

**Efficient Job Execution for Map Reduce Using Phase-Level
Scheduling Algorithm.**Nisha Shinde¹, Trupti Patil², Prajkta Shinde³, Himani Mavchi⁴^{1,2,3,4} Student Dept. of Computer Engineering., AISSMS's Institute Of information Technology,
Pune, Maharashtra, India

Abstract - The Rapid improvement in the computational speed of the technology has made our life easier. New technologies like parallel computing and distributed computing has made a significant improvement in the speed of computing and help us to solve many complicated issues. The map reduce which is used in parallel computing is one of the popular data model for high speed computation in computation technology. The Existing map reduce focuses on scheduling at the task-level. But unfortunately, the task-level scheduling leads to inefficient job schedules with low resource utilization and long job execution time. In this concept we divide the tasks into unequal parts called as phases and apply phase-level scheduling to these phases and achieve efficient resource usage.

A fine-grained, phase and resource-aware Map Reduce Scheduler was introduced that divides tasks into phases, where each phase has a constant resource usage profile, and performs scheduling at the phase level. System can provide accurate resource information that can be used by the scheduler so that it can take effective scheduling decisions and reduce the job execution time. The phase-level scheduling achieves higher resource utilization when compared to task level schedulers

Keywords- Map Reduce; Scheduling; Cloud Computing; Hadoop; Resource Allocation;

I. INTRODUCTION

From past few years, in almost every country across the world, significant financial resources have been allocated to the improvement of the processing speed of computer system. The two factors that triggered this shift are modern technology developments and also the latest need of faster systems. Most of the developed countries are facing currently the problem off middle and older aged marketplace from a largely youth-driven marketplace. Due to this trend there is a great need of improved processing speed systems. Due to this requirement there is a huge effort put by the computer industry through different ways.

Nowadays the high speed computation is a need of today's Critical computational tasks. Map reduce is the parallel programing model for a data intensive applications. The Map Reduce algorithm contains two important tasks, namely Map and Reduce. The Map task takes a set of data and converts it into another set of data, where individual elements are broken down into tuples (key-value pairs). The Reduce task takes the output from the Map as an input and combines those data tuples (key-value pairs) into a smaller set of tuples. The reduce task is always performed after the map job. A central component to a MapReduce system is its job scheduler. Its role is to create a schedule of Map and Reduce tasks, spanning one or more jobs, that minimizes job completion time and maximizes resource utilization. A schedule with too many concurrently running tasks on a single machine will result in heavy resource contention and long job completion time. Conversely, a schedule with too few concurrently running tasks on a single machine will cause the machine to have poor resource utilization.

There are two aspects that differentiate scheduling in Map Reduce from traditional cluster scheduling. The first aspect is the need for data locality, i.e., placing tasks on nodes that contain their input data. Locality is crucial for performance because the network bisection bandwidth in a large cluster is much lower than the aggregate bandwidth of the disks in the machines. Traditional cluster schedulers that give each user a fixed set of machines, like Torque, significantly degrade performance, because files in Hadoop are distributed across all nodes as in GFS. Grid schedulers like Condor support locality constraints, but only at the level of geographic sites, not of machines, because they run CPU-intensive applications rather than data-intensive workloads like Map Reduce. Even with a granular fair scheduler, we

@IJAERD-2017, All rights Reserved

found that locality suffered in two situations: concurrent jobs and small jobs. We address this problem through a technique called delay scheduling that can double throughput.

This Project aims at developing a scheduler For Hadoop map reduce to work at the phase-level instead of the task-level. Divide the tasks into unequal parts called as phases and apply phase-level scheduling to these phases and achieve efficient resource usage and better Execution speed. The motivation of our project was because of our interest in parallel processing with Hadoop and to develop Scheduler that can utilize the resources in better way and improve the performance of parallel processing. Now-days, there is no such scheduler which can work on phase level so that better utilization and efficiency is achieved. Every scheduler is a task level scheduler, so we decided to design a new scheduler. This Scheduler can help Hadoop to achieve better parallelism by utilizing all the resources that are available and by providing better efficiency.

