

On Scaling Dynamic Programming Problems with a Multithreaded Tabling System

Miguel Areias and Ricardo Rocha

*CRACS & INESC TEC, Faculty of Sciences, University of Porto
Rua do Campo Alegre, 1021/1055, 4169-007 Porto, Portugal
{miguel-areias, ricroc}@dcc.fc.up.pt*

Abstract

Tabling is a recognized and powerful implementation technique that improves the declarativeness and expressiveness of traditional Prolog systems in dealing with recursion and redundant computations. In a nutshell, tabling consists of storing intermediate answers for subgoals so that they can be reused when a variant subgoal appears. The tabling technique can thus be viewed as a natural tool to implement dynamic programming problems, where a general recursive strategy divides a problem in simple sub-problems that, often, are the same. When tabling is combined with multithreading, we have the best of both worlds, since we can exploit the combination of higher declarative semantics with higher procedural control. However, such combination for dynamic programming problems is very difficult to exploit in order to achieve execution scalability as we increase the number of running threads. To the best of our knowledge, no previous work showed to be able to scale the execution of multithreaded dynamic programming problems. In this work, we focus on two well-known dynamic programming problems, the Knapsack and the Longest Common Subsequence problems, and we discuss how we were able to scale their execution by taking advantage of the multithreaded tabling engine of the Yap Prolog system.

Keywords: Dynamic Programming, Multithreaded Tabling, Scalability

1. Introduction

Dynamic programming [1] is a general recursive strategy that consists in dividing a problem in simple sub-problems that, often, are the same. The idea behind dynamic programming is to reduce the number of computations: once an answer to a given sub-problem has been computed, it is memorized and the next time the same answer is needed, it is simply looked up. Dynamic programming is especially useful when the number of overlapping sub-problems grows exponentially as a function of the size of the input. Dynamic programming can be implemented using either a *bottom-up* or a *top-down* approach. In bottom-up, it starts

from the base sub-problems and recursively computes the next level sub-problems until reaching the answer to the given problem. On the other hand, the top-down approach starts from the given problem and uses recursion to subdivide a problem into sub-problems until reaching the base sub-problems. Answers to previously computed sub-problems are reused rather than being recomputed. An advantage of the top-down approach is that it might not need to compute all possible sub-problems.

Most of the proposals that can be found in the literature to parallelize dynamic programming problems follow the parallelization of a sequential bottom-up algorithm. All these proposals are usually based on a careful analysis of the sequential

algorithm in order to find the best way to minimize the data dependencies in the supporting data structure for memorization, often a matrix or an array, resulting in a parallelization that requires a synchronization mechanism before recursively computing the next level sub-problems. Alternatively, a generic proposal to parallelize top-down dynamic programming algorithms is Stivala et al.’s work [2], where a set of threads solve the entire dynamic program independently but with a randomized choice of sub-problems, i.e., each thread runs exactly the same function, but the randomization choice of sub-problems results in the threads diverging to compute different sub-problems while reusing the sub-problem’s results computed in the meantime by the other threads.

Tabling [3] is a recognized and powerful implementation technique that proved its viability and efficiency to overcome Prolog’s susceptibility to infinite loops and redundant computations. Tabling consists of saving and reusing the results of sub-computations during the execution of a program and, for that, the calls and the answers to tabled subgoals are memorized in a proper data structure called the *table space*. Tabling can thus be viewed as a natural tool to implement dynamic programming problems. When tabling is combined with multithreading, we have the best of both worlds, since we can exploit the combination of higher declarative semantics with higher procedural control. However, such combination for dynamic programming problems is very difficult to exploit in order to achieve execution scalability as we increase the number of running threads. To the best of our knowledge, XSB [4] and Yap [5] were the only Prolog systems that were able to combine both multithreading with tabling, but none of them showed until now to be able to scale the execution of multithreaded dynamic programming problems. This is a difficult task since we need to combine the explicit thread control required to launch, assign and schedule tasks to threads, with the built-in tabling evaluation mechanism, which is implicit and cannot be controlled by the user. This main motivation of this work is to show a simple multithreaded table space design that is efficient enough to obtain good a scalability. To do so, we focus on two well-known dynamic programming problems, the Knapsack and the Longest Common Subsequence (LCS) problems, and we discuss how we were able to scale their execution by taking advantage of the multithreaded tabling engine of the Yap Prolog sys-

tem. For each problem, we present a multithreaded tabled top-down and bottom-up approach. For the top-down approach, we use Yap’s mode-directed tabling support [6] that allows to aggregate answers by specifying pre-defined modes such as *min* or *max*. For the bottom-up approach, we use Yap’s standard tabling support [7]. To put our results in perspective, we also experimented with the state-of-the-art XSB Prolog system¹ using *thread shared tables* [4]. Our experiments on a 32-core AMD machine show that using Yap’s simple and efficient multithreaded table space design, we were able to scale the execution of both problems for both top-down and bottom-up approaches. Regarding the comparison with XSB and since the Yap’s results clearly outperform the results of obtained for XSB, we complete the comparison by outline some of the possible reasons for the big gap between both systems. The remainder of the paper is organized as follows. First, we briefly describe some background about Yap’s standard, mode-directed and multithreaded tabling support. Then, we discuss related work and briefly introduce XSB’s approach to multithreaded tabling. Next, for both Knapsack and LCS problems, we introduce the problem, present in detail our parallel implementations using either a top-down and bottom-up dynamic programming approach, and discuss the experimental results. At the end, we outline some conclusions.

2. Background

This section introduces some background needed for the following sections.

