

FILE COPY

2

SECURITY CLASSIFICATION OF THIS PAGE

DOCUMENTATION PAGE

Form Approved
OMB No. 0704-0188

AD-A216 472 IC

2b. DECLASSIFICATION/DOWNGRADING SCHEDULE

JAN 05 1990

4. PERFORMING ORGANIZATION REPORT NUMBER(S)

1b. RESTRICTIVE MARKINGS

3. DISTRIBUTION/AVAILABILITY OF REPORT

APPROVED FOR PUBLIC RELEASE;
DISTRIBUTION UNLIMITED

5. MONITORING ORGANIZATION REPORT NUMBER(S)

AFOSR-TR-89-1885

6a. NAME OF PERFORMING ORGANIZATION
DEPARTMENT OF PSYCHOLOGY
STANFORD UNIVERSITY

6b. OFFICE SYMBOL
(If applicable)
NL

7a. NAME OF MONITORING ORGANIZATION

(NL)
AIR FORCE OFFICE OF SCIENTIFIC RESEARCH

6c. ADDRESS (City, State, and ZIP Code)
BUILDING 420, JORDAN HALL
STANFORD, CA 94305-2130

7b. ADDRESS (City, State, and ZIP Code)
BOLLING AIR FORCE BASE, DC 20332-6448

8a. NAME OF FUNDING/SPONSORING
ORGANIZATION
AFOSR

8b. OFFICE SYMBOL
(If applicable)
NL

9. PROCUREMENT INSTRUMENT IDENTIFICATION NUMBER

AFOSR-89-0076

8c. ADDRESS (City, State, and ZIP Code)

BUILDING 410
BOLLING AFB, DC 20332-6448

10. SOURCE OF FUNDING NUMBERS

PROGRAM ELEMENT NO.	PROJECT NO.	TASK NO.	WORK UNIT ACCESSION NO.
61102F	2313	A4	

11. TITLE (Include Security Classification)

INDUCED PICTORIAL REPRESENTATIONS

12. PERSONAL AUTHOR(S)

TVERSKY, BARBARA

13a. TYPE OF REPORT
FIRST ANNUAL TECHNICAL

13b. TIME COVERED
FROM 11-1-88 TO 11-30-89

14. DATE OF REPORT (Year, Month, Day)
9 DECEMBER 5, 1989

15. PAGE COUNT
47

16. SUPPLEMENTARY NOTATION

17. COSATI CODES

FIELD	GROUP	SUB-GROUP
05	09	

18. SUBJECT TERMS (Continue on reverse if necessary and identify by block number)

MENTAL MODELS, VISUAL COGNITION, MENTAL MAPS,
DISCOURSE PROCESSING

19. ABSTRACT (Continue on reverse if necessary and identify by block number)

Researchers agree that mental representations of discourse are established at many levels, including a model of the situation described by the discourse. I describe two sets of studies investigating spatial properties of mental models induced by text. In the first set, Holly Taylor and I have found that descriptions written from different perspectives, route and survey, seem to induce the same perspective-free spatial models termed spatial frameworks. In the second set of studies, Nancy Franklin and later David Bryant and I have gathered detailed data on a spatial framework capturing a common situation, of an observer surrounded by objects. That spatial framework is not perception-like, but rather reflects conceptions of space. Extensions of both paradigms are discussed. This early research indicates that situation models constructed from text contain spatial properties, such as relative locations and directions, but are not perception-like or image-like. They are more general than a particular view, allow different perspectives, and have different

20. DISTRIBUTION/AVAILABILITY OF ABSTRACT

☒ UNCLASSIFIED/UNLIMITED ☐ SAME AS RPT. ☐ DTIC USERS

21. ABSTRACT SECURITY CLASSIFICATION

UNCLASSIFIED

22a. NAME OF RESPONSIBLE INDIVIDUAL

ALFRED R. FREGLY, PH.D.

22b. TELEPHONE (Include Area Code)

(202) 767-5021

22c. OFFICE SYMBOL

NL

90 01 04 147

19. ABSTRACT (Continued)

— tial access to different parts. *Figure 10 shows a typical cross-section of a*

Induced Pictorial Representations

Barbara Tversky

Stanford University

Send correspondence to:
Barbara Tversky
Department of Psychology, Bldg. 420
Stanford University
Stanford, California 94305-2130

0110
COPY
XEROX
10/10/89

Accession For	
NTIS CRA&I	☒
DTIC TAB	☐
Unannounced	☐
Justification	
By	
Distribution/	
Availability Codes	
Dist	Avail and/or Special
A-1	

This research was supported by AFOSR grant 89-0076 to Stanford University.

Abstract

Researchers agree that mental representations of discourse are established at many levels, including a model of the situation described by the discourse. I describe two sets of studies investigating spatial properties of mental models induced by text. In the first set, Holly Taylor and I have found that descriptions written from different perspectives, route and survey, seem to induce the same perspective-free spatial models termed spatial frameworks. In the second set of studies, Nancy Franklin and later David Bryant and I have gathered detailed data on a spatial framework capturing a common situation, of an observer surrounded by objects. That spatial framework is not perception-like, but rather reflects conceptions of space. Extensions of both paradigms are discussed. This early research indicates that situation models constructed from text contain spatial properties, such as relative locations and directions, but are not perception-like or image-like. They are more general than a particular view, allow different perspectives, and have differential access to different parts.

Induced Pictorial Representations

Barbara Tversky
Stanford University

There are many simple, everyday tasks, such as following road directions, using instructions to assemble a bicycle or reading a novel, that seem to entail constructing a spatial mental model from a description. In order to comprehend "Go straight till the first light, then turn left, go down about 3 blocks to Oak, and make a right," it is useful to have a spatial representation. Of course, the gist of the message could be remembered instead, but incorporating the instructions into a cognitive map helps, especially when things don't quite turn out the way expected, such as encountering a "No Left Turn" sign at the light. Indeed there is evidence that people do construct such spatial models. The nature of these models, and the conditions under which they are constructed are the topics of this proposal.

Ample research in memory and comprehension of text supports the assertion that listeners or readers form multiple representations of text, of sound or graphemic properties, of actual words or sentences, of gist, and finally of the situation described by the text (Bransford, Barclay, & Franks, 1972; Garnham, 1981; Johnson-Laird, 1983; van Dijk & Kintsch, 1983, among others). In order to make headway, rather than studying models constructed from abstract discourse, we have selected the domain of spatial environments. People have considerable experience converting spoken or written communications about environments into mental representations, and then acting on them. People then get feedback--they either get lost or they find their ways-- and can correct their models. In addition, there is a large body of data on how people learn and

remember environments from experience or from maps that can be compared to acquiring environments from descriptions.

We have developed two separate but related experimental paradigms to investigate mental spatial models constructed from text. In the first paradigm, we vary characteristics of the descriptions and observe the consequent mental models. This work has been done with Holly Taylor, a graduate student. In the second paradigm, we examine in great detail the spatial characteristics of a particular but very common situation, the one people are in most of the time, of having objects at different places around them. This work has been done with Nancy Franklin, a former graduate student, now an Assistant Professor at SUNY-Stony Brook. More recently, a current graduate student, David Bryant, has joined us. For each of the paradigms, I will first present the results we have obtained so far and then the research in progress.

