

The V*-Diagram: A Query-Dependent Approach to Moving KNN Queries

Sarana Nutanong^{†‡} Rui Zhang[†] Egemen Tanin^{†‡} Lars Kulik^{†‡}
 {sarana,rui,egemen,lars}@csse.unimelb.edu.au

[†]Department of Computer Science and Software Engineering
 University of Melbourne
 Victoria, Australia

[‡]NICTA Victoria Laboratory
 Australia

ABSTRACT

The moving k nearest neighbor (MkNN) query finds the k nearest neighbors of a moving query point continuously. The high potential of reducing the query processing cost as well as the large spectrum of associated applications have attracted considerable attention to this query type from the database community. This paper presents an incremental safe-region-based technique for answering MkNN queries, called the V*-Diagram. In general, a safe region is a set of points where the query point can move without changing the query answer. Traditional safe-region approaches compute a safe region based on the data objects but independent of the query location. Our approach exploits the current knowledge of the query point and the search space in addition to the data objects. As a result, the V*-Diagram has much smaller IO and computation costs than existing methods. The experimental results show that the V*-Diagram outperforms the best existing technique by two orders of magnitude.

1. INTRODUCTION

Current location-based services provide accurate position information with a high degree of temporal precision. Consider the following two scenarios. A driver in a GPS-equipped car issues a continuous query to find the nearest gas station while driving in a city. A tourist uses a location-aware mobile device to issue a continuous query for the nearest restaurant while walking to a museum. The queries are sent to a server that processes the queries and returns the answers. In these scenarios, the server has to continuously maintain the answer set which may change depending on the location of the query point. These queries are *location-based continuous spatial queries* [23] and the scenarios above are typical examples of *moving k nearest neighbor queries* (MkNN). A straightforward way to process a MkNN query is using a *sampling-based* method, which processes the MkNN query as a kNN query at sampled locations. This method does not provide answers between sampled locations. In order to provide an (almost) continuous answer to the query, a high sampling rate is required, which makes the method inefficient due to the frequent processing of kNN queries. The concept of the

safe region provides a more effective way to achieve continuous answers to location-based spatial queries. In a safe-region-based method, an answer is returned with a safe region. As long as the query point stays in the safe region, the answer remains the same. When the query point moves out of the safe region, another answer with its associated region is returned. Therefore, a safe-region-based method always (that is, continuously) provides accurate answers without the need for sampling. This approach also requires much less frequent communication between the mobile client and the server.

A classic example of safe-region-based techniques is the *Voronoi Diagram* [14]. The Voronoi Diagram is a well known space decomposition determined by distances to a given discrete set of objects, typically, a set of points. Specifically, the Voronoi Diagram of a set of points $P = \{p_1, p_2, \dots, p_n\}$ is defined as a set of cells where each cell $V(p_i)$ is a region of space that consists of all points of the data space that are closer to p_i than any other point in P . An example is given in Figure 1. Figure 1(a) is a set of points in a 2-dimensional (2D) space and Figure 1(b) is the Voronoi Diagram.

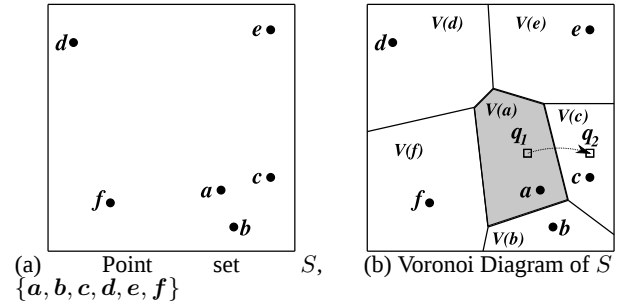


Figure 1: The Voronoi Diagram

Processing a 1NN query using the Voronoi Diagram involves: (i) locating which Voronoi cell the query point falls into; and (ii) identifying the associated object. In the above example, q_1 falls in $V(a)$ (the grey region) and therefore a is the nearest neighbor of q_1 . The answer remains valid as long as the query point stays in $V(a)$. When the query point moves across the boundary of $V(a)$ to $V(c)$, c becomes the NN.

The Voronoi Diagram can be generalized to the k^{th} -order Voronoi Diagram (kVD). In a kVD , each region is associated with the set of k nearest neighbors, termed kNN set or k NNs, rather than only the nearest one. The kVD can handle MkNN queries in the same manner as the basic first order VD handles 1NN queries.

Another useful generalization of the Voronoi Diagram is order sensitivity. The *ordered kVD* partitions the space into cells where each cell is associated to a particular ordering of the kNN set. The

Permission to copy without fee all or part of this material is granted provided that the copies are not made or distributed for direct commercial advantage, the VLDB copyright notice and the title of the publication and its date appear, and notice is given that copying is by permission of the Very Large Data Base Endowment. To copy otherwise, or to republish, to post on servers or to redistribute to lists, requires a fee and/or special permission from the publisher, ACM.

VLDB '08, August 24-30, 2008, Auckland, New Zealand
 Copyright 2008 VLDB Endowment, ACM 000-0-00000-000-0/00/00.

ordered k VD can be used to answer $MkNN$ queries that require the ranking of the k NNs by their distances to the query point.

The k VD has the following shortcomings:

1. **Expensive precomputation.** The k VD requires precomputing all the k VD cells and access to all the data points. Both computation and storage costs are high.
2. **No support for dynamically changing k values.** The k VD can only accommodate kNN queries with a specific k value; the ordered k VD can only accommodate kNN queries with k values no larger than the order of the diagram. As a result, the technique is not suitable for situations where the value of k is unknown in advance or can change dynamically.
3. **Inefficient update operations.** Many cells have to be recomputed for each insertion or deletion on the dataset.

The expensive precomputation is especially not justified if the query point is confined to a small region of the whole data space. For example, if a car is moving in a small part of a city, then it is unnecessary to compute the k VD for the entire city. Furthermore, one may require different k VDs for different needs. For example, a driver may need to find a gas station with a restroom facility while another driver needs one with a special type of fuel. Precomputing k VDs for all possible scenarios may be prohibitive.

In [23], Zhang et al. proposed an algorithm to locally compute a k VD cell, which mitigates the precomputation and update problems (shortcomings 1 and 3). However, the algorithm is still relatively expensive and does not address the problem of dynamically changing k values (shortcoming 2).

In this paper, a technique called the V^* -Diagram for $MkNN$ queries is proposed. The V^* -Diagram has the following key advantages:

1. It requires no precomputation.
2. It incrementally computes answers and therefore efficiently adapts to changes – such as insertions and deletions of objects, as well as, dynamically changing values of k .

The V^* -Diagram is based on the safe-region concept, but differs from any previous technique for $MkNN$ queries in the following aspect: previous safe-region-based techniques compute safe regions purely based on the data (for example, you can compute the k VD without referring to the query point); **the V^* -Diagram computes safe regions based on not only the data objects, but also the query point and the current knowledge of the search space.** This is one of the main novelties of the technique. By doing so, both computation and data retrieval of the V^* -Diagram are more economical than those of the other techniques.

The contributions of this paper are summarized as follows:

- We propose the V^* -Diagram technique and the associated algorithm, called V^*-kNN , to support efficient processing of $MkNN$ queries.
- We show how the V^* -Diagram technique can be applied to the domain of spatial networks.
- We perform an extensive experimental study with the results showing that the V^* -Diagram outperforms the best existing technique [23] by two orders of magnitude.

The rest of the paper is organized as follows: Section 2 describes the problem setup. Related work is discussed in Section 3. We formulate the V^* -Diagram in Section 4 and present the algorithm for $MkNN$ queries based on the V^* -Diagram in Section 5. In Section 6, we show how the V^* -Diagram technique can be applied to the domain of spatial networks. A performance analysis of the V^* -Diagram is given in Section 7. Section 8 presents experimental results and Section 9 concludes the paper.

2. PROBLEM SETUP

Let D be a set of data objects (represented by points) in a d -dimensional space and $dist(\cdot)$ be the distance function. The k

nearest neighbor (kNN) query is defined as follows: given D and a static query point q , find a set H that consists of k objects from D such that for any $p_1 \in H$ and any $p_2 \in D - H$, $dist(p_1, q) \leq dist(p_2, q)$.

The moving k nearest neighbor query ($MkNN$) is defined as follows: given D and a moving query point q , find the k NNs of q for every position of q .

Due to the nature of location-based applications, $MkNN$ queries are discussed in the context of two settings: (i) **Centralized processing paradigm.** Both the query issuer and the processor are on the same machine. Then the main performance measure is the query processing cost. (ii) **Client-server paradigm.** The query is issued by a client to a server through a wireless network (such as a mobile phone network). The performance measure involves both communication and query processing costs on the server side, and the former one is more important in delay-sensitive applications.

We assume an unknown trajectory which means that the location of q gets updated in a periodic manner. We also assume that no k VD is maintained but there is a generic spatial index such as the R-tree [4] built on the data objects, since it can be used for various query types and is efficient to maintain. This is also argued as a valid assumption in previous work [23].

3. RELATED WORK

3.1 KNN algorithms

Many kNN search algorithms have been proposed based on spatial hierarchical structures. One of their common features is the application of the branch-and-bound strategy on the tree structure. The tree can be traversed in a depth-first ($DF-kNN$) [16] or a best-first ($BF-kNN$) [5] manner. $BF-kNN$ can retrieve more nearest neighbors incrementally if k increases.

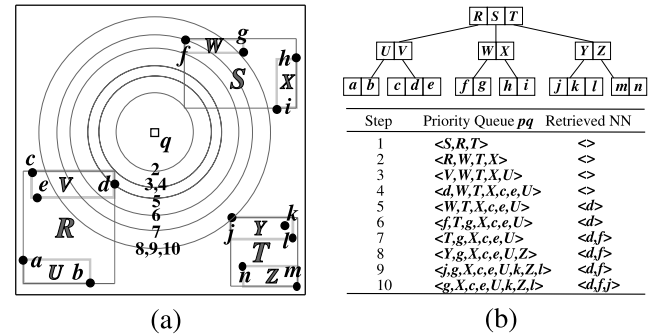


Figure 2: R-tree and $BF-kNN$

Figure 2 shows an example R-tree and the $BF-kNN$ algorithm. Figure 2(a) shows a set of points and how they can be grouped in an R-tree, where objects are indexed in a hierarchy of *minimum bounding rectangles* (MBRs). The corresponding R-tree is shown in the upper part of Figure 2(b). Starting from the root, $BF-kNN$ traverses all the entries in increasing order of their *mindist*, where the *mindist* of an entry is the minimum distance between the entry and the query point q . To do so, a priority queue is maintained to keep all the retrieved data points and *active* nodes. A node is active if its parent has been accessed but itself has not. The traversal stops if the first k elements retrieved from the priority queue are all data points.