2.1. Standard Tabling

The basic idea behind tabling is straightforward: programs are evaluated by storing answers for tabled subgoals in an appropriate data space, called the *table space*. Variant calls² to tabled subgoals are not re-evaluated against the program clauses, instead they are resolved by consuming the answers already stored in their table entries. During this process, as further new answers are found, they are stored in their tables and later returned to all variant calls.

¹To the best of our knowledge, Yap and XSB are the unique Prolog systems that combine multithreading with tabling support.

²Two terms are considered to be variant if they are the same up to variable renaming.

With these requirements, the design of the table space is critical to achieve an efficient implementation. Yap uses *tries* which is regarded as a very efficient way to implement the table space [8]. Tries are trees in which common prefixes are represented only once. The trie data structure provides complete discrimination for terms and permits look up and possibly insertion to be performed in a single pass through a term, hence resulting in a very efficient and compact data structure for term representation. Figure 1 shows the general table space organization for a tabled predicate in Yap.

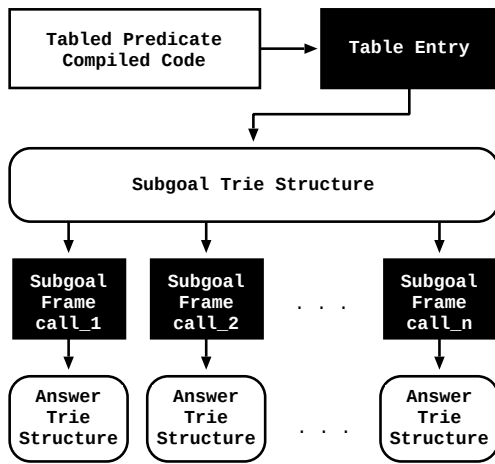


Figure 1: Yap's table space organization

At the entry point we have the *table entry* data structure. This structure is allocated when a tabled predicate is being compiled, so that a pointer to the table entry can be included in its compiled code. This guarantees that further calls to the predicate will access the table space starting from the same point. Below the table entry, we have the *subgoal trie structure*. Each different tabled subgoal call to the predicate at hand corresponds to a unique path through the subgoal trie structure, always starting from the table entry, passing by several subgoal trie data units, the *subgoal trie nodes*, and reaching a leaf data structure, the *subgoal frame*. The subgoal frame stores additional information about the subgoal and acts like an entry point to the *answer trie structure*. Each unique path through the answer trie data units, the *answer trie nodes*, corresponds to a different tabled answer to the entry subgoal.

2.2. Mode-Directed Tabling

In a traditional tabling system, all the arguments of a tabled subgoal call are considered when storing

answers into the table space. When a new answer is not a variant of any answer that is already in the table space, then it is always considered for insertion. Therefore, traditional tabling is very good for problems that require storing all answers. However, with dynamic programming, usually, the goal is to dynamically calculate optimal or selective answers as new results arrive. Writing dynamic programming algorithms can thus be a difficult task without further support.

Mode-directed tabling is an extension to the tabling technique that supports the definition of *modes* for specifying how answers are inserted into the table space. Within mode-directed tabling, tabled predicates are declared using statements of the form *'table p(m₁, ..., m_n)'*, where the *m_i*'s are *mode operators* for the arguments. The idea is to define the arguments to be considered for variant checking (the index arguments) and how variant answers should be tabled regarding the remaining arguments (the output arguments). In Yap, index arguments are represented with mode *index*, while arguments with modes *first*, *last*, *min*, *max*, *sum* and *all* represent output arguments [6]. After an answer is generated, the system tables the answer only if it is *preferable*, accordingly to the meaning of the output arguments, than some existing variant answer.

2.3. Multithreaded Tabling

Yap implements a SWI-Prolog compatible multithreading library [9]. Like in SWI-Prolog, Yap's threads have their own execution stacks and only share the code area where predicates, records, flags and other global non-backtrackable data are stored. For tabled evaluation, a thread views its tables as private but, at the engine level, we use a common table space, i.e., from the thread point of view, the tables are private but, from the implementation point of view, the tables are shared among all threads.

In previous work [5], he have proposed three designs for the common table space. This work uses the *Subgoal Sharing (SS)* design. In the SS design, the subgoal trie structure is shared among all threads and the leaf data structures representing each tabled subgoal call *C_i*, instead of pointing to a single subgoal frame, they point to a list of private subgoal frames, one per thread that is evaluating the call *C_i*. In the previous version of the SS design, threads would store the results of their

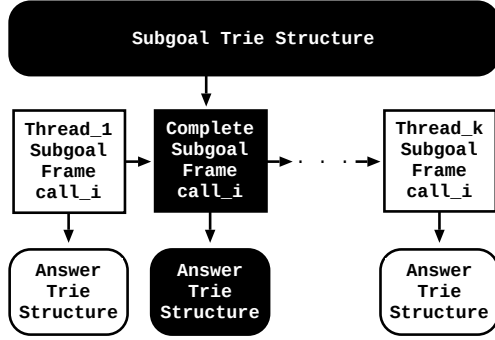


Figure 2: Subgoal sharing design

sub-computations in private structures and afterwards whenever they finished their execution, they would remove their private structures. As consequence no information about the results of sub-computations was shared among threads. In this work, we propose a new asynchronous version of the SS design, where the key idea is that a thread does not wait for the other threads to compute its sub-problem, but is able to use the result of the sub-problem, if another thread has already computed it. Thus, whenever a thread has to compute a sub-problem, it lookups if any other thread has already computed the result for that sub-problem. If it exists, then the thread uses the result, otherwise it computes the sub-problem itself in a private fashion (without concurrency) and afterwards, when the computation is completed it publishes the results, thus that they can be used by any other threads. Fig. 2 shows a small overview about the new asynchronous SS design. Concurrency among threads is restricted to the subgoal trie structure and to completed subgoal frames (black data structures in Fig. 2). All subgoal frames and answer tries are initially private to a thread (white data structures in Fig. 2). Later, when the first subgoal frame is completed, i.e., when we have found the full set of answers for it, it is marked as completed and put in the beginning of the list of private subgoal frames (configuration shown in Fig. 2). Following calls made by other threads to this subgoal call simply consume the answers from the completed subgoal frame, thus avoiding recomputing the subgoal call at hand. By sharing only completed answer tries, we avoid the problem of dealing with concurrent updates to the answer tries, the problem of managing the different set of answers that each thread has found, and, more importantly, the prob-

lem of dealing with concurrent deletes, as in the case of using mode-directed tabling.

