



Planning Guide

Abstract

This book provides recommendations for those considering the use of VoltDB.

Planning Guide

V14

Copyright © 2008-2024 Volt Active Data, Inc.

This document is published under copyright by Volt Active Data, Inc. All Rights Reserved.

The software described in this document is furnished under a license by Volt Active Data, Inc. Your rights to access and use VoltDB features are defined by the license you received when you acquired the software.

The VoltDB client libraries, for accessing VoltDB databases programmatically, are licensed separately under the MIT license.

Volt Active Data, VoltDB, and Active(N) are registered trademarks of Volt Active Data, Inc.

VoltDB software is protected by U.S. Patent Nos. 9,600,514, 9,639,571, 10,067,999, 10,176,240, and 10,268,707. Other patents pending.

This document was generated on September 29, 2025.

Table of Contents

Preface	vi
1. Organization of this Manual	vi
2. Other Resources	vi
1. The Planning Process	1
2. Proof of Concept	2
2.1. Effective Partitioning	2
2.2. Designing the Stored Procedures	2
3. Choosing Hardware	4
3.1. The Dimensions of Hardware Sizing	4
3.2. Sizing for Throughput	4
3.3. Sizing for Capacity	5
3.4. Sizing for Durability	7
4. Sizing Memory	9
4.1. Planning for Database Capacity	9
4.1.1. Sizing Database Tables	10
4.1.2. Sizing Database Indexes	11
4.1.3. An Example of Database Sizing	12
4.2. Distributing Data in a Cluster	12
4.2.1. Data Memory Usage in Clusters	13
4.2.2. Memory Requirements for High Availability (K-Safety)	13
4.3. Planning for the Server Process (Java Heap Size)	13
4.3.1. Attributes that Affect Heap Size	14
4.3.2. Guidelines for Determining Java Heap Size	14
5. Benchmarking	16
5.1. Benchmarking for Performance	16
5.1.1. Designing Your Benchmark Application	16
5.1.2. How to Measure Throughput	18
5.1.3. How to Measure Latency	19
5.2. Determining Sites Per Host	19

List of Figures

5.1. Determining Optimal Throughput and Latency	17
5.2. Determining Optimal Sites Per Host	20

List of Tables

3.1. Quick Estimates for Memory Usage By Datatype	6
4.1. Memory Requirements For Tables By Datatype	11

Preface

This book provides information to help those considering the use of VoltDB. Choosing the appropriate technology is more than deciding between one set of features and another. It is important to understand how the product fits within the existing technology landscape and what it requires in terms of systems, support, etc. This book provides guidelines for evaluating VoltDB, including sizing hardware, memory, and disks based on the specific requirements of your application and the VoltDB features you plan to use.

1. Organization of this Manual

This book is divided into 5 chapters:

- Chapter 1, *The Planning Process*
- Chapter 2, *Proof of Concept*
- Chapter 3, *Choosing Hardware*
- Chapter 4, *Sizing Memory*
- Chapter 5, *Benchmarking*

2. Other Resources

This book assumes you are already familiar with the VoltDB feature set. The choice of features, especially those related to availability and durability, will impact sizing. Therefore, if you are new to VoltDB, we encourage you to visit the VoltDB web site to familiarize yourself with the product options and features.

You may also find it useful to review the following books to better understand the process of designing, developing, and managing VoltDB applications:

- *VoltDB Tutorial*, a quick introduction to the product and is recommended for new users
- *Using VoltDB*, a complete reference to the features and functions of the VoltDB product
- *VoltDB Administrator's Guide*, information on managing VoltDB clusters

These books and more resources are available on the web from <http://www.voltactivedata.com/>.

Chapter 1. The Planning Process

Welcome to VoltDB, a best-in-class database designed specifically for high volume transactional applications. Since you are reading this book, we assume you are considering the use of VoltDB for an existing or planned database application. The goal of this book is to help you understand the impact of such a decision on your computing environment, with a particular focus on hardware requirements.

Technology evaluation normally involves several related but separate activities:

- **Feature Evaluation**

The goal of the feature evaluation is to determine how well a product's features match up to the needs of your application. For VoltDB, we strongly encourage you to visit our website and review the available product overviews and technical whitepapers to see if VoltDB is right for you. If you need additional information, please feel free to contact us directly.

- **Proof of Concept**

The proof of concept, or POC, is usually a small application that emulates the key business requirements using the proposed technology. The goal of the POC is to verify, on a small scale, that the technology performs as expected for the target usage.

- **Hardware Planning**

Once you determine that VoltDB is a viable candidate for your application, the next step is to determine what hardware environment is needed to run it. Hardware sizing requires an understanding of the requirements of the business application (volume, throughput, and availability needs) as well as the technology. The primary goal of this book is to provide the necessary information about VoltDB to help you perform this sizing exercise against the needs of your specific application.

- **Benchmarking**

Having determined the feasibility of the technology, the final activity is to perform benchmarking to evaluate its performance against the expected workload. Benchmarking is often performed against the proof of concept or a similar prototype application. Benchmarking can help validate and refine the hardware sizing calculations.

Let's assume you have already performed a feature evaluation, which is why you are reading this book. You are now ready to take the next step. The following chapters provide practical advice when building a proof of concept, sizing hardware, and benchmarking a solution with VoltDB.

Note that this book does *not* help with the detailed application design itself. For recommendations on application design we recommend the other books about VoltDB. In particular, *Using VoltDB*.

