

Seamless Integration of Diverse Data Types into Exploratory Visualization Systems

Anilkumar Patro,^{1†} Matthew O. Ward¹ and Elke A. Rundensteiner¹

Computer Science Department, Worcester Polytechnic Institute Worcester, Massachusetts, USA

Abstract

Visual data exploration involves presenting data in some graphical form and allowing the human to interactively gain insight into the data. Visual data exploration techniques have proven to be very useful for exploratory data analysis, especially for mining large databases. However, the lack of explicit content semantics within the framework of the visualization tool makes interpretation of the visualizations difficult at times. The aim of our research is the development of a comprehensive framework for integrating semantics of data into an existing visualization tool in a flexible yet minimally intrusive manner. The first stage of the framework involves a transformation process that enables the mapping between a wide range of data types to a form the visualization system can process. This stage exploits the strengths of XML for handling semantics extraction, flexible data modelling of domains and managing the data mapping rules to conform the data to the requirements of the visualization tool. The second stage of the framework enables the visualization tool to access and present the semantic information separate from the data display. We demonstrate and validate this framework by illustrating the extension of XmdvTool, a multivariate data visualization system, to handle varied data types such as categorical fields and unstructured text files. We believe the framework is sufficiently powerful to incorporate other forms of data, as well as to be used to extend other visualization tools, such as OpenDX.

1. Introduction

Visual data exploration is concerned with exploring data and information in such a way as to promote a deeper level of understanding and foster new insight into the data, relying on the humans' powerful ability to perform visual information processing [ODR*99]. In a simple sense, visualization is essentially a mapping process, taking data attributes and mapping them to representations that can be perceived. Humans can often easily detect patterns, trends and outliers in the underlying data when presented with visual depictions of the data, which may be more difficult to identify with automatic techniques.

The data is the key component of the visualization and it plays a large role in determining the effectiveness of the visualization tool. The usefulness of a data visualization tool might be limited by:

- The amount of data that can be handled by the tool.
- The preprocessing of the data that is required to convert the data from its original format to a form the visualization tool can process.
- The different types of data that can be handled by the tool.
- The limited presentation of data semantics within the visualization.

Even though visualization systems have tried to tackle the problem of displaying large datasets [And72] (for example, pixel-oriented displays [KKA95] and hierarchical parallel coordinates [FWR99]) effectively, the other issues have been given less attention. This paper tries to address the last two issues in an existing visualization system.

Most of the current software systems for graphically presenting data are not built to handle a wide range of data types and hence they tend to be special purpose. For example, Parallel Coordinates [ID90] for numeric variables, Mosaic Displays [Fri94] for nominal data, and ThemeScapes [WTP*95] for text corpus are few exemplars of many such techniques.

[†] This work is supported under NSF grant IIS-0119276.

The reason for the tools being special purpose could be that the data that the visualization system uses to generate graphics is primarily numeric. This is because numeric data can be easily mapped to graphical attributes, which are usually themselves ordinal in nature (e.g. size, position, intensity). However, a great deal of data can be of other types such as nominal / categorical data. The problem with visualizing this in traditional ways is that the ordinal nature of the graphical attribute may introduce errors in the interpretation of the visualization.

Data sets have structure, both in terms of the means of representation (syntax) and the types of interrelationships within a given record as well as between records (semantics). Semantics of the data may be as important as the data for particular domains. Most visualization systems can pictorially represent the structure, but fail to capture the semantics of the data. For example, the most common way to represent nominal variables is by way of including legends for the variables. Although this can be effective in some situations, it cannot be generalized to other data types and might not scale to the amount of data in the dataset. Inclusion of semantics might be important in situations where the user who provides the data is different from the user who analyzes the data. This scenario can make the process all the more difficult because the analyst may not know the true meaning of the data. And since there is no means of passing semantics to the visualization tool, the knowledge within the data is completely lost in the visual mapping process.

In this paper, we propose a framework for including heterogeneous data types and embedding data semantics into the structure of an existing visualization tool while attempting to minimize the modification of the tool itself. This would surely increase the utility value of the tool, enabling its use in various application domains. To realize this, one basically needs an embedding of the data semantics within the framework. This can be achieved by (a) capturing the description of the data, (b) storing the mapping between data attributes to numerical values and (c) driving the display of the particular visualization with this knowledge.

To capture the description of the data, we associate metadata with the original data. This metadata can provide information such as the format of the individual fields within the data records, which helps in its interpretation. It may also contain the base reference point from which some of the data fields are measured, the units used in the measurements, the symbol or number used to indicate a missing value, and the resolution at which measurements were acquired. Each of these may be important in selecting appropriate preprocessing operations and setting their parameters.

Storing the mapping information with the metadata itself can make the integration of different data types seamless within the system. This allows the framework to work with an existing visualization tool. The flexibility of having different mapping types will make the process more powerful,

and help to extend the tool to many application domains. One such domain is the area of text visualization, an emerging area, where a mapping from structured or unstructured text files to the visualization tool-compatible format would enable the same tool to handle this kind of data.

The data semantics can be displayed within the existing visualization by accommodating the metadata into the display. This would require minimal recoding of the input and display modules to integrate the proposed framework.

The remainder of the paper presents the framework as follows. Section 2 presents work by others that is related to ours. Section 3 presents the proposed framework. Section 4 provides details about the data mapping process that is one of the main components in the transformation process. Details about the other components of the framework are discussed in Section 5. The different encodings for the mappings are discussed in Section 6. We validate our ideas using two case studies in Section 7. Section 8 contains our conclusions and the future work.

2. Related Work

Visualization tools do exist that deal with different data types. But little attention has been given to the issue of handling data management and data semantics.

