



An Analysis of Budgeted Parallel Search on Conditional Galton–Watson Trees

David Avis^{1,2} · Luc Devroye³

Received: 13 April 2017 / Accepted: 22 October 2019 / Published online: 14 November 2019
© Springer Science+Business Media, LLC, part of Springer Nature 2019

Abstract

Recently Avis and Jordan have demonstrated the efficiency of a simple technique called budgeting for the parallelization of a number of tree search algorithms. The idea is to limit the amount of work that a processor performs before it terminates its search and returns any unexplored nodes to a master process. This limit is set by a critical budget parameter which determines the overhead of the process. In this paper we study the behaviour of the budget parameter on conditional Galton–Watson trees obtaining asymptotically tight bounds on this overhead. We present empirical results to show that this bound is surprisingly accurate in practice.

Keywords Parallel tree search · Random Galton–Watson tree · Probabilistic analysis of algorithms · Branching process

1 Introduction

The majority of algorithms have been designed, analyzed and implemented to run on single core processors. While multicore hardware is now ubiquitous these algorithms profit little, if any, from the additional processing power. On the other hand, the study of parallel algorithms also has a long history. Issues involved are complex and include architecture design, communication, data sharing, interrupts, deadlocks, load balancing, and the distinction between shared memory and distributed computing. For a comprehensive reference, see Mattson [23]. This complexity presents serious

✉ David Avis
avis@cs.mcgill.ca

Luc Devroye
luc@cs.mcgill.ca

¹ School of Informatics, Kyoto University, Kyoto, Japan

² School of Computer Science and GERAD, McGill University, Montreal, Quebec, Canada

³ School of Computer Science, McGill University, Montreal, Quebec, Canada

challenges to developing theoretical analyses to explain successful empirical results. Indeed many recent highly publicized computational success stories, such as Computer Go, Machine Learning and Deep Learning for AI, make massive use of parallel computation but have little or no theoretical foundation for the underlying algorithms.

In a recent paper, McCreesh and Prosser [24] attempt to fill this gap in the context of parallel branch and bound algorithms for the maximum clique problem. They show that the interaction of the search order and the most likely location of solutions is often the dominant consideration in load balancing for these types of problems. This leads to a new work-splitting technique which gives good scalability. In their conclusions they recall the following prescient remark of Hooker [13]:

Based on one's insight into an algorithm, for instance, one might expect good performance to depend on a certain characteristic. How to find out? Design a controlled experiment that checks how the presence or absence of this characteristic affects performance. Even better, build an explanatory mathematical model that captures the insight, as is done routinely in other empirical sciences, and deduce from it precise consequences that can be put to the test.

The purpose of this paper is to make a modest additional contribution along these lines by giving a theoretical analysis of the overhead produced by a simple and effective scaleable parallelization technique, described by Avis and Jordan [6,7], that is applicable to a certain class of tree search algorithms. This class includes reverse search which has been used for a large number of combinatorial and geometric enumeration problems. Trees in this class must be definable by an oracle which, given any node in the tree, gives the children of that node (if any) in an arbitrary but fixed order. Traversals of trees in this class can be performed in parallel by assigning independent subtrees to different processors. When the underlying tree is unbalanced, the main problem again becomes one of load balancing. The tree traversal problem differs significantly from branch and bound, where the goal is to explore as little of the tree as possible until a desired node is found. The bounding step removes subtrees from consideration and this step depends critically on what has already been discovered. Hence the order of traversal is crucial and the number of nodes evaluated varies dramatically depending on this order, as reported by McCreesh and Prosser [24]. Sharing of information is critical to the success of parallelization in branch and bound. These issues do not occur when the entire tree must be traversed, hence simpler parallelization techniques are possible.

Avis and Jordan approach the load balancing problem by using a simple technique called *budgeting*. A search of a subtree of the original tree terminates after a given number of nodes, called the *budget*, have been generated and returns all unexplored nodes to a master process which stores them on a job list. In a multiprocessor setting, the master process assigns nodes from the job list to available processors, called workers, that proceed in parallel, repopulating the list whenever their budget is reached. In their 2015 paper they apply this technique to parallelize *lrs*, a reverse search based code for the vertex/facet enumeration problem for convex polyhedra. Experimental results showed near linear speedups when using up to several hundred processors. In the second paper the authors describe a generic common wrapper that can be used with

a wide variety of legacy codes. They gave applications to other reverse search codes and, more recently, to satisfiability testers, again obtaining substantial speedups. The resulting software, called *mts*,¹ is freely available. Budgeting may be seen as a very simple form of *job stealing*, as described by Blumofe and Leiserson [9]. However in job stealing individual workers maintain their own job lists and poll each other as necessary to obtain additional work, requiring communications between workers, interrupt handling, deadlock avoidance and so on.

Budgeting has several very practical advantages. Firstly it does not require communication between workers and they do not need to process interrupts. This means that parallelism does not have to be added to the existing legacy code. Secondly large subtrees are automatically broken up, since they exceed the budget, whereas small subtrees are not, since they do not. Thirdly the master can dynamically control the size of the job list by varying the size of the budget when assigning work. Finally it is easy to checkpoint the process since all workers return to the master regularly due to the budget restriction. It is simply a matter of waiting for all workers to return and then outputting the remaining job list. On the other hand budgeting introduces two competing forms of overhead. One is the cost of restarting jobs on the master's list, which is directly proportional to the total size of the list. The other is the cost of idle workers if the list becomes empty. It is therefore very important to determine how the budget parameter affects the size of the job list.

The goal of this paper is to analyze the budgeting method when it is applied to critical Galton–Watson trees. These trees provide a challenge to efficient parallel search since they are very unbalanced: it is well-known, see, e.g., Flajolet and Odlyzko [12], that the height of such trees is of the order or the square root of their size. For this class of trees we study the critical budget parameter and how it determines the size and evolution of the job list. In Sect. 2 we formally introduce budgeting and give an explicit example of how it can be implemented in a depth first search setting. Next, in Sect. 3 we introduce Galton–Watson trees and observe that they include quite a large number of other classes of random trees. We then state our main result, a tight estimate on the overhead produced by budgeting for traversing trees in our family. In Sect. 4 we describe a random walk which is closely related to the evolution of the list of unexplored nodes produced by budgeting. This is followed in Sect. 5 by a proof of our main result. Finally in Sect. 6 we make use of the *mts* framework to empirically test our estimates on a sample of large Galton–Watson trees. We also give a discussion of how our main result can be used in a multiprocessor setting in order to choose an efficient budget parameter.

