



R*-Grove: Balanced Spatial Partitioning for Large-Scale Datasets

Tin Vu and Ahmed Eldawy*

Department of Computer Science and Engineering, University of California, Riverside, Riverside, CA, United States

OPEN ACCESS

Edited by:

Liyue Fan,
University at Albany, United States

Reviewed by:

Keke Chen,
Wright State University, United States
Arnab Bhattacharya,
Indian Institute of Technology
Kanpur, India
Jia Yu,
Arizona State University, United States

*Correspondence:

Ahmed Eldawy
eldawy@ucr.edu

Specialty section:

This article was submitted to
Data Mining and Management,
a section of the journal
Frontiers in Big Data

Received: 25 December 2019

Accepted: 24 July 2020

Published: 28 August 2020

Citation:

Vu T and Eldawy A (2020) R*-Grove:
Balanced Spatial Partitioning for
Large-Scale Datasets.
Front. Big Data 3:28.
doi: 10.3389/fdata.2020.00028

The rapid growth of big spatial data urged the research community to develop several big spatial data systems. Regardless of their architecture, one of the fundamental requirements of all these systems is to spatially partition the data efficiently across machines. The core challenges of big spatial partitioning are building high spatial quality partitions while simultaneously taking advantages of distributed processing models by providing load balanced partitions. Previous works on big spatial partitioning are to reuse existing index search trees as-is, e.g., the R-tree family, STR, Kd-tree, and Quad-tree, by building a temporary tree for a sample of the input and use its leaf nodes as partition boundaries. However, we show in this paper that none of those techniques has addressed the mentioned challenges completely. This paper proposes a novel partitioning method, termed R*-Grove, which can partition very large spatial datasets into high quality partitions with excellent load balance and block utilization. This appealing property allows R*-Grove to outperform existing techniques in spatial query processing. R*-Grove can be easily integrated into any big data platforms such as Apache Spark or Apache Hadoop. Our experiments show that R*-Grove outperforms the existing partitioning techniques for big spatial data systems. With all the proposed work publicly available as open source, we envision that R*-Grove will be adopted by the community to better serve big spatial data research.

Keywords: big spatial data, partitioning, R*-Grove, index optimization, query processing

1. INTRODUCTION

The recent few years witnessed a rapid growth of big spatial data collected by different applications such as satellite imagery (Eldawy et al., 2015b), social networks (Magdy et al., 2014), smart phones (Henke et al., 2016), and VGI (Goodchild, 2007). Traditional Spatial DBMS technology could not scale up to these petabytes of data which led to the birth of many big spatial data management systems such as SpatialHadoop (Eldawy and Mokbel, 2015), GeoSpark (Yu et al., 2015), Simba (Xie et al., 2016), LocationSpark (Tang et al., 2016), and Sphinx (Eldawy et al., 2017), to name a few.

Regardless of their architecture, all these systems need an essential preliminary step that partitions the data across machines before the execution can be parallelized. This is also known as *global indexing* (Eldawy and Mokbel, 2016). A common method that was first introduced in SpatialHadoop (Eldawy and Mokbel, 2015), is the sample-based STR partitioner. This method picks a small sample of the input to determine its distribution, packs this sample using the STR packing algorithm (Leutenegger et al., 1997), and then uses the boundaries of the leaf nodes to partition the entire data. **Figure 1A** shows an example of an STR-based partitioning where each data partition is depicted by a rectangle. The method was later generalized by replacing the STR bulk loading algorithm with other spatial indexes such as Quad-tree (Samet, 1984), Kd-Tree, and Hilbert R-trees

(Kamel and Faloutsos, 1994; Eldawy et al., 2015a). That STR-based partitioning was very attractive due to its simplicity and good load balancing which is very important for distributed applications. Its simplicity urged many other researchers to adopt it in their systems such as GeoSpark (Yu et al., 2015) and Simba (Xie et al., 2016) for in-memory distributed processing; Sphinx (Eldawy et al., 2017) for SQL big spatial data processing; HadoopViz (Eldawy et al., 2016; Ghosh et al., 2019) for scalable visualization of big spatial data; and in distributed spatial join (Sabek and Mokbel, 2017).

Despite their wide use, the existing partitioning techniques all suffer from one or more of the following three limitations. First, some partitioning techniques (STR, Kd-tree) prioritize load balance over spatial quality which results in suboptimal partitions. This is apparent in **Figure 1A** where the thin and wide partitions result in low overall quality for the partitions since square-like partitions are preferred for most spatial queries. Square-like partitions are preferred in indexing because they indicate that the index is not biased toward one dimension. Also, since most queries are shaped like a square or a circle, square-like partitions would minimize the overlap with the queries (Beckmann et al., 1990). Second, they could produce partitions that do not fill the HDFS blocks in which they are stored. Big data systems are optimized to process full blocks, i.e., 128 MB, to offset the fixed overhead in processing each block. However, the index structures used in existing partitioning techniques, e.g., R-trees, Kd-tree, Quad-tree, produce nodes with number of records in the range $[m, M]$, where $m \leq M/2$. In practice, m can be as low as $0.2M$ (Beckmann et al., 1990; Beckmann and Seeger, 2009). While those underutilized index nodes were desirable for disk indexing as they can accommodate future inserts, they result in underutilized blocks as depicted in **Figure 1A** where all blocks are $<80\%$ full. Moreover, this design might also produces poor load balance among partitions due to the wide range of partition sizes. Third, all existing partitioning techniques rely on a sample and try to balance the number of records per partition. This resembles traditional indexes where the index contains record IDs. However, in big spatial data partitioning, the entire record is written in each partition, not just its ID. When records are highly variant in size, all existing techniques end up with extremely unbalanced partitions.

This paper proposes a novel spatial partitioning technique for big data, termed R*-Grove, which completely addresses all of three aforementioned limitations. First, it produces high quality partitions by utilizing the R*-tree optimization techniques (Beckmann et al., 1990) which aim at minimizing the total area, overlap area, and margins. The key idea of the R*-Grove partitioning technique is to start with one partition that contains all sample points and then use the node split algorithm of the R*-tree to split it into smaller partitions. This results in compact square-like partitions as shown in **Figure 1B**. Second, in order to ensure that we produce full blocks and balanced partitions, R*-Grove introduces a new constraint that puts a lower bound on the ratio between the smallest and the largest block, e.g., 95%. This property is theoretically proven and practically validated by our experiments. Third, when the input records have variable sizes, R*-Grove combines a data size histogram with the sample

points to assign a weight for each sample point. These weights are utilized to guarantee that the size of each partition falls in a user-defined range.

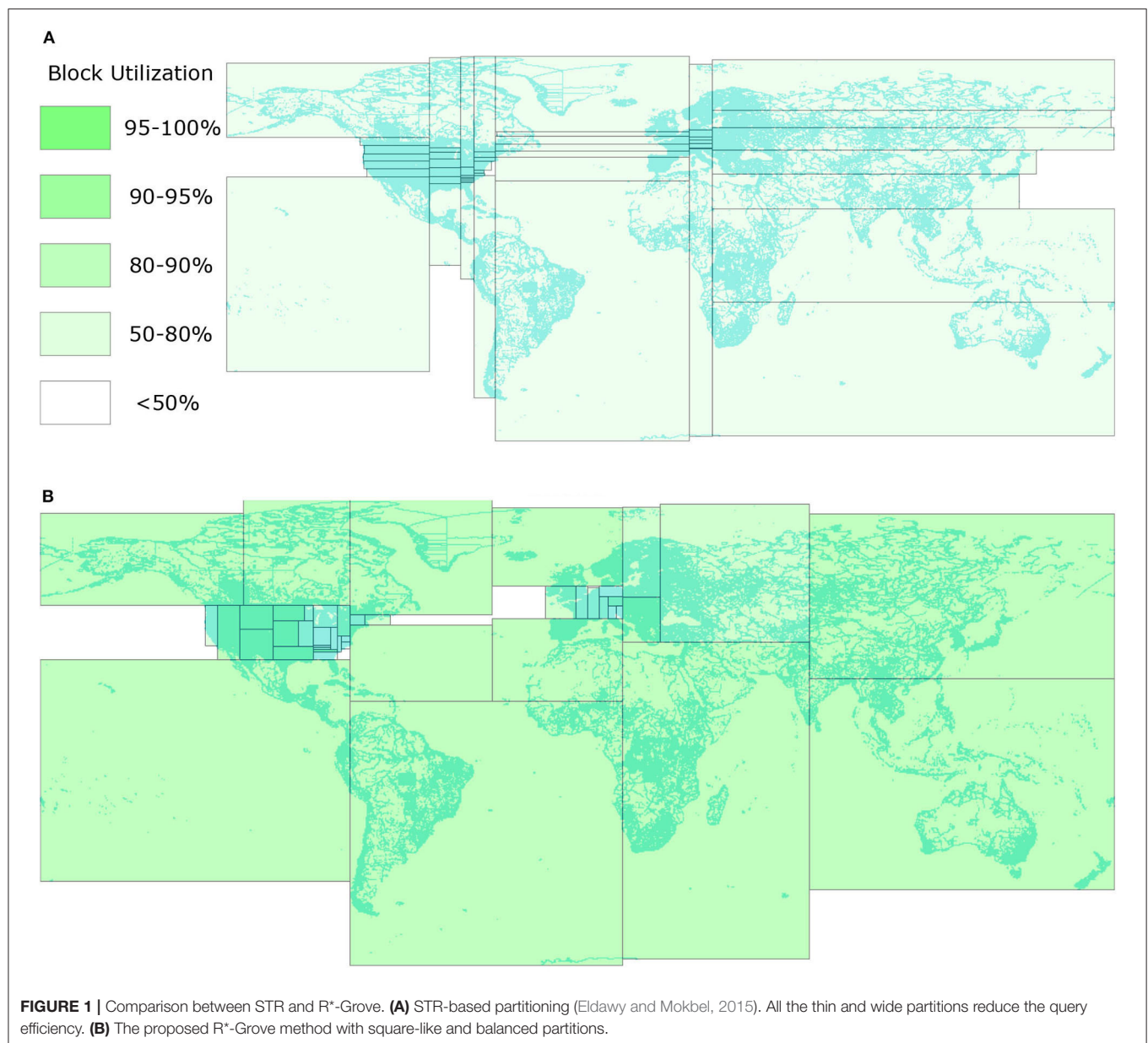
Given the wide adoption of the previous STR-based partitioner, we believe the proposed R*-Grove will be widely used in big spatial data systems. This impacts a wide range of spatial analytics and processing algorithms including indexing (Vo et al., 2014; Eldawy et al., 2015a), range queries (Eldawy and Mokbel, 2015; Yu et al., 2015), kNN queries (Eldawy and Mokbel, 2015), visualization (Eldawy et al., 2016; Ghosh et al., 2019), spatial join (Jacox and Samet, 2007), and computational geometry (Eldawy et al., 2013; Li et al., 2019). All the work proposed in this paper is publicly available as open source and supports both Apache Spark and Apache Hadoop. We run an extensive experimental evaluation with up-to 500 GB and 7 billion record datasets and up-to nine dimensions. The experiments show that R*-Grove consistently outperforms existing STR-based, Z-curve-based, Hilbert-Curve-based, and Kd-tree-based techniques in both partitions quality and query efficiency.

The rest of this paper is organized as follow. Section 2 describes the related works. Section 3 gives a background about big spatial data partitioning. Section 4 describes the proposed R*-Grove technique. Section 5 describes the advantages of R*-Grove in popular case studies of big spatial data systems. Section 6 gives a comprehensive experimental evaluation of the proposed work. Finally, section 7 concludes the paper.

2. RELATED WORK

This section discusses the related work in big spatial data partitioning. In general, distributed indexes for big spatial data are constructed in two levels, one global index that partitions the data across machines, and several local indexes that organize records in each partition. Previous work (Lu et al., 2014; Eldawy and Mokbel, 2015, 2016) showed that the global index provides far much improvement than local indexes. Therefore, in this paper we focus on global indexing and it can be easily combined with any of the existing local indexes. The work in global indexing can be broadly categorized into three approaches, namely, sampling-based methods, space-filling-curve (SFC)-based methods, and quad-tree-based methods.

The *sampling-based* method picks a small sample from the input data to infer its distribution. The sample is loaded into an in-memory index structure while adjusting the data page capacity, e.g., leaf node capacity, such that the number of data pages is roughly equal to the desired number of partitions. The order of sample objects does not affect the partition quality, since the sample is uniformly taken from the entire input dataset. Furthermore, most of algorithms sort the data as part of the partitioning process so the original order is completely lost. Some R-tree bulk-loading algorithms (STR Leutenegger et al., 1997 or OMT Lee and Lee, 2003) can also be used to speed up the tree construction time. Then, the minimum bounding rectangles (MBRs) of the data pages are used to partition the entire dataset. This method was originally proposed for spatial join and denoted the seeded-tree (Lo and Ravishankar, 1994). It was then used for



big spatial indexing in many systems including SpatialHadoop (Eldawy and Mokbel, 2015; Eldawy et al., 2015a), Scala-GiST (Lu et al., 2014), GeoSpark (Yu et al., 2015), Sphinx (Eldawy et al., 2017), Simba (Xie et al., 2016), and many other systems. This technique can be used with existing R-tree indexes but it suffers from two limitations, load imbalance and low quality of spatial partitions. Additionally, when there is a big variance in record sizes, the load imbalance is further amplified due to the use of the sample. We will further discuss these limitations in section 4.