Visualization systems such as OpenDX [Vis03], AVS [Adv03], and XmdvTool [War94], by default, make use of single precision floating-point values for the visualization of multivariate data sets. These are some of the systems that would surely benefit by the addition of data semantics within the tools.

Database oriented visualization systems such as TreeViz [Joh92], the VisDB system [KD94], the DeVise tool [LRB*97], Polaris [SH00] and SGI's MineSet [Rat96] do handle numeric and nominal types. But these systems just assign some order to the nominal variables and treat them as ordinal types. Statistical Data Analysis packages including SAS [CCS97], S Plus [KO02], and XGobi [SCB98] do allow one to specify and manipulate the mappings from nominal to numeric, but the mappings are completely embedded within the tool and not extensible to allow for other data types.

Research on providing general visualizations for abstract data have been underway for some time. APT [Mac86], TRIP [KK91], VISTA [SI94] and SAGE [RM90] systems are based on the process of translation from an application's data representation into pictorial representations. In this scheme, the data is represented by abstract objects and relations, which are mapped to graphical objects and graphical relations. Proximity Visualization [Bas01] is another system that visualizes the similarity between elements of abstract data collections. We have extended the ideas put forth by these systems and applied it to the creation of the proposed framework to facilitate the inclusion of various data types in general purpose visualization systems.

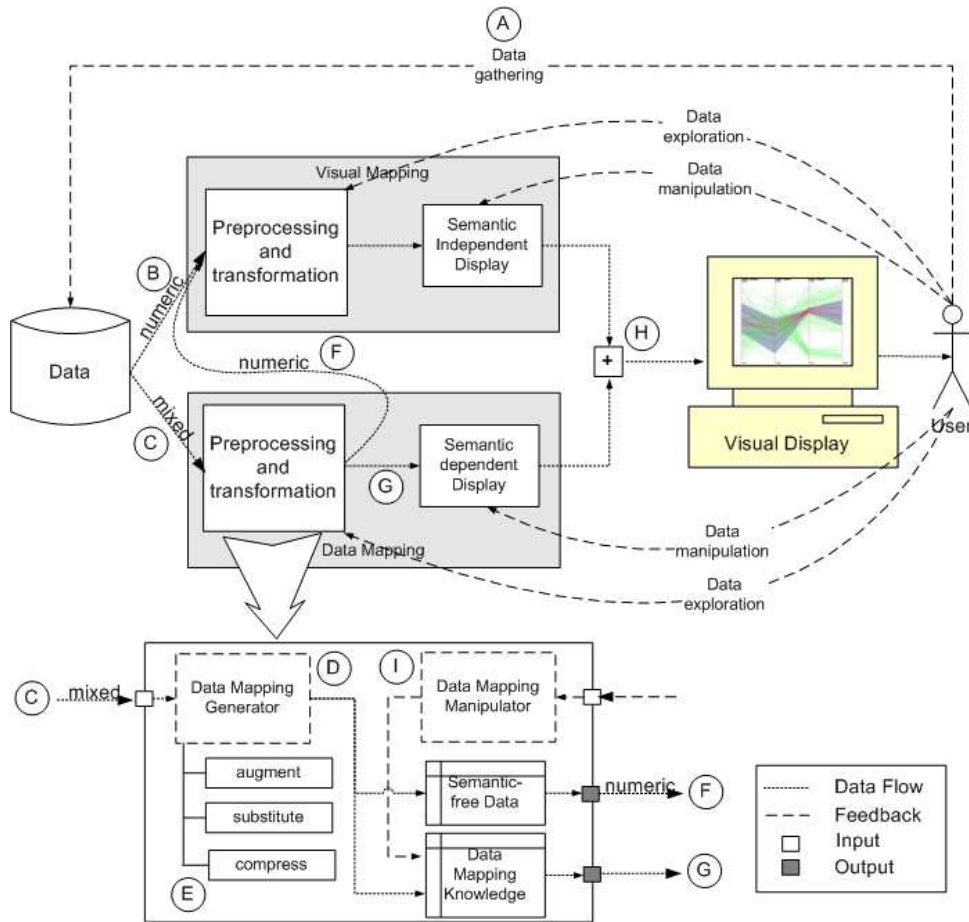


Figure 1: Block Diagram of the proposed framework.

3. Proposed Approach

There are four basic stages in the pipeline for a typical visualization tool (the visual mapping process), together with a number of feedback loops [War00]. They consist of:

- The collection and storage of data itself.
- The preprocessing designed to transform the data into a form that can be visualized.
- The graphics algorithms that produce an image on the screen.
- The human perceptual and cognitive system (the perceiver).

Figure 1 gives an overview of the framework proposed in this paper. Our framework introduces another mapping process, parallel to the visual mapping process, to include the semantics of the data within this pipeline of the visualization tool. The basic idea is that many visualization tools have a *semantically independent data display* and a *semantically dependent data display*. In current visualization systems, the semantically independent display is driven by the

visual mapping process and the semantically dependent display is normally ad hoc (indicate the numeric values instead of the actual values). Our framework hinges on this mapping process, that we term *data mapping*, that tries to capture the semantics of the data and then drive the semantically dependent display through this process.

The data mapping process outlined here involves two stages:

- Mapping Stage
 - This stage of the framework involves a transformation process that enables the mapping between a wide range of data types to a form the visualization system can process. This transformation process involves the creation, manipulation and storage of the mappings. The subcomponents of this process are:
 - Data Mapping Generator
 - This helps in the creation of mappings through the use of automatic / semi-automatic tools and process-

ing helper components. These are discussed in detail in Section 5.1 and Section 5.2.

