

*I thought the following four [rules] would be enough, provided that I made a firm and constant resolution not to fail even once in the observance of them. The first was never to accept anything as true if I had not evident knowledge of its being so. . . . The second, to divide each problem I examined into as many parts as was feasible, and as was requisite for its better solution. The third, to direct my thoughts in an orderly way. . . establishing an order in thought even when the objects had no natural priority one to another. And the last, to make throughout such complete enumerations and such general surveys that I might be sure of leaving nothing out.*

— René Descartes, *Discours de la Méthode* (1637)

*There are those who think that life has nothing left to chance  
A host of holy horrors to direct our aimless dance*

— Rush, “Freewill”, *Permanent Waves* (1980), lyrics by Neal Peart

*What is luck?  
Luck is probability taken personally.  
It is the excitement of bad math.*

— Penn Jillette (2001), quoting Chip Denman (1998)

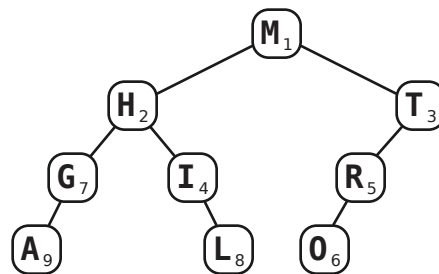
## 10 Randomized Binary Search Trees

In this lecture, we consider two randomized alternatives to balanced binary search tree structures such as AVL trees, red-black trees, B-trees, or splay trees, which are arguably simpler than any of these deterministic structures.

### 10.1 Treaps

#### 10.1.1 Definitions

A *treap* is a binary tree in which every node has both a *search key* and a *priority*, where the inorder sequence of search keys is sorted and each node's priority is smaller than the priorities of its children.<sup>1</sup> In other words, a treap is simultaneously a binary search tree for the search keys and a (min-)heap for the priorities. In our examples, we will use letters for the search keys and numbers for the priorities.



A treap. Letters are search keys; numbers are priorities.

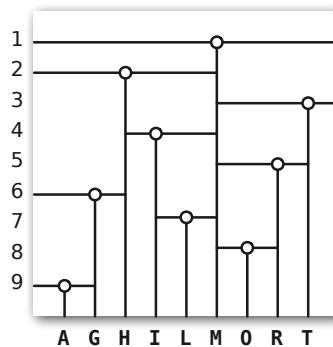
I'll assume from now on that all the keys and priorities are distinct. Under this assumption, we can easily prove by induction that the structure of a treap is completely determined by the

<sup>1</sup>Sometimes I hate English. Normally, 'higher priority' means 'more important', but 'first priority' is also more important than 'second priority'. Maybe 'posteriority' would be better; one student suggested 'unimportance'.

search keys and priorities of its nodes. Since it's a heap, the node  $v$  with highest priority must be the root. Since it's also a binary search tree, any node  $u$  with  $key(u) < key(v)$  must be in the left subtree, and any node  $w$  with  $key(w) > key(v)$  must be in the right subtree. Finally, since the subtrees are treaps, by induction, their structures are completely determined. The base case is the trivial empty treap.

Another way to describe the structure is that a treap is exactly the binary search tree that results by inserting the nodes one at a time into an initially empty tree, in order of increasing priority, using the standard textbook insertion algorithm. This characterization is also easy to prove by induction.

A third description interprets the keys and priorities as the coordinates of a set of points in the plane. The root corresponds to a T whose joint lies on the topmost point. The T splits the plane into three parts. The top part is (by definition) empty; the left and right parts are split recursively. This interpretation has some interesting applications in computational geometry, which (unfortunately) we won't have time to talk about.



A geometric interpretation of the same treap.

Treaps were first discovered by Jean Vuillemin in 1980, but he called them *Cartesian trees*.<sup>2</sup> The word 'treap' was first used by Edward McCreight around 1980 to describe a slightly different data structure, but he later switched to the more prosaic name *priority search trees*.<sup>3</sup> Treaps were rediscovered and used to build randomized search trees by Cecilia Aragon and Raimund Seidel in 1989.<sup>4</sup> A different kind of randomized binary search tree, which uses random rebalancing instead of random priorities, was later discovered and analyzed by Conrado Martínez and Salvador Roura in 1996.<sup>5</sup>

### 10.1.2 Treap Operations

The search algorithm is the usual one for binary search trees. The time for a successful search is proportional to the depth of the node. The time for an unsuccessful search is proportional to the depth of either its successor or its predecessor.

To insert a new node  $z$ , we start by using the standard binary search tree insertion algorithm to insert it at the bottom of the tree. At the point, the search keys still form a search tree, but the

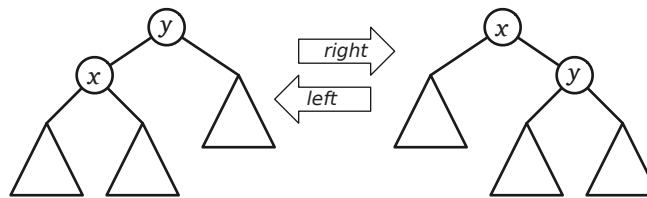
<sup>2</sup>J. Vuillemin, A unifying look at data structures. *Commun. ACM* 23:229–239, 1980.

