

PEOPLE'S DEMOCRATIC REPUBLIC OF ALGERIA
University of Sciences and Technology HOUARI
BOUMEDIENE

Faculty Informatics



Architecture of Apache Spark and Big Data Clustering with K-means

Presented by

- Mouhoub Massinissa
- Bouchair Aya
- Hannoun Amar Amine
- Merbaet Katia
- Zigadi Sadjji
- Mechidal Moncif

Contents

0.1	Introduction	1
1	Big Data and Database Systems	2
1.1	Introduction to Big Data	2
1.1.1	Types of Big Data	2
1.1.2	The 5 Vs of Big Data	3
1.1.3	Big Data Operations	3
1.2	SQL Databases and Database Management Systems	4
1.2.1	Definition of a DBMS	4
1.2.2	Types of DBMS	4
1.2.3	Functions of a DBMS	5
1.3	Introduction to SQL	6
1.3.1	Definition of SQL	6
1.3.2	Role of SQL in Data Manipulation and Querying	6
1.3.3	SQL and Different Database Management Systems	7
1.4	NoSQL Databases: Definitions and Various Implementations	9
1.4.1	Sample Example of NoSQL Database	10
1.4.2	Benefits of NoSQL Databases	11
1.4.3	Drawbacks of NoSQL database	11
1.4.4	Popular NoSQL Databases	12
2	Exploring Spark Architecture	16
2.1	Overview of Big Data Architectures	16
2.1.1	MapReduce: The Founding Model for Large-Scale Distributed Processing	16
2.1.2	Hadoop: The Open-Source Ecosystem for Big Data	17
2.1.3	Spark: A New Era for Unified and Performant Processing	18
2.2	Understanding the Spark Ecosystem	18
2.2.1	Spark Core	18
2.2.2	Spark SQL	18
2.2.3	Spark MLlib	19

2.2.4	Spark Streaming	20
2.2.5	Graph X	20
2.2.6	SparkR	21
2.3	Differences Between Big Data and Traditional Architectures .	22
2.3.1	Big Data Architecture	22
2.3.2	Benefits of big data architecture:	22
2.3.3	Challenges of big data architecture:	23
2.3.4	Traditional architecture	23
2.3.5	Key differences between big data and traditional architecture	24
3	Applications of Spark in Real Projects	25
3.1	Existing Projects and Implementations using Apache Spark . .	25
3.1.1	Real-Time Twitter Data Streaming ETL Pipeline . . .	25
3.1.2	Using Apache Spark for Real-Time Tweet Sentiment Analysis	27
3.2	Configuration of the Spark Ecosystem	30
3.2.1	Windows environment	30
3.2.2	Linux environment	34
3.2.3	MacOS environment	37
3.3	Our Use Case: Leveraging Spark for K-means clustering . . .	40
3.4	Conclusion	42

0.1 Introduction

The rapid growth of digital data has transformed the way information is stored, processed, and analyzed. Traditional database systems are often insufficient to handle the scale, variety, and speed of modern data, which has led to the emergence of Big Data technologies. Among these technologies, Apache Spark has become one of the most powerful and widely used frameworks for distributed computing, due to its in-memory processing capabilities, scalability, and support for diverse workloads such as batch processing, streaming, machine learning, and graph analytics.

This paper presents an overview of Big Data systems and database architectures, with a particular focus on Apache Spark and its ecosystem. It begins by introducing the main concepts of Big Data, relational and NoSQL databases, and the differences between traditional and Big Data architectures. It then explores the architecture of Apache Spark, describing its core components and key modules such as Spark SQL, Spark Streaming, MLlib, GraphX, and SparkR. Finally, the paper highlights real-world applications of Spark and presents a practical use case based on K-means clustering using Spark MLlib.

The objective of this work is to better understand how Spark can be used as a unified platform for efficient large-scale data processing and machine learning. By combining theoretical foundations with practical implementation, this study demonstrates the relevance of Spark in modern data-driven applications.

Chapter 1

Big Data and Database Systems

1.1 Introduction to Big Data

Big Data encompasses the vast expanse of information we interact with daily, comprising numerous zettabytes sourced from our computers, mobile devices, and various sensors.

Businesses leverage this data to inform decisions, refine processes and policies, and craft customer-centric products and services. Big Data earns its moniker not only due to its sheer volume but also its diversity and complexity, often surpassing the capabilities of traditional databases in capturing, managing, and processing such data.

1.1.1 Types of Big Data

1. **Structured Data:** This category encompasses easily organized and accessed data types such as financial records, demographic statistics, and machine logs. Think of an Excel spreadsheet with its neatly arranged columns and rows, facilitating straightforward visualization and categorization.
2. **Unstructured Data:** Here lie social media posts, audio files, images, and open-ended customer feedback. Unlike the orderly rows and columns of traditional databases, unstructured data often finds its home in data lakes, warehouses, and NoSQL databases due to its unconventional format.
3. **Semi-Structured Data:** Combining elements of both structured and

unstructured data, examples include emails containing unstructured message bodies alongside structured metadata like sender, recipient, and date. Similarly, devices employing geotags, timestamps, or semantic markers offer structured data embedded within unstructured content.

1.1.2 The 5 Vs of Big Data

In today's data-driven landscape, the 5 Vs of Big Data—Volume, Velocity, Variety, Veracity, and Value—define the challenges and opportunities inherent in managing vast amounts of information. From the sheer volume of data requiring advanced analytics to the real-time processing enabled by modern technologies, understanding these aspects is crucial for businesses aiming to derive actionable insights and competitive advantage from their data.

1. **Volume:** While not the sole defining factor, the sheer volume of data is undeniably a core aspect of Big Data. Harnessing this data necessitates advanced algorithms and AI-driven analytics, underpinned by robust storage, organization, and retrieval mechanisms.
2. **Velocity:** Modern Big Data technologies enable real-time processing and analysis, yielding actionable insights within milliseconds of data generation.
3. **Variety:** Beyond structured datasets, the true essence of Big Data lies in its amalgamation of structured, unstructured, and semi-structured data formats.
4. **Veracity:** Data is only valuable if it is accurate, relevant, and convenient. Yet, challenges such as human biases, social noise, and data provenance issues can undermine data quality.
5. **Value:** While the potential for fascinating insights is undeniable, the true value of Big Data lies in its ability to empower businesses with actionable intelligence, fostering competitiveness, resilience, and enhanced customer service.

1.1.3 Big Data Operations

- **Collecting Big Data:** Much of Big Data consists of huge sets of unstructured data, flowing from disparate and inconsistent sources. Traditional disk-based databases and data integration mechanisms simply

aren't up to managing all of this. Managing Big Data requires adopting In-Memory database solutions and specific Big Data software solutions.

- **Storing Big Data:** Big Data is, as the name suggests, voluminous. Many companies have on-premise storage solutions for their existing data and hope to save costs by reusing these repositories to process Big Data. However, Big Data is more efficient when not subject to size and memory constraints. Companies that don't integrate cloud storage solutions into their Big Data models from the start often regret it a few months later.[1]
- **Analyzing Big Data:** To make the most of Big Data, we need AI and Machine Learning for analysis. "Velocity," which is about getting insights fast, is crucial. For this, we rely on AI and modern databases to optimize processes and learn from experience.[1]

1.2 SQL Databases and Database Management Systems

1.2.1 Definition of a DBMS

A Database Management System (DBMS) is a software system that facilitates the creation, organization, management, and manipulation of databases.