An example run of $BF-kNN$ is shown in the lower part of Figure 2(b). At step 1, all entries in the root are inserted into the priority queue pq . Then the entries are dequeued and the corresponding nodes are retrieved in order. The first dequeued item is S and its

two child entries X and W are put back into pq . Then R is dequeued. Nodes U and V are put in pq and so on. At step 4, data point d is the head of pq . Data point d must be the first NN because all other entries have distances to q (which are bounded by $mindist$ of their nodes) larger than $dist(q, d)$. d can be retrieved as the first NN. If another NN is needed, the process continues until another data point is the head of pq . At step 6, f is discovered as the second NN. By this means, an arbitrary number of NNs can be incrementally obtained. If the value of k is fixed, an aggressive pruning can be performed on the nodes in pq to reduce the queue size, though the page access cost cannot be further reduced due to the search-space optimality of the algorithm [6]. Other techniques such as iDistance [10] has superior performance than these hierarchical structure based algorithms in high dimensional space. However, it cannot retrieve k NNs in an incremental manner.

3.2 Techniques for processing M_kNN queries

We have discussed two approaches for M_kNN queries in Section 1: one is sampling based and the other one is safe-region based.

SR- kNN . Song and Roussopoulos [18] introduced a method which will be referred to as SR- kNN in this paper. SR- kNN reduces the access costs in the sampling-based approach by retrieving redundant data entries and caching. It greatly reduces the cost of query reevaluation, but does not solve the problem of inaccurate answers between sampled locations. Thus, SR- kNN does not provide *truly continuous* answers.

The k^{th} -order Voronoi Diagram (kVD). Safe-region based techniques produce continuous answers and reduce processing and communication costs. The Voronoi Diagram [14] is a classic example and can be used to process MINN queries as described in Section 1. For M_kNN queries, the kVD is used. However, the kVD has shortcomings described in Section 1.

TP- kNN and C- kNN . Tao and Papadias [19] proposed the *time-parameterized kNN* (TP- kNN) query. Assuming a linear trajectory of the query point, a TP- kNN query finds (i) the current kNN set, (ii) a position on the trajectory where the kNN set changes and (iii) the objects that cause the change. This is done by finding the earliest point on the trajectory that has a different kNN set from the current one. This point is also known as the *influence point* (or equivalently *influence time* when the speed is known), which can be considered as the boundary of a safe region.

Tao and Papadias [20] also considered the *Continuous kNN* (C- kNN) query, which finds the kNN for every single point on a pre-defined linear trajectory. This is achieved by identifying all influence points on the trajectory. The main difference between C- kNN and TP- kNN is that C- kNN obtains all the influence points on the trajectory but TP- kNN finds only the first one. Both TP- kNN and C- kNN are limited to known linear trajectories.

RIS- kNN . Zhang et al. [23] proposed an algorithm called *Retrieve-Influence-Set kNN* (RIS- kNN) to locally compute kVD cells using a spatial index. RIS- kNN uses the TP- kNN query [19] to find each edge of a kVD cell, 360 degrees around the query point.

Figure 3 shows how RIS- kNN discovers all edges of $V(a)$ from the example in Figure 1(a). Initially, a is found to be the NN of q ; the cell is initialized to the whole data space, that is, rectangle $ABCD$. At step 1, a TPNN query is executed with the trajectory from q to the top left corner of the space ($\vec{qA}$) and d is returned as the object that changes the NN result along $\vec{qA}$. The perpendicular bisector of a and d , B_{ad} , contributes one edge to the cell. The cell is updated to the polygon $BCDEF$. At steps 2 and 3, two TPNN queries are executed with the trajectories from q to the new corners ($\vec{qE}$ and $\vec{qF}$). Two new edges are found according to B_{af} and B_{ae} . Since kVD cells for data points are convex polygons, this process continues until all corners of the cell have been checked

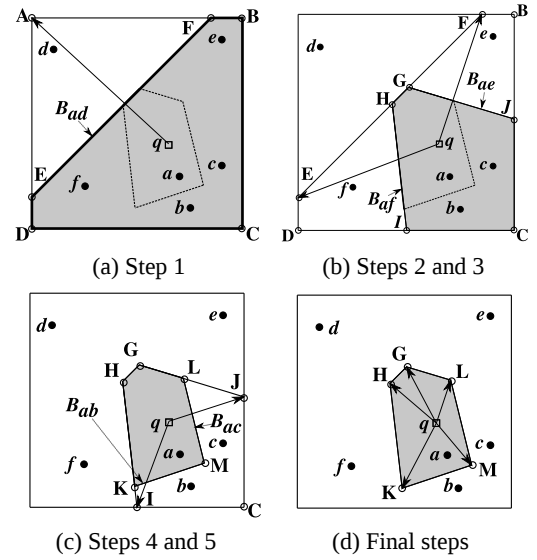


Figure 3: Computing a Voronoi cell locally

and they all have the same NN as q .

RIS- kNN mitigates the precomputation problem (shortcoming 1) of the kVD because it only accesses local data, but this algorithm is still expensive because it performs multiple (on average 12 [23]) TP- kNN queries, where each TP- kNN query involves a costly tree traversal. RIS- kNN does not solve the problem of dynamically changing k values (shortcoming 2); changing k to a larger value incurs recalculation of the kVD cell. The computation of the previous kVD cell cannot be reused and hence this algorithm is not *incremental*.

Due to the fact that only one kVD cell is maintained at a time, RIS- kNN handles dataset updates (insertions and deletions of objects) more efficiently than the traditional kVD technique. However, in a case where an update effects the kNN answers, the current kVD cell has to be recalculated.

IRU. Kulik and Tanin [13] introduced an algorithm called *incremental rank updates* (IRU) to compute regions where the ranking of all the objects (based on their distances) is the same. This is equivalent to computing the ordered kVD cell with $k = n$, where n is the total number of objects. Rather than computing the whole nVD , IRU incrementally computes a neighboring nVD cell from the current cell. In [13], an nVD cell is termed a *fixed-rank region* (FRR) since for any point in the region, the ranking of all the objects based on their distances is fixed. Based on the observation that only rank-adjacent objects can swap their ranks¹, defining the FRR of n objects requires at most $n - 1$ bisectors² of the $n - 1$ pairs of rank-adjacent objects. Continuous monitoring of the ranking of the objects is done by maintaining a rank-sorted list of objects and its corresponding list of bisectors of pairs of rank-adjacent objects (*rank-adjacent bisectors*). Each time a bisector is crossed by the query point, the ranks of the two corresponding objects are swapped and the list of rank-adjacent bisectors are updated.

An example is given in Figure 4, where the grey region is the current FRR that q_1 is in. Let us assume that q is the location of a moving query point which starts at q_1 and stops at q_2 . In Figure 4(a), q is at q_1 and the ranking is initially $\langle a, c, b, f, e, d \rangle$ and the corre-

¹In this paper, the rank of an object means the object's position in a list of objects sorted by their distances to some other object. We use "ranking" and "ordering" interchangeably due to the usage of both in the literature.

²The *bisector* of two objects a and b is the set of points where each point is equidistant to a and b .

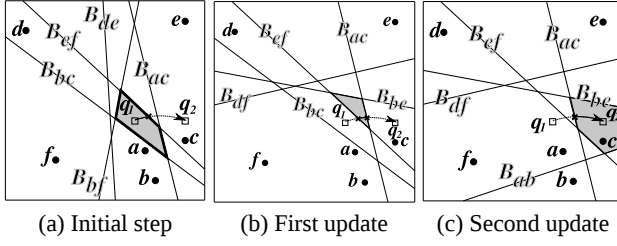


Figure 4: Incremental rank update

sponding list of bisectors is $\langle B_{ac}, B_{cb}, B_{bf}, B_{fe}, B_{ed} \rangle$. Then q crosses B_{ef} in Figure 4(b). This causes e and f to swap their ranks. Therefore B_{bf} and B_{de} are replaced by B_{be} and B_{df} , respectively. In Figure 4(c), B_{ac} is crossed. This causes a and c to swap their ranks and B_{bc} is replaced by B_{ab} . It is shown in [13] that only $\mathcal{O}(n)$ instead of $\mathcal{O}(n^2)$ bisectors are maintained during the iterations in the IRU algorithm. However, IRU still accesses all data objects to obtain the sorted list and checks $n - 1$ bisectors every time q moves.

Related spatial-network queries. Several k NN techniques for static query points were proposed in [8, 9, 15, 17]. There are also Mk NN techniques specific to the domain of spatial networks. Kolahdouzan et al. [12] proposed an algorithm that utilizes the *Voronoi Network Nearest Neighbor* (VN^3) [11]. Cho et al. [3] proposed a technique that issues static k NN queries at the intersection points on the query path.

Related moving-object queries. Hu et al. [7] proposed a safe-region-based technique for static window and k NN queries on moving objects. Each moving object maintains its own safe region and only report if its new location changes the results of any of the queries. Yu et al. [22] presented a query-indexing technique for monitoring k NN queries for moving objects and a given set of queries, where each query object maintain its own *critical region* to keep track of the k NN set. The similarity between [7, 22] and the Mk NN techniques [13, 14, 18–20, 23] including the V^* -Diagram is the use of the safe-region concept to reduce the access cost.

Benetis et al. [2] presented algorithms to process NN and reverse NN queries for moving objects with known trajectories. Result validity is thus expressed as a function of time. Similar to our work, Benetis et al. included methods to handle insertions and deletions of data points.

These techniques [2, 7, 22] focus on monitoring changes caused by location updates of data objects. The emphasis of the Mk NN techniques, on the other hand, is on the changes caused by location updates of the query point.

Summary. The Mk NN techniques are summarized in Table 1 on five features: providing **continuous** answer, **incremental** evaluation, **accessing** only **local** data (instead of all data), working on **unknown** query **path**, and providing **order-sensitive** k NNs. Only our proposed algorithm, V^* - k NN, has all these features.