II. RELATED WORK

The Hadoop was designed for large batch jobs. Hadoop uses Map Reduce to complete this batch job. Map Reduce uses scheduler to schedule the task. There are Several scheduler schemes are there for scheduling like FIFO scheduler ,fair scheduler ,capacity scheduler ,resource adaptive scheduler and Dominant resource fair scheduler. This job schedulers are integrated with the job tracker. The job tracker pulls the job from system and give it to the scheduler for scheduling. The simplest scheduling algorithm is a FIFO. It works on the FIFO manner of queue (first come first serve basis). The Fair scheduler is the default. Resources are shared evenly across pools and each user has its own pool by default. You can configure custom pools and guaranteed minimum access to pools to prevent starvation. This scheduler supports pre-emption. The goal of the fair scheduler is to give all users equitable access between pools to cluster resources. Each pool has configurable guaranteed capacity in slots. Each Pool is equal to the number of jobs. Jobs are placed in flat pools. The default is 1 pool per user.

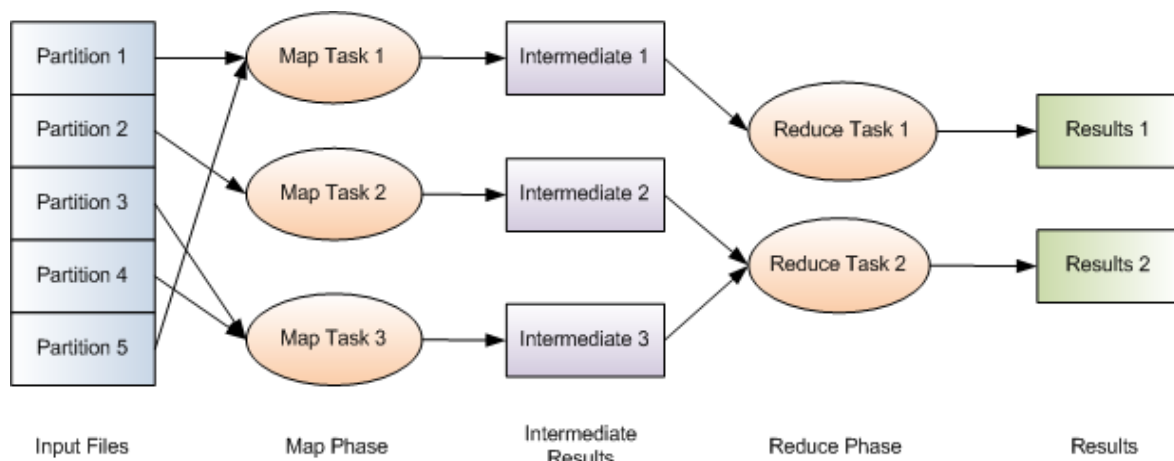


Figure 1. Typical MapReduce Job.

Hadoop also ships with the Capacity scheduler; here resources are shared across queues. You may configure hierarchical queues to reflect organizations and their weighted access to resources. You can also configure soft and hard capacity limits to users within a queue. Queues have ACLs to prevent rogues from accessing the queue. This scheduler supports resource-based scheduling and job priorities. The goal of the capacity scheduler is to give all queues access to cluster resources. Shares are assigned to queues as percentages of total cluster resources.

Although not a scheduler per se, Hadoop also supports the idea of provisioning virtual clusters from within larger physical clusters, called Hadoop on Demand (HOD). The HOD approach uses the Torque resource manager for node allocation based on the needs of the virtual cluster. With allocated nodes, the HOD system automatically prepares configuration files, and then initializes the system based on the nodes within the virtual cluster. Once initialized, the HOD virtual cluster can be used in a relatively independent way. HOD is also adaptive in that it can shrink when the

workload changes. HOD automatically de-allocates nodes from the virtual cluster after it detects no running jobs for a given time period. This behavior permits the most efficient use of the overall physical cluster assets.

So basically MapReduce work is to schedule the task in different levels. In a MapReduce technique, it is a collection of jobs and can be scheduled concurrently on multiple machines, resulting in reduction in job running time. Any companies refer MapReduce to process large volume of data. But they refer at task level to perform these data. Initially the task level performs two phases one is Mapper phase and other is reducer phase. In mapper phase, it takes data blocks Hadoop distributed file system and it maps, merge the data and stored in the multiple files. Then the second, reducer phase will fetch data from mapper output and shuffle, sort the data in a serialized manner. At the task level, performance and effectiveness become very important factor in day to day life. Task level scheduler is not capable to manage the available resources effectively. In the end it won't be able to give optimize performance.