These studies have several goals. First, to demonstrate that the mental models constructed from text with no special instructions to image nevertheless reflect spatial properties described in the text. Next, to discover which spatial properties are preserved, and how they are organized and accessed. At the same time, to investigate characteristics and organization of text in the generation of spatial mental models. Studies by Foos (1980), Mani and Johnson-Laird (1982), Ehrlich and Johnson-Laird (1982), and Perrig and Kintsch (1985) have shown that when descriptions are complete and coherent, readers' mental models preserve information about the spatial relations among objects in scene. Studies by Morrow, Greenspan and Bower (1987), Glenberg, Meyer and Lindem (1987) and Morrow, Bower and Greenspan (1989) indicate that some distance

information described in text is preserved in mental representations. Our first set of studies addresses the issue of the generality and perspective of spatial models. Specifically, are they like structural descriptions (e. g., Marr, 1982; Minsky, 1975; Palmer, 1977; Pinker, 1984), that is, perspective-free accounts of the spatial relations of parts of the scene that allow viewers to take different perspectives on them, or are they like images (e. g., Kosslyn, 1980; Shepard and Podgorny, 1978), that is, internalized perceptions, representing a scene from a particular viewpoint. Our second set of studies investigates representation and access of particular spatial relations from particular perspectives.

Part I: Survey and Route Descriptions.

When tourists visit a new place, they often buy guide books to let them know what is worth seeing and doing, and how to get there. An informal review of guide books reveals that they tend to adopt one of two perspectives on the place described. Some take the reader on a mental tour or *route* through the environment. A route description of the Smithsonian in Washington, D. C. might proceed: As you leave the Capitol going along the Mall, the first building you pass on your right is the East Wing of the National Gallery. Continuing on, you come to the main building of the National Gallery. On your left, across the Mall, you can see the Air and Space Museum....until you reach the Washington Monument." Another perspective commonly adopted is to give the reader a bird's eye view or *survey* of the place. A survey description of the same scene might proceed: At the east end of the Mall stands the Capitol and at the west end, the Washington Monument. Going from east to west along the north side of the Mall, the

first building is the East Wing of the National Gallery, then the National Gallery.... On the south side of the Mall, the eastern most building is the Air and Space Museum, directly south across the Mall from the National Gallery...."

Survey descriptions take a perspective from above and describe the locations of landmarks to one another in canonical direction terms, north, south, east, and west. Survey descriptions are hierarchical, and in general, begin with an overview of boundaries or large-scale regions, and become more specific. Route descriptions take the perspective of a moving observer in the environment, and describe the locations of landmarks relative to the observer's changing position in terms of observer's left, right, in front, and behind. Route descriptions are at a single level of analysis whose sequence is determined by the particular path. The initial question we asked is: Do route and survey descriptions lead to different mental representations? That is, do the representations generated by each perspective preserve that perspective, or are they perspective-free? This is a question of more generality than just spatial models as route and survey descriptions are appropriate for other topics as well, such as descriptions of time.

Experiment 1-1: Route vs. Survey Descriptions.

For our experiments we developed four fictitious environments, two large-scale, one county-sized and the other a small town, and two small-scale, a zoo and a convention center, containing from 11 to 15 landmarks each. Depictions of these environments are in Figures 1-4, but subjects in the initial experiments did not see these maps.

Insert Figure 1 about here

Insert Figure 2 about here

Insert Figure 3 about here

Insert Figure 4 about here

We wrote a survey and a route description of each environment. The survey descriptions took a perspective from above and used a hierarchical organization and canonical *directions terms to describe landmarks in relation to each other*. The route descriptions took a perspective from within the environment and used a sequential organization and egocentric direction terms to describe landmarks in relation to a moving ego. There is no widely-applicable measure of discourse coherence. Co-reference, that is, linking sentences in sequence by referring to the same thing, has sometimes been suggested (Johnson-Laird, 1983; van Dijk & Kintsch, 1983). Co-reference may be appropriate for route or sequential organizations, but not for hierarchical descriptions, where a new

descriptive part will refer back to the overview, but not to the previous sentence. Lacking an objective measure, we asked a group of pilot subject to evaluate the coherence of our texts, and they reported that the two types of descriptions were equally coherent. We also pretested our descriptions to make sure that readers could correctly place all landmarks in sketches. In addition to the locative information, each description contained non-locative information, for example, telling activities that could be performed in different parts of the environment, or giving elaborative details about landmarks. This information was identical for route and survey descriptions. The route and survey text for the town are presented in Figures 5 and 6.

Insert Figure 5 about here

Insert Figure 6 about here

We modeled our memory tasks on those of Perrig and Kintsch (1985) who tested a similar hypothesis but got inconclusive results, partly because their descriptions were too difficult and partly because their survey descriptions did not completely describe the environments. Subjects read two route and two survey descriptions, one large-scale and one small-scale environment for each description type. Over subjects, each environment was presented equally often as a route and as a survey description. Subjects could read each description up to 4 times each. Reading time was self-paced, and total times were

recorded. After reading each description, subjects were presented with statements to verify as true or false; reaction time and errors were recorded. Some statements tested the nonlocative information. Perspective should make no difference on performance on these questions. Other statements tested the locative information. These were either verbatim from the text or statements that could be inferred from the text and were from the same global perspective of the text, but that were not explicitly stated in the text. Half of both the verbatim and inference locative statements were from a route perspective and half from a survey perspective. Readers answered all questions regardless of perspective read. Thus, a verbatim statement from a different perspective was in effect an inference statement for that reader. Verbatim questions should be faster than inference questions if readers are using representations of the gist or propositions. In general, retrieval of linguistic or propositional information is faster than retrieval from images or mental models (Kosslyn, 1976). If the situation models readers construct depend on the particular perspective of the narrative, then readers should respond faster to inference statements from the perspective read than to inference and verbatim statements from the other perspective. If, however, readers construct the same spatial mental models irrespective of the perspective of the text, then there should be no differences on the inference questions that depend on perspective read. Following the questions, readers drew a map of each environment. This serves to check that readers were able to form an integrated and correct spatial model from the text.

Route maps took slightly but significantly longer to read.

Insert Figure 7 about here

Subjects made more map errors on route descriptions (1.31) than on survey descriptions (.68), but there were very few errors made on maps altogether. The data of primary interest are the reaction times and error rates to the different types of questions. These followed the same pattern, and are presented in Figures 8 and 9.

Insert Figure 8 about here

Insert Figure 9 about here

As in the case of the maps, overall performance was excellent. First, there were fewer errors and faster reaction times to verify the non-locative statements than locative statements, and perspective had no effect on performance on non-locative statements. We would not like to claim that non-locative information is generally easier than locative; surely one could write nonlocative statements that would be very difficult to remember. As for locative statements, the only differences to be found are on the verbatim statements. Subjects were faster and more accurate verifying statements that they had actually read than statements about inferences from information presented in the descriptions. Specifically, subjects were not faster or more accurate on inference

statements from the perspective read than from the other perspective, for both perspectives.

Characteristics of Spatial Representations from Text

The lack of any differences on verification of survey and route inference statements as a consequence of type of description read suggests that this information is verified from a situation model, and that the situation model constructed does not depend on the perspective of the text. Because readers are just as good taking a new perspective as taking a previous perspective, their situation models must be general enough to allow the taking of different perspectives with equal ease. What does such a representation look like? We would like to speculate that it doesn't look like anything that can be visualized. Rather, it is like an architect's 3-D model of a town; it can be viewed or visualized from many different perspectives, but it cannot be viewed or visualized as a whole. In fact, answering the locative questions required taking a particular perspective. Thus, this kind of representation differs from the representations proposed in the classic work on imagery (e. g., Finke & Shepard, 1986; Kosslyn, 1980; Shepard & Podgorny, 1978) which are perception-like and from a particular point of view.

Experiment 1-2: Drawing Order.