The *SFC-based* method builds a spatial index on top of an existing one-dimensional index by applying any space-filling curve, e.g., Z-curve or Hilbert curve. MD-HBase (Nishimura et al., 2013) builds Kd-tree-like and Quad-tree-like indexes on top of HBase by applying the Z-curve on the input data and

customizing the region split method in HBase to respect the structure of both indexes. GeoMesa (Fox et al., 2013) uses geohashing which is also based on the Z-curve to build spatio-temporal indexes on top of Accumulo. Unlike MD-HBase which only supports point data, GeoMesa can support rectangular or polygonal geometries by replicating a record to all overlapping buckets in the geohash. While this method can ensure a near-perfect load balance, it produces an even bigger spatial overlap between partitions as compared to the sampling-based approach described above. This drawback leads to the inefficient performance of spatial queries.

The *quad-tree-based* method relies heavily on the Quad-tree structure to build efficient and scalable Quad-tree index in Hadoop (Whitman et al., 2014). It starts by splitting the input

data into equi-sized chunks and building a partial Quad-tree for each split. Then, it combines the leaf nodes of the partial trees based on the Quad-tree structure to merge them into the final tree. While highly efficient, this method cannot generalize to other spatial indexes and is tightly tied to the Quad-tree structure. In addition, this Quad-tree-based partitioning tends to produce much more than the desired number of partitions which also leads to load imbalance.

Although there are several partitioning techniques for large-scale spatial data as mentioned above, sampling-based method is the most ubiquitous option, which is integrated in most of existing spatial data systems. Sampling-based methods are preferred as they are simple to implement and provide very good results. In this paper, we follow the sampling-based approach, and propose a method which utilizes R*-tree's advantages that were never used before for big spatial data partitioning. The proposed R*-Grove index has three advantages over the existing work. First, it inherits and improves the R*-tree index structure to produce high-quality partitions that are tailored to big spatial data. Second, the improved algorithm produces balanced partitions by employing a user-defined parameter, termed *balance factor*, α , e.g., 95%. In addition, it can produce spatially disjoint partitions which are necessary for some spatial analysis algorithms. Third, R*-Grove can couple a sample with a data size histogram to guarantee the desired load balance even when the input record sizes are highly variant. While R*-Grove is not the only framework for big spatial partitioning, it is the first one that is tailored for large-scale spatial datasets while existing techniques reuse traditional index structures, such as R-tree, STR, or Quad-tree, as black boxes.

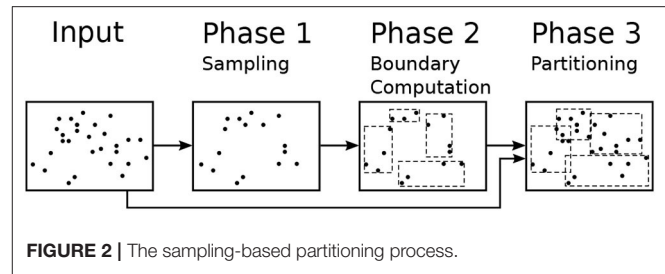
3. BACKGROUND

3.1. R*-Tree

The R*-tree (Beckmann et al., 1990) belongs to the R-tree family (Guttman, 1984) and it improves the insertion algorithm to provide high quality index. In R-tree, the number of children in each nodes has to be in the range $[m, M]$. By design, m can be at most $\lfloor M/2 \rfloor$ to ensure that splitting a node of size $M + 1$ is feasible. In this paper, we utilize and enhance two main functions of the R*-tree index, namely, CHOOSESUBTREE and SPLITNODE which are both used in the insertion process. For the CHOOSESUBTREE method, given the MBR of a record and a tree node, it chooses the best subtree to assign this record to. The SPLITNODE method takes an overflow node with $M + 1$ records and splits it into two nodes.

3.2. Sample-Based Partitioning Workflow

This section gives a background on the sampling-based partitioning technique (Vo et al., 2014; Eldawy and Mokbel, 2015; Eldawy et al., 2015a), just partitioning hereafter, that this paper relies on. **Figure 2** shows the workflow for the partitioning algorithm which consists of three phases, namely, sampling, boundary computation, and partitioning. The sampling phase (Phase 1) draws a random sample of the input records and converts each one to a point. Notice that sample points are picked from the entire file at no particular order so the order of



points does not affect the next steps. The boundary computation phase (Phase 2) runs on a single machine and processes the sample to produce partition boundaries as a set of rectangles. Given a sample S , the input size D , and the desired partition size B , this phase adjusts the capacity of each partition to contain $M = \lceil |S| \cdot B/D \rceil$ sample points which is expected to produce final partitions with the size of one block each. The final partitioning phase (Phase 3) scans the entire input in parallel and assigns each record to these partitions based on the MBR of the record and the partition boundaries. If each record is assigned to exactly one partition, the partitions will be spatially overlapping with no data replication. If each record is assigned to all overlapping partitions, the partitions will be spatially disjoint but some records can be replicated and duplicate handling will be needed in the query processing (Dittrich and Seeger, 2000). Some algorithms can only work if the partitions are spatially disjoint such as visualization (Eldawy et al., 2016) and some computational geometry functions (Li et al., 2019).

The proposed R*-Grove method expands Phase 1 by optionally building a histogram of storage size that assists in the partitioning algorithm at Phase 2. In Phase 2, it adapts R*-tree-based algorithms to produce the partition boundaries with desired level of load balance. In Phase 3, we propose a new data structure that improves the performance of that phase and allows us to produce spatially disjoint partitions if needed.

3.3. Quality Metrics

This paper uses the quality metrics defined in Eldawy et al. (2015a). Below, we redefine these metrics while accounting for the case of partitions that span multiple HDFS blocks. A single partition π_i is defined by two parameters, minimum bounding box mbb_i and size in bytes $size_i$. Given the HDFS block size B , e.g., 128 MB, we define the number of blocks for a partition π_i as $b_i = \lceil size_i/B \rceil$. Given a dataset that is partitioned into a set of l partitions, $\mathcal{P} = \{\pi_i\}$, we define the five quality metrics as follows.

DEFINITION 1 (Total Volume - Q_1). *The total volume is the sum of the volume of all partitions where the volume of a partition is the product of its side lengths.*

$$Q_1(\mathcal{P}) = \sum_{\pi_i \in \mathcal{P}} b_i \cdot \text{volume}(mbb_i)$$

We multiply by the number of blocks b_i because big spatial data systems process each block separately. Lowering the total volume is preferred to minimize the overlap with a query. Given the

popularity of the two-dimensional case, this is usually used under the term total area.

DEFINITION 2 (Total Volume Overlap - Q_2). This quality metric measures the sum of the overlap between pairs of partitions.

$$Q_2(\mathcal{P}) = \sum_{\pi_i, \pi_j \in \mathcal{P}, i \neq j} b_i \cdot b_j \cdot \text{volume}(mbb_i \cap mbb_j) + \sum_{\pi_i \in \mathcal{P}} \frac{b_i(b_i - 1)}{2} \cdot \text{volume}(mbb_i)$$

where $mbb_i \cap mbb_j$ is the intersection region between the two boxes. The first term calculates the overlaps between pairs of partitions and the second term accounts for self-overlap which treats a partition with multiple blocks as overlapping partitions. Lowering the volume overlap is preferred to keep the partitions apart.

DEFINITION 3 (Total Margin - Q_3). The margin of a block is the sum of its side lengths. The total margin is the sum of all margins as given below.

$$Q_3(\mathcal{P}) = \sum_{\pi_i \in \mathcal{P}} b_i \cdot \text{margin}(mbb_i)$$

Similar to Q_1 , multiplying by the number of blocks b_i treats each block as a separate partition. Lowering the total margin is preferred to produce square-like partitions.

DEFINITION 4 (Block Utilization - Q_4). Block utilization measures how full the HDFS blocks are.

$$Q_4(\mathcal{P}) = \frac{\sum_{\pi_i \in \mathcal{P}} \text{size}_i}{B \cdot \sum_{\pi_i \in \mathcal{P}} b_i}$$

The numerator $\sum \text{size}_i$ represents the total size of all partitions and denominator $B \sum b_i$ is the maximum amount of data that can be stored in all blocks used by these partitions. In big data applications, each block is processed in a separate task which has a setup time of a few seconds. Having full or near-full blocks minimize the overhead of the setup. The maximum value of block utilization is 1.0, or 100%.

DEFINITION 5 (Standard Deviation of Sizes).

$$Q_5(\mathcal{P}) = \sqrt{\frac{\sum_{\pi_i \in \mathcal{P}} (\text{size}_i - \overline{\text{size}})^2}{l}}$$

Where $\overline{\text{size}} = \sum \text{size}_i / l$ is the average partition size. Lowering this value is preferred to balance the load across partitions.

4. R*-GROVE PARTITIONING

This section describes the details of the proposed R*-Grove partitioning algorithm. R*-Grove employs three techniques that overcome the limitations of existing works. The first technique

adapts the R*-tree index structure for spatial partitioning by utilizing the CHOOSESUBTREE and SPLITNODE functions in the sample-based approach described in section 3. This technique ensures a high spatial quality of partitions. The second technique addresses the problem of load balancing by introducing a new constraint that guarantees a user-defined ratio between smallest and largest partitions. The third technique combines the sample points with its storage histogram to balance the sizes of the partitions rather than the number of records. This combination allows R*-Grove to precisely produce partitions with a desired block utilization, which cannot be achieved by any other partitioning techniques.

4.1. R*-Tree-Based Partitioning

This part describes how R*-Grove utilizes the R*-tree index structure to produce high quality partitions. It utilizes the SPLITNODE and CHOOSESUBTREE functions from the R*-tree algorithm in Phases 2 and 3 as described shortly. A naïve method (Vu and Eldawy, 2018) is to use the R*-tree as a blackbox in Phase 2 in **Figure 2** and insert all the sample points into an R*-tree. Then it emits the MBRs of the leaf nodes as the output partition boundaries. However, this technique was shown to be inefficient as it processes the sample points one-by-one and does not integrate the R*-tree index well in the partitioning algorithm. Therefore, we propose an efficient approach that runs much faster and produces higher quality partitions. It extends Phases 2 and 3 as follows.

Phase 2 computes partition boundaries by only using the SPLITNODE algorithm from the R*-tree index which splits a node with $M + 1$ records into two nodes with the size of each one in the range $[m, M]$. This algorithm starts by choosing the split axis, e.g., x or y , that minimizes the total margin. Then, all the points are sorted along the chosen axis and the split point is chosen as depicted in Algorithm 1. The CHOOSE SPLIT POINT algorithm simply considers all the split points and chooses the one that minimizes some cost function which is typically the total area of the two resulting partitions.

We set $M = \lceil |S| \cdot B / |D| \rceil$ as explained in section 3 and $m = 0.3M$ as recommended in the R*-tree paper. In particular, this phase starts by creating a single big tree node that has all the sample points S . Then, it recursively calls the SPLITNODE algorithm as long as the resulting node has more than M elements. This top-down approach has a key advantage over building the tree record-by-record as it allows the algorithm to look at all the records at the beginning and optimize for all of them. Furthermore, it avoids the FORCEDREINSERT algorithm which is known to slow down the R*-tree insertion process. Notice that this is different than the bulk loading algorithms as it does not produce a full tree. Rather, it just produces a set of boundaries that are produced as partitions. Phase 3 treats all the MBRs as leaf nodes in an R-tree and uses the CHOOSELEAF method from the R*-tree to assign an input record to a partition.

4.1.1. Run-Time Analysis

The SPLITNODE algorithm can be modeled as a recursive algorithm where each iteration sorts all the points and runs the linear-time splitting algorithm to produce two smaller partitions.

Algorithm 1 A simplified version of the traditional R*-tree splitting mechanism.

Inputs: P is the all sample records; m is the minimum size of a node.

Output: the optimal splitting position.

```

1: function CHOOSESPLITPOINT( $P, m$ )
2:   chosenK = -1; minCost =  $\infty$ 
3:   for  $k$  in  $[m, |P| - m]$  do
4:      $P_1 = P[1..k]$   $\triangleright P_1$  is the first  $k$  records of  $P$ 
5:      $P_2 = P[k + 1..|P|]$   $\triangleright P_2$  is all the remaining records
        $P - P_1$ 
6:     Calculate the cost of the partitions  $P_1$  and  $P_2$ 
7:     if the cost is smaller than minCost then
8:       Set chosenK =  $k$  and update minCost
9:   return chosenK

```

Algorithm 2 R*-tree-based split while ensuring valid partitions

Inputs: P is the all sample records; $[m, M]$ is the target range of sizes for final partitions.

Output: the optimal splitting position.

```

1: function CHOOSEVALIDSPLITPOINT( $P, m$ )
2:   for  $k$  in  $[m, |P| - m]$  do
3:     if either  $k$  or  $|P| - k$  is invalid then  $\triangleright$  Lemma 1
4:       Skip this iteration and continue
5:     Similar to Lines 4-8 in Algorithm 1
6:   return chosenK

```

The run-time can be expressed as $T(n) = T(k) + T(n - k) + O(n \log n)$, where k is the size of one group resulting from the partitioning, n is the number of records in the input partition. In particular, $T(k)$ and $T(n - k)$ are the running times to partition two partitions from splitting process. The term $O(n \log n)$ is the running time for the splitting part which requires sorting all the points. This recurrence relation has a worst case of $n^2 \log n$ if k is always $n - 1$. In order to guarantee a run-time of $O(n \log^2 n)$, we define a parameter $\rho \in [0, 0.5]$ which defines the minimum splitting ratio k/n . Setting this parameter to any non-zero fraction guarantees an $O(n \log^2 n)$ run-time. However, the restriction of k/n also limits the range of possible value of k . For example, if $n = 100$ and $\rho = 0.3$, k must be a number in the range $[30, 70]$. As this parameter gets closer to 0.5, the two partitions become closer in size and the run-time decreases but the quality of the index might also deteriorate due to the limited search space imposed by this parameter. To incorporate this parameter in the node-splitting algorithm, we call the CHOOSESPLITPOINT function with the parameters $(P, \max\{m, \rho \cdot |P|\})$, where $|P|$ is the number of points in the list P .

