



Dynamic memory-aware task-tree scheduling

Guillaume Aupy, Clément Brasseur, Loris Marchal

► To cite this version:

Guillaume Aupy, Clément Brasseur, Loris Marchal. Dynamic memory-aware task-tree scheduling. [Research Report] RR-8966, INRIA Grenoble - Rhone-Alpes. 2016. hal-01390107

HAL Id: hal-01390107

<https://inria.hal.science/hal-01390107v1>

Submitted on 31 Oct 2016

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

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



Dynamic memory-aware task-tree scheduling

Guillaume Aupy, Clément Brasseur, Loris Marchal

**RESEARCH
REPORT**

N° 8966

October 2016

Project-Team ROMA

ISRN INRIA/RR--8966--FR+ENG

ISSN 0249-6399



Dynamic memory-aware task-tree scheduling

Guillaume Aupy*, Clément Brasseur†, Loris Marchal†

Project-Team ROMA

Research Report n° 8966 — October 2016 — 38 pages

Abstract: Factorizing sparse matrices using direct multifrontal methods generates directed tree-shaped task graphs, where edges represent data dependency between tasks. This paper revisits the execution of tree-shaped task graphs using multiple processors that share a bounded memory. A task can only be executed if all its input and output data can fit into the memory. The key difficulty is to manage the order of the task executions so that we can achieve high parallelism while staying below the memory bound. In particular, because input data of unprocessed tasks must be kept in memory, a bad scheduling strategy might compromise the termination of the algorithm. In the single processor case, solutions that are guaranteed to be below a memory bound are known. The multi-processor case (when one tries to minimize the total completion time) has been shown to be NP-complete. We present in this paper a novel heuristic solution that has a low complexity and is guaranteed to complete the tree within a given memory bound. We compare our algorithm to state of the art strategies, and observe that on both actual execution trees and synthetic trees, we always perform better than these solutions, with average speedups between 1.25 and 1.45 on actual assembly trees. Moreover, we show that the overhead of our algorithm is negligible even on deep trees (10^5), and would allow its runtime execution.

Key-words: scheduling; memory; trees.

* University of Vanderbilt, Nashville TN, USA

† CNRS, LIP, École Normale Supérieure de Lyon, INRIA, France

Dynamic memory-aware task-tree scheduling

Résumé : La factorisation de matrices creuses à l'aide de méthodes directes multifrontales génère des arbres enracinés de tâches, où les arrêtes représentent des dépendences de données. Ce papier revisite l'exécution de ces arbres de tâches sur multi-processeurs avec mémoire bornée. Une tâche ne peut être exécutée que si ses entrées et sorties tiennent en mémoire. La difficulté principale tient en la gestion de l'ordre d'exécution des tâches pour atteindre un grand parallélisme tout en respectant la contrainte mémoire. En particulier, une fois calculées, les entrées d'une tâche non exécutée doivent être gardées en mémoire. De mauvais choix algorithmiques peuvent conduire à une sur-consommation mémoire et ainsi compromettre la terminaison du traitement de l'arbre.

Des solutions qui garantissent des bornes sur la consommation mémoire sont connues dans le cas mono-processeur. Le cas multi-processeur où l'on essaie de minimiser le temps d'exécution est NP-complet. Nous présentons dans ce papier une nouvelle solution heuristique à faible complexité et qui garantit la terminaison pour des bornes mémoires dans le cas multi-processeur. Nous nous comparons aux heuristiques les plus récentes et montrons la dominance de notre solution.

Mots-clés : ordonnancement; mémoire; arbres.

1 Introduction

Parallel workloads are often modeled as task graphs, where nodes represent tasks and edges represent the dependencies between tasks. There is an abundant literature on task graph scheduling when the objective is to minimize the total completion time, or makespan. However, with the increase of the size of the data to be processed, the memory footprint of the application can have a dramatic impact on the algorithm execution time, and thus needs to be optimized. This is best exemplified with an application which, depending on the way it is scheduled, will either fit in the memory, or will require the use of swap mechanisms or out-of-core. There are very few existing studies that take into account the memory footprint when scheduling task graphs, and even fewer of them targeting parallel systems.

In the present paper, we consider the parallel scheduling of rooted in-trees. The vertices of the trees represent computational tasks, and its edges represent the dependencies between these tasks, which are in the form of input and output data: each task requests for its processing all the data produced by its children tasks to be available in memory, and outputs a new data for its parent. We want to process the resulting task tree on a parallel system made of p computing units, also named processors, sharing a global memory of limited size M . At any time, the size of all the data currently produced but not yet consumed cannot exceed M . Our objective is to minimize the makespan, that is, the total time needed to process the whole task tree, under the memory constraint.

The motivation for this work comes from numerical linear algebra, and especially the factorization of sparse matrices using direct multifrontal methods [6]. During the factorization, the computations are organized as a tree workflow called the elimination tree, and the huge size of the data involved makes it absolutely necessary to reduce the memory requirement of the factorization. Note that we consider here that no numerical pivoting is performed during the factorization, and thus that the structure of the tree, as well as the size of the data are known before the computation really happens.

In this paper, we mainly build on two previous results. On the theoretical side, we have previously studied the complexity of the bi-criteria problem which considers both makespan minimization and peak memory minimization [7], and we have proposed a few heuristic strategies to schedule task trees under hard memory constraints. However, these strategies requires strong reduction properties on the tree. An arbitrary tree can be turned into a reduction tree, but this increases its memory footprint, which limits the performance of the scheduler under memory constraint. On the practical side, Agullo et al. [1] uses a simple activation strategy to ensure the correct termination of a multifrontal QR factorization, whose task graph is an in-tree. Both approaches have drawbacks: the first one artificially increases the peak memory of the tree, and the second one overestimates the memory booked to process a subtree. Our objective is to take inspiration from both to design a better scheduling algorithm.

Note that we are looking for a *dynamic* scheduling algorithm, that is, a strategy that dynamically reacts to task terminations to activate and schedule

new nodes. We suppose that only the tree structure and the data sizes are known before the execution, not the task processing times, so that one cannot rely on them to build a perfect static schedule. Finally, the scheduling complexity should be kept as low as possible, since scheduling decision need to be taken during the computation without delaying the task executions.

Our contributions are as follows:

- We provide a novel heuristic along with a proof of its termination for memory bounds.
- We provide data-structure optimizations to improve its computational complexity.
- We provide a thorough experimental study, both on actual and synthetic trees to show its dominance over state of the art algorithms.
- We propose a new makespan lower bound for memory-constrained parallel platforms.

The rest of the paper is organized as follows. We first present the problem, its notation and formalize our objective in Section 2. We then review related work and the two existing approaches listed above in Section 3. Next, we present our new scheduling algorithms, as well as the proof of its correctness in Section 4. Then, we propose a memory-aware makespan lower bound in Section 6. Finally, we present a set of comprehensive simulations to assess the benefit of the new algorithm 7, before presenting concluding remarks in Section 8

2 Model and objectives

2.1 Application model

Let T be a rooted in-tree (dependencies point toward the root) composed of n nodes, the tasks, denoted by their index $1, \dots, n$. A node i is characterized by its input data (one per child), its execution data (of size n_i), and its output data (of size f_i). When processing node i , all input, execution and output data must be allocated in memory. At the termination of node i , input and execution data are deallocated, and only the output data stays in memory. We denote by $Children(i)$ the set of children of node i , which is empty if i is a leaf. The memory needed for the processing of node i , illustrated on Figure 1, is given by:

$$MemNeeded_i = \left(\sum_{j \in Children(i)} f_j \right) + n_i + f_i. \quad (1)$$

2.2 Platform model

We consider a shared-memory parallel platform, composed of p homogeneous processors onto which each task can be computed. Those processors share a limited memory of size M .

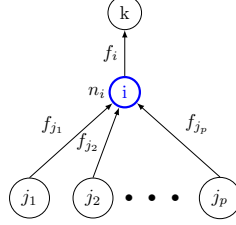


Figure 1: Input, execution and output data of a node i with children $j_1, \dots, j_p$ and parent k .

2.3 Objectives

Our objective here is to minimize the makespan, that is the total execution time, while keeping the size of the data stored in memory below the bound M . This problem is a variant of the bi-criteria problem which aims at minimizing both the makespan and the peak memory. Note that those two objective are antagonist: the best way to minimize the makespan is to parallelize as much as possible without regard to the memory used, while the best way to minimize the peak memory used is to execute the whole schedule on a single processor which would give the worse makespan. In a previous study [7], we have proven that this bi-criteria problem was NP-complete and inapproximable within constants factors of both the optimal memory and the optimal makespan. Our variant clearly inherits this complexity and in this study, we are mainly looking for heuristic solutions.

Note that the algorithms under considerations are natural candidates to replace the simple activation scheme of [1]. Thus, their runtime complexity is critical, as we want to take scheduling decisions very fast. While a notable, say quadratic, complexity is acceptable in the initial preprocessing phase, we are looking for $O(1)$ complexity at each task termination, or eventually logarithmic.

3 Related work and existing algorithms

Memory and storage have always been a limited parameter for large computations, as outlined by the pioneering work of Sethi and Ullman [13] on register allocation for task graphs. In the realm of sparse direct solvers, the problem of scheduling a tree so as to minimize peak memory has first been investigated by Liu [11] in the sequential case: he proposed an algorithm to find a peak-memory minimizing traversal of a task tree when the traversal is required to correspond to a postorder traversal of the tree. A follow-up study [10] presents an optimal algorithm to solve the general problem, without the postorder constraint on the traversal. Postorder traversals are known to be arbitrarily worse than optimal traversals for memory minimization [8]. However, they are very natural and straightforward solutions to this problem, as they allow to fully

process one subtree before starting a new one. Therefore, they are widely used in sparse matrix software like MUMPS [3, 4], and achieve good performance on actual elimination trees [8].

The problem of scheduling a task graph under memory or storage constraints also appears in the processing of scientific workflows whose tasks require large I/O files. Such workflows arise in many scientific fields, such as image processing, genomics, and geophysical simulations. The problem of task graphs handling large data has been identified by Ramakrishnan et al. [12] who propose some simple heuristics. In the context of quantum chemistry computations, Lam et al. [9] also consider task trees with data of large size.

Note that peak memory minimization is still a crucial question for direct solvers, as highlighted by Agullo et al. [2], who study the effect of processor mapping on memory consumption for multifrontal methods.

We now review the two scheduling strategies from the literature that target our problem.

3.1 Simple activation heuristic

Agullo et al.[1] use a simple activation strategy to ensure that a parallel traversal of a task tree will process the whole tree without running out of memory. The first step is to compute a postorder traversal, such as the memory-minimizing traversal of [11]. This postorder traversal, denoted by AO , will serve as an order to activate tasks. This solution requires that the available memory M is not smaller than the peak memory M_{AO} of the activation order. The strategy is summarized in Algorithm 1. The activation of a task i consists in allocating all the memory needed for this tasks. Then, only tasks that are both activated and whose dependency constraints are satisfied (i.e., all predecessors in the tree are already processed) are available for execution. Another scheduling heuristic may be used to choose which tasks the among available ones are executed: we denote by EO the order giving the priority of the tasks for execution. Note that *available nodes* are nodes whose children have been completed.