A serendipitous finding emerged from the first experiment which we have begun to chase. The experimenter noticed that subjects seemed to be drawing the landmarks of the maps in the order they had been mentioned in the descriptions. Half-way through the study, she began recording the drawing order, and analysis of those data confirmed the

hypothesis. To make absolutely sure there was no bias in this observation, two new experimenters and a new set of subjects were recruited. The experimenters were told to record the order of drawing the landmarks, but they did not know which type of description subjects read each trial nor the general hypothesis. Some of the descriptions were rewritten to make the route and survey orders as different as possible. The results of this third study confirmed those of the first. The correlation between description order and map drawing order was very high ($r = .72$), significantly higher than the correlation between other description order and map drawing order ($r = .22$). Approximately the same correlations were found in Experiment 3. The patterns of errors and reaction times found in Experiment 1 were replicated. That is, non-locative statements were faster and more accurate, and the only differences attributable to description perspective were in the verbatim statements; inference statements did not differ.

At first, finding correspondences between description order and map drawing order may seem contrary to the findings for speed and accuracy in verifying statements from the two types of text, where we found no differences except in verbatim or paraphrased statements. Now, for map drawings, we find a large difference due to description in a task that seems to depend on drawing a mental image or cognitive map. It is that latter assumption that we question. We argued earlier that the mental representations subjects construct are not like images, but rather like structural descriptions of the relations between the landmarks of the scene. As such, unlike a mental image or cognitive map, they cannot be visualized. They can be imagined from a particular perspective, but not in entirety. So, we conjecture, when subjects are asked to draw a map of the environments, they do so by reconstructing their mental models, and

they reconstruct in the same order as they originally constructed, that is, in the order of the description they read. Thus, this finding is indirect evidence in support of the contention that readers' mental models are not image- or map-like. If they were, there should be no differences in drawing order depending on description perspective; rather, drawing order should depend on characteristics of the image or map.

Experiment 1-3: Verbatim vs. Paraphrased Test Statements.

The superior performance on verbatim statements over inference statements suggests that these statements may be verified against memory for text per se. We performed a third study with a nearly identical design except for the addition of paraphrased verification statements. The paraphrased statements behaved just like verbatim statements, a common finding in memory for text (references?). It is widely accepted that text representations are for gist rather than exact wording. All of the other findings of the first experiment were replicated in this third study.

The studies completed so far seem to answer some questions, but raise others; the experiments in progress address those issues.

In Progress: Single Descriptions, No Maps.

It is possible that readers of route and survey descriptions form the same mental representations because they know they will be required to draw maps and because after one trial they know that they will verify both route and survey statements. We analyzed the first trial data of our subjects. At this point, subjects did not know ahead of time that they would be asked to verify statements from both route and survey perspectives.

Nevertheless, there were no differences in errors or reaction time to either route or survey statements as a function of the perspective of the narrative. These readers, however, knew that they would be drawing maps. Thus, we plan a short study on a large number of subjects, where each subject will be given general comprehension instructions and not informed about the map task; then they will read just one description, route or survey, and answer both types of questions. If there are still no differences between the two groups, then the conclusion that readers form perspective-free representations in comprehending text describing spatial relations among landmarks is strengthened.

In Progress: Studying Maps vs. Studying Description

As an addition to the first study, we tested a group of subjects who had studied maps of the environments rather than descriptions. Their speed and accuracy on the questions resembled that of the groups that had studied descriptions (see Figures 10 and 11).

Insert Figure 10 about here

Insert Figure 11 about here

Their order of drawing landmarks in maps contrasted to the drawing order of those who had read descriptions, whose orders corresponded to the order of mention in the

descriptions. Subjects who had studied maps tended to draw maps in a hierarchical fashion, beginning with borders or large entities and working inwards or smallwards. This pattern could be due to memory organization, that is spatial memory may be hierarchically organized (Stevens & Coupe, 1978; McNamara, 1986), or it could be due to demands of the drawing task, where bordering or larger elements set the scale for internal and smaller elements, or both. It could be that in writing descriptions of environments, subjects are implicitly aware that they are constructing mental models in the minds of readers, and that many of the constraints of model construction are similar to those of actual construction (Novick & Tversky, 1987). One way of separating these two explanations is to ask subjects to write descriptions of the maps. Of course, this is not a clean test of memory organization because the description task brings with it its own demands, namely, linearization. Linde and Labov (1975) studied people's descriptions of their apartments, and Levelt (1984) studied people's descriptions of map-like networks. Both groups found that subjects took readers on a mental tour of the environments. They attributed this to the linear characteristic of language as opposed to pictures and environments. Linde and Labov's and Levelt's findings and contentions notwithstanding, our subjects tended to organize their descriptions hierarchically, especially for the large-scale environments. There were more route descriptions for the convention center, one of the small-scale environments. In a study asking subjects to describe different arrangements of furniture, Ehrich and Koster (1983) found that when there was a natural hierarchical arrangement of the furniture, subjects' descriptions were hierarchical, but when the arrangement was random, subjects' descriptions were route-like.

We plan to run a larger study comparing organization of drawing and descriptions as a function of type of environment and task expectations, drawing or description. One of the large scale maps (the town) and one of the small scale maps (convention center) were selected for this experiment because they more easily lend themselves to both route and survey descriptions. Subjects will study just one map under one of two expectations: drawing a map or writing a description. They will be asked to perform both tasks (half the subjects in each order). The dependent variable will be the organization of their depictions or descriptions inferred in large part from the order in which they draw/describe different landmarks. The previous research and pilot work described above lead us to three predictions, about the maps, the expectations, and the task. Over all, hierarchical descriptions and depictions should predominate, but there should be relatively more route descriptions for the small-scale map, under description expectations, and for the description task.

Spatial Frameworks.

In the previous studies, we examined spatial mental models constructed from text in a very general way, and discovered that such models are quite general, that is, perspective-free, but allowing many possible perspectives. These were models of large or small geographic areas. Although we only examined the gross relative spatial positions of landmarks, investigation of more fine-grained and other relations is possible. In the next set of studies, we begin to examine the details of spatial mental models for specific situations and relations. In the first part, we suggested that spatial mental models derived from text are like structural descriptions, that is, a collection of parts and their

interrelationships. We will call these *spatial frameworks*. Now we begin to explore how parts of spatial frameworks are organized, what relations are preserved, and how information is accessed from them. To a certain extent, the structure of the mental models and the processes of retrieving information from them are inseparable (Anderson, 1978).

We begin with a simple and familiar situation, precisely because it is simple and familiar, that of a person surrounded by objects, and keeping track of those objects as the person's position changes. In the first set of studies, the movement is very simple as well, standing and turning in place, and lying down and turning in place. The locations of objects are kept simple as well, extending from the three body axes, at head or feet, in front or back, to the left or right. Later studies will complicate this situation in a variety of ways. In each case, we propose that in comprehending narratives describing spatial relations among objects, readers construct mental spatial frameworks. The spatial frameworks reflect the spatial relations described, and objects are attached to the appropriate places. The frameworks may be updated as the situation changes. In order to answer questions about objects or relations, people take the appropriate perspective in or on the framework, and retrieve the relevant information.

One-Place Spatial Framework: Experiments 2-1 to 2-5.

Keeping track of the objects around us, even those that cannot currently be seen, as we navigate through an environment, is a natural task that people seem to perform easily. In the first set of studies such environments are described in narratives, and then subjects are queried about objects in particular locations. In these studies, conducted

with Nancy Franklin (Franklin & Tversky, 1990), the texts describe naturalistic scenes in which the observer ("you" in the narrative) remained in one place and objects were located in the six canonical directions around the observer, above, below, in front, in back, to the left, and to the right. The protagonist was periodically reoriented to face a different object and queried about the objects in all possible locations. A three-dimensional task in which the observer is embedded in the scene was selected to encourage formation of a spatial model, rather than some other, for example, list-like representation.