Chapter 2. Proof of Concept

A proof of concept (POC) is a small application that tests the key requirements of the proposed solution. For database applications, the POC usually focuses on a few critical transactions, verifying that the database can support the proposed schema, queries, and, ultimately, the expected volume and throughput. (More on this in the chapter on Benchmarking.)

A POC is *not* a full prototype. Instead, it is just enough code to validate that the technology meets the need. Depending upon the specific business requirements, each POC emphasizes different database functionality. Some may focus primarily on capacity, some on scalability, some on throughput, etc.

Whatever the business requirements, there are two key aspects of VoltDB that must be designed correctly to guarantee a truly effective proof of concept. The following sections discuss the use of partitioning and stored procedures in POCs.

2.1. Effective Partitioning

VoltDB is a distributed database. The data is partitioned automatically, based on a partitioning column you, as the application developer, specify. You do not need to determine where each record goes — VoltDB does that for you.

However, to be effective, you must choose your partitioning columns carefully. The best partitioning column is not always the most obvious one.

The important thing to keep in mind is that VoltDB partitions both the data *and* the work. For best performance you want to partition the database tables and associated queries so that the most common transactions can be run in parallel. That is, the transactions are, in VoltDB parlance, "single-partitioned".

To be single-partitioned, a transaction must only contain queries that access tables based on a specific value for the partitioning column. In other words, if a transaction is partitioned on the EmpID column of the Employee table (and that is the partitioning column for the table), all queries in the transaction accessing the Employee table must include the constraint `WHERE Employee.EmpID = {value}`.

To make single-partitioned transactions easier to create, not all tables have to be partitioned. Tables that are not updated frequently can be replicated, meaning they can be accessed in any single-partitioned transaction, no matter what the partitioning key value.

When planning the partitioning schema for your database, the important questions to ask yourself are:

- Which are the critical, most frequent queries? (These are the transactions you want to be single-partitioned.)
- For each critical query, what database tables does it access and using what column?
- Can any of those tables be replicated? (Replicating smaller, less frequently updated tables makes joins in single-partitioned procedures easier.)

2.2. Designing the Stored Procedures

Designing the schema and transactions to be single-partitioned is one thing. It is equally important to make sure that the stored procedures operate in a way that lets VoltDB do its job effectively.

The first step is to write the transactions as stored procedures that are loaded into the schema. Do *not* write critical transactions as ad hoc queries to the database. VoltDB provides the `@AdHoc` system procedure

for executing arbitrary SQL queries, which can be helpful when building early prototypes to test queries or occasionally validate the content of the database. But `@AdHoc` queries often run as multi-partitioned transactions and, therefore, should not be used for critical or repetitive transactions.

The next step is to ensure that single-partitioned stored procedures are correctly identified as such in the schema using the `PARTITION ON` clause in the `CREATE PROCEDURE` statement and specifying the appropriate partitioning column.

Finally, when designing for maximum throughput, use asynchronous calls to invoke the stored procedures from within the POC application. Asynchronous calls allow VoltDB to queue the transactions (and their responses), avoiding any delays between when a transaction completes, the results are returned to the POC application, and the next procedure is invoked.

Chapter 5, *Benchmarking* later in this book provides additional suggestions for effectively designing and testing proof of concept applications.

Chapter 3. Choosing Hardware

VoltDB is designed to provide world class throughput on commodity hardware. You do not need the latest or most expensive hardware to achieve outstanding performance. In fact, a few low- to mid-range servers can easily outperform a single high end server, since VoltDB throughput tends to scale linearly with the number of servers in the cluster.

People often ask us at VoltDB "what type of servers should we use and how many?" The good news is that VoltDB is very flexible. It works well on a variety of configurations. The bad news is that the true answer to the question is "it depends." There is no one configuration that is perfect for all situations.

Like any technology question, there are trade offs to be made when choosing the "right" hardware for your application. This chapter explains what those trade offs are and provides some general rules of thumb that can help when choosing hardware for running a VoltDB database.

3.1. The Dimensions of Hardware Sizing

There are three key dimensions to sizing individual servers: the number and speed of the processors, the total amount of memory, and the size and type of disks available. When sizing hardware for a distributed database such as VoltDB, there is a fourth dimension to consider: the number of servers to use.

Each of these dimensions impacts different aspects of VoltDB performance. The number of processors affects how many partitions can be run on each server and, as a result, throughput. The available memory obviously impacts capacity, or the volume of data that can be stored. The size, number, and type of disks impacts the performance of availability features such as snapshots and command logging.

However, they also interact. The more memory per server, the longer it takes to write a snapshot or for a node to rejoin after a failure. So increasing the number of servers but reducing the amount of memory per server may reduce the impact of durability on overall database performance. These are the sorts of trade offs that need to be considered.

The following sections discuss hardware sizing for three key aspects of a VoltDB application:

- Throughput
- Capacity
- Durability

3.2. Sizing for Throughput

The minimum hardware requirements for running a VoltDB database server is a 64-bit machine with two or more processor cores. The more cores the server has, the more VoltDB partitions can potentially run on that server. The more unique partitions, the more throughput is possible with a well partitioned workload.

However, the number of processor cores is not the only constraint on throughput. Different aspects of the server configuration impact different characteristics of the database process.

For example, although the more physical cores a server has increases the number of partitions that server can potentially handle, at some point the number of transactions being received and data being returned exceeds the capacity of the network port for that server. As a consequence, going beyond 12-16 cores on a

single machine may provide little value to a VoltDB database, since the server's network adapter becomes the gating factor.

Rule of Thumb

VoltDB runs best on servers with between 4 and 16 cores.

