

The Application of Legibility Techniques to Enhance Information Visualisations

Rob Ingram and Steve Benford
Department of Computer Science
The University of Nottingham
Nottingham, NG7 2RD,UK

email: { rji, sdb } @cs.nott.ac.uk
<http://www.crg.cs.nott.ac.uk/>
fax: +44 115 951 4254
phone: +44 115 951 4225

Abstract

This paper shows how previous research into navigation through urban environments, which has emerged from the discipline of urban planning, can be adapted to enhance the design of information visualisations. The paper draws on Kevin Lynch's seminal work on the legibility of urban landscapes in order to propose a set of general techniques which can be applied to the task of information visualisation. It describes a specific instantiation of these techniques called LEADS, a legibility system which post-processes the output of a range of existing visualisation systems in order to enhance their legibility. The paper provides four examples of the application of LEADS to different information visualisations. Following this, it discusses experimental work, the conclusions of which provide some tentative support for the likely success of this approach. The outcomes of this work are both a recognition of the important relationship between the disciplines of urban planning and the design of information visualisations as well as more concrete algorithms to be used by the designers of such visualisations.¹

¹We would like to gratefully acknowledge the support of the UK's Engineering and Physical Sciences Research Council (EPSRC) in this work.

1. Introduction

This paper concerns the structure and navigation of information visualisations. More specifically, it provides some observations on how the design of such visualisations might affect people's ability to learn to navigate them. These observations are then translated into computational techniques (i.e. algorithms) for enhancing visualisations. In turn, these techniques are backed up with experimental results providing some early indications of their potential usefulness.

This work has arisen from the observation that research into the relationship between spatial form and navigation has not only been confined to the realm of information visualisation; there has been a wealth of previous research on this topic in the fields of urban planning and architecture. In many cases, this work has yielded concrete (!) results in terms of well defined design principles to support navigation in real-world spatial environments. The work described in this paper borrows directly from this previous research and applies it to the task of information visualisation. We focus in particular on the concept of the *legibility* of an urban landscape as initially developed by Kevin Lynch in the 1960s. Lynch's work focused on how an individual might develop

a so called cognitive map of an urban landscape as a result of repeated exposures to it. Following a series of experiments involving the navigation of real cities, he proposed the existence of five key features of an urban landscape which, if correctly designed, could enhance people's ability to form a cognitive map. These were: districts, edges, landmarks, paths and nodes.

This paper reviews Lynch's work and then goes on to develop a series of computational algorithms for *automatically* generating or emphasising these features in a range of information visualisations. It describes the implementation of these algorithms as a "legibility layer" called LEADS which post-processes the output of various information visualisations in order to provide additional navigation support. As such, this paper is not so much concerned with the core design of visualisation layouts as it is with how existing visualisations can be enhanced. Some evaluation experiments are then described, the results of which provide some initial justification for these algorithms and for this approach in general.

2. Goals and relationship to previous work

Before progressing further, it is first necessary to clearly state the problem that we are trying to address and to say something of its relationship to previous work on navigation and visualisation. Our work builds on two areas of research: navigation within large-scale information systems and related to this, the design of graphical information visualisations.

A problem which is often encountered in the use of large computer based information systems is that of getting lost. For example, considering the area of Hypertext and Hypermedia systems for a moment, Elm and Woods [1] define three categories of being lost:

- Not knowing where to go next;
- Knowing where to go but not how to get there;
- Not knowing a current location in relation to an overall context.

In addition other categories might be considered, such as mistaking the global or local frame of reference, erroneously believing one knows one's location or the path to the goal, or simply being disoriented.

Jacob Nielsen addressed this issue in a number of related papers [2][3][4] citing cases where users were unable to return to given locations within a hypertext system. A number of solutions have been proposed to these kinds of navigation problems, including guides, tours and various backtracking mechanisms. Indeed, variants of these have been integrated in the current generation of World Wide Web browsers (e.g. bookmarks and hot-lists). These techniques, however, introduce their own problems; for example, the passive movement mode in guided tours may inhibit learning of the system and reduces the hypertext to a linear form and although the information provided might help in some tasks a more active approach is preferred. As a result, Nielsen has proposed the use of overview diagrams (effectively maps) and fish-eye techniques to provide both local detail and global context. This naturally leads us on to the subject of visualisation.

Information visualisation techniques are intended, at least in part, to overcome some of these problems for a range of different information systems. So called "focus and context" techniques aim to situate a specific item of information within a representation of its surrounding context, thereby addressing the third of the above categories of being lost. Perhaps the best known examples of this approach are the Perspective Wall [5] and Cone Tree [6] visualisations from Xerox PARC.

More recently, researchers have begun to explore the potential of three dimensional

information visualisations which utilise interactive 3-D graphics to allow one or more users to situate themselves within a visualisation and to move about it or fly over it. Authors such as Thomas Erikson [7] and Kim Fairchild [8] have argued that three dimensional visualisation techniques may help increase the amount of information that people can meaningfully manage and, although the merits of the 3-D approach remain largely unproved, there is clearly a growing interest in this field. Indeed, there have already been numerous specific examples of three dimensional visualisations, many of the best known ones having been reported at the ACM's series of Computer Human Interaction (CHI) conferences.

The work described in the paper builds upon previous work concerning the navigation of information systems and the construction of information visualisations. Our specific aim is to introduce a framework for *automatically enhancing existing information visualisations* so as to aid users in more easily *learning to navigate* them. As such, we are not directly concerned with the design of specific visualisation layouts or indeed with application to specific kinds of information systems. Although the examples that we use to illustrate our work are largely based upon document repositories, the World Wide Web and three dimensional visualisations, our techniques are intended to be applicable across a wider range of visualisation styles and underlying information systems than these.

At this stage, we must also be careful to precisely state our interest in navigation. We are not primarily concerned with the mechanics of controlling the movement of viewpoints and embodiments (although we will briefly touch on this theme later on); nor are we primarily concerned with how people find their way around unfamiliar information environments (i.e. ones that they haven't seen before). Instead, our specific aim is to support people in learning to

navigate visualisations as a consequence of repeated exposures to them over a period of time. In other words, we are concerned with how people can be aided in gradually learning the structure of a graphical space. We anticipate that the techniques that we propose will be mostly applicable to visualisations which are persistent, which evolve relatively slowly and which are repeatedly visited (a visualisation of a region of the World Wide Web would be a good example).

Our approach to this particular navigation problem has been to apply legibility techniques, developed in the domain of urban planning, to the domain of information visualisation. This approach has also been recently adopted by other researchers. Rudy Darken and John Sibert [9] have considered a number of issues associated with navigation in virtual environments. They have evaluated a space which makes use of legibility features such as landmarks and districts (see section 3) to aid navigation tasks in a flat simulation of the sea containing ship objects. In developing the BEAD system Matthew Chalmers [10][11] has also considered the use of legibility features, among other techniques, in the production of useable information spaces and uses landmarks, edges and districts in a visualisation system based on a landscape metaphor.

This work presented here differs from this research in its development of techniques to *automatically* construct legibility features in the abstract spaces produced by a variety of existing or future visualisation systems. One of our main aims has been to accomplish this without requiring the users of visualisation systems to perform the placement of the features manually. Note also that because this research draws heavily on city planning literature this paper uses terminology from this area rather than the more traditional terminology of computer graphics.

In summary then, our aim has been to generate and evaluate techniques for automatically enhancing the legibility of information

visualisations in order to support users in more easily learning to navigate them over a period of time. The remainder of this paper is structured as follows. Section 3 reviews previous work on navigation in urban environments and introduces the key concept of legibility as proposed by Lynch. Section 4 adapts the concept of legibility to the task of information visualisation and describes the set of algorithms and techniques which have been developed for the automatic creation or enhancement of legibility features. Next, section 5 demonstrates the application of our algorithms to four existing visualisations: a network drawing tool, an interactive searching interface to a document repository, a scatter graph drawing tool and a landscape oriented document visualisation. Section 6 describes an initial experiment to assess the effectiveness of this implementation. Finally, section 7 offers by some reflections on specific problems encountered with this work and outlines possible next steps in its evolution.

3. The legibility of urban environments

Legibility, in the context of urban environments, is a term which has been used for many years in the discipline of City Planning. Work in this area has focused on the ease with which people are able to learn the layout of a city and then use this knowledge to perform wayfinding tasks. In his book “The Image of the City” [12] Kevin Lynch defines the legibility of a city as:

“-the ease with which its parts may be recognised and can be organised into a coherent pattern-”

Here, Lynch is referring to the formation of a *cognitive map* within the persons mind [13], a structure which is an internal representation of an environment which its inhabitants use as a reference when navigating to a destination. It is

proposed that cognitive maps fall into two types. Linear or sequential maps are based on movement through the space, or sequential images of a route, and what is observed on the journey. Spatial maps do not require this reference movement through the imagined space and so areas within the map do not need to be linked to routes through the area. A persons cognitive map may change over time. Typically, when we are new to an area our cognitive map will be linear but this will usually evolve to become increasingly spatial. Also, active exploration of an environment, as opposed to being guided thought it, encourages the formation of richer maps which are more likely to be spatial. For more detail on the formation of spatial knowledge Piaget and Inhelder [14] and Downs and Stea [15] are good starting points.

The Image of the City describes experiments carried out in a number of major US cities which suggest how cognitive maps are built up over time. The experiments involved obtaining information from long term inhabitants of the cities in the form of, for example, interviews, written descriptions of journeys through the city and drawn maps. By examining this data Lynch identified five major elements of urban landscapes which were identified by the inhabitants and used as the main building blocks of their cognitive maps. These features are:

- **Landmarks.** Static and recognisable objects which can be used to give a sense of location and bearing. Examples of landmarks might be prominent buildings or monuments or, on a smaller scale, recognisable shopfronts or roadside installations.
- **Districts.** Sections of the environment which have a distinct character which provides coherence, allowing the whole to be viewed as a single entity. Districts may be identifiable, for example, by the nature of the architecture of the buildings in the area or by their use.
- **Paths.** Major avenues of travel through the environment such as major roads or footpaths.