This simple procedure is efficient to schedule task trees without exceeding the available memory. However, it may book too much memory, and thus limit the available parallelism in the tree. Consider for example a chain of tasks $T_1 \rightarrow T_2 \rightarrow T_3$. Algorithm 1 will first book $n_1 + f_1$ for task T_1 , then $n_2 + f_2$ for T_2 and finally $n_3 + f_3$ for T_3 (assuming all this memory is available). However, no two tasks of this chain can be scheduled simultaneously because of their precedence order. Thus, it is not necessary to book n_1 , n_2 and n_3 at the same time, nor is it necessary to book memory for f_1 and f_3 : the memory used for T_1 can be later reused for the processing of T_2 and T_3 . By booking memory in a very conservative way, this heuristic may prevent nodes from other branches to be available for computation, and thus delay the processing of these nodes.

Algorithm 1: ACTIVATION(T, p, AO, EO, M)

```

1  $M_{Booked} \leftarrow 0$ 
2  $ACT \leftarrow \emptyset$ 
3 while the whole tree is not processed do
4   Wait for an event (task finishes or  $t = 0$ )
   // Free the memory booked by  $j$ 
5   foreach just finished node  $j$  do  $M_{Booked} \leftarrow M_{Booked} - n_j - \sum_{j \in Children(i)} f_j$ 
6    $continueActivation \leftarrow \text{true}$ 
7   while  $continueActivation$  do
8     // Activate the first node  $i$  of  $AO$  if possible
9      $i \leftarrow pop(AO)$ 
10    if  $M_{Booked} + n_i + f_i \leq M$  then
11       $M_{Booked} \leftarrow M_{Booked} + n_i + f_i$ 
12       $push(i, ACT)$ 
13    else
14       $push(i, AO)$ 
15       $continueActivation \leftarrow \text{false}$ 
16    // Process available nodes in  $ACT$  following priority order  $EO$ 
17    while there is an available processor  $P_u$  and there is an available node in  $ACT$  do
18      Let  $i$  be the available node in  $ACT$  with maximal priority in  $EO$ :
19       $Remove(i, ACT)$ 
20      Make  $P_u$  process  $i$ 

```

3.2 Booking strategy for reduction trees

In a previous publication [7], we have proposed a novel activation policy based on a refined memory booking strategy. However, our strategy is limited to special trees, also called reduction trees, who exhibit the following two properties:

- There is no execution data, i.e., $n_i = 0$ for each node i ;
- The size of the output of each node is smaller than the size of its inputs, that is, $f_i \leq \sum_{j \in Children(i)} f_j$.

Using these two properties, we are able to prove that the memory has been successfully booked for all the leaves of a subtree, all nodes in this subtree can be processed without additional memory. Moreover, we know how to compute the amount of booked memory that each completing node has to transmit to its parent.

Contrarily to the previous algorithm, this complex strategy allows to correctly predict the amount of memory that needs to be booked for a given subtree. However, it only applies to special trees, namely reduction trees. General trees may be transform into reduction trees by adding fictitious edges. However this increases the overall peak memory needed for any traversal, which limits its interest —we have indeed noticed that it does not give better performance than

Algorithm 1 on general trees—. Furthermore, in some cases it makes it a key limitation for general trees with limited memory as it does not always allow for the completion of those trees.

4 A dynamic fast algorithm

In this section, we propose a novel algorithm, named MEMBOOKING to schedule trees under limited memory, that overcomes the limitations of the previous two strategies. Similarly as in Algorithm 1, we rely on the activation of nodes, following an activation AO which is guaranteed to complete the whole tree within the prescribed memory in the case of a linear execution. However, activating a node does not correspond here to booking the exact memory $n_i + f_i$ needed for this node: some of this memory will be transferred by some of its descendant in the tree, and if needed, we only book what is absolutely needed. The core idea of the algorithm is the following: when a node completes its execution, we want (i) to reuse the memory that is freed on one of its ancestors and (ii) perform these transfers of booked memory in an As Late As Possible (ALAP) fashion. More precisely, the memory freed by a completed node j will only be transferred to one of its ancestors i if (a) all the descendants of i have enough memory to be executed (that is, they are activated), and (b) if this memory is necessary and cannot be obtained from another descendent of i that will complete its execution later. Finally, an execution order EO states which of the activated and available nodes should be processed whenever a processor is available.

In order to keep track of all nodes, we use *five states* to describe them (a node can only be in one state), which we present in reverse order of their use for a given node:

1. Finished (FN): This corresponds to nodes which have completed their execution.
2. Running (RUN): This corresponds to nodes being executed.
3. Activated (ACT): This corresponds to nodes for which we have booked enough memory (some of this memory might be booked by some descendant in the subtree).
4. Candidate for activation ($CAND$): This corresponds to nodes which are the next to be activated, that is, all their descendant have been activated but they are not activated yet. This is the initial state for all leaves.
5. Unprocessed (UN): This corresponds to nodes which have not yet been considered; it is the initial state for all interior nodes.

Because a node can only have be in one state at a given time, we write $j \in UN$ (resp. $CAND$, ACT , RUN , FN) if node j is in the corresponding state.

We now present the MEMBOOKING algorithm (Algorithm 2), as well as its proof of correctness. Some optimizations and data-structures used to reduce the

time complexity will be presented later (the full algorithm with optimizations is available in Appendix B).

Algorithm 2: MEMBOOKING(T, p, AO, EO, M)

```

// initialization of the memory
1 foreach node  $i$  do
2   |  $Booked[i] \leftarrow 0$ 
3   |  $BookedBySubtree[i] \leftarrow -1$ 
4  $M_{Booked} \leftarrow 0$ 
5  $UN \leftarrow T \setminus Leaves(T)$ 
6  $CAND \leftarrow Leaves(T)$ 
7  $ACT \leftarrow \emptyset$ 

// main loop
8 while the whole tree is not processed do
9   | Wait for an event (task finishes or  $t = 0$ )
10  | foreach just finished node  $j$  do
11    |  $DISPATCHMEMORY(j)$ 
12  |  $UPDATECAND-ACT(CAND, ACT)$ 
13  | while there is an available processor  $P_u$  and  $ACT \neq \emptyset$  do
14    | Let  $i$  be the available node in  $ACT$  with maximal priority in  $EO$ :
15    |  $Remove(i, ACT)$ 
16    | Make  $P_u$  process  $i$ 

```

Available nodes are nodes whose children have been completed.

At the beginning of the schedule, or each time a task completes, the MEMBOOKING algorithm performs these three consecutive operations:

1. Memory re-allocation: $DISPATCHMEMORY$ (Algorithm 3) reallocates the memory used by a node that just finished its execution. We present this algorithm in Section 4.1.
2. Node activation: $UPDATECAND-ACT$ (Algorithm 4) allocates the available memory following the activation order AO . We present this algorithm in Section 4.2.
3. Node scheduling: the schedule is done following the execution order EO amongst the nodes that are both activated and ready to be executed.

Note that while AO is a topological order, we do not have any constraint on EO .

To be able to keep track of the memory allocated to each node, we use two arrays of data that are updated during the computation, namely:

- $Booked[1..n]$, which contains the memory that is currently booked in order to process nodes 1 to n . We further call $M_{Booked} = \sum_i Booked[i]$;
- $BookedBySubtree[1..n]$, which sums the memory that is currently booked by the subtree rooted in $i \in \{1, \dots, n\}$.

4.1 Memory re-allocation

When a node finishes its computation, the memory that was used during its computation can be allocated to other nodes. Our memory allocation works in two steps:

1. First we free the memory that was used by the node that has just finished its execution. Note that we cannot free all the memory: if the node that finished is not the root of the tree, then its output needs to be saved. In that case we allocate this memory to its parent.
2. Then we allocate the memory freed to its ancestors in *ACT* following an As Late As Possible strategy (meaning that if there is already enough memory booked in the unfinished part of the subtree, we do not allocate it to the root of the subtree but keep it for later use). We thus compute the contribution $C_{i,j}$ of a terminated node j to its parent i as the difference between what is needed by node i and what can be provided later by its subtree.

Algorithm 3: DISPATCHMEMORY(j)

```

/* First we free the memory used by j */
1  $B = \text{Booked}[j]$ 
2  $\begin{cases} \text{Booked}[j] & \leftarrow 0 \\ M_{\text{Booked}} & \leftarrow M_{\text{Booked}} - B \\ \text{BookedBySubtree}[j] & \leftarrow \text{BookedBySubtree}[j] - B \end{cases}$ 
3  $i \leftarrow \text{parent}(j)$ 
4 if  $i \neq \text{NULL}$  then
5    $\begin{cases} \text{Booked}[i] & \leftarrow \text{Booked}[i] + f_j \\ M_{\text{Booked}} & \leftarrow M_{\text{Booked}} + f_j \\ \text{BookedBySubtree}[i] & \leftarrow \text{BookedBySubtree}[i] + f_j \end{cases}$ 
6    $B = B - f_j$ 
/* Then we dispatch the memory used by j between its ancestors which are in ACT, if it is necessary */
7 while  $i \neq \text{NULL}$  and  $i \in \text{ACT} \cup \text{RUN}$  and  $B \neq 0$  do
8    $C_{j,i} = \max(0, \text{MemNeeded}_i - (\text{BookedBySubtree}[i] - B))$ 
9    $\begin{cases} \text{Booked}[i] & \leftarrow \text{Booked}[i] + C_{j,i} \\ M_{\text{Booked}} & \leftarrow M_{\text{Booked}} + C_{j,i} \\ \text{BookedBySubtree}[i] & \leftarrow \text{BookedBySubtree}[i] - (B - C_{j,i}) \end{cases}$ 
10   $B = B - C_{j,i}$ 
11   $i \leftarrow \text{parent}(i)$ 

```

4.2 Node activation

Our second algorithm, *UPDATECAND-ACT*, updates both *ACT* and *CAND*. The key point of this sub-algorithm is that it is conceived such that nodes are activated following the *AO* order. We formally show this result in Section 5 (Lemma 1).

Algorithm 4: *UPDATECAND-ACT*(*CAND*, *ACT*)

```

1  WaitForMoreMem  $\leftarrow$  false
2  while  $\neg$ (WaitForMoreMem) and CAND  $\neq \emptyset$  do
3    Let i be the node of CAND with maximal priority in AO
4    MissingMemi =
       $\max \left( 0, \text{MemNeeded}_i - \left( \text{Booked}[i] + \sum_{j \in \text{Children}(i)} \text{BookedBySubtree}[j] \right) \right)$ 
5    if MBooked + MissingMemi  $\leq M$  then
      {
6        Booked[i]  $\leftarrow$  Booked[i] + MissingMemi
        MBooked  $\leftarrow$  MBooked + MissingMemi
        BookedBySubtree[i]  $\leftarrow$  Booked[i] +  $\sum_{j \in \text{Children}(i)} \text{BookedBySubtree}[j]$ 
7        remove(i, CAND); insert(i, ACT)
8        if  $\forall j \in \text{Children}(\text{parent}(i)), j \notin \text{UN} \cup \text{CAND}$  then
9          remove(parent(i), UN); insert(parent(i), CAND)
10       else
11         WaitForMoreMem  $\leftarrow$  true

```

5 Proof of correctness

In this section, we want to show the following result:

Theorem 1. *If T can be executed with a memory bound of M using the sequential schedule AO , then $\text{MEMBOOKING}(T, p, AO, EO, M)$ processes T entirely.*

To prove this theorem, we need to verify that the following conditions are respected:

1. The memory used never exceed the memory limit M ;
2. Each running task has enough memory to run;
3. No data is lost, that is, a result that was computed will not be overwritten until it has been used;
4. All tasks are executed.

A necessary condition to prove items 1, 2 and 3 using the *Booked*[1..*n*] introduced earlier is to prove the following results:

1. At all time, $\sum_i \text{Booked}[i] \leq M$;
2. If $i \in \text{RUN}$, then $\text{Booked}[i] = \text{MemNeeded}_i$;
3. For all $i \notin \text{FN}$, we cannot decrease the value of $\text{Booked}[i]$; finally, when i is moved to FN , its output f_i should be moved from $\text{Booked}[i]$ to its parent's memory usage: $\text{Booked}[\text{parent}(i)]$.

