

Spatial Data Structures*

Hanan Samet
Computer Science Department and
Institute of Advanced Computer Studies and
Center for Automation Research
University of Maryland
College Park, MD 20742

Abstract

An overview is presented of the use of spatial data structures in spatial databases. The focus is on hierarchical data structures, including a number of variants of quadtrees, which sort the data with respect to the space occupied by it. Such techniques are known as spatial indexing methods. Hierarchical data structures are based on the principle of recursive decomposition. They are attractive because they are compact and depending on the nature of the data they save space as well as time and also facilitate operations such as search. Examples are given of the use of these data structures in the representation of different data types such as regions, points, rectangles, lines, and volumes.

Keywords and phrases: spatial databases, hierarchical spatial data structures, points, lines, rectangles, quadtrees, octrees, R-tree, R^+ -tree image processing.

*This work was supported in part by the National Science Foundation under Grant IRI-9017393. Appears in *Modern Database Systems: The Object Model, Interoperability, and Beyond*, W. Kim, ed., Addison Wesley/ACM Press, Reading, MA, 1995, 361-385.

1 Introduction

Spatial data consists of spatial objects made up of points, lines, regions, rectangles, surfaces, volumes, and even data of higher dimension which includes time. Examples of spatial data include cities, rivers, roads, counties, states, crop coverages, mountain ranges, parts in a CAD system, etc. Examples of spatial properties include the extent of a given river, or the boundary of a given county, etc. Often it is also desirable to attach non-spatial attribute information such as elevation heights, city names, etc. to the spatial data. Spatial databases facilitate the storage and efficient processing of spatial and non-spatial information ideally without favoring one over the other. Such databases are finding increasing use in applications in environmental monitoring, space, urban planning, resource management, and geographic information systems (GIS) [Buchmann et al. 1990; Günther and Schek 1991].

A common way to deal with spatial data is to store it explicitly by parametrizing it and thereby obtaining a reduction to a point in a possibly higher dimensional space. This is usually quite easy to do in a conventional database management system since the system is just a collection of records, where each record has many fields. In particular, we simply add a field (or several fields) to the record that deals with the desired item of spatial information. This approach is fine if we just want to perform a simple retrieval of the data.

However, if our query involves the space occupied by the data (and hence other records by virtue of their proximity), then the situation is not so straightforward. In such a case we need to be able to retrieve records based on some spatial properties which are not stored explicitly in the database. For example, in a roads database, we may not wish to force the user to specify explicitly which roads intersect which other roads or regions. The problem is that the potential volume of such information may be very large and the cost of preprocessing it high, while the cost of computing it on the fly may be quite reasonable, especially if the spatial data is stored in an appropriate manner. Thus we prefer to store the data implicitly so that a wide class of spatial queries can be handled. In particular, we need not know the types of queries a priori.

Being able to respond to spatial queries in a flexible manner places a premium on the appropriate representation of the spatial data. In order to be able to deal with proximity queries the data must be sorted. Of course, all database management systems sort the data. The issue is which keys do they sort on. In the case of spatial data, the sort should be based on all of the spatial keys, which means that, unlike conventional database management systems, the sort is based on the space occupied by the data. Such techniques are known as *spatial indexing* methods.

One approach to the representation of spatial data is to separate it structurally from the nonspatial data while maintaining appropriate links between the two [Aref and Samet 1991a]. This leads to a much higher bandwidth for the retrieval of the spatial data. In such a case, the spatial operations are performed directly on the spatial data structures. This provides the freedom to choose a more appropriate spatial structure than the imposed non-spatial structure (e.g., a relational database). In such a case, a spatial processor can be used that is specifically designed for efficiently dealing with the part of the queries that involve proximity relations and search, and a relational database management system for the part of the queries that involve non-spatial data. Its proper functioning depends on the existence of a query optimizer to determine the appropriate processor for each part of the

query [Aref and Samet 1991b].

As an example of the type of query to be posed to a spatial database system, consider a request to “find the names of the roads that pass through the University of Maryland region”. This requires the extraction of the region locations of all the database records whose “region name” field has the value “University of Maryland” and build a map A . Next, map A is intersected with the road map B to yield a new map C with the selected roads. Now, create a new relation having just one attribute which is the relevant road names of the roads in map C . Of course, there are other approaches to answering the above query. Their efficiency depends on the nature of the data and its volume.

In the rest of this review we concentrate on the data structures used by the spatial processor. In particular, we focus on hierarchical data structures. They are based on the principle of recursive decomposition (similar to *divide and conquer* methods). The term *quadtree* is often used to describe many elements of this class of data structures. We concentrate primarily on region, point, rectangle, and line data. For a more extensive treatment of this subject, see [Samet 1990a; Samet 1990b].

Our presentation is organized as follows. Section 2 describes a number of different methods of indexing spatial data. Section 3 focusses on region data and also briefly reviews the historical background of the origins of hierarchical spatial data structures such as the quadtree. Sections 4, 5, and 6 describe hierarchical representations for point, rectangle, and line data, respectively, as well as give examples of their utility. Section 7 contains concluding remarks in the context of a geographic information system that makes use of these concepts.

2 Spatial Indexing

Each record in a database management system can be conceptualized as a point in a multi-dimensional space. This analogy is used by many researchers (e.g., [Hinrichs and Nievergelt 1983; Jagadish 1990]) to deal with spatial data as well by use of suitable transformations that map the spatial object (henceforth we just use the term *object*) into a point (termed a *representative point*) in either the same (e.g., [Jagadish 1990]), lower (e.g., [Orenstein and Merrett 1984]), or higher (e.g., [Hinrichs and Nievergelt 1983]) dimensional spaces. This analogy is not always appropriate for spatial data. One problem is that the dimensionality of the representative point may be too high [Orenstein 1989]. One solution is to approximate the spatial object by reducing the dimensionality of the representative point. Another more serious problem is that use of these transformations does not preserve proximity.

To see the drawback of just mapping spatial data into points in another space, consider the representation of a database of line segments. We use the term *polygonal map* to refer to such a line segment database, consisting of vertices and edges, regardless of whether or not the line segments are connected to each other. Such a database can arise in a network of roads, power lines, rail lines, etc. Using a representative point (e.g., [Jagadish 1990]), each line segment can be represented by its endpoints¹. This means that each line segment is represented by a tuple of four items (i.e., a pair of x coordinate values and a pair of y coordinate values). Thus, in effect, we have constructed a mapping from a two-dimensional

¹Of course, there are other mappings but they have similar drawbacks. We shall use this example in the rest of this section.

space (i.e., the space from which the lines are drawn) to a four-dimensional space (i.e., the space containing the representative point corresponding to the line).

This mapping is fine for storage purposes and for queries that only involve the points that comprise the line segments (including their endpoints). For example, finding all the line segments that intersect a given point or set of points or a given line segment. However, it is not good for queries that involve points or sets of points that are not part of the line segments as they are not transformed to the higher dimensional space by the mapping. Answering such a query involves performing a search in the space from which the lines are drawn rather than in the space into which they are mapped.

As a more concrete example of the shortcoming of the mapping approach suppose that we want to detect if two lines are near each other, or, alternatively, to find the nearest line to a given point or line. This is difficult to do in the four-dimensional space since proximity in the two-dimensional space from which the lines are drawn is not necessarily preserved in the four-dimensional space into which the lines are mapped. In other words, although the two lines may be very close to each other, the Euclidean distance between their representative points may be quite large.