- **Nodes.** Important points of interest along paths, e.g. road junctions or town squares. There is clearly a strong link between paths and nodes.
- **Edges.** Structures or features providing borders to districts or linear obstacles. Examples might be the waterfront in cities built on large rivers, or a major road. Note in the latter case that the road may have a dual nature, being a path for someone travelling in a car but an edge to the pedestrian.

Urban environments also contain other cues for navigation which we might consider including in virtual worlds. In particular we are interested in the use of signposts. In city planning there is some contention over the use of signs as they are sometimes considered to be, to some extent, an admission of failure in the initial design of a navigable space [16]. However, other researchers point out the extent with which signs are integrated with the everyday process of wayfinding and so consider them an essential element in the urban landscape. We will return to the subject of signposts later on.

4. Legibility techniques for information visualisation

In this section we consider how the general notion of legibility and the five specific legibility features identified by Lynch might be adapted for use in information visualisation systems. The algorithms proposed here were initially presented in [17]; we represent them here for completeness before progressing to discuss recent applications and experimental evaluation.

We have stated above that the formation of cognitive maps is a learning process, so the spaces to which we apply legibility enhancements must allow learning to happen. This leads us to define the following criteria for spaces to which our techniques might most usefully be applied. First,

the space, or the visualisation, ought to persist over a long period of time. This is necessary in order that enough time is available for learning of the space to take place. Second, the information in the space must be relatively stable when considered in terms of the ratio between the number of changes occurring and the overall volume of information present. Two problems may arise if the data is constantly in flux; it may be difficult to actually produce and place the legibility features in a useful way; and even if we could they would quickly become irrelevant, either because they no longer related to the data or because they moved too frequently to act as useful reference points. Finally, the space should be available for users to re-enter and use repeatedly. This will allow learning over time to take place naturally.

The usefulness of legibility features then obviously depends on the whole data space being used in a browsing mode rather than viewing only pieces of data which match a search query or travelling directly to such objects in the space. We believe that a useful aspect of visualisation techniques is the ability to present data in the context of related pieces of information and hence that browsing is an important part of using such a system. Chalmers [10] notes that large information spaces are often difficult to access because the user has difficulty recognising an effective starting point for browsing the data. He suggests that searches can provide this ‘way into’ the information. It is our hope that by enhancing the process of learning the layout of the space the use of legibility features may also help users establish an initial starting point for browsing.

The following sub-sections now describe how each of Lynch’s five legibility features might be introduced into a general information visualisation. This order in which they are presented (districts, edges, landmarks and then paths and nodes) is significant as the placement of some of the features depends on others which

have already been generated.

We also describe how these proposed techniques have been implemented in a prototype system called LEADS (LEgibility for Abstract Data Spaces). LEADS has been constructed as an example “legibility layer”, an independent sub-system which, following some configuration and basic integration, is able to post-process the output of other visualisation systems in order to enhance their legibility. LEADS is referred to throughout as a concrete example of how our techniques can actually be realised. However, we urge the reader to bear in mind that other implementations are possible.

Returning briefly to the nature of the information visualisation space, most such systems can be very broadly classified into two types: those which map attributes of the data directly onto the axes of the space, in the manner of a scatter graph, or those which use some incremental algorithm to draw similar objects closer together. We would like to point out that the techniques we are describing here may be, and have been, applied to both types of space with equal efficacy.

4.1. Districts

The discovery of districts within the data is a matter of finding groups of items which have strong similarities to each other which they do not have with other items in the space. Bearing in mind the goal of the work which is to accomplish the placement of features automatically, the discovery of these groups is not a trivial procedure. Some of the problems that must be contended with are: how many of these groups exist? how do we measure (dis)similarity between data items? when do we stop trying to divide the groups that we already have?

In order to identify districts within an arbitrary data space we can use techniques from the field of cluster analysis, for which a number of algorithms

have already been developed. Such algorithms take collections of data and analyse them according to the strengths of similarity of their different attributes. The output a clustering algorithm is a set of discrete clusters of data items which we directly map onto the idea of different districts in a visualisation of this data. Once districts have been identified, they need to be represented in the visualisation. For example, one might use the attributes of colour and shape to give each district a distinct character (see below for examples). Districts also provide the basis for creating the remaining legibility features.

It should be remembered however that the prime purpose of the representation of districts in the information visualisation systems is to provide a segmentation of the space which can be used as a reference for navigation. Therefore, while it is desirable for the districts to have a semantic relevance the strict accuracy of the classification method need not be an overriding concern. We accept that this is a major research issue in itself and that there are a number of issues which could be considered, such as fuzzy boundaries between clusters and hierarchical classifications, but our priority has been to establish a clustering tool which can be applied to a large variety of spaces with little or no modification. Scope exists to extend and explore the use of other aspects of classification methods as further research.

In choosing an algorithm for the initial LEADS implementation, we set the criteria that it should be simple, due to time constraints on the project, relatively computationally inexpensive and reasonably effective for a wide range of clusters. We have initially adopted Zahn’s Minimal Spanning Tree algorithm (MST) [18]. This algorithm first constructs a minimal spanning tree of the data such that the edge values are taken to be the euclidian distance, or other similarity measure, between the items. Clusters are then produced by identifying and eliminating inconsistent edges, which are defined as edges of

the spanning tree whose values are significantly greater than the nearby edge values. For example an edge might be eliminated if its length more than twice the mean length of other edges within a certain number of steps or if it is more than n times the standard deviation of those edges *longer* than the mean. Useful values for n are usually between 2.0 and 3.0.

The basic algorithm as described by Zahn is able to detect clusters of varying shapes and sizes so long as they are relatively distinct (some of the specific cases where it will not perform accurately are described in Zahn's original paper). A specific advantage for LEADS is that the use of minimal spanning trees lends itself to cluster analysis of so called network structured spaces where one can identify explicit links between data objects (e.g. hyper-media structures or generalised object models). The drawback to this is that in spaces where no such graph structure exists it is necessary to assume a fully-connected graph structure, where each object is conceptually connected to all others. This means that the MST algorithm must make use of a complete matrix of distance values between the objects, making this step of the algorithm quite computationally expensive. Despite this, and especially in the case of network spaces, when compared to other clustering algorithms this method is relatively inexpensive computationally. Fast algorithms for producing minimum spanning trees are well known and the identification of inconsistent edges is relatively simple.

A potential problem with the use of a clustering algorithm to discover districts within the data is the coherence of the clusters in the space. If the clustering is based on a similarity measure which uses all n dimensions of the data, where n may be much larger than three, then when the data is mapped onto the three dimensions of the space then it cannot be guaranteed that the objects within each cluster will form a coherent group. In practice here we are relying on the

efficacy of the underlying visualisation system to provide a layout for the space which provides a high enough degree of proximity between similar objects. The systems to which we have applied LEADS so far have shown no obvious signs of cluster fragmentation. We have also considered other methods of dealing with this potential problem, such as basing the clustering on a weighted combination of the similarity based on the data and euclidian distance in the space, or using other clustering algorithms based on more highly spatial criteria.

4.2. Edges

The next feature we need to cover is the edge. These are (usually) imposing features which section off one area from another. It seems sensible therefore to define edges as existing between adjoining districts. The main problems to be solved with edges are: between which districts should they be placed? what shape should they be? how big should they be and how should they be positioned and oriented?

We will consider three possible methods for defining edges. The first is a quick and dirty method which will not be effective for all shapes of clusters but which might work tolerably well for those which are generally spherical or cuboid in shape. This simple approach is to find the nearest neighbour data items between the clusters using the same similarity measure as was employed in the initial clustering process. The edge can then be placed between the clusters along the line connecting these two items, with an appropriate orientation. Provided that the objects do not excessively cut into the clusters or are positioned far from their logical joining point these methods should be effective in providing a reference point and defining the borders of clusters.

The second method involves finding the hull of each district and creating an edge just beyond this.

This has a number of drawbacks. Most importantly finding the hull of the cluster of objects in 3D space is not a trivial task and would work out to be computationally expensive. Considering that the system will already be carrying out a large number of computations in order to carry out the clustering to produce districts the trade-off between complexity and accurate positioning of the edge objects becomes quite important, especially if the situation of having to completely redo the whole legibility process after a large number of changes to the space looks likely to occur often.

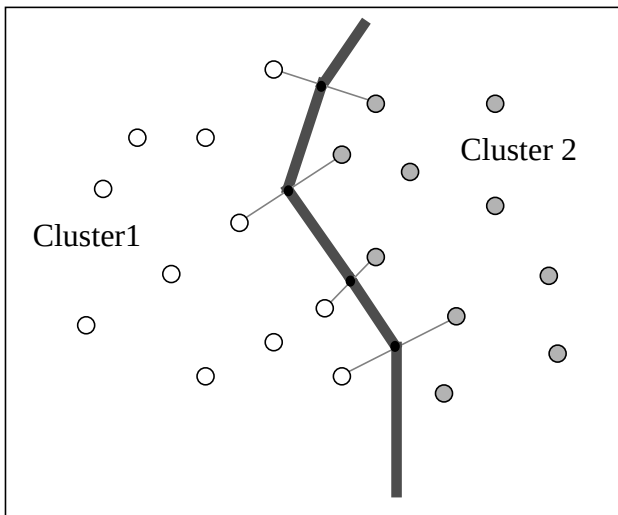


Figure 1: Formation of edge between districts by interpolation between pairs of adjacent data items

The third method is a refinement of this which again involves finding the hulls of the adjoining districts. Once this is complete the edge could be defined by interpolation between the points along adjoining edges of districts, as shown in figure 1. Sophisticated interpolation methods might produce a smoother edge topology at the expense of some time, as shown in Figure 1. One possible problem with this would be deciding which pairs of data items to use for the interpolation process. Items would need to be in the same general area of the space but problems might occur in making the choice, for example, where the items are densely

packed. It would also be necessary to consider whether objects may belong to more than one pair where the interpolation process is concerned.