In the proofs, we decompose time according to the different events (Algorithm 2, line 9): $t_0 = 0$ and when each task finishes. We know that to execute the tree T of n tasks, there are at most $n + 1$ events (each task can finish only once) that we denote $t_0 = 0 \leq t_1 \leq \dots \leq t_n$ (in case a node is never executed, we may write $t_i = +\infty$).

As our algorithm updates its variables describing the tasks' state and the amount of booked memory for each node after each event, we consider time-intervals $]t_i, t_{i+1}]$: this corresponds to the status of the algorithm after the modification of event t_i .

We start by proving that nodes are activated in the right order.

Lemma 1. *The nodes are moved from CAND to ACT following the AO order.*

Proof. First, we notice that nodes are inserted in *ACT* only in Algorithm 4, on line 7. Then one can verify that at all time, the element extracted from *CAND* is the first one with respect to the *AO* ordering, ensuring that the elements contained in *CAND* are extracted following *AO*.

To show that elements are inserted in *ACT* following *AO*, we then only need to prove that at all time, the next element according to *AO* is available in *CAND*. That is, for all $i \in \{1, \dots, n\}$, if for all $j < i$, $AO(j)$ has been extracted from *CAND* to *ACT*, then either $AO(i)$ has been extracted from *CAND* to *ACT*, or $AO(i) \in \text{CAND}$. We prove this by induction.

Because *AO* is a topological order, then $AO(1)$ is necessarily a leaf. Furthermore, *CAND* initially contains all leaves, hence $AO(1) \in \text{CAND}$.

Let $i > 1$. Assume that for all $j < i$, $AO(j)$ has been extracted from *CAND* to *ACT*. If $AO(i)$ has been extracted from *CAND* to *ACT* then the property is true. Otherwise, either $AO(i)$ is a leaf or an interior node. If it is a leaf, then $AO(i) \in \text{CAND}$ because leaves were inserted initially and they can only be removed from *CAND* when they are extracted to *ACT*. If $AO(i)$ is an interior node, let $AO(i_1), \dots, AO(i_m)$ be its children, because *AO* is a topological order, then for $j \leq m$, $i_j < i$. In particular, assume w.l.o.g that $i_1 < \dots < i_m$, then when $AO(i_m)$ was extracted from *CAND* to *ACT*, then $AO(i)$ was added to *CAND*, hence showing the result. $\square$

Then, we prove that until a node is activated, the only memory booked for this node are for the results of its completed children.

Lemma 2. *At any time t , if $i \in \text{UN} \cup \text{CAND}$, then*

$$\text{Booked}[i] = \sum_{j \in \text{Children}(i) \cap \text{FN}} f_j$$

Proof. If $i \in UN \cup CAND$, this means that until that time, i has always been either in UN or in $CAND$, and hence the only time when $\text{Booked}[i]$ could have been modified is by a call of `DISPATCHMEMORY` (Algorithm 3), more specifically on line 5. This happens only when a child of i completes, and hence is moved to FN , which proves the result. $\square$

Lemma 3. *At any time t , if $i \in ACT \cup RUN$, then*

- (1) $\sum_{j \in \text{Children}(i) \cap FN} f_j \leq \text{Booked}[i];$
- (2) $\text{MemNeeded}_i \leq \text{BookedBySubtree}[i]$
- (3) $\text{BookedBySubtree}[i] = \text{Booked}[i] + \sum_{j \in \text{Children}(i) \cap (ACT \cup RUN \cup FN)} \text{BookedBySubtree}[j].$

Proof. First one can notice that for $i \in ACT \cup RUN$, $\sum_{j \in \text{Children}(i) \cap FN} f_j \leq \text{Booked}[i]$: indeed, every time a child j of i finished its execution, `DISPATCHMEMORY` (Algorithm 3) added f_j to $\text{Booked}[i]$. Furthermore, until i is moved to FN , the algorithm never subtracts anything from $\text{Booked}[i]$.

Similarly, we have at all time $\text{MemNeeded}_i \leq \text{BookedBySubtree}[i]$: it is the case when i is moved from $CAND$ to ACT , and all further transformations of $\text{BookedBySubtree}[i]$ preserve this result.

To show the third result, we proceed by induction on time. For simplicity, let us denote by

$$\widetilde{\text{BBS}}(i) = \text{Booked}[i] + \sum_{j \in \text{Children}(i) \cap (ACT \cup RUN \cup FN)} \text{BookedBySubtree}[j]$$

Clearly the result is true at time $t_0 = 0$ since $ACT \cup RUN = \emptyset$. Assume the induction hypothesis: for all time $t \leq t_{i_0}$, if $i \in ACT \cup RUN$, $\text{BookedBySubtree}[i] = \widetilde{\text{BBS}}(i)$. We now prove that the result holds during the interval $]t_{i_0}, t_{i_0+1}]$ using a second induction on the tree structure, starting from leaves.

We first assume that $i \in ACT \cup RUN$ is a leaf. When i was moved from $CAND$ to ACT (using `UPDATECAND-ACT`), then the algorithm made sure that the property was respected, then we notice that neither $\text{Booked}[i]$ nor $\text{BookedBySubtree}[i]$ is modified for a leaf until that leaf is moved to FN , so the result still holds.

We now consider an interior node $i_c \in ACT \cup RUN$ such that the result holds for all its descendants until t_{i_0+1} and the result holds for that node until t_{i_0} (first induction hypothesis).

If i_c was moved to ACT at time t_{i_0} , then one can verify by looking at algorithm `UPDATECAND-ACT` that the result holds. Furthermore, if no descendant of i_c finished at t_{i_0} , then `DISPATCHMEMORY` has not modified the status of i_c , hence by induction hypothesis the result holds. Let us assume now that a descendant of i_c has finished its execution at t_{i_0} , let us call j_c the children of i_c on the path to that descendant.

If $j_c \in FN$ (meaning j_c is the descendant that finished at t_{i_0}), then one can see that B was subtracted from $\widetilde{BBS}(i_c)$. (DISPATCHMEMORY, line 2), then f_{j_c} was added (DISPATCHMEMORY, line 5), then C_{j_c, i_c} was added (DISPATCHMEMORY, line 9). Similarly, f_{j_c} was added to $\text{BookedBySubtree}[i]$ (DISPATCHMEMORY, line 5), then $B - C_{j_c, i_c}$ was subtracted (DISPATCHMEMORY, line 9) which keeps the value identical, showing the result. Finally, one can notice that at the end of the while iteration (line 11), the value B is now what was removed from $\text{BookedBySubtree}[i]$ compared to the beginning of the iteration.

Otherwise, $j_c \in ACT \cup RUN$ In that case, $\widetilde{BBS}(i_c)$ and $\text{BookedBySubtree}[i_c]$ can only be modified by the “while” loop of Algorithm DISPATCHMEMORY. More specifically, only $\text{BookedBySubtree}[j_c]$ and $\text{Booked}[i_c]$ will be modified. At the beginning of the iteration of the while-loop such that the index $i = i_c$, then B is exactly equal to what was subtracted from $\text{BookedBySubtree}[j_c]$. Furthermore, C_{j_c, i_c} is added to $\text{Booked}[i_c]$, and then $B - \text{Booked}[i_c]$ is subtracted from $\text{BookedBySubtree}[i_c]$ which shows the result. Finally, we notice again that at the end of the iteration, B is exactly equal to what was subtracted from $\text{BookedBySubtree}[i_c]$.

This proves the result for i_c , hence proving the wanted result. $\square$

Lemma 4. *At any time t , if $i \in ACT \cup RUN$, then*

$$\text{Booked}[i] \leq \text{MemNeeded}_i - \sum_{j \in \text{Children}(i) \cap (ACT \cup RUN)} f_j.$$

Proof. We prove this result by induction on time. At $t_0 = 0$, $ACT \cup RUN$ is empty, hence showing the result. Let us assume the result true for all node $i \in ACT \cup RUN$ at any time $t \leq t_{i_0}$ and let us that it also holds for $t \in]t_{i_0}, t_{i_0+1}]$.

We first start by the case where $i \in UN \cup CAND$ at t_{i_0} , and $i \in ACT \cup RUN$ in $]t_{i_0}, t_{i_0+1}]$. This means that node i was activated (moved to ACT) by Algorithm UPDATECAND-ACT at time t_{i_0} . Let us consider the missing memory as defined in UPDATECAND-ACT:

MissingMem_i

$$= \max \left(0, \text{MemNeeded}_i - \left(\text{Booked}[i] + \sum_{j \in \text{Children}(i)} \text{BookedBySubtree}[j] \right) \right),$$

Since i was in $CAND$ right before being activated, we had $\text{Booked}[i] = \sum_{j \in \text{Children}(i) \cap FN} f_j$ (Lemma 2).

Moreover, for all $j \in \text{Children}(i)$, $\text{BookedBySubtree}[j]$ does not change later when i is moved to ACT , hence according to Lemma 3 (we use item (3) on i

then item (2) on the set $Children(i) \cap (ACT \cup RUN)$,

$$\begin{aligned}
 \sum_{j \in Children(i)} \text{BookedBySubtree}[j] &= \sum_{j \in Children(i) \cap (ACT \cup RUN \cup FN)} \text{BookedBySubtree}[j] \\
 &\geq \sum_{j \in Children(i) \cap (ACT \cup RUN)} \text{BookedBySubtree}[j] \\
 &\geq \sum_{j \in Children(i) \cap (ACT \cup RUN)} \text{MemNeeded}_j \\
 &\geq \sum_{j \in Children(i) \cap (ACT \cup RUN)} f_j
 \end{aligned}$$

Hence we verify that:

- If $\text{MissingMem}_i = 0$, then when i is moved to ACT ,

$$\begin{aligned}
 \text{Booked}[i] &= \sum_{j \in Children(i) \cap FN} f_j \leq \text{MemNeeded}_i - \sum_{j \in Children(i)} \text{BookedBySubtree}[j] \\
 &\leq \text{MemNeeded}_i - \sum_{j \in Children(i) \cap (ACT \cup RUN)} f_j;
 \end{aligned}$$

- If $\text{MissingMem}_i \neq 0$, then when i is moved to ACT ,

$$\begin{aligned}
 \text{Booked}[i] &= \text{MemNeeded}_i - \sum_{j \in Children(i)} \text{BookedBySubtree}[j] \\
 &\leq \text{MemNeeded}_i - \sum_{j \in Children(i) \cap (ACT \cup RUN)} f_j.
 \end{aligned}$$

Let us now assume that i was already in $ACT \cup RUN$ at t_{i_0} . Let A_{i_0} be the value of $\text{Booked}[i]$ at t_{i_0} , and $B_{i_0} = \text{MemNeeded}_i - \sum_{j \in Children(i) \cap (ACT \cup RUN)} f_j$ at t_{i_0} . We have $A_{i_0} \leq B_{i_0}$ by induction hypothesis.

