

CSE 332: Data Structures & Parallelism

Lecture 27: Topological Sort & B Trees & Wrap Up

Ruth Anderson

Autumn 2025

Administrative

- EX12 – P/NP, last exercise! Due Fri Dec 5 (last day of class)
 - O.k. to use late days on EX12, will close Mon Dec 8
- Lecture on Fri Dec 5
 - Final Exam Review Session (similar to midterm review during lecture)
- Final Exam, Thursday Dec 11, 12:30-2:20pm in KNE 120
 - Topics and Old exams on [exams page](#)
- Please fill out Course Evaluations!
 - For Lecture and Section

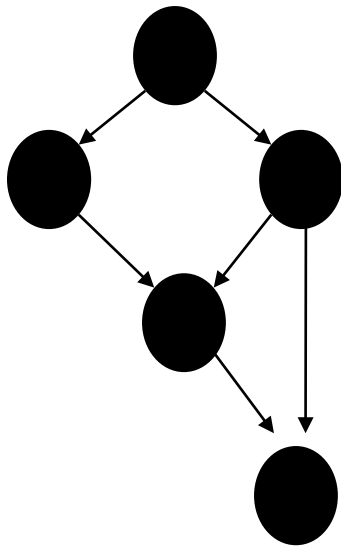
Today

- **Topological Sort:** a graph algorithm that can be used to schedule parallel tasks
- **B-Trees:** a dictionary that considers the memory hierarchy
- Wrap Up

Topological Sort

The DAG (Directed Acyclic Graph)

- A program execution using **fork** and **join** can be seen as a DAG
- [A DAG is a graph that is directed (edges have direction (arrows)), and those arrows do not create a cycle (ability to trace a path that starts and ends at the same node).]
 - **Nodes:** Pieces of work
 - **Edges:** Source must finish before destination starts



- A **fork** “ends a node” and makes two outgoing edges
 - New thread
 - Continuation of current thread
- A **join** “ends a node” and makes a node with two incoming edges
 - Node just ended
 - Last node of thread joined on

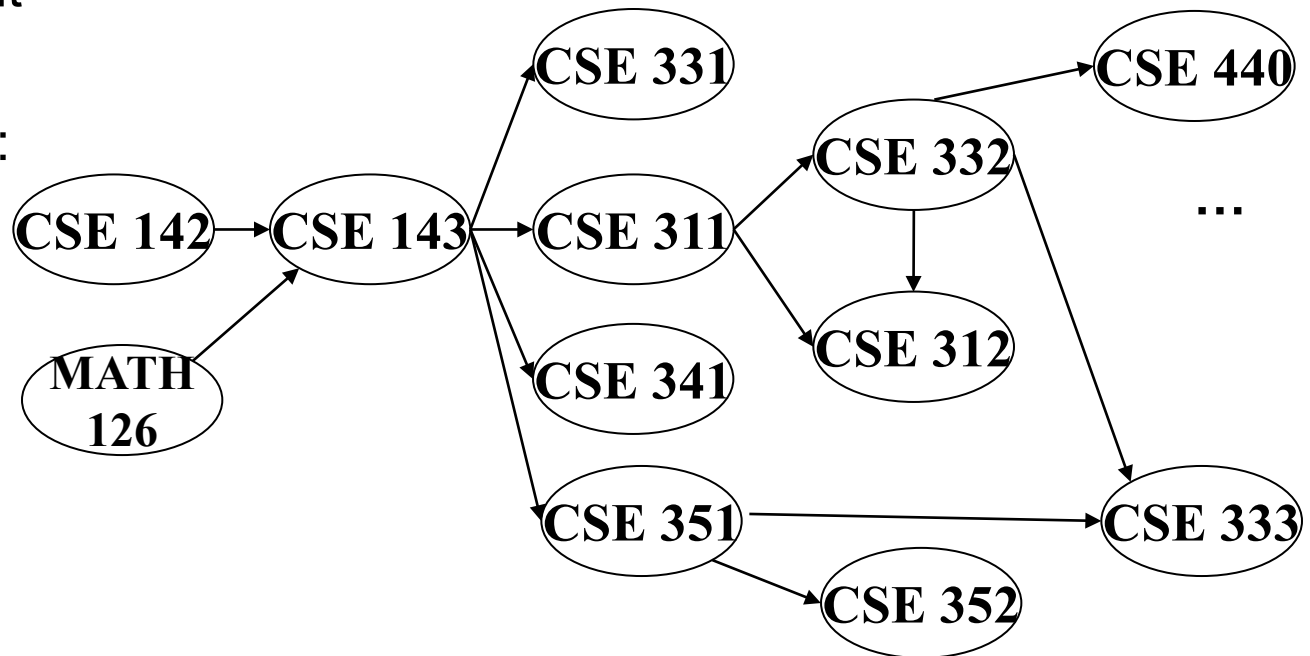
Connecting to performance

- Recall: T_P = running time if there are P processors available
- Work = T_1 = sum of run-time of all nodes in the DAG
 - That lonely processor does everything
 - **Any topological sort is a legal execution**
 - $O(n)$ for simple maps and reductions
- Span = T_∞ = sum of run-time of all nodes on the most-expensive path in the DAG
 - Note: costs are on the nodes not the edges
 - Our infinite army can do everything that is ready to be done, but still has to wait for earlier results
 - $O(\log n)$ for simple maps and reductions

Topological Sort

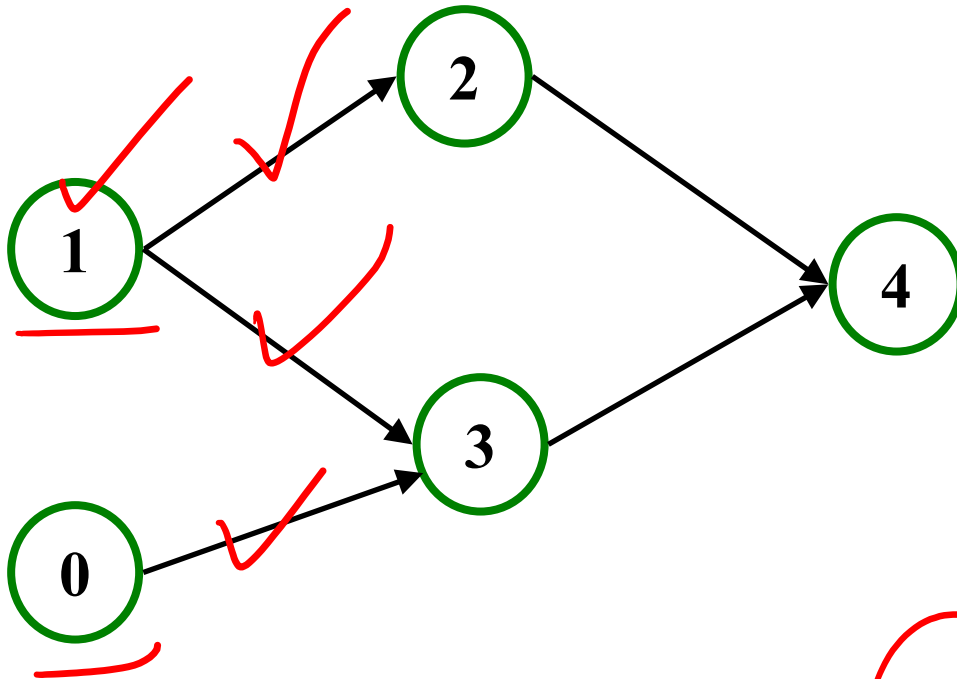
Problem: Given a DAG $G = (V, E)$, output all the vertices in order such that if no vertex appears before any other vertex that has an edge to it

Example input:



Example output:

142, 126, 143, 311, 331, 332, 312, 341, 351, 333, 440, 352



**Valid Topological
Sorts:**

0, 1, 2, 3, 4

0, 1, 3, 2, 4

1, 0, 3, 2, 4

1, 2, 0, 3, 4

1, 0, 2, 3, 4

Topological Sort Uses

- Scheduling instructions/tasks
- Figuring out how to finish your degree
- Computing the order in which to recompute cells in a spreadsheet
- Determining the order to compile files using a Makefile
- In general, taking a dependency graph and coming up with an order of execution

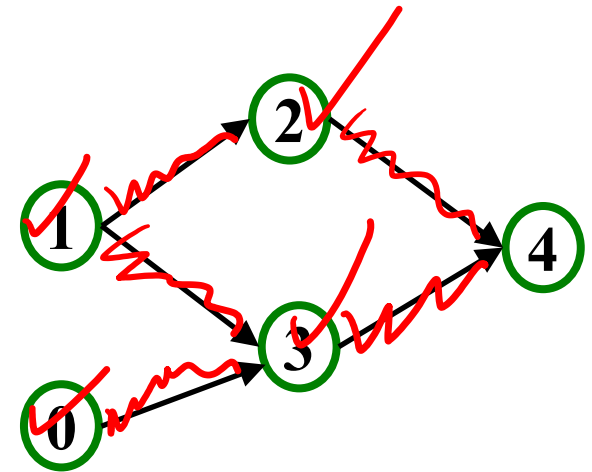
A First Algorithm for Topological Sort

1. Label ("mark") each vertex with its in-degree
 - Think "write in a field in the vertex"
 - Could also do this via a data structure (e.g., array) on the side
2. While there are vertices not yet output:
 - a) Choose a vertex v labeled with in-degree of 0
 - b) Output v and *conceptually* remove it from the graph
 - c) For each vertex w adjacent to v (i.e. w such that (v, w) in $\mathbf{E}$),

In-degree

0	0	✓	3	/
1	0	✓	2	→ 3
2	0	✓	4	/
3	2 0	✓	4	/
4	2 10	✓		

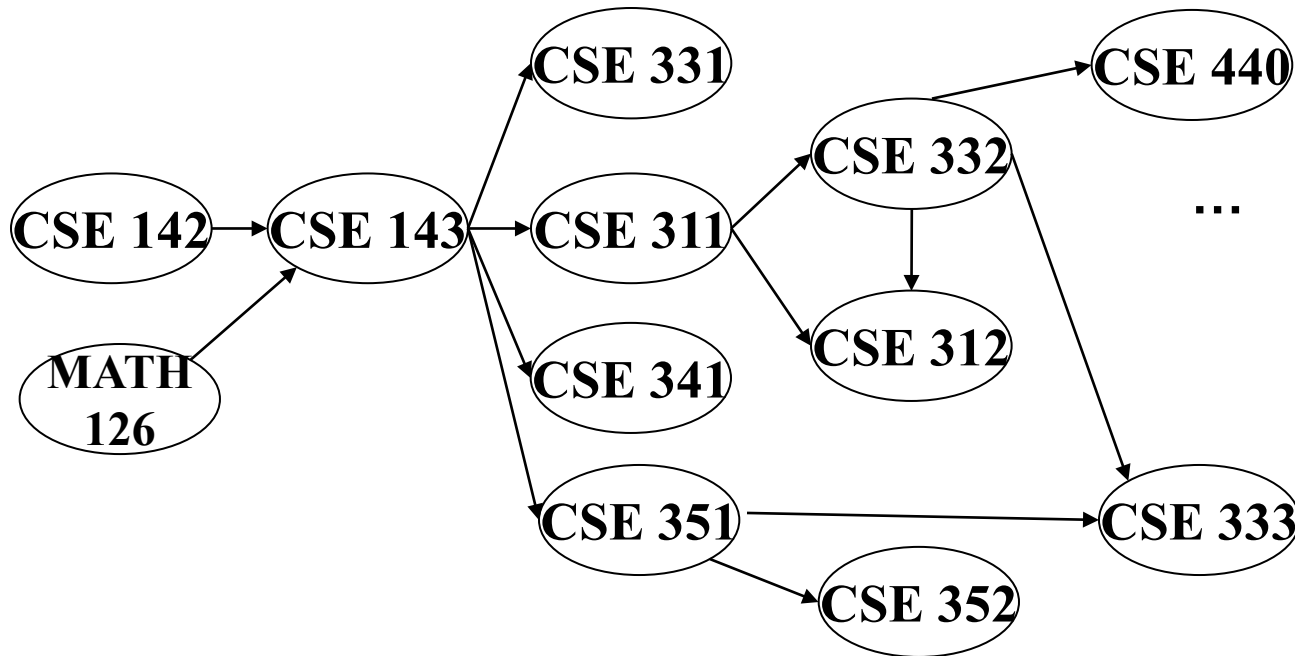
decrement the in-degree of w



Output: 0, 1, 2, 3, 4

Example

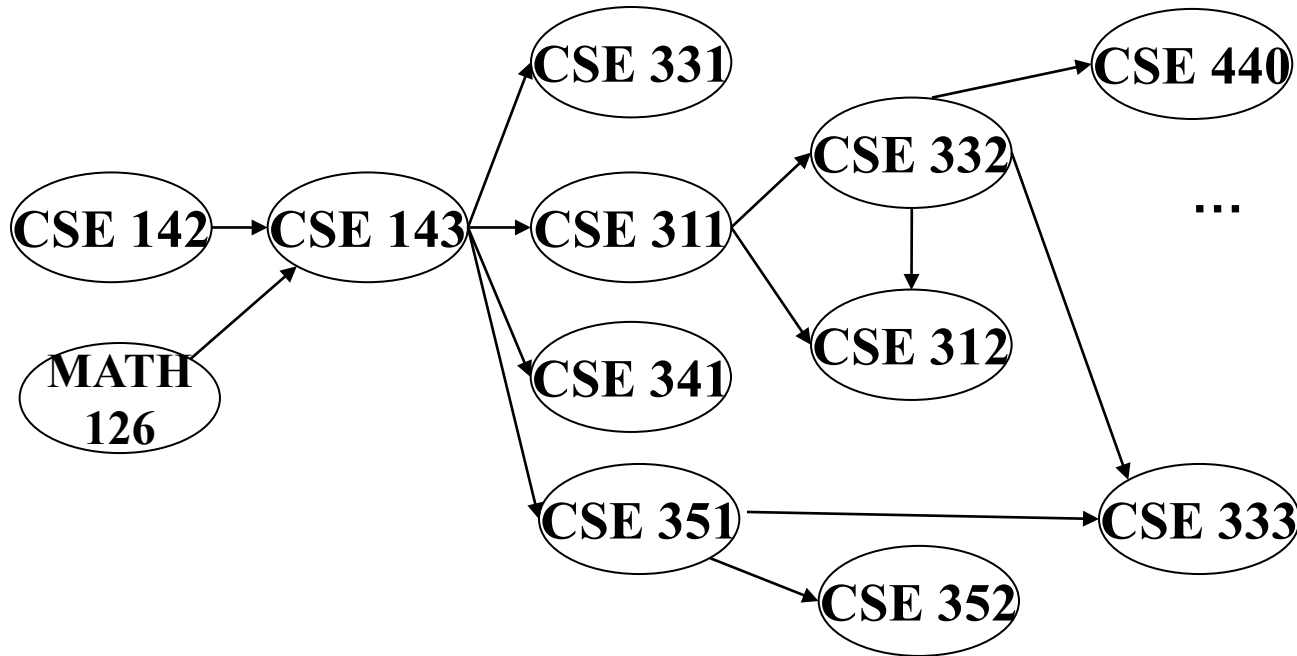
Output:



Node:	126	142	143	311	312	331	332	333	341	351	352	440
Removed?												
In-degree:	0	0	2	1	2	1	1	2	1	1	1	1

Example

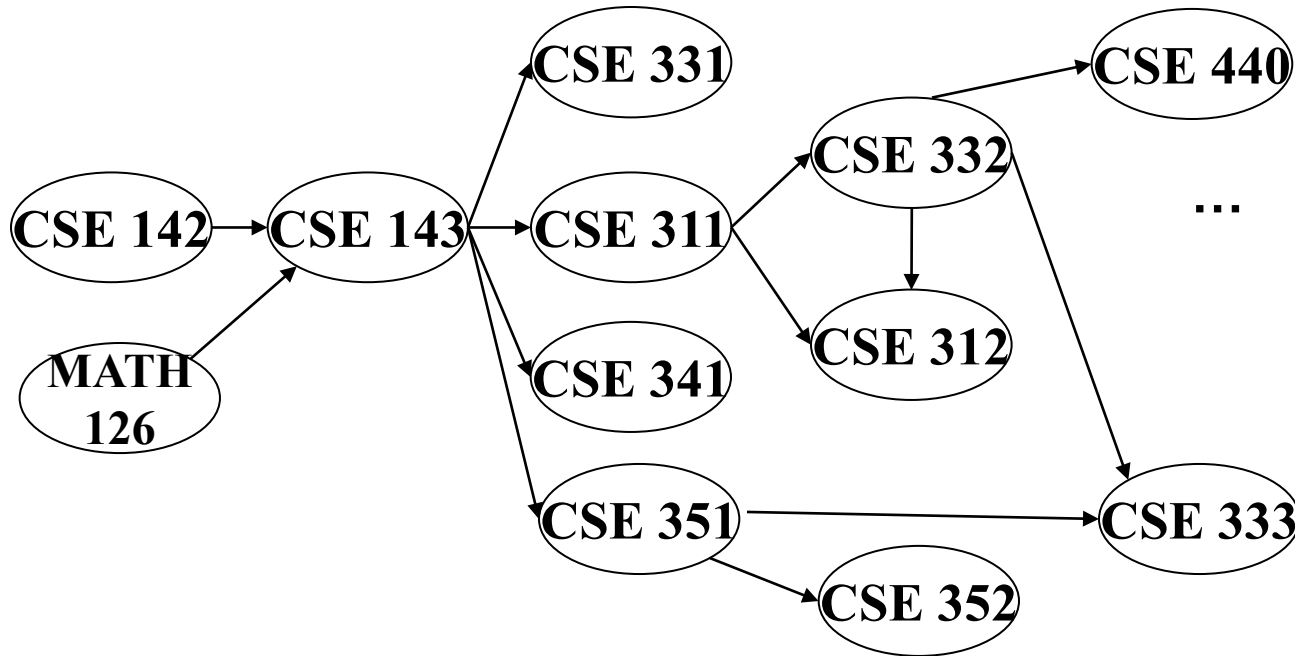
Output: 126



Node:	126	142	143	311	312	331	332	333	341	351	352	440
Removed?	x											
In-degree:	0	0	2	1	2	1	1	2	1	1	1	1
			1									