2 Budgeted Tree Search

We are given a tree T by its root, an upper bound Δ on the number of children of any node, and an adjacency oracle $Adj(v, j)$. The adjacency oracle takes a tree node v and an index $j \in \{0, 1, 2, \dots, \Delta - 1\}$. For each value of j in its range, $Adj(v, j)$ either returns *null* or gives a child of v . Each child is given exactly once. Our goal is to

¹ <http://cgm.cs.mcgill.ca/~avis/doc/tutorial.html>.

visit all of the nodes in T . For concreteness in this section we will describe depth first search (DFS) although the budgeting technique can be applied to breadth first search or other search strategies.

A budgeted DFS is initiated from tree node *start_vertex*, initially set to be the root, and proceeds in a depth first fashion outputting each node as it is generated. We use two stacks, *stack_v* to store vertices and *stack_j* to store indices. We also specify an integer budget parameter $b \geq 1$ which causes the tree search to be terminated when b nodes have been generated. At this point the most recently generated vertex v and all of the unexplored siblings along the return path from v to *start_vertex* are returned with the flag *unexplored* = *true*. All previously generated nodes were output with the flag *unexplored* = *false*. The pseudo-code is shown as Algorithm 1. Note that that *start_vertex* is not output so this should be done in the calling program. We exploit this property later.

Algorithm 1 Budgeted depth first search

```

1: procedure BDFS(start_vertex,  $\Delta$ , Adj,  $b$ )
2:    $j \leftarrow 0$    $v \leftarrow \text{start\_vertex}$    $\text{count} \leftarrow 0$    $\text{depth} \leftarrow 0$ 
3:   repeat
4:     unexplored  $\leftarrow$  false
5:     while  $j < \Delta$  and unexplored = false do
6:        $j \leftarrow j + 1$ 
7:       push(stack_v,  $v$ )
8:       push(stack_j,  $j$ )
9:        $v \leftarrow \text{Adj}(v, j)$ 
10:       $\text{depth} \leftarrow \text{depth} + 1$    $\text{count} \leftarrow \text{count} + 1$ 
11:      if  $\text{count} \geq b$  then ▷ budget is exhausted
12:        unexplored  $\leftarrow$  true
13:      end if
14:      output ( $v$ , unexplored)
15:      if  $\text{count} < b$  then ▷ continue down tree
16:         $j \leftarrow 0$ 
17:      end if
18:    end while
19:    if  $\text{depth} > 0$  then ▷ backtrack step
20:       $v \leftarrow \text{pop}(\text{stack\_v})$ 
21:       $j \leftarrow \text{pop}(\text{stack\_j})$ 
22:       $\text{depth} \leftarrow \text{depth} - 1$ 
23:    end if
24:  until  $\text{depth} = 0$  and  $j = \Delta$ 
25: end procedure

```

Consider the tree in Fig. 1 which has 25 nodes, $\Delta = 5$ and is rooted at vertex 0. For convenience the nodes are numbered 0, 1, ..., 24 in depth first search order but this is in no way essential. If we set the parameter b to be 25 or greater, all nodes are output in this order with *unexplored* = *false*. Suppose we set $b = 13$. Firstly nodes 1, ..., 12 are output in order with *unexplored* = *false*. Then nodes 13, 15, 16, 18, 22 are output with *unexplored* = *true*. On the other hand, if we set $b = 8$ then nodes 1, 2, ..., 7 are output with *unexplored* = *false* and nodes 8, 9, 10, 11, 15, 16, 18, 22 are output with *unexplored* = *true*.

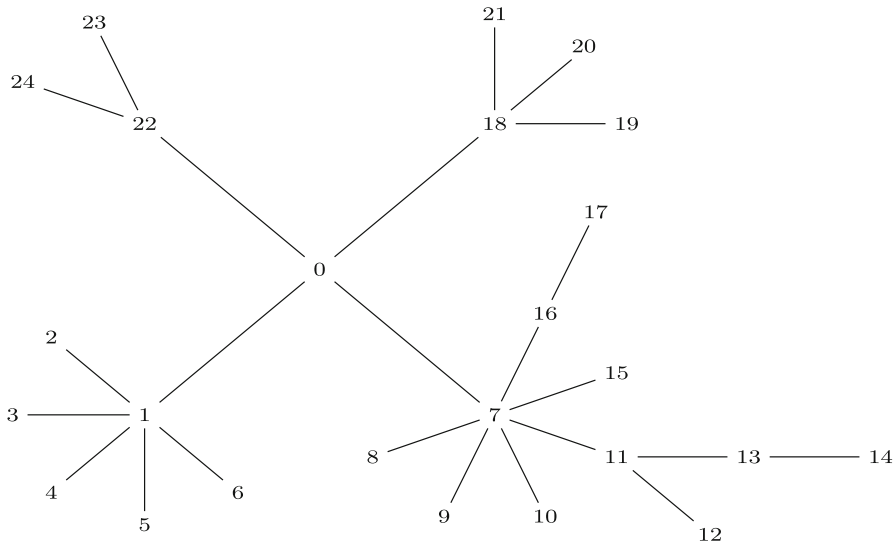


Fig. 1 Tree rooted at vertex 0 with 25 nodes and $\Delta = 6$

We will store the unexplored nodes in a list L which, for DFS, would be implemented as a stack. To complete the DFS of the tree we search the subtrees rooted by nodes in L using Algorithm 1 repeatedly. Note that each vertex in L now becomes a *start_vertex* and if a budget constraint b is applied additional nodes may be added to L . If L is empty when Algorithm 1 terminates the tree has been completely output. Since the root nodes are not output we will obtain a duplicate free list of all the nodes of the tree by concatenating outputs. Note that if the order of output is not important, the subtrees rooted by nodes in L can be explored in any order and in parallel. Furthermore, if $b = 1$ then all nodes in T will be returned to L . Alternatively, if $b > |T|$ then L remains empty as the tree is explored with a single call to Algorithm 1.