- If i is not an ancestor of the node that just finished, then during $]t_{i_0}, t_{i_0+1}]$, $\text{Booked}[i] = A_{i_0}$ and $\text{MemNeeded}_i - \sum_{j \in Children(i) \cap (ACT \cup RUN)} f_j = B_{i_0}$, hence the property holds.
- If i is the parent of the node j_c that just finished, then during $]t_{i_0}, t_{i_0+1}]$, $\text{MemNeeded}_i - \sum_{j \in Children(i) \cap (ACT \cup RUN)} f_j = B_{i_0} - f_{j_c}$. Furthermore, DISPATCHMEMORY performs the operation $\text{Booked}[i] \leftarrow A_{i_0} + f_{j_c} + C_{i,j_c}$, where either $C_{i,j_c} = 0$ in which case we can verify that the property holds, or $C_{i,j_c} \neq 0$, and in this case we have $\text{BookedBySubtree}[i] = \text{MemNeeded}_i$,

then by Lemma 3, we have:

$$\begin{aligned}
MemNeeded_i &= BookedBySubtree[i] \\
&= Booked[i] + \sum_{j \in Children(i) \cap (ACT \cup RUN)} BookedBySubtree[j] \\
&\geq Booked[i] + \sum_{j \in Children(i) \cap (ACT \cup RUN)} f_j
\end{aligned}$$

hence showing the result.

- If i is another ancestor of the node j_c that just finished, then during $[t_{i_0}, t_{i_0+1}]$, $MemNeeded_i - \sum_{j \in Children(i) \cap (ACT \cup RUN)} f_j = B_{i_0}$. Furthermore, DISPATCHMEMORY does: $Booked[i] \leftarrow A_{i_0} + C_{i,j_c}$, where either $C_{i,j_c} = 0$ in which case we can verify that the property holds, or $C_{i,j_c} \neq 0$, and

with an identical argument that above we obtain the result.

This concludes the proof. $\square$

Lemma 5. *If $i \in RUN$, then $Booked[i] = MemNeeded_i$, if $i \in FN$, then $BookedBySubtree[i] = 0$.*

Proof. We show this result by showing a stronger result: if $i \in RUN$, then $Booked[i] = MemNeeded_i = BookedBySubtree[i]$. Then, if $i \in FN$, we have $BookedBySubtree[i] = 0$ (this is a consequence of algorithm DISPATCHMEMORY).

The proof is an induction on the tree structure, starting from leaves. If i is a leaf, then by Lemma 3 and 4 we have:

$$Booked[i] \leq MemNeeded_i \leq BookedBySubtree[i] = Booked[i],$$

hence showing the result.

Let us now consider an interior node i and assume that the result is true for its descendants in the tree. When i is moved to RUN , then all its children are in FN , hence Lemma 3 (items (2) and (3)) and Lemma 4 can be written as:

$$\begin{aligned}
MemNeeded_i &\leq BookedBySubtree[i] \\
BookedBySubtree[i] &= Booked[i] + \sum_{j \in Children(i) \cap FN} BookedBySubtree[j] \\
Booked[i] &\leq MemNeeded_i
\end{aligned}$$

By induction hypothesis, $j \in Children(i) \cap FN, BookedBySubtree[j] = 0$, which shows the desired result. $\square$

We are now ready to show the final result:

Proof of Theorem 1. Let us remind what needs to be proven to show this result:

1. At all time, $\sum_i \text{Booked}[i] \leq M$;
2. If $i \in \text{RUN}$, then $\text{Booked}[i] = \text{MemNeeded}_i$;
3. For all $i \notin \text{FN}$, we cannot decrease the value of $\text{Booked}[i]$; finally, when i is moved to FN , its output f_i should be moved from $\text{Booked}[i]$ to its parent's memory usage: $\text{Booked}[\text{parent}(i)]$.
4. All tasks are executed.

We prove the item by checking how the three algorithms modify M_{Booked} . First, we notice that we ensure that at all time, $M_{\text{Booked}} = \sum_i \text{Booked}[i]$. DISPATCHMEMORY never increases M_{Booked} ; in UPDATECAND-ACT, when modified, M_{Booked} stays below M ; M_{Booked} is not directly modified in MEMBOOKING.

The second item is proven by Lemma 5. For the third item, we check that we never modify $\text{Booked}[i]$ for $i \in \text{FN}$ and that DISPATCHMEMORY correctly transfers f_j to the parent of j at its termination.

For the fourth and last item, we proved by contradiction: suppose that at some time t , all tasks are not processed and no more events happen. If not task termination happens, this means than no tasks are running ($\text{RUN} = \emptyset$). Let us call t_{i_0} the last event (we write $t_{i_0+1} = +\infty$).

- Consider the case where $\text{ACT} \neq \emptyset$. We know that all descendants of nodes in ACT are in $\text{ACT} \cup \text{RUN} \cup \text{FN}$, hence in particular there exists a node i in ACT such that all its descendants are in FN (or it has no descendants). Then, by Lemma 3, 4 and 5 we have during $]t_{i_0}, t_{i_0+1}[$: $\text{Booked}[i] = \text{MemNeeded}_i$, hence i should have been moved to RUN in the execution of MEMBOOKING at time t_{i_0} which contradicts $\text{RUN} = \emptyset$.
- Consider the case where $\text{ACT} = \emptyset$. According to Lemma 1, the i_0 nodes that were inserted in ACT are the i_0 first elements AO . Let k be the $i_0 + 1$ th node in AO . We have also seen in the proof of Lemma 1 that at t_{i_0} , k is the next node to be considered in CAND for activation. Hence the only reason why it was not moved to ACT by UPDATECAND-ACT is because $M_{\text{Booked}} + \text{MissingMem}_j > M$.

Because all tasks are either in FN or in $\text{UN} \cup \text{CAND}$, using Lemma 2 we can compute M_{Booked} :

$$\begin{aligned}
 M_{\text{Booked}} &= \sum_{i \in \text{FN}} \text{Booked}[i] + \sum_{i \in \text{UN} \cup \text{CAND}} \text{Booked}[i] \\
 &= \sum_{i \in \text{FN}} 0 + \sum_{i \in \text{UN} \cup \text{CAND}} \sum_{j \in \text{Children}(i) \cap \text{FN}} f_j \\
 &= \sum_{i, \text{AO}[i] > \text{AO}[k]} \sum_{j \in \text{Children}(i) \cap \text{FN}} f_j
 \end{aligned}$$

We now evaluate the missing memory:

$$\begin{aligned}
& \text{MissingMem}_k \\
&= \text{MemNeeded}_k - \left(\text{Booked}[k] + \sum_{j \in \text{Children}(k)} \text{BookedBySubtree}[j] \right) \\
&= \text{MemNeeded}_k - \text{Booked}[k] \quad \text{thanks to Lemma 5} \\
&= \text{MemNeeded}_k - \sum_{j \in \text{Children}(k) \cap FN} f_j \quad \text{thanks to Lemma 2.}
\end{aligned}$$

In the end, we have

$$M_{\text{Booked}} + \text{MissingMem}_k = \text{MemNeeded}_k + \sum_{i, AO[i] > AO[k]} \sum_{j \in \text{Children}(i) \cap FN} f_j$$

which is exactly the memory used by the sequential schedule AO when processing k and hence is smaller than M , contradicting the previous hypothesis.

In the end, we were able to show that if $RUN = \emptyset$, then $ACT \cup UN \cup CAND = \emptyset$, hence showing that $FN = T$ and the whole tree was processed. $\square$

5.1 Complexity analysis

In this section we give a complexity analysis of the algorithm presented in the previous section. We have chosen to separate the idea of the main algorithm from the optimizations used to lower the execution cost so that the proof of Theorem 1 is more understandable. In this section we now detail those optimizations. A complete version of the Algorithm is available in Appendix B.

Theorem 2. *Let T a tree with n nodes, and H be its height, AO an activation order, EO an execution order, M a memory bound and p processors, then $\text{MEMBOOKING}(T, p, AO, EO, M)$ runs in $O(n(H + \log n))$.*

Proof. First let us define some data structures that we use and update during the execution but that we did not disclose in the presentation of the algorithm to simplify its proof of correctness.

First we introduce some informative arrays:

- We keep an array of size n , **ChNotAct**, such that at all time:

$$\forall i, \text{ChNotAct}[i] = |\{j \in \text{Children}(i) | j \in UN \cup CAND\}|.$$

This array keeps track of the number of children of each nodes that are still in UN or $CAND$. At the beginning of the execution, this array is initiated for all nodes with their number of children. Then it is updated in constant time by **UPDATECAND-ACT** when a node i is moved from $CAND$ to ACT (line 7): $\text{ChNotAct}[\text{parent}(i)] \leftarrow \text{ChNotAct}[\text{parent}(i)] - 1$. Hence the time complexity of updating this table throughout execution is $O(n)$.

- We keep an array of size n , **ChNotFin**, such that at all time:

$$\forall i, \text{ChNotFin}[i] = |\{j \in \text{Children}(i) | j \notin FN\}|.$$

This array keeps track of the number of children of each nodes that are not finished. At the beginning of the execution, this array is initiated for all nodes with their number of children. Then it is updated in constant time by **MEMBOOKING** when a node i finishes (line 10): $\text{ChNotFin}[\text{parent}(i)] \leftarrow \text{ChNotFin}[\text{parent}(i)] - 1$. Hence the time complexity of updating this table throughout execution is $O(n)$.

- We keep an array of size n , **NotUnCand**, such that at all time:

$$\text{NotUnCand}[i] \Leftrightarrow i \notin UN \cup CAND$$

At the beginning of execution, this array is initiated to **false**. It is then updated in constant time by **UPDATECAND-ACT** when a node i is moved from *CAND* to *ACT* (line 7): $\text{NotUnCand}[i] \leftarrow \text{true}$. Hence the time complexity of updating this table throughout execution is $O(n)$.

Now we introduce the main structures used in the computation:

- We implement *CAND* as a heap whose elements are sorted according to the activation order (*AO*). All elements are inserted and removed (with complexity $O(\log n)$) at most once from *CAND*, hence a time complexity of $O(n \log n)$. Furthermore, in **UPDATECAND-ACT**, extracting i from *CAND* (on line 3) is done in constant time.
- In practice belonging to *ACT* can be verified by checking if $\text{MemNeeded}_i \leq \text{BookedBySubtree}[i]$ (Lemma 3, item 2). Hence we do not need a data structure for *ACT*.
- Nevertheless, we use a data structure to remember the elements of *ACT* whose children have all finished their execution. Hence instead of implementing *ACT*, we use:

$$\text{ACTf} = \text{ACT} \cap \{i | \text{ChNotFin}[i] = 0\}.$$

This is implemented as a heap whose elements are sorted according to the activation order (*EO*). In order to keep this property at all time, in **UPDATECAND-ACT** when a node i is moved from *CAND* to *ACT* (line 7), we check if $\text{ChNotFin}[i] = 0$, in which case i is inserted in *ACTf*. Similarly, in **MEMBOOKING** when a node i finishes (line 10), if $\text{ChNotFin}[\text{parent}(i)]$ is set to 0, then we insert $\text{parent}(i)$ in *ACTf* in $O(\log n)$. Note that all elements are inserted and removed at most once from *ACTf*, hence the time complexity of elements going through *ACTf* is $O(n \log n)$. Finally, in **MEMBOOKING**, extracting i from *ACTf* (line 14) is done in constant time.

If we detail all operations contained in **MEMBOOKING**:

- DISPATCHMEMORY is called exactly once per node (every time a node finishes). For a given node j of depth h_j , it does at most $O(h_j)$ operations (note that the test $i \in ACT \cup RUN$ on line 7 can be done in constant time by checking if $\text{NotUnCand}[i] = \text{true}$), which gives a total cumulative time complexity of $O(nH)$.
- UPDATECAND-ACT is called for each event. Note that we have already accounted for removing or inserting all elements from the different sets. Finding the element to remove from CAND on line 3 is done in constant time because CAND is a heap sorted according to AO. Similarly, the test on line 8 can be done in constant time by simply testing whether $\text{ChNotAct}[\text{parent}(i)] = 0$.

