

第5章Spark SQL

(源自:<https://biglab.site>)

(版本:Ver1.7-20230828)

学习目标

- 了解Spark SQL基本概念
- 掌握Spark SQL与Shell交互
- 掌握创建DataFrame对象的方法
- 掌握DataFrame查看数据的方法
- 掌握DataFrame的查询及输出操作

任务背景

基于大数据技术对房价进行分析和预测，是科学精准调控，促进房地产市场平稳健康发展，实现“房住不炒”的重要手段。现有一份房屋销售数据文件house.csv，记录了某地区的房屋销售情况，包含销售价格、销售日期、房屋评分等共14个数据字段，字段说明如下。（1英尺=0.3048米）

字段名称	说明
selling_price	销售价格（单位：美元）
bedrooms_num	卧室数
bathroom_num	浴室数
housing_area	房屋面积（单位：平方英尺）
parking_area	停车区面积（单位：平方英尺）
floor_num	楼层数
sales_data	销售日期
built_area	建筑面积（单位：平方英尺）
basement_area	地下室面积（单位：平方英尺）
year_bulit	修建年份
year_repair	修复年份
latitude	纬度
longitude	经度
housing_rating	房屋评分

在进行房价数据分析之前，由于无法直接判断出各个数据字段之间的关系，因此需要先对数据进行基础的探索，探索各个数据字段间的关系并加以分析。

本章将使用Spark SQL即席查询框架解决房价数据探索分析的问题。

1. 首先介绍Spark SQL框架及其编程数据模型DataFrame的基本概念，并对Spark SQL相关环境进行配置。
2. 其次介绍DataFrame的查询、输出操作API的用法。
3. 最后结合房价数据探索分析实例，帮助读者掌握DataFrame的基础操作。

学习重点

理论学习

- (1) 认识Spark SQL。
- (2) Spark SQL CLI配置。
- (3) Spark SQL与Shell交互。
- (4) DataFrame基础操作方法

实验学习

- (1) 配置Spark SQL CLI。
- (2) 掌握DataFrame基础操作。
- (3) 探索分析房屋售价数据。
- (4) 统计分析各季度房屋销量和销售额。
- (5) 探索分析房屋评分。

数据准备

上传spark自带的example到hdfs

```
1 hdfs dfs -mkdir -p /user/myname/sparkSql
2 hdfs dfs -put /usr/local/spark/examples/src/main/resources/users.parquet
  /user/myname/sparkSql/
3 hdfs dfs -put /usr/local/spark/examples/src/main/resources/people.json
  /user/myname/sparkSql/
4 hdfs dfs -put /usr/local/spark/examples/src/main/resources/people.txt
  /user/myname/sparkSql/
```

下载测试数据文件spark_data.tar.gz（若前面已下载过，可跳过）

```
1 mkdir /root/spark/
2 cd /root/spark/
3 wget http://biglab.site/b59510spark/file/spark_data.tar.gz
4 tar -xvzf ./spark_data.tar.gz
5 hdfs dfs -mkdir /user/myname
```

上传测试数据到hdfs

```
1 cd /root/spark
2 hdfs dfs -put ./spark_data/movies.dat /user/myname/sparkSql/
3 hdfs dfs -put ./spark_data/users.dat /user/myname/sparkSql/
4 hdfs dfs -put ./spark_data/ratings.dat /user/myname/sparkSql/
```

上传实训数据到hdfs

```
1 wget -c https://biglab.site/b46488/file/train5data.tar.gz
2 tar -xzf ./train5data.tar.gz
3 hdfs dfs -put -d ./train5data /user/myname/sparkSql/
4 hdfs dfs -ls /user/myname/sparkSql/train5data
```

环境准备

版本配套(了解)

Hadoop 3.3.1

```
1 [hadoop@master jars]$ hadoop version
2 Hadoop 3.3.1
3 Source code repository https://github.com/apache/hadoop.git -r
a3b9c37a397ad4188041dd80621bdeefc46885f2
4 Compiled by ubuntu on 2021-06-15T05:13Z
5 Compiled with protoc 3.7.1
6 From source with checksum 88a4ddb2299aca054416d6b7f81ca55
7 This command was run using /usr/local/hadoop-
3.3.1/share/hadoop/common/hadoop-common-3.3.1.jar
```

Spark 3.1.3

```
1 [hadoop@master ~]$ spark-shell
2 22/08/16 14:30:44 WARN NativeCodeLoader: Unable to load native-hadoop library
for your platform... using builtin-java classes where applicable
3 22/08/16 14:31:02 INFO SignalUtils: Registering signal handler for INT
4 22/08/16 14:35:16 INFO SparkContext: Running Spark version 3.1.3
```

Scala 2.12.16

```
1 [hadoop@master bin]$ scala -version
2 scala code runner version 2.12.16 -- Copyright 2002-2022, LAMP/EPFL and
Lightbend, Inc.
```

Hive 3.1.3

```
1 [hadoop@master bin]$ scala -version
2 Scala code runner version 2.12.16 -- Copyright 2002-2022, LAMP/EPFL and
   Lightbend, Inc.
3 [hadoop@master bin]$ hive -version
4 which: no hbase in
   (/usr/local/hive/bin:/usr/local/bin:/usr/bin:/usr/local/sbin:/usr/sbin:/usr/lo
   cal/hadoop-3.3.1/sbin:/usr/local/hadoop-3.3.1/bin:/usr/lib/jvm/java-1.8.0-
   openjdk-1.8.0.322.b06-
   1.e17_9.x86_64/jre/bin:/usr/local/scala/bin:/usr/local/spark/bin:/usr/local/sp
   ark/sbin:/home/hadoop/.local/bin:/home/hadoop/bin)
5 Hive Session ID = 1551baf4-0917-496f-987f-57b7c21ae61a
6
7 Logging initialized using configuration in jar:file:/usr/local/hive/lib/hive-
   common-3.1.3.jar!/hive-log4j2.properties Async: true
```

软件包准备

```
1 cd /usr/local/spark/jars/
2 wget http://home.hddly.cn:90/soft/mysql/mysql-connector-java-5.1.45-bin.jar
3 (或 wget http://bigdata.hddly.cn/b46488/file/chap6/mysql-connector-java-
   5.1.45-bin.jar)
4 scp /usr/local/spark/jars/mysql-connector-java-5.1.45-bin.jar
   slave1:/usr/local/spark/jars/
5 scp /usr/local/spark/jars/mysql-connector-java-5.1.45-bin.jar
   slave2:/usr/local/spark/jars/
```

(备注: 若有指定ssh端口号, 可参考:

scp -P 8022 /usr/local/spark/jars/mysql-connector-java-5.1.45-bin.jar c24:/usr/local/spark/jars/)

任务5.1认识Spark SQL

任务描述

使用Spark SQL探索分析房价数据前, 需要先了解Spark SQL是什么, 有什么作用。

本节任务如下。

1. 了解Spark SQL框架、Spark SQL的编程模型DataFrame和Spark SQL的运行过程。
2. 对Spark SQL相关环境进行配置。

Spark SQL简介

- 是一个用来处理结构化数据的Spark组件
- 是分布式的SQL查询引擎, 并且提供一个DataFrame可编程抽象数据模型
- 前身是Shark, 由于Shark需要依赖Hive, Spark SQL摆脱了对Hive依赖

- Spark SQL在数据兼容性好
 - 可直接处理RDD
 - 可处理Parquet文件或JSON文件
 - 可处理外部数据库数据，及Hive中存在的表
 - 可统一处理RDD和关系表
- Spark SQL最核心的
编程抽象是DataFrame
 - DataFrame是一个分布式的Row对象的数据集合
 - 它本身实现了RDD绝大多数功能
- 基本处理流程
 - Spark SQL通常从外部数据源加载数据为DataFrame，
 - 然后通过DataFrame上丰富的API进行查询、转换，
 - 最后将结果进行展现或存储

Mysql元数据库安装和配置

删除已有mysql安装新mysql

```

1 rm -rf /usr/lib/mysql
2 rm -rf /usr/include/mysql
3 rm -rf /etc/my.cnf
4 rm -rf /var/lib/mysql
5 rm -rf /usr/share/mysql
6 cd /root
7 yum remove -y mariadb-libs
8 yum install -y net-tools
9 yum install -y perl
10 yum remove mysql mysql-server
11 yum install -y mariadb-libs
12 rpm -ivh mysql-community-* --force --nodeps
13 cp /root/mysql-connector-java-5.1.32-bin.jar /usr/local/hive/lib/

```

启mysql服务

```

1 systemctl start mysqld
2 systemctl enable mysqld
3 service mysqld restart

```

查看临时密码

```

1 [root@master ~]# more /var/log/mysqld.log |grep password |grep generated
2 2023-11-08T03:46:49.898572Z 1 [Note] A temporary password is generated for
   root@localhost: K6UdV+XTmoAB
3 [root@master ~]

```

使用临时密码登陆:

```
1 | mysql -uroot -p'K6UdV+XTmoAB'
```

-p后的临时密码需要替换成本机查询出来的密码。

修改mysql的root密码

进入mysql命令后, 修改mysql的root密码为:123456

```
1 | alter user 'root'@'localhost' identified by '@Root_123456';
2 | set global validate_password_policy=0;
3 | set global validate_password_length=6;
4 | set password for 'root'@'localhost'=password('123456');
5 | GRANT ALL PRIVILEGES ON *.* TO 'root'@'%' IDENTIFIED BY '123456' WITH GRANT
   | OPTION;
6 | FLUSH PRIVILEGES;
7 | CREATE SCHEMA `hive` DEFAULT CHARACTER SET utf8mb4 COLLATE utf8mb4_bin ;
8 | grant all on *.* to hive@'%' identified by 'hive';
9 | grant all on *.* to hive@'localhost' identified by 'hive';
10 | grant all on *.* to hive@'master' identified by 'hive';
11 | grant all privileges on *.* to 'hive'@'%' identified by 'hive' with grant
   | option;
12 | grant all privileges on *.* to 'hive'@'master' identified by 'hive' with
   | grant option;
13 | flush privileges;
14 | create database hive;
15 | show databases;
16 | exit;
```

Spark Hive 安装

在master主机上安装hive:

```
1 | cd /root
2 | wget https://repo.huaweicloud.com/apache/hive/hive-3.1.3/apache-hive-3.1.3-
   | bin.tar.gz --no-check-certificate
3 | tar -xvzf ./apache-hive-3.1.3-bin.tar.gz
4 | mv ./apache-hive-3.1.3-bin /usr/local/hive
```

Spark SQL配置

软件包准备

下载mysql-connector-java-5.1.49-bin.jar到 /usr/local/spark/jars目录下

```
1 wget https://repo.huaweicloud.com/mysql/Downloads/Connector-J/mysql-connector-  
java-5.1.49.tar.gz  
2 tar -xzf ./mysql-connector-java-5.1.49.tar.gz  
3 cp ./mysql-connector-java-5.1.49/mysql-connector-java-5.1.49-bin.jar  
/usr/local/spark/jars/  
4 cd /usr/local/spark/jars/  
5 ls mysql-connector-java *
```

如果ls有显示mysql-connector-java-5.1.49-bin.jar文件，说明mysql包已完成准备。

(还可以从mysql官网下载connector: <https://downloads.mysql.com/archives/c-j/>)

替换hive-site.xml

```
1 sed -i 's/home.hddly.cn:53306/master:3306/g' \  
/usr/local/hive/conf/hive-site.xml  
3 ln -s /usr/local/hive/conf/hive-site.xml \  
/usr/local/spark/conf/hive-site.xml
```

核对hive-site.xml

```
1 vi /usr/local/hive/conf/hive-site.xml
```

```
1 <property>  
2 <name>javax.jdo.option.ConnectionURL</name>  
3 <value>jdbc:mysql://master:3306/hive?  
useSSL=false&useUnicode=true&characterEncoding=UTF-8</value>  
4 </property>  
5 <property>  
6 <name>javax.jdo.option.ConnectionDriverName</name>  
7 <value>com.mysql.jdbc.Driver</value>  
8 <description>Driver class name for a JDBC metastore</description>  
9 </property>  
10 <property>  
11 <name>javax.jdo.option.ConnectionUserName</name>  
12 <value>hive</value>  
13 <description>Username to use against metastore database</description>  
14 </property>  
15 <property>  
16 <name>javax.jdo.option.ConnectionPassword</name>  
17 <value>hive</value>  
18 <description>password to use against metastore database</description>  
19 </property>  
20 <property>  
21 <name>hive.metastore.event.db.notification.api.auth</name>  
22 <value>false</value>  
23 </property>  
24 <property>  
25 <name>hive.metastore.schema.verification</name>  
26 <value>false</value>  
27 </property>
```