<sup>3</sup>E. M. McCreight. Priority search trees. *SIAM J. Comput.* 14(2):257–276, 1985.

<sup>4</sup>R. Seidel and C. R. Aragon. Randomized search trees. *Algorithmica* 16:464–497, 1996.

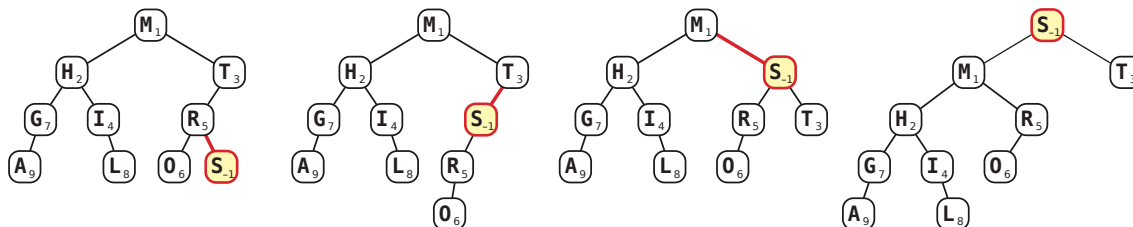
<sup>5</sup>C. Martínez and S. Roura. Randomized binary search trees. *J. ACM* 45(2):288–323, 1998. The results in this paper are virtually identical (including the constant factors!) to the corresponding results for treaps, although the analysis techniques are quite different.

priorities may no longer form a heap. To fix the heap property, as long as  $z$  has smaller priority than its parent, perform a *rotation* at  $z$ , a local operation that decreases the depth of  $z$  by one and increases its parent's depth by one, while maintaining the search tree property. Rotations can be performed in constant time, since they only involve simple pointer manipulation.



A right rotation at  $x$  and a left rotation at  $y$  are inverses.

The overall time to insert  $z$  is proportional to the depth of  $z$  before the rotations—we have to walk down the treap to insert  $z$ , and then walk back up the treap doing rotations. Another way to say this is that the time to insert  $z$  is roughly twice the time to perform an unsuccessful search for  $key(z)$ .



Left to right: After inserting  $S$  with priority  $-1$ , rotate it up to fix the heap property.

Right to left: Before deleting  $S$ , rotate it down to make it a leaf.

To delete a node, we just run the insertion algorithm backward in time. Suppose we want to delete node  $z$ . As long as  $z$  is not a leaf, perform a rotation at the child of  $z$  with smaller priority. This moves  $z$  down a level and its smaller-priority child up a level. The choice of which child to rotate preserves the heap property everywhere except at  $z$ . When  $z$  becomes a leaf, chop it off.

We sometimes also want to *split* a treap  $T$  into two treaps  $T_<$  and  $T_>$  along some pivot key  $\pi$ , so that all the nodes in  $T_<$  have keys less than  $\pi$  and all the nodes in  $T_>$  have keys bigger than  $\pi$ . A simple way to do this is to insert a new node  $z$  with  $key(z) = \pi$  and  $priority(z) = -\infty$ . After the insertion, the new node is the root of the treap. If we delete the root, the left and right sub-treaps are exactly the trees we want. The time to split at  $\pi$  is roughly twice the time to (unsuccessfully) search for  $\pi$ .

Similarly, we may want to *join* two treaps  $T_<$  and  $T_>$ , where every node in  $T_<$  has a smaller search key than any node in  $T_>$ , into one super-treap. Merging is just splitting in reverse—create a dummy root whose left sub-treap is  $T_<$  and whose right sub-treap is  $T_>$ , rotate the dummy node down to a leaf, and then cut it off.

The cost of each of these operations is proportional to the depth of some node  $v$  in the treap.

- **Search:** A successful search for key  $k$  takes  $O(\text{depth}(v))$  time, where  $v$  is the node with  $key(v) = k$ . For an unsuccessful search, let  $v^-$  be the inorder predecessor of  $k$  (the node whose key is just barely smaller than  $k$ ), and let  $v^+$  be the inorder successor of  $k$  (the node whose key is just barely larger than  $k$ ). Since the last node examined by the binary search is either  $v^-$  or  $v^+$ , the time for an unsuccessful search is either  $O(\text{depth}(v^+))$  or  $O(\text{depth}(v^-))$ .