The most time consuming event is the computation of MissingMem_i on line 4 which could be computed up to $O(n)$ times for a given node if the condition on M_{Booked} and M on line 5 is not satisfied (hence giving a time complexity of $O(n^2)$). To avoid this case, we will extend the property (3) proven in Lemma 3 to a node i of CAND such that $\text{BookedBySubtree}[i]$ has already been computed once. To do this, we first initialized $\text{BookedBySubtree}[j]$ to -1, and replace line 4 in UPDATECAND-ACT by:

If $\text{BookedBySubtree}[i] = -1$, **then**

$$\text{BookedBySubtree}[i] \leftarrow \text{Booked}[i] + \sum_{j \in \text{Children}(i)} \text{BookedBySubtree}[j]$$

$$\text{MissingMem}_i \leftarrow \max(0, \text{MemNeeded}_i - \text{BookedBySubtree}[i])$$

We also extend the “while” loop in DISPATCHMEMORY (line 7) from $i \in ACT \cup RUN$ to $i \in \{i | \text{BookedBySubtree}[i] \neq -1\}$.

- Finally, the last “while” loop of MEMBOOKING (line 13) is entered once per element contained in ACT, that is exactly n times. Furthermore, because $ACTf$ is a heap sorted according to EO, removing one element is done in $O(\log n)$, which gives a cumulative complexity of $O(n \log n)$ for this last loop.

Finally accounting for all operations, the total time complexity of this optimized algorithm is $O(n(\log(n) + H))$. $\square$

6 A new makespan lower bound

It is usual in scheduling problems to compare the makespan of proposed algorithms to lower bounds, as the optimal makespan is usually out of reach (NP-complete). The classical lower bound for scheduling task graphs considers the maximum between the average workload (total computation time divided

by the number of processors) and the longest path in the graph. In a memory-constrained environment, the memory bound itself may prevent the simultaneous execution of too many tasks. We propose here a new lower bound that takes this into account.

Theorem 3. *Let $C_{\max}$ be the makespan of any correct schedule of a tree whose peak memory is at most the memory bound M , and t_i the processing time of task i . Then*

$$C_{\max} \geq \frac{1}{M} \sum_i \text{MemNeeded}_i \times t_i.$$

Proof. Consider a task i as described in the model of Section 2: its processing requires a memory of MemNeeded_i (see Equation (1)). As stated in the theorem, we denote by t_i its processing time. Consider the total memory usage of a schedule, that is, the sum over all time instants t of the memory used by this schedule. Then, task i contributes to at least $\text{MemNeeded}_i \times t_i$ to this total memory usage. For a schedule of makespan $C_{\max}$, the total memory usage cannot be larger than $C_{\max} \times M$, where M is the memory bound. Thus, $\sum_i \text{MemNeeded}_i \times t_i \leq C_{\max} \times M$ which concludes the proof. $\square$

We have noticed in the simulations described in the next section that with eight processors, this new lower bound improved the classical lower bound in 22% of the case for on actual assembly trees, and in these cases the average increase in the bound was 46%. For the simulations on synthetic trees, it has improved the lower bound in 33% of the cases, with an average improvement of 37%. Contrarily to the previous lower bound, this new lower bound does not depend on the number of processors, hence the improvement is even greater with more processors.

It is important to understand that the more precise the lower bound, the more information is available for a possible improvement of the considered heuristics.

7 Simulations

We report here the results of the simulations that we performed to compare our new booking strategy (MEMBOOKING) to the two other scheduling heuristics presented above: the basic ACTIVATION policy [1] presented in Section 3.1 and the booking strategy [7] for reduction trees, denoted MEMBOOKINGREDTREE, from Section 3.2. In the latter, the tree is first transformed into a “reduction tree” [7] by adding some fictitious nodes and edges before the scheduling strategy can be applied.

7.1 Data sets

The trees used for the simulations come from two data sets, which we briefly describe below.

The first data set, also called *assembly trees* are trees corresponding to the multifrontal direct factorization of a set of sparse matrices obtained from the University of Florida Sparse Matrix Collection (<http://www.cise.ufl.edu/research/sparse/matrices/>). This data set is taken from [7], where more information can be found on multifrontal factorization and on how the trees are constructed. This data set consists in 608 trees which contains from 2,000 to 1,000,000 nodes. Their height ranges from 12 to 70,000 and their maximum degree ranges from 2 to 175,000.

The second data set is synthetic. The node degree is taken randomly in $[1; 5]$, with a higher probability for small values to avoid very large and short trees, on which we already observed with the previous data set that our algorithm outperforms other strategies. The following table gives the precise node degree probabilities:

X	$Pr(\delta = X)$
1	0.58
2	0.17
3,4,5	0.08

Edges weights follow a truncated exponential distribution of parameter 1. More precisely, we first generate a random number from an exponential distribution of parameter 1, which is multiplied it by 100 and then truncated to fit in the interval $[10; 10.000]$. The processing size of a node is 10% of its outgoing edge weight and its processing time is proportional to its outgoing edge degree. We generated 50 synthetic trees of 1.000, 10.000 and 100.000 nodes, which results in trees of respective average height of 63, 95 and 131.

7.2 Simulation setup

All three strategies were implemented in C, with special care to avoid complexity issues. These strategies have been applied to the two tree families described above, with the following parameters:

- We tested 5 different number of processors (2,4,8,16,32). The results were quite similar, expect for the extreme case (too large or too small parallelism). We mainly report the results for 8 processors but also mention results for other number of processors.
- For each tree, we first computed the postorder traversal that minimizes the peak memory. This gives the minimum amount of memory needed for both ACTIVATION and MEMBOOKING (MEMBOOKINGREDTREE is likely to use more memory as it works on a transformed tree). The heuristics are then tested with a factor of this minimal memory, which we call below *normalized memory bound*. We only plot an average result when a given strategy was able to schedule at least 95% of the trees within the memory bound.

- The previous postorder was used as input for both the activation order *AO* and the execution order *EO* for ACTIVATION and MEMBOOKING in general. We also tested other orders for activation and execution, such as other postorders, critical path ordering, or even optimal (non-postorder) ordering for peak memory [10]. As seen later, this only results in slightly noticeable change in performance.

During the simulations of the parallel executions, we reported the makespan (total completion time), which is normalized by the maximum of the classical lower bound and our new memory related lower bound (see Section 6). We also reported the peak memory of the resulting schedules, as well as the time needed to compute the schedule. This scheduling time does not include the computation of the activation or execution order, which may be done beforehand.

7.3 Results on the assembly trees

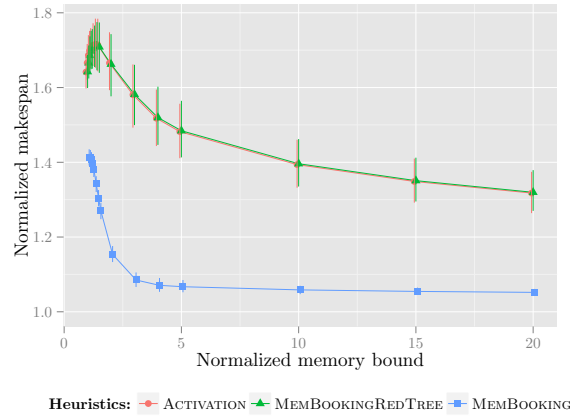


Figure 2: Makespan of assembly trees with all heuristics depending on the memory bound

Figure 2 plots the average normalized makespan of all strategies on various memory constraints, while Figure 3 shows the speedup of MEMBOOKING over ACTIVATION. We notice that for a memory bound twice the minimum memory, MEMBOOKING is 1.4 faster than ACTIVATION on average. However, even this particular speedup spans a wide interval (between 1 and 6) due to the large heterogeneity of the assembly trees. Note that the two heuristics from the literature give very similar results: this is explained by the fact that MEMBOOKINGREDTREE first transforms the trees before applying a smart booking strategy: on these trees, adding fictitious edges has the same effect than booking to much memory (as ACTIVATION does) and hinders the benefit of the booking strategy. We also note that MEMBOOKING is able to take advantage of

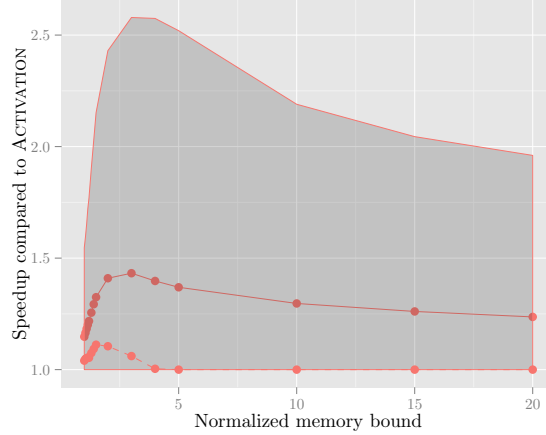


Figure 3: Speedup of MEMBOOKING compared to ACTIVATION on assembly trees. Plain lines represent the average speedup, dotted lines connect the median speedup while the ribbon depicts the results between the first and ninth decile (the maximum speedup always lies between 5 and 6).

very scarce memory conditions: as soon as the available memory increases from its minimum value, its makespan drops and reaches only 10% above the lower bound for 3x the minimum memory, leaving very little room to hope for better algorithms. This is also illustrated by Figure 4, which plots the real use of the memory by the heuristics: while ACTIVATION and MEMBOOKINGREDTREE are very conservative, MEMBOOKING is able to use larger fraction of the available memory when it is limited.

Discussion on execution time: From the complexity analysis (Section 5.1) we expect our algorithm to add significant overhead when trees are very deep (nH term of the worst-case complexity). We first study the cumulative running time in Figure 5 of the various strategy as a function of the number of nodes in the trees. All strategies have similar running times, except on a subset of trees for which our heuristic is much slower (≈ 10 s). We verify that those running times indeed depend on the height of the tree in Figure 6. Another noticeable fact from this figure is that overall the average overhead for each node remains negligible (below 1ms per node with height $H = 10^5$!).

As future work, it may be interesting to get rid of this height factor in the complexity of the algorithm especially for cases when $H = O(n)$.

