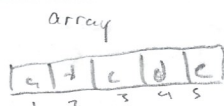


Peak Finding

One Dimensional



a-i are #'s.

Position a is a peak if $a \geq a-1$ and $a \geq a+1$

Basic algo goes through $1 \dots n-1$ $O(n)$

Better: middle $\rightarrow$ compare either side $\rightarrow$ go to middle of side that's higher
 to keep splitting array in 2. you can do this (worst case) $\log_2(n)$ times
 $\hookrightarrow$ each of these $\log_2(n)$ are of length 1 i.e. $O(1)$
 $\Rightarrow O(\log_2(n))$

2 Dimensional

could go through each one ... slow $O(mn)$

Better: pick middle column $j = m/2$, find global max on column j , say at (i, j)

compare $(i, j-1), (i, j), (i, j+1)$
 Pick left col if $(i, j-1) > (i, j)$ similar for right.
 else (i, j) is a peak.

if not peak you now have same problem w half as many rows.
 $\hookrightarrow$ get down to base case of one column, find global max, done.

$$T(n, m) = T(n, m/2) + O(n)$$

$$T(n, 1) = O(n)$$

$$\Rightarrow T(n, m) = O(n) + \dots + O(n) = O(n \log_2 m)$$

2) Random Access Machine

- Rand. Access Memory (RAM)
 $\hookrightarrow$ modelled by big array

- in $O(1)$ time can:

- load $O(1)$ words
- do $O(1)$ computation
- store $O(1)$ words

- $O(1)$ registers.

- word is w bits

(should be $\log_2(\text{size of memory})$ or rep'd w bits.)



Document Distance Problem, document = sequence of words

word = string of alphanumeric chars

think of documents as vector, $D[w] = \# \text{ occurrences of } w \text{ in } D$

$D_1 = \text{"the cat"}$

$D_2 = \text{"the dog"}$

how similar are these documents

$\hookrightarrow$ inner product

$$D(D_1, D_2) = D_1 \cdot D_2 = \sum_w D_1[w] \cdot D_2[w]$$

$\hookrightarrow$ scale invariant, 100 words identical ≤ 100000 words of eg. a.b.

Better would be angle between $\Rightarrow D(D_1, D_2) = \arccos\left(\frac{D_1 \cdot D_2}{|D_1| |D_2|}\right)$

Now we need to do 3 things

- 1) split docs into words
- 2) compute word frequencies
- 3) dot product.

how do you do these

- 1) scan through alpha numerics
- 2) for word in doc: count[word] += 1

$O(n)$
 $O(2 | \text{words} |)$

2) Pointer Machine

- dynamically allocated objects
- object has $O(1)$ fields
- field = word (ex: int)
 pointer to object or None.

3) Python Model

1) "list" = array $L[i] = L[i] + s$ $O(1)$ time

2) object with $O(1)$ attributes, $x = x.next \rightarrow O(1)$ time

Ex $L.append(x) \rightarrow$ table doubling (lecture 9)

$\hookrightarrow O(1)$ time

$$L = L_1 + L_2$$

(think about how to implement)

$$O(1 + |L_1| + |L_2|) \begin{cases} O(|L_1|) \{ \text{for } x \text{ in } L_1: L.append(x) \} \\ O(|L_2|) \{ \text{for } x \text{ in } L_2: L.append(x) \} \end{cases}$$

$len(L)$ $O(1)$, length stored w each list so just look it up

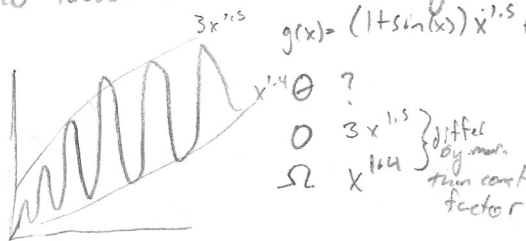
$L.sort(L) \rightarrow O(|L| \log |L|)$ time to compare (lecture 3)

dict: $D[key] = val \rightarrow O(1)$ more in later lects.

Tours: $x + y$ $O(|x| + |y|)$, $x \neq y$ $O((|x| + |y|)^{1/3})$

Asymptotic complexity Θ means you have upper and lower bound, can only use Θ if upper and lower bound differ by a constant factor (or less than)

ex $g(x) = x \cdot (1 + \sin(x))$
 upper $f(x) = 2x$
 lower $f(x) = 0$
 can't use Θ



$g(x) = 1.1x^2 + (10 + \sin(x+15))x^{1.5} + 600 = \Theta(x^2)$
 $\checkmark$ $1.2x^2$ x^2
 $O(1.2x^2)$ } get rid of const factor for the form
 $\Omega(x^2)$
 $O(x^2), \Omega(x^2) \Rightarrow \Theta(x^2)$

Worse
 $O(1)$
 $O(\log n)$
 $O(n)$
 $O(n \log n)$
 $O(n^2)$

ex $\log(n^{100}) \sim \Theta(\log n)$
 $\log_{100}(\log n)$

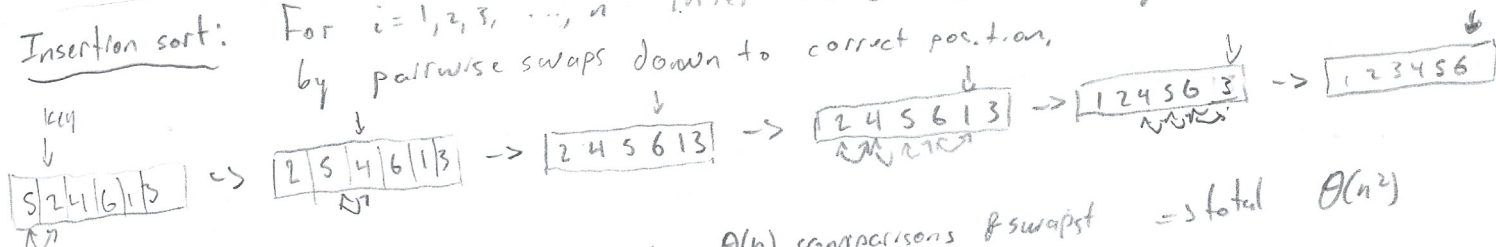
$\log_5(n) \sim \Theta(\log n)$
 $\rightarrow \frac{\log(n)}{\log(5)}$

$\log(ab) = \log a + \log b$

2) Doc distance need to get rid of punctuation and capitals to prevent error

Recitation

3) Insertion sort: For $i = 1, 2, 3, \dots, n$ Insert $A[i]$ into sorted array $[0; i-1]$ by pairwise swaps down to correct position.



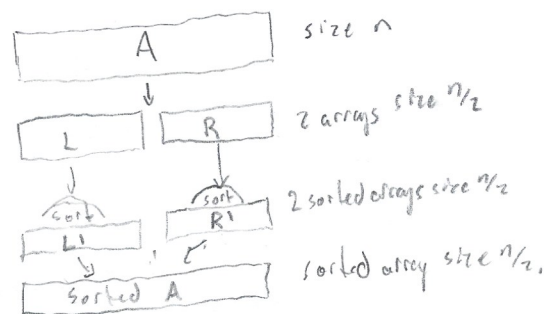
$\Theta(n)$ steps (key positions) Each step has $\Theta(n)$ comparisons & swaps $\Rightarrow$ total $\Theta(n^2)$

Binary Insertion: IF you only care about #comparisons $\rightarrow$ compare $\Rightarrow$ swaps, binary search $A[0; i-1]$ already sorted $\rightarrow$ search now $\Theta(\log n)$ $\rightarrow \Theta(n)$ steps $\rightarrow \Theta(n \log n)$ comparisons (still $\Theta(n^2)$ swaps)

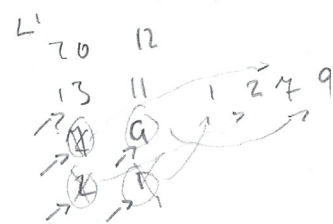
Merge Sort Divide and conquer.