A successful parallelization of Algorithm 1 depends on controlling the job list L . On one hand, if the budget parameter b is too small, then L tends to become large causing considerable overhead in restart costs. If b is too large, then L becomes smaller, reducing overhead, but may become empty before the entire tree has been searched. This causes processors to become idle and reduces the speedup that can be obtained. We are interested in how L evolves over time.

Let L_b be the job list generated by repeatedly applying Algorithm 1 to a tree T for a given budget b and let R_n denote the total number of nodes added to the list during a complete search of T . We can observe the size of L_b each time Algorithm 1 is called. For example, again referring to Fig. 1, Algorithm 1 is called a total of 6 times if we keep $b = 13$ throughout. $|L_{13}| = 5$ and the size of L_{13} is respectively 0,4,3,2,1,0 for these 6 times.

For a larger example consider Fig. 2. This is based on visiting all the nodes of the tree T_{10} which has $\Delta = 10$ and 264,933,315 nodes.² Its construction is described in

² All computational results in the paper were obtained on mai20 at Kyoto University: 2x Xeon E5-2690 (10-core 3.0GHz), 20 cores, 128GB memory, 3TB hard drive.

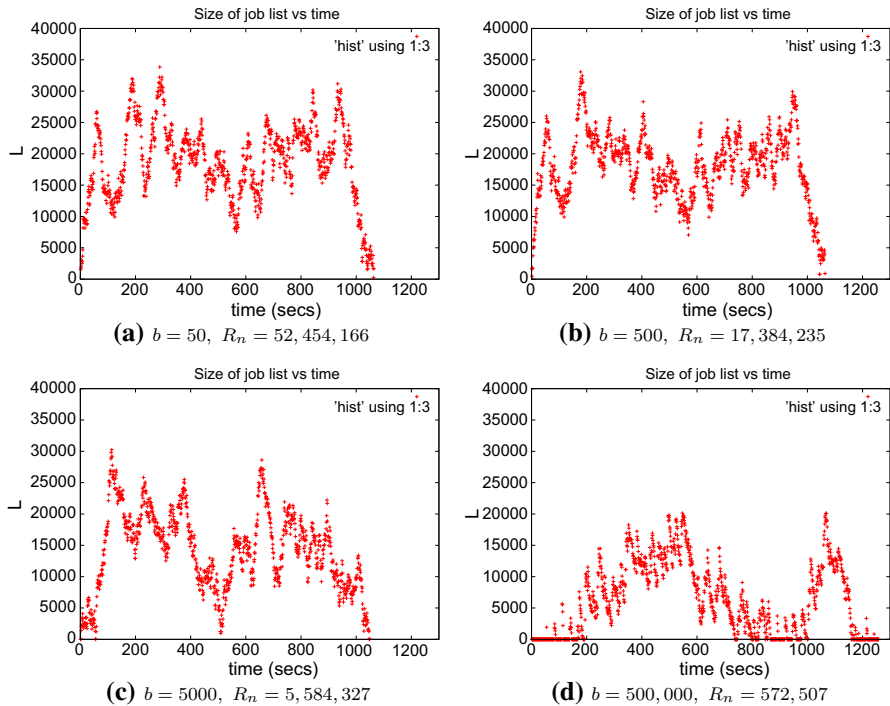


Fig. 2 Evolution of the job list for T_{10} : budget = b , total restarts = R_n , *mai20*, 16 cores

detail in Sect. 6. The plots were obtained by running Algorithm 1 in parallel on 16 cores using *mts* version 0.1 and different budget sizes b .

The first thing to observe is that, as expected, the number R_n of jobs returned to L is inversely proportional to b . By comparing the results for $b = 50, 5000, 500,000$ we see that as the budget scales up by a factor of 100, the size of R_n scales down by roughly a factor of 10, the square root of the scale factor. In the following sections we show that this is typical behaviour for Galton–Watson trees, of which T_{10} is a sample.

For $b = 50$ and $b = 500$ we have the desired behaviour for L : it quickly reaches a reasonable size (compared to the total tree size) and remains non-empty until the end of the run, implying all processors are busy. For $b = 5000$ the job list comes dangerously close to being empty before the run ends and for the extreme case $b = 500,000$ it is empty for several time intervals. All things being equal a small budget would seem optimum. However in practical applications, such as vertex enumeration, the startup cost for a job taken from L can be very significant. A budget small enough to keep a non-empty job list yet large enough to prevent too many restarts leads to optimum performance.

Before proceeding we note two important points. Firstly the total size of the job list is independent of the number of workers assigned. Indeed a single worker repeatedly executing Algorithm 1 produces the same number of jobs as a thousand workers in parallel. Secondly, if the job list stays well populated for the duration of the run then workers are never idle. If the jobs on the list are independent, as they are in a Galton–

Watson tree, then each worker behaves in the same way in probability. It is therefore sufficient to study the behaviour of a single worker applying Algorithm 1 repeatedly until the job list is empty.

3 The Probabilistic Model and Our Main Result

A Galton–Watson (or Galton–Watson–Bienaymé) tree (see [5]) is a rooted random ordered tree. Each node independently generates a random number of children drawn from a fixed offspring distribution ξ . The distribution of ξ defines the distribution of T , a random Galton–Watson tree. In what follows, we are mainly interested in critical Galton–Watson trees, i.e., those having $E\{\xi = 1\}$, and $P\{\xi = 1\} < 1$. In addition, we assume that the variance of ξ is finite (and hence, nonzero).

Moon [26] and Meir and Moon [25] defined the *simply generated trees* as ordered labelled trees of size n that are all equally likely given a certain pattern of labeling for each node of a given degree. The most important examples include the Catalan trees (equiprobable binary trees), the equiprobable k -ary trees, equiprobable unary-binary trees (ordered trees with up to two children), random Motzkin trees, random planted plane trees (equiprobable ordered trees of unlimited degrees) and Cayley trees (equiprobable unordered rooted trees). It turns out that all these trees can be represented as critical Galton–Watson trees conditional on their size, n , a fact first pointed out by Kennedy [16], and further developed by Kolchin [18,19] and others. So, let T_n be a Galton–Watson tree conditional on its size being n . For example, when ξ is 0 or 2 with probability 1/4, and 1 with probability 1/2, we obtain the uniform binary (Catalan) tree. Uniformly random full binary trees are obtained by setting $P\{\xi = 0\} = P\{\xi = 2\} = 1/2$. A uniformly random k -ary tree has its offspring distributed as a binomial $(k, 1/k)$ random variable. A uniform planted plane tree is obtained for the geometric law $P\{\xi = i\} = 1/2^{i+1}$, $i \geq 0$. When ξ is Poisson of parameter 1, one obtains (the shape of) a random rooted labeled (or Cayley) tree. For ξ uniform on $\{0, 1, 2, \dots, k\}$, T_n is like a uniform ordered tree with maximal degree of k . All such trees can be dealt with at once in the Galton–Watson framework.

