

A New Model of Graph and Visualization Usage

J. Gregory Trafton
Naval Research Laboratory
trafton@itd.nrl.navy.mil

Susan B. Trickett
George Mason University
strickett@gmu.edu

Abstract

We propose that current models of graph comprehension do not adequately capture how people use graphs and complex visualizations. To investigate this hypothesis, we examined 3 sessions of scientists using an *in vivo* methodology. We found that in order to obtain information from their graphs, scientists not only read off information directly from their visualizations (as current theories predict), but they also used a great deal of mental imagery (which we call spatial transformations). We propose an extension to the current model of visualization comprehension and usage to account for this data.

Introduction

If a person looks at a standard stock market graph or a meteorologist is examining a complex meteorological visualization, how is information extracted from these graphs? The most influential research on graph and visualization comprehension is Bertin's (1983) task analysis that suggests three main processes in graph and visualization comprehension:

1. Encode visual elements of the display: For example, identify lines and axes. This stage is influenced by pre-attentive processes and is affected by the discriminability of shapes.
2. Translate the elements into patterns: For example, notice that one bar is taller than another or the slope of a line. This stage is affected by distortions of perception and limitations of working memory.
3. Map the patterns to the labels to interpret the specific relationships communicated by the graph. For example, determine the value of a bar graph.

Most of the work done on graph comprehension has examined the encoding, perception, and representation of graphs. Cleveland and McGill, for example, have examined the psychophysical aspects of graphical perception (Cleveland & McGill, 1984, 1986). Similarly, Pinker's theory of graph comprehension, while quite broad, focuses on the encoding and understanding of graphs (Pinker, 1990). Kosslyn's work emphasizes the cognitive processes that make a graph more or less difficult to read. Kosslyn's syntactic and semantic (and to a lesser degree pragmatic) level of analysis focuses on encoding, perception, and representation of graphs (Kosslyn, 1989). Recent work by Carpenter and Shah (1998)

shows that people switch between looking at the graph and the axes in order to comprehend the visualization.

This scheme seems to work very well when the graph contains all the information the user needs (i.e., when the information is explicitly represented in one form or another). Thus, when an undergraduate is asked to extract specific information from a bar-graph, the above process seems to hold. However, graph usage outside the laboratory is probably not simply a series of information extractions. For example, when looking at a stock market graph, the goal may not be just to determine the current or past price of the stock, but perhaps to determine what the price of the stock will be sometime in the future. A weather forecaster looking at a meteorological visualization is frequently trying to predict what the weather will be in the future, as well as what the current visualization shows (Trafton, Kirschenbaum, Tsui, Miyamoto, Ballas, & Raymond, 2000). A scientist examining results from a recent experiment can not always display the available information in a way that perfectly shows the answer to her hypotheses.

How do current theories of graph comprehension hold up when a graph or visualization does not contain the exact information needed? Unfortunately, the theories do not say anything about this situation. In fact, there are no specifications in any theory of graph comprehension about how information could or would be extracted from a visualization where that information is not represented in some form. If a graph does not contain the information needed by the user, the graph is often labeled "bad" or "useless" (Kosslyn, 1989; Pinker, 1990).

Current graph comprehension theories do not have a great deal to say about what to do when a graph does not explicitly show the needed information for a variety of reasons. The main reason is probably that most graph comprehension studies have used fairly simple graphs for which no particular domain knowledge is required (e.g., Carter, 1947; Lohse, 1993; Pinker, 1990). However, in real-world situations, people use complex visualizations that require a great deal of domain knowledge, and all the needed information would probably not be explicitly represented in the graph. This study will thus try to answer two questions about graph comprehension. Do expert users of visualizations ever need information that is not on a specific graph they are using? If so, how do they extract that information from the graph?