4.2. Load Balancing for Partitions With Equal-Size Records

In this section, we focus on balancing the number of records in partitions assuming equal-size records. We further extend this in the next section to support variable-size records. The method in section 4.1 does a good job in producing high-quality

partitions similar to what the R*-tree provides. However, it does not address the second limitation, that is, balancing the sizes of the partitions. Recall that the R-tree index family requires the leaf nodes to have sizes in the range $[m, M]$, where $m \leq M/2$. With the R*-tree algorithm explained earlier, some partitions might be 30% full which reduces block utilization and load balance. We would like to be able to set m to a larger value, say, $m = 0.95M$. Unfortunately, if we do so, the SPLITNODE algorithm would just fail because it will face a situation where there is no valid partitioning.

To illustrate the limitation of the SPLITNODE mechanism, consider the following simple example. Let us assume we choose $m = 9$ and $M = 10$ while the list P contains 28 points. If we call the SPLITNODE algorithm on the 28 points, it might produce two partitions with 14 records each. Since both contain more than $M = 10$ points, the splitting method will be called again on each of them which will produce an incorrect answer since there is no way to split 14 records into two groups while each of them contain between 9 and 10 records. A correct splitting algorithm would produce three partitions with sizes 9, 9, and 10. Therefore, we need to introduce a new constraint to the splitting algorithm so that it always produces partitions with sizes in the range $[m, M]$.

4.2.1. The Final Finding

The SPLITNODE algorithm can be minimally modified to guarantee final leaf partitions in the range $[m, M]$ by satisfying the following validity constraint:

$$\lceil S_i/M \rceil \leq \lfloor S_i/m \rfloor, i \in \{1, 2\},$$

where S_1 and S_2 are the sizes of the two resulting partitions of the split. Algorithm 2 depicts the main changes to the algorithm that introduces a new constraint test in Line 3 that skips over invalid partitioning. The rest of this section provides the theoretical proof that this simple constraint guarantees the algorithm termination with leaf partitions in the range $[m, M]$. We start with the following definition.

DEFINITION 6. Valid Partition Size: An integer number S is said to be a valid partition size with respect to a range $[m, M]$ if there exists a set of integers $X = \{x_1, \dots, x_n\}$ such that $\sum x_i = S$ and $x_i \in [m, M] \forall i \in [1, n]$. In words, if we have S records, there is at least one way to split them such that each split has between m and M records.

For example, if $m = 9$ and $M = 10$, the sizes 14, 31, and 62, are all invalid while the sizes 9, 27, and 63, are valid. Therefore, to produce balanced partitions, the SPLITNODE algorithm should keep the invariant that the partition sizes are always valid according to the above definition. Going back to the earlier example, if $S = 28$, the answer $S_1 = S_2 = 14$ will be rejected because $S_1 = 14$ is invalid. Rather, the result of the first call to the SplitNode algorithm will result in two partitions with sizes $\{10, 18\}$ or $\{9, 19\}$. The following lemma shows how to test a size for validity in constant time.

LEMMA 1. *Validity Test:* An integer S is a valid partition size w.r.t a range $[m, M]$ iff $L \leq U$ in which L (lower bound) and U (upper bound) are computed as:

$$L = \lceil S/M \rceil$$

$$U = \lfloor S/m \rfloor$$

PROOF. First, if S is valid then, by definition, there is a partitioning of S into n partitions such that each partition is in the range $[m, M]$. It is easy to show that $L \leq U$ and we omit this part for brevity. The second part is to show that if the inequality $L \leq U$ holds, then there is at least one valid partitioning. Based on the definition of L and U , we have:

$$U = \lfloor S/m \rfloor \Rightarrow U \leq S/m \Rightarrow S \geq m \cdot U \Rightarrow S \geq m \cdot L \Rightarrow S - m \cdot L \geq 0 \quad (1)$$

$$L = \lceil S/M \rceil \Rightarrow L \geq S/M \Rightarrow S \leq M \cdot L \Rightarrow S - m \cdot L \leq (M - m) \cdot L \quad (2)$$

Based on Inequalities 1 and 2, we can make a valid partitioning as follows:

1. Start with L empty partitions. Assign m records to each partition. The remaining number of records is $S - m \cdot L \geq 0$. This is satisfied due to Inequality 1.
2. Since each partition now has m records, it can receive up-to $M - m$ additional records in order to keep its validity. Overall, L partitions of size m can accommodate up-to $(M - m) \cdot L$ records to keep a valid partitioning. But the remaining number of records $S - m \cdot L$ is not larger than the upper limit of what the partitions can accommodate, $(M - m) \cdot L$ as shown in Inequality 2. Therefore, this condition is satisfied as well.

In conclusion, it follows that if the condition $L \leq U$ holds, we can always find a valid partitioning scheme for S records which completes the proof.

If we apply this test for the example above, we find that 28 is valid because $L = \lceil 28/10 \rceil = 3 \leq U = \lfloor 28/9 \rfloor = 3$ while 62 is invalid because $L = \lceil 62/10 \rceil = 7 > U = \lfloor 62/9 \rfloor = 6$. This approach works fine as long as the initial sample size S is valid but how do we guarantee the validity of S ? We show that this is easily guaranteed if the size S is above some threshold S^* as shown in the following lemma.

LEMMA 2. Given a range $[m, M]$, any partition of size $S \geq S^*$ is valid where S^* is defined by the following formula:

$$S^* = \left\lceil \frac{m}{M-m} \right\rceil \cdot m \quad (3)$$

PROOF. Following Definition 6, we will prove that for any partition size $S \geq S^*$, there exists a way to split it into k groups such that the size of each group is in the range $[m, M]$.

First, let $i = \left\lceil \frac{m}{M-m} \right\rceil$, we have:

$$S \geq S^* = \left\lceil \frac{m}{M-m} \right\rceil \cdot m = i \cdot m \quad (4)$$

$$\Rightarrow S = i \cdot m + X, X \geq 0. \text{ Let } X = a \cdot m + b, a \geq 0, 0 \leq b < m \quad (5)$$

$$\Rightarrow S = i \cdot m + (a \cdot m + b) = (i + a) \cdot m + b, a \geq 0, 0 \leq b < m \quad (6)$$

Second, since $b < m$, we have:

$$\frac{b}{M-m} < \frac{m}{M-m} \leq i \Rightarrow \frac{b}{i} < M-m \quad (7)$$

From Equations (6) and (7), we can make a valid partitioning for a partition size S as follows:

1. Start with $i + a$ empty partitions. Assign m records to each partition. The remaining number of records is b . This step is based on Equation (6).
2. Equation (7) means that we can split b records over i groups such that each group receives at most $M - m$ records. Since we already have $i + a$ groups each of size m , adding $M - m$ to i groups out of them will increase their sizes to M which still keeps them in the valid range $[m, M]$. The remaining groups will still have m records making them valid too.

This completes the proof of Lemma 2.

Based on Lemma 2, a question is raised as how large the size of sample points S should be to ensure that a good block utilization is achievable. As we mentioned from beginning, R*-Grove allows us to configure a parameter $\alpha = m/M$, that called *balance factor*, is computed as the ratio between minimum and maximum number of records of a leaf node in the tree. α should be close to 1 to guarantee a good block utilization. Let's assume that $0 < r \leq 1$ is the sampling ratio and p is the storage size of a single point. The maximum number of records M is computed in the section 4.1 as:

$$M = \left\lceil \frac{|S| \cdot B}{D} \right\rceil \Rightarrow M = \left\lceil \frac{|S| \cdot p}{D} \cdot \frac{B}{p} \right\rceil \Rightarrow M = \left\lceil \frac{r \cdot B}{p} \right\rceil \quad (8)$$

From Equation (8), we can rewrite Lemma 2 as:

$$\begin{aligned} |S| \geq S^* &= \left\lceil \frac{m}{M-m} \right\rceil \cdot m \Rightarrow |S| \geq \left\lceil \frac{\alpha}{1-\alpha} \right\rceil \cdot \alpha \cdot \left\lceil \frac{r \cdot B}{p} \right\rceil \\ &\Rightarrow |S| \cdot p \geq \left\lceil \frac{\alpha}{1-\alpha} \right\rceil \cdot \alpha \cdot \lceil r \cdot B \rceil \end{aligned} \quad (9)$$

Therefore, assume that we want to configure the *balance factor* as $\alpha = 0.95$, sample ratio $r = 1\%$, and block size $B = 128$ MB, then the term $|S| \cdot p$ in Equation (9) would be computed as 23 MB. In other words, if the storage size of sample points $|S| \cdot p \geq 23$ MB, it will be guaranteed to produce a valid partitioning. This is a reasonable size that can be stored in main memory and processed in a single machine.

Algorithm 3 Choose the split point with weights

Inputs: P is the all sample records; w is an array of weights of corresponding records in P ; $[m, M]$ is the target range of sizes for final partitions.

Output: the optimal splitting position.

```

1: function CHOOSEWEIGHTEDSPLITPOINT( $P, w, m$ )
2:    $W = \sum_{1 \leq i \leq |P|} w_i$ 
3:   for  $k$  in  $[m, |P| - m]$  do
4:      $W_1 = \sum_{1 \leq i \leq k} w_i$ 
5:     if either  $W_1$  or  $W - W_1$  is invalid then  $\triangleright$  Lemma 1
6:       Skip this iteration and continue
7:     Similar to Lines 4-8 in Algorithm 1
8:   return chosenK

```

4.3. Load Balancing for Datasets With Variable-Size Records

The above two approaches can be combined to produce high-quality and balanced partitions in terms of number of records. However, the partitioning technique needs to write the actual records in each partition and often these records are of variable sizes. For example, the sizes of records in the OSM-Objects dataset (allobjects, 2019) range from 12 bytes to 10 MB per record. Therefore, balancing the number of records can result in a huge variance in the partition sizes in terms of number of bytes.

To overcome this limitation, we combine the sample points with a *storage size histogram* of the input as follows. The storage size histogram is used to assign a weight to each sample point that represents the total size of all records in its vicinity. To find these weights, Phase 1 computes, in addition to the sample, a storage size histogram of the input. This histogram is created by overlaying a uniform grid on the input space and computing the total size of all records that lie in each grid cell (Chasparis and Eldawy, 2017; Siddique et al., 2019). This histogram is computed on the full dataset not the sample, therefore, it catches the actual size of the input. After that, we count the number of *sample* points in each grid cell. Finally, we divide the total weight of each cell among all sample points in this cell. For example, if a cell has a weight of 1,000 bytes and contains five sample points, the weight of each point in this cell becomes 200 bytes.

In Phase 2, the SPLITNODE function is further improved to balance the *total weight* of the points in each partition rather than the number of points. This also requires modifying the value of M to be $M = \lceil \sum w_i / N \rceil$, where w_i is the assigned weight to the sample point p_i , and N is the desired number of partitions. Algorithm 3 shows how the algorithm is modified to take the weights into account. Line 4 calculates the weight of each partitioning point which is used to test the validity of this split point as shown in Line 5.

Unfortunately, if we apply this change, the algorithm is no longer guaranteed to produce balanced partitions. The reason is that the proof of Lemma 1 is no longer valid. That proof assumed that the partition sizes are defined in terms of number of records which makes all possible partition sizes part of the search space in the **for-loop** in Line 2 of Algorithm 2. However, when the size

of each partition is the sum of the weights, the possible sizes are limited to the weights of the points. For example, let us assume a partition with five points all of the same weight $w_i = 200$ while $m = 450$ and $M = 550$. The condition in Definition 6 suggests that the total weight 1,000 is valid because $L = \lceil 1,000/550 \rceil = 2 \leq U = \lfloor 1000/450 \rfloor = 2$. However, given the weights $w_i = 200$ for $i \in [1, 5]$, there is no valid partitioning, i.e., there is no way to make two partitions each with a total weight in the range $[450, 550]$.

To overcome this issue, this part further improves the SPLITNODE algorithm so that it still guarantees a valid partitioning even for the case described above. The key idea is to make minimal changes to the weights to ensure that the algorithm will terminate with a valid partitioning; we call this process *weight correction*. For example, the case described earlier will be resolved by changing the weights of two points from 200 and 200 to 100 and 300. This will result in the valid partitioning $\{200, 200, 100\}$ and $\{300, 200\}$ which is valid. Keep in mind that these weights are approximate anyway as they are based on the sample and histogram so these minimal changes would not hugely affect the overall quality, yet, they ensure that the algorithm will terminate correctly. The following part describes how these weight changes are applied while ensuring a valid answer.

First of all, we assume that the points are already sorted along the chosen axis as explained in section 4.1. Further, we assume that Algorithm 3 failed by not finding any valid partitions, i.e., return -1 . Now, we make the following definitions to use them in the weight update function.

DEFINITION 7. Point position: Let p_i be point $\#i$ in the sort order and its weight is w_i . We define the position of the point i as $pos_i = \sum_{j \leq i} w_j$.