III. PROPOSED SYSTEM

In the proposed system we present the fine-grained resource-aware scheduler, which performs scheduling at phase-level. It allows the job owners to specify the phase-level requirements. The architecture states that there are majorly three components: a scheduler called the phase-based scheduler which is located at the master node, local node managers that coordinate phase transitions with the scheduler and a Resource manager that is indeed used resource information at the phase-level. For scheduling whenever a task needs to be scheduled, the scheduler replies with a message with the task scheduling request. Then the node manager assigns that task. Each time a task finishes executing a phase it notifies and asks permission of the node manager to go to the next phase. The node manager then forwards the permission request to the scheduler through the regular message. If sufficient resources are available the scheduler decides and informs its decision to the local node manager whether it can proceed or wait the execution of the next phase. Finally, if the task is given permission to execute the next phase, the node manager grants the task to continue its duty. Once the task is completed, the task status is forwarded to the node manager and then forwarded again to the scheduler.

We suggest that it would be more efficient if we make the scheduler to work at the phase-level instead of the task-level. The reason is because the task demands a lot of requirements during its lifetime. In this concept we divide the tasks into unequal parts called as phases and apply phase-level scheduling to these phases and achieve efficient resource usage

3.1. Working Principle

The system functions can be described as follows:

When a User wants to perform the parallel computation the system must request for a job from user to perform parallel computation. And be sure that job is fit for parallel execution. So node manager has to send a message for request to schedule the job and whenever scheduler is ready it has to response to that request message.

Scheduling of job is the main functionality of this system, scheduling is the method by which work specified by some means is assigned to resources that complete the work. Here we first allocate the job to the scheduler and then divide this job into smaller task. After that This Smaller task are assign to node manager.

After that we have to Assign the resources to the job in the effective manner so that overall performance should be improve. At the end system should calculate the time and space to measure the performance improvement of system over traditional system.

3.2 Phase level scheduling algorithm.

The main steps in scheduling algorithm are as follows

- 1) The system is modeled as $S = \{s, e, X, Y, Jf, Jp, Jn\}$ Where, s is the initial state and the user will input the data (D).
- 2) X = Set of inputs in the system $X = \{D, k, tm\}$ where $D = \{d1, d2, d3, ..., dn\}$ and Y = Set of outputs Y
- 3) Jp = Job performance and Jf = job fairness.
- 4) Allocation of job to multiple task. Reduce the delay of particular task.
- 5) If $isPendingMapTask()$ is true then Calculate number of map task and Running Map Task.
- 6) Else if $isShuffleTask()$ is true then Calculate number of Shuffle task Running Map Task.
- 7) Else Calculate no of reduce task and Running Reduce Task
- 8) At the end Calculate Jf, Jp .
- 9) End.

IV. System Architecture

In the Phase Level Scheduling algorithm we will going the design the Scheduler which will work on Hadoop MapReduce which will improve the performance of parallel computation by better Resource utilization. As you can see in diagram our system mainly has three module as scheduler, Resource manager and node manager.