At its core, a DBMS serves as an interface between users and databases, providing efficient methods for storing, retrieving, updating, and managing data. It acts as a centralized repository for data storage, ensuring data integrity, security, and accessibility.

The primary role of a DBMS is to facilitate the efficient and effective management of large volumes of structured and unstructured data.

1.2.2 Types of DBMS

1. **Relational DBMS (RDBMS):** Relational DBMS organizes data into tables with rows and columns, and establishes relationships between tables using keys. It follows the principles of relational algebra and SQL for data manipulation. Examples include MySQL, Oracle Database, and PostgreSQL. RDBMS is suitable for transactional applications, data warehousing, and online analytical processing (OLAP).
2. **Object-Oriented DBMS (OODBMS):** Object-Oriented DBMS stores data as objects, encapsulating data and behavior together. It sup-

ports inheritance, encapsulation, abstraction and polymorphism, enabling complex data modeling and manipulation. Examples include db4o (dataBase For Objects) and ObjectStore. OODBMS is suitable for applications with complex data structures and object-oriented programming paradigms.

3. **Hierarchical DBMS:** Hierarchical DBMS organizes data in a tree-like structure with parent-child relationships. It represents data hierarchically, where each record has one parent and multiple children. Examples include IBM Information Management System (IMS) and Windows Registry. Hierarchical DBMS is suitable for applications with hierarchical data structures, such as file systems and organizational charts.
4. **NoSQL (Not Only SQL) DBMS:** NoSQL DBMS encompasses various non-relational database systems designed to handle large volumes of unstructured or semi-structured data. It offers flexible schema design, horizontal scalability, and high availability. Examples include MongoDB, Cassandra, and Redis. NoSQL databases are suitable for real-time analytics, content management, and distributed data storage.

1.2.3 Functions of a DBMS

- **Data Definition:** DBMS enables users to define the structure of the database, including creating tables, specifying data types, defining constraints, and establishing relationships between entities. This function ensures data integrity and consistency by enforcing predefined rules and standards.
- **Data Manipulation:** DBMS allows users to manipulate data stored in the database through operations such as insertion, deletion, modification, and retrieval. Users can execute queries to extract specific information from the database, perform calculations, and generate reports.
- **Data Storage:** DBMS manages the physical storage of data on storage devices such as hard disks or solid-state drives. It optimizes storage allocation, manages disk space, and ensures efficient data retrieval by organizing data into storage structures such as tables, indexes, and files.
- **Data Retrieval:** DBMS provides mechanisms for retrieving data from the database based on user queries and criteria. It supports various

retrieval operations, including simple and complex queries, sorting, filtering, and aggregation, to retrieve relevant information efficiently.

- **Security:** DBMS implements security mechanisms to protect data from unauthorized access, manipulation, and disclosure. It supports authentication, authorization, encryption, and access control mechanisms to ensure data confidentiality, integrity, and availability.

1.3 Introduction to SQL

1.3.1 Definition of SQL

Structured Query Language (SQL) is a standard programming language used for managing and manipulating relational databases. It was first developed by IBM in the 1970s and has since become the go to language for interacting with databases. SQL provides a powerful and standardized way to define, access, and manipulate data stored in a relational database management system.

1.3.2 Role of SQL in Data Manipulation and Querying

SQL is essential for extracting valuable insights from databases through querying and analysis. It enables users to retrieve and manipulate large volumes of data efficiently, perform complex data operations, and generate meaningful reports and visualizations.

SQL's flexibility and power make it a fundamental tool for data management and analysis in various industries and applications. SQL allows users to:

- **Retrieve data:** Using SELECT statements, which enable users to extract specific information from one or more tables in the database.
- **Manipulate data:** Using INSERT, UPDATE, and DELETE statements, allowing users to add, modify, or remove records from the database.
- **Modify the database schema:** Using CREATE, ALTER, and DROP statements, enabling users to define or modify the structure of tables and other database objects.
- **Control access to data:** Using GRANT and REVOKE statements, allowing users to grant or revoke permissions to access or manipulate data.

1.3.3 SQL and Different Database Management Systems

MySQL:

It's an open-source relational database management system (RDBMS) known for its speed, reliability, and ease of use. It is widely used in web development, e-commerce, and other applications requiring a robust and scalable database solution. MySQL supports various features, including:

1. **Data types:** MySQL supports a wide range of data types, including integers, strings, dates, and timestamps.
2. **Indexing:** MySQL allows users to create indexes on columns to improve query performance and data retrieval.
3. **Transactions:** MySQL supports transactional operations, ensuring data integrity and consistency in multi-user environments.
4. **Replication:** MySQL supports replication, allowing users to create replicas of databases for scalability and high availability.

Oracle:

Oracle Database is a proprietary relational database management system (RDBMS) developed by Oracle Corporation. It is known for its scalability, high availability, and advanced features for enterprise-level applications. Oracle Database supports various features, including:

1. **Scalability:** Oracle Database is designed to handle large volumes of data and users, making it suitable for enterprise-scale applications.
2. **High availability:** Oracle Database provides built-in features for data replication, failover, and recovery, ensuring continuous availability and reliability.
3. **Advanced security options:** Oracle Database offers advanced security features, including encryption, access controls, and auditing, to protect sensitive data.

Oracle Database is widely used in industries such as finance, telecommunications, and healthcare, where reliability, scalability, and security are critical requirements. Companies such as Amazon, Walmart, and Bank of America rely on Oracle Database for mission-critical applications.

PostgreSQL:

PostgreSQL is an open-source object-relational database management system (ORDBMS) known for its extensibility, standards compliance, and advanced features. It supports various features, including:

1. **Support for JSON:** PostgreSQL natively supports storing and querying JSON data, making it suitable for applications with semi-structured data.
2. **Full-text search:** PostgreSQL provides built-in support for full-text search, enabling users to perform advanced text searches on large volumes of textual data.
3. **Advanced data types:** PostgreSQL offers a rich set of data types, including arrays, geometric types, and user-defined types, allowing users to model complex data structures.

PostgreSQL is commonly used in industries such as healthcare, research, and government, where data integrity, flexibility, and scalability are paramount. Companies such as Instagram, Spotify, and GitHub rely on PostgreSQL for managing their data efficiently.[2]

1.4 NoSQL Databases: Definitions and Various Implementations

A NoSQL database is a type of database that does not use the traditional tabular relations found in relational databases. Instead, NoSQL databases use a variety of data models to store and retrieve data, such as document-oriented, key-value, column-family, and graph databases.

