

Apache Hadoop® Disaggregation and Tiering

Big Analytics Drive Big Value

Highlights

Analytics applications are rapidly becoming the key applications for Big Data workloads. These applications address large datasets coming from transactional processing, Internet-of-Things (IoT) devices, Web properties, and other sources to find patterns in the data that can be leveraged by data scientists to take decisive action in a fast-moving marketplace.

Solution

- Provide cost-effective and ultra-speed ingest capabilities
- Minimize server sprawl by tuning storage and compute separately
- Reduce overall costs by requiring fewer servers
- Reduce time from ingest to result, improving responsiveness

Analytics is an umbrella term used to describe specific workloads that are widely deployed within financial services companies, Web properties, manufacturing industries, and other large organizations. These workloads are needed to cope with the data tsunami that is hitting firms of all types. Sensors scattered throughout a factory floor need to have their output processed quickly to ensure maximum uptime. Customer patterns need to be considered in real time for fashion-forward retailers. Financial services companies need to process historical and real-time trade data to guarantee the maximum return on their investments.

Hadoop Scale-out Woes

Apache Hadoop® is a leading software component that allows users to identify key data points in extremely large datasets. Built on top of the base Hadoop services, applications such as Apache HBase™ (a massive-scale distributed Big Data store) and Apache Spark™ (a general-purpose compute engine focusing on streaming and in-memory Big Data workloads) make it easier to produce meaningful results. By leveraging these technologies, companies faced with a big data onslaught can effectively produce insight from the raw data.

Hadoop is well-known for its ability to scale out across multiple, identically configured servers. At smaller scales, using Hadoop simplifies a data center manager's job, as fewer server configurations need to be supported. If more compute or storage is needed, identical servers are added. Unfortunately, as datasets increase in size and longer time periods need to be examined, this simplicity comes at a high cost: server sprawl.

Disaggregation to the Rescue

Hadoop was designed in a world where gigabit networking was state of the art, hard drive sizes were measured in gigabytes, and flash memory was slow and sold in megabytes. Keeping data close to compute was essential to attaining high performance. In today's world where flash can deliver gigabytes per second of bandwidth and 10-gigabit networks are ubiquitous, combining storage and compute ends up wasting space, power, and IT dollars. Massive server sprawl is eating into operational budgets when nodes are rolled out simply to extend storage space or deal with peak data ingest loads. This approach doesn't give IT administrators the flexibility they need to perform their jobs, nor provide the speed that data scientists require to do their jobs effectively.

Disaggregation—separating and tiering storage from the Hadoop compute infrastructure—can help solve these problems. By using commodity networking, high-density SSD servers, and ultra-capacity storage servers to implement the Hadoop storage tier, companies are free to tune the Hadoop compute tier to their own needs.

Pain Point: Compute Cores Idling, Waiting on Disk I/O

Moving the bulk storage out of the compute nodes in a Hadoop cluster enables the datacenter architect to configure each compute node with a small amount of local flash. This flash can store the operating system, greatly reducing deployment and start-up times. It can also provide temporary storage for Hadoop processing and its infamous "shuffle" phase, where data from multiple compute nodes is rearranged by a smaller subset of nodes. Enterprise-class Ultrastar products provide multiple performance and attachment options for flash storage, including add-in cards and 2.5" SFF form factors, to allow deployment in a wide variety of high-density compute nodes.

Pain Point: Storage Node Sprawl

Storage nodes consisting of a standard server and a high-density, SAS-attached JBOD provide massive capacity with a minimum footprint in space, power, and cost. In just a 5U height using high-capacity 14TB Ultrastar HDDs, a storage administrator can attach a single 1U storage server to a single 4U JBOD that houses 60 Ultrastar DC HC520 HDDs and stores 840 terabytes. Attaining such capacity in a typical Hadoop cluster with front-loaded disk drives would require at least 16U of rack space!

Pain Point: Ingesting a Data Tsunami

In the most demanding of applications, big data arrives in the cluster in massive, high-volume bursts. These bursts need to be processed in near-real time, but once processed they can be stored away for less urgent needs. For these use cases, an all-flash storage tier with NVMe™-based Ultrastar SSDs can significantly improve performance and shrink the number of servers required to manage this data. With capacities of 7.68 terabytes or more per 2.5" SFF NVMe SSD, a single 2-U all-flash storage server can contain over 360 terabytes of microsecond latency flash.

Summary

Hadoop, HBase and Spark are the foundations of big data analytics in many companies. You can minimize the expense and size of the infrastructure needed to implement them with Ultrastar SSDs. You can also disaggregate the storage and implement a high-density hard drive storage tier, a local flash tier directly tied to the compute CPUs, and, where applicable, an all-flash storage tier for maximum results and minimum expenditure.

	NVMe SSD	NVMe SSD	HelioSeal HDD
Pain Point	Ultrastar DC SN630	Ultrastar DC SN200	Ultrastar DC HC500 Series
Compute cores idling, waiting on disk I/O	• •		
Server node sprawl			• • •
Ingesting a data tsunami	• •	• • •	
Legend: • = Good •• = Better ••• = Best			

Western Digital.

5601 Great Oaks Parkway
San Jose, CA 95119, USA
US (Toll-Free): 800.801.4618
International: 408.717.6000

www.westerndigital.com

© 2017-2019 Western Digital Corporation or its affiliates. All rights reserved. Produced 4/17. Rev. 4/19. Western Digital, the Western Digital logo, HelioSeal and Ultrastar are trademarks or registered trademarks of Western Digital Corporation or its affiliates in the US and/or other countries. The NVMe™ and NVM Express word marks are trademarks of NVM Express, Inc. Apache®, Apache Hadoop, Apache HBase, Apache Spark, and Hadoop® are either registered trademarks or trademarks of The Apache Software Foundation in the United States and/or other countries. All other marks are the property of their respective owners. One gigabyte (GB) is equal to 1,000MB (one billion bytes) and one terabyte (TB) is equal to 1,000GB (one trillion bytes) when referring to solid-state capacity. Accessible capacity will vary from the stated capacity due to formatting and partitioning of the drive, the computer's operating system, and other factors.