- **Insert/Delete:** Inserting a new node with key  $k$  takes either  $O(\text{depth}(v^+))$  time or  $O(\text{depth}(v^-))$  time, where  $v^+$  and  $v^-$  are the predecessor and successor of the new node. Deletion is just insertion in reverse.
- **Split/Join:** Splitting a treap at pivot value  $k$  takes either  $O(\text{depth}(v^+))$  time or  $O(\text{depth}(v^-))$  time, since it costs the same as inserting a new dummy root with search key  $k$  and priority  $-\infty$ . Merging is just splitting in reverse.

In the worst case, the depth of an  $n$ -node treap is  $\Theta(n)$ , so each of these operations has a worst-case running time of  $\Theta(n)$ .

### 10.1.3 Random Priorities

A *randomized treap* is a treap in which the priorities are *independently and uniformly distributed continuous random variables*. Whenever we insert a new search key into the treap, we generate a random real number between (say) 0 and 1 and use that number as the priority of the new node. The only reason we're using real numbers is so that the probability of two nodes having the same priority is zero, since equal priorities make the analysis slightly messier. (In practice, we could just choose random integers from a large range like 0 to  $2^{31} - 1$  and break ties arbitrarily; occasional ties have almost no practical effect on the performance of the data structure.) Also, since the priorities are independent, each node is equally likely to have the smallest priority.

The cost of all the operations we discussed—search, insert, delete, split, join—is proportional to the depth of some node in the tree. Here we'll see that the *expected* depth of *any* node is  $O(\log n)$ , which implies that the expected running time for any of those operations is also  $O(\log n)$ .

Let  $x_k$  denote the node with the  $k$ th smallest search key. To simplify notation, let us write  $i \uparrow k$  (read “ $i$  above  $k$ ”) to mean that  $x_i$  is a proper ancestor of  $x_k$ . Since the depth of  $v$  is just the number of proper ancestors of  $v$ , we have the following identity:

$$\text{depth}(x_k) = \sum_{i=1}^n [i \uparrow k].$$

(Again, we're using Iverson bracket notation.) Now we can express the *expected* depth of a node in terms of these indicator variables as follows.

$$E[\text{depth}(x_k)] = \sum_{i=1}^n E[[i \uparrow k]] = \sum_{i=1}^n \Pr[i \uparrow k]$$

(Just as in our analysis of matching nuts and bolts, we're using linearity of expectation and the fact that  $E[X] = \Pr[X = 1]$  for any zero-one variable  $X$ ; in this case,  $X = [i \uparrow k]$ .) So to compute the expected depth of a node, we just have to compute the probability that some node is a proper ancestor of some other node.

Fortunately, we can do this easily once we prove a simple structural lemma. Let  $X(i, k)$  denote either the subset of treap nodes  $\{x_i, x_{i+1}, \dots, x_k\}$  or the subset  $\{x_k, x_{k+1}, \dots, x_i\}$ , depending on whether  $i < k$  or  $i > k$ . The order of the arguments is unimportant; the subsets  $X(i, k)$  and  $X(k, i)$  are identical. The subset  $X(1, n) = X(n, 1)$  contains all  $n$  nodes in the treap.

**Lemma 1.** *For all  $i \neq k$ , we have  $i \uparrow k$  if and only if  $x_i$  has the smallest priority among all nodes in  $X(i, k)$ .*

**Proof:** There are four cases to consider.

If  $x_i$  is the root, then  $i \uparrow k$ , and by definition, it has the smallest priority of *any* node in the treap, so it must have the smallest priority in  $X(i, k)$ .

On the other hand, if  $x_k$  is the root, then  $k \uparrow i$ , so  $i \not\uparrow k$ . Moreover,  $x_i$  does not have the smallest priority in  $X(i, k)$  —  $x_k$  does.

On the gripping hand<sup>6</sup>, suppose some other node  $x_j$  is the root. If  $x_i$  and  $x_k$  are in different subtrees, then either  $i < j < k$  or  $i > j > k$ , so  $x_j \in X(i, k)$ . In this case, we have both  $i \not\uparrow k$  and  $k \not\uparrow i$ , and  $x_i$  does not have the smallest priority in  $X(i, k)$  —  $x_j$  does.

Finally, if  $x_i$  and  $x_k$  are in the same subtree, the lemma follows from the inductive hypothesis (or, if you prefer, the Recursion Fairy), because the subtree is a smaller treap. The empty treap is the trivial base case.  $\square$

Since each node in  $X(i, k)$  is equally likely to have smallest priority, we immediately have the probability we wanted:

$$\Pr[i \uparrow k] = \frac{[i \neq k]}{|k - i| + 1} = \begin{cases} \frac{1}{k - i + 1} & \text{if } i < k \\ 0 & \text{if } i = k \\ \frac{1}{i - k + 1} & \text{if } i > k \end{cases}$$

To compute the expected depth of a node  $x_k$ , we just plug this probability into our formula and grind through the algebra.

$$\begin{aligned} E[\text{depth}(x_k)] &= \sum_{i=1}^n \Pr[i \uparrow k] = \sum_{i=1}^{k-1} \frac{1}{k - i + 1} + \sum_{i=k+1}^n \frac{1}{i - k + 1} \\ &= \sum_{j=2}^k \frac{1}{j} + \sum_{i=2}^{n-k+1} \frac{1}{i} \\ &= H_k - 1 + H_{n-k+1} - 1 \\ &< \ln k + \ln(n - k + 1) - 2 \\ &< 2 \ln n - 2. \end{aligned}$$

In conclusion, every search, insertion, deletion, split, and join operation in an  $n$ -node randomized binary search tree takes  $O(\log n)$  expected time.

