

Graph Algorithms

Diego Rivera Correa

Slides owned by Mirek Riedewald



This work is licensed under the Creative Commons Attribution 4.0 International License. To view a copy of this license, visit <http://creativecommons.org/licenses/by/4.0/>.

Key Learning Goals

- Write the pseudo-code for breadth-first search (BFS) in MapReduce and in Spark.
- Write the pseudo-code for single-source shortest path in MapReduce and in Spark.
- Write the pseudo-code for PageRank (without dangling pages) in MapReduce and in Spark.
- How can the single-source shortest path algorithm detect that no more iterations are needed?
- How can the PageRank algorithm detect that no more iterations are needed?

Key Learning Goals

- Why is Spark better suited for BFS and BFS-based algorithms than MapReduce?
- How can we handle dangling pages in PageRank?

Introduction

- Graphs are very general and a large variety of real-world problems can be modeled naturally as graph analysis or mining problems. Data occurring as graphs include:
 - Social interactions, e.g., Facebook friendships, Twitter followers, email flows, and phone call patterns
 - Transportation networks, e.g., roads, bus routes, and flights
- Since graphs are so general, many graph problems are inherently complex—a perfect target for distributed data-intensive computation.

What is a Graph?

- A graph consists of 2 main components:
 - **Vertices:** The things being connected.
 - **Edges:** The connections between those things.
- Edges can be one-way or two-way (directed or undirected).
- Graphs might contain cycles. If a graph does not contain any cycle, it is acyclic.

What is a Graph?



Graph Problems

- Graph search and path planning
 - Find driving directions from point A to point B.
 - Recommend new friends in a social network.
 - Find the best route for IP packets or delivery trucks.
- Graph clustering
 - Identify user communities in social networks based on graph structure, e.g., connectedness.
 - Partition a large graph into homogeneous partitions to parallelize graph processing.
- Bipartite graph matching
 - Match nodes on the “left” with nodes on “right” side. For example, match job seekers and employers, singles looking for dates, or papers with qualified reviewers.
- Maximum flow
 - Determine the greatest possible traffic between a source and a sink node, e.g., to optimize transportation networks.

GRAPH LANDSCAPE

2026 EDITION

Graph Databases

Property Graphs	RDF	Multi-model

Industry-specific Graph Apps

Cybersecurity	Law Enforcement/ Financial Crime
Other 	

Graph Processing Engines

Graph Viz Applications

Consulting

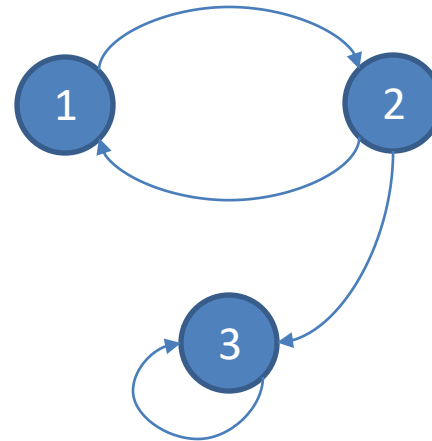
Graph Representations

- Graphs are usually represented in one of these three formats:
 - Adjacency matrix
 - Adjacency list
 - Set of individual edges
- We will take a quick look at each of them.

Adjacency-Matrix Representation

- The adjacency matrix M represents a graph in a matrix of size $|V|$ by $|V|$. Its size is quadratic in the number of vertices. Entry $M(i,j)$ contains the weight of the edge from vertex i to vertex j ; or 0 if there is no edge.
 - The adjacency matrix of an undirected graph is symmetric along the diagonal.

	1	2	3
1	0	1	0
2	1	0	1
3	0	0	1



Two-Hop Paths in Linear Algebra

- Consider entry (1, 3) in matrix $M \times M$, which is marked in green. Its value 1 is caused by $M(1, 2)=1$ (marked orange) and $M(2, 3)=1$ (marked black). Intuitively, the two-step path from node 1 to node 3 is made up of edges (1, 2) and (2, 3).

0	1	0
1	0	1
0	0	1

M

 $\times$

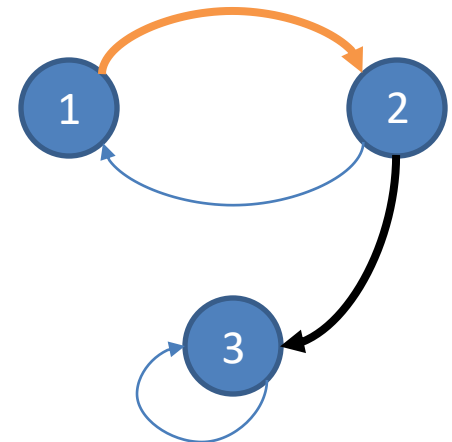
0	1	0
1	0	1
0	0	1

M

 $=$

1	0	1
0	1	1
0	0	1

$M \times M$



Adjacency-List Representation

- The adjacency-list representation only stores the existing *outlinks* for each vertex. Hence its size is linear in the number of edges.

Adjacency matrix

	1	2	3
1	0	1	0
2	1	0	1
3	0	0	1

Adjacency list for the same graph

1: 2

2: 1, 3

3: 3

Adjacency-List Properties

- Pro: Compared to the adjacency matrix, the adjacency list is much more space-efficient for sparse graphs.
- Pro: It is still easy to compute over the outlinks of a vertex, because they are stored in the adjacency list for that vertex.
- Con: Computation over the inlinks of a vertex is more challenging compared to the adjacency matrix. Instead of scanning through a column, each of the different lists must be searched.

Set-of-Edges Representation

- The set-of-edges approach stores each edge of the graph explicitly.
 - The adjacency list format is set-of-edges, grouped by the “start” vertex of each edge.

Adjacency matrix

	1	2	3
1	0	1	0
2	1	0	1
3	0	0	1

Set of edges for the same graph

(1,2), (2,1), (2,3), (3,3)

Set-of-Edges Properties

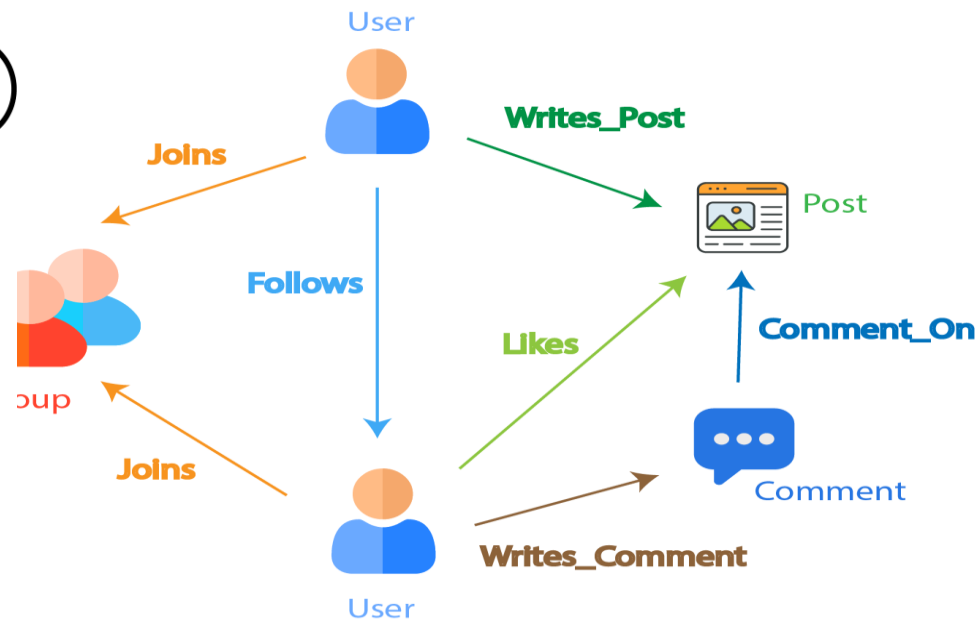
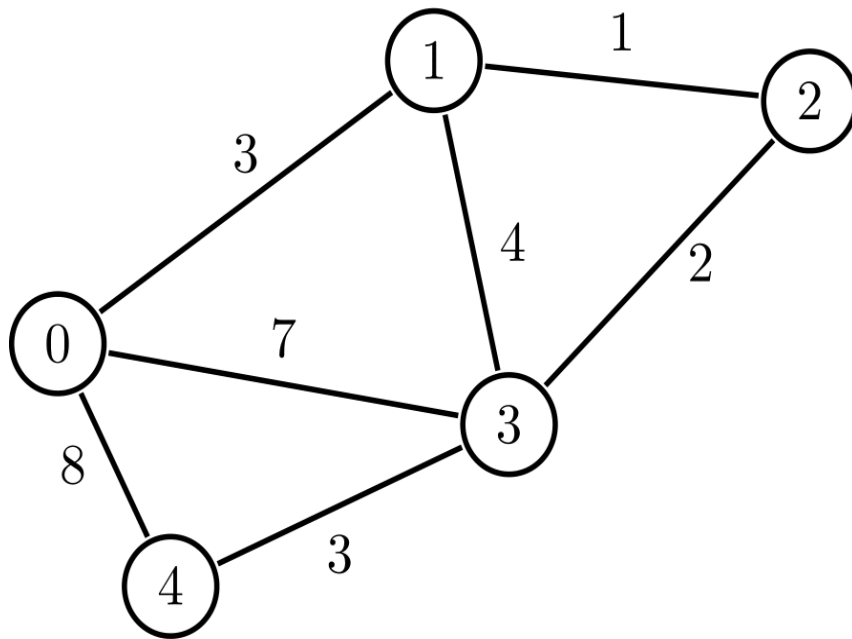
- Pro: This representation is still space-efficient for sparse graphs and it is easy to perform per-edge manipulations, e.g., inverting each edge.
- Pro: The uniform record format—each record is a pair of vertices, in contrast to adjacency lists, whose sizes can vary—can simplify programming and data processing.

Set-of-Edges Properties

- Con: Out of the three, this is the most challenging format for computations over inlinks or outlinks of a vertex. The records corresponding to a given vertex' inlinks and outlinks might be scattered all over the file storing the set of edges.

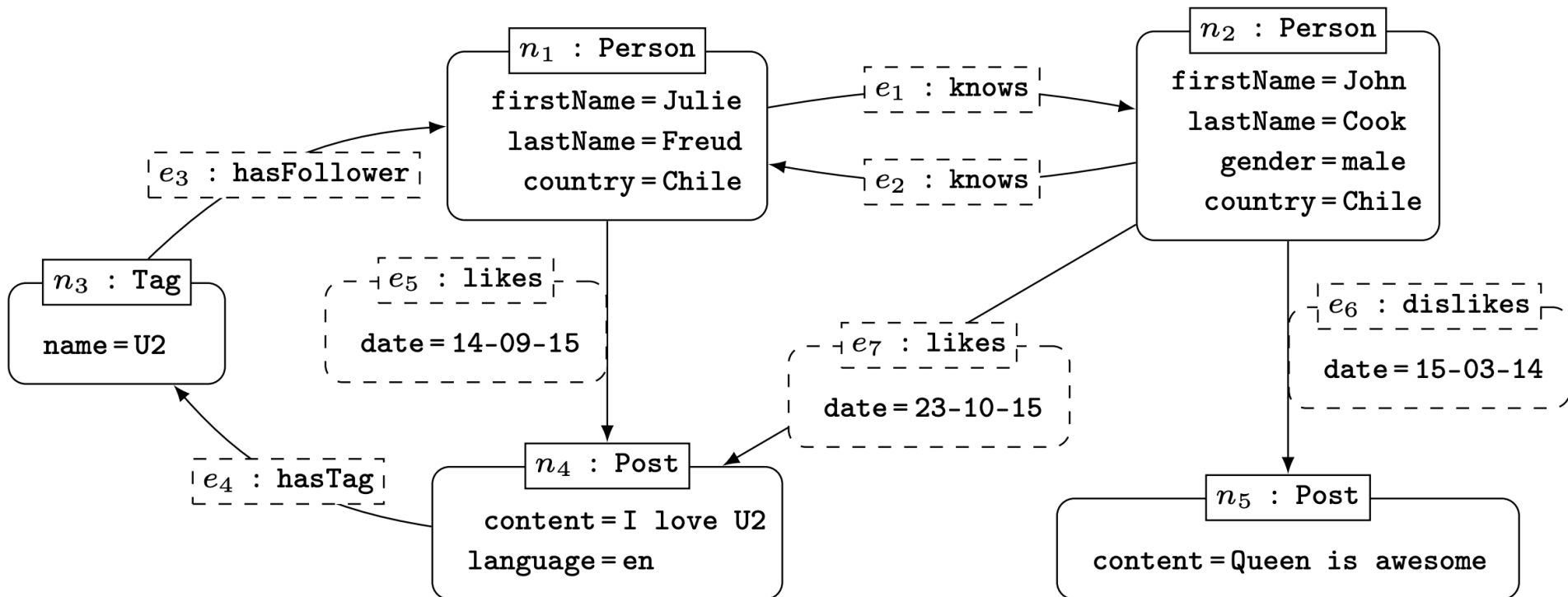
Other Graph Representations

- There are also richer ways of modeling graphs. The common extensions are adding weights and labels to edges.



Newer Graph Representations

- In the graph-structured data literature, the most popular "new" format for storing data is called the property graph.
- The idea is that both vertices and edges contain labels and attributes (properties).



Let us first look at the classic breadth-first algorithm for exploring the k -hop neighborhood of a given source vertex.

This algorithm forms the basis of two others we discuss later.

Breadth-First Search (BFS)

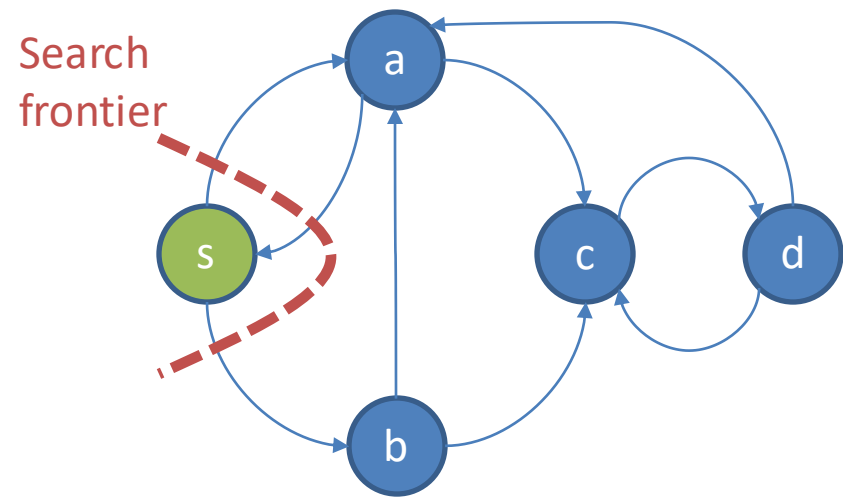
- Starting from a given source vertex s , breadth-first search first reaches all 1-hop neighbors of s , then all 2-hop neighbors, and so on.
- A vertex will be encountered repeatedly if and only if it can be reached through different paths.

