

Introduction

Spark was designed for massive parallel computation on large-scale distributed cluster. It is currently the most popular open-source project of Apache Software Foundation, and the first choice of “big-data engine” for cluster-computation.

Although Spark Core learns from Hadoop MapReduce framework at lower level, with extensive optimizations and design choices, it quickly outperformed its predecessor, providing faster, more resilient, more convenient cluster-computation experience. Also, with extensive modules and libraries like Spark SQL, Spark Streaming, GraphX, and MLlib, it is a powerful platform to cater different computation needs, and different environments.

These powerful capabilities of Spark, understanding its key features and advantages is important for anyone who works in data-driven fields.

Objective

Our goal is to provide an overview of Apache Spark ecosystem and several key Spark modules. Including lower-level interfaces to machine and higher-level interfaces to users.

Also, we will compare Spark with its predecessor Hadoop, and review its advantages. The intention is to provide Spark intermediate learners with a better understanding of this powerful platform and offer a different perspective.

Spark Ecosystem

Extended Modules

Spark Core

Resource
Management

Distributed File
System

Modules and Innovations