Since a treap is exactly the binary tree that results when you insert the keys in order of increasing priority, a randomized treap is the result of inserting the keys in *random* order. So our analysis also automatically gives us the expected depth of any node in a binary tree built by random insertions (without using priorities).

#### 10.1.4 Randomized Quicksort (Again!)

We've already seen two completely different ways of describing randomized quicksort. The first is the familiar recursive one: choose a random pivot, partition, and recurse. The second is a less familiar iterative version: repeatedly choose a new random pivot, partition whatever subset contains it, and continue. But there's a third way to describe randomized quicksort, this time in terms of binary search trees.

<sup>6</sup>See Larry Niven and Jerry Pournelle, *The Gripping Hand*, Pocket Books, 1994.

**RANDOMIZED QUICKSORT:**

$T \leftarrow$  an empty binary search tree  
 insert the keys into  $T$  in random order  
 output the inorder sequence of keys in  $T$

Our treap analysis tells us that this algorithm will run in  $O(n \log n)$  expected time, since each key is inserted in  $O(\log n)$  expected time.

Why is this quicksort? Just like last time, all we've done is rearrange the order of the comparisons. Intuitively, the binary tree is just the recursion tree created by the normal version of quicksort. In the recursive formulation, we compare the initial pivot against everything else and then recurse. In the binary tree formulation, the first "pivot" becomes the root of the tree without any comparisons, but then later as each other key is inserted into the tree, it is compared against the root. Either way, the first pivot chosen is compared with everything else. The partition splits the remaining items into a left subarray and a right subarray; in the binary tree version, these are exactly the items that go into the left subtree and the right subtree. Since both algorithms define the same two subproblems, by induction, both algorithms perform the same comparisons.

We even saw the probability  $1/(|k - i| + 1)$  before, when we were talking about sorting nuts and bolts with a variant of randomized quicksort. In the more familiar setting of sorting an array of numbers, the probability that randomized quicksort compares the  $i$ th largest and  $k$ th largest elements is exactly  $2/(|k - i| + 1)$ . The binary tree version of quicksort compares  $x_i$  and  $x_k$  if and only if  $i \uparrow k$  or  $k \uparrow i$ , so the probabilities are exactly the same.

## 10.2 Skip Lists

*Skip lists*, first discovered by Bill Pugh in the late 1980's,<sup>7</sup> have many of the usual desirable properties of balanced binary search trees, but their superficial structure is very different.

At a high level, a skip list is just a sorted linked list with some random shortcuts. To do a search in a normal singly-linked list of length  $n$ , we obviously need to look at  $n$  items in the worst case. To speed up this process, we can make a second-level list that contains roughly half the items from the original list. Specifically, for each item in the original list, we duplicate it with probability  $1/2$ . We then string together all the duplicates into a second sorted linked list, and add a pointer from each duplicate back to its original. Just to be safe, we also add sentinel nodes at the beginning and end of both lists.

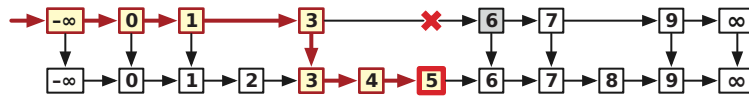


A linked list with some randomly-chosen shortcuts.

Now we can find a value  $x$  in this augmented structure using a two-stage algorithm. First, we scan for  $x$  in the shortcut list, starting at the  $-\infty$  sentinel node. If we find  $x$ , we're done. Otherwise, we reach some value bigger than  $x$  and we know that  $x$  is not in the shortcut list. Let  $w$  be the largest item less than  $x$  in the shortcut list. In the second phase, we scan for  $x$  in the original list, starting from  $w$ . Again, if we reach a value bigger than  $x$ , we know that  $x$  is not in the data structure.

Since each node appears in the shortcut list with probability  $1/2$ , the expected number of nodes examined in the first phase is at most  $n/2$ . Only one of the nodes examined in the second

<sup>7</sup>William Pugh. Skip lists: A probabilistic alternative to balanced trees. *Commun. ACM* 33(6):668–676, 1990.

