

Branching List Scheduling Algorithms for the Identical Parallel Machine Scheduling Problem

Hakim Hadj-Djilani, Louis-Claude Canon, Laurent Philippe, Julien Bernard

Contacts: hakim.hadj-djilani@inria.fr / surname.name@umlp.fr

Abstract

This paper proposes new heuristics to the classic non-preemptive scheduling problem of assigning n jobs on m identical parallel machines with the objective to minimize the makespan. Starting from the List Scheduling method (LS), used for example in LPT [Graham, 1969] or SLACK [Della Croce & Scatamacchia, 2020] heuristics, we derive a branching strategy that also considers assigning a job to the second best machine, in addition to the first one, the only one usually dealt with in literature. Along the exploration of solutions, we keep only the best solution. This strategy leads to two heuristics, thought in an effort to speed up the solution search. The first, named Branch & Parallelize LS (BPLS), parallelizes the two alternatives considered on each job. The second, named BBLS, applies a Branch & Bound method to widely prune the solution tree. As a trade-off between computation time and solution quality, these heuristics are parameterized in order to assign only a subset of jobs according to the branching strategy. Out of this subset, the assignment is made by the classic LS, that is, each time on the first available machine. We show on literature instances that our heuristics outperform many of well-known algorithms on the vast majority of the considered instances. We also investigate the subspace of instances for which our heuristics are not able to beat all of these algorithms. Finally, we propose an ad hoc heuristic named MULTI-BBLS (MBBLS) which consists of multiple calls to BBLS that permit to rank first even on this instance subspace.

Keywords: identical parallel machines scheduling, heuristics, branch and bound, list scheduling

1 Introduction

Let us consider the identical parallel machine scheduling problem, denoted $P||C_{\max}$ according to Graham's three-field notation [Graham et al., 1979]: given m identical parallel machines $\{1 \leq i \leq m\}$ and a larger number of jobs $\{1 \leq j \leq n\}$ whose processing times are denoted $\{p_j\}_{j \in \{1, 2, \dots, n\}}$, we have to find a non-preemptive schedule such that the makespan $C_{\max}$ is minimum. More formally, the makespan is defined by $C_{\max} = \max_i \{C_i, 1 \leq i \leq m\}$, $C_i = \sum_{k_i} p_{k_i}$ being the i -th machine completion time for its assigned jobs k_i . $C_{\max}$ is hence the completion time for the execution of all jobs. Most of the time, as we do in this paper, the p_j values are considered as a set of positive integers¹. Finding the optimal solution, denoted $C_{\max}^*$, constitutes a strongly NP-hard problem in combinatorial optimization as described in [Garey & Johnson, 1979]. This problem has practical utility in industry and engineering and has been the focus of extensive research for more than half a century. Several approximation algorithms have been proposed in the literature.

A first family of algorithms is based on the List Scheduling algorithm (LS). These algorithms all work in two stages: the jobs are first sorted under a certain order then the LS algorithm is applied. The LS heuristic rule is to assign each job to the current least loaded machine, that is the first one available. The most famous of these algorithms is the Longest Processing Time First rule (LPT) due to Graham (see [Graham, 1969]), which sorts the jobs in non-increasing order of p_j 's. More recently, another LS algorithm, named SLACK, was published in [Della Croce & Scatamacchia, 2020]. SLACK proceeds in a bit more complicated way than LPT for the sorting. Firstly, it sorts also the jobs in non-increasing order of p_j 's. But then it splits the list in tuples of size m , completing the last tuple, if necessary, with null processing time jobs. Next, the tuples are sorted in non-increasing order of their slack, which is the absolute difference between the smallest and the largest p_j 's of the tuple. Finally, LS is applied to the job list formed by concatenating the sorted tuples.

Another family of heuristics developed to tackle the $P||C_{\max}$ problem is based on the bin-packing problem (BPP). In this problem, n objects of different finite sizes have to fit into bins of finite capacity or length L . The objective is to minimize the number of bins used. The First Fit Decreasing algorithm (FFD) [Coffman et al., 1978]

¹But note that this is made without loss of generality because floating/fixed-point numbers can be bijectively converted to integers, the $P||C_{\max}$ problem solved and then the resulting schedule converted back to floating/fixed-point numbers.

Table 1: Approximation ratio or bounds & Time complexity

Algorithm	Value / bounds	Time complexity
LPT [Graham, 1969]	$r^{LPT} = \frac{4}{3} - \frac{1}{3m}$	$O(n \log n)$
SLACK [Della Croce & Scatamacchia, 2020]	N.D.	$O(n \log n)$
MULTIFIT [Coffman et al., 1978]	$r^{MULTIFIT} \leq 1.22 + 2^{-k}$	$O(n \log n + k n \log m)$
COMBINE [Lee & Massey, 1988]	$r^{COMBINE} \leq r^{MULTIFIT}, r^{LPT}$	$O(n \log n + k n \log m)$
LISTFIT [Gupta & Ruiz-Torres, 2001]	$r^{LISTFIT} = 13/11 + 2^{-k}$	$O(n^2 \log n + n^2 k \log m)$
Hochbaum & Shmoys	$r^{PTAS} = 1 + \varepsilon$	$O((n/\varepsilon)^{\lceil 1/\varepsilon^2 \rceil})$
PTAS [Hochbaum & Shmoys, 1987]		
LDM [Michiels et al., 2003]	$\frac{4}{3} - \frac{1}{3(m-1)} \leq r^{LDM} \leq \frac{4}{3} - \frac{1}{3m}$ for $m \geq 3$ and $r^{LDM} = 7/6$ for $m = 2$	$O(n \log n)$

gives a bin-packing solution in two steps. First, it sorts the objects in non-increasing order of their size. Then, starting with only one open bin, it assigns the objects according to this order, one by one to the first bin it can fit into. If there is not enough space in all bins, then a new bin is opened. A form of duality exists between BPP and $P||C_{\max}$. Indeed, given a solution of BPP, as a number of m' bins, the capacity L is an upper bound to $C_{\max}^*$ for the $P||C_{\max}$ instance composed of as many jobs as objects in the BPP instance, with processing times p_j equal to the object sizes and with an assignment on $m \geq m'$ machines. This duality is the base of the MULTIFIT algorithm [Coffman et al., 1978]: a binary search of the $P||C_{\max}$ instance makespan, identified to the BPP capacity L , by iterating on the FFD algorithm to find out if it is an upper bound, i.e. $m \geq m'$. The COMBINE algorithm [Lee & Massey, 1988] came next as a combination of LPT and MULTIFIT. It first uses the former and, if its solution is not guaranteed to be optimal, it tries MULTIFIT with a search upper bound set to the makespan found by LPT. About a decade later, Gupta and Ruiz-Torres proposed LISTFIT in [Gupta & Ruiz-Torres, 2001]. Their approach is also based on FFD but tries $4n$ different specially forged orders of p_j 's as input of FFD. Thanks to its wider solution exploration, LISTFIT provides very often better solutions than MULTIFIT.

Of course myriad of other algorithms exist for $P||C_{\max}$. We can, for instance, cite the Polynomial Time Approximation Scheme proposed in [Hochbaum & Shmoys, 1987], which has some similarities with MULTIFIT in that it also uses the duality with BPP and a binary search, but turns to dynamic programming for the problem solving. The Largest Differencing Method (LDM) is another heuristic that allows solving a $P||C_{\max}$ equivalent partitioning problem ([Michiels et al., 2003], [Karmarkar & Karp, 1982]). Many metaheuristics as genetic algorithms [Min & Cheng, 1999], tabu search, simulated annealing [Glass et al., 1994], swarm optimization [Kashan & Karimi, 2009], harmony search [Chen et al., 2012] or ant colony optimization [Yibao et al., 2002], have also largely been used to search for local optima of $P||C_{\max}$ solutions.

Because they produce suboptimal solutions, $P||C_{\max}$ heuristics, are evaluated according to their approximation ratio: $r^A \geq \frac{C_{\max}^A(I)}{C_{\max}^*(I)}$, with A an algorithm to be evaluated on any instance I , $C_{\max}^A$ the algorithm makespan and $C_{\max}^*$ the optimal solution. For any algorithm A , we have $r^A \geq 1$. When possible, a worst-case analysis is made to define exact value or upper bounds, as tight as possible, for an algorithm approximation ratio. Table 1 lists known approximation ratios and computational complexities of several well-known algorithms. Alternatively, evaluating and comparing the performance of algorithms is made through empirical protocols using pseudo-random instances defined in the literature.

The main contribution of this paper is the Branch & Bound List Scheduling algorithm (BBLS). Because LS-based heuristics concentrate on the job list order and restricts the assignment of jobs to the first available machine, we propose with BBLS to consider assigning jobs not only to the first available machine but also to the next ones. The Branch & Bound part of BBLS is highly linked to the research introduced in [Dell'Amico & Martello, 1995] which aims at producing exact/optimal $P||C_{\max}$ solutions considering basically all the m machines for a job assignment. On the other hand BBLS is more about searching approximate solutions in a smaller amount of time and hence limits the number of considered machines. To this end, it leverages many elements of optimization that are developed and assessed in this paper.

This article is organized in five sections. In Section 2, we examine the LS heuristic and describe the branching strategy of a first LS variant, called BLS for Branch LS. BLS is a first step toward BBLS. It represents the alternative job assignments on the first available machines as a tree and then searches for the best solution. We show how this strategy can easily be parallelized into an algorithm named BPLS (for Branch & Parallelize LS). Then, in Section 3, pursuing an effort to speed up these variants, we derive BBLS that is able to prune the tree of solutions by testing their makespan lower bounds as in [Dell'Amico & Martello, 1995]. Besides,

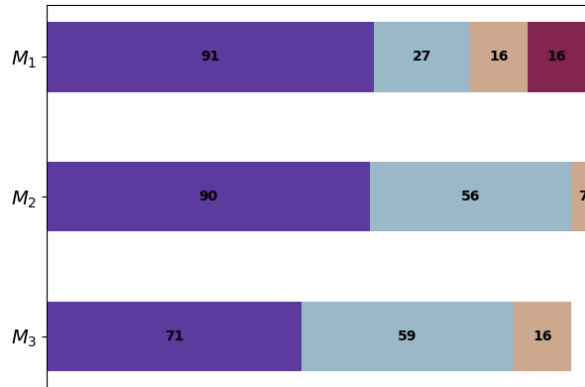
we introduce several optimization properties and parameters to speed up the tree pruning. Next, in Section 4 we define MULTI-BBLS (or MBBLs) methods that mainly consist in calling BBLs several times in parallel in order to enhance the quality of solutions. Finally, in Section 5, we set up an empirical protocol based on literature instances [Gupta & Ruiz-Torres, 2001], [Della Croce & Scatamacchia, 2020] to show how BPLS and BBLs can outperform the renowned algorithms presented before. Section 6 comes as a complement to see how the MBBLs approaches can be combined together to beat all the tested algorithms for subgroups of instances on which BBLs and BPLS alone are not able to dominate.

2 The BLS and BPLS algorithms

As previously explained, the LS algorithm consists in assigning each job to the current least loaded machine. In many cases however, it can be observed that some of the machines are not much more loaded than the first one. Hence it is meaningful to assess the impact of also considering several least loaded machines instead of only the first one. On the other hand, it is worth noticing that the exploration with several machines at each step of the assignment quickly leads to a combinatorial explosion. Furthermore, tests made by considering the third least loaded machine showed to be less efficient when adding the pruning process². We hence only explore alternative of the second least loaded machine in the remainder of the paper.

Let us start with a simple example, before explaining how an algorithm based on this idea works. Consider the following instance $I = \{P = (91, 90, 71, 59, 56, 27, 16, 16, 16, 7), m = 3\}$ as input of LPT. The sorting stage of P is already done in non-increasing order. Then, with LS, comes the assignment of the jobs whose processing times are listed in P to obtain the schedule shown in Figure 1. If we introduce an exception to the LS rule, by allowing once to assign a job to the second least loaded machine, we obtain the assignment illustrated in Figure 2, which happens to be an optimal solution for I . This short example shows that exploring solutions with the two least loaded machines has the potential to find better results than LS or even optimal solutions for certain instances.

Figure 1: LPT schedule for instance $I = \{P = (91, 90, 71, 59, 56, 27, 16, 16, 16, 7), m = 3\}$ the solution makespan is $C_{\max}^{LPT} = 153$



We can now define the Branch LS algorithm (BLS). Basically, it is a recursive algorithm that concurrently tries, for each job, an assignment on the two least loaded machines. Once the last job is assigned, the two job assignment alternatives are compared, at each level of recursion, according to their solution makespans in order to keep the best solution. This approach is simply a solution search structured in a binary tree. The left child node represents an assignment on the first least loaded machine, while the right one is an assignment on the second least loaded machine. Notice that there is no interest to assess the two alternatives for the m first jobs that are always assigned one on each machine. The last job of the list is likewise assigned according to LS because this is always the best choice.

Unfortunately, this approach has an exponential computational cost of $C(n, m) = \Theta(m2^{n-m})$. A first step for reducing this cost is to consider that not only the m first jobs can be excluded from the binary tree but also several other jobs since, as can be seen on the example of Figure 2, several jobs are assigned to the first least loaded

²We give more details about this question and our related tests in 5.2.1. The pruning process is discussed in 3.2

Figure 2: An optimal schedule for instance in Fig. 1. The makespan is $C_{\max} = 150$. The black framed job represents an assignment on the second least loaded machine, other assignments are all made on the first available machine (wrt LS)



machine in the optimal solution. Knowing which jobs to include seems complicated but limiting to a parameter, denoted N , the number of consecutive jobs included in the binary tree can achieve a linear asymptotic complexity since 2^N is a constant. The BLS complexity is then defined by Equation (1).

$$C(n, m) = O(2^N m) + O(n - N) = O(n) \quad (1)$$

Note that, even if 2^N is just a multiplicative constant, it should not be underestimated as it is an exponential term.

The BLS heuristic, given as Algorithm 1, recursively builds the schedule, denoted σ for an input instance defined by $(P = \{p_j\}_{1 \leq j \leq n}, m)$. It starts (line 3) from an empty σ , assigns the first m jobs $j \in \{1, \dots, m\}$ according to LS and makes a recursive call (line 7) to enter in the branching part of the algorithm. Once σ is initialized the two least loaded machines, i_1 and i_2 ³, are identified (lines 9, 10). The two following conditions (lines 11, 12) assert that at least two jobs remain to assign (remember that the machines are not challenged for the last jobs) and less than $N + m$ jobs are already assigned into σ . The algorithm then takes the first of the remaining jobs, whose processing time is p_1 , assigns it alternatively on i_1 and i_2 by updating respective schedules σ'_1 and σ'_2 (lines 13, 14). Then (lines 15, 16) it makes two BLS recursive calls responsible for assigning the remaining jobs into σ'_1 and σ'_2 according to the same recursive process. On each recursive call N is decremented and P loses its first processing time. Going ahead, N becomes zero or P becomes a singleton. In both cases BLS switches back to LS algorithm for assigning the remaining jobs: jobs $m + 1 + N$ to $n - 1$, lines 20-22 and last job, line 24. Afterward, on lines 17 and 18, the makespans of the two alternative schedules σ'_1 and σ'_2 are compared, keeping the best makespan solution. This process continues toward parent calls until the recursion level of the first job assigned is reached back to finally return the BLS overall best solution found.

It is worth noticing that, to keep it simple, the branching strategy is here applied starting from job $m + 1$ to job $m + N$. For more generality, the N jobs could have been picked in the remainder of the job list⁴. Note that parameter N should not be greater than $(n - m - 1)$ because, as already mentioned, the m first jobs and the last one are assigned according to LS.

2.1 BPLS: a parallelization of BLS

Due to the BLS algorithmic complexity it is worth proposing solutions to improve its running time on large instances. Since BLS is considering many alternative solutions, it is possible to parallelize the search. Note that this is not the case for LPT and SLACK that produce a single solution for which there is no possible parallelization regarding the dependency between the algorithm steps. A basic parallelization can be done on the two alternatives (i_1, i_2 choices) using two processes, each one computing one of the assignments made according to the branching strategy. We name BPLS the parallelization variant of BLS (for Branch & Parallelize LS). An example is given in Figure 3 for a parallelization using a total of four processes $\text{Proc}_{k \in \{1, \dots, 4\}}$ to produce a schedule when $N = 2$. The number of processes must be limited to at most 2^N , which is the number of leaves of the binary tree, above what there is no interest to parallelize. There is likewise and obviously no point to

³Notice that we should respectively name $i_{1,j}$ and $i_{2,j}$ the two least loaded machines at the assignment time of the j -th job but, when there is no ambiguity, we only denote them i_1 and i_2 , without any mention of the next job to assign.

⁴We work toward this variation with BBLS and its parameter S in 3.3

Algorithm 1: BLS(P, m, σ, N)

Input:
 m : integer
// p_j : processing time of job j
 $P \leftarrow \{p_j\}_{j \in \{1, \dots, n\}}$
// $1 \leq \sigma[i][k] \leq n$: k -th job on machine i
 $\sigma \leftarrow \{()\}_{i \in \{1, \dots, m\}}$: schedule
// default complexity limit
Optional parameter: $N \leftarrow \infty$: integer

Data:
 i_1, i_2 : integer

```

1 begin
2   if empty( $\sigma$ ) then
3      $\sigma \leftarrow \text{assign}(\sigma, 1, 1)$ 
4      $\sigma \leftarrow \text{assign}(\sigma, 2, 2)$ 
5     ...
6      $\sigma \leftarrow \text{assign}(\sigma, m, m)$ 
7      $\sigma' \leftarrow \text{BLS}(P - \{p_1, \dots, p_m\}, m, \sigma, N)$ 
8     return  $\sigma'$ 
9    $i_1 \leftarrow \arg \min_{i \in \{1, \dots, m\}} \sum_{1 \leq k \leq \text{size}(\sigma[i])} (p_{\sigma[i][k]})$ 
10   $i_2 \leftarrow \arg \min_{i \in \{1, \dots, m\} - \{i_1\}} \sum_{1 \leq k \leq \text{size}(\sigma[i])} (p_{\sigma[i][k]})$ 
11  if size( $P$ ) > 1 then
12    if  $N > 0$  then
13       $\sigma'_1 \leftarrow \text{assign}(\sigma, i_1, 1)$ 
14       $\sigma'_2 \leftarrow \text{assign}(\sigma, i_2, 1)$ 
15       $\sigma'_1 \leftarrow \text{BLS}(P - \{p_1\}, m, \sigma'_1, N - 1)$ 
16       $\sigma'_2 \leftarrow \text{BLS}(P - \{p_1\}, m, \sigma'_2, N - 1)$ 
17      if makespan( $\sigma'_2$ ) < makespan( $\sigma'_1$ ) then return  $\sigma'_2$ 
18      else return  $\sigma'_1$ 
19    else
20       $\sigma' \leftarrow \text{assign}(\sigma, i_1, 1)$ 
21       $\sigma' \leftarrow \text{BLS}(P - \{p_1\}, m, \sigma', 0)$ 
22      return  $\sigma'$ 
23  else
24    return  $\text{assign}(\sigma, i_1, 1)$ 
```

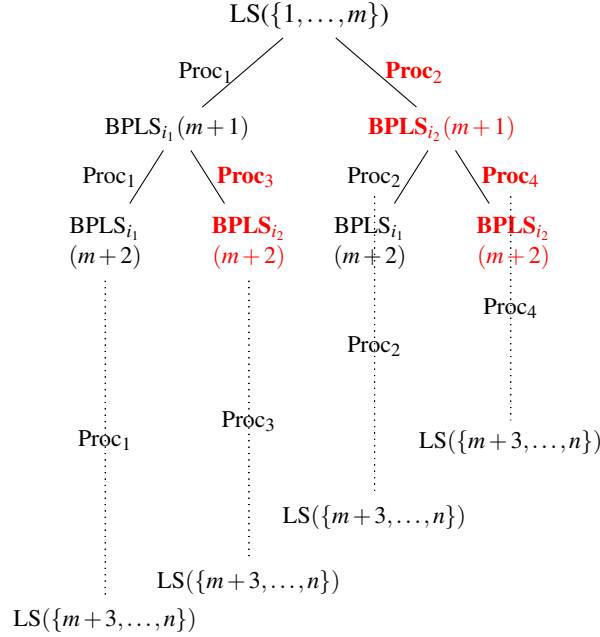
parallelize using more processes than the number of threads the CPU used is capable to run efficiently in parallel or concurrently. When the number of threads T used is lower than 2^N , the basic strategy is to evenly assign the tree branches to the threads, that is, with $\frac{2^N}{T}$ branches per thread. A basic experiment has, for example, shown that running BPLS with at most 2^7 processes on instances defined as $n = 2^{12}$, $m = 2, 4, 8, 16$ and $N = 10$ can generate a median speedup of 28 (from a time of approximately 12.5 sec for BLS to only 0.44 sec for BPLS⁵). More elements about both the computation time and performance of the makespan minimization are presented in Section 5. About the performance, we can already state that BLS and BPLS cannot produce worse solutions than LS because the LS solution is included in the solutions evaluated by the branching strategy and kept as final solution if it is the best.

2.2 A branching strategy for all LS-based algorithms

A primary interest of this branching strategy used in BLS and BPLS algorithms is that it can be used in any LS algorithm. Thereby we can easily write a B-SLACK or a BP-SLACK algorithm or likewise a B-LPT or a BP-LPT algorithm. The job assignment stage, in the LS algorithm, has only to be replaced by one of the branching

⁵On an Intel(R) Xeon(R) Platinum 9242 equipped of 48 cores, that is 96 threads

Figure 3: *Parallelization tree of BPLS ($N = 2$) using 4 processes $Proc_k$. Left and right BPLS nodes represent respectively i_1 and i_2 alternative machine assignments performed each time on two separate processors. The classic LS heuristic is used for the first jobs indexed from 1 to m and also for the last jobs indexed from $m + 3$ to n . The two jobs in between are assigned using BPLS.*



LS algorithms. This is also true for the BBLS algorithm presented in the next section which has been used to write a BB-SLACK variant.