$A[0; n] \rightarrow L = A[0; \frac{n}{2}-1]$
 $R = A[\frac{n}{2}; n]$

Complexity $T(n) = C_1 + 2T(\frac{n}{2}) + C_2 n$
 divide recursion merge



Merge: 2 sorted arrays as in two finger algo



Complexity of merging $\Theta(n)$

Cn comparisons on each level,
 $\log n$ levels.

Space: had to create new elements (copies of arrays)

$\Rightarrow \Theta(n)$ auxiliary space

Note: In place sorting $\Rightarrow \Theta(1)$ auxiliary space

$\rightarrow$ For insertion sort, $\Theta(1)$ is temporary variable created when you make a swap.

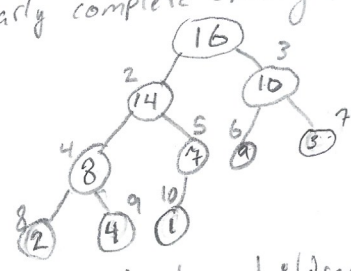
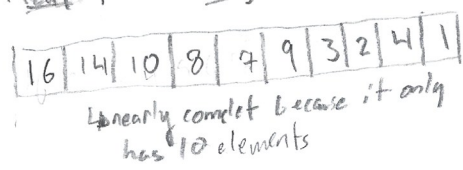
Time complexity: Cn per merge, $\log n$ merges

$\boxed{O(n \log n)}$

4) Priority Queue: Implements a set S of elements, each element associated w a key.

- Insert(S, v): insert element v into set S
- Max(S): return element of S w largest key
- extract-max(S): " " " " " " and remove from S
- increase-key(S, x, k): increase value of x 's key to new value k

Heap: An array visualized as a nearly complete binary tree



Heap as tree
 root of tree: first element ($i=1$)
 parent(i) = $i/2$
 left(i) = $2i$ right(i) = $2i+1$

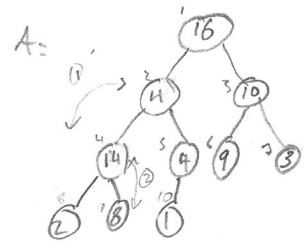
Max heap property: The key of a node $\geq$ the key of its children (there is also a min heap property)

How do we maintain this property?
 Little value
 This is a property you want to maintain as you modify the heap

Heap operations:

- build-max-heap: produces a max heap from an unordered array.
- max-heapify: correct a single violation of the heap property in a subtree's root. (Key assure)
- max-heapify(A, i): ensure that trees rooted at left(i) and right(i) are max heaps.

max-heapify($A, 2$):
 heap size (A) = 10



- ① Exchange $A[2], A[4]$
- Call Max-Heapify($A, 4$)
- ① Exchange $A[4], A[9]$

Exchanges called node w the largest of its children
 Simple, could run all the way down to the leaves.
 $\rightarrow O(\log n)$

Convert $A[1 \dots n]$ into a max heap:

Build-max-heap(A):
 for $i = n/2$ down to 1
 do max-heapify(A, i)

$n/2 + 1 \rightarrow n$ are all leaves
 Build your way up the tree so that you know everything below it satisfies max heap property.
 Kind of dragging biggest elements up.

Complexity?

Observe: Max-heapify takes $O(k)$ for nodes that are one level above leaves.
 and in general, order k time for nodes k levels above leaves.
 Look like $(n \log n)$ But...

Total work: $n/4(1c) + n/8(2c) + n/16(3c) + \dots + 1(\log n c)$
 $= c 2^k (\frac{1}{2^0} + \frac{2}{2^1} + \frac{3}{2^2} + \dots + \frac{(k+1)}{2^k})$
 $= 2^k \rightarrow (n/4) \rightarrow O(n)$
 convergent series = bounded by a constant.

Set $\frac{n}{4} = 2^k$
 constants go away for asymptotics

Algorithm for sorting $O(n \log n)$

- 1) Build-max-heap from unordered array
- 2) Find max element $A[1]$
- 3) Swap elements $A[1]$ and $A[n]$, now max is at end of array.
- 4) Discard node n from heap, decrementing heap size
- 5) New root may violate max heap property, but children are max heap so run max-heapify (shift down $\log n$ rows)

build up max heap
 Move biggest element to end & discard
 now put subtrees at all max heaps so run max-heapify
 repeat 2-5 n times
 total $n \times (n \log n)$

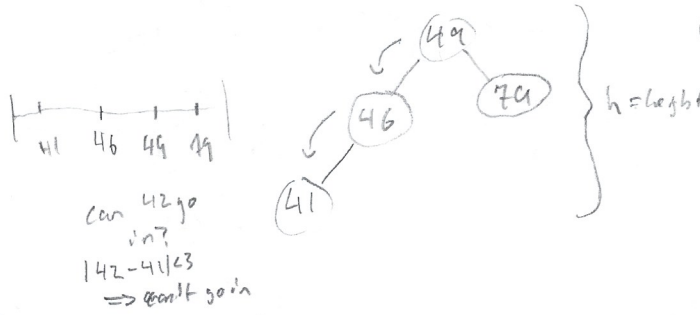
5) Scheduling and Binary Search Tree

Want: Fast insertion into a sorted array
 no data structure we have gone over can do this.
 Binary Search Tree; Invariant: $\forall$ node x , if y is in the left subtree of $x \Rightarrow key(y) \leq key(x)$, if y is in the right subtree of $x \Rightarrow key(y) \geq key(x)$

array can search but shift in $O(n)$
 Impl/ls: shift pointers
 can't bin. search

node: $key(x)$
 Pointers: $parent(x)$, $leftchild(x)$, $rightchild(x)$
 (unlike a heap)

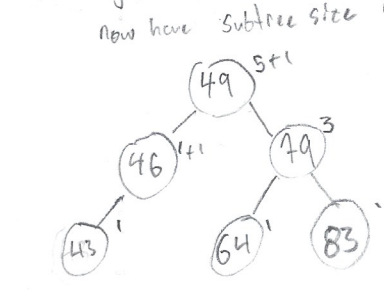
Application: Airport runway reservation system.
 - Airport has 1 runway
 - Reservations for future landings
 - Reserve request specifies landing time t .
 - Add t to set R if no other landings scheduled within k minutes (case $k=3$)



Find min(): go left until you hit a leaf $O(h)$
 next largest: $O(h)$ time.

* New requirement: How many planes are scheduled to land at times $\leq t$

Augment BST structure.



now have subtree size associated w each node, note, insert/delete will modify size.
 insert 43, add one to 'size' of every node you touch along the way, and assign 1 to 43.
 delete more complicated
 * How many $\leq t$.
 1) Walk down to find desired time
 2) Add in nodes that are smaller
 3) Add subtree's size to the left. $\rightarrow$ whenever you go right, all nodes in subtree to the left are smaller

6) AVL trees on the topic of balancing

height of BST: length of longest path down to leaf,
 height of node: length of longest path from it down to leaf
 $h = \max(\text{height of left child, height of right child}) + 1$
 children of leaves, null, defined to have depth -1



AVL trees require heights of left & right children of every node to differ by at most ± 1
 AVL trees are balanced: worst case is when right subtree has height 1 more than left $\forall$ node.



