

論文 / 著書情報
Article / Book Information

題目(和文)	スケーラブルなデータストアにおけるクエリ適応的な多次元索引に関する研究
Title(English)	Query-Adaptive Multidimensional Indexing on Scalable Data Stores
著者(和文)	西村祥治
Author(English)	Shoji Nishimura
出典(和文)	学位:博士(工学), 学位授与機関:東京工業大学, 報告番号:甲第10874号, 授与年月日:2018年3月26日, 学位の種別:課程博士, 審査員:横田 治夫,宮崎 純,権藤 克彦,吉瀬 謙二,金子 晴彦
Citation(English)	Degree:Doctor (Engineering), Conferring organization: Tokyo Institute of Technology, Report number:甲第10874号, Conferred date:2018/3/26, Degree Type:Course doctor, Examiner:,,,,
学位種別(和文)	博士論文
Category(English)	Doctoral Thesis
種別(和文)	論文要旨
Type(English)	Summary

論文要旨

THESIS SUMMARY

専攻： 計算工学 専攻
Department of
学生氏名： 西村 祥治
Student's Name

申請学位 (専攻分野)： 博士 (工学)
Academic Degree Requested Doctor of
指導教員 (主)： 横田 治夫
Academic Supervisor(main)
指導教員 (副)：
Academic Supervisor(sub)

要旨 (英文 800 語程度)
Thesis Summary (approx.800 English Words)

Recently, massive data management plays an increasingly important role in data analytics because data access is a major bottleneck. Scalable data stores are widely applied as a data management infrastructure. The data stores should have the unlimited scalability on the amount of data and data insertion throughput by just adding servers. The data stores also should provide efficient simple operations such as key lookup and key range access. However, they cannot natively support multidimensional queries, which are fundamental functions of data analytics.

Multidimensional indexing is a solid method to achieve efficient multidimensional query processing. In the context of scalable data stores, the index itself should be scalable. Moreover, the scalable data stores distribute data among multiple servers and enlarge management unit of data due to huge volume of data. Therefore, data partitioning and data skipping are critical to maximize query performance.

This doctoral dissertation thus proposes a solution to introduce multidimensional indexing layer on top of scalable data stores, which includes optimization techniques of data partitioning and data skipping for prioritized workloads. The dissertation discovers notable properties of space-filling curves and highly utilizes them to realize efficient multidimensional queries on top of scalable data stores.

The dissertation first introduces MD-HBase, a method to construct a multidimensional indexing layer over range-partitioned data stores. The indexing layer provides equivalent functionality of the kd-tree, a traditional multidimensional index. The method applies Z-curve, a space-filling curve, to translate a multidimensional data to a key. The method emulates the same space-splitting as the kd-tree by letting the underlying data store choose the longest common prefix of keys within the page as a range-partitioning pivot. The proposed method thus realizes a scalable multidimensional index layer by basic operations provided by the underlying data stores. The dissertation also proposes efficient algorithms for multidimensional range query and k-nearest neighbor query on top of the index layer. The algorithms eliminate any false positive scans in the index layer by subdividing the input query recursively into sub-queries exactly as the space is partitioned by the index structure.

The dissertation then proposes a data partitioning framework named QUILTS to achieve near-optimal range partitioning based on well-designed space-filling curves. Optimal data partitioning minimizes the number of page accesses during query processing. However, it is impossible to achieve optimal against any query patterns. Moreover, it is known as an NP-hard problem to achieve optimal against a specific query pattern and data distribution. The dissertation thus proposes an approximate method to reduce the number of page accesses for given query patterns over the range partitioning based on space filling curve. The problem can be interpreted as selecting a curve to achieve near-optimal partitioning. The dissertation proposes a model to measure how a space-filling curve fits a given query pattern. The key insight is that the measurement should change in the case of the sparse or dense distribution. Through the analysis on the model, the C-curve, also known as the composite index, and the Z-curve have opposite properties on a given query pattern and data distribution. The dissertation proposes a bit-merging curve family, which is a generalization of the C- and Z-curves, to synthesize a hybrid curve that inherits the both curve properties. Finally, the dissertation proposes a heuristic method to design a query-aware and skew-tolerant curve that achieves near-optimal data partitioning.

The dissertation also proposes a data skipping method named multidimensional range filter to achieve efficient pre-computation with small foot print. The core problem is to answer the data existence within a given range. There is a trade-off between space and false-positive ratio to give a correct answer. The dissertation analyzes the cause of high false-positive ratio in the multi dimension case and then identifies two reasons: the data sparsity and query pattern awareness. The dissertation introduces a special handling mechanism for sparse region to the data structure. For the query pattern awareness, the above curve design method can be applied to improve the false-positive ratio. The dissertation implements a noble range

management mechanism that controls granularity of range management adapting to query patterns and query frequency.

The dissertation prototypes the multidimensional index layer on top of a range-partitioned key-value store. The prototype system achieves efficient multidimensional range query and k-nearest neighbor query without discouraging the scalability of the underlying data store such as insert throughput. The dissertation also implements two optimization techniques on top of the prototype systems and conducts extensive experiments for data warehousing (DWH) and geographic information systems (GIS) applications with real-world data. The proposed partitioning method reduces the number of page accesses by an order of magnitude. The proposed method also reduces the false-positive ratio to 10-20% on the both applications while a traditional method achieves at least 60%.

The works in the dissertation contribute directly for designing the scalable data management system that supports efficient multidimensional queries and leading to the data analytics infrastructure for the Internet scale data.

備考：論文要旨は、和文 2000 字と英文 300 語を 1 部ずつ提出するか、もしくは英文 800 語を 1 部提出してください。

Note：Thesis Summary should be submitted in either a copy of 2000 Japanese Characters and 300 Words (English) or 1copy of 800 Words (English).

注意：論文要旨は、東工大リサーチリポジトリ(T2R2)にてインターネット公表されますので、公表可能な範囲の内容で作成してください。

Attention: Thesis Summary will be published on Tokyo Tech Research Repository Website (T2R2).