3 A Branch & Bound LS algorithm

As previously mentioned, BLS and BPLS algorithms explore the whole binary tree for assigning N jobs, each time on the two first available machines. On the other hand, there is no interest to continue exploring toward a node if the best makespan found so far in the exploration is a lower bound for the subset of solutions this node is encoding. That remark leads us straight right to branch-and-bound (BB) algorithms. In fact, BLS is half way to be a BB algorithm because it already provides the branching part but not the bounding one. The principle for a Branch and Bound LS algorithm (BBLS) is, as for BLS, that any node in the tree represents a partial schedule and child nodes encode the solution subsets of the parent node. One job is added in the partial schedule at each level of the tree whether on machine i_1 or on machine i_2 . In addition to that process, a lower bound of the subset of solutions represented by a node can be computed on the fly to find out if this solution has a chance to give a better makespan than the current best one. So the algorithm starts by exploring the tree with the LS solution as initial best makespan and, for each node, it computes a lower bound for the corresponding solution. If this lower bound is greater than the current best makespan, then there is no interest to go further on this branch, since all child solutions are bounded with the same lower limit, and the branch is pruned out of the tree. Otherwise the algorithm continues exploring and, every time it reaches a leaf, it compares the makespan of this particular schedule to the current best makespan. If the new solution has a lower makespan it becomes the new best solution for the next of the exploration. Such a method can really speed up the tree exploration by widely pruning the tree.

3.1 $P||C_{\max}$ Lower bounds

To implement BBLS, especially the pruning strategy discussed above, we need lower bounds as tight as possible to the optimal makespan. Simple lower bounds can be established on any $P||C_{\max}$ instance, whatever algorithm is used. We present here the lower bounds already described in [Dell'Amico & Martello, 1995]. The simplest

one, L_0 , is obtained by a preemptive relaxation of the $P||C_{\max}$ instance:

$$L_0 = \frac{1}{m} \sum_{j=1}^n p_j$$

L_0 formula means that this is not possible to obtain a smaller makespan than the one obtained by considering that the jobs can be preempted. Indeed, if the jobs can be split, the resulting fragments can always be spread equitably to the machines, such that the completion time is the same on all machines: $C_i = C_{i+1}, \forall 1 \leq i < m$. This completion time is the optimal makespan.

A second lower bound comes from the non-preemptive constraint of the problem which implies that the makespan of any $P||C_{\max}$ instance cannot be lower than any p_j . This lower bound is

$$L_1 = \max(L_0, \max\{p_j\})$$

A third bound, more precise when not equal to L_1 or L_0 , is obtained by another problem relaxation which consists in reducing the instance to the $(m+1)$ jobs of largest processing times. Assume that the processing times $(p_j)_{j \in 1, \dots, n}$ are sorted in non-increasing order, then the jobs of processing times p_m and p_{m+1} must be assigned to the same machine to get an optimal solution for the relaxed instance. Indeed, the machine running the job m will be the first available machine at the moment of assigning job $(m+1)$. Therefore we have the following bound:

$$L_2 = \max(L_1, p_m + p_{m+1})$$

Another lower bound, denoted L_3 , is proposed in [Dell'Amico & Martello, 1995]. To introduce this bound we recall that the bin-packing problem (BPP) is dual to $P||C_{\max}$. Dell'Amico and Martello limit the number of bins of the BPP obtained by a dual reformulation of a $P||C_{\max}$ instance I into an instance I' of BPP. We denote $m'(L, \{p_j\}_{1 \leq j \leq n})$, a solution of the bin packing problem, i.e. a number of bins for I' with L the bin capacity and p_j the item sizes, that come directly from I . If an optimal solution $m^*(L, \{p_j\})$ is such that $m^*(L, \{p_j\}) > m$, m being the number of machines in I , then L is obviously a lower bound of the optimal makespan of I . Since the p_j 's are integers, it is also clear that $L+1$ is also a valid lower bound for I . That is what Dell'Amico and Martello did using two possible lower bounds of the optimal solution of I' with a certain capacity L to prove that $L+1$ is a lower bound for I . The two lower bounds are denoted B_α and B_β . The first one was introduced in [Martello & Toth, 1990] and the second one in [Dell'Amico & Martello, 1995]. Their definitions are based on a subset of processing times $\{\bar{p}\} = \{p_j \leq \frac{L}{2}\}$ and other subsets denoted J_1 , J_2 and J_3 . These three subsets form together a partition of all jobs whose processing times are greater or equal to a given $\bar{p}$. They are defined by:

$$J_1 = \{j | L - \bar{p} < p_j\} \quad (2)$$

$$J_2 = \{j | L/2 < p_j \leq L - \bar{p}\} \quad (3)$$

$$J_3 = \{j | \bar{p} \leq p_j \leq L/2\} \quad (4)$$

The two bounds are then defined as follows for a given $\bar{p}$ and a given bin capacity L :

$$B_\alpha(L, \bar{p}) = |J_1| + |J_2| + \max(0, \left\lceil \frac{\sum_{j \in J_3} p_j - (L|J_2| - \sum_{i \in J_2} p_i)}{L} \right\rceil) \quad (5)$$

$$B_\beta(L, \bar{p}) = |J_1| + |J_2| + \max(0, \left\lceil \frac{|J_3| - \sum_{j \in J_2} \left\lfloor \frac{L - p_j}{\bar{p}} \right\rfloor}{\left\lfloor \frac{L}{\bar{p}} \right\rfloor} \right\rceil) \quad (6)$$

The idea of these bounds is that no jobs in J_1 or J_2 can be in the same bin because they all exceed the processing time $\frac{L}{2}$. So one bin is used for each one of these jobs. The remaining J_3 jobs cannot fit in bins that already contain a J_1 job. Since this type of job has a p_j greater than $L - \bar{p}$, this does not leave enough space to let a J_3 job fit into the bin. So the bounds B_α and B_β , in their right 'max' term, take both account of the minimum number of new bins that are necessary to make all the J_3 jobs fit into the bins and form a full bin packing. Precisely, some of the J_3 jobs might fit into bins that already contain one J_2 job, if there is such. But if other J_3 jobs remain then the opening of at least one new bin is needed. In B_α the J_3 jobs are considered preemptively while in B_β the fit is non-preemptive but the J_3 jobs processing times are all reduced to $\bar{p}$. Finally, it is stated that, consistently with the dual relationship between $P||C_{\max}$ and BPP, if

$$B_\alpha(L, \bar{p}) > m \text{ or } B_\beta(L, \bar{p}) > m \quad (7)$$

then $(L+1)$ is a lower bound for I . In [Dell'Amico & Martello, 1995] the bound named L_3 is the maximum capacity that verifies the property (7) but in this paper we denote L_3 any bound that verifies this property because, as detailed in the next subsection, it is enough to justify a pruning.

For a single value L , it costs up to $O(n^2)$ to compute if $L + 1$ is a lower bound for a $P||C_{\max}$ instance. This is a considerable cost but, hoping that the tree of solutions will be massively pruned in BBLS, the lower bound computation could be worth it. This is empirically verified in Section 5. Besides, in Appendix J, we give new properties or recall already known ones, all purposed for optimization of bounds calculation. However, it should be noticed, as summed up in Table 2, that not all of those properties have been practically efficient in our experimental framework and implementations.

3.2 The BBLS algorithm and its pruning process

Algorithm 2: Function *bb_node_instance* Computes a tree node transformed instance.

- P is the p_j 's of jobs yet to assign,
 - σ is the node partial schedule.
-

```

1 Function bb_node_instance( $\sigma, P, m$ )
2 begin
3   //  $P$  does not contain any  $p_j$  of job already assigned in  $\sigma$ 
4   for  $j \leftarrow 1, \dots, m$  do
5      $p'_j \leftarrow \sum_{1 \leq k \leq \text{size}(\sigma[j])} (p_{\sigma[j][k]})$ 
6   for  $j \leftarrow 1, \dots, \text{size}(P)$  do  $p'_{m+j} \leftarrow P[j]$ 
7   return  $(p'_j)_{1 \leq j \leq m + \text{size}(P)}$ 

```

The condition (7) now allows to complete the bounding part of the BBLS algorithm envisioned in the beginning of this section, i.e. pruning nodes for which the best makespan found so far, denoted $C_{\max}^b$, is larger than the node lower bound. Considering a given node, we can proceed, as in [Dell'Amico & Martello, 1995], by transforming the problem instance I to another instance I' whose solutions are precisely the subset of solutions encoded by this node. The *bb_node_instance* function, given as Algorithm 2, builds the transformed instance of a BBLS tree node.

The jobs previously assigned to the m machines are replaced by m jobs where p'_j ($1 \leq j \leq m$), the processing time of job j in the transformed instance I' , is equal to the sum of the processing times of the jobs already assigned to the j -th machine. These jobs are those assigned by the parent nodes that led to the considered node. From the way they are created, we name these m jobs the *aggregated jobs*. These jobs processing times are calculated in line 4 of *bb_node_instance*. The jobs that were not assigned yet to any machine are simply added as is in the transformed instance (line 6).

Algorithm 3: Function *is_L3_bound*

Tests if L is a L_3 lower bound for $I' = (P, m)$

```

1 Function is_L3_bound( $L, P, m$ )
2 begin
3    $L \leftarrow L - 1$ 
4   // use (7) on  $I' = (P, m)$  for all  $0 < \bar{p} \leq L/2$ 
5   for  $\bar{p} \in \{p_j | j \in \{1, \dots, \text{size}(P)\}, p_j \leq L/2\}$  do
6      $J_1 \leftarrow \{j | L - \bar{p} < p_j\}$ 
7      $J_2 \leftarrow \{j | L/2 < p_j \leq L - \bar{p}\}$ 
8      $J_3 \leftarrow \{j | \bar{p} \leq p_j \leq L/2\}$ 
9      $B_\alpha(L, \bar{p}) \leftarrow |J_1| + |J_2| + \max(0, \left\lceil \frac{\sum_{j \in J_3} p_j - (L|J_2| - \sum_{j \in J_2} p_j)}{L} \right\rceil)$ 
10     $B_\beta(L, \bar{p}) \leftarrow |J_1| + |J_2| + \max(0, \left\lceil \frac{|J_3| - \sum_{j \in J_2} \left\lfloor \frac{L - p_j}{\bar{p}} \right\rfloor}{\left\lfloor \frac{L}{\bar{p}} \right\rfloor} \right\rceil)$ 
11    if  $B_\alpha(L, \bar{p}) > m \vee B_\beta(L, \bar{p}) > m$  then return true
12  return false

```

Next, property (7) has to be tested to verify that the current best bound $C_{\max}^b$ is a lower bound of the transformed instance. The function *is_L3_bound*, given as Algorithm 3, is used in this purpose. If this is the case,

then no better solution can be found by considering this instance node and further child nodes. The node is hence pruned out. On the contrary, if $C_{\max}^b$ is not a lower bound for the transformed instance then the exploration continues on the child nodes. Note that $C_{\max}^b$ can be tested first against L_2 which is less expensive to compute than L_3 . If $C_{\max}^b > L_2$ then *is_L3_bound* is used to test if $C_{\max}^b$ is a L_3 lower bound. Besides, since L_2 was tested before, the set of $\bar{p}$ tested in *is_L3_bound* can be limited to the jobs that verify the condition $\bar{p} \leq p_{m+2}$ as proved in [Dell’Amico & Martello, 1995]. It implies however to sort the p_j ’s in non-increasing order.

At this point BBLS, given as Algorithm 4, is almost a complete Branch & Bound algorithm. As mentioned before, the classic LS solution, that assigns each task to the first available machine, is used for the initialization. In the exploration tree it means that, for building the first complete solution, the left child node is always chosen until the leftmost leaf is reached. This is made very naturally because the solution search is a depth-first search. For this reason, in Algorithm 4, line 12, the left child node is never ignored as long as $C_{\max}^b = \infty$ (its initial value). Notice that Algorithm 4 is very similar to Algorithm 1. Indeed, only a few parts change compared to BLS, mainly for the pruning implementation integrated in lines 12 to 30. In this block, precisely lines 12 and 17, $C_{\max}^b$ is tested as a lower bound of the transformed instance corresponding to each alternative assignment of the current job (to machine i_1 or i_2). Hence, if $C_{\max}^b$ is confirmed to be a lower bound, the node is pruned out, that is no further recursive call is made for the corresponding partial schedule (σ'_1 or σ'_2) which is marked as pruned out using the empty set $\emptyset$. This backtracking process comes as a complement to picking the best solution of the two alternatives as already mentioned for BLS. In the following subsection is described the parameter S which is an optimization of BBLS that is not present in BLS.

3.3 A shift parameter S indexing the first job to start BBLS on

In BLS (Algorithm 1) and also BBLS (Algorithm 4), the N parameter defines the number of jobs on which the branching is performed, other jobs being assigned according to the classic LS. This parameter however does not indicate the job of the list where to start the branching. The only constraints induced by N are that the job start index is strictly greater than m and smaller than or equal to $(n - N)$. For this reason, we introduce another BBLS parameter, named S , a positive integer that defines a shift of job index on which the branching strategy is started. Hence, BBLS uses the job index $m + S + 1$ to start the branching strategy for N jobs (line 3, jobs $j = 1$ to $j = m + S$ are assigned according to the classic LS). Note that a meaningful S must be such that $0 \leq S \leq n - N - m - 1$. Regarding the computational cost, the following properties indicate that there is an interest into picking a larger value of S . The Proposition 3.1 shows that a greater S allows a smaller computational cost for *is_L3_bound* which represents the quadratic dominant cost of BBLS.

Proposition 3.1. *The computational cost of is_L3_bound is a decreasing function of S .*

Proof. At the right child node of the BB tree root, the size of BBLS first transformed instance I' produced by *bb_node_instance* is $n' = n - S - 1$. The dominant term of *is_L3_bound* complexity being n'^2 , this is enough to conclude. More precisely, note that with a larger S and same bound L , the number of $\bar{p}$ candidates to evaluate, that is the subset $\{p'_j \leq L/2 | L = C_{\max}^b - 1\}$ can only stay equal or shrink. Furthermore, this subset for a child node cannot be larger than the parent’s, because already assigned jobs are aggregated in the transformed instance jobs. $\square$

To test our assertion about the speedup obtained with a greater S used in BBLS we have defined a protocol introduced in Appendix I. Complementary properties about wider tree pruning in relation to larger S are discussed in Appendix F.

4 MULTI-BBLS: parallel runs of BBLS

In Section 2 we present BLS and its parallel version BPLS. Then, in Section 3, we introduce BBLS, an improvement of BLS that uses pruning to reduce the solution search tree. Doing so we free computational power that we can use to enhance the BBLS solutions at a small additional cost. We name this basic idea MULTI-BBLS (or MBBLs) which consists generally in running BBLS multiple times, mostly in parallel.

We devise three specific strategies in this purpose. They are described in Sections 4.1, 4.2, 4.3. Each one of these strategies is not necessarily more efficient than a simple BBLS for all instances. Interesting use cases are presented in Section 6 where it is shown that combining these multiple calls of BBLS allows to outperform other heuristics on instances for which all our tested configurations of a single BBLS have been unable to do so (see Section 5).

Algorithm 4: BBLS(P, m, σ, N, S)

Input: $P \leftarrow \{p_j\}_{j \in \{1, \dots, n\}}$ m : integer $\sigma \leftarrow \{()\}_{i \in \{1, \dots, m\}} : \text{schedule} \ // \ 1 \leq \sigma[i][k] \leq n$: k -th job on machine i *Optional parameter:* $N \leftarrow \infty$: integer // default value (complexity limit)*Optional parameter:* $S \leftarrow 0$: integer // job index shift for branching strategy start**Data:** $C_{\max}^b \leftarrow \infty$ // global variable i_1, i_2 : integer**1 begin****2 if** *empty*(σ) **then****3** $\sigma \leftarrow LS(\{p_1, \dots, p_{m+S}\}, m)$ **4** $\sigma' \leftarrow \text{BBLS}(P - \{p_1, \dots, p_{m+S}\}, m, \sigma, N)$ **5** **return** σ' **6 else****7** $i_1 \leftarrow \arg \min_{i \in \{1, \dots, m\}} \sum_{1 \leq k \leq \text{size}(\sigma[i])} (P\sigma[i][k])$ **8** $i_2 \leftarrow \arg \min_{i \in \{1, \dots, m\} - \{i_1\}} \sum_{1 \leq k \leq \text{size}(\sigma[i])} (P\sigma[i][k])$ **9 if** *size*(P) > 1 **then****10 if** $N > 0$ **then****11** $\sigma'_1 \leftarrow \text{assign}(\sigma, i_1, 1)$ **12 if** $C_{\max}^b = \infty \vee \neg \text{is_L3_bound}(C_{\max}^b, \text{bb_node_instance}(\sigma'_1, P, m), m)$ **then****13** $\sigma'_1 \leftarrow \text{BBLS}(P - \{p_1\}, m, \sigma'_1, N - 1)$ **14 else****15** $\sigma'_1 \leftarrow \emptyset$ // elimination of the i_1 node**16** $\sigma'_2 \leftarrow \text{assign}(\sigma, i_2, 1)$ **17 if** $\neg \text{is_L3_bound}(C_{\max}^b, \text{bb_node_instance}(\sigma'_2, P, m), m)$ **then****18** $\sigma'_2 \leftarrow \text{BBLS}(P - \{p_1\}, m, \sigma'_2, N - 1)$ **19 else****20** $\sigma'_2 \leftarrow \emptyset$ // elimination of the i_2 node**21** // Keep the best solution between σ'_1 and σ'_2 or $\emptyset$
22 // if the two nodes were pruned out**23 if** $\sigma'_1 = \emptyset \wedge \sigma'_2 = \emptyset$ **then return** $\emptyset$ **24 else if** $\sigma'_1 = \emptyset$ **then return** σ'_2 **25 else if** $\sigma'_2 = \emptyset$ **then return** σ'_1 **26 else****27** **if** *makespan*(σ'_2) $<$ *makespan*(σ'_1) **then return** σ'_2 **28** **else return** σ'_1 **29 else****30** $\sigma' \leftarrow \text{assign}(\sigma, i_1, 1)$ **31** $\sigma' \leftarrow \text{BBLS}(P - \{p_1\}, m, \sigma', 0)$ **32** **return** σ' **33 else****34** $\sigma' \leftarrow \text{assign}(\sigma, i_1, 1)$ **35 if** *makespan*(σ') $<$ $C_{\max}^b$ **then** $C_{\max}^b \leftarrow \text{makespan}(\sigma')$ **36 return** σ'

Table 2: : Propositions in this paper

Proposition	Contribution	Used in implementation
3.1, F.1-F.3	New	Yes, as arguments for a larger S .
J.1	[Martello & Toth, 1990]	Yes
J.2, J.3	Extension of [Martello & Toth, 1990]	Yes for J_2 , no for J_3 .
J.4	New	Yes but no significant speedup observed.
Bound (9)	New	No, might slow down BBLS.

4.1 MBBS1: trying different orders of the job list

As mentioned previously, BBLS, as any other LS algorithm, assigns jobs in a specific order. However, trying different orders does not give the same result not only for BBLS but also for a classic LS like LPT or SLACK. LISTFIT solutions are also more efficient than MULTIFIT's because LISTFIT tries much more orders of the n jobs than MULTIFIT, which only tries one. The idea for MBBS1 is then to leverage the power of a multi-core CPU to run BBLS on different job list orders. Because these runs are parallel it should not take more time than only one BBLS run, except if we add more than one supplementary order per CPU core. If needed one can decide to try more orders to pursue better solutions but it goes obviously with an increase of the computational cost.

A question remains about what orders should be tried. We denote $\mathcal{O}$ the number of orders we want to try. We already know that the LPT and SLACK orders produce good solutions for many instances. So, if $\mathcal{O} \geq 2$, we always pick both of them. If $\mathcal{O} > 2$ we simply try to pick additional orders that are as different as possible because we do not have any information about which order is more efficient for a given instance. Indeed, it does not seem very wise to try very similar orders. Following this idea, we start from the LPT order and additionally try, as input of BBLS, different job permutations of this initial order. More precisely, let $\mathcal{P}_{LPT}$, the ordered set of all the permutations of jobs j , whose first permutation is the LPT order, we denote $d = \frac{|\mathcal{P}_{LPT}|}{\mathcal{O}-1}$ the index difference between two permutations we want to try as input of BBLS (including the LPT order). We then define $\pi(k)$, the function that gives the k -th permutation in $\mathcal{P}_{LPT}$ sorted in lexicographic order. The LPT order is the permutation $\pi(1)$, defined by $\pi(1) = 1, 2, \dots, n$, assuming that $p_1 \geq p_2 \geq \dots \geq p_n$. The other permutations picked are $\pi(\lfloor d \rfloor), \pi(\lfloor 2 * d \rfloor), \dots, \pi((\mathcal{O}-1)d = |\mathcal{P}_{LPT}|)$. We hence have a number of $\mathcal{O}$ evenly spaced permutations from the ordered set $\mathcal{P}_{LPT}$. Because of the lexicographic order, $\pi((\mathcal{O}-1)d)$, the last permutation picked is the so called SPT order (for Shortest Processing Time first), $\pi((\mathcal{O}-1)d) = n, n-1, \dots, 1$. Note that it does not need to build the entire set $\mathcal{P}_{LPT}$ to extract a number of $(\mathcal{O}-1)$ permutations. The permutation extraction complexity is considered $O(1)$.

This approach is written in Algorithm 5, in which three functions must be explained.

- $P \mapsto P_{order} = \text{order}(P, \text{order})$ is the function that takes a list of processing times P and returns P_{order} , the list ordered in a specific *order* (e.g. 'LPT', 'SLACK' or 'SPT').
- $\text{parallel}(\text{function}(\text{args}))$ is a non-blocking function that performs the parallel execution, as a thread or a process, of a *function* configured with arguments *args*.
- $\text{waitall}()$ blocks the current program until all *parallel* calls made before has ended their execution.