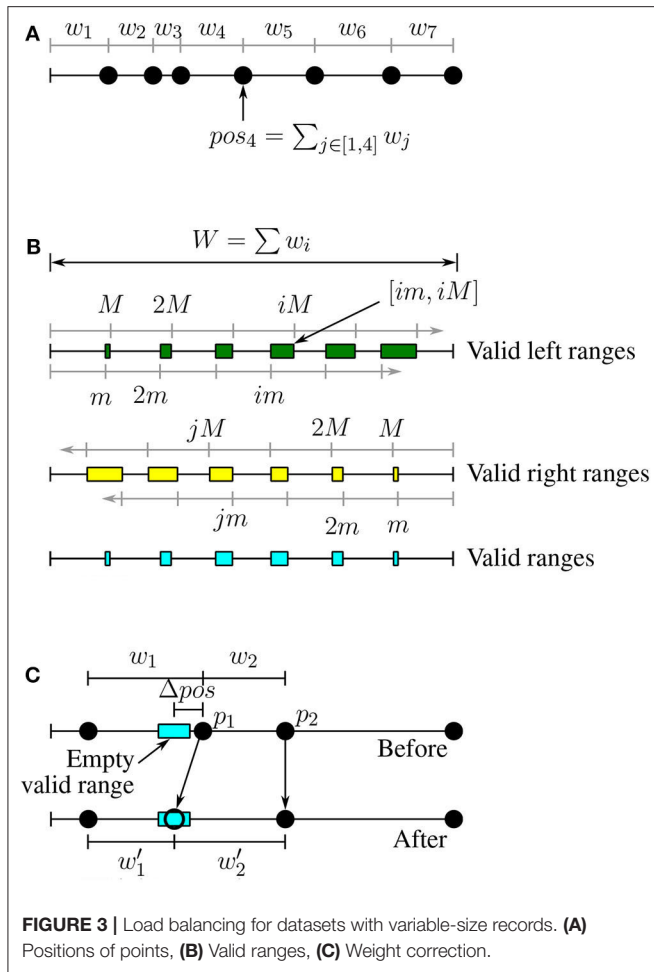
Based on this definition, we can place all the points on a linear scale based on their position as shown in **Figure 3A**.

DEFINITION 8. Valid left range: A range of positions $VL = [vl_s, vl_e]$ is a valid left range if for all positions $vl \in VL$ the value vl is valid w.r.t. $[m, M]$. All the valid left ranges can be written in the form $[im, iM]$ where i is a natural number and they might overlap for large values of i (see **Figure 3B**).

DEFINITION 9. Valid right range: A range of positions $VR = [vr_s, vr_e]$ is a valid right range if for all positions $vr \in VR$ the value $W - vr$ is valid w.r.t. $[m, M]$. Similar to valid left ranges, all valid right ranges can be written in the form $[W - jM, W - jm]$, where $W = \sum w_i$ (see **Figure 3B**).

DEFINITION 10. Valid range: A range of positions $V = [v_s, v_e]$ is valid if for all positions $v \in V$, v belongs to at least one valid left range and at least one valid right range. In other words, the valid ranges are the intersection of the valid left ranges and valid right ranges.

Figure 3B illustrates the valid left, valid right, and valid ranges. If we split a partition around a point with a position in a valid left range, the first partition will be valid. Similarly for valid



right positions the second partition (on the right) will be valid. Therefore, we would like to split a partition around a point in one of the valid ranges (intersection of left and right).

LEMMA 3. Empty valid ranges: If Algorithm 3 fails by returning -1 , then none of the point positions in P falls in a valid range.

PROOF. By contradiction, let a point p_i has a position pos_i that falls in a valid range. In this case, the partitions $P_1 = \{p_k : k \leq i\}$ and $P_2 = \{p_l : l > i\}$ are both valid partitions because the total weight of P_1 is equal to the position pos_i which is valid because pos_i falls in a valid left range. Similarly, the total weight of P_2 is valid because pos_i falls in a valid right range. In this case, Algorithm 3 should have found this partitioning as a valid partitioning because it tests all the points which is a contradiction.

A corollary to Lemma 3 is that when Algorithm 3 fails by returning -1 , then all valid ranges are empty.

As a result, we would like to slightly modify the weights of some points in the sample points in order to enforce some points to fall in valid ranges. We call this the *weight correction* process. This process is described in the following lemma:

LEMMA 4. Weight correction: Given any empty valid range $[v_s, v_e]$, we can modify the weight of only two points such that the position of one of them will fall in the range.

PROOF. Figure 3C illustrates the proof of this lemma. Given an empty valid range, we modify the two points with positions that follow the empty valid range, p_1 and p_2 , where $pos_1 < pos_2$. We would like to move the point p_1 to the new position $pos'_1 = (v_s + v_e)/2$ which is in the middle of the empty valid range. To do that, we reduce the weight w_1 by $\Delta pos = pos_1 - pos'_1$. The updated weight $w'_1 = w_1 - \Delta pos$. To keep the position of p_2 and all the following points intact, we have to also increase the weight of p_2 by Δpos ; that is, $w'_2 = w_2 + \Delta pos$.

We do the weight correction process for *all* empty valid ranges to make them non-empty and then we repeat Algorithm 3 to choose the best one among them.

The only remaining part is how to enumerate all the valid ranges. The idea is to simply find a valid left range, an overlapping valid right range, and compute their intersection, all in constant time. Given a natural number i , the valid left range is in the form $[im, iM]$. Assume that this range overlaps a valid right range in the form $[W - jM, W - jm]$. Since they overlap, the following two inequalities should hold:

$$W - jm \geq im \Rightarrow j < \frac{W - im}{m}$$

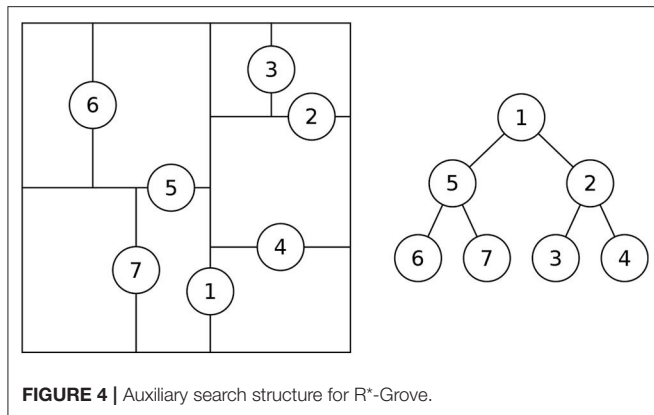
$$W - jM \leq iM \Rightarrow j > \frac{W - iM}{M}$$

Therefore, the lower bound of j is $j_1 = \lceil \frac{W - iM}{M} \rceil$ and the upper bound of j is $j_2 = \lfloor \frac{W - im}{m} \rfloor$. If $j_1 \leq j_2$, then there is a solution to these inequalities which we use to generate the bounds of the valid range $[v_s, v_e]$. Notice that if there is more than one valid solution to j , all of them should be considered to generate all the valid ranges but we omit this special case for brevity.

4.4. Implementation Considerations

4.4.1. Optimization of Phase 3

The CHOOSESUBTREE operation in R*-tree chooses the node that results in the least overlap increase with its siblings (Beckmann et al., 1990). A straight-forward implementation of this method is $O(n^2)$ as it needs to compute the overlap between each candidate partition and all other partitions. In the R*-tree index, this cost is limited due to the limited size of each node. However, this step can be too slow as the number of partitions in R*-Grove can be extremely large. To speed up this step, we use a K-d-tree-like auxiliary search structure as shown in Figure 4. This index structure is generated during Phase 2 as the partition boundaries are computed. Each time the NODESPLIT operation completes, the search structure is updated by adding a corresponding split in the direction of the chosen split axis. This auxiliary search structure is stored in memory and replicated to all nodes. It will be used in Phase 3, when we physically store the input records to the partitions. Given a spatial record, it will be assigned to the corresponding partition using a search algorithm which is



similar to the K-d-tree's point search algorithm (Bentley, 1975). Based on this similarity, we can estimate the running time to choose a partition to be $O(\log(n))$. Notice that this optimization is not applicable in traditional R*-trees as the partitions might be overlapping while in R*-Grove we utilize the fact that we only partition points which guarantees disjoint partitions.

Since the partition MBRs in Phase 2 are computed from sample objects, there will be objects which do not fall in any partition in Phase 3. R*-Grove addresses this problem in two ways. First, if no disjoint partitions are desired, it chooses a single partition based on the CHOOSELEAF method in original R*-tree. In short, an object will be assigned to the partition in which the enlarged area or margin is minimal. Second, if disjoint partitions are desired, R*-Grove uses the auxiliary data structure, which covers the entire space, to assign this record to all overlapping partitions.

4.4.2. Disjoint Indexes

Another advantage of using the auxiliary search structure described above, is that it allows for building a disjoint index. This search structure naturally provides disjoint partitions. To ensure that the partitions cover the entire input space, we assume that input region is infinite, that is, starts from $-\infty$ and ends at $+\infty$ in all dimensions. Then, Phase 3 replicates each record to all overlapping partitions by directly searching in this k-d-tree-like structure with range search algorithm, which has the $O(\sqrt{n})$ running time (Lee and Wong, 1977). This advantage was not possible with the black-box R*-tree implementation as it is not guaranteed to provide disjoint partitions.

5. CASE STUDIES

This section describes three case studies where the R*-Grove partitioning technique can improve big spatial data processing. We consider three fundamental operations, namely, indexing, range query, and spatial join.

5.1. Indexing

Spatial data indexing is an essential component in most big spatial data management systems. The state-of-the-art global indexing techniques rely on reusing existing index structures

with a sample which are shown to be inefficient in terms of quality and load balancing (Vo et al., 2014; Eldawy et al., 2015a; Yu et al., 2015).

R*-Grove partitioning can be used for the *global indexing* step which partitions records across machines. In big spatial data indexing, the global index is the most crucial step as it ensures load balancing and efficient pruning when the index is used. If only the number of records needs to be balanced or if the records are roughly equi-sized, then the techniques described in sections 4.1 and 4.2 can be used. If the records are of a variable size and the total sizes of partitions need to be balanced, then the histogram-based step in section 4.3 can be added to ensure a higher load balance. Notice that the index would hugely benefit from the balanced partition size as it reduces the total number of blocks in the output file which improves the performance of all Spark and MapReduce queries that create one task per file block.

5.2. Range Query

Range query is a popular spatial query, which is also the building block of many other complex spatial queries. Previous studies found a strong correlation between the performance of range queries and the performance of other queries such as spatial join (Hoel and Samet, 1994; Eldawy et al., 2015a). Therefore, the performance of range query could be considered as a good reflection about the quality of a partitioning technique. A good partitioning technique allows the query processor to make two optimization techniques. First, it can prune the partitions that are completely outside the query range. Second, it can directly write to the output the partitions that are completely contained in the query range without further processing (Eldawy et al., 2017). For very small ranges, most partitioning techniques will behave similarly as it is most likely that the small query overlaps one partition and no partitions are completely contained (Eldawy et al., 2015a). However, as the query range increases, the differences between the partitioning techniques become apparent. Since most range queries are expected to be square-like, the R*-Grove partitioning is expected to perform very well as it minimizes the total margin which produces square-like partitions. Furthermore, the balanced load across partitions minimizes the straggler effect where one partition takes significantly longer time than all other partitions.

5.3. Spatial Join

Spatial join is another important spatial query that benefits from the improved R*-Grove partitioning technique. In spatial join, two big datasets need to be combined to find all the overlapping pairs of records. To support spatial join on partitioned big spatial data, each dataset is partitioned independently. Then, a spatial join runs between the partition boundaries to find all pairs of overlapping partitions. Finally, these pairs of partitions are processed in parallel. An existing approach (Zhou et al., 1998) preserves spatial locality to reduce the processing jobs. However, it still relies on traditional index like R-Tree, which also inherited its limitations. The R*-Grove partitioning has two advantages for the spatial join operation. First, it is expected to reduce the number of partitions by increasing the load balance which

TABLE 1 | Datasets for experiments.

Dataset	Type	Dimensions	Size (GB)	# Records
(1) OSM-Nodes	Point	2	500	7.4 billions
(2) OSM-Roads	Line segments	2	20	59 millions
(3) OSM-Parks	Polygon	2	7.2	10 millions
(4) OSM-Objects	Polygon	2	96	264 millions
(5) NYC-Taxi	Point	4,5,7	46	173 millions
(6) Diagonal points	Point	3,4,5,9	100	80 millions

reduces the total number of pairs. Second, it produces square-like partitions which is expected to overlap with fewer partitions of the other dataset as compared to the very thin and wide partitions that the STR or other partitioning techniques produce. These advantages allows R*-Grove to significantly outperform other partitioning techniques in spatial join query performance. We will validate these advantages in the section 6.5.2.

6. EXPERIMENTS

In this section, we carry out an extensive experimental study to validate the advantages of R*-Grove over widely used partitioning techniques, such as bulk loading STR, Kd-tree, Z-Curve, and Hilbert curve. We will show how R*-Grove addresses the current limitations of those techniques, leads to a better performance in big spatial data processing. In addition, we also show other capabilities of R*-Grove in the context of big spatial data, for example, how it works with large or multi-dimensional datasets. The experimental results in this section provide an evidence to the spatial community to start using R*-Grove if they would like to improve the system performance of their spatial applications.