Thus we need different representations for spatial data. One way to overcome these problems is to use data structures that are based on spatial occupancy. Spatial occupancy methods decompose the space from which the data is drawn (e.g., the two-dimensional space containing the lines) into regions called *buckets*. They are also commonly known as *bucketing methods*. Traditionally, bucketing methods such as the grid file [Nievergelt et al. 1984], BANG file [Freeston 1987], LSD trees [Henrich et al. 1989], buddy trees [Seeger and Kriegel 1990], etc. have always been applied to the transformed data (i.e., the representative points). In contrast, we are applying the bucketing methods to the space from which the data is drawn (i.e., two-dimensions in the case of a collection of line segments).

There are four principal approaches to decomposing the space from which the data is drawn. One approach buckets the data based on the concept of a minimum bounding (or enclosing) rectangle. In this case, objects are grouped (hopefully by proximity) into hierarchies, and then stored in another structure such as a B-tree [Comer 1979]. The R-tree (e.g., [Beckmann et al. 1990; Guttman 1984]) is an example of this approach.

The R-tree and its variants are designed to organize a collection of arbitrary spatial objects (most notably two-dimensional rectangles) by representing them as d -dimensional rectangles. Each node in the tree corresponds to the smallest d -dimensional rectangle that encloses its son nodes. Leaf nodes contain pointers to the actual objects in the database, instead of sons. The objects are represented by the smallest aligned rectangle containing them.

Often the nodes correspond to disk pages and, thus, the parameters defining the tree are chosen so that a small number of nodes is visited during a spatial query. Note that the bounding rectangles corresponding to different nodes may overlap. Also, an object may be spatially contained in several nodes, yet it is only associated with one node. This means that a spatial query may often require several nodes to be visited before ascertaining the presence or absence of a particular object.

The basic rules for the formation of an R-tree are very similar to those for a B-tree. All leaf nodes appear at the same level. Each entry in a leaf node is a 2-tuple of the form

(R,O) such that R is the smallest rectangle that spatially contains object O . Each entry in a non-leaf node is a 2-tuple of the form (R,P) such that R is the smallest rectangle that spatially contains the rectangles in the child node pointed at by P . An R-tree of order (m,M) means that each node in the tree, with the exception of the root, contains between $m \leq \lceil M/2 \rceil$ and M entries. The root node has at least two entries unless it is a leaf node.

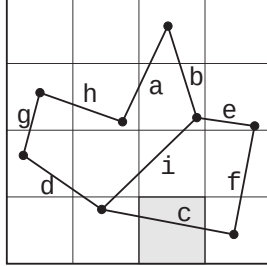


Figure 1: Example collection of line segments embedded in a 4×4 grid.

For example, consider the collection of line segments given in Figure 1 shown embedded in a 4×4 grid. Let $M = 3$ and $m = 2$. One possible R-tree for this collection is given in Figure 2a. Figure 2b shows the spatial extent of the bounding rectangles of the nodes in Figure 2a, with broken lines denoting the rectangles corresponding to the subtrees rooted at the non-leaf nodes. Note that the R-tree is not unique. Its structure depends heavily on the order in which the individual line segments were inserted into (and possibly deleted from) the tree.

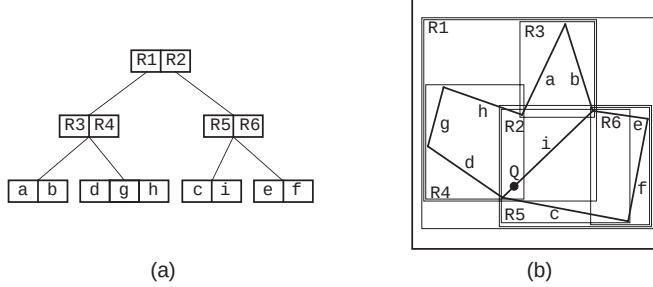


Figure 2: (a) R-tree for the collection of line segments in Figure 1, and (b) the spatial extents of the bounding rectangles.

The drawback of these methods is that they do not result in a disjoint decomposition of space. The problem is that an object is only associated with one bounding rectangle (e.g., line segment i in Figure 2 is associated with rectangle $R5$, yet it passes through $R1$, $R2$, $R4$, and $R5$). In the worst case, this means that when we wish to determine which object is associated with a particular point (e.g., the containing rectangle in a rectangle database, or an intersecting line in a line segment database) in the two-dimensional space from which the objects are drawn, we may have to search the entire database.

For example, suppose we wish to determine the identity of the line segment in the collection of line segments given in Figure 2 that passes through point Q . Since Q can be in either of $R1$ or $R2$, we must search both of their subtrees. Searching $R1$ first, we find that Q could only be contained in $R4$. Searching $R4$ does not lead to the line segment that contains Q even though Q is in a portion of bounding rectangle $R4$ that is in $R1$. Thus, we must search

R2 and we find that Q can only be contained in R5. Searching R5 results in locating i, the desired line segment.

The other approaches are based on a decomposition of space into disjoint cells, which are mapped into buckets. Their common property is that the objects are decomposed into disjoint subobjects such that each of the subobjects is associated with a different cell. They differ in the degree of regularity imposed by their underlying decomposition rules and by the way in which the cells are aggregated. The price paid for the disjointness is that in order to determine the area covered by a particular object, we have to retrieve all the cells that it occupies. This price is also paid when we want to delete an object. Fortunately, deletion is not so common in these databases. A related drawback is that when we wish to determine all the objects that occur in a particular region we often retrieve many of the objects more than once. This is particularly problematic when the result of the operation serves as input to another operation via composition of functions. For example, suppose we wish to compute the perimeter of all the objects in a given region. Clearly, each object's perimeter should only be computed once. Eliminating the duplicates is a serious issue (see [Aref and Samet 1992] for a discussion of how to deal with this problem in a database of line segments).

The first method based on disjointness partitions the objects into arbitrary disjoint subobjects and then groups the subobjects in another structure such as a B-tree. The partition and the subsequent groupings are such that the bounding rectangles are disjoint at each level of the structure. The R^+ -tree [Sellis et al. 1987] and the cell tree [Günther 1988] are examples of this approach. They differ in the data with which they deal. The R^+ -tree deals with collections of objects that are bounded by rectangles, while the cell tree deals with convex polyhedra.

The R^+ -tree is an extension of the k-d-B-tree [Robinson 1981]. The R^+ -tree is motivated by a desire to avoid overlap among the bounding rectangles. Each object is associated with all the bounding rectangles that it intersects. All bounding rectangles in the tree (with the exception of the bounding rectangles for the objects at the leaf nodes) are non-overlapping². The result is that there may be several paths starting at the root to the same object. This may lead to an increase in the height of the tree. However, retrieval time is sped up.

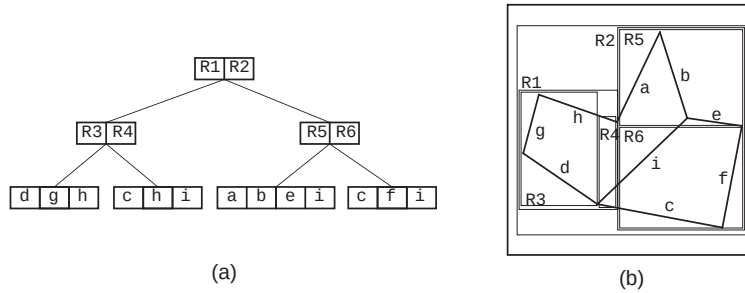


Figure 3: (a) R^+ -tree for the collection of line segments in Figure 1 and (b) the spatial extents of the bounding rectangles.