4.2 MBBS2: one BBLS per subgroup of jobs

Another way to run multiple BBLS for a $P||C_{\max}$ instance $I = (P = \{p_j\}_{j \in \{1, \dots, n\}}, m)$ is to divide the set of n jobs into a partition of K job subsets of sizes $n_1, \dots, n_K$ such that $n_1 + \dots + n_K = n$. Then BBLS is performed on each subset so that we obtain K different sub-schedules on m machines. These BBLS executions can be totally parallel if K is lower or equal to the number of CPU cores. Finally, these sub-schedules are grouped together to obtain a full schedule for the original instance I . This approach is presented as Algorithm 6. For the grouping part (lines 15-16), we consider I' , another $P||C_{\max}$ instance, formed from the K sub-schedules. Each sub-schedule gives m jobs. One job $p'_{k,i}$ is obtained by summing the processing times p_j of the jobs assigned to a machine i . Hence I' has a number of $K \times m$ jobs. A last BBLS execution is performed on I' (line 19) then the resulting schedule is transformed back to the original p_j 's of I by splitting the $p'_{k,i}$ accordingly (lines 21-24) in order to reverse the summing process made to form I' . This conversion is done by the *replace* function. It replaces the merger job of index (k, i) of instance I' set in schedule σ' by the corresponding jobs of original instance I whose p_j were

Algorithm 5: MBBL1($P, m, N, S, \mathcal{O}$)

Input:

$P \leftarrow \{p_j\}_{j \in \{1, \dots, n\}}$

$m : \text{integer}$

Optional parameter : $N \leftarrow \infty : \text{integer}$

Optional parameter : $S \leftarrow 0 : \text{integer}$

Optional parameter : $\mathcal{O} \leftarrow 1 : \text{integer}$

Data:

$C_{\max}[\mathcal{O}] : \text{integer}$

$\sigma[\mathcal{O}] \leftarrow \{()\}_{i \in \{1, \dots, m\}} : \text{schedule}$

$k \leftarrow 2 : \text{integer}$

```
1 begin
2    $P_{LPT} \leftarrow \text{order}(P, 'LPT')$ 
3    $\sigma[1], C_{\max}[1] \leftarrow \text{parallel}(BBL1(P_{LPT}, m, \sigma[1], N, S))$ 
4    $d \leftarrow \frac{|P_{LPT}|}{\mathcal{O}-1}$ 
5   while  $k < \mathcal{O}$  do
6     if  $k = \mathcal{O} - 1$  then  $\pi_k \leftarrow \pi(|P_{LPT}|)$ 
7
8     else  $\pi_k \leftarrow \pi(\lfloor d * (k - 1) \rfloor)$ 
9
10     $P_k \leftarrow P_{LPT}[\pi_k[1]], \dots, P_{LPT}[\pi_k[n]]$ 
11     $\sigma[k], C_{\max}[k] \leftarrow \text{parallel}(BBL1(P_k, m, \sigma[k], N, S))$ 
12     $k \leftarrow k + 1$ 
13   $P_{SLACK} \leftarrow \text{order}(P, 'SLACK')$ 
14   $\sigma[\mathcal{O}], C_{\max}[\mathcal{O}] \leftarrow \text{parallel}(BBL1(P_{SLACK}, m, \sigma[\mathcal{O}], N, S))$ 
15   $\text{waitall}()$ 
16   $k_{\text{best}} \leftarrow \arg \min_k (C_{\max})$ 
17  return  $\sigma[k_{\text{best}}]$ 
```

Algorithm 6: MBBS2(P, m, N, S, K)

Input:
 $P \leftarrow \{p_j\}_{j \in \{1, \dots, n\}}$
 $m : \text{integer}$
Optional : $N \leftarrow \infty, S \leftarrow 0, K \leftarrow 1 : \text{integer}$

Data:
 $\sigma_K[K], \sigma', \sigma \leftarrow \{()\}_{i \in \{1, \dots, m\}} : \text{schedule}$
 $S_k, N_k, S', N', n' \leftarrow K \times m : \text{integer}$
 $n_1, \dots, n_K, q, r, k \leftarrow 0, j \leftarrow 1 : \text{integer}$

```

1 begin
2    $q \leftarrow \lfloor n/K \rfloor$ 
3    $r \leftarrow n - Kq$ 
4    $n_1, n_2, \dots, n_K \leftarrow q$ 
5   for  $j \leftarrow 1$  to  $r$  do  $n_j \leftarrow n_j + 1$ 
6    $j \leftarrow 1$ 
7   while  $j \leq n$  do
8      $k \leftarrow k + 1$ 
9      $P_k \leftarrow \{P[j], P[j+1], \dots, P[j+n_k-1]\}$ 
10     $j \leftarrow j + n_k$ 
11     $N_k \leftarrow \max(\min(n_k - m - 1, N), 0)$ 
12     $S_k \leftarrow \max(\min(S, n_k - N_k - m - 1), 0)$ 
13     $\sigma_K[k] \leftarrow \text{parallel}(BBS(P_k, m, \sigma_K[k], N_k, S_k))$ 
14  waitall()
15  for  $k \leftarrow 1, \dots, K$  do
16    for  $i \leftarrow 1, \dots, m$  do  $p'_{k,i} \leftarrow p_{\sigma_K[k][i][1]} + p_{\sigma_K[k][i][2]} + \dots + p_{\sigma_K[k][i][\text{size}(\sigma_K[k][i])]}$ 
17   $N' \leftarrow \max(\min(N, n' - m - 1), 0)$ 
18   $S' \leftarrow \max(0, \min(S, n' - N' - m - 1))$ 
19   $\sigma' \leftarrow BBS((p'_{1,1}, \dots, p'_{1,m}, p'_{2,1}, \dots, p'_{2,m},$ 
20     $\dots, p'_{K,1}, \dots, p'_{K,m}), m, \sigma', N', S')$ 
21  for  $k \leftarrow 1, \dots, K$  do
22    for  $i \leftarrow 1, \dots, m$  do
23       $\sigma \leftarrow \text{replace}(\sigma', (k, i), \sigma_K[k], (i, 1),$ 
24         $\dots, (i, \text{size}(\sigma_K[k])))$ 
25  return  $\sigma$ 
```

summed to obtain $p'_{k,i}$. The schedule σ thus built is the $P||C_{\max}$ solution of I . The functions *parallel*, *waitall* are the same as in Section 4.1.

Using MBBS2, makes sense particularly on instances where n is very large. Indeed, it explores combinatorial possibilities on different subsets of the jobs instead of just one subset (defined by N and S in BBS). The branching strategy is working at different levels of the job list and it is done in parallel, so the computation time does not have to be more important to get more different solutions and potentially a better schedule. Finally, when the sub-schedules are added together to form one full schedule, we know it has been more refined than what BBS alone is capable of. The same idea lays behind MBBS3, described in the following, but for this heuristic the sub-problems are defined on subsets of machines instead of subsets of jobs.

4.3 MBBS3: one BBS per subgroup of machines

Yet another way to run multiple BBS is to pose sub-problems that reduce the number of machines used. If m is the number of machines of the original instance I , we defined a number of K sub-instances I_k in which the numbers of machines are $m_1, m_2, \dots, m_K$. That is only interesting if m is large enough (in our tests we considered it is worth it if $m > 5$). Indeed, the larger m is, the less BBS is exploring the field of possibilities, because the algorithm is not allowed to try assignments on the third up to the m -th least loaded machines. On the other hand, increasing the tree arity was not found to be an efficient way to proceed (as explained and empirically elaborated

in Section 5.2.1).

The strategy of splitting the m machines by groups is an attempt to enlarge the spanning of the exploration. Algorithm 7 follows this strategy. It works in three times:

Algorithm 7: MBBLS3(P, m, N, S, K)

Input:

$P \leftarrow \{p_j\}_{j \in \{1, \dots, n\}}$

m : integer

Optional parameter : $N \leftarrow \infty$: integer

Optional parameter : $S \leftarrow 0$: integer

Optional parameter : $K \leftarrow 1$: integer

Data:

$\sigma_K[K] \leftarrow \{()_{i \in \{1, \dots, m_k\}}\}$: schedule

$\sigma' \leftarrow \{()_{i \in \{1, \dots, K\}}\}$: schedule

$\sigma \leftarrow \{()_{i \in \{1, \dots, m\}}\}$: schedule

$k \leftarrow 0$: integer

S_k, N_k : integer

q, r : integer

$n_1, \dots, n_K$: integer

$m_1, \dots, m_K$: integer

```

1 begin
2    $q \leftarrow \lfloor m/K \rfloor$ 
3    $r \leftarrow m - Kq$ 
4   if  $r > 0$  then
5     error('m must be evenly divided by K')
6    $\sigma' \leftarrow \text{BBLs}(P, K, \sigma', N, S)$ 
7   for  $k \leftarrow 1, \dots, K$  do
8      $m_k \leftarrow q$ 
9      $n_k \leftarrow \text{size}(\sigma'[k])$ 
10     $P_k \leftarrow (p_{\sigma'[k][1]}, \dots, p_{\sigma'[k][n_k]})$ 
11     $N_k \leftarrow \max(\min(n_k - m_k - 1, N), 0)$ 
12     $S_k \leftarrow \max(n_k - N_k - m_k - 1, 0)$ 
13     $\sigma_K[k] \leftarrow \text{parallel}(\text{BBLs}(P_k, m_k, \sigma_K[k], N_k, S_k))$ 
14  waitall()
15   $\sigma \leftarrow (\sigma_K[1][1], \dots, \sigma_K[1][m_1], \dots, \sigma_K[K][1],$ 
16     $\dots, \sigma_K[K][m_k])$ 
17  return  $\sigma$ 
```

1. Split the n jobs into K subgroups as balanced as possible. This partition problem is equivalent to a $P||C_{\max}$ problem on K machines so we can use BBLs in that purpose too (line 6). Of course K must stay small to be consistent with the idea to really reduce the number of machines. In our experiments we limit m to at most 25 machines. Another constraint is that all the m_k must be equal to each other to avoid unbalanced partition of jobs.
2. Solve corresponding K sub-instances of $P||C_{\max}$ whose numbers of machines are $m_1, \dots, m_K$ (line 13).
3. Concatenate together the K (line 15) sub-schedules obtained to solve $P||C_{\max}$ on the initial instance I .

An additional refinement, of the local search kind, precisely a neighborhood search, is possible before the last step, by proceeding to job interchanges between sub-schedules in order to reduce the possible unbalance between two sub-schedules. A simple rule would be: let I_{k_1}, I_{k_2} the sub-instances for which the makespans obtained are $C_{\max, k_1}, C_{\max, k_2}$, such that they are respectively the greatest and the smallest of all sub-schedules. Then if $\delta = C_{\max, k_1} - C_{\max, k_2}$ is such that it exists a pair of jobs $(p_1, p_2) \in I_{k_1} \times I_{k_2}$, with p_1 being assigned to the machine of greatest completion time $C_{\max, k_1}$, and $0 < p_1 - p_2 < \delta$ then swapping these two jobs between sub-schedules of I_{k_1} and I_{k_2} is lowering the overall $C_{\max}$ for I . We do not have to assess all pairs (p_1, p_2) , only taking the two first largest can be enough. This interchanging process can hence be repeated in a time complexity of at most $O(n)$. It is not included in Algorithm 7 but could easily be inserted on the end.

Table 3: Experiment instances

Instances ⁶	Distribution	m	n	[a,b]
Della Croce and Scatamacchia Instances				
DCU	Uniform	5, 10, 25	10, 50, 100, 500, 1000	[1, 100], [1, 1000] [1, 10000]
DCNU	Non-uniform	5, 10, 25	10, 50, 100, 500, 1000	[1, 100], [1, 1000] [1, 10000]
Gupta et Ruiz-Torres Instances				
GR1	Uniform	3, 4, 5	2m, 3m, 5m	[1, 20], [20, 50]
GR2 ⁷	Uniform	2, 3, 4, 5, 6, 8, 10	10, 30, 50, 100	[100, 800]
GR3	Uniform	3, 5, 8, 10	3m+1, 3m+2, 4m+1, 4m+2, 5m+1, 5m+2	[1, 100], [100, 200]

5 Experimental evaluation of BBLS and BPLS

In this section we mainly provide a global evaluation of BBLS and BPLS. They are compared to other algorithms mentioned in the introduction. More precisely, we present benchmarks for BB-SLACK and BP-SLACK, variants of SLACK [Della Croce & Scatamacchia, 2020], already introduced in 2.2. We did not integrate BB-LPT and BP-LPT in our tests because LPT is less effective than SLACK on most of the instances considered as shown in [Della Croce & Scatamacchia, 2020]. For reproducibility, all the source code on which the experiments are based is available online [Hadj-Djilani, 2025].

Before the benchmarks we describe the sets of instances used. They are all taken as references from the literature.

5.1 The instances

We selected different sets of instances from the literature, they are listed in Table 3. The first group of instances comes from [Della Croce & Scatamacchia, 2020] but originally from [França et al., 1994] for uniform instances (DCU) and from [Frangioni et al., 2004] for non-uniform instances (DCNU). The second group of instances is composed of three subgroups denoted GR1, GR2 and GR3 in reference to experiences E1, E2 and E3 made in [Gupta & Ruiz-Torres, 2001]. We chose those instances: (i) firstly because, as variants of SLACK, BB-SLACK and BP-SLACK should be evaluated according to the same protocol used in [Della Croce & Scatamacchia, 2020], (ii) secondly because, in our experiments, LISTFIT delivers good performance but comparing SLACK and LISTFIT using one or another set of instances did not give the same results, hence the need to use several set of instances like the one used to evaluate LISTFIT in [Gupta & Ruiz-Torres, 2001]. The instances are generated as follows: for each combination of m, n and interval $[a, b]$ in the lines of the Table 3, such that $m < n$, 20 instances are drawn with processing times picked randomly in interval $[a, b]$ according to the instance group distribution (uniform or non-uniform). The non-uniform distribution is such that 98% of the processing times are integers uniformly distributed in $[0.9(b - a), b]$ while the remaining ones are uniformly distributed in $[a, 0.2(b - a)]$. All these instances sum up to 780 DCU instances and as many DCNU instances, 360 GR1 instances, 540 GR2 instances and 960 GR3 instances for a total of 3420 instances.

5.2 Benchmarking BBLS and BPLS against well-known heuristics

5.2.1 Protocol and algorithms

We present now an experiment established in order to evaluate the performance of BBLS and BPLS. This evaluation takes into account not only the quality of the solutions but also the computational time. The tests are performed on BB-SLACK (BBLS) and BP-SLACK (BPLS). We compare those algorithms to other well-known algorithms:

- LS-based algorithms with LPT and SLACK,
- Bin-packing based algorithms like MULTIFIT, LISTFIT for which the number of iterations is always $k = 7$. This is the configuration advised for MULTIFIT in [Coffman et al., 1978] and also applied in [Gupta & Ruiz-Torres, 2001].

⁶Exact instances used in 5.2 are available online [Hadj-Djilani, 2023a]

⁷Although the value $m = 5$ is not used in the original E2 experiment ([Gupta & Ruiz-Torres, 2001]), it is used in GR2 instance set.

- COMBINE which is a mixed strategy between the two previous kinds (MULTIFIT being called with the same number of iterations, $k = 7$),
- We also chose to evaluate LDM because it can be very efficient. Besides, it implements yet another strategy [Michiels et al., 2003].

BP-SLACK and BB-SLACK are both configured with $N = 10$ and $N = 15$ in two separate runs. Varying the N value allows us to get an insight on how the compromise between solution quality and computation time evolves with respect to N . However the shifting parameter S is always zero for BP-SLACK while a value of S near to the maximum ($n - N - m - 1$) is used for BB-SLACK. Recall that S matters only in a pruning scenario, that is why BP-SLACK (BPLS) does not use it. We executed these algorithms on all the instances described in Table 3. The exact instances used are available online [Hadj-Djilani, 2023a].

The algorithm implementations were all written in Python using Numpy (<https://numpy.org>). For the BBLS implementation, to optimize the computation of the L_3 bound, we used Numba (<https://numba.pydata.org>). Optimization properties J.1, J.2, J.3 were enabled for BBLS. Table 2 sums up all the properties considered and their activation or deactivation in experiment implementation.

We carry out the experiment on the same CPU⁸ for all algorithms. Except for BP-SLACK, LISTFIT and MBLS (Section 6), which are parallelized, all other heuristics are run sequentially, including BB-SLACK. BP-SLACK uses at most 2^7 parallel processes on this same CPU as described Section 2.1.

The LISTFIT implementation has also been parallelized. Indeed, the $4n$ MULTIFIT calls needed by LISTFIT can easily be spread on several processors/cores as threads or processes. A quick experiment allowed to verify that the parallelization of LISTFIT is worth it on the CPU used⁸ if the number of jobs is greater or equal to 100 (limit under what we used a sequential LISTFIT). These details matter in the computation time measurement made further in this section.

Note that three sets of experiments, done to widen the scope of our investigations, are not presented here. First, we ran the algorithms on perfect matching instances (PMI's, instances where the optimal solution is known) but it did not give any major difference in the evaluation of our algorithms. Please refer to Appendix A for more details on these results and how we built PMI's. Second, we did not include our experiments about BBLS configured with ternary trees because it turns out to be less efficient for the same computational cost as binary tree BBLS. We tested a ternary tree based BBLS and limited the number of nodes to get the same cost than that obtained with binary BBLS. Indeed, it is possible to reduce the value of N to achieve a comparable cost whatever is the basis of the exponential in Equation (1). Precisely, if $N = N_2$ in basis 2 and $N = N_3$ in basis 3 and we choose an arbitrary N_2 , we can adjust the value of N_3 to obtain in the induced ternary tree an equal or lower number of nodes than in the binary tree induced by N_2 . The formula of this constraint is simply: $N_3 = \lfloor \log_3(2^{N_2+2} - 1) - 1 \rfloor$. This implies that less jobs are assigned by the branching strategy for ternary BBLS than for binary BBLS. In practice, we chose $N_2 \in \{10, 15\}$ for binary BBLS and $N_3 \in \{6, 9\}$ for ternary BBLS. We noticed that for more than 90% of instances listed in Table 3, binary BBLS produced equal or better solutions than ternary BBLS while the latter conversely produced better or equal solutions to binary BBLS for only 70% of the instances. Furthermore, ternary BBLS was twice more computationally expensive than binary BBLS even if the full ternary tree of the former was smaller, in number of nodes, than that of binary BBLS. It suggests that the tree pruning is far less efficient with arity 3.

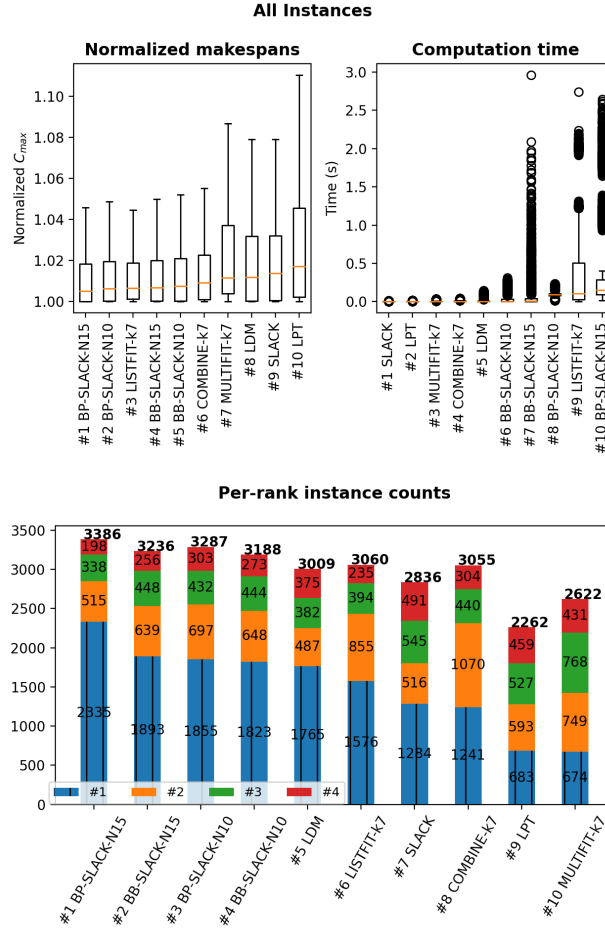
Finally, we also compared our heuristics against exact solvers considering again the instances in Table 3. Indeed, it is an interesting way to evaluate both the computation time and solution quality. We did not compare to B&B exact method from [Dell'Amico & Martello, 1995] because of its very higher worst case computational time complexity (in this method, contrary to our heuristics, all machines are considered for each job assignment). We referred instead to a mixed integer linear programming (MILP) formulation and a cutting plane solving method similar to the one developed in [Mokotoff, 2004]. We limited the CPU⁸ computation time up to an upper bound of 300s in order to make the experiment feasible in a reasonable time and because we are more particularly interested in viewing on which instances an exact solver can make our heuristics computationally less interesting. Hence using the Python-MIP package (<https://www.python-mip.com/>) as a frontend of Coin-OR state-of-the-art solver⁹, we observed, as shown in Appendix E that:

- The MILP method is about 100 times slower,
- The MILP method allows to find only 64% of guaranteed optimal solutions within the time constraint of 300s, while our heuristics managed to find 52% of optimal solutions (a number that cannot anyway be larger than MILP's because it is the method used here to know a solution is optimal). For the comparison, only 28% of LPT and 48% of LISTFIT solutions are optimal.
- In the same time, the results show that for about 75 to 79% of the instances our BBLS/BPLS solutions are better or equal to MILP solutions, including those that MILP method cannot ensure they are optimal, typically instances for which the MILP solver reached the time limit.

⁸An Intel(R) Xeon(R) Platinum 9242 equipped of 48 cores, that is 96 threads.

⁹Python-MIP is also compatible to Gurobi solver according to its documentation

Figure 4: Performance of the branching heuristics relatively to well-known algorithms for all instances Table 3 (DCU, DCNU, GR1, GR2, GR3)



5.2.2 The results

In this subsection we show, with the experimental results obtained, how BPLS and BBLS are competitive relatively to well-known heuristics. They managed to rank first for the vast majority of instances, whether considering the median makespans or counting the number of instances for which they obtain the best makespan. We show also that these results are performed in a reasonable median computation time comparable to what reference algorithms provide. We finally assess the question of instance subsets for which LISTFIT and LDM produce better makespans and envision the idea behind the use of MBBLs in next section to do as good as those heuristics with our list scheduling branching strategy.