Table 1: Comparison of Mk NN techniques

Technique	Continuous	Incremental	Local access	Unknown path	Order-sensitive
SB- k NN [18]	×	✓	✓	✓	✓
k VD [14]	✓	×	×	✓	×
Ordered k VD [14]	✓	×	×	✓	✓
TP k NN [19]	✓	×	✓	×	×
C k NN [20]	✓	×	✓	×	×
RIS- k NN [23]	✓	×	✓	✓	×
IRU [13]	✓	✓	×	✓	✓
V^* - k NN	✓	✓	✓	✓	✓

4. THE V^* -DIAGRAM

We formulate the V^* -Diagram in this section. The V^* -Diagram is a safe-region-based technique. Previous techniques compute safe regions purely based on the data. The V^* -Diagram computes safe regions based on not only the data, but also the query point and the current knowledge of the search space.

The V^* -Diagram assumes a metric space and a spatial hierarchical index on the dataset. Hence the BF- k NN algorithm can be used to retrieve NNs incrementally as discussed in Section 3.1.

In the V^* -Diagram, the $(k + x)$ NNs of the moving query point q are maintained, where x is the number of auxiliary objects to help the V^* -Diagram work effectively (more analysis on x is in Section 5.1). The V^* -Diagram comprises two types of regions, which are discussed in Section 4.1 and Section 4.2. These regions are then put together to form a combined safe region, discussed in Section 4.3. Commonly used symbols are summarized in Table 2.

Table 2: Symbols

Symbol	Meaning
n	The number of objects in the database.
k	The number of requested nearest neighbors.
x	The number of auxiliary objects.
q	The moving k NN query point.
q_b	The position where the latest BF- k NN call is made.
q'	The current position of the query point.
p	A data object.
p_k	The current k^{th} NN of q .
z	The $(k + x)^{th}$ NN of q_b when q is at q_b .
S_k	The safe region with regard to p_k .
$B_{p_i p_j}$	The bisector of two objects p_i and p_j .

4.1 Safe region with regard to a data object

To help explain the concept of a *safe region with regard to a data object*, the notions of *search sphere*, *known region* and *reliable region* are first introduced. Recall the BF- k NN algorithm in Section 3.1. Each object/node retrieved from the priority queue corresponds to an implicit *search sphere* (centered at the query point), which delimits the current search coverage, and the sphere expands gradually as more nodes/objects are accessed. Numbers are assigned to those spheres in Figure 2(a) according to the steps in Figure 2(b) that access the corresponding nodes. For example, sphere 2 corresponds to step 2 where node S is retrieved; sphere 5 corresponds to step 5 where object d is retrieved. Intuitively, the search sphere denotes the region we have full knowledge of, because all the objects in the sphere are already retrieved.

In the V^* -Diagram, BF- k NN is called repeatedly to help maintain $(k + x)$ NNs. Let q_b be the position of q where the latest BF- k NN call is made. Let z be the $(k + x)^{th}$ NN to q_b . The search sphere corresponding to z centered at q_b is the latest one (since BF- k NN stops when z is obtained), and we call this search sphere the *known region*, denoted by $W(q_b, z)$. We highlight z because it determines the boundary of the known region. Figure 5 gives an example. $W(q_b, z)$ is actually a sphere centered at q_b with the radius $dist(q_b, z)$. Point p is one of the $(k + x - 1)$ NNs of q_b and other objects in $W(q_b, z)$ are not shown.

Next, we formulate a region within which q can move while p remains one of the $(k + x)$ NNs of q . Let q' denote a later position of q after q_b . Suppose q' is at the position as shown in Figure 5. We extend the line segment $q_b q'$ and it exits $W(q_b, z)$ at χ . Let $sph(v, l)$ denote a sphere with center v and radius l . As long as p is in $sph(q', dist(q', \chi))$, it is one of the $(k + x)$ NNs of q' . This is because any object outside $sph(q', dist(q', \chi))$ must be farther to q' than p to q' and there are at most $(k + x)$ objects inside $sph(q', dist(q', \chi))$. Since any object in $sph(q', dist(q', \chi))$ remains one of the $(k + x)$ NNs of q' , we call $sph(q', dist(q', \chi))$

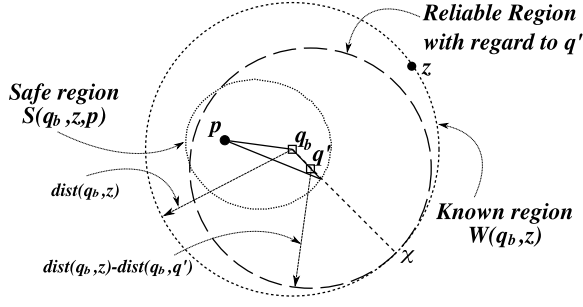


Figure 5: The known, reliable, and safe regions

the *reliable region* with regard to q ³ and any object in the reliable region a *reliable object*. If p is a reliable object, p is said to be *reliable*; otherwise, it is said to be *unreliable*. Mathematically, p being in the reliable region with regard to q' is expressed as

$$\text{dist}(q', p) \leq \text{dist}(q_b, z) - \text{dist}(q_b, q'), \quad (1)$$

where $\text{dist}(q_b, z) - \text{dist}(q_b, q')$ is the length of $\overline{q'z}$.

If we view q' as a variable, then Equation (1) actually describes all the possible positions of q' that guarantee p remaining among the $(k+x)$ NNs. Consequently, we can formulate the safe region with regard to p as follows.

DEFINITION 1 (SAFE REGION WITH REGARD TO A POINT). Given a known region $W(q_b, z)$ and a data point p within $W(q_b, z)$, the safe region with regard to p is:

$$S(q_b, z, p) = \{q' : \text{dist}(q', p) + \text{dist}(q_b, q') \leq \text{dist}(q_b, z)\}.$$

For Euclidean distance in a 2D space, the boundary of $S(q_b, z, p)$ is an ellipse as illustrated in Figure 5. The two foci of the ellipse are q_b and p ; the sum of the distances from q_b and p to any point on the ellipse is $\text{dist}(q_b, z)$. We further have the following results.

COROLLARY 1. In Euclidean space, the safe region with regard to z , $S(q_b, z, z)$, is the line segment $\overline{q_b z}$.

PROOF. For any q' in $S(q_b, z, z)$, $\text{dist}(q', z) + \text{dist}(q_b, q') \leq \text{dist}(q_b, z)$. By the triangular inequality, we have $\text{dist}(q', z) + \text{dist}(q_b, q') \geq \text{dist}(q_b, z)$. The set of points that satisfies both inequalities is $\{q' : \text{dist}(q', z) + \text{dist}(q_b, q') = \text{dist}(q_b, z)\}$. In Euclidean space, this is the line segment $\overline{q_b z}$. $\square$

COROLLARY 2. p is unreliable iff q' is outside of $S(q_b, z, p)$.

The proof is straightforward and we omit it.

4.2 Fixed-rank region

As discussed in Section 3.2, the fixed-rank region (FRR) was introduced in [13] to denote a set of possible query-point locations that share a specific ranking of all objects. For the V*-Diagram, the FRR is applied to a subset of all objects, specifically, the $(k+x)$ NN. The IRU algorithm is used to incrementally maintain the FRR.

Given a list L of objects, $\langle p_1, p_2, \dots, p_m \rangle$, the FRR of L is the set of all points such that for any point v in the set, $p_1, p_2, \dots, p_m$ are already in sorted (ascending) order by their distances to v . Let $H_{p_i p_j}$ be defined as $\{v \in DS : \text{dist}(v, p_i) \leq \text{dist}(v, p_j)\}$, where DS is the data space. An FRR is a function of a list and is formulated as follows.

DEFINITION 2 (FIXED-RANK REGION).

$$F\langle p_1, p_2, \dots, p_m \rangle = \bigcap_{i=1}^{m-1} H_{p_i p_{i+1}}.$$

³We omit “with regard to q ” when the context is clear.

$F\langle p_1, p_2, \dots, p_m \rangle$ may be written in the compact format of $F(L)$.

In Figure 4(a), the ranking of the objects according to their distances to q_1 is $L = \langle a, c, b, f, e, d \rangle$. The FRR $F(L)$ is defined as $H_{ac} \cap H_{cb} \cap H_{bf} \cap H_{fe} \cap H_{ed}$. The boundary of $F(L)$ is defined by five bisectors, B_{ac} , B_{cb} , B_{bf} , B_{fe} and B_{ed} .

We use the IRU algorithm described in Section 3.2 to incrementally compute $F(L)$ in which q currently resides for the $(k+x)$ maintained objects. An FRR is represented as (i) a list B of the $(k+x-1)$ rank-adjacent bisectors, $B_{p_i p_{i+1}}$, for $i = 1, 2, \dots, k+x-1$, and (ii) a reference point (which could be any point in $F(L)$, q_b in our case). The FRR is incrementally maintained by (i) checking whether q crosses a bisector in B , and (ii) if yes, performing updates accordingly.

The purpose of maintaining the FRR using the IRU algorithm is to keep the $(k+x)$ objects sorted according to their distances to q . One alternative solution to IRU is to computing distances between the $(k+x)$ objects and q for every position of q , which is sampling-based. Although both methods have the same complexity, the FRR is important to the formulation of a region where the (order-sensitive) k NN does not change.

4.3 Integrated safe region

We are now ready to formulate the safe region for the MkNN query, called the *integrated safe region* (ISR). The ISR is the intersection of the current FRR of the $(k+x)$ maintained objects and the safe regions with regard to the k nearest objects. We will first define the ISR formally and then prove that ISR satisfies the MkNN safe-region requirements.

Let O denote the $(k+x)$ NN set of q_b , L be the list of these $(k+x)$ objects sorted by their distances to q , and z still be the farthest retrieved object to q_b , which is p_{k+x} . The ISR is then formulated as:

$$F(L) \cap \left(\bigcap_{i=1}^k S(q_b, z, p_i) \right) \quad (2)$$

The computation of the ISR can be greatly reduced based on Lemma 1 and Theorem 3 below.

LEMMA 1. $F\langle p_i, p_j \rangle \cap S(q_b, z, p_j) \cap S(q_b, z, p_i) = F\langle p_i, p_j \rangle \cap S(q_b, z, p_j)$.