In a typical experiment, subjects read 10 different scenarios containing 5 critical objects each. One scenario described an evening at the Opera, where the objects were: plaque, loudspeaker, sculpture, lamp, and bouquet. Other scenes were in a hotel lobby, at a Halloween party, in a lagoon, on a ship, and more, each with appropriate objects. Objects were selected so that they could reasonably appear in any position, and positions were determined randomly. Subjects read part of the scenario from a piece of paper, and then turned to a computer screen to read the continuation of the narrative which oriented and reoriented readers in front of specific objects, and presented the questions. Between questions was more filler material to eliminate any priming effects.

Questions were one word terms, such as "left" or "front." It should be noted that "head" and "feet" were used as questions instead of (or as well as) "up" and "down." That way, all the questions were with respect to the central being or object and not with respect to the environment. This also allowed us to ask unambiguous questions about reclining observers (in parallel experiments, one asking "up" or "down" and the other

asking "head" or "feet", identical patterns of data were obtained). Subjects were instructed to press a button to see the alternatives as soon as they were sure which object was in the indicated direction. These R1 reaction times were the ones which entered the analysis. Subjects then selected the correct alternative from a randomized list of all alternatives; these R2 reaction times did not vary systematically, indicating that subjects followed instructions about R1 times.

Three classes of models to account for the response times were considered. According to the *equiavailability* model, all locations are equally available to the observer, as they would be in a picture, or if there were no bias toward one direction or another (Levine, Jankovic, & Palij, 1982). According to the *mental transformation* model, the observer imagines him or herself facing the stipulated direction, and then mentally turning to face the cued direction to verify the object. It is as if the reader were viewing the scene and turning to inspect the cued direction to see what object is there. This model is closest to the classic models of imagery (e. g. Kosslyn, 1980; Finke & Shepard, 1986) that are based on internalized perception. In this case, reaction times should increase with the objects' angular disparity from straight-ahead. This "mental rotation" is similar to that studied by Shepard and Cooper (1982) in that the observer imagines him/herself perceiving an environment, but different in that the observer imagines him/herself turning rather than imagining an object rotating. Just as mental rotation has been demonstrated to be an analog process, one might expect mental turning to have that same characteristic. On the other hand, Hintzman, O'Dell and Arndt (1981) asked subjects to point to real or imagined objects arrayed in a circle around the subject. They found correlations between degree of rotation and reaction time for a real scene but

not for an imagined scene, concluding that the imaginal representations are "not strictly holistic, but consist of orientation-specific representations, and --at least in part--of relational propositions specific to object pairs(p. 149)."

Neither the equiavailability nor the mental transformation model was supported by the data. There were reliable differences in reaction times depending on the direction cuing the object, but these differences did not correlate with angle of disparity from straight-ahead. Instead, the pattern of reaction times in the 5 experiments conformed to what we called the *spatial framework* model, to a conceptual spatial model based on perceptual experience, but not based on the perception of experiencing.

According to the one-place spatial framework model (influenced by writings of Clark, 1973; Levelt, 1984; Shepard and Hurwitz, 1984; Talmy, 1973, amongst others), the canonical position of the observer is upright, and the canonical world of the observer can be described by one vertical and two horizontal dimensions. The vertical dimension is correlated with gravity, an important asymmetrizing factor in the world; furthermore, in canonical movement, vertical spatial relations generally remain constant with respect to the observer, but horizontal spatial relations change. Whereas there are environmental reference points for the vertical dimension, the sky and the ground, for example, the reference points for the two horizontal dimensions are more arbitrary, often the prominent dimensions of the observer's body. Thus, for the upright observer, the vertical dimension is predominant. Of the two horizontal dimensions, the front/back dimension predominates over the left/right. The former is asymmetric perceptually and functionally: the observer can more readily see, attend to , and move toward the *front*

than the *back*. The left/right dimension, in contrast, is derived from the front/back, and has no salient asymmetries. Thus, for the upright observer, this model predicts that the vertical dimension should be fastest, followed by front/back, followed by left/right. In addition, it predicts that front should be faster than back. In three experiments, exactly this pattern of data was obtained (see Table 1).

Insert Table 1 about here

When the observer/protaganist is reclining, the vertical dimension no longer corresponds to any of the axes of the body. In that case, according to the spatial framework model, the relative salience of the body axis determines the speed of accessibility. Clearly, the left/right axis is least salient, having no asymmetries and being dependent for definition on the front/back axis. Both the front/back and the head/feet axes have asymmetries; however, given the reclining position, and the difficulty moving, the front/back axis which is the site of most of the important perceptual apparatus is more salient than the head/feet axis. Thus, the spatial model predicts that for the reclining observer, accessing objects along the front/back axis should be fastest, followed by the head/feet, and left/right last. Again, the obtained results of two experiments match that pattern (see Table 1). Overall, reaction times for the reclining observer were longer than those for the upright observer, indicating greater difficulty imagining the reclining case.

These experiments demonstrated that a complex mental spatial framework can be constructed from descriptions, that that framework for representing locations of objects in

space is biased and not equipotential, that the source of bias can be understood in terms of a perceptual analysis of the characteristics of the perceptual and functional apparatus of the observer in relation to the world. Unlike the classic models of imagery, this spatial framework is not based on internalized perception, but rather is a construction based on our understanding of the spatial world.

Extensions of the One-Place Spatial Framework.

How general and robust is this spatial framework? David Bryant, Nancy Franklin and I have explored the generality of this model in two directions. The first is alternative perspectives. The first five experiments used second-person narratives to draw readers into the text. We wondered whether readers would adopt the perspectives of a third person and of an object. The one-place spatial framework is presumed to be organized around body axes, and readers were probed for objects with direction cues corresponding to the presumed axes. Would the spatial framework pattern emerge if readers were cued with objects for directions?

Experiment 2-6: Third-Person Perspective.

We rewrote the narratives of the fifth experiment, using both upright and reclining protagonists, so that they were in the third person, half around a male protagonist and half around a female protagonist. The same pattern of data appeared in this study as in the previous five studies, as is clear from Table 2.

Insert Table 2 about here

Although some readers reported that it was easier to identify with a same-sex protagonist, analysis of reaction times did not support this assertion.

Experiment 2-7: Object-centered Narratives.

Again, we rewrote the narratives, this time using an object in the center, and having an invisible agent reorienting the object. For example, there is a saddle in the middle of the barn; in front of it is a feedbag, to its left is a shovel; it is being turned from side to side in order to be cleaned. All questions were from the point of view of the central object, and "top" and "bottom" were substituted for "head" and "feet" because this is how those parts are referred to. When the object was upright, readers were able to take the perspective of the object quite easily; the upright pattern of data resembles that of previous studies (see Table 3).

Insert Table 3 about here

When the object was reclined, however, readers encountered difficulties, especially with "top" and "bottom." Another study is underway instructing readers to take the point of view of the central object to see if this will improve reclining performance.

Experiments 2-8 and 2-9: Object Cues.

According to the spatial framework model, readers construct a mental scaffolding in the canonical directions around themselves and hang objects in the appropriate places. If the mental scaffolding is primary, then giving object cues for directions should yield the same pattern of data as giving direction probes for objects, as in the previous studies. Spontaneous descriptions of environments follow this order, first location, then object (Ehrich and Koster, 1983). In solving geometric analogies where order of operations is not constrained, subjects prefer to first ascertain location of figures, and then nature of figures. Moreover, when subjects are required to solve the analogies in the opposite order, they make more errors and take more time (Novick & Tversky, 1987). Thus there is some reason to expect that direction cues should yield faster reaction times than object cues.