Figure 4 presents the makespans obtained with each algorithm as well as their execution times. In order to compare the results for different instances in the same figure the makespans are normalized using the lower bound L_1 (see Section 3.1). Hence the ranks according to the median normalized makespans are indicated on the figure along the algorithm names which are sorted according to their rank from left to right. Two algorithms might have the same rank in case of tie. The makespan boxplot is useful as it gives an insight of the algorithm result distribution over the full set of instances. However, medians should be read carefully because it is totally possible, for two algorithms A_1 and A_2 , to obtain makespan medians m_1 and m_2 such that $m_1 \leq m_2$ but with a number of better results far more important for A_2 . For that reason, the stacked bar chart at the bottom of Figure 4, displays the number of instances for which an algorithm ranks first, with the best makespan, or second and so on (with possible ties here too). The ranks span from #1 to #4 since the last ranks (#5 to #10) were not shown in the figure for readability reasons. Besides, we are more interested in first ranks even if looking more carefully almost all instances are accounted in the figures. Note that on the top of each bar is printed the total number of instances for all accounted ranks. Likewise, over each sub-bar is displayed the number of instances counted for the corresponding algorithm rank listed in the legend.

All instances grouped: Figure 4 shows the results obtained with instances from all different groups (DC and GR) combined. A first general remark is that all heuristics perform very well with median normalized makespans very close to one, that is median makespans close to the lower bound L_1 . For any tested algorithm the difference between the median makespan and L_1 is smaller than 2% of L_1 . The algorithms are very competitive to each other. Indeed they all have been already proven to be effective in the literature. However it is clear that the spanning of the makespans over the instances can vary significantly from one algorithm to another (e.g. about 5% of L_1 for LISTFIT to more than 10% for LPT¹⁰). Besides, depending on the jobs processing times, even 1% of the makespan can be interesting for the performance of a schedule.

Furthermore, even if the makespan changes just a bit from an algorithm to another for a given instance, we are still interested in enhancing the performance toward the optimal solution. Anyway, out of makespan value considerations, counting the number of #1 ranks, we see significant variations among algorithms. We can see in the "per-rank instance counts" bar chart of Figure 4 that BP-SLACK and BB-SLACK give the best makespan for a number of instances significantly greater than all other algorithms. A larger value of N giving obviously the best results. A value of $S = 0$ for BP-SLACK runs shows also a significant advantage over BB-SLACK runs that almost maximize S . This advantage however comes with a larger computational cost. This can be seen in the "Computation time" boxplots, where BP-SLACK ranks 9 and 10, while BB-SLACK ranks 6 and 7, even if there are outliers with larger times for BB-SLACK configured with $N = 15$. These outliers are most likely due to a poor pruning of the corresponding solution trees by BBLS but we note that, for the most part, they do not go as high as BP-SLACK set also to $N = 15$.

Looking at the "Normalized makespans" boxplots, we see again that BP-SLACK and BB-SLACK algorithms are top ranked, from 1 to 5. LISTFIT ranks 3, just after the BP-SLACK algorithm and before the BB-SLACK algorithm. From the bar chart we can nevertheless verify that LISTFIT has less instances than BB-SLACK ($N = 15$) for which it obtains the best makespan. On the other hand, considering the makespans of all instances, the median for LISTFIT is better than that of BB-SLACK ($N = 10, 15$).

In order to evaluate furthermore the statistical significance of the results obtained in Figure 4 we ran a Wilcoxon statistical test (implemented using Scipy's provided utility, <https://scipy.org>). The test purpose is to reject the null hypothesis according to which two algorithms would have the same distribution of normalized makespans. The results shown in Appendix D are pretty clear (Tables 5, 6): all the p-values (the confidence indices) are quasi-zero which means that the null hypothesis is strongly rejected. The second greatest p-value is $1.9e - 05$ and happens between BP-SLACK-N10 and BB-SLACK-N15 distributions. It makes sense because those two algorithms are similar. The overall greatest p-value, $7.3e - 02$, is between LISTFIT and BB-SLACK-N10, which are not algorithms of the same family. However, because LISTFIT is one of the best competitor of BPLS/BBLS algorithms, it is not surprising that its makespan distribution can be closer (indeed, in Figure 4, LISTFIT rank is 3, just between BPLS and BBLS algorithms). Still, the two distributions at hand are quite different.

In the following, we assess the heuristics over the different subgroups of instances. In particular we show for which instance subgroups LISTFIT manages to be better than the BB/BP-SLACK algorithms.

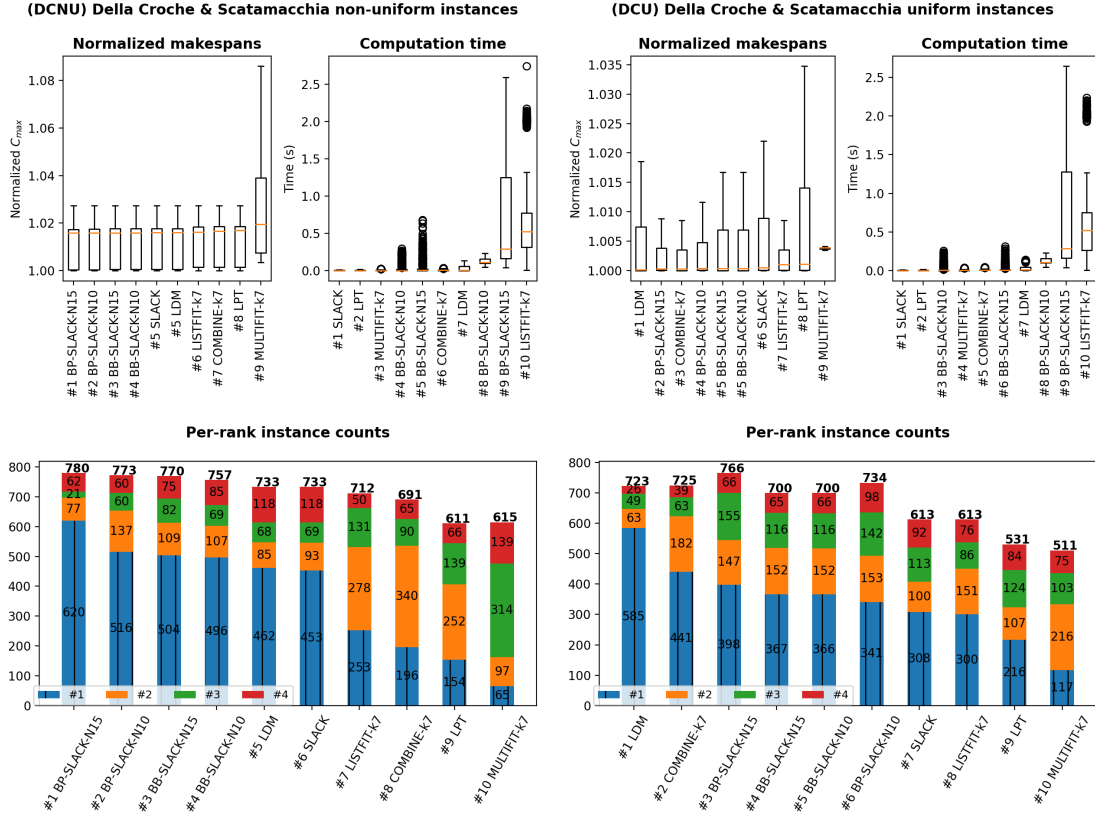
DC instances: The results obtained with DC instances are presented in Figure 5. Note that we also ran the experiment on equivalent instances from University of Bologna and obtained similar results (see Appendix B, Figure 17).

Starting with the DCNU instances, that is the non-uniform instances, it is pretty obvious that BB-SLACK, and even more BP-SLACK, outperform all other algorithms. Whether we look at the bar chart or at the boxplots, they accumulate a clearly larger number of instances on which they ranked at the top, and reach smaller median makespans than all other algorithms.

Regarding the computation time, LISTFIT occupies the last position (rank 10) just behind BP-SLACK (ranks 8 and 9) and surprisingly BB-SLACK ($N=10$) managed to do faster than COMBINE and LDM. It is likely that the tree pruning was very efficient on this configuration and set of instances. According to this experiment, if we were to restrain ourselves to only BB-SLACK $N = 10$ against all other algorithms (excluding those of our contributions), we could, for the DCNU instances, have the best solutions to $P||C_{\max}$ and be ranked 4 out of 7 regarding the computation time. In this figure we also notice, as in [Della Croce & Scatamacchia, 2020], that SLACK is very efficient for that kind of instances compared

¹⁰The figures for the normalized makespans do not show the outliers for the sake of readability but we noticed outliers were not making a difference between one algorithm to each other. The greatest outlier makespan was about 30% above L_1 for all algorithms.

Figure 5: Benchmark on DCNU and DCU instances (see Table 3)



to COMBINE, MULTIFIT or LISTFIT (bin-packing algorithms in general).

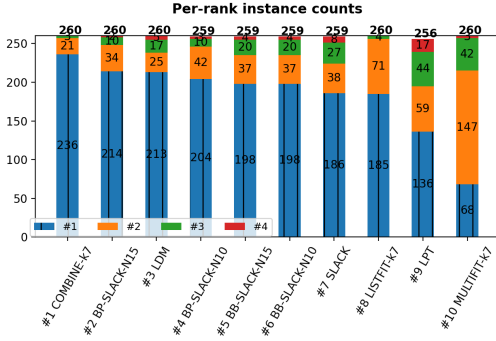
For uniform instances (DCU), we observe that BB-SLACK and BP-SLACK do quite good with a median normalized makespan, at rank 2 for BP-SLACK $N = 15$ and rank 3 for its number of best makespan instances. This configuration of BP-SLACK ranks just after LDM (rank 1) and just before COMBINE (rank 3) for the median makespan, while LDM is first, COMBINE second and BP-SLACK $N = 15$ third for the number of best makespan instances. Note however that the latter manages to accumulate the most important number of instances for which it is ranked in the top 4. BP-SLACK $N = 10$ produces less good results than LDM and COMBINE too and is ranked 4 regarding the median makespan and 6 for the number of best makespan instances, behind BB-SLACK $N = 15$.

This one ranks better than BB-SLACK $N = 10$ but with almost the same performance because the former is the best for 367 instances, while the latter is for 366 instances. It is however worth noting that, on this set of DCU instances, BB-SLACK $N = 10$ median computation time is smaller than that of MULTIFIT, COMBINE, LDM and LISTFIT. According to the median time BB-SLACK $N = 15$ is faster than LDM and LISTFIT. However, BP-SLACK configurations, $N = 15$ and $N = 10$ (with $S = 0$), are slower than all other algorithms except LISTFIT which is the slowest of all.

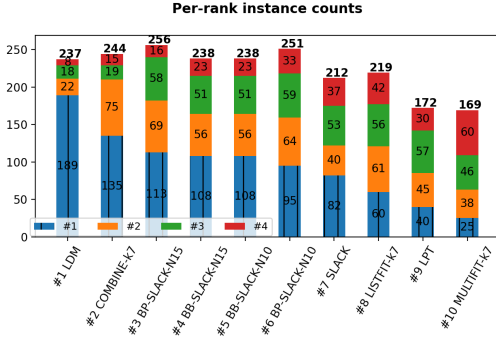
In Figure 6, we see that LDM becomes more and more dominant on DCU instances, when b grows from 100 to 10000. It is also noticeably true when n grows as shown in Table 4. An explanation for these two facts is to find in the LDM strategy. First, LDM subdivides the jobs into smaller packets of jobs, the partial schedules. When new assignments of jobs are made by merging two partial schedules, the whole set of n jobs is taken into account (by comparing the differences of machine loads among partial schedules). On the other side, BBS and BPLS only consider the loads of the two least loaded machines when assigning the next job with no consideration at all to other jobs yet to assign. Going blindly on the remaining jobs to assign, is a risk that is more and more important as n grows. Similarly, when b grows, the difference between two p_j 's has a greater chance to be larger so that a wrong choice of assignment might be more penalizing on the machine load balancing. Besides, for a same n , the availability of smaller jobs becomes less probably when b is larger (remember that here we consider uniform instances).

Figure 6: Benchmark on DCU instances for different $[a, b]$ (see Table 3)

(DCU) Della Croce & Scatamacchia uniform instances $ab=1-100$



(DCU) Della Croce & Scatamacchia uniform instances $ab=1-1000$



(DCU) Della Croce & Scatamacchia uniform instances $ab=1-10000$

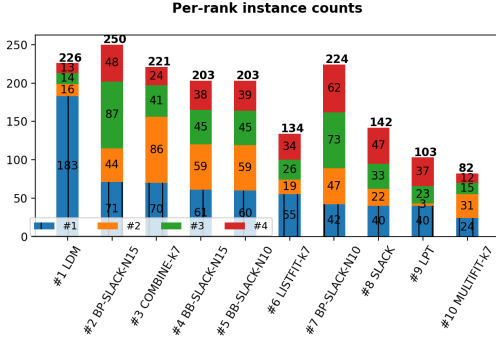


Table 4: LDM versus BP/BB-SLACK wins w.r.t. n

n	number of instances	LDM	BP-SLACK $N = 15$	BP-SLACK $N = 10$	BB-SLACK $N = 15$	BB-SLACK $N = 10$
10	60	54	55	55	55	55
50	180	76	86	131	84	84
100	180	100	35	47	43	44
500	180	175	70	70	78	78
1000	180	180	95	95	106	106

But smaller jobs are precisely what is needed to balance differences of machine loads implied by larger jobs. In Section 6 we show how MBLS2, introduced in Section 4.2, can mitigate those penalties on BLS, by working on smaller packets of jobs first, as LDM does, and then grouping the sub-schedules obtained into a complete one. For GR instances the rankings are also subtly different. First we consider GR1 and GR2 on

Figure 7. It is clear that both BP-SLACK and BB-SLACK outperform all other algorithms for all our metrics, except for the computation time. However BB-SLACK ($N = 10, 15$) is faster than LISTFIT which produces significantly less good solutions, especially for GR2 instances. For GR3 instances, in Figure 8, the interesting point is that LISTFIT outperforms most of our algorithms on median makespan and all of them on number of best makespan instances. However its median execution time is larger than BB-SLACK's. We zoomed on specific subgroups of GR3 instances to understand where exactly LISTFIT achieves to be the best algorithm.

GR instances: In Figures 9 to 12 we clearly see that LISTFIT gets its advantage when the interval of processing times is $[a = 100, b = 200]$ (see bottom subplots), while the interval $[a = 1, b = 100]$ (see top subplots) fits better to BP-SLACK and BB-SLACK best performance. When $m \geq 8$, there is a slight difference because BP/BB algorithms do less well when m becomes larger. Indeed, with $m = 10$, LISTFIT ranks first, even if $a = 1$ (with 84 instances for which LISTFIT is the best and 80 for BP-SLACK $N = 15$).

The fact that our algorithms are less efficient when $a = 100$ is to find in the other LS-based algorithms that seem to all follow this tendency, contrary to FFD based algorithms that behave well in this configuration, LISTFIT being at the top. Preventing any smaller p_j puts LS-based algorithms in difficult situation to balance machines loads because, as mentioned before in LDM performance comparison, the smaller jobs are necessary in this purpose. About the effect of the m value, it is clear that, when m is greater, there is much of a chance that the optimal choice of assignment for each job is another machine than the first or second available ones. Therefore the likelihood to get less efficient solutions for BP-SLACK and BB-SLACK is larger when m grows. MBBLS1 and MBBLS2 introduced in Sections 4.1 and 4.2 are used in Section 6 to overcome these LS-based algorithm disadvantages.

Figure 7: Benchmark on GR1 and GR2 instances (see Table 3).

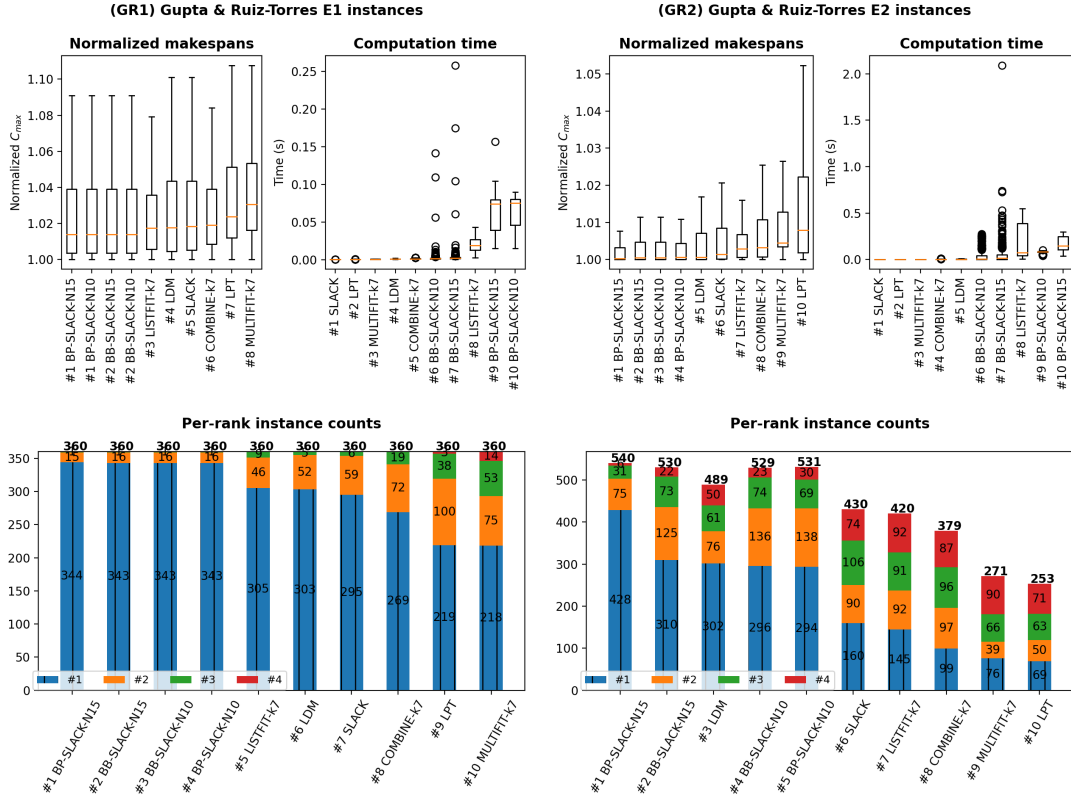


Figure 8: Benchmark on GR3 instances (see Table 3).

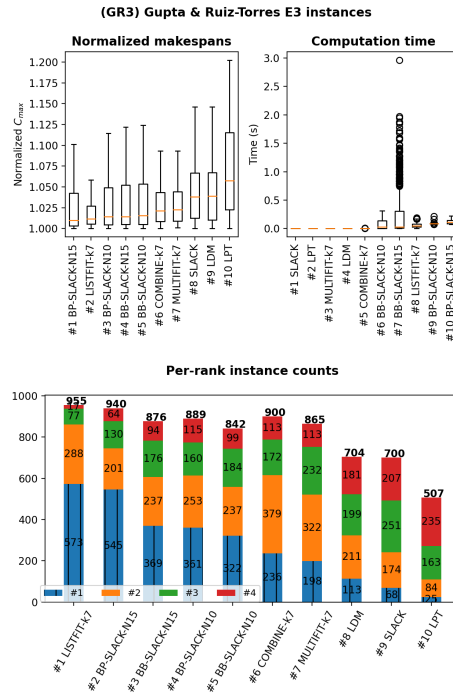


Figure 9: Focus on GR3 instances, $m = 3$ (see Table 3)

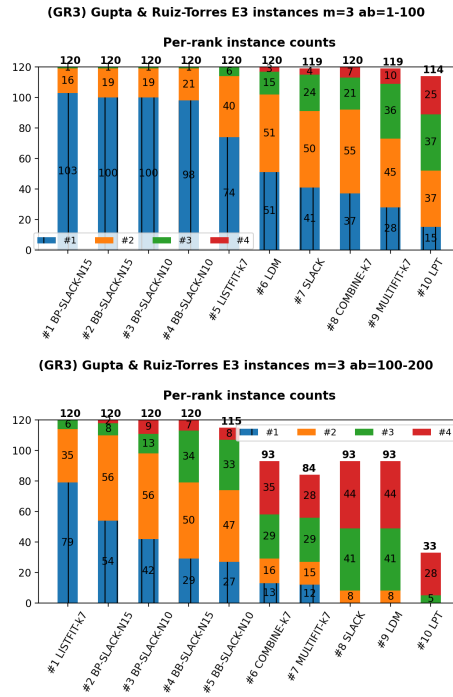


Figure 10: Focus on GR3 instances, $m = 5$ (see Table 3)

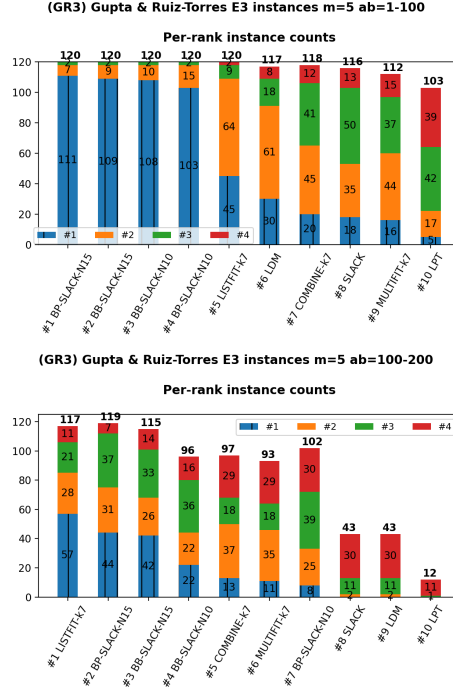


Figure 11: Focus on GR3 instances, $m = 8$ (see Table 3)

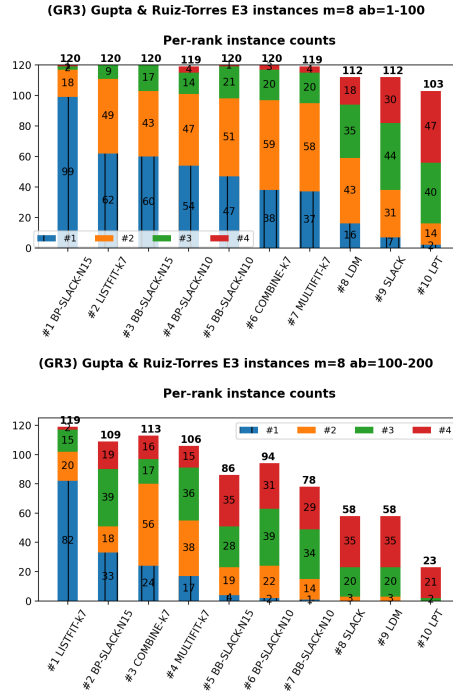
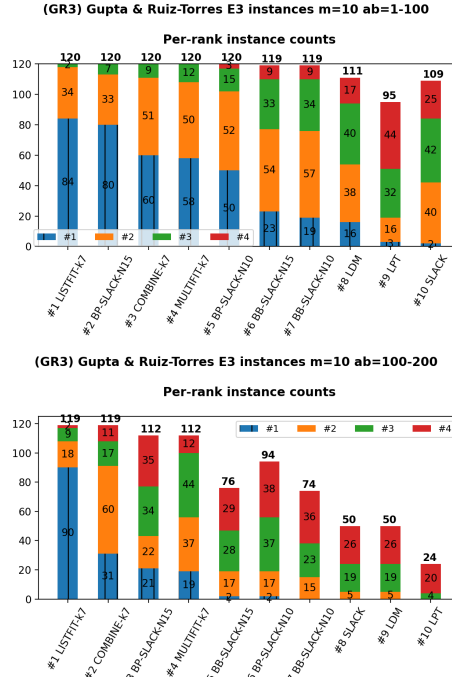


Figure 12: Focus on GR3 instances, $m = 10$ (see Table 3)



6 An ad hoc MBBLS

In Section 5.2 we show that LISTFIT or LDM can be more efficient on some subsets of instances. To take this into account, we can of course, as a direct alternative, combine the LS branching with other heuristics, noticeably LDM or LISTFIT, but the goal pursued is to push the LS branching strategy to its best capabilities. Therefore, we propose an ad hoc MBBLS heuristic using the three heuristics introduced in Sections 4.1, 4.2, 4.3. The flowchart, presented in Figure 13, shows how this heuristic works. Basically, it calls MBBLS1, MBBLS2, MBBLS3 or just BB-SLACK in order to maximize the quality of solutions for all the tested instances. Indeed, BB-SLACK performs very well on non-uniform instances (DCNU), so it was maintained for those instances, but has clear weaknesses on other instances.