- Data Mapping Manipulator
The data exploration process will also involve the manipulation of these mappings to gain better insight into the data. These subcomponents help in the manipulation of the mappings within the framework of the visualization system. Section 5.3 presents the Data Mapping Manipulators in detail.
- Data Mapping Encodings
Efficient and extensible storage of the mappings is an issue within the process. Section 6 provides details on the various encodings that are possible for storing the mappings.
- Display Stage
Semantic information associated with the data is captured by the first stage of the process. But this would be of no use if this information were not to be displayed on the screen. The integration of the mapping knowledge with the display of the tool would enable the presentation of this semantic information, either separate from or integrated with the actual display. This should involve minimal recoding of the display of the visualization tool so as to incorporate the different types of data that can be provided by the user.

The flow of data within the entire framework can now be conceptualized (alphabets in parentheses refer to the corresponding circles in Figure 1) as follows:

- Data is acquired by the user through some source (A).
- If the data is numeric in nature, then it can be directly fed to the visual mapping process (B).
- Otherwise the data is directed to the data mapping process (C) to convert it to a form that the visual mapping process can utilize.
- A Data Mapping Generator (D) creates the set of mappings and the numeric representation of the data. A set of processing helper components (E) assist in further refining the mappings. The numeric data (semantic-free data) (F) created by the Generator is directed to the visual mapping process and the mapping information is stored within the data mapping process in the data mapping knowledge repository (G).
- The visual mapping process and the data mapping process can both drive the entire display (H) to provide an effective visualization of the data. If the user is not satisfied with the visualization, he may want to explore the possibilities of manipulating the the numeric data or the mappings. The mappings can be modified through the Data Mapping Manipulator (I).

The idea of data mapping can be best understood with the help of a simple example. Consider we want to represent the grades of students for a particular class. In presenting the data to the visualization system in numeric form, the data

would have to be converted to values such as 4.0 for 'A', 3.0 for 'B', and so on. This converted data can be easily visualized by the system, but this mapping cannot be represented within the tool. Our data mapping process can store these mappings and as part of the framework drive the display to show the mapped values. Section 4.3 elaborates on this small example presented here.

4. Data Mapping

Techniques for the visualization of data are divided into two distinct groups [RW97]:

- **Non-transformational techniques:** These techniques simply map the data directly onto graphical primitives. The location of the graphical primitives on the screen is determined by some other numeric data. For example, we can use colors or the shapes of the points to represent nominal values. However, this may not work for many numeric displays such as ScatterPlot Matrices and Parallel Coordinates if the colors or shapes already have pre-defined purposes.
- **Transformational techniques:** These techniques transform the data into numeric values. For example, the most natural transformation we might apply to a set of nominal values is a frequency-based analysis. Given n nominal data points, each with m possible discrete values, we create a new table that counts the frequency of each value. In this example, we've transformed our nominal data into numeric values, and now it's a straightforward job to map these numeric values to the appropriate kinds of graphical entities (e.g. columns of a histogram).

The Data Mapping process that we outline here is a transformational technique that is minimally intrusive within the framework of an existing Visualization system. The need for such a mapping process was pointed out in the introduction. This block is the main system component responsible for storing the knowledge of the mappings.

4.1. Definition of Data Mapping

The main idea of the data mapping process is to include many diverse data types into an existing visualization system through association of metadata with the data to drive the visualization. This is a critical step of the knowledge discovery process within a semantics-based visualization system. It involves transforming data into a suitable format for a particular analysis tool. A good data transformation would facilitate better accuracy of the results of exploratory data analysis. Thus, Data Mapping can be defined as *a way to specify the binding of semantics of the data to a representation of the data that the visualization system can process*.

A rich and flexible data mapping process is the central component of the visualization framework. This process will serve to organize the information found in the complex, heterogeneous data sets obtained from different sources. The

need is for a process that describes the semantic content of the data rather than just reducing it to a collection of numerical data structures.

Since the framework that we are proposing needs to work with an existing visualization tool, the data mapping process has to work in tandem with the existing environment of the tool. The semantics and the mappings must serve as additions to the already existing numerical to graphical attribute mapping within the tool.

4.2. Data Mapping Knowledge

The data mapping knowledge or the metadata carries two categories of meta knowledge:

1. **Domain Description:** This is the description that characterizes the data types within the data. The inclusion of a large variety of data types within the framework will allow the visualization tool to process data gathered from many sources. Data types that would immediately benefit the visualization system if incorporated within the tool are:

- **Simple data types:** These data types are common in most data sets. Simple data types include types such as real and integer for interval data, strings for categorical data and the uniform resource identifiers (URI) for document collections. Other simple data types such as date and time can also be included for time-varying data. Our framework also allows the specification of domain constraints for the simple data types to validate the data that is being fed to the visualization system. A simple example of this is an enumeration of strings, where the value of a particular data entry can only be one of the values in the enumeration. This allows the integrity of the data set to be maintained and in some cases better efficiency for the storage of these simple data types may result.

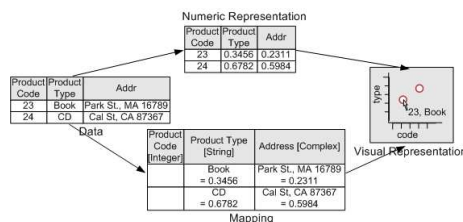


Figure 2: Data Mapping Example.

- **Complex data types:** These data types are essentially a collection of simple data types. One example of such a data type would be an address composed of several simple data types such as string for the street address, state and country information and integer for the zip code.