Considering the problems that would come with attempting to find the hulls of all the clusters for the initial implementation of the LEADS system we decided to use the first method described here, positioning the edges between cluster nearest neighbours. This highlighted the problem of deciding the orientation and size of the planes used to represent the edges. In orientating edges the method used is to simply align the object's shortest axis with the line joining the spatially nearest neighbours. This provides a reasonable orientation in the majority of cases. The size of the edge is dictated by the size of the clusters which it separates. The philosophy used is to make the edge some proportion of the size of the smaller cluster in the direction of the two major axes of the edge. The proportion currently used is 80%, but this may be defined by a system resource. This means that the edge can never overwhelm the smallest cluster of the pair and that its size is dependent on its orientation to some extent.

4.3. Landmarks

Landmarks need to be stable points within the data space that can be used as a common reference point for navigation. We will illustrate three possible methods which we have devised for defining landmark positions.

The methods described here are based on the definitions of the districts and will therefore rely on these being quite stable so that the landmarks are not constantly moving as the database changes. This should not be a problem in the relatively slow moving type of space for which the system has been designed, although it must be expected that the landmarks will be subject to some drift over time as the database profile evolves.

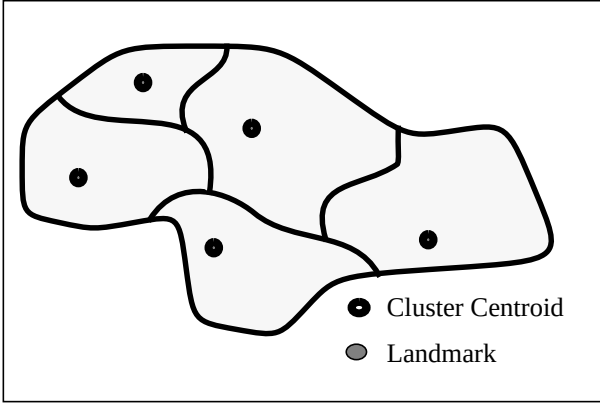


Figure 2: Cluster centroids as landmarks

Method 1: Cluster centroids (see figure 2). Most partitional clustering algorithms will define the centroid of a group while it is being formed. This centroid is a virtual data element that best describes the cluster as a whole. Our first method is to simply define landmarks to be the cluster centroids as defined within the clustering algorithm. If the algorithm used does not define centroids it should be possible to simply apply a centroid generation method similar to that which would be used with a clustering algorithm to the finished clusters. A slight variation of this method might be to place the landmark at a cluster's geometric centre. This method will therefore define a single landmark for each of the clusters which could act as a beacon for navigation to the heart of the district

The main question to consider is if this is the most suitable place for a landmark. The centre of the cluster may be densely populated with data items but on the other hand it could just as easily be rather sparse. In large clusters these landmarks may easily become lost in the vast group. We suspect that this method of placing landmarks might not take account of the concentration of clusters and data and so will not necessarily position landmarks in areas where external references were most useful. We also feel that the districts may already be sufficiently well defined

by their own character and edges so that reference points at their centres might not be essential.

Method 2: Cluster intersections (see figure 3). This second method places landmarks wherever more than two districts intersect. This is a simple method that requires identifying which districts are adjacent and finding a midpoint between the closest members of the districts. This method results in fewer landmarks than simply using the centroids but those landmarks which are created appear in the areas of the space where a number of clusters meet. This implies that these are areas where there are a number of clusters bordering on each other. It is in this sort of area where a stable reference point might help in navigation and so is a better choice for the placement of landmarks.

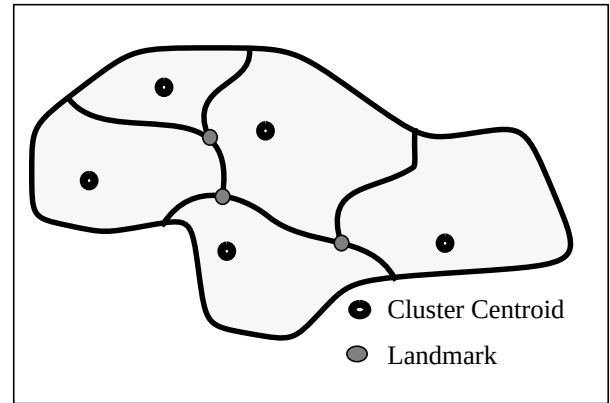


Figure 3: Cluster intersections as landmarks

Method 3: Cluster centroid triangulation. The final method again uses the centroids of the districts but will place landmarks at the centre of the triangles formed by the centroids of any group of three adjacent clusters. This is illustrated in figure 4 and will produce approximately the same number of landmarks as the cluster intersection method. The position of the landmarks, although seemingly quite similar to that produced by Method 2 is significantly different in that the landmark is placed in a central position between the points which most typify the districts rather than relying solely on the geometric point of intersection. We suspect that this will give a more

even balance to the placement of the landmarks than simply positioning them at the intersection points and will make the landmark positions more responsive to the shapes of the clusters as well as their positions. This is the method that we have implemented in the LEADS system.

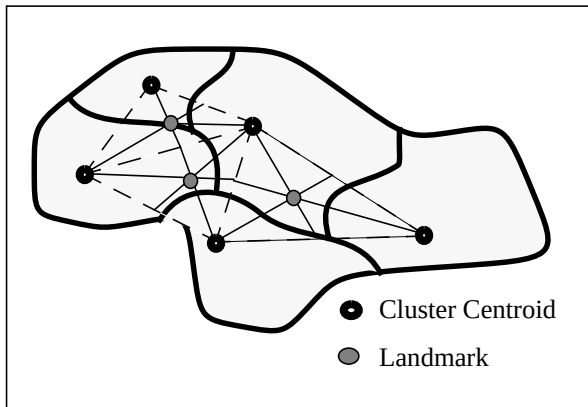


Figure 4: Landmark positions calculated using cluster centroid triangulation

It is anticipated that this final method might give rise to some drift in the position of landmark objects over time. However, if the space satisfies the criterion of relative stability over time this movement should not be enough to render the landmarks ineffective as reference points.

4.4. Nodes and paths

Nodes and paths are strongly inter-related and will therefore be considered together. We propose that paths should be composed of links between individual elements of the visualisation. Eventually, we intend paths to evolve as a function of use of the space with information about movement of users and database access being recorded for this purpose. Two main issues therefore need to be considered. The first is which, if any, elements will be used as initial nodes and paths, before any usage information has been gathered. The second is how the path layout will evolve with use. So far, we have implemented algorithms for the first of these and have

proposals for the second.

The choice of initial nodes and paths in an unused information space is a difficult problem. On one hand it will be useful to have these features available to aid initial legibility and imageability of the space while on the other hand we must be careful not to colour the usage of the space too much. Presuming that absolutely no usage information has been recorded, it will be necessary to make educated guesses at data elements which might possibly become nodes through usage. We have implemented the following approach in LEADS:

- The nearest neighbour elements between districts might represent good initial choices as these are the most similar items across district boundaries. We shall call such nodes gateways.
- An additional main node for each district might be defined as that data element closest to the centroid of the cluster as this might be seen as being the item most typical of the district. In network type spaces there may be further data available to identify such nodes automatically, such as connectivity information (valency) or measures of amounts of data stored, both of which imply importance.

The current LEADS implementation first identifies the gateway nodes and forms paths between nearest neighbour pairs in adjoining districts and then within districts all the gateway nodes are linked together with a spanning tree.

Considering dynamic usage information, we propose that the positions of node, path and gateway features should evolve in response to the way in which they are accessed. The type of statistics that will be used will be concerned with both the frequency of access to individual items and the sequence in which the items are examined. For example:

- The ability of a data item to potentially become a node might depend on the number of queries which were made to the item within a certain time period. For example, within each district

this metric might be the top $n\%$ of the frequently most used nodes. In network spaces we might also consider the number of connections which lead to and from the item, its *valency*, as this may also be considered as an indicator of the importance of the object in the space.

- The definition of a path link between two elements might depend on the number of times those elements were accessed in sequence. For example, we might define the set of path links as the top $n\%$ of possible links where n will be found by experimentation.

In LEADS, we produce two levels of path which reflect different levels of use within the space. Within the set of path edges some of those most used will be further highlighted to show the most popular routes between gateways. Special care must be taken however to evolve slowly from the set of initial paths and nodes. We do not wish to eradicate any legibility features that we have defined too quickly, but it is also desirable that paths and nodes do not remain in the space when they do not accurately reflect its usage. It will also be important to ensure the connectivity of paths to all gateway nodes.

Other features

Text is an essential feature of many visualisations. In information spaces, text may be used as an identification sign of an area or element, as a sign giving directions to other areas or as an annotation. In his book *Wayfinding in Architecture*, Passini talks at length about the nature, placement and content of signs [13] and we have adopted his philosophy where possible. This book contains an in depth examination of the wayfinding process and concludes that signs should be placed at points in the space where wayfinding decisions are necessary and also at regular intervals between these point where the distance is large as a reassurance to users. In an

urban environment a decision point would generally correspond to an intersection on a path so we are instantly provided with a possibility for sign positions in the information space.

In three dimensional information visualisations it is often difficult to maintain orientation, and even more difficult to apply this orientation to the data items in the space to recognise how their attributes are mapped onto the major axes. A useful addition to the space might be an optional representation of the axes which will display their orientation with respect to the user and a textual annotation indicating which attributes are mapped onto the axes in information spaces where this is relevant.

Finally, a simple mechanism for realigning the viewpoint with the major axes can easily be added to the VR visualisation system to allow users to easily correct unintentional rotations and to regain a familiar orientation.

In summary then, this section has proposed a number of techniques for generating or enhancing legibility features in information visualisations and has described how these have been implemented in the LEADS system. The following section now describes three example applications of LEADS to different pre-existing information visualisations.