Example

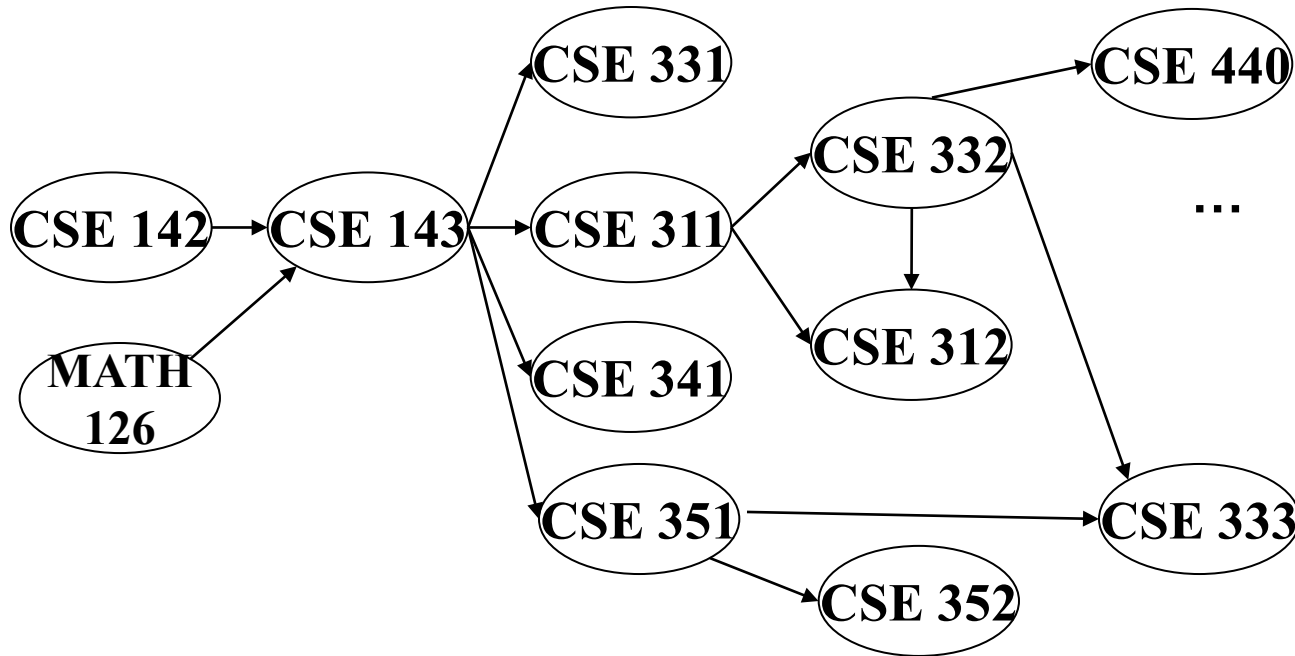
Output: 126
142



Node:	126	142	143	311	312	331	332	333	341	351	352	440
Removed?	x	x										
In-degree:	0	0	2	1	2	1	1	2	1	1	1	1
			1									
			0									

Example

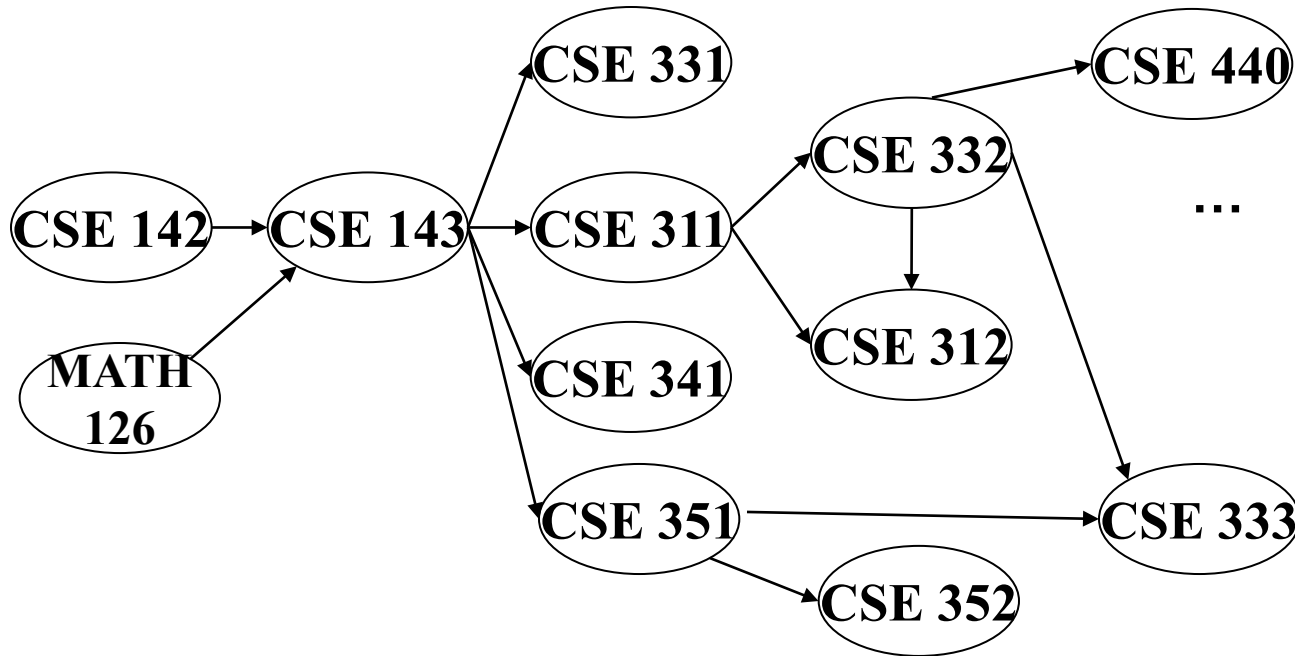
Output: 126
142
143



Node:	126	142	143	311	312	331	332	333	341	351	352	440
Removed?	x	x	x									
In-degree:	0	0	2	1	2	1	1	2	1	1	1	1
			1	0		0			0	0		
			0									

Example

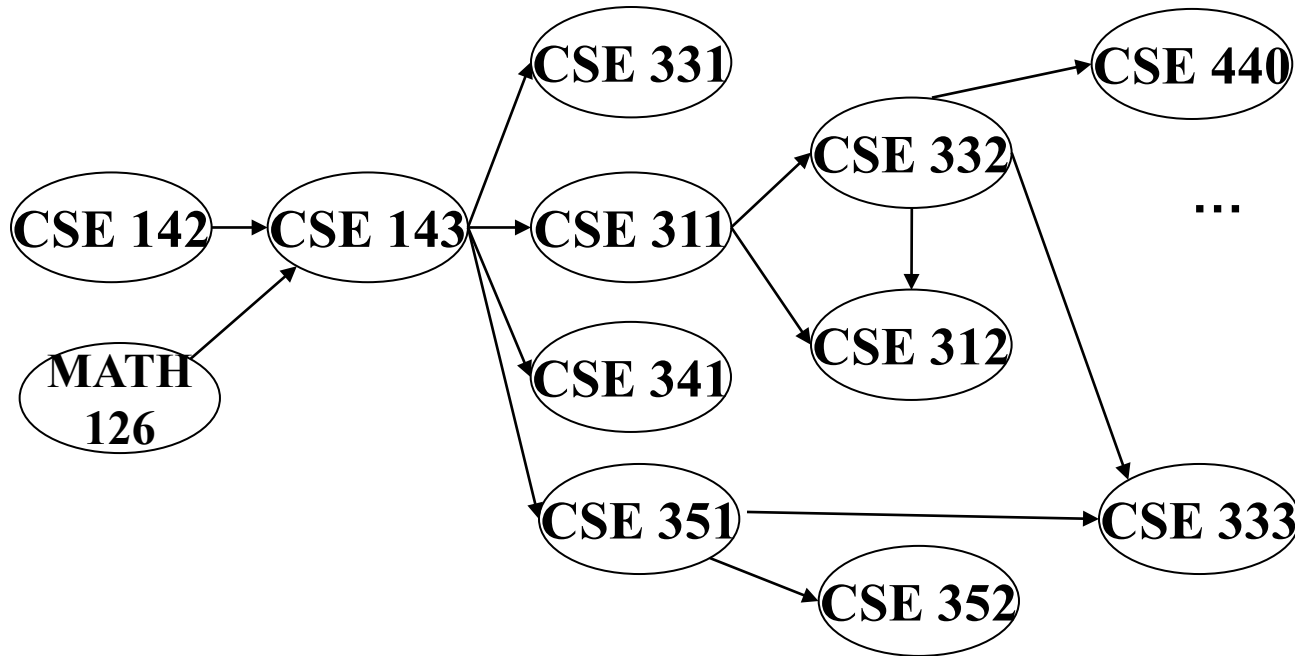
Output: 126
142
143
311



Node:	126	142	143	311	312	331	332	333	341	351	352	440
Removed?	x	x	x	x								
In-degree:	0	0	2	1	2	1	1	2	1	1	1	1
			1	0	1	0	0		0	0		
			0									

Example

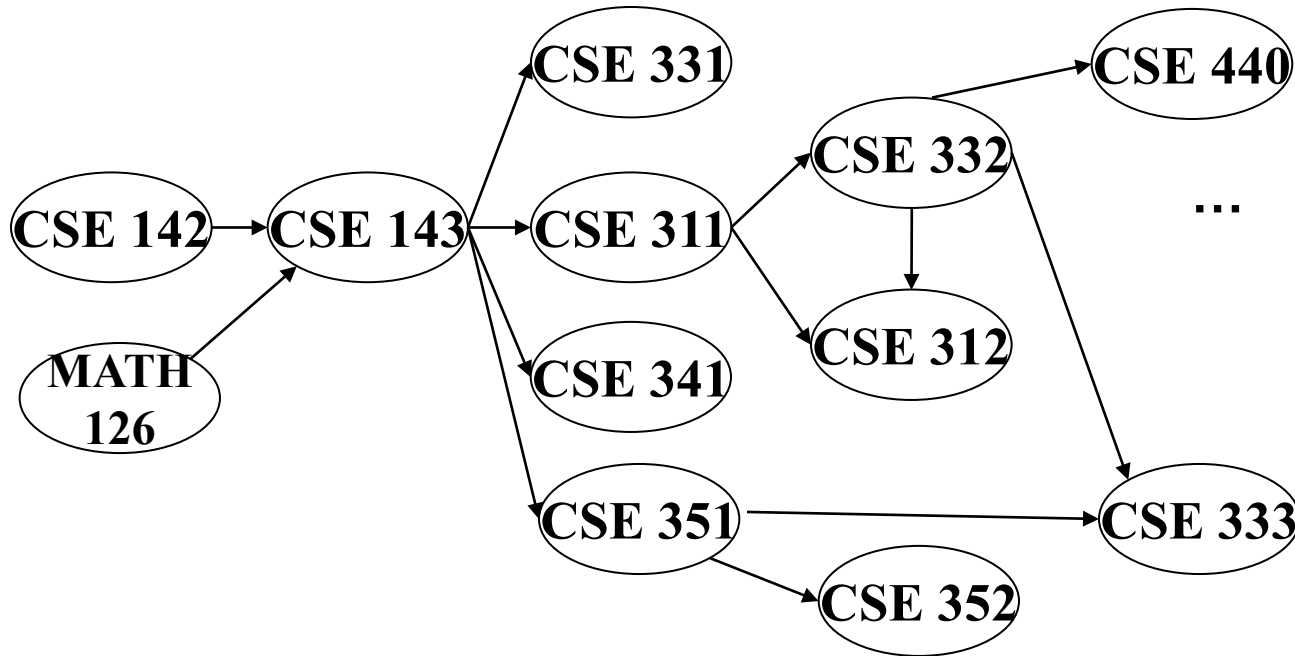
Output: 126
142
143
311
331



Node:	126	142	143	311	312	331	332	333	341	351	352	440
Removed?	x	x	x	x		x						
In-degree:	0	0	2	1	2	1	1	2	1	1	1	1
			1	0	1	0	0		0	0		
			0									

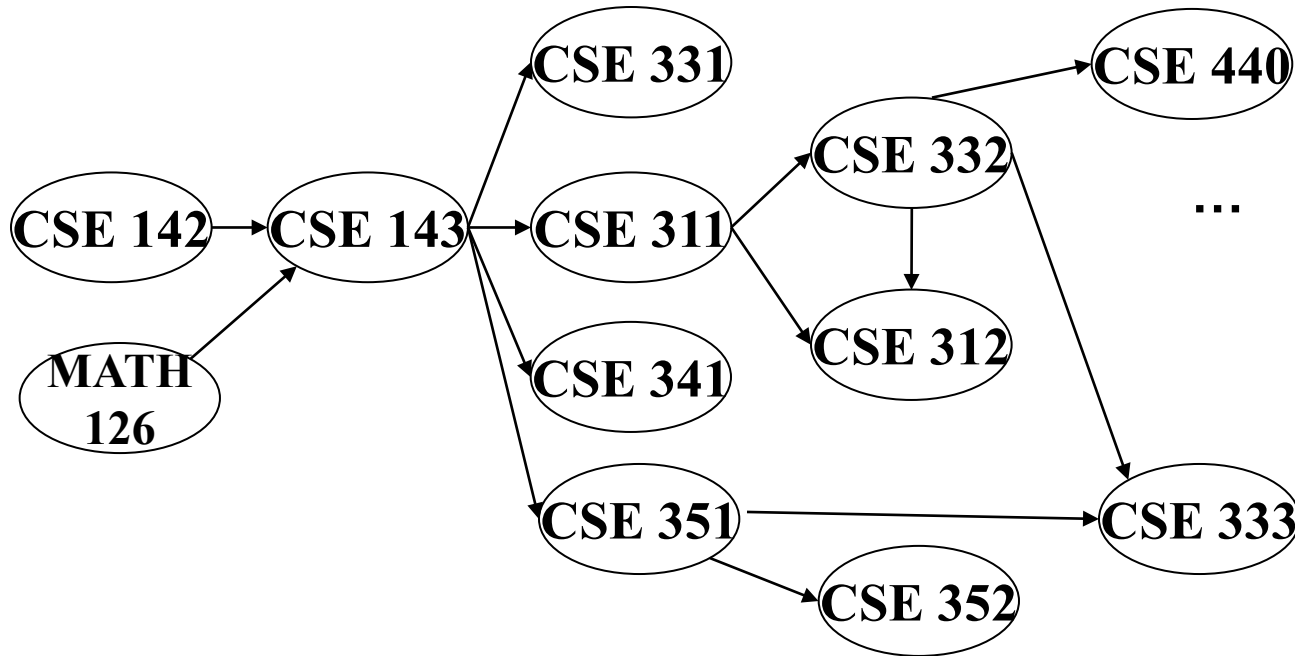
Example

Output: 126
142
143
311
331
332



Node:	126	142	143	311	312	331	332	333	341	351	352	440
Removed?	x	x	x	x		x	x					
In-degree:	0	0	2	1	2	1	1	2	1	1	1	1
			1	0	1	0	0	1	0	0		0
			0		0							