On GR3 instances BB-SLACK does not perform as well as LISTFIT when m grows so, for $m > 5$, MBBLS calls MBBLS3 which divides the $P||C_{\max}$ problem into sub-problems that use a smaller number of machines. On the same subset of instances, we noticed that LISTFIT is also better if the lower bound a , chosen to draw the p_j 's, is greater or equal to 100, hence, in that case, we call MBBLS1. Two exceptions should be noticed in the flowchart. First, when $m = 5$, even if $a \geq 100$, rather to call MBBLS1, MBBLS calls BB-SLACK. The reason for this choice is simply that BB-SLACK does better on this subset. For the same reason, a second exception is made about MBBLS3. It is not called when $\frac{n}{m} = 2$, even if $m > 5$. Among tested instances the only precise case that verifies this condition is when $m = 25$ and $n = 50$. Indeed, on these instances, we observed that BB-SLACK is more efficient than MBBLS3. We also use MBBLS2 on cases where $n \geq 100$ since that is the point of this heuristic to run on instances that contain a large number of jobs. That choice was made to manage outperforming LDM on DCU instances.

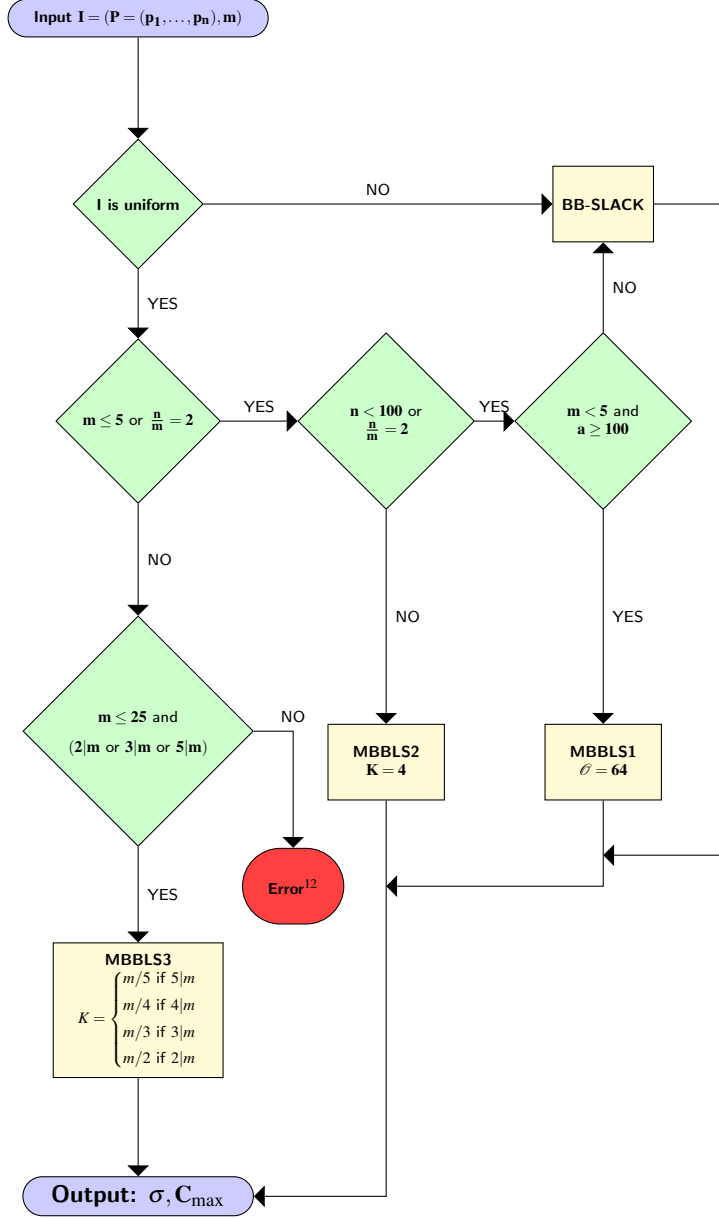
By running MBBLS as defined by the flowchart, we were able to outperform LDM and LISTFIT on instances for which BB-SLACK alone was not able to do so (as shown in Section 5.2.2). Figures 14 and 15¹¹ show the results respectively for the GR3 and DCU instances. In Figure 14, apart from our heuristics, only LISTFIT is considered, because that is the challenging algorithm identified in Section 5.2.2 for GR3 instances (see Figure 8). LDM is likewise included in Figure 15 regarding the results obtained for DCU instances (see Figure 5). Another point that requires attention, is the increase of N up to 16 for MBBLS. The reason for this configuration is the greater number of wins for LDM compared to MBBLS $N = 15$. Note however that, even if MBBLS $N = 15$ has less wins than LDM, the former more often ranks in the two first algorithms than the latter.

Anyway, MBBLS $N = 16$ is enough to produce more wins than LDM. Finally, note that although LDM is

¹¹We also ran the experiment on equivalent instances from University of Bologna website and obtained similar results (see Appendix B, Figure 18)

ranked first according to its median normalized makespan (Figure 15) the spanning of its makespans, is quite larger than that obtained with any configuration of MBBS (even $N = 10$). MBBS $N = 17$, while not present in the figure, was also ran and gave a median normalized $C_{\max}$ lower than that of LDM. This allows us to conclude that the MBBS approach, with increasing values of N can outperform LDM.

Figure 13: Ad hoc MBBS flowchart

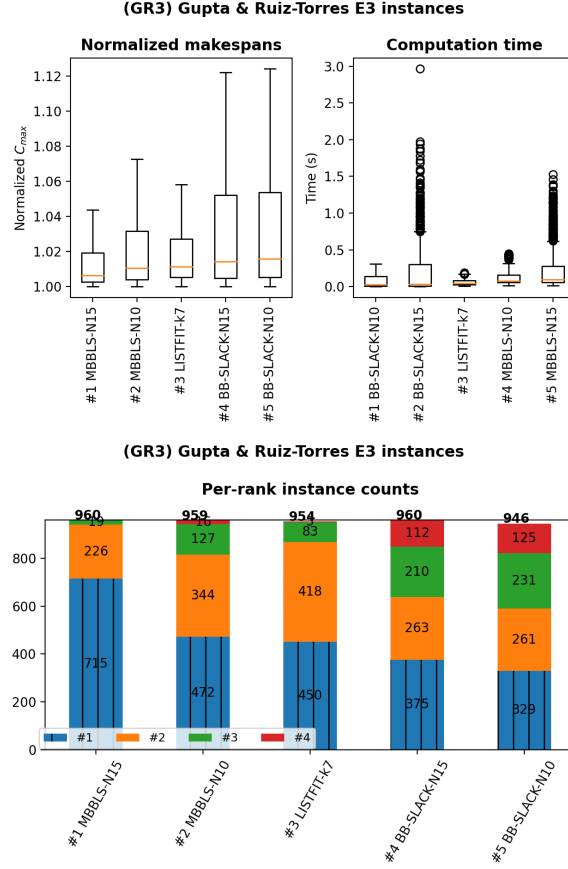


This complementary experiment has been performed on the same CPU as in Section 5.2. BB-SLACK is still run sequentially but MBBS1, MBBS2 and MBBS3 are parallelized. MBBS1 runs $\ell = 64$ parallel processes, one by tested job order while the two other heuristics create K parallel processes, one by subproblem, and run one extra BBS call sequentially. Figure 13 exhibits the different values of K used.

Conclusion

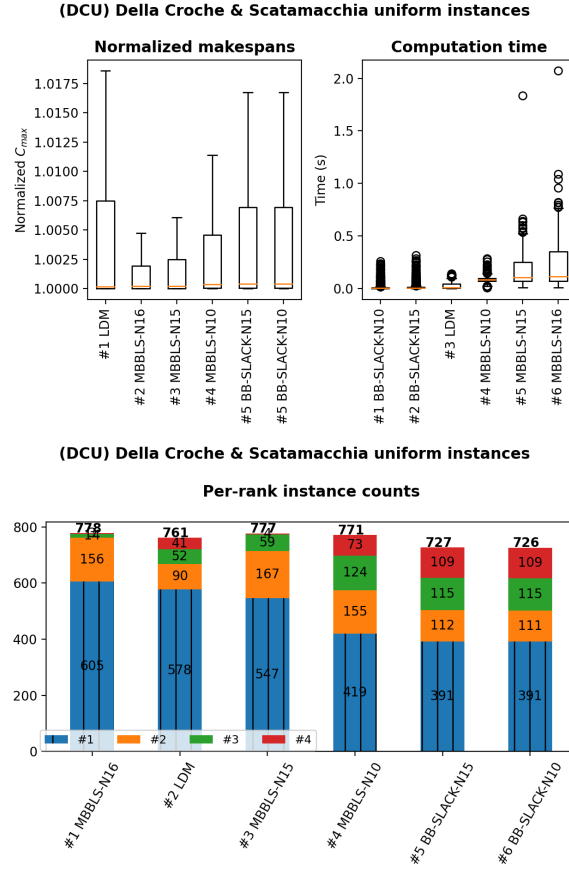
In this paper we show how a modification in the LS algorithm, considering an assignment on other machines than the first available one (actually only the second one), can improve the makespan obtained for a $P||C_{\max}$ problem. We show that branching strategies can easily be derived, as in BLS algorithm, although the cost is

Figure 14: MBBLS versus LISTFIT on GR3 instances (see table 3)



exponential in the number of jobs n . A parameter N , that limits the number of jobs treated by BLS, the others being handled by the classic LS, is our first step to speed up the method. This parameter N is used without losing the algorithm interest regarding the quality of solutions. Then we show, with BPLS, that a basic parallelization of the alternatives encoded in the solution search tree allows to speed up significantly the branching process. Besides, it makes possible to scale up to the available computational power of a multi-core CPU or a computation grid, adjusting the quality of solutions through the parameter N . Next, we develop a Branch & Bound algorithm, named BBLS, on the same principle introduced in [Dell'Amico & Martello, 1995] by allowing to prune out from the tree, subsets of solutions for which equivalent transformed instances indicate lower bounds that are not worthy relatively to the best solution so far. A shifting parameter S , added to BBLS, allows to speed up the computation of lower bounds and enhance the pruning effect of the algorithm. We provide additional Propositions in the goal to optimize BBLS. To assess the performance of the branching heuristics, benchmarks include many known algorithms as challengers. They show clearly, on literature instances, that BB-SLACK and BP-SLACK can really outperform the other algorithms on the vast majority of these instances. On the other hand, we also identify where they are not so efficient compared to LISTFIT (on instances whose job times start from 100 and above or when m grows) or LDM (on instances for which n or job times are larger). Making multiple calls to BBLS, varying the order of the jobs, or relying on sub-problems with a smaller number of machines or jobs, we have been able to define an ad hoc heuristic, named MBBLS, purposed to beat LDM and LISTFIT on the problematic subsets of instances. Appendices are joined to the paper in order to confirm or complete the results shown in the body of the paper on other instances, as perfect matching instances (Appendix A), instances found online (Appendix B) or instances created through independent generators (Appendix C). Future works are envisioned as benchmarking our heuristics on other instance families, as bin-packing instances. Evaluating these heuristics against the metaheuristics mentioned in introduction should also be interesting.

Figure 15: MBBLS versus LDM on DCU instances (see table 3)



Acknowledgements

We gratefully acknowledge the support of the Centre Blaise Pascal’s IT test platform at ENS de Lyon (Lyon, France) for computing facilities. The platform operates the SIDUS solution [Emmanuel Quemener, 2014]. We also thank anonymous reviewers for their valuable comments.

References

- [Chen et al., 2012] Chen, J., Pan, Q.-K., Wang, L., & Li, J.-Q. (2012). A hybrid dynamic harmony search algorithm for identical parallel machines scheduling. *Engineering Optimization*, 44(2), 209–224.
- [Coffman et al., 1978] Coffman, Jr, E. G., Garey, M. R., & Johnson, D. S. (1978). An application of bin-packing to multiprocessor scheduling. *SIAM Journal on Computing*, 7(1), 1–17.
- [Della Croce & Scatamacchia, 2020] Della Croce, F. & Scatamacchia, R. (2020). The longest processing time rule for identical parallel machines revisited. *Journal of Scheduling*, 23(2), 163–176.
- [Dell’Amico & Martello, 1995] Dell’Amico, M. & Martello, S. (1995). Optimal scheduling of tasks on identical parallel processors. *ORSA Journal on Computing*, 7(2), 191–200.
- [Emmanuel Quemener, 2014] Emmanuel Quemener, M. C. (2014). Sidus, the solution for extreme deduplication of an operating system. *The Linux Journal*.
- [França et al., 1994] França, P. M., Gendreau, M., Laporte, G., & Müller, F. M. (1994). A composite heuristic for the identical parallel machine scheduling problem with minimum makespan objective. *Computers & operations research*, 21(2), 205–210.
- [Frangioni et al., 2004] Frangioni, A., Necciari, E., & Scutella, M. G. (2004). A multi-exchange neighborhood for minimum makespan parallel machine scheduling problems. *Journal of Combinatorial Optimization*, 8, 195–220.

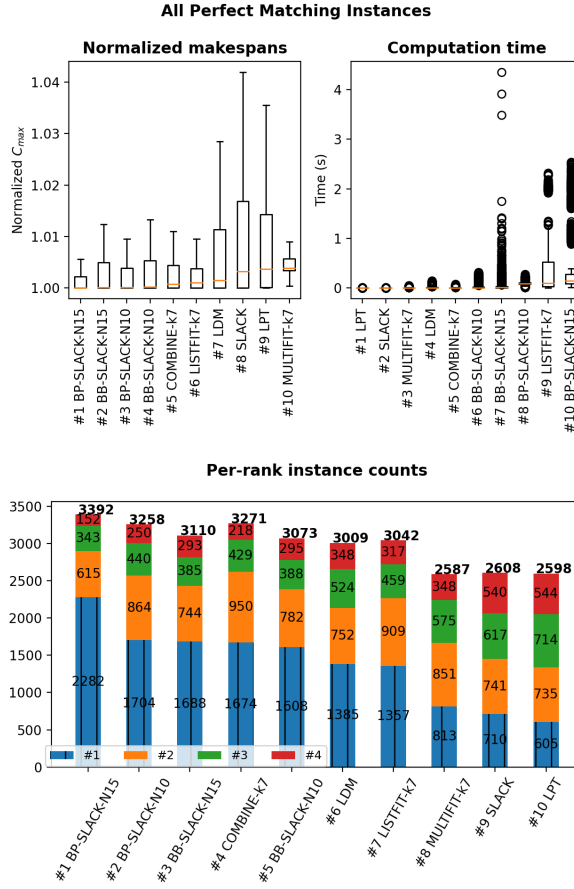
- [Garey & Johnson, 1979] Garey, M. R. & Johnson, D. S. (1979). Computers and intractability wh freeman and company. *New York*, (pp. 209–210).
- [Glass et al., 1994] Glass, C. A., Potts, C., & Shade, P. (1994). Unrelated parallel machine scheduling using local search. *Mathematical and Computer Modelling*, 20(2), 41–52.
- [Graham, 1969] Graham, R. L. (1969). Bounds on multiprocessing timing anomalies. *SIAM journal on Applied Mathematics*, 17(2), 416–429.
- [Graham et al., 1979] Graham, R. L., Lawler, E. L., Lenstra, J. K., & Kan, A. R. (1979). Optimization and approximation in deterministic sequencing and scheduling: a survey. 5, 287–326.
- [Gupta & Ruiz-Torres, 2001] Gupta, J. N. & Ruiz-Torres, A. J. (2001). A listfit heuristic for minimizing makespan on identical parallel machines. *Production Planning & Control*, 12(1), 28–36.
- [Hadj-Djilani, 2023a] Hadj-Djilani, H. (2023a). Dc & gr instances. Dataset on Zenodo, <https://doi.org/10.5281/zenodo.10203754>.
- [Hadj-Djilani, 2023b] Hadj-Djilani, H. (2023b). Gera instances. Dataset on Zenodo, <https://doi.org/10.5281/zenodo.10203980>.
- [Hadj-Djilani, 2023c] Hadj-Djilani, H. (2023c). Perfect matching instances. Dataset on Zenodo, <https://doi.org/10.5281/zenodo.10400559>.
- [Hadj-Djilani, 2023d] Hadj-Djilani, H. (2023d). Pimsgen instances. Dataset on Zenodo, <https://doi.org/10.5281/zenodo.10204011>.
- [Hadj-Djilani, 2025] Hadj-Djilani, H. (2025). B*ls exp code. Software on Zenodo, <https://doi.org/10.5281/zenodo.15433697>.
- [Hochbaum & Shmoys, 1987] Hochbaum, D. S. & Shmoys, D. B. (1987). Using dual approximation algorithms for scheduling problems theoretical and practical results. *Journal of the ACM (JACM)*, 34(1), 144–162.
- [Karmarkar & Karp, 1982] Karmarkar, N. & Karp, R. M. (1982). The differencing method of set partitioning.
- [Kashan & Karimi, 2009] Kashan, A. H. & Karimi, B. (2009). A discrete particle swarm optimization algorithm for scheduling parallel machines. *Computers & Industrial Engineering*, 56(1), 216–223.
- [Lee & Massey, 1988] Lee, C.-Y. & Massey, J. D. (1988). Multiprocessor scheduling: combining lpt and multi-fit. *Discrete applied mathematics*, 20(3), 233–242.
- [Martello & Toth, 1990] Martello, S. & Toth, P. (1990). Lower bounds and reduction procedures for the bin packing problem. *Discrete applied mathematics*, 28(1), 59–70.
- [Michiels et al., 2003] Michiels, W., Korst, J., Aarts, E., et al. (2003). Performance ratios for the karmarkar-karp differencing method. *Electronic Notes in Discrete Mathematics*, 13, 71–75.
- [Min & Cheng, 1999] Min, L. & Cheng, W. (1999). A genetic algorithm for minimizing the makespan in the case of scheduling identical parallel machines. *Artificial Intelligence in Engineering*, 13(4), 399–403.
- [Mokotoff, 2004] Mokotoff, E. (2004). An exact algorithm for the identical parallel machine scheduling problem. *European Journal of Operational Research*, 152(3), 758–769.
- [Yibao et al., 2002] Yibao, C., Jianchu, Y., & Yifang, Z. (2002). Ant system based optimization algorithm and its applications in identical parallel machine scheduling. *Journal of Systems Engineering and Electronics*, 13(3), 78–85.

A Results on perfect matching instances

This appendix presents the results obtained with perfect matching instances respecting the exact same (m, n) combinations as instances presented in Table 3 but with a different law for picking the p_j 's. Indeed, a perfect matching instance (PMI) is an instance for which it exists a solution schedule that is such that $C_{\max} = \frac{\sum_{1 \leq j \leq n} p_j}{m} = L_0$ which is also the completion time of all machines. The main interest of this kind of instances is that the optimal makespan is easily computable. Our PMI's were built as follows:

1. Generate an instance $I = (\{p_j\}_{1 \leq j \leq n}, m)$ as in Table 3.
2. Solve it with LPT and save $C_{\max}^{LPT}$. The goal is to obtain a reference makespan for our PMI.
3. For building the PMI start by creating the m blocks representing the machine completion times all equal to $C_{\max}^{LPT}$. Then cut these blocks in random positions until n sub-blocks are obtained. The n sub-blocks identify the job processing times of the PMI.