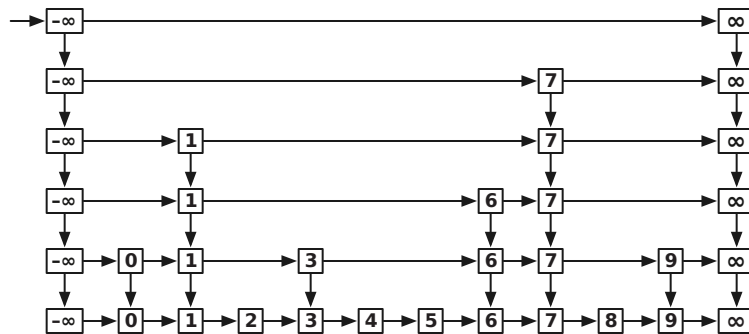


Searching for 5 in a list with shortcuts.

phase has a duplicate. The probability that any node is followed by  $k$  nodes without duplicates is  $2^{-k}$ , so the expected number of nodes examined in the second phase is at most  $1 + \sum_{k \geq 0} 2^{-k} = 2$ . Thus, by adding these random shortcuts, we've reduced the cost of a search from  $n$  to  $n/2 + 2$ , roughly a factor of two in savings.

### 10.2.1 Recursive Random Shortcuts

Now there's an obvious improvement—add shortcuts to the shortcuts, and repeat recursively. That's exactly how skip lists are constructed. For each node in the original list, we repeatedly flip a coin until we get tails. Each time we get heads, we make a new copy of the node. The duplicates are stacked up in levels, and the nodes on each level are strung together into sorted linked lists. Each node  $v$  stores a search key ( $key(v)$ ), a pointer to its next lower copy ( $down(v)$ ), and a pointer to the next node in its level ( $right(v)$ ).



A skip list is a linked list with recursive random shortcuts.

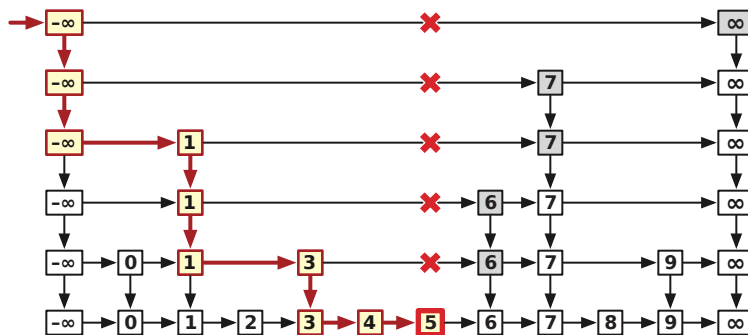
The search algorithm for skip lists is very simple. Starting at the leftmost node  $L$  in the highest level, we scan through each level as far as we can without passing the target value  $x$ , and then proceed down to the next level. The search ends when we either reach a node with search key  $x$  or fail to find  $x$  on the lowest level.

```

SKIPLISTFIND( $x, L$ ):
   $v \leftarrow L$ 
  while ( $v \neq \text{NULL}$  and  $key(v) \neq x$ )
    if  $key(right(v)) > x$ 
       $v \leftarrow down(v)$ 
    else
       $v \leftarrow right(v)$ 
  return  $v$ 

```

Intuitively, since each level of the skip lists has about half the number of nodes as the previous level, the total number of levels should be about  $O(\log n)$ . Similarly, each time we add another level of random shortcuts to the skip list, we cut the search time roughly in half, except for a constant overhead, so  $O(\log n)$  levels should give us an overall expected search time of  $O(\log n)$ . Let's formalize each of these two intuitive observations.



Searching for 5 in a skip list.

### 10.2.2 Number of Levels

The actual values of the search keys don't affect the skip list analysis, so let's assume the keys are the integers 1 through  $n$ . Let  $L(x)$  be the number of levels of the skip list that contain some search key  $x$ , not counting the bottom level. Each new copy of  $x$  is created with probability  $1/2$  from the previous level, essentially by flipping a coin. We can compute the expected value of  $L(x)$  recursively—with probability  $1/2$ , we flip tails and  $L(x) = 0$ ; and with probability  $1/2$ , we flip heads, increase  $L(x)$  by one, and recurse:

$$E[L(x)] = \frac{1}{2} \cdot 0 + \frac{1}{2}(1 + E[L(x)])$$

Solving this equation gives us  $E[L(x)] = 1$ .

In order to analyze the expected worst-case cost of a search, however, we need a bound on the *number of levels*  $L = \max_x L(x)$ . Unfortunately, we can't compute the average of a maximum the way we would compute the average of a sum. Instead, we derive a stronger result: **The depth of a skip list storing  $n$  keys is  $O(\log n)$  with high probability.** “High probability” is a technical term that means the probability is at least  $1 - 1/n^c$  for some constant  $c \geq 1$ ; the hidden constant in the  $O(\log n)$  bound could depend on  $c$ .