5. Four example applications

We have applied LEADS to four different visualisation systems, three of which were locally available at the start of this research and a fourth which was developed externally. These are:

- The **FDP-Grapher** tool for visualising three dimensional network structures [19];
- The **VR-VIBE** system for interactive visualisation and searching of document databases and the World Wide Web [20];
- The **Q-PIT** 3-D scatter-graph visualiser [20]; and

- The **BEAD** document space visualisation system [10][11];

Between them, these four examples cover a range of different visualisation layouts which are broadly representative of the different approaches described in the literature. They also include examples of fully 3-D visualisations where the dimension of height is used in the same way as depth and width and also a landscape style visualisation where information is arranged mostly in a plane and the dimension of height enables the user to fly over the scene and to zoom in and out.

We will now focus on each of these applications in turn, describing the underlying visualisation system and demonstrating the effects of applying LEADS to its output. We will give a particularly detailed description of the FDP-Grapher example, showing images of the various different legibility features in isolation, as this was the example used for experimental work. We will provide somewhat shorter descriptions of the other examples, limiting images to more general before and after shots.

5.1. FDP-Grapher

FDP-Grapher is a 3-D tool for visualising arbitrary network structures. The underlying visualisation approach is based on the Force Directed Placement (FDP) technique where the nodes of the network are treated as masses and the links as springs [21]. Initially the nodes are placed randomly and the system then passes through repeated cycles of repositioning the nodes, based on the tension in the links, until a relatively stable formation is found. In addition, each link in the network can be given a weight which will alter the way in which the tension value is calculated. The resulting visualisation shows the network drawn in 3-D space such that densely inter-linked groups of nodes are positioned closely together.

FDP-Grapher might have many applications.

One of our current applications is to visualise regions of the World Wide Web; in other words, to produce the kind of overview diagrams proposed by Nielsen (see section 2). Users of this application are able to:

- see an overview of up to a hundred or so linked Web nodes defined by an initial position (specified by a WWW URL) and a link adjacency distance;
- navigate the resulting visualisation with six degrees of freedom;
- select nodes in order to either obtain summary information or to launch the Mosaic browser in order to inspect their contents; and
- grab nodes and reposition them in order to stretch out the visualisation;

The WWW is seen as an ideal underlying database for this example because it satisfies all of the criteria for an appropriate space as defined in section 4. The WWW is highly persistent and is re-used on a very regular basis by a large number of users. Despite its perceived chaotic nature, the WWW as a whole also satisfies our criterion for stability, as the number of significant changes made to its contents over short periods of time are very small with respect to the overall size of the database and the number of retrievals.

Figure 5 shows a number of different screenshots from FDP-Grapher being used to visualise a network of 239 nodes (the example network actually used for experimentation). Each screenshot is taken from the same point of view but shows different combinations of legibility features (LEADS allows the user to interactively switch different features on or off). Figure 5 (a) shows the basic visualisation with no additional legibility features. Figure 5 (b) shows districts distinguished by the colour and shape of their constituent nodes. Figure 5 (c) shows the placement of a number of landmarks within the visualisation and figure 5 (d) shows how those links representing key pathways are emphasised through changes in colour and thickness. Figure 5

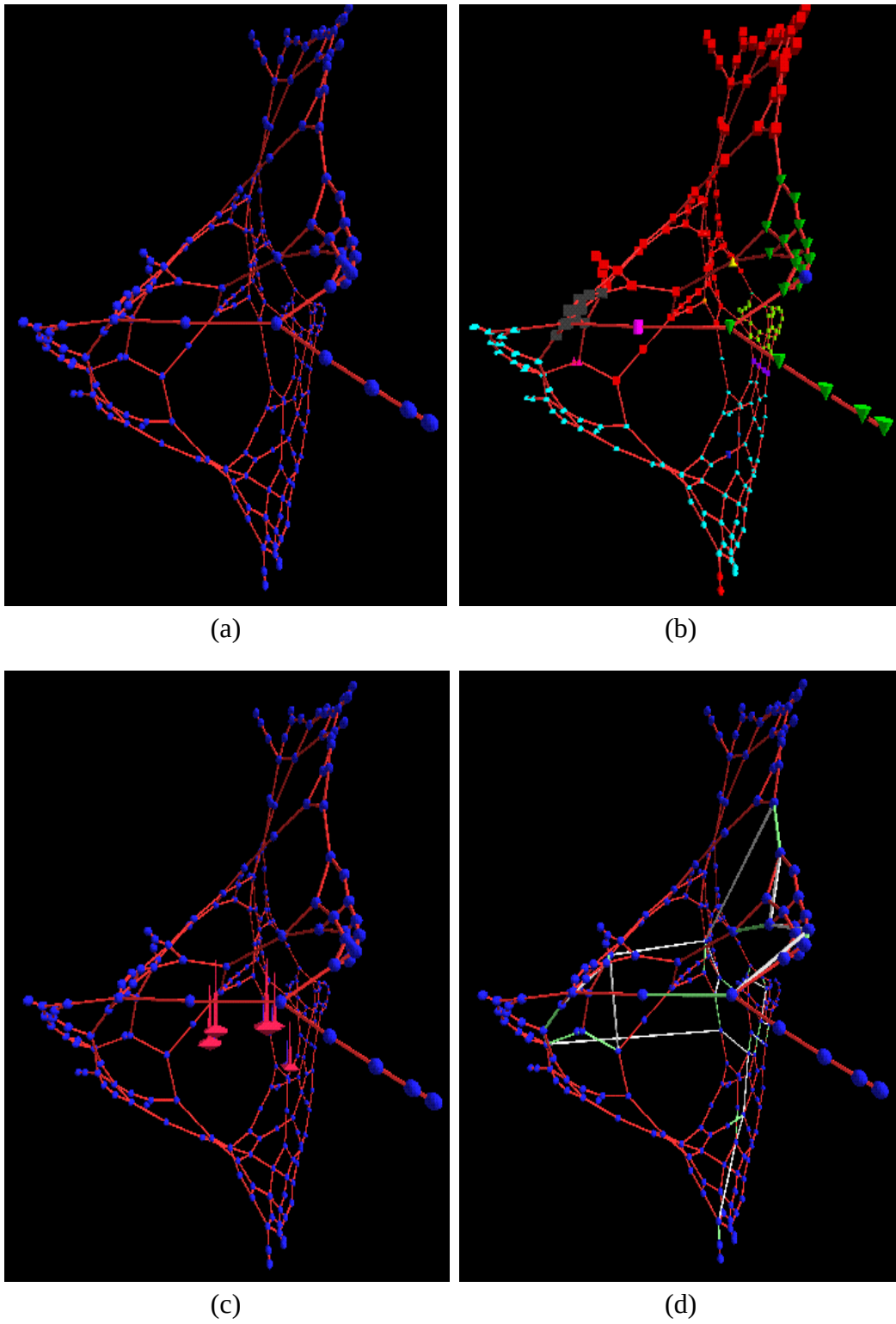
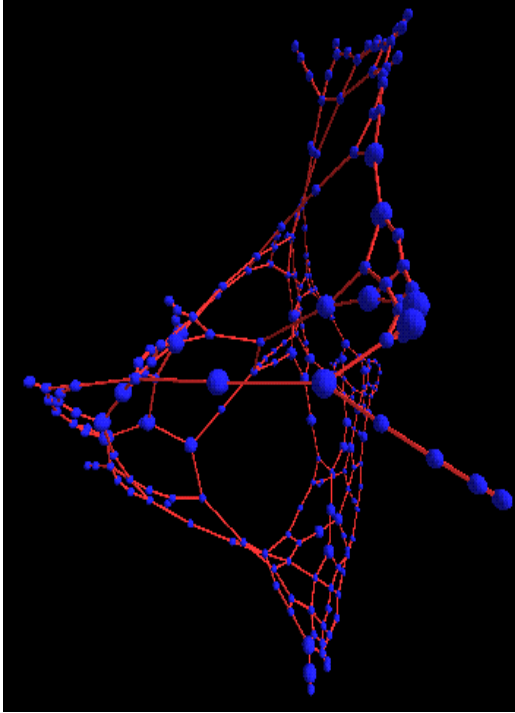
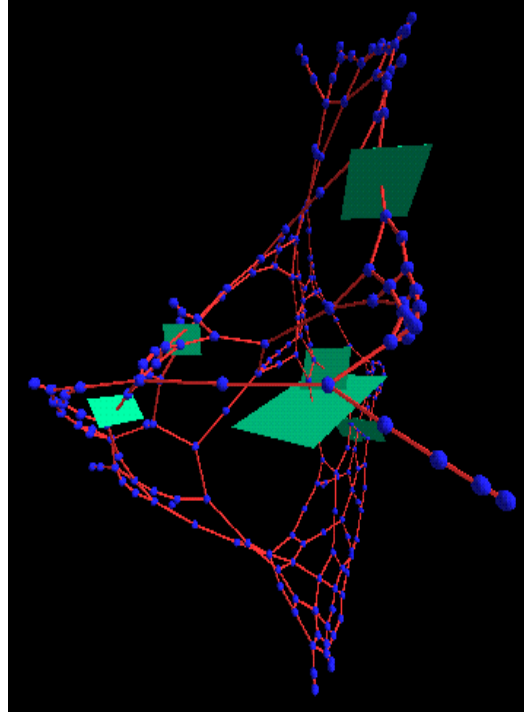


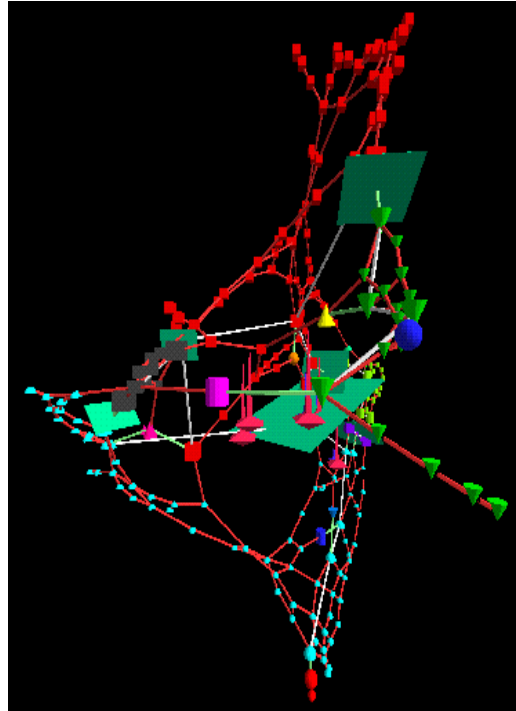
Figure 5: LEADS applied to FDP-Grapher showing individual legibility features:
(a) no features, (b) districts, (c) landmarks, (d) paths



(e)



(f)



(g)

Figure 5:(cont.) Spaces with individual legibility features: (e) nodes (indicated by increased size of data items), (f) edges, (g) all features together

(e) shows the emphasis of key nodes by increasing their size and figure 5 (f) includes a number of additional edge objects between the major districts. Finally, figure 5 (g) shows the visualisation with all of the features turned on simultaneously. When studying these images, it is important to bear in mind that they represent a distant perspective view and that the user is able to fly right into the centre of the visualisation, in which case they are surrounded by the nodes and links. Thus, what might appear to be a somewhat cluttered image from this distance becomes more open from closer in or even inside.