Recall that R_n is the number of restarts generated when Algorithm 1 is used repeatedly to explore a tree T with n nodes using a budget b . Our main result is:

Theorem 1 *If T_n is a Galton–Watson tree of size n determined by ξ , where $E\{\xi\} = 1$ and $0 < V\{\xi\} \stackrel{\text{def}}{=} \sigma^2 < \infty$, then*

$$\frac{R_n}{n} \rightarrow \frac{1}{\mu_b}$$

in probability as $n \rightarrow \infty$, where

$$\mu_b \stackrel{\text{def}}{=} E\{\min(|T|, b)\},$$

and T is an unconditional Galton–Watson tree for the same ξ . In addition, as $b \rightarrow \infty$,

$$\mu_b \sim \sqrt{\frac{8b}{\pi\sigma^2}}.$$

4 Preliminary Results: The Random Walk View

Let ξ be a random variable representing the number of children of the root in a critical Galton–Watson tree: $E\{\xi\} = 1$, $P\{\xi = 0\} > 0$. Set $X = \xi - 1$, and let $X_1, X_2, \dots$ be a sequence of i.i.d. random variables distributed as X . Define the partial sums $Q(0) = 1$, $Q(t) = 1 + \sum_{i=1}^t X_i$. There is a well-known depth first or preorder construction of a random Galton–Watson tree which lends itself well to the study of all properties of randomly selected nodes (see, e.g., Le Gall [20], or Aldous [1–3]). Nodes in an ordered tree can be encoded with a vector of child numbers. The root corresponds to the empty vector. Its children have encodings $1, 2, \dots, \xi$. More generally, the $i - th$ child of a node with label j is encoded (j, i) . A preorder listing of the nodes is nothing but a lexicographic listing of the node vectors.

We can traverse a random Galton–Watson tree by visiting nodes in preorder, starting at the root. To do so, a list V of nodes to be visited is kept, which is initially of size one (V just contains the root). When node u is visited, we consider ξ_u , the number of children of u , and remove u from V . Thus, V increases by $\xi_u - 1$. The next node in lexicographic order is taken from V , and the process continues until V is empty.

We denote by N the size of a random Galton–Watson tree. The size of V after t nodes have been processed is precisely the Q process introduced above, $Q(t)$. Thus, $Q(0) = 1$, and

$$Q(t) = 1 + X_1 + \dots + X_t,$$

where $X_t = \xi_t - 1$, and indexing of the nodes is by their lexicographic rank. An example of the Q process derived from a tree is shown in Fig. 3. The path in the grid from vertex $(0,1)$ to vertex $(9,0)$ is derived from a preorder traversal of the tree to the right. Except for vertex $(9,0)$, each vertex of the path is of the form $(x, Q(x))$ where x is the tree node to be explored next and $Q(x)$ is the number of unexplored tree nodes we have discovered so far. We start at $(0,1)$ as V contains only the root 0 and so V has size $Q(0) = 1$. Since the root has two children the list of unexplored nodes becomes $\{1, 8\}$ which has length 2 so we move to vertex $(1,2)$ in the grid. Tree node 1 has 3 children so the list becomes $\{2, 3, 7, 8\}$ with length 4, we move to $(2,4)$ and so on. Finally we explore node 8 which has no children, the unexplored list is empty and we terminate the Q process.

We have the identity

$$[N = n] = [Q(1) > 0, Q(2) > 0, \dots, Q(n-1) > 0, Q(n) = 0].$$

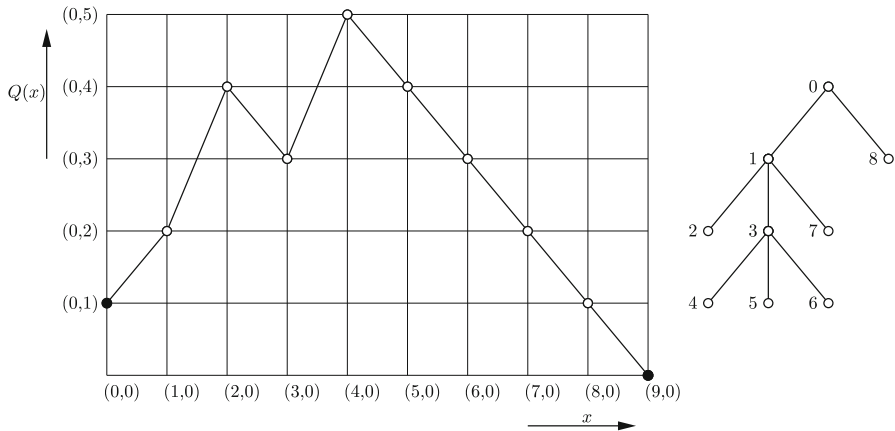


Fig. 3 The Q -process derived from a tree

The standard circular symmetry argument for random walks (see, e.g., Dwass [11]) shows that

$$\begin{aligned} P\{N = n\} &= P\{Q(0) = 1, Q(1) > 0, Q(2) > 0, \dots, Q(n-1) > 0, Q(n) = 0\} \\ &= \frac{1}{n} P\{Q(n) = 0\} \\ &= \frac{1}{n} P\{X_1 + \dots + X_n = -1\}. \end{aligned}$$