Figure 3 is an example of one possible R^+ -tree for the collection of line segments in

²From a theoretical viewpoint, the bounding rectangles for the objects at the leaf nodes should also be disjoint. However, this may be impossible (e.g., when the objects are line segments where many line segments intersect at a point).

Figure 1. This particular tree is of order (2,3) although in general it is not possible to guarantee that all nodes will always have a minimum of 2 entries. In particular, the expected B-tree performance guarantees are not valid (i.e., pages are not guaranteed to be m/M full) unless we are willing to perform very complicated record insertion and deletion procedures. Notice that line segments *c* and *h* appear in two different nodes, while line segment *i* appears in three different nodes. Of course, other variants are possible since the R^+ -tree is not unique.

Methods such as the R^+ -tree and the cell tree (as well as the R^* -tree [Beckmann et al. 1990]) have the drawback that the decomposition is data-dependent. This means that it is difficult to perform tasks that require composition of different operations and data sets (e.g., set-theoretic operations such as overlay). In contrast, the remaining two methods, while also yielding a disjoint decomposition, have a greater degree of data-independence. They are based on a regular decomposition. The space can be decomposed either into blocks of uniform size (e.g., the uniform grid [Franklin 1984]) or adapt the decomposition to the distribution of the data (e.g., a quadtree-based approach such as [Samet and Webber 1985]). In the former case, all the blocks are of the same size (e.g., the 4×4 grid in Figure 1). In the latter case, the widths of the blocks are restricted to be powers of two, and their positions are also restricted.

The uniform grid is ideal for uniformly distributed data, while quadtree-based approaches are suited for arbitrarily distributed data. In the case of uniformly distributed data, quadtree-based approaches degenerate to a uniform grid, albeit they have a higher overhead. Both the uniform grid and the quadtree-based approaches lend themselves to set-theoretic operations and thus they are ideal for tasks which require the composition of different operations and data sets. In general, since spatial data is not usually uniformly distributed, the quadtree-based regular decomposition approach is more flexible. The drawback of quadtree-like methods is their sensitivity to positioning in the sense that the placement of the objects relative to the decomposition lines of the space in which they are embedded effects their storage costs and the amount of decomposition that takes place. This is overcome to a large extent by using a bucketing adaptation that decomposes a block only if it contains more than n objects.

All of the spatial occupancy methods discussed above are characterized as employing spatial indexing because with each block the only information that is stored is whether or not the block is occupied by the object or part of the object. This information is usually in the form of a pointer to a descriptor of the object. For example, in the case of a collection of line segments in the uniform grid of Figure 1, the shaded block only records the fact that a line segment crosses it or passes through it. The part of the line segment that passes through the block (or terminates within it) is termed a *q-edge*. Each q-edge in the block is represented by a pointer to a record containing the endpoints of the line segment of which the q-edge is a part [Nelson and Samet 1986]. This pointer is really nothing more than a spatial index and hence the use of this term to characterize this approach. Thus no information is associated with the shaded block as to what part of the line (i.e., q-edge) crosses it. This information can be obtained by clipping [Foley et al. 1990] the original line segment to the block. This is important for often the precision necessary to compute these intersection points is not available.

3 Region Data

A region can be represented either by its interior or by its boundary. In this section we focus on the representations of regions by their interior, while the use of a boundary is discussed in Section 6 in the context of collections of line segments as found, for example, in polygonal maps. The most common region representation is the image array. In this case, we have a collection of picture elements (termed *pixels*). Since the number of elements in the array can be quite large, there is interest in reducing its size by aggregating similar (i.e., homogeneous or equal-valued) pixels. There are two basic approaches. The first approach breaks up the array into $1 \times m$ blocks [Rutovitz 1968]. This is a row representation and is known as a *runlength code*. A more general approach treats the region as a union of maximal square blocks (or blocks of any other desired shape) that may possibly overlap. Usually the blocks are specified by their centers and radii. This representation is called the *medial axis transformation (MAT)* [Blum 1967].

When the maximal blocks are required to be disjoint, to have standard sizes (squares whose sides are powers of two), and to be at standard locations (as a result of a halving process in both the x and y directions), the result is known as a *region quadtree* [Klinger 1971]. It is based on the successive subdivision of the image array into four equal-size quadrants. If the array does not consist entirely of 1s or entirely of 0s (i.e., the region does not cover the entire array), it is then subdivided into quadrants, subquadrants, etc., until blocks are obtained (possibly 1×1 blocks) that consist entirely of 1s or entirely of 0s. Thus, the region quadtree can be characterized as a variable resolution data structure.

As an example of the region quadtree, consider the region shown in Figure 4a which is represented by the $2^3 \times 2^3$ binary array in Figure 4b. Observe that the 1s correspond to pixels that are in the region and the 0s correspond to pixels that are outside the region. The resulting blocks for the array of Figure 4b are shown in Figure 4c. This process is represented by a tree of degree 4.

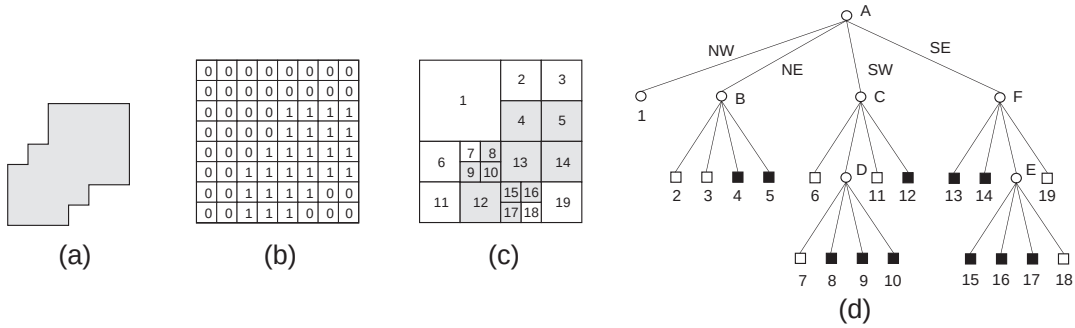


Figure 4: (a) Sample region, (b) its binary array representation, (c) its maximal blocks with the blocks in the region being shaded, and (d) the corresponding quadtree.

In the tree representation, the root node corresponds to the entire array. Each son of a node represents a quadrant (labeled in order NW, NE, SW, SE of the region represented by that node). The leaf nodes of the tree correspond to those blocks for which no further subdivision is necessary. A leaf node is said to be BLACK or WHITE, depending on whether its corresponding block is entirely inside or entirely outside of the represented region. All non-leaf nodes are said to be GRAY. The quadtree representation for Figure 4c is shown in

Figure 4d. Of course, quadtrees can also be used to represent non-binary images. In this case, the same merging criteria is applied to each color. For example, in the case of a landuse map, simply merge all wheat growing regions, and likewise for corn, rice, etc. [Samet et al. 1984].

The term *quadtree* is often used in a more general senses to describe a class of hierarchical data structures whose common property is that they are based on the principle of recursive decomposition of space. They can be differentiated on the following bases:

1. the type of data that they are used to represent,
2. the principle guiding the decomposition process, and
3. the resolution (variable or not).

Currently, they are used for points, rectangles, regions, curves, surfaces, and volumes (see the remaining sections for further details on the adaptation of the quadtree to them). The decomposition may be into equal parts on each level (termed a *regular decomposition*), or it may be governed by the input. The resolution of the decomposition (i.e., the number of times that the decomposition process is applied) may be fixed beforehand or it may be governed by properties of the input data.