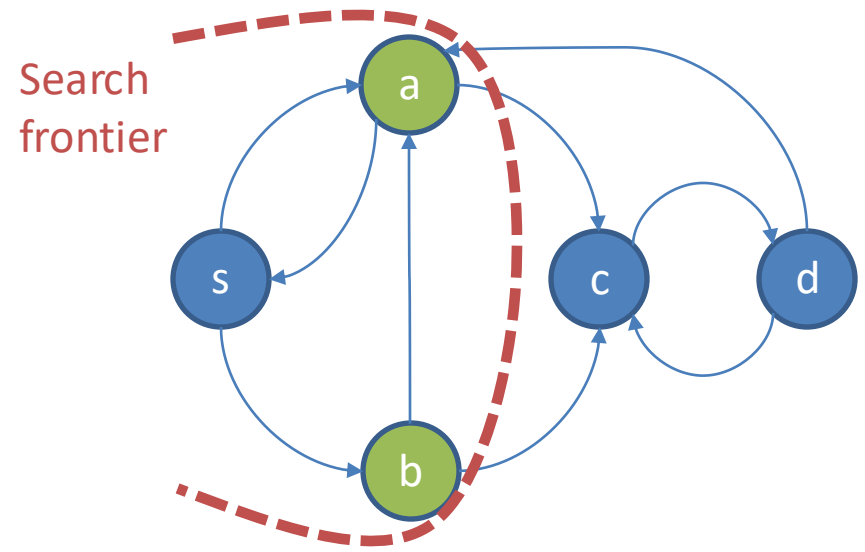
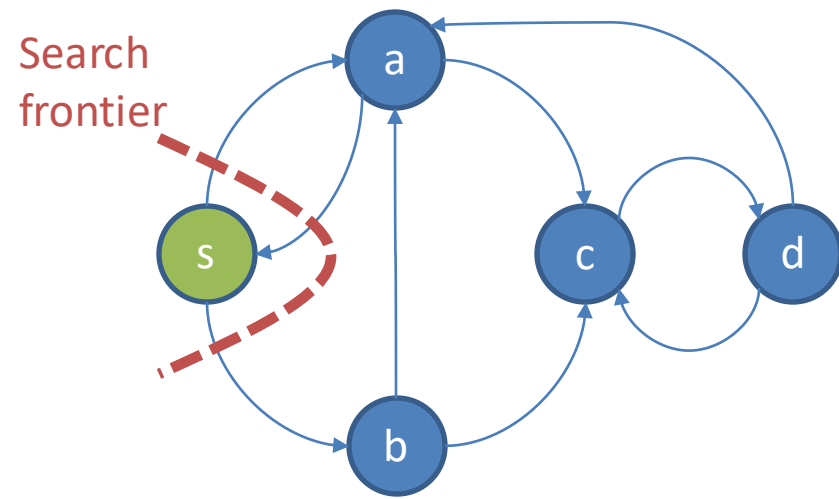
Breadth-First Search (BFS)

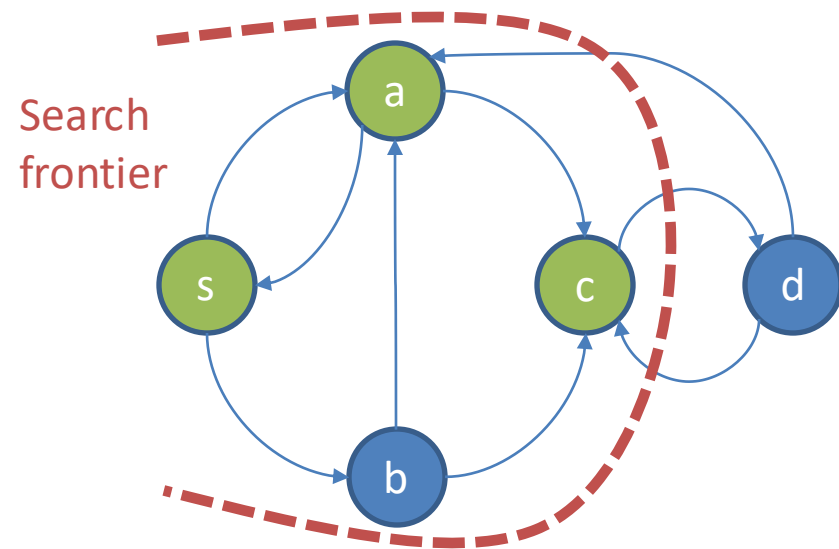
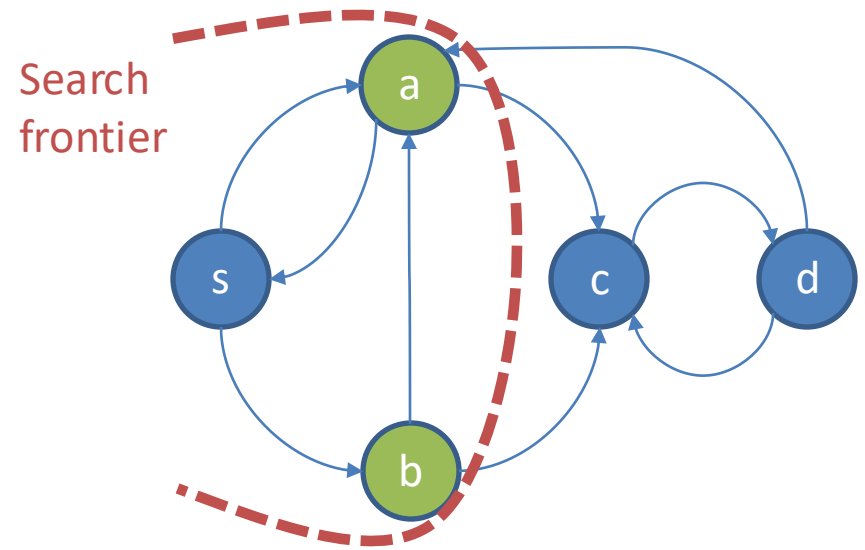
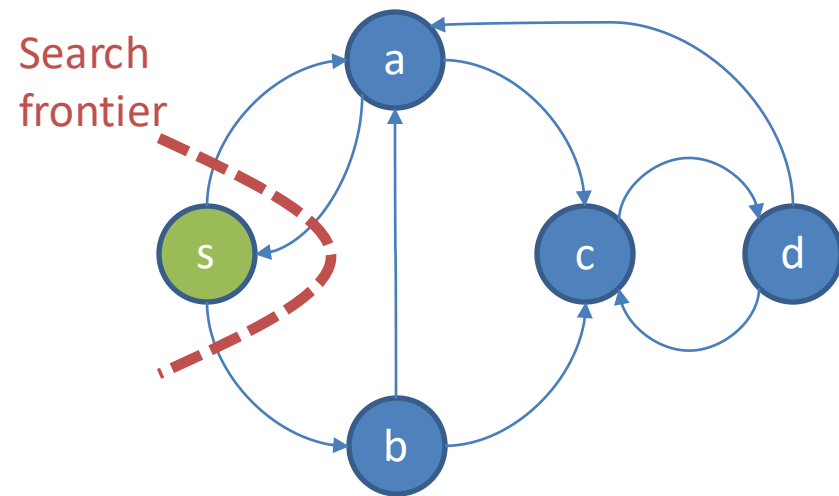
- If the graph has cycles, then the search process can continue indefinitely, repeatedly traversing the cycles.
- In an acyclic graph, breadth-first search will terminate after at most $|V|-1$ iterations—this is the length of the longest path possible in any acyclic graph.

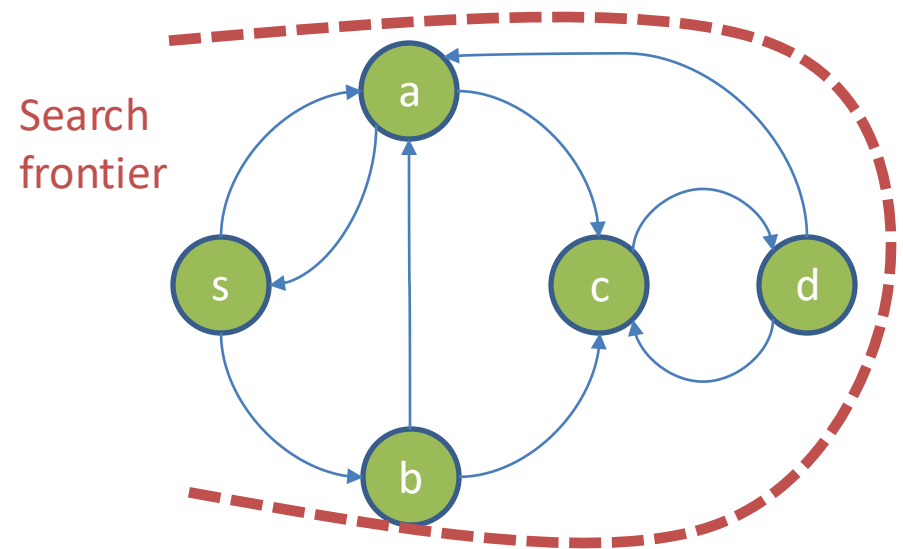
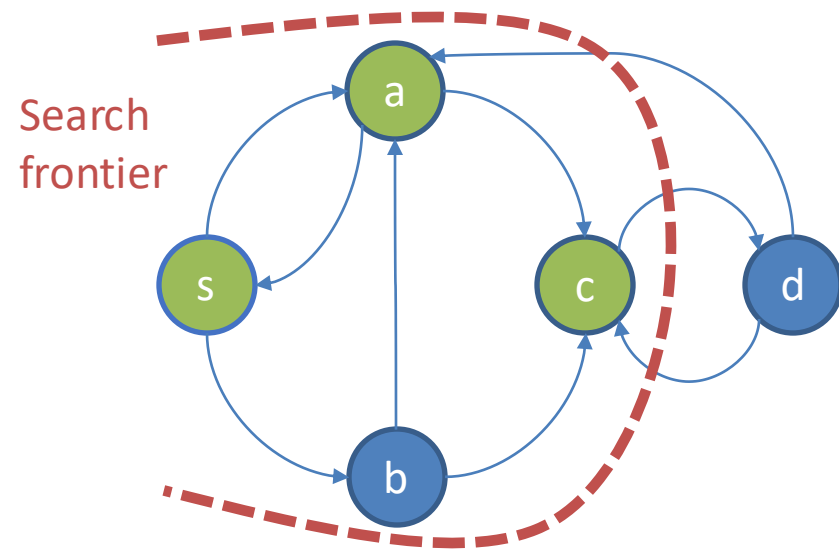
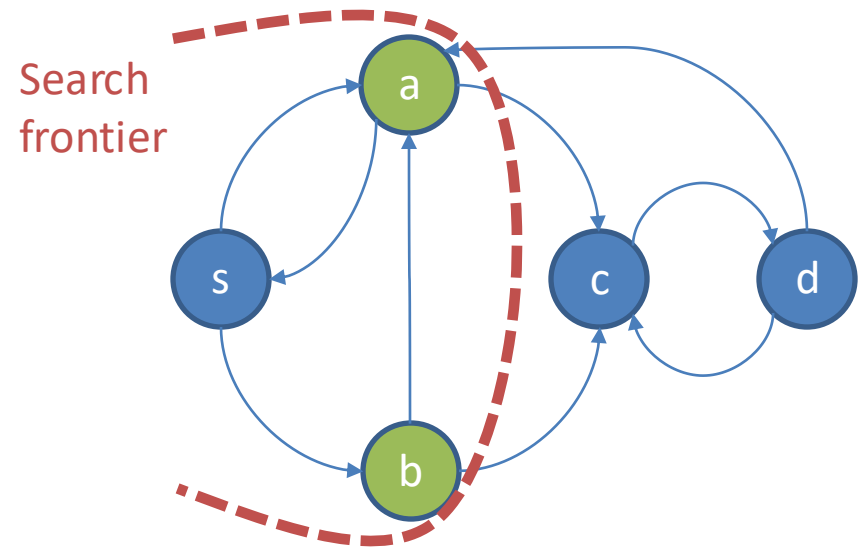
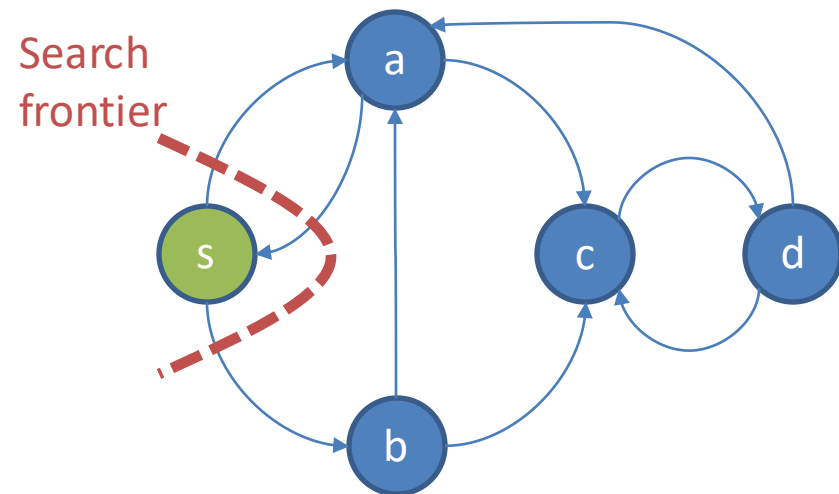
BFS Example

- The following images illustrate how breadth-first traversal reaches vertices in each iteration.
 - Green color highlights all vertices visited in the corresponding iteration. Notice how some vertices are visited repeatedly.
 - Due to the particular structure of the example graph, after iteration 3, each vertex will be reached in each of the following iterations.
- The **search frontier** illustrates how breadth-first traversal reaches new vertices.
 - Iteration 1: Vertices a and b are reached in exactly one hop from s.
 - Iteration 2: Vertex s is reached again (from a), and so is a (from b). Vertex c is reached for the first time.
 - Iteration 3: All vertices are reached, including d for the first time.









Parallel BFS

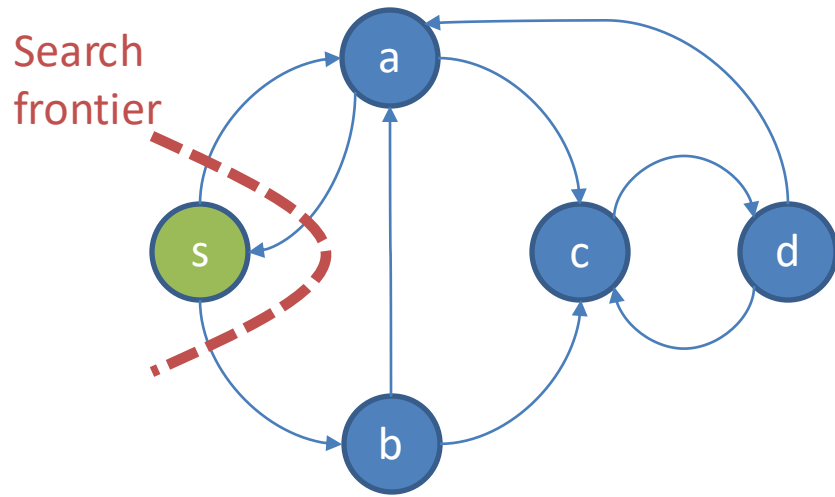
- We work with a graph represented in **adjacency-list** format. Parallel BFS relies on the following observation: If vertex v was reached in iteration i , then in iteration $(i+1)$ the algorithm must explore all adjacent vertices, i.e., all vertices in v 's adjacency list.
 - Each vertex can be processed independently. In the beginning, only source vertex s is “reached,” but in later iterations there could be many reached vertices, enabling effective parallelization.

Parallel BFS

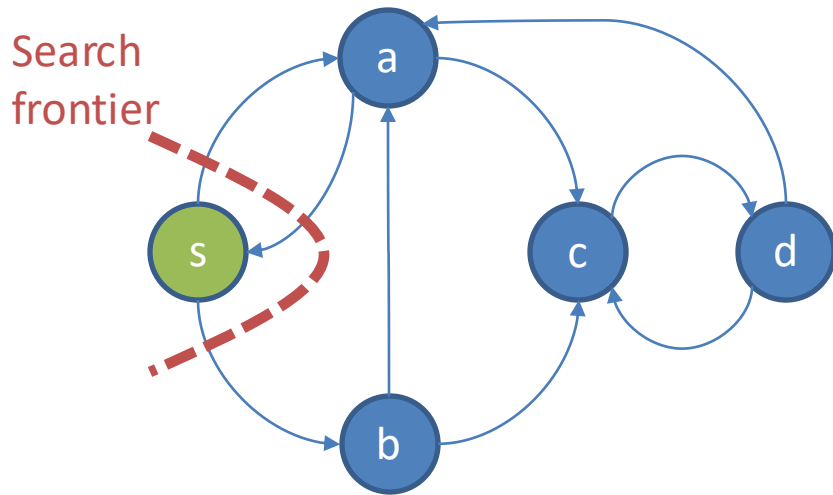
- To turn this idea into an algorithm, let $\text{active}(i)$ denote the set of vertices reached in iteration i . The next iteration replaces each active vertex by the vertices in its adjacency list as follows:
- Each task receives a partition of $\text{active}(i)$. For each vertex v in that partition, and for each vertex w in v 's adjacency list, emit w . Then $\text{active}(i+1)$ is the union of the emitted vertices.

Parallel BFS

- This algorithm has two problems:
 - A vertex may be emitted multiple times, requiring shuffling to eliminate duplicates.
 - How can the task access an adjacency list? For large graphs, it is infeasible to copy all adjacency lists (i.e., the entire graph) to all tasks. If a task receives only some adjacency lists, then it can only process those vertices!
- Let us return to our example to see what happens. Assume there are two tasks.



Initially, only source node s is active. The driver informs all tasks in iteration 1 about this.



Initially, only source node s is active. The driver informs all tasks in iteration 1 about this.