3. Related Work

In this section, we briefly introduce XSB’s alternative approach to multithreaded tabling. Remember that, to the best of our knowledge, Yap and XSB are the unique Prolog systems that combine multithreading with tabling support.

The XSB system supports two types of models for the combination of multithreading with tabling: *private tables* and *shared tables* [4, 10]. On the private tables model, each thread keeps its own copy of the table space. On one hand, this avoids concurrency over the tables but, on the other hand, the same table can be repeatedly computed by several threads, thus increasing the memory usage necessary to represent the table space. Moreover, since no information is shared between threads, a thread cannot reuse the results being tabled by another thread to reduce execution time in a concurrent environment.

For shared tables, the running threads store only once the same table, even if multiple threads use it. This model can be viewed as a variation of the *table-parallelism* proposal [11], where a tabled computation can be decomposed into a set of smaller sub-computations, each being performed by a different thread. Each sub-computation is computed independently by the first thread calling it, the *generator thread*, and each generator is the sole responsible for fully exploiting and obtaining the complete set of answers for a sub-computation. Variant sub-computations by other threads are resolved by consuming the answers stored by the generator thread. When a set of sub-computations being computed by different threads is mutually dependent, then a *usurpation operation* [12] synchronizes threads and a single thread assumes the computation of all sub-computations, turning the remaining threads into consumer threads. This maintains the correctness of the table space in a concurrent environment, but has a major disadvantage: it restricts the potential of concurrency to non-mutually dependent sub-computations. In particular, as our experiments will show, this severely constraints the goal of scaling multithreaded dynamic programming problems.

4. The 0-1 Knapsack Problem

The Knapsack problem [13] is a well-known problem in combinatorial optimization that can be found in many domains such as logistics, manufacturing, finance or telecommunications. Given a set of items, each with a weight and a profit, the goal is to determine the number of items of each kind to include in a collection so that the total weight is equal or less than a given capacity and the total profit is as much as possible. The most common variant of the problem is the *0-1 Knapsack problem*, which restricts the number of copies of each kind of item to be zero or one. In what follows, we will focus on this variant. The 0-1 Knapsack problem can be formulated as follows. Given a set of items $i \in \{1, \dots, n\}$, each with a weight $w_i \in \mathbb{N}^*$ and a profit $p_i \in \mathbb{N}^*$, and a Knapsack with capacity $C \in \mathbb{N}^*$, the following formulas define the Knapsack problem (KS) and the restriction (KS_R) that avoids any trivial solution, by insuring that each item fits into the Knapsack and that the total weight of all items exceeds the Knapsack capacity.

$$KS = \begin{cases} \max \sum_{i=1}^n p_i \cdot x_i, \\ \text{s.t.} \sum_{i=1}^n w_i \cdot x_i \leq C, \\ x_i \in \{0, 1\}, i \in \{1, \dots, n\}. \end{cases}$$

$$KS_R = \begin{cases} \forall i \in \{1, \dots, n\}, w_i \leq C, \\ \sum_{i=1}^n w_i > C. \end{cases} \quad (1)$$

4.1. Top-Down Approach

We first introduce a standard top-down approach that solves the Knapsack problem using mode-directed tabling. Figure 3 shows our Yap's implementation adapted from [14] to include the dimension of profitability.

The table directive declares that predicate *ks/3* is to be tabled using modes (*index, index, max*), meaning that the third argument (the profit) should store only the maximal answers for the first two arguments (the index of the number of items being considered and Knapsack's capacity). The remaining part of the program implements a recursive top-down definition of the Knapsack problem. The first clause is the base case and defines that the empty set is a solution with profit 0. The second clause excludes the current item from the solution set and

```
% table declaration
:- table ks(index, index, max).
% base case
ks(0, C, 0).
% exclude case
ks(I, C, P) :-
    I > 0, ks_exc(I, C, P, 1).
% include case
ks(I, C, P) :-
    I > 0, ks_inc(I, C, P, 1).

% exclude N items starting from I
ks_exc(I, C, P, N) :-
    J is I - N, ks(J, C, P).
% include I and exclude the next
% N-1 items
ks_inc(I, C, P, N) :-
    item(I, Ci, Pi), Cj is C - Ci,
    Cj >= 0, J is I - N,
    ks(J, Cj, Pj), P is Pi + Pj.
```

Figure 3: A top-down approach for the Knapsack problem with mode-directed tabling

the third includes the current item in the solution if its inclusion does not overcome the current capacity of the Knapsack. For simplicity of integration with the parallel approach presented next, we are already using two auxiliary predicates, *ks_exc/4* and *ks_inc/4*, as a way to implement the exclude and include cases. These auxiliary predicates take an extra argument N (fourth argument) that represents the number of items to jump (or exclude) in the recursion procedure. Here, for the sequential version of the problem, N is always 1, i.e, we always move to the next item.