6.1. Experimental Setup

6.1.1. Datasets

Table 1 summarizes the datasets will be used in our experiments. We use both real world and synthetic datasets for our experiments: (1) Semi-synthetic OpenStreetMap (OSM-Nodes) dataset with 7.4 billion points and a total size of 500 GB. This is a semi-synthetic dataset which represents all the points in the world. The points in this dataset are generated within a pre-specified distance from original points from OSM-Nodes dataset; (2) OSM-Roads with size 20 GB; and (3) OSM Parks with size 7.2 GB, which contain line segments and polygons for spatial join experiments. (4) OSM-Objects dataset with size 92 GB, which contains many variable-size records. (5) NYC-Taxi dataset with size 41.7 GB with up-to seven dimensions. All of those datasets are available online on UCR-STAR (ucrstar, 2019)—our public repository for spatial data; (6) Synthetic multi-dimensional diagonal points, with the number of dimensions are 3, 4, 5, and 9. This synthetic dataset is generated using our open source Spatial Data Generator (Vu et al., 2019). Dataset (5) and (6) allow us to show the advantages of R*-Grove in multi-dimensional datasets.

6.1.2. Parameters and Performance Metrics

In the following experiments, we partition the mentioned datasets with different datasets size $|D|$ in different techniques then we measure: (1) partition quality metrics, namely, total partition area, total partition margin, total partition overlap, block utilization (maximum is 1.0, i.e. 100%), standard deviation of partition size in MB (load balance). Notice that unit is not relevant for area, margin and overlap metric; (2) total partitioning time (in seconds), (3) for range queries, we measure the number of processed blocks and query running time, (4) for spatial join, we measure the number of processed blocks and total running time. We fix the balance factor $\alpha = 0.95$ and HDFS block size at 128 MB.

6.1.3. Machine Specs

All the experiments are executed on a cluster of one head node and 12 worker nodes, each having 12 cores, 64 GB of RAM, and a 10 TB HDD. They run CentOS 7 and Oracle Java 1.8.0_131. The cluster is equipped with Apache Spark 2.3.0 and Apache Hadoop 2.9.0. The proposed indexes are available for running in both Spark and Hadoop. Unless otherwise mentioned, we use Spark by default. The source code is available at <https://bitbucket.org/tvu032/beast-tv/src/rsgrove/>. The implementation for R*-Grove (RSGrovePartitioner) is located at *indexing* package.

6.1.4. Baseline Techniques

We compare R*-Grove to K-d Tree, STR, Z-curve, and Hilbert curve (denoted H-Curve thereafter) which are widely used in existing big spatial data systems (Eldawy and Mokbel, 2016). Z-Curve is adopted in some systems under the name Geohash which behaves in the same way.

6.2. Effectiveness of the Proposed Improvements in R*-Grove

In this experiment, we compare the three following variants of R*-Grove: (1) *R*-tree-black-box* is the application of the method in section 4.1. Simply, it uses the basic R*-tree algorithm to compute high-quality partition but it does not guarantee a high block utilization or load balance. (2) *R*-tree-gray-box* applies the improvements in sections 4.1 and 4.2. In addition to the high-quality partition, this method can also guarantee a high block utilization in terms of number of records per partition but it does not perform well if records have highly-variable sizes since it does not include the size adjustment technique in section 4.3. (3) *R*-Grove* applies all the three improvements at sections 4.1, 4.2, and 4.3. It has the advantage of producing high-quality partitions and can also guarantee a high block utilization in terms of storage size even when the record sizes are highly variable.

In **Figure 5**, we partition the OSM-Objects dataset, which contains variable-size records to validate our proposed improvements. Overall, R*-Grove outperforms R*-tree-black-box and R*-tree-gray-box in all of spatial quality metrics. Especially, R*-Grove provides excellent load balance between partitions as shown in **Figure 5D**, which is the standard deviation of partition size in OSM-Objects dataset. Given the HDFS block size is 128 MB, R*-Grove has the standard deviation of partition size 5 – 6 times smaller than R*-Tree-gray-box and

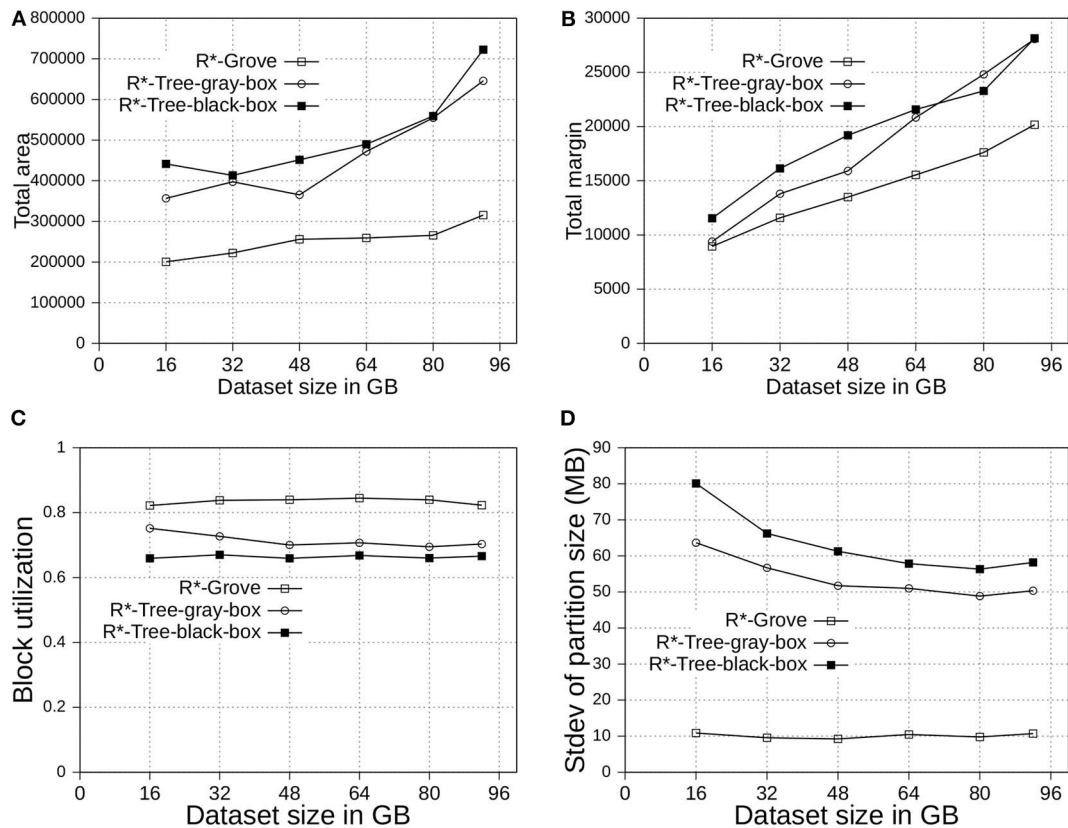


FIGURE 5 | Partition quality with variable-record-size dataset in R*-Grove and its two variants, R*-tree-black-box and R*-tree-gray-box. **(A)** Total area, **(B)** Total margin, **(C)** Block utilization, **(D)** Load balance.

R*-Tree-black-box. Since then, the following experiments will evaluate the performance of R*-Grove with existing widely-used spatial partitioning techniques.

6.3. Results Overview

Figure 6 shows an overview of the advantages of R*-Grove over other partitioning techniques for indexing, range query, and spatial join. In this experiment, we compare to four popular baseline techniques, namely, STR, Kd-Tree, Z-Curve, and H-Curve. We use OSM-Nodes dataset (ucrstar, 2019) for this experiment. The numbers on the y - axis are normalized to the largest number for a better representation except for block utilization which is reported as-is. Except for block utilization, the lower the value in the chart the better it is. The first two groups, total area and total margin, show that index quality of R*-Grove is clearly better than other baselines in both measures. For block utilization, on average, a partition in R*-Grove occupy around 90%, while other techniques could only utilize 60 – 70% storage capacity of an HDFS block. R*-Grove also has a better load balance when compared to other techniques. The last two groups indicate that R*-Grove significantly outperforms other partitioning techniques in terms of range query and spatial join query performance. We will go into further details in the rest of this section.

6.4. Partition Quality

This section shows the advantages of R*-Grove for indexing big spatial data when compared to other partitioning techniques. We use OSM-Nodes and OSM-Objects dataset with size up to 200 and 92 GB, respectively. We compare five techniques, namely, R*-Grove, STR, Kd-Tree, Z-Curve and H-Curve. We implemented those techniques on Spark with sampling-based partitioning mechanism. **Figures 7A, 8A** show that there is no significant difference of indexing performance between different techniques. This result is expected since the main difference between them is in Phase 2 which runs on a single machine on a sample of a small size (and a histogram in case of R*-Grove). Typically, Phase 2 takes only a few seconds to finish. These results suggest that the proposed R*-Grove algorithm requires the same computational resources as the baseline techniques. Meanwhile, it provides a better query performance by providing a higher partition quality as detailed next.

6.4.1. Total Area and Total Margin

Figures 7B, 8B show the total area of indexed datasets when we vary the OSM-Nodes and OSM-Objects dataset size from 20 to 200 GB and 16 to 92 GB, respectively. R*-Grove is the winner since it minimizes the total area of all partitions. While H-Curve performs generally better than Z-Curve, they are both doing bad since they do not take partition area into account in their

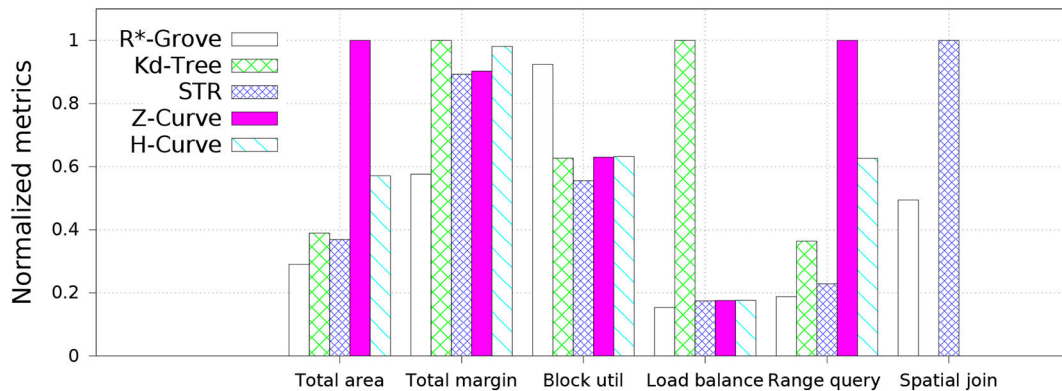


FIGURE 6 | The advantages of R*-Grove when compared to existing partitioning techniques.

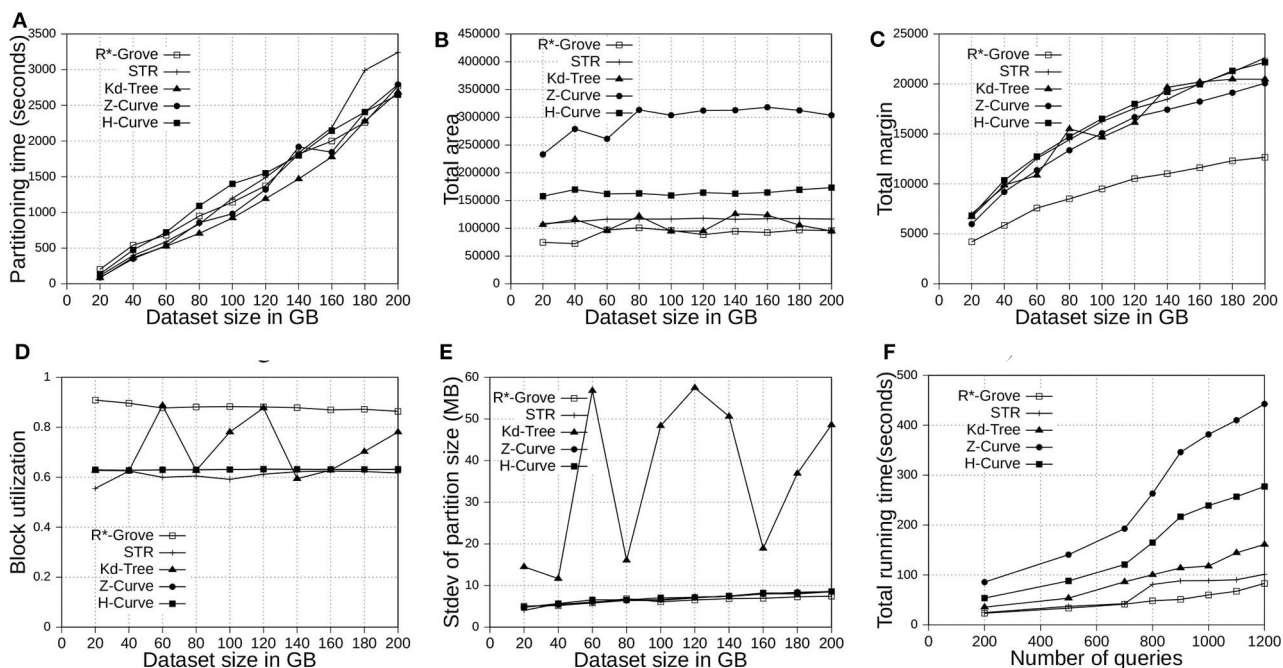


FIGURE 7 | Indexing performance and partition quality of R*-Grove and other partitioning techniques in OSM-Nodes datasets with similar-size records. (A) Partitioning time, (B) Total area, (C) Total margin, (D) Block utilization, (E) Load balance, (F) Range query performance.

optimization criteria. Specially, **Figure 8B** strongly validates the advantages of R*-Grove in non-point datasets. **Figures 7C, 8C** report the total margin for the same experiment. R*-Grove is the clear winner because it inherits the splitting mechanism of R*-Tree, which is the only one among all those that tries to produce square-like partitions. As the input size increases, more partitions are generated which causes the total margin to increase.