5.2. VR-VIBE

The VR-VIBE visualisation [20] supports information retrieval from electronic document repositories and is a three dimensional and interactive extension of the original 2-D VIBE as reported by Olsen *et al.* [22]. Unlike traditional text retrieval systems which only allow users to run a single keyword search at a time, VR-VIBE allows users to explore the effects of comparing and dynamically manipulating multiple simultaneous keyword searches. In essence, a number of keyword searches are defined, each of which consists of one or more text keywords. These are then positioned in a virtual space to form a spatial framework of queries. Document icons are positioned within this framework according to the strengths of their relative attractions to each query (i.e. the more strongly an individual document matches an individual query, the closer it is placed to it).

VR-VIBE users may dynamically interact with the visualisation in a number of ways: selecting documents displays summary details or launches Mosaic to view the document source if stored in the WWW (links are maintained from the local document repository into the Web); raising a relevance filter removes all documents whose overall score falls below a threshold value from the

display; grabbing and dropping queries dynamically re-arranges the space into a new configuration; switching queries on and off also changes the space and, finally, any number of new queries may be defined and positioned on the fly.

Figure 6 shows example screenshots of VR-VIBE before and after the application of LEADS. The first image shows a visualisation of approximately 1500 document references stretched out within a framework of five queries, currently positioned at the corners of an inverted pyramid. The second image shows the same visualisation from the same viewpoint but with the addition of legibility features.

5.3. Q-PIT

The Q-PIT visualiser draws three dimensional scatter graphs of tabular data. Q-PIT takes data which consists of records, each of which has a fixed number of typed fields, and maps these onto different display attributes of a three dimensional visualisation. These include the three spatial dimensions as well as representational attributes such as colour, shape, size and spin-speed. Q-PIT users are able to specify different mappings between the underlying data attributes and the visualisation attributes via a configuration file each time they launch the visualisation. Once the visualisation has been created, they are then able to inspect the contents of objects and also modify them, in which case they may see an animated movement of any changed objects to new locations in space. As with 2-D scatter graphs, there might be many potential uses of a Q-PIT style visualisation involving comparison and correlation of different combinations of attributes belonging to a set of data objects. Figure 7 shows before and after images of applying LEADS to Q-PIT (in this case a small database of some 130 personnel records).

An interesting issue raised by this example is the likelihood of LEADS overloading visual

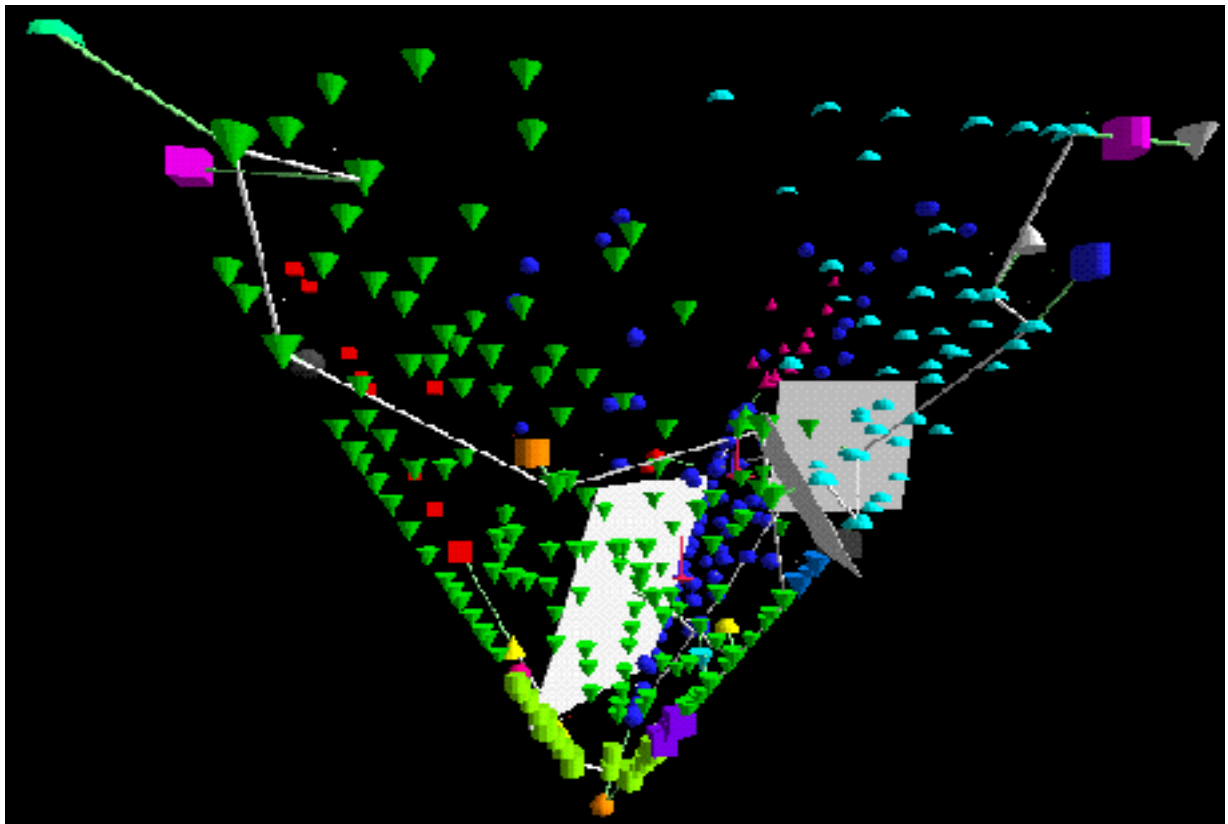
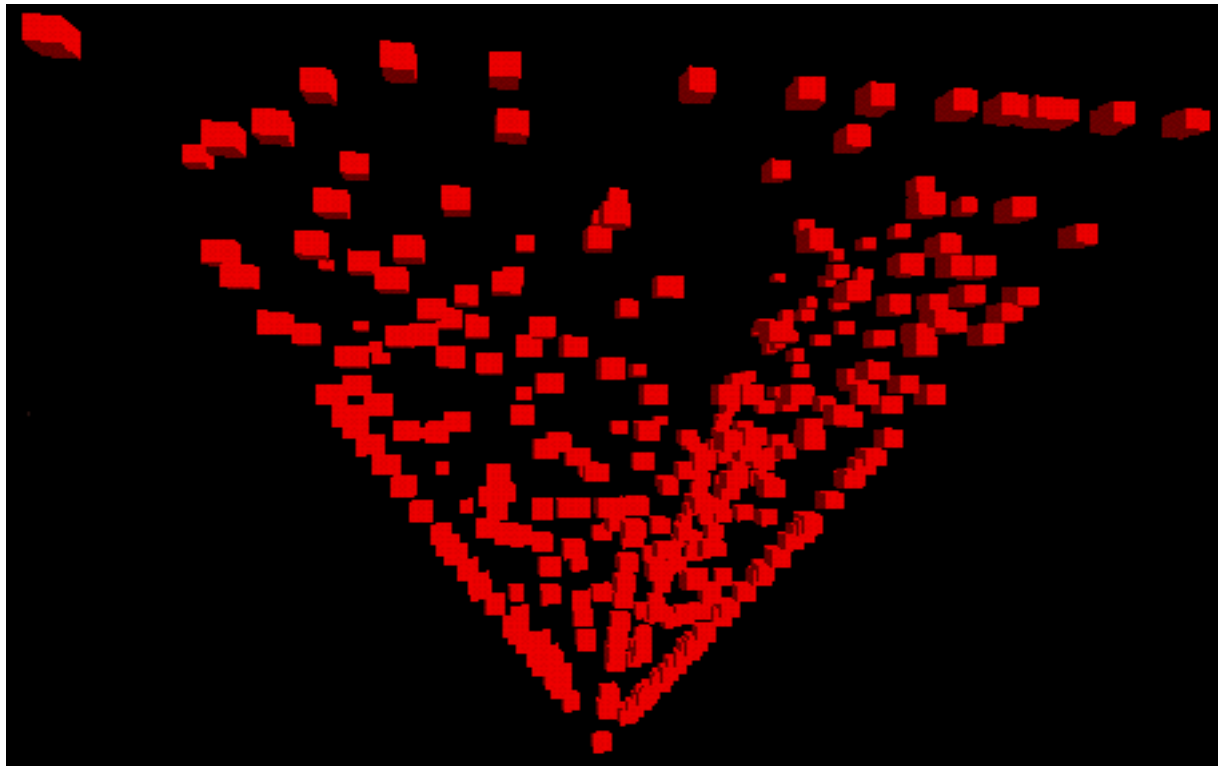