The results of two experiments, the first using only object probes, and the second using both object and direction probes but in separate narratives, support the first hypothesis, that the one-place spatial framework pattern of data emerges for object probes (see Table 4).

Insert Table 4 about here

Overall reaction times were slower in the first of these studies, using only object probes, lending some support for the proposal that direction probes should be faster than object probes. However, no differences were found between direction and object probes in the

study comparing these within subject (see Table 5).

Insert Table 5 about here

Thus, the spatial framework is organized around directions, but direction probes for objects do not seem to be more efficient than object probes for directions.

In Progress: External Perspective.

Thus far, the observer's view point has been in the center of an array of objects surrounding the observer. Two sets of studies are in progress giving the observer an external point of view, and asking questions from this point of view. In the first, the observer is looking onto a cube-like array of objects. The narratives describe each object with respect to some other object using the terms, front/back, left/right, above/below. Thus, the perspective of the observer is an outsider's perspective, and "in front of the car" doesn't mean "in front of the car's front" but rather "toward the observer from the car." Predictions from the spatial framework are for the most part the same: the fastest dimension should be above/below followed by front/back and then left/right, for the same reasons as before. Front/back, however, may lose its asymmetry in this case. Both objects are in front of the observer with an outsider's perspective; the front one is closer than the back one, but this may not be sufficient to create the large advantage to front in reaction time with an insider's perspective.

Another way to get an outsider's perspective is to view what was the observer in the paradigm studies. This essentially entails taking the third-person narratives, and rewriting them so that front/back and left/right are with respect to the outsider rather than with respect to the third person. This should induce readers to adopt an outsider's perspective, and the spatial framework predictions are modified as for the cube study. Readers must alternately take the insider and outsider perspectives to answer questions. Both patterns of reaction times are of interest, as well as the time to switch perspectives.

In Progress: Construction of Framework.

The experiments done and proposed so far have been directed at retrieving information from the spatial framework. Nancy Franklin and I are investigating construction of the spatial framework by looking at reading times for the initial description. Sentence length (by syllable) and sentence structure are standardized, and order of presenting information counterbalanced. Will the pattern of reading times resemble the pattern of retrieval times? That is, will upright times be fastest for head/feet, then front/back, then left/right, and reclining times fastest for front/back, then head/feet, then left/right? And will times be faster for upright than for reclining? If so, this would indicate that the construction of the spatial framework is subject to the same constraints as retrieval of information from it.

Summary and New Directions.

The research described is aimed at uncovering characteristics of mental models formed from text. Mental models of spatial descriptions were selected because on the one hand, they are familiar and well-practiced, and on the other hand, relatively easy to operationalize. In the first set of experiments, subjects read route or survey descriptions of four environments, and verified verbatim and inference statements about those environments both from the read and the other perspective. Subjects were equally fast and accurate in verifying inference statements from the read perspective and the other perspective. This lead us to the conclusion that subjects' mental models capture the spatial relations described in the text, but not from any particular perspective. Rather, unlike images which have been compared to internalized perceptions, they are perspective-free and allow the taking of many perspectives at information retrieval. We termed such models *spatial frameworks*; like structural descriptions, they contain information about the parts of a scene or environment and the relations between the parts.

The second set of studies examined information updating and retrieval from a particular environment, one that is simple and common, that of an observer surrounded by objects. We found that retrieval times to report what objects lie at 6 canonical directions from the observer (at head or feet, to the left or right, in front or back) differ reliably and systematically depending on the direction and the posture of the observer. The particular relations among the retrieval times were explained by an analysis of how space is conceived.

This research is being extended in many directions: to the analysis of other environments, especially dynamic ones; to the construction of mental models; to the construction and use of mental models that are not spatial, and more.

References

- Anderson, J. R. (1978). Arguments concerning representations for mental imagery. *Psychological Review*, 85, 249-277.
- Bransford, J. D., Barclay, J. R., & Franks, J. J. (1972). Sentence memory: A constructive versus interpretive approach. *Cognitive Psychology*, 3, 193-209.
- Clark, H. H. (1973). Space, time, semantics and the child. In T. E. Moore (Ed.), *Cognitive development and the acquisition of language* (pp. 27-63). N. Y.: Academic Press.
- Ehrich, V., & Koster, C. (1983). Discourse organization and sentence form: The structure of room descriptions in Dutch. *Discourse Processes*, 6, 169-195.
- Ehrlich, K. & Johnson-Laird, P. N. (1982). Spatial descriptions and referential continuity. *Journal of Verbal Learning and Verbal Behavior*, 21, 296-306.
- Finke, R. A., & Shepard, R. N. (1986) Visual functions of mental imagery. In K. R. Boff, L. Kaufman & J. P. Thomas (Eds.), *Handbook of perception and human performance*. N. Y.: Wiley.
- Foos, P. W. (1980). Constructing cognitive maps from sentences. *Journal of Experimental Psychology: Human Learning and Memory*, 6, 25-38.

- Franklin, N. & Tversky, B. (1990). Searching imagined environments. *Journal of Experimental Psychology: General*, 119,
- Garnham, A. (1981). Mental models as representations of text. *Memory and Cognition*, 9, 560-565.
- Glenberg, A. M., Meyer, M., & Lindem, K. (1987). Mental models contribute to foregrounding during text comprehension. *Journal of Memory and Language*, 26, 69-83.
- Hintzman, D. L., O'Dell, C. S., & Arndt, D. R. (1981). Orientation in cognitive maps. *Cognitive Psychology*, 13, 149-206.
- Johnson-Laird, P. N. (1983). *Mental models*. Cambridge, Mass: Harvard.
- Kosslyn, S. M. (1976). Using imagery to retrieve semantic information: A developmental study. *Child Development*, 47, 434-444.
- Kosslyn, S. M. (1980). *Image and mind*. Cambridge, MA: Harvard University Press.
- Levelt, W. (1984). Some perceptual limitations on talking about space. In A. J. van Doorn, W. A. de Grind, & J. J. Koenderink (Eds.), *Limits in perception*. Utrecht, The Netherlands: VNU Science Press.
- Levine, M., Jankovic, I., & Palij, M. (1982). Principles of spatial problem solving. *Journal of Experimental Psychology: General*, 111, 157-175.

- Linde, C., & Labov, W. (1975). Spatial structures as a site for the study of language and thought. *Language*, 51, 924-939.
- Mani, K., & Johnson-Laird, P. N. (1982). The mental representation of spatial descriptions. *Memory and Cognition*, 10, 181-187.
- Marr, D. (1982). *Vision: A computational investigation into the human representation and processing of visual information*. N. Y.: Freeman.
- McNamara, T. P. (1986). Mental representations of spatial relations. *Cognitive Psychology*, 18, 87-121.
- Minsky, M. (1975). A framework for representing knowledge. In P. H. Winston (Ed.), *The psychology of computer vision*. New York: McGraw Hill.
- Morrow, D. G., Greenspan, S., & Bower, G. H. (1987). Accessibility and situation models in narrative comprehension. *Journal of Memory and Language*, 26, 165-187.
- Morrow, D. G., Bower, G. H., & Greenspan, S. (1989). Updating situation models during narrative comprehension. *Journal of Memory and Language*, 28, 292-312.
- Novick, L. R. and Tversky, B. (1987). Cognitive constraints on ordering operations: The case of geometric analogies. *Journal of Experimental Psychology: General*, 116, 50-67.