Report Documentation Page			Form Approved OMB No. 0704-0188		
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE 2001		2. REPORT TYPE		3. DATES COVERED 00-00-2001 to 00-00-2001	
4. TITLE AND SUBTITLE A New Model of Graph and Visualization Usage			5a. CONTRACT NUMBER		
			5b. GRANT NUMBER		
			5c. PROGRAM ELEMENT NUMBER		
6. AUTHOR(S)			5d. PROJECT NUMBER		
			5e. TASK NUMBER		
			5f. WORK UNIT NUMBER		
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Naval Research Laboratory, Navy Center for Applied Research in Artificial Intelligence (NCARAI), 4555 Overlook Avenue SW, Washington, DC, 20375			8. PERFORMING ORGANIZATION REPORT NUMBER		
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)			10. SPONSOR/MONITOR'S ACRONYM(S)		
			11. SPONSOR/MONITOR'S REPORT NUMBER(S)		
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES The proceedings of the twenty-third annual conference of the cognitive science society, 1-4 Aug 2001, Edinburgh, Scotland					
14. ABSTRACT We propose that current models of graph comprehension do not adequately capture how people use graphs and complex visualizations. To investigate this hypothesis, we examined 3 sessions of scientists using an in vivo methodology. We found that in order to obtain information from their graphs, scientists not only read off information directly from their visualizations (as current theories predict), but they also used a great deal of mental imagery (which we call spatial transformations). We propose an extension to the current model of visualization comprehension and usage to account for this data.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 6	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

There are several possible things that users could do when trying to extract information from a graph. In the simplest case, the information is explicitly available, and they can simply read off the information from the visualization. What do they do when information they need is not available on the visualization? They could create a completely new visualization that does show the information. They could also collect more data or consult another source. They could create an explicit plan to look for more data or run another experiment.

What do they do when the visualization is all they have to work with? What kind of mental operations could users perform on graphs and visualizations in order to extract information that is not explicit? One possibility is that people use some sort of visual imagery to extract information that is not explicitly represented on a graph or visualization. For example, a weather forecaster may mentally imagine a front moving east over the next several days (Trafton et al., 2000), or a stock analyst may mentally extend a line on a graph and think that a stock will continue to rise. We have developed a framework for coding and working with these kinds of graphs and visualizations called *Spatial Transformations* that will be used to investigate these issues. We will argue that spatial transformations are a fundamental aspect of complex visualization usage.

Spatial Transformations are cognitive operations that a scientist performs on a visualization. Sample spatial transformations are mental rotation (e.g., Shepard & Metzler, 1971), creating a mental image, modifying that mental image by adding or deleting features to or from it, animating an aspect of a visualization (Hegarty, 1992) time series progression prediction, mentally moving an object, mentally transforming a 2D view into a 3D view (or vice versa), comparisons between different views (Kosslyn, Sukel, & Bly, 1999; Trafton, Trickett, & Mintz, 2001), and anything else a scientist mentally does to a visualization in order to understand it or facilitate problem solving. Also note that a spatial transformation can be done on either an internal (i.e., mental) image or an external image (i.e., a scientific visualization on a computer-generated image). What all spatial transformations have in common is that they involve the use of mental imagery. A more complete description of spatial transformations can be found at <http://iota.gmu.edu/users/trafton/405st.html>.

We will examine the number of times that users needed information from a visualization. If all or most of the information is available explicitly on the visualization, we should see primarily read-offs (Kosslyn, 1989; Pinker, 1990). If, however, a particular visualization does not explicitly display particular information that a scientist wants, we will examine how the scientist goes about obtaining that information. We expect that in complex visualizations, there is a great deal of information that is needed in addition to what is displayed, and we expect scientists to use spatial transformations to retrieve that information.

Method

In order to investigate the issues discussed above, we have adapted Dunbar's in vivo methodology (Dunbar, 1995, 1996; Trickett, Trafton, & Schunn, 2000b). This approach offers several advantages. First, it allows the observation of experts, who are thus able to use their domain knowledge to guide their strategy selection. Second, it allows the collection of "on-line" measures of thinking, which allow the investigation of the scientists' reasoning as it occurs (Ericsson & Simon, 1993). Finally, the tasks (experiment design, data analysis, etc.) conducted by the scientists, as well as the tools they use, are fully authentic.