初始化元数据库

```
1 | cd /usr/local/hive/bin
2 | ./schematool -dbType mysql -initSchema
```

运行后会出现如下completed的内容，说明元数据库初始化完成

```
1 | Initialization script completed
2 | schemaTool completed
3 | [root@master bin]#
```

#####

启动hive测试下

```
1 | hive
```

配置spark

修改配置spark-env.sh,添加或修改行

```
1 | export SPARK_CLASSPATH=/usr/local/spark/jars/mysql-connector-java-5.1.45-
   | bin.jar
```

分发到从机

```
1 | cd /usr/local/spark/conf
2 | scp ./* slave1:/usr/local/spark/conf/
3 | scp ./* slave2:/usr/local/spark/conf/
```

修改日志级别

```
1 | echo '' > /usr/local/hive/conf/log4j.properties
2 | echo 'log4j.rootLogger = ERROR,stdout' >>
   | /usr/local/hive/conf/log4j.properties
3 | echo 'log4j.appender.stdout=org.apache.log4j.ConsoleAppender' >>
   | /usr/local/hive/conf/log4j.properties
4 | echo 'log4j.appender.stdout.Target=System.out' >>
   | /usr/local/hive/conf/log4j.properties
5 | echo 'log4j.appender.stdout.layout=org.apache.log4j.PatternLayout' >>
   | /usr/local/hive/conf/log4j.properties
6 | echo 'log4j.appender.stdout.layout.ConversionPattern=%d{yyyy/MM/dd
   | HH:mm:ss,SSS}- %c{1}: %m%n' >> /usr/local/hive/conf/log4j.properties
```

核对log4j.properties 修改后如：


```
1 log4j.rootLogger = ERROR,stdout
2 log4j.appender.stdout=org.apache.log4j.ConsoleAppender
3 log4j.appender.stdout.Target=System.out
4 log4j.appender.stdout.layout=org.apache.log4j.PatternLayout
5 log4j.appender.stdout.layout.ConversionPattern=%d{yyyy/MM/dd HH:mm:ss,SSS}-
  %c{1}: %m%n
```

分发日志配置

```
1 \cp /usr/local/hive/conf/log4j.properties
  /usr/local/spark/conf/log4j.properties
```

启动Spark集群

```
1 /usr/local/spark/sbin/start-all.sh
```

启动Spark客户端

观察从机jps是否有worker进程，如没有则运行从机上的worker程序，如果有则不用运行：

```
1 /usr/local/spark/sbin/start-worker.sh master:7077
```

启动spark-sql

```
1 spark-sql
```

使用HQL创建表

```
1 show databases;
2 create database myname;
3 use myname;
4 create table students(id int,name string,score double, classes string);
```

Spark SQL与Shell交互

特性

- Spark SQL已经集成在spark-shell中
- 在Spark-shell中执行SQL需使用SQLContext对象来调用sql()方法
- Spark SQL 查询分成SQLContext和HiveContext
- SQLContext和HiveContext区别
 - HiveContext继承了SQLContext
 - 过时，已被SparkSession代替
 - HiveContext对象
 - val hiveContext=new org.apache.spark.sql.hive.HiveContext(sc)

- SQLContext只支持SQL语法解析器
 - 过时，已被SparkSession代替
 - SQLContext对象
 - `val sqlContext= new org.apache.spark.sql.SQLContext(sc)`
- HiveContext同时支持SQL语法解析器和HiveQL语法解析器
 - 过时，已被SparkSession代替
- SparkSession
 - spark2后SQLContext和HiveContext合并为SparkSession,启动spark-shell后为spark对象
 - spark-shell启动过程会初始化SparkSession对象为spark,它同时支持SQL语法解析器和HiveQL语法解析器

启动spark-shell

- spark-shell
- SparkSession对象
 - 进入spark-shell后可直接使用spark对象：
Spark session available as 'spark'.

使用spark测试:

```
1 val dbs= spark.sql("show databases;")
2 dbs.collect;
```

运行结果:

```
1 scala> val dbs= spark.sql("show databases;");
2 dbs: org.apache.spark.sql.DataFrame = [namespace: string]
3
4 scala> dbs.collect;
5 res0: Array[org.apache.spark.sql.Row] = Array([default], [huangmm], [law],
6 [limm], [myname], [zhangmm])
7 scala>
```

任务5.2掌握DataFrame基础操作

任务描述

- DataFrame由SchemaRDD发展来
- Spark1.3.0开始，SchemaRDD更名为DataFrame
- 不是直接继承自RDD
- 它实现了RDD的绝大多数功能
- 可以理解为一个分布式Row对象的数据集合

创建DataFrame对象

通过结构化数据文件，Hive中的表，外部数据库，Spark计算过程生成RDD

结构化数据文件创建DataFrame

load()方法将HDFS上的格式化文件转换为DataFrame

load默认导入的文件格式是Parquet

Parquet文件

上传测试文件(如果前面数据准备已完成，这步可忽略)：

```
1 hdfs dfs -mkdir -p /user/myname/sparkSql
2 hdfs dfs -put /usr/local/spark/examples/src/main/resources/users.parquet
  /user/myname/sparkSql/
3 hdfs dfs -put /usr/local/spark/examples/src/main/resources/people.json
  /user/myname/sparkSql/
4 hdfs dfs -put /usr/local/spark/examples/src/main/resources/people.txt
  /user/myname/sparkSql/
```

进入spark-shell, 创建DataFrame

```
1 val dfUsers=spark.read.load("/user/myname/sparkSql/users.parquet")
2 dfUsers.collect;
```

运行结果如：

```
1
2 scala> val dfUsers=spark.read.load("/user/myname/sparkSql/users.parquet")
3 dfUsers: org.apache.spark.sql.DataFrame = [name: string, favorite_color:
  string ... 1 more field]
4
5 scala> dfUsers.collect;
6 res0: Array[org.apache.spark.sql.Row] = Array([Alyssa,null,wrappedArray(3, 9,
  15, 20)], [Ben,red,wrappedArray()])
7
8 scala>
```

通过format加载Json

文件并创建DataFrame

```
1 val dfPeople
  =spark.read.format("json").load("/user/myname/sparkSql/people.json")
2 dfPeople.collect;
```

运行结果如：

```

1 scala> val dfPeople
  =spark.read.format("json").load("/user/myname/sparkSql/people.json")
2 dfPeople: org.apache.spark.sql.DataFrame = [age: bigint, name: string]
3
4 scala> dfPeople.collect;
5 res1: Array[org.apache.spark.sql.Row] = Array([null,Michael], [30,Andy],
  [19,Justin])
6
7 scala>

```

通过json方法将json

文件创建DataFrame

```

1 val dfPeople=spark.read.json("/user/myname/sparkSql/people.json")
2 dfPeople.collect;

```

运行结果如：

```

1
2 scala> val dfPeople=spark.read.json("/user/myname/sparkSql/people.json")
3 dfPeople: org.apache.spark.sql.DataFrame = [age: bigint, name: string]
4
5 scala> dfPeople.collect;
6 res2: Array[org.apache.spark.sql.Row] = Array([null,Michael], [30,Andy],
  [19,Justin])

```

外部数据创建DataFrame

还可以从外部数据库(mysql,oracle)中创建DataFrame

需通过JDBC或ODBC连接的方式访问数据库

将mysql中的test中people表数据转换成DataFrame

```

1 val url="jdbc:mysql://home.hddly.cn:53306/test?useSSL=false"
2 val jdbcDF=spark.read.format("jdbc").options(Map("url"->url,"user"-
  >"test","password"->"test","dbtable"->"class")).load()
3 jdbcDF.show()

```

运行结果如：

```

1 scala> val url="jdbc:mysql://home.hddly.cn:53306/test?useSSL=false"
2 url: String = jdbc:mysql://home.hddly.cn:53306/test?useSSL=false
3
4 scala> val jdbcDF=spark.read.format("jdbc").options(Map("url"->url,"user"-
  >"test","password"->"test","dbtable"->"class")).load()
5 jdbcDF: org.apache.spark.sql.DataFrame = [id: int, name: string ... 2 more
  fields]
6

```

```

7 scala> jdbcDF.show()
8 +---+-----+-----+-----+-----+
9 | id| name|                                text| stud|
10 +---+-----+-----+-----+-----+
11 | 1|tipdm|大数据产品 - 泰迪科技-专注于大...| 张三|
12 | 2|tipdm|大数据产品 - 泰迪科技-专注于大...| 张三|
13 | 3|tipdm|大数据产品 - 泰迪科技-专注于大...| 张三|
14 | 4|tipdm|大数据产品 - 泰迪科技-专注于大...| 张三|
15 | 5|tipdm|大数据产品 - 泰迪科技-专注于大...| 张三|
16 | 6|tipdm|大数据产品 - 泰迪科技-专注于大...| 张三|
17 | 7|tipdm|大数据产品 - 泰迪科技-专注于大...|包思瑶|
18 | 8|tipdm|大数据产品 - 泰迪科技-专注于大...|薛伟杰|
19 | 9|tipdm|大数据产品 - 泰迪科技-专注于大...| 张三|
20 |10|tipdm|大数据产品 - 泰迪科技-专注于大...|汪睿鹏|
21 |11|tipdm|大数据产品 - 泰迪科技-专注于大...| 林维|
22 |12|tipdm|大数据产品 - 泰迪科技-专注于大...| 张三|
23 |13|tipdm|大数据产品 - 泰迪科技-专注于大...| 林维|
24 |14|tipdm|大数据产品 - 泰迪科技-专注于大...| 林维|
25 |15|tipdm|大数据产品 - 泰迪科技-专注于大...|汪睿鹏|
26 |16|tipdm|大数据产品 - 泰迪科技-专注于大...|包思瑶|
27 |17|tipdm|大数据产品 - 泰迪科技-专注于大...|汪睿鹏|
28 |18|tipdm|大数据产品 - 泰迪科技-专注于大...| 李洪|
29 |19|tipdm|大数据产品 - 泰迪科技-专注于大...| 张三|
30 |20|tipdm|大数据产品 - 泰迪科技-专注于大...|陈建良|
31 +---+-----+-----+-----+-----+
32 only showing top 20 rows
33

```

RDD创建DataFrame

利用反射机制推断RDD

例子

```

1 case class Person(name:String,age:Int)
2 val data=sc.textFile("/user/myname/sparkSql/people.txt").map(_.split(","))
3 val people=data.map(p=>Person(p(0),p(1).trim.toInt)).toDF
4 people.collect
5 people.show()

```

运行结果如：

```

1 scala> case class Person(name:String,age:Int)
2 defined class Person
3
4 scala> val
data=sc.textFile("/user/myname/sparkSql/people.txt").map(_.split(","))
5 data: org.apache.spark.rdd.RDD[Array[String]] = MapPartitionsRDD[24] at map
at <console>:24
6
7 scala> val people=data.map(p=>Person(p(0),p(1).trim.toInt)).toDF
8 people: org.apache.spark.sql.DataFrame = [name: string, age: int]

```

```

9
10 scala> people.collect
11 res4: Array[org.apache.spark.sql.Row] = Array([Michael,29], [Andy,30],
    [Justin,19])
12
13 scala> people.show()
14 +-----+----+
15 |  name|age|
16 +-----+----+
17 |Michael| 29|
18 |  Andy| 30|
19 | Justin| 19|
20 +-----+----+
21
22
23 scala>

```

采用编程方式指定Schema的方式将RDD转换为DataFrame

1,从原来的RDD创建一个元组或列表的RDD

2,用StructType创建一个和步骤1中创建的RDD

中元组或列表的结构相匹配的Schema

3,通过spark提供的createDataFrame方法

将Schema应用到RDD上

例如

```

1 val people=sc.textFile("/user/myname/sparkSql/people.txt")
2 val schemaString ="name age"
3 import org.apache.spark.sql.Row
4 import org.apache.spark.sql.types.{StructType,StructField,StringType}
5 val schema=StructType(schemaString.split("
    ").map(fieldName=>StructField(fieldName,StringType,true)))
6 val rowRDD=people.map(_.split(",")).map(p=>Row(p(0),p(1).trim))
7 val peopleDataFrame=spark.createDataFrame(rowRDD,schema)
8 peopleDataFrame.show