- Palmer, S. E. (1977). Hierarchical structure in perceptual representation. *Cognitive Psychology*, 9, 441-474.
- Pinker, S. (1984). Visual cognition: An introduction. *Cognition*, 18, 1-63.
- Perrig, W., & Kintsch, W. (1985). Propositional and situational representations of text. *Journal of Memory and Language*, 24, 503-518.
- Shepard, R. N., & Cooper, L. A. (1982). *Mental images and their transformations*. Cambridge, MA: MIT Press/Bradford Books.
- Shepard, R. N. & Hurwitz, S. (1984). Upward direction, mental rotation, and discrimination of left and right turns in maps. *Cognition*, 18, 161-194.
- Shepard, R. N. & Podgorny, P. (1978). Cognitive processes that resemble perceptual processes. In W. K. Estes (Ed.), *Handbook of learning and cognitive processes*. Hillsdale, N. J.: Erlbaum.
- Stevens, A. L., & Coupe, P. (1978) Distortions in judged spatial relations. *Cognitive Psychology*, 10, 422-437.
- Talmy, L. (1983). How language structures space. In H. L. Pick & L. P. Acredolo (Eds.), *Spatial orientation: Theory, research, and application*. N. Y.: Plenum Press.
- van Dijk, T. A., & Kintsch, W. (1983). *Strategies of discourse comprehension*. N. Y.: Academic Press.

Figure Captions.

Figure 1. Map of Resort Area

Figure 2. Map of Town

Figure 3. Map of Zoo

Figure 4. Map of Convention Center

Figure 5. Route Text of Town

Figure 6. Survey Text of Town

Figure 7. Study Times for Route and Survey Texts in Experiment 1-1

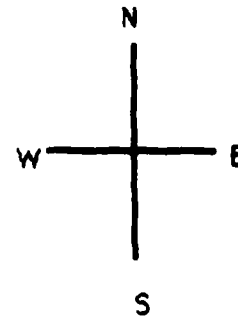
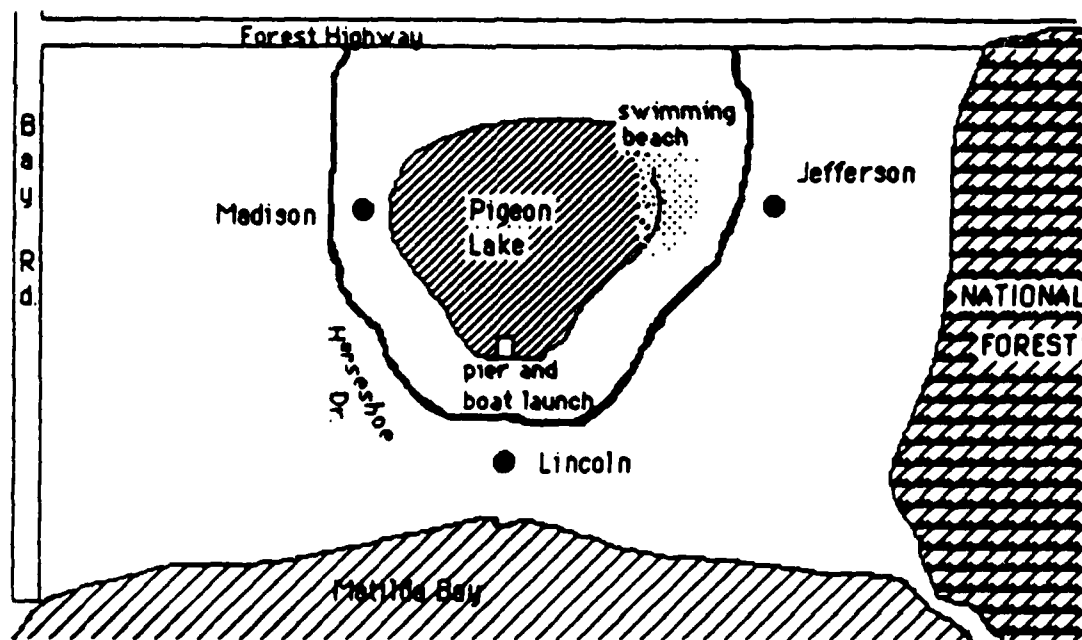
Figure 8. Reaction Times per Syllable in Experiment 1-1

Figure 9. Errors in Experiment 1-1

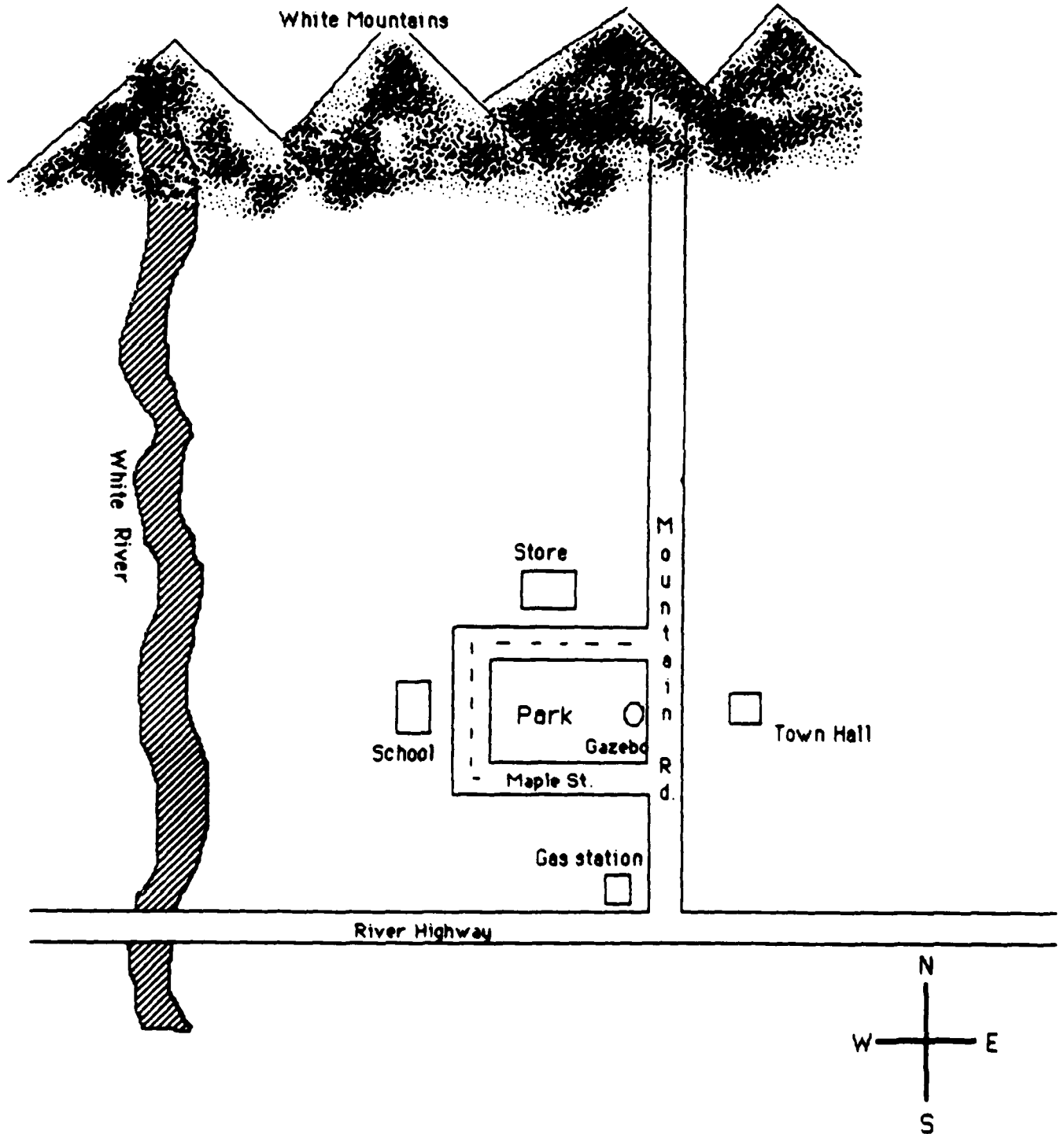
Figure 10. Reaction Times per Syllable in Experiment 1-2

