
INSIGHTFULSAGA

CERTIFICATION PREPARATION SERIES

PYSPARK

CERTIFICATION PREPARATION

Complete Practice Workbook

Practice Tests Comprehensive Question Bank Verified Explanations

"Prepare Smarter. Practice With Purpose."

WORKBOOK SCOPE & COVERAGE

12 Complete Practice Tests covering Questions 1 to 250 (250 total questions)

Core architectural concepts, production scenarios, performance, and security questions

Verified answer keys with concise technical explanations and architectural rationales

Formatted as a high-density, printable technical reference workbook for exam preparation

TABLE OF CONTENTS

Practice Test 01: Foundation Concepts	Questions 1 20 (20 Qs)
Practice Test 02: DataFrame Operations and Transformations	Questions 21 40 (20 Qs)
Practice Test 03: Data Engineer Review Notes	Questions 41 60 (20 Qs)
Practice Test 04: Window Functions and Analytical Processing	Questions 61 80 (20 Qs)
Practice Test 05: Production Project Walkthrough Scenarios	Questions 81 100 (20 Qs)
Practice Test 06: Code Review Scenarios	Questions 101 120 (20 Qs)
Practice Test 07: Find The Problem In The Code	Questions 121 140 (20 Qs)
Practice Test 08: Predict The Output Questions	Questions 141 160 (20 Qs)
Practice Test 09: Complete the Logic / Best Coding Solution	Questions 161 180 (20 Qs)
Practice Test 10: Debug The ETL Pipeline	Questions 181 200 (20 Qs)
Practice Test 11: Spark Optimization Scenarios	Questions 201 220 (20 Qs)
Practice Test 12: Principal Engineer and Architecture Scenarios	Questions 221 250 (30 Qs)

PYSPARK CERTIFICATION PRACTICE TEST 01

Foundation Concepts Questions 1 20 (20 Questions)

QUESTION 1

A newly joined Data Engineer at a retail company is asked to process 2 TB of sales data every night. The company requires distributed processing across multiple machines instead of relying on a single server.

What is the PRIMARY reason PySpark is being used?

- A. Distributed Data Processing
- B. Dashboard Development
- C. Data Warehousing
- D. API Development

ANSWER

A. Distributed Data Processing

EXPLANATION

Explanation

PySpark distributes workload across multiple executors, enabling large-scale data processing.

QUESTION 2

A banking company creates a Spark application.

Which component acts as the central coordinator of the application?

- A. Driver
- B. Executor
- C. Worker

D. Partition

ANSWER

A. Driver

QUESTION 3

A newly hired Data Engineer executes:

```
"""python
df.filter(col("salary") > 50000)
"""
```

No output appears and no job runs.

What Spark concept explains this behavior?

- A. Lazy Evaluation
- B. Broadcast Join
- C. Caching
- D. Checkpointing

ANSWER

A. Lazy Evaluation

QUESTION 4

A healthcare company processes patient records using DataFrames.

Which Spark abstraction is generally preferred for modern ETL development?

- A. DataFrame
- B. RDD Only
- C. Java Collections
- D. CSV Files

ANSWER

A. DataFrame

QUESTION 5

A retail company stores customer transactions in a DataFrame.

The engineer wants to display the first 20 rows.

Which action will trigger execution?

- A. show()
- B. select()
- C. filter()
- D. withColumn()

ANSWER

A. show()

QUESTION 6

A newly joined engineer asks:

"What component performs the actual computation on worker nodes?"

- A. Executor
- B. Driver
- C. SparkSession
- D. DAG

ANSWER

A. Executor

QUESTION 7

A PySpark application contains:

```
"""python
df.filter(col("amount") > 1000)
.select("customer_id")
"""
```

These operations are examples of:

- A. Transformations
- B. Actions
- C. Checkpoints
- D. Partitions

ANSWER

A. Transformations

QUESTION 8

A telecom company executes:

```
""python  
df.count()  
""
```

What type of operation is count()?

- A. Action
- B. Transformation
- C. Partitioning
- D. Caching

ANSWER

A. Action

QUESTION 9

A manufacturing company loads data into Spark.

The dataset is automatically divided into smaller units before processing.

These units are called:

- A. Partitions
- B. Executors
- C. Stages
- D. Queues

ANSWER

A. Partitions

QUESTION 10

A newly hired engineer wants to start working with Spark DataFrames.

Which object is usually created first?

- A. SparkSession
- B. Executor
- C. Partition
- D. Broadcast

ANSWER

A. SparkSession

QUESTION 11

A retail workflow contains multiple transformations.

Spark builds an execution plan before processing begins.

This execution plan is known as:

- A. DAG
- B. Cluster
- C. Cache
- D. Queue

ANSWER

A. DAG

QUESTION 12

A Data Engineer runs:

```
""python  
df.collect()  
""
```

What is the result?

- A. Data is returned to the Driver
- B. Data is repartitioned
- C. Cache is created
- D. DAG is deleted

ANSWER

A. Data is returned to the Driver

QUESTION 13

A banking company processes billions of records.

Why should collect() be used carefully?

- A. It may overload Driver memory
- B. It increases partitions
- C. It creates executors
- D. It improves caching

ANSWER

A. It may overload Driver memory

QUESTION 14

A healthcare company reads a CSV file.

Which command is commonly used?

```
""python  
spark.read.csv(...)  
""
```

This operation creates:

- A. DataFrame
- B. RDD Only
- C. Executor
- D. DAG

ANSWER

A. DataFrame

QUESTION 15

A newly joined engineer wants SQL capabilities on top of Spark DataFrames.

What Spark component provides this foundation?

- A. Spark SQL
- B. Spark Streaming
- C. MLlib
- D. GraphX

ANSWER

A. Spark SQL

QUESTION 16

A company performs:

```
""python  
df.filter(...)  
""
```

followed by

```
""python  
df.show()  
""
```

When does actual execution begin?

- A. During show()
- B. During filter()
- C. During DataFrame creation
- D. During partition creation

ANSWER

A. During show()

QUESTION 17

A logistics company processes petabytes of shipment data.
What key Spark characteristic makes this possible?

- A. Distributed Computing
- B. Local Execution
- C. Single-thread Processing
- D. Manual Scaling

ANSWER

A. Distributed Computing

QUESTION 18

A newly joined engineer hears:
"The Driver creates tasks and sends them to Executors."
Which statement is TRUE?

- A. Driver coordinates execution
- B. Driver stores all source files
- C. Driver replaces executors
- D. Driver performs all processing

ANSWER

A. Driver coordinates execution

QUESTION 19

A PySpark application contains several transformations followed by count().
How many jobs are typically triggered by count()?

- A. One Action-Triggered Job
- B. One Job per Transformation
- C. One Job per Partition
- D. Zero Jobs

ANSWER

A. One Action-Triggered Job

QUESTION 20

A newly hired PySpark Engineer is asked:
"What is the primary benefit of PySpark?"
Choose the BEST answer.

- A. Scalable Distributed Data Processing
- B. Dashboard Reporting
- C. Data Visualization
- D. Office Automation

ANSWER

A. Scalable Distributed Data Processing

EXPLANATION

Explanation

PySpark enables distributed computation across clusters, allowing organizations to process massive datasets efficiently.

title: PySpark Certification Practice Test 02

description: DataFrame Operations and Transformations

PySpark Certification Practice Test 02

PYSPARK CERTIFICATION PRACTICE TEST 02

DataFrame Operations and Transformations Questions 21 40 (20 Questions)

QUESTION 21

A newly joined Data Engineer at an e-commerce company receives a customer dataset containing more than 50 columns.

Business users only need:

- customer_id
- customer_name
- country

The engineer wants to reduce unnecessary data processing as early as possible.

Which PySpark operation should be used?

- A. select()
- B. filter()
- C. distinct()
- D. union()

ANSWER

A. select()

EXPLANATION

Explanation

select() allows engineers to choose only required columns.

QUESTION 22

A retail company wants to identify transactions where:

sales_amount > 10000

Which DataFrame operation is MOST appropriate?

- A. filter()
- B. select()
- C. withColumn()
- D. drop()

ANSWER

A. filter()

QUESTION 23

A banking company stores customer information.

Engineers want to create a new column:

is_premium_customer

based on account balance conditions.

Which operation should be used?

- A. withColumn()
- B. select()
- C. distinct()
- D. limit()

ANSWER

A. withColumn()

QUESTION 24

A healthcare company receives source files containing duplicate patient records. Engineers need unique patient rows before loading downstream systems. Which operation is most appropriate?

- A. distinct()
- B. orderBy()
- C. filter()
- D. cast()

ANSWER

A. distinct()

QUESTION 25

A telecom company wants to sort customers from highest monthly usage to lowest. Which DataFrame operation should be applied?

- A. orderBy()
- B. drop()
- C. filter()
- D. agg()

ANSWER

A. orderBy()

QUESTION 26

A newly joined Engineer notices a column: customer_temp_field is no longer needed. What is the BEST operation?