It should be noted that the preceding discussion refers to physical processor cores. Some servers support hyperthreading. *Hyperthreading* is essentially the virtualization of processor cores, doubling the reported number of cores. For example, a system with 4 cores and hyperthreading acts like an 8 core machine. These virtualized cores can improve VoltDB performance, particularly for servers with a small (2 or 4) number of physical cores. However, the more physical cores the system has, the less improvement is seen in VoltDB performance. Therefore, hyperthreading is not recommended for VoltDB servers with more than 8 physical cores.

The alternative to adding processors to an individual server for improving throughput is to add more servers. If a single 4-core server can handle 3 partitions, two such servers can handle 6 partitions, three can handle 9, etc. This is how VoltDB provides essentially linear scaling in throughput.

But again, there are limits. For peak performance it is key that network latency and disruption between the cluster nodes be kept to a minimum. In other words, all nodes of the cluster should be on the same network switch. Obviously, the capacity of the network switch constrains the number of nodes that it can support. (A 32 port switch is not uncommon.)

Rule of Thumb

Best performance is achieved with clusters of 2-32 servers connected to the same network switch.

It is possible to run a VoltDB cluster across switches. (For example, this is almost always the case in cloud environments.) However, latency between the cluster nodes will have a negative impact on performance and may ultimately limit overall throughput. In these situations, it is best to benchmark different configurations to determine exactly what performance can be expected.

Finally, it should be noted that the speed of the processor cores may not have a significant impact on overall throughput. Processor speed affects the time it takes to execute individual transactions, which may be only a small percentage of overall throughput. For workloads with very compute-intensive transactions, faster processors can improve overall performance. But for many small or simple transactions, improved processor speed will have little or no impact.

3.3. Sizing for Capacity

The second aspect of database sizing is capacity. Capacity describes the maximum volume of data that the database can hold.

Since VoltDB is an in-memory database, the capacity is constrained by the total memory of all of the nodes in the cluster. Of course, one can never size servers too exactly. It is important to allow for growth over time and to account for other parts of the database server that use memory.

Chapter 4, *Sizing Memory* explains in detail how memory is assigned by the VoltDB server for database content. Use that chapter to perform accurate sizing when you have a known schema. However, as a rough estimate, you can use the following table to approximate the space required for each column. By adding up the columns for each table and index (including index pointers) and then multiplying by the expected number of rows, you can determine the total amount of memory required to store the database contents.

Table 3.1. Quick Estimates for Memory Usage By Datatype

Datatype	Bytes in Table	Bytes in Index
TINYINT	1	1
SMALLINT	2	2
INTEGER	4	4
BIGINT	8	8
DOUBLE	8	8
DECIMAL	16	16
TIMESTAMP	8	8
VARCHAR ^a or VARBINARY (less than 64 bytes)	length + 1	length + 1
VARCHAR ^a or VARBINARY (64 bytes or greater)	length	8
index pointers	n/a	40

^aFor VARCHAR columns declared in characters, rather than in bytes, the length is calculated as four bytes for every character. In other words, for storage calculations a string column declared as VARCHAR(16) has the same length as a column declared as VARCHAR(64 BYTES).

You must also account for the memory required by the server process itself. If you know how many tables the database will contain and how many sites per host will be used, you can calculate the server process memory requirements using the following formula:

$$384\text{MB} + (10\text{MB} \times \text{number of tables}) + (128\text{MB} \times \text{sites per host})$$

This formula assumes you use K-safety, which is recommended for all production environments. If the cluster is also an XDCR database, you should increase the multiplier for sites per host from 128 to 256 megabytes:

$$384\text{MB} + (10\text{MB} \times \text{number of tables}) + (256\text{MB} \times \text{sites per host})$$

If you do not know how many tables the database will contain or how many sites per host you expect to use, you can use 2 gigabytes as a rough estimate for the server process size for moderately sized databases and servers. But be aware that you may need to increase that estimate once your actual configuration is defined.

Finally, your estimate of the memory required for the server overall is the combination of the memory required for the content and the memory for the server process itself, plus 30% as a buffer.

$$\text{Server memory} = (\text{content} + \text{server process}) + 30\%$$

When sizing for a cluster, where the content is distributed across the servers, the calculation for the memory required for content on each server is the total content size divided by the number of servers, plus some percentage for replicated tables. For example, if 20% of the tables are replicated, a rough estimate of the space required for each server is given by the following equation:

$$\text{Per server memory} = ((\text{content} / \text{servers}) + 20\% + \text{server}) + 30\%$$

When sizing memory for VoltDB servers, it is important to keep in mind the following points:

- Memory usage includes not only storage for the data, but also temporary storage for processing transactions, managing queues, and the server processes themselves.
- Even in the best partitioning schemes, partitioning is never completely balanced. Make allowances for variations in load across the servers.

- If memory usage exceeds approximately 70% of total memory, the operating system can start paging and swapping, severely impacting performance.

Rule of Thumb

Keep memory usage per server within 50-70% of total memory.

Memory technology and density is advancing so rapidly, (similar to the increase in processor cores per server), it is feasible to configure a small number of servers with extremely large memory capacities that provide capacity and performance equivalent to a larger number of smaller servers. However, the amount of memory in use can impact the performance of other aspects of database management, such as snapshots and failure recovery. The next section discusses some of the trade offs to consider when sizing for these features.

