

BITS WILP Cloud Computing Solved Paper - BDS - 2024 - EC2 Makeup

..... crackbitswilp.in

Q 1.1 [03 Marks]

What are the desirable characteristics of a Big Data System?

Explanation

- Application does not need to bother about sharding and replication.
- Developers focus on application logic rather than data management.
- Flexible schema.
- Treat data as immutable where possible.
- Keep timestamped versions.
- Avoid destroying good copies of data.
- Provide application-specific consistency models.
- Handle volume, velocity and variety.
- Eventual consistency support.
- Distributed scale-out architecture.
- Low-cost long-term storage (HDFS).
- Computation close to data (MapReduce).

..... crackbitswilp.in

Q 1.2 [04 Marks]

A 2-tier application uses a 3 node application cluster and a 2 node DB cluster. The application works only when both tiers are available. The application tier is in an active-active load-balanced configuration with the given nodes. But the database tier is in a cold standby mode where it takes 12 hours to switch for bringing a passive node online. If an application node fails every 10 days and a DB node fails every 100 days, find the following:

**PREMIUM SOLVED PREVIOUS QUESTION
PAPERS WITH EXPLANATIONS**

 **+91 8660242841**

WhatsApp Us · Quick Response

**TOPIC WISE FREE NOTES & CONCEPTS FOR
EASY LEARNING**

crackbitswilp.in

Free Preview · All Subjects

- (a) MTTF of the application tier
- (b) MTTF of the database tier
- (c) Availability of the database tier
- (d) Overall availability of the 2-tier system, assuming MTTR of the application tier is negligible.

Explanation

- (a) MTTF of an application node = 10 days. Therefore,

$$\text{MTTF (Application Tier)} = 10 \times 3 = 30 \text{ days}$$

- (b) MTTF of a DB node = 100 days. Therefore,

$$\text{MTTF (DB Tier)} = 100 \text{ days}$$

since it is an active-passive configuration.

- (c) Availability of the DB tier:

$$\text{Availability} = \frac{\text{MTTF}}{\text{MTTF} + \text{MTTR}} = \frac{100}{100 + 0.5} = 0.995 \approx 99.5\%$$

- (d) Overall availability of the 2-tier system:

$$\begin{aligned} \text{Availability (System)} &= \text{Availability (Application Tier)} \times \text{Availability (DB Tier)} \\ &= 1 \times 0.995 = 0.995 \approx 99.5\% \end{aligned}$$

Note: Availability of the application tier is

$$\frac{\text{MTTF}}{\text{MTTF} + \text{MTTR}}$$

where MTTR is assumed to be 0 (or negligible) in an active-active configuration.

(4 × 1 mark awarded if the formula/reasoning is clear in addition to the numerical answer.)

..... crackbitswilp.in

Q 1.3 [4 Marks]

In a distributed computing system consisting of multiple nodes, data is accessed by an application from other nodes. The data is cached in a memory of size 10 MB on the local node, and it can be delivered only from this cache to the application.

Calculate the average access time of a block of data consisting of 1 MB fetched from one node and processed sequentially from memory in another node of the cluster if the hit rate is 80%. You may

ignore the advantages in the access time contributed by the filesystem/network caching.

Typical timing values are given below:

- Read 1 MB sequentially from disk = 20,000 μ s
- Read 1 MB sequentially from memory = 250 μ s
- Send 1 MB over a 1 Gbps network = 100 μ s

You need to create a suitable relationship between the parameters and calculate the average access time using the formula created by you.

Explanation

Let,

- T_{avg} = Average access time
- h = Cache hit rate
- T_d = Disk read time
- T_m = Memory read/write time
- T_n = Network transport time

The average access time is given by

$$T_{avg} = h \times T_m + (1 - h) \times (T_d + T_n + T_m)$$

Substituting the given values,

$$T_{avg} = 0.8 \times 250 + 0.2 \times (20000 + 100 + 250)$$

$$= 200 + 0.2 \times 20350$$

$$= 4070 \mu s$$

Required Data:

- Read 1 MB sequentially from disk = 20,000,000 ns
- Read 1 MB sequentially from memory = 250,000 ns
- Send 2 KB over a 1 Gbps network = 20,000 ns

..... crackbitswilp.in

Q 1.4 [4 Marks]

You are configuring a NoSQL cluster with 7 nodes. The nodes of the cluster can be configured with servers (1) hosted within the same Rack, (2) hosted on different Racks within the same datacenter, or (3) servers hosted on remote datacenters.

How do you select the right combination of servers for the nodes of the cluster for maximizing:

- (a) Consistency
- (b) Availability
- (c) Partition Tolerance

Which consistency level will be selected by you for:

- (a) Write-intensive applications
- (b) Read-intensive applications

Explanation

(a) **Configuration for maximizing consistency:**

Select a replication factor of 1 so that the data is not replicated. With only one copy of the data, there are no consistency issues. If a replication factor greater than 1 is selected, the configuration should introduce minimum replication lag. Therefore, place all 7 nodes on servers hosted within the same rack.

(b) **Configuration for maximizing availability:**

Distribute the 7 nodes across diverse groups. At least one node should be placed on a different rack to tolerate rack failures, and at least one node should be placed in a different datacenter to tolerate failures affecting the primary datacenter.

(c) **Configuration for maximizing partition tolerance:**

Partitions are caused by network link failures. If only partition tolerance is to be addressed, all 7 nodes can be hosted on the same rack. The probability of network failures within the same rack is very small.