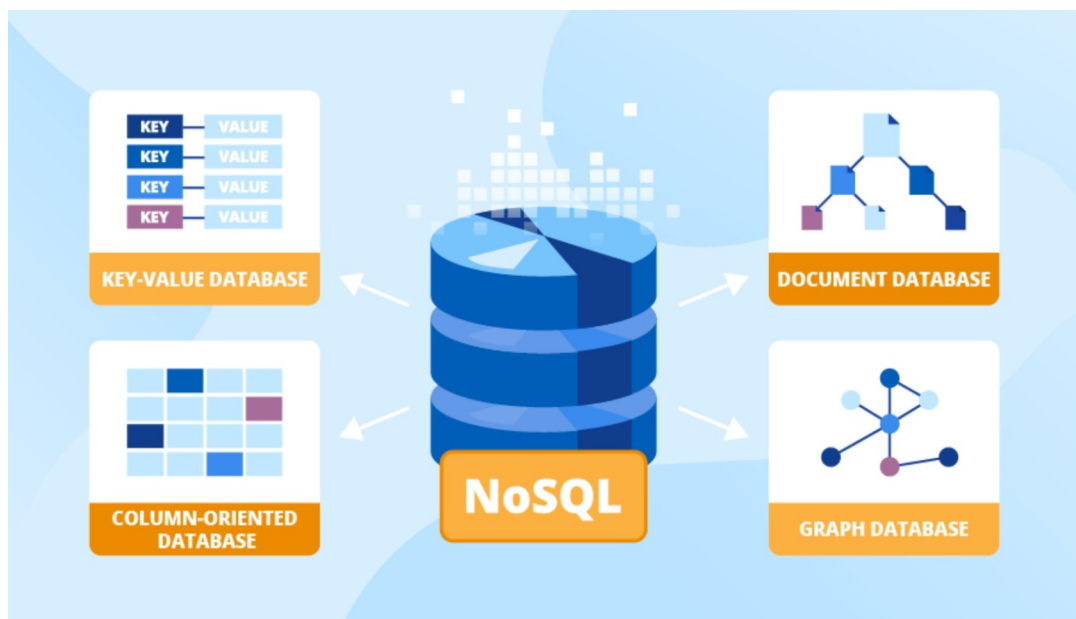


Figure 1.1: NoSQL.

One of the main advantages of NoSQL databases is their ability to scale horizontally, meaning that they can handle large amounts of data and traffic by adding more servers to a cluster. They are also more flexible than relational databases, as they allow for the storage of unstructured or semi-structured data, such as JSON or XML documents.

1.4.1 Sample Example of NoSQL Database

Here's a sample example of a NoSQL database using MongoDB: Let's say we want to store information about books in a MongoDB database. We could create a database called “**library**” and a collection called “**books**”. Each document in the “**books**” collection would represent a single book, and could contain the following fields:

```
1 {
2   "title": "The Great Gatsby",
3   "author": "F. Scott Fitzgerald",
4   "published_year": 1925,
5   "publisher": "Scribner",
6   "genres": ["Classic", "Fiction"],
7   "rating": 4.5,
8   "reviews": [
9     {
10      "user": "john_doe",
11      "rating": 4,
12      "comment": "Great book, highly recommended!"
13    },
14    {
15      "user": "jane_smith",
16      "rating": 5,
17      "comment": "One of the best books I've ever read."
18    }
19  ]
20 }
```

Figure 1.2: NoSQL code example.

In this example, each **book document** has a *title*, *author*, *published year*, *publisher*, *genres*, *rating*, and *reviews*. The “**genres**” field is an array of strings, while the “**reviews**” field is an array of sub-documents, each of which contains information about a single review.

With this document-oriented structure, we can easily retrieve all books by a certain **author**, filter books by **genre**, or search for books with a certain **rating**. We can also add or remove fields as needed without worrying about maintaining a strict schema.

1.4.2 Benefits of NoSQL Databases

NoSQL databases offer several benefits over traditional relational databases, including:

1. **Scalability:** NoSQL databases are designed to scale horizontally, meaning that they can handle large amounts of data and traffic by adding more servers to a cluster. This makes them well-suited for handling large-scale applications that need to handle high volumes of data.
2. **Flexibility:** NoSQL databases can handle different types of data and do not require a fixed schema. This means that they are well-suited for handling unstructured or semi-structured data, such as JSON or XML documents.
3. **High Performance:** NoSQL databases are designed to optimize read and write performance, making them ideal for use in applications that require low-latency data access.
4. **Cost-effectiveness:** NoSQL databases are often open-source, meaning that they can be used for free or at a low cost. Additionally, because they are designed to scale horizontally, they can often be run on commodity hardware, further reducing costs.
5. **Availability and fault-tolerance:** NoSQL databases are often designed to be highly available and fault-tolerant, meaning that they can continue to operate even if individual nodes fail. This makes them ideal for use in distributed systems that require high levels of uptime.

1.4.3 Drawbacks of NoSQL database

While NoSQL databases offer many benefits, there are also some drawbacks that should be considered when choosing whether to use a NoSQL database. These drawbacks include:

1. **Lack of ACID compliance:** NoSQL databases often sacrifice ACID compliance (Atomicity, Consistency, Isolation, Durability) to achieve high scalability and performance. This can make them less suitable for applications that require strict data consistency or transactional integrity.
2. **Limited query functionality:** NoSQL databases often lack the sophisticated query functionality provided by SQL-based relational databases. This can make it more difficult to perform complex data queries or joins.

3. **Limited tooling and community support:** NoSQL databases often have fewer development tools and less community support than SQL-based databases. This can make it more difficult to find help or to integrate the database with other technologies.
4. **Complexity:** NoSQL databases can be more complex to design and manage than SQL-based databases, especially if you are working with distributed databases or complex data models.
5. **Learning curve:** NoSQL databases often require learning a new set of skills and query languages, which can be challenging for developers who are used to working with SQL-based databases.

1.4.4 Popular NoSQL Databases

NoSQL databases are commonly used in modern web and mobile applications, where the data is often unstructured, and the application needs to handle large volumes of traffic and user-generated content. Some examples of popular NoSQL databases include MongoDB, Cassandra, Redis, and Neo4j.

MongoDB:

MongoDB is a popular open-source NoSQL document-oriented database system that was first released in 2009. It is designed to be highly scalable, flexible, and performant, making it well-suited for handling large volumes of data and traffic. Some of its key features of MongoDB include:

1. **Document-oriented data model:** MongoDB stores data in flexible JSON-like documents, allowing for more dynamic and expressive data modeling compared to traditional relational databases.
2. **High availability and scalability:** MongoDB supports automatic failover and load balancing, allowing it to handle large volumes of data and traffic with high availability.
3. **Rich query language:** MongoDB supports a powerful query language that includes support for a wide range of query operations and aggregation pipelines.
4. **Indexing:** MongoDB provides extensive indexing options, allowing you to create indexes on any field in a collection, including indexes that support text search and geospatial queries.

5. **Built-in sharding:** MongoDB supports automatic sharding, allowing you to scale your database horizontally across multiple servers or clusters.
6. **Open-source community:** MongoDB has a large and active open-source community that provides ongoing support, documentation, and tools.

Overall, MongoDB is a popular choice for many types of applications, including e-commerce, content management systems, and big data analytics. Its flexible data model, powerful query language, and built-in scaling capabilities make it a good fit for a wide range of use cases.

Cassandra:

Apache Cassandra is a popular open-source NoSQL distributed database system that was first released in 2008. It was originally developed by Facebook and is now maintained by the Apache Software Foundation.

Cassandra is designed to be highly scalable, fault-tolerant, and performant, making it well-suited for handling large volumes of data and traffic. Some of its key features of Cassandra include:

1. **Distributed architecture:** Cassandra is designed to run on multiple servers, with each node in the cluster being equal and independent. This allows for easy scalability and fault tolerance.
2. **No single point of failure:** Cassandra is designed to be highly fault-tolerant, with no single point of failure in the system. Data is replicated across multiple nodes to ensure that it is always available.
3. **Tunable consistency:** Cassandra provides tunable consistency, allowing you to choose the level of consistency that you need for your data.
4. **Column-family data model:** Cassandra stores data in a column-family data model, which allows for flexible data modeling and supports semi-structured and unstructured data.
5. **CQL:** Cassandra provides a powerful query language called CQL (Cassandra Query Language) that is similar to SQL, making it easy for developers to learn and use.