Figure 6: VR-VIBE before and after the application of LEADS

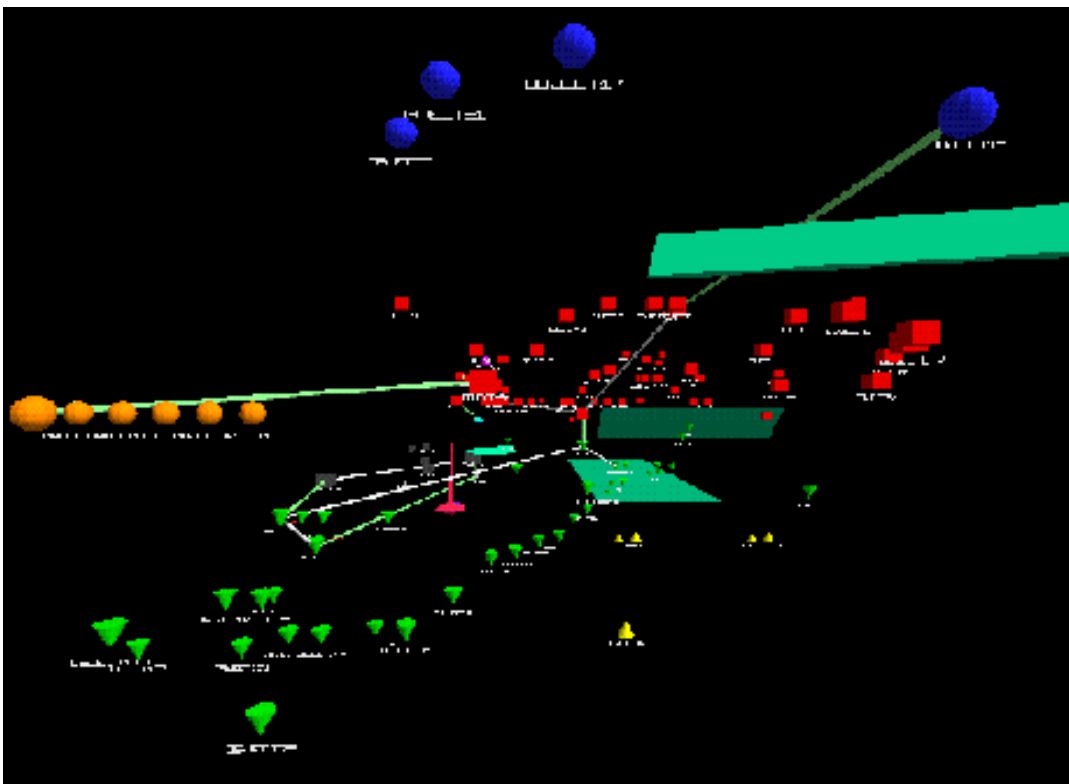
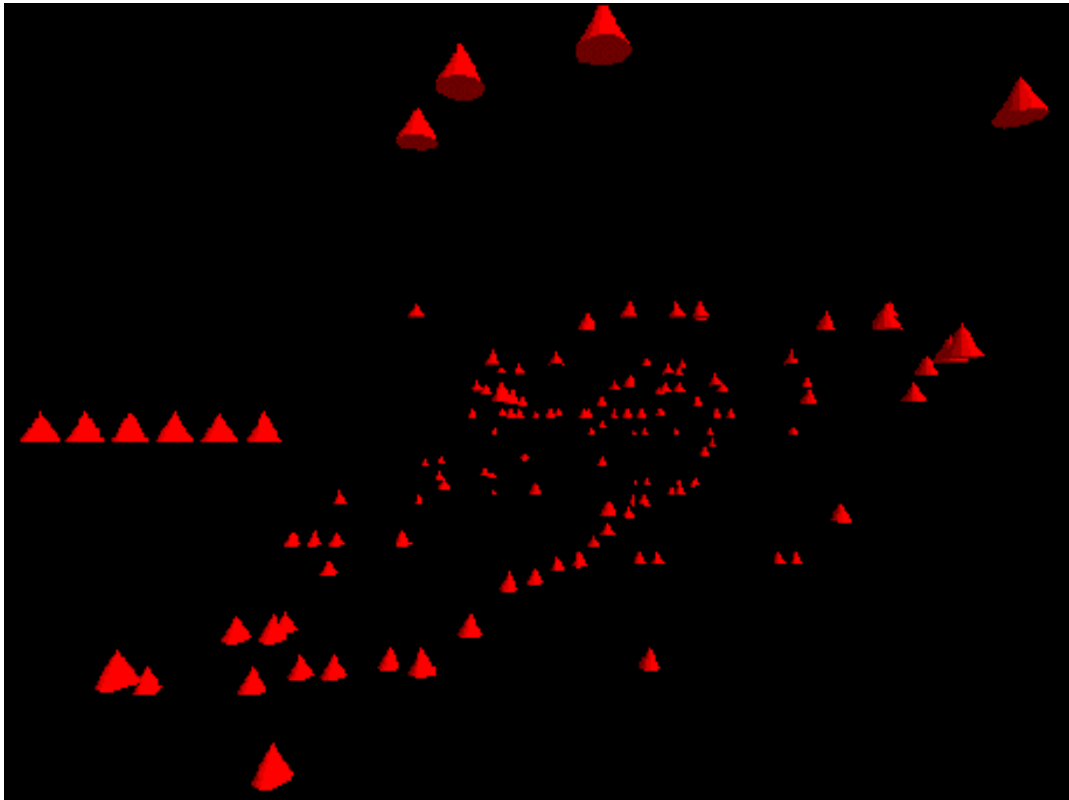


Figure 7: Q-PIT before and after the application of LEADS

display attributes currently used by Q-PIT resulting in possible confusion (e.g. LEADS using colour to distinguish districts when Q-PIT has already assigned it some other intrinsic meaning). Unlike FDP-Grapher and VR-VIBE, Q-PIT does not fix the meanings of its display attributes until run-time. LEADS therefore cannot guarantee that the attributes that it uses to emphasise districts, nodes and paths, or the objects that it introduces to represent edges and landmarks will not be confused with representations of the underlying data. This suggests that more careful consideration must be given when applying LEADS type systems to visualisations that make extensive and dynamic use of different display attributes. It also suggests that visualisations might usefully make such mappings explicitly visible (as Q-PIT does in its configuration file).

5.4 BEAD

BEAD is a visualisation system developed by Matthew Chalmers at Rank Xerox EuroPARC and continuing at the Union Bank of Switzerland. It was initially designed for the visualisation of large document stores but has recently been applied to other data such as groups of time series representing stock performance. The algorithms developed by Chalmers for the placement of data items in the space are based on annealing methods. The goal of the algorithm is the reduction of the total energy of the system by incremental movements of objects, where the energy value is provided by attractive/repulsive forces between items. The level of the force is calculated using a similarity measure between the items, for example word co-occurrence in document spaces, and their distance from each other, with similar objects being drawn towards each other and dissimilar ones forced apart.

A major feature of BEAD is that it uses the metaphor of a landscape to represent the space. To

achieve this the forces in the placement algorithm are skewed so that the space converges towards the x - z plane. When the energy of the system reaches a sufficiently low level the positions of the objects are used as the data points for a De Launay triangulation which generates the polygons of the landscape. The result of this is a single, well defined 'island' of data whose shoreline can be considered as an edge, in the manner of Lynch, and where similar objects should be drawn together into district like groups.

The main difference in the application of LEADS to BEAD spaces is that in this case the presence of the landscape object which is generated by the system provides scope for the representation of legibility features which is not present in the other, more generally three dimensional, visualisations considered so far. We took advantage of this to improve the cohesion and visual definition of the districts being displayed in the space. The previous examples have shown districts which were distinguished from each other visually by changing the shape and, in particular, the colours of the objects they contain. For the BEAD system the districts were shown by colouring the polygons of the landscape itself rather than the objects. As figure 8 shows this results in extremely good definition of the area of the districts. The polygons which fall between districts (as is inevitable when using a De Launay triangulation) are given a neutral colour. In this way the colouring system also implicitly defines the edges which emphasise the district borders.

6. Experimental Evaluation

In this section we describe an initial experiment which was carried out to evaluate the effectiveness of the LEADS implementation of our proposed legibility techniques. The aim of the experiment was to test the way in which its