- A. drop()
- B. select()
- C. filter()
- D. cache()

ANSWER

A. drop()

QUESTION 27

A retail company stores product prices as strings. Engineers must perform mathematical calculations. Which operation should be applied first?

- A. cast()
- B. filter()
- C. distinct()
- D. repartition()

ANSWER

A. cast()

QUESTION 28

An insurance company wants to classify policyholders as:
- High Risk
- Medium Risk
- Low Risk
based on risk scores. Which PySpark function is most suitable?

- A. when()
- B. filter()
- C. drop()
- D. union()

ANSWER**A. when()**

QUESTION 29

A logistics company must return only shipments delivered successfully.
Which operation should be applied?

- A. filter()
- B. drop()
- C. cast()
- D. union()

ANSWER**A. filter()**

QUESTION 30

A company receives customer data with:
First_Name
Business standards require:
first_name
What operation is MOST appropriate?

- A. withColumnRenamed()
- B. filter()
- C. distinct()
- D. count()

ANSWER**A. withColumnRenamed()**

QUESTION 31

A Data Engineer wants to know how many records exist after filtering.
Which operation will trigger execution?

- A. count()
- B. select()
- C. filter()
- D. withColumn()

ANSWER**A. count()**

QUESTION 32

A retail company needs the first 100 records for testing.
Which operation should be considered?

- A. limit()
- B. cast()
- C. union()
- D. repartition()

ANSWER**A. limit()**

QUESTION 33

Two regional sales datasets have identical schemas.
Business users want both datasets combined into one.
Which operation should be used?

- A. union()
- B. distinct()
- C. drop()
- D. cast()

ANSWER

A. union()

QUESTION 34

A healthcare company needs the average patient age.
Which operation category should be used?

- A. Aggregation
- B. Filtering
- C. Ordering
- D. Casting

ANSWER

A. Aggregation

QUESTION 35

A banking company wants to determine the number of unique customers.
What combination is MOST appropriate?

- A. distinct() + count()
- B. filter() + show()
- C. cast() + union()
- D. drop() + limit()

ANSWER

A. distinct() + count()

QUESTION 36

A manufacturing company loads machine readings.
Engineers create five new columns using withColumn().
When are these transformations actually executed?

- A. During an Action
- B. Immediately
- C. During DataFrame Creation
- D. During Executor Startup

ANSWER

A. During an Action

EXPLANATION

Explanation
PySpark transformations are lazily evaluated.

QUESTION 37

A retail analyst requests all customer records where:
Country = 'India'
AND
Sales > 50000
Which operation should be applied?

- A. filter()
- B. drop()
- C. cast()
- D. union()

ANSWER

A. filter()

QUESTION 38

A telecom company wants to replace NULL values in a usage column with 0.
Which operation is most appropriate?

- A. fillNa()
- B. filter()
- C. distinct()

D. orderBy()

ANSWER

A. fillna()

QUESTION 39

A newly joined engineer notices some columns contain unexpected leading and trailing spaces. Which function should be evaluated?

- A. trim()
- B. union()
- C. repartition()
- D. collect()

ANSWER

A. trim()

QUESTION 40

A global retailer processes billions of records daily.

Management asks:

"Why are operations like select(), filter(), and withColumn() called transformations?"

Choose the BEST answer.

- A. They define data processing logic without immediately executing it
- B. They trigger jobs immediately
- C. They create executors
- D. They increase cluster size

ANSWER

A. They define data processing logic without immediately executing it

EXPLANATION

Explanation

Transformations build the logical execution plan. Spark executes them only when an action such as count(), show(), or collect() is triggered.

title: PySpark Certification Practice Test 03

description: Data Engineer Review Notes

PySpark Certification Practice Test 03

PYSPARK CERTIFICATION PRACTICE TEST 03

Data Engineer Review Notes Questions 41 60 (20 Questions)

QUESTION 41

A newly joined Data Engineer reviews the following requirement:

"We have customer data and customer address data.

Only customers existing in both datasets should appear in the output."

Which join type should be recommended?

- A. Inner Join
- B. Left Join
- C. Right Join
- D. Full Join

ANSWER

A. Inner Join

EXPLANATION

Review Note

Inner Join returns only matching records from both datasets.

QUESTION 42

Review Requirement:

"All customers must appear in the output.

Address information should appear only when available."

Which join should be chosen?

- A. Inner Join
- B. Left Join
- C. Right Join
- D. Cross Join

ANSWER

B. Left Join

QUESTION 43

Review Requirement:

"Return all records from both systems even when no matching key exists."

Which join is appropriate?

- A. Inner Join
- B. Left Join
- C. Full Outer Join
- D. Semi Join

ANSWER

C. Full Outer Join

QUESTION 44

A retail company wants total sales amount by store.

Which operation should be performed first?

- A. groupBy()
- B. filter()
- C. drop()
- D. cast()

ANSWER

A. groupBy()

QUESTION 45

Review Requirement:

"Calculate total revenue for each product."

Which aggregation function is most appropriate?

- A. sum()
- B. avg()
- C. min()
- D. max()

ANSWER

A. sum()

QUESTION 46

A banking company wants average account balance by branch.

Which aggregation should be recommended?

- A. avg()
- B. sum()
- C. count()
- D. max()

ANSWER

A. avg()

QUESTION 47

Review Requirement:

"Determine how many transactions occurred for each customer."

Which aggregation fits best?

- A. count()
- B. avg()
- C. max()
- D. first()

ANSWER

A. count()

QUESTION 48

A healthcare company wants to know the highest insurance claim amount submitted by each patient.

Which function should be used?

- A. max()
- B. min()
- C. avg()
- D. count()

ANSWER

A. max()

QUESTION 49

A telecom company wants the lowest monthly usage recorded for each customer.

Which aggregate is appropriate?

- A. min()
- B. max()
- C. sum()
- D. count()

ANSWER

A. min()

QUESTION 50

Review Requirement:

"Count distinct customers who placed orders."

Which solution best fits?

- A. countDistinct()
- B. count()
- C. avg()
- D. sum()

ANSWER

A. countDistinct()

QUESTION 51

A retail company has:

sales_df

customers_df

The review note states:

"Keep only customer records that have sales."

Which join type should be evaluated?

- A. Left Semi Join
- B. Full Join
- C. Cross Join
- D. Right Join

ANSWER

A. Left Semi Join

EXPLANATION

Review Note

Semi Join returns rows from the left table when a match exists.

QUESTION 52

Review Requirement:

"Identify customers that do NOT have any orders."

Which join is most suitable?

- A. Left Anti Join
- B. Inner Join
- C. Right Join
- D. Full Join

ANSWER

A. Left Anti Join

QUESTION 53

A Data Engineer sees:

```
""python
```

```
df.groupBy("department").agg(sum("salary"))
```

```
""
```

What business question is being answered?

- A. Total salary by department
- B. Employee count
- C. Maximum salary
- D. Duplicate records

ANSWER

A. Total salary by department

QUESTION 54

Review Requirement:

"Find average monthly sales by region."

Which operation sequence is most appropriate?

- A. groupBy() + avg()
- B. distinct() + count()
- C. filter() + orderBy()
- D. drop() + cast()

ANSWER

A. groupBy() + avg()

QUESTION 55

A logistics company wants shipment statistics by country.

Multiple metrics are required:

- Total Shipments
- Average Cost
- Maximum Cost

Which feature should be used?

- A. agg()
- B. distinct()
- C. fillna()
- D. union()

ANSWER

A. agg()

QUESTION 56

A newly hired engineer joins two datasets.

Output row count is much larger than expected.

The review document shows:

No join condition specified.

What is the MOST likely issue?

- A. Cross Join
- B. Broadcast Join
- C. Left Join
- D. Semi Join

ANSWER

A. Cross Join

QUESTION 57

Review Requirement:

"Display top-selling stores first."

Which operation should be added after aggregation?

- A. orderBy(desc())
- B. drop()
- C. cast()
- D. fillna()

ANSWER

A. orderBy(desc())

QUESTION 58

A manufacturing company wants total machine downtime by plant.

Which operation sequence is most appropriate?

- A. groupBy() + sum()
- B. filter() + distinct()
- C. cast() + union()
- D. drop() + orderBy()

ANSWER

A. groupBy() + sum()

QUESTION 59

A banking company aggregates billions of transaction records.

Only 300 branch-level records are returned.

What explains this behavior?

- A. Aggregation reduces data volume
- B. Data was deleted
- C. Partitions were removed
- D. Executors failed

ANSWER

A. Aggregation reduces data volume

QUESTION 60

A Senior Data Engineer reviews an ETL design:

1. Join customer data
2. Join transaction data
3. Aggregate spending
4. Publish customer metrics

What is the PRIMARY business purpose of these operations?