$|h_l - h_r| \leq 1$
 $N_h = \min \# \text{ nodes in AVL tree w height } h$

$N_h = 1 + N_{h-1} + N_{h-2}$

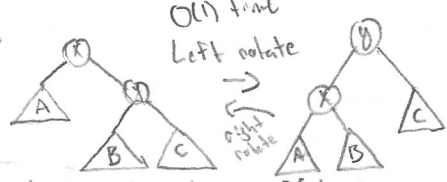
looks like Fibonacci, $N_h = F_h = \frac{\phi^h - \psi^h}{\sqrt{5}} \Rightarrow h < 1.44 \log n$

could take worst case n rotations to balance a tree.

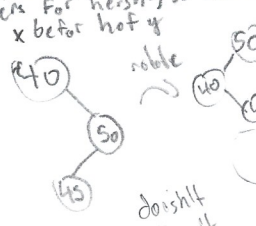
AVL Insert

1) simple BST insert
 2) Fix AVL property from the changed node up
 could unbalance every node above insertion. Rotation is $O(1)$ but could have to rotate up to $O(\log n)$

Rotation



AVL sort: Insert n items $O(n \log n)$
 2) in-order traversal $O(n)$



Since we right heavy: 2 cases
 if x 's right child is right heavy (RR)
 if x 's right child is left heavy (RL)

do shift fix it
 rotate so for then rotate

Good for insert/delete, min/max, successor/predecessor

L7

Comparison Model

- all input items are black boxes (ADT)
- only operations allowed are comparisons ($<, \leq, >, \geq, =$)
- time cost = #comparisons

Decision Tree: any comparison algorithm can be viewed as a tree of all possible comparisons and their outcomes and resulting answers.

For any particular n .

Decision tree

- internal node
- leaf
- root-to-leaf path length
- height of tree

algorithm

- binary decision
- Found answer
- algorithm execution
- running time.
- Worst case running time

Searching Lower Bound: n preprocessed items, finding a given item among them in comparison model requires $\Omega(\lg n)$ in worst case.

Why? decision tree is binary & must have $\geq n$ leaves, one for each answer. $\Rightarrow$ height $\geq \lg n$
could have multiple paths leading to the same answer

Sorting Lower Bound

Decision trees where answers (leaves) are possible sorted orders, for n items possible orders is $n!$
#leaves $\geq n! \Rightarrow$ height $\geq \lg(n!)$ now using Stirling's or sum of $\ln()$

we get: $\Omega(n \lg n)$

These analyses were both done by thinking about them as decision trees

RAM

Linear Time Sorting (integer sorting)

- Assume n keys sorting are integers in range $\{0, 1, \dots, k-1\}$ (and each fits in a word)
- Can do a lot more than comparisons
- For k (not too big) can sort in $O(n)$ time

Counting Sort

L = array of k empty lists

for j in range(n):

$L[\text{key}(A[j])].\text{append}(A[j])$

output = $[\]$

for i in range(k):

output.extend($L[i]$).

$O(k)$

$O(1) \rightarrow O(n)$

$O(L_i + 1) \rightarrow O(n + k)$

Radix Sort

• Imagine each integer in base b (break it into a bunch of columns)

$\Rightarrow$ # of digits = $d = \log_b k + 1$

• sort ints by least sig digit

• sort by most significant

example: sort in excel by the least important column first, then... finally the most important last
 $\hookrightarrow$ sorted by all in order of significance.

$\hookrightarrow$ each sort use counting sort by digit $\Rightarrow O(n + b)$

Total time: $O((n + b) \cdot d)$

minimized when $b = \Theta(n)$ $\hookrightarrow O((n + b) \log_b k)$
 $O(n \log n \log k)$

if $k \leq n^c$ then $O(n \log n)$

Hashing In Python, Dictionary: Abstract data type, maintain set of items, each w a key
 - insert(item), - delete(item), - search(key)
 ↳ return item w given key or report doesn't exist.
 ↳ overwrites an existing key
 $O(\log n)$ via AVL, but we can do better $\rightarrow O(1)$

Python: dict: $D[key] \sim$ search, $D[key] = val \sim$ insert, $del D[key] \sim$ delete
 $\rightarrow item = (key, value)$.

Motivation: docids, databases, compilers & interpreters, network router
 • substring search, string commonalities, file/directory synchronization, cryptography



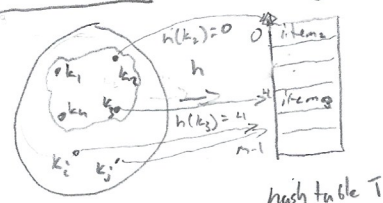
Simple Approach: Direct access table, store items in an array indexed by key

Problems: (1) Key may not be non-neg integers
 (2) Huge memory hog

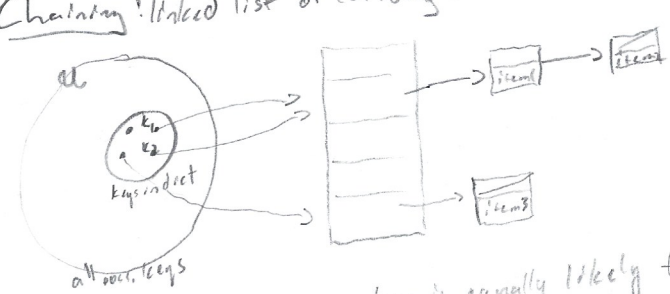
Solution to (1): prehash, - maps keys to non-neg integers

- in theory, keys are finite and discrete (string of bits)
 - in python, $hash(x)$ is the prehash of x , $hash(' \backslash 0 B ') = hash(' \backslash 0 \backslash 0 C ') = 64$
 Ideally: $hash(x) = hash(y) \Leftrightarrow x = y$
 -- hash: if you want to tell it what to do
 ↳ pre hash should not change our time.

Solution to (2): hashing: reduce U of all keys (integers) down to reasonable size, m , for the table
 $h: U \rightarrow \{0, 1, \dots, m-1\}$
 if $h(k_i) = h(k_j)$ but $k_i \neq k_j$
 $\Rightarrow$ collision
 ↳ guaranteed to happen by pigeonhole principle.



Chaining: linked list of colliding elements in each slot of hash table.



• worst case everything goes to 1 slot and you just end up w 1 big linked list.
 $\Rightarrow$ worst case $O(n)$

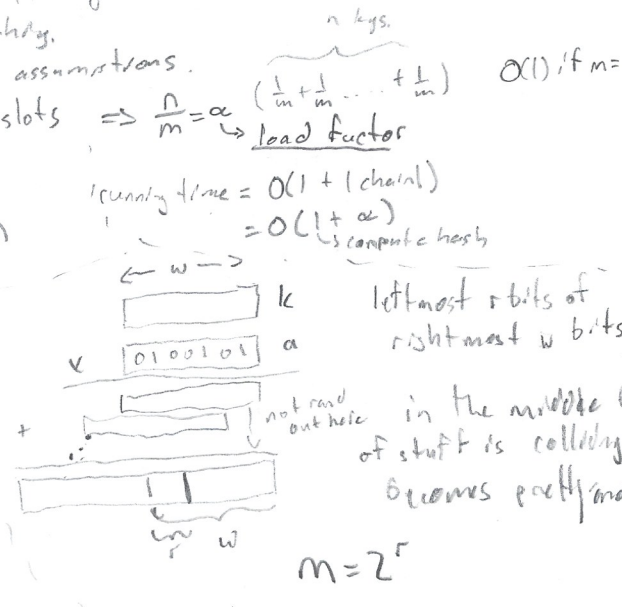
Simple Uniform Hashing: each key is equally likely to be hashed to any slot of the table
 independent of whole other key hashing.
 ↳ uniformly random and independent assumptions.