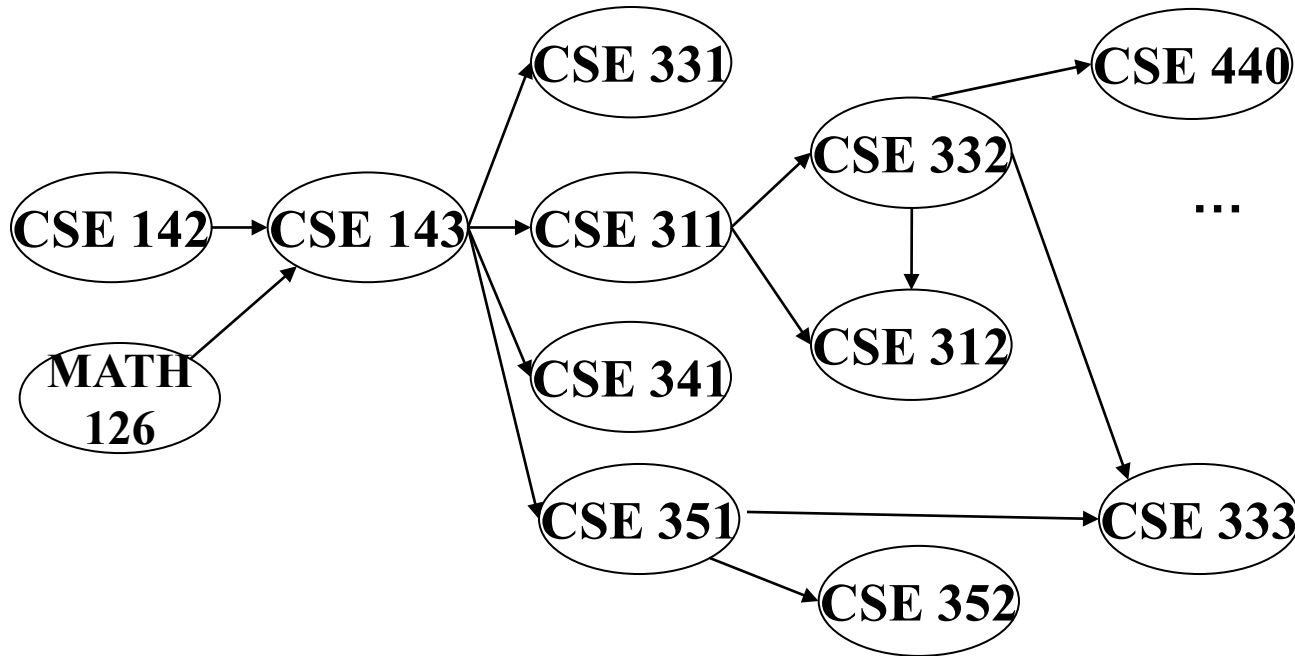
Example



Output: 126
142
143
311
331
332
312

Node:	126	142	143	311	312	331	332	333	341	351	352	440
Removed?	x	x	x	x	x	x	x					
In-degree:	0	0	2	1	2	1	1	2	1	1	1	1
			1	0	1	0	0	1	0	0		0
			0		0							

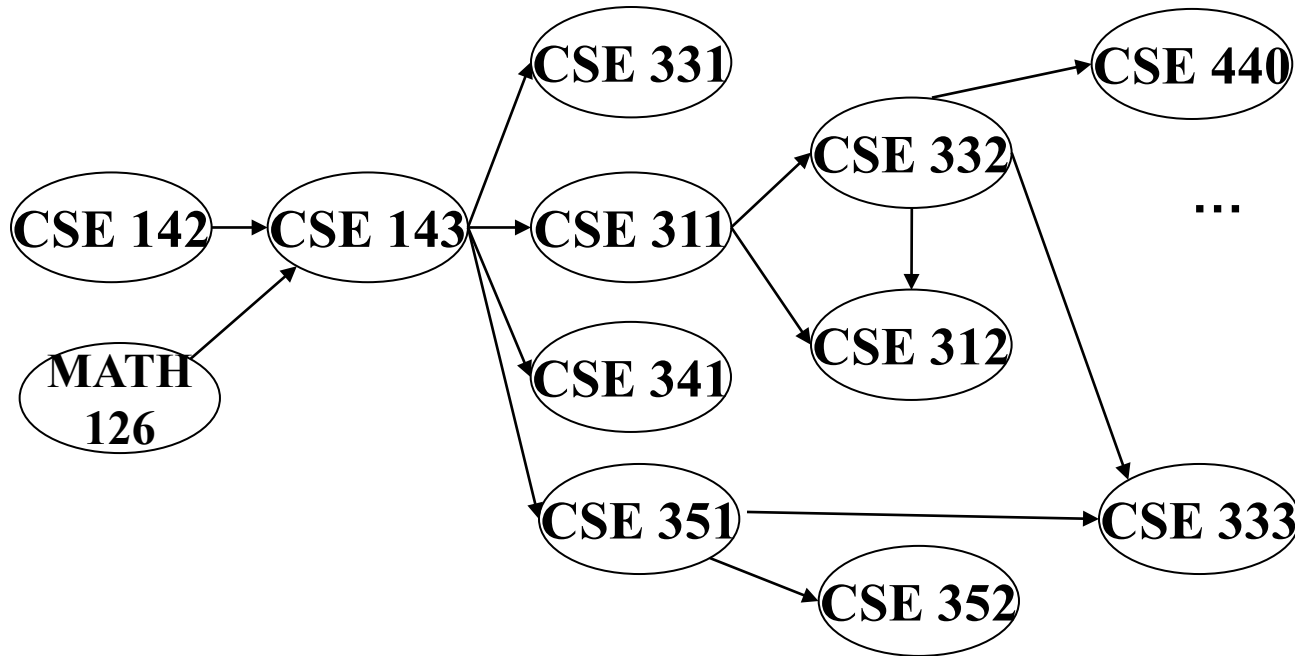
Example



Output: 126
142
143
311
331
332
312
341

Node:	126	142	143	311	312	331	332	333	341	351	352	440
Removed?	x	x	x	x	x	x	x		x			
In-degree:	0	0	2	1	2	1	1	2	1	1	1	1
			1	0	1	0	0	1	0	0		0
			0		0							

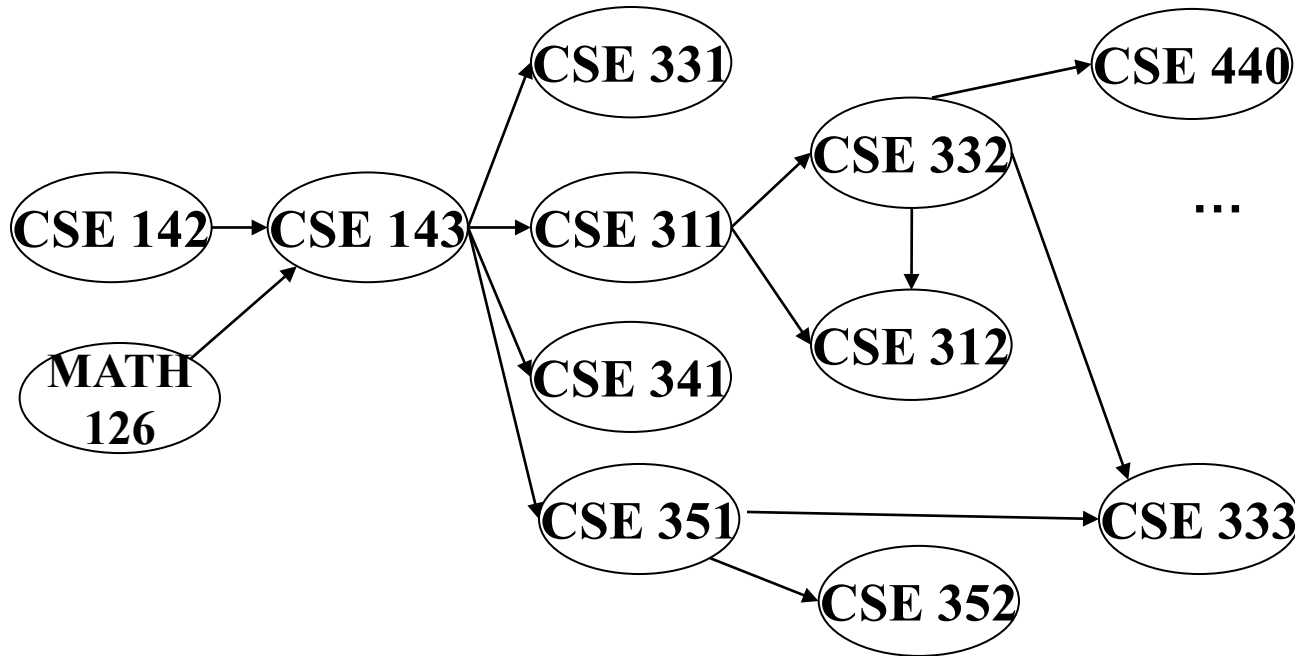
Example



Output: 126
142
143
311
331
332
312
341
351

Node:	126	142	143	311	312	331	332	333	341	351	352	440
Removed?	x	x	x	x	x	x	x		x	x		
In-degree:	0	0	2	1	2	1	1	2	1	1	1	1
			1	0	1	0	0	1	0	0	0	0
			0		0			0				

Example



Output: 126
142
143
311
331
332
312
341
351
333
352
440

Node:	126	142	143	311	312	331	332	333	341	351	352	440
Removed?	x	x	x	x	x	x	x	x	x	x	x	x
In-degree:	0	0	2	1	2	1	1	2	1	1	1	1
			1	0	1	0	0	1	0	0	0	0
			0		0			0				

A couple of things to note

- Needed a vertex with in-degree of 0 to start
 - No cycles
- Ties between vertices with in-degrees of 0 can be broken arbitrarily
 - Potentially many different correct orders

Topological Sort: Running time?

```
→ labelEachVertexWithItsInDegree();  
for(ctr=0; ctr < numVertices; ctr++){  
    v = findNewVertexOfDegreeZero();  
    put v next in output  
    for each w adjacent to v  
        w.indegree--;  
}
```

$O(V+E)$
 V times
 $O(V)$
 $O(1)$
 $O(d)$
 $O(V^2 + E)$

Topological Sort: Running time?

```
labelEachVertexWithItsInDegree();  
for(ctr=0; ctr < numVertices; ctr++){  
    v = findNewVertexOfDegreeZero();  
    put v next in output  
    for each w adjacent to v  
        w.indegree--;  
}
```

- What is the worst-case running time?
 - Initialization $O(|V| + |E|)$ (assuming adjacency list)
 - Sum of all find-new-vertex $O(|V|^2)$ (because each $O(|V|)$)
 - Sum of all decrements $O(|E|)$ (assuming adjacency list)
 - So total is $O(|V|^2 + |E|)$ – not good for a sparse graph!

Doing better

The trick is to avoid searching for a zero-degree node every time!

- Keep the “pending” zero-degree nodes in a list, stack, queue, box, table, or something
- Order we process them affects output but not correctness or efficiency provided add/remove are both $O(1)$

Using a queue:

1. Label each vertex with its in-degree, enqueue 0-degree nodes
2. While queue is not empty
 - a) $\mathbf{v} = \text{dequeue}()$
 - b) Output $\mathbf{v}$ and remove it from the graph
 - c) For each vertex $\mathbf{w}$ adjacent to $\mathbf{v}$ (i.e. $\mathbf{w}$ such that $(\mathbf{v}, \mathbf{w})$ in $\mathbf{E}$), decrement the in-degree of $\mathbf{w}$, if new degree is 0, enqueue it

Topological Sort(optimized): Running time?

```
labelAllAndEnqueueZeros();  
for(ctr=0; ctr < numVertices; ctr++){  
    v = dequeue();  
    put v next in output  
    for each w adjacent to v {  
        w.indegree--;  
        if (w.indegree==0)  
            enqueue(w);  
    }  
}
```

Topological Sort(optimized): Running time?

```
labelAllAndEnqueueZeros();  
for(ctr=0; ctr < numVertices; ctr++){  
    v = dequeue();  
    put v next in output  
    for each w adjacent to v {  
        w.indegree--;  
        if(w.indegree==0)  
            enqueue(w);  
    }  
}
```

- What is the worst-case running time?
 - Initialization: $O(|V|+|E|)$ (assuming adjacency list)
 - Sum of all enqueues and dequeues: $O(|V|)$
 - Sum of all decrements: $O(|E|)$ (assuming adjacency list)
 - So total is $O(|E| + |V|)$ – much better for sparse graph!

B-Trees

B-Trees

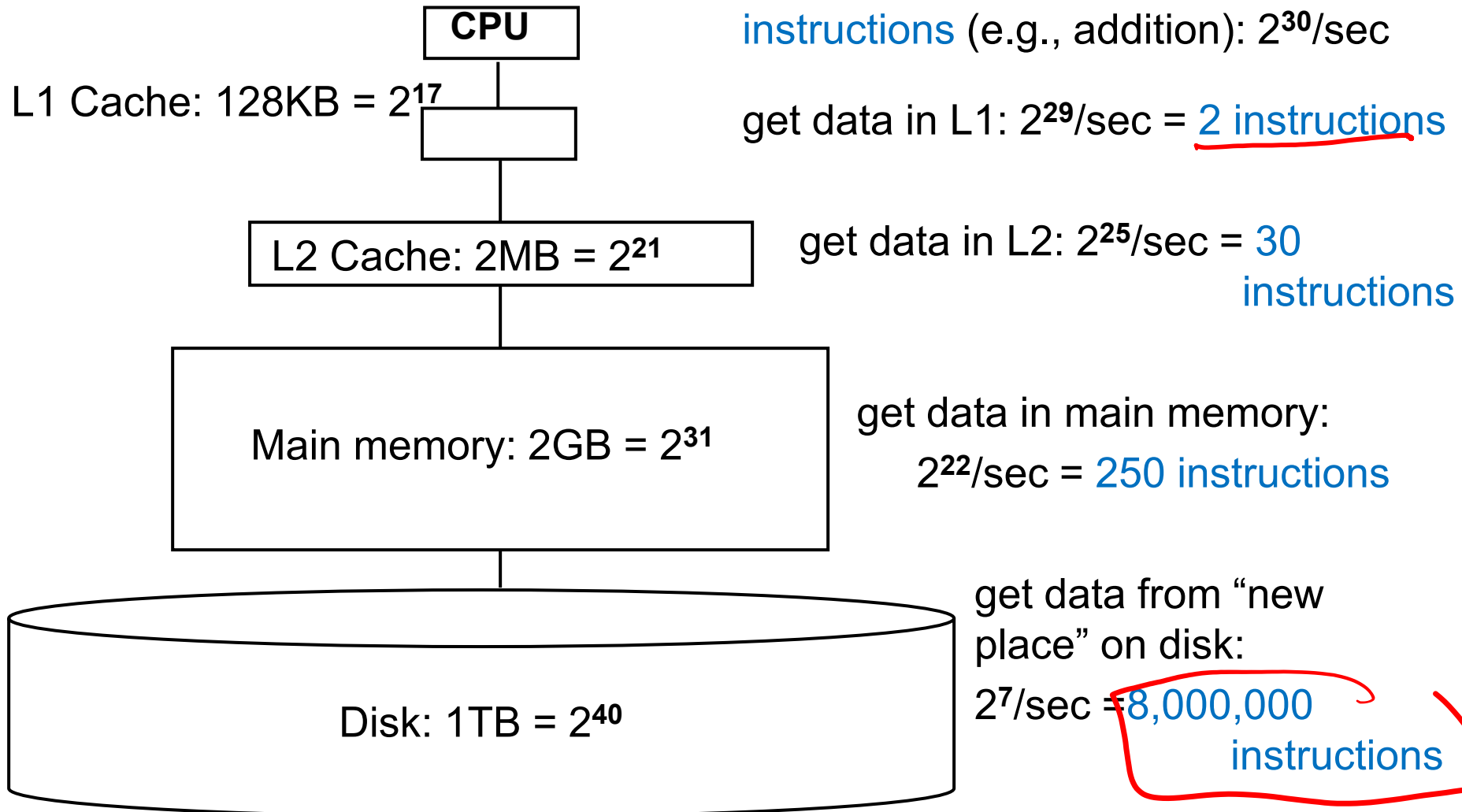
- We have a data structure for the dictionary ADT (AVL tree) that has worst-case $O(\log n)$ behavior
- B-Trees are a dictionary that considers the memory hierarchy
- To motivate why B trees are better for really large dictionaries (say, over 1GB = 2^{30} bytes), need to understand some memory-hierarchy basics
 - Don't always assume "every memory access has an unimportant $O(1)$ cost"
 - Learn more in CSE351/333/471, focus here on relevance to data structures and efficiency

Why do we need to know about the memory hierarchy?

- One of the assumptions that Big-Oh makes is that all operations take the same amount of time.
- Is that really true?

A typical hierarchy

“Every desktop/laptop/server is different” but here is a plausible configuration these days



Morals

It is much faster to do:

Than:

5 million arithmetic ops

1 disk access

2500 L2 cache accesses

1 disk access

400 main memory accesses

1 disk access

Why are computers built this way?

- Physical realities (speed of light, closeness to CPU)
- Cost (price per byte of different technologies)
- Disks get much bigger not much faster
 - Spinning at 7200 RPM accounts for much of the slowness and unlikely to spin faster in the future
- Speedup at higher levels (e.g. a faster processor) makes lower levels *relatively slower*
- Later in the course: more than 1 CPU!

“Fuggedaboutit”, usually

The hardware automatically moves data into the caches from main memory for you

- Replacing items already there
- So algorithms much faster if “data fits in cache” (often does)

Disk accesses are done by software (e.g., ask operating system to open a file or database to access some data)

So most code “just runs” but sometimes it’s worth designing algorithms / data structures with knowledge of memory hierarchy

- And when you do, you often need to know one more thing...

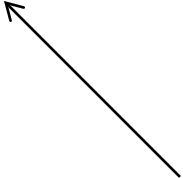
How does data move up the hierarchy?