Two sets of scientists were videotaped while conducting their own research. All the scientists were experts, having earned their Ph.D.s more than 6 years previously. In the first set, two astronomers, one a tenured professor at a university, the other a fellow at a research institute, worked collaboratively to investigate computer-generated visual representations of a new set of observational data. At the time of this study, one astronomer had approximately 20 publications in this general area, and the other approximately 10. The astronomers have been collaborating for some years, although they do not frequently work at the same computer screen and the same time to examine data.

In the second dataset, a physicist with expertise in computational fluid dynamics worked alone to inspect the results of a computational model he had built and run. Two related sessions were recorded with this scientist over consecutive days. He works as a research scientist at a major U.S. scientific research facility, and had earned his Ph.D. over 20 years previously. He had inspected the data previously but had made some adjustments to the physics parameters underlying the model and was therefore revisiting the data.

Both sets of scientists were instructed to carry out their work as though no camera were present and without explanation to the experimenter (Ericsson & Simon, 1993). The relevant part of the astronomy session lasted about 53 minutes, and the two physics sessions each lasted approximately 15 minutes. All utterances were later transcribed and segmented according to complete thought. All segments were coded by 2 coders as on-task (pertaining to data analysis) or off-task (e.g., jokes, phone call interruptions, etc.). Inter-rater reliability for this coding was more than 95%. Off-task segments were excluded from further analysis. On-task segments ($N = 649$ for the astronomy dataset and $N = 189$ for the first physics dataset and $N = 176$ for the second physics dataset) were further coded as described below.

The Tasks and the Data

Astronomy The astronomical data under analysis were optical and radio data of a ring galaxy. The astronomers' high-level goal was to understand its evolution and structure by understanding the flow of gas in the galaxy. In order to understand the flow of gas, the astronomers must make inferences about the velocity field, represented by

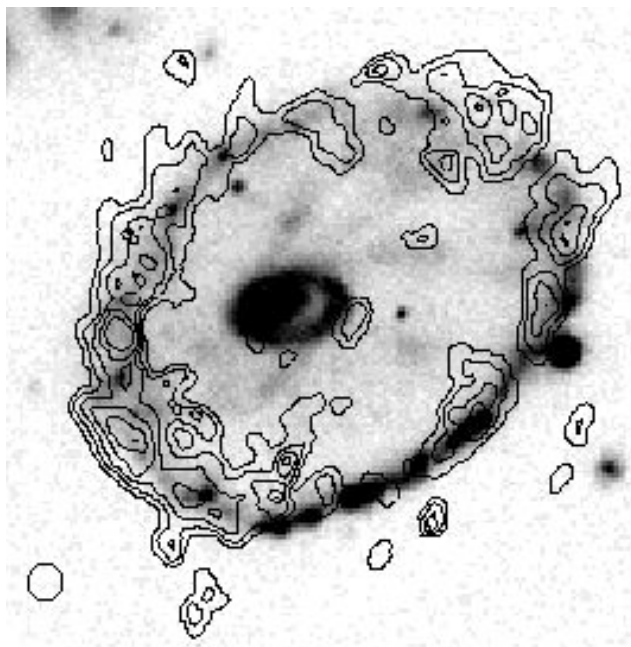


Figure 1: An example of the kind of visualizations examined by the astronomers.

contour lines on the 2-dimensional display. The astronomers' task was made difficult by two characteristics of their data. First, the data were one- or at best 2-dimensional, whereas the structure they were attempting to understand is 3-dimensional. Second, the data were noisy, and there was no easy way to distinguish between noise and real phenomena. Figure 1 shows a screen snapshot of the type of data the astronomers were examining. In order to make their inferences, the astronomers used different types of image, representing different phenomena (e.g., different forms of gas), which represent different information about the structure and dynamics of the galaxy. Some of these images could be overlaid on each other. In addition, the astronomers could choose from images created by different processing algorithms, each with advantages and disadvantages (e.g., more or less resolution). Finally, they could adjust different features of the display, such as contrast or false color. A more complete description of this dataset can be found in Trickett, Fu, Schunn, and Trafton (2000a) and Trickett, Trafton, and Schunn (2000b).