2. **Mapping Description:** This is the description that characterizes the mapping of the original data to assigned numerical values. The semantically dependent display uses these mappings to show the semantics. Mappings that could be specified are:

- **Enumerated Mappings:** This simply lists out all the mapping pairs between the data and its numerical representation within the metadata. This represents a finite number of enumerable domains possible for a dataset.
- **Function-driven Mappings:** This presents a function within the data mapping knowledge to map a particular data value to its numerical representation on demand. This is a much more flexible representation of the mapping because this does not depend on the number of data values within the dataset.

4.3. Example

Consider the example of a purchase transaction dataset from Amazon.com that could have many variables such as product code (numerical values), product type (book, CD, etc), and addresses. This example incorporates the use of simple data types such as integer for the code, strings or enumerations for the product type and complex data types such as addresses. A Data Mapping Generator can preprocess this data set and generate numeric mappings for the product type and addresses. The mappings, for example, can specify that the product type - book - will have a value of 0.3456 and that the product type - CD - may have a value of 0.6782. The metadata associated with the data will specify the data type of each dimension and the mappings associated with the values. The numeric data will contain these values to specify book or CD for the product type dimension. The visualization tool is capable of presenting the actual data using the numeric representation of the raw data and identifying the actual values using the stored mappings. The display section of the tool can be modified to make use of these mappings in order to present the visualization in a much more user-friendly way, as will be shown in the case studies (Section 7). Figure 2 gives a better view of this example. Note that the product code does not have any mappings associated with it since it's already in numeric form.

5. Data Mapping Helpers

The data mapping process is a simple process and as such can be handled manually. But, it would become cumbersome to generate these mappings for a large amount of data. Also, the feedback from the visualization might be used to modify the mappings from the original in an interactive manner. Thus, we propose a system of mapping helpers that assist the user to create and manipulate the mappings in this process. These may be automatic or semi-automatic tools that execute within the system or as stand-alone processes.

5.1. Data Mapping Generator

Converting from one data representation to another is both time-consuming and labor-intensive. The user is already involved getting the raw data into a form the visualization tool can interpret. Adding another conversion process for generating the mappings would not be acceptable for large amounts of data. Generators are a set of analytic tools that aid in the creation of the metadata needed for the data mapping process, either automatically or semi-automatically (i.e., with user input). The data mapping generator is not only responsible for generating the numeric form of the data for the visualization tool, but also for creating the mappings from the particular data type to the numeric values.

For example, to display nominal variables using numeric displays, one way is to map the nominal values to numbers, i.e. assign order and spacing to the nominal values. This could be driven by a data mapping generator that deals with nominal values. One example of such a data mapping generator is described in [MH99].

5.2. Data Processing Helpers

Processing Helpers are analytic tools that facilitate the importation of heterogeneous data types into a visualization system. To date we have identified three classes of Helpers:

1. **Substitution:** This type of helper does not change the number of records or dimensions, but substitutes values of one type with values that are interpretable by the visualization tool. Examples of this type of helper is converting nominal variables into numeric types, applying scaling operators, and transforming data from one coordinate system into another [MH99, RRBW03].
2. **Augmentation:** This helper adds to the dimensionality of a data set by performing some analytic process that generates derived values or classifications. Examples include extracting the first three principal components and adding the values in as new dimensions or performing clustering [JW88] on the data and labelling each data record according to the cluster to which it belongs.
3. **Compression:** this type of helper can reduce the number of records, the number of dimensions, or both. For example, we can implement helpers that eliminate outliers, remove redundant dimensions, or replace each record with its equivalent position in a lower dimension space using a process such as multidimensional scaling [JW88].

It is critical that each Helper augment the metadata associated with a data set to properly reflect to the user the nature of the mappings. Without this additional information it would be relatively easy for the user to misinterpret what is seen in the visualization.

5.3. Data Mapping Manipulator

After the metadata has been derived from the mapping process, it can be interpreted from the visualization much more

efficiently by the analyst. The analyst may further want to modify the mappings in order to derive a possible improvement. The manual modification of the metadata might require expert knowledge about the data and metadata. The framework includes a set of manipulator helper components that free the user from manually modifying the associations. The manipulators could include expert system tools that generate suggestions on how to enhance the mappings.

In order to create suggestions for the improvement of the mappings, a base for the decisions leading to these suggestions has to be provided by gathering statistical data about the mapping process. For example, given a nominal to numeric mapping, a manipulator might be used to reorder and re-quantify the nominal mappings and statistics can be displayed to show the "goodness" factor [MH99] of the new ordering.

A persistent mapping scheme would also be beneficial in the interchange of semantics for the same set of data. This would mean that the user can try various semantic associations with the dataset in order to get an effective visualization. For example, the user could manipulate the data mappings through the Data Mapping Manipulator and store this new set of mappings in a separate file. The user can then associate the newly created metadata with the data and can then study the effect of the manipulations.

6. Data Mapping Encodings

There is a need for the mapping knowledge to be persistent within the process because we do not want to modify the visualization tool in order to maintain these mappings. The mappings can be encoded in a variety of ways, including

- **ASCII Encoding**
A simple and fast solution would be to define a simple ascii file format with attribute-value pairs for the purpose of storing the mappings and to specify the type association. Such a method works for small to medium size volumes of data, and can be extremely efficient when coded correctly. However, it does require newcomers to learn the syntax of the data encoding format and is also difficult to validate. Basically, the problem of ascii encoding is the tight integration of the format with the code, making it hard to comprehend, extend and validate.
- **Binary Encoding**
Binary encoding of the mappings is possible and one could also associate a binary value for each type possible within the system. Binary encoding will, in most cases, result in compact files but has the same problems as of ASCII encodings along with the problem of being platform dependent.
- **XML Encoding**
XML is the eXtensible Markup Language, a World Wide Web Consortium (W3C) standard for the representation of information as structured documents, and is an emerging