Unfortunately, the term *quadtree* has taken on more than one meaning. The region quadtree, as shown above, is a partition of space into a set of squares whose sides are all a power of two long. A similar partition of space into rectangular quadrants is termed a *point quadtree* [Finkel and Bentley 1974]. It is an adaptation of the binary search tree to two dimensions (which can be easily extended to an arbitrary number of dimensions). It is primarily used to represent multidimensional point data where the rectangular regions need not be square. The quadtree is also often confused with the pyramid [Tanimoto and Pavlidis 1975]. The pyramid is a multiresolution representation which is an exponentially tapering stack of arrays, each one-quarter the size of the previous array. In contrast, the region quadtree is a variable resolution data structure.

The distinction between a quadtree and a pyramid is important in the domain of spatial databases, and can be easily seen by considering the types of spatial queries. There are two principal types [Aref and Samet 1990]. The first is location-based. In this case, we are searching for the nature of the feature associated with a particular location or in its proximity. For example, “what is the feature at location X?”, “what is the nearest city to location X?”, or “what is the nearest road to location X?” The second is feature-based. In this case, we are probing for the presence or absence of a feature, as well as seeking its actual location. For example, “does wheat grow anywhere in California?”, “what crops grow in California?”, or “where is wheat grown in California?”

Location-based queries are easy to answer with a quadtree representation as they involve descending the tree until finding the object. If a nearest neighbor is desired, then the search is continued in the neighborhood of the node containing the object. This search can also be achieved by unwinding the process used to access the node containing the object. On the other hand, feature-based queries are more difficult. The problem is that there is no indexing by features. The indexing is only based on spatial occupancy. The goal is to process the query without examining every location in space. The pyramid is useful for

such queries since the nodes that are not at the maximum level of resolution (i.e., at the bottom level) contain summary information. Thus we could view these nodes as feature vectors which indicate whether or not a feature is present at a higher level of resolution. Therefore, by examining the root of the pyramid (i.e., the node that represents the entire image) we can quickly tell if a feature is present without having to examine every location.

For example, consider the block decomposition of the non-binary image in Figure 5a. Its truncated pyramid is given in Figure 5b. The values of a nonleaf node p in the truncated pyramid indicate if the feature is present in the subtrees of p . In the interest of saving space, the pyramid is not shown in its entirety here.

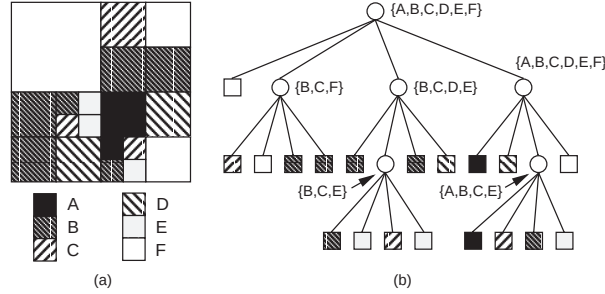


Figure 5: (a) Sample non-binary image, and (b) its corresponding truncated pyramid.

Quadtree-like data structures can also be used to represent images in three dimensions and higher. The octree [Hunter 1978; Meagher 1982] data structure is the three-dimensional analog of the quadtree. It is constructed in the following manner. We start with an image in the form of a cubical volume and recursively subdivide it into eight congruent disjoint cubes (called octants) until blocks are obtained of a uniform color or a predetermined level of decomposition is reached. Figure 6a is an example of a simple three-dimensional object whose raster octree block decomposition is given in Figure 6b and whose tree representation is given in Figure 6c.

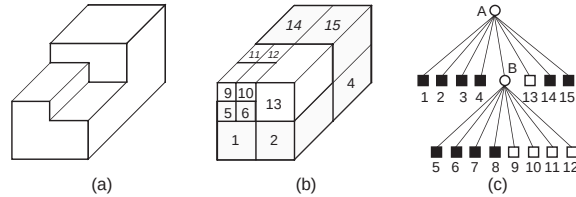


Figure 6: (a) Example three-dimensional object; (b) its octree block decomposition; and (c) its tree representation.

The quadtree is particularly useful for performing set operations as they form the basis of most complicated queries. For example, to “find the names of the roads that pass through the University of Maryland region,” we will need to intersect a region map with a line map. For a binary image, set-theoretic operations such as union and intersection are quite simple to implement [Hunter and Steiglitz 1979].

In particular, the intersection of two quadtrees yields a BLACK node only when the corresponding regions in both quadtrees are BLACK. This operation is performed by simultaneously traversing three quadtrees. The first two trees correspond to the trees being

intersected and the third tree represents the result of the operation. If any of the input nodes are WHITE, then the result is WHITE. When corresponding nodes in the input trees are GRAY, then their sons are recursively processed and a check is made for the mergibility of WHITE leaf nodes. The worst-case execution time of this algorithm is proportional to the sum of the number of nodes in the two input quadtrees, although it is possible for the intersection algorithm to visit fewer nodes than the sum of the nodes in the two input quadtrees.

Performing the set operations on an image represented by a region quadtree is much more efficient than when the image is represented by a boundary representation (e.g., vectors) as it makes use of global data. In particular, to be efficient, a vector-based solution must sort the boundaries of the region with respect to the space which they occupy, while in the case of a region quadtree, the regions are already sorted.

One of the motivations for the development of hierarchical data structures such as the quadtree is a desire to save space. The original formulation of the quadtree encodes it as a tree structure that uses pointers. This requires additional overhead to encode the internal nodes of the tree. In order to further reduce the space requirements, two other approaches have been proposed. The first treats the image as a collection of leaf nodes where each leaf is encoded by a pair of numbers. The first is a base 4 number termed a *locational code*, corresponding to a sequence of directional codes that locate the leaf along a path from the root of the quadtree (e.g., [Gargantini 1982]). It is analogous to taking the binary representation of the x and y coordinates of a designated pixel in the block (e.g., the one at the lower left corner) and interleaving them (i.e., alternating the bits for each coordinate). The second number indicates the depth at which the leaf node is found (or alternatively its size).

The second, termed a *DF-expression*, represents the image in the form of a traversal of the nodes of its quadtree [Kawaguchi and Endo 1980]. It is very compact as each node type can be encoded with two bits. However, it is not easy to use when random access to nodes is desired. For a static collection of nodes, an efficient implementation of the pointer-based representation is often more economical spacewise than a locational code representation [Samet and Webber 1989]. This is especially true for images of higher dimension.

Nevertheless, depending on the particular implementation of the quadtree we may not necessarily save space (e.g., in many cases a binary array representation may still be more economical than a quadtree). However, the effects of the underlying hierarchical aggregation on the execution time of the algorithms are more important. Most quadtree algorithms are simply preorder traversals of the quadtree and, thus, their execution time is generally a linear function of the number of nodes in the quadtree. A key to the analysis of the execution time of quadtree algorithms is the *Quadtree Complexity Theorem* [Hunter 1978] which states that the number of nodes in a quadtree region representation is $O(p + q)$ for a $2^q \times 2^q$ image with perimeter p measured in pixel-widths. In all but the most pathological cases (e.g., a small square of unit width centered in a large image), the q factor is negligible and thus the number of nodes is $O(p)$.