6. **Linear scalability:** Cassandra is designed to scale linearly, allowing it to handle large volumes of data and traffic without sacrificing performance.

Overall, Cassandra is a popular choice for many types of applications, including those that require high scalability and fault tolerance, such as e-commerce, social media, and big data analytics. Its distributed architecture, tunable consistency, and powerful query language make it a good fit for a wide range of use cases.

Redis:

Redis is a popular open-source in-memory data structure store that was first released in 2009. It is designed to be highly performant, with low latency and high throughput, making it well-suited for handling high-speed data and real-time applications. Some key features of Redis include:

1. **In-memory data store:** Redis stores data in memory, which allows for extremely fast read and write operations.
2. **Data structures:** Redis supports a wide range of data structures, including strings, hashes, lists, sets, and sorted sets. This allows for more expressive and flexible data modeling compared to traditional key-value stores.
3. **Pub/sub messaging:** Redis provides a pub/sub messaging system, allowing multiple clients to subscribe to channels and receive real-time updates when new data is published.
4. **Transactions:** Redis supports transactions, allowing multiple commands to be executed atomically.
5. **Lua scripting:** Redis provides a Lua scripting engine, allowing you to write custom scripts to perform complex operations.
6. **Persistence:** Redis supports persistence, allowing data to be written to disk for durability and backup purposes.

Overall, Redis is a popular choice for many types of applications, including real-time analytics, gaming, and chat applications. Its in-memory data store, and flexible data structures, make it a good fit for high-speed, real-time applications.

Neo4j:

Neo4j is a popular open-source graph database management system that was first released in 2007. It is designed to store and process highly interconnected data, making it well-suited for use cases such as social networking, recommendation engines, and knowledge graphs. Some key features of Neo4j include:

1. **Graph data model:** Neo4j stores data in a graph data model, which allows for complex relationships between entities to be modeled and queried.
2. **ACID compliance:** Neo4j is ACID-compliant, meaning that it ensures data consistency and reliability.
3. **Cypher query language:** Neo4j provides a powerful query language called Cypher, which is specifically designed for querying graph data.
4. **High availability and scalability:** Neo4j supports clustering and automatic failover, allowing it to handle large volumes of data and traffic with high availability.
5. **Spatial data:** Neo4j supports spatial data, allowing you to store and query geospatial data.
6. **Integrations:** Neo4j integrates with a wide range of tools and frameworks, including Python, Java, and GraphQL.

Overall, Neo4j is a popular choice for many types of applications, including social networking, recommendation engines, and knowledge graphs. Its graph data model, powerful query language, and support for spatial data make it a good fit for use cases that require modeling complex relationships between entities.

As a conclusion, implementing a NoSQL database requires careful consideration of your specific needs, as well as a solid understanding of the NoSQL database type that you have chosen to use.

Chapter 2

Exploring Spark Architecture

2.1 Overview of Big Data Architectures

The advent of Big Data has brought a multitude of challenges and opportunities for businesses and organizations.

Traditional data processing systems were not equipped to handle the massive volume, velocity, and variety of data generated in the modern era. This led to the emergence of new technologies and frameworks, such as Apache Spark, that are specifically designed to address the complexities of Big Data processing.

By offering faster processing, simpler development, enhanced scalability, and cost-efficiency, Spark has become a cornerstone of modern Big Data pipelines across diverse industries.

2.1.1 MapReduce: The Founding Model for Large-Scale Distributed Processing

MapReduce is a programming model and an associated implementation for processing and generating large data sets with a parallel, distributed algorithm on the cluster. It consists of two main phases:

- **Map:** The input data is broken down into smaller chunks, and each chunk is processed independently by a mapper function.
- **Reduce:** The outputs of the mapper functions are combined into a final output by a reducer function.

MapReduce is not well-suited for tasks that require multiple passes over the data, such as iterative algorithms or interactive data analysis [1]. Addition-

ally, it can be inefficient for small data sets due to its overhead, and it requires data to be stored in a specific format, which can add complexity to the data processing pipeline.

2.1.2 Hadoop: The Open-Source Ecosystem for Big Data

Hadoop is an open-source software framework for distributed storage and processing of large data sets. It consists of four main components:

1. **Hadoop Distributed File System (HDFS):** A distributed file system that stores data across multiple nodes in a cluster.
2. **YARN:** A resource manager that schedules and manages jobs on a cluster.
3. **MapReduce:** implementation as a programming model for processing large data sets in parallel on a Hadoop cluster. It is particularly suited for tasks that can be broken down into independent units of work, such as log processing, data filtering, data aggregation, and basic data transformations.
4. **Hadoop Common:** A set of utilities and libraries that are common to all Hadoop components.

Hadoop's scalable and cost-effective platform for storing and processing large data sets has been instrumental in the rise of Big Data. It has empowered a wide range of organizations, including businesses, governments, and research institutions, to analyze massive datasets and extract valuable insights for a variety of purposes.

2.1.3 Spark: A New Era for Unified and Performant Processing

Spark is a unified analytics engine for large-scale data processing. It can be used for batch processing, streaming, machine learning, and graph processing. Spark is faster than MapReduce [3] because it uses in-memory computing and a more efficient execution model. Spark can be integrated into the Hadoop Ecosystem, to provide a unified platform for Big Data processing.

2.2 Understanding the Spark Ecosystem

Apache Spark is a powerful open-source distributed computing framework designed for processing large-scale data across clusters of computers. Its ecosystem consists of various components that offer different functionalities for data processing, analytics, and machine learning. Let's delve into the core components of the Spark ecosystem.

2.2.1 Spark Core

All the functionalities being provided by Apache Spark are built on the top of Spark Core. It delivers speed by providing in-memory computation capabilities. Thus Spark Core is the foundation of parallel and distributed processing of huge datasets.

The key features of Apache Spark Core are:

- It is in charge of essential I/O functionalities.
- Significant in programming and observing the role of the Spark cluster.
- Task dispatching.
- Fault recovery.
- It overcomes the snag of MapReduce by using in-memory computation.
- It's embedded with a special collection called RDD (resilient distributed dataset).

2.2.2 Spark SQL

The Spark SQL component is a distributed framework for structured data processing. Using Spark SQL, Spark gets more information about the structure of data and the computation. With this information, Spark can perform

extra optimization. It uses the same execution engine while computing an output. It does not depend on API/ language to express the computation.

Spark SQL works to access structured and semi-structured information. It also enables powerful, interactive, analytical application across both streaming and historical data.

Spark SQL is a Spark module for structured data processing. Thus, it acts as a distributed SQL query engine.

Features of Spark SQL include:

- Cost based optimizer.
- Mid query fault-tolerance; this is done by scaling thousands of nodes and multi-hour queries using the Spark engine.
- DataFrames and SQL provide a common way to access a variety of data sources. It includes Hive, Avro, Parquet, ORC, JSON, and JDBC.
- Full compatibility with existing Hive data.
- Provision to carry structured data inside Spark programs, using either SQL or a familiar Data Frame API.