- A. Convert detailed records into meaningful business insights
- B. Increase partitions
- C. Create executors
- D. Reduce cluster memory

ANSWER**A. Convert detailed records into meaningful business insights****EXPLANATION**

Review Note

Joins combine related datasets, while aggregations summarize data into business-friendly metrics.

title: PySpark Certification Practice Test 04

description: Window Functions and Analytical Processing

PySpark Certification Practice Test 04

PYSPARK CERTIFICATION PRACTICE TEST 04

Window Functions and Analytical Processing Questions 61 80 (20 Questions)

QUESTION 61

An interviewer says:

"A retail company wants the latest order for every customer. The source table contains multiple orders per customer."

Which function would most likely be used?

- A. row_number()
- B. count()
- C. distinct()
- D. sum()

ANSWER**A. row_number()****EXPLANATION**

Explanation

row_number() is commonly used with partitionBy(customer_id) and ordered by order_date desc to identify the latest record.

QUESTION 62

An interviewer observes:

"The finance team wants employees ranked by salary within each department. If two employees have the same salary, they should receive the same rank, and the next rank should be skipped."

Which function matches this requirement?

- A. rank()
- B. dense_rank()
- C. row_number()
- D. lag()

ANSWER**A. rank()**

QUESTION 63

Observation:

"A telecom company wants rankings without gaps even when values tie."

Which function should be recommended?

- A. dense_rank()
- B. rank()
- C. row_number()
- D. count()

ANSWER**A. dense_rank()**

QUESTION 64

Observation:

"A healthcare analyst wants to compare today's patient count with yesterday's patient count."

Which function is most appropriate?

- A. lag()
- B. lead()
- C. rank()
- D. distinct()

ANSWER

A. lag()

QUESTION 65

Observation:

"A retail company wants to view the next scheduled delivery date for each shipment."

Which function should be considered?

- A. lead()
- B. lag()
- C. row_number()
- D. avg()

ANSWER

A. lead()

QUESTION 66

An interviewer asks:

"Which function would you use to find the most recent transaction per customer?"

- A. row_number()
- B. dense_rank()
- C. sum()
- D. max()

ANSWER

A. row_number()

QUESTION 67

Observation:

"A bank wants a running account balance after every transaction."

Which window function pattern is typically used?

- A. sum() over Window
- B. distinct()
- C. dropDuplicates()
- D. collect()

ANSWER

A. sum() over Window

QUESTION 68

Observation:

"A sales manager wants cumulative sales as each day passes."

Which approach should be used?

- A. Running Total Window
- B. Left Join
- C. Union
- D. Distinct

ANSWER

A. Running Total Window

QUESTION 69

Observation:

"A customer can have multiple records. We only want the newest record."

Which pattern is most commonly used?

- A. row_number() = 1
- B. count() = 1
- C. distinct()
- D. collect()

ANSWER

A. row_number() = 1

QUESTION 70

Observation:

"A company wants the Top 3 products by revenue within each country."

Which capability best supports this?

- A. Window + rank()
- B. dropDuplicates()
- C. fillna()
- D. union()

ANSWER

A. Window + rank()

QUESTION 71

An interviewer reviews a solution:

```
""python
```

```
Window.partitionBy("department")
```

```
""
```

What business requirement is most likely being addressed?

- A. Ranking within groups
- B. Reading CSV files
- C. Creating partitions in storage
- D. Broadcasting data

ANSWER

A. Ranking within groups

QUESTION 72

Observation:

"A logistics company wants to compare current shipment cost with previous shipment cost."

Which function is most suitable?

- A. lag()
- B. lead()
- C. count()
- D. max()

ANSWER

A. lag()

QUESTION 73

Observation:

"A manufacturing company wants the next machine maintenance date for every machine."

Which function would help?

- A. lead()
- B. lag()
- C. sum()
- D. avg()

ANSWER

A. lead()

QUESTION 74

Observation:

"A company wants to identify the first transaction for every customer."

Which function is commonly used?

- A. row_number()
- B. union()
- C. distinct()
- D. collect()

ANSWER

A. row_number()

QUESTION 75

Observation:

"A retailer wants to know how much sales changed compared to the previous day."

Which function is most likely used?

- A. lag()
- B. dense_rank()
- C. countDistinct()
- D. cache()

ANSWER

A. lag()

QUESTION 76

Observation:

"Employees must be ranked by performance score. Employees with equal scores should receive consecutive rankings without gaps."

Which function fits best?

- A. dense_rank()
- B. rank()
- C. lead()
- D. lag()

ANSWER

A. dense_rank()

QUESTION 77

Observation:

"A fraud analytics team wants the previous transaction amount available in every record."

Which function is appropriate?

- A. lag()
- B. lead()
- C. row_number()
- D. distinct()

ANSWER

A. lag()

QUESTION 78

Observation:

"A business wants the next upcoming customer renewal date visible in the current row."

Which function should be applied?

- A. lead()
- B. lag()
- C. rank()
- D. sum()

ANSWER**A. lead()**

QUESTION 79

Observation:

"A company wants the Top 5 customers by yearly spending."

Which window-based approach is most common?

- A. rank()/dense_rank() + filter
- B. union()
- C. fillna()
- D. collect()

ANSWER**A. rank()/dense_rank() + filter**

QUESTION 80

A Senior PySpark Interviewer makes the following observation:

"Most real-world analytics problems involving latest records, ranking, comparisons, trends, Top-N reporting, and cumulative calculations are solved using Window Functions."

How should this statement be evaluated?

- A. Correct
- B. Incorrect

ANSWER**A. Correct****EXPLANATION**

Explanation

Window Functions are among the most important PySpark features for analytical workloads and are heavily used in enterprise reporting, customer analytics, finance, telecom, retail, and healthcare projects.

title: PySpark Certification Practice Test 05

description: Production Project Walkthrough Scenarios

PySpark Certification Practice Test 05

PYSPARK CERTIFICATION PRACTICE TEST 05Production Project Walkthrough Scenarios Questions 81 100 (20 Questions)

QUESTION 81

A newly joined Data Engineer at a retail company receives daily sales files from 2,000 stores.

The data contains duplicate transactions because some stores resend files after network failures.

Before loading into the Gold layer, management wants only unique transactions.

What should be implemented first?

- A. Deduplication Logic
- B. Broadcast Join
- C. Repartition
- D. Cache

ANSWER**A. Deduplication Logic**

QUESTION 82

A banking company receives customer account updates every day.

Business users must see both:

- Current Account Details
- Historical Account Changes

Which design pattern should be implemented?

- A. SCD Type 2
- B. SCD Type 1
- C. Full Refresh
- D. Truncate and Load

ANSWER

A. SCD Type 2

QUESTION 83

A healthcare company processes 500 million claim records daily.

Only 10,000 records change each day.

Management wants to reduce processing time.

What is the BEST recommendation?

- A. Incremental Processing
- B. Full Reload
- C. Collect Dataset
- D. Single Partition

ANSWER

A. Incremental Processing

QUESTION 84

A telecom company joins a 2 TB transaction table with a 5 MB reference table.

Engineers report poor performance.

Which optimization should be considered?

- A. Broadcast Join
- B. Cross Join
- C. Collect()
- D. Union

ANSWER

A. Broadcast Join

QUESTION 85

A retail organization reports sales metrics taking 40 minutes to generate.

The same DataFrame is reused by multiple downstream transformations.

What should engineers evaluate?

- A. Cache/Persist
- B. Collect
- C. DropDuplicates
- D. Coalesce(1)

ANSWER

A. Cache/Persist

QUESTION 86

A manufacturing company processes machine sensor data.

Business users want alerts whenever temperature exceeds safety limits.

Which transformation approach best supports this?

- A. withColumn() + when()
- B. collect()
- C. distinct()
- D. union()

ANSWER

A. withColumn() + when()

QUESTION 87

A retail company stores customer information in a Bronze table.

The Silver layer requires:

- Valid customer IDs
- No duplicates
- Standardized formats

What is the PRIMARY goal of the Silver layer?

- A. Data Cleansing and Validation
- B. Dashboard Reporting
- C. Archival Storage
- D. Cluster Management

ANSWER

A. Data Cleansing and Validation

QUESTION 88

An insurance company notices some records are rejected because mandatory columns contain NULL values.

What should be introduced?

- A. Data Quality Validation
- B. Cross Join
- C. Broadcast Hint
- D. Collect

ANSWER

A. Data Quality Validation

QUESTION 89

A banking workflow identifies customers with unusually large transactions.

The business wants:

- Normal
- High Value
- Risk

categories.

Which PySpark function is commonly used?

- A. when()
- B. union()
- C. repartition()
- D. collect()

ANSWER

A. when()

QUESTION 90

A healthcare company loads data from multiple hospitals.

Schema differences occasionally occur.

Engineering leadership wants early detection of schema changes.

What should be implemented?

- A. Schema Validation Checks
- B. More Executors
- C. More Partitions
- D. Coalesce

ANSWER

A. Schema Validation Checks

QUESTION 91

A logistics company wants to identify records arriving multiple times.
Which operation is most commonly used?