To give some hindsight on the importance of this factor, we decided to study the speedup of MEMBOOKING compared to that of ACTIVATION as a function of the tree height. In particular, one can see from the correlations between Figures 5 and 6 that the trees of large height are trees where $H = O(n)$. We report these results in Figure 7 where we show the relation between achievable speedup and tree height: very deep trees usually correspond to thin ones and

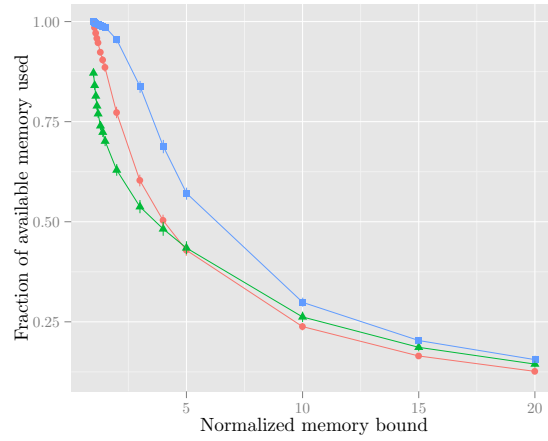


Figure 4: Fraction of memory used by all heuristics on assembly trees (same legend as Figure 2)

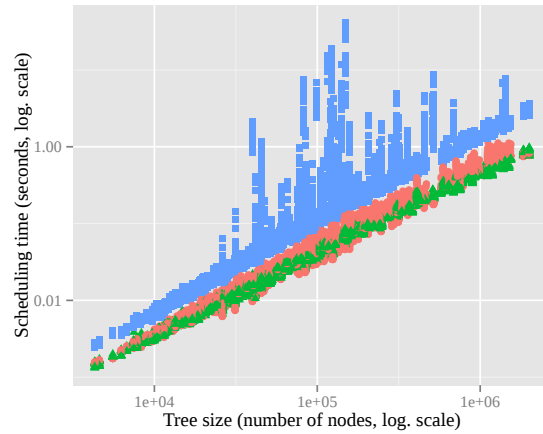


Figure 5: Running times of the heuristics on assembly trees (same legend as Figure 2)

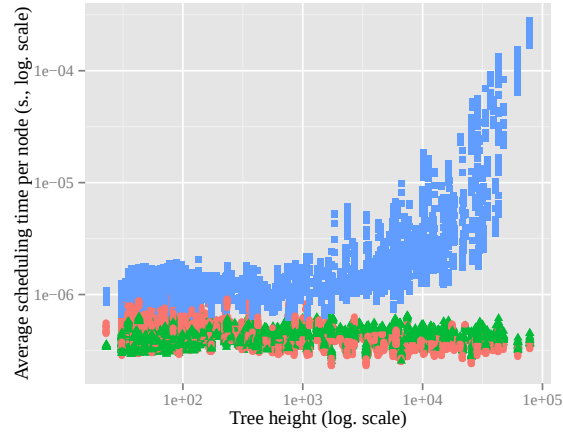


Figure 6: Running times of the heuristics on assembly trees (same legend as Figure 2)

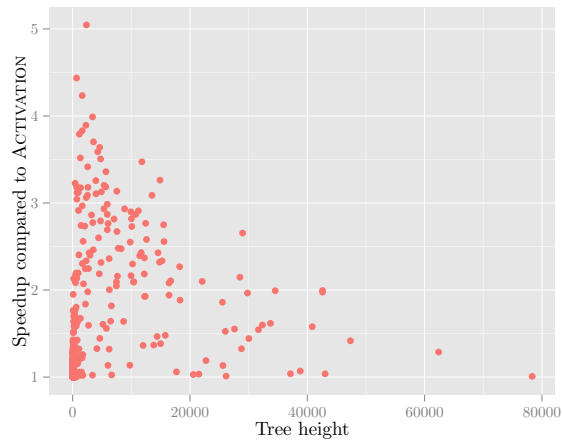


Figure 7: Speedup of MEMBOOKING compared to ACTIVATION on assembly trees when the normalized memory bound is 2 for all 608 trees.

do not provide big opportunities for increasing the parallelism. This is why our strategy achieves best speedups on shallow trees. With this in mind, an interesting study would be to derive a good measure on trees which may hint whether the use of a sophisticated strategy such as MEMBOOKING is needed. Finding such a good measure would need a particular research effort and is out of scope of this paper.

7.3.1 Changing the activation or execution order

On Figure 8, we present the average makespan of the ACTIVATION heuristic for various activation and execution order. This figure is similar to Figure 2 for the

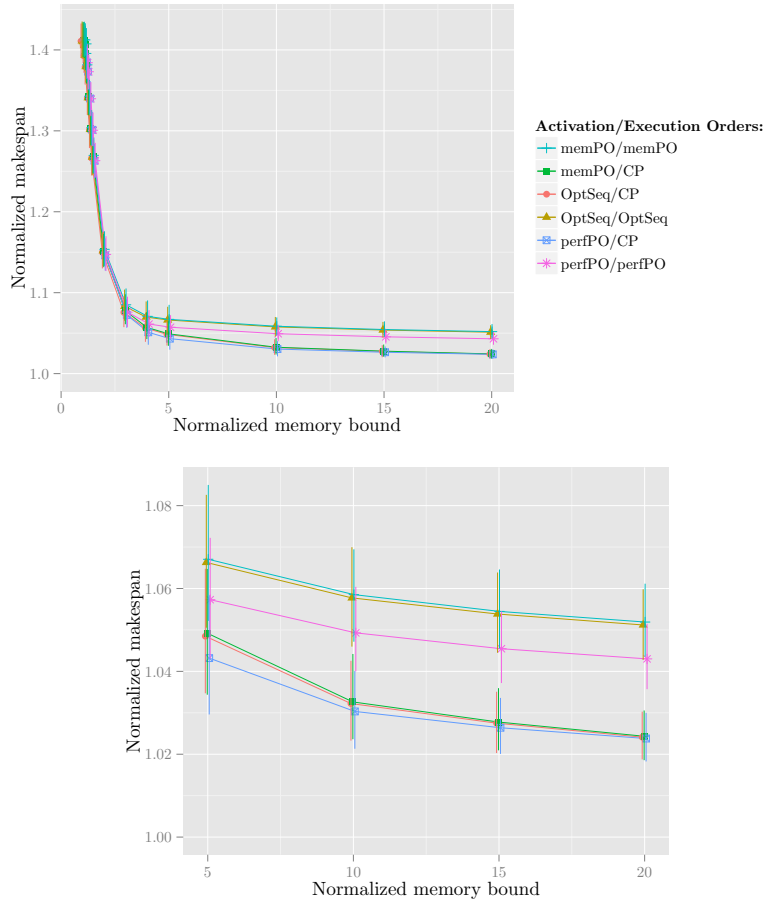


Figure 8: Makespan of assembly trees for the proposed MEMBOOKING strategy using different activation and execution order. The bottom plot corresponds to a zoom on the part where the memory is less limited.

other scheduling heuristics. The activation and execution order used are the following:

- memPO (memory PostOrder): the sequential postorder traversal that minimizes the peak memory (NB: this is the order chosen activation and execution order of both ACTIVATION and MEMBOOKING in all other plots of this section);
- CP (Critical Path): nodes orders by decreasing bottom-level;
- OptSeq (Optimal Sequential): the sequential (non postorder) traversal that minimizes the peak memory, computed as in [10];
- perfPO (performance PostOrder): another postorder traversal, designed for parallel performance (subtrees with larger critical path are scheduled first, which, in a parallel execution, is supposed to give higher priority to nodes with large critical path).

We notice that the results of using different orders for activation and/or execution slightly change the results: using CP as an execution order always gives a small but noticeable improvement over the other strategies. On the contrary, the choice of the activation order has little effect on the final makespan. The same effects can be seen on ACTIVATION (when changing the activation/execution orders) and MEMBOOKING (when changing its priority order). However, the gap between the performance of different orders is much smaller than the gap of using different scheduling strategy: changing the activation/execution order does not change the ranking of the scheduling policies

7.3.2 Results on other numbers of processors

Figure 9 depicts the makespan of all strategies for increasing number of processors. We notice that the gain of the proposed MEMBOOKING strategy increases with the number of available processors, as it also increased the potential parallelism in the tree.

7.4 Results on the synthetic trees

The simulations on synthetic trees show the same general trends as what we notices on assembly trees, and thus we only review briefly their results.

Figure 10 and 11 shows that MEMBOOKING is once again able to schedule trees faster in a memory-constrained environnement. Synthetic trees are more regular and homogeneous than assembly trees, so that the speedup of MEMBOOKING over ACTIVATION is more regular. It reaches an average 1.3 when the memory is twice the bound.

Finally, the scheduling time of all strategies is always below 0.1 seconds due to the smaller height of the trees. Note that MEMBOOKINGREDTREE is not able to schedule most trees in a very constrained memory environment: when the memory bound is smaller than 1.4 times the minimum memory peak of a sequential postorder processing, MEMBOOKINGREDTREE is unable to schedule 33% of the trees (or more) and thus is not included in the plot.

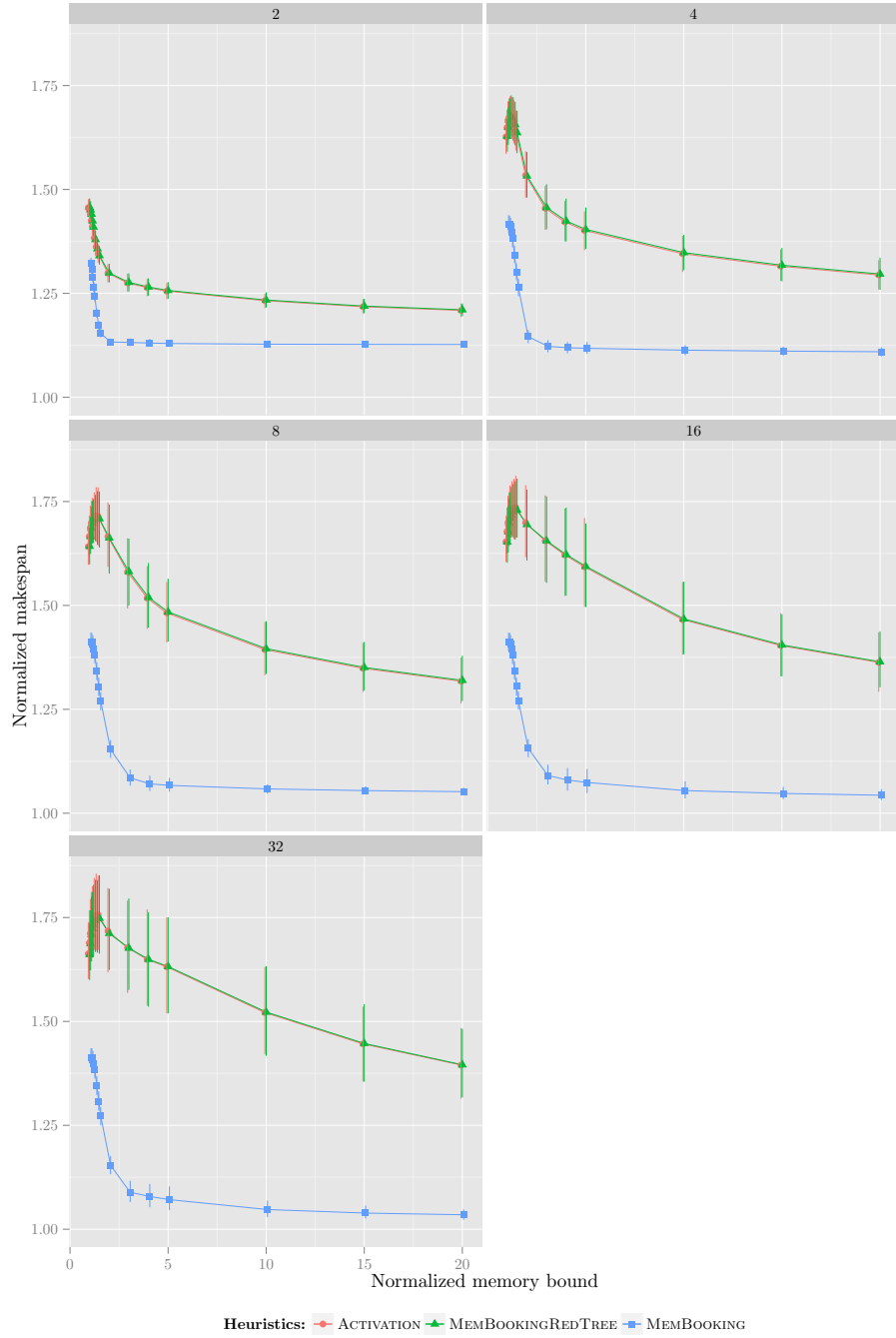


Figure 9: Makespan of assembly trees for all strategies on various number of processors.