To parallelize top-down dynamic programming algorithms, we followed Stivala's et al. work [2] where a set of threads solve the entire program independently but with a randomized choice of the sub-problems. For the Knapsack problem, we have two sub-problems, the exclude and include cases. We can thus consider two alternative execution choices at each step: (i) exclude first and include next (as in the sequential version presented in Fig. 3), or (ii) include first and exclude next. The randomized choice of sub-problems results in the threads diverging to compute different sub-problems simultaneously while reusing the sub-problem's results computed in the meantime by the other threads. Since the number of overlapping sub-problem is usually

high in these kind of problems, it is expected that the available set of sub-problems to be computed will be evenly divided by the number of available threads resulting in less computation time required to reach the final result.

For the parallel version of the Knapsack problem, we have implemented two alternative versions. The first version simply follows Stivala's et al. original random approach. The second version extends the first one with an extra step where the computation is first moved forward using a random displacement of the number of items to be excluded and only then the computation is performed for the next item, as usual. By doing this, it is expected that the sub-problems closer to the base cases are computed earlier, meaning that their subgoal frames are also marked as completed earlier, which avoids re-computation when other threads call the same sub-problems. Figure 4 shows the implementation. The difference between the two versions is that the first version does not consider the first extra clause in the *aux_exc/4* and *aux_inc/4* auxiliary predicates.

4.2. Bottom-Up Approach

A straightforward method to solve the Knapsack problem bottom-up is for a fixed capacity c , to consider all 2^n possible subsets of the n items and choose the one that maximizes the profit. The recursive application of this algorithm to increasing capacities $c \in \{1, \dots, C\}$, yields a Knapsack of maximum profit for the given capacity C [15]. The bottom-up characteristic comes from the fact that, given a Knapsack with capacity c and using i items, $i < n$, the decision to include the next item j , $j = i + 1$, leads to two situations: (i) if j is not included, the Knapsack profit is unchanged; (ii) if j is included, the profit is the result of the maximum profit of the Knapsack with the same i items but with capacity $c - w_j$ (the capacity needed to include the weight w_j of item j) increased by p_j (the profit of the item j being included). The algorithm then decides whether or not to include an item based on which choice leads to maximum profit. Figure 5 shows the $KS[n, C]$ matrix. The rows define the items and the columns define the Knapsack capacities. The first column and row are filled with zeros, which are the initial profit for the Knapsacks with no items or no capacity.

The sequential version of the algorithm can be constructed row by row or column by column. The computation of each sub-problem $KS[j, c]$ considers the maximum profitability obtained between

```
% table declaration
:- table ks(index, index, max).
% base case
ks(0, C, 0).
% random choice
ks(I, C, P) :-
  I > 0, random(2, maxRandom, N),
  R is N mod 2,
  ( R == 0 ->
    aux_exc(I, C, P, N)
  ;
    aux_inc(I, C, P, N)).
% try exclude first and include next
aux_exc(I, C, P, N) :-
  ks_exc(I, C, P, N).
aux_exc(I, C, P, _) :-
  ks_exc(I, C, P, 1).
aux_exc(I, C, P, _) :-
  ks_inc(I, C, P, 1).
% try include first and exclude next
aux_inc(I, C, P, N) :-
  ks_inc(I, C, P, N).
aux_inc(I, C, P, _) :-
  ks_inc(I, C, P, 1).
aux_inc(I, C, P, _) :-
  ks_exc(I, C, P, 1).
```

Figure 4: A top-down parallel version of the Knapsack problem with mode-directed tabling

$KS[j - 1, c]$ and $KS[j - 1, c - w_{j-1}] + p_j$. When all sub-problems are computed, $KS[n, C]$ holds the best profitability for the full problem. Figure 6 shows Yap's implementation. For simplicity of presentation, we are omitting the predicate that implements the main loop used to recursively traverse the matrix and launch the computation for each sub-problem.

The table directive declares that predicate *ks/3* is to be tabled using standard tabling. Since here a sub-problem can be computed from the results of its sub-problems, standard tabling is enough and there is no need for mode-directed tabling. The first two clauses of *ks/3* are the base cases and define that the Knapsacks with no items or no capacity have profit 0. The third clause deals with the cases where an item's weight exceeds the Knapsack capacity and the fourth clause is the one that implements the main case discussed above.

Filling cells in subsequent rows requires accessing two cells from the previous row: one from the

	0	...	c-w _j	...	c	...	C
0	0	0	0	0	0	0	0
⋮	0						
i	0		⋮	⋮	⋮		
j	0				↓		
⋮	0						
n	0						KS [n, C]

Figure 5: Knapsack bottom-up matrix

```
% table declaration
:- table ks/3.
% base cases
ks(0, _, 0). ks(_, 0, 0).
% item I exceeds capacity C
ks(I, C, P) :-
  I > 0, item(I, Ci, Pi), Ci > C,
  J is I - 1, ks(J, C, P).
% item I fits in capacity C
ks(I, C, P) :-
  I > 0, item(I, Ci, Pi), Ci ≤ C,
  Cj is C - Ci, Cj ≥ 0,
  J is I - 1,
  ks(J, Cj, Pj), P1 is Pj + Pi,
  ks(J, C, P2), max(P1, P2, P).
```

Figure 6: A bottom-up approach for the Knapsack problem with standard tabling

same column and one from the column offset by the weight of the current item. Thus, computing a row i depends only on the sub-problems at row $i - 1$. A possible parallelization is, for each row, to divide the computation of the C columns between the available threads and then wait for all threads to complete in order to synchronize before computing the next row.

Here, since we want to take advantage of the built-in tabling mechanism, which is implicit and cannot be controlled by the user, we want to avoid this kind of synchronization between iterations. Hence, when a sub-problem in the previous row was not computed yet (i.e., marked as completed in one of the subgoal frames for the given call), instead of waiting for the corresponding result to be computed by another thread, the current thread starts also its computation and for that it can recursively call many other sub-problems not computed yet. Despite this can lead to redundant sub-computations,

it avoids synchronization. In fact, as we will see, this strategy showed to be very effective.