2.2.3 Spark MLlib

MLlib is a powerful machine learning library within Apache Spark designed for scalability and speed. It offers a wide range of high-quality algorithms for tasks like clustering, regression, classification, and collaborative filtering. Originally based on RDDs (Resilient Distributed Datasets), MLlib has transitioned to a DataFrame-based API in Spark Version 2.0, providing a more user-friendly approach and leveraging Spark's optimized execution engine.

MLlib simplifies machine learning workflows by providing Java, Scala, and Python APIs for distributed implementations of common learning algorithms like linear models, naive Bayes, decision tree ensembles, and k-means clustering. It also includes utilities for convex optimization, distributed linear algebra, statistical analysis, and feature extraction.

Additionally, MLlib benefits from the broader Spark ecosystem. Spark SQL integration supports data preprocessing and structured data processing, while Spark Streaming enables the development of online learning algorithms for processing live data streams. MLlib's `spark.ml` package further simplifies multi-stage learning pipelines with high-level APIs, leveraging Spark's

extensive capabilities for transforming and processing large-scale datasets efficiently.

2.2.4 Spark Streaming

Spark Streaming is an extension of the core Spark API that allows for processing of live data streams in a scalable, fault-tolerant manner.

One key feature of Spark Streaming is micro-batching, where the incoming data stream is divided into small batches for processing. This approach ensures fault tolerance and allows Spark to efficiently handle real-time streaming data.

Spark Streaming works in three phases:

- **Gathering:** Data is collected from various sources, including basic sources like file systems and socket connections, as well as advanced sources like Kafka, Flume, and Kinesis.
- **Processing:** The gathered data is processed using complex algorithms expressed with high-level functions such as map, reduce, join, and window operations.
- **Data Storage:** The processed data is then stored or delivered to file systems, databases, or live dashboards.

Spark Streaming also provides a high-level abstraction called a discretized stream, or DStream. A DStream represents a continuous stream of data and can be created either from external sources like Kafka and Flume or through high-level operations on existing DStreams. Internally, a DStream is a sequence of Resilient Distributed Datasets (RDDs), which are the fundamental data structure in Spark.

2.2.5 Graph X

GraphX in Spark is an essential API tailored for graph processing and parallel execution within the Spark framework. It serves as a robust network graph analytics engine and data store, enabling various operations such as clustering, classification, traversal, searching, and pathfinding on graphs.

One of the key features of GraphX is its extension of the Spark RDD (Resilient Distributed Dataset) by introducing a new Graph abstraction. This

abstraction represents a directed multigraph, where properties can be attached to each vertex and edge. This structure provides a flexible and efficient way to model and analyze complex relationships within graphs.

GraphX also offers optimizations for handling primitive data types when representing vertices and edges. This optimization enhances the performance of graph computations. Additionally, GraphX supports fundamental graph operators such as subgraph, join Vertices, and aggregate Messages. It also includes an optimized variant of the Pregel API, a widely used programming model for graph processing. These features make GraphX a powerful tool for graph analytics and processing within the Spark ecosystem.

2.2.6 SparkR

SparkR made its debut with the Apache Spark 1.4 release, introducing the SparkR DataFrame as its cornerstone. As a critical data structure in R for processing data, DataFrames find counterparts in other languages like Python with libraries such as Pandas. The primary objective of SparkR was to merge R's usability with Spark's scalability, presenting an R package that serves as a lightweight interface for utilizing Apache Spark directly within R.

SparkR offers several advantages:

- **Data Sources API:** By leveraging Spark SQL's data sources API, SparkR seamlessly reads data from diverse sources like Hive tables, JSON files, and Parquet files.
- **Data Frame Optimizations:** SparkR DataFrames benefit from optimizations made to the computation engine, including enhanced code generation and memory management.
- **Scalability:** Operations executed on SparkR DataFrames are distributed across all available cores and machines in the Spark cluster. This scalability enables SparkR to handle vast datasets, scaling from terabytes of data to clusters comprising thousands of machines.

In summary, SparkR empowers R users to harness the scalability and performance capabilities of Apache Spark for efficient large-scale data processing tasks, seamlessly integrating R's data manipulation and graphical capabilities with Spark's distributed computing prowess.

Apache Spark enhances existing Big Data tools for analysis without reinventing the wheel. Its popularity stems from the robust Apache Spark

Ecosystem Components, making it a preferred choice over other Big Data frameworks. Spark serves as a versatile platform for various data processing tasks, including real-time data analytics, structured data processing, and graph processing.

With its growing momentum, Apache Spark emerges as a promising alternative for supporting ad-hoc queries and iterative processing, replacing traditional MapReduce approaches. It facilitates interactive code execution through Python and Scala REPL (Read-Eval-Print Loop), offering flexibility for writing and compiling applications in Scala, Java, or other supported languages.

2.3 Differences Between Big Data and Traditional Architectures

Big data architecture and traditional architecture are two different approaches to data management and processing [4]. Big data architecture is designed to handle the large volumes of data that are generated by modern applications, while traditional architecture is designed for smaller datasets [5].

2.3.1 Big Data Architecture

It's a new approach to data management and processing that is designed to handle the large volumes of data that are generated by modern applications. Big data architecture is typically based on a distributed file system, such as Hadoop, and a distributed processing framework, such as Spark.

2.3.2 Benefits of big data architecture:

- **Decision-making:** Big data architecture can help businesses make better decisions by providing them with insights into their data. [6]
- **Increased efficiency:** Big data architecture can help businesses improve their efficiency by automating processes and reducing the time it takes to analyze data.
- **New product and service development:** Big data architecture can help businesses develop new products and services by providing them with insights into customer needs.

- **Reduced costs:** Big data architecture can help businesses reduce costs by improving efficiency and reducing the need for manual data processing.

2.3.3 Challenges of big data architecture:

- **Complexity:** Big data architecture can be complex to implement and manage.
- **Cost:** Big data architecture can be expensive to implement and maintain.
- **Security:** Big data architecture can pose security risks if not properly implemented.
- **Skills:** Big data architecture requires specialized skills that may not be available in-house.

2.3.4 Traditional architecture

Traditional architecture is the traditional way of storing and processing data. It is typically based on a relational database management system (RDBMS). RDBMSs are designed for structured data, which is data that is organized in a table format.

2.3.4.1 Advantages of traditional architecture:

- It is well-understood and a mature technology.
- It is relatively easy to implement and manage.
- It is scalable to a certain extent.
- It is relatively secure.

2.3.4.2 Disadvantages of traditional architecture:

- It is not designed for big data.
- It is not designed for real-time data processing.
- It can be expensive to scale.

2.3.5 Key differences between big data and traditional architecture

- **Data volume:** Big data architecture is designed to handle large volumes of data, while traditional architectures are designed for smaller datasets.
- **Data variety:** Big data architecture is designed to handle a variety of data types, including structured, semi-structured, and unstructured data [7]. Traditional architecture is typically designed for structured data.
- **Data velocity:** Big data architecture is designed to handle data that is generated in real time or near real time. Traditional architecture is typically designed for batch processing of data.[8]
- **Scalability:** Big data architecture is designed to be scalable to handle increasing volumes of data. Traditional architecture is typically not as scalable.
- **Cost:** Big data architecture can be more expensive to implement and maintain than traditional architecture.
- **Complexity:** Big data architecture can be more complex to implement and manage than traditional architecture.