3.4. Sizing for Durability

Durability refers to the ability of a database to withstand — or recover from — unexpected events. VoltDB has several features that increase the durability of the database, including K-Safety, snapshots, command logging, and database replication

K-Safety replicates partitions to provide redundancy as a protection against server failure. Note that when you enable K-Safety, you are replicating the unique partitions across the available hardware. So the hardware resources — particularly servers and memory — for any one copy are being reduced. The easiest way to size hardware for a K-Safe cluster is to size the initial instance of the database, based on projected throughput and capacity, then multiply the number of servers by the number of replicas you desire (that is, the K-Safety value plus one).

Rule of Thumb

When using K-Safety, configure the number of cluster nodes as a whole multiple of the number of copies of the database (that is, $K+1$).

K-Safety has no real performance impact under normal conditions. However, the cluster configuration can affect performance when recovering from a failure. In a K-Safe cluster, when a failed server rejoins, it gets copies of all of its partitions from the other members of the cluster. The larger (in size of memory) the partitions are, the longer they can take to be restored. Since it is possible for the restore action to block database transactions, it is important to consider the trade off of a few large servers that are easier to manage against more small servers that can recover in less time.

Two of the other durability features — snapshots and command logs — have only a minimal impact on memory and processing power. However, these features do require persistent storage on disk.

Most VoltDB disk-based features, such as snapshots, export overflow, network partitions, and so on, can be supported on standard disk technology, such as SATA drives. They can also share space on a single disk, assuming the disk has sufficient capacity, since disk I/O is interleaved with other work.

Command logging, on the other hand, is time dependent and must keep up with the transactions on the server. The chapter on command logging in *Using VoltDB* discusses in detail the trade offs between asynchronous and synchronous logging and the appropriate hardware to use for each. But to summarize:

- Use fast disks (such as battery-backed cache drives) for synchronous logging
- Use SATA or other commodity drives for asynchronous logging. However, it is still a good idea to use a dedicated drive for the command logs to avoid concurrency issues between the logs and other disk activity.

Rule of Thumb

When using command logging, whether synchronous or asynchronous, use a dedicated drive for the command logs. Other disk activity (including command log snapshots) can share a separate drive.

Finally, database replication (DR) does not impact the sizing for memory or processing power of the servers. But it does require duplicates of the initial hardware for each additional cluster. When using cross datacenter replication (XDCR) you will need one complete copy for each of the clusters participating in the XDCR relationship.

Rule of Thumb

When using database replication, multiply the number of servers needed by the number of clusters involved — two or more to match the number of clusters in a XDCR environment.

Chapter 4. Sizing Memory

An important aspect of system sizing is planning for the memory required to support the application. Because VoltDB is an in-memory database, allocating sufficient memory is vital.

Section 3.3, “Sizing for Capacity” provides some simple equations for estimating the memory requirements of a prospective application. If you are already at the stage where the database schema is well-defined and want more precise measurements of potential memory usage, this chapter provides details about how memory gets allocated.

For VoltDB databases, there are three aspects of memory sizing the must be considered:

- Understanding how much memory is required to store the data itself; that is, the contents of the database
- Evaluating how that data is distributed across the cluster, based on the proportion of partitioned to replicated tables and the K-safety value
- Determining the memory requirements of the server process

The sum of the estimated data requirements per server and the java heap size required by the server process provide the total memory requirement for each server.

4.1. Planning for Database Capacity

To plan effectively for database capacity, you must know in advance both the structure of the data and the projected volume. This means you must have at least a preliminary database schema, including tables, columns, and indexes, as well as the expected number of rows for each table.

It is often useful to write this initial sizing information down, for example in a spreadsheet. Your planning may even allow for growth, assigning values for both the initial volume and projected long-term growth. For example, here is a simplified example of a spreadsheet for a database supporting a flight reservation system:

Name	Type	Size	Initial Volume	Future Volume
Flight	Replicated table		5,000	20,000
- FlightByID	Index		5,000	20,000
- FlightByDepartTime	Index		5,000	20,000
Airport	Replicated table		10,000	10,000
- AirportByCode	Index		10,000	10,000
Reservation	Table		100,000	200,000
- ReservByFlight	Index		100,000	200,000
Customer	Table		200,000	1,000,000
- CustomerByID	Index		200,000	1,000,000
- CustomerByName	Index		200,000	1,000,000

Using the database schema, it is possible to calculate the size of the individual table records and indexes, which when multiplied by the volume projections gives a good estimate the the total memory needed to store the database contents. The following sections explain how to calculate the size column for individual table rows and indexes.

4.1.1. Sizing Database Tables

The size of individual table rows depends on the number and datatype of the columns in the table. For fixed-size datatypes, such as INTEGER and TIMESTAMP, the column is stored inline in the table record using the specified number of bytes. Table 4.1, “Memory Requirements For Tables By Datatype” specifies the length (in bytes) of fixed size datatypes.

For variable length datatypes, such as VARCHAR and VARBINARY, how the data is stored and, consequently, how much space it requires, depends on both the actual length of the data and the maximum possible length. If the maximum length is less than 64 bytes, the data is stored inline in the tuple as fixed-length data consuming the maximum number of bytes plus one for the length. So, for example, a VARCHAR(32 BYTES) column takes up 33 bytes, no matter how long the actual data is.

Note that VARCHAR columns can be declared in characters (the default) or bytes. For storage calculations, variable-length strings declared in characters are considered to consume 4 bytes for every character. In other words, a variable length string declared as VARCHAR(8) in characters consume the same amount of space as a string declared as VARCHAR(32 BYTES).

If the maximum length is 64 bytes or more, the data is stored in pooled memory rather than inline. To do this, there is an 8-byte pointer stored inline in the tuple, a 24-byte string reference object, and the space required to store the data itself in the pool. Within the pool, the data is stored as a 4-byte length, an 8-byte reverse pointer to the string reference object, and the data.