```

运行结果:

```

1 scala> val people=sc.textFile("/user/myname/sparkSql/people.txt")
2 people: org.apache.spark.rdd.RDD[String] = /user/myname/sparkSql/people.txt
    MapPartitionsRDD[1] at textFile at <console>:24
3
4 scala> val schemaString ="name age"
5 schemaString: String = name age
6
7 scala> import org.apache.spark.sql.Row
8 import org.apache.spark.sql.Row
9
10 scala> import org.apache.spark.sql.types.
    {StructType,StructField,StringType}
11 import org.apache.spark.sql.types.{StructType, StructField, StringType}
12

```

```

13 scala> val schema=StructType(schemaString.split("
14 schema: org.apache.spark.sql.types.StructType =
15   StructType(StructField(name,StringType,true),
16   StructField(age,StringType,true))
17
18 scala> val rowRDD=people.map(_._split(",")).map(p=>Row(p(0),p(1).trim))
19 rowRDD: org.apache.spark.rdd.RDD[org.apache.spark.sql.Row] =
20 MapPartitionsRDD[3] at map at <console>:27
21
22 scala> val peopleDataFrame=spark.createDataFrame(rowRDD,schema)
23 peopleDataFrame: org.apache.spark.sql.DataFrame = [name: string, age:
24 string]
25
26 scala> peopleDataFrame.show
27
28 +-----+----+
29 |  name | age |
30 +-----+----+
31 | Michael | 29 |
32 |  Andy  | 30 |
33 |  Justin | 19 |
34 +-----+----+

```

Hive中的表创建DataFrame

查询Hive中的表,转换成DataFrame

在hive中准备表, 在hive中执行:

```

1 show databases;
2 create database myname;
3 use myname;
4 create table students(id int,name string,score double, classes string);
5 insert into students(id,name,score,classes) values(1,'limm',90,'class1');

```

(参考常见问题1处理insert失败的问题)

在spark-shell中执行

```

1 spark.sql("use myname")
2 val people=spark.sql("select * from students")
3 people.show

```

运行结果:

```

1 scala> spark.sql("use myname")
2 res0: org.apache.spark.sql.DataFrame = []
3
4 scala> val people=spark.sql("select * from students")

```

```
5 | people: org.apache.spark.sql.DataFrame = [id: int, name: string ... 2 more
   | fields]
6 |
7 | scala> people.show
8 | +---+-----+-----+-----+
9 | | id |  name|score|classes|
10 | +---+-----+-----+-----+
11 | | 10|limm10| 90.0| class1|
12 | | 12|limm12| 90.0| class1|
13 | +---+-----+-----+-----+
14 |
15 |
16 | scala>
```

查看DataFrame数据

查看及获取数据的常用函数或方法

方法	描述
printSchema	打印数据模式
show	查看数据
first/head/take/takeAsList	获取若干行数据
collect/collectAsList	获取所有数据

数据准备

上传测试文件(如果前面数据准备已完成，这步可忽略)：

```
1 | cd /root/spark
2 | hdfs dfs -put ./spark_data/movies.dat /user/myname/sparkSql/
3 | hdfs dfs -put ./spark_data/users.dat /user/myname/sparkSql/
4 | hdfs dfs -put ./spark_data/ratings.dat /user/myname/sparkSql/
```

DataFrame常用操作函数或方法

创建DataFrame对象movies

```
1 | case class Movie(movieId:Int,title:String,Genres:String)
2 | val data=sc.textFile("/user/myname/sparkSql/movies.dat").map(_.split(":"))
3 | val movies=data.map(m=>Movie(m(0).trim.toInt,m(1),m(2))).toDF()
```

运行结果如：


```

1 scala> case class Movie(moveieId:Int,titile:String,Genres:String)
2 defined class Movie
3
4 scala> val
data=sc.textFile("/user/myname/sparkSql/movies.dat").map(_.split(":"))
5 data: org.apache.spark.rdd.RDD[Array[String]] = MapPartitionsRDD[15] at map
at <console>:24
6
7 scala> val movies=data.map(m=>Movie(m(0).trim.toInt,m(1),m(2))).toDF()
8 movies: org.apache.spark.sql.DataFrame = [moveieId: int, titile: string ...
1 more field]
9
10 scala>

```

printSchema:打印数据模式

```
1 movies.printSchema
```

运行结果:

```

1 scala> movies.printSchema
2 root
3 |-- moveieId: integer (nullable = false)
4 |-- titile: string (nullable = true)
5 |-- Genres: string (nullable = true)
6
7
8 scala>

```

show查看数据

方法	解释
show()	显示前20条记录
show(numRows:Int)	显示numRows条
show(truncate:Boolean)	是否最多只显示20个字符，默认为true
show(numRows:Int,truncate:Boolean)	显示numRows条记录并设置过长字符串的显示格式

```

1 movies.show()
2 movies.show(false)

```

运行结果:

```

1 scala> movies.show()
2 +-----+-----+-----+-----+
3 |moveieId|          titile|          Genres|
4 +-----+-----+-----+-----+
5 |         1| Toy Story (1995)|Animation|Childre...|

```

```

6 | 2| Jumanji (1995)|Adventure|Childre...|
7 | 3|Grumpier Old Men ...| Comedy|Romance|
8 | 4|Waiting to Exhale...| Comedy|Drama|
9 | 5|Father of the Bri...| Comedy|
10 | 6| Heat (1995)|Action|Crime|Thri...|
11 | 7| Sabrina (1995)| Comedy|Romance|
12 | 8| Tom and Huck (1995)|Adventure|Children's|
13 | 9| Sudden Death (1995)| Action|
14 | 10| GoldenEye (1995)|Action|Adventure|...|
15 | 11|American Presiden...|Comedy|Drama|Romance|
16 | 12|Dracula: Dead and...| Comedy|Horror|
17 | 13| Balto (1995)|Animation|Children's|
18 | 14| Nixon (1995)| Drama|
19 | 15|Cutthroat Island ...|Action|Adventure|...|
20 | 16| Casino (1995)| Drama|Thriller|
21 | 17|Sense and Sensibi...| Drama|Romance|
22 | 18| Four Rooms (1995)| Thriller|
23 | 19|Ace Ventura: when...| Comedy|
24 | 20| Money Train (1995)| Action|

```

```

25 +-----+-----+-----+-----+

```

```

26 only showing top 20 rows

```

```

29 scala> movies.show(false)

```

```

30 +-----+-----+-----+-----+
31 |moveieId|titile                                     |Genres
32 +-----+-----+-----+-----+
33 |1        |Toy Story (1995)                                |Animation|Children's|Comedy
34 |2        |Jumanji (1995)                                  |Adventure|Children's|Fantasy|
35 |3        |Grumpier Old Men (1995)                         |Comedy|Romance
36 |4        |Waiting to Exhale (1995)                       |Comedy|Drama
37 |5        |Father of the Bride Part II (1995)             |Comedy
38 |6        |Heat (1995)                                     |Action|Crime|Thriller
39 |7        |Sabrina (1995)                                  |Comedy|Romance
40 |8        |Tom and Huck (1995)                             |Adventure|Children's
41 |9        |Sudden Death (1995)                             |Action
42 |10       |GoldenEye (1995)                                |Action|Adventure|Thriller
43 |11       |American President, The (1995)                 |Comedy|Drama|Romance
44 |12       |Dracula: Dead and Loving It (1995)             |Comedy|Horror

```

```

45 |13      |Balto (1995)                |Animation|Children's
46 |14      |Nixon (1995)                |Drama
47 |15      |Cutthroat Island (1995)     |Action|Adventure|Romance
48 |16      |Casino (1995)               |Drama|Thriller
49 |17      |Sense and Sensibility (1995)|Drama|Romance
50 |18      |Four Rooms (1995)           |Thriller
51 |19      |Ace Ventura: When Nature Calls (1995)|Comedy
52 |20      |Money Train (1995)           |Action
53 +-----+-----+-----+-----+
54 |
55 |
56 |
57 |

```

only showing top 20 rows

scala>

first/head/take/takeAsList获取若干行数据

```

1 | movies.first()
2 | movies.head(3)
3 | movies.take(3)
4 | movies.takeAsList(3)

```

运行结果：

```
1 scala> movies.first()
2 res6: org.apache.spark.sql.Row = [1,Toy Story
  (1995),Animation|Children's|Comedy]
3
4 scala> movies.head(3)
5 res7: Array[org.apache.spark.sql.Row] = Array([1,Toy Story
  (1995),Animation|Children's|Comedy], [2,Jumanji
  (1995),Adventure|Children's|Fantasy], [3,Grumpier Old Men
  (1995),Comedy|Romance])
6
7 scala> movies.take(3)
8 res8: Array[org.apache.spark.sql.Row] = Array([1,Toy Story
  (1995),Animation|Children's|Comedy], [2,Jumanji
  (1995),Adventure|Children's|Fantasy], [3,Grumpier Old Men
  (1995),Comedy|Romance])
9
10 scala> movies.takeAsList(3)
11 res9: java.util.List[org.apache.spark.sql.Row] = [[1,Toy Story
  (1995),Animation|Children's|Comedy], [2,Jumanji
  (1995),Adventure|Children's|Fantasy], [3,Grumpier Old Men
  (1995),Comedy|Romance]]
12
13 scala>
```

collect/collectAsList获取所有数据

```
1 movies.collect()
2 movies.collectAsList()
```

运行结果：

```

1 scala> movies.collect()
2 res10: Array[org.apache.spark.sql.Row] = Array([1,Toy Story
(1995),Animation|Children's|Comedy], [2,Jumanji
(1995),Adventure|Children's|Fantasy], [3,Grumpier Old Men
(1995),Comedy|Romance], [4,Waiting to Exhale (1995),Comedy|Drama], [5,Father
of the Bride Part II (1995),Comedy], [6,Heat (1995),Action|Crime|Thriller],
[7,Sabrina (1995),Comedy|Romance], [8,Tom and Huck
(1995),Adventure|Children's], [9,Sudden Death (1995),Action], [10,GoldenEye
(1995),Action|Adventure|Thriller], [11,American President, The
(1995),Comedy|Drama|Romance], [12,Dracula: Dead and Loving It
(1995),Comedy|Horror], [13,Balto (1995),Animation|Children's], [14,Nixon
(1995),Drama], [15,Cutthroat Island (1995),Action|Adventure|Romance],
[16,Casino (1995),Drama|Thriller], [17,Sense and Sensibil...
3
4 scala> movies.collectAsList()
5 res11: java.util.List[org.apache.spark.sql.Row] = [[1,Toy Story
(1995),Animation|Children's|Comedy], [2,Jumanji
(1995),Adventure|Children's|Fantasy], [3,Grumpier Old Men
(1995),Comedy|Romance], [4,Waiting to Exhale (1995),Comedy|Drama], [5,Father
of the Bride Part II (1995),Comedy], [6,Heat (1995),Action|Crime|Thriller],
[7,Sabrina (1995),Comedy|Romance], [8,Tom and Huck
(1995),Adventure|Children's], [9,Sudden Death (1995),Action], [10,GoldenEye
(1995),Action|Adventure|Thriller], [11,American President, The
(1995),Comedy|Drama|Romance], [12,Dracula: Dead and Loving It
(1995),Comedy|Horror], [13,Balto (1995),Animation|Children's], [14,Nixon
(1995),Drama], [15,Cutthroat Island (1995),Action|Adventure|Romance],
[16,Casino (1995),Drama|Thriller], [17,Sense and Sens...
6
7 scala>

```

掌握DataFrame查询操作

1, 方法一将DataFrame注册成为临时表, 然后通过SQL语句进行查询

```

1 val people=sc.textFile("/user/myname/sparkSql/people.txt")
2 val schemaString ="name age"
3 import org.apache.spark.sql.Row
4 import org.apache.spark.sql.types.{StructType,StructField,StringType}
5 val schema=StructType(schemaString.split("
").map(fieldName=>StructField(fieldName,StringType,true)))
6 val rowRDD=people.map(_.split(",")).map(p=>Row(p(0),p(1).trim))
7 val peopleDataFrame=spark.createDataFrame(rowRDD,schema)
8 peopleDataFrame.registerTempTable("peopleTempTab")
9 val personsRDD=spark.sql("select name,age from peopleTempTab where age
>20").rdd
10 personsRDD.collect

```

运行结果如:

```

1 scala> val people=sc.textFile("/user/myname/sparkSql/people.txt")
2 people: org.apache.spark.rdd.RDD[String] = /user/myname/sparkSql/people.txt
MapPartitionsRDD[42] at textFile at <console>:24