Physics The physicist was working to evaluate how deep into a pellet a laser light will go before being reflected. His high-level goal was to understand the fundamental physics underlying the reaction, an understanding that hinged on an understanding of the relative importance and growth rates of different modes. The physicist had built a model of the reaction; other scientists had independently conducted experiments in which lasers were fired at pellets and the reactions recorded. A close match between model and empirical data would indicate a good

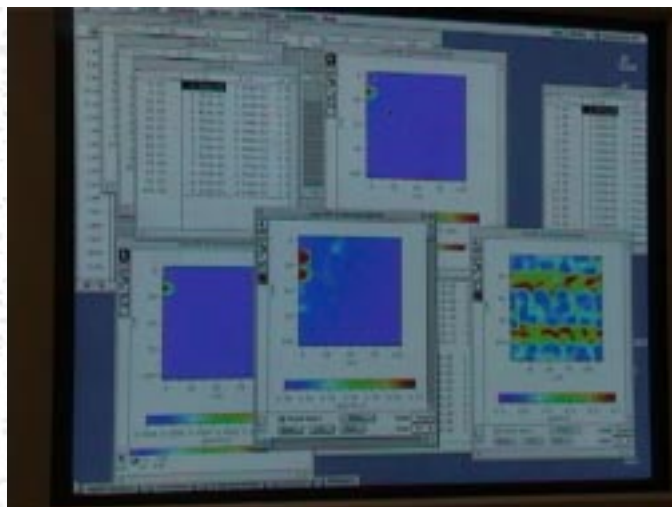


Figure 2: An example of the kind of visualizations examined by the physicist.

understanding of the underlying theory. Although the physicist had been in conversation with the experimentalist, he had not viewed the empirical data, and in this session he was investigating only the results of his computational model. However, he believed the model to be correct (i.e., he had strong expectations about what he would see), and in this sense, this session may be considered confirmatory.

The data consisted of two different kinds of representation of the different modes, shown over time (nanoseconds). The physicist was able to view either a Fourier decomposition of the modes or a representation of the "raw" data.

Figure 2 shows an example of the physicist's data. He could choose from black-and-white or a variety of color representations, and could adjust the scales of the displayed image, as well as some other features. He was able to open numerous views simultaneously.

Coding Scheme

Our goals in this research are first, to determine if complex visualizations contain all the information needed by the scientists, and, if not, to investigate what happens when they do not have all the information they need. We propose that spatial transformations are a major portion of extracting information from a visualization when the data is not explicitly represented. Consequently, we identified every situation where a scientist wanted to extract information from a visualization. Next, we coded what the scientist did to extract information, including reading off the information directly from the graph, spatial transformations, changing the visualization, plans or discussions about getting more data, and abandoning their attempt to get the information. We now describe and provide examples of this coding scheme in detail.

Example	Explanation
After all, it is ten to the minus six...	Scientist is looking at a line and extracting the y-axis value
I mean, the fact you see such a strong concentration of gas in the ring, um...	Scientist is reading off the amount of gas in the ring
That's about 220 km/sec, which is the velocity spread of a normal galaxy.	Scientist is reading off the velocity spread

Table 1: Examples of information that is read off the visualizations.

Spatial Transformation	Example	Explanation
Create Mental Image	I mean, in a perfect, in a perfect world, in a perfect sort of spider diagram...	Scientist is creating a mental image of a spider diagram; there is no spider diagram displayed.
Modify Image	So that [line] would be below the black line	Scientist is adding a new (hypothesized) line to a current visualization
Modify Image	If there was no streaming motion or sort of piling of gas	Scientist has imaged a previous mental image and is now removing the streaming motions from his mental image
Comparison:	Maybe it's a projection effect, although if that's true, there should be a very large velocity dispersion.	Scientist is comparing a current image to a previously created mental image.