- A. dropDuplicates()
- B. count()
- C. union()
- D. collect()

ANSWER

A. dropDuplicates()

QUESTION 92

A retail company creates Gold-level sales metrics.
Executives consume the results in Power BI.
What is the PRIMARY purpose of the Gold layer?

- A. Business Reporting
- B. Raw Storage
- C. Landing Zone
- D. Temporary Processing

ANSWER

A. Business Reporting

QUESTION 93

A telecom company receives late-arriving records.
Business users still want historical accuracy.
Which data engineering concept becomes important?

- A. Late Arriving Data Handling
- B. Cache Management
- C. Driver Scaling
- D. Repartitioning Only

ANSWER

A. Late Arriving Data Handling

QUESTION 94

A financial institution audits an ETL pipeline.
Auditors request:
- Records Processed
- Start Time
- End Time
- Status

What should the engineering team maintain?

- A. Audit Logging
- B. Broadcast Tables
- C. Caching
- D. Union Logic

ANSWER

A. Audit Logging

QUESTION 95

A large retailer loads daily inventory files.
One corrupted file causes the entire workflow to fail.
Management wants resilience.
What should be added?

- A. Error Handling Framework
- B. More Workers
- C. More Schemas
- D. collect()

ANSWER**A. Error Handling Framework**

QUESTION 96

A banking company tracks customer status changes.

Customers move through:

- Prospect
- Active
- Premium

The business wants to analyze transitions over time.

Which design is most suitable?

- A. Historical Tracking (SCD Type 2)
- B. Full Refresh
- C. Distinct Only
- D. Cache

ANSWER**A. Historical Tracking (SCD Type 2)**

QUESTION 97

A healthcare ETL pipeline must recover automatically after temporary source outages.

What production feature is most important?

- A. Retry Strategy
- B. More Columns
- C. More Aggregations
- D. More Unions

ANSWER**A. Retry Strategy**

QUESTION 98

A telecom company discovers different pipelines calculate customer revenue differently.

Executives lose trust in reports.

What should be standardized?

- A. Business Rules
- B. Executor Count
- C. Partition Count
- D. Cache Usage

ANSWER**A. Business Rules**

QUESTION 99

A newly joined Senior Data Engineer reviews a PySpark ETL flow.

Requirements include:

- Scalability
- Reusability
- Monitoring
- Data Quality

What architecture characteristic is being emphasized?

- A. Production Readiness
- B. Local Processing
- C. Manual Operations
- D. Single-Node Design

ANSWER**A. Production Readiness**

QUESTION 100

A Principal Data Engineer is asked:

"What is the primary objective of a PySpark-based enterprise ETL platform?"

- A. Convert raw large-scale data into trusted business-ready information
- B. Create maximum partitions
- C. Use the largest cluster possible
- D. Generate the longest code

ANSWER

A. Convert raw large-scale data into trusted business-ready information

EXPLANATION

Explanation

The goal of enterprise PySpark platforms is not just processing data, but delivering reliable, scalable, high-quality business information that supports decision-making.

title: PySpark Certification Practice Test 06

description: Code Review Scenarios

PySpark Certification Practice Test 06

PYSPARK CERTIFICATION PRACTICE TEST 06

Code Review Scenarios Questions 101 120 (20 Questions)

QUESTION 101

A developer submits:

```
"""python
large_df.collect()
"""
```

The dataset contains 500 million records.

What is the BEST review comment?

- A. Avoid collect() on very large datasets because it may crash the Driver
- B. Use collect() more frequently
- C. Increase the number of columns
- D. Use additional joins

ANSWER

A

QUESTION 102

Code Review Note:

A 2 TB transaction table is joined with a 20 MB lookup table.

The developer uses a regular join.

What is the BEST review recommendation?

- A. Consider a Broadcast Join
- B. Use collect()
- C. Convert to CSV first
- D. Repartition to one partition

ANSWER

A

QUESTION 103

A developer writes:

```
"""python
df.repartition(1)
"""
```

before writing a 1 TB dataset.

What is the BEST review comment?

- A. This may create a performance bottleneck
- B. This improves parallelism
- C. This eliminates shuffles
- D. This improves scalability

ANSWER

A

QUESTION 104

A developer applies the same expensive transformation chain three times.

What is the BEST recommendation?

- A. Consider caching the intermediate DataFrame
- B. Add more columns
- C. Use collect()
- D. Create a Cross Join

ANSWER

A

QUESTION 105

A code review finds:

```
"""python
df.filter(...)
.show()
.show()
.show()
"""
```

What is the BEST observation?

- A. Multiple actions may trigger repeated computation
- B. show() improves performance
- C. show() creates partitions
- D. show() removes shuffles

ANSWER

A

QUESTION 106

A developer uses:

```
"""python
df.withColumn(...)
.withColumn(...)
.withColumn(...)
"""
```

25 times.

What is the BEST review feedback?

- A. Large chains may impact readability and maintainability
- B. Spark cannot support withColumn()
- C. withColumn always causes failure
- D. Remove all transformations

ANSWER

A

QUESTION 107

A code review shows:

```
"""python
df.count()
"""
```

used repeatedly for debugging.

What is the BEST recommendation?

- A. Minimize unnecessary actions
- B. Count more frequently
- C. Replace with repartition()
- D. Convert to RDD

ANSWER

A

QUESTION 108

A developer joins customer and transaction data.

Duplicate customer records appear.

What is the BEST review question?

- A. Was the join key unique?
- B. Add additional columns
- C. Increase partitions
- D. Create another join

ANSWER

A

QUESTION 109

A developer hardcodes:

```
""python  
"/mnt/dev/customer/"  
""
```

inside ETL logic.

What is the BEST review feedback?

- A. Externalize environment-specific configuration
- B. Add more hardcoded paths
- C. Convert to broadcast
- D. Use collect()

ANSWER

A

QUESTION 110

A review identifies no validation checks before writing Gold tables.

What is the BEST recommendation?

- A. Add Data Quality Validation
- B. Remove Aggregations
- C. Use Cross Join
- D. Increase Cluster Size

ANSWER

A

QUESTION 111

A developer uses:

```
""python  
dropDuplicates()  
""
```

without specifying business keys.

What is the BEST review comment?

- A. Confirm deduplication business logic
- B. Use more partitions
- C. Convert to JSON
- D. Add collect()

ANSWER

A

QUESTION 112

A code review reveals transformations are repeated across many notebooks.
What is the BEST engineering recommendation?

- A. Create reusable functions/modules
- B. Duplicate code further
- C. Increase workers
- D. Use more joins

ANSWER**A**

QUESTION 113

A developer writes:

```
"""python  
df.write.mode("overwrite")  
"""
```

for a production customer dimension.

What is the BEST review question?

- A. Are we intentionally replacing all existing data?
- B. Increase executors
- C. Add more partitions
- D. Use collect()

ANSWER**A**

QUESTION 114

A review identifies no audit columns such as:

- created_date
- updated_date
- load_timestamp

What is the BEST recommendation?

- A. Add audit metadata
- B. Remove timestamps entirely
- C. Add more joins
- D. Increase partitions

ANSWER**A**

QUESTION 115

A developer processes 800 million records using a Python UDF.

Performance is very poor.

What is the BEST review recommendation?

- A. Prefer built-in Spark functions when possible
- B. Use more UDFs
- C. Convert to collect()
- D. Add more show()

ANSWER**A**

QUESTION 116

A review shows:

```
"""python  
df.cache()  
"""
```

but the DataFrame is only used once.

What is the BEST comment?

- A. Cache may be unnecessary
- B. Cache everything
- C. Increase memory only
- D. Convert to RDD

ANSWER**A**

QUESTION 117

A developer handles failed records by dropping them silently.

What is the BEST review feedback?

- A. Implement error tracking and logging
- B. Ignore bad data
- C. Remove validations
- D. Add more partitions

ANSWER**A**

QUESTION 118

A code review reveals there is no retry or failure handling in a production ingestion process.

What is the BEST recommendation?

- A. Improve resiliency and recovery logic
- B. Disable monitoring
- C. Use collect()
- D. Create additional joins

ANSWER**A**

QUESTION 119

A developer created an ETL process that works for current data volume but may struggle if volume grows 10x.

What is the BEST review concern?

- A. Scalability
- B. Variable names
- C. Column ordering
- D. Notebook color theme

ANSWER**A**

QUESTION 120

A Principal Data Engineer reviews a PySpark project.

Which review comment provides the MOST enterprise value?

- A. Optimize for reliability, scalability, maintainability, and data quality
- B. Maximize line count
- C. Increase notebook cells
- D. Add more comments than logic

ANSWER**A****EXPLANATION**

Explanation

Enterprise PySpark development is about building scalable, reliable, maintainable, and high-quality data systems rather than just making code run successfully.

title: PySpark Certification Practice Test 07

description: Find The Problem In The Code

PySpark Certification Practice Test 07

PYSPARK CERTIFICATION PRACTICE TEST 07Find The Problem In The Code Questions 121 140 (20 Questions)

QUESTION 121

A developer writes:

```
"""python
df = spark.read.parquet("/sales")
result = df.collect()
print(len(result))
"""
```

The dataset contains 2 billion records.