We next introduce our generic multithreaded scheduler used to load balancing the access to a set of concurrent tasks. We assume that the number of tasks is known before execution starts and that tasks are numbered incrementally starting at 1. For the Knapsack problem, we will consider that the number of tasks is the number of capacities $c \in \{1, \dots, C\}$ (alternatively, we could have considered the number of items $i \in \{1, \dots, n\}$). In a nutshell, the scheduler uses a user-level *mutex* to protect a concurrent *queue* that stores the indices of the available tasks. In fact, since tasks are numbered incrementally, the queue simply needs to store the index of the next available task. When a thread gets access to the queue of tasks, it picks a chunk of consecutive tasks and updates the queue's stored index accordingly. Figure 7 shows the Prolog code that implements the main execution loop of each thread.

```
% initialize mutex
:- mutex_create(queueLock).
% initialize queue
:- set_value(queueIndex, 0).

do_work(NumberOfTasks, ChunkSize) :-
  mutex_lock(queueLock),
  get_value(queueIndex, Current),
  ( Current = NumberOfTasks ->
    % terminate execution
    mutex_unlock(queueLock)
  );
  First is Current + 1,
  Last is Current + ChunkSize,
  set_value(queueIndex, Last),
  mutex_unlock(queueLock),
  compute_tasks(First, Last),
  % get more work
  do_work(NumberOfTasks, ChunkSize)
).
```

Figure 7: The generic execution loop of each thread for the bottom-up approach

The top declarations initialize the *queueLock* mutex and the *queueIndex* queue. The predicate *do_work/2* implements the main execution loop of each thread and is recursively executed until no more tasks exist in the queue. It receives two arguments: the total number of tasks in the problem (*NumberOfTasks*); and the chunk size to be

considered when retrieving tasks from the queue (*ChunkSize*). In each loop, a thread starts by gaining access to the mutex and then it checks the queue. If the queue is empty, case in which the test $Current = NumberOfTasks$ succeeds (In order to avoid low-level details which are not relevant to this work, the reader can assume that *NumberOfTasks* is a multiple of *ChunkSize*) the mutex is released and the thread terminates execution. Otherwise, the thread picks a new chunk of consecutive tasks and updates the queue's stored index accordingly. Variables *First* and *Last* define the lower and upper bounds of the chunk of tasks obtained. The tasks are then evaluated using the *compute_tasks/2* predicate, which calls the *ks/3* predicate for the set of Knapsack sub-problems associated with the task. After the *compute_tasks/2* finishes, the *do_work/2* predicate is called again to get more tasks from the queue. The process repeats until no more tasks exist.

5. The Longest Common Subsequence Problem

The problem of computing the length of the Longest Common Subsequence (LCS) is representative of a class of dynamic programming algorithms for string comparison that are based on getting a similarity degree. A good example is the sequence alignment, which is a fundamental technique for biologists to investigate the similarity between species. The LCS problem can be defined as follows. Given a finite set of symbols S and two sequences $U = \langle u_1, u_2, \dots, u_n \rangle$ and $V = \langle v_1, v_2, \dots, v_m \rangle$ such that $\forall_{i \in 1, \dots, n}, u_i \in S$ and $\forall_{i \in 1, \dots, m}, v_i \in S$, we say that U has a common subsequence with V of length k if there are indices $i_1, i_2, \dots, i_k, j_1, j_2, \dots, j_k : 1 \leq i_1 < i_2 < \dots < i_k \leq n$ and $1 \leq j_1 < j_2 < \dots < j_k \leq m$ such that $\forall_{l \in 1, \dots, k}, u_{i_l} = v_{j_l}$. The length k is considered to be the longest common subsequence if it is maximal.

5.1. Top-Down Approach

We next introduce a standard top-down approach that solves the LCS problem using mode-directed tabling. Figure 8 shows Yap's implementation adapted from [14].

The first two clauses of *lcs/3* are the base cases defining that for empty sequences the LCS (third argument) is 0. The third clause deals with the cases where the current symbols in both sequences

```
% table declaration
:- table lcs(index, index, max).
% base cases
lcs(_, 0, 0). lcs(0, _, 0).
% matched case
lcs(Iu, Iv, L) :-
    Iu > 0, Iv > 0, symbol_u(Iu, S),
    symbol_v(Iv, S), Ju is Iu - 1,
    Jv is Iv - 1, lcs(Ju, Jv, Lj),
    L is Lj + 1.
% sequence U case
lcs(Iu, Iv, L) :-
    Iu > 0, Iv > 0,
    lcs_u(Iu, Iv, L, 1).
% sequence V case
lcs(Iu, Iv, L) :-
    Iu > 0, Iv > 0,
    lcs_v(Iu, Iv, L, 1).
% jump N symbols in sequence U
lcs_u(Iu, Iv, L, N) :-
    symbol_u(Iu, Su), symbol_v(Iv, Sv),
    Su \= Sv, Ju is Iu - N,
    lcs(Ju, Iv, L).
% jump N symbols in sequence V
lcs_v(Iu, Iv, L, N) :-
    symbol_u(Iu, Su), symbol_v(Iv, Sv),
    Su \= Sv, Jv is Iv - N,
    lcs(Iu, Jv, L).
```

Figure 8: A top-down approach for the LCS problem with mode-directed tabling

match (arguments I_u and I_v represent, respectively, the current indices in sequences U and V to be considered). The fourth and fifth clauses represent the opposite case, where the symbols do not match, and each clause moves one of the sequences to the next symbol (note that recursion is done in descending order until reaching index 0). Again, for simplicity of integration with the parallel approach presented next, we are already using two auxiliary predicates, *lcs_u/4* and *lcs_v/4*, as a way to implement the unmatched cases. As for the Knapsack problem, these two auxiliary predicates take an extra argument N (fourth argument) that represents the number of symbols to jump in the recursion procedure. For the sequential version of the problem, N is always 1, meaning that we always move to the next symbol.