In order for a search key  $x$  to appear on level  $\ell$ , it must have flipped  $\ell$  heads in a row when it was inserted, so  $\Pr[L(x) \geq \ell] = 2^{-\ell}$ . The skip list has at least  $\ell$  levels if and only if  $L(x) \geq \ell$  for at least one of the  $n$  search keys.

$$\Pr[L \geq \ell] = \Pr[(L(1) \geq \ell) \vee (L(2) \geq \ell) \vee \dots \vee (L(n) \geq \ell)]$$

Using the **union bound** —  $\Pr[A \vee B] \leq \Pr[A] + \Pr[B]$  for any random events  $A$  and  $B$  — we can simplify this as follows:

$$\Pr[L \geq \ell] \leq \sum_{x=1}^n \Pr[L(x) \geq \ell] = n \cdot \Pr[L(x) \geq \ell] = \frac{n}{2^\ell}.$$

When  $\ell \leq \lg n$ , this bound is trivial. However, for any constant  $c > 1$ , we have a strong upper bound

$$\Pr[L \geq c \lg n] \leq \frac{1}{n^{c-1}}.$$

We conclude that **with high probability, a skip list has  $O(\log n)$  levels.**

This high-probability bound indirectly implies a bound on the *expected* number of levels. Some simple algebra gives us the following alternate definition for expectation:

$$E[L] = \sum_{\ell \geq 0} \ell \cdot \Pr[L = \ell] = \sum_{\ell \geq 1} \Pr[L \geq \ell]$$

Clearly, if  $\ell < \ell'$ , then  $\Pr[L(x) \geq \ell] > \Pr[L(x) \geq \ell']$ . So we can derive an upper bound on the expected number of levels as follows:

$$\begin{aligned} E[L(x)] &= \sum_{\ell \geq 1} \Pr[L \geq \ell] = \sum_{\ell=1}^{\lg n} \Pr[L \geq \ell] + \sum_{\ell \geq \lg n+1} \Pr[L \geq \ell] \\ &\leq \sum_{\ell=1}^{\lg n} 1 + \sum_{\ell \geq \lg n+1} \frac{n}{2^\ell} \\ &= \lg n + \sum_{i \geq 1} \frac{1}{2^i} \quad [i = \ell - \lg n] \\ &= \lg n + 2 \end{aligned}$$

So in expectation, a skip list has *at most two* more levels than an ideal version where each level contains exactly half the nodes of the next level below. Notice that this is an *additive* penalty over a perfectly balanced structure, as opposed to treaps, where the expected depth is a constant *multiple* of the ideal  $\lg n$ .

### 10.2.3 Logarithmic Search Time

It's a little easier to analyze the cost of a search if we imagine running the algorithm backwards. `UPLEFTFIND` takes the output from `SKIPLISTFIND` as input and traces back through the data structure to the upper left corner. Skip lists don't really have up and left pointers, but we'll pretend that they do so we don't have to write ' $v \leftarrow \text{up}(v)$ ' or ' $v \leftarrow \text{left}(v)$ '.<sup>8</sup>

```

UPLEFTFIND(v):
  while (level(v) ≠ L)
    if up(v) exists
      v ← up(v)
    else
      v ← left(v)

```

Now for *every* node  $v$  in the skip list,  $\text{up}(v)$  exists with probability  $1/2$ . So for purposes of analysis, `UPLEFTFIND` is equivalent to the following algorithm:

```

FLIPWALK(v):
  while (v ≠ L)
    if COINFLIP = HEADS
      v ← up(v)
    else
      v ← left(v)

```

<sup>8</sup> Leonard de Vinci wrote all his notes using mirror-writing, but not because he wanted to keep his discoveries secret. He just had really bad arthritis in his right hand!

Obviously, the expected number of heads is exactly the same as the expected number of TAILS. Thus, the expected running time of this algorithm is twice the expected number of upward jumps. But we already know that the number of upward jumps is  $O(\log n)$  with high probability. It follows the running time of FLIPWALK is  $O(\log n)$  with high probability (and therefore in expectation).

## Exercises

1. Prove that a treap is exactly the binary search tree that results from inserting the nodes one at a time into an initially empty tree, in order of increasing priority, using the standard textbook insertion algorithm.
2. Consider a treap  $T$  with  $n$  vertices. As in the notes above, identify nodes in  $T$  by the ranks of their search keys; thus, 'node 5' means the node with the 5th smallest search key. Let  $i$ ,  $j$ , and  $k$  be integers such that  $1 \leq i \leq j \leq k \leq n$ .
  - (a) Prove that the expected number of proper descendants of any node in a treap is exactly equal to the expected depth of that node.
  - (b) The *left spine* of a binary tree is a path starting at the root and following only left-child pointers. What is the expected number of nodes in the left spine of  $T$ ?
  - (c) What is the expected number of leaves in  $T$ ? [Hint: What is the probability that node  $k$  is a leaf?]
  - (d) What is the expected number of nodes in  $T$  with two children?
  - (e) What is the expected number of nodes in  $T$  with exactly one child?
  - \* (f) What is the expected number of nodes in  $T$  with exactly one grandchild?
  - (g) Define the *priority rank* of a node in  $T$  to be one more than the number of nodes with smaller priority. For example, the root of  $T$  always has priority rank 1, and one of the children of the root has priority rank 2. What is the expected priority rank of node  $i$ ?
  - (h) What is the expected priority rank of the left child of the root (given that such a node exists)?
  - \* (i) What is the expected priority rank of the leftmost grandchild of the root (given that such a node exists)?
  - \* (j) What is the expected priority rank of a node with depth  $d$ ?
  - (k) What is the *exact* probability that node  $j$  is a common ancestor of node  $i$  and node  $k$ ?
  - (l) What is the *exact* expected length of the unique path in  $T$  from node  $i$  to node  $k$ ?
3. Recall that a *priority search tree* is a binary tree in which every node has both a *search key* and a *priority*, arranged so that the tree is simultaneously a binary search tree for the keys and a min-heap for the priorities. A **heater** is a priority search tree in which the *priorities* are given by the user, and the *search keys* are distributed uniformly and independently at random in the real interval  $[0, 1]$ . Intuitively, a heater is a sort of anti-treap.<sup>9</sup>