What is the biggest problem?

- A. Driver Out Of Memory Risk
- B. Missing Cache
- C. Missing Broadcast
- D. Missing Repartition

ANSWER**A. Driver Out Of Memory Risk****EXPLANATION**

```
### Explanation
collect() moves all data to the Driver.
```

QUESTION 122

A developer writes:

```
"""python
large_df.join(small_df, "customer_id")
"""
```

Fact Table = 5 TB

Lookup Table = 10 MB

What optimization opportunity is missing?

- A. Broadcast Join
- B. Cross Join
- C. Full Join
- D. Union

ANSWER**A. Broadcast Join**

QUESTION 123

Review the code:

```
"""python
df.repartition(1)
.write
.parquet("/output")
"""
```

Output size = 800 GB

What is the main concern?

- A. Single Partition Bottleneck
- B. Duplicate Records
- C. Missing Aggregation
- D. Missing Join

ANSWER**A. Single Partition Bottleneck**

QUESTION 124

A developer writes:

```
"""python
for i in range(100):
print(df.count())
"""
```

What is the problem?

- A. Repeated Expensive Actions
- B. Missing Partition
- C. Missing Join
- D. Wrong Schema

ANSWER

A

QUESTION 125

Review:

```
"""python
df.show()
df.show()
df.show()
"""
```

What is the concern?

- A. Multiple Job Executions
- B. Driver Failure
- C. Schema Drift
- D. Partition Skew

ANSWER

A

QUESTION 126

A Data Engineer writes:

```
"""python
df1.crossJoin(df2)
"""
```

Both tables contain 50 million rows.

What is the likely issue?

- A. Massive Data Explosion
- B. Missing Cache
- C. Missing Cast
- D. Missing Filter

ANSWER

A

QUESTION 127

Review the code:

```
"""python
df.withColumn("new_col1", ...)
.withColumn("new_col2", ...)
.withColumn("new_col3", ...)
...
"""
```

More than 60 withColumn statements exist.

What is the primary concern?

- A. Maintainability and Optimization Complexity
- B. Duplicate Records
- C. Broadcast Failure

D. Shuffle Elimination

ANSWER

A

QUESTION 128

Developer writes:

```
""python
df.cache()
""
```

The DataFrame is used only once.

What issue exists?

- A. Unnecessary Cache Usage
- B. Missing Partition
- C. Missing Join
- D. Wrong Window Function

ANSWER

A

QUESTION 129

Review:

```
""python
df.write.mode("overwrite")
""
```

Target table = Production Customer Table

What question should be asked?

- A. Are we intentionally replacing all existing data?
- B. Is Spark Installed?
- C. How Many Columns Exist?
- D. What Is The Cluster Name?

ANSWER

A

QUESTION 130

A developer creates:

```
""python
df.dropDuplicates()
""
```

Business key is:

(customer_id, order_id)

What concern exists?

- A. Deduplication Criteria Not Verified
- B. Missing Cache
- C. Missing Broadcast
- D. Missing Count

ANSWER

A

QUESTION 131

Review:

```
""python
customers.join(
transactions,
customers.id == transactions.id
)
""
```

Output records are unexpectedly higher.

Most likely cause?

- A. Non-Unique Join Keys
- B. Missing Cast
- C. Executor Failure
- D. Missing Repartition

ANSWER

A

QUESTION 132

Developer writes:

```
""python
df.filter(col("amount") > 100)
""
```

but never stores the result.

What is the issue?

- A. Transformation Result Discarded
- B. Partition Skew
- C. Driver Overflow
- D. Broadcast Issue

ANSWER

A

QUESTION 133

Review:

```
""python
df.collect()
for row in rows:
    process(row)
""
```

Used for 500 million rows.

What is the main problem?

- A. Processing Distributed Data On Driver
- B. Missing Window Function
- C. Missing Aggregation
- D. Missing Sort

ANSWER

A

QUESTION 134

Developer writes:

```
""python
from pyspark.sql.functions import udf
""
```

A UDF is applied to 2 billion records even though Spark built-in functions exist.
Concern?

- A. Performance Degradation
- B. Duplicate Records
- C. Missing Cache
- D. Missing Cast

ANSWER

A

QUESTION 135

Review:

```
"""python  
df.coalesce(1)  
"""
```

File size = 1.5 TB.

Issue?

- A. Excessive Single Partition Processing
- B. Missing Aggregation
- C. Missing Join
- D. Missing Window

ANSWER

A

QUESTION 136

Developer performs:

```
"""python  
large_df.join(large_df2)  
"""
```

No filters.

No partitions.

No broadcast.

Performance is terrible.

What should be investigated first?

- A. Shuffle Cost
- B. Variable Names
- C. Column Order
- D. Notebook Layout

ANSWER

A

QUESTION 137

Review:

```
"""python  
df.write.mode("append")  
"""
```

Business users report duplicate records daily.

Most likely concern?

- A. Duplicate Load Risk
- B. Cache Failure
- C. Wrong Partition Count
- D. Window Function Issue

ANSWER

A

QUESTION 138

Developer writes:

```
"""python  
df.count()  
"""
```

solely for troubleshooting.

Table contains 5 billion records.

Concern?

- A. Expensive Full Dataset Scan
- B. Missing Schema
- C. Missing Window

D. Missing Cache

ANSWER

A

QUESTION 139

Review:

```
"""python
df.orderBy("customer_name")
"""
```

Table has 3 billion rows.

What operation should engineers expect?

- A. Large Shuffle Operation
- B. Broadcast Join
- C. Data Loss
- D. Duplicate Records

ANSWER

A

QUESTION 140

A Principal Data Engineer reviews a notebook containing:

- collect()
- repartition(1)
- repeated count()
- unnecessary cache()
- Python UDFs

What is the overall feedback?

- A. Significant Performance and Scalability Risks
- B. Excellent Production Design
- C. Perfect Optimization
- D. No Improvements Needed

ANSWER

A

EXPLANATION

Explanation

The code may work on small datasets but will struggle significantly in production-scale environments.

title: PySpark Certification Practice Test 08

description: Predict The Output Questions

PySpark Certification Practice Test 08

PYSPARK CERTIFICATION PRACTICE TEST 08

Predict The Output Questions Questions 141 160 (20 Questions)

QUESTION 141

Input Data:

```
| id |
```

```
|----|
```

```
| 1 |
```

```
| 2 |
```

```
| 3 |
```

Code:

```
"""python
df.count()
```

""

What will be returned?

- A. 3
- B. 2
- C. 1
- D. Error

ANSWER

A. 3

QUESTION 142

Input Data:

```
| country |  
|-----|  
| India |  
| India |  
| USA |
```

Code:

```
""python  
df.select("country").distinct().count()  
""
```

Output?

- A. 1
- B. 2
- C. 3
- D. 4

ANSWER

B. 2

QUESTION 143

Input Data:

```
| id | salary |  
|---|-----|  
| 1 | 1000 |  
| 2 | 5000 |  
| 3 | 2000 |
```

Code:

```
""python  
df.filter(col("salary") > 2000).count()  
""
```

Output?

- A. 1
- B. 2
- C. 3
- D. 0

ANSWER

A. 1

QUESTION 144

Input Data:

```
| id |  
|---|  
| 1 |  
| 1 |  
| 2 |
```

Code:

```
""python  
df.dropDuplicates().count()
```

""

Output?

- A. 2
- B. 3
- C. 1
- D. Error

ANSWER

A. 2

QUESTION 145

Input Data:

| amount |

|-----|

| 100 |

| 200 |

| 300 |

Code:

""python

df.agg(sum("amount"))

""

Result?

- A. 100
- B. 200
- C. 600
- D. 300

ANSWER

C. 600

QUESTION 146

Input Data:

| amount |

|-----|

| 100 |

| 200 |

| 300 |

Code:

""python

df.agg(avg("amount"))

""

Result?

- A. 100
- B. 200
- C. 300
- D. 600

ANSWER

B. 200

QUESTION 147

Input Data:

| id |

|----|

| 1 |

| 2 |

| 3 |

Code:

""python

df.limit(2).count()

""

Output?

- A. 3
- B. 2
- C. 1
- D. Error

ANSWER

B. 2

QUESTION 148

Input Data:

| value |

|-----|

| NULL |

| 100 |

Code:

```
""python
```

```
df.fillna(0)
```

""

NULL value becomes?

- A. NULL
- B. 100
- C. 0
- D. Empty String

ANSWER

C. 0

QUESTION 149

Input Data:

sales_df

| id |

|---|

| 1 |

| 2 |

customer_df

| id |

|---|

| 2 |

| 3 |

Inner Join Row Count?

- A. 0
- B. 1
- C. 2
- D. 4

ANSWER

B. 1

QUESTION 150

Same Data

Left Join Row Count?