Figure 16: Overall performance on PMI's of our branching heuristics relatively to well-known algorithms



Note that a simpler method is to compute a LPT schedule and then fills the gaps to C_{max}^{LPT} for each machine by introducing new jobs or by enlarging pre-existing ones. However this method is kind of biased because it tends to advantage LPT by design for much of the instances.

The Figure 16 presents the results obtained with PMI's running the same experiment as in Section 5.2. The exact instances used to produce this figures are available online [Hadj-Djilani, 2023c]. The main change we notice is that COMBINE is a way more efficient with these PMI's. Still BB-SLACK and BP-SLACK algorithms stay in the top five, as in Section 4, according to both median normalized makespans and counts of first ranks. Because we use PMI's we are able to see that the gap between solutions and optima is about 1% of $L_0 = C_{max}^*$ for BP-SLACK runs and a little more for BB-SLACK's. For comparison this gap can go up to 4% in SLACK solutions.

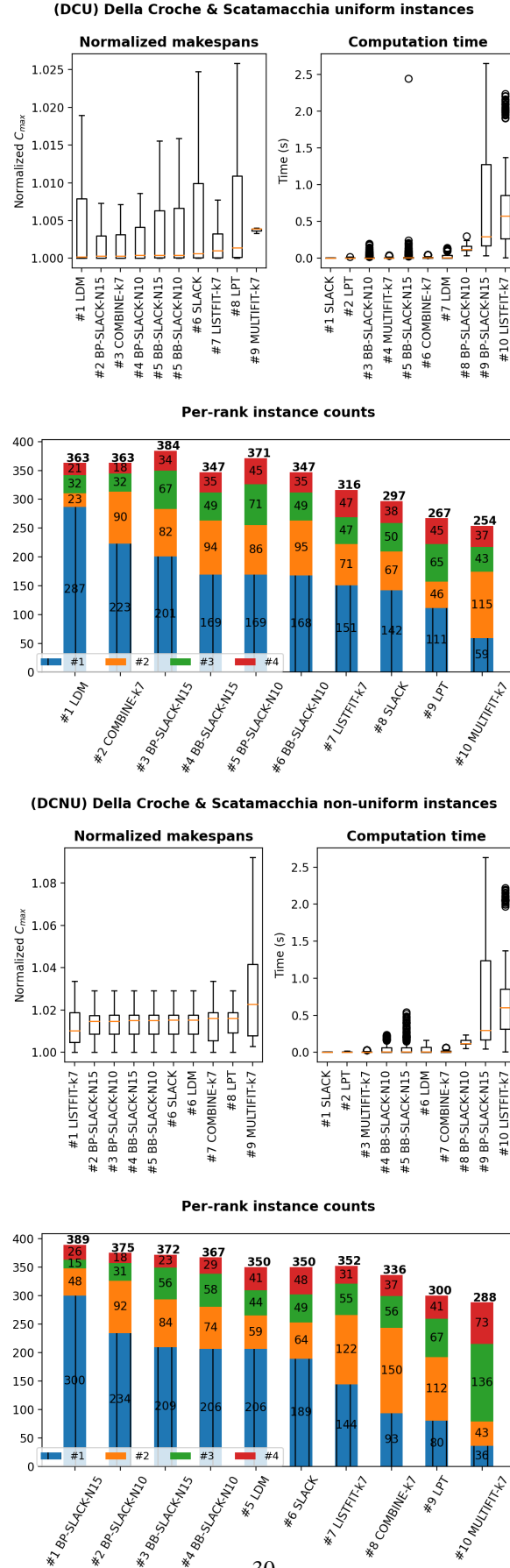
B Results on instances from University of Bologna

The figures 17 and 18 are the results obtained with instances found on the website of the University of Bologna in this page:

<https://site.unibo.it/operations-research/en/research/library-of-codes-and-instances-1>¹³

¹³In the section named "Library of Instances on the P—Cmax Problem". The instances are located in the archived linked here <https://site.unibo.it/operations-research/en/research/library-of-codes-and-instances-1/cmax.zip/@download/file/cmax.zip>.

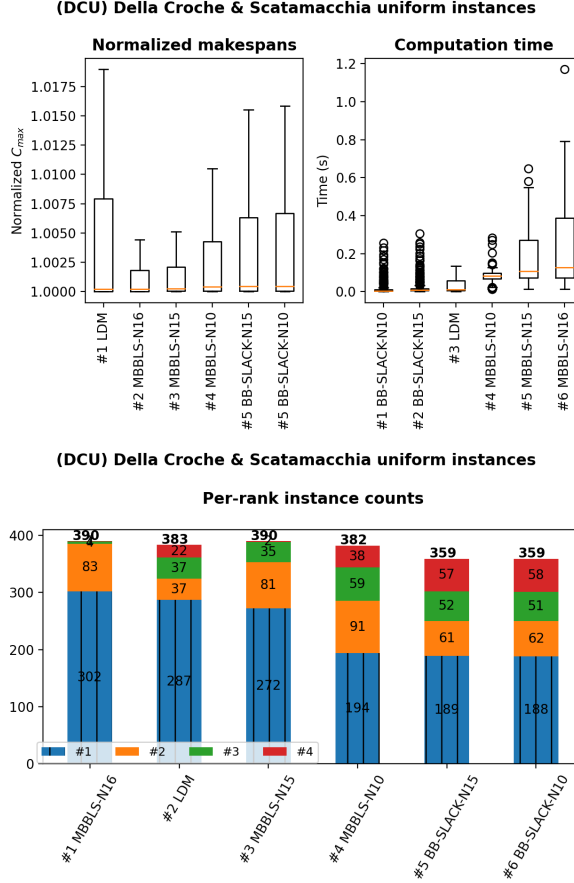
Figure 17: Benchmark on DCU and DCNU instances from University of Bologna (see Table 3)



We used them to run the same experiment as in Figures 5 and 15 discussed respectively in Sections 5.2.2 and 6. Note that in our experiments we used 20 instances per valid combination of instance parameters while for the instances here only 10 instances per combination were generated. That’s why we have twice more instances in our paper experiments. The results are quite similar to what obtained in the body of the paper. We only note that LISTFIT manages to make the best median normalized makespan for non-uniform instances in Figure 17 while it was only ranked 6 with our set of instances as shown in Figure 5.

Nevertheless it doesn’t change its rank in “Per-rank instance counts” bar chart which is 7 both in Figure 17 and Figure 5.

Figure 18: MBBLS versus LDM on DCU instances from University of Bologna (see Table 3)



C Results on Gera and PMSGGen instances

Gera and PMSGGen are generators of $P||C_{max}$ instances provided by CommaLAB, University of Pisa ¹⁴.

Gera generates 10 random instances for each valid combination of $(m, n) \in \{5, 10, 25\} \times \{10, 50, 100, 500, 1000\}$. The p_j 's are drawn in $[a = 1, b]$ according to different distributions:

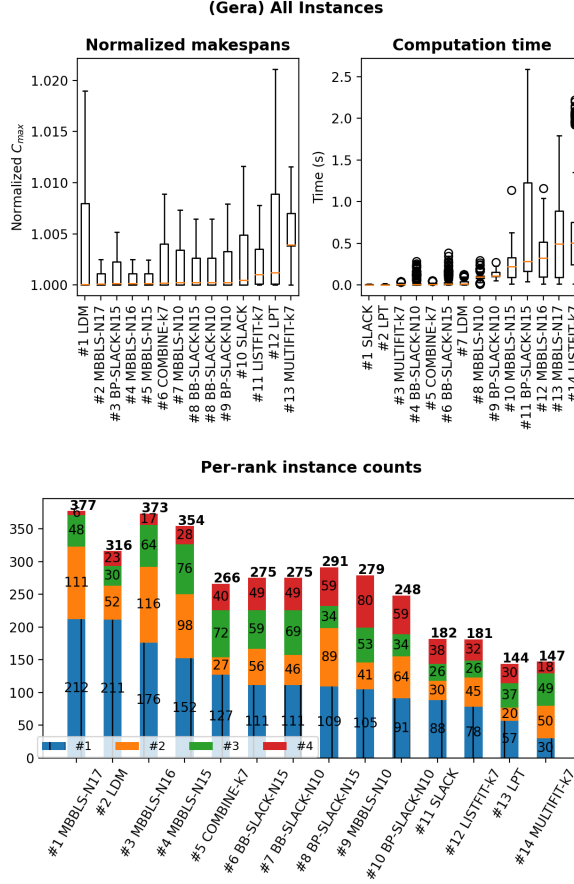
- A uniform law for which $b = 10000$ (geraun program).
- A normal law for which $b = 10000$ (geran program).
- An exponential law for which $b = 1000$ (gerae program).

It provides a total of 390 instances per run (130 per distribution). We performed the same experiment as defined in Section 5.2.1 using these instances. We grouped the results on all instances in the Figure 19.

Although the benchmark gives results that differ from what obtained in Section 5.2.2 and 6 we can see that MBBLS $N = 17$ manages to be at the top of all heuristics. Indeed, the number of first ranks obtained on instances

¹⁴At time of writing, the generators are downloaded here: <https://commalab.di.unipi.it/datasets/MS/>

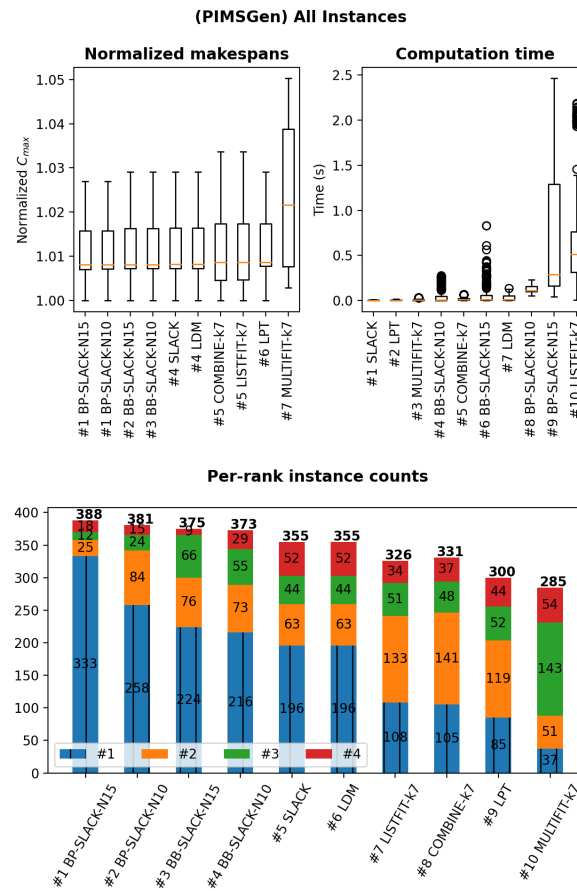
Figure 19: All Gera instances (normal, exponential and uniform laws) are grouped in this figure to show the overall performance of our branching heuristics relatively to well-known algorithms. The exact instances used to produce this figure are available online [Hadj-Djilani, 2023b].



is the greatest of all tested algorithms. Besides, even if LDM has the smallest median normalized makespan, the spanning of its makespans is much larger than that of MBBL5 $N = 17$ and other configurations of our heuristics. Finally, that performance is made possible in a reasonable execution time. Notably we see in the computation time boxplot that MBBL5 is faster than LISTFIT according to the median time.

PIMSGen generates 10 random instances for each of the 39 valid combinations $(m, n, b) \in \{5, 10, 25\} \times \{10, 50, 100, 500, 1000\} \times \{100, 1000, 10000\}$. It gives a total of 390 instances. The p_j 's of an instance are integers drawn uniformly in $[0.9(b - a), b]$ for 99% of them while the remaining ones (1% of the instances) are picked in $[a, 0.2(b - a)]$ according to a uniform distribution too. Note that this instance generation law is almost the same as the one defined for DCNU instances from the Table 3 (the only difference is that only 98% of the p_j 's are picked in the largest subinterval for DCNU instances and not 99% as PIMSGen does). Hence Figure 20 is a quite good confirmation of the results shown in Figure 5.

Figure 20: Benchmark including all PIMSGen instances. Exact instances used to produce this figure are available online [Hadj-Djilani, 2023d]



D Wilcoxon tests for comparison of algorithms' normalized makespan distributions

This appendix provides Wilcoxon statistical test p-values obtained for all instances considered in experiment of Section 5 and separately on each family of instances considered in Table 3.

	LPT	SLACK	LDM	COMBINE	MULTIFIT	LISTFIT
BB-SLACK-N15	0	1.2e-267	3.3e-73	7.6e-29	6.5e-121	3.6e-04
BB-SLACK-N10	0	1.4e-260	2.5e-65	1.3e-21	3.7e-107	7.3e-02
BP-SLACK-N15	0	0	1.4e-137	8.8e-75	4.0e-194	2.8e-30
BP-SLACK-N10	0	7.3e-303	1.3e-86	5.9e-39	3.5e-143	3.0e-07

Table 5: **Wilcoxon test p-values** comparing BPLS/BBLS makespan distributions to those of other algorithms for **all instances** (see Figure 4 and Table 3)

	BB-SLACK-N15	BB-SLACK-N10	BP-SLACK-N15	BP-SLACK-N10
BB-SLACK-N15		2.3e-09	1.1e-68	1.9e-05
BB-SLACK-N10	2.3e-09		6.6e-88	1.0e-11
BP-SLACK-N15	1.1e-68	6.6e-88		4.7e-114
BP-SLACK-N10	1.9e-05	1.0e-11	4.7e-114	

Table 6: **Wilcoxon test p-values** comparing BPLS/BBLS makespan distributions between them for **all instances** (see Figure 4 and Table 3))

	LPT	SLACK	LDM	COMBINE	MULTIFIT	LISTFIT
BB-SLACK-N15	4.9e-69	1.1e-40	3.8e-15	4.3e-11	1.6e-32	2.2e-08
BB-SLACK-N10	4.9e-69	1.1e-40	2.2e-15	4.3e-11	1.6e-32	2.2e-08
BP-SLACK-N15	1.2e-77	4.6e-54	1.6e-07	1.5e-07	2.4e-50	1.6e-16
BP-SLACK-N10	5.6e-75	1.1e-48	8.1e-13	4.6e-12	3.2e-43	3.6e-12

Table 7: **Wilcoxon test p-values** comparing BPLS/BBLS makespan distributions to those of other algorithms for **DCU instances** (see Figure 5 and Table 3)

	BB-SLACK-N15	BB-SLACK-N10	BP-SLACK-N15	BP-SLACK-N10
BB-SLACK-N15		8.9e-01	2.0e-09	4.0e-02
BB-SLACK-N10	8.9e-01		2.0e-09	4.0e-02
BP-SLACK-N15	2.0e-09	2.0e-09		2.0e-24
BP-SLACK-N10	4.0e-02	4.0e-02	2.0e-24	

Table 8: **Wilcoxon test p-values** comparing BPLS/BBLS makespan distributions between them for **DCU instances** (see Figure 5 and Table 3)

	LPT	SLACK	LDM	COMBINE	MULTIFIT	LISTFIT
BB-SLACK-N15	9.2e-120	4.9e-27	1.1e-24	2.2e-45	2.9e-88	7.0e-21
BB-SLACK-N10	9.2e-120	3.7e-23	6.6e-21	7.4e-45	4.6e-88	7.1e-21
BP-SLACK-N15	1.2e-119	2.7e-38	6.5e-36	1.0e-47	1.2e-88	7.4e-21
BP-SLACK-N10	1.6e-119	6.1e-30	1.1e-27	3.7e-47	2.8e-88	8.6e-21

Table 9: **Wilcoxon test p-values** comparing BPLS/BBLS makespan distributions to those of other algorithms for **DCNU instances** (see Figure 5 and Table 3)

	BB-SLACK-N15	BB-SLACK-N10	BP-SLACK-N15	BP-SLACK-N10
BB-SLACK-N15		1.4e-04	3.5e-17	3.7e-07
BB-SLACK-N10	1.4e-04		2.1e-20	2.3e-09
BP-SLACK-N15	3.5e-17	2.1e-20		3.0e-19
BP-SLACK-N10	3.7e-07	2.3e-09	3.0e-19	

Table 10: **Wilcoxon test p-values** comparing BPLS/BBLS makespan distributions between them for **DCNU instances** (see Figure 5 and Table 3)

	LPT	SLACK	LDM	COMBINE	MULTIFIT	LISTFIT
BB-SLACK-N15	3.0e-25	3.7e-06	7.4e-05	5.2e-10	1.9e-21	1.4e-03
BB-SLACK-N10	3.0e-25	3.7e-06	7.4e-05	5.2e-10	1.9e-21	1.4e-03
BP-SLACK-N15	1.7e-25	2.5e-06	5.1e-05	2.6e-10	8.8e-22	1.0e-03
BP-SLACK-N10	3.0e-25	3.7e-06	7.7e-05	3.9e-10	8.8e-22	1.3e-03

Table 11: **Wilcoxon test p-values** comparing BPLS/BBLs makespan distributions to those of other algorithms for **GR1 instances**(see Figure 7 and Table 3)

	BB-SLACK-N15	BB-SLACK-N10	BP-SLACK-N15	BP-SLACK-N10
BB-SLACK-N15		1.0e+00	9.2e-01	1.0e+00
BB-SLACK-N10	1.0e+00		9.2e-01	1.0e+00
BP-SLACK-N15	9.2e-01	9.2e-01		9.2e-01
BP-SLACK-N10	1.0e+00	1.0e+00	9.2e-01	

Table 12: **Wilcoxon test p-values** comparing BPLS/BBLs makespan distributions between them for **GR1 instances** (see Figure 7 and Table 3)

	LPT	SLACK	LDM	COMBINE	MULTIFIT	LISTFIT
BB-SLACK-N15	5.3e-81	3.6e-68	3.6e-10	1.2e-51	4.4e-64	2.9e-33
BB-SLACK-N10	5.5e-81	5.6e-68	7.3e-09	1.4e-51	4.6e-64	8.0e-33
BP-SLACK-N15	4.6e-85	1.4e-72	6.4e-20	7.0e-67	1.4e-74	9.6e-50
BP-SLACK-N10	3.2e-83	2.3e-69	6.0e-08	5.2e-60	1.2e-70	9.0e-40

Table 13: **Wilcoxon test p-values** comparing BPLS/BBLs makespan distributions to those of other algorithms for **GR2 instances** (see Figure 7 and Table 3)

	BB-SLACK-N15	BB-SLACK-N10	BP-SLACK-N15	BP-SLACK-N10
BB-SLACK-N15		1.6e-01	1.3e-17	7.3e-01
BB-SLACK-N10	1.6e-01		2.7e-20	7.5e-01
BP-SLACK-N15	1.3e-17	2.7e-20		3.5e-31
BP-SLACK-N10	7.3e-01	7.5e-01	3.5e-31	

Table 14: **Wilcoxon test p-values** comparing BPLS/BBLs makespan distributions between them for **GR2 instances** (see Figure 7 and Table 3)

	LPT	SLACK	LDM	COMBINE	MULTIFIT	LISTFIT
BB-SLACK-N15	1.0e-156	3.6e-128	2.8e-102	1.7e-01	8.7e-01	3.9e-41
BB-SLACK-N10	1.3e-156	2.5e-126	5.1e-98	7.7e-04	3.4e-02	2.0e-55
BP-SLACK-N15	4.0e-158	5.1e-150	3.1e-137	3.5e-08	8.5e-11	8.4e-10
BP-SLACK-N10	2.1e-157	3.6e-143	2.3e-116	9.9e-01	2.5e-01	4.4e-38

Table 15: **Wilcoxon test p-values** comparing BPLS/BBLs makespan distributions to those of other algorithms for **GR3 instances** (see Figure 8 and Table 3)

	BB-SLACK-N15	BB-SLACK-N10	BP-SLACK-N15	BP-SLACK-N10
BB-SLACK-N15		3.6e-10	1.4e-36	2.0e-02
BB-SLACK-N10	3.6e-10		5.6e-53	4.4e-08
BP-SLACK-N15	1.4e-36	5.6e-53		1.3e-53
BP-SLACK-N10	2.0e-02	4.4e-08	1.3e-53	

Table 16: **Wilcoxon test p-values** comparing BPLS/BBLs makespan distributions between them for **GR3 instances** (see Figure 8 and Table 3)

E MILP comparison experiments

Algo.	Normed $C_{\max}$ median	All instances (see Figure 4 and Table 3)		CPU time (s) mean/median
		#opt/#total lower bound	#wins/#ties vs MILP	
LPT	1.017278	0.28	0.46	0.003/0.001
SLACK	1.013887	0.38	0.62	0.002/0.001
LDM	1.011864	0.41	0.66	0.017/0.003
COMBINE	1.009256	0.38	0.62	0.007/0.003
MULTIFIT	1.011609	0.26	0.38	0.004/0.002
LISTFIT	1.006538	0.43	0.68	0.346/0.106
BB-SLACK-N15	1.006904	0.49	0.75	0.086/0.008
BB-SLACK-N10	1.007477	0.48	0.74	0.039/0.005
BP-SLACK-N15	1.005063	0.51	0.78	0.440/0.152
BP-SLACK-N10	1.006416	0.49	0.75	0.105/0.090
MBBLS-N10	1.006166	0.49	0.75	0.093/0.082
MBBLS-N15	1.004025	0.51	0.78	0.329/0.102
MBBLS-N16	1.003928	0.52	0.79	0.528/0.105
MBBLS-N17	1.003839	0.52	0.79	0.885/0.105
MILP	1.008062	0.64		122.234/13.526

Table 17: **Comparison of results to exact MILP method** (performed with a time limitation of 300 seconds to find an optimal solution). The normed $C_{\max}$ median, column 2, are normalized using the lower bound L_1 (see Section 3.1). The “#opt/#total” column is the lower bound of the rate of optimal solutions found according to MILP guaranty. “#wins/#ties” column is the ratio of instances for which the considered algorithm found an equal or better solution than MILP. In last column are the median and mean computation times for each algorithm.