To complicate the calculation somewhat, data stored in pooled memory is not stored as arbitrary lengths. Instead, data is incremented to the smallest appropriate "pool size", where pool sizes are powers of 2 and intermediary values. In other words, pool sizes include 2, 4, 6 (2+4), 8, 12 (8+4), 16, 24 (8+16), 32 and so on up to a maximum of 1 megabyte for data plus 12 bytes for the pointer and length. For example, if the LastName column in the Customer table is defined as VARCHAR(32) (that is, a maximum length of 128 bytes) and the actual content is 95 bytes, the column consumes 160 bytes:

	8	Inline pointer
	24	String reference object
4		Data length
8		Reverse pointer
<u>95</u>		Data
107	128	Pool total / incremented to next pool size
	160	Total

Note that if a variable length column is defined with a maximum length greater than or equal to 64 bytes, it is not stored inline, even if the actual contents is less than 64 bytes. Variable length columns are stored inline only if the maximum length is less than 64 bytes.

Table 4.1, “Memory Requirements For Tables By Datatype” summarizes the memory requirements for each datatype.

Table 4.1. Memory Requirements For Tables By Datatype

Datatype	Size (in bytes)	Notes
TINYINT	1	
SMALLINT	2	
INTEGER	4	
BIGINT	8	
DOUBLE	8	
DECIMAL	16	
TIMESTAMP	8	
VARCHAR (<64 bytes)	maximum size + 1	Stored inline
VARBINARY (<64 bytes)	maximum size + 1	Stored inline
VARCHAR (>= 64 bytes)	32 + (actual size + 12 + padding)	Pooled resource. Total size includes an 8-byte inline pointer, a 24-byte reference pointer, plus the pooled resource itself.
VARBINARY (>=64 bytes)	32 + (actual size + 12 + padding)	Same as VARCHAR.

For tables with variable length columns less than 64 bytes, memory usage can be sized very accurately using the preceding table. However, for tables with variable length columns greater than 64 bytes, sizing is approximate at best. Besides the variability introduced by pooling, any sizing calculation must be based on an estimate of the *average length of data in the variable columns*.

For the safest and most conservative estimates, you can use the maximum length when calculating variable length columns. If, on the other hand, there are many variable length columns or you know the data will vary widely, you can use an estimated average or 90th percentile figure, to avoid over-estimating memory consumption.

4.1.2. Sizing Database Indexes

Indexes are sized in a way similar to tables, where the size and number of the index columns determine the size of the index.

VoltDB uses tree indexes. You can calculate the size of the individual index entries by adding up the size for each column in the index plus 40 bytes for overhead (pointers and lengths). The size of the columns are identical to the sizes when sizing tables, as described in Table 4.1, “Memory Requirements For Tables By Datatype”, with the exception of non-inlined binary data. For variable length columns equal to or greater than 64 bytes in length, the index only contains an 8-byte pointer; the data itself is not replicated.

So, for example, the CustomerByName index on the Customer table, which is an index containing the VARCHAR(32) fields LastName and FirstName, has a length of 56 bytes for each entry:

8	Pointer to LastName
8	Pointer to FirstName
40	Overhead
56	Total

The following equation summarizes how to calculate the size of an index.

$$(\text{sum-of-column-sizes} + 8 + 32) * \text{rowcount}$$

4.1.3. An Example of Database Sizing

Using the preceding formulas it is possible to size the sample flight database mentioned earlier. For example, it is possible to size the individual rows of the Flight table based on the schema columns and datatypes. The following table demonstrates the sizing of the Flight table.

Column	Datatype	Size in Bytes
FlightID	INTEGER	4
Carrier	VARCHAR(32)	160
DepartTime	TIMESTAMP	8
ArrivalTime	TIMESTAMP	8
Origin	VARCHAR(3 BYTES)	4
Destination	VARCHAR(3 BYTES)	4
Destination	VARCHAR(3 BYTES)	4
Total:		192

The same calculations can be done for the other tables and indexes. When combined with the expected volumes (described in Section 4.1, “Planning for Database Capacity”), you get a estimate for the total memory required for storing the database content of approximately 500 megabytes, as shown in the following table.

Name	Type	Size	Final Volume	Total Size
Flight	Replicated table	184	20,000	3,840,000
- FlightByID	Index	36	20,000	1,040,008
- FlightByDepartTime	Index	48	20,000	960,000
Airport	Replicated Table	484	10,000	4,840,000
- AirportByCode	Index	44	10,000	440,000
Reservation	Table	243	200,000	48,600,000
- ReservByFlight	Index	36	200,000	10,400,008
Customer	Table	324	1,000,000	324,000,000
- CustomerByID	Index	36	1,000,000	52,000,008
- CustomerByName	Index	56	1,000,000	56,000,000
Total:				502,120,024

4.2. Distributing Data in a Cluster

In the simplest case, a single server, the sum of the sizing calculations in the previous section gives you an accurate estimate of the memory required for the database content. However, VoltDB scales best in a cluster environment. In that case, you need to determine how much of the data will be handled by each server, which is affected by the number of servers, the number and size of partitioned and replicated tables, and the level of availability, or K-safety, required.

The following sections explain how to determine the distribution (and therefore sizing) of partitioned and replicated tables in a cluster and the impact of K-safety.

4.2.1. Data Memory Usage in Clusters

Although it is tempting to simply divide the total memory required by the database content by the number of servers, this is not an accurate formula for two reasons:

- Not all data is partitioned. Replicated tables (and their indexes) appear on all servers.
- Few if any partitioning schemes provide perfectly even distribution. It is important to account for some variation in the distribution.