- A. 1
- B. 2
- C. 3
- D. 4

ANSWER

B. 2

QUESTION 151

Input Data:

```
| salary |  
|-----|  
| 1000 |  
| 1000 |  
| 2000 |
```

Code:

```
""python  
rank()  
""
```

Ordered by salary descending.

Ranks?

- A. 1,1,2
- B. 1,1,3
- C. 2,2,3
- D. 1,2,3

ANSWER

B. 1,1,3

QUESTION 152

Using same input:

```
""python  
dense_rank()  
""
```

Output?

- A. 1,1,2
- B. 1,1,3
- C. 2,2,3
- D. 1,2,3

ANSWER

A. 1,1,2

QUESTION 153

Input Data:

```
| salary |  
|-----|  
| 1000 |  
| 2000 |  
| 3000 |
```

Code:

```
""python  
lag("salary")  
""
```

For salary = 2000

Returned value?

- A. 1000
- B. 2000
- C. 3000
- D. NULL

ANSWER

A. 1000

QUESTION 154

Input Data:

| salary |

|-----|

| 1000 |

| 2000 |

| 3000 |

Code:

""python

lead("salary")

""

For salary = 2000

Returned value?

A. 1000

B. 2000

C. 3000

D. NULL

ANSWER**C. 3000**

QUESTION 155

Input Data:

| value |

|-----|

| A |

| B |

| C |

Code:

""python

row_number()

""

Order By value

Output?

A. 1,2,3

B. 0,1,2

C. 1,1,1

D. 3,2,1

ANSWER**A. 1,2,3**

QUESTION 156

Input Data:

| amount |

|-----|

| 100 |

| 200 |

| 300 |

Running Total Output?

A. 100,300,600

B. 100,200,300

C. 600,600,600

D. 100,100,100

ANSWER**A. 100,300,600**

QUESTION 157

Input Data:

```
| id |  
|----|  
| 1 |  
| 2 |
```

Code:

```
""python  
df.union(df)  
""
```

Row Count?

- A. 2
- B. 3
- C. 4
- D. 1

ANSWER**C. 4**

QUESTION 158

Input Data:

```
| id |  
|----|  
| 1 |  
| 2 |  
| 3 |
```

Code:

```
""python  
df.filter(col("id") > 1)  
""
```

Returned ids?

- A. 1
- B. 2,3
- C. 1,2
- D. 3

ANSWER**B. 2,3**

QUESTION 159

Input Data:

```
| amount |  
|-----|  
| 100 |  
| 200 |  
| 300 |
```

Code:

```
""python  
df.orderBy(desc("amount"))  
""
```

First row?

- A. 100
- B. 200
- C. 300
- D. NULL

ANSWER**C. 300**

QUESTION 160

Input Data:

```
| customer_id | amount |  
|-----|-----|  
| C1 | 100 |  
| C1 | 200 |  
| C2 | 300 |
```

Code:

```
""python  
df.groupBy("customer_id").agg(sum("amount"))  
""
```

How many output rows?

- A. 1
- B. 2
- C. 3
- D. 4

ANSWER**B. 2****EXPLANATION**

Explanation

Aggregation returns one row per group. Since there are two unique customers (C1 and C2), the output contains two rows.

PySpark Certification Practice Test 09

Questions 161-180

Complete the Logic / Best Coding Solution

PYSPARK CERTIFICATION PRACTICE TEST 09Complete the Logic / Best Coding Solution Questions 161 180 (20 Questions)

QUESTION 161

A retail company wants the latest customer record.

Input:

```
customer_id | city | update_date  
-----|-----|-----  
1 | Bangalore | 2024-01-01  
1 | Chennai | 2024-02-01  
2 | Mumbai | 2024-01-15
```

Which window function is the BEST choice?

- A. row_number()
- B. count()
- C. avg()
- D. min()

ANSWER**A. row_number()**

QUESTION 162

A banking company wants the Top 5 customers by balance per branch.

Which approach is most suitable?

- A. Window + row_number()
- B. distinct()
- C. union()
- D. collect()

ANSWER**A. Window + row_number()**

QUESTION 163

A telecom company wants:

Current Month Usage

Previous Month Usage

Difference

Which function should be used?

- A. lag()
- B. lead()
- C. count()
- D. sum()

ANSWER

A. lag()

QUESTION 164

Find duplicate records based on:

customer_id

order_id

Best solution?

- A. `groupBy().count()`
- B. `collect()`
- C. `cache()`
- D. `coalesce()`

ANSWER

A. `groupBy().count()`

QUESTION 165

A dataset contains:

salary

Requirement:

Find the 3rd highest salary.

Best window function?

- A. `dense_rank()`
- B. `lag()`
- C. `sum()`
- D. `lead()`

ANSWER

A. `dense_rank()`

QUESTION 166

Input:

customer_id

transaction_date

amount

Requirement:

Running Total.

Best approach?

- A. `Window + sum()`
- B. `groupBy()`
- C. `collect()`
- D. `distinct()`

ANSWER

A. `Window + sum()`

QUESTION 167

A healthcare company receives late-arriving records.

Requirement:

Keep latest version per patient.

Most suitable pattern?

- A. row_number + desc timestamp
- B. count()
- C. min()
- D. union()

ANSWER

A

QUESTION 168

Input:

employee_id

salary

department

Requirement:

Highest-paid employee in each department.

Best solution?

- A. row_number partitioned by department
- B. collect()
- C. full join
- D. cache()

ANSWER

A

QUESTION 169

A table contains:

customer_id

event_date

Requirement:

Find consecutive events.

Primary functions?

- A. lag() + datediff()
- B. sum()
- C. avg()
- D. union()

ANSWER

A

QUESTION 170

Requirement:

Remove duplicates but keep the latest timestamp row.

Most suitable logic?

- A. row_number over timestamp desc
- B. distinct()
- C. count()
- D. collect()

ANSWER

A

QUESTION 171

Input:

sales_date

sales_amount

Requirement:

7-Day Moving Average.

Best solution?

- A. Window Between
- B. distinct()
- C. left join
- D. repartition()

ANSWER

A

QUESTION 172

A logistics company wants:

Current Shipment Cost

Previous Shipment Cost

Cost Difference

Which pattern should be used?

- A. lag()
- B. dense_rank()
- C. count()
- D. cache()

ANSWER

A

QUESTION 173

Requirement:

Identify customers who placed orders on 3 consecutive days.

Main concepts?

- A. lag + datediff + window
- B. union
- C. cache
- D. broadcast

ANSWER

A

QUESTION 174

A retail company wants:

Top Product per Category

Best coding approach?

- A. row_number partitionBy(category)
- B. collect
- C. repartition(1)
- D. distinct

ANSWER

A

QUESTION 175

Requirement:

Find customers whose current order amount is greater than previous order amount.

Best function?

- A. lag()
- B. lead()

- C. rank()
- D. avg()

ANSWER

A

QUESTION 176

A bank wants:

Monthly Balance Trend

with month-over-month comparison.

Which function is most appropriate?

- A. lag()
- B. collect()
- C. union()
- D. broadcast()

ANSWER

A

QUESTION 177

Input:

transaction_date

amount

Requirement:

Cumulative yearly revenue.

Best solution?

- A. Window + sum()
- B. distinct()
- C. min()
- D. collect()

ANSWER

A

QUESTION 178

A manufacturing company wants anomalies where:

today_temperature >

previous_temperature * 2

Best window function?

- A. lag()
- B. count()
- C. union()
- D. cache()

ANSWER

A

QUESTION 179

Find customers with no orders.

Best join type?

- A. Left Anti Join
- B. Inner Join
- C. Cross Join
- D. Right Join

ANSWER

A

QUESTION 180

A senior PySpark interviewer asks:

"Which PySpark concepts are most important for advanced analytics problems?"

- A. Window Functions
- B. print()
- C. show()
- D. limit()

ANSWER

A

title: PySpark Certification Practice Test 10

description: Debug The ETL Pipeline

PySpark Certification Practice Test 10

PYSPARK CERTIFICATION PRACTICE TEST 10

Debug The ETL Pipeline Questions 181 200 (20 Questions)

QUESTION 181

Problem:

```
"""python
df.filter(col("salary") > 10000)
"""
```

Expected Output:

100 rows

Actual Output:

1000 rows

Most likely issue?

- A. Result not assigned back to DataFrame
- B. Cache issue
- C. Broadcast issue
- D. Partition issue

ANSWER

A

EXPLANATION

```
### Fix
"""python
df = df.filter(col("salary") > 10000)
"""
```

QUESTION 182

Problem:

```
"""python
df.join(df2)
"""
```

Output row count is extremely large.

Most likely fix?

- A. Add join condition
- B. Add cache
- C. Add collect
- D. Add count

ANSWER

A

QUESTION 183

Problem:

```
"""python
df.collect()
"""
```

Job fails with Driver OOM.

Most likely fix?