Similarly to Knapsack's problem, to parallelize the LCS sequential top-down approach, we have implemented two alternative versions. The first ver-

sion follows Stivala's et al. original random approach. The second version extends the first one with an extra step where the computation is first moved forward using a random displacement of the number of symbols to jump and only then the computation is performed for the next symbol, as usual. Figure 9 shows the implementation. The difference between the two versions is that the first version does not consider the first extra clause in the *aux_u/4* and *aux_v/4* auxiliary predicates.

```
% table declaration
:- table lcs(index, index, max).
% base cases
lcs(_, 0, 0). lcs(0, _, 0).
% matched case
lcs(Iu, Iv, L) :-
    Iu > 0, Iv > 0,
    symbol_u(Iu, S), symbol_v(Iv, S),
    Ju is Iu - 1, Jv is Iv - 1,
    lcs(Ju, Jv, Lj), L is Lj + 1.
% random choice
lcs(Iu, Iv, L) :-
    Iu > 0, Iv > 0,
    random(2, maxRandom, N),
    R is N mod 2,
    ( R == 0 ->
        aux_u(Iu, Iv, L, N)
    ;
        aux_v(Iu, Iv, L, N) ).
% try sequence U first and V next
aux_u(Iu, Iv, L, N) :-
    lcs_u(Iu, Iv, L, N).
aux_u(Iu, Iv, L, _) :-
    lcs_u(Iu, Iv, L, 1).
aux_u(Iu, Iv, L, _) :-
    lcs_v(Iu, Iv, L, 1).
% try sequence V first and U next
aux_v(Iu, Iv, L, N) :-
    lcs_v(Iu, Iv, L, N).
aux_v(Iu, Iv, L, _) :-
    lcs_v(Iu, Iv, L, 1).
aux_v(Iu, Iv, L, _) :-
    lcs_u(Iu, Iv, L, 1).
```

Figure 9: A top-down parallel version of the LCS problem with mode-directed tabling

5.2. Bottom-Up Approach

We now introduce our bottom-up approach to the LCS problem, which is based on [15]. In a nutshell,

the bottom-up characteristic comes from the fact that, the maximum length of a common subsequence between two sequences U and V is: (i) if the initial symbols of both sequences match, then they are part of the longest common subsequence and the length of the longest common subsequence can be incremented by one; (ii) if the initial symbols do not match then two situations arise: the longest common subsequence may be obtained from U and V without the initial symbol or from V and U without the initial symbol. Since we want the longest subsequence, the maximum of these two must be selected. The following formula formalizes the LCS problem as described above:

$$LCS[j, l] = \begin{cases} LCS[j-1, l-1] + 1, & \text{if } u_j = v_l. \\ \max \{LCS[j, l-1], LCS[j-1, l]\}, & \text{otherwise.} \end{cases} \quad (2)$$

Figure 10 shows the LCS matrix that represents the bottom-up approach. The rows define the indices to be considered in sequence U and the columns define the indices in sequence V . The first column and the first row are filled with zeros, meaning that for empty sequences the LCS is 0. The sequential version of the algorithm can be constructed row by row or column by column, since the computation of each sub-problem $LCS[j, l]$ only depends on the sub-computations done for the preceding row and column. At the end, $LCS[n, m]$ holds the LCS for the problem.

	0	...	k	l	...	m
0	0	0	0	0	0	0
⋮	0					
i	0					
j	0					
⋮	0					
n	0					LCS [n, m]

Figure 10: LCS bottom-up matrix

Figure 11 shows Yap's implementation. Again, for simplicity of presentation, we are omitting the predicate that implements the main loop used to recursively traverse the matrix and launch the computation for each sub-problem.

The table directive declares that predicate *lcs/3*

```

% table declaration
:- table lcs/3.
% base cases
lcs(_, 0, 0). lcs(0, _, 0).
% matched case
lcs(Iu, Iv, L) :-
    Iu > 0, Iv > 0, symbol_u(Iu, S),
    symbol_v(Iv, S), Ju is Iu - 1,
    Jv is Iv - 1, lcs(Ju, Jv, Lj),
    L is Lj + 1.
% unmatched case
lcs(Iu, Iv, L) :-
    Iu > 0, Iv > 0, symbol_u(Iu, Su),
    symbol_v(Iv, Sv), Su \= Sv,
    Ju is Iu - 1, Jv is Iv - 1,
    lcs(Ju, Iv, L1), lcs(Iu, Jv, L2),
    max(L1, L2, L).

```

Figure 11: A bottom-up approach for the LCS problem with standard tabling

is to be tabled using standard tabling. The first two clauses of *lcs/3* are the base cases and the third and fourth clauses deal with the cases where the initial symbols of both sequences match and do not match, respectively.

Concerning the parallelization of the matrix, a possible approach is, for each row, divide the computation of the m columns between the available threads or, for each column, divide the computation of the n rows between the available threads. Here, we will follow the same approach as for the Knapsack problem and we will use the generic multithreaded scheduler that implements the thread execution loop presented in Fig. 7. The number of concurrent tasks to be considered is the size of sequence U (alternatively, we could have considered the size of sequence V) and the evaluation of the *compute_tasks/2* predicate calls the *lcs/3* predicate for the set of LCS sub-problems associated with a given task.