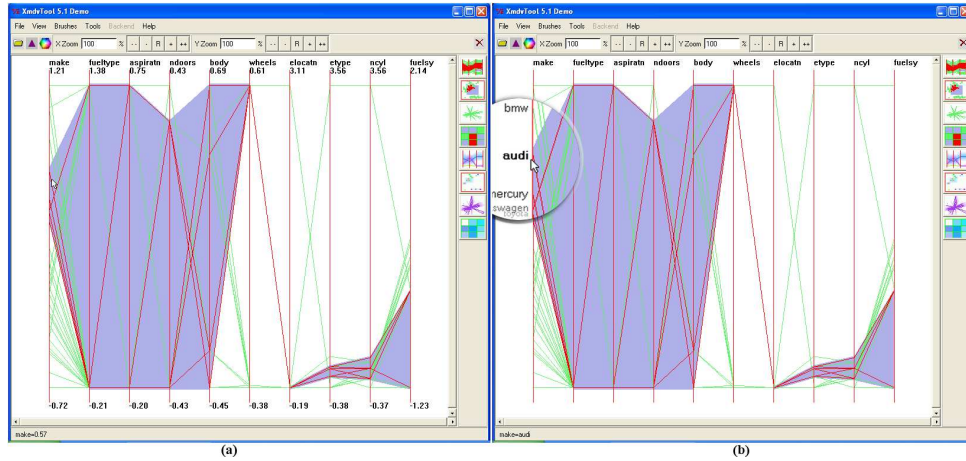


Figure 3: Nominal Data within XmdvTool. (a) Without annotation (b) With annotation (The labels have been highlighted by zooming into the region of the mouse cursor).

standard for digital information markup. It allows information (data, metadata, documentation, resources of any kind) to be expressed in a structured, flexible, and easily parsable manner. XML allows for contents-based tagging of any information resource, and consequently it allows for powerful, focused, and efficient contents-based search and retrieval of information. For these reasons, XML Encoding provides a more elegant solution to the above specified problems in other encoding systems. The advantages are:

- XML is self-describing, i.e., its tags are meaningful allowing us to semantically markup the document.
- XML is extensible in many ways, including extensible tag library, extensible document structure, and extensible document elements.
- XML is very suitable for creation and manipulation by a programming language due to the widespread availability of parsers.
- XML can be easily validated.

We chose the Data Mapping encodings to be in XML format due to the above mentioned advantages.

7. Case Studies

The development of the framework presented in this article was driven by our goal of adding support for embedding semantics within XmdvTool. We show that the ideas presented here are general enough to be applied to other visualization systems by providing an implementation for OpenDX.

XmdvTool is a public-domain software package designed for the interactive visual exploration of multivariate data [War94]. The tool provides four methods for displaying both flat (non-hierarchical) form data and hierarchically clustered data, namely scatterplots, star glyphs, parallel coordinates,

and dimensional stacking. XmdvTool also supports a variety of interaction modes and tools, including brushing in screen, data, and structure spaces, zooming, panning, and distortion techniques, and the masking and reordering of dimensions. More details can be found at <http://davis.wpi.edu/xmdv/>.

OpenDX is a widely used general purpose data visualization software package that runs on most major Unix and Windows platforms. Data Explorer (OpenDX) employs a dataflow client-server execution model that allows one to author visualization applications by selecting modules of appropriate functionality from a large library, and then to describe the flow of data through those modules. The selection of modules and dataflow can be specified through a visual program editor (VPE) that provides a point - click - connect graphic interface for application authoring. The functionality of the provided module library can be extended by writing new modules following well-defined guidelines.

7.1. XmdvTool

While to date, XmdvTool has addressed visualization of numeric data, the inclusion of our framework enables the same tool to visualize categorical data. Unlike numeric data, categorical values do not have an order. So, this is problematic for typical visualization techniques, since categorical values need to be mapped to the numeric values to specify the graphical attributes. Since the semantics of the data has changed, it would require special coding to provide the special displays required for categorical data within XmdvTool. This is also true for other types of data such as text document collections.

There are several independent approaches to visualizing nominal variables. One can use displays that are specifically designed for nominal variables: sieve diagrams

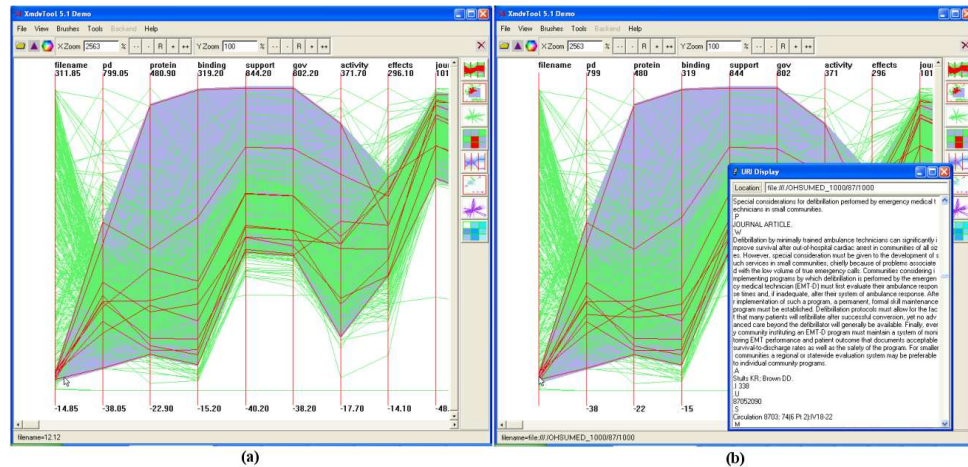


Figure 4: Text Visualization within XmdvTool (a) Without annotation (b) With annotation (along with a uri display window showing the complete text for the selected document)