Table 2: Examples of spatial transformations.

Desire to extract information A scientist would frequently want to extract some amount of information from a visualization. Comments varied from the very general (“What do we see?”) to the very specific (“Let’s see, how does oh-three versus three-oh [look]?”).

Read-Off A scientist would be able to read-off information directly from the graph. Information that was read off a visualization was explicitly on the graph and the scientist simply had to read-off a particular value. For every utterance, we evaluated whether a value was read off the visualization. Table 1 shows several examples of information that was read off of the visualization.

Spatial Transformations As discussed earlier, spatial transformations are cognitive operations that a scientist performs on a visualization. For every utterance in each protocol we evaluated whether there was a spatial transformation. Spatial transformations were further coded as Create Image, Modify Image, or Comparison. Table 1 shows examples of each category of spatial transformation (note that these utterances are independent of one another and do not represent a sequence). Table 2 shows several examples of spatial transformations that were used by the scientists.

Changing the Visualization The scientists were using their own tools and were able to change the visualization to a completely different representation. For example, a scientist could change the data display from the raw data

to a Fourier mode display). Alternatively, the scientists could “tweak” the current representation (from black and white to color, for example). We coded the visualization changes where the scientists were looking for additional information. If they simply made a mistake and tweaked the visualization, we did not count that visualization change. For example, while looking at a particularly compressed visualization, one of the scientists said “Where’s three-oh at? Don’t see three [oh]. That’s what I figured, I was gonna get spaghetti. Let’s do a re-plot.” and then replotted the data with a reduced dataset.

Plans to gather more data Occasionally, the scientists wanted or needed to gather more data. We coded every time they made a plan to gather more data. For example, one scientist said “So that means that this guy is in fact between him and him, which is exactly what the experimentalist believes he saw. Now, somewhere along the line I have to get their results.”

Abandoning their attempt to get information Sometimes the scientists either could not decide what data to get or simply abandoned their quest for a specific information. We coded every time the scientists abandoned their attempt to get information. For example, one scientist, unable to explain a particular feature after extensive investigation of the image, said “Yeah well, [let’s] gloss over it.”

Results

Our two goals in this paper are to explore whether scientists are able to directly extract the amount of information they need from the visualizations they examine and if not, to explore how they do get the information that is needed.

How often is needed information directly available?

Of the 1014 total utterances in the three sessions, almost half (481) involved some form of information gathering.

As Figure 3 shows, approximately half of those information gathering instances were read-off, suggesting that the scientist did use the visualization a great deal to extract information. However, there were many times when the scientists needed information from a visualization but it was not available directly from the visualization. Thus, the visualizations seem to be good, but far from perfect from an information gathering point of view.

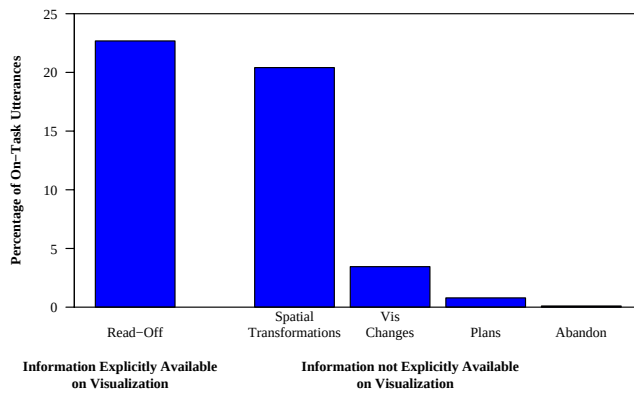


Figure 3: The number of read-offs, spatial transformations, visualization changes, plans to collect future data, and decisions to abandon the attempt to get more data for all datasets.