```

```

3
4 scala> val schemaString ="name age"
5 schemaString: String = name age
6
7 scala> import org.apache.spark.sql.Row
8 import org.apache.spark.sql.Row
9
10 scala> import org.apache.spark.sql.types.{StructType, StructField, StringType}
11 import org.apache.spark.sql.types.{StructType, StructField, StringType}
12
13 scala> val schema=StructType(schemaString.split("
14 ").map(fieldName=>StructField(fieldName,StringType,true)))
15 schema: org.apache.spark.sql.types.StructType =
16   StructType(StructField(name,StringType,true),
17   StructField(age,StringType,true))
18
19 scala> val rowRDD=people.map(_.split(",")).map(p=>Row(p(0),p(1).trim))
20 rowRDD: org.apache.spark.rdd.RDD[org.apache.spark.sql.Row] =
21   MapPartitionsRDD[44] at map at <console>:27
22
23 scala> val peopleDataFrame=spark.createDataFrame(rowRDD,schema)
24 peopleDataFrame: org.apache.spark.sql.DataFrame = [name: string, age:
25   string]
26
27 scala> peopleDataFrame.registerTempTable("peopleTempTab")
28 warning: there was one deprecation warning (since 2.0.0); for details,
29 enable `:setting -deprecation` or `:replay -deprecation`
30
31 scala> val personsRDD=spark.sql("select name,age from peopleTempTab where
32   age >20").rdd
33 personsRDD: org.apache.spark.rdd.RDD[org.apache.spark.sql.Row] =
34   MapPartitionsRDD[49] at rdd at <console>:25
35
36 scala> personsRDD.collect
37 res13: Array[org.apache.spark.sql.Row] = Array([Michael,29], [Andy,30])
38
39 scala>

```

2, 方法二直接在DataFrame对象上进行查询

DataFrame提供了很多查询的方法, 是Transformation操作,是懒操作

方法	描述
where	条件查询
select/selectExpr/col/apply	查询指定字段的数据信息
limit	查询前n行记录
order by	排序查询
group by	分组查询
join	连接查询

创建DataFrame对象rating和用户

```

1 case class Rating(userId:Int, movieId:Int, rating:Int, timestamp:Long)
2 var
  ratingData=sc.textFile("/user/myname/sparkSql/ratings.dat").map(_.split(":"))
3 val
  rating=ratingData.map(r=>Rating(r(0).trim.toInt, r(1).trim.toInt, r(2).trim.toIn
t, r(3).trim.toLong)).toDF()
4 case class User(userId:Int, gender:String, age:Int, occupation:Int, zip:String)
5 val userData=sc.textFile("/user/myname/sparkSql/users.dat").map(_.split(":"))
6 val
  user=userData.map(u=>User(u(0).trim.toInt, u(1), u(2).trim.toInt, u(3).trim.toInt
, u(4))).toDF()

```

运行结果:

```

1 scala> case class Rating(userId:Int, movieId:Int, rating:Int, timestamp:Long)
2 defined class Rating
3
4 scala> var
  ratingData=sc.textFile("/user/myname/sparkSql/ratings.dat").map(_.split(":")
))
5 ratingData: org.apache.spark.rdd.RDD[Array[String]] = MapPartitionsRDD[52]
at map at <console>:26
6
7 scala> val
  rating=ratingData.map(r=>Rating(r(0).trim.toInt, r(1).trim.toInt, r(2).trim.to
Int, r(3).trim.toLong)).toDF()
8 rating: org.apache.spark.sql.DataFrame = [userId: int, movieId: int ... 2
more fields]
9
10 scala> case class
  User(userId:Int, gender:String, age:Int, occupation:Int, zip:String)
11 defined class User
12
13 scala> val
  userData=sc.textFile("/user/myname/sparkSql/users.dat").map(_.split(":"))
14 userData: org.apache.spark.rdd.RDD[Array[String]] = MapPartitionsRDD[56] at
map at <console>:26
15

```

```

16 | scala> val
    | user=userData.map(u=>User(u(0).trim.toInt,u(1),u(2).trim.toInt,u(3).trim.toI
    | nt,u(4))).toDF()
17 | user: org.apache.spark.sql.DataFrame = [userId: int, gender: string ... 3
    | more fields]
18 |
19 | scala>

```

查询条件:where查询&filter

where查询

```

1 | val userWhere=user.where("gender='F' and age=18")
2 | userWhere.show(3)

```

结果:

```

1 | scala> val userWhere=user.where("gender='F' and age=18")
2 | userWhere: org.apache.spark.sql.Dataset[org.apache.spark.sql.Row] = [userId:
  | int, gender: string ... 3 more fields]
3 |
4 | scala> userWhere.show(3)
5 | +-----+-----+---+-----+-----+
6 | |userId|gender|age|occupation| zip|
7 | +-----+-----+---+-----+-----+
8 | |    18|    F| 18|          3|95825|
9 | |    34|    F| 18|          0|02135|
10 | |    38|    F| 18|          4|02215|
11 | +-----+-----+---+-----+-----+
12 | only showing top 3 rows
13 |
14 |
15 | scala>

```

filter

```

1 | val userFilter=user.filter("gender='F' and age=18")
2 | userFilter.show(3)

```

结果:

```

1 |
2 | scala> val userFilter=user.filter("gender='F' and age=18")
3 | userFilter: org.apache.spark.sql.Dataset[org.apache.spark.sql.Row] =
  | [userId: int, gender: string ... 3 more fields]
4 |
5 | scala> userFilter.show(3)
6 | +-----+-----+---+-----+-----+
7 | |userId|gender|age|occupation| zip|
8 | +-----+-----+---+-----+-----+
9 | |    18|    F| 18|          3|95825|
10 | |    34|    F| 18|          0|02135|
11 | |    38|    F| 18|          4|02215|
12 | +-----+-----+---+-----+-----+

```



```
13 | only showing top 3 rows
14 |
15 |
16 | scala>
```

查询指定字段的数据信息select&selectExpr&col&apply

select:获取指定字段值

```
1 | val userSelect=user.select("userId","gender")
2 | userSelect.show
```

运行结果:

```
1 | scala> val userSelect=user.select("userId","gender")
2 | userSelect: org.apache.spark.sql.DataFrame = [userId: int, gender: string]
3 |
4 | scala> userSelect.show
5 | +-----+-----+
6 | |userId|gender|
7 | +-----+-----+
8 | |    1|    F|
9 | |    2|    M|
10 | |    3|    M|
11 | |    4|    M|
12 | |    5|    M|
13 | |    6|    F|
14 | |    7|    M|
15 | |    8|    M|
16 | |    9|    M|
17 | |   10|    F|
18 | |   11|    F|
19 | |   12|    M|
20 | |   13|    M|
21 | |   14|    M|
22 | |   15|    M|
23 | |   16|    F|
24 | |   17|    M|
25 | |   18|    F|
26 | |   19|    M|
27 | |   20|    M|
28 | +-----+-----+
29 | only showing top 20 rows
30 |
31 |
32 | scala>
```

selectExpr:对指定字段进行特殊处理

```
1 spark.udf.register("replace", (x:String)=>
2   {x match{
3     case "M"=>0
4     case "F"=>1
5   }})
6 val userSelectExpr=user.selectExpr("userId","replace(gender) as sex","age")
7 userSelectExpr.show(3)
```

运行结果:

```
1 scala> val userSelectExpr=user.selectExpr("userId","replace(gender) as
sex","age")
2 userSelectExpr: org.apache.spark.sql.DataFrame = [userId: int, sex: int ...
1 more field]
3
4 scala> userSelectExpr.show(3)
5 +-----+-----+
6 |userId|sex|age|
7 +-----+-----+
8 |    1| 1|  1|
9 |    2| 0| 56|
10 |    3| 0| 25|
11 +-----+-----+
12 only showing top 3 rows
13
14
15 scala>
```

col/apply

获取DataFrame指定字段,但是只能获取一个字段,并且返回对象为Column类型

```
1 val userCol=user.col("zip")
2 user.select(userCol).collect
3 val userApply=user.apply("zip")
4 user.select(userApply).collect
```

运行结果:

```

1 scala> val userCol=user.col("zip")
2 userCol: org.apache.spark.sql.Column = zip
3
4 scala> user.select(userCol).collect
5 res19: Array[org.apache.spark.sql.Row] = Array([48067], [70072], [55117],
6 [02460], [55455], [55117], [06810], [11413], [61614], [95370], [04093],
7 [32793], [93304], [60126], [22903], [20670], [95350], [95825], [48073],
8 [55113], [99353], [53706], [90049], [10023], [01609], [23112], [19130],
9 [14607], [33407], [19143], [06840], [19355], [55421], [02135], [02482],
10 [94123], [66212], [02215], [61820], [10543], [15116], [24502], [60614],
11 [98052], [94110], [75602], [94305], [92107], [77084], [98133], [10562],
12 [72212], [96931], [56723], [55303], [60440], [30350], [30303], [55413],
[72118], [95122], [98105], [54902], [53706], [55803], [57706], [60181],
[53706], [02143], [53703], [95008], [55122], [53706], [94530], [01748],
[55413], [15321], [98029], [98103], [49327], [606...

6
7 scala> val userApply=user.apply("zip")
8 userApply: org.apache.spark.sql.Column = zip
9
10 scala> user.select(userApply).collect
11 res20: Array[org.apache.spark.sql.Row] = Array([48067], [70072], [55117],
[02460], [55455], [55117], [06810], [11413], [61614], [95370], [04093],
[32793], [93304], [60126], [22903], [20670], [95350], [95825], [48073],
[55113], [99353], [53706], [90049], [10023], [01609], [23112], [19130],
[14607], [33407], [19143], [06840], [19355], [55421], [02135], [02482],
[94123], [66212], [02215], [61820], [10543], [15116], [24502], [60614],
[98052], [94110], [75602], [94305], [92107], [77084], [98133], [10562],
[72212], [96931], [56723], [55303], [60440], [30350], [30303], [55413],
[72118], [95122], [98105], [54902], [53706], [55803], [57706], [60181],
[53706], [02143], [53703], [95008], [55122], [53706], [94530], [01748],
[55413], [15321], [98029], [98103], [49327], [606...

```

limit

获取指定DataFrame的前n行记录

```

1 val userLimit=user.limit(3)
2 userLimit.show

```

运行结果:

```

1 scala> val userLimit=user.limit(3)
2 userLimit: org.apache.spark.sql.Dataset[org.apache.spark.sql.Row] = [userId:
int, gender: string ... 3 more fields]
3
4 scala> userLimit.show
5 +-----+-----+---+-----+-----+
6 |userId|gender|age|occupation| zip|
7 +-----+-----+---+-----+-----+
8 |      1|      F|  1|          10|48067|

```

```

9 |      2 |      M | 56 |      16 | 70072 |
10 |      3 |      M | 25 |      15 | 55117 |
11 | +-----+-----+-----+-----+
12 |
13 |
14 | scala>

```

orderBy/sort

orderBy排序方法

```

1 | val userOrderBy1=user.orderBy(desc("UserId"))
2 | userOrderBy1.show(3)
3 | val userOrderBy2=user.orderBy($"UserId".desc)
4 | userOrderBy2.show(3)
5 | val userOrderBy3=user.orderBy(-user("userId"))
6 | userOrderBy3.show(3)

```

运行结果:

```

1 | scala> val userOrderBy1=user.orderBy(desc("UserId"))
2 | userOrderBy1: org.apache.spark.sql.Dataset[org.apache.spark.sql.Row] =
  [userId: int, gender: string ... 3 more fields]
3 |
4 | scala> userOrderBy1.show(3)
5 | +-----+-----+-----+-----+
6 | |userId|gender|age|occupation| zip|
7 | +-----+-----+-----+-----+
8 | | 6041 |      M | 25 |      7 | 11107 |
9 | | 6040 |      M | 25 |      6 | 11106 |
10 | | 6039 |      F | 45 |      0 | 01060 |
11 | +-----+-----+-----+-----+
12 | only showing top 3 rows
13 |
14 |
15 | scala> val userOrderBy2=user.orderBy($"UserId".desc)
16 | userOrderBy2: org.apache.spark.sql.Dataset[org.apache.spark.sql.Row] =
  [userId: int, gender: string ... 3 more fields]
17 |
18 | scala> userOrderBy2.show(3)
19 | +-----+-----+-----+-----+
20 | |userId|gender|age|occupation| zip|
21 | +-----+-----+-----+-----+
22 | | 6041 |      M | 25 |      7 | 11107 |
23 | | 6040 |      M | 25 |      6 | 11106 |
24 | | 6039 |      F | 45 |      0 | 01060 |
25 | +-----+-----+-----+-----+
26 | only showing top 3 rows
27 |
28 |
29 | scala> val userOrderBy3=user.orderBy(-user("userId"))
30 | userOrderBy3: org.apache.spark.sql.Dataset[org.apache.spark.sql.Row] =
  [userId: int, gender: string ... 3 more fields]
31 |
32 | scala> userOrderBy3.show(3)