[Fri99], mosaic displays [Fri99, HK81], maps from Correspondence Analysis [Gre93], fourfold displays [Fri99], treemaps [KW01], dimensional stacking [LWW90] and Cat-Trees [KW01]. Unfortunately, several of these approaches are either special-purpose, not readily available in common data analysis software [Fri99], or cannot handle high cardinality nominal variables well. So, the goal was to include the capability to analyze categorical data using existing visualization techniques within XmdvTool.

```
<okcmeta>
<dimensions value="10" />
<datapoints value="205" />

<dimension name="make" type="word">
<value scale="0.62750">honda</value>
<value scale="0.53840">mitsubishi</value>
<value scale="0.42860">plymouth</value>
.
.
</dimension>

<dimension name="fueltype" type="word">
<value scale="0.14135">gas</value>
<value scale="1.30745">diesel</value>
</dimension>
.
.
</okcmeta>
```

Figure 5: XML Snippet for the Automobile DataSet

For the following discussion, we are using the Automobile Data Set, which has the following nominal variables - make, fuel type, aspiration, number of doors, body type, wheels, engine location, engine type, number of cylinders and fuel system. A Data Mapping Generator for nominal data is used to generate the ordering and spacing attributes through the use of a data-driven, multivariate, scalable and distance-preserving method [RRBW03]. This generates the

numerical data required by XmdvTool as well as the data mapping knowledge required by the framework.

Figure 3 (a) shows how the Automobile Data Set would be represented if the data mapping process is not used. The difference is clearly visible in Figure 3 (b) where the semantically dependent part of the display has been recoded. The display now makes use of the data mapping knowledge and presents the visualization in a more user-friendly way. The statusbar and the annotated labels indicate the nominal values instead of the numerical values when the mouse hovers over a particular data entry. Figure 5 shows a snippet of the XML mapping produced by the Data Mapping Generator and stored by the system. This gives an example of how the semantics can be attached for the purpose of nominal data analysis within XmdvTool.

We have also used the framework for Text Visualization in XmdvTool. Text Visualization is primarily concerned with the presentation of a corpus of unstructured textual information so as to aid in the process of finding interesting or useful patterns in the document collection. Establishing a set of meaningful concepts for a text document collection allows one to look at hierarchical ordering of the concepts, and then mine for relationships in between documents and between concepts.

There are many text visualization tools available. Bead [CC92] presents similarity relationships between documents by 3-dimensional "particle" clouds. Using a similar visualization approach, Starlight [RMDT96] aims to visualize the content of multimedia databases. Closely related to these approaches is Galaxies [WTP*95] which yields a simple 2D scatter plot. VxInsight [DHJ*98] visualizes the document distribution density by means of a mountain terrain metaphor. These approaches are focused on the display of

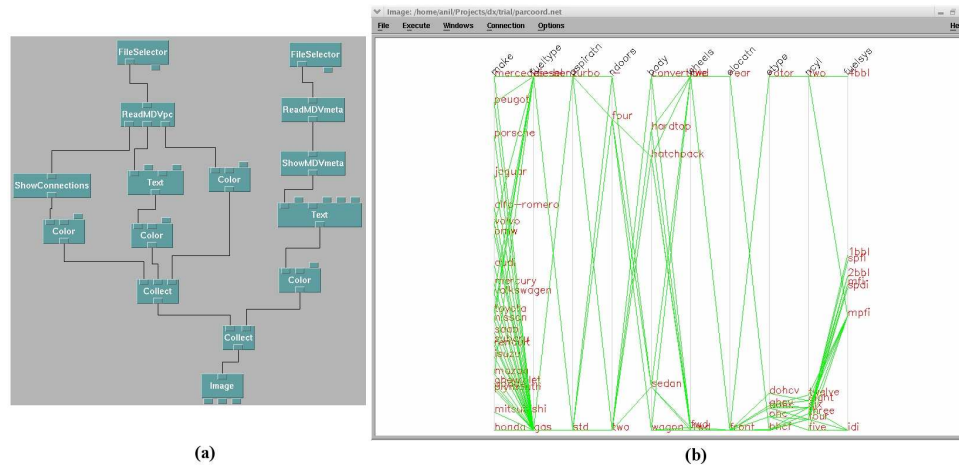


Figure 6: Nominal Data within OpenDX. (a) Visual program for the visualization (b) Annotated Parallel Coordinate visualization in OpenDX.

similarity between individual documents. The scaling techniques used try to optimize the distances between documents with respect to the given proximity measure. However, directly using a 2D or 3D space results in a relatively high information loss and can only very coarsely represent the original documents' relationships.

Our approach is to use a multivariate data visualization tool for the purpose of unstructured text mining. The hope is that the visualization of the document collection using a large number of dimensions would help gain better insight into the collection.

The Data Mapping Generator for the text visualization system is responsible for text extraction, pruning and dimension reduction [HWR03]. An available public domain tool, rainbow [McC96], was employed for text extraction. The text was tokenized by the most used tokenization options; the words from the SMART stoplist (524 common words), such as "the" and "of", are neglected before tokenization; the Porter stemming algorithm was applied for all words before they are counted. After tokenization, a document-term matrix was acquired and was processed by dimension reduction algorithms such as Multidimensional Scaling, Self-Organizing Maps and Principal Component Analysis.