Task 0:

active: s

Adjacency lists:

s : a, b

b : a, c

d : a, c

for each active vertex v

for each vertex w in v 's adjacency list

emit w

Task 1:

active: s

Adjacency lists:

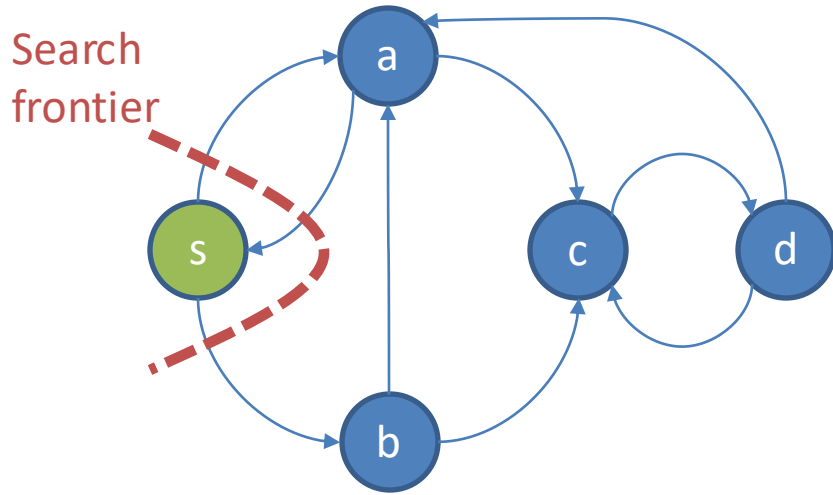
a : s, c

c : d

for each active vertex v

for each vertex w in v 's adjacency list

emit w



In iteration 1, all nodes in the adjacency list of s are emitted. This is done by task 0. Task 1 has no active node and hence does not emit anything. The output is sent to the task that owns the corresponding adjacency list.

Task 0:

active: s

Adjacency lists:

s : a, b

b : a, c

d : a, c

for each active vertex v

for each vertex w in v 's adjacency list

emit w

Task 1:

active: s

Adjacency lists:

a : s, c

c : d

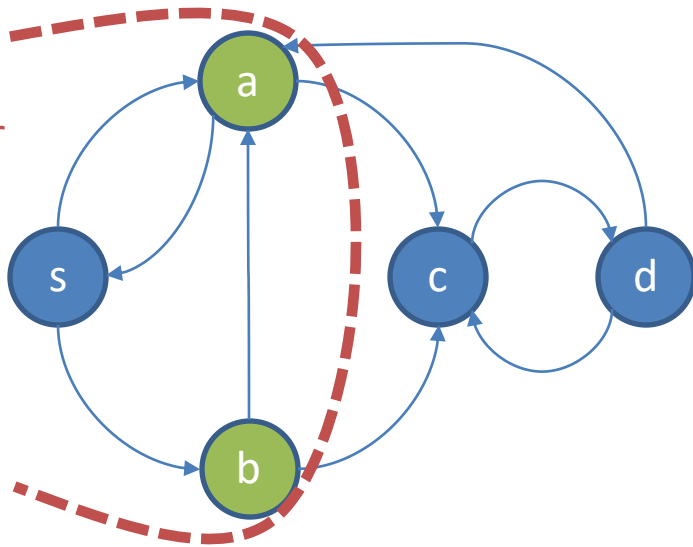
for each active vertex v

for each vertex w in v 's adjacency list

emit w



Search
frontier



Now vertices a and b are active. Notice that task 0 only knows about “its” vertex b, while task 1 only knows about a.

Task 0:

active: **b**

Adjacency lists:

s: a, b

b: a, c

d: a, c

for each active vertex v

for each vertex w in v 's adjacency list

emit w

Task 1:

active: **a**

Adjacency lists:

a: s, c

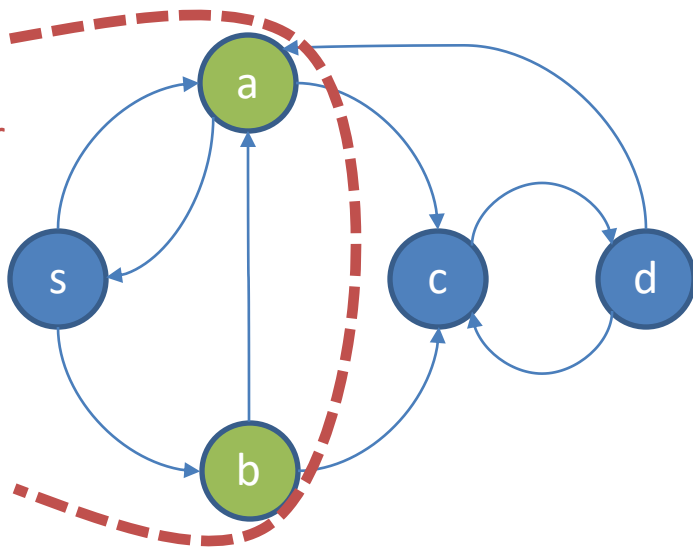
c: d

for each active vertex v

for each vertex w in v 's adjacency list

emit w

Search
frontier



In iteration 2, all nodes in the adjacency lists of a and b are emitted. Since a and b are “owned” by different tasks, this work is done in parallel.

Note that c is received twice by task 1. It is straightforward to deal with such duplicates locally.

Task 0:

active: b

Adjacency lists:

s: a, b

b: a, c

d: a, c

for each active vertex v

for each vertex w in v's adjacency list
emit w

Task 1:

active: a

Adjacency lists:

a: s, c

c: d

for each active vertex v

for each vertex w in v's adjacency list
emit w

a, c

c

s

BFS in MapReduce

- The communication of the emitted active vertices requires **shuffling**. In MapReduce, this means that each iteration is a full MapReduce job.
- The parallel algorithm requires a partition of the adjacency list to be kept in the memory of a task across iterations.

BFS in MapReduce

- This is *not possible* in MapReduce. Instead, each Map phase must read the adjacency lists again in every iteration.
- Hence both, newly active vertices and adjacency list data must be read, shuffled, and emitted in every iteration.

BFS in MapReduce

- To traverse the graph, what do we need to know? Given a vertex, we need:
 - (1) The vertex's adjacency list.
 - (2) An indicator to determine if the vertex is part of the frontier (i.e., active).
- We can create the appropriate object to implement this idea.

BFS in MapReduce

```
// Map processes a vertex object,  
// which contains an ID, adjacency list,  
// and active status.  
map(vertex N) {  
    // Pass along the graph structure  
    emit(N.id, N)  
  
    // If N is active, mark each vertex in  
    // its adjacency list as active  
    if (N.isActive)  
        for all id m in N.adjacencyList do  
            emit(m, "active")  
}
```

The driver program repeatedly calls this MapReduce job, each time passing the previous job's output directory as the input to the next.

```
// Reduce receives the vertex object for M and  
// between zero or more "active" flags.  
reduce(id m, [o1, o2,...]) {  
    isReached = false  
    M = NULL  
  
    for all o in [o1,o2,...] do  
        if isVertexObject(o) then  
            // The vertex object was found: recover graph structure  
            M = o  
        else  
            // An active flag was found: m should be active  
            isReached = true  
  
    // Update active status of vertex m  
    M.setActive( isReached )  
    emit(M)  
}
```

Avoiding Vertex Revisits

- The BFS algorithm will revisit vertices that are reachable through paths of different lengths from the start vertex.
- What if we only want to activate newly-reached vertices?
 - We need another flag `wasReached` to keep track of vertices that had been encountered in an earlier iteration.
 - A vertex becomes active only if it is newly reached.

Avoiding Vertex Revisits: Map

```
// Map processes a vertex object, which contains vertex id,  
// adjacency list, wasReached flag, and active flag.  
// Initially, only the start vertex satisfies wasReached and active.  
map(vertex N) {  
    // Pass along the graph structure only for vertices not reached yet  
    if (not N.wasReached) then  
        emit(N.id, N)  
  
    // If N is active, reset to inactive and explore its neighbors  
    if (N.active) {  
        N.active = false  
  
        // Set all neighbors to “reached” status  
        for all m in N.adjacencyList do  
            emit(m, true) // true indicates reached status for outlink destination  
    }  
}
```

Avoiding Vertex Revisits: Reduce

```
// Reduce receives the node object for node M and
// possibly "true" flags if the node was reached from an active node
reduce(id m, [d1,d2,...]) {
  isReached = false; M = NULL

  for all d in [d1,d2,...] do
    if isNode(d) then // The vertex object was found: recover graph structure
      M = d
    else // A "true" was found, indicating M was reached from an active vertex
      isReached = true

  // Skip if no vertex object was encountered, i.e., M had been reached earlier
  if (M is not NULL) { // This guarantees that M.wasReached = false
    // Update reached and active status of M if newly reached
    if (isReached) then {
      M.wasReached = true; M.active = true
    }
    emit(vertex M)
  }
}
```

Improving the MapReduce Algorithm

- Can we avoid the expensive cycle of reading the graph in the Map phase, sending it from Mappers to Reducers, and then writing it in the Reduce phase, which repeats **in every iteration**?
- Recall that distributed-file-system accesses are expensive, and the network is a precious shared resource.

Improving the MapReduce Algorithm

- One could attempt to use the file cache for this purpose.
 - However, it is very limited in the sense that it makes the same data available for each Mapper or Reducer but does not support managing different partitions on different machines.
- In general, MapReduce lacks mechanisms to exploit the repetitive structure of iterative computations. Spark's RDD abstraction addresses exactly this issue.

BFS in Spark

- We discuss the Spark Scala pseudo-code for BFS with RDDs. The version with DataSet is left as a voluntary challenge.
- The Spark program separates static data (the graph structure) from evolving data (the active vertex set).
 - This leverages cached RDD partitions and avoids the shuffling of the adjacency list data we saw in the MapReduce implementation.

BFS in Spark

- On the downside, the new active-status information must be joined with the graph data, so that tuples of type (vertexID, adjacencyList) and (vertexID, activeStatus) are joined on vertexID.
- This join normally requires shuffling, but careful co-partitioning avoids shuffling (we will see this in a future lecture).

Spark Implementation

```
// Assume the input file contains in each line a vertex ID and its adjacency list.
// Function getAdjListPairData returns a pair of vertex ID and its adjacency list, creating a pair RDD.
graph = sc.textFile(...).map( line => getAdjListPairData(line) )

// Add some code here to make sure that graph has a Partitioner. This is needed for avoiding shuffling
// in the join below.
[...]

// Ask Spark to try and keep this pair RDD around in memory for efficient re-use
graph.persist()

// The active vertex set initially only contains source vertex s, which needs to be passed through the
// context. Create a pair RDD for s. Value 1 is a dummy value.
activeVertices = sc.parallelize( ("s", 1) )

// Function extractVerticesAsPair returns each vertex id m in the adjacency list as (m, 1),
// where 1 is a dummy value.
for (iterationCount <- 1 to k) {
  activeVertices = graph.join( activeVertices )
    .flatMap( (id, adjList, dummy) => extractVerticesAsPair(adjList) )
    .reduceByKey( (x, y) => x ) // Remove duplicate vertex occurrences
}
```

Real Spark Code

- For an example of BFS, look at the transitive closure program, here the version from the Spark 2.4.0 distribution
 - <http://khoury.northeastern.edu/home/mirek/code/SparkTC.scala>
- It first generates a random graph and converts it to a pair RDD.
- Note the use of join to extend each path.

Next, we explore two important graph algorithms: single-source shortest path (SSSP) and PageRank (PR), starting with SSSP.

The parallel versions of both are built on BFS.

Single-Source Shortest Path

- Consider the well-known problem of finding the shortest path from a source vertex s to all other vertices in a graph.
- The length of a path is equivalent to the total weight of all edges belonging to the path.
- Let all edges in the graph have *non-negative weights*.

Single-Source Shortest Path

- We will first discuss a classic sequential algorithm called **Dijkstra's algorithm**, which finds the solution very efficiently.
- Then we will explore how to solve the problem in parallel.

Dijkstra's Algorithm

1. Let $d[v]$ denote the distance of vertex v from source s . Initialize $d[v]$ by setting $d[s]=0$ and $d[v]=\infty$ for all $v \neq s$.
2. Insert all graph vertices into a priority queue sorted by distance. (Initially s will be first, while all other vertices appear in some arbitrary order.)

Dijkstra's Algorithm

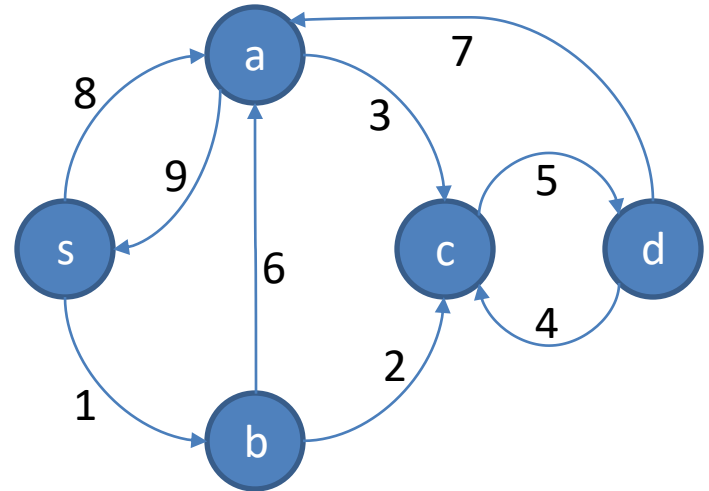
3. Repeat until the queue is empty
 1. Remove the first vertex u from the queue.
Output $(u, d[u])$. (The shortest path to u was found.)
 2. For each vertex v in u 's adjacency list do
 1. If v is in the queue and $d[v] > d[u] + \text{weight}(u, v)$, then set $d[v]$ to $d[u] + \text{weight}(u, v)$. (A shorter path to v was found through u .)

Dijkstra's Algorithm Example

Output:

Priority queue: $(s, 0)$, (d, ∞) , (a, ∞) , (b, ∞) , (c, ∞)

Vertex being processed:



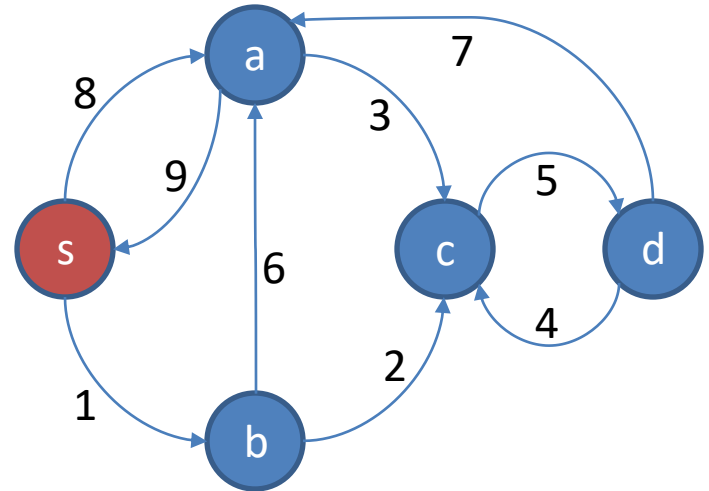
Initially all vertices and their current distance values are in the priority queue.

Dijkstra's Algorithm Example

Output: $(s, 0)$

Priority queue: $(\cancel{s, 0}), (d, \infty), (a, \infty), (b, \infty), (c, \infty)$

Vertex being processed: $s: (a, 8), (b, 1)$



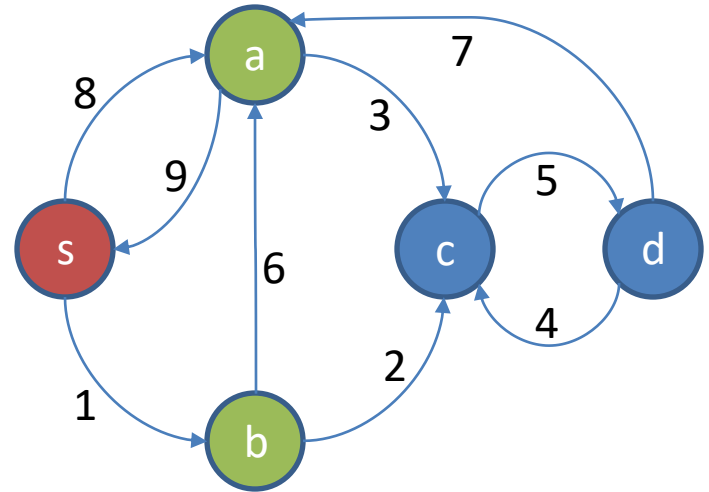
The first vertex, s , is removed and output.

Dijkstra's Algorithm Example

Output: (s,0)

Priority queue: (b,1), (a,8), (d,∞), (c,∞)

Vertex being processed: s: (a,8), (b,1)



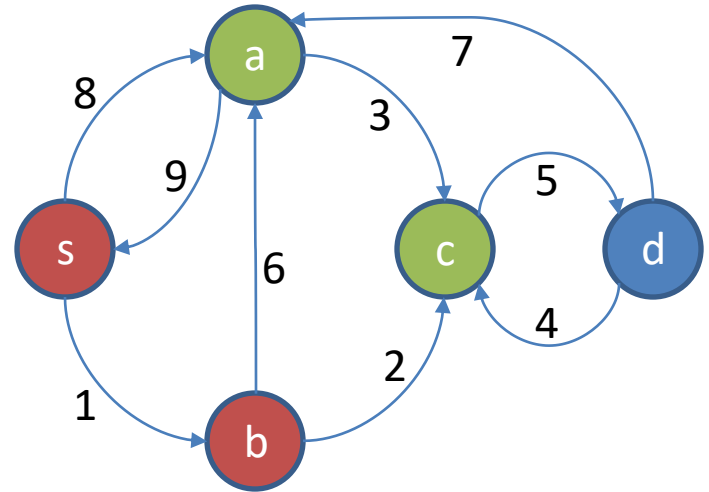
For all vertices in the adjacency list of the removed vertex s , the distances are updated. For example, since $d[s]=0$ and entry $(a,8)$ appears in the adjacency list, there exists a path to vertex a of length $0+8 = 8$.