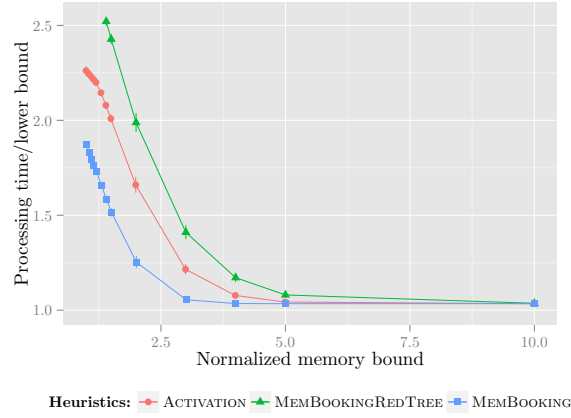


Figure 10: Makespan of synthetic trees with all heuristics depending on the memory bound

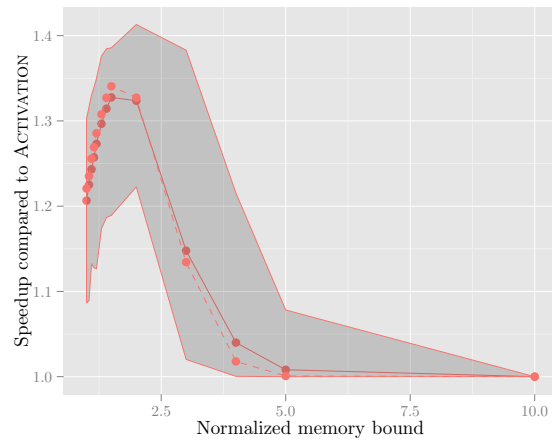


Figure 11: Speedup of MEMBOOKING compared to ACTIVATION on synthetic trees. Plain lines represent the average speedup, dotted lines connect the median speedup while the ribbon depicts the minimum/maximum results.

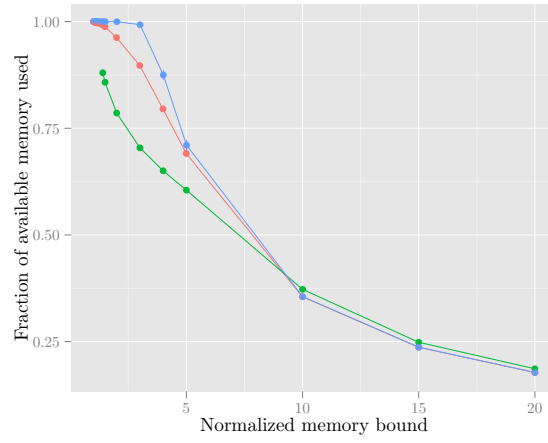


Figure 12: Fraction of memory used by all heuristics on synthetic trees (same legend as Figure 10)

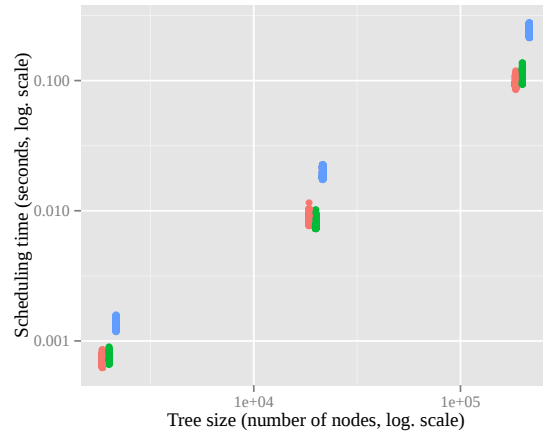


Figure 13: Running times of the heuristics on synthetic trees (same legend as Figure 10)

7.4.1 Changing the activation or execution order

On Figure 14, we present the average makespan of the ACTIVATION heuristic for various activation and execution order. The activation/execution orders used in these simulations are the same as the one described in Section 7.3.1.

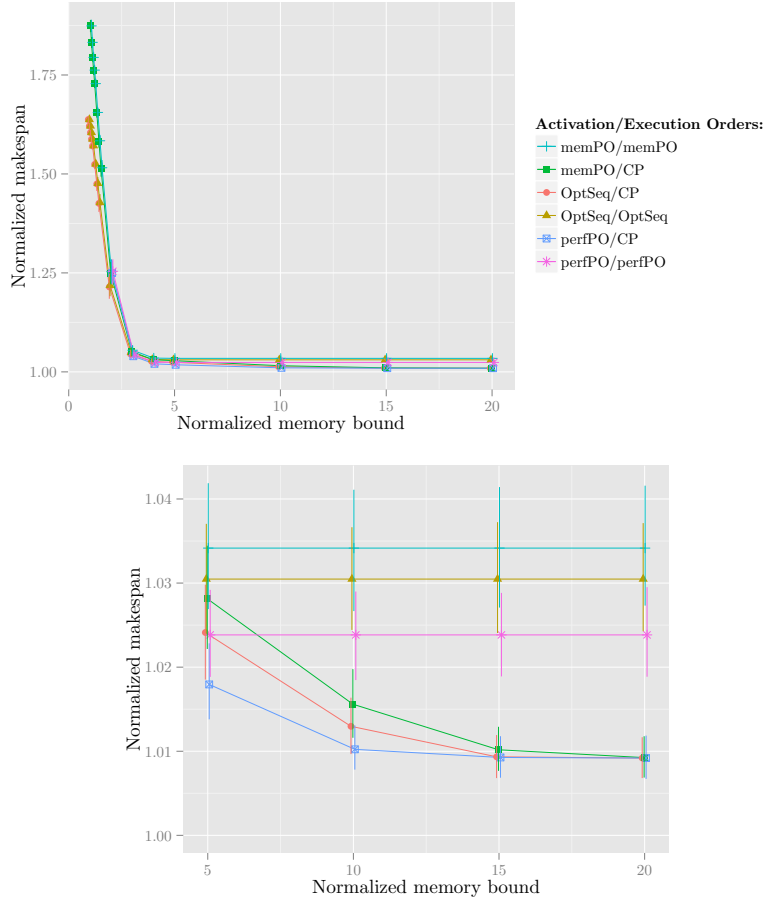


Figure 14: Makespan of assembly trees for the proposed MEMBOOKING strategy using different activation and execution order. The bottom plot corresponds to a zoom on the part where the memory is less limited.

Once again, we notice that an improvement when using CP (Critical Path) as an execution order, and that the improvement is small enough to keep the ranking of the different scheduling policies.

7.4.2 Results on other numbers of processors

Figure 15 depicts the makespan of all strategies for increasing number of processors for synthetic trees. Once again, we notice that the gain of the proposed MEMBOOKING strategy increases with the number of available processors: it is hardly noticeable for 2 or 4 processors, but becomes important for 8, 16 and 32 processors.

8 Conclusion

In this paper, we have proposed a novel algorithm for scheduling task trees on computing platforms with bounded shared memory. The proposed algorithm carefully allocates memory to activated nodes, and accurately predicts how much memory can be recycled from the processing of a node’s ancestors. We have shown that it is always able to schedule a tree under an admissible memory bound, and that its complexity is sufficiently small to allow its implementation is actual runtime schedulers. By performing simulations on both actual assembly trees of sparse direct multifrontal solvers and on a broader class of synthetic trees, we have proved that it outperforms its two competitors from the literature, especially when memory is a scarce resource. Incidentally, we have proposed a new makespan lower bound that takes into account a bound on the shared memory, which, to the best of our knowledge, is the first of its kind.

This study is the first step in the design of realistic schedulers for task trees handling large data, such as assembly trees. One major extension would be to consider parallel tasks rather than only sequential ones. To this goal, one would need to make several adaptations to cope with the extra memory needed for a parallel processing, and to solve the unavoidable trade-off between allocating many processors to big tasks (and losing on tree parallelism) and allocating many tasks in parallel (and threatening the memory bound). Nevertheless, we are confident that the algorithm presented in this paper (or its adaptation) would still provide an improvement over the classical ACTIVATION algorithm. Another necessary extension would be to consider distributed memories, or even a mix of distributed/shared memory (as in clusters of cores sharing a dedicated memory).

Acknowledgments

This work was supported by the ANR SOLHAR project funded by the French National Research Agency.

References

- [1] Emmanuel Agullo, Alfredo Buttari, Abdou Guermouche, and Florent Lopez. “Multifrontal QR Factorization for Multicore Architectures over Runtime

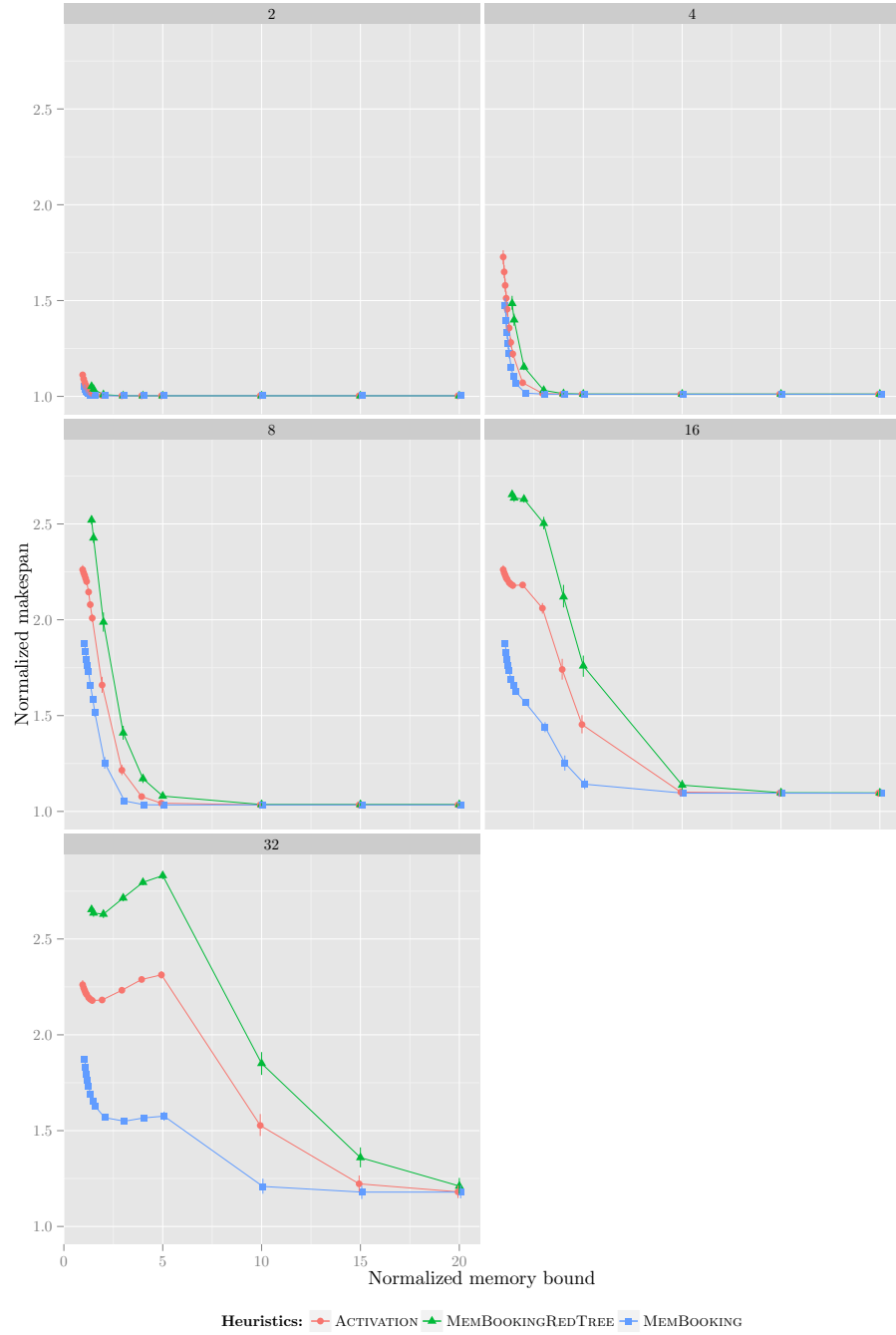


Figure 15: Makespan of all strategies for synthetic trees on various number of processors.