<sup>9</sup>If you choose not to decide, you still have made a choice.

The following problems consider an  $n$ -node heater  $T$  whose priorities are the integers from 1 to  $n$ . We identify nodes in  $T$  by their *priorities*; thus, ‘node 5’ means the node in  $T$  with *priority* 5. For example, the min-heap property implies that node 1 is the root of  $T$ . Finally, let  $i$  and  $j$  be integers with  $1 \leq i < j \leq n$ .

- (a) Prove that in a random permutation of the  $(i+1)$ -element set  $\{1, 2, \dots, i, j\}$ , elements  $i$  and  $j$  are adjacent with probability  $2/(i+1)$ .
  - (b) Prove that node  $i$  is an ancestor of node  $j$  with probability exactly  $2/(i+1)$ . [Hint: Use part (a)!]
  - (c) What is the exact probability that node  $i$  is a *descendant* of node  $j$ ? [Hint: Don't use part (a)!]
  - (d) What is the *exact* expected depth of node  $j$ ?
  - (e) Describe and analyze an algorithm to insert a new item into a heater. Express the expected running time of the algorithm in terms of the rank of the newly inserted item.
  - (f) Describe an algorithm to delete the minimum-priority item (the root) from an  $n$ -node heater. What is the expected running time of your algorithm?
- \*4. In the usual theoretical presentation of treaps, the priorities are random real numbers chosen uniformly from the interval  $[0, 1]$ . In practice, however, computers have access only to random *bits*. This problem asks you to analyze an implementation of treaps that takes this limitation into account.

Suppose the priority of a node  $v$  is abstractly represented as an infinite sequence  $\pi_v[1.. \infty]$  of random bits, which is interpreted as the rational number

$$\text{priority}(v) = \sum_{i=1}^{\infty} \pi_v[i] \cdot 2^{-i}.$$

However, only a finite number  $\ell_v$  of these bits are actually known at any given time. When a node  $v$  is first created, *none* of the priority bits are known:  $\ell_v = 0$ . We generate (or “reveal”) new random bits only when they are necessary to compare priorities. The following algorithm compares the priorities of any two nodes in  $O(1)$  expected time:

```

LARGERPRIORITY( $v, w$ ):
  for  $i \leftarrow 1$  to  $\infty$ 
    if  $i > \ell_v$ 
       $\ell_v \leftarrow i$ ;  $\pi_v[i] \leftarrow \text{RANDOMBIT}$ 
    if  $i > \ell_w$ 
       $\ell_w \leftarrow i$ ;  $\pi_w[i] \leftarrow \text{RANDOMBIT}$ 
    if  $\pi_v[i] > \pi_w[i]$ 
      return  $v$ 
    else if  $\pi_v[i] < \pi_w[i]$ 
      return  $w$ 

```

Suppose we insert  $n$  items one at a time into an initially empty treap. Let  $L = \sum_v \ell_v$  denote the total number of random bits generated by calls to LARGERPRIORITY during these insertions.