To accurately calculate the memory usage per server in a cluster, you must account for all replicated tables and indexes plus each server's portion of the partitioned tables and indexes.

Data per server = replicated tables + (partitioned tables/number of servers)

Using the sample sizing for the Flight database described in Section 4.1.3, “An Example of Database Sizing”, the total memory required for the replicated tables and indexes (for the Flight and Airport tables) is only approximately 12 megabytes. The memory required for the remaining partitioned tables and indexes is approximately 490 megabytes. Assuming the database is run on two servers, the total memory required for data on each server is approximately 256 megabytes:

	12	Replicated data
2		Number of servers
490	<u>245</u>	Partitioned data total / per server
	256	Total

Of course, no partitioning scheme is perfect. So it is a good idea to provide additional space (say 20% to 30%) to allow for any imbalance in the partitioning.

4.2.2. Memory Requirements for High Availability (K-Safety)

The features you plan to use with a VoltDB database also impact capacity planning, most notably K-Safety. K-Safety improves availability by replicating data across the cluster, allowing the database to survive individual node failures.

Because K-Safety involves replication, it also increases the memory requirements for storing the replicated data. Perhaps the easiest way to size a K-Safe cluster is to size a non-replicated cluster, then multiply by the K-Safety value plus one.

For example, let's assume you plan to run a database with a K-Safety value of 2 (in other words, three copies) on a 6-node cluster. The easiest way to determine the required memory capacity per server is to calculate the memory requirements for a 2-node (non K-Safe) cluster, then create three copies of that hardware and memory configuration.

4.3. Planning for the Server Process (Java Heap Size)

The preceding sections explain how to calculate the total memory required for storing the database content and indexes. You must also account for the database process itself, which runs within the Java heap. Calculating the memory required by the server process both helps you define the total memory needed for planning purposes but also identifies the Java heap size that you will need to set in production when starting the database.

It is impossible to define an exact formula for the optimal heap size. But the following basic guidelines can be used to accommodate differing hardware configurations and database designs.

4.3.1. Attributes that Affect Heap Size

The database features that have the most direct impact on the server process and, therefore, the Java heap requirements are:

- Schema size, in terms of number tables and stored procedures
- The number of sites per host
- The features in use, specifically K-safety and/or database replication

The schema size affects the base requirements for the server process. The more tables the schema has and the more stored procedures it contains, the more heap space it will take up. In particular, it is important to provide enough heap so the schema can be updated, no matter what other features are enabled.

The general rule of thumb is a base Java heap size of 384MB, plus 10MB for every table in the schema. Stored procedures don't impact the heap size as much as the number of tables do. However, if you have lots of stored procedures (that is, hundreds or thousands of them) it is a good idea to add additional space to the base heap size for good measure.

Beyond the base heap size, use of K-safety and database replication each increases the requirements for Java heap space, with the increase proportional to the number of sites per host. In general each feature requires an additional 128MB for every site per host.

For example, a K-safe cluster with 4 sites per host requires an additional 512MB, while a K-safe cluster with 8 sites per host requires an extra gigabyte. If that cluster is also in an XDCR relationship, those extra heap requirements are doubled to 1GB and 2GB, respectively.

Note that the Java heap requirements for features are based solely on the sites per host, not the number of nodes in the cluster or the K-safety value. Any K-safety value greater than zero has the same requirements, in terms of the server process requirements.

4.3.2. Guidelines for Determining Java Heap Size

The recommended method for determining the appropriate heap size for a VoltDB cluster node is the following:

Step #1 Calculate the base Java heap requirement using the following formula:

$$384\text{MB} + (10\text{MB} \times \text{number of tables}) = \text{base Java heap size}$$

Be sure to allow for growth if you expect to add tables in the future. Also, if you expect to have large numbers of stored procedures (200 or more), increment the base heap size accordingly. Note that where memory is available, additional Java heap space beyond the minimum settings may provide performance benefits during operational events, such as schema updates and node rejoins.

Step #2 Based on the hardware to be used in production (specifically cores per server), and performance testing of the proposed schema, determine the optimal number of sites per host. 8 sites per host is the default. Setting the sites per host greater than 24 is not recommended.

Step #3 Determine which database features will be used in production. K-safety, network partition detection, and command logging are recommended for all production environments. Database replication (DR) is an optional feature that provides additional durability.

If K-safety is in use, but DR is not, multiply the number of sites per host by 128MB to determine the feature-specific Java heap requirements. If K-safety is enabled and the cluster is performing database replication (DR), multiply the number of sites per host by 256MB.

- Step #4 Add the base Java heap requirements defined in Step #1 to the feature-specific requirements in Step #3 to determine the recommended Java heap size for the production servers.

Chapter 5. Benchmarking

Benchmarking is the process of evaluating performance against a known baseline. For database applications, you can benchmark against other technologies, against a target metric (such as a specific number of transactions per second), or you can use the same application on different hardware configurations to determine which produces the best results.

For VoltDB applications, benchmarking is useful for establishing metrics with regards to:

- Optimum throughput
- Optimum number of sites per host

5.1. Benchmarking for Performance

When establishing performance criteria for database applications, there are usually two key metrics:

- **Throughput** — how many transactions can be completed at one time, usually measured in transactions per second, or TPS
- **Latency** — how long each individual transaction takes, usually measured as the average or percentile of the latency for a sample number of transactions

Note that neither of these metrics is exact. For many databases, throughput can vary depending upon the type of transaction; whether it is a read or write operation. One of the advantages of VoltDB is that throughput does not change significantly for write versus read operations. However, VoltDB throughput *does* change when there are many multi-partitioned transactions versus a single-partitioned transaction workload. This is why it is important to design your schema and stored procedures correctly when benchmarking a VoltDB database.