How was needed information extracted if it was not simply read off? As Figure 3 shows, the vast majority of information that was not read off was gathered by using spatial transformations. In fact, there was no statistical difference between the number of times that the scientists read off information directly from the graph and the number of spatial transformations, $\chi^2(1) = 1.21, p > .20$.

Additionally, scientists chose to use a spatial transformation to get needed information from a visualization rather than changing the visualization, $\chi^2(1) = 122.25, p < .001$, making plans to gather more data $\chi^2(1) = 184.19, p < .001$, or abandoning their attempt to answer their question, $\chi^2(1) = 204.02, p < .001$.¹

¹All χ^2 's used the Bonferroni adjustment.

General Discussion

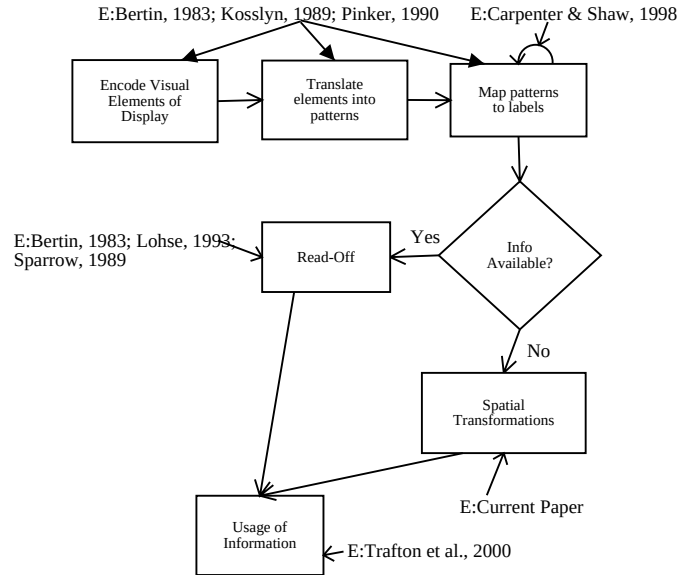


Figure 4: Our current theoretical model of complex visualization usage. The “E:” shows evidence for each stage of the model.

We have conducted a detailed analysis of expert scientists at work in their own laboratories, analyzing data that they have collected themselves. Our results show that these scientists do extract a great deal of information from the visualizations. However, these visualizations do not provide the scientists with all the information they need to answer their questions. We found that when they needed information that was not explicitly provided by the visualization, they tended to perform spatial transformations to answer their questions.

It is interesting that the scientists did not simply change the visualization more frequently to get the needed information. There was some evidence in the protocols that it was not easy to create new visualizations. For example, some of the visualizations had to be re-done because of an error that was made in the display (i.e., needed data was not included in the plot or the plot was not presented logarithmically when it should have been). However, this problem did not seem to have prevented the scientists from trying to make the changes: there were no instances of a scientist saying the visualization tool was too complicated or difficult to work with (though these tools could no doubt be improved). Thus, the scientists’ use of spatial transformations do not seem to be a substitute for “bad” graphs, but rather a strategy to understand the data more thoroughly.

As suggested earlier, current theories of graph comprehension can not account for this pattern of results. Current theories (e.g., Bertin, 1983; Kosslyn, 1989; Pinker, 1990) deal primarily with how users extract information that is explicitly available on a graph or vi-

sualization. In this study, we have shown that users do not simply extract information that is explicitly shown on a visualization; rather, they extract information and use mental imagery to create similar visualizations, modify those mental images, and compare their mental results to on-screen results. These spatial transformation seem to be used for a variety of reasons, including hypothesis testing and understanding their own mental representation through a process of aligning various mental images (Trafton et al., 2001).

How can we integrate these new results into current theories? We believe that the current theoretical model should be expanded to include spatial transformations as part of the cognitive processes that users go through to interpret and use visualizations. Figure 4 shows our current model of graph comprehension, along with evidence that supports each stage of this model.

We believe, as Figure 4 shows, that when people use graphs or visualizations, they initially go through a process to understand the graph itself. Then, when they need to extract information, they can either read off that information directly from the visualization or, if that information is not available, perform a spatial transformation to get the needed information.² Finally, that information is actually used by the user.