The asymptotics for this probability distribution are well-known, and will be given below. Let $0 < \sigma^2 \stackrel{\text{def}}{=} V\{\xi\} > 0$ (which implies $P\{\xi = 1\} < 1$), and let the span d be the greatest common divisor of all $i > 0$ for which $P\{\xi = i\} > 0$. A Galton–Watson tree with span d can only have sizes that are $1 \bmod d$. From Petrov [27, p. 197] or Kolchin [19, p. 16, p. 105], we recall that

$$\lim_{n \rightarrow \infty} \sup_{k=1 \bmod d} \left| \sigma \sqrt{n} P\{X_1 + \dots + X_n = k\} - \frac{d}{\sqrt{2\pi}} e^{-\frac{k^2}{2\sigma^2 n}} \right| = 0,$$

while, clearly, $P\{X_1 + \dots + X_n = k\} = 0$ if $k \not\equiv 1 \pmod{d}$. Thus, along $n \in \mathcal{N}$,

$$P\{N = n\} \sim \frac{d}{\sigma \sqrt{2\pi n^{3/2}}},$$

and, as $n \rightarrow \infty$,

$$P\{N \geq n\} \sim \sqrt{\frac{2}{\pi n \sigma^2}}.$$

Lemma 1 Under the conditions of Theorem 1, as $b \rightarrow \infty$,

$$E\{\min(N, b)\} \sim \sqrt{\frac{8b}{\pi\sigma^2}}.$$

Proof We have

$$\begin{aligned} E\{\min(N, b)\} &= \sum_{t=1}^{\infty} tP\{\min(N, b) = t\} \\ &= \sum_{t=1}^b tP\{N = t\} + bP\{N > b\} \\ &= \sum_{t=1}^b P\{X_1 + \cdots + X_t = -1\} + bP\{N > b\} \\ &= \sum_{t \leq b: t \in \mathcal{N}} \frac{d + o(1)}{\sigma\sqrt{2\pi t}} + (1 + o(1))\sqrt{\frac{2b}{\pi\sigma^2}} \\ &\sim \int_1^b \frac{1}{\sigma\sqrt{2\pi t}} dt + \sqrt{\frac{2b}{\pi\sigma^2}} \\ &\sim \sqrt{\frac{8b}{\pi\sigma^2}}. \end{aligned}$$

□

Consider the Q -process conditional on visiting (x, y) , with $0 < x < n$ and $y > 0$. Define

$$S(x, y) = \min\{t > x : Q(t) = Q(x) - 1 \mid Q(x) = y\} - x,$$

where we recall that $Q(t) = 1 + X_1 + \cdots + X_t$. In other words, $S(x, y)$ is the size of the subtree of node x in depth first search order given that $Q(x) = y$. Referring back to the example in Fig. 3 we see that $S(0, 1) = 9$, $S(1, 2) = 7$, $S(2, 4) = 1$, $S(3, 3) = 4$, $S(4, 5) = 1$, $S(5, 4) = 1$, $S(6, 3) = 1$, $S(7, 2) = 1$, $S(8, 1) = 1$.

Note that $S(x, y) = 1 \bmod d$. In addition, since $Q(0) = 1$, we must have $t + Q(t) = 1 \bmod d$, so that the only values (x, y) of interest to us have $x + y = 1 \bmod d$. We consider two positive constants a, c , with $a < 1/2$, $c < 1$, and let $\mathcal{S}(a, c)$ denote the collection of all random variables $S(x, y)$ with $an \leq x \leq (1 - a)n$, $c\sqrt{n} \leq y \leq (1/c)\sqrt{n}$. Lemma 2 below shows that within $\mathcal{S}(a, c)$, all random variables are close to N , the size of an unconditional Galton–Watson tree. This is an explicit version of a well-known theorem due to Aldous [1–3] (see also Janson [15]), which states that if T_n^* is the subtree of a uniform random node of T_n , then T_n^* tends in distribution to an unconditional Galton–Watson tree T as $n \rightarrow \infty$.

Lemma 2 For any fixed $0 < a < 1/2$, $0 < c < 1$ and integer $b > 0$,

$$\lim_{n \rightarrow \infty} \sup_{1 \leq i \leq b; i \equiv 1 \pmod d} \sup_{(x,y): \begin{cases} S(x,y) \in \mathcal{S}(a,c) \\ x+y \equiv 1 \pmod d \end{cases}} \left| \frac{P\{S(x,y) = i\}}{P\{N = i\}} - 1 \right| = 0.$$

Furthermore,

$$\lim_{n \rightarrow \infty} \sup_{(x,y): \begin{cases} S(x,y) \in \mathcal{S}(a,c) \\ x+y \equiv 1 \pmod d \end{cases}} |E\{\min(b, S(x,y))\} - E\{\min(b, N)\}| = 0.$$

Proof The last part follows from the first one, so we only consider $P\{S(x,y) = i\}$ for $1 \leq i \leq b, i \equiv 1 \pmod d$. We have, for $x + y \equiv 1 \pmod d, n \equiv 1 \pmod d$,

$$\begin{aligned} & P\{Q(x) = y, Q(t) > 0 \text{ for all } x < t < n, Q(n) = 0\} \\ &= P\{Q(x) = y\} \times \frac{1}{n-x} P\{X_1 + \cdots + X_{n-x} = -y\} \\ &= P\{Q(x) = y\} \times \frac{d + o(1)}{(n-x)\sqrt{2\pi(n-x)\sigma^2}} \exp\left(-\frac{y^2}{2(n-x)\sigma^2}\right). \end{aligned}$$

Also,

$$\begin{aligned} & P\{Q(x) = y, S(x,y) = i, Q(t) > 0 \text{ for all } x+i < t < n, Q(n) = 0\} \\ &= P\{Q(x) = y\} \times P\{N = i\} \times \frac{1}{n-x-i} P\{X_1 + \cdots + X_{n-x-i} = -(y-1)\} \\ &= P\{Q(x) = y\} \times P\{N = i\} \times \frac{d + o(1)}{(n-x-i)\sqrt{2\pi(n-x-i)\sigma^2}} \\ &\quad \exp\left(-\frac{(y-1)^2}{2(n-x-i)\sigma^2}\right). \end{aligned}$$