Analysis: expected length of a chain, n keys stored in table w m slots $\Rightarrow \frac{n}{m} = \alpha$ load factor

Hash Functions

- 1) Division method: $h(k) = k \bmod m$ (bad if k & m have common factor)
 ↳ overall bad still
- 2) Multiplication method: $h(k) = [(a \cdot k) \bmod 2^w] \gg (w-r)$
 ↳ w bit $\rightarrow$ word size w bits, r bit

(Good): 3) Universal Hashing $h(k) = [(ak + b) \bmod P] \bmod m$
 random $a \in \{0, 1, \dots, P-1\}$, prime $P > 1M$ $\rightarrow$ bring # to $[0, m-1]$
 For worst case keys $k_1 \neq k_2: P, \{h(k_1) = h(k_2)\} = 1/m$
 ↳ collision worst case



$m = 2^r$

9) How to choose m ? want $m = O(n)$, $\alpha = O(1)$ $\#$ scale m w n *

Idea: start w small m , ex $m=8$, then grow/shrink as necessary.

If $n > m$, grow table: $m \rightarrow m'$, reallocate memory; make table size m' , build new hash h'
 rehash: for item in T : $\left. \begin{matrix} T'.insert(item) \\ \end{matrix} \right\} O(m+n)$ visit every key in old table every item in linked lists new table

how big do we make m ?

double it $\rightarrow m' = 2m$, only have to rebuild every $2m$

$= O(1+2+4+8+\dots+n) = O(n)$ operation is not constant time?

Amortization: operation takes " $T(n)$ amortized" $\Rightarrow$ if k operations, take $\leq k \cdot T(n)$ time
 nothing as " $T(n)$ on average" where average over all operations

Table doubling, k inserts, take $O(k)$ time $\rightarrow O(1)$ inserts.
 only care about total not time of each step for most operations.

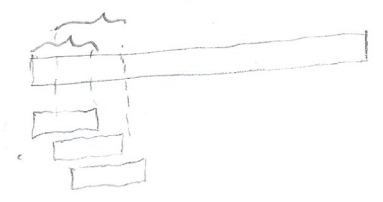
Also: k inserts & deletes take $O(k)$

Deletion if $m = \frac{n}{4} \rightarrow$ shrink to $\frac{n}{2}$, $\Rightarrow$ amortized time $= O(1)$
 $m = 8 \ 9 \ 9 \ 8 \ 8$
 $m = 8 \ 8 \ 16 \ 16 \ 8$ } alternate by and forth w linear time for each operation.
 prevents what would happen if shrink to half

String Matching given 2 strings s & t ; does s occur as a substring of t ?
 simple would be loop through + checking each string of t (len(s))
 expensive to compare strings want to view hash of these strings.

$$O(|s| \cdot (|t| - |s| + 1)) = O(|s| \cdot |t|)$$

Rolling Hash: computing hash of each string wouldn't save you any time
 but the next string differs by only 1 character.



$r.append(c)$: add char c to end of r
 $r.skip(c)$: delete first char of r (assuming its c)

r maintains a string x - r maintains a string x

Karp-Rabin Algorithm: $rs()$ rolling hash, $rt()$ rolling hash

for c in s : $rs.append(c)$
 for i in $1 : \text{len}(s)$:
 $rt.append(c)$
 if $rs() == rt()$: ...
 for i in range($\text{len}(s), \text{len}(t)$):
 $rt.skip(t[i - \text{len}(s)])$
 $rt.append(t[i])$
 if $rs() == rt()$
 check whether $s == t[i - \text{len}(s) : i + 1]$
 yes: if equal: Found match
 else: happens w prob $\leq \frac{1}{|s|}$
 $\Rightarrow O(1)$ expected time.
 $O(|s| + |t| + \# \text{ matches} \cdot |s|)$ expected time

How do we implement? which hash?

division method: $h(k) = k \bmod m$ random prime $\geq |s|$
 treating string x as multi/digit #, u , in base a $\rightarrow$ alphabet size

look at $r.append(c)$:

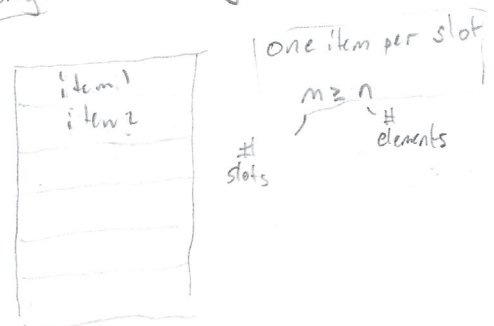


$u \rightarrow u \cdot a + \text{ord}(c)$
 $r \rightarrow r \cdot a + \text{ord}(c) \bmod m$
 $r.skip$ $u \rightarrow u - c \cdot a$

how does it change when adding a single character.

Simplest and most important thing to remember $\rightarrow$ open addressing

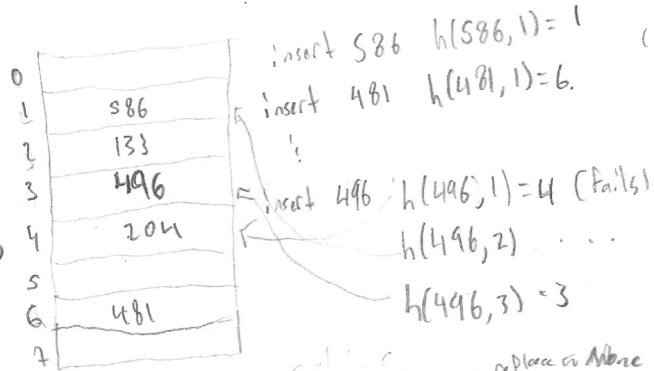
Open Addressing: no chaining



Probing: Hash function specifies order of slots to probe for a key (for insert/search/delete)

$$h: \mathcal{U} \times \{0, 1, \dots, m-1\} \rightarrow \{0, 1, \dots, m-1\}$$

try 1: $h(k, 1)$, try 2: $h(k, 2)$, ..., try $m-1$: $h(k, m-1)$
 arbitrary key k trial count try $m-1$
 $\rightarrow$ this should be a permutation of $0, 1, \dots, m-1$
 (for a specific key k) use all slots in hash table.



Now if I delete 586, a search for 496 would fail incorrectly.

searching? None = empty slot (Flag)
 Insert: keep probing till empty slot, then insert.
 Search(k): As long as slots encountered are occupied by keys $\neq k$. Keep probing until you encounter k or find an empty slot.

Delete: replace deleted item w/ different flag DeleteMe (different from None, prevents incorrect for)

$\Rightarrow$ Insert treats DeleteMe the same as None
 $\Rightarrow$ Search keep going (different than None)

Probing Strategies

Linear Probing: $h(k, i) = (h'(k) + i) \bmod n$ ordinary hash fun
 , satisfying permutation, too much cluster.
 $\rightarrow$ Clusters: consecutive groups of occupied slots, which keep getting longer.
 $\rightarrow$ causes you to lose average constant time lookup.

Double Hashing: $h(k, i) = (h_1(k) + i h_2(k)) \bmod m$

$\rightarrow$ IF $h_2(k)$ is relatively prime to $m \rightarrow$ permutation.
 $\rightarrow m = 2^r$, $h_2(k) \forall k$ is odd, construct h_2 to only produce odd #'s.
 adding a different amount with second hash helps separate keys.

LS

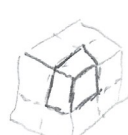
13 Graph Search | 'Explore' a graph.