- (a) Prove that  $E[L] = \Theta(n)$ .
  - (b) Prove that  $E[\ell_v] = \Theta(1)$  for any node  $v$ . [Hint: This is equivalent to part (a). Why?]
  - (c) Prove that  $E[\ell_{\text{root}}] = \Theta(\log n)$ . [Hint: Why doesn't this contradict part (b)?]
5. Prove the following basic facts about skip lists, where  $n$  is the number of keys.
- (a) The expected number of nodes is  $O(n)$ .
  - (b) A new key can be inserted in  $O(\log n)$  time with high probability.
  - (c) A key can be deleted in  $O(\log n)$  time with high probability.
6. Suppose we are given two skip lists, one storing a set  $A$  of  $m$  keys, and the other storing a set  $B$  of  $n$  keys. Describe and analyze an algorithm to merge these into a single skip list storing the set  $A \cup B$  in  $O(n + m)$  expected time. Do *not* assume that every key in  $A$  is smaller than every key in  $B$ ; the two sets could be arbitrarily intermixed. [Hint: Do the obvious thing.]
7. Any skip list  $\mathcal{L}$  can be transformed into a binary search tree  $T(\mathcal{L})$  as follows. The root of  $T(\mathcal{L})$  is the leftmost node on the highest non-empty level of  $\mathcal{L}$ ; the left and right subtrees are constructed recursively from the nodes to the left and to the right of the root. Let's call the resulting tree  $T(\mathcal{L})$  a *skip list tree*.
- (a) Show that any search in  $T(\mathcal{L})$  is no more expensive than the corresponding search in  $\mathcal{L}$ . (Searching in  $T(\mathcal{L})$  could be *considerably* cheaper—why?)
  - (b) Describe an algorithm to insert a new search key into a skip list tree in  $O(\log n)$  expected time. Inserting key  $x$  into the tree  $T(\mathcal{L})$  should produce *exactly* the same tree as inserting  $x$  into the skip list  $\mathcal{L}$  and then transforming  $\mathcal{L}$  into a tree. [Hint: You need to maintain some additional information in the tree nodes.]
  - (c) Describe an algorithm to delete a search key from a skip list tree in  $O(\log n)$  expected time. Again, deleting key  $x$  from  $T(\mathcal{L})$  should produce *exactly* the same tree as deleting  $x$  from  $\mathcal{L}$  and then transforming  $\mathcal{L}$  into a tree.
8. Consider the following “loose” variant of treaps. Instead of generating priorities uniformly at random, for each node, we flip an independent fair coin until it comes up heads, and define the priority of the node to be the number of flips. Thus, for every positive integer  $k$ , a node has priority  $k$  with probability  $2^{-k}$ . In addition, we invert the heap property, by requiring that the priority of any node is *not* larger than the priority of its parent.
- This method creates many nodes with the same priority; in particular, about half the nodes will have priority 1. Thus, a single set of search keys and priorities may be consistent with many different loose treaps.
- Prove that the expected depth of any node in an  $n$ -node “loose treap” is  $O(\log n)$ . To be pedantic, the expectation is over the random choice of priorities, but for each choice of priorities, we consider the worst possible loose treap.

[Hint: This is almost exactly the same as the previous question.]

9. A **meldable priority queue** stores a set of keys from some totally-ordered universe (such as the integers) and supports the following operations:

- **MAKEQUEUE**: Return a new priority queue containing the empty set.
- **FINDMIN**( $Q$ ): Return the smallest element of  $Q$  (if any).
- **DELETEMIN**( $Q$ ): Remove the smallest element in  $Q$  (if any).
- **INSERT**( $Q, x$ ): Insert element  $x$  into  $Q$ , if it is not already there.
- **DECREASEKEY**( $Q, x, y$ ): Replace an element  $x \in Q$  with a smaller key  $y$ . (If  $y > x$ , the operation fails.) The input is a pointer directly to the node in  $Q$  containing  $x$ .
- **DELETE**( $Q, x$ ): Delete the element  $x \in Q$ . The input is a pointer directly to the node in  $Q$  containing  $x$ .
- **MELD**( $Q_1, Q_2$ ): Return a new priority queue containing all the elements of  $Q_1$  and  $Q_2$ ; this operation destroys  $Q_1$  and  $Q_2$ .

A simple way to implement such a data structure is to use a heap-ordered binary tree, where each node stores a key, along with pointers to its parent and two children. **MELD** can be implemented using the following randomized algorithm:

```

MELD( $Q_1, Q_2$ ):
  if  $Q_1$  is empty return  $Q_2$ 
  if  $Q_2$  is empty return  $Q_1$ 
  if  $\text{key}(Q_1) > \text{key}(Q_2)$ 
    swap  $Q_1 \leftrightarrow Q_2$ 
  with probability  $1/2$ 
     $\text{left}(Q_1) \leftarrow \text{MELD}(\text{left}(Q_1), Q_2)$ 
  else
     $\text{right}(Q_1) \leftarrow \text{MELD}(\text{right}(Q_1), Q_2)$ 
  return  $Q_1$ 

```

- Prove that for any heap-ordered binary trees  $Q_1$  and  $Q_2$  (not just those constructed by the operations listed above), the expected running time of **MELD**( $Q_1, Q_2$ ) is  $O(\log n)$ , where  $n = |Q_1| + |Q_2|$ . [Hint: What is the expected length of a random root-to-leaf path in an  $n$ -node binary tree, where each left/right choice is made with equal probability?]
- Prove that **MELD**( $Q_1, Q_2$ ) runs in  $O(\log n)$  time with high probability.
- Show that each of the other meldable priority queue operations can be implemented with at most one call to **MELD** and  $O(1)$  additional time. (This implies that every operation takes  $O(\log n)$  time with high probability.)