```

```

33 | +-----+-----+-----+-----+
34 | |userId|gender|age|occupation| zip|
35 | +-----+-----+-----+-----+
36 | | 6041|      M| 25|          7|11107|
37 | | 6040|      M| 25|          6|11106|
38 | | 6039|      F| 45|          0|01060|
39 | +-----+-----+-----+-----+
40 | only showing top 3 rows
41 |
42 |
43 | scala>

```

sort方法

```

1 | val userSort1=user.sort(asc("userId"))
2 | userSort1.show(3)
3 | val userSort2=user.sort($"userId".asc)
4 | userSort2.show(3)
5 | val userSort3=user.sort(user("userId"))
6 | userSort3.show(3)

```

运行结果:

```

1 | scala> val userSort1=user.sort(asc("userId"))
2 | userSort1: org.apache.spark.sql.Dataset[org.apache.spark.sql.Row] = [userId:
  | int, gender: string ... 3 more fields]
3 |
4 | scala> userSort1.show(3)
5 | +-----+-----+-----+-----+
6 | |userId|gender|age|occupation| zip|
7 | +-----+-----+-----+-----+
8 | |    1|      F|  1|          10|48067|
9 | |    2|      M| 56|          16|70072|
10 | |    3|      M| 25|          15|55117|
11 | +-----+-----+-----+-----+
12 | only showing top 3 rows
13 |
14 |
15 | scala> val userSort2=user.sort($"userId".asc)
16 | userSort2: org.apache.spark.sql.Dataset[org.apache.spark.sql.Row] = [userId:
  | int, gender: string ... 3 more fields]
17 |
18 | scala> userSort2.show(3)
19 | +-----+-----+-----+-----+
20 | |userId|gender|age|occupation| zip|
21 | +-----+-----+-----+-----+
22 | |    1|      F|  1|          10|48067|
23 | |    2|      M| 56|          16|70072|
24 | |    3|      M| 25|          15|55117|
25 | +-----+-----+-----+-----+
26 | only showing top 3 rows
27 |
28 |
29 | scala> val userSort3=user.sort(user("userId"))

```

```

30 userSort3: org.apache.spark.sql.Dataset[org.apache.spark.sql.Row] = [userId:
   int, gender: string ... 3 more fields]
31
32 scala> userSort3.show(3)
33 +-----+-----+-----+-----+
34 |userId|gender|age|occupation| zip|
35 +-----+-----+-----+-----+
36 |    1|    F|  1|         10|48067|
37 |    2|    M| 56|         16|70072|
38 |    3|    M| 25|         15|55117|
39 +-----+-----+-----+-----+
40 only showing top 3 rows
41
42

```

groupBy

两种分组方法:传入String类型和传入Column类型

groupBy返回的是GroupData对象:

方法	描述
max(colNames:String)	获取分组中指定字段或者所有的数值类型字段的最大值
min(colNames:String)	获取分组中指定字段或者所有的数字类型字段的最小值
mean(colNames:String)	获取分组中指定字段或者所有的数字类型字段的平均值
sum(colNames:String)	获取分组中指定字段或者所有的数字类型字段和值
count()	获取分组中的元素个数

```

1 val userGroupBy1 = user.groupBy("gender")
2 val userGroupBy2 = user.groupBy(user("gender"))
3 val userGroupByCount = user.groupBy(user("gender")).count
4 userGroupByCount.show()

```

运行结果:

```

1 scala> val userGroupBy1 = user.groupBy("gender")
2 userGroupBy1: org.apache.spark.sql.RelationalGroupedDataset =
   RelationalGroupedDataset: [grouping expressions: [gender: string], value:
   [userId: int, gender: string ... 3 more fields], type: GroupBy]
3
4 scala> val userGroupBy2 = user.groupBy(user("gender"))
5 userGroupBy2: org.apache.spark.sql.RelationalGroupedDataset =
   RelationalGroupedDataset: [grouping expressions: [gender: string], value:
   [userId: int, gender: string ... 3 more fields], type: GroupBy]
6 scala> val userGroupByCount = user.groupBy(user("gender")).count
7 userGroupByCount: org.apache.spark.sql.DataFrame = [gender: string, count:
   bigint]
8

```

```

9  scala> userGroupByCount.show()
10  +-----+-----+
11  |gender|count|
12  +-----+-----+
13  |      F| 1709|
14  |      M| 4332|
15  +-----+-----+
16
17
18  scala>

```

join

DataFrame提供了三种join方法用于连接两个表

方法	描述
join(right:DataFrame)	两个表做笛卡尔积
join(right:DataFrame,joinExprs:Column)	根据两表中相同的某个字段进行连接
join(right:DataFrame,joinExprs:Column,joinType:String)	根据两表中相同的某个字段进行连接并指定连接类型

连接查询

```

1  val dfJoin = user.join(rating,"userId")
2  dfJoin.show(3)
3  val dfJoin2 = dfJoin.join(user,Seq("userId","gender"),"left_outer")
4  dfJoin2.show(3)

```

运行结果

```

1  scala> val dfJoin = user.join(rating,"userId")
2  dfJoin: org.apache.spark.sql.DataFrame = [userId: int, gender: string ... 6
  more fields]
3
4  scala> dfJoin.show(3)
5  +-----+-----+-----+-----+-----+-----+-----+-----+
6  |userId|gender|age|occupation|  zip|movieId|rating|timestamp|
7  +-----+-----+-----+-----+-----+-----+-----+-----+
8  |   148|    M|  50|          17|57747|  2987|    5|979576038|
9  |   148|    M|  50|          17|57747|  2989|    4|977333611|
10 |   148|    M|  50|          17|57747|   647|    3|977352859|
11 +-----+-----+-----+-----+-----+-----+-----+-----+
12 only showing top 3 rows
13
14 scala> val dfJoin2 = dfJoin.join(user,Seq("userId","gender"),"left_outer")
15 dfJoin2: org.apache.spark.sql.DataFrame = [userId: int, gender: string ... 9
  more fields]
16
17 scala> dfJoin2.show(3)

```

```

18 | +-----+-----+-----+-----+-----+-----+-----+-----+-----+
19 | |userId|gender|age|occupation| zip|movieId|rating|timestamp|age|occupation|
20 | |zip|
21 | +-----+-----+-----+-----+-----+-----+-----+-----+-----+
21 | | 71| M| 25| 14|95008| 648| 4|977875582| 25|
21 | 14|95008|
22 | | 71| M| 25| 14|95008| 1250| 4|977875378| 25|
22 | 14|95008|
23 | | 71| M| 25| 14|95008| 2997| 3|977875912| 25|
23 | 14|95008|
24 | +-----+-----+-----+-----+-----+-----+-----+-----+-----+
24 | +-----+
25 | only showing top 3 rows
26 |
27 |
28 | scala>

```

掌握DataFrame输出操作

方法一

1, 创建Map对象, 存储数据, 如保存路径, json文件头信息

```
val saveOptions=Map("header"->"true","path"->"/user/myname/sparkSql/copyOfUser.json")
```

2, 从DataFrame对象中选择出userId,gender和age

```
val copyOfUser=user.select("userId","gender","age")
```

3, 调用save函数保存2中的

DataFrame数据到文件中

```
import org.apache.spark.sql.SaveMode
```

```
copyOfUser.write.format("json").mode(SaveMode.Overwrite).options(saveOptions).save()
```

4, 查看保存结果

hdfs上查看目录: /user/myname/sparkSql/copyOfUser.json

源码

```

1 | import org.apache.spark.sql.SaveMode
2 | case class User(userId:Int,gender:String,age:Int,occupation:Int,zip:String)
3 | val userData=sc.textFile("/user/myname/sparkSql/users.dat").map(_.split("::"))
4 | val
   user=userData.map(u=>User(u(0).trim.toInt,u(1),u(2).trim.toInt,u(3).trim.toInt
   ,u(4))).toDF()
5 | val saveOptions=Map("header"->"true","path"-
   >"/user/myname/sparkSql/copyOfUser.json")
6 | val copyOfUser=user.select("userId","gender","age")
7 | copyOfUser.write.format("json").mode(SaveMode.Overwrite).options(saveOptions).
   save()

```


运行结果

```
1 scala> import org.apache.spark.sql.SaveMode
2 import org.apache.spark.sql.SaveMode
3
4 scala> case class
User(userId:Int,gender:String,age:Int,occupation:Int,zip:String)
5 defined class User
6
7 scala> val
userData=sc.textFile("/user/myname/sparkSql/users.dat").map(_.split(":"))
8 userData: org.apache.spark.rdd.RDD[Array[String]] = MapPartitionsRDD[32] at
map at <console>:25
9
10 scala> val
user=userData.map(u=>User(u(0).trim.toInt,u(1),u(2).trim.toInt,u(3).trim.toI
nt,u(4))).toDF()
11 user: org.apache.spark.sql.DataFrame = [userId: int, gender: string ... 3
more fields]
12
13 scala> val saveOptions=Map("header"->"true","path"-
>"/user/myname/sparkSql/copyOfUser.json")
14 saveOptions: scala.collection.immutable.Map[String,String] = Map(header ->
true, path -> /user/myname/sparkSql/copyOfUser.json)
15
16 scala> val copyOfUser=user.select("userId","gender","age")
17 copyOfUser: org.apache.spark.sql.DataFrame = [userId: int, gender: string
... 1 more field]
18
19 scala>
copyOfUser.write.format("json").mode(SaveMode.Overwrite).options(saveOptions
).save()
20
21 scala>
```

方法二

1, 直接调用save方法

```
1 val dfPeople =spark.read.json("/user/myname/sparkSql/people.json")
2 dfPeople.show()
3 import org.apache.spark.sql.SaveMode
4 dfPeople.write.mode(SaveMode.Overwrite).json("/user/myname/sparkSql/people_cop
y")
```

说明:

DataFrameReader.json读文件,返回DataFrame , 所以dfPeople是DataFrame

dfPeople.write, 返回DataFrameWriter对象, DataFrameWriter.json保存文件

运行结果:

```

1 scala> import org.apache.spark.sql.SaveMode
2 import org.apache.spark.sql.SaveMode
3
4 scala> case class
User(userId:Int,gender:String,age:Int,occupation:Int,zip:String)
5 defined class User
6
7 scala> val
userData=sc.textFile("/user/myname/sparkSql/users.dat").map(_.split(":"))
8 userData: org.apache.spark.rdd.RDD[Array[String]] = MapPartitionsRDD[32] at
map at <console>:25
9
10 scala> val
user=userData.map(u=>User(u(0).trim.toInt,u(1),u(2).trim.toInt,u(3).trim.toI
nt,u(4))).toDF()
11 user: org.apache.spark.sql.DataFrame = [userId: int, gender: string ... 3
more fields]
12
13 scala> val saveOptions=Map("header"->"true","path"-
>"/user/myname/sparkSql/copyOfUser.json")
14 saveOptions: scala.collection.immutable.Map[String,String] = Map(header ->
true, path -> /user/myname/sparkSql/copyOfUser.json)
15
16 scala> val copyOfUser=user.select("userId","gender","age")
17 copyOfUser: org.apache.spark.sql.DataFrame = [userId: int, gender: string
... 1 more field]
18
19 scala>
copyOfUser.write.format("json").mode(SaveMode.Overwrite).options(saveOptions
).save()
20
21 scala>

```

2, saveAsTable方法保存到表

```

1 case class User(userId:Int,gender:String,age:Int,occupation:Int,zip:String)
2 val userData=sc.textFile("/user/myname/sparkSql/users.dat").map(_.split(":"))
3 val
user=userData.map(u=>User(u(0).trim.toInt,u(1),u(2).trim.toInt,u(3).trim.toInt
,u(4))).toDF()
4
5 val copyOfUser=user.select("userId","gender","age")
6 copyOfUser.write.saveAsTable("copyUser")
7 spark.sql("select * from copyUser").show(5)