6. Performance Analysis

The environment for our experiments was a machine with 32-core AMD Opteron (tm) Processor 6274 @ 2.2 GHz with 32 GBytes of main memory and running the Linux kernel 3.8.3-1.fc17.x86_64 with Jemalloc 3.1.0 [16]. We used Yap Prolog, version 6.3.2, with the SS design and the memory allocator [17]. To put our results in perspective, we also experimented with XSB Prolog version 3.4.0,

using the shared tables model. To verify the correctness of the benchmarks we have confirmed that the final results of the computations were correct on both Prolog systems and on all of the strategies that we have implemented.

For the Knapsack problem, we fixed the number of items and capacity, respectively, 1600 and 3200. For the LCS problem, we used both sequences with a fixed size of 3200 symbols each. Then, for each problem we created three different datasets, D_{10} , D_{30} and D_{50} , meaning that the values for the weights/profits for the Knapsack problem and the symbols for LCS problem were randomly generated in an interval between 1 and 10%, 30% and 50% of the total number of items/symbols, respectively. For the top-down approaches, we only experimented with Yap since XSB does not support mode-directed tabling. We tested Yap with Stivala's et al. original version (YAP_{TD_1}) and with the extended version using the extra random displacement clause (YAP_{TD_2}). For both Knapsack and LCS problems, we used a *maxRandom* value corresponding to 10% of the total number of items/symbols in the problem. For the bottom-up approaches, we experimented with Yap (YAP_{BU}) and XSB (XSB_{BU}) and we used a *ChunkSize* value of 5.

Table 1 and Table 2 show the results obtained, respectively, for the Knapsack and LCS problems for both top-down and bottom-up approaches using the Yap and XSB Prolog systems. In particular, both tables show the execution time, in milliseconds, for one thread (column **Time** (T_1)) and the corresponding speedup, for the execution with 8, 16, 24 and 32 threads (columns **Speedup** (T_1/T_p)). Results in bold highlight the best speedup obtained for each system/dataset configuration. All results are the average of 10 runs.

Analyzing the general picture of both tables, one can observe that for both problems, the top-down YAP_{TD_2} and bottom-up YAP_{BU} approaches have the best results with excellent speedups for 8, 16, 24 and 32 threads. In particular, for 32 threads, they obtain speedups around 19 and 18, respectively, for the Knapsack and LCS problems. The results for the top-down YAP_{TD_1} approach are not so interesting, regardless of the fact that it can slightly scale for the Knapsack problem up to 16 threads.

Regarding the base execution times with one thread, YAP_{TD_2} clearly pays the cost of the extra clause with an average execution time around 1.4 to 1.5 times slower than YAP_{TD_1} and YAP_{BU} . As

Table 1: Execution time, in milliseconds, for one thread and corresponding speedup, for the execution with 8, 16, 24 and 32 threads, for the top-down and bottom-up approaches of the **Knapsack** problem using the Yap and XSB Prolog systems

System/Dataset		# Threads (p)				
		Time (T_1)	Speedup (T_1/T_p)			
		1	8	16	24	32
Top-Down Approaches						
YAP_{TD1}	D₁₀	18,319	1.96	2.10	2.01	1.89
	D₃₀	17,664	3.41	3.96	3.83	3.62
	D₅₀	17,828	4.72	6.12	6.21	6.07
YAP_{TD2}	D₁₀	23,816	6.78	11.95	14.81	16.79
	D₃₀	25,049	7.39	13.63	16.85	19.35
	D₅₀	24,866	7.38	13.67	16.78	19.23
Bottom-Up Approaches						
YAP_{BU}	D₁₀	17,054	7.25	13.32	17.12	19.60
	D₃₀	17,005	7.22	13.47	17.29	19.64
	D₅₀	16,550	7.16	13.29	17.04	19.60
XSB_{BU}	D₁₀	37,338	0.81	0.79	0.73	0.54
	D₃₀	38,245	0.82	0.75	0.75	0.56
	D₅₀	39,100	0.82	0.79	0.73	0.54

a complementary information, note that the execution time with one thread for the clean top-down approach without randomization for the D_{10} , D_{30} and D_{50} datasets is, respectively, 11.844, 11.706 and 12.151 for the Knapsack problem and 23.722, 23.471 and 23.374 for the LCS problem. In this regard, the execution times for YAP_{TD2} and YAP_{BU} are quite different although their similar average speedups. For example, consider the D_{50} dataset of the Knapsack problem with 32 threads, while the speedup 19.23 of YAP_{TD2} corresponds to an execution time of 1.293 seconds, the speedup 19.60 of YAP_{BU} corresponds to only 0.844 seconds. Similarly for the LCS problem, if considering the D_{50} dataset with 32 threads, while the speedup 18.33 of YAP_{TD2} corresponds to 2.325 seconds, the speedup 18.07 of the YAP_{BU} corresponds to only 1.500 seconds.

Regarding the comparison with XSB, Yap’s results clearly outperform those of XSB. For the execution time with one thread, XSB shows higher times than all Yap’s approaches (around two times the execution times for YAP_{TD1} and YAP_{BU}). For the parallel execution of the Knapsack problem, XSB shows no speedups and for the parallel execution of the LCS problem we have no results available (*n.a.*) since we got *segmentation fault* execution errors. From our point of view, these results that were obtained for the XSB’s shared tables model were a consequence of the *usurpation operation*, since it constrained the concurrency to non-mutual

dependent sub-computations, which consequently constrained the full potentiality of the parallelism. As the parallel algorithms implemented in this performance analysis for the Knapsack and LCS problems, create mutual dependent sub-computations which are executed in different threads, the XSB is actually unable to execute the benchmarks in a parallel fashion. By other works, even if we launch an arbitrary large number of threads on those benchmarks, the system would only use one thread to evaluate all the computations.