- A. Avoid collect on huge datasets
- B. Increase columns
- C. Add distinct
- D. Add orderBy

ANSWER

A

QUESTION 184

Problem:

```
"""python
df.write.mode("append")
"""
```

Duplicates appear every day.

Best fix?

- A. Implement incremental/deduplication logic
- B. Use collect
- C. Cache data
- D. Use coalesce(1)

ANSWER

A

QUESTION 185

Problem:

A join takes 2 hours.

Tables:

```
"""text
```

Transactions = 5 TB

Reference = 20 MB

```
"""
```

Fix?

- A. Broadcast Reference Table
- B. Full Outer Join
- C. collect()
- D. coalesce(1)

ANSWER

A

QUESTION 186

Problem:

```
"""python
df.repartition(1)
"""
```

Output size:

```
"""text
```

600 GB

```
"""
```

Best fix?

- A. Remove single partition design
- B. Add collect()

- C. Use more filters
- D. Use count()

ANSWER

A

QUESTION 187

Problem:

Pipeline runtime increased suddenly.

Logs show:

```
""python
```

```
Python UDF
```

```
""
```

processing billions of rows.

Best recommendation?

- A. Use Spark Native Functions
- B. Add more print()
- C. Add union()
- D. Add count()

ANSWER

A

QUESTION 188

Problem:

Pipeline loads customer data.

Business finds outdated records.

Most likely fix?

- A. Latest Record Logic using row_number
- B. Use collect
- C. Add repartition
- D. Drop NULLs

ANSWER

A

QUESTION 189

Problem:

```
""python
```

```
dropDuplicates()
```

```
""
```

still leaves duplicates.

Most likely issue?

- A. Business keys not defined properly
- B. Cache issue
- C. Driver issue
- D. Executor issue

ANSWER

A

QUESTION 190

Problem:

```
""python
```

```
orderBy()
```

```
""
```

takes very long.

Dataset:

```
""text
```

```
3 Billion Rows
```

"""

Most likely explanation?

- A. Large Global Shuffle
- B. Broadcast Failure
- C. Missing Cache
- D. Missing Union

ANSWER

A

QUESTION 191

Problem:

Job succeeds.

Gold table row count:

"""text

0

"""

Most likely next step?

- A. Validate upstream filters
- B. Increase executors
- C. Increase cluster size
- D. Use cache

ANSWER

A

QUESTION 192

Problem:

SCD Type 2 table has multiple active rows.

Expected:

"""text

One active record

"""

Most likely fix?

- A. Correct end-date/current-flag logic
- B. Add repartition
- C. Add cache
- D. Add collect

ANSWER

A

QUESTION 193

Problem:

Incremental load misses records.

Watermark:

"""python

last_load_date

"""

Most likely area to review?

- A. Watermark Logic
- B. Cache Storage
- C. Broadcast Join
- D. Window Function

ANSWER

A

QUESTION 194

Problem:

```
""python  
df.count()  
""
```

used repeatedly for debugging.

Pipeline becomes slow.

Best recommendation?

- A. Remove unnecessary actions
- B. Add more count()
- C. Add more partitions
- D. Add unions

ANSWER

A

QUESTION 195

Problem:

Customer table contains:

```
""text  
1 customer  
5 duplicate orders  
""
```

Join multiplies rows.

Best fix?

- A. Verify data cardinality before join
- B. Add cache
- C. Add repartition
- D. Add count

ANSWER

A

QUESTION 196

Problem:

Source schema changed.

Pipeline starts failing.

Best engineering solution?

- A. Schema Validation Framework
- B. collect()
- C. count()
- D. Cache Everything

ANSWER

A

QUESTION 197

Problem:

Driver CPU constantly at 100%.

Logs show:

```
""python  
collect()  
foreach()  
local loops  
""
```

Most likely root cause?

- A. Large Processing on Driver
- B. Missing Broadcast
- C. Missing Filter

D. Missing Cache

ANSWER

A

QUESTION 198

Problem:

A skewed join causes one task to run much longer than others.

Most likely technique?

- A. Salting / AQE
- B. collect()
- C. union()
- D. count()

ANSWER

A

QUESTION 199

Problem:

Business reports inconsistent revenue numbers.

Two ETL pipelines calculate revenue differently.

Best correction?

- A. Standardize Business Logic
- B. Add Executors
- C. Repartition Data
- D. Cache More

ANSWER

A

QUESTION 200

Problem:

A PySpark ETL solution works perfectly in QA but fails at production scale.

What is the MOST likely missing consideration?

- A. Scalability Testing
- B. Column Naming
- C. Notebook Comments
- D. Print Statements

ANSWER

A

EXPLANATION

Explanation

Many ETL solutions work on millions of rows but fail on billions because scalability, skew, shuffles, memory usage, and execution plans were not evaluated properly.

title: PySpark Certification Practice Test 11

description: Spark Optimization Scenarios

PySpark Certification Practice Test 11

PYSPARK CERTIFICATION PRACTICE TEST 11

Spark Optimization Scenarios Questions 201 220 (20 Questions)

QUESTION 201

A PySpark job joins:

- Sales Table = 4 TB
- Product Lookup = 15 MB

Spark UI shows:

- Large shuffle
- Long join stage

Best optimization?

- A. Broadcast Join
- B. Cross Join
- C. Collect()
- D. Coalesce(1)

ANSWER

A. Broadcast Join

EXPLANATION

Explanation

Broadcasting the small lookup table avoids a costly shuffle.

QUESTION 202

Spark UI shows:

- 200 Tasks
- 199 finish quickly
- 1 task runs for 45 minutes

Most likely optimization?

- A. Salting for Data Skew
- B. Cache
- C. collect()
- D. Union

ANSWER

A

QUESTION 203

A DataFrame is used in:

- 5 joins
- 3 aggregations
- 4 transformations

Spark recomputes it repeatedly.

Best optimization?

- A. Cache/Persist
- B. count()
- C. orderBy()
- D. collect()

ANSWER

A

QUESTION 204

A write operation creates:

```
"""text
10,000 small files
"""
```

Performance issues appear downstream.

Best optimization?

- A. Coalesce
- B. collect()
- C. Window Functions
- D. Cross Join

ANSWER**A**

QUESTION 205

A job contains multiple joins.

Spark 3.x is available.

The team wants Spark to optimize join strategies automatically.

Best feature?

- A. AQE (Adaptive Query Execution)
- B. collect()
- C. Union
- D. Cache Removal

ANSWER**A**

QUESTION 206

A table is partitioned by:

```
""text
```

```
year
```

```
month
```

```
""
```

Query:

```
""sql
```

```
WHERE year = 2025
```

```
""
```

Which optimization helps?

- A. Partition Pruning
- B. Cross Join
- C. count()
- D. collect()

ANSWER**A**

QUESTION 207

Spark UI shows:

```
""text
```

```
Shuffle Read = Extremely High
```

```
""
```

during a join.

What should be investigated first?

- A. Join Strategy
- B. Variable Names
- C. Comments
- D. Notebook Title

ANSWER**A**

QUESTION 208

A dataset contains:

```
""text
```

```
5 Billion Rows
```

```
""
```

Output needs only:

```
""text
```

```
100 Sample Records
```

```
""
```

Best approach?

- A. limit(100)
- B. collect()
- C. count()
- D. orderBy()

ANSWER

A

QUESTION 209

A job repeatedly scans the same source files.
Which optimization may help?

- A. Cache
- B. Union
- C. distinct()
- D. collect()

ANSWER

A

QUESTION 210

A Data Engineer writes:

```
"""python
repartition(5000)
"""
```

for a table containing only 100,000 rows.
What should be reviewed?

- A. Partition Overhead
- B. Window Functions
- C. Delta Tables
- D. NULL Handling

ANSWER

A

QUESTION 211

A Spark job performs:

```
"""python
orderBy()
"""
```

on 2 billion rows.
What major operation should be expected?

- A. Global Shuffle
- B. Broadcast
- C. Cache
- D. Partition Pruning

ANSWER

A

QUESTION 212

A dimension table is used in nearly every ETL job.
Its size:

```
"""text
5 MB
"""
```

Best optimization?

- A. Broadcast
- B. Collect
- C. Repartition(1)

D. Full Join

ANSWER

A

QUESTION 213

A Spark application experiences memory pressure.

The cached dataset is rarely reused.

Best action?

- A. Remove Unnecessary Cache
- B. Add More collect()
- C. Increase Actions
- D. More withColumn()

ANSWER

A

QUESTION 214

Spark UI shows:

“text

Thousands of tiny tasks

“

with minimal processing per task.

Which area should be investigated?

- A. Excessive Partition Count
- B. Broadcast Hints
- C. UDF Logic
- D. Window Functions

ANSWER

A

QUESTION 215

A job fails because one partition contains most of the data.

Which issue is occurring?

- A. Data Skew
- B. Cache Failure
- C. Watermark Failure
- D. Schema Drift

ANSWER

A

QUESTION 216

A company stores data in Delta format.

Queries always filter by:

“text

customer_id

“

What should engineers evaluate?