PROOF. For any point $v \in F\langle p_i, p_j \rangle$, v satisfies

$$\text{dist}(v, p_i) \leq \text{dist}(v, p_j). \quad (3)$$

For any point $v \in S(q_b, z, p_j)$, by definition v satisfies

$$\text{dist}(v, p_j) + \text{dist}(q_b, v) \leq \text{dist}(q_b, z). \quad (4)$$

For any $v \in F\langle p_i, p_j \rangle \cap S(q_b, z, p_i)$, it satisfies inequalities (3) and (4). By adding the two inequalities,

$$\text{dist}(v, p_i) + \text{dist}(q_b, v) \leq \text{dist}(q_b, z). \quad (5)$$

Inequality (5) shows that $v \in S(q_b, z, p_i)$, given that $v \in F\langle p_i, p_j \rangle \cap S(q_b, z, p_j)$. It can be concluded that: $F\langle p_i, p_j \rangle \cap S(q_b, z, p_j) \subseteq S(q_b, z, p_i)$ and hence $S(q_b, z, p_i)$ can be discarded in $F\langle p_i, p_j \rangle \cap S(q_b, z, p_j) \cap S(q_b, z, p_i)$. $\square$

An example is given in Figure 6. The grey region $F\langle a, c, b, f \rangle \cap S(q_1, f, c) \cap S(q_1, f, a)$ is exactly the same as $F\langle a, c, b, f \rangle \cap S(q_1, f, c)$.

THEOREM 3. For $k \geq 2$,

$$F\langle p_1, p_2, \dots, p_k \rangle \cap \left(\bigcap_{i=1}^k S(q_b, z, p_i) \right) = F\langle p_1, p_2, \dots, p_k \rangle \cap S(q_b, z, p_k)$$

PROOF. Lemma 1 shows the case of $k = 2$, that is, $F\langle p_1, p_2 \rangle \cap S(q_b, z, p_2) \cap S(q_b, z, p_1) = F\langle p_1, p_2 \rangle \cap S(q_b, z, p_2)$. If the theorem holds for $k = l$, that is,

$$F\langle p_1, p_2, \dots, p_l \rangle \cap \left(\bigcap_{i=1}^l S(q_b, z, p_i) \right) = F\langle p_1, p_2, \dots, p_l \rangle \cap S(q_b, z, p_l),$$

then the theorem can be verified for $k = l + 1$ as follows.

$$\begin{aligned} & F\langle p_1, p_2, \dots, p_{l+1} \rangle \cap \left(\bigcap_{i=1}^{l+1} S(q_b, z, p_i) \right) \\ &= F\langle p_1, p_2, \dots, p_l \rangle \cap F\langle p_l, p_{l+1} \rangle \\ & \quad \cap \left(\bigcap_{i=1}^l S(q_b, z, p_i) \right) \cap S(q_b, z, p_{l+1}) \\ &= F\langle p_1, p_2, \dots, p_l \rangle \cap F\langle p_l, p_{l+1} \rangle \cap S(q_b, z, p_l) \\ & \quad \cap S(q_b, z, p_{l+1}). \end{aligned}$$

By applying Lemma 1, $S(q_b, z, p_l)$ can be removed from the above expression. Hence, we obtain a final result of

$$\begin{aligned} & F\langle p_1, p_2, \dots, p_l \rangle \cap F\langle p_l, p_{l+1} \rangle \cap S(q_b, z, p_{l+1}) \\ &= F\langle p_1, p_2, \dots, p_{l+1} \rangle \cap S(q_b, z, p_{l+1}). \end{aligned}$$

The theorem therefore holds for any integer value of k greater than or equal to 2. $\square$

Since $F(L) = F\langle p_1, p_2, \dots, p_k \rangle \cap F\langle p_k, p_{k+1}, \dots, p_{k+x} \rangle$, based on Theorem 3, expression (2) can be reduced to $F(L) \cap S(q_b, z, p_k)$. Therefore, the ISR can be defined as follows.

DEFINITION 3 (INTEGRATED SAFE REGION (ISR)). Let O be the $(k+x)$ NN set of q_b , L be the list of these $(k+x)$ objects sorted by their distances to q , z be the farthest retrieved object to q_b , and p_k be the k^{th} object in L . The integrated safe region with respect to q_b , z , p_k and L is defined as:

$$I(q_b, z, p_k, L) = F(L) \cap S(q_b, z, p_k) \quad (6)$$

Next, we prove that the ISR defined above satisfies the requirements of being a safe region for the Mk NN query, that is, the k NNs as well as their order do not change when q remains in the ISR.

THEOREM 4. If the ISR $I(q_b, z, p_k, L)$ is not an empty set, every point q' in $I(q_b, z, p_k, L)$ has the same order-sensitive k NNs.

PROOF. According to Definition 3, (1) since $I \subseteq F(L)$ (parameters of I omitted), the ranking of the $(k+x)$ objects is fixed for all points in I , which satisfies the order-sensitivity requirement; (2) every point q' in I is also in the safe regions with regard to the first k objects in L . As a result, there can be no object outside $W(q_b, z)$ nearer to q' than any of the first k objects in L . Therefore, for any q' in I , q' has the same order-sensitive k NNs as q_b . $\square$

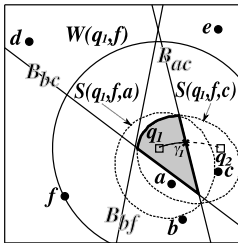


Figure 6: Integrated safe region example ($k = 2, x = 2$)

As exemplified in Figure 6, four objects retrieved by a 4NN query ($k = 2$ and $x = 2$) at q_1 are $\{a, c, b, f\}$. Point q_1 is the most

recent point where $(k+x)$ NNs are retrieved from the database. As long as q' remains in $F\langle a, c, b, f \rangle \cap S(q_1, f, c)$ (the grey region), then: (i) no object outside $W(q_1, f)$ is nearer to the two objects: a and c ; and (ii) the ranking of $\langle a, c, b, f \rangle$ is unchanged.

The diagram that contains the information used in computing the ISR is called the **V*-Diagram**. It consists of: (i) the bisectors of the rank-adjacent pairs in L ; and (ii) the boundary of the safe region with regard to p_k . We may also use the V*-Diagram to refer generally to the whole technique based on it, including the algorithms. Although the V*-Diagram in the example of Figure 6 only computes a single ISR, it actually allows incremental computation of new ISRs, which is further discussed in Section 5.

5. ALGORITHMS

In this section, we present V*- k NN, an algorithm for Mk NN queries based on the V*-Diagram, followed by a discussion on the effect of x , the number of auxiliary objects. We also present the algorithms to handle insertions/deletions and dynamically changing k values.

The V*- k NN algorithm uses the following data structures and variables to compute and maintain the ISR.

1. L : a list of $(k+x)$ objects always sorted in ascending order by their distances to q ; these objects are the $(k+x)$ NNs retrieved at q_b .
2. z : the farthest retrieved object in the known region when q is at q_b .
3. p_k : the k^{th} object in L .
4. S_k : the safe region with regard to p_k .
5. B : a list of bisectors of pairs of rank-adjacent objects (*rank-adjacent bisectors*) in the order corresponding to L .

We do not explicitly maintain $F(L)$ because it is represented by B , and checking whether q moves out of the current $F(L)$ is also done by checking whether q crosses any bisector in B .

V*- k NN produces answers continuously as shown in Algorithm 1. It has an initialization part (lines 1 to 3) and a continuous processing part (lines 4 to 19). The initialization part calls the algo-

Algorithm 1: V*- k NN(q_0, k, x)

```

1  $q_b \leftarrow q_0$ 
2  $(L, z, S_k, B, ISR) \leftarrow \text{Compute-V}^*(q_b, k, x)$ 
3 ReportResult( $L.\text{Head}(k)$ )
4 while ( $\text{Event} \leftarrow \text{GetEvent}()$ ) do
5    $q \leftarrow \text{Event.Position}$ 
6   switch Event.Type do
7     case RankUpdate
8        $\text{Bisector} \leftarrow \text{Event.Bisector}$ 
9        $L.\text{OrderSwap}(\text{Bisector.Index})$ 
10       $B.\text{Update}(L, \text{Bisector.Index})$ 
11      if  $\text{Bisector.Index} \leq k$  then
12        ReportResult( $L.\text{Head}(k)$ )
13      if  $\text{Bisector.Index} \in [k-1, k]$  then
14         $p_k \leftarrow L.\text{Item}(k)$ 
15         $S_k \leftarrow S(q_b, z, p_k)$ 
16       $ISR \leftarrow \text{ConstructISR}(S_k, B, q)$ 
7     case ReliabilityUpdate
18        $q_b \leftarrow q$ 
19        $(L, z, S_k, B, ISR) \leftarrow \text{Compute-V}^*(q_b, k, x)$ 
```

rithm *Compute-V** (Algorithm 2) to compute the initial ISR using the starting point q_0 of the trajectory as q_b . Then the continuous processing part starts.

Algorithm Compute-V* (Algorithm 2) runs as follows. It first calls the BF- k NN algorithm to retrieve $(k + x)$ objects with q_b as the query point; with the retrieved objects, it sets z and p_k accordingly; then, it computes $S(q_b, z, p_k)$ and rank-adjacent bisectors based on L and assigns them to S_k and B , respectively; finally, the current ISR is computed. To determine the correct half plane for each bisector in B , q_b is used as the reference point. Readers may notice that symbol z is explained differently here from Table 2. Object z is used to determine the known region and it is the $(k + x)^{th}$ NN of q_b when q is at q_b . When deletion is taken into account, if the $(k + x)^{th}$ NN of q_b gets deleted, we still use the deleted object to determine the known region. Therefore, we make it more accurate here than in Table 2 when we did not consider deletions.

The continuous processing part of V*- k NN is event driven. It basically maintains the ISR as q moves. An event is triggered when q exits the current ISR. There are two types of events with this regard, *RankUpdate* and *ReliabilityUpdate*. These events are generated by a separate inexpensive process that constantly checks the current position of q against the ISR. When an event is generated, it is associated with a timestamp and the corresponding query position.