The Quadtree Complexity Theorem holds for three-dimensional data [Meagher 1980] where perimeter is replaced by surface area, as well as for objects of higher dimensions d for which it is proportional to the size of the $(d - 1)$ -dimensional interfaces between these objects.

The Quadtree Complexity Theorem also directly impacts the analysis of the execution time of algorithms. In particular, most algorithms that execute on a quadtree representation of an image instead of an array representation have an execution time that is proportional to the number of blocks in the image rather than the number of pixels. In its most general case, this means that the application of a quadtree algorithm to a problem in d -dimensional space executes in time proportional to the analogous array-based algorithm in the $(d - 1)$ -dimensional space of the surface of the original d -dimensional image. Therefore, quadtrees act like dimension-reducing devices.

4 Point Data

Multidimensional point data can be represented in a variety of ways. The representation ultimately chosen for a specific task is influenced by the type of operations to be performed on the data. Our focus is on dynamic files (i.e., the amount of data can grow and shrink at will) and on applications involving search. In Section 3 we briefly mentioned the point quadtree. In higher dimensions (i.e., greater than 3) it is preferable to use the k-d tree [Bentley 1975] as every node has degree 2 since the partitions cycle through the different attributes.

There are many different representations for point data. Most of them are some variants of the bucket methods discussed in Section 2. These include the grid file and EXCELL which are described in Section 6. For more details, see [Samet 1990b]. In this section we present the PR quadtree (P for point and R for region) [Orenstein 1982; Samet 1990b] as it is based on a regular decomposition. It is an adaptation of the region quadtree to point data which associates data points (that need not be discrete) with quadrants. The PR quadtree is organized in the same way as the region quadtree. The difference is that leaf nodes are either empty (i.e., WHITE) or contain a data point (i.e., BLACK) and its coordinate values. A quadrant contains at most one data point. For example, Figure 7 is a PR quadtree corresponding to some point data.

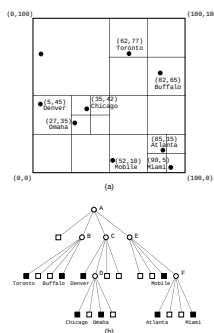


Figure 7: A PR quadtree.

The shape of the PR quadtree is independent of the order in which data points are inserted into it. The disadvantage of the PR quadtree is that the maximum level of decomposition depends on the minimum separation between two points. In particular, if two points are very close, then the decomposition can be very deep. This can be overcome by viewing the blocks or nodes as buckets with capacity c and only decomposing a block when it contains more than c points. Of course, bucketing methods such as the R-tree and the R^+ -tree can also be used.

PR quadtrees, as well as other quadtree-like representations for point data, are especially attractive in applications that involve search. A typical query is one that requests the determination of all records within a specified distance of a given record - e.g., all cities within 100 miles of Washington, DC. The efficiency of the PR quadtree lies in its role as a pruning device on the amount of search that is required. Thus many records will not need to be examined. For example, suppose that in the hypothetical database of Figure 7 we wish to find all cities within 8 units of a data point with coordinates (84,10). In such a case, there is no need to search the NW, NE, and SW quadrants of the root (i.e., (50,50)). Thus we can restrict our search to the SE quadrant of the tree rooted at root. Similarly, there is no need to search the NW, NE, and SW quadrants of the tree rooted at the SE quadrant (i.e., (75,25)). Note that the search ranges are usually orthogonally defined regions such as rectangles, boxes, etc. Other shapes are also feasible as the above example demonstrated (i.e., a circle).

5 Rectangle Data

The rectangle data type lies somewhere between the point and region data types. Rectangles are often used to approximate other objects in an image for which they serve as the minimum rectilinear enclosing object. For example, bounding rectangles can be used in cartographic applications to approximate objects such as lakes, forests, hills, etc. In such a case, the approximation gives an indication of the existence of an object. Of course, the exact boundaries of the object are also stored; but they are only accessed if greater precision is needed. For such applications, the number of elements in the collection is usually small, and most often the sizes of the rectangles are of the same order of magnitude as the space from which they are drawn.

Rectangles are also used in VLSI design rule checking as a model of chip components for the analysis of their proper placement. Again, the rectangles serve as minimum enclosing objects. In this application, the size of the collection is quite large (e.g., millions of components) and the sizes of the rectangles are several orders of magnitude smaller than the space from which they are drawn.

The representation that is used depends heavily on the problem environment. If the environment is static, then frequently the solutions are based on the use of the plane-sweep paradigm [Preparata and Shamos 1985], which usually yields optimal solutions in time and space. However, the addition of a single object to the database forces the re-execution of the algorithm on the entire database. We are primarily interested in dynamic problem environments. The data structures that are chosen for the collection of the rectangles are differentiated by the way in which each rectangle is represented.

One representation discussed in Section 2 reduces each rectangle to a point in a higher dimensional space, and then treats the problem as if we have a collection of points [Hinrichs and Nievergelt 1983]. Each rectangle is a Cartesian product of two one-dimensional intervals where each interval is represented by its centroid and extent. Each set of intervals in a particular dimension is, in turn, represented by a grid file [Nievergelt et al. 1984].

The grid file is a two-level grid for storing multidimensional points. It uses a grid directory (a two-dimensional array of grid blocks for two-dimensional point data) on disk indicating the address of the bucket (i.e., page) that contains the data associated with the

grid block. A set of linear scales (actually a pair of one-dimensional arrays in the case of two-dimensional data) are kept in core. The linear scales access the grid block in the grid directory (on disk) that is associated with a particular point. The grid file guarantees access to any record with two disk operations – that is, one for each level of the grid. The first access is to the grid block, while the second access is to the grid bucket. The linear scales are necessary because the grid lines in the grid directory can be in arbitrary positions.

In contrast, EXCELL [Tamminen 1981] also guarantees access to any record with two disk operations but makes use of regular decomposition. This means that the linear scales are not necessary. However, a grid partition results in doubling the size of the grid directory.

The second representation is region-based in the sense that the subdivision of the space from which the rectangles are drawn depends on the physical extent of the rectangle - not just one point. Representing the collection of rectangles, in turn, with a tree-like data structure has the advantage that there is a relation between the depth of node in the tree and the size of the rectangle(s) that are associated with it. Interestingly, some of the region-based solutions make use of the same data structures that are used in the solutions based on the plane-sweep paradigm.

There are three types of region-based solutions currently in use. The first two solutions adapt the R-tree and the R⁺-tree (discussed in Section 2) to store rectangle data (i.e., in this case the objects are rectangles instead of line segments as in Figures 2 and 3). The third is a quadtree-based approach and uses the MX-CIF quadtree [Kedem 1982].

In the MX-CIF *quadtree* each rectangle is associated with the quadtree node corresponding to the smallest block which contains it in its entirety. Subdivision ceases whenever a node's block contains no rectangles. Alternatively, subdivision can also cease once a quadtree block is smaller than a predetermined threshold size. This threshold is often chosen to be equal to the expected size of the rectangle [Kedem 1982]. For example, Figure 8 is the MX-CIF quadtree for a collection of rectangles. Note that rectangle **F** occupies an entire block and hence it is associated with the block's father. Also rectangles can be associated with both terminal and non-terminal nodes.