Thus, as

$$\begin{aligned} & P\{S(x,y) = i\} \\ &= \frac{P\{Q(x) = y, S(x,y) = i, Q(t) > 0 \text{ for all } x+i < t < n, Q(n) = 0\}}{P\{Q(x) = y, Q(t) > 0 \text{ for all } x < t < n, Q(n) = 0\}}, \end{aligned}$$

we have

$$\begin{aligned} & \frac{P\{S(x,y) = i\}}{P\{N = i\}} \\ &= (1 + o(1)) \left(\frac{n-x}{n-x-i}\right)^{3/2} \exp\left(\frac{y^2}{2(n-x)\sigma^2} - \frac{(y-1)^2}{2(n-x-i)\sigma^2}\right) \end{aligned}$$

$$\begin{aligned}
&= (1 + o(1)) \exp \left(\frac{y^2}{2(n-x)\sigma^2} - \frac{(y-1)^2}{2(n-x-i)\sigma^2} \right) \\
&= (1 + o(1)) \exp \left(\frac{y^2 - (y-1)^2}{2(n-x)\sigma^2} \right) \\
&= (1 + o(1)) \exp \left(\frac{2y-1}{2(n-x)\sigma^2} \right) \\
&= 1 + o(1),
\end{aligned}$$

where all the $o(1)$ terms are uniform over all i, x, y within the given ranges. $\square$

The Q process can be viewed as a continuous curve on $[0, n]$ by using linear interpolation between $(t, Q(t))$ and $(t+1, Q(t+1))$. When properly rescaled, it converges in distribution to Brownian excursion. Brownian excursion is Brownian motion $B(t)$ on $[0, 1]$ conditioned to be positive and to take the value 0 at time 1. Alternatively, it is a Brownian bridge process conditioned to be positive. Another representation of a Brownian excursion $e(t)$ due to Paul Lévy (and noted by Ito and McKean [14]) is in terms of the last time τ_- that Brownian motion $B(t)$ hits zero before time 1 and the first time τ_+ that Brownian motion $B(t)$ hits zero after time 1:

$$\{e(t) : 0 \leq t \leq 1\} \stackrel{\mathcal{L}}{=} \left\{ \frac{|B((1-t)\tau_- + t\tau_+)|}{\sqrt{\tau_+ - \tau_-}} : 0 \leq t \leq 1 \right\}.$$

Since B , and thus e , are continuous processes, we note the following:

- (i) $\max_{0 \leq t \leq 1} e(t)$ is a random variable. In fact, it has the theta distribution, as proved by Rényi and Szekeres [28], Kennedy [17], Flajolet and Odlyzko [12] and others.
- (ii) For every $\epsilon \in (0, 1/2)$, $\min_{\epsilon \leq t \leq 1-\epsilon} e(t)$ is a random variable without an atom at zero (for otherwise, it would contradict Lévy's representation).

The convergence of partial sums of i.i.d. zero mean random variables with a finite variance to Brownian motion is well-known. As a consequence, partial sums conditioned on being positive and attaining 0 at the start and end, when suitably rescaled, converge to Brownian excursion. In particular,

$$\left(\frac{Q(tn)}{\sigma\sqrt{n}}, 0 \leq t \leq 1 \right) \xrightarrow{\mathcal{L}} (e(t), 0 \leq t \leq 1)$$

in the sup norm metric. In other words, there exists a sequence of Brownian excursions $e_1, e_2, \dots$, such that

$$\sup_{0 \leq t \leq 1} \left| \frac{Q(tn)}{\sigma\sqrt{n}} - e_n(t) \right| \rightarrow 0$$

in probability as $n \rightarrow \infty$. The work on this convergence of the Q process goes back to Le Gall [20,21], Aldous [1–4], Bennes and Kersting [8], Marckert and Mokkadem [22], Duquesne [10], and others. A consequence of this and remarks (i) and (ii) above

is that for given $\epsilon > 0$ and $\delta > 0$, we can find $\theta > 0$ such that with probability at least $1 - \delta$,

$$\theta\sqrt{n} \leq \min_{\epsilon n \leq t \leq (1-\epsilon)n} Q(t) \leq \max_{\epsilon n \leq t \leq (1-\epsilon)n} Q(t) \leq \frac{1}{\theta}\sqrt{n}.$$

Call this event $G(\epsilon, \theta)$.

5 Proof of Theorem 1

We are now ready to prove Theorem 1. Let us denote the sizes of the trees explored by the depth first search steps by $N_1, N_2, N_3, \dots$, where we note that $1 \leq N_i \leq b$ for all i . The indices of the nodes at which the depth first search operations are started are denoted by $T_1 < T_2 < \dots$, with $T_1 = 0, T_{i+1} = T_i + N_i$ for all i . We have

$$R_n = R = \min\{r > 0 : T_r = n - 1\}.$$

This forces $Q(T_R) = 1$, and thus $Q(T_R + 1) = 0$. In total, the driving stack has held R_n nodes.

By duality,

$$[R \geq r] \equiv [N_1 + \dots + N_r < n].$$

Recalling the event $G(\epsilon, \theta)$ from the previous section, and denoting by $I(\epsilon)$ the interval of indices $[\epsilon n, r - \epsilon n]$, we have

$$P\{R \geq r\} \leq P\left\{\sum_{i \in I(\epsilon)} N_i < n\right\} = P\left\{\sum_{i \in I(\epsilon)} (N_i - EN_i) < n - \sum_{i \in I(\epsilon)} EN_i\right\}.$$

Note that for $i \in I(\epsilon)$, we have $T_i \geq i \geq \epsilon n$, and $T_i \leq r - \epsilon n$. Also, for any i and j , by simple conditioning, we have

$$E\{(N_i - EN_i)(N_j - EN_j)\} = 0.$$

By Chebyshev's inequality, and the fact that $V\{N_i\} \leq b^2$, we have

$$P\{R \geq r\} \leq \frac{b^2|I(\epsilon)|}{\left(\sum_{i \in I(\epsilon)} EN_i - n\right)_+^2},$$

where $(y)_+ = \max(y, 0)$. Let E_i be the event

$$\left[\theta\sqrt{n} \leq Q(T_i) \leq \frac{1}{\theta}\sqrt{n}, \epsilon n \leq T_i \leq (1 - \epsilon)n\right].$$