Similarly, latency can vary, both in how it is measured and what impacts it. You can measure latency as the time from when the client issues a transaction request until the response is received, or from when the database server receives the request until it queues the response.

The former measurement, latency from the client application's perspective, is perhaps the most accurate "real world" metric. However, this metric includes both database latency and any network latency between the client and the server. The latter measurement, latency from the database perspective, is a more accurate measurement of the technology's capability. This metric includes the time required to process the transaction itself (that is, the stored procedures and the database queries it contains) as well as time the request spends in the queue waiting to be executed.

5.1.1. Designing Your Benchmark Application

There is a relationship between throughput, latency, and system configuration. Throughput is a combination of the amount of time it takes to execute a transaction (which is part of latency) and the number of transactions that can be run in parallel (that is, the percentage of single-partitioned transactions plus the number of unique partitions, which is a combination of sites per host and number of servers).

This is why it is important that benchmarking for performance be done in conjunction with benchmarking for server configuration (as discussed later). Different configurations will result in different values for throughput and latency.

Your benchmark results are also affected by the design of the application itself.

5.1.1.1. Good Application Design

As mentioned before, throughput and latency are not abstract numbers. They are a consequence of the schema, stored procedures, application, and server configuration in use. As with any VoltDB solution, to affectively benchmark performance you must start with a well-designed application. Specifically:

- Partition all tables (except those that are small and primarily read-only)
- Make all frequent transactions single-partitioned
- Use asynchronous procedure calls

See Chapter 2, *Proof of Concept* for more information on writing effective VoltDB applications.

It is common practice to start benchmarking with the proof of concept, since the POC includes an initial schema and the key transactions required by the solution. If you decide to proceed with development, it is also a good idea to benchmark the application periodically throughout the process to make sure you understand how the additional functions impact your overall performance targets.

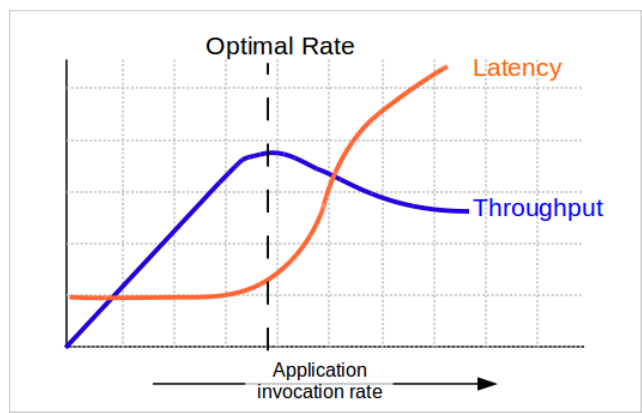
However, before benchmarking the POC, it is important to decide how you are going to measure performance. The following sections provide useful information for determining what and how to measure performance in a VoltDB application.

5.1.1.2. Rate Limiting

Latency and throughput are related. To measure throughput, your test application can make stored procedure calls as fast as it can and then measure how many are completed in a second (TPS). However, if the test application invokes more transactions than the Volt client API can transmit in the time frame, or the server can complete in that time frame, the additional invocations must wait in a queue before being processed. The time these transactions wait in the queue will result in the latency measurement suddenly spiking.

There is an optimal throughput for any given application and system configuration. As the rate of invocations from the client application increases, there is a direct increase in TPS while latency stays relatively flat. However, as the invocation rate approaches and surpasses the optimal throughput (the limit of what the current configuration can process), latency increases dramatically and TPS may even drop, as shown in Figure 5.1, “Determining Optimal Throughput and Latency”.

Figure 5.1. Determining Optimal Throughput and Latency



If your test application “fire hoses” the database server — that is, it sends invocations as fast as it can — all you can measure is the misleading throughput and latency on the right side of the preceding chart.

To determine the optimal rate, you need to be able control the rate at which your benchmark application submits transaction requests. This process is called *rating limiting*.

At its simplest, rate limiting is constraining the number of invocations issued by the application. For example, the following program loop, constrains the application to invoking the `SignIn` stored procedure a maximum of 10 times per millisecond, or 10,000 times per second.

```
boolean done = false;
long maxtxns = 10;
while (!done) {
    long txns = 0;
    long millisecs = System.currentTimeMillis();
    while (millisecs == System.currentTimeMillis()) {
        if (txns++ < maxtxns) {
            myClient.callProcedureAsync("SignIn", id, millisecs)
                .whenComplete((resp, th) -> signInComplete(resp, th));
        }
    }
}
```

You could use a command line argument to parameterize the rate limit value `maxtxns` and then use multiple runs of the application to create a graph similar to Figure 5.1, “Determining Optimal Throughput and Latency”. An even better approach is to use a limit on latency to automatically control the invocation rate and let the application close in on the optimal throughput value.

How rate limiting based on latency is done is to have a variable for the target latency as well as a variable for maximum allowable throughput (such as `maxtxns` in the preceding example). The application measures both the average throughput and latency for a set period (every 1 to 5 seconds, for example). If the average latency exceeds the goal, reduce the maximum transactions per second, then repeat. After the same period, if the latency still exceeds the goal, reduce the maximum transaction rate again. If the latency *does* meet the goal, incrementally *increase* the maximum transaction rate.