Algo.	Normed $C_{\max}$ median	DCU instances (see Figures 4, 5 and Table 3)		
		#opt/#total lower bound	#wins/#ties vs MILP	CPU time (s) mean/median
LPT	1.001148	0.38	0.78	0.006/0.002
SLACK	1.000558	0.44	0.89	0.004/0.002
LDM	1.000182	0.47	0.95	0.035/0.011
COMBINE	1.000334	0.46	0.93	0.011/0.007
MULTIFIT	1.003891	0.21	0.51	0.009/0.005
LISTFIT	1.001070	0.42	0.87	0.659/0.520
BB-SLACK-N15	1.000386	0.47	0.95	0.022/0.007
BB-SLACK-N10	1.000386	0.47	0.95	0.016/0.004
BP-SLACK-N15	1.000306	0.47	0.96	0.826/0.289
BP-SLACK-N10	1.000371	0.46	0.94	0.134/0.109
MBBLS-N10	1.000345	0.47	0.95	0.079/0.084
MBBLS-N15	1.000217	0.48	0.97	0.163/0.108
MBBLS-N16	1.000196	0.48	0.98	0.224/0.117
MBBLS-N17	1.000180	0.48	0.98	0.335/0.129
MILP	1.003537	0.49		166.877/300.154
Algo.	Normed $C_{\max}$ median	DCNU instances (see Figures 4, 5 and Table 3)		
		#opt/#total lower bound	#wins/#ties vs MILP	CPU time (s) mean/median
LPT	1.017059	0.23	0.61	0.006/0.002
SLACK	1.016017	0.28	0.75	0.004/0.002
LDM	1.016017	0.28	0.75	0.034/0.011
COMBINE	1.016703	0.24	0.70	0.013/0.010
MULTIFIT	1.019522	0.11	0.21	0.007/0.004
LISTFIT	1.016343	0.25	0.72	0.669/0.526
BB-SLACK-N15	1.015977	0.28	0.75	0.045/0.006
BB-SLACK-N10	1.015978	0.28	0.75	0.032/0.004
BP-SLACK-N15	1.015853	0.29	0.76	0.826/0.293
BP-SLACK-N10	1.015936	0.28	0.76	0.135/0.111
MBBLS-N10	1.016159	0.26	0.67	0.104/0.101
MBBLS-N15	1.016102	0.26	0.67	0.687/0.307
MBBLS-N16	1.016081	0.26	0.67	1.221/0.433
MBBLS-N17	1.016052	0.26	0.67	2.235/0.596
MILP	1.014667	0.30		227.958/312.498

Table 18: **Comparison of results to exact MILP method for DCU & DCNU instances.** Refer to Table 17 about meaning of columns.

Algo.	Normed $C_{\max}$ median	GR1 instances (see Figures 4, 7 and Table 3)		CPU time (s) mean/median
		#opt/#total lower bound	#wins/#ties vs MILP	
LPT	1.023701	0.62	0.62	0.000/0.000
SLACK	1.018350	0.81	0.81	0.000/0.000
LDM	1.017726	0.83	0.83	0.001/0.001
COMBINE	1.019185	0.76	0.76	0.001/0.001
MULTIFIT	1.030489	0.60	0.60	0.001/0.001
LISTFIT	1.017391	0.84	0.84	0.021/0.019
BB-SLACK-N15	1.013975	0.92	0.92	0.004/0.002
BB-SLACK-N10	1.013975	0.92	0.92	0.003/0.001
BP-SLACK-N15	1.013845	0.93	0.93	0.062/0.074
BP-SLACK-N10	1.013845	0.92	0.92	0.062/0.075
MBBLS-N10	1.013825	0.93	0.93	0.043/0.019
MBBLS-N15	1.013825	0.93	0.93	0.052/0.019
MBBLS-N16	1.013785	0.93	0.93	0.050/0.018
MBBLS-N17	1.013785	0.93	0.93	0.050/0.017
MILP	1.013785	1.00		4.920/4.762

Algo.	Normed $C_{\max}$ median	GR2 instances (see Figures 4, 7 and Table 3)		CPU time (s) mean/median
		#opt/#total lower bound	#wins/#ties vs MILP	
LPT	1.007949	0.25	0.29	0.001/0.001
SLACK	1.001410	0.49	0.60	0.001/0.001
LDM	1.000587	0.57	0.71	0.004/0.003
COMBINE	1.003210	0.36	0.43	0.003/0.003
MULTIFIT	1.004486	0.29	0.33	0.002/0.001
LISTFIT	1.002880	0.44	0.51	0.150/0.076
BB-SLACK-N15	1.000482	0.68	0.84	0.051/0.014
BB-SLACK-N10	1.000509	0.66	0.82	0.037/0.008
BP-SLACK-N15	1.000319	0.71	0.90	0.155/0.148
BP-SLACK-N10	1.000575	0.67	0.84	0.081/0.083
MBBLS-N10	1.000969	0.56	0.75	0.085/0.070
MBBLS-N15	1.000526	0.59	0.79	0.390/0.103
MBBLS-N16	1.000500	0.60	0.80	0.678/0.104
MBBLS-N17	1.000398	0.62	0.82	1.178/0.106
MILP	1.001420	0.79		71.833/12.458

Algo.	Normed $C_{\max}$ median	GR3 instances (see Figures 4, 8 and Table 3)		CPU time (s) mean/median
		#opt/#total lower bound	#wins/#ties vs MILP	
LPT	1.057320	0.12	0.13	0.001/0.001
SLACK	1.037968	0.20	0.21	0.001/0.001
LDM	1.039006	0.23	0.25	0.002/0.002
COMBINE	1.021402	0.32	0.37	0.002/0.002
MULTIFIT	1.022659	0.30	0.34	0.001/0.001
LISTFIT	1.011432	0.45	0.51	0.061/0.046
BB-SLACK-N15	1.014228	0.41	0.45	0.220/0.030
BB-SLACK-N10	1.015715	0.39	0.43	0.077/0.026
BP-SLACK-N15	1.010084	0.46	0.54	0.117/0.114
BP-SLACK-N10	1.014159	0.41	0.46	0.087/0.087
MBBLS-N10	1.010655	0.50	0.57	0.119/0.079
MBBLS-N15	1.006331	0.55	0.67	0.243/0.097
MBBLS-N16	1.006174	0.55	0.68	0.307/0.098
MBBLS-N17	1.006105	0.55	0.68	0.382/0.097
MILP	1.011559	0.81		72.405/8.821

Table 19: **Comparison of results to exact MILP method for GR1, GR2, GR3 instances.** Refer to Table 17 about meaning of columns.

F Complementary propositions about greater pruning efficiency in relation to S

The Propositions F.1, F.2 give an insight on why a greater $|J_1|$ and more generally larger aggregate jobs, in formula of bounds $B_\alpha(L, \bar{p})$, $B_\beta(L, \bar{p})$, are better to obtain a wider pruning and finally the Proposition F.3 shows that a larger S facilitates a larger $|J_1|$.

In order to introduce the Proposition F.1 we present some notations and the related context. Let $(k-1)$ and k the previous and next jobs assigned by BBLS:

- their processing times p_{k-1} and p_k ,
- their respective BBLS node transformed instances I'_{k-1} and I'_k ,
- with job subsets $J_{i \in \{1,2,3\}, k-1}(L, \bar{p})$ and $J_{i \in \{1,2,3\}, k}(L, \bar{p})$ as defined in Section 3.1.

k is assigned to machine i_k whose previous load is C_{i_k} . This load corresponds to a job in I'_{k-1} which is denoted a_{i_k} . The job in I'_k formed by aggregation of k and a_{i_k} is denoted a_k and its processing time is $(C_{i_k} + p_k)$.

Proposition F.1. *If the next conditions are fulfilled for all $\bar{p}$ then I'_k node will not be pruned out by BBLS.*

- (i) $k, a_{i_k} \in J_{3, k-1}(L, \bar{p})$,
- (ii) $a_k \in J_{3, k}(L, \bar{p})$,
- (iii) or alternatively to (i), (ii): $p_k + C_{i_k} < \bar{p}$,
- (iv) $|J_{1, k}(L, \bar{p})| = |J_{1, k-1}(L, \bar{p})|$.

Proof. The proof is obtained by intersecting lemma Propositions G.1, G.2 of Appendix G. Besides, note that the parent node (instance I'_{k-1}) has not been pruned out or the child node (instance I'_k) wouldn't even be considered. $\square$

Additionally, an example of instance and assignment verifying conditions (i, ii, iv) of Proposition F.1 is given in Figure 21. The original instance of this example is $I = (P = (p_1, \dots, p_n) = (15, 14, 11, 7, 5, 3, 3, 1, 19, 18, 16, 16), m = 4)$ in SLACK order. Here $S = 0$, the first 4 jobs are assigned according to LS. For the assignment of the job $k = 6$ on the first machine available $i_1 = 4$, with $p_{k=6} = 3$, the transformed instance is $I'_{k=6} = (P' = (19, 18, 16, 16, 16, 15, 14, 10, 3, 1), m = 4)$. Given the current best makespan $C_{\max}^b = 33$ and $L = 32$ with $C_{i_k} = p_{a_{i_k}} = 7$, $p_{a_k} = 10$, $\bar{p} = 1$, we can verify that the proposition conditions are verified:

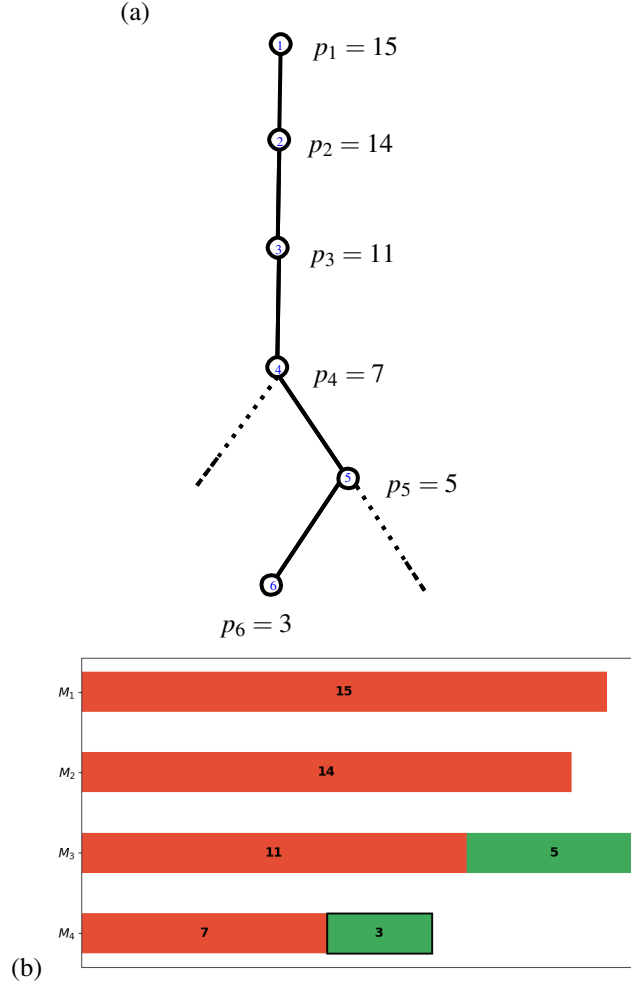


Figure 21: Proposition F.1 example verifying conditions (i, ii, iv) for one $\bar{p}$. (a) is the partial BBLs tree up to assignment of the job $k = 6$. (b) is a representation of the corresponding partial schedule in which the job is framed in black. The bars indicate the processing times p_j .

- Condition (i) is verified because $\bar{p} = 1$ and $L/2 = 16$, $p_k = 3$, $p_{a_k} = C_k = 7$ so $k, a_k \in J_{3,k-1}(L, \bar{p} = 1)$.
- Condition (ii) is verified because $p_{a_k} = C_k + p_k = 7 + 3 = 10$ and so $a_k \in J_{3,k}(L, \bar{p} = 1)$.
- Condition (iv) is verified because $L - \bar{p} = 31$, hence $J_{1,k-1} = J_{1,k} = \emptyset$.

Hence, we conclude that the node for this partial solution/schedule must not be pruned out because of this $\bar{p}$ and this consists with conditions (7) ($B_\alpha(L = 32, \bar{p} = 1) = 4 \leq m = 4$, $B_\beta(L = 32, \bar{p} = 1) = 2 \leq m = 4$).

Another example for condition (iii, iv) is given Figure 22 of Appendix H. This proposition has been also verified in practice on all instances (and implied BBLs trees) of our experiment framework shown in Table 3.

Proposition F.2. Let j , the next job to assign by BBLs, p_j its processing time and I'_j the corresponding node transformed instance. If all the m aggregated jobs of I'_j are in $J_1(L, \bar{p})$ and it exists a p_j such that $p_j \geq \bar{p}$ then the node will be pruned out.

Proof. All the m machines are already occupied, each one with a job of processing time greater than $L - \bar{p}$, then no more job in $J_3(L, \bar{p})$ can fit into any machine without overpassing the L capacity. Condition (7) will be verified and the node pruned out. $\square$

Proposition F.3. Let I'_1 the first transformed instance of a BBLs tree and the new aggregated jobs $j'_a \in I'_1$ not initially in the problem instance I . The number of these new aggregated jobs is $\#\{j'_a\} = \min(S + 1, m)$. Each one of their processing times $p_{j'_a}$ is an increasing function of S .

Proof. The root node of a BBLs tree has no transformed instance since the corresponding job is assigned according to LS before BBLs start. The first left child node has no transform instance either because BBLs needs to initialize $C_{\max}^b$ using LS. The first transformed instance I'_1 is computed for the root node's right child. In I'_1 , there is a number of $\min(S+1, m)$ new jobs formed by aggregation because the first m jobs of the list are assigned with LS and S supplementary jobs are also assigned with LS on the machines to form at most m new aggregated jobs in I'_1 . By construction, the greater is S the larger are the m aggregated jobs processing times. $\square$

Proposition F.3 pushes us to increase the value of parameter S to increase $|J_1|$ and $|J_2|$ in order to obtain a pruning out of a node as soon as possible in BBLs tree evaluation.

The preceding propositions are experimented through the protocol described in Appendix F.

G Lemma propositions for Proposition F.1

The propositions in this section allow to prove Proposition F.1 by intersecting their respective predicates. Indeed, the following propositions permit to infer that a child node will not be pruned out according to only one bound. The other bound can nevertheless infer a pruning of the node and conversely. Precisely, Proposition G.1 tells that a child node will not be pruned out according to bound B_α (5) and likewise Proposition G.2 implies that no pruning will occur for a node according to bound B_β (6).

Proposition G.1. *Let $(k-1)$ and k the previous and next jobs assigned by BBLs:*

- *their processing times p_{k-1} and p_k ,*
- *their respective BBLs node transformed instances I'_{k-1} and I'_k ,*
- *with job subsets $J_{i \in \{1,2,3\}, k-1}(L, \bar{p})$ and $J_{i \in \{1,2,3\}, k}(L, \bar{p})$ as defined in Section 3.1.*

k is assigned to machine i_k whose previous load is C_{i_k} . This load corresponds to a job in I'_{k-1} which is denoted a_{i_k} . The job in I'_k formed by aggregation of k and a_{i_k} is denoted a_k and its processing time is $(C_{i_k} + p_k)$.

Therefore, because the parent node has not been pruned out, if the next conditions are fulfilled for all $\bar{p}$ then I'_k node is such that $m \geq B_{\alpha, k-1}(L, \bar{p}) \geq B_{\alpha, k}(L, \bar{p})$.

- (i) $k \in J_{3, k-1}(L, \bar{p})$,
- (ii) $a_{i_k} \in J_{2, k-1}(L, \bar{p}) \cup J_{3, k-1}(L, \bar{p})$,
- (iii) or alternatively to (i), (ii): $p_k + C_{i_k} < \bar{p}$,
- (iv) $|J_{1, k}(L, \bar{p})| = |J_{1, k-1}(L, \bar{p})|$.

Proof. First, (iii) implies that $B_{\alpha, k-1}(L, \bar{p}) = B_{\alpha, k}(L, \bar{p})$. Second, consider the space left, for $J_{3, k}$ jobs after the assignment of jobs in J_1 and J_2 . Note that this space verification is made, according to $B_\alpha(L, \bar{p})$, as if the J_3 jobs were assigned preemptively.

$$\text{space}_{J_{3, k}} = mL - |J_{1, k}|L - \sum_{j \in J_{2, k}} p_j$$

Considering conditions (i, ii, iv), it comes that:

- $\text{space}_{J_{3, k}} = \text{space}_{J_{3, k-1}} - p_k$ if $a_k \in J_{2, k}, a_{i_k} \in J_{2, k-1}$,
- $\text{space}_{J_{3, k}} = \text{space}_{J_{3, k-1}} - p_k - C_{i_k}$ if $a_k \in J_{2, k}, a_{i_k} \in J_{3, k-1}$,
- $\text{space}_{J_{3, k}} = \text{space}_{J_{3, k-1}}$ if $a_k \in J_{3, k}, a_{i_k} \in J_{3, k-1}$.

Thus, because jobs a_{i_k}, k in I'_{k-1} of times C_{i_k}, p_k are replaced in I'_k by one aggregated job a_k in $J_{3, k} \cup J_{2, k}$ whose time is $(C_{i_k} + p_k)$, it is left exactly the same space as in parent node for jobs $J_{3, k-1} - \{p_k, a_{i_k}\}$. Therefore, because the parent node has not been pruned out, we have: $m \geq B_{\alpha, k-1}(L, \bar{p}) \geq B_{\alpha, k}(L, \bar{p})$. $\square$

Proposition G.2. *In the same context as Proposition G.1 (without conditions (i-iv)), if the following conditions are fulfilled then $m \geq B_{\beta, k-1}(L, \bar{p}) \geq B_{\beta, k}(L, \bar{p})$.*

- (a) $a_{i_k} \in J_{3, k-1}(L, \bar{p})$,
- (b) $a_k \in J_{3, k}(L, \bar{p})$,
- (c) or alternatively to (i), (ii): $p_k + C_{i_k} < \bar{p}$,
- (d) $|J_{1, k}(L, \bar{p})| = |J_{1, k-1}(L, \bar{p})|$.

Proof. First, (c) implies that $B_{\beta,k-1}(L, \bar{p}) = B_{\beta,k}(L, \bar{p})$. Second, note that $J_{1,k}(L, \bar{p}) = J_{1,k-1}(L, \bar{p})$, $J_{2,k}(L, \bar{p}) = J_{2,k-1}(L, \bar{p})$, and the replacement of jobs k, a_k of processing times p_k, C_{i_k} (from I'_{k-1}) by the job a_k of processing time $(C_{i_k} + p_k)$ verifies $|J_{3,k}(L, \bar{p})| \leq |J_{3,k-1}(L, \bar{p})|$. Then, because parent node has not been pruned out, $m \geq B_{\beta,k-1}(L, \bar{p}) \geq B_{\beta,k}(L, \bar{p})$. $\square$

Let us discuss now why the two previous propositions are each one specific to only one bound. The Proposition G.1 does not apply to $B_{\beta,k}(L, \bar{p})$ because it is possible to build counterexamples for which $B_{\beta,k-1}(L, \bar{p}) < B_{\beta,k}(L, \bar{p})$ even if conditions (i, ii, iv) are true:

- If $a_{i_k} \in J_{3,k-1}(L, \bar{p})$ and $a_k \in J_{2,k}$, we have $|J_{3,k}| = |J_{3,k-1}| - 2$ and $|J_{2,k}| = |J_{2,k-1}| + 1$. Now suppose that $p_k + C_{i_k} = L - \bar{p}$ (maximum time of a J_2 job), that is $\left\lfloor \frac{L - (p_k + C_{i_k})}{\bar{p}} \right\rfloor = 1$, so $B_{\beta,k} = |J_{1,k-1}| + |J_{2,k-1}| + 1 + \max(0, \left\lfloor \frac{|J_{3,k-1}| - (\sum_{j \in J_{2,k-1}} \left\lfloor \frac{L - p_j}{\bar{p}} \right\rfloor)}{\left\lfloor \frac{L}{\bar{p}} \right\rfloor} \right\rfloor - \frac{3}{\left\lfloor \frac{L}{\bar{p}} \right\rfloor})$. Finally, denote $R = \frac{|J_{3,k-1}| - (\sum_{j \in J_{2,k-1}} \left\lfloor \frac{L - p_j}{\bar{p}} \right\rfloor)}{\left\lfloor \frac{L}{\bar{p}} \right\rfloor}$ if $3 < \left\lfloor \frac{L}{\bar{p}} \right\rfloor$ and $(R - \lfloor R \rfloor) > \frac{3}{\left\lfloor \frac{L}{\bar{p}} \right\rfloor}$ it comes that $B_{\beta,k} = B_{\beta,k-1} + 1 > B_{\beta,k-1}$.

Likewise, Proposition G.2 does not apply to $B_{\alpha,k}(L, \bar{p})$ because it is possible to build counterexamples for which $B_{\alpha,k-1}(L, \bar{p}) < B_{\alpha,k}(L, \bar{p})$ even if conditions (a, b, d) are true:

- First, assume that $p_k < \bar{p}$, that is p_k was not taken into account in parent node calculation of $B_{\alpha,k-1}(L, \bar{p})$. Second, add the constraint

$$\text{space}_{J_{3,k-1}} - \sum_{j \in J_{3,k-1}} p_j < p_k$$

then because $\sum_{j \in J_{3,k}} = \sum_{j \in J_{3,k-1}} p_j + p_k$ it comes that $mL < |J_{1,k}|L + \sum_{j \in J_{2,k}} p_j + \sum_{j \in J_{3,k}} p_j$, so there is not enough space for $J_{3,k}$ jobs to fit into m machines/bins of capacity L and $B_{\alpha,k} > m \geq B_{\alpha,k-1}$.