(d) **Consistency levels:**

- Write-intensive applications: Consistency Level 0 (None) **or** ONE (Primary only).
- Read-intensive applications: Consistency Level ONE.

Award 1 additional mark for justification of the selected configuration and 1 mark each for the listed configurations.

..... crackbitswilp.in

Q1.5 [4 Marks]

You have to design a microblogging application where a user can have various types of interactions – e.g., signup, connect to users as friends, read posts from friends and others, create a new post, reply to an existing post or reply to a comment on a post. The application needs to run at a large scale with users across the globe.

Discuss your choice of database (e.g., RDBMS, NoSQL) and the type(s) of configuration you will choose as per the CAP Theorem. Given the example user interactions listed above, it is possible that

different interaction scenarios need unique configurations.

Explanation

Answer:

- (a) Considering the large-scale nature of the application and the absence of strict multi-data-item transactional consistency requirements, a **NoSQL database** is the preferred choice.

Suitable reasons include:

- Better horizontal scalability.
- High availability for globally distributed users.
- Flexible schema suitable for social media workloads.

- (b) Most of the user interactions do not require strong consistency among multiple readers and writers.

- Sign-up and friend connection requests are independent operations.
- Reading and writing posts do not require strict consistency.
- Replies to posts or comments require only that the referenced post/comment already exists.
- Latest posts and comments need not be immediately visible to every user.

Therefore, according to the CAP Theorem, an **AP (Availability + Partition Tolerance)** configuration is the preferred choice for the database.

(2 marks if NoSQL is selected with at least one valid reason. 2 marks if AP configuration is selected with appropriate justification.)

..... crackbitswilp.in

Q 1.6 [4 Marks]

The employee data of a company with the following schema is given to you in the form of a CSV file. The file has 1 million records. The 2 instances given are examples.

Employee_ID	Name	Age	Gender	Salary
1201	Gopal	45	Male	50000
1202	Manisha	40	Female	48000

Write pseudocode for a MapReduce program to find out the total number of Male and Female employees having the same age. You need to output Age, No. of Males with this age, No. of Females with this age. The number of output records should be equal to the number of distinct values of Age of employees.

Explanation**Answer:****Mapper:**

- Read each employee record.
- Extract the Age and Gender fields.
- If Gender is Male, emit:
 $(\text{Age}, (1, 0))$
- If Gender is Female, emit:
 $(\text{Age}, (0, 1))$

Reducer:

- For each Age, initialize MaleCount = 0 and FemaleCount = 0.
- For every value received:
 - Add the first value to MaleCount.
 - Add the second value to FemaleCount.
- Output:
 $(\text{Age}, \text{MaleCount}, \text{FemaleCount})$

Sample Output:

Age	Male Count	Female Count
40	0	1
45	1	0

(Award marks for a correct Mapper, Reducer, and output format that produces one record for each distinct age.)

..... crackbitswilp.in

Q 1.7 [4 Marks]

What are the basic design principles of a Hadoop cluster?

Explanation**Answer:**

- Scale-out:** Scale-up architectures are rarely used, and scale-out is the standard approach in big data processing.
- Share nothing:** Communication and dependencies are bottlenecks. Individual components should be as independent as possible so that processing can continue even if other components fail.

- (c) **Expect failure:** Components will fail at inconvenient times, so the system should be designed to tolerate failures.
- (d) **Smart software, dumb hardware:** Push intelligence into the software, which is responsible for managing and allocating generic hardware resources.
- (e) **Move processing, not data:** Perform processing locally on the data. Only program binaries and status reports should be transferred over the network, as they are much smaller than the actual data.
- (f) **Build applications, not infrastructure:** Instead of focusing on data movement and processing, rely on the Hadoop framework for job scheduling, error handling, and coordination, and focus on application logic and its implementation.

..... crackbitswilp.in

Q 1.8 [4 Marks]

In a Hadoop cluster, jobs can be processed in 4 different queues. Fair Share Scheduler is plugged into the YARN Resource Manager with queues Q₀, Q₁, Q₂, and Q₃. The scheduler is configured to preempt containers from other queues if required.

The weights allocated for each of the queues are given below:

- Q₀ – Weight 0.00
- Q₁ – Weight 0.50
- Q₂ – Weight 0.30
- Q₃ – Weight 0.20

What will be the percentage of total capacity allocated to the jobs at run time in each of the following scenarios?

- (a) Job₁ is submitted to Q₃ and it is the only job running on the cluster.
- (b) While Job₁ was running in Q₃, Job₂ is submitted to Q₂.
- (c) While Job₁ was running in Q₃ and Job₂ was running in Q₂, Job₃ is submitted to Q₁.
- (d) When Job₃ was the only job running in Q₁, Job₄ is submitted to Q₀.

Explanation

Answer:

- (a) 100.0% of the total capacity is allocated to Job₁.
- (b) 60.0% of the total capacity is allocated to Job₂ and 40.0% is allocated to Job₁.
- (c) 50.0% of the total capacity is allocated to Job₃, 30.0% to Job₂, and 20.0% to Job₁.
- (d) Job₄ will start running only after Job₃ is completed, and it will receive 100.0% of the cluster resources.

(1 mark for each correct answer.)

TOPIC WISE FREE NOTES & CONCEPTS FOR
EASY LEARNING

crackbitswilp.in

Free Preview · All Subjects

PREMIUM SOLVED PREVIOUS QUESTION
PAPERS WITH EXPLANATIONS

 **+91 8660242841**

WhatsApp Us · Quick Response