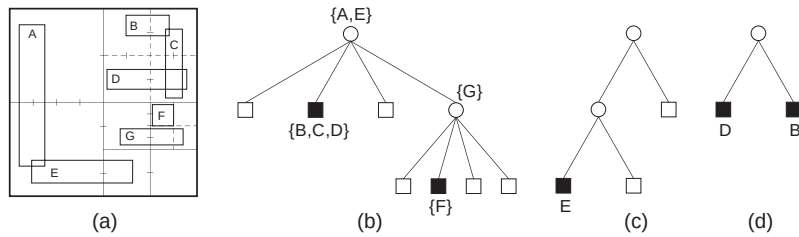


Figure 8: (a) Collection of rectangles and the block decomposition induced by the MX-CIF quadtree; (b) the tree representation of (a); the binary trees for the y axes passing through the root of the tree in (b), and (d) the NE son of the root of the tree in (b).

It should be clear that more than one rectangle can be associated with a given enclosing block and, thus, often we find it useful to be able to differentiate between them. This is done in the following manner [Kedem 1982]. Let P be a quadtree node with centroid (CX, CY) , and let S be the set of rectangles that are associated with P . Members of S are organized into two sets according to their intersection (or collinearity of their sides) with the lines passing through the centroid of P 's block – that is, all members of S that intersect the line

$x = CX$ form one set and all members of S that intersect the line $y = CY$ form the other set.

If a rectangle intersects both lines (i.e., it contains the centroid of P 's block), then we adopt the convention that it is stored with the set associated with the line through $x = CX$. These subsets are implemented as binary trees (really tries), which in actuality are one-dimensional analogs of the MX-CIF quadtree. For example, Figure 8c and Figure 8d illustrate the binary trees associated with the y axes passing through the root and the NE son of the root of the MX-CIF quadtree of Figure 8c, respectively. Interestingly, the MX-CIF quadtree is a two-dimensional analog of the interval tree [Edelsbrunner 1980], which is a data structure that is used to support optimal solutions based on the plane-sweep paradigm to some rectangle problems.

6 Line Data

Section 3 was devoted to the region quadtree, an approach to region representation that is based on a description of the region's interior. In this section, we focus on a representation that specifies the boundaries of regions. The simplest representation is the polygon consisting of vectors which are usually specified in the form of lists of pairs of x and y coordinate values corresponding to their start and end points. The vectors are usually ordered according to their connectivity. One of the most common representations is the chain code [Freeman 1974] which is an approximation of a polygon's boundary by use of a sequence of unit vectors in the four principal directions. Using such representations, given a random point in space, it is very difficult to find the nearest line to it as the lines are not sorted. Nevertheless, the vector representation is used in many commercial systems (e.g., ARC/INFO [Peuquet and Marble 1990]) on account of its compactness.

In this section we concentrate on the use of bucketing methods. There are a number of choices (see [Hoel and Samet 1992] for an empirical comparison). The first two are the R-tree and the R⁺-tree which have already been explained in Section 2. The third uses regular decomposition to adaptively sort the line segments into buckets of varying size. There is a one-to-one correspondence between buckets and blocks in the two-dimensional space from which the line segments are drawn. There are a number of approaches to this problem [Samet 1990b]. They differ by being either vertex-based or edge-based. Their implementations make use of the same basic data structure. All are built by applying the same principle of repeatedly breaking up the collection of vertices and edges (making up the polygonal map) into groups of four blocks of equal size (termed *brothers*) until obtaining a subset that is sufficiently simple so that it can be organized by some other data structure. This is achieved by successively weakening the definition of what constitutes a legal block, thereby enabling more information to be stored in each bucket.

The PM quadtree family [Samet and Webber 1985] is vertex-based. We illustrate the PM₁ quadtree. It is based on a decomposition rule stipulating that partitioning occurs as long as a block contains more than one line segment unless the line segments are all incident at the same vertex which is also in the same block (e.g., Figure 9). A similar representation has been devised for three-dimensional images (e.g., [Ayala et al. 1985]). The decomposition criteria are such that no node contains more than one face, edge, or vertex unless the faces all meet at the same vertex or are adjacent to the same edge. This representation is quite useful since its space requirements for polyhedral objects are significantly smaller than those

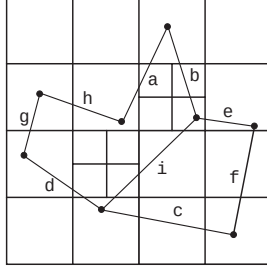


Figure 9: PM₁ quadtree for the collection of line segments of Figure 1.

of a conventional octree.

The PMR quadtree [Nelson and Samet 1986; Nelson and Samet 1987] is an edge-based variant of the PM quadtree (see also edge-EXCELL [Tamminen 1981]). It makes use of a probabilistic splitting rule. A block is permitted to contain a variable number of line segments. The PMR quadtree is constructed by inserting them one-by-one into an initially empty structure consisting of one block. Each line segment is inserted into all of the blocks that it intersects or occupies in its entirety. During this process, the occupancy of each affected block is checked to see if the insertion causes it to exceed a predetermined *splitting threshold*. If the splitting threshold is exceeded, then the block is split *once*, and only once, into four blocks of equal size. The rationale is to avoid splitting a node many times when there are a few very close lines in a block. In this manner, we avoid pathologically bad cases. For more details, see [Nelson and Samet 1986].

A line segment is deleted from a PMR quadtree by removing it from all the blocks that it intersects or occupies in its entirety. During this process, the occupancy of the block and its siblings (the ones that were created when its predecessor was split) is checked to see if the deletion causes the total number of line segments in them to be less than the splitting threshold. If the splitting threshold exceeds the occupancy of the block and its siblings, then they are merged and the merging process is recursively reapplied to the resulting block and its siblings. Notice the asymmetry between splitting and merging rules.

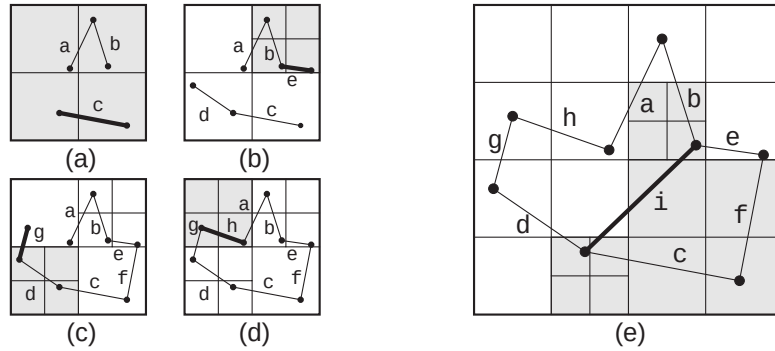


Figure 10: PMR quadtree for the collection of line segments of Figure 1. (a) – (e) illustrate snapshots of the construction process with the final PMR quadtree given in (e).

Figure 10e is an example of a PMR quadtree corresponding to a set of 9 edges labeled a – i inserted in increasing order. Observe that the shape of the PMR quadtree for a given polygonal map is not unique; instead it depends on the order in which the lines are inserted

into it. In contrast, the shape of the PM_1 quadtree is unique. Figure 10a–e shows some of the steps in the process of building the PMR quadtree of Figure 10e. This structure assumes that the splitting threshold value is two. In each part of Figure 10a–e, the line segment that caused the subdivision is denoted by a thick line, while the gray regions indicate the blocks where a subdivision has taken place. The insertion of line segments c , e , g , h , and i cause the subdivisions in parts a, b, c, d, and e, respectively, of Figure 10. The insertion of line segment i causes three blocks to be subdivided (i.e., the SE block in the SW quadrant, the SE quadrant, and the SW block in the NE quadrant). The final result is shown in Figure 10e. Note the difference from the PM_1 quadtree in Figure 9 – that is, the NE block of the SW quadrant is decomposed in the PM_1 quadtree while the SE block of the SW quadrant is not decomposed in the PM_1 quadtree.