Algorithm 2: Compute-V*(q_b, k, x)

```

1  $L \leftarrow \text{BF-}k\text{NN}(q_b, k + x)$ 
2  $z \leftarrow L.\text{Item}(k + x)$ 
3  $p_k \leftarrow L.\text{Item}(k)$ 
4  $S_k \leftarrow S(q_b, z, p_k)$ 
5  $B \leftarrow \text{CreateBisectorList}(L)$ 
6  $\text{ISR} \leftarrow \text{ConstructISR}(S_k, B, q_b)$ 
7 return ( $L, z, S_k, B, \text{ISR}$ )
```

Given that the query trajectory is unknown, the query positions are updated discretely and checking for new events has to be done based on these discrete updates. To provide accurate answers, the checking should be performed at a high frequency, which is acceptable because of the low cost. Note that this is different from *processing the query* based on sampling, which requires frequent tree searches instead of event checking. The answer we provide is continuous, not based on sampled locations. Since the query positions are updated discretely, the events could happen anytime between two consecutive query updates. To compute the exact time (position) the events happen, we assume a linear trajectory between two consecutive query positions.

We describe the two event types, *RankUpdate* and *ReliabilityUpdate* below and discuss how to handle them with reference to Algorithm 1 (in the algorithm, q is used to denote the position of the query point when an event happens).

RankUpdate This event is triggered when q exits the current $F(L)$, that is, crossing a rank-adjacent bisector. Besides the timestamp and query position, a *RankUpdate* event also contains the information of the bisector crossed by q (line 8). For this event, the ranks of the two objects corresponding to the bisector are swapped (line 9) and the bisector list B is updated accordingly (line 10) as explained in the IRU algorithm (Section 3.2). If the event affects the rank of any of the k NNs (line 11), then the new k NNs are reported. Moreover, if the rank update changes p_k (line 13), S_k also needs to be updated (lines 14 to 15). Changes are reported to the user if at least one of the k nearest objects is affected.

ReliabilityUpdate This event is triggered when q is leaving S_k (that is, on the boundary of S_k). It means that the k^{th} NN is about to become unreliable and hence the number of reliable objects is about to become less than k . Therefore, the *ReliabilityUpdate* event calls Compute-V* using the event position, q , to obtain x new auxiliary objects so that all the

$(k + x)$ maintained objects are reliable again. The new ISR is constructed accordingly. This event does not cause result update because neither the k NN set nor their ordering changes. If another object really becomes nearer than the k^{th} NN, a *RankUpdate* will be triggered, which will update the result.

Next, we give a running example of the algorithm. Recall the example in Figure 6. At the starting point q_1 , 4 NNs are retrieved in the order of $\langle a, c, b, f \rangle$. The ISR is $F\langle a, c, b, f \rangle \cap S(q_1, f, c)$. On the trajectory of q (starting at q_1), two events happen when q

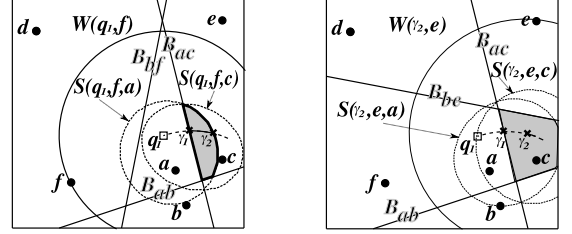


Figure 7: Example for Algorithm 1, ($k = 2, x = 2$)

crosses B_{ac} at γ_1 and exits $S(q_1, f, a)$ at γ_2 . Figure 7 shows the effects of these two events.

Figure 7(a) shows how the ISR changes after q crosses B_{ac} . At the instant that q is crossing B_{ac} at γ_1 , a *RankUpdate* event is triggered, which causes a and c to swap their ranks. The list L becomes $\langle c, a, b, f \rangle$, and this causes both $F(L)$ and S_k to change. Now a becomes p_k (2^{nd} NN), and hence the ISR becomes $F\langle c, a, b, f \rangle \cap S(q_1, f, a)$ (the grey region). The current 2 NNs, c and a , are reported to the user in that order.

Figure 7(b) shows how the ISR changes after q exits $S(q_1, f, a)$. At the instant that q is exiting $S(q_1, f, a)$ at γ_2 , a *ReliabilityUpdate* event is triggered, which calls Compute-V* to retrieve more objects. The new $(k + x)$ NNs are $\langle c, a, b, e \rangle$, with the corresponding ISR, $F\langle c, a, b, e \rangle \cap S(\gamma_2, e, a)$ (the grey region).

5.1 On the number of auxiliary objects

Auxiliary objects are an important part of the V*-Diagram technique. They allow q to move away from q_b while retaining the current k NNs by providing the knowledge beyond the coverage of the search sphere of the original k NNs. This makes it possible to continuously evaluate the MkNN query. In this subsection, we discuss possible values of x , the number of auxiliary objects. Generally, we find that x should not assume the values of 0 and 1, which are explained below.

Having x equal to 0 implies that z and the k^{th} object in L , p_k , are the same object. According to Corollary 1, in Euclidean space, the safe region with regard to z (which is S_k) is the line segment $\overline{q_b z}$. Unless q moves along $\overline{q_b z}$, q exits S_k as it starts moving. Probabilistically, it is highly unlikely that q moves along $\overline{q_b z}$ since $\overline{q_b z}$ is just one direction among the infinite possible directions q can move towards. Therefore, it is probable that q always exits S_k as it moves and it triggers the *ReliabilityUpdate* event infinitely. As a result, x should not be set to 0 for Euclidean space.

When x is a positive integer, the problem of infinitely triggering the *ReliabilityUpdate* event does not happen except under the coincidence described in the next paragraph. Therefore, any integer greater than 1 is a valid value for x . The effect of the value of x on performance is further investigated in Sections 7 and 8.

In theory, the problem of S_k being a line segment may happen with any value of x when the last $(x + 1)$ objects in L have the same distance to q_b , which is $\text{dist}(q_b, z)$. In this case, $\text{dist}(q_b, p_k)$ equals to $\text{dist}(q_b, z)$. By definition, $S_k = \{q' : \text{dist}(q', p_k) + \text{dist}(q_b, q') \leq \text{dist}(q_b, z)\}$. If we replace $\text{dist}(q_b, z)$ by $\text{dist}(q_b, p_k)$ in the inequality of the definition, we

get $S_k = \{q' : \text{dist}(q', p_k) + \text{dist}(q_b, q') \leq \text{dist}(q_b, p_k)\}$, which is a line segment in Euclidean space. To completely avoid this problem, we can check whether $\text{dist}(q_b, p_k)$ is equal to $\text{dist}(q_b, z)$ after we retrieve $(k + x)$ objects by a BF- k NN call. If they are equal, then we increase the value of x until $\text{dist}(q_b, p_k)$ is different from $\text{dist}(q_b, z)$.

In general, a larger value of x provides a larger S_k , and hence the less frequent we need to retrieve new objects from the database. Having the value of x too small will result in highly frequent BF- k NN calls. On the other hand, a too large x value also incurs the overhead of retrieving more objects in every BF- k NN call and more computation for maintaining them.

5.2 Insertions and deletions of objects

Algorithm 3: DatasetUpdate($q, p, \text{Operation}$)

```

1 if  $p \in W(q_b, z)$  then
2   if  $\text{Operation} = \text{Insertion}$  then
3      $L \leftarrow \text{Insert}(L, p, q)$ 
4   else
5      $L \leftarrow \text{Delete}(L, p)$ 
6    $B.\text{Update}(L)$ 
7   if  $k > L.\text{Length}()$  then
8      $q_b \leftarrow q$ 
9      $(L, z, S_k, B, \text{ISR}) \leftarrow \text{Compute-V}^*(q_b, k, x)$ 
10     $\text{ReportResult}(L.\text{Head}(k))$ 
11   else if  $\text{dist}(q, p) \leq \text{dist}(q, p_k)$  then
12      $p_k \leftarrow L.\text{Item}(k)$ 
13      $S_k \leftarrow S(q_b, z, p_k)$ 
14     if  $q \notin S_k$  then
15        $q_b \leftarrow q$ 
16        $(L, z, S_k, B, \text{ISR}) \leftarrow \text{Compute-V}^*(q_b, k, x)$ 
17     else
18        $\text{ISR} \leftarrow \text{ConstructISR}(S_k, B, q)$ 
19        $\text{ReportResult}(L.\text{Head}(k))$ 
20   else
21      $\text{ISR} \leftarrow \text{ConstructISR}(S_k, B, q)$ 

```

In this subsection, we describe the algorithm to perform updates (that is, insertions and deletions) to the dataset for V*- k NN. The algorithm is called *DatasetUpdate* and is presented in Algorithm 3. In this algorithm, q denotes the position of the query point when the update happens and it is passed in as an input. Let p be the object to be inserted or deleted. First, the algorithm checks whether the p is in $W(q_b, z)$. If not, the update can be safely ignored because it cannot affect the ISR. Otherwise, an insertion/deletion of p into/from L is performed and B is updated accordingly (lines 2 to 6): insertion of p needs q for computing distances from q and the maintained objects to find the correct insertion slot in L ; deleting p from L requires only a simple lookup operation. After the bisector update, the ISR and L could be in one of the following three cases:

- (i) *The length of L becomes smaller than k as a result of a deletion (line 7).* In this case, q_b is set to q and Compute-V* is called to retrieve more objects and compute the new ISR accordingly (lines 8 and 9). The new result is reported (line 10).
- (ii) *The length of L is still greater than k but the update affects the k NN set (line 11).* We update p_k and S_k (lines 12-13) and check if q is inside the new S_k (line 14). If q is not inside the new S_k , then Compute-V* is called (line 16). Otherwise, the ISR is updated to reflect the changes in B and S_k . Since the k NNs have changed, the new result is reported to the user (line

19).

- (iii) *The update has no effects to the k NN set (line 20).* Only the ISR is updated to reflect the change in B .

5.3 M k NN with dynamically changing k values

The ability to gracefully handle changes to the value of k is crucial for the *distance browsing* functionality [6]. For static k NN queries, distance browsing is a feature that allows NNs to be incrementally retrieved without having to specify the value of k in advance. In this paper, we allow the value of k to be changed without incurring heavy computation.

Algorithm *KUpdate* (Algorithm 4) shows how V*- k NN handles dynamically changing k values. The algorithm has two inputs: the current location q of the query point and the new k value. We first check if the new k is greater than the length of L . If yes, q_b is set to q and Compute-V* is called. Otherwise, p_k and S_k are updated for the new k (lines 5 and 6). If q is not inside the new S_k , q_b is set to q and Compute-V* is called. Otherwise, only the ISR has to be updated to incorporate the new S_k (line 11). Finally, the new k NNs are reported (line 12). As we can see, the V*- k NN algorithm can easily accommodate dynamically changing k values due to its incremental nature.