7. Conclusions and Further Work

Starting from two well-known dynamic programming problems, the Knapsack and the Longest Common Subsequence (LCS) problems, we have discussed how we were able to scale their execution by taking advantage of the multithreaded tabling engine of the Yap Prolog system. We have presented multithreaded tabled top-down and bottom-up approaches using, respectively, Yap’s mode-directed tabling support and Yap’s standard tabling support. Our experiments, on a 32-core AMD machine, showed that using either top-down or bottom-up techniques, we were able to scale the execution of both problems by taking advantage of the state-of-the-art multithreaded tabling engine of the Yap Prolog system. Further work will include studying other dynamic programming

Table 2: Execution time, in milliseconds, for one thread and corresponding speedup, for the execution with 8, 16, 24 and 32 threads, for the top-down and bottom-up approaches of the **LCS** problem using the Yap and XSB Prolog systems

System/Dataset		# Threads (p)				
		Time (T_1)	Speedup (T_1/T_p)			
		1	8	16	24	32
Top-Down Approaches						
YAP _{<i>TD1</i>}	D ₁₀	30,708	1.53	1.45	1.40	1.29
	D ₃₀	30,817	1.53	1.46	1.38	1.28
	D ₅₀	30,707	1.52	1.44	1.39	1.27
YAP _{<i>TD2</i>}	D ₁₀	42,556	7.25	13.13	16.26	18.32
	D ₃₀	42,511	7.21	13.24	16.19	18.34
	D ₅₀	42,631	7.21	13.15	16.27	18.33
Bottom-Up Approaches						
YAP _{<i>BU</i>}	D ₁₀	27,253	6.97	10.78	14.88	17.91
	D ₃₀	27,045	6.88	11.20	14.74	17.92
	D ₅₀	27,102	6.97	11.91	14.51	18.07
XSB _{<i>BU</i>}	D ₁₀	68,255	n.a.	n.a.	n.a.	n.a.
	D ₃₀	69,700	n.a.	n.a.	n.a.	n.a.
	D ₅₀	70,100	n.a.	n.a.	n.a.	n.a.

problems and explore the impact of applying multithreaded tabling to other application domains.

Acknowledgments

This work is partially funded by the ERDF (European Regional Development Fund) through the COMPETE Programme and by FCT (Portuguese Foundation for Science and Technology) within project SIBILA (NORTE-07-0124-FEDER-000059). Miguel Areias is funded by the FCT grant SFRH/BD/69673/2010.

- [1] R. Bellman, Dynamic Programming, Princeton University Press, 1957.
- [2] A. Stivala, P. Stuckey, M. G. de la Banda, M. Hermenegildo, A. Wirth, Lock-Free Parallel Dynamic Programming, Journal of Parallel and Distributed Computing 70 (8) (2010) 839–848.
- [3] W. Chen, D. S. Warren, Tabled Evaluation with Delaying for General Logic Programs, Journal of the ACM 43 (1) (1996) 20–74.
- [4] R. Marques, T. Swift, Concurrent and Local Evaluation of Normal Programs, in: International Conference on Logic Programming, no. 5366 in LNCS, Springer, 2008, pp. 206–222.
- [5] M. Areias, R. Rocha, Towards Multi-Threaded Local Tabling Using a Common Table Space, Journal of Theory and Practice of Logic Programming, International Conference on Logic Programming, Special Issue 12 (4 & 5) (2012) 427–443.
- [6] J. Santos, R. Rocha, On the Efficient Implementation of Mode-Directed Tabling, in: International Symposium on Practical Aspects of Declarative Languages, no. 7752 in LNCS, Springer, 2013, pp. 141–156.
- [7] V. Santos Costa, R. Rocha, L. Damas, The YAP Prolog System, Journal of Theory and Practice of Logic Programming 12 (1 & 2) (2012) 5–34.
- [8] I. V. Ramakrishnan, P. Rao, K. Sagonas, T. Swift, D. S. Warren, Efficient Access Mechanisms for Tabled Logic Programs, Journal of Logic Programming 38 (1) (1999) 31–54.
- [9] J. Wielemaker, Native Preemptive Threads in SWI-Prolog, in: International Conference on Logic Programming, no. 2916 in LNCS, Springer, 2003, pp. 331–345.
- [10] T. Swift, D. S. Warren, XSB: Extending Prolog with Tabled Logic Programming, Theory and Practice of Logic Programming 12 (1 & 2) (2012) 157–187.
- [11] J. Freire, R. Hu, T. Swift, D. S. Warren, Exploiting Parallelism in Tabled Evaluations, in: International Symposium on Programming Languages: Implementations, Logics and Programs, no. 982 in LNCS, Springer, 1995, pp. 115–132.
- [12] R. Marques, T. Swift, J. C. Cunha, A Simple and Efficient Implementation of Concurrent Local Tabling, in: International Symposium on Practical Aspects of Declarative Languages, no. 5937 in LNCS, Springer, 2010, pp. 264–278.
- [13] S. Martello, P. Toth, Knapsack Problems: Algorithms and Computer Implementations, John Wiley and Sons, 1990.
- [14] H.-F. Guo, G. Gupta, Simplifying Dynamic Programming via Mode-directed Tabling, Software Practice and Experience 38 (1) (2008) 75–94.
- [15] V. Kumar, Introduction to Parallel Computing, 2nd Edition, Addison-Wesley, 2002.
- [16] J. Evans, A Scalable Concurrent malloc(3) Implementation for FreeBSD, in: The Technical BSD Conference, 2006.
- [17] M. Areias, R. Rocha, An Efficient and Scalable Memory Allocator for Multithreaded Tabled Evaluation of Logic Programs, in: International Conference on Parallel and Distributed Systems, IEEE Computer Society, 2012, pp. 636–643.