- A. Z-Ordering
- B. collect()
- C. union()
- D. count()

ANSWER

A

QUESTION 217

Spark UI indicates:

```
"""text
```

Task Serialization Time

```
"""
```

is unusually high.

First review?

- A. Large Objects Sent to Executors
- B. More Joins
- C. More Aggregations
- D. More Partitions

ANSWER

A

QUESTION 218

A DataFrame is used only once.

A developer proposes:

```
"""python
```

```
cache()
```

```
"""
```

Best recommendation?

- A. Caching may not provide value
- B. Cache everything
- C. Always persist
- D. Convert to RDD

ANSWER

A

QUESTION 219

A Spark SQL query scans an entire table even though a filter is applied.

What optimization should be verified?

- A. Predicate Pushdown
- B. Cross Join
- C. collect()
- D. count()

ANSWER

A

QUESTION 220

A Senior Data Engineer is reviewing a slow PySpark job.

Which area generally provides the largest performance gains?

- A. Reducing Shuffles and Data Movement
- B. Renaming Columns
- C. Adding Comments
- D. Rearranging Notebook Cells

ANSWER

A

EXPLANATION

```
### Explanation
```

In large-scale Spark environments, performance improvements typically come from reducing shuffles, optimizing joins, handling skew, pruning partitions, and minimizing unnecessary data movement.

```
---
```

```
---
```

```
title: PySpark Certification Practice Test 12
```

```
description: Principal Engineer and Architecture Scenarios
```

```
---
```

```
# PySpark Certification Practice Test 12
```

PYSPARK CERTIFICATION PRACTICE TEST 12Principal Engineer and Architecture Scenarios Questions 221-250 (30 Questions)

QUESTION 221

A retail company currently performs full refreshes of a 10 TB sales table every night.

Only 0.5% of records change daily.

What is the BEST architectural improvement?

- A. Incremental Processing
- B. More Executors
- C. More Partitions
- D. Full Reload Twice Daily

ANSWER**A. Incremental Processing**

QUESTION 222

A banking company wants to track all customer address changes historically.

What design pattern should be implemented?

- A. SCD Type 2
- B. SCD Type 1
- C. Truncate and Load
- D. Full Refresh

ANSWER**A. SCD Type 2**

QUESTION 223

A healthcare organization receives streaming patient-monitoring data.

Business users require updates within seconds.

Which architecture is most appropriate?

- A. Structured Streaming
- B. Daily Batch Processing
- C. Full Refresh
- D. CSV Export

ANSWER**A. Structured Streaming**

QUESTION 224

A telecom company has hundreds of pipelines calculating revenue differently.

Executives no longer trust reports.

What should be implemented first?

- A. Common Business Rules Framework
- B. More Clusters
- C. More Partitions
- D. More Spark Jobs

ANSWER**A**

QUESTION 225

A manufacturing organization wants:

- Bronze Layer
- Silver Layer
- Gold Layer

What architecture pattern is being adopted?

- A. Medallion Architecture
- B. Lambda Architecture
- C. Kappa Architecture
- D. Star Schema

ANSWER

A

QUESTION 226

A company wants every pipeline execution to capture:

- Start Time
- End Time
- Record Count
- Status

What should be implemented?

- A. Audit Framework
- B. Broadcast Join
- C. Cache
- D. Z-Ordering

ANSWER

A

QUESTION 227

A retail company receives duplicate files frequently.

Leadership wants prevention before data reaches reporting tables.

Where should duplicate detection happen first?

- A. Silver Layer
- B. Gold Layer Only
- C. Dashboard Layer
- D. Reporting Tool

ANSWER

A

QUESTION 228

A financial institution requires:

- Reliability
- Replay Capability
- Historical Tracking

Which design principle becomes critical?

- A. Data Lineage and Auditability
- B. More Executors
- C. More Notebooks
- D. Larger CSV Files

ANSWER

A

QUESTION 229

A logistics company wants a single reusable framework used by all teams.

What should architects prioritize?

- A. Standardized ETL Framework
- B. Independent Development
- C. Separate Coding Styles
- D. Separate Business Logic

ANSWER

A

QUESTION 230

A PySpark platform serves:

- BI Teams
- Data Science
- Reporting
- AI Applications

What principle is most important?

- A. Shared Trusted Data Foundation
- B. Data Duplication
- C. Team-Specific Data Copies
- D. Local Metrics

ANSWER

A

QUESTION 231

A company must process 50 TB daily today and 500 TB within three years.

What architectural principle is MOST important?

- A. Scalability by Design
- B. Static Configuration
- C. Single Node Processing
- D. Fixed Clusters

ANSWER

A

QUESTION 232

A newly hired Architect wants to identify downstream impact before changing schemas.

What capability should be prioritized?

- A. Data Lineage
- B. Cache Management
- C. Repartitioning
- D. Broadcast Hints

ANSWER

A

QUESTION 233

A retail company wants trustworthy analytics datasets available to all departments.

What concept best fits?

- A. Data Products
- B. Excel Exports
- C. Independent Copies
- D. Manual Requests

ANSWER

A

QUESTION 234

A healthcare company requires validation of:

- Mandatory Columns
- Null Checks
- Reference Integrity

What should be embedded into ETL?

- A. Data Quality Framework
- B. More Workers
- C. More Partitions
- D. More Joins

ANSWER

A

QUESTION 235

A banking company processes CDC events.

Operations include:

“text

Insert

Update

Delete

“

What design capability is required?

- A. Merge/Upsert Framework
- B. Full Refresh
- C. Broadcast Join
- D. Cache

ANSWER**A**

QUESTION 236

A global telecom company struggles with ownership confusion.

Incidents take hours to resolve.

What governance improvement is needed?

- A. Defined Data Ownership
- B. More Executors
- C. More Storage
- D. More Notebooks

ANSWER**A**

QUESTION 237

A company wants pipeline failures automatically reported before users notice.

What capability becomes critical?

- A. Monitoring and Alerting
- B. Repartition
- C. Union
- D. Cache

ANSWER**A**

QUESTION 238

A retail company wants to reduce cloud costs while maintaining performance.

What should engineers monitor?

- A. Cost Attribution and Resource Usage
- B. Notebook Count
- C. Column Count
- D. File Names

ANSWER**A**

QUESTION 239

A company frequently acquires other organizations.

What architecture principle reduces onboarding effort?

- A. Standardized Data Models
- B. Separate Standards
- C. Different Schemas Per Team

D. Manual Mapping Only

ANSWER

A

QUESTION 240

A platform team wants new developers onboarded quickly.
What investment provides the biggest benefit?

- A. Documentation and Reusable Templates
- B. More Clusters
- C. More Databases
- D. More Views

ANSWER

A

QUESTION 241

A company wants all production deployments tracked and auditable.
What should exist?

- A. CI/CD and Deployment Audit Trail
- B. Manual Production Changes
- C. Shared Credentials
- D. Local Scripts

ANSWER

A

QUESTION 242

A PySpark environment supports hundreds of engineers.
Executives want consistent coding practices.
What should be established?

- A. Engineering Standards
- B. Personal Preferences
- C. Independent Frameworks
- D. Separate Rules

ANSWER

A

QUESTION 243

A business wants trusted datasets discoverable without IT tickets.
What capability should be introduced?

- A. Data Catalog
- B. Shared Spreadsheets
- C. Manual Requests
- D. Local Documentation

ANSWER

A

QUESTION 244

A streaming pipeline processes financial transactions.
Business requires:

“text

Exactly Once Processing

“

What should architects prioritize?

- A. Reliable Streaming Design
- B. More Partitions
- C. More Columns

D. Larger Files

ANSWER

A

QUESTION 245

A company wants to move from reactive support to proactive operations.
What capability should mature first?

- A. Observability
- B. Cache Usage
- C. Distinct Logic
- D. show()

ANSWER

A

QUESTION 246

A platform serves:

- Analytics
- AI
- Operations
- Compliance Teams

What architectural principle dominates?

- A. Governed Shared Platform
- B. Independent Copies
- C. Manual Data Exchange
- D. Team Isolation

ANSWER

A

QUESTION 247

A global organization wants architecture decisions reviewed before implementation.
What process should exist?

- A. Architecture Review Board
- B. Open Deployments
- C. Direct Production Changes
- D. Manual Emails

ANSWER

A

QUESTION 248

A Principal Engineer wants to reduce dependency on tribal knowledge.
What should improve?

- A. Knowledge Management
- B. Cluster Size
- C. Cache Usage
- D. File Naming

ANSWER

A

QUESTION 249

A company wants success measured by outcomes rather than technical metrics alone.
What should leadership define?

- A. Business-Focused KPIs
- B. Partition Counts
- C. Executor Counts
- D. Notebook Counts

ANSWER

A

QUESTION 250

Spark UI shows:

""text

Shuffle Read = Extremely High

""

during a join.

What should be investigated first?

- A. Join Strategy
- B. Variable Names
- C. Comments
- D. Notebook Title

ANSWER

A