- Moving data up the memory hierarchy is slow because of *latency* (think distance-to-travel)
 - Since we're making the trip anyway, may as well carpool
 - Get a block of data in the same time it would take to get a byte
 - Sends nearby memory because:
 - It's easy
 - And likely to be asked for soon (think fields/arrays)
- Side note: Once a value is in cache, may as well keep it around for awhile; accessed once, a particular value is more likely to be accessed again in the **near future** (more likely than some random other value)

Spatial Locality



Temporal locality

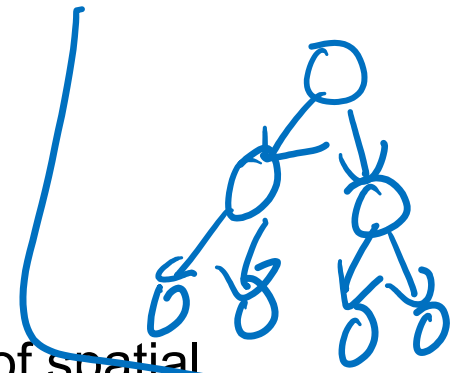


Locality

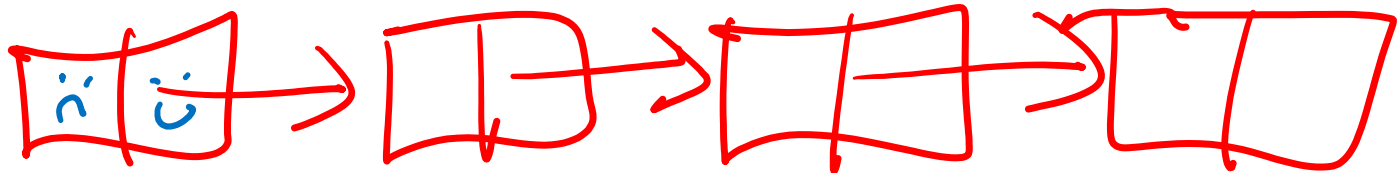
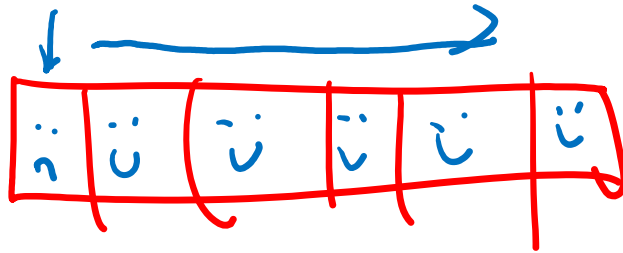
Temporal Locality (locality in **time**) – If an address is referenced, ***it*** will tend to be referenced again soon.

Spatial Locality (locality in **space**) – If an address is referenced, ***addresses that are close by*** will tend to be referenced soon.

Arrays vs. Linked lists



- Which has the potential to best take advantage of spatial locality?



Block/line size

- The amount of data moved from **disk** into **memory** is called the “**block**” size or the “**page**” size
 - Not under program control, determined in hardware and OS
- The amount of data moved from **memory** into **cache** is called the cache “**line**” size
 - Not under program control, determined in hardware

Connection to data structures

- An **array** benefits more than a **linked list** from block moves
 - Language (e.g., Java) implementation can put the list nodes anywhere, whereas array is typically contiguous memory
- Suppose you have a queue to process with 2^{23} items of 2^7 bytes each on disk and the block size is 2^{10} bytes
 - An **array** implementation needs 2^{20} disk accesses
 - If “perfectly streamed”, > 4 seconds
 - If “random places on disk”, 8000 seconds (> 2 hours)
 - A **linked list** implementation in the worst case needs 2^{23} “random” disk accesses (> 16 hours) – probably not that bad
- Note: “array” doesn’t necessarily mean “good”
 - Binary heaps “make big jumps” to percolate (different block)

BSTs?

- Looking things up in balanced binary search trees is $O(\log n)$, so even for $n = 2^{39}$ (512GB) we need not worry about minutes or hours
- Still, number of disk accesses matters:
 - Pretend for a minute we had an AVL tree of height 55
 - The total number of nodes could be? _____
 - Most of the nodes will be on disk: the tree is shallow, but it is still many gigabytes big so the entire *tree* cannot fit in memory
 - Even if memory holds the first 25 nodes on our path, we still potentially need 30 disk accesses if we are traversing the entire height of the tree.

Note about numbers; moral

- **Note:** All the numbers in this lecture are “ballpark” “back of the envelope” figures
- **Moral:** Even if they are off by, say, a factor of 5, the moral is the same:

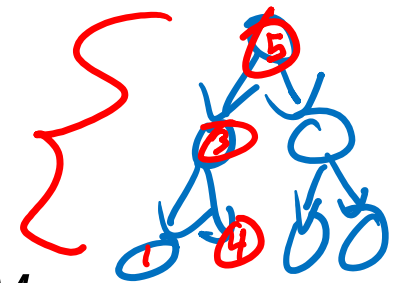
***If your data structure is mostly on disk,
you want to minimize disk accesses***

- A better data structure in this setting would exploit the block size and relatively fast memory access to ***avoid disk accesses...***

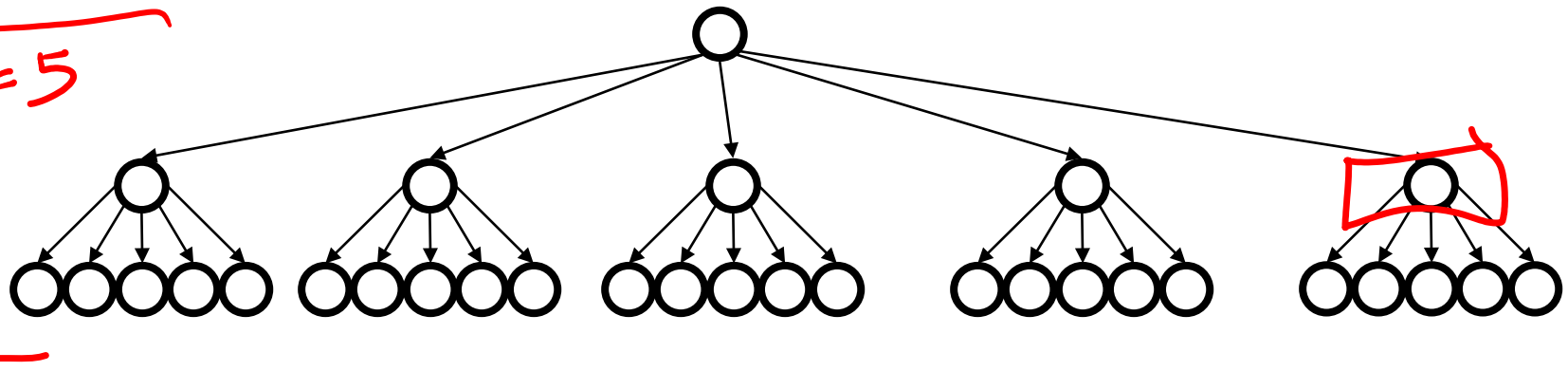
Our goal

- **Problem:** A dictionary with so much data *most of it is on disk*
- **Desire:** A balanced tree (logarithmic height) that is even shallower than AVL trees so that we can minimize disk accesses and exploit disk-block size
- **A key idea:** Increase the branching factor of our tree

M-ary Search Tree



- Build some sort of search tree with branching factor M :
 - Have an array of sorted children (**Node** [])
 - Choose M to fit snugly into a disk block (1 access for array)



Perfect tree of height h has $(M^{h+1}-1)/(M-1)$ nodes (textbook, page 4)

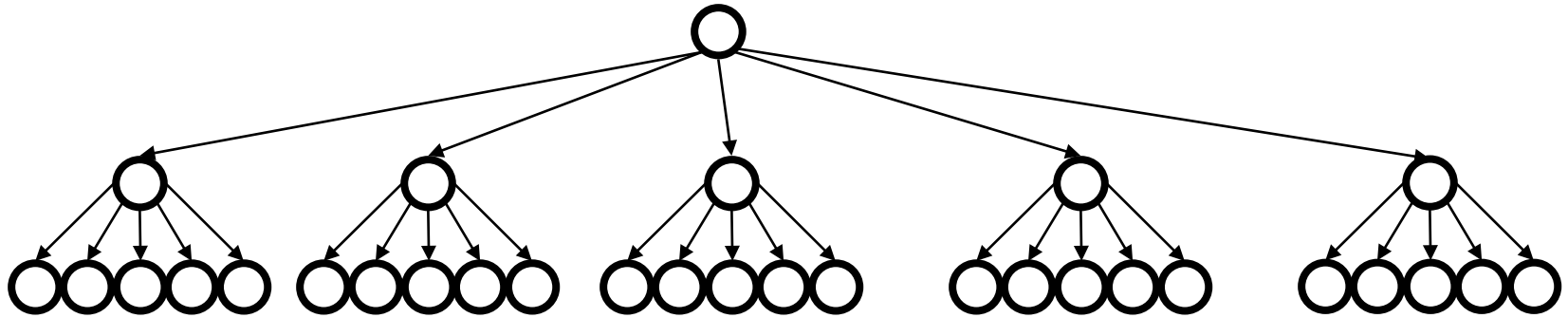
What is the height of this tree? $\log_m N$

What is the worst case running time of find?

$$O(\log_m N)$$

$$O(\log N)$$

M-ary Search Tree



- # hops for `find`?
 - If we have a balanced M-ary tree:
 - Approx. $\log_M n$ hops instead of $\log_2 n$ (for balanced BST)
 - Example: $M = 256 (=2^8)$ and $n = 2^{40}$ that's 5 hops instead of 40 hops
- Sounds good, but how do we decide which branch to take?
 - Binary tree: Less than/greater than node value?
 - M-ary: In range 1? In range 2? In range 3?... In range M?
- Runtime of `find` if balanced: $O(\log_2 M \log_M n)$
 - $\log_M n$ is the height we traverse.
 - $\log_2 M$: At each step, find the correct child branch to take using binary search among the M options!

Questions about *M*-ary search trees

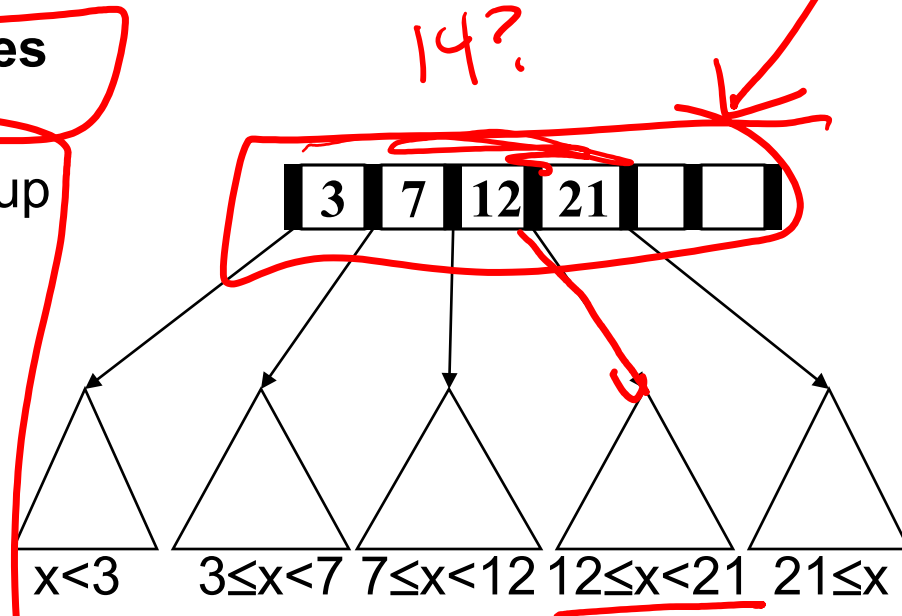
- What should the **order** property be?
- How would you **rebalance** (ideally without more disk accesses)?
- Storing **real data** at inner-nodes (like we do in a BST) seems kind of wasteful...
 - To access the node, will have to load the **data** from disk, even though most of the time we won't use it!!
 - Usually we are just “passing through” a node on the way to the value we are actually looking for.

So let's use the branching-factor idea, but for a **different kind of balanced tree**:

- **Not** a binary *search tree*
- But still logarithmic height for any $M > 2$

B+ Trees (we and the book say “B Trees”)

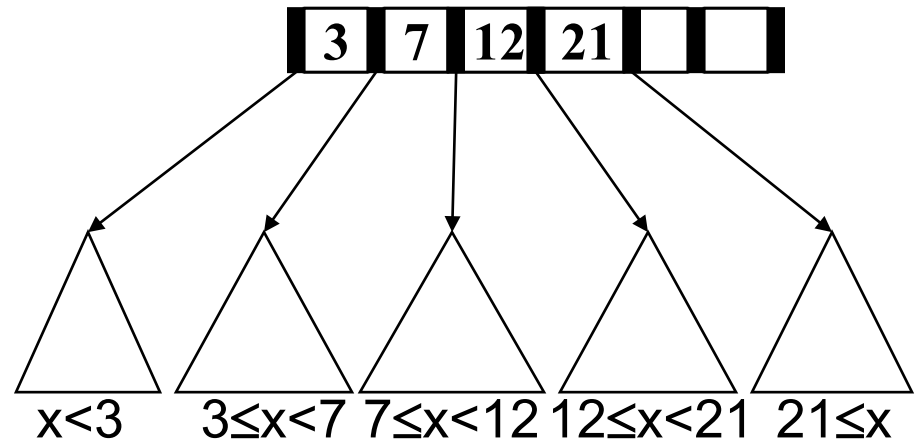
- Two types of nodes: **internal nodes** & **leaves**
- Each **internal node** has room for up to $M-1$ keys and M children
 - No other data; **all data at the leaves!**
- **Order property:**
Subtree **between** keys a and b contains only data that is $\geq a$ and $< b$ (notice the $\geq$)
- **Leaf** nodes have up to L sorted data items
- As usual, we'll ignore the “along for the ride” data in our examples
 - Remember no data at non-leaves



Remember:

- **Leaves** store data
- **Internal nodes** are ‘signposts’

Find



- Different from BST in that we don't store data at internal nodes
- But **find** is still an easy root-to-leaf recursive algorithm
 - At each internal node do binary search on (up to) $M-1$ keys to find the branch to take
 - At the leaf do binary search on the (up to) L data items
- But to get logarithmic running time, we need a balance condition...

Structure Properties

- **Root** (special case)
 - If tree has $\leq L$ items, root is a leaf (occurs when starting up, otherwise unusual)
 - Else has between 2 and M children
- **Internal nodes**
 - Have between $\lceil M/2 \rceil$ and M children, i.e., **at least half full**
- **Leaf nodes**
 - All leaves at the same depth
 - Have between $\lceil L/2 \rceil$ and L data items, i.e., **at least half full**

Any $M > 2$ and L will work, but:

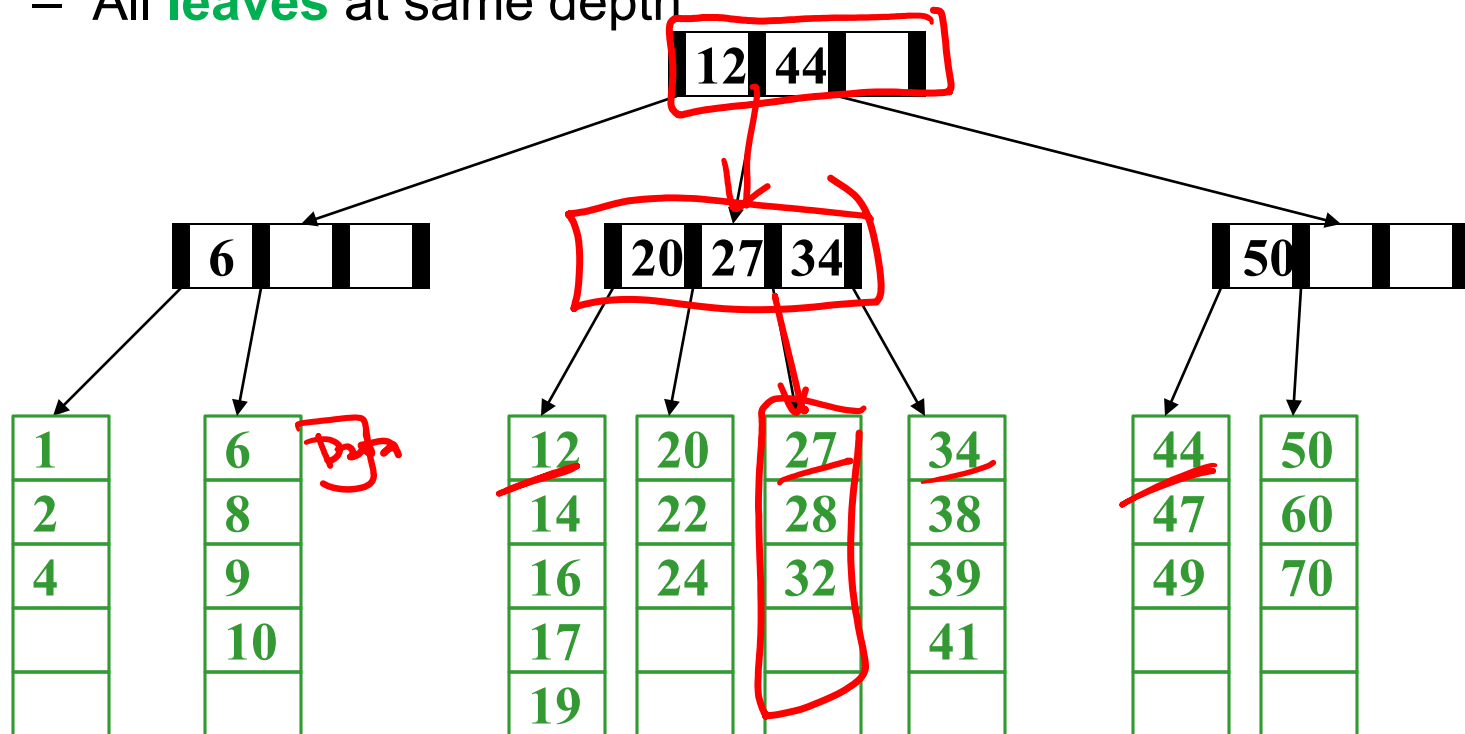
We pick M and L **based on disk-block size**

