

Distributed Big Table for Relational Database

Taseem Nazir¹, Khalil Ahmed², Aaqib Hassan Bulla³, Junaid Farooq War⁴, ⁵Haider Mehraj

^{1,2}Department of Computer Science and Engineering, SoET, BGSSBU

^{3, 4, 5} Department of Electronics and Communication Engineering, SoET, BGSSBU

Abstract—In order to change the underlying file system storage and distribute the data among numerous independent machines, big tables are widely used. The main aim of the paper is to provide dynamic control to clients over data layout and design and to provide efficient storage and retrieval of data without disturbing the availability and performance.

Keywords: Big table, Dynamic Control, Availability, Performance, Storage

I. INTRODUCTION

Big table is a storage system to manage the data that is widely distributed amongst multiple servers. The tasks at MNC'S that work with data store data in Big table, including web indexing. These applications place very different demands on Big table, both in terms of data size and latency requirements. Despite these varied demands, Big table has successfully provided high-performance solution for all. Big table works on the principle of partition and replication of data .In case of partition, data is bifurcated on the basis of a key and stored on multiple servers and in case of replication, data is replicated among multiple servers. Big table is different from relational database management system where in spite of rows, columns are stored in sequential storage locations. A Big table holds numerous tables with each table comprising of further set of tablets, and each tablet contains all data associated with a row range. At first, each table consists of just one tablet. When the size of the table increases, it is automatically split into further more tablets of fixed size. In contrast to database management system, big table depends on cluster management system to meet all the functionalities of resource management and failures.

II. LOCKS AND IMPLEMENTATION

Big table depends upon the availability and consistency of lock services. Locks are used to ensure –there is at most one transaction writing any data in a tablet, reads and writes in a tablet are atomic .The Big table implementation has three major components: a library that is associated with every client, A master server and numerous tablet servers. In order to accommodate changes in workload, we can add or remove the tablet servers. Balancing the server load, changes in the schema, garbage collection in the file system are done by the master server. Bifurcation of large tables, handling read and write requests and management of set of tables is done by the tablet server. The communication overhead between the client and master is quite low leading to less overhead of master servers and thus helps in efficiency.

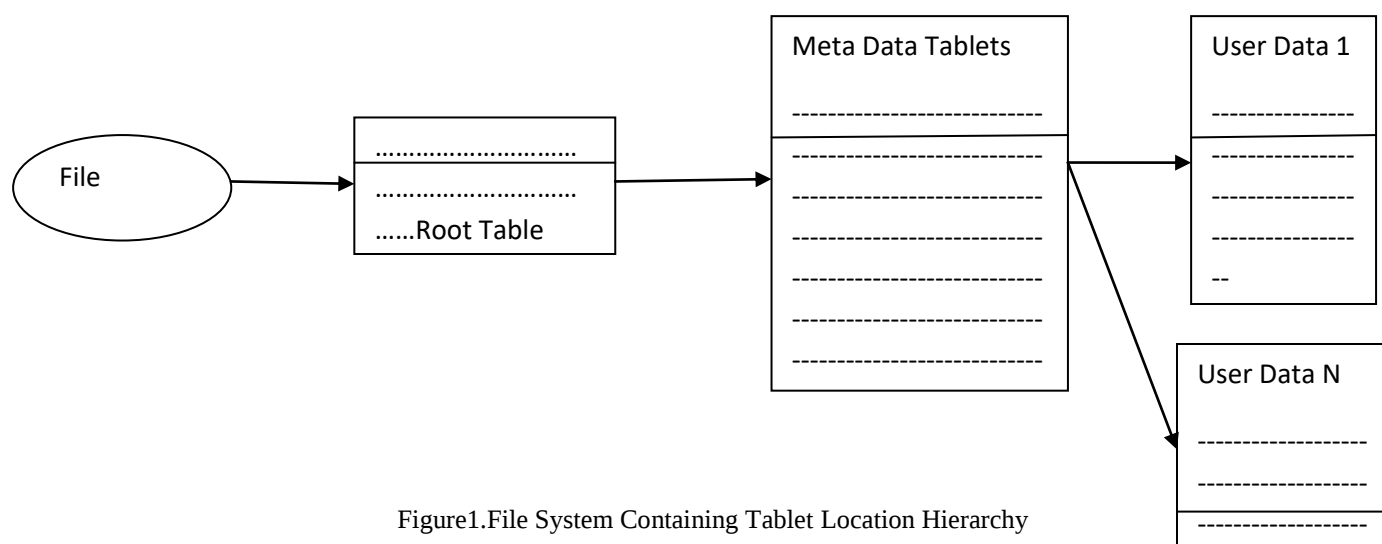


Figure1.File System Containing Tablet Location Hierarchy

III. REFINEMENTS

The implementation described in the previous section required a number of refinements to achieve the high performance, availability, and reliability required by users. The refinements could be done on various basis like:

A. Club The Data Locally

Clients can group multiple column families together into a *group*. A separate SSTable is generated for each locality group in each tablet. Segregating column families that are not typically accessed together into separate locality groups enables more efficient reads. For example, page metadata in Webtable (like that of a language and checksums) can be in one locality group, and the contents of the page can be in a different group.