Figure 11. Errors in Experiment 1-2

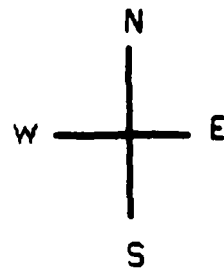
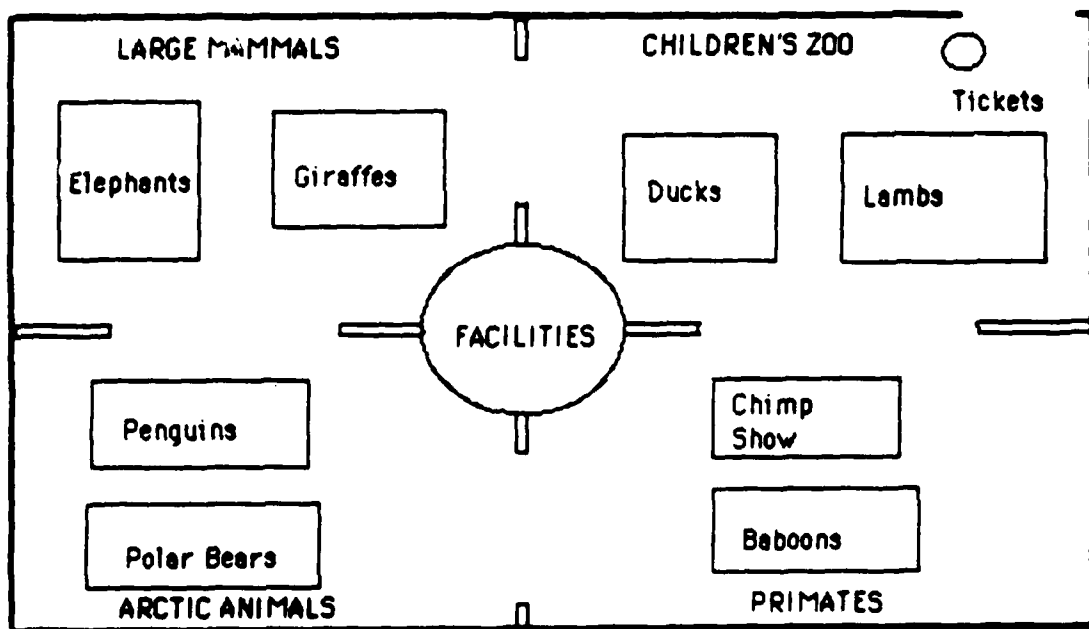
RESORT AREA



TOWN



ZOO



CONVENTION CENTER

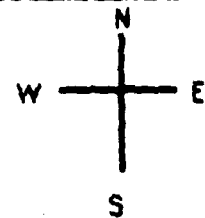
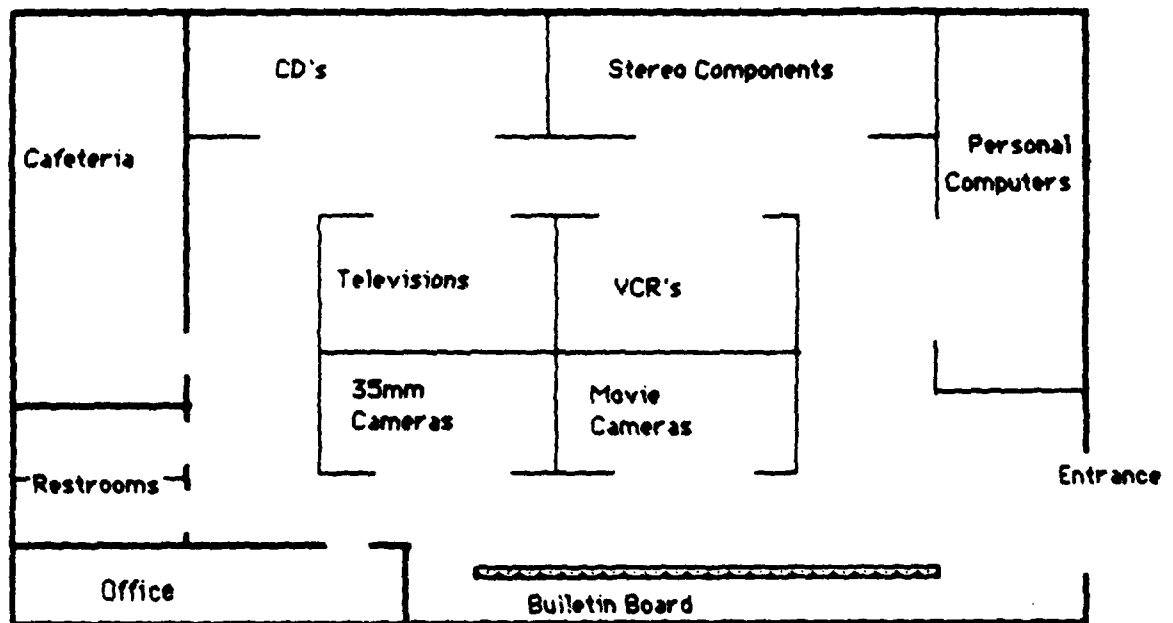