Recall: graph = $G = (V, E)$, V set of vertices, E set of edges
 $E = \{v, w\}$ unordered pairs, or directed graph $e = (v, w)$ ordered pairs

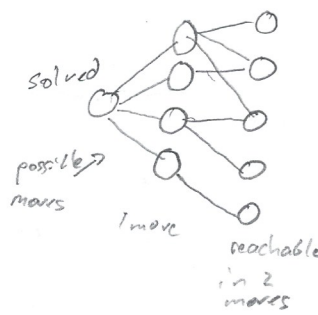
Applications: • web crawling • social networking • network broadcast • garbage collection
 • model checking • checking math conjectures • solving puzzles & games.

2x2x2 Configuration graph

• vertex for each possible state of the cube
 • edge for each possible move
 2x2x2: 11
 3x3x3: 20
 $n \times n \times n: \theta(\frac{n^2}{\lg n})$



vertices = $8! = 40320$
 3 directions for each permute the 8



Graph Representation:

Adjacency list:
 • Array Adj of $|V|$ linked lists
 • $\forall$ vertex $u \in V$, Adj[u] stores its neighbors
 $\hookrightarrow$ set $\{v \in V \mid (u, v) \in E\}$
 ex graph 2) Adj[b] = {a, c}, Adj[a] = {c}, Adj[c] = {b}
 implemented as:
 Adj:

b	a	c
a	c	
c	b	

 or

a	c
c	b

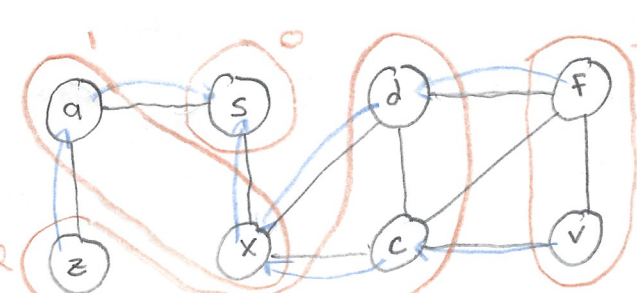
 better if discussing multiple graphs

Object oriented fashion:
 $v.neighbors = Adj[v]$

Implicit representation:
 - Adj(u) is a function } not stored just called when you want to compute } good for sub's cube
 - v.neighbors() is a method } less space

Breadth-First Search

- visit all nodes reachable from given $s \in V$ (vertex) $\hookrightarrow$ connected component.
- achieve $O(V+E)$ time
- look at nodes reachable in 0 moves, 1 move, 2 moves, ...
 $\{s\}$ Adj[s]
- careful to avoid duplicates:



parents always point all the way back to s.
 Forms a tree, w/ root s.
 shortest path back to s, (next page)

at 3, next = $\emptyset$
 $\hookrightarrow$ that the end frontier-shoot
 $\hookrightarrow$ frontier becomes $\emptyset$
 $\hookrightarrow$ while ends.

BFS(s, Adj):
 level = $\{s: \emptyset\}$ # given vertex and its adjacent
 parent = $\{s: None\}$ # $\{x: i\}$ can reach x in i moves
 i = 1
 frontier = [s] # what we've just reached in i
 while frontier:
 next = []
 for u in frontier: # scan frontier
 for v in Adj[u]:
 if v not in level:
 level[v] = i
 parent[v] = u
 next.append(v)
 frontier = next
 i += 1 # track current level

Shortest paths

- $v \leftarrow \text{parent}[v]$
- $\leftarrow \text{parent}[\text{parent}[v]]$
- $\leftarrow \dots$
- $\leftarrow s$

is a shortest path from s to v of length $\text{level}[v]$

↳ may not be unique.

↳ this is interesting, level is shortest # steps to get to root
↳ to get there just follow parents $\text{level}[v]$ times.

Depth first search

good example Recitation 14:20:00.

1) $\text{parent} = \{s: \text{none}\}$

DFS-Visit($v, \text{Adj}[s]$):

for v_i in $\text{Adj}[s]$:

if v not in parent :

$\text{parent}[v] = s$

DFS-Visit($v, \text{Adj}[v]$)

iterate over outgoing edges

- recursively explore the graph, backtracking as necessary

- careful not to repeat vertices

↳ are you in parent dictionary already?

↳ if you visit something you've already seen, just skip it.

- will only visit vertices reachable from s .

↳ not necessarily the whole graph.

2) check each vertex, visit everything reachable from it

↳ check next vertex visit everything reachable that hasn't already been visited.

Ensures you search entire graph, because you check each starting point
↳ finds all strongly connected components of graph.

2) DFS(v, Adj):

$\text{parent} = \{\}$

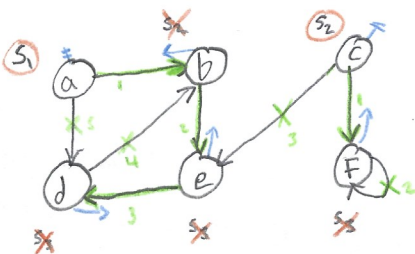
for s in V :

if s not in parent :

$\text{parent}[s] = \text{None}$

DFS-Visit($v, \text{Adj}[s]$)

iterate over choices of
if we don't know where
start is not strongly
or disconnected.



Analysis $\Theta(V+E)$, linear time.

- visit each vertex once in DFS alone $O(V)$

- DFS-Visit(s, v) called at most once per vertex v
↳ pay $|\text{Adj}[v]| \Rightarrow O(\sum_{v \in V} |\text{Adj}[v]|) = O(E)$ → directed.
2E in undirected.

Edge Classification

- tree edge (parent pointer): visit new vertex via edge, green edges in pic
- forward edges: node to descendant in tree ex $a \rightarrow d$
- backward edges: node to ancestor in tree ex $d \rightarrow b$
- cross edges: between two non-ancestor related subtrees.

In an undirected graph, can only have tree edges, and backward edges.

Cycle detection: G has a cycle $\Leftrightarrow$ DFS has a back edge.

↳ find cycle; start at backedge and follow tree edges.

- can't have forward edge cause it would have been found as a back edge ex $d \rightarrow a$
- can't have cross edge, edge would have been explored when at that node initially.

generally need to track
nodes discovered, but
still need to visit, Q
↳ use a queue

track and query if node
has been visited already
↳ hash table

Topological Sort Job scheduling example: given dir. acyclic graph (DAG) order vertices st. all edges point low $\rightarrow$ high.
 Use: Topological Sort $\rightarrow$ run DFS and return reverse of finishing times for vertices

Example intuition: Rec 14 4:00:00

Shortest Path $G(E, V, W)$
 G : graph, E : edges, V : vertices, W : weights $W: E \rightarrow \mathbb{R}$

path $p: \langle v_0, v_1, \dots, v_k \rangle$
 $(v_i, v_{i+1}) \in E$ for $0 \leq i < k$

$W(p) = \sum_{i=0}^{k-1} w(v_i, v_{i+1})$, find p w min weight
 Exponential # of paths, algs need to be smart...

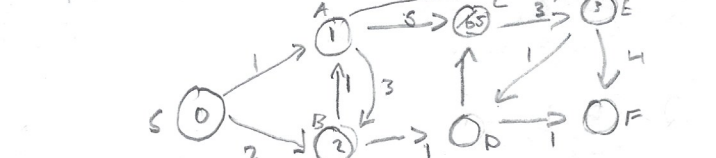
Weighted Graphs

$v_0 \xrightarrow{p} v_k$, (v_0) is a path from $v_0 \rightarrow v_0$ w weight 0.