Dijkstra's Algorithm Example

Output: (s,0), (b,1)

Priority queue: (c,3), (a,7), (d,∞)

Vertex being processed: b: (a,6), (c,2)



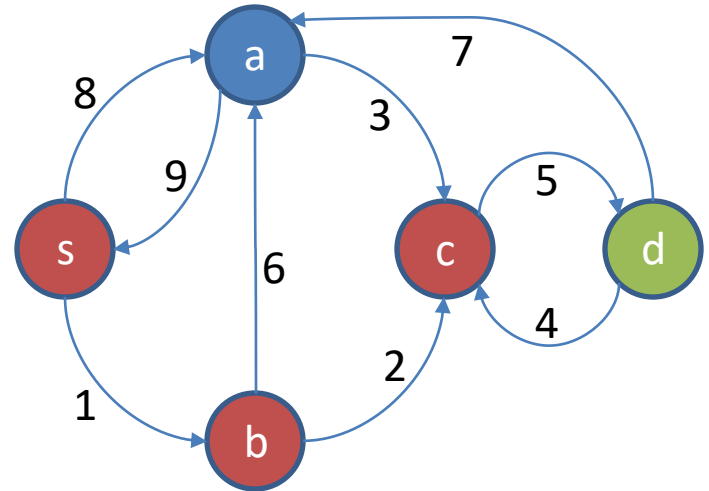
Now vertex b is removed and output, followed by updating of distances of vertices a and c.

Dijkstra's Algorithm Example

Output: (s,0), (b,1), (c,3)

Priority queue: (a,7), (d,8)

Vertex being processed: c: (d,5)



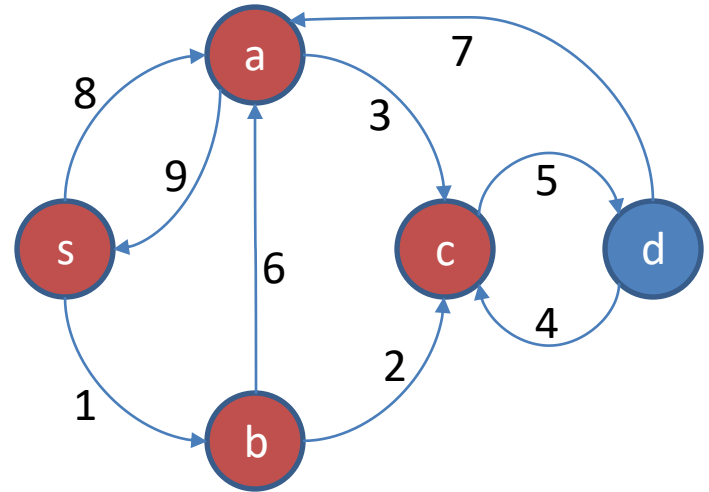
Now vertex c is removed and output, followed by updating of the distance of vertex d.

Dijkstra's Algorithm Example

Output: (s,0), (b,1), (c,3), (a,7)

Priority queue: (d,8)

Vertex being processed: a: (c,3), (s,9)



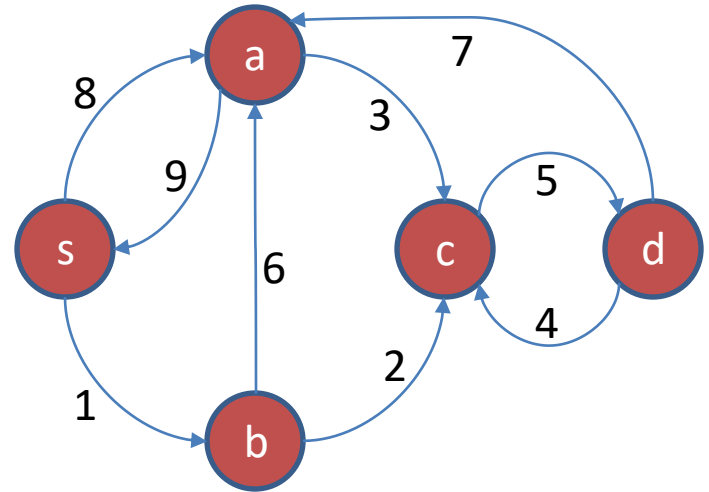
Now vertex a is removed and output. Since vertices c and s in a's adjacency list are not in the queue, no distance updates are performed.

Dijkstra's Algorithm Example

Output: (s,0), (b,1), (c,3), (a,7), (d,8)

Priority queue:

Vertex being processed: d: (c,4), (a,7)



Finally vertex d is output and the algorithm terminates.

Parallel Single-Source Shortest Path

- Dijkstra's algorithm is elegant and very efficient for sequential execution, but difficult to adapt for parallel execution: vertices are removed from the priority queue one-by-one.
- One cannot remove multiple vertices at once for parallel processing, without risking incorrect results.

Parallel Single-Source Shortest Path

- Recall the example where after processing vertex s , entries $(b,1)$ and $(a,8)$ were the first entries in the queue.
- Removing both and processing them in parallel would not have worked, because the shortest path to a is going through b .

Parallel Single-Source Shortest Path

- Example from the deck: After processing s , the queue is: $(b, 1), (a, 8), \dots$
- It may look tempting to process both b and a in parallel. But that is unsafe. Why? $s \rightarrow a = 8, s \rightarrow b = 1, b \rightarrow a = 6$
 - So the path through b is better: $s \rightarrow b \rightarrow a = 1 + 6 = 7$
- If we process a too early, we treat distance 8 as final, even though the true shortest distance is 7.

Parallel Single-Source Shortest Path

- While removing multiple vertices at once jeopardizes correctness in *any* parallel environment, the use of a single priority queue represents a particular challenge for MapReduce and Spark.
- It is not clear how to best implement such a shared data structure in a system without an efficiently implemented shared-memory abstraction.

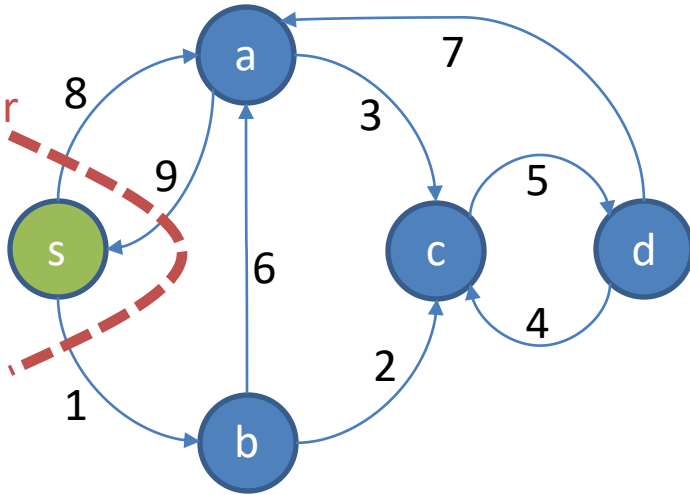
Shortest Path using BFS

- Even without the priority queue, which is the key component of Dijkstra's algorithm, the parallel solution can still rely on the following property:
- If there exists a path of length $d[u]$ to some vertex u , then for each vertex v in u 's adjacency list there exists a path of length $d[u] + \text{weight}(u, v)$.

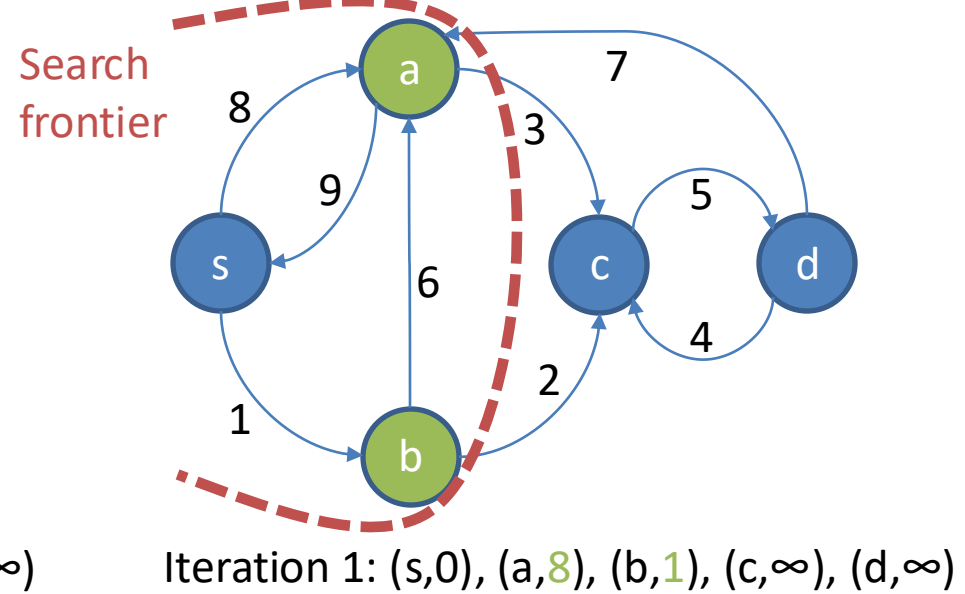
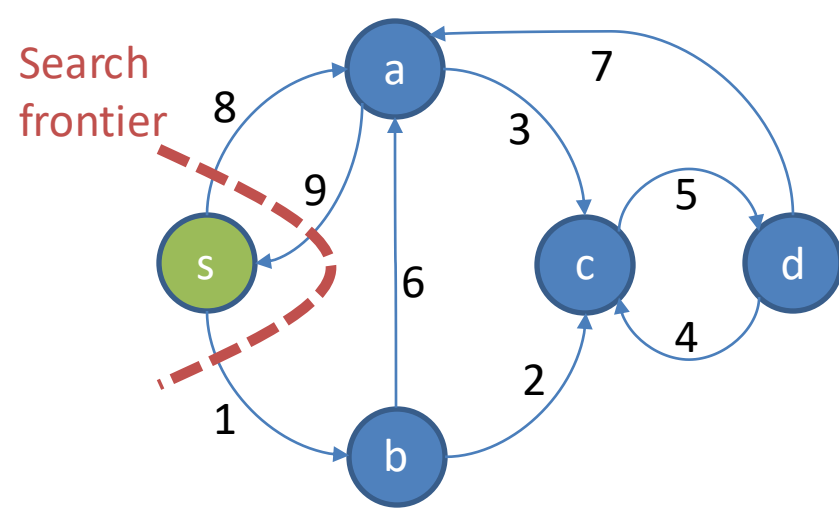
Shortest Path using BFS

- We design an algorithm that exploits this property by systematically exploring the graph using BFS:
 - In the first iteration, find all vertices reachable from s in exactly one hop and update their distances.
 - In the second iteration, find all vertices reachable from s in exactly two hops and update their distances. And so on.
 - Continue this process until the shortest path to each vertex was found.
- We illustrate the algorithm next for the same example graph.

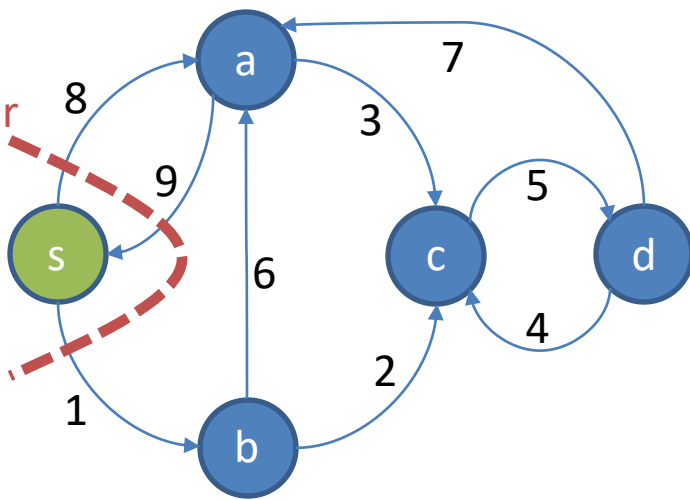
Search
frontier



Initial state: $(s, 0)$, (a, ∞) , (b, ∞) , (c, ∞) , (d, ∞)

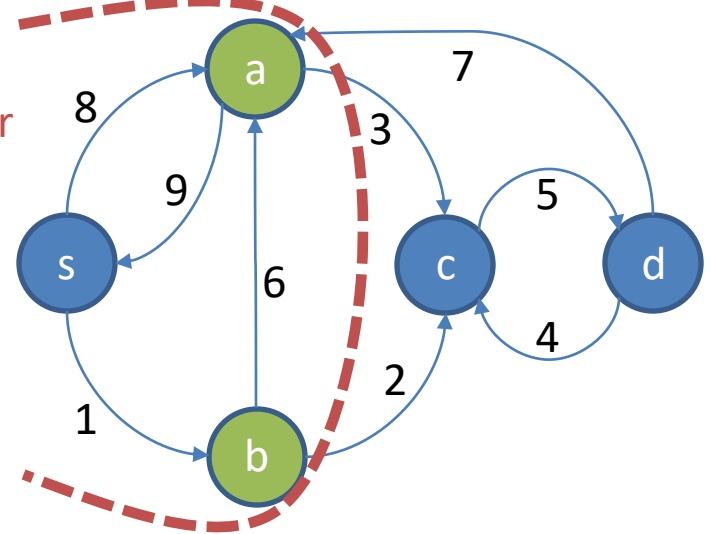


Search
frontier



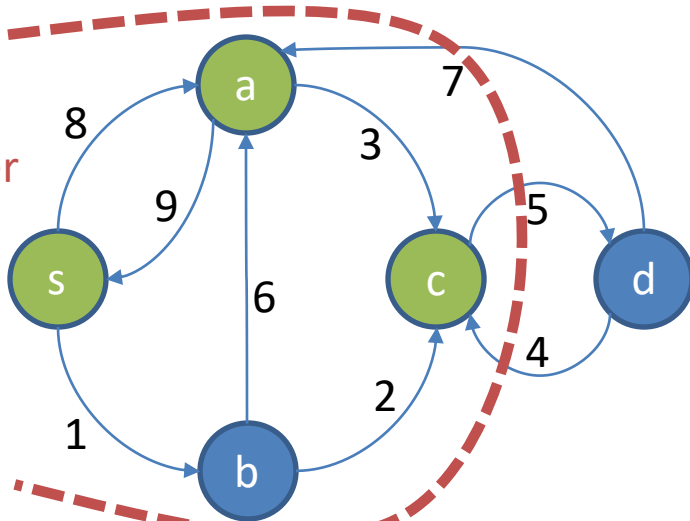
Initial state: $(s,0)$, (a,∞) , (b,∞) , (c,∞) , (d,∞)