By using this mix of rate limiting and automated adjustment based on a latency goal, the test application will eventually settle on an optimal throughput rate for the current configuration. This is the method used by the sample applications (such as `voter`).

5.1.2. How to Measure Throughput

Normally for benchmarking it is necessary to “instrument” the application. That is, add code to measure and report on the benchmark data. Although it is possible (and easy) to instrument a VoltDB application, it is not necessary.

The easiest way to measure throughput for a VoltDB database is to monitor the database while the benchmark application is running. You can monitor a database using the VoltDB Management Console, which is available from any VoltDB server and provides a graphical display of the overall throughput for the database cluster.

For more detailed information, you can instrument your application by using a variable to track the number of completed transactions (incrementing the variable in the asynchronous procedure callback). You then periodically report the average TPS by dividing the number of transactions by the number of seconds since the last report. This approach lets you configure the reporting to whatever increment you choose.

See the `voter` sample application for an example of an instrumented benchmarking application. The README for the `voter` application explains how to customize the reporting through the use of command line arguments.

5.1.3. How to Measure Latency

It is also possible to get a sense of the overall latency without instrumenting your application using the VoltDB Management Console. The latency graph provided by the Management Console shows the average latency, measured every few seconds, as reported by the server. At times, this method can produce a dramatic sawtooth graph with extreme highs and lows. In this case, the best way to interpret the overall latency of the application is to imagine a line drawn across the high points of the graph.

For a more accurate benchmark, you can use latency metrics built into the VoltDB system procedures and Java client interface. To instrument your client application, you can use the `ClientResponse` interface returned by stored procedures invocations to measure latency. Part of the `ClientResponse` class are the `getClientRoundTrip()` and `getClusterRoundTrip()` methods. Both methods return an integer value representing the number of milliseconds required to process the transaction. The `getClientRoundTrip()` method reports total latency, including the network round trip from the client application to the server. The `getClusterRoundTrip()` method reports an estimate of the latency associated with the database only.

It is easy to use these methods to collect and report custom latency information. For example, the following code fragment captures information about the minimum, maximum, and combined latency during each procedure completion callback.

```
public void requestComplete(ClientResponse response, Throwable th) {  
  
    txncount++;  
    int latency = response.getClusterRoundTrip();  
    if (latency < minlatency) minlatency = latency;  
    if (latency > maxlatency) maxlatency = latency;  
    sumlatency += latency;  
    .  
    .  
    .  
}
```

The benchmarking application can then use this information to periodically report specific latency values, like so:

```
if (System.currentTimeMillis() >= nextreport ) {  
    // report latency  
    printf("Min latency: %d\n" +  
        "Max latency: %d\n" +  
        "Average latency: %d\n",  
        minlatency, maxlatency, sumlatency/txncount);  
    // reset variables  
    txncount=0;  
    minlatency = 5000;  
    maxlatency = 0;  
    sumlatency = 0;  
    // report every 5 seconds  
    nextreport = System.currentTimeMillis() + 5000;  
}
```

5.2. Determining Sites Per Host

Another important goal for benchmarking is determining the optimal number of sites per host. Each VoltDB server can host multiple partitions, or "sites". The ideal number of sites per host is related to the number

of processor cores available on the server. However, it is not an exact one-to-one relationship. Usually, the number of sites is slightly lower than the number of cores.

The equation becomes even more complex with hyperthreading, which "virtualizes" multiple processors for each physical core. Hyperthreading can improve the number of sites per host that a VoltDB server can support. But again, not in direct proportion to a non-hyperthreaded server.

Important

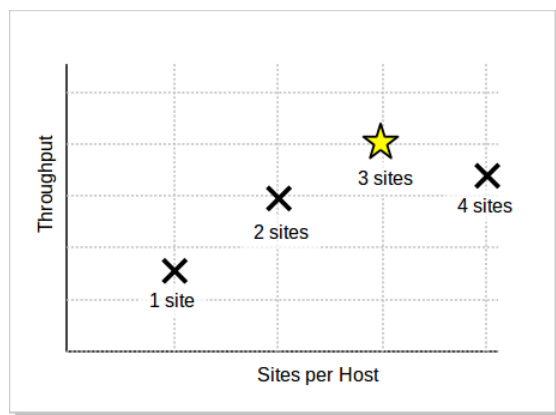
In general, VoltDB performs best with between 4 and 16 sites per host. However, you should never exceed 24 sites per host, even if the number of processor cores might support more, because the processing required to manage so many sites begins to conflict with the data processing.

The easiest way to determine the optimal number of sites per host is by testing, or benchmarking, against the target application. The process for determining the correct number of sites for a specific hardware configuration is as follows:

1. Create a benchmark application that measures the optimal throughput, as described in Section 5.1, "Benchmarking for Performance".
2. Run the benchmark application multiple times, each time increasing the number of sites per host for the database.
3. Make note of the optimal throughput for each run and graph the optimal TPS against the number of sites per host.

As the number of sites per host increases, the optimal throughput increases as well. However, at some point, the number of sites exceeds the number of threads that the hardware can support, at which point the throughput levels out, or even decreases, as contention occurs in the processor. Figure 5.2, "Determining Optimal Sites Per Host" shows the results graph of a hypothetical benchmark of sites per host for a quad core server. In this case, the optimal number of sites per host turned out to be three.

Figure 5.2. Determining Optimal Sites Per Host



By graphing the relationship between throughput and partitions using a benchmark application, it is possible to maximize database performance for the specific hardware configuration you will be using.