Shortest path weight from u to v as:
 $\delta(u, v) = \begin{cases} \min \{ W(p) : u \xrightarrow{p} v \} & \exists \text{ any such path.} \\ \infty & \text{otherwise} \end{cases}$

$\rightarrow$ min weight if path exists, ∞ otherwise.

Think of all weights as delay as since you haven't found any get, want to reduce them down.



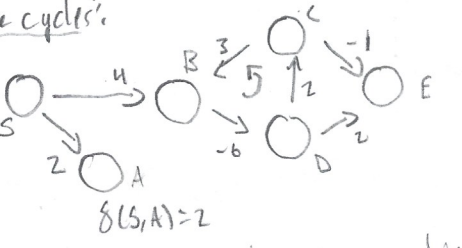
$d(v)$: current weight (# decide node)
 $\pi[v]$: predecessor on current best path to v
 $\rightarrow \pi[s] = \text{nil}$

When $d(v) = \delta(s, v) \Rightarrow \text{done}$

Negative Weights

why? course talks, social networks etc...

Negative cycles:
 $\delta(s, s) = 0$
 $\rightarrow$ can't get into neg cycle

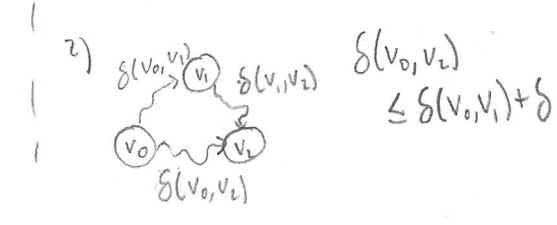
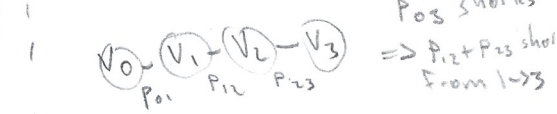


For B, C, D, E, can keep running through $B \rightarrow D \rightarrow C$ and keep getting more negative
 $\rightarrow$ need a termination condition that accounts for

General Structure: shortest path algs. (no negative cycles)
 Initialize, for $u \in V$, $d[u] \leftarrow \infty$, $\pi[u] = \text{NIL}$, $d[s] \leftarrow 0$
 Repeat: select some edge (u, v) [somehow]
 "Relax" edge (u, v) : if $d[v] > d[u] + w(u, v)$:
 $d[v] = d[u] + w(u, v)$
 $\pi[v] \leftarrow u$
 until all edges have $d[v] \leq d[u] + w(u, v)$

Optimal substructure.

1) Subpaths of a shortest path, are shortest paths.



Pos shortest $\Rightarrow P_{12} + P_{23}$ shorter from 1 to 3

$\delta(v_0, v_2) \leq \delta(v_0, v_1) + \delta(v_1, v_2)$

16 Dijkstra

Relaxation is safe: maintains invariant $d[v] \geq \delta(s, v)$

Relax(u, v, w)

if $d[v] > d[u] + w(u, v)$

$d[v] = d[u] + w(u, v)$

$\pi[v] = u$

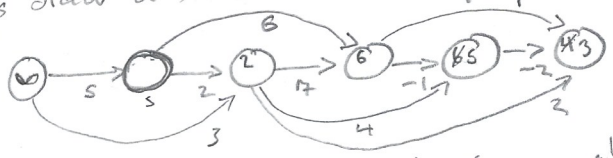
Lemma: The relaxation operation maintains the invariant that $d[v] \geq \delta(s, v) \forall v \in V$

For DAG (no res cycles)

use topological sort (Whenever DAG, try topol. sort is a great option)

- 1) Topological sort the DAG. Path from u to v implies u is before v in the ordering
- 2) One pass over vertices in topologically sorted order, relax each edge that leaves each vertex

$O(V+E)$ Can always draw a sorted DAG linearly, everything to left of source will be infinity



go node-by-node, only moving to next once all vertices have been relaxed

or, No neg cycles

Dijkstra(G, w, s)

Initialize(G, s)

While $Q \neq \emptyset$:

$u \leftarrow \text{Extract min}(Q)$ % delete u from Q

$S' \leftarrow S' \cup \{u\}$

for each vertex $v \in \text{Adj}[u]$

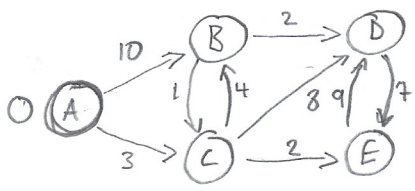
Relax(u, v, w)

$S \leftarrow \emptyset$ $Q \leftarrow V[G]$

every vertex you can reach from u

good to read this proof

priority queue
priorities are d values
first is s cause $d(s)=0$



$S = \{ \}$

$Q = \{A, B, C, D, E\}$
0 ∞ ∞ ∞ ∞

$S = \{A\}$

$Q = \{B, C, D, E\}$
10 3 ∞ ∞

$S = \{A, C\}$

$Q = \{B, D, E\}$
7 11 5

$S = \{A, C, E\}$

$Q = \{B, D\}$
7 11

Complexity

$O(V)$ inserts into priority Queue

$O(V)$ Extract-min ops.

$O(E)$ Decrease key ops
↳ Relax()

Implement queues as:

Arrays

$O(V)$ For extract min

$O(1)$ For decrease key

$\Rightarrow \theta(V \cdot V + E \cdot 1) \approx \theta(V^2)$

Binary Min-Heap

$O(\lg V)$ For ext. min

$O(\lg V)$ For decrease key

$\Rightarrow \theta(V \lg V + E \lg V)$

$\approx \theta(V \lg V + E)$

1 w Fibonacci he (61016)

17 Bellman-Ford

Generic S.P. Algo:

Used if neg cycles are a possibility

Initialize for $v \in V$ $d[v] \leftarrow \infty$, $\pi[v] \leftarrow nil$
 $d[s] \leftarrow 0$

Main: repeat select edge [somehow]
 Relax(u, v, w)

Until you can't relax anymore

Bellman-Ford (G, w, s) $d[s] = 0$, others ∞

Initialize()

for $i = 1$ to $|V| - 1$

for each edge $(u, v) \in E$

Relax(u, v, w)

for each edge $(u, v) \in E$:

if $d[v] > d[u] + w(u, v)$

then report -ve cycle exists.

this loop
 $O(E)$ →

would be right
 here if no -ve
 cycles →

check for
 -ve cycles.

then just relax
 every edge $|V|-1$ times

Relax(u, v, w):

if $d[v] > d[u] + w(u, v)$:

$d[v] = d[u] + w(u, v)$

$\pi[v] = u$

You relax every edge, in every pass.
 You make $|V|-1$ passes.

$O \rightarrow O \rightarrow O \rightarrow O \rightarrow O \rightarrow O \dots O(2^{n/2})$

2) Will not terminate if $\exists$ a -ve cycle
 reachable from the source.

Theorem: If $G(V, E)$ contains no -ve weight cycles, then:
 after B-F executes, $d[v] = \delta(s, v) \forall v \in V$.

Corollary: If a value $d[v]$ fails to converge after $|V|-1$ passes, then,
 $\exists$ a -ve weight cycle reachable from s .

Proofs in notes

you have topol. sorted vertices

guarantee to have $\delta(v_0, v_1)$ after first loop

then to have $\delta(v_0, v_2)$ after second loop

$\vdots$

$\delta(v_0, v_k)$ after k passes.

if $k = |V|$, get k after $|V|-1$ iterations

essentially expanding the frontier by 1 every time and will have
 found shortest simple path

other condition is to check that there hasn't been a negative cycle
 updating a bunch.

you will catch it because no paths should be updating after $|V|-1$ passes

order edges are relaxed each loop doesn't matter.