Search
frontier

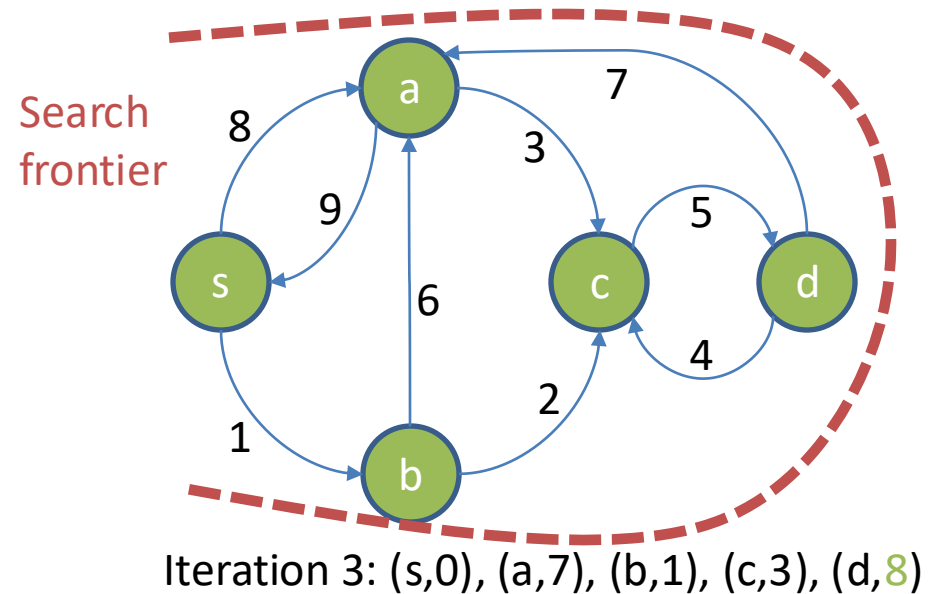
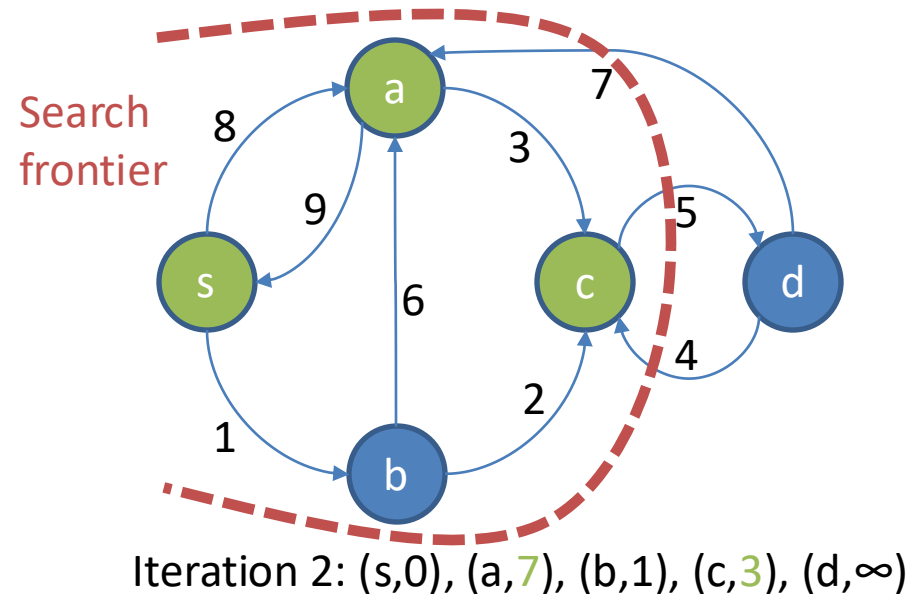
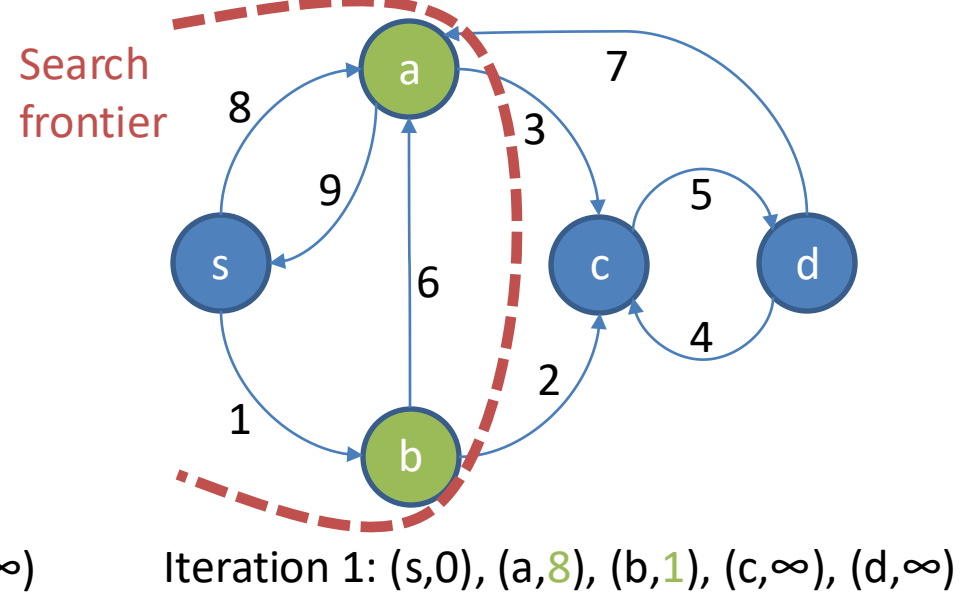
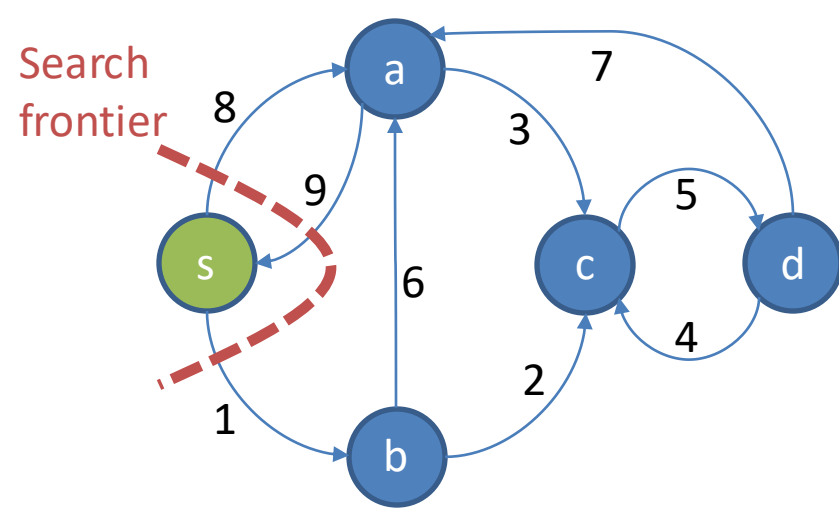


Iteration 1: $(s,0)$, $(a,8)$, $(b,1)$, (c,∞) , (d,∞)

Search
frontier



Iteration 2: $(s,0)$, $(a,7)$, $(b,1)$, $(c,3)$, (d,∞)



When to Stop Iterating

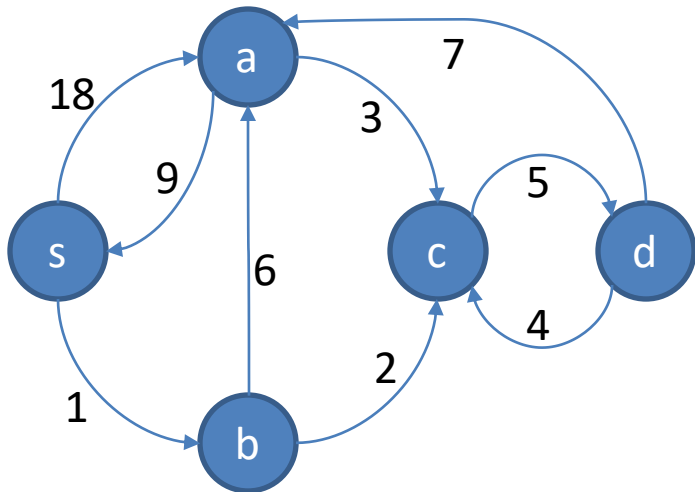
- How does the algorithm know when all shortest paths were found?

When to Stop Iterating

- How does the algorithm know when all shortest paths were found?
- If all edges have the same weight, stop iterating as soon as no vertex has distance ∞ any more: Later iterations can only find longer paths. This can be detected using a global counter or accumulator.

When to Stop Iterating

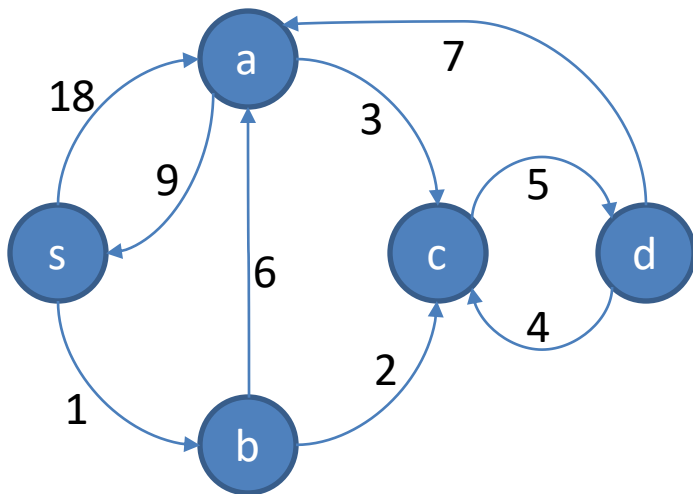
- If edges have **different weights**, then a “detour” path can have a lower total weight than a more “direct” connection and hence we cannot stop as soon as all vertex distances are finite. Instead, iterations must continue until no shortest distance changes.



This is an example for a graph where a “detour” path consisting of four edges is shorter than the direct path. Path (s, a) has length 18, while path (s, b, c, d, a) has length 15. Detecting this shortest path requires 4 iterations.

When to Stop Iterating

- In a graph with **cycles of negative weight** the algorithm never terminates: each traversal of the cycle reduces the distances of all its vertices indefinitely.



This is an example for a graph where a “detour” path consisting of four edges is shorter than the direct path. Path (s, a) has length 18, while path (s, b, c, d, a) has length 15. Detecting this shortest path requires 4 iterations.

Single-Source Shortest Path in MapReduce

- The algorithm is virtually identical to BFS, but also keeps track of the shortest distance found for each vertex so far:
 - Map processes a single vertex u , emitting $d[u] + \text{weight}(u, v)$ for each vertex v in u 's adjacency list.
 - Reduce collects the newly computed distances for all inlinks of vertex v and determines if any of them is shorter than its currently known shortest distance.
- The driver program repeatedly calls the MapReduce program as many times as necessary, exploring ever longer paths.

MapReduce Code for a Single Iteration

```
// The vertex object now also contains the
// current min distance and edges in the
// adjacency list have a weight
map(vertex N) {
    // Pass along the graph structure
    emit(N.id, N)

    // Compute the distance for each outlink
    for all m in N.adjacencyList do
        emit( m, N.distance + weight(n,m) )
}
```

```
// Reduce receives the vertex object for vertex m and
// the newly computed distances for m's inlinks
reduce(id m, [d1, d2,...]) {
    dMin =  $\infty$ 
    M = NULL

    for all d in [d1, d2,...] do
        if isVertex(d) then
            // The vertex object was found: recover graph structure
            M = d
        else
            // A distance value for an inlink was found: keep track
            // of the minimum.
            if d < dMin then dMin = d

    // Update distance of vertex m if necessary.
    if dMin < M.distance then M.distance = dMin
    emit( m, M )
}
```

MapReduce Algorithm Analysis

- Each iteration of the algorithm performs a large amount of work:
 - How much data is transferred from Map to Reduce?
 - How many reduce function calls are made?

MapReduce Algorithm Analysis

- Each iteration of the algorithm performs a large amount of work:
 - The entire graph is read from the distributed file system, transferred from Mappers to Reducers, and then, with updated distance values, written to the distributed file system.
 - For every vertex u , no matter if it potentially lies on a shorter path to another vertex v or not, distance $d[u] + \text{weight}(u, v)$ is computed and sent to the Reducers.
 - For every vertex v , no matter if its shortest path was already found in previous iterations or not, the Reduce function is executed to re-compute the shortest distance.

MapReduce Algorithm Analysis

- This brute-force approach performs many irrelevant computations:
 - What happens to nodes far away from the source in iteration 1?
 - What if we already found the shortest distance?

MapReduce Algorithm Analysis

- This brute-force approach performs many irrelevant computations:
 - In early iterations, Map computes distances for vertices that have not yet been reached, therefore still have infinity distance.
 - In later iterations, the program keeps re-computing paths for vertices whose shortest path was already found.
- Contrast this with the elegance of Dijkstra's algorithm, which avoids this irrelevant computation, but needed the priority queue in order to achieve this.

Improving the MapReduce Algorithm: Avoiding Useless Work

- Can we avoid processing vertices that do not improve any distances in the iteration?
 - If a vertex u has distance $d[u]=\infty$, then it cannot help reduce the distance for any of the vertices in its adjacency list.
 - Assume vertex u had the same distance $d[u]=x$ in iterations i and $(i+1)$. For any vertex v in its adjacency list, the Map function call for u would emit the same value $x+\text{weight}(u,v)$ in both iterations. Since this value was already included in the Reduce computation in iteration i , it cannot result in improvements in iteration $(i+1)$.

Improving the MapReduce Algorithm: Avoiding Useless Work

- Like BFS, the program distinguishes between “active” and “inactive” vertices. Active vertices are those that could potentially help reduce the distance for another vertex.
- We define a vertex to be **active** if and only if its distance value changed in the previous iteration.
 - The only exception to this rule is source vertex s , which is set to “active” before the first iteration. Note that a vertex that was active in one iteration could become inactive in the next, and vice versa.

Improving the MapReduce Algorithm:

Avoiding Useless Work

- It is easy to prove that a vertex whose distance reached the final value, i.e., the shortest-path distance from s , will remain inactive afterwards.
- Hence the algorithm can stop iterating as soon as all vertices become inactive.

Improved MapReduce Code for a Single Iteration

```
// The vertex object now also keeps track
// if the vertex is active.
map( vertex N ) {
    // Pass along the graph structure
    emit(N.id, N)

    // Compute the distance for each outlink
    // of an active vertex
    if N.isActive {
        N.setActive(false)

        for all m in N.adjacencyList do
            emit( m, N.distance + weight(n,m) )
    }
}
```

```
// Reduce receives the vertex object for vertex m and
// the newly computed distances for m's inlinks
reduce(id m, [d1, d2,...]) {
    dMin =  $\infty$ ; M = NULL

    for all d in [d1, d2,...] do
        if isVertex(d) then
            // The vertex object was found: recover graph structure
            M = d
        else
            // A distance value for an inlink was found: keep track
            // of the minimum.
            if d < dMin then dMin = d

    // Update distance of vertex m if necessary.
    if dMin < M.distance then {
        M.distance = dMin
        // The distance change can affect the distance for
        // nodes in the adjacency list, hence set status to active
        M.setActive(true)
    }
    emit( m, M )
}
```

Spark Implementation

```
// Each line in the input file contains a vertex ID and its adjacency list. Function getAdjListPairData returns
// a pair (vertex ID, adjacency list), creating a pair RDD. An entry in the adjacency list is a pair of destination
// vertex id and edge weight. The source node's adjacency list has a special "source" flag for setting of
// the initial distances.
graph = sc.textFile(...).map( line => getAdjListPairData(line) )

// Add code here to make sure the graph has a Partitioner. This avoids shuffling in the join below.

// Tell Spark to try and keep this pair RDD in memory for efficient re-use
graph.persist()

// Create the initial distances. mapValues ensures that the same Partitioner is used as for the graph RDD.
distances = graph.mapValues( (id, adjList) => hasSourceFlag(adjList) match {
                                case true => 0
                                case _ => infinity }

// Function extractVertices returns each vertex id m in n's adjacency list as (m, distance(n)+w),
// where w is the weight of edge (n, m). It also returns n itself as (n, distance(n))
for (iterationCount <- 1 to k) {    // Use Accumulator instead to determine when last iteration is reached
    distances = graph.join( distances )
        .flatMap( (n, adjList, currentDistanceOfN) => extractVertices(adjList, currentDistanceOfN) )
        .reduceByKey( (x, y) => min(x, y) )    // Remember the shortest of the distances found
}
```

We now explore the famous PageRank (PR) algorithm.

Compare and contrast the parallel version to SSSP.

PageRank

- PageRank was popularized by Google as a measure for evaluating the importance of a Web page. Intuitively it assigns greater importance to pages that are linked from many other important pages.
- PageRank captures the probability of a “random Web surfer” landing on a given page. The random Web surfer can reach a page by jumping to it or by following the link from another page pointing to it.

PageRank

- PageRank helps identify the most relevant results for a keyword query.
- For example, consider query “Northeastern University.” A person entering these keywords will most likely be looking for the Northeastern.edu homepage.
- If a search engine only considers traditional measures of importance such as term frequency, it might highly rank a spam page containing term “Northeastern” many times.
 - Assuming that Northeastern.edu will be linked from many more important pages than the spam site, taking into account the PageRank value can help boost its rank in the result list.

PageRank Definition

- The importance of a Web page is measured by the probability that a random Web surfer will land on it. They can reach a page either by typing a random URL into the browser's address field (aka perform a **random jump**) or by following a random **link** from the page they currently visit. The probability of doing the latter is α .
- Formally, $P(n) = (1 - \alpha) \frac{1}{|V|} + \alpha \sum_{m \in L(n)} \frac{P(m)}{C(m)}$, where
 - $|V|$ is the number of pages (vertices) in the Web graph considered.
 - α is the probability of the surfer following a link (vs. $1-\alpha$ for a random jump).
 - $L(n)$ is the set of all pages in the graph linking to n .
 - $P(m)$ is the PageRank of another page m .
 - $C(m)$ is the out-degree of page m , i.e., the number of links on that page.
- The 2 main terms in the formula correspond to the 2 ways of reaching n :
 - The random surfer reaches it via a random jump if they (1) decide to make a random jump (probability $1-\alpha$) and (2) choose exactly this 1 out of $|V|$ possible Web pages.
 - They reach it via a link if they (1) decide to follow a link (probability α), (2) were visiting another page m (which they happen to be currently visiting with probability $P(m)$) that links to n , and (3) selected the 1 out of $C(m)$ links on m pointing to n .