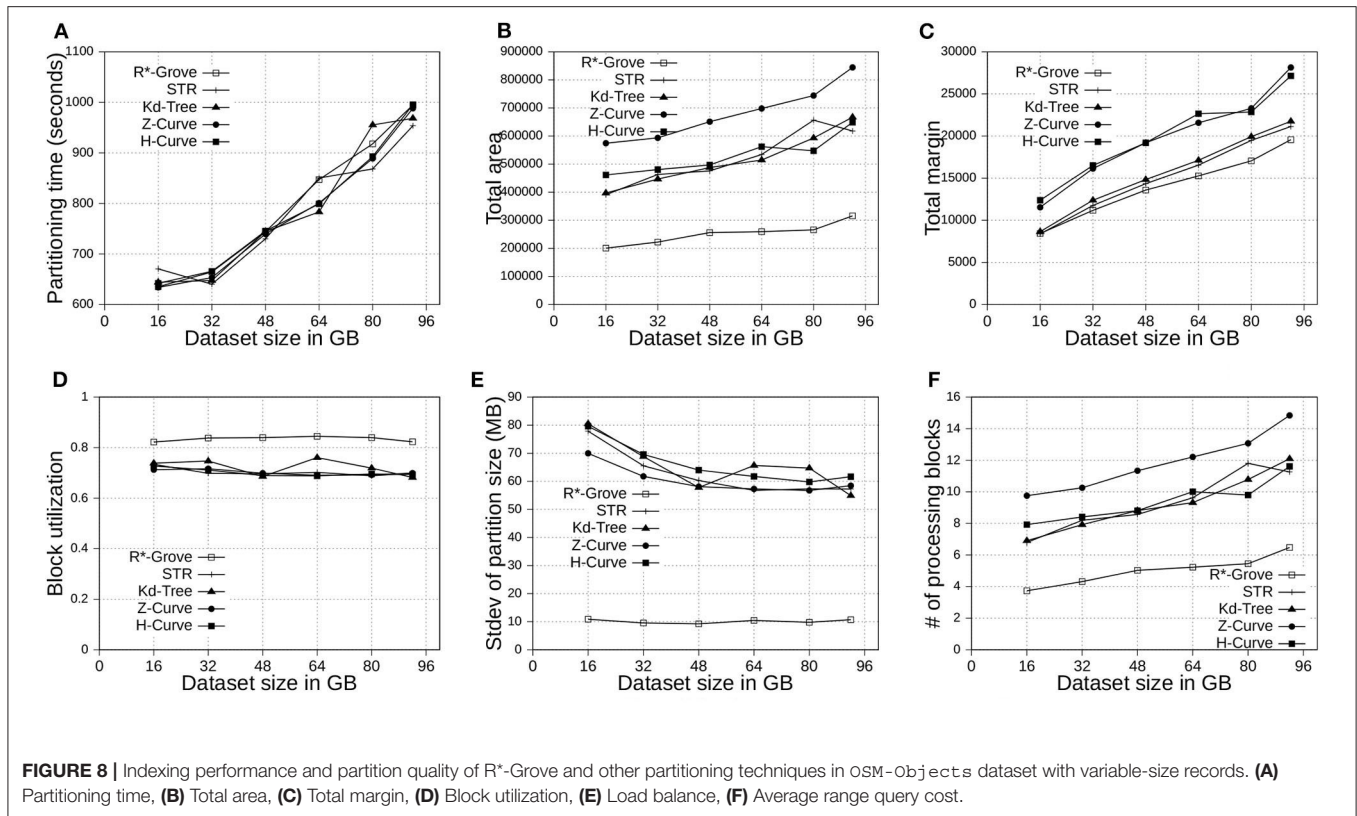
6.4.2. Block Utilization

Figures 7D, 8D show the block utilization as the input size increases. R*-Grove outperforms other partitioning techniques due to the proposed improvements in sections 4.2 and 4.3 specifically improve block utilization. Using R*-Grove, each

partition almost occupies a full block in HDFS which increases the overall block utilization. Z-Curve and H-Curve perform similarly since they produce equi-sized partition by creating split points along the curve. The high variability of the Kd-tree is due to the way it partitions the space at each iteration. Since it always partitions the space along the median, it only works perfectly if the number of partitions is a power of two; otherwise, it could be very inefficient. This occasionally results in partitions of high block utilization but they could be highly variable in size.

6.4.3. Load Balance

Figures 7E, 8E show the standard deviation of partition size in MB for the OSM-Nodes and OSM-Objects datasets,



respectively. Note that the HDFS block size is set to 128 MB. A smaller standard deviation indicates a better load balance. In **Figure 7E**, the dataset OSM-Nodes contains records of almost the same size so R*-Grove performs only slightly better than Z-Curve, H-Curve, and STR even though these three techniques try to primarily balance the partition sizes. In **Figure 8E**, the OSM-Objects dataset contains highly variable record sizes. In this case, R*-Grove is way better than all other techniques as it is the only one that employs the storage histogram to balance variable size records. In particular, we observe that the standard deviation of partition size on STR, Kd-Tree, Z-Curve, and H-Curve is about 50–60% of the HDFS block size. Meanwhile, R*-Grove maintains a value around 10 MB, which is only 8% of the block size.

6.4.4. Effect of Sampling Ratio

Since the proposed R*-Grove follows the *sampling-based* partitioning mechanism, a valid question is how the sampling ratio affects partition quality and performance? In this experiment, we execute several partitioning operations using R*-Grove in OSM-Objects datasets. All the partitioning parameters are kept fixed, except the sampling ratio, which is varying from 0.001 to 3%. For each sampling ratio value, we execute the partition operation three times, then compute the average and standard deviation of quality measures and partition construction time. Partition construction is the process that compute partition MBRs from the sample. **Figure 9** uses the

average values to plot the line and the standard deviation for the error bars. First, **Figure 9A** shows that the higher sampling ratio requires higher time for the partition construction process. This is expected due to the number of sample records which the partitioner has to use to compute partition MBRs. **Figures 9B–E** show the downward trend of total area, total margin, total overlap, and standard deviation of partition size. **Figure 9F** shows the upward trend of block utilization when the sampling ratio increases. In addition, the standard deviation of small sampling ratios is much higher than that for high sampling ratios. These results indicate that the higher sampling ratios promise better partition quality. However, an important observation is that the partition quality measures start stabilizing for sampling ratios larger than 1%. This behavior was also validated in a previous work (Eldawy et al., 2015a). In short, this work shows that a sample ratio of 1% dataset is enough to achieve virtually a same partition quality as sample ratio 100%. In the following experiments, we choose 1% as the default sampling ratio for all partitioning techniques.

6.4.5. Effect of Minimum Split Ratio

In section 4.1, we introduced parameter ρ , namely minimum splitting ratio, to speed up the running time of SPLITNODE algorithm used in Phase 2, boundary computation. In this experiment, we verify how the minimum splitting ratio impacts the partition quality and performance. We also use OSM-Objects dataset with R*-Grove partitioning as the

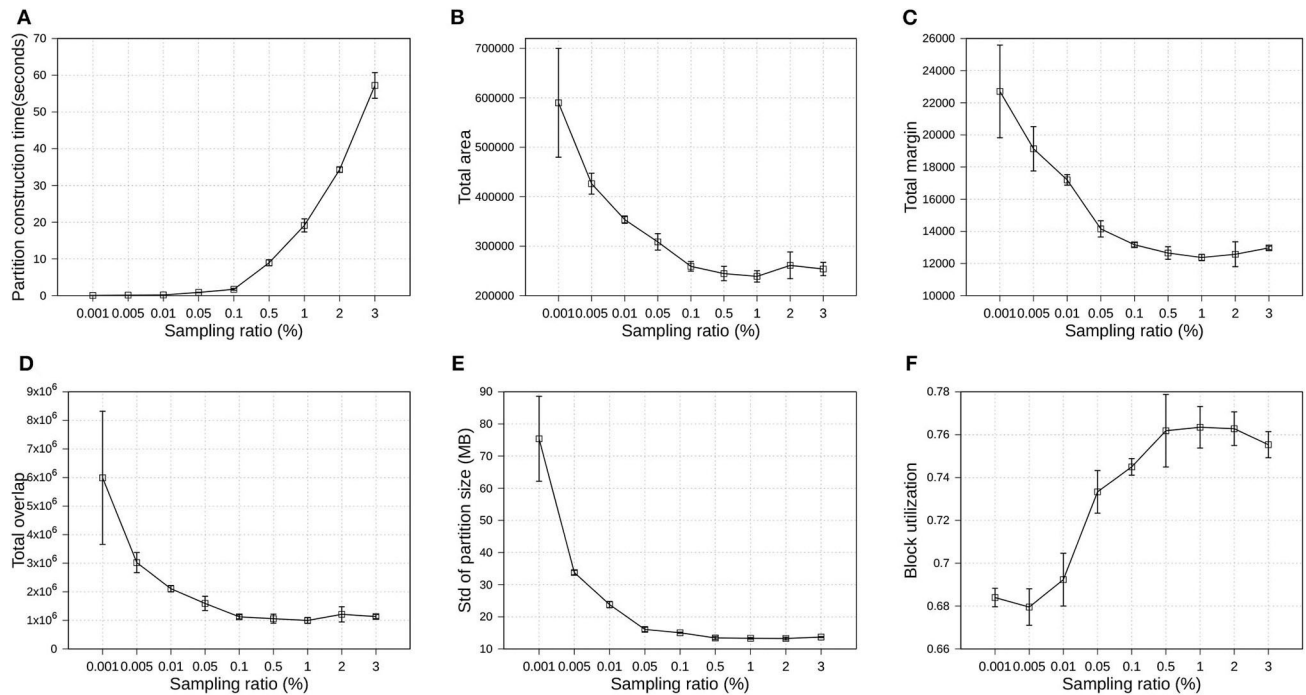


FIGURE 9 | Indexing performance and partition quality of R*-Grove in OSM-Objects datasets with different sampling ratios. **(A)** Partition construction time, **(B)** Total area, **(C)** Total margin, **(D)** Total overlap, **(E)** Load balance, **(F)** Block utilization.

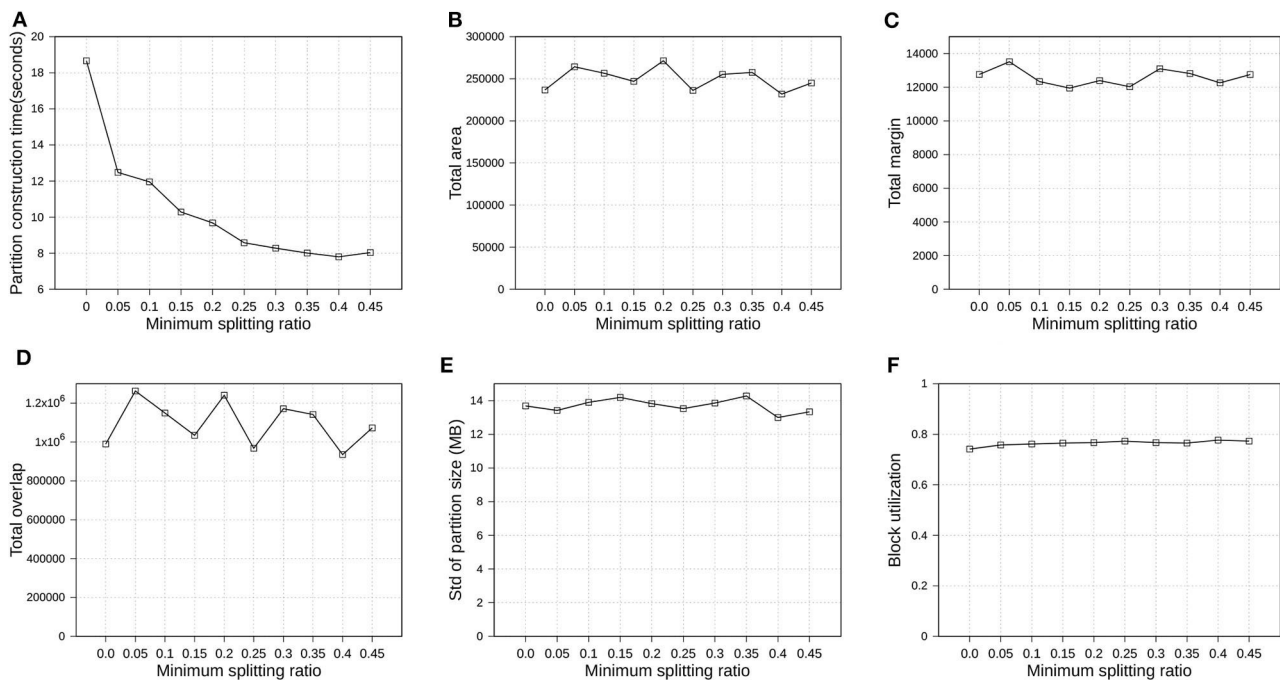


FIGURE 10 | Indexing performance and partition quality of R*-Grove in OSM-Objects datasets with different minimum splitting ratios. **(A)** Partition construction time, **(B)** Total area, **(C)** Total margin, **(D)** Total overlap, **(E)** Load balance, **(F)** Block utilization.

previous experiment in section 6.4.4. We vary ρ from 0 to 0.45. **Figure 10** shows the overview of the experimental results. First, **Figure 10A** shows that the running time of Phase 2, boundary computation, decreases as ρ increases which is expected due to the balanced splitting in the recursive algorithm which causes it to terminate earlier. According to the run-time analysis in section 4.1, a larger value of ρ reduces the depth of the recursive formula which results in a lower running time. However, this minimum splitting ratio also shrinks the search space for optimal partitioning scheme. Fortunately, the number of records in the 1% sample is usually large enough such that the boundary computation algorithm could still find a good partitioning scheme even for high value of ρ . In the following experiments, we choose $\rho = 0.4$ as the default value for R*-Grove partitioning.