Algorithm 4: KUpdate(q, k)

```

1 if  $k > L.\text{Length}()$  then
2    $q_b \leftarrow q$ 
3    $(L, z, S_k, B, \text{ISR}) \leftarrow \text{Compute-V}^*(q_b, k, x)$ 
4 else
5    $p_k \leftarrow L.\text{Item}(k)$ 
6    $S_k \leftarrow S(q, z, p_k)$ 
7   if  $q \notin S_k$  then
8      $q_b \leftarrow q$ 
9      $(L, z, S_k, B, \text{ISR}) \leftarrow \text{Compute-V}^*(q_b, k, x)$ 
10  else
11     $\text{ISR} \leftarrow \text{ConstructISR}(S_k, B, q)$ 
12   $\text{ReportResult}(L.\text{Head}(k))$ 

```

6. THE V*-DIAGRAM IN SPATIAL NETWORKS

When the movement of q is constrained by network connectivity, NN problems should be solved based on the network distance. For example, a car is travelling on a road and it keeps track of the k nearest gas stations based on the road network distance.

In this section, we show how the V*-Diagram technique can be applied to the domain of spatial networks, which also satisfies our metric space assumption. Due to the space limitation, we could not fully elaborate the application of V*-Diagram in the spatial-network model, but only present the essential components of the technique.

The essence of the V*-Diagram is the ISR, which consists of two key components: safe regions with regard to data objects and the fixed-rank region (FRR). Hence, we focus our discussions on how to determine these two components in spatial networks, while the algorithms to compute the ISR and process M k NN queries can be reused. To avoid an exhaustive discussion, we use an example to illustrate the main idea.

A spatial network is usually represented as a set of vertices and a set of edges, where an edge is defined by two vertices. Given points p_1 and p_2 in the network, $\text{dist}(p_1, p_2)$ is the length of the shortest path between p_1 and p_2 . Figure 8 shows a spatial network of eight nodes, a to h . Distance $\text{dist}(q_b, p)$ is 2 using the path via b .

We use $pnt(p_1, p_2, l)$ to denote a point on the edge (p_1, p_2) with distance l to p_1 , and $seg(p_1, p_2, l)$ to denote a section on the edge (p_1, p_2) with the starting point p_1 and length l . For example, z is at $pnt(a, e, 3)$, and all points between a and z can be represented as $seg(a, e, 3)$.

Consider an $MkNN$ query with $k = 1$. Suppose x is 2 for the V^* -Diagram and q_b is at $pnt(a, b, 3)$. The 3 NNs for q_b are: p , r and z . We can use any kNN technique for spatial networks (such as those described in [8, 9, 15, 17]) to retrieve the $(k + x)$ NNs at q_b . The algorithm also returns the edges in the range of $dist(q_b, z)$ to q_b (that is, $W(q_b, z)$), (a, b) , (a, c) , (a, e) , (b, d) , (b, f) , (d, c) , (d, g) , and the distances from the vertices of these edges to q_b .

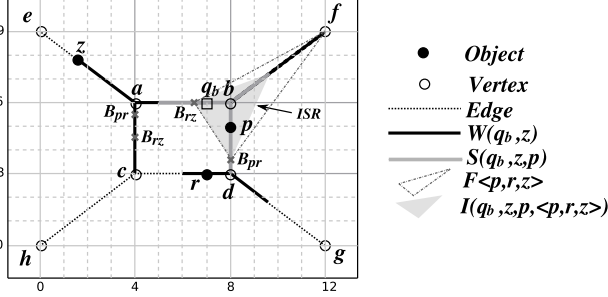


Figure 8: V^* -Diagram in a spatial network ($k = 1$ and $x = 2$)

Next, we determine the safe region for p_k (which is p in this example). The safe region $S(q_b, z, p)$ is the region where for any point q' in the region, the sum of $dist(q_b, q')$ and $dist(p, q')$ is less than or equal to $dist(q_b, z)$, which is 6. We need to explore all the edges in $S(q_b, z, p)$ to identify its boundary. Since $W(q_b, z)$ encloses $S(q_b, z, p)$, and we already know the edges that are in $W(q_b, z)$ via the kNN query executed at q_b , we only need to consider those edges. We use edge (b, f) as an example. For a point q' on (b, f) to be in $S(q_b, z, p)$, $dist(q_b, q') + dist(p, q')$ has to be less than 6 according to the definition. Since $dist(q_b, q') + dist(p, q')$ is equal to $dist(q_b, b) + dist(p, b) + 2dist(b, q')$, and both $dist(q_b, b)$ and $dist(p, b)$ are 1, q' is in $S(q_b, z, p)$ when $dist(b, q')$ is less than or equal to 2. Consequently, $pnt(b, f, 2)$ forms part of the boundary of $S(q_b, z, p)$. The boundary points on other edges can be computed in a similar manner. After obtaining all boundary points, we get $S(q_b, z, p)$, which consists of edge (b, d) , $seg(b, a, 3)$ and $seg(b, f, 2)$, plotted as grey thick segments in the figure.

As discussed in Section 4.2, the FRR is determined by the rank-adjacent bisectors the maintained objects. In a spatial network, a bisector reduces to points on edges. For example, B_{pr} has two bisecting points, one on $pnt(b, d, 2.5)$ and the other on $pnt(a, c, 0.5)$. The FRR $F(p, r, z)$ is the region with B_{pr} and B_{rz} as the boundary and containing q_b . It consists of (b, f) , $seg(b, a, 1.5)$ and $seg(b, d, 2.5)$, shown as the segments in the dashed triangle.

As a result, the ISR ($S(q_b, z, p) \cap F(p, r, z)$) consists of $seg(b, a, 1.5)$, $seg(b, d, 2.5)$ and $seg(b, f, 2)$. It is shown as the segments in the grey region. According to Algorithm 1, exiting the ISR via $pnt(b, a, 1.5)$ (B_{rz}) or $seg(b, d, 2.5)$ (B_{pr}) triggers a *RankUpdate* event, and exiting the ISR via $pnt(b, f, 2)$ triggers a *ReliabilityUpdate* event.

7. PERFORMANCE ANALYSIS

Among all techniques listed in Table 1, only RIS- kNN by Zhang et al. [23] and V^* - kNN provide continuous answers for the $MkNN$ query with unknown query trajectory and without accessing all the data. Therefore, this section focuses on a comparative performance analysis on RIS- kNN and V^* - kNN in terms of tree node accesses (that is, IO cost). We assume a 2D space, although the analysis can

be extended to higher dimensional spaces. We also assume that the data objects are uniformly distributed.

V^* - kNN only has node accesses in the calls to BF- kNN , which are triggered by the *ReliabilityUpdate* event (line 19 of Algorithm 1). We analyze the frequency of *ReliabilityUpdate* events, f_b , as follows. A *ReliabilityUpdate* event happens when q exits $S(q_b, z, p_k)$. We denote the point of exit as q_e . The *ReliabilityUpdate* event happens only once during the process of q moving away from q_b until reaching q_e . Then a BF- kNN is performed and q_e becomes the new q_b . A *ReliabilityUpdate* event will happen again when the next time q exits $S(q_b, z, p_k)$. Therefore, f_b is inversely proportional to the distance q travels from q_b to q_e . In the worst case, q moves in a straight line, and f_b is inversely proportional to $dist(q_b, q_e)$. The expected value of $dist(q_b, q_e)$ is obtained as follows. When q moves to q_e , p_k is on the boundary of the reliable region (with regard to q_e). For a better understanding, imagine in Figure 5, p is p_k , q' is q_e and it is on the boundary of $S(q_b, z, p)$. We can see that $dist(q_b, q_e)$ equals $dist(q_b, \chi) - dist(q_e, \chi)$. Note that $dist(q_b, \chi)$ is the radius of $W(q, z)$, which is a sphere that contains $(k + x)$ points; $dist(q_e, \chi)$ is the radius of the reliable region with regard to q_e , which is a sphere that contains k points.

According to [21], the distance between the query point q_b and the k^{th} NN is $2/C_v(1 - \sqrt{1 - \sqrt{k/n}})$, where C_v is the vicinity constant [21]. Therefore $dist(q_b, q_e)$, which is $dist(q_b, \chi) - dist(q_e, \chi)$, can be expressed as

$$\frac{2}{C_v}(\sqrt{1 - \sqrt{k/n}} - \sqrt{1 - \sqrt{(k+x)/n}}).$$

Since f_b is inversely proportional to $dist(q_b, q_e)$,

$$\mathcal{O}(f_b) = \mathcal{O}(1/(\sqrt{1 - \sqrt{k/n}} - \sqrt{1 - \sqrt{(k+x)/n}})).$$

The expression of $\mathcal{O}(f_b)$ can be relaxed as follows:

$$\begin{aligned} \frac{1}{\sqrt{1 - \sqrt{k/n}} - \sqrt{1 - \sqrt{(k+x)/n}}} &= \frac{\sqrt{1 - \sqrt{k/n}} + \sqrt{1 - \sqrt{(k+x)/n}}}{\sqrt{(k+x)/n} - \sqrt{k/n}} \\ &\leq \frac{2\sqrt{1 - \sqrt{k/n}}}{\sqrt{(k+x)/n} - \sqrt{k/n}} \leq \frac{2}{\sqrt{(k+x)/n} - \sqrt{k/n}} \\ &= 2\frac{\sqrt{(k+x)/n} + \sqrt{k/n}}{(k+x)/n - k/n} = 2\frac{\sqrt{(k+x)/n} + \sqrt{k/n}}{x/n} \leq 4\frac{\sqrt{(k+x)/n}}{x/n}. \end{aligned}$$

Therefore $\mathcal{O}(f_b)$ is $\mathcal{O}(\sqrt{\frac{(k+x)n}{x^2}})$. Typically, x is comparable to k and $(k + x)$ is much smaller than xk . Thus, we obtain that f_b is $\mathcal{O}(\sqrt{\frac{kn}{x}})$. Let C_{nn} be the cost of a BF- kNN call. Then we obtain

the total IO cost of V^* - kNN , $C_{nn}\mathcal{O}(\sqrt{\frac{kn}{x}})$.