Note on notation: Inner nodes drawn horizontally, leaves vertically to distinguish. Include empty cells

Example

Suppose $M=4$ (max # pointers in **internal node**)
and $L=5$ (max # data items at **leaf**)

- All **internal nodes** have at least 2 children
- All **leaves** have at least 3 data items (only showing keys)
- All **leaves** at same depth



Balanced enough

Not hard to show height h is logarithmic in number of data items n

- Let $M > 2$ (if $M = 2$, then a list tree is legal – no good!)
- Because all nodes are at least half full (except root may have only 2 children) and all leaves are at the same level, the minimum number of data items n for a height $h > 0$ tree is...

$$n \geq \underbrace{2 \lceil M/2 \rceil^{h-1}}_{\text{minimum number of leaves}} \underbrace{\lceil L/2 \rceil}_{\text{minimum data per leaf}}$$

Example: B-Tree vs. AVL Tree

Suppose we have 100,000,000 items

- **Maximum height of AVL tree?**
 - Recall $S(h) = 1 + S(h-1) + S(h-2)$
 - lecture8.xlsx reports: **37**
- **Maximum height of B tree** with $M=128$ and $L=64$?
 - Recall $(2 \lceil M/2 \rceil^{h-1}) \lceil L/2 \rceil$
 - lecture9.xlsx reports: **5** (and 4 is more likely)
 - Also not difficult to compute via algebra

Disk Friendliness

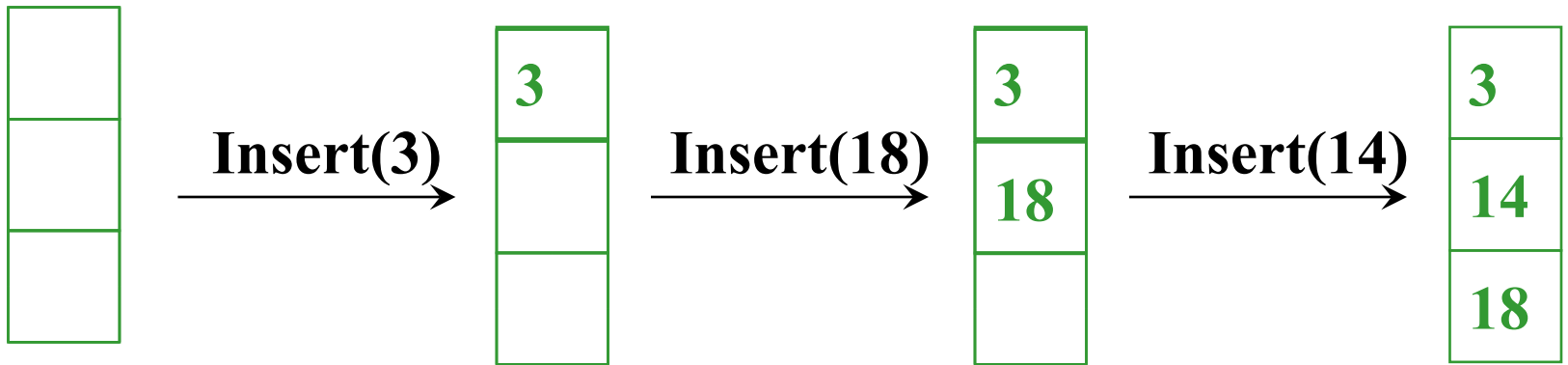
What makes B trees so disk friendly?

- Many keys stored in one **internal node**
 - All brought into memory in one disk access
 - *IF* we pick M wisely
 - Makes the binary search over $M-1$ keys totally worth it (insignificant compared to disk access times)
- **Internal nodes** contain only keys
 - Any **find** wants only one data item; wasteful to load unnecessary items with internal nodes
 - So only bring one **leaf** of data items into memory
 - Data-item size doesn't affect what M is

Maintaining balance

- How do we implement the other dictionary operations yet
 - **insert**
 - **delete**
- As with AVL trees, the hard part is maintaining structure properties
 - Example: for **insert**, there might not be room at the correct leaf

Building a B-Tree (insertions)

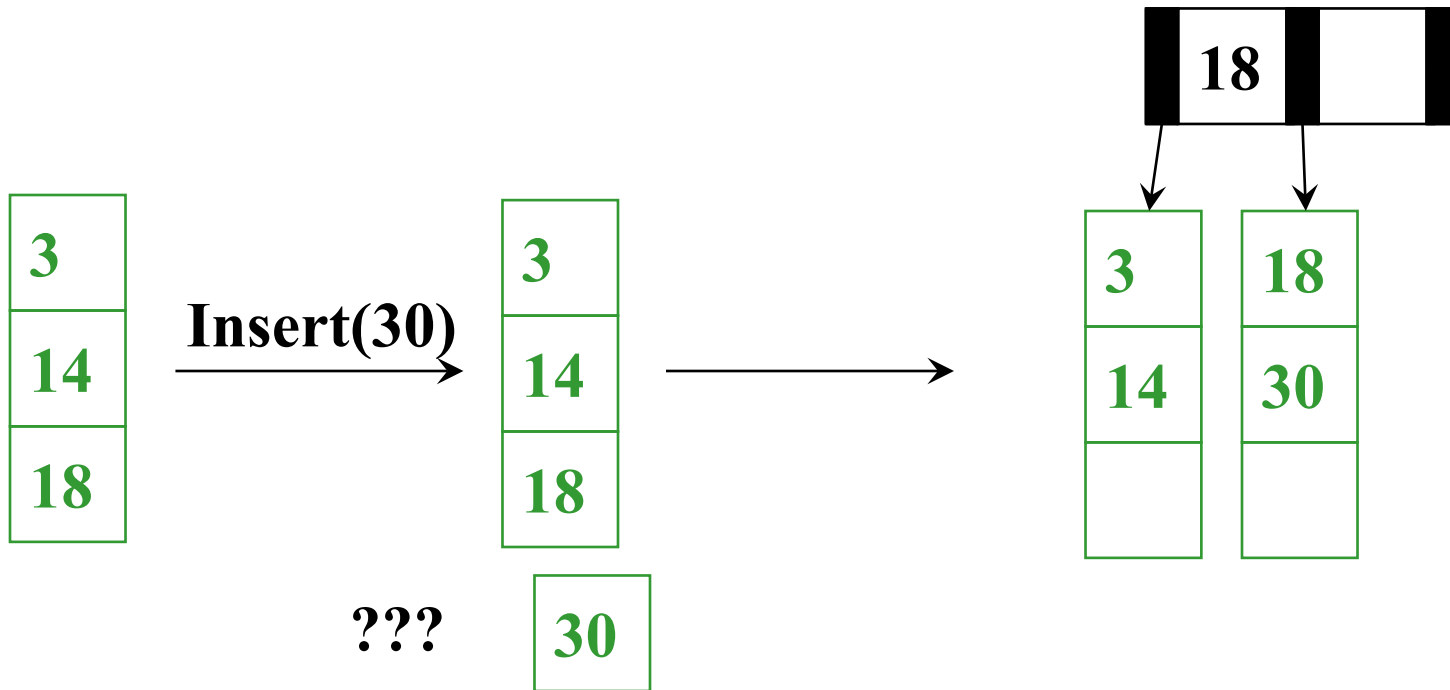


The empty B-Tree (the **root** will be a leaf at the beginning)

Just need to keep data in order

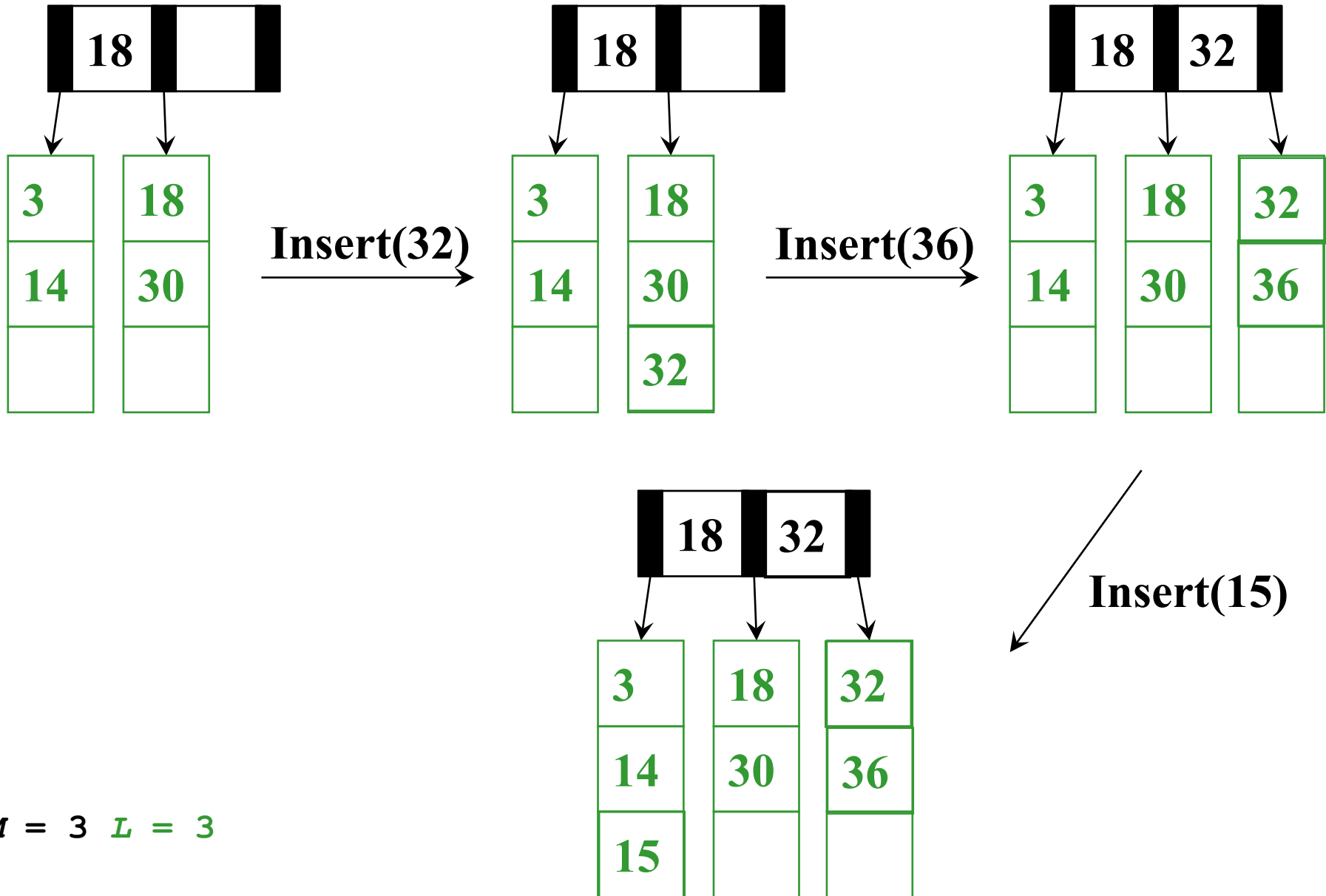
$$M = 3 \quad L = 3$$

$M = 3$ $L = 3$

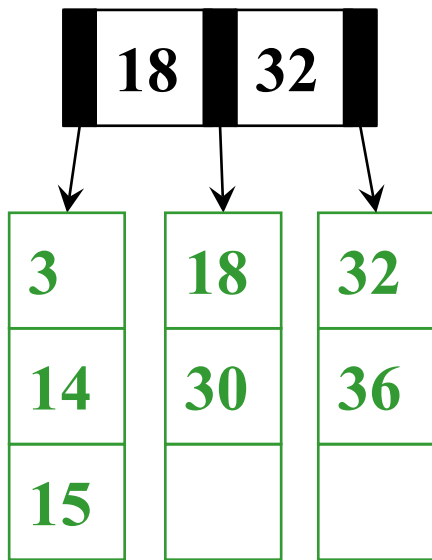


- When we ‘overflow’ a **leaf**, we split it into 2 leaves
- Parent gains another child
- If there is no parent (like here), we create one; how do we pick the key shown in it?
 - Smallest element in right tree

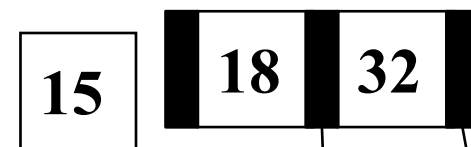
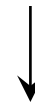
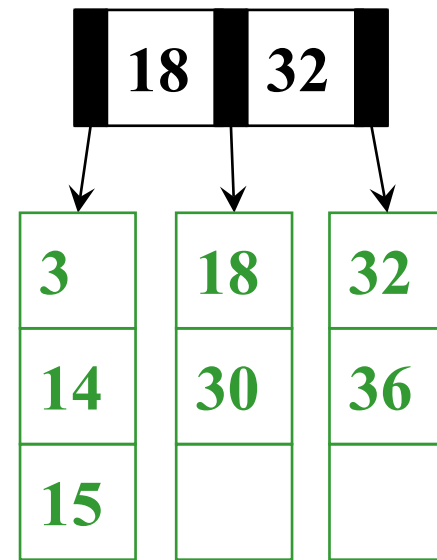
Split **leaf** again



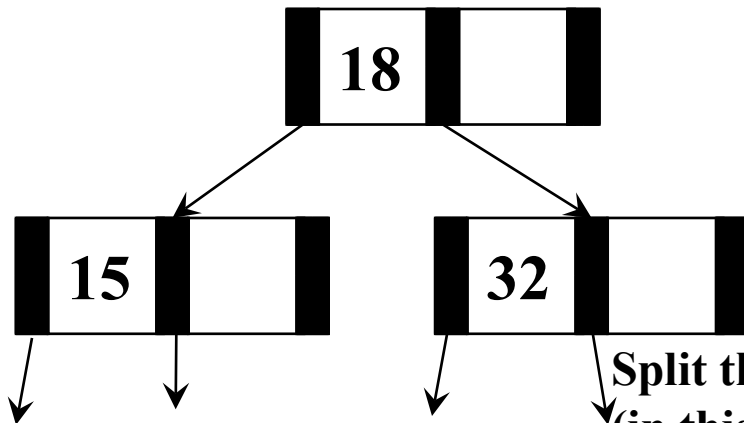
$M = 3$ $L = 3$



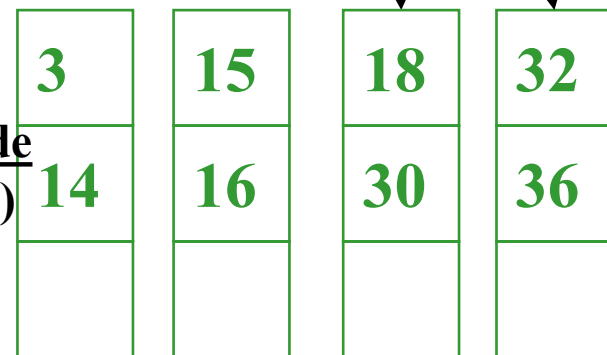
Insert(16)

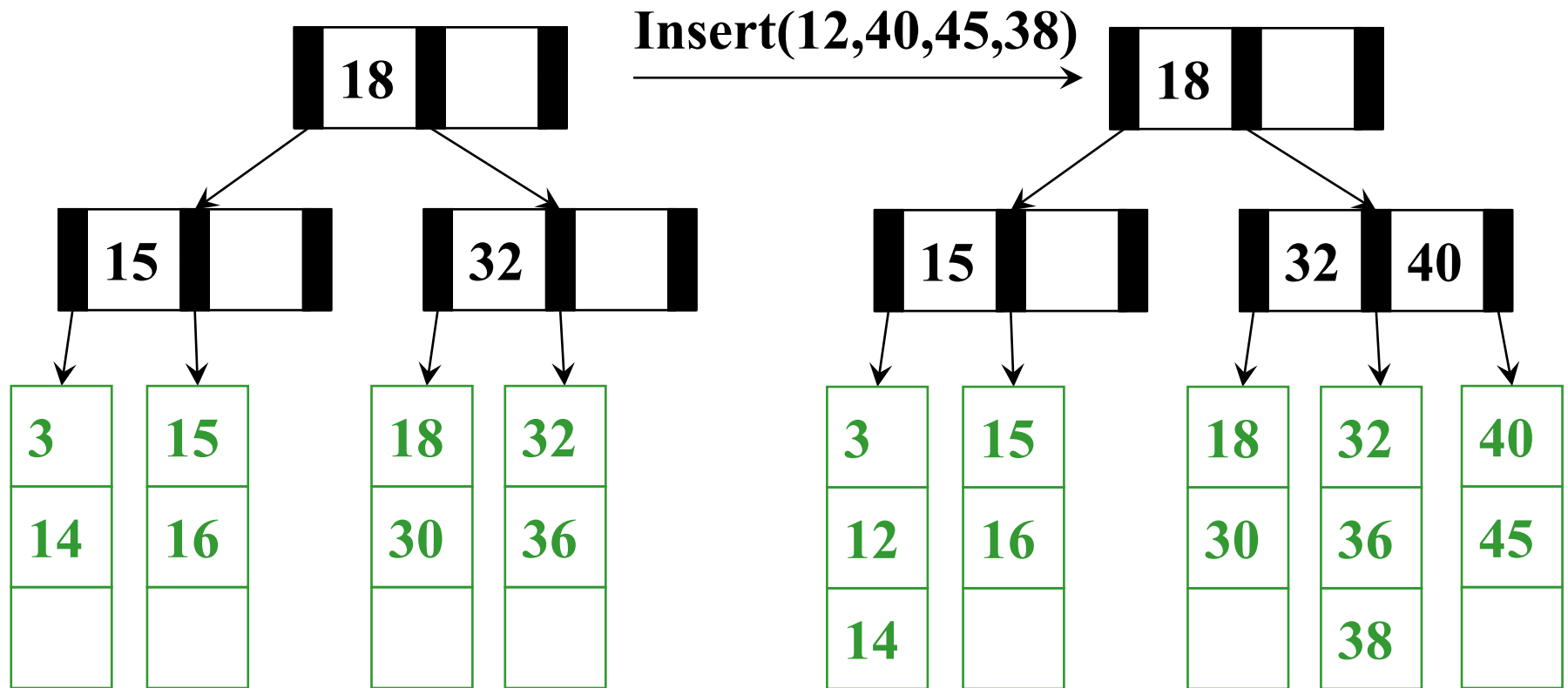


What now?