6.5. Spatial Query Performance

6.5.1. Range Query

Figure 7F shows the performance of range query on the OSM-Nodes dataset with size 200 GB. For partitioned OSM-Nodes dataset, we run a number of range queries (from 200 to 1,200) all with the same range query size which is 0.01% of the area covered by the entire input. All the queries are sent in one batch to run in parallel to put the cluster at full utilization. It is clear that R*-Grove outperforms all other techniques, especially when we run a large number of queries. This is the result of the high-quality and load-balanced partitions which minimize the number of blocks needed to process for each query. **Figure 8F** shows the average cost of a range query on the OSM-Objects dataset in terms of number of blocks that need to be processed, the lower the better. This value is also computed for a range query with size 0.01% of space area. This result further confirms that R*-Grove provide a better query performance for variable-size records datasets.

6.5.2. Spatial Join

In this experiment, we split OSM-Parks and OSM-Roads datasets to get multiple datasets as follows: Parks1, Park2 with sizes 3.6 and 7.2 GB; Roads1 and Roads2 with sizes 10 and 20 GB, respectively. This allows us to study the effect of the input size on the spatial join query while keeping the input data characteristics the same, i.e., distribution and geometry size. We

compare to STR since it is the best competitor of R*-Grove in previous experiments. **Figure 11** shows the performance of the spatial join query. In general, R*-Grove significantly outperforms STR in all query instances.

Figure 11A shows the number of accessed blocks for each spatial join query over the datasets which are partitioned by R*-Grove and STR. We can notice that R*-Grove needs to access 40–60% fewer blocks than STR for two reasons. First, the better load balance in R*-Grove reduces the overall number of blocks in each dataset. Second, the higher partition quality in R*-Grove results in fewer overlapping partitions between the two datasets. The number of accessed blocks is an indicator to estimate the actual performance of spatial join queries. Indeed, this is further verified in **Figure 11B**, which shows actual running time for those queries. As we described, STR does not produce high quality partitions, thus the compound effect will even make it worst for spatial join query, which always relates to multiple STR partitioned datasets. On the other hand, R*-Grove addresses the limitations of STR so it can significantly improve the performance of spatial join query.

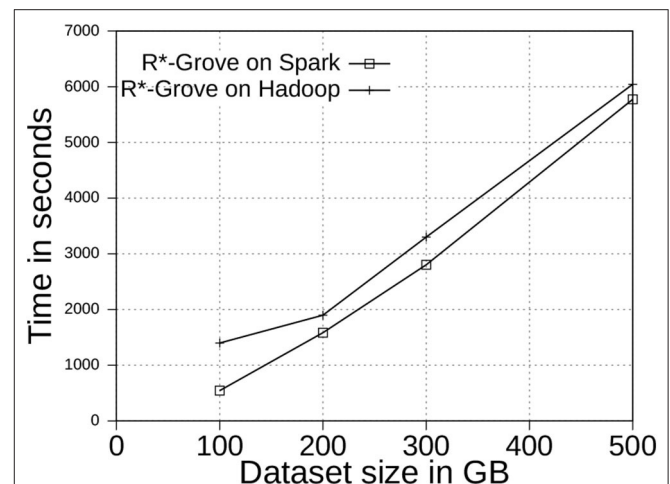


FIGURE 12 | The scalability of R*-Grove partitioning in Spark and Hadoop.

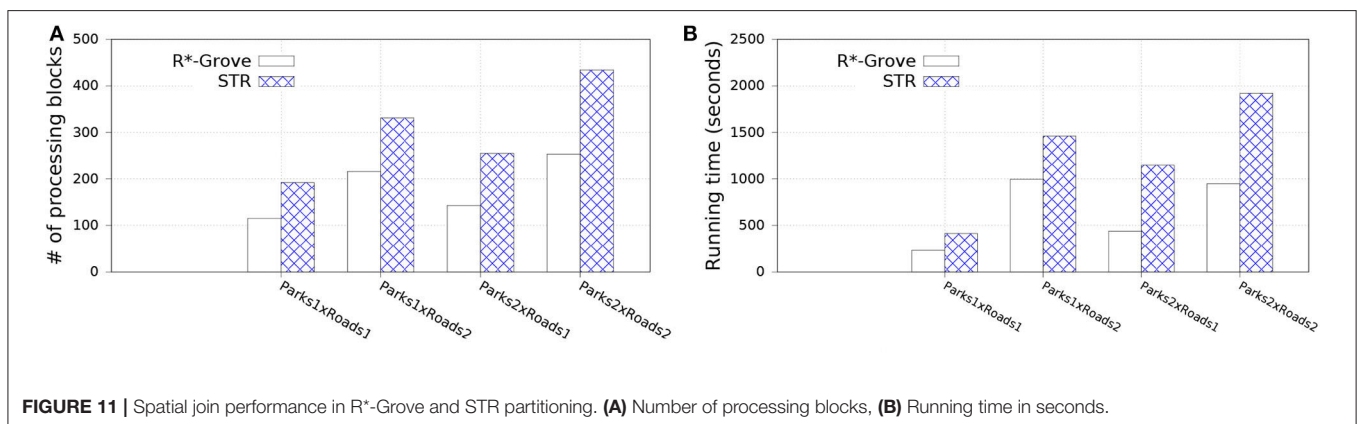
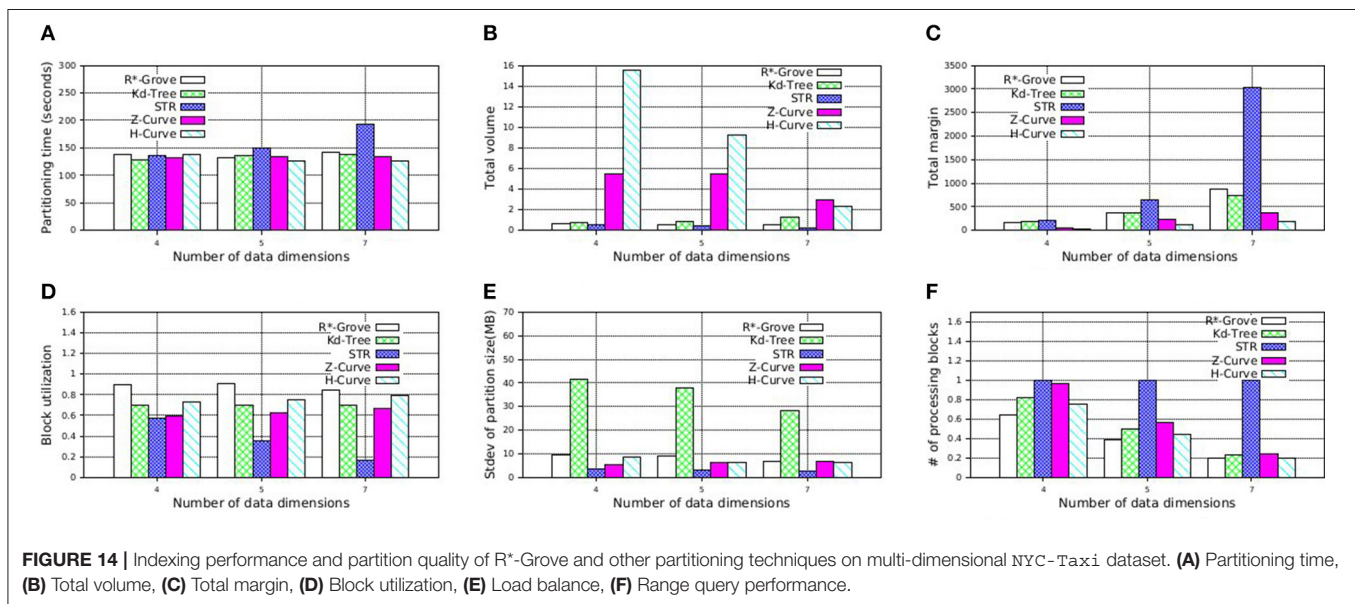
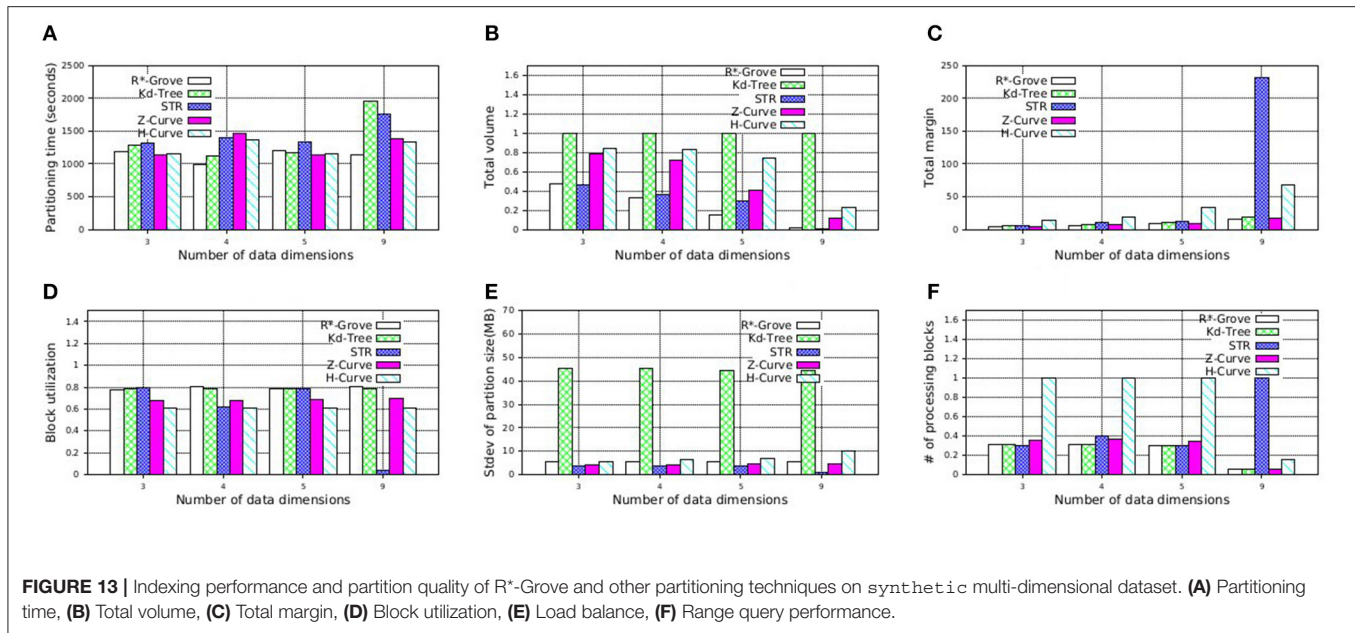


FIGURE 11 | Spatial join performance in R*-Grove and STR partitioning. (A) Number of processing blocks, (B) Running time in seconds.



6.6. Performance on Larger Datasets and Multi-Dimensional Data

6.6.1. Scalability

Figure 12 shows the indexing time for two-dimensional OSM-Nodes dataset with sizes 100, 200, 300, and 500 GB. We executed the same indexing jobs in both Spark and Hadoop to see how the processing model affects the indexing performance. We observed that Spark outperforms Hadoop in terms of total indexing time. This experiment also demonstrates that R*-Grove is ready to work with large volume datasets on both Hadoop and Spark. We also observe that the gap between Hadoop and Spark decreases as the input size increases as Spark starts to spill more data to disk.

6.6.2. Multi-Dimensional Datasets

In this experiment, we study the quality of R*-Grove on multi-dimensional datasets. Inspired by Beckmann and Seeger (2008), we use four synthetic datasets with number of dimensions 3, 4, 5, and 9. We measure the running time and the quality of the five partitioning techniques: R*-Grove, STR, Kd-Tree, Z-Curve, and H-Curve. Figure 13A shows that R*-Grove is mostly the fastest technique to index the input dataset due to the best load balance among partitions. Figure 13B shows that R*-Grove significantly reduces total area of partitions. Figure 13C shows the total margin of all the techniques. While the total margin varies with the number of dimensions since they are different datasets, the techniques maintain the same order in terms of

quality from best to worst, i.e., R*-Grove, Z-Curve, Kd-tree, STR, and H-Curve, except the last group, where H-Curve is better than STR. This experiment indicates that R*-Grove could maintain its characteristics for multi-dimensional datasets. **Figures 13D,E** report the block utilization and standard deviation of partition size, respectively. R*-Grove is the best technique that keeps both measures good. **Figure 13F** depicts the normalized range query performance of different techniques, which affirms the advantages of R*-Grove. Notice that this is the only experiment where Z-Curve performs better than H-Curve. The reason is that the generated points are generated close to a diagonal line in the d -dimension. Since the Z-Curve just interleaves the bits of all dimensions, it will result in sorting these points along the diagonal line which results in a good partitioning. However, the way H-Curve rotates the space with each level will cause it to jump across the diagonal.