The RIS- kNN processes the $MkNN$ query as follows. Every time q exits the current kVD cell, RIS- kNN is executed to obtain the new kVD cell and the corresponding kNN . The total cost depends on the frequency of crossing kVD cells and the cost of each RIS- kNN run. In the worst case, q moves along a straight line. The frequency of q crossing kVD cells is proportional to the average linear density⁴ of the kVD cells. The number of the kVD cells is $\mathcal{O}(kn - k)$ in 2D space [14]. We assume that k is much smaller than n . Thus, the density of kVD cells is $\mathcal{O}(kn)$, which corresponds to a linear density of $\mathcal{O}(\sqrt{kn})$. Each RIS- kNN run requires 12 TPkNN queries on average [23]. Let C_{tpnn} be the cost of a TPkNN call. Then the total IO cost of RIS- kNN for the $MkNN$ query is $12C_{tpnn}\mathcal{O}(\sqrt{kn})$.

⁴The number of kVD cells crossed per unit length along a straight line.

We now compare the costs of V^*-kNN and $RIS-kNN$. For the frequency component, V^*-kNN is smaller than $RIS-kNN$ especially when x is large. For the cost-per- kNN -algorithm-call component, the first $TPkNN$ call of an $RIS-kNN$ run is always more expensive than a $BF-kNN$ call because $TPkNN$ has to retrieve at least the kNN and it needs to access more nodes to obtain the influence object. In addition, there are 11 subsequent $TPkNN$ calls of an $RIS-kNN$ run, where each call is much more expensive than a $BF-kNN$ call in practice. Consequently, the IO cost of V^*-kNN is much lower than that of $RIS-kNN$.

In terms of the CPU cost, V^*-kNN maintains the rank of $(k + x)$ objects, which causes more computation than $RIS-kNN$ on the client side. However, given today's mobile devices and trends (e.g., phones with multimedia and graphical functionalities), we argue that the CPU power of these devices is adequate to check $(k + x - 1)$ bisectors at a reasonably high frequency (e.g., once every second) for practical values of k and x . In many realistic settings, the benefits from the communication-cost reduction will outweigh this overhead on the client side.

On the server side, V^*-kNN has a lower CPU cost than $RIS-kNN$. It is because $RIS-kNN$ uses the $TP-kNN$ query which incurs a much higher computational cost than $BF-kNN$.

8. EXPERIMENTAL STUDY

This section presents the results of our experimental study. In our implementation, the R^* -tree [1] is used to index the data objects. The page size is 1 KB, which has a node capacity of 50 entries.

We used both synthetic and real datasets in our experiments. All datasets span the space of $10,000 \times 10,000$ square units. We generated synthetic datasets with uniform (U) and Zipfian (Z) distributions with the default cardinality of 25,000 data points. The real datasets are 65,743 and 119,897 postal addresses from California (C) and North-Eastern USA (N), respectively.

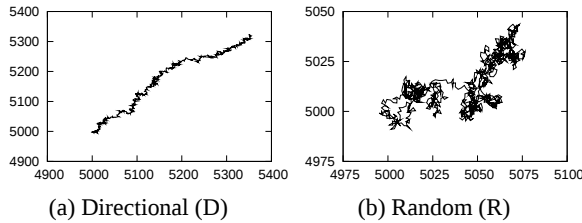


Figure 9: Trajectory types

We generated two different types of query trajectories, random (R) and directional (D), as shown in Figure 9. Each trajectory consists of 1,001 points. Between two points, the trajectory is assumed to be a straight line segment. The length of each segment is 1 unit. For each type of trajectory, we generated 20 different trajectories. We run these 20 trajectories as a query set for each experiment and present the average result. We measured both the cumulative total response time and the cumulative number of page accesses for a whole trajectory as the performance metric.

8.1 Choosing the value of x

In the first set of experiments, we study the impact of the value of x on the performance of V^*-kNN . We varied x from 3 to 36 with k set to 20 and no buffer space. We did not use the values less than 3 for x because too small values of x do not yield a reasonable size of S_k to make V^*-kNN effective. Figure 10 shows result of the response time and number of page accesses as functions of x for both query trajectory types. The response time first decreases as x increases but becomes more constant as x keeps increasing. There is a small increase when x becomes large (around 30). The num-

ber of page accesses always decreases as x increases. This is because as x becomes larger, S_k becomes larger and hence $BF-kNN$ is called less frequently. This reduces both CPU time and IO cost. When x becomes too large, the computational overhead of maintaining more objects becomes more significant and may outweigh the savings in CPU time. Therefore, the CPU time increases for a large x value. The number of page accesses mainly depends on the frequency of $BF-kNN$ calls and hence always decreases. These results confirm our discussion in Section 5.1 and the analysis in 7. In all these experiment, the x value of 9 provides a good performance, so 9 is used as the default value of x for the rest of the experiments. As we can see, some variations around the value of 9 do not affect the performance much.

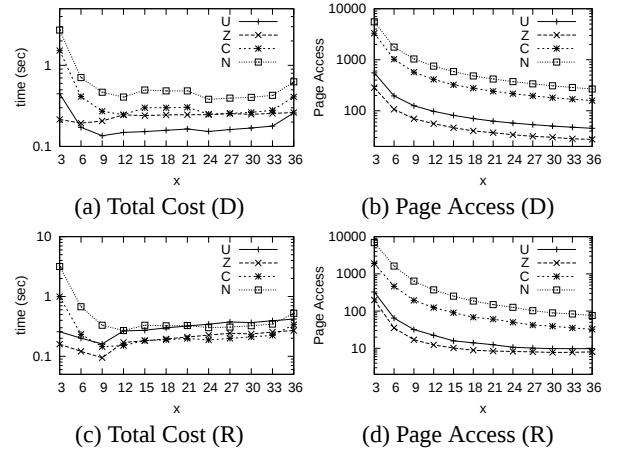


Figure 10: Effect of x

8.2 Comparative study: centralized

Among all the techniques discussed in the related work, $RIS-kNN$ by Zhang et al. [23] is the only algorithm that is comparable to our work. Therefore, we perform a comparative experimental study on V^*-kNN and $RIS-kNN$. The next four sets of experiments compare these two techniques using different experimental parameters.

The effect of the buffer size. In this set of experiments, we use the buffer sizes of 0, 8, 16, 24 and 32 pages. Figure 11 shows the results for two synthetic datasets with the default dataset size and the two real datasets. V^*-kNN outperforms the $RIS-kNN$ in all settings in terms of both total response time and number of page accesses. In most cases, the improvement factor is two orders of magnitude. For both methods, the page access cost decreases as the buffer size increases as expected.

The effect of the number of query location updates. In this set of experiments, we vary the number of location updates in the query trajectory from 0 to 1000. Figure 12 shows the response time on the four datasets. In all experiments, V^*-kNN outperforms $RIS-kNN$ and the improvement factor is two orders of magnitude in most cases. The results of the number of page accesses have very similar behavior as those of the total response time. Therefore we do not present them for the remaining experiments due to space limitation. For both techniques, the total response time increases as the number of location updates increases.

The effect of the dataset size. In this set of experiments, we vary the number of objects in the dataset from 5,000 to 25,000 for the synthetic datasets. Figure 13 shows the response time results. Again, V^*-kNN outperforms $RIS-kNN$ in all settings and the improvement factor is two orders of magnitude in most cases. For both techniques, the total response time increases as the number of objects increases.