The PMR quadtree is very good for answering queries such as finding the nearest line to a given point [Hoel and Samet 1991]. It is preferred over the PM_1 quadtree as it results in far fewer subdivisions. In particular, in the PMR quadtree there is no need to subdivide in order to separate line segments that are very “close” or whose vertices are very “close,” which is the case for the PM_1 quadtree. This is important since four blocks are created at each subdivision step. Thus when many subdivision steps occur, many empty blocks are created and thus the storage requirements of the PMR quadtree are considerably lower than those of the PM_1 quadtree. Generally, as the splitting threshold is increased, the storage requirements of the PMR quadtree decrease while the time necessary to perform operations on it will increase. Another advantage of the PMR quadtree over the PM_1 quadtree is that by virtue of being edge based, it can easily deal with nonplanar graphs.

Observe that although a bucket in the PMR quadtree can contain more line segments than the splitting threshold, this is not a problem. In fact, it can be shown [Samet 1990b] that the maximum number of line segments in a bucket is bounded by the sum of the splitting threshold and the depth of the block (i.e., the number of times the original space has been decomposed to yield this block).

7 Concluding Remarks

The use of hierarchical data structures in spatial databases enables the focussing of computational resources on the interesting subsets of data. Thus, there is no need to expend work where the payoff is small. Although many of the operations for which they are used can often be performed equally as efficiently, or more so, with other data structures, hierarchical data structures are attractive because of their conceptual clarity and ease of implementation.

When the hierarchical data structures are based on the principle of regular decomposition, we have the added benefit of a spatial index. All features, be they regions, points, rectangles, lines, volumes, etc., can be represented by maps which are in registration. This means that a query such as “finding the names of the roads that pass through the University of Maryland region” can be executed by simply overlaying the region and road maps even though they represent data of different types. The overlay performs an intersection operation where the common feature is the area spanned by the University of Maryland and the roads that pass through it (i.e., a spatial join).

In fact, such a system, known as QUILT, has been built [Shaffer et al. 1990] for representing geographic information using the quadtree variants described here. In this case, the

quadtree is implemented as a collection of leaf nodes where each leaf node is represented by its locational code, which is a spatial index. The collection is in turn represented as a B-tree. There are leaf nodes corresponding to region, point, and line data.

The disadvantage of quadtree methods is that they are shift sensitive in the sense that their space requirements are dependent on the position of the origin. However, for complicated images the optimal positioning of the origin will usually lead to little improvement in the space requirements. The process of obtaining this optimal positioning is computationally expensive and is usually not worth the effort [Li et al. 1982].

The fact that we are working in a digitized space may also lead to problems. For example, the rotation operation is not generally invertible. In particular, a rotated square usually cannot be represented accurately by a collection of rectilinear squares. However, when we rotate by 90° , then the rotation is invertible. This problem arises whenever one uses a digitized representation. Thus, it is also common to the array representation.

8 Acknowledgements

The assistance of Erik G. Hoel in the preparation of the figures is greatly appreciated.

References

- AREF, W. G. AND SAMET, H. 1990. Efficient processing of window queries in the pyramid data structure. In *Proceedings of the 9th ACM SIGACT-SIGMOD-SIGART Symposium on Principles of Database Systems (PODS)*, (Nashville, TN, Apr.), pp. 265–272.
- AREF, W. G. AND SAMET, H. 1991a. Extending a DBMS with spatial operations. In GÜNTHER, O. AND SCHEK, H.-J., EDS., *Advances in Spatial Databases*, pp. 299–318. Lecture Notes in Computer Science 525. Springer-Verlag, Berlin.
- AREF, W. G. AND SAMET, H. 1991b. Optimization strategies for spatial query processing. In LOHMAN, G., ED., *Proceedings of the Seventeenth International Conference on Very Large Databases (VLDB)*, (Barcelona, Spain, Sept.), pp. 81–90.
- AREF, W. G. AND SAMET, H. 1992. Uniquely reporting spatial objects: yet another operation for comparing spatial data structures. In *Proceedings of the Fifth International Symposium on Spatial Data Handling*, volume 1, (Charleston, SC, Aug.), pp. 178–189.
- AYALA, D., BRUNET, P., JUAN, R., AND NAVAZO, I. 1985. Object representation by means of nonminimal division quadtrees and octrees. *ACM Transactions on Graphics*, 4, 1 (Jan.), 41–59.
- BECKMANN, N., KRIEGEL, H. P., SCHNEIDER, R., AND SEEGER, B. 1990. The R*-tree: an efficient and robust access method for points and rectangles. In *Proceedings of the SIGMOD Conference*, (Atlantic City, NJ, June), pp. 322–331.
- BENTLEY, J. L. 1975. Multidimensional binary search trees used for associative searching. *Communications of the ACM*, 18, 9 (Sept.), 509–517.
- BLUM, H. 1967. A transformation for extracting new descriptors of shape. In WATHEN-DUNN, W., ED., *Models for the Perception of Speech and Visual Form*, pp. 362–380. MIT Press, Cambridge, MA.
- BUCHMANN, A., GÜNTHER, O., SMITH, T. R., AND WANG, Y.-F., EDS. 1990. *Design and Implementation of Large Spatial Databases*. Lecture Notes in Computer Science 409. Springer-Verlag, Berlin.

- COMER, D. 1979. The ubiquitous B-tree. *ACM Computing Surveys*, 11, 2 (June), 121–137.
- EDELSBRUNNER, H. 1980. Dynamic rectangle intersection searching. Institute for Information Processing Report 47, Institute for Information Processing, Technical University of Graz, Graz, Austria, (Feb.).
- FINKEL, R. AND BENTLEY, J. 1974. Quad trees: a data structure for retrieval on composite keys. *Acta Informatica*, 4, 1, 1–9.
- FOLEY, J. D., VAN DAM, A., FEINER, S. K., AND HUGHES, J. F. 1990. *Computer Graphics: Principles and Practice*. Addison-Wesley, Reading, MA, second edition.
- FRANKLIN, W. R. 1984. Adaptive grids for geometric operations. *Cartographica*, 21, 2 & 3 (Summer & Autumn), 160–167.
- FREEMAN, H. 1974. Computer processing of line-drawing images. *ACM Computing Surveys*, 6, 1 (Mar.), 57–97.
- FREESTON, M. 1987. The BANG file: a new kind of grid file. In *Proceedings of the SIGMOD Conference*, (San Francisco, May), pp. 260–269.
- GARGANTINI, I. 1982. An effective way to represent quadtrees. *Communications of the ACM*, 25, 12 (Dec.), 905–910.
- GÜNTHER, O. 1988. *Efficient structures for geometric data management*. PhD thesis, Electronics Research Laboratory, College of Engineering, University of California at Berkeley, Berkeley, CA. (Lecture Notes in Computer Science 337, Springer-Verlag, Berlin, 1988).
- GÜNTHER, O. AND SCHEK, H.-J., EDS. 1991. *Advances in Spatial Databases*. Lecture Notes in Computer Science 525. Springer-Verlag, Berlin.
- GUTTMAN, A. 1984. R-trees: a dynamic index structure for spatial searching. In *Proceedings of the SIGMOD Conference*, (Boston, June), pp. 47–57.
- HENRICH, A., SIX, H. W., AND WIDMAYER, P. 1989. The LSD tree: spatial access to multidimensional point and non-point data. In APERS, P. M. G. AND WIEDERHOLD, G., EDS., *Proceedings of the Fifteenth International Conference on Very Large Databases (VLDB)*, (Amsterdam, Aug.), pp. 45–53.
- HINRICHS, K. AND NIEVERGELT, J. 1983. The grid file: a data structure designed to support proximity queries on spatial objects. In NAGL, M. AND J. PERL, E., EDS., *Proceedings of the WG'83 (International Workshop on Graphtheoretic Concepts in Computer Science)*, (Truner Verlag, Linz, Austria,), pp. 100–113.
- HOEL, E. G. AND SAMET, H. 1991. Efficient processing of spatial queries in line segment databases. In GÜNTHER, O. AND SCHEK, H.-J., EDS., *Advances in Spatial Databases*, pp. 237–256. Lecture Notes in Computer Science 525. Springer-Verlag, Berlin.
- HOEL, E. G. AND SAMET, H. 1992. A qualitative comparison study of data structures for large line segment databases. In *Proceedings of the SIGMOD Conference*, (San Diego, June), pp. 205–214.
- HUNTER, G. M. 1978. *Efficient computation and data structures for graphics*. PhD thesis, Department of Electrical Engineering and Computer Science, Princeton University, Princeton, NJ.
- HUNTER, G. M. AND STEIGLITZ, K. 1979. Operations on images using quad trees. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 1, 2 (Apr.), 145–153.