```

运行结果:

```

1 scala> case class
User(userId:Int,gender:String,age:Int,occupation:Int,zip:String)
2 defined class User
3

```

```

4  scala> val
   userData=sc.textFile("/user/myname/sparkSql/users.dat").map(_.split(":"))
5  userData: org.apache.spark.rdd.RDD[Array[String]] = MapPartitionsRDD[6] at
   map at <console>:24
6
7  scala> val
   user=userData.map(u=>User(u(0).trim.toInt,u(1),u(2).trim.toInt,u(3).trim.toI
   nt,u(4))).toDF()
8  user: org.apache.spark.sql.DataFrame = [userId: int, gender: string ... 3
   more fields]
9
10 scala>
11     | val copyOfUser=user.select("userId","gender","age")
12 copyOfUser: org.apache.spark.sql.DataFrame = [userId: int, gender: string
   ... 1 more field]
13
14 scala> copyOfUser.write.saveAsTable("copyUser")
15
16 scala> spark.sql("select * from copyUser").show(5)
17 +-----+-----+----+
18 |userId|gender|age|
19 +-----+-----+----+
20 |    1|    F|  1|
21 |    2|    M| 56|
22 |    3|    M| 25|
23 |    4|    M| 45|
24 |    5|    M| 25|
25 +-----+-----+----+
26 only showing top 5 rows
27
28

```

任务5.3探索分析房屋售价数据

任务描述

从购房者角度出发，影响用户购房的主要因素有房屋价格、房屋属性。从地产公司角度出发，除了房屋本身的居住属性外，地产公司也会考虑什么样的房屋属性在市场上更具有价值。本节的任务如下。

使用Spark SQL实现房价数据的探索分析，从数据中的房屋售价、房屋评分和房屋销售日期等维度进行探索分析，获得房价波动的内在原因。

获取数据

创建表并导入数据

准备数据

```
1 hdfs dfs -put /root/spark/spark_data/house.csv /user/myname/sparkSql/
2 hdfs dfs -ls /user/myname/sparkSql/house.csv
```

打开hive,运行如下创建表和load数据

```
1 create database house;
2 use house;
3 create table king_county_house(
4     selling_price double,
5     bedrooms_num double,
6     bathroom_num double,
7     housing_area double,
8     parking_area double,
9     floor_num double,
10    housing_rating double,
11    built_area double,
12    basement_area double,
13    year_built int,
14    year_repair int,
15    latitude double,
16    longitude double,
17    sale_data string
18 )
19 row format delimited fields terminated by ',';
20 load data local inpath '/root/spark/spark_data/house.csv' overwrite into
    table king_county_house;
```

运行结果:

```
1 hive> use house;
2 OK
3 Time taken: 0.089 seconds
4 hive> create table king_county_house(
5     > selling_price double,
6     > bedrooms_num double,
7     > bathroom_num double,
8     > housing_area double,
9     > parking_area double,
10    > floor_num double,
11    > housing_rating double,
12    > built_area double,
13    > basement_area double,
14    > year_built int,
15    > year_repair int,
16    > latitude double,
17    > longitude double,
18    > sale_data string
19    > )
20    > row format delimited fields terminated by ',';
21 OK
22 Time taken: 1.354 seconds
```

```

23 | hive> load data local inpath '/root/spark/spark_data/house.csv' overwrite
    | into table king_county_house;
24 | Loading data to table house.king_county_house
25 | 79774 [cd973e9f-ac65-4d34-96d5-409feef08b62 main] WARN
    | org.apache.hadoop.hive.metastore.ObjectStore -
    | datanucleus.autoStartMechanismMode is set to unsupported value null .
    | Setting it to value: ignored
26 | 80836 [cd973e9f-ac65-4d34-96d5-409feef08b62 main] WARN
    | org.apache.hadoop.hive.metastore.ObjectStore -
    | datanucleus.autoStartMechanismMode is set to unsupported value null .
    | Setting it to value: ignored
27 | OK
28 | Time taken: 1.528 seconds
29 | hive>

```

在spark-shell中读取数据

```

1 | val house = spark.table("house.king_county_house")
2 | house.take(3)

```

运行结果:

```

1 | scala> val house = spark.table("house.king_county_house")
2 | house: org.apache.spark.sql.DataFrame = [selling_price: double, bedrooms_num:
    | double ... 12 more fields]
3 |
4 | scala> house.take(3)
5 | res6: Array[org.apache.spark.sql.Row] =
    | Array([null,null,null,null,null,null,null,null,null,null,null,null,sale_d
    | ata],
    | [545000.0,3.0,2.25,1670.0,6240.0,1.0,8.0,1240.0,430.0,1974,0,47.6413,-122.113,
    | 20200302],
    | [785000.0,4.0,2.5,3300.0,10514.0,2.0,10.0,3300.0,0.0,1984,0,47.6323,-122.036,2
    | 0200211])
6 |
7 | scala>

```

任务实现

源码实现:

```

1 | package main.java
2 |
3 | import org.apache.spark.sql.types.DataTypes
4 | import org.apache.spark.sql.{ColumnName, DataFrame, SparkSession}
5 |
6 | object process {
7 |   def main(args: Array[String]): Unit = {
8 |     val spark = SparkSession.builder()
9 |       .master("local")
10 |      .appName("process")
11 |      .config("spark.some.config.option", "some-value")

```

```

12     .enableHiveSupport()
13     .getOrCreate()
14     spark.sparkContext.setLogLevel("WARN")
15     import org.apache.spark.sql.functions._
16     import spark.implicits._
17     // 读取数据集
18     val house = spark.table("house.king_county_house")
19     val dataColumnName = house.columns.toList
20     for (i <- dataColumnName) {
21         if (i == "sale_data" || i == "year_built" || i == "year_repair") {
22             println(i + ":")
23             null_count(house, i)
24             println("-" * 20)
25         }
26         else {
27             max_min_mean_std(house, i)
28         }
29     }
30     // 转换时间格式
31     val houseDate = house.na.drop()
32     .withColumn("date", to_date(col("sale_data"), "yyyyMMdd"))
33     // 划分季度
34     val houseQuarter = houseDate.withColumn("quarter", quarter(col("date")))
35     houseQuarter.show()
36     // 统计每个季度的房屋销量
37     houseQuarter.groupBy("quarter").count().sort("quarter").show()
38     // 统计每个季度的房屋销售总额
39     houseQuarter.groupBy(
40         "quarter").sum("selling_price").sort("quarter").show()
41     houseQuarter.selectExpr("mean(selling_price)").show(100)
42
43     // 房屋评分分布统计
44
45     houseQuarter.groupBy("housing_rating").count().sort(desc("count")).show()
46
47     // 房屋评分与单位售价的关系
48     houseQuarter.groupBy("housing_rating").agg(
49         avg(col("selling_price") / col("housing_area"))).sort(
50         "housing_rating").show()
51     val houseAgeTitleUdf = udf((x: Int) => AgeTitle(x))
52     val houseRepair = houseQuarter.filter(
53         "year_repair != 0").withColumn(
54         "houseAge", -(col("year_built").cast(DataTypes.IntegerType) - 2020))
55     val houseAgeTitle = houseRepair.withColumn(
56         "houseAgeTitle", houseAgeTitleUdf(col("houseAge")))
57     houseAgeTitle.groupBy("houseAgeTitle").count().show()
58
59     // 求字段的最大值、最小值、均值及标准差
60     def max_min_mean_std(data: DataFrame, columnName: String): Unit = {
61         println(columnName + ":")
62         data.selectExpr("max(" + columnName + ") as max")
63             .foreach(line => println("max:" + line.toString()))
64         data.selectExpr("min(" + columnName + ") as min")
65             .foreach(line => println("min:" + line.toString()))

```

```

66     data.selectExpr("mean(" + columnName + ") as mean")
67     .foreach(line => println ("mean:" + line.toString()))
68     data.selectExpr("stddev(" + columnName + ") as std")
69     .foreach(line => println("std:" + line.toString()))
70     null_count(data, columnName)
71     println("-" * 20)
72 }
73
74 def null_count(data: DataFrame, columnName: String): Unit = {
75     println(columnName + "缺失值数: " + (data.count() -
data.na.drop().count()))
76 }
77
78 def AgeTitle(age: Int) = {
79     if (age >= 30) {
80         "old"
81     }
82     else if (age <= 10) {
83         "new"
84     }
85     else {
86         "middle"
87     }
88 }
89 }

```

打包上传

使用sftp上传word.jar，例如：其中D:\yuxm\wordcount\out\artifacts\word要根据本机情况修改

```

1 | lcd D:\yuxm\wordcount\out\artifacts\word
2 | cd /root/spark/
3 | put word.jar

```

结果如：

```

1 | sftp> lcd D:\yuxm\wordcount\out\artifacts\word
2 | sftp> cd /root/spark/
3 | sftp> put word.jar
4 | Uploading word.jar to /root/spark/word.jar
5 | 100% 29KB 29KB/s 00:00:00
6 | D:/yuxm/wordcount/out/artifacts/word/word.jar: 30327 bytes transferred in 0
seconds (29 KB/s)
7 | sftp>
8 |

```

运行结果：

切换到shell命令行执行脚本

```

1 | spark-submit --master local --class main.java.process /root/spark/word.jar

```

运行结果：

```
1
2 [root@master ~]# spark-submit --master local --class main.java.process
   /root/spark/word.jar
3 SLF4J: Class path contains multiple SLF4J bindings.
4 SLF4J: Found binding in [jar:file:/usr/local/spark/jars/slf4j-log4j12-
   1.7.30.jar!/org/slf4j/impl/StaticLoggerBinder.class]
5 SLF4J: Found binding in [jar:file:/usr/local/hadoop-
   3.3.1/share/hadoop/common/lib/slf4j-log4j12-
   1.7.30.jar!/org/slf4j/impl/StaticLoggerBinder.class]
6 SLF4J: See http://www.slf4j.org/codes.html#multiple_bindings for an
   explanation.
7 SLF4J: Actual binding is of type [org.slf4j.impl.Log4jLoggerFactory]
8 selling_price:
9 max:[6885000.0]
10 min:[75000.0]
11 mean:[542874.9288]
12 std:[372925.76587803545]
13 selling_price缺失值数: 1
14 -----
15 bedrooms_num:
16 max:[10.0]
17 min:[0.0]
18 mean:[3.3676]
19 std:[0.8931685255322745]
20 bedrooms_num缺失值数: 1
21 -----
22 bathroom_num:
23 max:[7.75]
24 min:[0.0]
25 mean:[2.1168]
26 std:[0.7740995950252791]
27 bathroom_num缺失值数: 1
28 -----
29 housing_area:
30 max:[9890.0]
31 min:[390.0]
32 mean:[2082.4884]
33 std:[922.8789159010951]
34 housing_area缺失值数: 1
35 -----
36 parking_area:
37 max:[1651359.0]
38 min:[572.0]
39 mean:[15352.7339]
40 std:[45776.22980321974]
41 parking_area缺失值数: 1
42 -----
43 floor_num:
44 max:[3.5]
45 min:[1.0]
46 mean:[1.50285]
47 std:[0.5436418233940561]
48 floor_num缺失值数: 1
49 -----
50 housing_rating:
```



```

51 max:[13.0]
52 min:[3.0]
53 mean:[7.6648]
54 std:[1.173873399424226]
55 housing_rating缺失值数: 1
56 -----
57 built_area:
58 max:[8860.0]
59 min:[390.0]
60 mean:[1791.4749]
61 std:[829.449437281676]
62 built_area缺失值数: 1
63 -----
64 basement_area:
65 max:[4820.0]
66 min:[0.0]
67 mean:[291.0135]
68 std:[446.64133907009403]
69 basement_area缺失值数: 1
70 -----
71 year_built:
72 year_built缺失值数: 1
73 -----
74 year_repair:
75 year_repair缺失值数: 1
76 -----
77 latitude:
78 max:[47.7776]
79 min:[47.1593]
80 mean:[47.5606288899999]
81 std:[0.13857015670482747]
82 latitude缺失值数: 1
83 -----
84 longitude:
85 max:[-121.315]
86 min:[-122.519]
87 mean:[-122.21584520000049]
88 std:[0.13973856419611894]
89 longitude缺失值数: 1
90 -----
91 sale_data:
92 sale_data缺失值数: 1
93 -----
94 +-+-----+-----+-----+-----+-----+-----+
95 |selling_price|bedrooms_num|bathroom_num|housing_area|parking_area|floor_num|
96 |housing_rating|built_area|basement_area|year_built|year_repair|latitude|lon-
97 |gitude|sale_data|date|quarter|
98 +-+-----+-----+-----+-----+-----+-----+
99 |545000.0|3.0|2.25|1670.0|6240.0|
100 |1.0|8.0|1240.0|430.0|1974|0|
101 |47.6413|-122.113|20200302|2020-03-02|1|