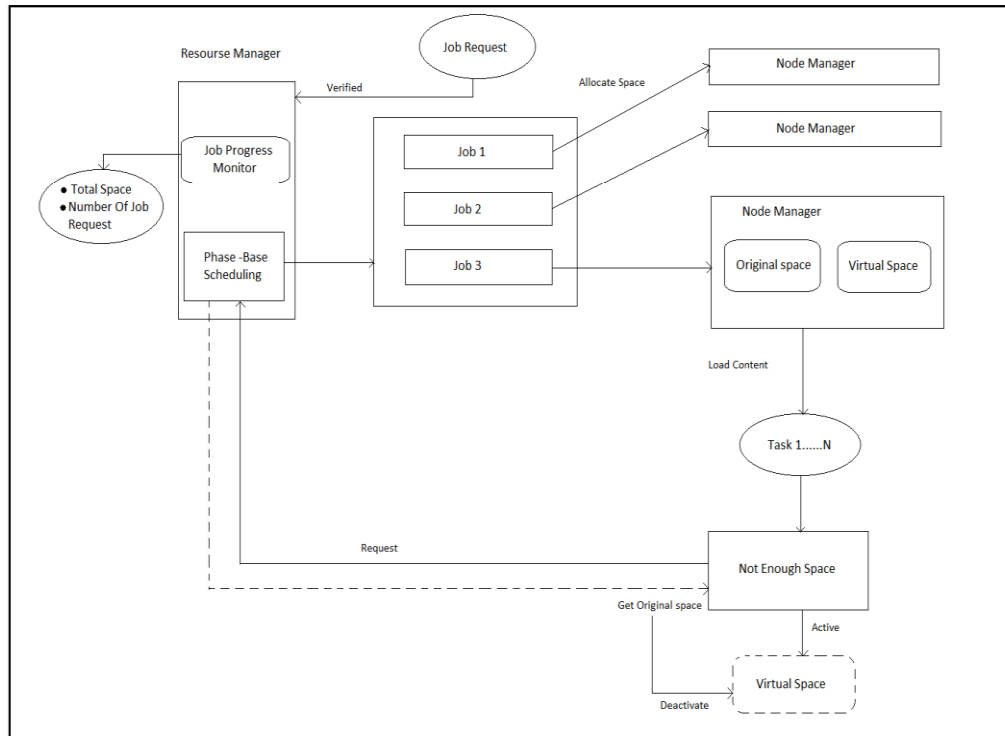


Figure 2. Phase Level scheduling architecture.

Also the other modules are users, job requester and tasks. The basic work flow of the system is User first create the resources and jobs. After creating the job requester Request for a job. Then the Job is allocated to the scheduler. Scheduler Divide job into smaller tasks and Assign each task to the node scheduler and also Schedule the Resources. Then the node scheduler Performs Computation and calculate the time and space. While proceeding in a phase level, phase-based scheduler send message to node manager. Upon receiving heartbeat message from node manager reporting resource availability on node, the scheduler must select which phase should be scheduled on node. For each job J consists of two types of tasks: map task M and reduce task R. We define the Utility function with machine n and assigning phase I as shown in equation. In utilization, PRISM is able to achieve shorter results and is able to achieve shorter job running time while maintaining high resource utilization for large workloads containing a mixture of jobs, which are same cluster.

IV. CONCLUSION

MapReduce is a popular data model for data related high speed computing. However, despite recent efforts toward designing resource-efficient MapReduce schedulers, existing work mainly focuses on designing task-level schedulers, they work fine but Still Provides Sub-Optimal Job Performance Because Varying Resource Requirements of tasks. Also task level scheduler does not able to utilize the resources effectively and it effects the performance of the system.

It is oblivious to the fact that the execution of each task can be divided into phases with drastically different resource consumption characteristics. That's when phase level scheduler comes into picture. The phase level scheduler which works on phase instead of task. Phase-Level consist of how run-time resources can be used and how it varies over the long life time. In short they utilize resources effectively. And it in the result performance gets improved.

REFERENCES

- [1] Qi Zhang, Mohamed Faten Zhani, Yuke Yang and Raouf Boutaba, "PRISM: Fine-Grained Resource-Aware Scheduling for MapReduce", IEEE Transaction on cloud computing, VOL. 3, NO. 2, Page Number, APRIL/JUNE 2015.
- [2] Suryakant S. Bhalke, "Allocation of Phase-Based Scheduler for MapReduce Job Scheduling", IJARCCCE, ISO 3297:2007 Certified Vol. 5, July 2016.
- [3] Dhanya Sudhakaran and Shini Renjith," Phase Based Resource Aware Scheduler with Job Profiling of Map Reduce." International Journal of Latest Trend in Engineering and Technology, July 2015.
- [4] Ms.Savitri.D.H and Narayana H.M, "PRISM: allocation of resources in phase-level using map-reduce in Hadoop", International Journal of Research in Science & Engineering e-ISSN: 2394-8299 Volume: 1 Special Issue: 2, May 2014.