Acknowledgments: This research was supported in part by grant 55-7850-00 to the first author from the Office of Naval Research. We would also like to thank Wai-Tat Fu, Anthony Harrison, and William Liles for comments on a previous draft of this paper.

References

- Bertin, J. (1983). *Semiology of graphs*. Madison, WI: University of Wisconsin Press.
- Carpenter, P. A., & Shah, P. (1998). A model of the perceptual and conceptual processes in graph comprehension. *Journal of Experimental Psychology: Applied*, 4(2), 75–100.
- Carter, L. F. (1947). An experiment on the design of tables and graphs used for presenting numerical data. *Journal of Applied Psychology*, 31, 640–650.
- Cleveland, W. S., & McGill, R. (1984). Graphical perception: theory, experimentation, and application to the development of graphical method. *Journal of the American Statistical Association*, 79, 531–553.
- Cleveland, W. S., & McGill, R. (1986). An experiment in graphical perception. *International Journal of Man-Machine Studies*, 25, 491–500.
- Dunbar, K. (1995). How scientists really reason: Scientific reasoning in real-world laboratories. In R. J. Sternberg, & J. E. Davidson (Eds.), *The nature of insight*, (pp. 365–395). Cambridge, MA: MIT Press.
- Dunbar, K. (1996). How scientists think: Online creativity and conceptual change in science. In T. B. Ward, S. M. Smith, & S. Vaid (Eds.), *Creative thought: An Investigation of Conceptual Structures and Processes*, (pp. 461–493). Washington, DC: APA Press.
- Ericsson, K. A., & Simon, H. A. (1993). *Protocol analysis: Verbal reports as data* (Revised edition). Cambridge, MA: MIT Press.
- Hegarty, M. (1992). Mental animation: Inferring motion from static displays of mechanical systems. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 18(5), 1084–1102.
- Kosslyn, S. M. (1989). Understanding charts and graphs. *Applied Cognitive Psychology*, 3, 185–226.
- Kosslyn, S. M., Sukel, K. E., & Bly, B. M. (1999). Squinting with the mind's eye: Effects of stimulus resolution on imaginal and perceptual comparisons. *Memory and Cognition*, 27(2), 276–287.
- Lohse, G. L. (1993). A cognitive model for understanding graphical perception. *Human Computer Interaction*, 8, 353–388.
- Pinker, S. (1990). A theory of graph comprehension. In R. Freedle (Ed.), *Artificial intelligence and the future of testing*, (pp. 73–126). Hillsdale, NJ: Lawrence Erlbaum Associates, Inc.
- Shepard, R. N., & Metzler, J. (1971). Mental rotation of three-dimensional objects. *Science*, 171, 701–703.
- Trafton, J. G., Kirschenbaum, S. S., Tsui, T. L., Miyamoto, R. T., Ballas, J. A., & Raymond, P. D. (2000). Turning pictures into numbers: Extracting and generating information from complex visualizations. *International Journal of Human Computer Studies*, 53(5), 827–850.
- Trafton, J. G., Trickett, S. B., & Mintz, F. E. (2001). Overlaying images: Spatial transformations of complex visualizations. Paper to be presented at Model-Based Reasoning: Scientific Discovery, Technological Innovation, Values in Pavia, Italy.
- Trickett, S. B., Fu, W., Schunn, C. D., & Trafton, J. G. (2000a). From dippy-doodles to streaming motions: Changes in representation in the analysis of visual scientific data. In *Proceedings of the Twenty Second Annual Conference of the Cognitive Science Society*.
- Trickett, S. B., Trafton, J. G., & Schunn, C. D. (2000b). Blobs, dippy-doodles and other funky things: Framework anomalies in exploratory data analysis. In *Proceedings of the Twenty Second Annual Conference of the Cognitive Science Society*.

²This model does not show the other ways to gather information because it did not show up as a major component in our current datasets.