H Proposition F.1 secondary example

In Figure 22, an example for Proposition F.1 is shown in which the original instance is $I = (P = (p_1, \dots, p_n) = (13, 6, 1, 18, 16, 15, 20, 19, 19), m = 3)$ w.r.t. SLACK order. Here $S = 0$, the first 4 jobs are assigned according to LS. For the assignment of the job $k = 6$, with $p_{k=6} = 15$, the transformed instance is $I'_{k=6} = (P' = (20, 19, 19, 29, 24, 16), m = 3)$. Given the current best makespan $C_{\max}^b = 47$ and $L = 46$ with $C_{i_k} = p_{a_{i_k}} = 1$, $p_{a_k} = C_{i_k} + p_k = 16$, $\bar{p} = 19$, we can verify that the proposition conditions are verified:

- Condition (iii) is verified because $p_{a_k} = C_{i_k} + p_k = 15 + 1 = 16 \leq \bar{p} = 19$.
- Condition (iv) is verified because $L - \bar{p} = 27$, hence $J_{1,k-1} = J_{1,k} = \{j' = 4\}$, with $p_{j'=4} = 29$.

Hence, we conclude that the node for this partial solution/schedule must not be pruned out because of this $\bar{p}$ and this is consistently verified with conditions (7) ($B_{\alpha}(L = 46, \bar{p} = 19) = 3 \leq m = 3, B_{\beta}(L = 46, \bar{p} = 19) = 3 \leq m = 3$).

I Testing the parameter S of BBLs

The goal of the experiment described here is to verify empirically the speedup effect of a greater S parameter as discussed in Section 3.3 in particular with Proposition 3.1.

The instance family tested is a subset of GR2 instances (see Table 3). Note that other instance subsets of Table 3 would allow to make similar observations on computation times and BBLs tree pruning. So we generate 200 pseudo-random instances for each pair (m, n) with $m \in \{4, 8\}$ and $n = 100$, that is a total of 400 instances. The processing times are picked in $\{100, \dots, 800\}$ according to a uniform law.

For our tests, we use BB-SLACK and BP-SLACK algorithms, that is BBLs and BPLs. In these runs N is always set to 10 but the value of S varies. BB-SLACK is ran two times, the first one with $S = 0$ and the second one with $S = 80$. For BP-SLACK, S is not used so it is equivalent to the configuration $S = 0$. The algorithm is here parallelized with at most 8 processes¹⁵. We intend to show that a larger value of S is interesting not only by speeding up the computation of BPP bounds (5), (6) but also by pruning widely the BB trees. BP-SLACK plays the role of the baseline to evaluate BB-SLACK performance in both cases $S = 0$ and $S = 80$. Indeed, remember that BBLs first objective remains to work faster than BPLs, without risking to produce poorer solutions. This is only possible when $S > 0$, otherwise BBLs (or BB-SLACK) produces the exact same solutions as BPLs (or BP-SLACK) as we verify in the following. We also verify that increasing S for a same N does not necessarily imply a loss of quality of the solutions.

¹⁵Precisely, 8 cores of a AMD Ryzen 5 2500U CPU

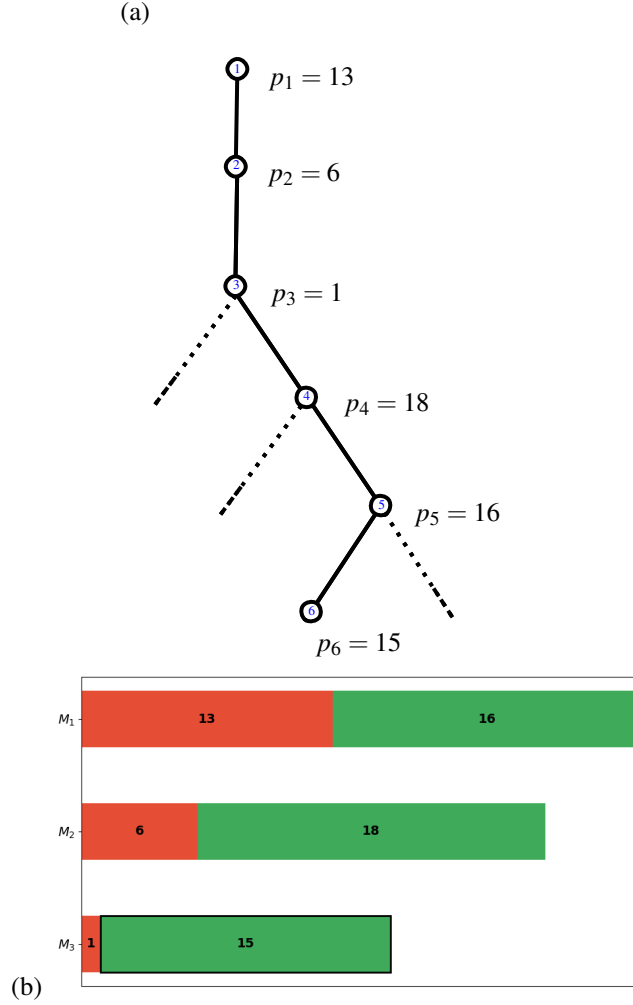


Figure 22: Proposition F.1 example verifying conditions (iii, iv) for one $\bar{p}$. (a) is the partial BBLs tree up to assignment of the job $k = 6$. (b) is a representation of the corresponding schedule in which the job is framed in black. The bars indicate the processing times p_j .

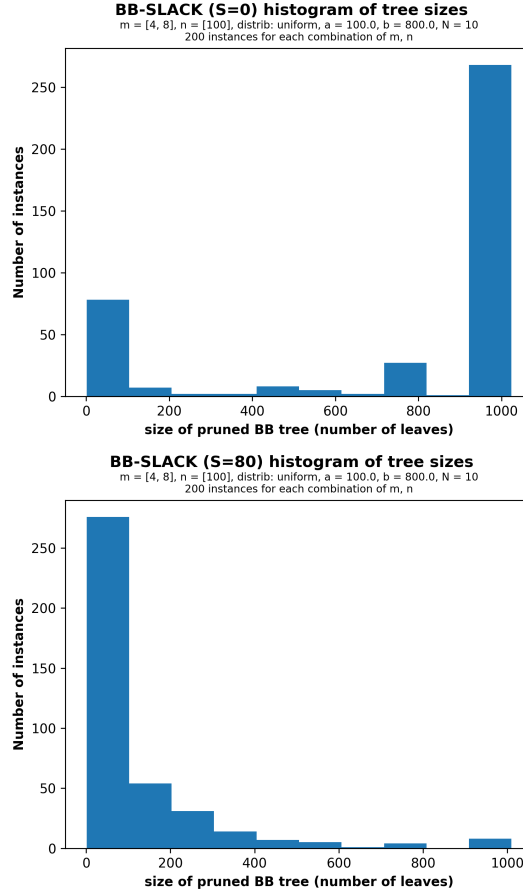
We present the results of this experiment in Figures 23, 24 and 25, where the top subplot is the case $S = 0$ for BB-SLACK while the bottom one is the case $S = 80$. In Figure 23 the histogram shows that $S = 80$ results in a wider pruning of the trees. Indeed, for most of the instances (more than 250 out of 400), the pruning is such that the number of leaves is lower than 100, while it can go up to $2^N = 1024$ with no pruning. Besides, no tree of BB-SLACK with $S = 80$ is full, with 1024 leaves, the greatest tree size being 1010 leaves. The histogram for $S = 0$ shows, on the contrary, that for most of the instances (about 300 over 400) there is almost no pruning in BB-SLACK run (35 trees are full with 1024 leaves). We can therefore verify that increasing S is efficient to achieve a wider pruning for many instances.

The speedup of BB-SLACK, compared to BP-SLACK, according to the size of the pruned tree is shown in Figure 24. We can observe that, even for a same size of tree, BB-SLACK is faster when $S = 80$. This experimentally confirms Proposition 3.1. For example, for a size of 10 leaves, BB-SLACK with $S = 0$ reached a speedup about 13, while BB-SLACK with $S = 80$ allows a speedup about 45.

With Figure 25 we can evaluate how S can affect the quality of the $P||C_{\max}$ solutions for the considered instance family. In the two subplots shown in this figure, the results are partitioned in 3 groups:

1. The left bars count for instances for which BB-SLACK pruned the tree and was faster than BP-SLACK.
2. The middle bars count the instances for which BB-SLACK pruned the tree but was slower than BP-SLACK

Figure 23: Histograms for the comparison of BB-SLACK tree sizes for $S = 0$ (on the top) and $S = 80$ (on the bottom)



3. The right bars count cases for which BB-SLACK was not able to prune the tree at all and hence was slower than BP-SLACK.

Note that cases 2 and 3 did not happen for $S = 80$, so there is no bar at all.

In the top subplot, $S = 0$, so that the same pair (N, S) is used for BB-SLACK and BP-SLACK. Hence their solutions are exactly the same. Indeed, the sizes of the red bars are equal to these of the blue bars, which means that, for all the instances, BB-SLACK (BBLs) and BP-SLACK (BPLs) returned the same makespans. We can also see that, for $S = 0$, on the most part of the instances BB-SLACK is slower than BP-SLACK.

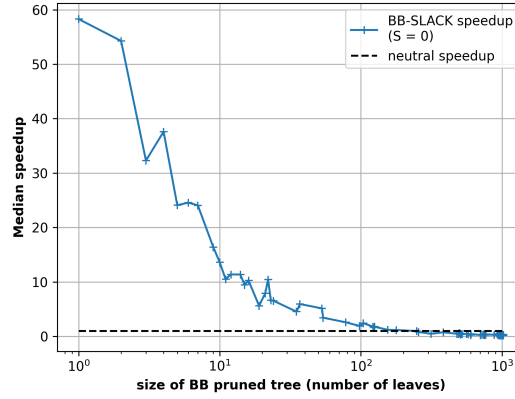
In the bottom subplot, where $S = 80$, the instances for which the makespans of BB-SLACK and BP-SLACK are equal are accounted in red, the instances for which BB-SLACK is better than BP-SLACK are in green and the ones for which BP-SLACK does better are in orange. The green part is the largest for this instance family, which means that increasing S does not lead to poorer performance in that case. It is noteworthy in this figure that again, for all of the instances, BB-SLACK set with $S = 80$ is faster than BP-SLACK running on the same instances, even if the tree is not largely pruned. Indeed, up to a tree of 1010 leaves BB-SLACK does faster than BP-SLACK while when $S = 0$ the tree needs to be smaller than 178 leaves for BB-SLACK to do faster. We see in Section 5.2 that S can nevertheless impact the performance of BB-SLACK, in a sense that for a same N , the smaller S the better the solutions even if that is not always true for specific subgroups of instances. This tendency is globally observed on the considered instances.

J Other properties for optimization of the pruning process

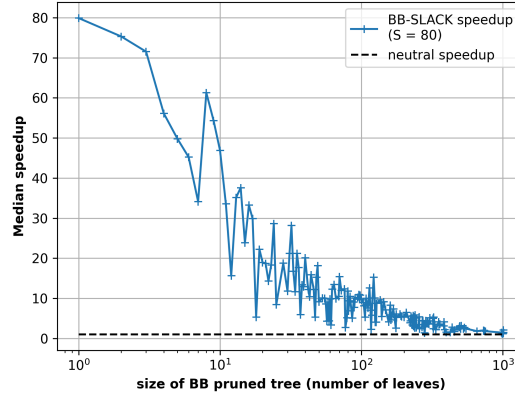
According to our tests, using the parameter S , along with the parameter N , is the most efficient optimization of the pruning process we experienced. We present nevertheless several other properties to optimize BBLs and in particular *is.L3.bound*. The simple Proposition J.1 was already known in [Martello & Toth, 1990], ideas behind Propositions J.2 and J.3 also but were limited to B_α and decreasingly sorted p_j 's. Here we propose a

Figure 24: BB-SLACK vs BP-SLACK speedups obtained w.r.t. the tree size for $S = 0$ (on the top) and $S = 80$ (on the bottom)

BB-SLACK ($S=0$) Speedup vs BP-SLACK (parallelized with 2^3 procs)



BB-SLACK ($S=80$) Speedup vs BP-SLACK (parallelized with 2^3 procs)



complement of those properties to arbitrary order of p_j 's. We also derive a relaxed bound from B_β , named B'_β , that is eligible to the optimization brought by Propositions J.2 and J.3 contrary to B_β . We also explain why despite this optimization, this derived bound should not be used in the pruning process.

Finally, we introduce Proposition J.4 for which no significant speedup has been observed experimentally.

Proposition J.1. *For a given L , $J_1 \cup J_2$ and the sum $|J_1| + |J_2|$ are the same for all $\bar{p}$.*

Proof. $J_1 \cup J_2 = \{p_j | p_j > \frac{L}{2}\}$, this set does not depend on $\bar{p}$. As $\bar{p}$ increases J_2 shrinks and J_1 widens but their union does not change. Besides, by definition $J_1 \cap J_2 = \emptyset$. $\square$

The previous property allows to precompute the sum $|J_1| + |J_2|$ once and for all $\bar{p}$ without having to build J_1 for each $\bar{p}$ in the loop that computes B_α and B_β . Still J_2 has to be built separately on each iteration as it depends on $\bar{p}$. The next proposition mitigates this need.

Proposition J.2. *For a given L and any pair $(\bar{p}_{j_1}, \bar{p}_{j_2})$ such that $0 < \bar{p}_{j_1} \leq \bar{p}_{j_2} \leq \frac{L}{2}$, we have $J_2(L, \bar{p}_{j_2}) \subseteq J_2(L, \bar{p}_{j_1})$. Likewise, $J_3(L, \bar{p}_{j_2}) \subseteq J_3(L, \bar{p}_{j_1})$.*

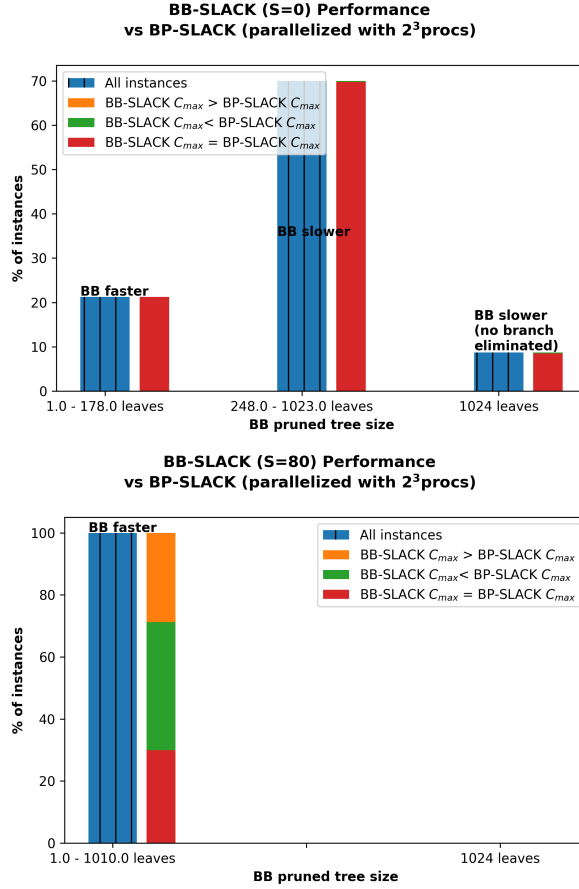
Proof. It is obvious by definition of J_2 and J_3 . $\square$

Corollary J.3. *Let $\bar{p}_{sup}$ be the greatest $\bar{p} \leq \frac{L}{2}$, it comes that $\forall \bar{p} \leq \frac{L}{2}, J_2(L, \bar{p}_{sup}) \subseteq J_2(L, \bar{p})$ and $J_3(L, \bar{p}_{sup}) \subseteq J_3(L, \bar{p})$.*

Proof. For any $\bar{p}$, take $\bar{p}_{j_2} = \bar{p}_{sup}$ and $\bar{p}_{j_1} = \bar{p}$ then refer to Proposition J.2. $\square$

If the $\bar{p}$'s are considered in non-increasing order in the loop of *is L3.bound* then B_α can be calculated iteratively from a greater $\bar{p}$ to a smaller $\bar{p}$ because according to Proposition J.2:

Figure 25: Performance of BB-SLACK compared to BP-SLACK, in case $S = 0$ (on the top) and $S = 80$ (on the bottom)



$$\sum_{j \in J_2(L, \bar{p}_{j_1})} p_j = \sum_{j \in J_2(L, \bar{p}_{j_2})} p_j + \sum_{j \in (J_2(L, \bar{p}_{j_1}) \setminus J_2(L, \bar{p}_{j_2}))} p_j$$

$$\sum_{j \in J_3(L, \bar{p}_{j_1})} p_j = \sum_{j \in J_3(L, \bar{p}_{j_2})} p_j + \sum_{j \in (J_3(L, \bar{p}_{j_1}) \setminus J_3(L, \bar{p}_{j_2}))} p_j$$

Hence, along iterations on $\bar{p}$'s, keeping track of sums corresponding to $J_2(L, \bar{p})$, $J_3(L, \bar{p})$ it is possible to compute $B_\alpha(L, \bar{p})$ by incrementing the previous sums instead of computing the whole sums for each $\bar{p}$. As argued in [Martello & Toth, 1990] it allows to move from a $O(n^2)$ to a $O(n)$ complexity for the calculation of B_α for all $\bar{p}$. The same article provides further improvement which consists of verifying that the current $B_\alpha(L, \bar{p})$ is not already the largest. It might spare the calculation of other B_α bounds that depend on remaining $\bar{p}$'s. Let $\bar{p}_{\inf}$, the smallest possible $\bar{p}$ and corresponding J_3 sum $S_{3, \sup}(J_3(L, \bar{p}_{\inf})) = \sum_{j|p_j \leq L/2} p_j$ which is the greatest one for all J_3 subsets defined on the same L . Then if the following inequality is true for the current $B_\alpha(L, \bar{p})$ and corresponding J_1, J_2 , it comes that this is the greatest bound for one L :

$$|J_1| + |J_2| + \left\lceil \frac{S_{3, \sup} - (L|J_2| - \sum_{j \in J_2} p_j)}{L} \right\rceil \leq B_\alpha(L, \bar{p}) \quad (8)$$

Indeed, if the $\bar{p}$'s are considered in non-increasing order, the quantity $(L|J_2| - \sum_{j \in J_2} p_j)$ can only increase, as well as the $\sum_{j \in J_3} p_j$ sums, which are all bounded superiorly by $S_{3, \sup}$. Conclude that if (8) is true for the current $B_\alpha(L, \bar{p})$, then this is the greatest bound possible. Notice that the S_3 term of the inequality can be precomputed and other terms of the left member are already computed within B_α formula.

On the other hand, if the $\bar{p}$'s are sorted in arbitrary order, this is still possible to enhance computation time. Corollary J.3 implies that in $B_\alpha(L, \bar{p})$ the sum terms that correspond to the subsets $J_2(L, \bar{p}_{\sup})$, $J_3(L, \bar{p}_{\sup})$ can be precomputed once and for all $\bar{p}$ and incremented by the sums corresponding to the specific parts of $J_2(L, \bar{p})$, $J_3(L, \bar{p})$ to compute the whole formula for one $\bar{p}$:

$$\begin{aligned}\sum_{j \in J_2(L, \bar{p})} p_j &= \sum_{j \in J_2(L, \bar{p}_{\text{sup}})} p_j + \sum_{j \in (J_2(L, \bar{p}) \setminus J_2(L, \bar{p}_{\text{sup}}))} p_j \\ \sum_{j \in J_3(L, \bar{p})} p_j &= \sum_{j \in J_3(L, \bar{p}_{\text{sup}})} p_j + \sum_{j \in (J_3(L, \bar{p}) \setminus J_3(L, \bar{p}_{\text{sup}}))} p_j\end{aligned}$$

Let us now consider if previous results about B_α also apply to B_β . Indeed, B_β was not present in [Martello & Toth, 1990] because it was introduced later in [Dell'Amico & Martello, 1995]. It is easily understandable that the term $\left\lfloor \frac{L-p_j}{\bar{p}} \right\rfloor$ makes inevitable the recomputation of the whole sum for each $\bar{p}$. Indeed the $\bar{p}$ denominator cannot be factored because of the enclosing floor function. To make the precomputation and sum update possible it is necessary to introduce a relaxed bound B'_β , very similar to B_β with the only difference that the floor function is applied on the whole sum instead of each term.

$$B'_\beta(L, \bar{p}) = |J_1| + |J_2| + \max(0, \left\lceil \frac{|J_3| - \left\lfloor \frac{1}{\bar{p}} \sum_{j \in J_2} L - p_j \right\rfloor}{\left\lfloor \frac{L}{\bar{p}} \right\rfloor} \right\rceil) \quad (9)$$

Now the sum of $(L - p_j)$ can be updated along the growth of J_2 for the different $\bar{p}$'s considered as explained previously for B_α . It is worth mentioning that the relaxation of B_β into B'_β is, more precisely, a preemptive relaxation of the bound, in which the J_3 jobs can be split and then assigned by bits to the J_2 bins. Note that B_α is also a preemptive bound. Naturally, this relaxation implies that $B'_\beta \leq B_\beta$ because $\left\lfloor \frac{1}{\bar{p}} \sum_{j \in J_2} L - p_j \right\rfloor \geq \sum_{j \in J_2} \left\lfloor \frac{L - p_j}{\bar{p}} \right\rfloor$. Therefore even if the relaxed bound can be computed more rapidly it might lead to a less efficient pruning of the BBLS tree and finally a slower method. This risk has been verified experimentally.

As a complement, the following proposition allows to spare some calculations for B_β bound by reasoning on J_2 and J_3 cardinals. Of note, no significant speedup has been observed experimentally using all instances in Table 3. Nevertheless, we noticed, in those tests, that about 43% of `is_L3_bound` calls were concerned by this proposition.

Proposition J.4. *If $|J_3| \leq |J_2|$ then $B_\beta(L, \bar{p}) = |J_1| + |J_2|$.*

Proof. With $j \in J_2$ we have that $p_j \leq L - \bar{p}$ then $L - p_j \geq \bar{p}$. Hence $\sum_{j \in J_2} \left\lfloor \frac{L - p_j}{\bar{p}} \right\rfloor \geq |J_2|$ so if $|J_3| \leq |J_2|$ we obtain that $\max(0, \left\lceil \frac{|J_3| - \sum_{j \in J_2} \left\lfloor \frac{L - p_j}{\bar{p}} \right\rfloor}{\left\lfloor \frac{L}{\bar{p}} \right\rfloor} \right\rceil) = 0$ □