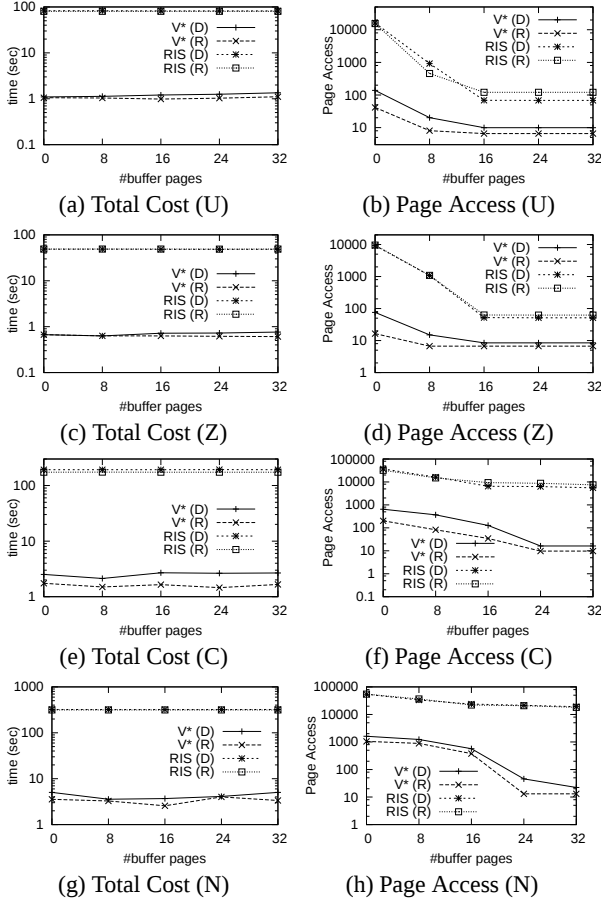


Figure 11: Effect of buffer size

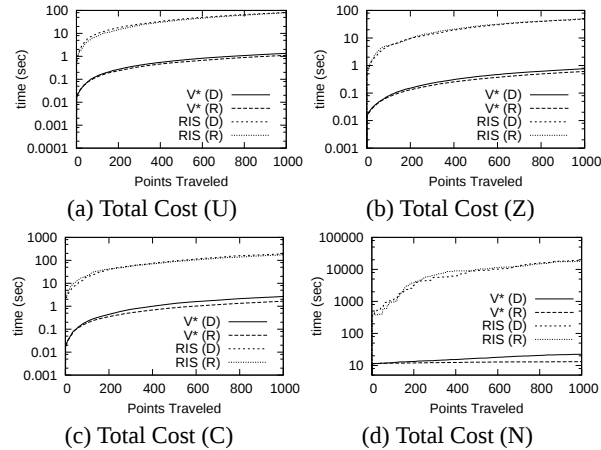


Figure 12: Effect of the number of location updates

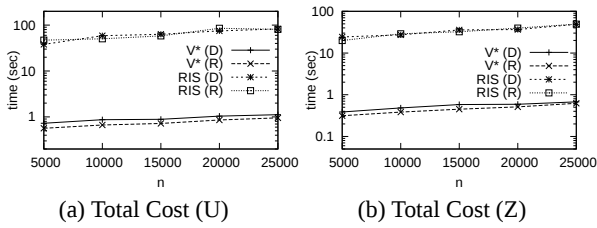


Figure 13: Effect of dataset size

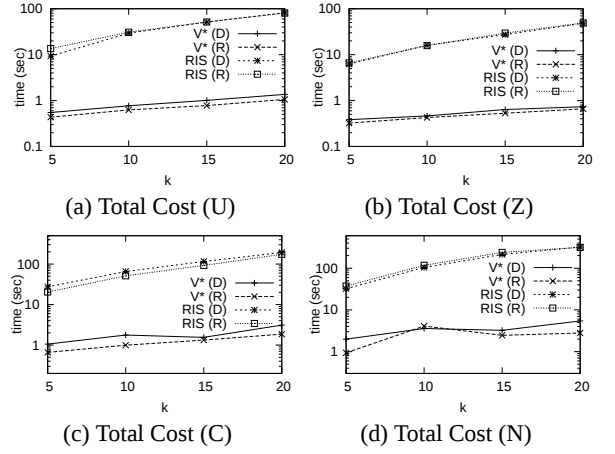


Figure 14: Effect of k

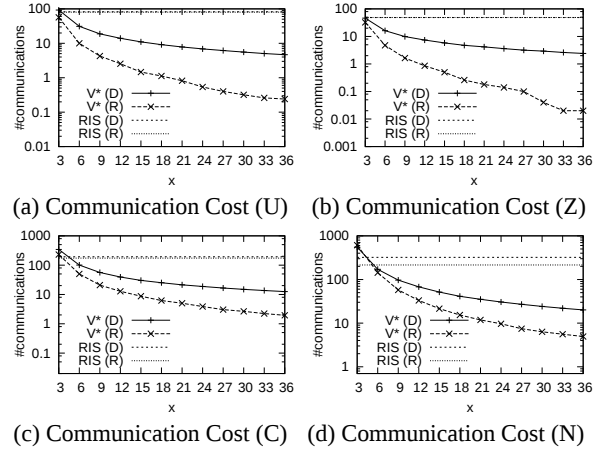


Figure 15: Communication cost: effect of x

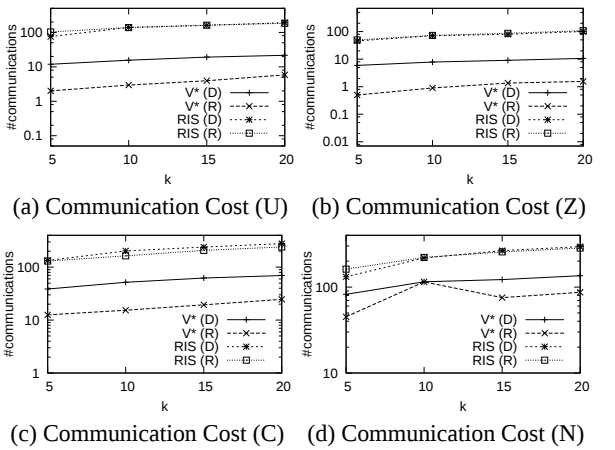


Figure 16: Communication cost: effect of k

The effect of k . In this set of experiments, we vary the value of k from 5 to 20 for the four datasets. Figure 14 shows the response time results. We observe similar results as in previous experiments. V^*-kNN outperforms $RIS-kNN$ in all settings and the improvement factor is two orders of magnitude in most cases. For both techniques, the total response time increases as the value of k increases, but the total response time of V^*-kNN increases slower than that of $RIS-kNN$.

8.3 Comparative study: client-server

In a high-latency client-server setting, communication costs between the mobile client and the server dominate the other costs. The following experiments compare the communication costs of V^*-kNN and $RIS-kNN$. The number of times a client has to communicate with the server is used as the performance measure.

For V^*-kNN , the communication cost is measured by the number of times the $BF-kNN$ query is executed because other operations are local. For $RIS-kNN$, the communication cost is the number of times the query itself is executed, i.e., the number of kVD cells crossed.

Figure 15 shows the communication costs with the increasing x values in the four datasets and two trajectory types. The default buffer size of 16 pages is used. Since the parameter x does not apply to $RIS-kNN$, its communication cost is constant as x changes. For all datasets, we can see that the communication cost of V^*-kNN decreases as the value of x increases. This conforms with the cost analysis in Section 7, which states that the retrieval cost of the V^*-kNN with respect to x is $x^{-1/2}$. The V^*-kNN starts to outperform the $RIS-kNN$ algorithm when x is: 3 for the uniform dataset (U), 3 for the Zipfian dataset (Z), 6 for the California dataset (C), and 6 for the North-Eastern USA dataset (N).

The difference between the communication and total response time costs (Figure 10) is notable. Unlike the total response time, the communication cost measure disregards the CPU cost. Only the database-access count is considered, and thus there is no penalty for large values of x .

The next experiment is a study on how the communication cost changes as k is varied from 5 to 20. For all datasets, x is set to 9 and the buffer size is 16 pages. Since a larger value of k produces denser Voronoi cells in the space, the cost of $RIS-kNN$ increases as k increases due to more frequent cell crossings. The communication cost of V^*-kNN also has a positive correlation with k as suggested by the analysis in Section 7. V^*-kNN still outperforms $RIS-kNN$ in all settings, as shown in Figure 16. The effect of k on the communication and total costs are very similar. This is because k has positive correlations with: communication, tree-traversal and computation costs.

8.4 Summary

The V^*-kNN algorithm consistently outperforms the $RIS-kNN$ algorithm for all settings using the measures of total response time, number of page accesses and communication costs. V^*-kNN is also more scalable with increasing k values. There is a tradeoff between CPU time and data retrieval costs for V^*-kNN , controlled by the value of x . In other words, x provides the ability to tune the query performance for different application domains.

9. CONCLUSIONS

In this paper, the V^* -Diagram and the associated algorithm V^*-kNN have been introduced to efficiently process moving k nearest neighbor queries ($MkNN$). A key difference between the V^* -Diagram and previous safe-region-based techniques for $MkNN$ queries is that, previous techniques only utilizes the knowledge on the data while the V^* -Diagram exploits also the query location and

the knowledge of the current search space. As a result, the V^* -Diagram is more economical in terms of both IO and CPU costs. We also showed that the V^* -Diagram can be applied to other useful domains such as spatial networks. We performed an extensive experimental study and the results show that our algorithm outperforms the best existing technique by two orders of magnitude.

10. ACKNOWLEDGMENTS

We would like to thank the anonymous reviewers for their comments that improved our paper. This work is supported under the Australian Research Council's Discovery funding scheme (project number DP0880215).

11. REFERENCES

- [1] N. Beckmann, H. Kriegel, R. Schneider, and B. Seeger. The R*-tree: an efficient and robust access method for points and rectangles. In *SIGMOD*, pages 322–331, 1990.
- [2] R. Benetis, C. S. Jensen, G. Karciuskas, and S. Saltenis. Nearest and reverse nearest neighbor queries for moving objects. *VLDB J.*, 15(3):229–249, 2006.
- [3] H. J. Cho and C. W. Chung. An efficient and scalable approach to CNN queries in a road network. In *VLDB*, pages 865–876, 2005.
- [4] A. Guttman. R-trees: a dynamic index structure for spatial searching. In *SIGMOD*, pages 47–57, 1984.
- [5] G. R. Hjaltason and H. Samet. Ranking in spatial databases. In *Symposium on Large Spatial Databases*, pages 83–95, 1995.
- [6] G. R. Hjaltason and H. Samet. Distance browsing in spatial databases. *ACM Trans. Database Syst.*, 24(2):265–318, 1999.
- [7] H. Hu, J. Xu, and D. L. Lee. A generic framework for monitoring continuous spatial queries over moving objects. In *SIGMOD*, pages 479–490, 2005.
- [8] X. Huang, C. S. Jensen, H. Lu, and S. Saltenis. S-GRID: A versatile approach to efficient query processing in spatial networks. In *SSTD*, pages 93–111, 2007.
- [9] X. Huang, C. S. Jensen, and S. Saltenis. The islands approach to nearest neighbor querying in spatial networks. In *SSTD*, pages 73–90, 2005.
- [10] H. V. Jagadish, B. C. Ooi, K.-L. Tan, C. Yu, and R. Zhang. iDistance: An adaptive B⁺-tree based indexing method for nearest neighbor search. *ACM Trans. Database Syst.*, 30(2):364–397, 2005.
- [11] M. R. Kolahdouzan and C. Shahabi. Voronoi-based k nearest neighbor search for spatial network databases. In *VLDB*, pages 840–851, 2004.
- [12] M. R. Kolahdouzan and C. Shahabi. Alternative solutions for continuous k nearest neighbor queries in spatial network databases. *GeoInformatica*, 9(4):321–341, 2005.
- [13] L. Kulik and E. Tanin. Incremental rank updates for moving query points. In *GIScience*, pages 251–268, 2006.
- [14] A. Okabe, B. Boots, and K. Sugihara. *Spatial tessellations: concepts and applications of Voronoi diagrams*. John Wiley & Sons, Inc., 1992.
- [15] D. Papadias, J. Zhang, N. Mamoulis, and Y. Tao. Query processing in spatial network databases. In *VLDB*, pages 802–813, 2003.
- [16] N. Roussopoulos, S. Kelley, and F. Vincent. Nearest neighbor queries. In *SIGMOD*, pages 71–79, 1995.
- [17] H. Samet, J. Sankaranarayanan, and H. Alborzi. Scalable network distance browsing in spatial databases. In *SIGMOD*, pages 43–55, 2008.
- [18] Z. Song and N. Roussopoulos. k -nearest neighbor search for moving query point. In *SSTD*, pages 79–96, 2001.
- [19] Y. Tao and D. Papadias. Time-parameterized queries in spatio-temporal databases. In *SIGMOD*, pages 334–345, 2002.
- [20] Y. Tao, D. Papadias, and Q. Shen. Continuous nearest neighbor search. In *VLDB*, pages 287–298, 2002.
- [21] Y. Tao, J. Zhang, D. Papadias, and N. Mamoulis. An efficient cost model for optimization of nearest neighbor search in low and medium dimensional spaces. *IEEE Trans. Knowl. Data Eng.*, 16(10):1169–1184, 2004.
- [22] X. Yu, K. Q. Pu, and N. Koudas. Monitoring k -nearest neighbor queries over moving objects. In *ICDE*, pages 631–642, 2005.
- [23] J. Zhang, M. Zhu, D. Papadias, Y. Tao, and D. L. Lee. Location-based spatial queries. In *SIGMOD*, pages 443–454, 2003.