Figure 4 (a) shows the visualization of document collection of medical abstracts (in 298 files, and taking 215 terms into account) without the data mapping framework. This figure shows the importance of including semantic meaning within the visualization system so that the data analyst can analyze the data more effectively. For example, an analyst would be completely lost in numerical data when he wants the association between the terms that are being visualized to the set of documents he has selected. Figure 4 (b) changes the same visualization in (a) to a more comprehensible vi-

sualization by presenting the document when the user clicks over a particular data entry.

7.2. OpenDX

OpenDX employs the data flow paradigm. The data is stored as a stream of bytes, transferred between the modules that comprise a data-flow "map" (commonly referred to as "network"). Typically, a data set is read from a file using one module and passed as a stream of bytes to another module where it is transformed into a visualization, and then passed to the rendering module for display.

OpenDX has a flexible data structure that describes the complex data sets that currently exist. Although this data structure is very accommodating, the data model is structure-based rather than knowledge-based. When the modules in the OpenDX map are executed, the data sets simply flow through the modules in the map losing the semantics of the data set. The framework presented here can help use the existing features of OpenDX to integrate and corroborate the use of various data types within the visualization. For this purpose, we implemented a sample set of modules to illustrate the idea.

A Parallel Coordinate module was developed which reads in numerical multivariate data sets in XmdvTool format and creates positions, connections, and axes for displaying the data. The network for displaying the parallel coordinates simply breaks the components out (data, border, axes), colors them, sets the data field to only Show Connections (an existing module in OpenDX), and combines the results into an Image. To use the features of the data mapping framework, two modules were implemented in OpenDX - the meta-input module and the meta-display module. The input module reads the mappings and the display module uses the

meta library to display the various data values as annotation driving the semantically dependent display. Figure 6 (a) shows the network and Figure 6 (b) presents the resultant image for the automobile car dataset. So, a simple addition of the data mapping pipeline transformed a numerical visualization to any data type driven visualization.

7.3. Implementation

The implementation of the data mapping process is in a separate library. This framework has been implemented in C++ using object oriented concepts so that it is easy to extend the framework for other data types. Once the framework was ready, it just took 2-3 man weeks time to integrate each of the nominal and text visualization capabilities into Xmdv-Tool and OpenDX. The data input module and the data display modules for each of the visualization techniques needed to be modified to integrate loading and displaying of type information within the tool.

We believe that it would take more or less the same amount of time for integrating the implemented framework within other visualization tools depending on the complexity of the implementation of the system.

8. Conclusion and Future Work

The key contribution of this work is the development of a general framework to enable heterogeneous data types to be included within the environment of an existing and future visual data exploration tools and to enable capturing and integration of data semantics into the system. The main component of this framework consists of a data mapping process, which enables the mapping between varied data types to a form that the visualization tool can interpret. This enables the framework to present the semantic information either separate from or integrated with the display. We have a working prototype implementation for a publicly available multivariate visualization tool and another general purpose visualization system which validates the applicability of the framework.

Future research plans include incorporating this framework within other visualization tools. This effort will help validate the generality of the framework and demonstrate the impact that this idea can have in the development of future visualization systems. Even though the framework includes many of the data types found in most data sets, the extension to other data types will be explored to verify that the system is sufficiently flexible. The inclusion of function-driven mappings in addition to the currently implemented enumerated mappings within the data mapping knowledge will certainly reduce the storage requirements of the metadata. Currently, the mapping helper tools - data mapping generators and manipulators - are a separate system of tools and utilities from the system. We envision the complete system for the entire

data mapping process could be executed through a integrated graphical user interface.

Acknowledgments

We gratefully acknowledge our colleagues in the XmdvTool group at WPI for their contributions to this research. The research funds from NSF and NSA are gratefully appreciated.

References

- [Adv03] ADVANCED VISUAL SYSTEMS INC.: Avs/express visualization edition. http://www.avs.com/software/soft_t/avsxps.html, 2003.
- [And72] ANDREWS D.: Plots of high dimensional data.
- [Bas01] BASALAJ W.: Proximity visualization of abstract data. <http://www.pavis.org/>, 2001.
- [CC92] CHALMERS M., CHITSON P.: Bead: Explorations in information visualization. In *Research and Development in Information Retrieval* (1992), pp. 330–337.
- [CCS97] CODY R. P., CODY R., SMITH J.: *Applied Statistics and the SAS Programming Language*, 4th ed. Prentice Hall, 1997.
- [DHJ*98] DAVIDSON G. S., HENDRICKSON B., JOHNSON D. K., MEYERS C. E., WYLIE B. N.: Knowledge mining with vxinsight: Discovery through interaction. In *Journal of Intelligent Information Systems*, Vol. 11, No. 3 (1998), pp. 259–285.
- [Fri94] FRIENDLY M.: Mosaic displays for multi-way contingency tables. 190–200.
- [Fri99] FRIENDLY M.: Visualizing categorical data. In *Cognition and Survey Research*, Sirken M. G., Herrmann D. J., Schechter S., Schwarz N., Tannur J. M., Tourangeau R., (Eds.). John Wiley & Sons, Inc., New York, 1999, pp. 319–348.
- [FWR99] FUA Y., WARD M., RUNDENSTEINER E.: Hierarchical parallel coordinates for exploration of large datasets. In *Proc. of Visualization '99*, p. 43-50 (San Francisco, California, USA, Oct. 1999).
- [Gre93] GREENACRE M. J.: *Correspondence Analysis in Practice*. Academic Press, London, 1993.
- [HK81] HARTIGAN J., KLEINER B.: Mosaics for contingency tables. In *Computer Science and Statistics: Proceedings of the 13th Symposium on the Interface* (New York, 1981), Eddy W. F., (Ed.), Springer-Verlag, pp. 268–273.