We have for $i \in I(\epsilon)$,

$$\begin{aligned} EN_i &\leq E\{N_i 1_{[G(\epsilon, \theta)]}\} + bP\{G^c(\epsilon, \theta)\} \\ &\leq \mu_b(1 + o(1))P\{G(\epsilon, \theta)\} + b\delta \\ &\leq \mu_b(1 + o(1)) + b\delta, \end{aligned}$$

where the $o(1)$ refers to the behavior as $n \rightarrow \infty$ for fixed b, ϵ, θ . Similarly,

$$\begin{aligned} EN_i &\geq E\{N_i 1_{[E_i]}\} \\ &\geq \mu_b(1 + o(1))P\{E_i\} \\ &\geq \mu_b(1 + o(1))(1 - P\{G^c(\epsilon, \theta)\}) \\ &\geq \mu_b(1 + o(1))(1 - \delta). \end{aligned}$$

Therefore,

$$\begin{aligned} \mu_b(1 + o(1))(1 - \delta)(r - 2\epsilon n) &\leq \mu_b(1 + o(1))(1 - \delta)|I(\epsilon)| \\ &\leq \sum_{i \in I(\epsilon)} EN_i \\ &\leq (\mu_b(1 + o(1)) + b\delta)|I(\epsilon)|. \\ &\leq (\mu_b(1 + o(1)) + b\delta)r. \end{aligned}$$

For n large enough,

$$\sum_{i \in I(\epsilon)} EN_i - n \geq \mu_b(1 - 2\delta)(r - 2\epsilon n) - n.$$

To show that for any small $\gamma > 0$, $P\{R \geq n(1 + \gamma)/\mu_b\} \rightarrow 0$, we choose first $\epsilon = \gamma/(4\mu_b)$, then $\delta = \gamma/(4 + 4\gamma)$, and then $\theta > 0$ so small that $P\{G^c(\epsilon, \theta)\} \leq \delta$. With those choices,

$$\begin{aligned} \mu_b(1 - 2\delta)(\mu_b^{-1}(n + n\gamma) - 2\epsilon n) - n &= n(1 - 2\delta)(1 + \gamma/2) - n \\ &= n \frac{\gamma^2}{4 + 4\gamma} \\ &\stackrel{\text{def}}{=} n\gamma'. \end{aligned}$$

So, for n large enough, by Chebyshev's inequality invoked above,

$$P\{R \geq n(1 + \gamma)/\mu_b\} \leq \frac{b^2 r}{(n\gamma')^2} \rightarrow 0.$$

Table 1 Test results

Budget b	R_n T_2	T_3	T_5	T_{10}	T_{20}	T_{40}
50	36, 164, 525	70, 861, 643	51, 881, 667	52, 454, 166	118, 688, 002	51, 764, 031
500	10, 656, 708	21, 596, 567	16, 331, 233	17, 384, 235	40, 853, 820	18, 621, 338
5000	3, 293, 174	6, 764, 700	5, 181, 648	5, 584, 327	13, 295, 349	6, 111, 917
50,000	1, 029, 919	2, 121, 621	1, 643, 866	1, 758, 032	4, 242, 028	1, 994, 306
500,000	321, 993	675, 660	520, 932	572, 507	1, 340, 837	620, 958
n	440, 480, 727	722, 813, 312	393, 789, 282	264, 933, 315	434, 056, 761	145, 049, 024
Seed	1, 486, 862, 923	2, 053, 907, 278	2, 018, 277, 783	1, 671, 700, 366	1, 992, 061, 106	2, 121, 539, 929

Table 2 Accuracy of Theorem 1 estimate

Budget b	$R_n/\sigma n$ T_2	T_3	T_5	T_{10}	T_{20}	T_{40}	Estimate $\sqrt{\pi/8b}$
50	.10055	.09804	.09316	.09333	.08872	.08082	.08862
500	.02963	.02988	.02933	.03093	.03054	.02907	.02802
5000	.00916	.00936	.00936	.00994	.00994	.00954	.00886
50,000	.00286	.00294	.00295	.00313	.00317	.00311	.00280
500,000	.00090	.00093	.00094	.00102	.00100	.00097	.00089

The other side is dealt with in the same manner. First,

$$\begin{aligned}
 P\{R < r\} &\leq P\left\{\sum_{i \in I(\epsilon)} N_i \geq n\right\} \\
 &= P\left\{\sum_{i \in I(\epsilon)} (N_i - EN_i) \geq n - \sum_{i \in I(\epsilon)} EN_i\right\} \\
 &\leq \frac{b^2 r}{\left(n - \sum_{i \in I(\epsilon)} EN_i\right)_+^2}.
 \end{aligned}$$

Note that for n large enough,

$$n - \sum_{i \in I(\epsilon)} EN_i \geq n - (\mu_b + 2b\delta)r.$$

To show that for any small $\gamma > 0$, $P\{R < n(1 - \gamma)/\mu_b\} \rightarrow 0$, we choose first $\delta = \gamma\mu_b/(2b)$, and then $\theta > 0$ so small that $P\{G^c(\epsilon, \theta)\} \leq \delta$. With those choices,

$$n - (\mu_b + 2b\delta) \frac{n(1 - \gamma)}{\mu_b} = n\gamma^2.$$

So, for n large enough, by Chebyshev's inequality invoked above,

$$P\{R < n(1 - \gamma)/\mu_b\} \leq \frac{b^2 r}{n^2 \gamma^4} \rightarrow 0.$$

6 Empirical Results and Discussion

In the spirit of the remark by Hooker quoted in the introduction, we put the result of Theorem 1 to the test experimentally. Except for $\Delta = 2$ we considered critical Galton–Watson trees with $p_i = 1/i\Delta$, $i = 1, 2, \dots, \Delta$ and $p_0 = 1 - \sum_{i=1}^{\Delta} p_i$. For this family of trees we have that $\sigma^2 = (\Delta - 1)/2$. For $\Delta = 2$ we used the distribution $p = (1/3, 1/3, 1/3)$ which has $\sigma^2 = 2/3$. For each $\Delta = 2, 3, 5, 10, 20, 40$ we generated random Galton–Watson trees until we found one with at least 10^8 nodes. In Table 1 we show the results of searching these trees with various budgets b . As before we denote the number of restarts by R_n . According to Theorem 1, $R_n/\sigma n$ should asymptotically approach $\sqrt{\pi/8b}$ as b gets large. The results in Table 2 show good empirical agreement with this asymptotic bound even for small values of b .