As a conclusion, big data architecture and traditional architecture are two different approaches to data management and processing.

Big data architecture is designed to handle the large volumes of data that are generated by modern applications, while traditional architecture is designed for smaller datasets. The best approach for a particular organization will depend on its specific needs and requirements.

Chapter 3

Applications of Spark in Real Projects

3.1 Existing Projects and Implementations using Apache Spark

We will be talking about some projects

3.1.1 Real-Time Twitter Data Streaming ETL Pipeline

Summary

This project explores the use of Apache Spark for sentiment analysis of tweets, highlighting its efficiency and flexibility in processing large volumes of real-time data. Apache Spark serves as a central pillar in the collection, processing, and analysis of data, thereby illustrating a powerful use case of Spark in the field of social media data analysis.

Introduction

Sentiment analysis on Twitter provides valuable insights for various applications, ranging from monitoring brand reputation to making data-driven marketing decisions. This project utilizes Apache Spark to analyze sentiments expressed in tweets in real-time, thus demonstrating Spark's power in data stream processing applications.

Context and Justification

Apache Spark was chosen for this project due to its speed, ability to handle real-time data streams, and easy integration with other technologies such as Apache Kafka and Delta Lake. This report justifies the choice of Spark by its in-memory processing capabilities and built-in tools for data analysis.

Methodology

First, *ETL* stands for Extract, Transform, Load. It's a common process in data engineering where data is extracted from a source, transformed into a usable format, and loaded into a destination like a database.

• Components used

- **Tweepy** A Python library for accessing Twitter's API (to get Twitter data).
- **Apache Kafka** A software platform for streaming data. It's used to handle and organize the incoming Twitter data.
- **Apache Spark** A powerful analytics engine and framework. It's used for processing and analyzing Twitter data.
- **Delta Lake** A storage layer for data lakes, used to store the processed data.

The data flow is described, starting from extraction (Twitter API) to transformation (Spark) and loading (Delta Lake).

1. **Data extraction:** Data is extracted in real-time using the Twitter API and directed to a Kafka messaging management system.
2. **Data Transformation with Spark:** The core of this project lies in using Spark to transform and analyze tweets. Data from Kafka is consumed by Spark Streaming, which processes it to detect and classify sentiments.
 - **Preprocessing:** Cleaning and tokenization of tweets.
 - **ML Pipeline:** Use of MLlib to create a text processing pipeline, including TF-IDF and logistic regression for sentiment classification.
3. **Data Loading:** Transformed data are loaded into Delta Lake for storage and further analysis, taking advantage of transactional management and scalability of Delta Lake.

System Architecture

The system utilizes a distributed architecture where each component plays a key role in processing and analyzing the data. Kafka acts as a hub for real-time data, Spark as the processing and analysis engine, and Delta Lake as the persistence layer.

Technical Details

- **Configuration of Spark Streaming:** Spark Streaming is configured to read data from the Kafka topic. A structured schema is applied to the data to facilitate manipulation and processing.
- **Tweet Processing and Sentiment Prediction:** Tweets are cleaned, tokenized, and transformed using a TF-IDF model followed by logistic regression to predict sentiment.
- **Recording Results:** Results are stored in Delta Lake, optimized for both batch and streaming processing, providing a robust platform for managing output data.

Performance Analysis

The model achieves an accuracy of 83.2% on the test dataset. This section also discusses performance metrics used to evaluate the model.

Conclusion

This Twitter sentiment analysis project has highlighted the effectiveness of Apache Spark as a central platform for real-time data processing. With its ability to handle large volumes of data quickly and efficiently, Spark has proven to be an indispensable tool for meeting the demands of this ambitious project.

3.1.2 Using Apache Spark for Real-Time Tweet Sentiment Analysis

Summary

This Spark project aims to analyze and compare the performance of different regions in managing the COVID-19 pandemic. Using relevant metrics and statistics, the aim is to identify which regions were most successful in containing the spread of the virus and minimizing its impact. The project will

also focus on real-time monitoring of COVID-19 data, analyzing the popularity and volume of online discussions on the subject, and identifying the hashtags most used to describe the pandemic.

Introduction

The COVID-19 pandemic has had a profound impact on the whole world, shattering lives and economies. While countries have implemented different strategies and measures to combat the virus, it is crucial to assess the effectiveness of these approaches and identify best practices. This Spark project is part of this effort, using data analysis and natural language processing tools to provide valuable insights into pandemic management.

Context and Justification

Understanding regional disparities in the management of COVID-19 is essential for several reasons. Firstly, it helps identify the factors that contribute to the success or failure of different strategies. Secondly, it can inform policy-makers and public health officials in their efforts to develop more effective interventions. Thirdly, it can help to target resources and support efforts towards the most troubled regions.

Methodology

The Spark project will use a combination of data analysis and natural language processing (NLP) techniques to achieve its objectives. The methodology can be divided into three main sub-sections:

1. Data collection and preparation

- Identify reliable and recognized sources of COVID-19 data, such as international public health organizations and national governments.
- Collect data on key performance indicators, such as infection rates, death rates, vaccination rates and other relevant measures.
- Clean and pre-process data to ensure quality and consistency.
- Integrate data from different sources into a unified format to facilitate analysis.

2. COVID-19 data analysis

- Apply statistical analysis techniques to identify trends and patterns in COVID-19 data.
- Compare the performance of different regions in terms of key performance indicators.
- Identify factors contributing to the success or failure of different COVID-19 management strategies.
- Visualize analysis results to facilitate communication and understanding

3. Analysis of sentiment and trends in online discussions

- Collect data on online discussions about COVID-19 from sources such as social networks, online forums and press articles.
- Apply natural language processing techniques to analyze discussion sentiment and trends.
- Identify emerging topics, popular hashtags and dominant opinions on COVID-19.
- Track the evolution of public perceptions and concerns over time.

By combining these three sub-sections, the Spark project can provide a comprehensive and insightful analysis of COVID-19 management around the world.

Technical Details

The Spark project will use a variety of tools and technologies to carry out its analyses. These tools may include :

- Databases and data warehouses to store and manage COVID-19 data
- Statistical analysis and visualization tools to identify trends and patterns in the data
- Natural language processing (NLP) platforms to analyze online discussions on COVID-19
- Data mining and machine learning tools to discover hidden information in the data

Conclusion

This Spark project aims to provide a comprehensive and insightful analysis of COVID-19 management around the world.

By identifying the best-performing regions and understanding the factors that contribute to their success, the project can help improve global efforts to combat the pandemic.

In addition, analysis of online discussions about COVID-19 can provide valuable insights into public perceptions and concerns, which can help inform communication and awareness strategies.

3.2 Configuration of the Spark Ecosystem

3.2.1 Windows environment

Step 1 — Prerequisites

Before installing Apache PySpark, it is essential to ensure that certain prerequisites are met:

- **Java Development Kit (JDK):** It is necessary to have JDK version 8 or higher installed on the system.
- **Python:** The latest version of Python must also be installed.

Step 2 — JDK installation

1. Download the JDK installer file from the official website.
2. Run the installation file and follow the instructions.
3. Choose to specify a specific directory for the Java installation, usually in the C:/Java/jdk folder.