- Systems.” In: *Euro-Par 2013 Parallel Processing - 19th International Conference*. 2013, pp. 521–532.
- [2] Emmanuel Agullo, Patrick R. Amestoy, Alfredo Buttari, Abdou Guermouche, Jean-Yves L’Excellent, and François-Henry Rouet. “Robust Memory-Aware Mappings for Parallel Multifrontal Factorizations.” In: *SIAM J. Scientific Computing* 38.3 (2016).
- [3] P. R. Amestoy, I. S. Duff, J. Koster, and J.-Y. L’Excellent. “A fully asynchronous multifrontal solver using distributed dynamic scheduling.” In: *SIAM Journal on Matrix Analysis and Applications* 23.1 (2001), pp. 15–41.
- [4] P. R. Amestoy, A. Guermouche, J.-Y. L’Excellent, and S. Pralet. “Hybrid scheduling for the parallel solution of linear systems.” In: *Parallel Computing* 32.2 (2006), pp. 136–156.
- [5] Guillaume Aupy, Clément Brasseur, and Loris Marchal. *Dynamic memory-aware task-tree scheduling*. Research Report 8966. France: INRIA, Oct. 2016.
- [6] Timothy A. Davis. *Direct Methods for Sparse Linear Systems*. Fundamentals of Algorithms. Philadelphia: Society for Industrial and Applied Mathematics, 2006.
- [7] Lionel Eyraud-Dubois, Loris Marchal, Oliver Sinnén, and Frédéric Vivien. “Parallel Scheduling of Task Trees with Limited Memory.” In: *TOPC* 2.2 (2015), p. 13.
- [8] Mathias Jacquelin, Loris Marchal, Yves Robert, and Bora Ucar. “On Optimal Tree Traversals for Sparse Matrix Factorization.” In: *Proceedings of the 25th IEEE International Parallel and Distributed Processing Symposium (IPDPS’11)*. Los Alamitos, CA, USA: IEEE Computer Society, 2011, pp. 556–567.
- [9] Chi-Chung Lam, Thomas Rauber, Gerald Baumgartner, Daniel Cociorva, and P. Sadayappan. “Memory-optimal evaluation of expression trees involving large objects.” In: *Computer Languages, Systems & Structures* 37.2 (2011), pp. 63–75.
- [10] Joseph W. H. Liu. “An application of generalized tree pebbling to sparse matrix factorization.” In: *SIAM J. Algebraic Discrete Methods* 8.3 (1987), pp. 375–395.
- [11] Joseph W. H. Liu. “On the storage requirement in the out-of-core multifrontal method for sparse factorization.” In: *ACM Transaction on Mathematical Software* (1986).
- [12] Arun Ramakrishnan, Gurmeet Singh, Henan Zhao, Ewa Deelman, Rizos Sakellariou, Karan Vahi, Kent Blackburn, David Meyers, and Michael Samidi. “Scheduling Data-Intensive Workflows onto Storage-Constrained Distributed Resources.” In: *Proceedings of the IEEE Symposium on Cluster Computing and the Grid (CCGrid’07)*. Los Alamitos, CA, USA: IEEE Computer Society, 2007, pp. 401–409.

- [13] Ravi Sethi and J.D. Ullman. “The Generation of Optimal Code for Arithmetic Expressions.” In: *J. ACM* 17.4 (1970), pp. 715–728.
- [14] Wayne E. Smith. “Various optimizers for single-stage production.” In: *Naval Research Logistics Quarterly* 3.1-2 (1956), pp. 59–66. ISSN: 1931-9193.

A Sequential post-order minimizing average memory

We define the average memory of a schedule as follows:

$$AvgMem = \frac{1}{C_{\max}} \int_{t=0}^{C_{\max}} mem(t) dt,$$

where $mem(t)$ is the memory used by the schedule at time t and $C_{\max}$ is the makespan.

Theorem 4. *A post-order traversal that minimizes the average memory is obtained by processing subtrees by non-increasing T_i/f_i value, where T_i is the total time needed to process the subtree rooted in i .*

Proof. Let us consider a tree rooted at node r , with k children $c_1, \dots, c_k$. If the subtrees are processed in this order, the average memory of the processing of the whole subtree is:

$$AvgMem(r) = \frac{1}{T_r} \left(\underbrace{\sum_{i=1}^k \left(\left(\sum_{j=1}^{i-1} f_j \right) + AvgMem(c_i) \right) \cdot T_{c_i}}_{\text{processing of } c_i} + \underbrace{\left(\sum_{j=1}^k f_j \right) t_r}_{\text{processing of } r} \right)$$

Let us now rewrite this expression:

$$T_r \cdot AvgMem(r) = \left(\sum_{j=1}^k f_j \right) t_r + \sum_{i=1}^k AvgMem(c_i) T_{c_i} + \sum_{i=1}^k \left(\sum_{j=1}^i f_j \right) T_{c_i} - \sum_{i=1}^k f_i T_{c_i}$$

The only term that depends on the subtree ordering is the third and penultimate one, namely $\sum_{i=1}^k \left(\sum_{j=1}^i f_j \right) \times T_{c_i}$. As this term does not depend on the average memory of the subtrees, it is easy to see that an optimal post-order for a tree can be obtained by optimizing the average memory of its subtrees (to minimize the second term) and carefully ordering them (to minimize the third term).

We then rewrite the third term using $T_{c_i} = w_i$ and $f_i = p_i$ and obtain $\sum_{i=1}^k w_i \left(\sum_{j=1}^i p_j \right)$ which we identify as the weighted sum flow of independant tasks on a single processor (problem 1|| $\sum w_i C_i$ using Graham’s notation). Following the classical Smith’s rule [14], we know that an ordering with minimal

sum flow is obtained by processing the tasks by non-increasing w_i/p_i . Thus, an optimal ordering of the subtrees for the average memory is obtained by processing them by non-increasing T_{c_i}/f_i . $\square$

B Complete and optimized algorithm

Algorithm 5: INIT(T, AO, EO)

```

// Initialization of the data structures
CAND           $\leftarrow$  AO-sorted heap (init: Leaves( $T$ ); fun: AO-insert,
                                AO-remove, AO-min)
ACTf           $\leftarrow$  EO-sorted heap (init: empty; fun: EO-insert,
                                EO-remove, EO-min)
1 ChNotAct[1..n]  $\leftarrow$  array of size  $n$  (init:  $\forall i, \text{ChNotAct}[i] \leftarrow |\text{Children}(i)|$ )
  ChNotFin[1..n]  $\leftarrow$  array of size  $n$  (init:  $\forall i, \text{ChNotFin}[i] \leftarrow |\text{Children}(i)|$ )
  NotUnCand[1..n]  $\leftarrow$  array of size  $n$  (init:  $\forall i, \text{NotUnCand}[i] = \text{false}$ )
  Booked[1..n]  $\leftarrow$  array of size  $n$  (init:  $\forall i, \text{Booked}[i] = 0$ )
  BookedBySubtree[1..n]  $\leftarrow$  array of size  $n$  (init:  $\forall i, \text{BookedBySubtree}[i] = -1$ )

2  $M_{\text{Booked}} \leftarrow 0$ 

```

Algorithm 6: MEMBOOKING(T, p, AO, EO, M)

```

1 INIT( $T, AO, EO$ )

2 while the whole tree is not processed do
3   Wait for an event (task finishes or  $t = 0$ )
4   foreach just finished node  $j$  do
5      $B = \text{Booked}[j]$ 
6      $\begin{cases} \text{Booked}[j] & \leftarrow 0 \\ M_{\text{Booked}} & \leftarrow M_{\text{Booked}} - B \\ \text{BookedBySubtree}[j] & \leftarrow 0 \end{cases}$ 
7      $i \leftarrow \text{parent}(j)$ 
8     if  $i \neq \text{NULL}$  then
9        $\text{ChNotFin}[i] \leftarrow \text{ChNotFin}[i] - 1$ 
10      if  $\text{ChNotFin}[i] = 0$  and  $\text{BookedBySubtree}[i] \geq \text{MemNeeded}_i$  then
11         $\text{EO-insert}(i, \text{ACTf})$ 
12         $\begin{cases} \text{Booked}[i] & \leftarrow \text{Booked}[i] + f_j \\ M_{\text{Booked}} & \leftarrow M_{\text{Booked}} + f_j \\ B & \leftarrow B - f_j \end{cases}$ 
13      while  $i \neq \text{NULL}$  and  $\text{BookedBySubtree}[i] \neq -1$  and  $B \neq 0$  do
14         $C_{j,i} = \min(B, \max(0, \text{MemNeeded}_i - (\text{BookedBySubtree}[i] - B)))$ 
15         $\begin{cases} \text{Booked}[i] & \leftarrow \text{Booked}[i] + C_{j,i} \\ M_{\text{Booked}} & \leftarrow M_{\text{Booked}} + C_{j,i} \\ \text{BookedBySubtree}[i] & \leftarrow \text{BookedBySubtree}[i] - (B - C_{j,i}) \\ B & \leftarrow B - C_{j,i} \\ i & \leftarrow \text{parent}(i) \end{cases}$ 
16       $\text{WaitForMoreMem} \leftarrow \text{false}$ 
17      while  $!(\text{WaitForMoreMem})$  and  $\text{CAND} \neq \emptyset$  do
18         $i \leftarrow \text{AO-min}(\text{CAND})$ 
19        if  $\text{BookedBySubtree}[i] = -1$  then
20           $\text{BookedBySubtree}[i] \leftarrow \text{Booked}[i] + \sum_{j \in \text{Children}(i)} \text{BookedBySubtree}[j]$ 
21         $\text{MissingMem}_i = \max(0, \text{MemNeeded}_i - \text{BookedBySubtree}[i])$ 
22        if  $M_{\text{Booked}} + \text{MissingMem}_i \leq M$  then
23           $\begin{cases} \text{Booked}[i] & \leftarrow \text{Booked}[i] + \text{MissingMem}_i \\ M_{\text{Booked}} & \leftarrow M_{\text{Booked}} + \text{MissingMem}_i \\ \text{BookedBySubtree}[i] & \leftarrow \text{Booked}[i] + \sum_{j \in \text{Children}(i)} \text{BookedBySubtree}[j] \end{cases}$ 
24           $\text{AO-remove}(i, \text{CAND});$  if  $\text{ChNotFin}[i] = 0$  then  $\text{EO-insert}(i, \text{ACTf})$ 
25           $\text{ChNotAct}[\text{parent}(i)] \leftarrow \text{ChNotAct}[\text{parent}(i)] - 1$ 
26          if  $\text{ChNotAct}[\text{parent}(i)] = 0$  then  $\text{AO-insert}(\text{parent}(i), \text{CAND})$ 
27        else
28           $\text{WaitForMoreMem} \leftarrow \text{true}$ 
29      while there is an available processor  $P_u$  and  $\text{ACTf} \neq \emptyset$  do
30         $i \leftarrow \text{EO-min}(\text{ACTf}); \text{EO-remove}(i, \text{ACTf})$ 
31        Make  $P_u$  process  $i$ 

```



**RESEARCH CENTRE
GRENOBLE – RHÔNE-ALPES**

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

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

ISSN 0249-6399