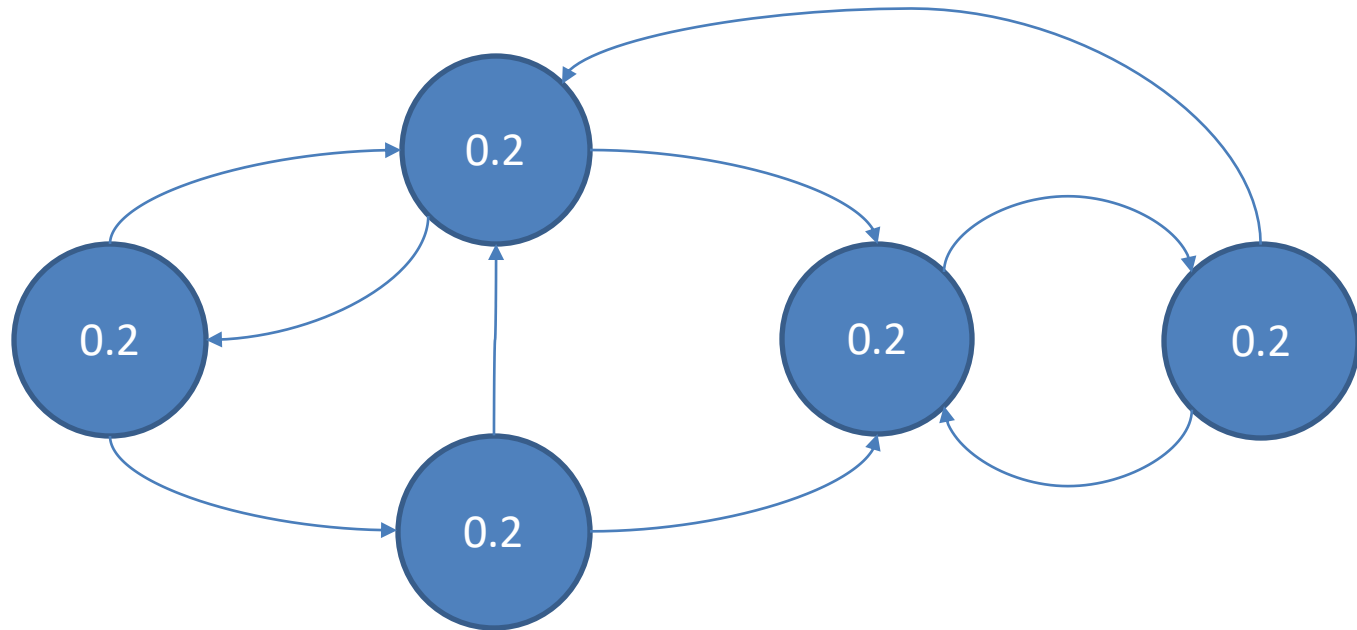
Iterative PageRank Computation

- Notice that the definition of PageRank creates a “chicken-and-egg problem”.
- To compute the PageRank of page n , we need to know the PageRank of all other pages linking to it, which in turn might depend on n 's PageRank.

Iterative PageRank Computation

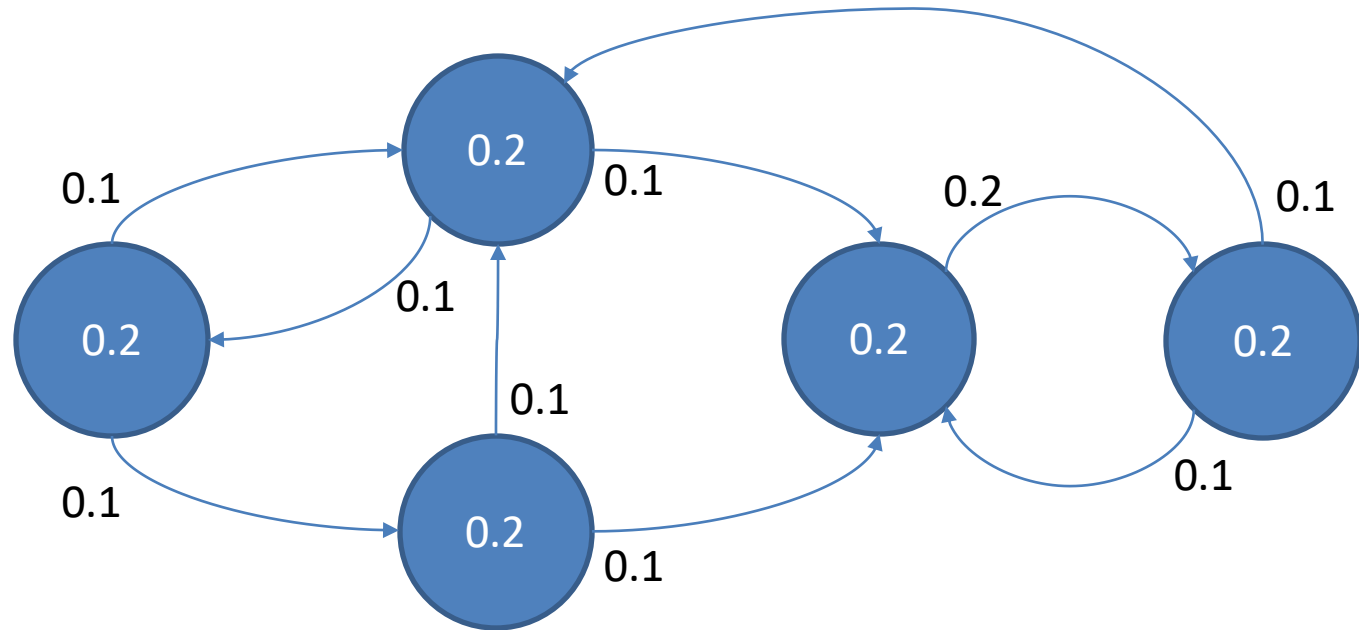
- Fortunately, this recursive definition admits an iterative algorithm for computing all PageRank values in the graph.
- Starting with some initial values, each iteration computes the new PageRank for all pages. This process continues until a **fixpoint** is reached, meaning that the value for every page does not change any more.

Initial State



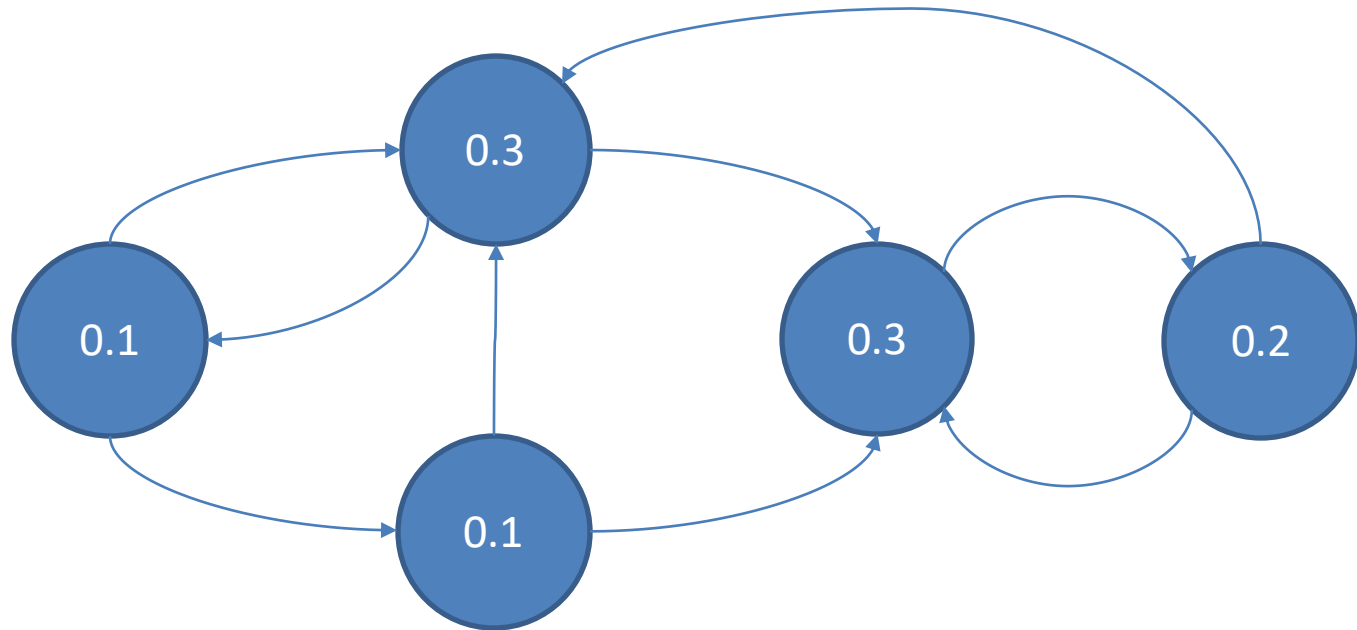
Assume that all PageRank values are initialized to 0.2. For simplicity we set $\alpha=1$, i.e., the random surfer only follows links.

Iteration 1: PageRank Transfer



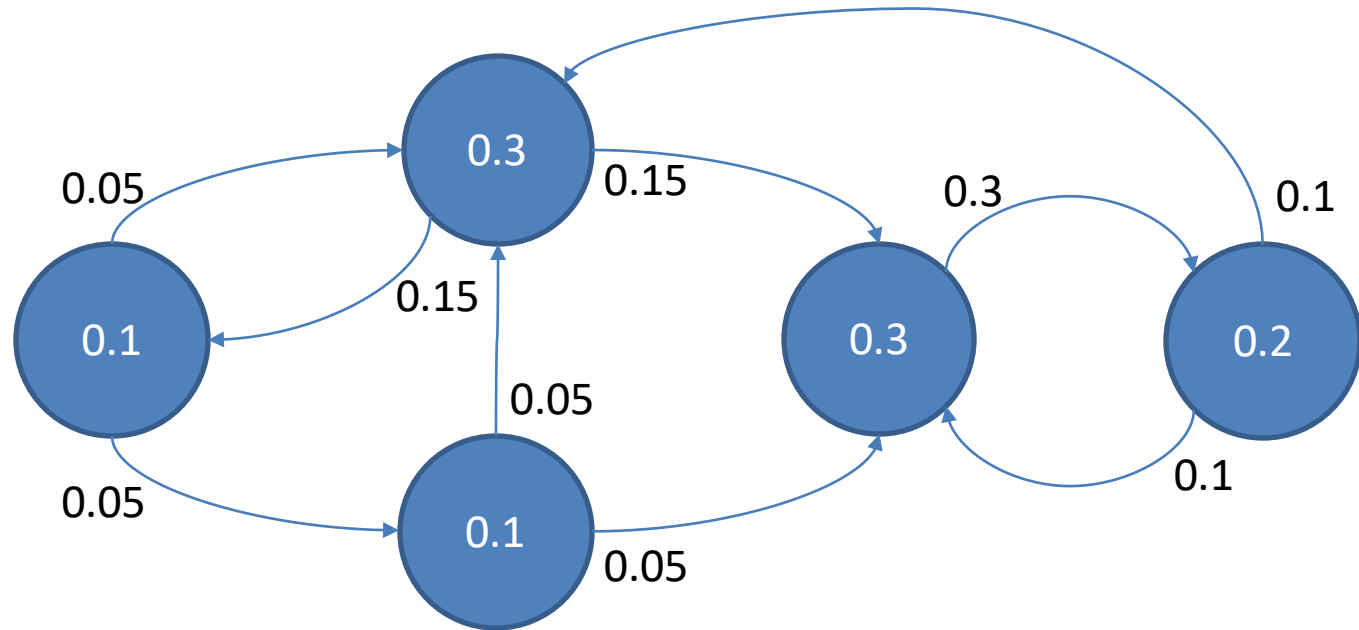
Since $\alpha=1$, each page passes on its full PageRank value, distributed equally over the outgoing links.

Iteration 1: Updated PageRank



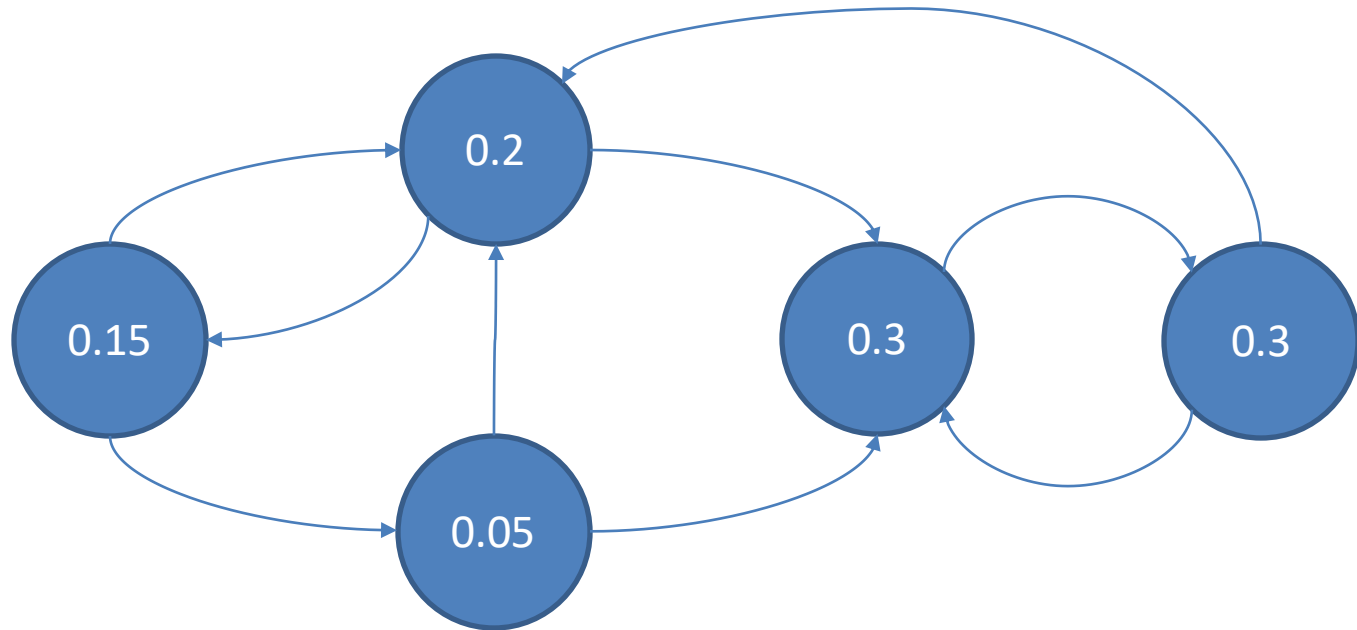
It is clearly visible how some pages receive more weight than others.

Iteration 2: PageRank Transfer



Since $\alpha=1$, each page again passes on its full PageRank value along the outgoing links.

Iteration 2: Updated PageRank



Already after two iterations, major properties of the algorithm show. First, from iteration to iteration, the PageRank of a page can oscillate between higher and lower values. E.g., the leftmost page changed from 0.2 to 0.1, then to 0.15. Over time, these changes become smaller as the values converge. Second, despite the oscillations, the general tendency is for some pages to accumulate larger values, while others drop.

PageRank using BFS

- Observing the steps during an iteration of PageRank, it becomes clear that there are two phases.
- In the first phase of the iteration, a page **sends out** fractions of its current PageRank along its *outgoing* edges. In the second phase, each page **sums up** the PageRank contributions along all its *incoming* edges.

PageRank using BFS

- Given page n 's current value $P(n)$ and adjacency list, one can compute all its outgoing contributions.
- Adding the incoming contributions for page m requires re-shuffling.
 - The term $(1-\alpha)/|V|$ remains constant throughout the computation. Both α and $|V|$ can be shared with all tasks.
- Interestingly, the computation pattern again matches BFS.

PageRank in MapReduce

- The algorithm is virtually identical to BFS, but an iteration must also keep track of the current PageRank value of each page:
 - Map processes a single vertex n , emitting the **PageRank of n , divided by n 's outdegree**, for each vertex m in n 's adjacency list.
 - Reduce collects PageRank contributions from all inlinks of vertex m and then applies the formula.

PageRank in MapReduce

- The driver program repeatedly calls the MapReduce program until (near) convergence, i.e., when all PageRank values stabilize.