```

98		785000.0	4.0	2.5	3300.0	10514.0
		2.0	10.0	3300.0	0.0	1984
		47.6323	-122.036	20200211	2020-02-11	1
99		765000.0	3.0	3.25	3190.0	5283.0
		2.0	9.0	3190.0	0.0	2007
		47.5534	-122.002	20200107	2020-01-07	1
100		720000.0	5.0	2.5	2900.0	9525.0
		2.0	9.0	2900.0	0.0	1989
		47.5442	-122.138	20191103	2019-11-03	4
101		449500.0	5.0	2.75	2040.0	7488.0
		1.0	7.0	1200.0	840.0	1969
		47.7289	-122.172	20190603	2019-06-03	2
102		248500.0	2.0	1.0	780.0	10064.0
		1.0	7.0	780.0	0.0	1958
		47.4913	-122.318	20200506	2020-05-06	2
103		675000.0	4.0	2.5	1770.0	9858.0
		1.0	8.0	1770.0	0.0	1971
		47.7382	-122.287	20200305	2020-03-05	1
104		730000.0	2.0	2.25	2130.0	4920.0
		1.5	7.0	1530.0	600.0	1941
		47.573	-122.409	20190701	2019-07-01	3
105		311000.0	2.0	1.0	860.0	3300.0
		1.0	6.0	860.0	0.0	1903
		47.5496	-122.279	20190807	2019-08-07	3
106		660000.0	2.0	1.0	960.0	6263.0
		1.0	6.0	960.0	0.0	1942
		47.6646	-122.202	20191204	2019-12-04	4
107		435000.0	2.0	1.0	990.0	5643.0
		1.0	7.0	870.0	120.0	1947
		47.6802	-122.298	20200227	2020-02-27	1
108		350000.0	3.0	1.0	1240.0	10800.0
		1.0	7.0	1240.0	0.0	1959
		47.5233	-122.185	20190904	2019-09-04	3
109		385000.0	3.0	2.25	1630.0	1598.0
		3.0	8.0	1630.0	0.0	2008
		47.6904	-122.347	20190902	2019-09-02	3
110		235000.0	2.0	1.0	930.0	10505.0
		1.0	6.0	930.0	0.0	1930
		47.4337	-122.329	20200413	2020-04-13	2
111		350000.0	3.0	1.0	1300.0	10236.0
		1.0	6.0	1300.0	0.0	1971
		47.5028	-121.77	20190930	2019-09-30	3
112		1350000.0	4.0	1.75	2000.0	3728.0
		1.5	9.0	1820.0	180.0	1926
		47.643	-122.299	20200507	2020-05-07	2
113		459900.0	3.0	1.75	2580.0	11000.0
		1.0	7.0	1290.0	1290.0	1951
		47.5646	-122.181	20190530	2019-05-30	2
114		430000.0	6.0	3.0	2630.0	8800.0
		1.0	7.0	1610.0	1020.0	1959
		47.7166	-122.293	20190723	2019-07-23	3
115		718000.0	5.0	2.75	2930.0	7663.0
		2.0	9.0	2930.0	0.0	2013
		47.5308	-122.184	20191003	2019-10-03	4

```

116 |      220000.0|      4.0|      1.0|      1440.0|      8250.0|
      1.0|      7.0|      1440.0|      0.0|      1959|      0|
47.4325| -122.291| 20191215|2019-12-15|      4|
117 +-----+-----+-----+-----+-----+-----+
      +-----+-----+-----+-----+-----+-----+
      +-----+-----+-----+-----+-----+-----+
118 only showing top 20 rows
119
120 +-----+-----+
121 |quarter|count|
122 +-----+-----+
123 |      1| 1906|
124 |      2| 3139|
125 |      3| 2759|
126 |      4| 2196|
127 +-----+-----+
128
129 +-----+-----+
130 |quarter|sum(selling_price)|
131 +-----+-----+
132 |      1| 1.006865855E9|
133 |      2| 1.770422075E9|
134 |      3| 1.474271351E9|
135 |      4| 1.177190007E9|
136 +-----+-----+
137
138 +-----+-----+
139 |mean(selling_price)|
140 +-----+-----+
141 |      542874.9288|
142 +-----+-----+
143
144 +-----+-----+
145 |housing_rating|count|
146 +-----+-----+
147 |      7.0| 4162|
148 |      8.0| 2828|
149 |      9.0| 1212|
150 |      6.0|  925|
151 |     10.0|  514|
152 |     11.0|  192|
153 |      5.0|  100|
154 |     12.0|   45|
155 |      4.0|   14|
156 |     13.0|    6|
157 |      3.0|    2|
158 +-----+-----+
159
160 +-----+-----+
161 |housing_rating|avg((selling_price / housing_area))|
162 +-----+-----+
163 |      3.0|      289.3034825870647|
164 |      4.0|      421.6765988891822|
165 |      5.0|      244.872904522352|
166 |      6.0|      274.3597415143096|

```

```

167 |          7.0|          252.72629268682468|
168 |          8.0|          260.1858918866025|
169 |          9.0|          274.36099140045263|
170 |         10.0|          311.8355441252322|
171 |         11.0|          337.7554628229812|
172 |         12.0|          407.72963099333407|
173 |         13.0|          495.34880449365454|
174 +-----+
175
176 +-----+
177 |houseAgeTitle|count|
178 +-----+
179 |          old|  425|
180 |        middle|   2|
181 +-----+
182
183 [root@master ~]#

```

实训作业

统计分析某公司每年的产品销售量及销售额

训练要点

- 1, RDD转化为DataFrame的方法
- 2, DataFrame的查询操作

需求说明

以某公司的交易数据，该数据包含3个.txt文档数据，分别为日期数据tbData.txt，订单表头数据tbStock.txt和订单明细tbStockDetail.txt.

将tbData.txt,tbStock.txt,tbStockDetail.txt上传到 HDFS,通过sparkContext读取该数据得到一个RDD,将该RDD数据转换成DataFrame进行查询分析。

分析目标如：

- 1, 计算所有订单中每年的销售单数，销售总额
- 2, 计算所有订单每年最大金额订单的销售额
- 3, 计算所有订单中每年最畅销货品

实现思路及步骤

- 1, 将数据上传到HDFS
- 2, 使用SparkContext读取HDFS上的数据
- 3, 将2中得到的RDD数据转为DataFrame，并注册成为临时表
- 4, 针对分析目标1，使用join连接3个临时表进行查询
- 5, 针对分析目标2，首先求出每份订单的销售额以其发生时间，然后以第1步的查询结果和tbData连接，求出每年最大金额订单的销售订单的销售额
- 6, 针对分析目标3，首先求出每年每个货品的销售利息额，然后求出每年单品销售的最大金额，最后求出每年与销售额最大相符的货品就是最畅销货品

作业要求

- 1, 提供分析过程的所有脚本
- 2, 提供分析目标1,2,3的运行过程截图
- 3, 提供分析目标1,2,3的运行结果截图

参考

1,准备数据

上传实训数据到hdfs

```
1 wget -c http://bigdata.hddly.cn/b46488/file/train5data.tar.gz
2 tar -xvzf ./train5data.tar.gz
3 hdfs dfs -put -d ./train5data /user/myname/sparkSql/
4 hdfs dfs -ls /user/myname/sparkSql/train5data
```

2,加载数据

```
1 case class tbStock(ordernumber:String,locationid:String,dateid:String)
  extends Serializable
2   val tbStockRdd =
  spark.sparkContext.textFile("/user/myname/sparkSql/train5data/tbStock.txt")
3   val tbStockDS = tbStockRdd.map(_._split(",")).map(attr=>
  tbStock(attr(0),attr(1),attr(2))).toDS
4   tbStockDS.show
5
6   case class tbStockDetail(ordernumber:String, rownum:Int, itemid:String,
  number:Int, price:Double, amount:Double) extends Serializable
7   val tbStockDetailRdd =
  spark.sparkContext.textFile("/user/myname/sparkSql/train5data/tbStockDetail.
  txt")
8   val tbStockDetailDS = tbStockDetailRdd.map(_._split(",")).map(attr=>
  tbStockDetail(attr(0),attr(1).trim().toInt,attr(2),attr(3).trim().toInt,attr
  (4).trim().toDouble, attr(5).trim().toDouble)).toDS
9   tbStockDetailDS.show()
10
11  case class tbDate(dateid:String, years:Int, theyear:Int, month:Int,
  day:Int, weekday:Int, week:Int, quarter:Int, period:Int, halfmonth:Int)
  extends Serializable
12  val tbDateRdd =
  spark.sparkContext.textFile("/user/myname/sparkSql/train5data/tbDate.txt")
13  val tbDateDS = tbDateRdd.map(_._split(",")).map(attr=>
  tbDate(attr(0),attr(1).trim().toInt,
  attr(2).trim().toInt,attr(3).trim().toInt, attr(4).trim().toInt,
  attr(5).trim().toInt, attr(6).trim().toInt, attr(7).trim().toInt,
  attr(8).trim().toInt, attr(9).trim().toInt)).toDS
14  tbDateDS.show()
```

3,注册表

```
1 tbStockDS.createOrReplaceTempView("tbStock")
2 tbDateDS.createOrReplaceTempView("tbDate")
3 tbStockDetailDS.createOrReplaceTempView("tbStockDetail")
```

4,查询统计

计算所有订单中每年的销售单数、销售总额

```
1 spark.sql("SELECT c.theyear, COUNT(DISTINCT a.ordernumber), SUM(b.amount) FROM
tbStock a JOIN tbStockDetail b ON a.ordernumber = b.ordernumber JOIN tbDate c
ON a.dateid = c.dateid GROUP BY c.theyear ORDER BY c.theyear").show
```

计算所有订单每年最大金额订单的销售额

```
1 spark.sql("SELECT a.dateid, a.ordernumber, SUM(b.amount) AS SumOfAmount FROM
tbStock a JOIN tbStockDetail b ON a.ordernumber = b.ordernumber GROUP BY
a.dateid, a.ordernumber").show
2 spark.sql("SELECT theyear, MAX(c.SumOfAmount) AS SumOfAmount FROM (SELECT
a.dateid, a.ordernumber, SUM(b.amount) AS SumOfAmount FROM tbStock a JOIN
tbStockDetail b ON a.ordernumber = b.ordernumber GROUP BY a.dateid,
a.ordernumber ) c JOIN tbDate d ON c.dateid = d.dateid GROUP BY theyear ORDER
BY theyear DESC").show
```

计算所有订单中每年最畅销货品

```
1 spark.sql("SELECT c.theyear, b.itemid, SUM(b.amount) AS SumOfAmount FROM
tbStock a JOIN tbStockDetail b ON a.ordernumber = b.ordernumber JOIN tbDate c
ON a.dateid = c.dateid GROUP BY c.theyear, b.itemid").show
2
3 spark.sql("SELECT d.theyear, MAX(d.SumOfAmount) AS MaxOfAmount FROM (SELECT
c.theyear, b.itemid, SUM(b.amount) AS SumOfAmount FROM tbStock a JOIN
tbStockDetail b ON a.ordernumber = b.ordernumber JOIN tbDate c ON a.dateid =
c.dateid GROUP BY c.theyear, b.itemid ) d GROUP BY d.theyear").show
4
5 spark.sql("SELECT DISTINCT e.theyear, e.itemid, f.maxofamount FROM (SELECT
c.theyear, b.itemid, SUM(b.amount) AS sumofamount FROM tbStock a JOIN
tbStockDetail b ON a.ordernumber = b.ordernumber JOIN tbDate c ON a.dateid =
c.dateid GROUP BY c.theyear, b.itemid ) e JOIN (SELECT d.theyear,
MAX(d.sumofamount) AS maxofamount FROM (SELECT c.theyear, b.itemid,
SUM(b.amount) AS sumofamount FROM tbStock a JOIN tbStockDetail b ON
a.ordernumber = b.ordernumber JOIN tbDate c ON a.dateid = c.dateid GROUP BY
c.theyear, b.itemid ) d GROUP BY d.theyear ) f ON e.theyear = f.theyear AND
e.sumofamount = f.maxofamount ORDER BY e.theyear").show
```