B. Compression

Clients can control whether or not the SSTables for a locality group are compressed, and if so, which compression format is used. The user-specified compression format is applied to each SSTable block where size is controlled by local group parameters. Although we lose some space by compressing each block separately, we benefit in that small portions of an SSTable can be read without decompressing the entire file. Many clients use a two-pass custom compression scheme. The first pass uses Bentley and McIlroy's scheme [3], which compresses long common strings across a large window. The second pass uses a fast compression algorithm that looks for repetitions in a small 16 KB window of the data. Both compression passes are very fast. They encode at 100-200 MB/s, and decode at 400-1000 MB/s on modern machines. Even though we emphasized speed instead of space reduction when choosing our compression algorithms, this two-pass compression scheme does surprisingly well.

C. Caching for read performance

To improve read performance, tablet servers use two levels of caching. The Scan Cache is a higher-level cache that caches the key-value pairs returned by the SSTable interface to the tablet server code. The Block Cache is a lower-level cache that caches SSTable blocks that were read from GFS. The Scan Cache is most useful for applications that tend to read the same data repeatedly. The Block Cache is useful for applications that tend to read data that is close to the data they recently read.

IV. PERFORMANCE

N-Tablet servers are used to measure the performance. The configuration of servers made use of 1GB of memory, 1500 machines. The same number of client machines were used in order to ensure that there is no starvation in terms of client tablets. The tablet servers and master, test clients, and GFS servers all ran on the same set of machines. Every machine ran a GFS server. Some of the machines also ran either a tablet server, or a client process, or processes from other jobs that were using the pool at the same time as these experiments. R is the distinct number of Bigtable row keys involved in the test. R was chosen so that each benchmark read or wrote approximately 1 GB of data per tablet server. The *sequential write* benchmark used row keys with names 0 to $R - 1$. This space of row keys was partitioned into 10N equal-sized ranges. These ranges were assigned to the N clients by a central scheduler that assigned the next available range to a client as soon as the client finished processing the previous range assigned to it. This dynamic assignment helped mitigate the effects of performance variations caused by other processes running on the client machines. A string under single row was written, each string was generated with the help of random function and was therefore uncompressible. In order to avoid any cross-row compression, strings under different row key were distinct. However, performance does not increase linearly. For most benchmarks, there is a significant drop in per-server throughput when going from 1 to 50 tablet servers. This drop is caused by imbalance in load in multiple server configurations, often due to other processes contending for CPU and network. Our load balancing algorithm attempts to deal with this imbalance, but cannot do a perfect job for two main reasons: rebalancing is throttled to reduce the number of tablet movements (a tablet is unavailable for a short time, typically less than one second, when it is moved), and the load generated by our benchmark shifts around as the benchmark progresses.

V. CONCLUSION

C-Store and Bigtable share many characteristics: both systems use a shared-nothing architecture and have two different data structures, one for recent writes, and one for storing long-lived data, with a mechanism for moving data from one form to the

other. The systems differ significantly in their API: C-Store behaves like a relational database, whereas Bigtable provides a lower level read and write interface and is designed to support many thousands of such operations per second per server. C-Store is also a .read-optimized relational DBMS., whereas Bigtable provides good performance on both read-intensive and write-intensive applications. Big table's load balancer has to solve some of the same kinds of load and memory balancing problems faced by shared-nothing databases[1][2]. Our problem issomewhat simpler:

- A. There is no consideration for the possibility of multiple copies of the same data, possibly in alternate forms due to views or indices;
- B. Let the user describe what data belongs in memory and what data should stay on disk, rather than trying to determine this dynamically;
- C. There is no complex queries to execute or optimize.

REFERENCES

- [1] STONEBRAKER, M., AOKI, P. M., DEVINE, R., LITWIN, W., AND OLSON, M. A. Mariposa: A new architecture for distributed data. In *Proc. of the Tenth ICDE*(1994), IEEE Computer Society, pp. 54.65.
- [2]. COPELAND, G. P., ALEXANDER, W., BOUGHTER, E. E., AND KELLER, T. W. Data placement in Bubba. In *Proc. of SIGMOD* (1988), pp. 99.108.
- [3]. BENTLEY, J. L., AND MCILROY, M. D. Data compression using long common strings. In *Data Compression Conference* (1999), pp. 287.295.
- [4]. ABADI, D. J., MADDEN, S. R., AND FERREIRA, M. C. Integrating compression and execution in column oriented database systems. *Proc. of SIGMOD* (2006).
- [5]. Bigtable: A Distributed Storage System for Structured Data Fay Chang, Jeffrey Dean, Sanjay Ghemawat, Wilson C. Hsieh, Deborah A. Wallach Mike Burrows, Tushar Chandra, Andrew Fikes, Robert E. Gruber
- [6] Surabhi Mittal and Gaurav Tripathi, "Role of Mobile Phone Technology in Improving Small Farm Productivity", *Agricultural Economics Research Review*, Vol. 22, pp. 451-459, 2009.