Split the internal node
(in this case, the **root**)





$M = 3$ $L = 3$

Note: Given the **leaves** and the structure of the tree, we can always fill in internal node keys; ‘the smallest value in my right branch’

Insertion Algorithm

1. Insert the data in its **leaf** in sorted order
2. If the **leaf** now has $L+1$ items, *overflow!*
 - Split the **leaf** into two nodes:
 - Original **leaf** with $\lceil (L+1) / 2 \rceil$ smaller items
 - New **leaf** with $\lfloor (L+1) / 2 \rfloor = \lceil L/2 \rceil$ larger items
 - Attach the new child to the parent
 - Adding new key to parent in sorted order
3. If step (2) caused the parent to have $M+1$ children, *overflow!*
 - ...

Insertion algorithm continued

3. If an **internal node** has $M+1$ children
 - Split the **node** into **two nodes**
 - Original **node** with $\lceil (M+1) / 2 \rceil$ smaller items
 - New **node** with $\lfloor (M+1) / 2 \rfloor = \lceil M/2 \rceil$ larger items
 - Attach the new child to the parent
 - Adding new key to parent in sorted order

Splitting at a node (step 3) could make the parent overflow too

- *So repeat step 3 up the tree until a node doesn't overflow*
- If the **root** overflows, make a new **root** with two children
 - This is the only case that increases the tree height

Efficiency of insert

- Find correct leaf: $O(\log_2 M \log_M n)$
- Insert in leaf: $O(L)$
- Split leaf: $O(L)$
- Split parents all the way up to root: $O(M \log_M n)$

Total: $O(L + M \log_M n)$

But it's not that bad:

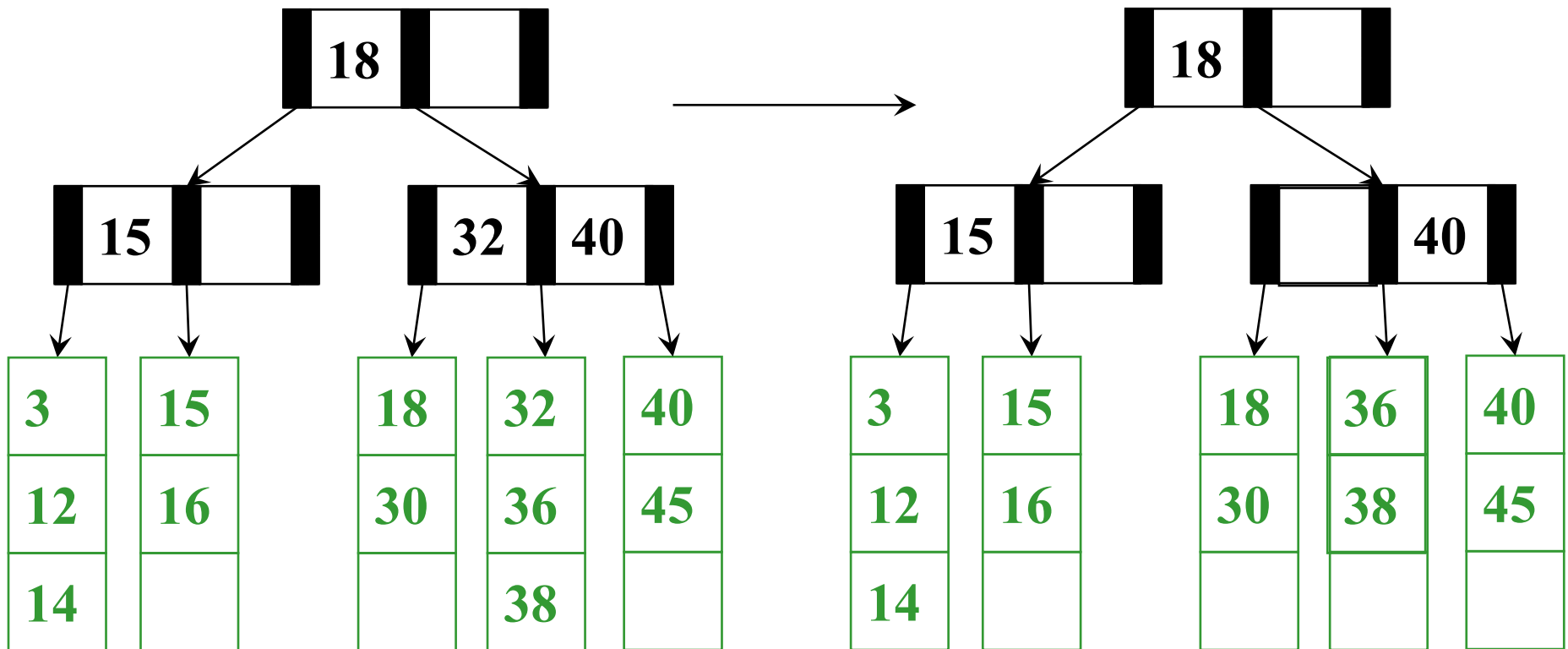
- Splits are not that common (only required when a node is FULL, M and L are likely to be large, and after a split, will be half empty)
- Splitting the **root** is extremely rare
- Remember disk accesses were the name of the game:
 $O(\log_M n)$

B-Tree Reminder: Another dictionary

- Before we talk about deletion, just keep in mind overall idea:
 - Large data sets won't fit entirely in memory
 - Disk access is slow
 - Set up tree so we do one disk access per node in tree
 - Then our goal is to keep tree shallow as possible
 - Balanced binary tree is a good start, but we can do better than $\log_2 n$ height
 - In an M-ary tree, height drops to $\log_M n$
 - Why not set M really really high? Height 1 tree...
 - Instead, set M so that each node fits in a disk block

And Now for Deletion...

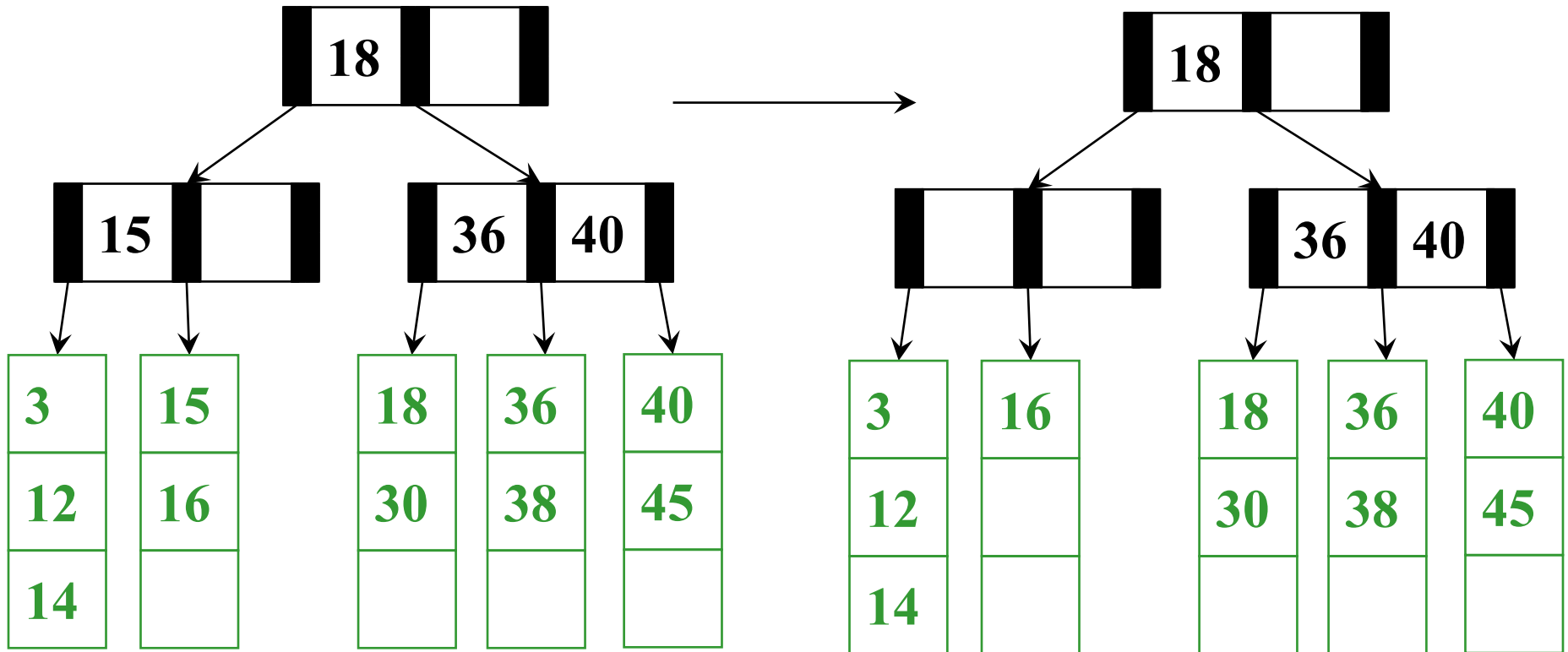
Delete(32)



Easy case: Leaf still has enough data; just remove

$M = 3$ $L = 3$

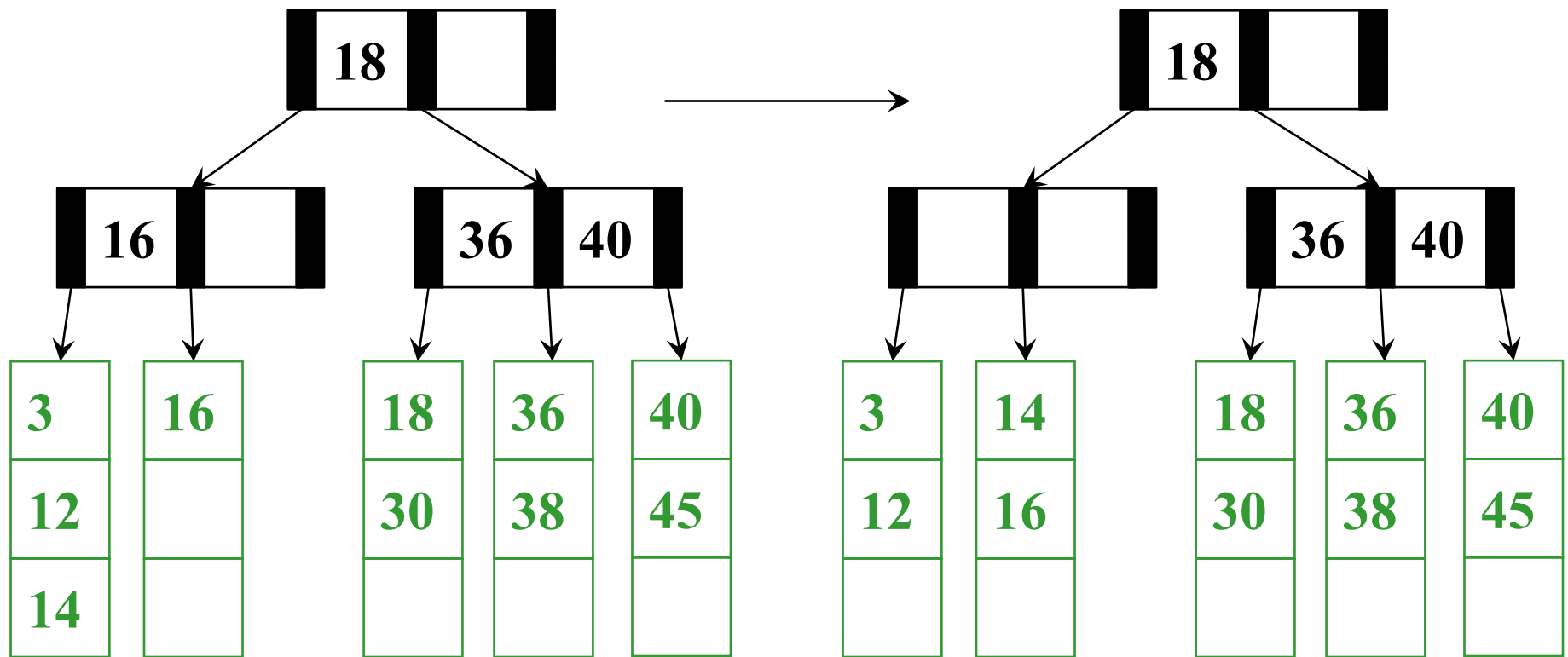
Delete(15)



Is there a problem?

$M = 3$ $L = 3$

12/03/2025

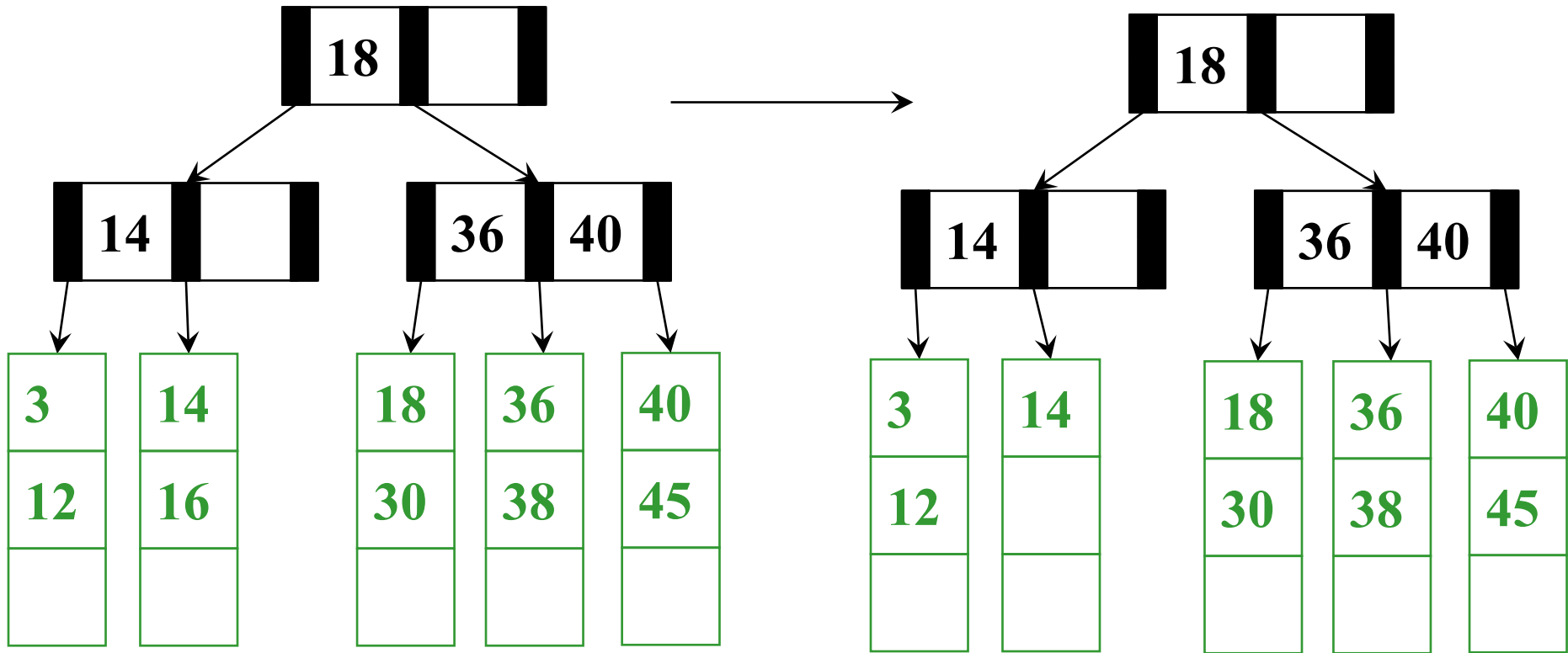


$M = 3$ $L = 3$

12/03/2025

Adopt from neighbor!