- No known algorithm for finding shortest simple path if -ve cycle exists,
- so is longest path

Lecture 18 | Speeding Up Dijkstra

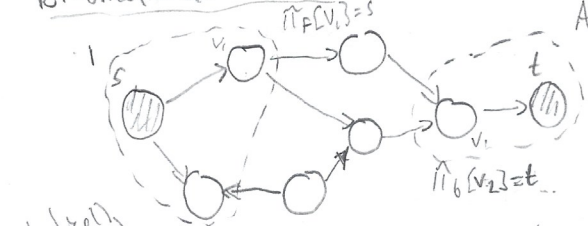
doesn't change worst case, improves "average" case.

Dijkstra: Initialize $d[s] = 0$ $d[u \neq s] = \infty$
 while $Q \neq \emptyset$
 do $u \leftarrow \text{Extract.Min}(Q)$
 for each vertex $v \in \text{Adj}[u]$
 do $\text{RELAX}(u, v, w)$

$Q \leftarrow V \setminus \{s\}$
 $c \rightarrow t$
 stop if $u = t$

DS needs edges able to be traversed both ways.

Bidirectional Search



Alternate: Forward search from s
 • backwards search from t (following edges backwards)

π_f : forward
 π_b : backward.

$d_f[u]$: distances for forward search
 $d_b[u]$: distances for backward search.

Priority Queues: Q_f : forward
 Q_b : backwards.

frontiers meeting

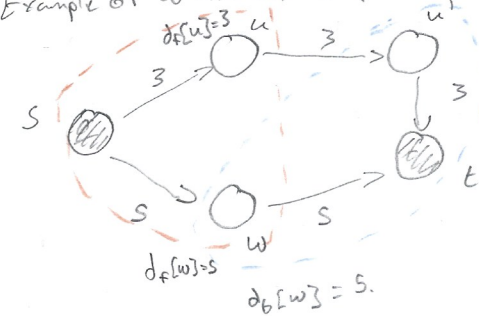
2 different priorities:
 $d_f[s] = 0$ $d_b[t] = 0$

Termination Condition: Some vertex has been processed both in the forward and backward search and has been deleted from Q_f and Q_b .

Shortest Path: If w was processed first from both Q_f, Q_b
 find S.P. using π_f from s to w
 find S.P. using π_b from t to w (backwards).

Not quite: w may not be on S.P.

Example of w not on SP.



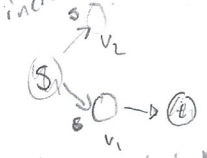
$s(s, x)$	find $d_f[x]$	Q_f	$s(t, x)$	find $d_b[x]$	Q_b
1) $\{s\}$	$\{u, u', w, t\}$	1) $\{t\}$	$\{u, u', w, s\}$		
3) $\{s, u\}$	$\{u', w, t\}$	4) $\{t, u'\}$	$\{u, w, s\}$		
5) $\{s, u, w\}$	$\{u', t\}$	6) $\{t, u', w\}$	$\{u, s\}$		

Terminate! w processed for both forwards and backwards.
 But incorrect length of 10 (should be 6)

Do a little more work...

Find an x (possibly different from w) w minimum value of $d_f[x] + d_b[x]$ (look in Q_f and Q_b)
 $d_f[u] + d_b[u] = 9$ ✓
 $d_f[u'] + d_b[u'] = 9$ ✓

increase potential (uphill)



how to get it to guess v_1 instead of v_2 .

decrease potential (downhill)

Goal-Directed Search

Modify edge weights with potential functions:

$$\bar{w}(u, v) = w(u, v) - \lambda(u) + \lambda(v)$$

Correctness: $\bar{w}(p) = w(p) - \lambda_t(s) + \lambda_t(t)$ - any path from $s \rightarrow t$ shifted by same amount. shortest path before will still be shortest path.

landmark $l \in V$ } $\lambda_t(u) = \delta(u, l) - \delta(t, l)$

precompute $\delta(u, l)$
 this increases speed, have a point you know you need to pass through.

1) Dynamic Programming (DP) - general powerful algo. design technique.
 DP $\approx$ careful brute force, subproblems and reuse.

Fibonacci Numbers: goal: compute F_n
 $F_1 = F_2 = 1$ $F_n = F_{n-1} + F_{n-2}$

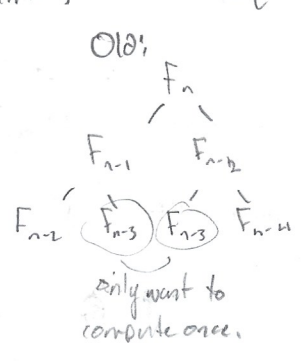
Naive recursive algo: $Fib(n)$:
 if $n \leq 2$: $F = 1$
 else: $F = fib(n-1) + fib(n-2)$
 return F

Exponential: $T(n) = T(n-1) + T(n-2) + \Theta(1)$
 $\begin{cases} \geq F_n \approx e^n \\ T(n) \geq 2 \cdot T(n-2) \\ = \Theta(2^{n/2}) \end{cases}$

Memoized DP algorithm: whenever you compute a fibonacci number, store it in a dictionary.

memo = {}
 $fib(n)$:
 if n in memo: return memo[n]
 if $n \leq 2$: $F = 1$
 else: $F = fib(n-1) + fib(n-2)$
 memo[n] = F
 return F

same as above



$Fib(k)$ only recurses the first time it's called, $\forall k$

- memoized calls cost $\Theta(1)$
- number of non-memoized calls (first time you call $Fib(a)$ or $Fib(b)$) is n : $Fib(1), Fib(2), \dots, Fib(n)$
- non recursive work per call is constant. $\Rightarrow$ time = $\Theta(n)$ $\rightarrow$ n calls, each costing $\Theta(1)$

DP $\approx$ recursion + memorization.
 - memoize (remember, or write down rough work), and reuse solutions to subproblems that help solve the problem

$\Rightarrow \text{time} = \underbrace{\# \text{subproblems}}_n \cdot \underbrace{\text{time/subproblem}}_{\Theta(1)}$

$\rightarrow$ don't count memoized occurrences.

Bottom-Up DP algorithm:

$Fib = \{ \}$
 for k in range(1, $n+1$):
 if $k \leq 2$: $F = 1$
 else: $F = Fib[k-1] + Fib[k-2]$
 $Fib[k] = F$
 return $Fib[n]$

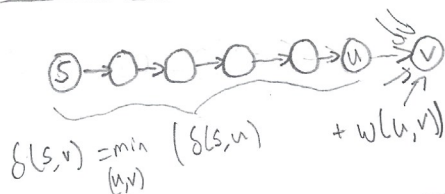
in general $\rightarrow$

- bottom up does exactly same computation as Memoized
 - topological sort of subproblem dependency DAG
-
- can often save space (only need to store last 2 values)

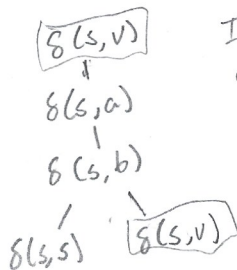
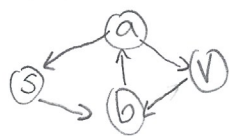
In all dynamic programs you can always go bottom-up like here in fibonacci

Shortest paths : $\delta(s, v) \forall v$

Guessing: don't know the answer? guess. try all guesses, take the best one.



DP $\approx$ guessing + recursion + memoization



Infinite time on graphs w cycles.