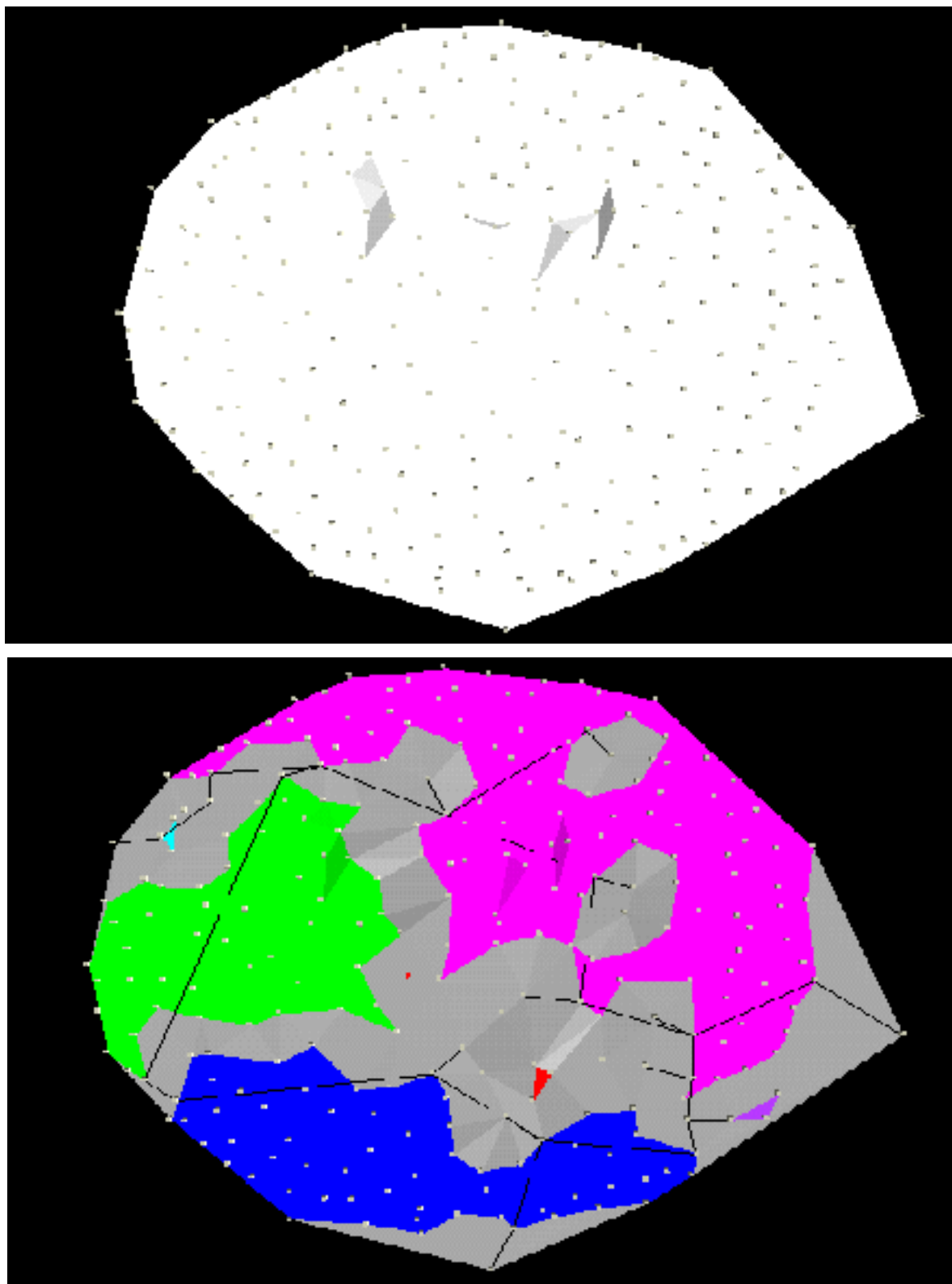


Figure 8: BEAD before and after the application of LEADS

subjects learned the layout of an information space which they entered and used repeatedly. To facilitate this the subjects were asked to complete a search task a number of times in the same space. Separate experiments were carried out in spaces with and without the LEADS legibility enhancements. The experiment only considered the addition of all five of the legibility features together and not the effects of individual features or smaller combinations of features.

The search task was to find five named objects in a graphical space of approximately 240 nodes. The graph used was randomly generated but not so random that objects of similar numeric value are spread throughout the space. This meant that the search task was not entirely one of finding a needle in a haystack. The subjects were required to complete the search task three times in the space of two days. The data items for the search were chosen to be spread quite evenly through the space.

From this we can identify that the independent variable in the experiment was the level of legibility information added to the basic space. The dependent variable was the time taken to complete the search task, which was dependent on how well the positions of the objects were learned.

Due to the nature of the environment within which the work was carried out the best source of subjects for the experiment was the students within the Department of Computer Science at the University of Nottingham. Apart from the simple availability of people this did have another advantage in that it limited the amount of variation between the subjects quite significantly. They all had a similar standard of education and of experience with computers and databases. The subjects were divided into two groups and allocated randomly to each of the two tasks.

Two types of result were gathered from the experiments: statistical, from the times taken to find each object; and anecdotal from observation

of the subjects and through questionnaires completed after the final attempt at the task. Because the time available for initial experimentation was limited and hence the size of the sample was small (six people in total) the statistical results cannot be taken as statistically significant. However, in conjunction with the other results we believe that they can provide a useful initial evaluation to guide further development and more detailed evaluation.

6.1. Measured results

These results are based on the times taken by subjects to find each of the five target objects and hence the time taken to complete the task as a whole, if achieved. Non completion of the task was due to an absolute limit of forty-five minutes on each attempt which was seen as a reasonable time to avoid undue boredom and frustration in the subjects.

Table 1 shows the absolute times taken for subjects to find each object (in the correct order) during the trials. The numbers in the left hand column indicate the number of the subject and of their attempt at the task. Subjects S1, S2 and S3 were using the space without any legibility enhancement. The remaining subjects used the space with all LEADS features added. The remaining columns each represent one of the five target objects. Dark shading of the table cell indicates that the target was not found. The main observation from the data is that the rate of task completion on the second and third attempts amongst users of the space with enhancement was higher than for those without. Users of the space with enhancements seemed to complete the task almost trivially on the third attempt while those in the raw space still had problems. This would seem to imply that greater learning of object positions did take place with the aid of legibility features.

Trial	T1	T2	T3	T4	T5
S1.1					
S1.2	9:00	12:30	25:00	39:00	
S1.3	19:00	22:00	41:00	42:30	
S2.1	32:00				
S2.2	34:00				
S2.3	9:00	18:00			
S3.1	8:0				
S3.2	1:30	4:00	15:00	21:00	32:00
S3.3	10:00	17:00	24:00	25:30	29:00
S4.1	16:00	37:00	44:00		
S4.2	2:00	4:30	12:00	15:00	31:00
S4.3	0:30	5:00	9:00	34:00	
S5.1	2:00	7:00	17:00	21:00	34:00
S5.2	0:30	2:00	2:30	10:00	14:00
S5.3	0:30	6:30	8:30	12:00	18:00
S6.1	13:00	33:00	34:30	45:00	
S6.2	4:00	6:30	7:00	10:00	11:00
S6.3	13:00	14:00	15:00	16:00	17:00

Table 1: Absolute times taken by subjects to find search objects

Another interesting point is that the mean time taken to find individual objects was consistently smaller on *all* attempts for the users of the space with legibility enhancements. These users seemed to gain some immediate advantage from the legibility features, without the chance for learning to take place. One explanation for this might be that the existing partition of the space into districts lead to the use of a more structured searching technique even before the space has been learned.

6.2. Results from questionnaires and observation

One of the first things that becomes clear from viewing the subjects performance was the importance of the way in which the volunteers moved through the space. There seemed to be a great variety in this despite the small sample used and the way they moved seemed to have an influence on the search methods used. The most effective search technique seemed to be one of

flying through strips of the space and gathering information on the numbers of the objects contained there. In contrast to this one subject adopted a technique of moving directly to individual objects, looking in the immediate vicinity and then pulling right back from the space before moving to another object in a different section of the space. This seemed to waste a great deal of time and did not allow him to observe a sufficient number of objects to be effective for the task, and this was reflected in his performance.

In general the subjects using the space with legibility enhancements seemed to remember the positions of the objects more effectively. They commented that the main aid to remembering the position of the objects was the colour (and to a lesser extent shape) of the districts. Although these were the main feature used as a memory aid from the observation it was clear that other features were also being used. One example was an object which was near to the largest Edge, which appeared as a green plane in the visualisation. When searching for this target the subjects often checked the objects along the line of this plane. Subjects using the space without legibility enhancements would often show signs of remembering certain attributes of the area of space the object was in but not the definite position.

The questionnaire contained 15 closed and five open questions which fell into five areas designed to gather information on different aspects of the user's experience of the space. These categories were: orientation, (learning) position of target objects, movement, system performance and use of other features (e.g. text labels.) Table 2 lists the closed questions and the answers given.

The responses to the section regarding orientation in the space are somewhat ambiguous. On the general question of whether the subject felt disoriented in the space the surprising answer is that the group without legibility enhancements felt less disoriented. This may indicate that the

extra information added to the space could possibly add to initial overload of information. However, on the related question of whether the users felt most disorientation within the group of objects or outside it (i.e. gaining an overview) the subjects from the group with enhancements seemed more comfortable with the latter, a view of the whole space. We might surmise then that the divisions presented by the clustering into districts are of most use on the global scale of the space and that more effective local cues for orientation are necessary.

The greatest agreement among the volunteers came in the section on the use of text. All subjects had text visible almost all of the time. The subjects also had strong agreement that the ability to align the labels towards the current viewpoint position was very useful. This seemed to be reflected in the number of times that the feature was used during the experiments.

The most important section of the questionnaire is that on the way in which the subjects learned the positions of target objects. In only one of these ("During the third visit I had no impression of the overall structure of the space.") did the subjects show significant agreement in the answers given. The responses of both groups indicated that they had in general gained a good impression of the overall structure of the space.

However, the questions where the answers did differ are more interesting. The first of these was asking if the subjects knew where the search targets were when they attempted the trial for the third and final time. The subjects using the enhanced space were much more positive that they had learned the absolute locations of the objects than those without enhancements. The other two questions were testing whether the subjects knew the positions of the targets relative to each other. The subjects with legibility features were again more positive that they knew the relative direction of their next target object and that they had learned routes between the objects.

Statement	Group 1			Group 2		
I did not feel disoriented in the virtual space	2	1	4	4	4	2
I made use of the axis object to check my orientation (1=Never, 5=Always)	3	2	1	3	3	2
I considered space to have a fixed top and bottom	5	4	2	2	5	3
I felt less disorientated outside the group of objects then when in the centre of the group	2	3	5	1	2	2
During the third visit to the space I knew where the search objects were	4	3	5	3	2	2
Navigation in the world was straightforward	2	2	4	2	3	1
During the third visit I had no impression of the overall structure of the space	4	2	4	4	5	4
I easily learned routes between different search objects	4	4	5	2	3	2
During the third visit I did not know in which direction to move to reach the next search object	4	2	2	5	4	5
I preferred moving freely through the space to following links	1	1	5	3	1	1
The response time of the system was adequate for my participation	4	4	2	3	4	2
I found it useful to be able to switch between the two modes of movement	2	1	2	1	1	5
I often wished to move more quickly through the space	3	2	2	2	2	4
I rarely aligned the text towards my position in the space	5	4	5	4	5	5
I often had the text visible	(5)	2	1	1	1	1
Were you using the space with or without different coloured areas? (1=With, 2=Without)	2	2	2	1	1	1

Table 2: Results from closed questions. Group 1 contained the subjects using the space without legibility enhancement, Group2 the subjects using the space with all features added. Answer 1 = strong agreement, 5 = strong disagreement, unless stated.