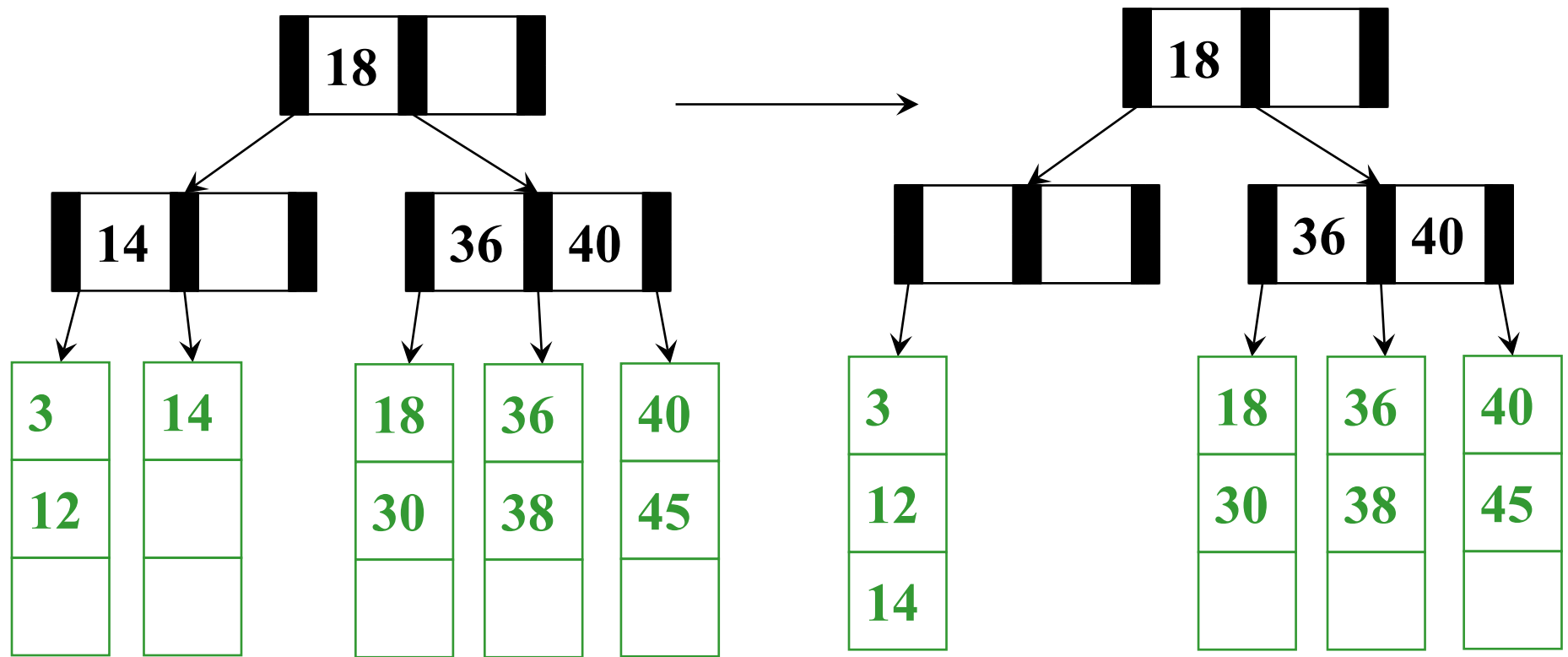
Delete(16)



Is there a problem?

$M = 3$ $L = 3$

12/03/2025



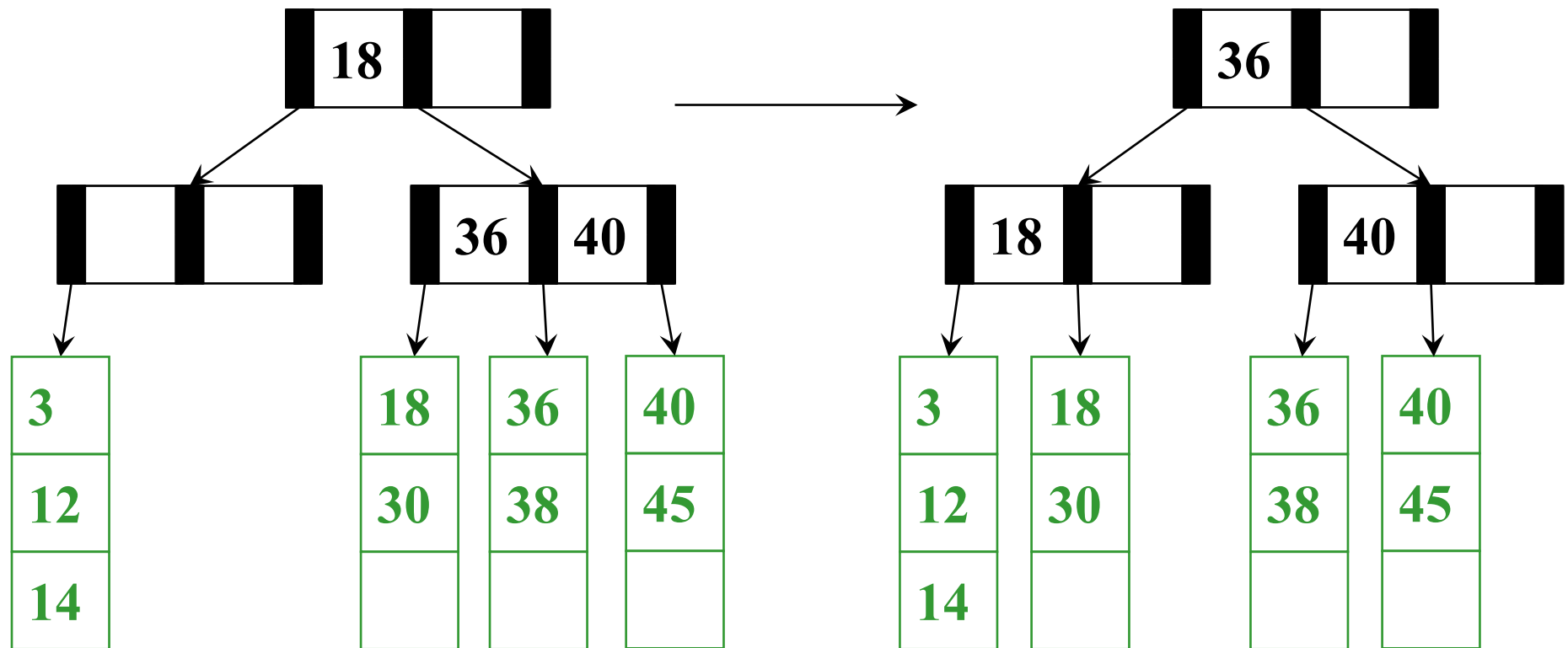
Merge with neighbor!

$M = 3$ $L = 3$

12/03/2025

But hey, Is there a problem?

66

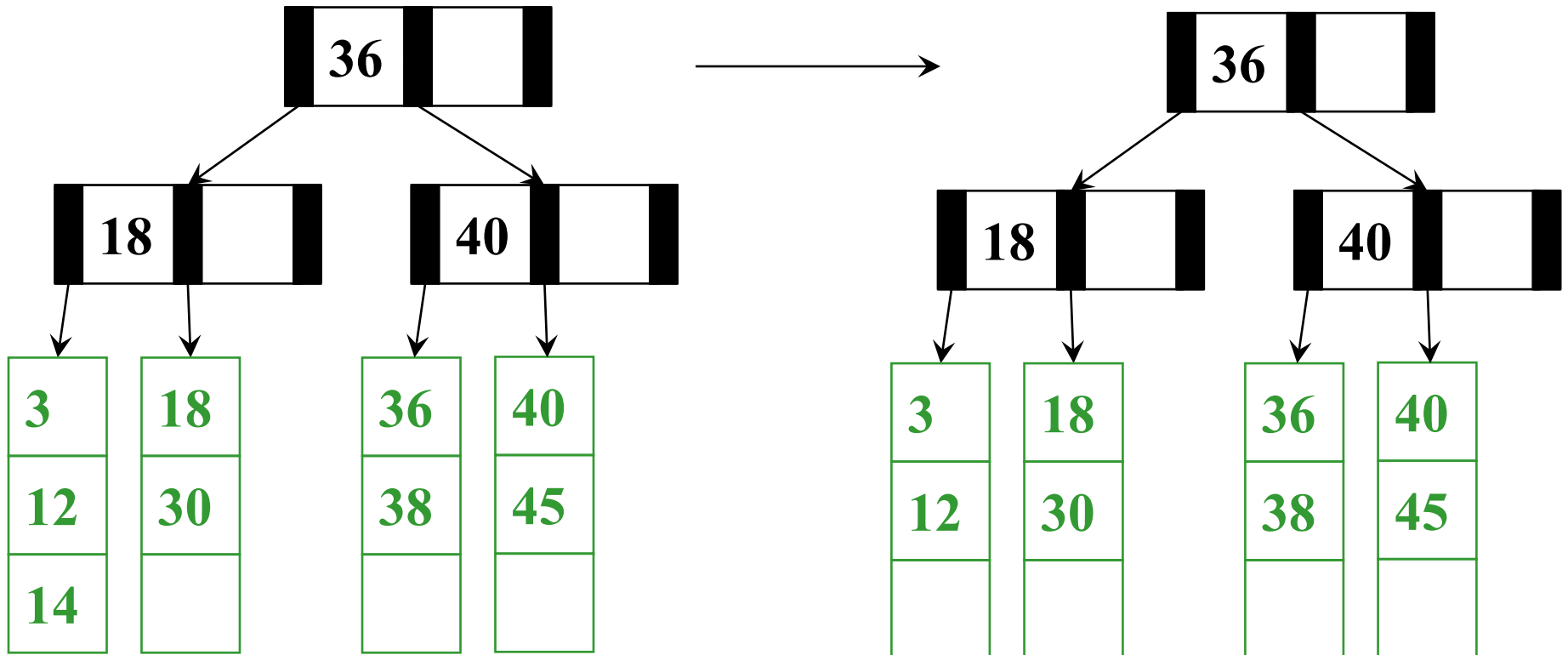


$M = 3$ $L = 3$

12/03/2025

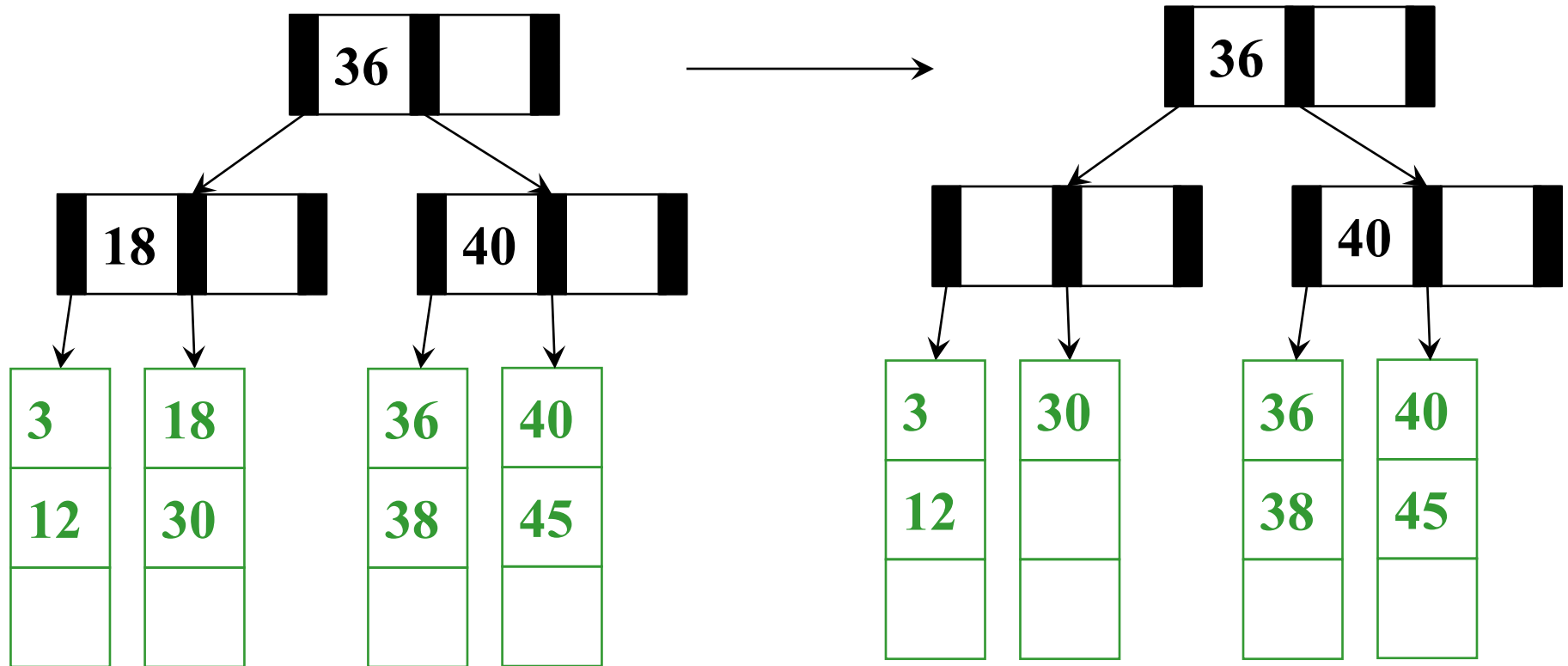
Adopt from neighbor!

Delete(14)



$M = 3$ $L = 3$

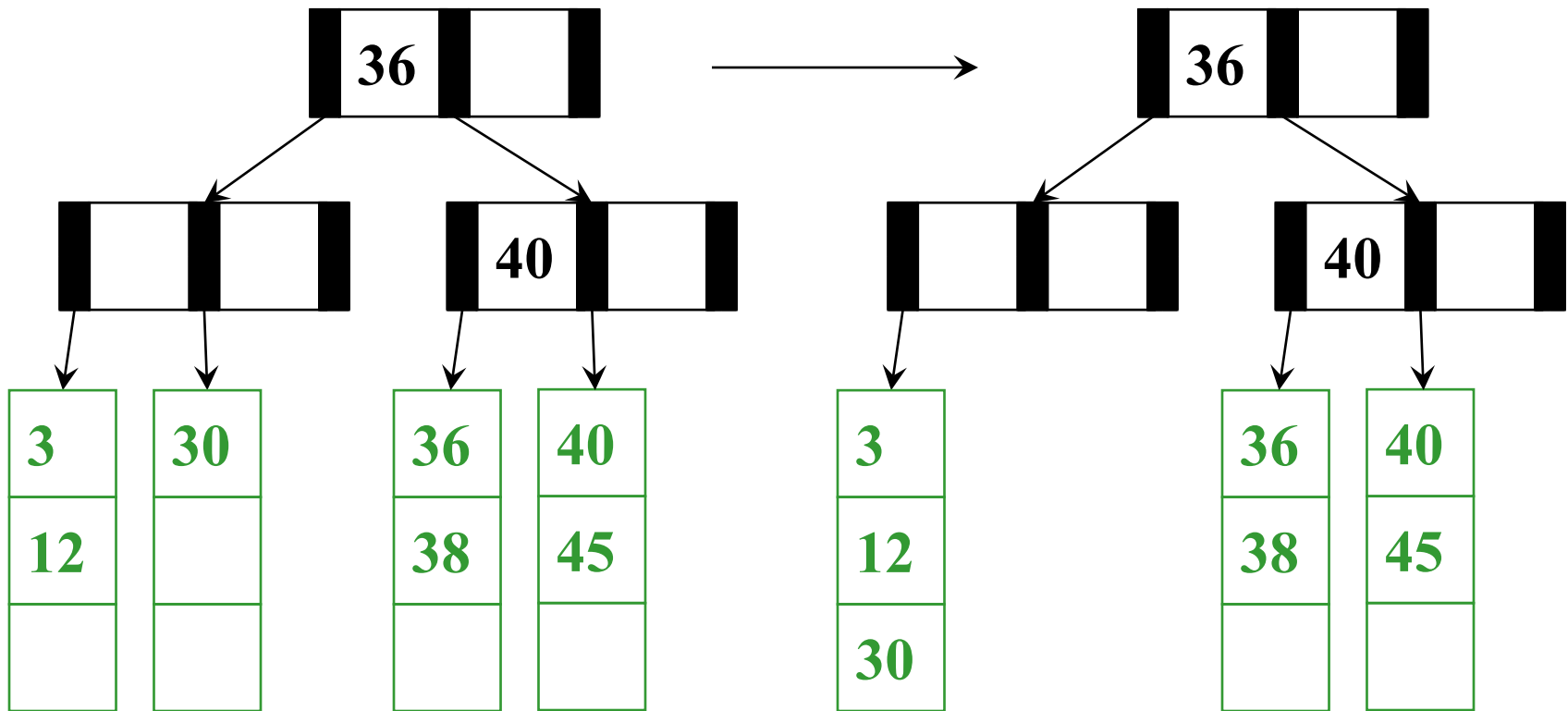
Delete(**18**)



$M = 3$ $L = 3$

12/03/2025

Is there a problem?