- JAGADISH, H. V. 1990. On indexing line segments. In MCLEOD, D., SACKS-DAVIS, R., AND SCHEK, H., EDS., *Proceedings of the Sixteenth International Conference on Very Large Databases (VLDB)*, (Brisbane, Australia, Aug.), pp. 614–625.
- KAWAGUCHI, E. AND ENDO, T. 1980. On a method of binary picture representation and its application to data compression. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2, 1 (Jan.), 27–35.
- KEDEM, G. 1982. The quad-CIF tree: a data structure for hierarchical on-line algorithms. In *Proceedings of the Nineteenth Design Automation Conference*, (Las Vegas, June), pp. 352–357.
- KLINGER, A. 1971. Patterns and search statistics. In RUSTAGI, J. S., ED., *Optimizing Methods in Statistics*, pp. 303–337. Academic Press, New York.
- LI, M., GROSCH, W. I., AND JAIN, R. 1982. Normalized quadtrees with respect to translations. *Computer Graphics and Image Processing*, 20, 1 (Sept.), 72–81.
- MEAGHER, D. 1980. Octree encoding: a new technique for the representation, the manipulation, and display of arbitrary 3-d objects by computer. Technical Report IPL-TR-80-111, Electrical and Systems Engineering Technical Report, Rensselaer Polytechnic Institute, Troy, NY, (Oct.).
- MEAGHER, D. 1982. Geometric modeling using octree encoding. *Computer Graphics and Image Processing*, 19, 2 (June), 129–147.
- NELSON, R. C. AND SAMET, H. 1986. A consistent hierarchical representation for vector data. *Computer Graphics*, 20, 4 (Aug.), 197–206. (also *Proceedings of the SIGGRAPH'86 Conference*, Dallas, August 1986).
- NELSON, R. C. AND SAMET, H. 1987. A population analysis for hierarchical data structures. In *Proceedings of the SIGMOD Conference*, (May), pp. 270–277.
- NIEVERGELT, J., HINTERBERGER, H., AND SEVCIK, K. C. 1984. The grid file: an adaptable, symmetric multikey file structure. *ACM Transactions on Database Systems*, 9, 1 (Mar.), 38–71.
- ORENSTEIN, J. A. 1982. Multidimensional tries used for associative searching. *Information Processing Letters*, 14, 4 (June), 150–157.
- ORENSTEIN, J. A. 1989. Redundancy in spatial databases. In *Proceedings of the SIGMOD Conference*, (Portland, OR, June), pp. 294–305.
- ORENSTEIN, J. A. AND MERRETT, T. H. 1984. A class of data structures for associative searching. In *Proceedings of the Third ACM SIGACT-SIGMOD Symposium on Principles of Database Systems*, (Waterloo, Canada, Apr.), pp. 181–190.
- PEUQUET, D. J. AND MARBLE, D. F. 1990. ARC/INFO: An example of a contemporary geographic information system. In PEUQUET, D. J. AND MARBLE, D. F., EDS., *Introductory Readings In Geographic Information Systems*, pp. 90–99. Taylor & Francis, London.
- PREPARATA, F. P. AND SHAMOS, M. I. 1985. *Computational Geometry: An Introduction*. Springer-Verlag, New York.
- ROBINSON, J. T. 1981. The k-d-b-tree: a search structure for large multidimensional dynamic indexes. In *Proceedings of the SIGMOD Conference*, (Ann Arbor, MI, Apr.), pp. 10–18.
- RUTOVITZ, D. 1968. Data structures for operations on digital images. In ET AL., G. C. C., ED., *Pictorial Pattern Recognition*, pp. 105–133. Thompson Book Co., Washington, DC.
- SAMET, H. 1990a. *Applications of Spatial Data Structures: Computer Graphics, Image Processing, and GIS*. Addison-Wesley, Reading, MA.

- SAMET, H. 1990b. *The Design and Analysis of Spatial Data Structures*. Addison-Wesley, Reading, MA.
- SAMET, H., ROSENFELD, A., SHAFFER, C. A., AND WEBBER, R. E. 1984. A geographic information system using quadrees. *Pattern Recognition*, 17, 6 (November/December), 647–656.
- SAMET, H. AND WEBBER, R. E. 1985. Storing a collection of polygons using quadrees. *ACM Transactions on Graphics*, 4, 3 (July), 182–222. (also *Proceedings of Computer Vision and Pattern Recognition 83*, Washington, DC, June 1983, 127–132; and University of Maryland Computer Science TR-1372).
- SAMET, H. AND WEBBER, R. E. 1989. A comparison of the space requirements of multi-dimensional quadtree-based file structures. *Visual Computer*, 5, 6 (Dec.), 349–359. (also University of Maryland Computer Science TR-1711).
- SEEGER, B. AND KRIEGEL, H. P. 1990. The buddy-tree: an efficient and robust access method for spatial data base systems. In MCLEOD, D., SACKS-DAVIS, R., AND SCHEK, H., EDS., *Proceedings of the Sixteenth International Conference on Very Large Databases (VLDB)*, (Brisbane, Australia, Aug.), pp. 590–601.
- SELLIS, T., ROUSSOPOULOS, N., AND FALOUTSOS, C. 1987. The R^+ -tree: a dynamic index for multi-dimensional objects. In STOCKER, P. M. AND KENT, W., EDS., *Proceedings of the Thirteenth International Conference on Very Large Databases (VLDB)*, (Brighton, England, Sept.), pp. 507–518.
- SHAFFER, C. A., SAMET, H., AND NELSON, R. C. 1990. QUILT: A geographic information system based on quadrees. *International Journal of Geographical Information Systems*, 4, 2 (April–June), 103–131.
- TAMMINEN, M. 1981. The EXCELL method for efficient geometric access to data. *Acta Polytechnica Scandinavica*, Mathematics and Computer Science Series, No. 34, Helsinki, Finland.
- TANIMOTO, S. AND PAVLIDIS, T. 1975. A hierarchical data structure for picture processing. *Computer Graphics and Image Processing*, 4, 2 (June), 104–119.