DAGs: $O(V+E)$

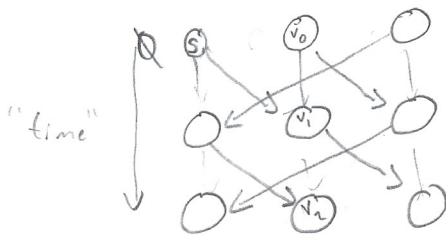
time for subprob $\delta(s, v) = \text{indegree}(v) + 1$
 $\Rightarrow \text{total time} = \sum_{v \in V} \text{indegree}(v) + 1 = O(E) + V$

subproblem dependencies should be acyclic (otherwise infinite)
 that's why we topologically sort DAG for bottom up.

If cycle need to make it acyclic. How?

Explod graph into multiple layers.

Make every edge go down to next layer



$\delta_k(s, v)$ = weight of shortest $s \rightarrow v$ path that uses $\leq k$ edges

$$\delta_k(s, v) = \min_{(u,v) \in E} (\delta_{k-1}(s, u) + w(u, v))$$

subproblems = V^2

20) 5 steps for dynamic programming

- 1) define subproblems $\rightarrow$ # subproblems
- 2) guess (part of solution) $\rightarrow$ # choices for guess
- 3) relate subproblem solutions (recurrence) $\rightarrow$ time per subproblem
- 4) recurse and memoize $\left\{ \begin{array}{l} \text{check subproblem recurrence} \\ \text{on build DP table bottom up} \end{array} \right\}$ is cyclic $\rightarrow$ has topological sort
- 5) solve original problem

cf Fibonacci:

- 1) F_k for $k=1, \dots, n$ n of them
- 2) no guess
- 3) $F_k = F_{k-1} + F_{k-2}$

Text Justification: split text into "good" lines

text = list of words

$$\text{badness}(i, j) = \begin{cases} \infty & \text{if they don't fit} \\ (\text{page width} - \text{total width})^3 & \text{otherwise} \end{cases}$$

use words $[i:j]$ as line.

- 1) subproblems: suffixes words $[i:]$

- # subproblems: n

- 2) guess: where to start 2nd line.

- # choices $\leq n - i = O(n)$

- 3) recurrence: $DP(i) = \min_{j \in \text{range}(i+1, n+1)} (DP(j) + \text{badness}(i, j))$

$\begin{matrix} i \\ j \end{matrix} \sim \sim \sim$ time per subprob
 $\sim \sim \sim$ $O(n)$
 $DP(n) = \emptyset \rightarrow$ no words after n .

- 4) topo. order: $i = n, n-1, \dots, 0$

total time $O(n^2)$

- 5) original problem

$DP(0)$

Parent pointers: remember which guess is best.

$$\text{parent}[i] = \text{argmin}(\dots) = j \text{ value.}$$

If you ignore the first i words, what's the best outcome?

Best outcome if you chose to leave $i+1$ or $i+2$ or $\dots$

Each DP you solve for the best case of the subproblems, then take the min of current badness + best case. Keep working your way back up to right answer.

21) Subproblems for strings/sequences.

- suffixes $x[i:]$ $\forall i$

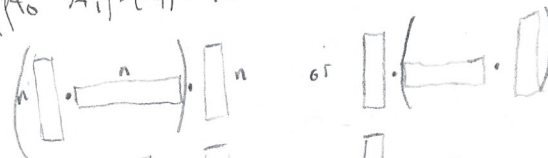
- prefixes $x[:i]$ $\forall i$

- substrings $x[i:j]$

never use both $\rightarrow$ induced subseq

Parenthesization: optimal evaluation of associative expression.

$(A_0 \cdot A_1) \cdot (\dots) \cdot A_n$ reduce cost of multiplication.



$O(n^2)$



$O(n)$

time/subproblem = $O(n)$

④ time = $O(n^3)$

topological order: increasing substring size

- 2) guess the outermost/last multiplication

$$(A_0 \dots A_{k-1}) \cdot (A_k \dots A_{j-1})$$

$$(A_0 \dots A_k) \cdot (A_{k+1} \dots A_{j-1})$$

not prefixes

can't use subseq

- 1) subproblem = optimal evaluation $A_i \dots A_{j-1}$

- 3) recurrence $DP(i, j) = \min_{k \in \text{range}(i+1, j)} (DP(i, k) + DP(k, j) + \text{cost of } (A_i:k) \cdot (A_k:j))$

Edit distance: given 2 strings x and y . what's the cheapest possible sequence of character edits to turn $x \rightarrow y$.
 char edits: insert c , delete c , replace $c \rightarrow c'$ each operation and c has different cost.

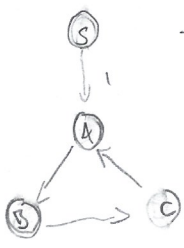
longest common subsequence: **HEROGLYPHOLOGH, MICHAELANGELO $\Rightarrow$ HELLO**

cost of insert/delete = 1
 replace = $\begin{cases} 0 & \text{if } c=c' \\ \infty & \text{else} \end{cases}$

① subproblem = edit distance on $x[i:]$ and $y[j:] \forall i, j \rightarrow \# \text{subproblems} = \theta(|x| \cdot |y|)$

② guess $\begin{matrix} x \\ y \end{matrix}$

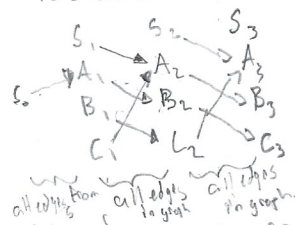
Cycles



The minimum distance calculation will call its own distance. This will lead to an infinite number of calls and never solved. Need to incorporate time.

$$\delta_k(s, v) = \delta_{k-1}(u, v) + w_{uv}$$

let's build this graph



longest simple path is $V-1$, so we need 3 layers

$$V' = (V-1) \cdot V + 1 \approx O(V^2)$$

$$E' = O(VE)$$

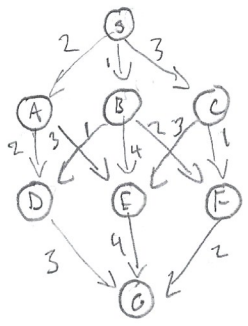
$$O(V' + E') = O(VE)$$

This is Belman Ford and is the actual intuition behind why it works

All dynamic programming problems can be thought of as a graph. Solve it smartly w/o building a graph

Optimal substructure, you are solving smaller parts of the problem and building on them (shortest paths). This is shown below. Find shortest path for subproblem then choose which one is best slightly bigger problem.

Example DP helped intuition from recursion.



base case is $\delta(s, s) = 0$

$$\delta(s, D) = \min \begin{pmatrix} \delta(s, A) + w_{AD} \\ \delta(s, B) + w_{BD} \end{pmatrix} = 2$$

$$\delta(s, E) = \min \begin{pmatrix} \delta(s, A) + w_{AE} \\ \delta(s, B) + w_{BE} \\ \delta(s, C) + w_{CE} \end{pmatrix} = 5$$

$$\delta(s, F) = \min \begin{pmatrix} \delta(s, B) + w_{BF} \\ \delta(s, C) + w_{CF} \end{pmatrix} = 3$$

$$\delta(s, G) = \min \begin{pmatrix} \delta(s, D) + w_{DG} \\ \delta(s, E) + w_{EG} \\ \delta(s, F) + w_{FG} \end{pmatrix} = 5$$

This is precisely what turns the exponential number of calculations into polynomial.

Just store your results after you compute them.

Key takeaway is that having $\delta(s, E)$ calculated prevents it needed to be recalculated again. Further having $\delta(s, A), \delta(s, B), \delta(s, F)$ calculated prevents them being recalculated for $\delta(s, D), \delta(s, E), \delta(s, F), \delta(s, G)$.