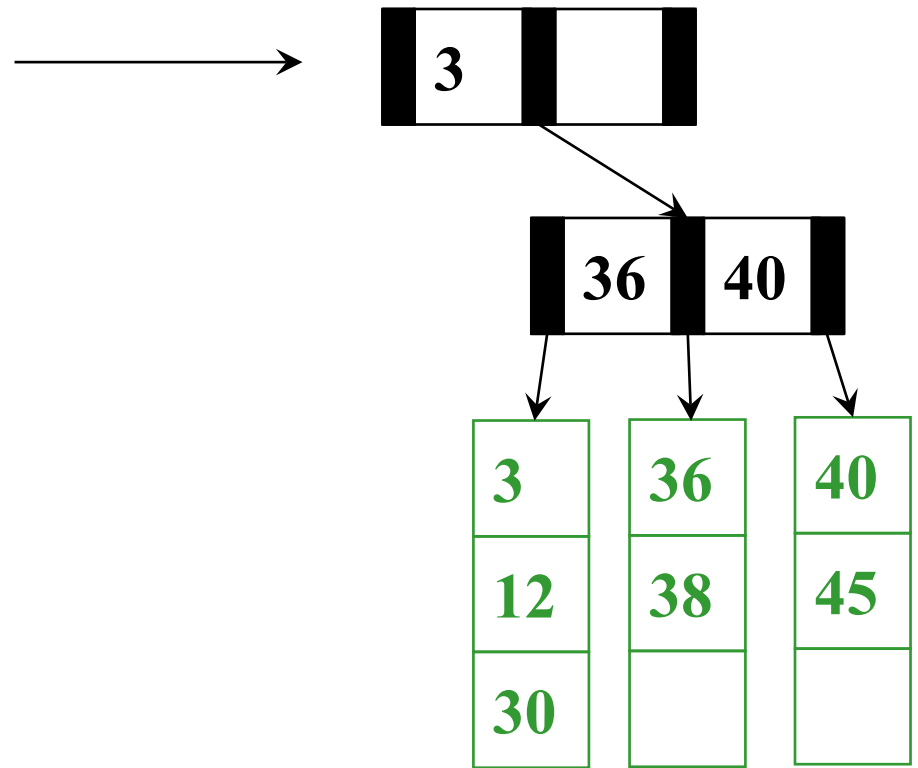
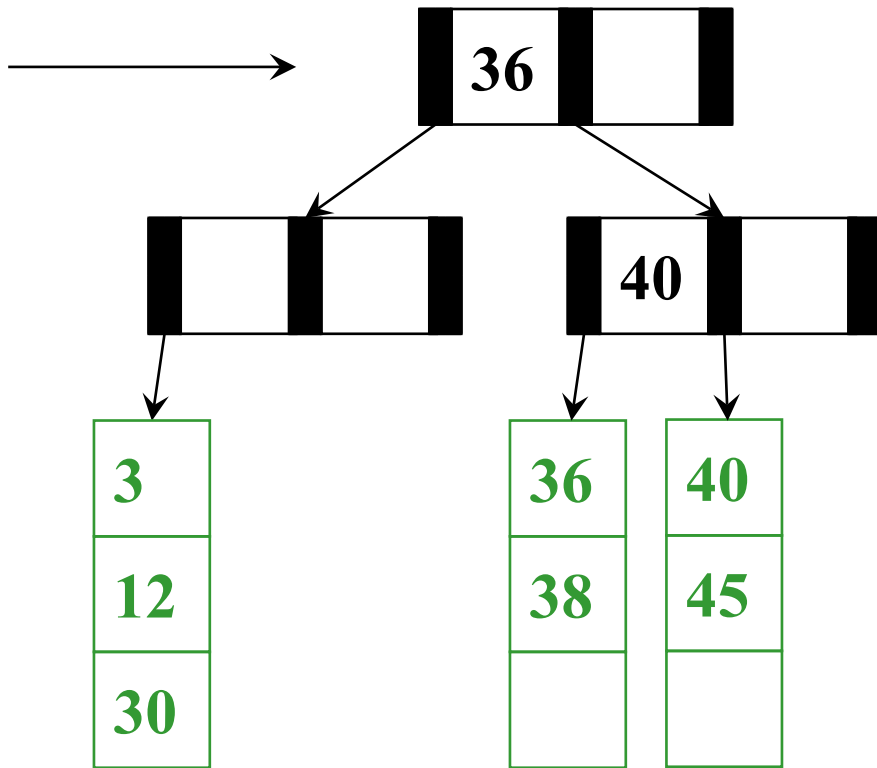
$M = 3$ $L = 3$

12/03/2025

Merge with neighbor!

But hey, Is there a problem?

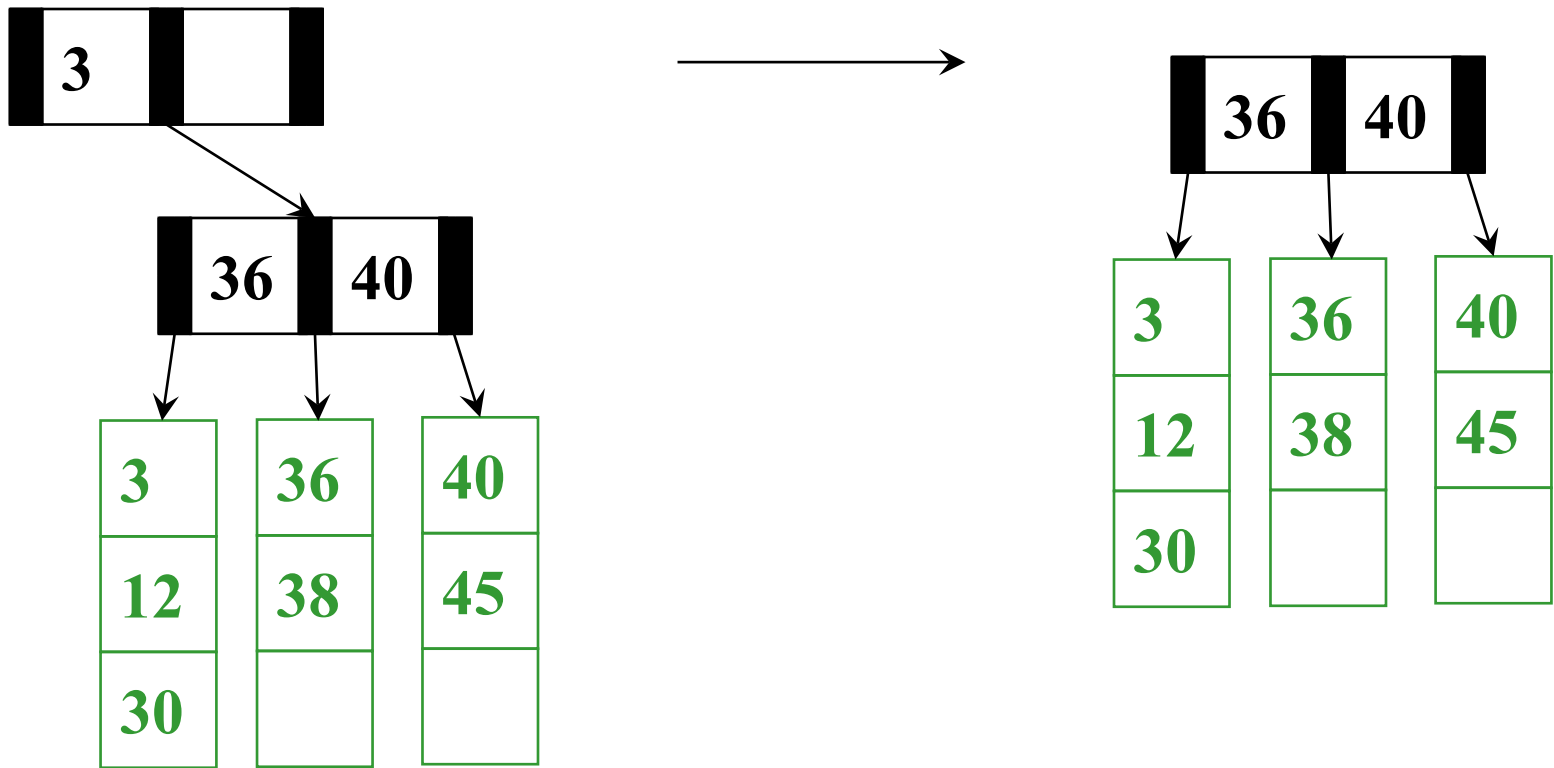
70



$M = 3$ $L = 3$

Merge with neighbor!

But hey, Is there a problem?



$M = 3$ $L = 3$

Pull out the root!

Deletion Algorithm, part 1

1. Remove the data from its leaf
2. If the leaf now has $\lceil L/2 \rceil - 1$, *underflow!*
 - If a neighbor has $> \lceil L/2 \rceil$ items, *adopt* and update parent
 - Else *merge* node with neighbor
 - Guaranteed to have a legal number of items
 - Parent now has one less node
3. If step (2) caused the parent to have $\lceil M/2 \rceil - 1$ children, *underflow!*
 - ...

Deletion algorithm (continued)

3. If an internal node has $\lceil M/2 \rceil - 1$ children
 - If a neighbor has $> \lceil M/2 \rceil$ items, *adopt* and update parent
 - Else *merge* node with neighbor
 - Guaranteed to have a legal number of items
 - Parent now has one less node, may need to continue up the tree

If we merge all the way up through the root, that's fine unless the root went from 2 children to 1

- In that case, delete the root and make child the root
- This is the only case that decreases tree height

Worst-Case Efficiency of Delete

- Find correct leaf: $O(\log_2 M \log_M n)$
- Remove from leaf: $O(L)$
- Adopt from or merge with neighbor: $O(L)$
- Adopt or merge all the way up to root: $O(M \log_M n)$

Total: $O(L + M \log_M n)$

But it's not that bad:

- Merges are not that common
- Disk accesses are the name of the game: $O(\log_M n)$

Insert vs delete comparison

Insert

- Find correct leaf: $O(\log_2 M \log_M n)$
- Insert in leaf: $O(L)$
- Split leaf: $O(L)$
- Split parents all the way up to root: $O(M \log_M n)$

Delete

- Find correct leaf: $O(\log_2 M \log_M n)$
- Remove from leaf: $O(L)$
- Adopt/merge from/with neighbor leaf: $O(L)$
- Adopt or merge all the way up to root: $O(M \log_M n)$

B Trees in Java?

For most of our data structures, we have encouraged writing high-level, reusable code, such as in Java with generics

It is worthwhile to know enough about “how Java works” to understand why this is probably a bad idea for B trees

- If you just want a balanced tree with worst-case logarithmic operations, no problem
 - If $M=3$, this is called a 2-3 tree
 - If $M=4$, this is called a 2-3-4 tree
- Assuming our goal is efficient number of disk accesses
 - Java has many advantages, but it wasn't designed for this

The key issue is extra *levels of indirection*...

Naïve approach in Java

Even if we assume data items have `int` keys, you cannot get the data representation you want for “really big data”

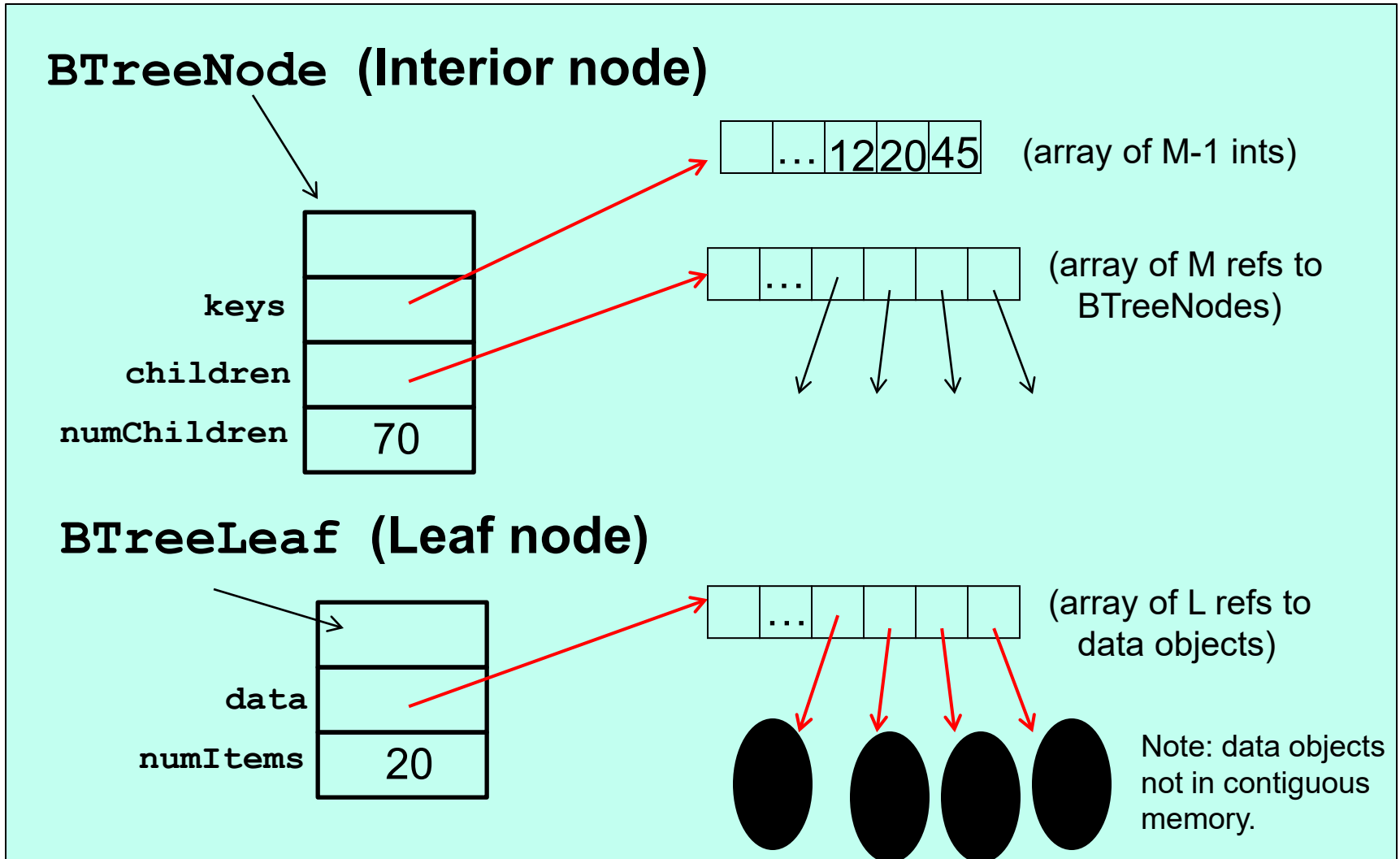
```
interface Keyed {
    int getKey();
}

class BTreeNode<E implements Keyed> {
    static final int M = 128;
    int[] keys = new int[M-1];
    BTreeNode<E>[] children = new BTreeNode[M];
    int numChildren = 0;
    ...
}

class BTreeLeaf<E implements Keyed> {
    static final int L = 32;
    E[] data = (E[])new Object[L];
    int numItems = 0;
    ...
}
```

All the **red** references indicate “unnecessary” indirection that might be avoided in another programming language.

What that looks like in Java



The moral

- The whole idea behind B trees was to keep related data in contiguous memory
- But that's “the best you can do” in Java
 - Again, the advantage is generic, reusable code
 - But for your performance-critical web-index, not the way to implement your B-Tree for terabytes of data
- Other languages (e.g., C++) have better support for “flattening objects into arrays”
- Levels of indirection matter!

Wrap Up!

What Have We Done This Quarter?

- Data Structures
 - Classic structures (hash tables, balanced BSTs, binary heaps, etc.)
 - You now understand these deeply, since you implemented many of them!
 - And you know how to analyze them with O , Ω , Θ .
 - And the ADTs
 - Analyze tradeoffs – there's not just one “right” answer!
- Algorithms
 - Sorting algorithms (examples of Divide & Conquer)
 - Graph algorithms (examples of Greedy Algorithms)
 - Effectively using our data structures to solve problems more efficiently (e.g. we used data structures in graph algorithms)

What Else Have We Done This Quarter?

- Parallelism & Concurrency
 - Exploit multiple processors
 - Fork-Join patterns (maps, reduces, prefixes, packs)
 - Concurrency: share resources safely
- Mixing theory and practice
 - Big- O analysis is the starting point, and you have more tools now (recurrences, worst- vs. best-case, amortization,...)
 - But in-practice constant factors, cache behavior, and a bunch of other things also matter (and are not captured by big- O)

What's Next?

- If you enjoyed shortest paths and MSTs
 - Take CSE 421: Algorithms
- If you loved P vs. NP or the sorting lower bound (what can't we do?)
 - Take CSE 431: Theory of Computation
- If you liked parallelism & concurrency
 - CSE 451: Operating Systems for how scheduling & locks actually work
 - CSE 452: Distributed Systems for hard concurrency problems.
- If your favorite part was writing and debugging code
 - CSE 331 has lots of code, CSE 333 does too.

Thank You!

- Thank you for a great quarter!
- Thanks to the CSE 332 Staff for a great quarter!
- We hope to see you in the future!