The results from this final set of questions strongly indicate that the volunteers who were using the space with legibility enhancements had greater success in learning the positions of the objects and relating this to their current position in the space. This meant that they knew both the location of the objects and how to reach them. While the subjects using the space had an equally good impression of the structure of the data they were not so confident of the positions of the objects. This therefore provides some support for the hypothesis that the LEADS legibility enhancements are useful in helping users in learning the layout of abstract data spaces.

The open questions highlighted the problems of using text in three dimensional spaces, such as occlusion by other objects and alignment. This seemed to be the major obstacle to efficient use of the system.

Subjects using the space with legibility enhancements also emphasised again the effectiveness of the districts in helping them to quickly return to objects they had discovered previously.

In summary then, the experiment seems to provide some initial evidence for the effectiveness of the legibility enhancements for allowing the users to learn the space. While the emphasis was on the use of districts the other features may have played a supporting and reinforcing role to increase the overall effect.

7. Conclusions, further reflections and future work

This paper has proposed a number of general techniques for improving the legibility of three dimensional information visualisations so that their users might more easily learn to navigate them. This work has adapted well established techniques from the discipline of urban planning and, we argue, points in a more general way to a

potentially useful relationship between these different fields. The primary goal of our work has been to develop a set of algorithms for automatically creating or enhancing legibility features within information visualisations. These include:

- the use of clustering algorithms to create districts;
- the creation of edge objects separating districts;
- the placement of landmark objects at a central point between districts; and
- emphasising nearest neighbour node objects within districts and creating paths between them.

We have implemented these techniques in a prototype system called LEADS which is intended to provide an additional legibility layer sitting on top of current information visualisations. So far, we have applied LEADS to four existing and contrasting information visualisation tools. The results of initial evaluation work suggest that this approach is promising and that it warrants further investigation.

We now conclude this paper by offering some additional observations which have arisen from the process of designing and implementing the LEADS system. We will also propose some broader directions for future research. We begin with some specific technical issues.

First, the idea of introducing legibility features such as paths seems to sit uncomfortably with six degrees of freedom navigation. Paths in a city are actually travelled along and the environment is therefore experienced from the perspective of the path. We suspect that users of our legible virtual environments should also directly experience paths as part of navigation. Thus, we have recently incorporated a simple interface technique to support navigation via paths, so that selecting a path causes the user to traverse it to its destination.

Second, automatically determining an

appropriate scale and appearance for legibility features has not always been easy. In particular, features such as landmarks and edges must be visible without being intrusive. Creating useful edges has proved to be a particularly difficult task as an edge should ideally follow the contours of a district. In a 3-D space, determining a sensible size for edges has been difficult. At one extreme an edge might be a hull completely surrounding a district. At the other it might be a thin flat surface dividing two districts. The former is likely to be visually intrusive; the latter may provide an insufficient sense of separation between districts.

Third, other features and tools are clearly needed to help people navigate. For example, the use of textual information in the form of signposts is an important part of navigating conventional urban environments. In the field of city planning signposts are often considered to be something of an admission of failure [16], that the planner was unable to produce a sufficiently legible environment, but we concur with the view of Passini that they are in fact a useful, even invaluable tool during wayfinding tasks. For this reason we are endeavouring to include signposts in our environments where possible. Currently these appear attached to the paths near to the node objects and identify the item to which the path leads. However, this gives rise to the problem of how to name districts and landmarks. More specifically, given that districts and landmarks are automatically created by LEADS, we are left with the problem of automatic name generation or alternatively the automatic creation of symbolic identifiers on signposts.

Fourth, we note again the problem of overloading visual attributes as discussed in section five. Ideally, LEADS ought to be able to query a visualisation application to find out what mappings or representations it already uses, in order to choose non-conflicting ones for legibility information. In turn, this suggests that visualisations ought to make their mappings

available to other applications in some explicit form.

Fifth, we need to be careful that by adding additional objects to an information visualisation, we do not increase the rendering overhead thereby degrading system performance. However, we suspect that some LEADS extensions might lead to improved system performance. For example, LEADS might enable the development of a distancing technique for data spaces whereby all of the individual data items in a district would be replaced by a single object when viewed from a distance.

We also need to instigate a more detailed program of experimental work in order to more rigorously test our hypotheses and also to tease apart the effects on navigation of different legibility features and of different techniques for enhancing or creating them.

Adopting a more general perspective, this paper provides one example of how the field of urban planning can inform the design of three dimensional information visualisations. However, Lynch's work on legibility, although very well known, is not the only research in this field that might be adapted to this purpose. In particular, work on *the social logic of space* [23] has considered how urban form relates to both navigation and social action. Thus, as opposed to focusing on the relationship between a single user and an information visualisation, we might instead consider how to support shared and cooperative access to such visualisations as proposed in recent work on Collaborative Virtual Environments and Populated Information Terrains [20]. Some work into this topic has recently been carried out by John Bowers as part of the European COMIC project resulting in a program called the Virtual City Builder which constructs virtual cities based on underlying principles proposed by Hillier and Hanson. One of the main aims of our future research is to explore the integration of this work with our own

Legibility based techniques in order to gain a more complete view of how urban planning techniques can support the design of future visualisation systems. An initial discussion of the relationship between these approaches can be found in [24].

8. References

- [1] Elm, W. C. and Woods, D. D. (1985) Getting Lost: A Case Study in Interface Design. *Proceedings of the Human Factors Society*. pp. 927-931.
- [2] Nielsen, J. (1990) The Art of Navigating Hypertext. *Communications of the ACM*, **33**, 3. 297-310.
- [3] Nielsen, J. (1990) *Hypertext and Hypermedia*, Academic Press, San Diego, CA. ISBN 0-12-518410-7.
- [4] Nielsen, J. (1995) *Multimedia and Hypertext: The Internet and Beyond*. Academic Press, San Diego, CA. ISBN 0-12-518408-5.
- [5] Mackinlay, J. D., Robertson, G.G. and Card, S. K. (1991) Perspective Wall: Detail and Context Smoothly Integrated, In *Proc. CHI'91, Human Factors in Computing Systems*. ACM SIGCHI, May 1991, pp 173-179.
- [6] Card, S. K., Robertson, G. G. and Mackinlay, J. D. (1991) The Information Visualiser, an Information Workspace. In *Proc. CHI'91, Human Factors in Computing Systems*. ACM SIGCHI. May, 1991. pp 181-188.
- [7] Erikson, T. (1993) Artificial Realities as Data Visualisation Environments: Problems and Prospects. In Alan Wexelblat (ed.) *Virtual Reality: Applications and Explorations*. Academic Press Professional, Cambridge, MA. ISBN 0-12-745045-9. pp. 3-22.
- [8] Fairchild, K. M. (1993) Information Management Using Virtual Reality Based Visualisations,. In Alan Wexelblat (ed.), *Virtual Reality: Applications and Explorations*. Academic Press Professional, Cambridge, MA, ISBN 0-12-745045-9, pp. 45-74.
- [9] Darken, R. P. and Sibert, J. L. (1993) A Toolset for Navigation in Virtual Environments. In *Proceedings of UIST'93*, November 3-5, pp. 157-165.
- [10] Chalmers, M. (1993) Using a Landscape Metaphor to Represent a Corpus of Documents. In *Proc. European Conference on Spatial Information Theory*, Elba, September. Published as: Frank, A. and Campari, I. (eds.), *Spatial Information Theory*. Springer Verlag. pp. 377-90.
- [11] Chalmers, M. (1995) Design perspectives in visualising complex information. In *Proc. IFIP 3rd Visual Databases Conference*, Lausanne, Switzerland, Chapman and Hall.
- [12] Lynch, K. (1960) *The Image of the City*. MIT Press, Cambridge, MA.
- [13] Passini, R. (1992) *Wayfinding in Architecture*. Van Nostrand Reinhold.
- [14] Piaget, J. and Inhelder, B. (1969) *The Psychology of the Child*. Routledge and Kegan Paul. First published in French as *La Psychologie de L'enfant*. Presses Universitaires, Paris. 1966
- [15] Downs, R. M. and Stea, D. (eds.) (1973) *Image and Environment: Cognitive Mapping and Spatial Behaviour*. Aldine Publishing Company, Chicago.
- [16] Canter, D. (1984) Way-finding and Signposting: Pennance or Prosthesis. In Easterby and Zwaga (eds.), *Nato Conference on Visual Presentation of Information, Information Design*. Wiley.
- [17] Ingram, R. J. and Benford, S. D., Improving the Legibility of Information Visualisations. In *Proc. IEEE Viz'95*. Atlanta, US. November 1995, IEEE Press.
- [18] Zahn, C.T., Graph-theoretical Methods for Detecting and Describing Gestalt Clusters. *IEEE Transactions on Computers*. **20**, 68-86.
- [19] Benford, S., Bowers, J., Fahlén, L. E., et al. (1995) Networked Virtual Reality and Co-operative Work. *Presence: Teleoperators and Virtual Environments*, **4**, 4, 364-386.

- [20] Benford, S., Snowdon, D., Greenhalgh, C., et al. (1995) VR-VIBE: A Virtual Environment for Co-operative Information Retrieval. *Computer Graphics Forum*. **14**, 3. 349-360.
- [21] Fruchterman, T. M. J. and Reingold, E. M., (1991) Graph Drawing by Force Directed Placement. *Software Practice and Experience*, **21**, 11, 1129-1164.
- [22] Olsen, K. A., Korfhage, R. R., Sochats, K. M. et al. (1993) Visualisation of a Document Collection: The VIBE System. *Information Processing and Management*, **29**, 1, 69-81.
- [23] Hillier, B. and Hanson, J. (1984) *The Social Logic of Space*. Cambridge University Press, Cambridge.
- [24] Ingram, R. J., Benford, S. D. and Bowers, J. M. (1996) Building Virtual Cities: Applying Urban Planning Principles to the Design of Virtual Environments In *Proc. ACM Conference on Virtual Reality Software and Technology (VRST'96)*. Hong Kong, July 1996. pp 83-91. ACM Press.