We now discuss how results like Theorem 1 are useful in practice. Recall from the introduction that budgeting causes two competing forms of overhead: restart cost and idle time. Let us model a unit of computation time by a single node evaluation in Algorithm 1. We can model the restart cost for a job as a constant r units of time, i.e., r node evaluations. For example a node evaluation in *mplrs* involves pivoting an m by n matrix. The restart cost in this case is approximately $r = n$, as the input matrix has to be pivoted up to n times to reach a restart cobasis. Note that the restart cost is born by an individual worker and so the total restart cost is independent of the number of workers. As a proportion of the total work required it is precisely the ratio R_n/n that is estimated in Theorem 1.

If the job list becomes empty then any worker returning for work remains idle until some new jobs arrive. The cost is the number of idle workers for each time unit when the list is empty. Unlike restart costs, these costs scale up as the number of workers increases and can greatly reduce the parallelization efficiency achieved. To minimize these costs the job list should always be well populated.

In practice the budget is not fixed for the duration of a run. The master monitors the size of job list and if it drops too low scales the budget down and if it gets too large scales the budget up. Theorem 1 tells us how this scaling will work for Galton–Watson trees: scaling down (up) by a factor of 100 increases (decreases) the work returned on average by a factor of 10. The dependence of the job list size on the square root of the budget explains why relatively small changes in the budget do not greatly effect performance. This can be observed in Fig. 2 for example.

The empirical results obtained by Avis and Jordan [6,7] come from a variety of practical applications. Whilst these trees are not generated according to the Galton–Watson model, the behaviour observed is surprisingly similar to that predicted by Theorem 1. It will be interesting to see if similar theoretical results hold for other classes of trees.

Acknowledgements We gratefully acknowledge helpful conversations with Louigi Addario-Berry and Charles Jordan. The research of Avis was supported by a JSPS Grant No. (24300002) Kakenhi Grant and a Grant-in-Aid for Scientific Research on Innovative Areas, ‘Exploring the Limits of Computation (ELC)’. The research of Devroye was supported by the Natural Sciences and Engineering Research Council of Canada.

References

1. Aldous, D.: The continuum random tree. II. An overview. In: *Stochastic Analysis* (Durham, 1990), vol. 167, pp. 23–70, Cambridge University Press, Cambridge (1991)
2. Aldous, D.: The continuum random tree. I. *Ann. Probab.* **19**, 1–28 (1991)
3. Aldous, D.: Asymptotic fringe distributions for general families of random trees. *Ann. Appl. Probab.* **1**, 228–266 (1991)
4. Aldous, D.: The continuum random tree. III. *Ann. Probab.* **21**, 248–289 (1993)
5. Athreya, K.B., Ney, P.E.: *Branching Processes*. Springer, Berlin (1972)
6. Avis, D., Jordan, C.: mplrs: a scaleable parallel vertex/facet enumeration code. [arXiv:1511.06487](https://arxiv.org/abs/1511.06487) (2015)
7. Avis, D., Jordan, C.: A parallel framework for reverse search using mts. [arXiv:1610.07735](https://arxiv.org/abs/1610.07735) (2016)
8. Bennes, J., Kersting, G.: A random walk approach to Galton–Watson trees. *J. Theor. Probab.* **13**, 777–803 (2000)
9. Blumofe, N., Leiserson, C.: Scheduling multithreaded computations by work stealing. *J. ACM* **46**, 720–748 (1999)
10. Duquesne, T.: A limit theorem for the contour process of conditioned Galton–Watson trees. *Ann. Probab.* **31**, 996–1027 (2003)
11. Dwass, M.: The total progeny in a branching process. *J. Appl. Probab.* **6**, 682–686 (1969)
12. Flajolet, P., Odlyzko, A.: The average height of binary trees and other simple trees. *J. Comput. Syst. Sci.* **25**, 171–213 (1982)
13. Hooker, J.N.: Testing heuristics: we have it all wrong. *J. Heuristics* **1**, 33–42 (1995)
14. Ito, K., McKean, H.P.: *Diffusion Processes and Their Sample Paths*. Springer, Berlin (1974)
15. Janson, S.: Simply generated trees, conditioned Galton–Watson trees, random allocations and condensation. *Probab. Surv.* **9**, 103–252 (2012)
16. Kennedy, D.P.: The Galton–Watson process conditioned on the total progeny. *J. Appl. Probab.* **12**, 800–806 (1975)
17. Kennedy, D.P.: The distribution of the maximum Brownian excursion. *J. Appl. Probab.* **13**, 371–376 (1976)
18. Kolchin, V.F.: Branching processes and random trees. In: *Problems in Cybernetics, Combinatorial Analysis and Graph Theory* (in Russian), pp. 85–97, Nauka, Moscow (1980)
19. Kolchin, V.F.: *Random Mappings*. Optimization Software Inc., New York (1986)
20. Le Gall, J.-F.: Marches aléatoires, mouvement Brownien et processus de branchement. In: *Séminaire de Probabilités XXIII*, edited by Azéma, J., Meyer, P.A., Yor, M. vol. 1372, pp. 258–274. *Lecture Notes in Mathematics*, Springer-Verlag, Berlin (1989)
21. Le Gall, J.-F.: Random trees and applications. *Probab. Surv.* **2**, 245–311 (2005)
22. Marckert, J.F., Mokkadem, A.: The depth first processes of Galton–Watson trees converge to the same Brownian excursion. *Ann. Probab.* **31**, 1655–1678 (2003)
23. Mattson, T., Saunders, B., Massingill, B.: *Patterns for Parallel Programming*. Addison-Wesley, Boston (2004)
24. McCreesh, C., Prosser, P.: The shape of the search tree for the maximum clique problem and the implications for parallel branch and bound. *ACM Trans. Parallel Process.* **2**, 8:1–8:27 (2015)
25. Meir, A., Moon, J.W.: On the altitude of nodes in random trees. *Can. J. Math.* **30**, 997–1015 (1978)
26. Moon, J.W.: *Counting Labelled Trees*. Canadian Mathematical Congress, Montreal (1970)
27. Petrov, V.V.: *Sums of Independent Random Variables*. Springer, Berlin (1975)
28. Rényi, A., Szekeres, G.: On the height of trees. *J. Aust. Math. Soc.* **7**, 497–507 (1967)

Publisher’s Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.