MapReduce Code for a Single Iteration

```
// The vertex object now also stores
// the current PageRank value
map( vertex N ) {
    // Pass along the graph structure
    emit(N.id, N)

    // Compute contributions to send
    // along outgoing links
    p = N.pageRank / N.adjacencyList.size()
    for all m in N.adjacencyList do
        emit( m, p )
}
```

```
// Reduce receives the vertex object for vertex m and
// the PageRank contributions for all m's inlinks
reduce(id m, [p1, p2,...]) {
    s = 0
    M = NULL

    for all p in [p1, p2,...] do
        if isVertex(p) then
            // The vertex object was found: recover graph structure
            M = p
        else
            // A PageRank contribution from an inlink was found:
            // add it to the running sum.
            s += p

    M.pageRank = (1-alpha)/|V| + (alpha)s
    emit( m, M )
}
```

MapReduce Algorithm Analysis

- A careful look reveals that this program is structurally almost identical to the one for single-source shortest path. Hence it shares the same weaknesses. In each iteration:
 - The entire graph is read from the distributed file system.
 - The entire graph is transferred from Mappers to Reducers.
 - The entire graph, with updated PageRank values, is written to the distributed file system.

MapReduce Algorithm Analysis

- On the other hand, the PageRank program does not perform irrelevant computation.
- In contrast to single-source shortest path, **all** PageRank values must be recomputed in every iteration.

PageRank in Spark

- The Spark program separates static data (the graph structure) from evolving data (the current PageRank of each vertex).
 - This leverages cached RDD partitions and avoids MapReduce's shuffling of the adjacency list data.
 - On the downside, the new PageRank data of type (vertexID, PageRank) must be joined with the graph data of type (vertexID, adjacencyList). This join on vertexID normally requires shuffling, but careful co-partitioning avoids that.

Spark Program

```
// Assume that our neighbor list was saved as a Spark objectFile
val links = sc.objectFile[(String, Seq[String])]("links").partitionBy(new HashPartitioner(100))
    .persist()
val numPages = links.count().toDouble

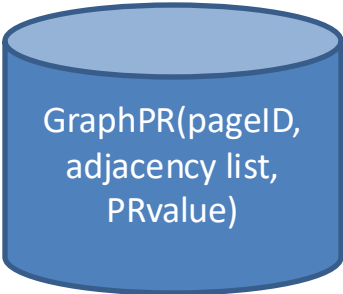
// Initialize each page's rank
var ranks = links.mapValues(v => 1.0 / numPages)

// Run 10 iterations of PageRank
for (i <- 0 until 10) {
  val contributions = links.join(ranks).flatMap {
    case (pageId, (links, rank)) => links.map(dest => (dest, rank / links.size))
  }
  ranks = contributions.reduceByKey((x, y) => x + y).mapValues(v => 0.15 / numPages + 0.85*v)
}

// Write out the final ranks
ranks.saveAsTextFile("ranks")
```

Let us look at the key difference between Spark and MapReduce.

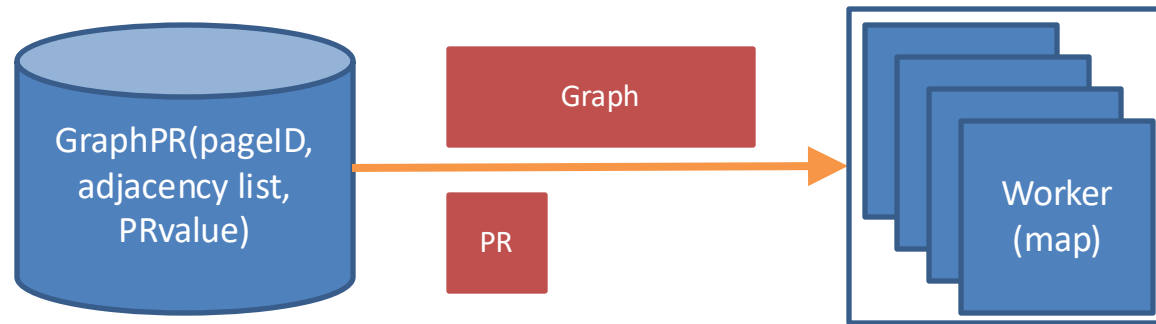
Typical iterative job (PageRank) in Hadoop MapReduce:



GraphPR(pageID,
adjacency list,
PRvalue)

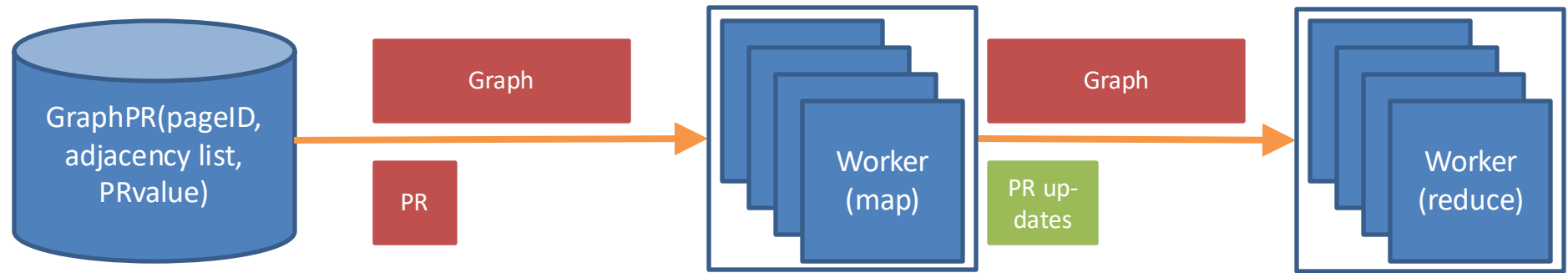
In MapReduce, each input record in GraphPR contains both the adjacency list and the PageRank value for each page. The PageRank algorithm iterates through the entire graph to update these PRvalues.

Typical iterative job (PageRank) in Hadoop MapReduce:



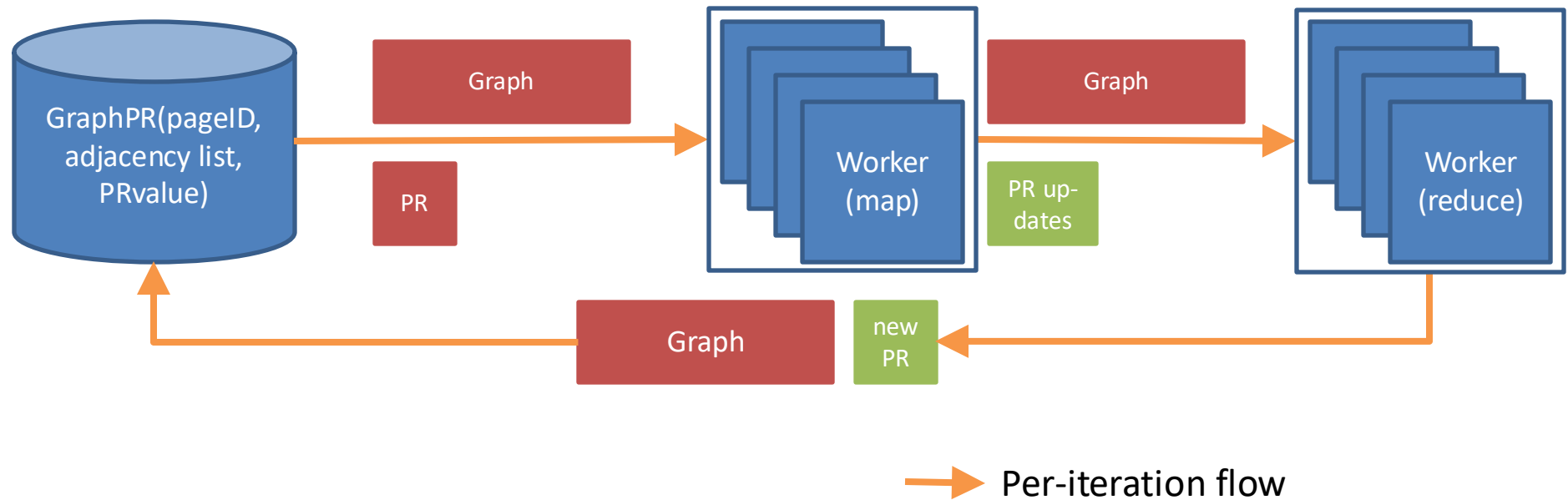
The MapReduce program first loads the Graph with its PRvalues into the Mappers. (Graph and PRvalue data are shown separately for better comparability to the Spark approach.)

Typical iterative job (PageRank) in Hadoop MapReduce:



Each Mapper forwards graph structure and update info for the PRvalues to the corresponding Reducers.

Typical iterative job (PageRank) in Hadoop MapReduce:

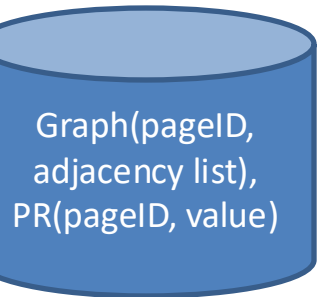


The Reducers compute the new PRvalues and emit the updated graph to the distributed file system.

This cycle repeats in the next iteration.

Spark manages data that changes in each iteration (PR: the PageRank value of a page) separate from data that does not (Graph: the graph structure, i.e., the pages and their adjacency lists of links.) This enables it to exploit in-memory data and dramatically reduce data movement.

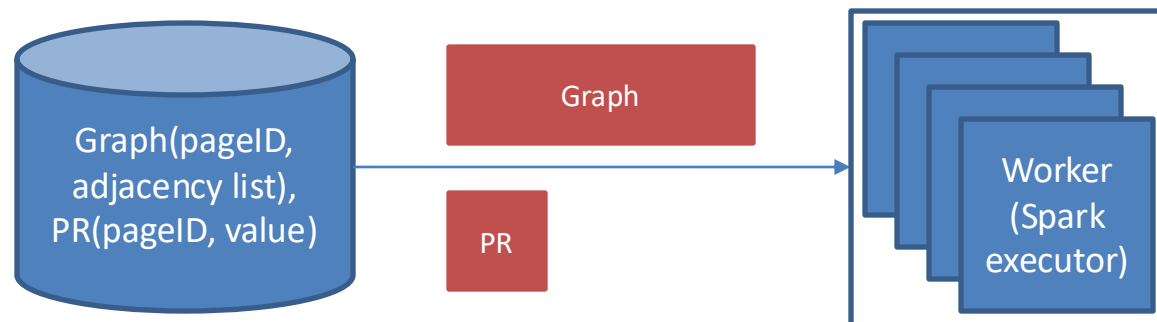
Typical iterative job (PageRank) in Spark:



First, both Graph and PR are loaded into the Spark worker tasks. Ideally (when the graph fits into the combined memory of all tasks) this initial loading step happens only once, not in each iteration.



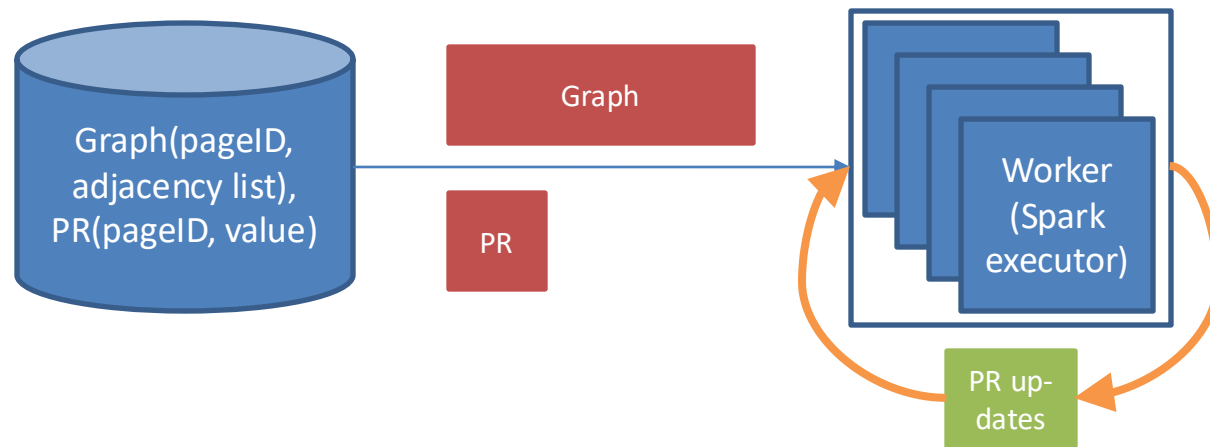
Typical iterative job (PageRank) in Spark:



Since the workers can keep the Graph data as an RDD or DataSet in memory, Graph does not need to be loaded and saved repeatedly. Instead, each iteration only passes the PR value updates to the task processing the corresponding node.



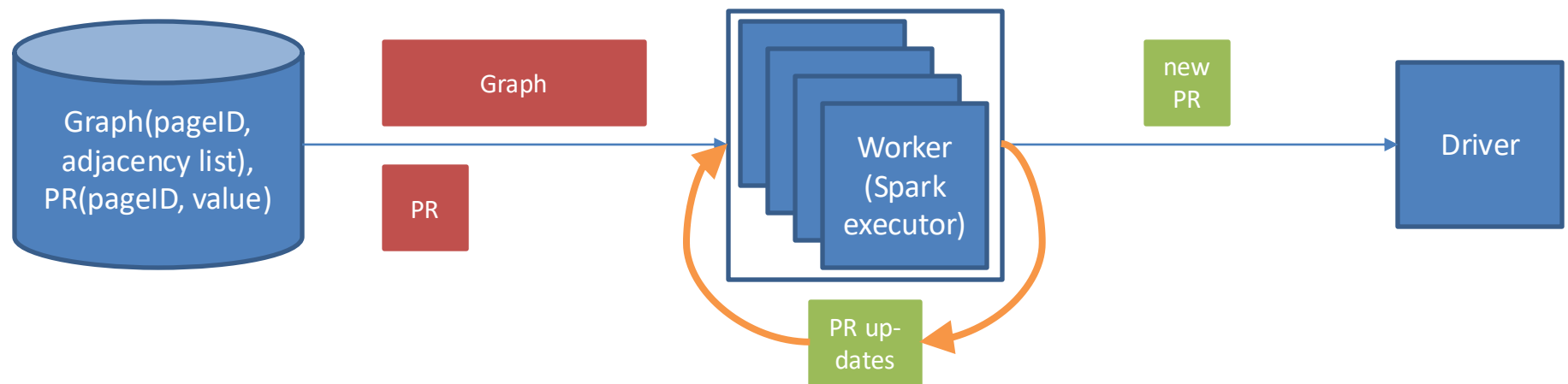
Typical iterative job (PageRank) in Spark:



In the very end, the final PR values are collected by the driver. This again is a one-time data transfer, not one happening in each iteration.



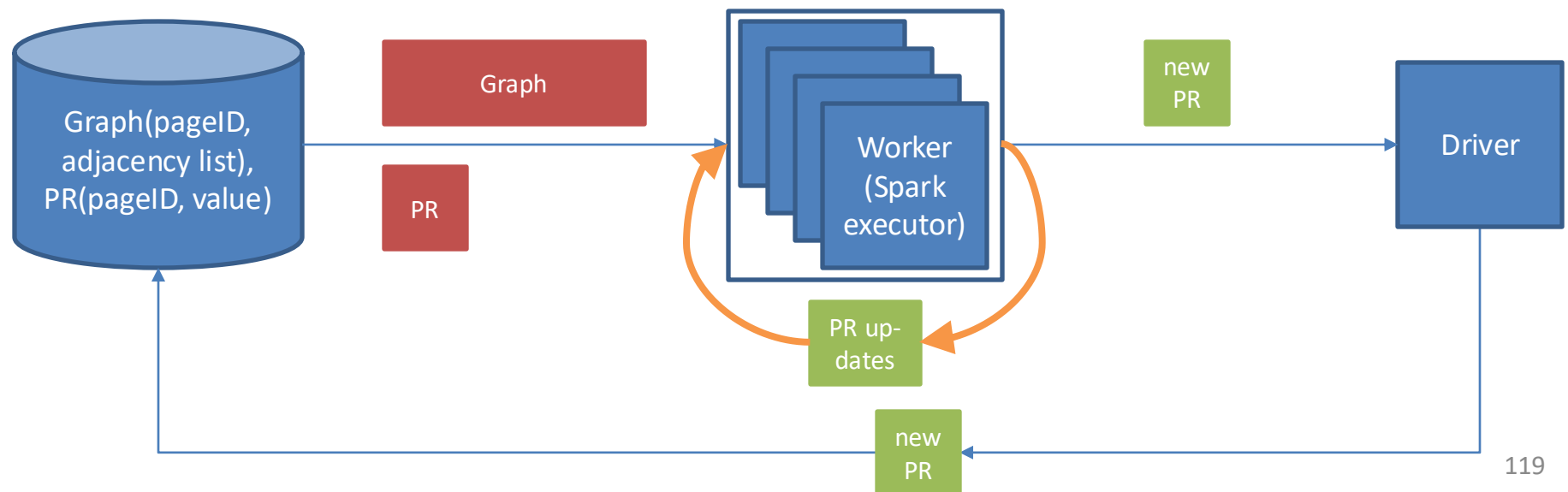
Typical iterative job (PageRank) in Spark:



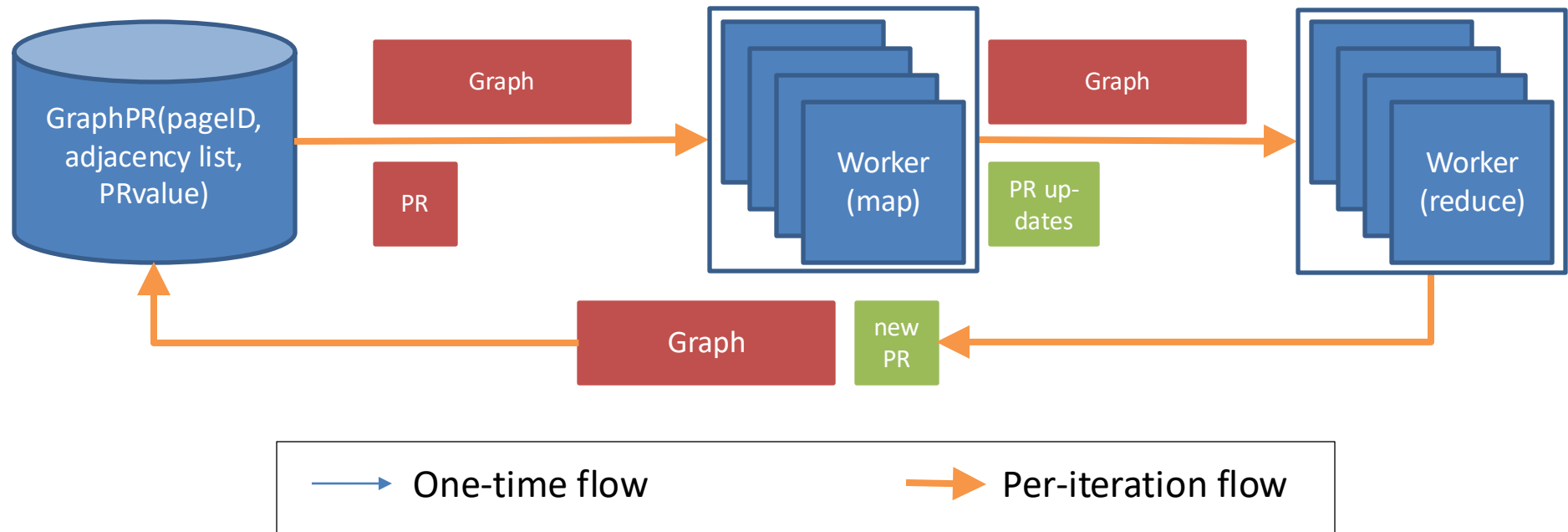
The driver then saves the final PR values to the distributed file system.



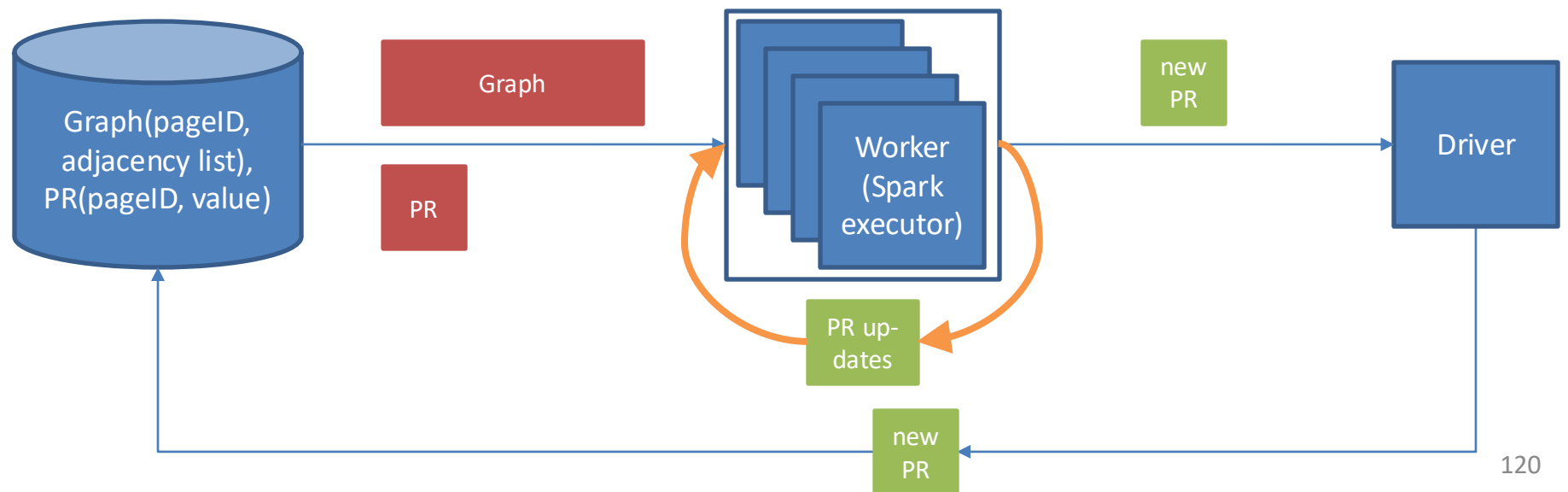
Typical iterative job (PageRank) in Spark:



PageRank in Hadoop MapReduce:



PageRank in Spark:



Dangling Pages

- A dangling page is one that has no outgoing links.
- These pages cause a PageRank "leak". Consider a page x with $\text{PageRank}(x)$.
- During an iteration, $\text{PageRank}(x)$ is not passed along other pages. Hence, a mass of $\alpha * \text{PageRank}(x)$ disappears.
 - The PageRank sum of the graph decreases.

Dangling Pages

- How do we address this? If a random surfer reaches a dangling page, it cannot follow links and must randomly jump.
- We can model this mathematically by conceptually adding imaginary links from x to every page in the graph, including itself.
- We call this total PageRank mass over all dangling pages δ .

Dangling Pages

$$P(n) = (1 - \alpha) \frac{1}{|V|} + \alpha \left(\frac{\delta}{|V|} + \sum_{m \in L(n)} \frac{P(m)}{C(m)} \right)$$

- Since PageRank values of dangling pages may change, we must compute the new value of δ in each iteration.

Dangling-Pages Solution in MapReduce

- The formula computing the new PageRank in Reduce must be modified to include the term with variable δ .
- Unfortunately, it is not possible to compute some value and also read it out in the *same* MapReduce job. We next discuss multiple possible solutions in more detail.

Dangling-Pages Solution in MapReduce

- Solution 1: add a separate job to each iteration to compute δ .
 - During an iteration, first execute a MapReduce job that computes δ . This is a simple global aggregation job, summing up PageRank values for all dangling pages.
 - Then pass the newly computed δ as a parameter to the modified MapReduce program that updates all PageRanks using the new formula with δ .
- Can we avoid the extra job in each iteration?

Alternative Solutions

- Solution 2: global counter.
 - We can compute δ “for free” during the Map phase. When Map processes a node with an empty adjacency list, it adds the node’s PageRank to a global counter.
 - By the end of the Map phase, this counter’s value is δ . If Reducers can read the counter value, they can compute the correct new PageRank in the Reduce function.
 - Otherwise, the global counter is read in the driver after the job completed.

Alternative Solutions

- Solution 3: dummy page.
 - For a dangling page n , the Map function emits (dummy, n 's PageRank). The Reduce call for the dummy node then adds up all contributions. The driver reads this output and passes it into the next job's context.

Number of Iterations

- In theory, the computation terminates when the **fixpoint** is reached, i.e., none of the PageRank values changes from one iteration to the next. In practice approximate numbers suffice, therefore iterations can stop as soon as all PageRank values “barely change.”
 - It is often reasonable to stop when no page’s PageRank changes by more than 1%. This threshold for *relative* change is computed for each page n as $| \text{new}(n) - \text{old}(n) | / \text{old}(n)$, where $\text{old}(n)$ and $\text{new}(n)$ are the PageRank of n in previous and current iteration, respectively.

Number of Iterations

- The easiest way to check for convergence is via a global counter or Accumulator that determines for how many pages the PageRank change was greater than the threshold.
- The counter is updated in the Reduce function (or the corresponding Spark function), where both old and new PageRank of a page are available. The driver program then checks the counter value to decide about another iteration.

Number of Iterations

- Depending on graph size, structure, and convergence threshold, the number of iterations can vary widely.
- The paper that proposed PageRank reported convergence after 52 iterations for a graph with 322 million edges.

Summary

- Large graphs tend to be sparse and hence are often stored in adjacency-list or set-of-edges format. This representation enables per-vertex computation in a single round, which can pass information along **outgoing** edges to all direct neighbors.
 - It is possible to extend this capability by pre-computing other data structures, e.g., the list of neighbors within a certain distance, to push information directly to nodes farther away.
- Computation along **incoming** edges requires shuffling.

Summary

- For some problems, the most efficient or most elegant sequential algorithms rely on a “centralized” data structure. These algorithms are difficult to parallelize directly.
- Iterative algorithms are common in practice and can be implemented easily in MapReduce and Spark.
 - However, MapReduce’s lack of support for persistent in-memory data results in inefficiencies due to costly data transfer.

Summary (Cont.)

- Each iteration of the presented algorithms has linear time complexity in graph size, scaling to large graphs.
- For algorithms that push data along graph edges, shuffle cost in Spark depends on the number of edges “across partitions.” (We discuss this on the next pages.)
 - A good partitioning therefore keeps densely connected regions of the graph together, cutting through sparse ones. Unfortunately, finding a balanced min-cut-style graph partitioning is hard.

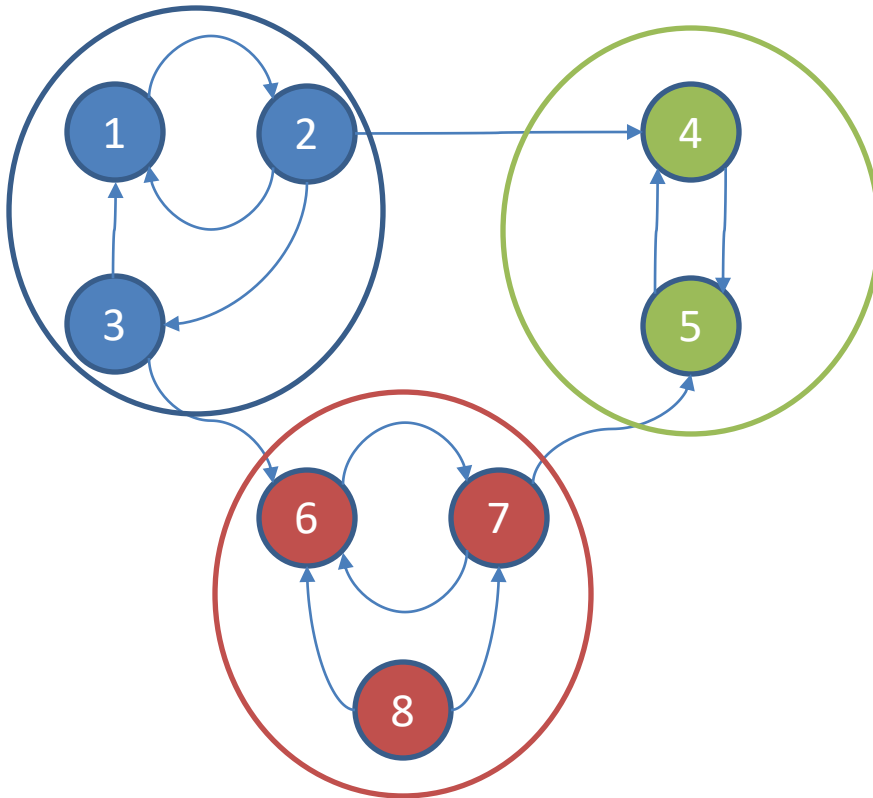
Summary (Cont.)

- Big data raises the issue of **numerical stability**:
In a billion-vertex graph an individual PageRank could *underflow* floating-point precision.

Side-Note: Graph Partitions versus Communication

- Consider the PageRank algorithm in Spark, where the graph RDD is partitioned over the different tasks. After loading the graph for the first iteration, no further graph-data movement occurs.
- However, PageRank values still change and need to be passed along outgoing edges to be aggregated along incoming edges.

Partitions versus Communication



Graph partitions:

Partition A:

node 1: 2
node 2: 1, 3, 4
node 3: 1, 6

Partition B:

node 4: 5
node 5: 4

Partition C:

node 6: 7
node 7: 5, 6
node 8: 6, 7

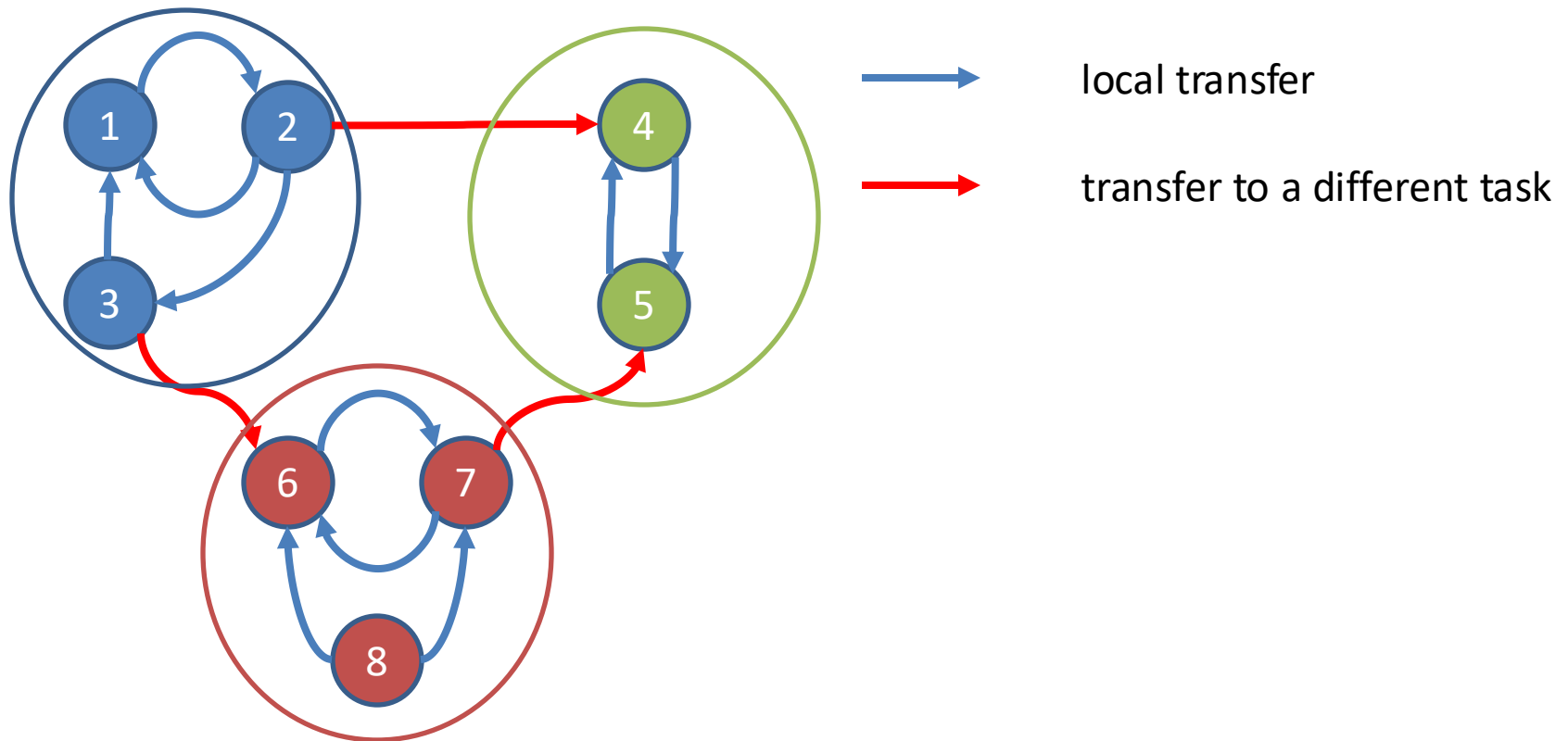
Partitions versus Communication

- During an iteration of the PageRank computation, each vertex sends its contributions along the outgoing edges.
- For destination vertices in the same partition, data transfer is local (indicated by a blue arrow on the next page).

Partitions versus Communication

- To avoid costly network transfer, the graph should be partitioned such that the number of edges connecting vertices in different partitions is minimized.
- In the example, only three edges (shown in red) require data transfer to a different task. After receiving the incoming contributions, for each vertex the new PageRank can be computed.

Efficient Iterative Processing



References

- S. Brin and L. Page. The anatomy of a large-scale hypertextual web search engine. Computer Networks and ISDN Systems, Elsevier, pages 107-117, 1998.
 - https://scholar.google.com/scholar?cluster=16140207188598694220&hl=en&as_sdt=0,22&as_vis=1