Additionally, STR becomes very bad as the number dimensions increases. This can be explained by the way STR computes the number of partitions given a sample data points. The existing STR implementation always creates a tree with a fixed node degree n and d levels where d is the number of dimensions. This configuration results in n^d leaf nodes or partitions. It computes the node degree n as the smallest integer that satisfies $n^d \geq P$ where P is the number of desired partitions. For example, for an input dataset of 100 GB, $d = 9$ dimensions, and a block size of $B = 128$ MB, the number of desired partitions $P = 100 \cdot 1,024/128 = 800$ partitions and $n = 3$. This results in a total of $3^9 = 19,683$ partitions. Obviously, as d increases, the gap between the ideal number of partitions P and the actual number of partitions n^d increases which results in a very poor block utilization as shown in this experiment. Finally, **Figure 13F** shows the average cost of a range query in terms of number of processed blocks, which indicates that R*-Grove is the winner when we want to speed up spatial query processing.

To further support our findings, we also execute similar experiment on NYC-Taxi dataset, which contains up to seven dimensions as follows: *pickup_latitude*, *pickup_longitude*, *dropoff_latitude*, *dropoff_longitude*, *pickup_datetime*, *trip_time_in_secs*, *trip_distance*. These attribute values are normalized in order to avoid the dominance of some columns. We decide to partition this dataset using multiple attributes which is picked in the aforementioned order with size 4, 5, and 7. **Figure 14** shows that R*-Grove balances all the different quality metrics. Specially, **Figure 14F** indicates that R*-Grove is the winner when compared to other techniques in terms of spatial query performance. We also notice that H-Curve performs better than Z-Curve with this real dataset. We conclude that R*-Grove is a better option for indexing multi-dimensional spatial data since it outperforms

or got an equivalent performance with other indexes in all metrics.

7. CONCLUSION

This paper proposes R*-Grove, a novel partitioning technique which can be widely used in many big spatial data processing systems. We highlighted three limitations in existing partitioning techniques such as STR, Kd-Tree, Z-Curve, and Hilbert Curve. These limitations are the low quality of the partitions, the imbalance among partitions, and the failure to handle variable-size records. We showed that R*-Grove overcomes these three limitations to produce high quality partitions. We showed three case studies in which R*-Grove can be used to facilitate big spatial indexing, range query, and spatial join. An extensive experimental evaluation was carried out on big spatial datasets and showed that R*-Grove is scalable and speeds up all the operations in the case studies. We believe that R*-Grove promises to be a good replacement to existing big spatial data partitioning techniques in many systems. In the future, we will further study the proposed technique for in-memory and streaming applications to see how it behaves under these architectures.

DATA AVAILABILITY STATEMENT

The datasets generated for this study are available on the UCR Spatio-temporal Active Repository (UCR-STAR, <https://star.cs.ucr.edu/>) or on request to the corresponding author. In particular, we used OSM2015/all_nodes, OSM2015/roads, OSM2015/parks, OSM2015/all_objects, NYCTaxi. For the diagonal points dataset, we generated them using the spatial data generator (Vu et al., 2019) with following parameters: dataset size $|D| = 80$ million points; number of dimensions $d = 3, 4, 5, 9$; the percentage (ratio) of the points that are exactly on the line $perc = 0.05$; the size of the buffer around the line where additional points are scattered $buf = 0.1$.

AUTHOR CONTRIBUTIONS

AE and TV worked on the theoretical proofs, the design, and implementation of the algorithms. TV wrote the manuscript and carried out the experimental evaluation with the guidance of AE. All authors contributed to the article and approved the submitted version.

FUNDING

This work was supported in part by the National Science Foundation (NSF) under grants IIS-1838222 and CNS-1924694.

REFERENCES

allobjects (2019). *Openstreetmap All Objects Dataset*. Available online at: http://star.cs.ucr.edu/#dataset=OSM2015/all_objects

Beckmann, N., Kriegel, H., Schneider, R., and Seeger, B. (1990). "The R*-tree: an efficient and robust access method for points and rectangles," in *SIGMOD* (Atlantic City, NJ), 322–331. doi: 10.1145/93597.98741

- Beckmann, N., and Seeger, B. (2008). A benchmark for multidimensional index Structures.
- Beckmann, N., and Seeger, B. (2009). "A revised R*-tree in comparison with related index structures," in *SIGMOD* (Providence, RI), 799–812. doi: 10.1145/1559845.1559929
- Bentley, J. L. (1975). Multidimensional binary search trees used for associative searching. *Commun. ACM* 18, 509–517. doi: 10.1145/361002.361007
- Chasparis, H., and Eldawy, A. (2017). "Experimental evaluation of selectivity estimation on big spatial data," in *Proceedings of the Fourth International ACM Workshop on Managing and Mining Enriched Geo-Spatial Data* (Chicago, IL), 8:1–8:6.
- Dittrich, J., and Seeger, B. (2000). "Data redundancy and duplicate detection in spatial join processing," in *Proceedings of the 16th International Conference on Data Engineering* (San Diego, CA), 535–546. doi: 10.1109/ICDE.2000.839452
- Eldawy, A., et al. (2017). "Sphinx: empowering impala for efficient execution of SQL queries on big spatial data," in *SSTD* (Arlington, VA), 65–83. doi: 10.1007/978-3-319-64367-0_4
- Eldawy, A., Alarabi, L., and Mokbel, M. F. (2015a). Spatial partitioning techniques in spatial hadoop. *PVLDB* 8, 1602–1605. doi: 10.14778/2824032.2824057
- Eldawy, A., Li, Y., Mokbel, M. F., and Janardan, R. (2013). "Cg_hadoop: computational geometry in mapreduce," in *SIGSPATIAL* (Orlando, FL), 284–293. doi: 10.1145/2525314.2525349
- Eldawy, A., and Mokbel, M. F. (2015). "Spatialhadoop: a mapreduce framework for spatial data," in *ICDE* (Seoul), 1352–1363. doi: 10.1109/ICDE.2015.7113382
- Eldawy, A., and Mokbel, M. F. (2016). The era of big spatial data: a survey. *Found. Trends Databases* 6, 163–273. doi: 10.1561/9781680832259
- Eldawy, A., Mokbel, M. F., Al-Harathi, S., Alzaidy, A., Tarek, K., and Ghani, S. (2015b). "SHADED: A mapreduce-based system for querying and visualizing spatio-temporal satellite data," in *ICDE* (Seoul), 1585–1596. doi: 10.1109/ICDE.2015.7113427
- Eldawy, A., Mokbel, M. F., and Jonathan, C. (2016). "Hadoopviz: A mapreduce framework for extensible visualization of big spatial data," in *ICDE* (Helsinki), 601–612. doi: 10.1109/ICDE.2016.7498274
- Fox, A. D., Eichelberger, C. N., Hughes, J. N., and Lyon, S. (2013). "Spatio-temporal indexing in non-relational distributed databases," in *Big Data* (Santa Clara, CA), 291–299. doi: 10.1109/BigData.2013.6691586
- Ghosh, S., Eldawy, A., and Jais, S. (2019). "Aid: An adaptive image data index for interactive multilevel visualization," in *ICDE* (Macau). doi: 10.1109/ICDE.2019.00150
- Goodchild, M. F. (2007). Citizens as voluntary sensors: Spatial data infrastructure in the world of web 2.0. *IJSDIR* 2, 24–32.
- Guttman, Bughin, J., Chui, M., Manyika, J., Saleh, T., Wiseman, B., A. (1984). "R-trees: A dynamic index structure for spatial searching," in *SIGMOD* (Boston, MA), 47–57. doi: 10.1145/602259.602266
- Henke, N., Bughin, J., Chui, M., Manyika, J., Saleh, T., Wiseman, B., et al. (2016). *The Age of Analytics: Competing in a Data-Driven World*. Technical Report, McKinsey Global Institute.
- Hoel, E. G., and Samet, H. (1994). "Performance of data-parallel spatial operations," in *Vldb'94, Proceedings of 20th International Conference on Very Large Data Bases* (Santiago de Chile), 156–167.
- Jacox, E. H., and Samet, H. (2007). Spatial join techniques. *ACM Trans. Database Syst.* 32:7. doi: 10.1145/1206049.1206056
- Kamel, I., and Faloutsos, C. (1994). "Hilbert R-tree: an improved r-tree using fractals," in *Vldb* (Santiago de Chile), 500–509.
- Lee, D.-T., and Wong, C. (1977). Worst-case analysis for region and partial region searches in multidimensional binary search trees and balanced quad trees. *Acta Inform.* 9, 23–29. doi: 10.1007/BF00263763
- Lee, T., and Lee, S. (2003). "Omt: Overlap minimizing top-down bulk loading algorithm for R-tree," in *CAISE Short Paper Proceedings*, Vol. 74, 69–72.
- Leutenegger, S. T., Lopez, M. A., and Edgington, J. (1997). "STR: A simple and efficient algorithm for r-tree packing," in *ICDE* (Birmingham), 497–506. doi: 10.1109/ICDE.1997.582015
- Li, Y., Eldawy, A., Xue, J., Knorozova, N., Mokbel, M. F., and Janardan, R. (2019). Scalable computational geometry in MapReduce. *Vldb J.* 28, 523–548. doi: 10.1007/s00778-018-0534-5
- Lo, M., and Ravishanker, C. V. (1994). "Spatial joins using seeded trees," in *SIGMOD* (Minneapolis, MN), 209–220. doi: 10.1145/191839.191881
- Lu, P., Chen, G., Ooi, B., Vo, H., and Wu, S. (2014). ScalaGiST: scalable generalized search trees for MapReduce systems. *PVLDB*, 7, 1797–1808. doi: 10.14778/2733085.2733087
- Magdy, A., Alarabi, L., Al-Harathi, S., Musleh, M., Ghanem, T. M., Ghani, S., et al. (2014). "Taghreed: a system for querying, analyzing, and visualizing geotagged microblogs," in *SIGSPATIAL* (Dallas, TX; Fort Worth, TX), 163–172. doi: 10.1145/2666310.2666397
- Nishimura, S., Das, S., Agrawal, D., and El Abbadi, A. (2013). MD-hbase: design and implementation of an elastic data infrastructure for cloud-scale location services. *Distrib. Parallel Databases* 31, 289–319. doi: 10.1007/s10619-012-7109-z
- Sabek, I., and Mokbel, M. F. (2017). "On spatial joins in mapreduce," in *SIGSPATIAL* (Redondo Beach, CA), 21:1–21:10. doi: 10.1145/3139958.3139967
- Samet, H. (1984). The quadtree and related hierarchical data structures. *ACM Comput. Surveys* 16, 187–260. doi: 10.1145/356924.356930
- Siddique, A. B., Eldawy, A., and Hristidis, V. (2019). Comparing synopsis techniques for approximate spatial data analysis. *Proc. VLDB Endow.* 12, 1583–1596. doi: 10.14778/3342263.3342635
- Tang, M., Yu, Y., Malluhi, Q. M., Ouzzani, M., and Aref, W. G. (2016). LocationSpark: a distributed in-memory data management system for big spatial data. *PVLDB* 9, 1565–1568. doi: 10.14778/3007263.3007310
- UCRSTAR (2019). *The UCR Spatio-Temporal Active Repository (UCR-STAR)*. Available online at: <https://star.cs.ucr.edu/>
- Vo, H., Aji, A., and Wang, F. (2014). "SATO: a spatial data partitioning framework for scalable query processing," in *SIGSPATIAL* (Dallas, TX; Fort Worth, TX), 545–548. doi: 10.1145/2666310.2666365
- Vu, T., and Eldawy, A. (2018). "R-Grove: growing a family of R-trees in the big-data forest," in *SIGSPATIAL* (Seattle, WA), 532–535. doi: 10.1145/3274895.3274984
- Vu, T., Migliorini, S., Eldawy, A., and Belussi, A. (2019). "Spatial data generators," in *1st ACM SIGSPATIAL International Workshop on Spatial Gems (SpatialGems 2019)*, 7.
- Whitman, R. T., Park, M. B., Ambrose, S. M., and Hoel, E. G. (2014). "Spatial indexing and analytics on Hadoop," in *SIGSPATIAL* (Dallas, TX; Fort Worth, TX), 73–82. doi: 10.1145/2666310.2666387
- Xie, D., Li, F., Yao, B., Li, G., Zhou, L., and Guo, M. (2016). "Simba: Efficient in-memory spatial analytics," in *SIGMOD* (San Francisco, CA), 1071–1085. doi: 10.1145/2882903.2915237
- Yu, J., Wu, J., and Sarwat, M. (2015). "Geospark: a cluster computing framework for processing large-scale spatial data," in *SIGSPATIAL* (Bellevue, WA), 70:1–4. doi: 10.1145/2820783.2820860
- Zhou, X., Abel, D. J., and Truffet, D. (1998). Data partitioning for parallel spatial join processing. *GeoInformatica* 2, 175–204. doi: 10.1023/A:1009755931056

Conflict of Interest: The authors declare that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

Copyright © 2020 Vu and Eldawy. This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY). The use, distribution or reproduction in other forums is permitted, provided the original author(s) and the copyright owner(s) are credited and that the original publication in this journal is cited, in accordance with accepted academic practice. No use, distribution or reproduction is permitted which does not comply with these terms.