Step 3 — Python Installation

1. Download Python from python.org
2. Execute the installation file and select "Add python.exe to path"

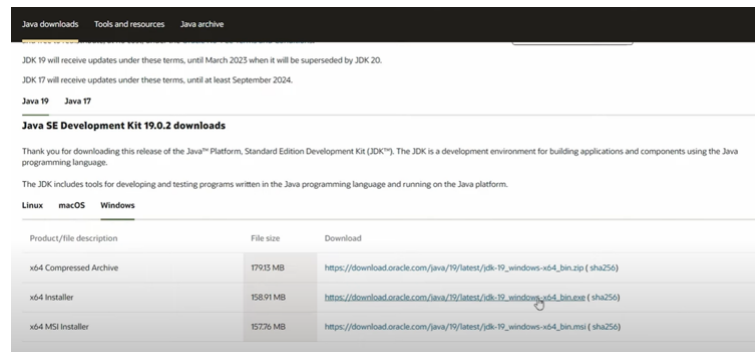


Figure 3.1: JDK Website

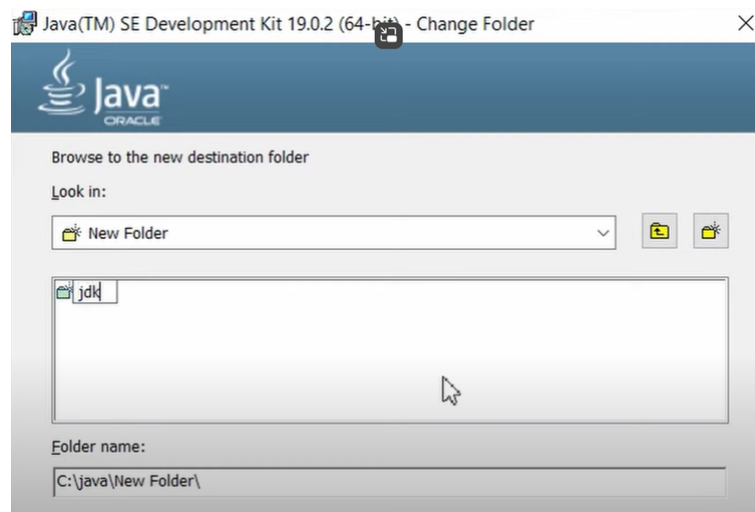


Figure 3.2: JDK Installation

Step 4 — Download and installation of Apache Spark

1. Go to the official Spark website (spark.apache.org) and select the appropriate version.
2. Download the corresponding tar file.
3. Extract the contents of the tar file to a specific directory (for example, C:/spark).

Step 5 — WinUtils download

1. Download the WinUtils.exe file from a Git repository.



Figure 3.3: python.org

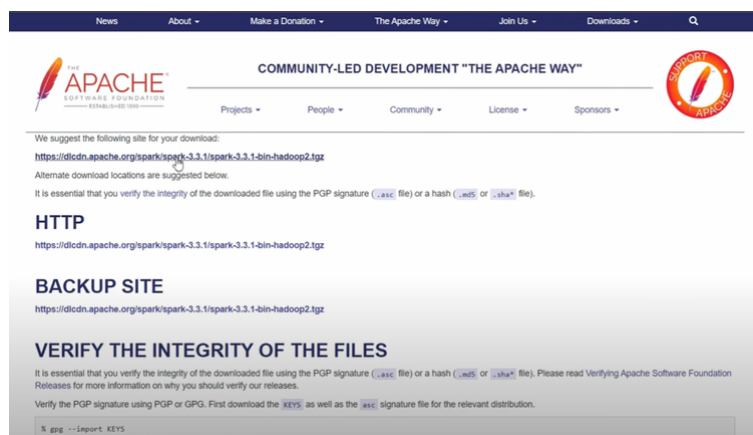


Figure 3.4: Apache Spark official website

2. Place the WinUtils.exe file in the bin directory of the Hadoop folder (for example, C:/Hadoop/bin).

Step 6 — Configuring environment variables

Access system environment variables through Windows Settings. Add the following variables:

- **JAVA_HOME:** Path to the Java installation directory (for example, C:/Java/jdk).
- **HADOOP_HOME:** Path to the Hadoop installation directory (for example, C:/Hadoop).

- **SPARK_HOME:** Path to the Spark installation directory (for example, C:/spark).
- Path to the Python installation directory (for example, C:/Users/UserName/AppData/Local/Programs/Python/Python3x-32/Python3x-32/).
- Also add the corresponding bin directory paths for JAVA_HOME, HADOOP_HOME and SPARK_HOME.

3.2.1.1 Step 7 — Download Pyspark

1. Open a terminal and launch the following command:

```
curl https://bootstrap.pypa.io/get-pip.py -o get-pip.py
```

2. Then launch this command:

```
python get-pip.py
```

3. After installing *pip* you get verify it by launching this command:

```
python -m pip help
```

4. Now install *pyspark* using the following command:

```
pip install pyspark
```

Step 8 — Verifying the installation

1. Open Command Prompt and check the versions of Java and Python.
2. Launch the Spark shell by typing “spark-shell” in the command prompt.
3. Also verify that PySpark is running by typing ”pyspark” into the command prompt.

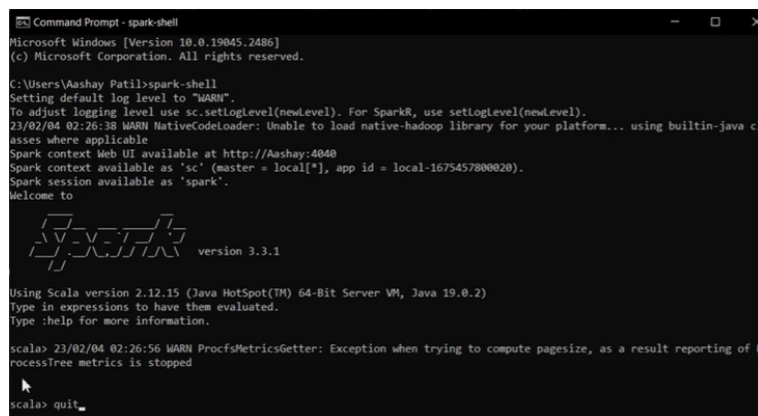


Figure 3.5: Spark-shell

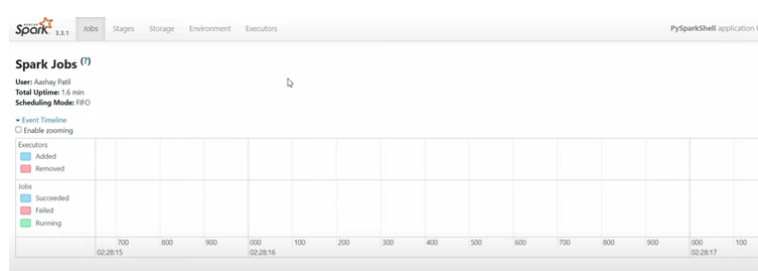


Figure 3.6: Apache Spark Jobs tab

Step 9 — Spark UI Visualization

1. Open a web browser and access the Spark user interface by entering the URL provided in the video.
2. Explore different UI features to monitor running Spark applications.