- [HWR03] HUANG S., WARD M. O., RUNDENSTEINER E. A.: *Exploration of Dimensionality Reduction for Text Visualization*. Technical Report TR-03-14, Worcester Polytechnic Institute, Computer Science Department, 2003.
- [ID90] INSELBERG A., DIMSDALE B.: Parallel coordinates: A tool for visualizing multidimensional geometry. In *Proc. of Visualization '90*, p. 361-78 (1990).
- [Joh92] JOHNSON B.: Treemap visualization of hierarchically structured information. In *Proceedings of the SIGCHI conference on Human factors in computing systems* (1992), ACM Press, pp. 369-370.
- [JW88] JOHNSON R. A., WICHERN D. W.: *Applied Multivariate Statistical Analysis*, second ed. Prentice Hall International, Inc., 1988.
- [KD94] KEIM D., DRIEGEL H.: Visdb: Database exploration using multidimensional visualization. In *Computer Graphics & Applications*, Sept. 1994, p. 40-49 (1994).
- [KK91] KAMADA T., KAWAI S.: A general framework for visualizing abstract objects and relations. In *ACM Trans. Gr.* (1991), vol. 10, pp. 1-39.
- [KKA95] KEIM D., KRIEGEL H., ANKERST M.: Recursive pattern: a technique for visualizing very large amounts of data. In *Proc. of Visualization '95*, p. 279-86 (San Francisco, CA, USA; 27 Oct.-1 Nov. 1995, 1995), ACM; New York, NY, USA.
- [KO02] KRAUSE A., OLSON M.: *The Basics of S-PLUS*, 3rd ed. Springer Verlag, 2002.
- [KW01] KOLATCH E., WEINSTEIN B.: Cat-trees: Dynamic visualization of categorical data using treemaps. http://www.cs.umd.edu/class/spring2001/cmsc838b/Project/Kolatch_Weinstein, 2001.
- [LRB*97] LIVNY M., RAMAKRISHNAN R., BEYER K. S., CHEN G., DONJERKOVIC D., LAWANDE S., MYLLYMAKI J., WENGER R. K.: DEVise: Integrated querying and visualization of large datasets. In *Proc ACM SIGMOD Intl Conf on Management of Data*, May 13-15, 1997, Tucson, Arizona, USA (1997), ACM Press, pp. 301-312.
- [LWW90] LEBLANC J., WARD M., WITTELS N.: Exploring n-dimensional databases. In *Proc. of Visualization '90*, p. 230-7 (San Francisco, CA, USA; 23-26 Oct. 1990, 1990), IEEE Comput. Soc. Press; Los Alamitos, CA, USA.
- [Mac86] MACKINLAY J.: Automating the design of graphical presentations of relational information. In *ACM Transactions on Graphics* (1986), vol. 5, pp. 110-141.
- [McC96] MCCALLUM A. K.: Bow: A toolkit for statistical language modeling, text retrieval, classification and clustering. <http://www.cs.cmu.edu/~mccallum/bow>, 1996.
- [MH99] MA S., HELLERSTEIN J.: Ordering categorical data to improve visualization, 1999.
- [ODR*99] OWEN G. S., DOMIK G., RHYNE T.-M., BRODLIE K. W., SANTOS B. S.: Hypervis - teaching scientific visualization using hypermedia. <http://www.siggraph.org/education/materials/HyperVis/hypervis.htm>, 1999.
- [Rat96] RATHJENS D.: *MineSet User's Guide*. Silicon Graphics, Inc., 1996.
- [RM90] ROTH S. F., MATTIS J.: Data characterization for intelligent graphics presentation. In *Proceedings, ACM Conference on Human Factors in Computing Systems* (1990), pp. 193-200.
- [RMDT96] RISCH J. S., MAY R. A., DOWSON S. T., THOMAS J. J.: A virtual environment for multimedia intelligence data analysis. pp. 33-41.
- [RRBW03] ROSARIO G. E., RUNDENSTEINER E. A., BROWN D. C., WARD M. O.: *Mapping Nominal Values to Numbers for Effective Visualization*. Technical Report TR-03-11, Worcester Polytechnic Institute, Computer Science Department, 2003.
- [RW97] RESNICK R., WARD M. O.: Visualizing nominal data. http://www.cs.wpi.edu/~matt/courses/cs563/talks/nominal_data/index.html, 1997.
- [SCB98] SWAYNE D. F., COOK D., BUJA A.: Xgobi: Interactive dynamic data visualization in the x window system. *Journal of Computational and Graphical Statistics* 7, 1 (1998).
- [SH00] STOLTE C., HANRAHAN P.: Polaris: A system for query, analysis and visualization of multidimensional relational databases. *Information Visualization 2000* (2000).
- [SI94] SENAY H., IGNATIUS E.: A knowledge-based system for visualization design. In *IEEE Computer Graphics and Applications* (November 1994), pp. 36-47.
- [Vis03] VISUALIZATION AND IMAGERY SOLUTIONS, INC.: Opendx: The open source software project based on ibm's visualization data

- explorer. <http://www.opendx.org/index2.php>, 2003.
- [War94] WARD M.: Xmdvtool: Integrating multiple methods for visualizing multivariate data. In *Proc. of Visualization '94* (Washington, DC, USA; 17-21 Oct. 1994, 1994), IEEE Comput. Soc. Press; Los Alamitos, CA, USA, pp. 326–33.
- [War00] WARE C.: *Information Visualization: Perception for Design*. Harcourt Publishers Ltd, 2000.
- [WTP*95] WISE J. A., THOMAS J. J., PENNOCK K., LANTRIP D., POTTIER M., SCHUR A., CROW V.: Visualizing the non-visual: Spatial analysis and interaction with information from text documents. In *Proc. of Information Visualization '1995*, p. 51-58 (1995).