常见问题

问题1: hive 插入数据, insert后报错 FAILED: Execution Error, return code 2 from org.apache.hadoop.hive.ql.exec.mr.MapRedTask

hive 插入数据的时候, 不能直接运行, 报错

问题现象

```
1  hive> insert into students1(id,name,score,classes)
   values(10,'limm10',90,'class1');
2  401282 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN
   org.apache.hadoop.hive.ql.session.SessionState - METASTORE_FILTER_HOOK will
   be ignored, since hive.security.authorization.manager is set to instance of
   HiveAuthorizerFactory.
3  401475 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN
   org.apache.hadoop.hive.metastore.ObjectStore -
   datanucleus.autoStartMechanismMode is set to unsupported value null .
   Setting it to value: ignored
4  401931 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN
   org.apache.hadoop.hive.metastore.ObjectStore -
   datanucleus.autoStartMechanismMode is set to unsupported value null .
   Setting it to value: ignored
5  413409 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN
   org.apache.hadoop.hive.metastore.ObjectStore -
   datanucleus.autoStartMechanismMode is set to unsupported value null .
   Setting it to value: ignored
6  414744 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN
   org.apache.hadoop.hive.ql.Driver - Hive-on-MR is deprecated in Hive 2 and
   may not be available in the future versions. Consider using a different
   execution engine (i.e. spark, tez) or using Hive 1.X releases.
7  Query ID = root_20231007070737_e1b391bf-7584-43fd-a78f-5037068cb77f
8  Total jobs = 3
9  Launching Job 1 out of 3
10 Number of reduce tasks determined at compile time: 1
11 In order to change the average load for a reducer (in bytes):
12   set hive.exec.reducers.bytes.per.reducer=<number>
13 In order to limit the maximum number of reducers:
14   set hive.exec.reducers.max=<number>
15 In order to set a constant number of reducers:
16   set mapreduce.job.reduces=<number>
17 420632 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN
   org.apache.hadoop.mapreduce.JobResourceUploader - Hadoop command-line
   option parsing not performed. Implement the Tool interface and execute your
   application with ToolRunner to remedy this.
18 Starting Job = job_1696262549114_0004, Tracking URL =
   http://master:8088/proxy/application_1696262549114_0004/
19 Kill Command = /usr/local/hadoop-3.3.1/bin/mapred job -kill
   job_1696262549114_0004
20 Hadoop job information for Stage-1: number of mappers: 0; number of
   reducers: 0
21 499586 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN mapreduce.Counters
   - Group org.apache.hadoop.mapred.Task$Counter is deprecated. Use
   org.apache.hadoop.mapreduce.TaskCounter instead
22 2023-10-07 07:09:18,189 Stage-1 map = 0%, reduce = 0%
```

```

23 499621 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN mapreduce.Counters
    - Group org.apache.hadoop.mapred.Task$Counter is deprecated. Use
    org.apache.hadoop.mapreduce.TaskCounter instead
24 Ended Job = job_1696262549114_0004 with errors
25 499674 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] ERROR
    org.apache.hadoop.hive ql.exec.Task - Ended Job = job_1696262549114_0004
    with errors
26 Error during job, obtaining debugging information...
27 499687 [Thread-545] ERROR org.apache.hadoop.hive ql.exec.Task - Error
    during job, obtaining debugging information...
28 FAILED: Execution Error, return code 2 from
    org.apache.hadoop.hive ql.exec.mr.MapRedTask
29 499956 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] ERROR
    org.apache.hadoop.hive ql.Driver - FAILED: Execution Error, return code 2
    from org.apache.hadoop.hive ql.exec.mr.MapRedTask
30 MapReduce Jobs Launched:
31 500076 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN mapreduce.Counters
    - Group FileSystemCounters is deprecated. Use
    org.apache.hadoop.mapreduce.FileSystemCounter instead
32 Stage-Stage-1: HDFS Read: 0 HDFS Write: 0 FAIL
33 Total MapReduce CPU Time Spent: 0 msec

```

错误原因:

namenode内存空间不够, JVM剩余内存空间不够新job运行所致

解决办法:

在hive> 下, 输入

```
1 | set hive.exec.mode.local.auto=true;
```

通过命令修改的话在重启hive后会失效, 要长期有效, 修改文件:

```
1 | vi /usr/local/hive/conf/hive-site.xml
```

修改hive.exec.mode.local.auto的value为true,如下:

```

1 | <property>
2 |   <name>hive.exec.mode.local.auto</name>
3 |   <value>true</value>
4 |   <description>Let Hive determine whether to run in local mode
    automatically</description>
5 | </property>

```

再次验证

再次运行insert语句, 结果如:

```

1 | hive> set hive.exec.mode.local.auto=true;
2 | hive> insert into students1(id,name,score,classes)
    values(10,'limm10',90,'class1');

```



```

3 1351659 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN
org.apache.hadoop.hive.metastore.ObjectStore -
datanucleus.autoStartMechanismMode is set to unsupported value null .
Setting it to value: ignored
4 Automatically selecting local only mode for query
5 1361460 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN
org.apache.hadoop.hive.q1.Driver - Hive-on-MR is deprecated in Hive 2 and
may not be available in the future versions. Consider using a different
execution engine (i.e. spark, tez) or using Hive 1.X releases.
6 Query ID = root_20231007072328_132fb4ba-e5bd-40af-8b4c-84998b9da947
7 Total jobs = 3
8 Launching Job 1 out of 3
9 Number of reduce tasks determined at compile time: 1
10 In order to change the average load for a reducer (in bytes):
11   set hive.exec.reducers.bytes.per.reducer=<number>
12 In order to limit the maximum number of reducers:
13   set hive.exec.reducers.max=<number>
14 In order to set a constant number of reducers:
15   set mapreduce.job.reduces=<number>
16 1363677 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN
org.apache.hadoop.metrics2.impl.MetricsSystemImpl - JobTracker metrics
system already initialized!
17 1363722 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN
org.apache.hadoop.mapreduce.JobResourceUploader - Hadoop command-line
option parsing not performed. Implement the Tool interface and execute your
application with ToolRunner to remedy this.
18 1365196 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN
org.apache.hadoop.fs.FileUtil - Command 'ln -s /usr/local/hadoop-
3.3.1/tmp/mapred/local/job_local1065285107_0001_dc86f705-a9c4-4b66-84e2-
276eb725b6b7/libjars /root/libjars/*' failed 1 with: ln: failed to create
symbolic link '/root/libjars/*': No such file or directory
19
20 1365196 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN
org.apache.hadoop.mapred.LocalDistributedCacheManager - Failed to create
symlink: /usr/local/hadoop-
3.3.1/tmp/mapred/local/job_local1065285107_0001_dc86f705-a9c4-4b66-84e2-
276eb725b6b7/libjars <- /root/libjars/*
21 Job running in-process (local Hadoop)
22 1371900 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN mapreduce.Counters
- Group org.apache.hadoop.mapred.Task$Counter is deprecated. Use
org.apache.hadoop.mapreduce.TaskCounter instead
23 2023-10-07 07:23:50,513 Stage-1 map = 0%, reduce = 0%
24 1373201 [LocalJobRunner Map Task Executor #0] WARN
org.apache.hadoop.hive.serde2.objectinspector.StandardStructObjectInspector
- Trying to access 5 fields inside a list of 6 elements: [Long, 10, 10, 0,
48 4c 4c a0 01 01 83 c5 91 8e 03, null]
25 1373202 [LocalJobRunner Map Task Executor #0] WARN
org.apache.hadoop.hive.serde2.objectinspector.StandardStructObjectInspector
- ignoring similar errors.
26 1373202 [LocalJobRunner Map Task Executor #0] WARN
org.apache.hadoop.hive.serde2.objectinspector.StandardStructObjectInspector
- Trying to access 6 fields inside a list of 7 elements: [String, 6, 6, 1,
0, 48 4c 4c a0 01 01 c1 ae df ef 05, null]

```

```

27 1373202 [LocalJobRunner Map Task Executor #0] WARN
    org.apache.hadoop.hive.serde2.objectinspector.StandardStructObjectInspector
    - ignoring similar errors.
28 1373202 [LocalJobRunner Map Task Executor #0] WARN
    org.apache.hadoop.hive.serde2.objectinspector.StandardStructObjectInspector
    - Trying to access 5 fields inside a list of 6 elements: [Double, 90.0,
    90.0, 0, 48 4c 4c a0 01 01 82 9b 88 e1 06, null]
29 1373202 [LocalJobRunner Map Task Executor #0] WARN
    org.apache.hadoop.hive.serde2.objectinspector.StandardStructObjectInspector
    - ignoring similar errors.
30 1373437 [pool-19-thread-1] WARN
    org.apache.hadoop.metrics2.impl.MetricsSystemImpl - JobTracker metrics
    system already initialized!
31 2023-10-07 07:23:52,581 Stage-1 map = 100%, reduce = 100%
32 Ended Job = job_local1065285107_0001
33 Stage-4 is selected by condition resolver.
34 Stage-3 is filtered out by condition resolver.
35 Stage-5 is filtered out by condition resolver.
36 Moving data to directory
    hdfs://master:9864/user/hive/warehouse/myname.db/students1/.hive-
    staging_hive_2023-10-07_07-23-29_497_8368754191379137565-1/-ext-10000
37 Loading data to table myname.students1
38 1374415 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN
    org.apache.hadoop.hive.metastore.ObjectStore -
    datanucleus.autoStartMechanismMode is set to unsupported value null .
    Setting it to value: ignored
39 1375519 [c1fc6cdd-cfa6-4959-963d-ac46ea866976 main] WARN
    org.apache.hadoop.hive.metastore.ObjectStore -
    datanucleus.autoStartMechanismMode is set to unsupported value null .
    Setting it to value: ignored
40 MapReduce Jobs Launched:
41 Stage-Stage-1: HDFS Read: 0 HDFS Write: 659769448 SUCCESS
42 Total MapReduce CPU Time Spent: 0 msec
43 OK
44 Time taken: 27.088 seconds
45 hive>

```

查看表信息，可查记录数：

```

1  hive> desc formatted students1;
2  OK
3  # col_name          data_type          comment
4  id                  int
5  name                string
6  score               double
7  classes             string
8
9  # Detailed Table Information
10 Database:           myname
11 OwnerType:          USER
12 Owner:              root
13 CreateTime:         Sat Oct 07 04:56:20 EDT 2023
14 LastAccessTime:     UNKNOWN
15 Retention:          0

```

```

16 Location:
   hdfs://master:9864/user/hive/warehouse/myname.db/students1
17 Table Type:          MANAGED_TABLE
18 Table Parameters:
19     numFiles          2
20     spark.sql.create.version    3.1.3
21     spark.sql.sources.schema.numParts    1
22     spark.sql.sources.schema.part.0 {"type":"struct","fields":
   [{"name":"id","type":"integer","nullable":true,"metadata":{}},
   {"name":"name","type":"string","nullable":true,"metadata":{}},
   {"name":"score","type":"double","nullable":true,"metadata":{}},
   {"name":"classes","type":"string","nullable":true,"metadata":
   {}}]}
23     totalSize          44
24     transient_lastDdlTime    1696678532
25
26 # Storage Information
27 SerDe Library:        org.apache.hadoop.hive.serde2.lazy.LazySimpleSerDe
28
29 InputFormat:          org.apache.hadoop.mapred.TextInputFormat
30 OutputFormat:
   org.apache.hadoop.hive ql.io.HiveIgnoreKeyTextOutputFormat
31 Compressed:           No
32 Num Buckets:          -1
33 Bucket Columns:      []
34 Sort Columns:         []
35 Storage Desc Params:
   serialization.format    1
36 Time taken: 0.127 seconds, Fetched: 33 row(s)
37 hive>

```

由结果看出：numFiles：2 表示已录入两条。

版本历史

- Ver1.1-20221007
 - 更新脚本sqlcontent为spark
 - 增加随堂练习
 - 增加实训作业
- Ver1.2-20221011
 - 先更新/usr/local/spark/conf/hive-site.xml为链接方式，后删除了软链接,因为spark 和hive有着天然的集成关系
- Ver1.3-20221015
 - 增加随堂和实训的视频学习
- Ver1.4-20221016
 - 增加理论视频学习
- Ver1.5-20221030
 - 增加理论视频学习2
- Ver1.7-20230828
 - 增加markdown版本