3.2.2 Linux environment

Step 1 — Check prerequisites

Check if Java is installed using:

```
java --version
```

or:

which java


```
wget https://d1cdn.apache.org/spark/spark-3.5.1/
spark-3.5.1-bin-hadoop3.tgz
```

2. Extract the file:

```
tar -xvf spark_file_name
```

3. Remove the tar folder

```
sudo rm -r spark_file_name
```

4. Rename the extracted folder

```
sudo mv spark_extracted_folder Spark
```

5. Add spark to your path by adding the line below to you shell file (ex: `.zshrc`):

```
export PATH="\$PATH:/opt/spark/Spark/bin"
```

3.2.2.1 Step 4 — Pyspark installation

Launch the following command:

```
pip install pyspark
```

If you have any problems install it using the following command on Debian:

```
sudo apt get python-pyspark
```

In Arch you are going to need *yay*:

```
git clone https://aur.archlinux.org/yay.git
cd yay
makepkg -si
yay -S python-pyspark
yay -S python-py4j
```

3.2.2.2 Step 5 — Use Spark

You verify the installation by executing this command:

```
spark-shell
```

3.2.3 MacOS environment

Step 1 — Install Homebrew

To install *Homebrew* execute this line below:

```
/bin/bash -c "$(curl -fsSL  
https://raw.githubusercontent.com/Homebrew/install/HEAD/install.sh)"
```

3.2.3.1 Step 2 — Install Java (openjdk8)

Install *JDK* using the command below:

```
brew install --cask homebrew/cask-versions/adoptopenjdk8
```

3.2.3.2 Step 3 — Set \$JAVA_HOME environment variable

1. Find your \$JAVA_HOME using the command below:

```
ls /Library/Java/JavaVirtualMachines/adoptopenjdk-8.jdk/  
Contents/Home
```

2. Add it to your shell file (`/.zshrc` for example):

```
export JAVA_HOME=/Library/Java/JavaVirtualMachines/adoptopenjdk-  
8.jdk/Contents/Home
```

3. Source your shell file:

```
source ~/.zshrc
```

3.2.3.3 Step 4 — Install Spark

1. Download Latest Version from <https://spark.apache.org/downloads.html>
2. Unzip the file:

```
tar -xf spark-3.4.1-bin-hadoop3.tgz
```

3.2.3.4 Step 5 — Set \$SPARK_HOME environment variable

Where this file lives, that is your \$SPARK_HOME, go ahead and write that to the necessary environment variable in your `/.zshrc` file as you did with \$JAVA_HOME



Figure 3.9: Spark Home

Set this to the environment variable:

```
export SPARK_HOME=~/.spark-3.4.1-bin-hadoop3
```

Step 6 — Test Spark

To launch the Scala Based Spark, open a terminal and type:

```
spark-shell
```

You typically should work with virtual environments so that you don't face problems with other packages.

Install *py4j*:

```
pip install py4j
```

Create a virtual environment in the directory that you want to work in.

```
python3 -m venv env
```

To activate the virtual environment in that folder run

3.3 Our Use Case: Leveraging Spark for K-means clustering

For our use case we choosed to implement the k-means algorithm that ships in with Spark MLlib in order to do a clustering of some data that we have generated and then evaluate this clustering.

Step 1 — Imports

First we start by simply importing all the *pyspark* libraries that we will be needing:

```
from pyspark.sql import SparkSession
from pyspark.ml.clustering import KMeans
from pyspark.ml.feature import VectorAssembler
from pyspark.ml.evaluation import ClusteringEvaluator
```

- **SparkSession:** We need this class from the *pyspark.sql* sub-package in order to manipulate our spark session.
- **KMeans:** From the *MLlib* of *pyspark* we import the KMeans class in order to do the clustering.
- **VectorAssembler:** It will permit us to do the proper preprocessing operations in order to apply machine learning algorithms to the data.
- **ClusteringEvaluator:** This class will permit us to evaluate the clustering that we did.

Step 2 — Spark session creation and reading data

In this step we create the spark session using the line below:

```
spark = SparkSession.builder.appName("KMeansExample").getOrCreate()
```

Next we read the data from a certain file:

```
data = spark.read.csv("data.csv", header=True, inferSchema=True)
```

After that we prepare the data to be processed by the machine learning algorithm:

```
feature_cols = data.columns
assembler = VectorAssembler(inputCols=feature_cols, outputCol="features")
data = assembler.transform(data)
```

Step 3 — K-means

Next we apply the k-means algorithm on the data:

```
kmeans = KMeans(k=3, seed=123)
model = kmeans.fit(data)
```

We can show the clusters of each point:

```
data_clusterd = model.transform(data)
data_clusterd.show()
```

Step 4 — Clustering evaluation

Finally we can evaluate our clustering by using the `ClusteringEvaluator` class:

```
evaluator = ClusteringEvaluator()
silhouette = evaluator.evaluate(data_clusterd)

print("Silhouette score = " + str(silhouette))
```

3.4 Conclusion

In this paper, we explored the fundamental concepts of Big Data, database systems, and Apache Spark architecture. We showed that traditional data processing approaches are no longer sufficient for the volume and complexity of data generated today, and that distributed frameworks such as Spark provide an effective solution for large-scale analytics. Through the study of Spark's ecosystem, we highlighted its flexibility, performance, and ability to support multiple data processing paradigms within a single framework. The practical part of the paper demonstrated how Spark MLlib can be used to implement the K-means clustering algorithm on distributed data. This use case confirmed the usefulness of Spark for machine learning tasks, especially when working with large datasets that require fast and scalable processing. It also illustrated how Spark simplifies the development of data-driven applications by providing high-level abstractions and optimized execution. Overall, Apache Spark stands out as a versatile and efficient platform for Big Data processing and analysis. Its growing adoption in industry and research shows its importance in addressing the challenges of modern data management. Future work could extend this study by exploring other Spark-based machine learning algorithms, real-time streaming applications, or more advanced clustering and classification techniques.

Bibliography

- [1] Spark: Cluster Computing with Working Sets. Consulté le : 01-04-2024. URL: https://www.usenix.org/legacy/event/hotcloud10/tech/full_papers/Zaharia.pdf.
- [2] Database with PostgreSQL. Consulté le : 16-04-2024. URL: <https://www.oracle.com/cloud/postgresql/>.
- [3] Apache Spark Officially Sets a New Record in Large-Scale Sorting. Consulté le : 01-04-2024. URL: <https://www.databricks.com/blog/2014/11/05/spark-officially-sets-a-new-record-in-large-scale-sorting.html>.
- [4] Learn Why Data Scientists and Developers Need More Than a Data Lake. Consulté le : 05-04-2024. URL: <https://www.actian.com/blog/cloud-data-warehouse/learn-why-data-scientists-and-developers-need-more-than-a-data-lake/>.
- [5] Big data analytics. Consulté le : 05-04-2024. URL: <https://www.ibm.com/analytics/big-data-analytics>.
- [6] Big Data Analytics. What it is and why it matters. Consulté le : 05-04-2024. URL: https://www.sas.com/en_us/insights/analytics/big-data-analytics.html.
- [7] What is a Data Lake? Consulté le : 05-04-2024. URL: <https://azure.microsoft.com/en-us/resources/cloud-computing-dictionary/what-is-a-data-lake>.
- [8] Difference between Traditional Data and Big Data. Consulté le : 05-04-2024. URL: <https://www.javatpoint.com/difference-between-traditional-data-and-big-data>.