

Cleaning with PySpark Cheat Sheet

Defining Schema

```
from pyspark.sql.types import *
Schema = StructType([
    StructField('Store', StringType(), nullable=True),
    StructField('Store Type', StringType(), nullable=True),
    StructField('Assessment', StringType(), nullable=True),
    StructField('Competition Distance', FloatType(), nullable=True),
    StructField('Competition Open Since Month', IntegerType(), nullable=True),
    StructField('Competition Open Since Year', IntegerType(), nullable=True),
    StructField('Promo 2', IntegerType(), nullable=True),
    StructField('Promo 2 Since Week', IntegerType(), nullable=True),
    StructField('Promo 2 Since Year', IntegerType(), nullable=True),
    StructField('Promo Interval', StringType(), nullable=True)
])

df = spark.read.option("header", True).schema(Schema).csv('source.csv')
# We can drop invalid rows while reading the dataset by setting the read mode as "DROPMALFORMED"
df_1 = spark.read.option("header", True).option("mode", 'DROPMALFORMED').csv('source.csv')
df.show()
```

Spark does not detect schema itself properly, so we need to define the schema as well for the data set.

PySpark DataTypes

Type	Size (Byte)	Default	Range (Digits)
byte	1	0	3 Ints
short	2	0	5
int	4	0	10
long	8	0	Lots
floats	4	0.0f	Lots floats
double	8	0.0d	Lots
DecimalType	32	0.0	Lots

String Data Types

StringType

Varchar
A variant of StringType which has a length limitation. Data writing will fail if the input string exceeds the length limitation

CharType
Reading column of type CharType -
pe(n) always returns string values of length n. Char type column comparison will pad the short one to the longer length.

Complex Data Types

ArrayType -
nts values comprising a sequence of elements
elementType, containsNull)

MapType -
Represents values comprising a set of key-value pairs. The data type of keys is described by keyType and the data type of values is described by valueType. For a MapType value, keys are not allowed to have null values. valueContainsNull is used to indicate if values of a MapType value can have null values.

StructType
-
Represents values with the structure described by a sequence of Struct Fields

Filtering Data

Adding, renaming and removing columns

(fields)



```
voter_df.filter(voter_df['name'].isNotNull())add-withColumn
  ORvoter_ df.w it hCo lum n(' -
voter_ df.w here(~ voter_ df._ cl.is Null())year',
voter_ df.d at e.year)
voter_ df.f il ter (vo ter _df.da te.yrearnaming> -
withCo lum nRe named
1800)
test_d f.w ith -
voter_ df.w he re( vot er_ df[ '_c 0']Col.co umn- Ren ame
d(' Gen der',
nta ins ('V OTE'))
#Multiple Conditions
drop
whereDF =
flatte nDF.wh ere ((c ol( " voterfir - df.d ro p(' unu sed
_co -
stN ame ") == " xia ngr ui") |
(col("flumn') irs -
tNa me") ==
" mic hae l")).so rt( asc from("la
pyspar-
k.sql import
stN ame "))
functions as F
whereD F.s how (tr unc ate =Fa
lse)
add_n = udf(lambda x, y: x + y,
Intege rTy pe())
#Unique Values
voter_df =
df.sel ect (df ["VOTER
NAME"]).di sti nct()
# Show the rows with 10 highest IDs inand
the we cast the id column to an
Integer type.
voter_ df.o rd erB y(v ote r_d f.R OWdf ID=.
d f-.wit hCo lum n(' id_ -
es c() ).s how(10)
```

User Defined Functions

1. Define a Python method

```
def
revers eSt rin g(m ystr):
return mystr[ ::-1]
```

2. Wrap the function and

```
store as a variable
udfRev ers eString =
udf(re -
ver seS tring,
String Type())
```

3. Use with Spark

```
user_df =
user_d f.w ith Col -
umn ('R eve rse Name',
udfRev ers eSt rin g(u se
r - _df.Name))
```

```
df.cre ate OrR epl ace Tem pVi e
w( " tab le1 ")
```

```
df2 = spark.s ql ("SELECT field1,
field2 FROM table1 ")
```

```
.when(<if condition>, <then x>)
df.sel ect (df.Name,
df.Age, .when( df.Age >= 18,
" Adu lt")
.when( df.Age < 18,
" Min or"))
.other wise() is like else
df.sel ect (df.Name,
df.Age,
```

```
.when( df.Age >= 18,
" Adu lt")
.other wis e("M ino r")
```

Remove duplicate rows & replace values

```
dropDuplicates()
test_d f_n o_dup =
tes t_d f.s ele ct( 'Us
er_ -
ID' ,'G ender',
'Age').dr opD -
upl ica tes()
fillna()
```

used to replace null value with
any other value

```
df.fil lna (va lue ==
9 9,s -
ubset=
```

```
parsed_df =
```

```
spark.read.parquet('parsed_data.parquet')
company_df = spark.r ea d.p arg uet ('c -
omp ani es.p ar quet')
verifi ed_df =
parsed _df.jo in( com - pan y_df,
parsed _df.co mpany ==
compan y_d f.c ompany)
# This automa tically removes any rows
with a company not in the valid_df !
```

View data/a ctions:

```
printS che ma(), head(), show(), count(),
columns and describe()
```

show() - Displa ys/ Prints a number of rows
in a tabular format. By default it displays 20
rows and to change the default number, you
can pass a value to show(n).

where as take(n) returns first n rows as
Array of row objects. It is an alias for
first().

count() - total rows

Published 3rd September, 2022.

Last updated 12th September, 2022.

Page 2 of 3.

Remove duplicate rows & replace values (cont)

```
> ["Promo2_SinceYear", "Promo2_SinceYear"]\
  .show()
Promo2_SinceYear
df.withColumn("grate_right",
  when(df.CombinationDistance == 2000, 1)
  .otherwise(0))\
  .alias('value_desc').show()
```



By **datamansam**

cheatography.com/datamansam/

Published 3rd September, 2022.

Last updated 12th September, 2022.

Page 3 of 3.

Sponsored by **Readable.com**

Measure your website readability!

<https://readable.com>