Figure 5

ROUTE DESCRIPTION: TOWN

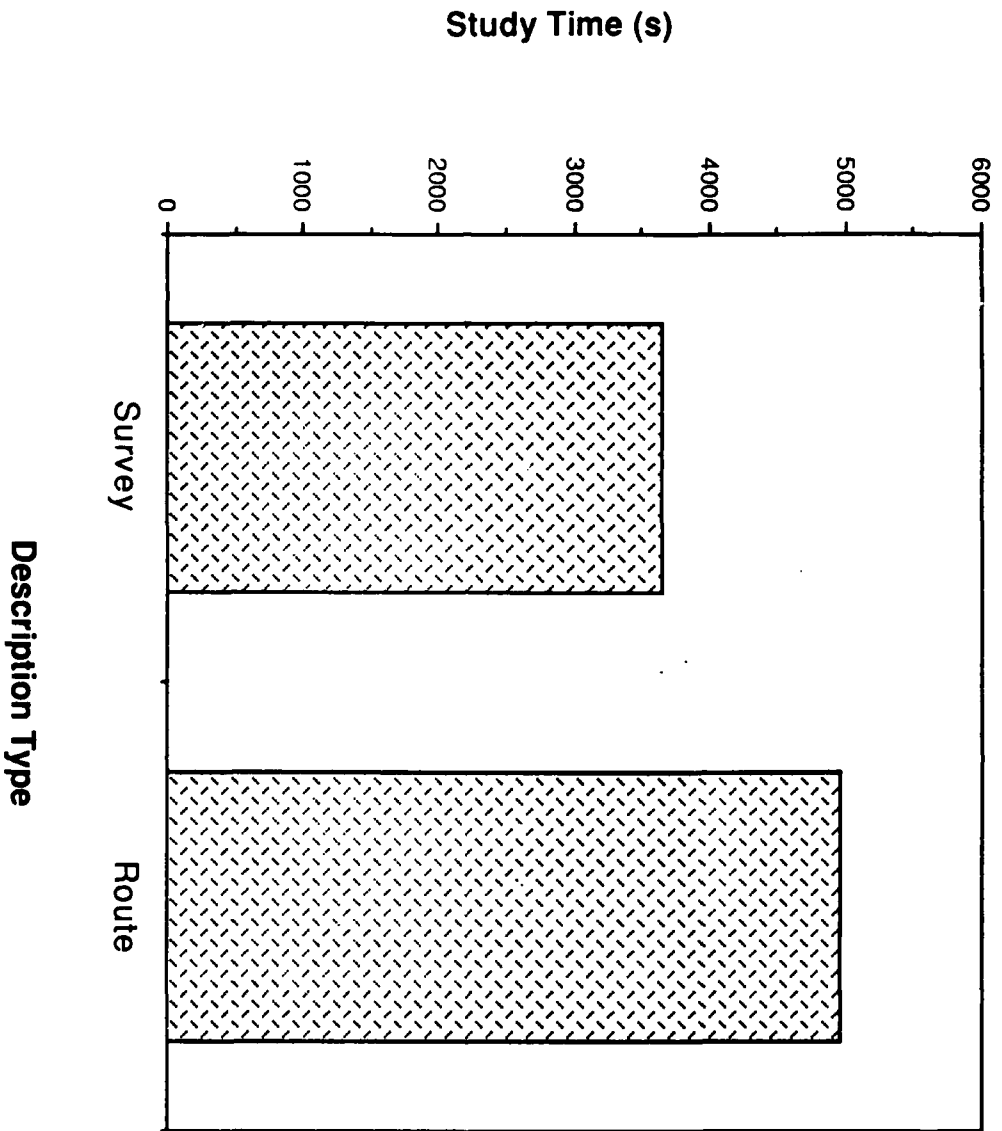
One of the largest town fairs and pumpkin festivals in the United States is held each year in the town of Etna. Etna is a typical small New England town. The lay-out of the town has not changed much since it was founded in the 1700's. To reach Etna, drive east along the River Highway to where the highway crosses the White River. Continuing on the River Highway, for another half mile past the river you come to, on your left, Mountain Rd. You have reached the town of Etna. As you turn left onto Mountain Rd. from the River Highway, you see, on your immediate left, the Gas Station. One of the mechanics from the Gas Station sits in front of the station office and waves to all the cars that drive past. Straight ahead, you can see the road disappearing into the distant White Mountains. You drive on Mountain Rd. a block past the Gas Station, and come to, on your left, Maple St. Turning left onto Maple St., you see that the street is lined with large maple trees. These maples, when they come alive with color in the fall, are an attraction for many tourists. After turning left onto Maple St. from Mountain Rd., you see, on your right, the Town Park--a central feature of Etna. You travel a block on Maple St. and are forced to make a right turn. On your left, about a half a block after you turn off of Maple St., is the School. The little red, one-roomed schoolhouse is the original school built when the town was founded. Continuing along this street for another half a block, you are again forced to make a right turn. You turn and drive a half a block where you see, on your left, the Store. People often gather at the Store to find out the latest town news. This road continues for another half a block where it dead-ends into Mountain Rd. After you make a right turn onto Mountain Rd., you drive about a half a block to where you see, on your left, the Town Hall. The Town Hall is the oldest structure in the town and one of the buildings around which the town was built. From your position with the Town Hall on your left, you see, on your right, a white Gazebo near the edge of the park. The Gazebo is used to house the town band during afternoon concerts. You return to where Mountain Rd dead-ends into the River Highway. You turn left from Mountain Rd. and leave the town of Etna by taking the River Highway.

Figure 6

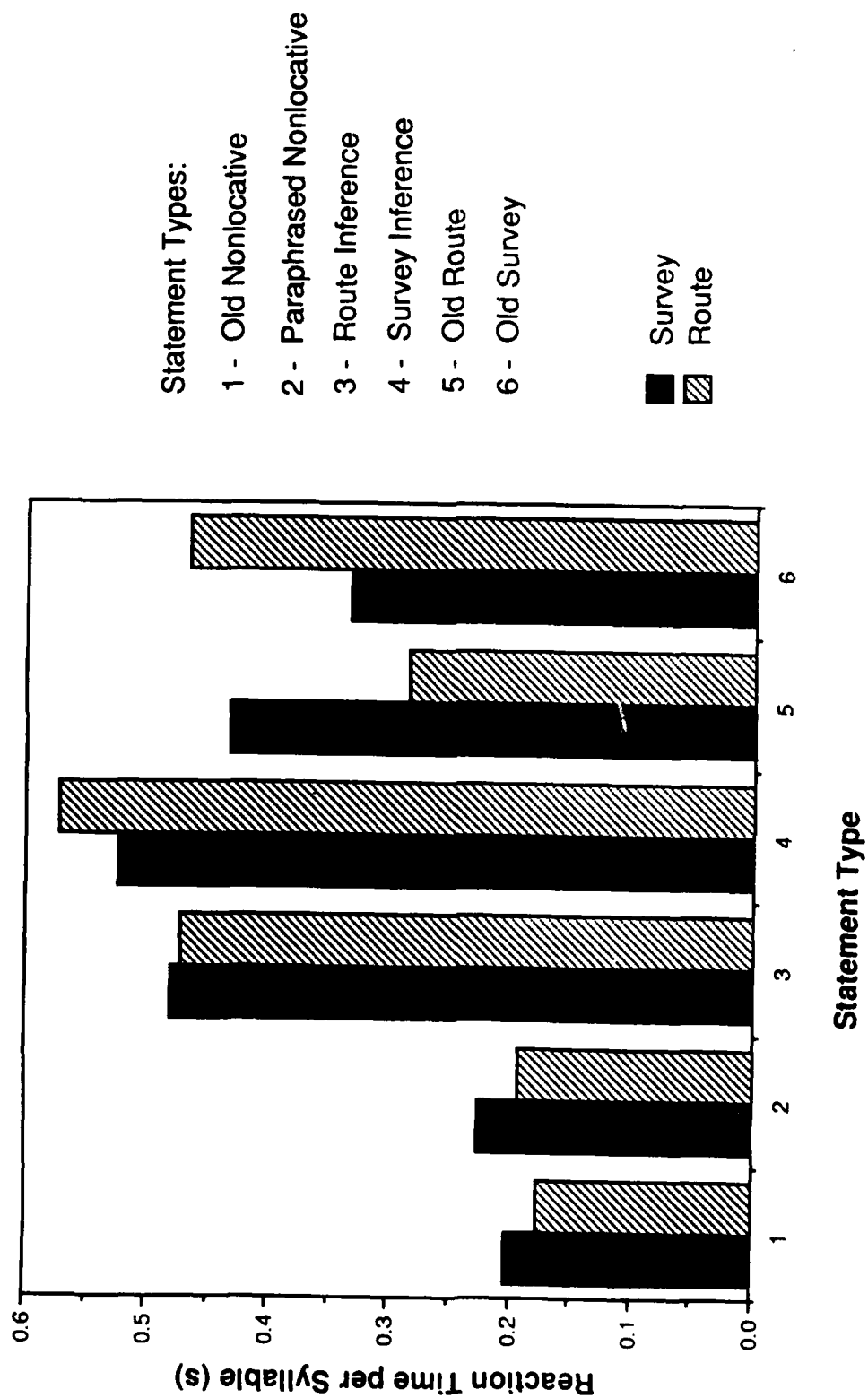
SURVEY DESCRIPTION: TOWN

One of the largest town fairs and pumpkin festivals in the United States is held each year in the town of Etna. Etna is a typical small New England town. The lay-out of the town has not changed much since it was founded in the 1700's. Etna and its surrounding areas are bordered by four major landmarks: the White Mountains, the White River, the River Highway, and Mountain Rd. The northern border is made up of the White Mountain Range. Running north-south along the western border of this region is the White River. The southern border is made up of the River Highway. Along the eastern border, connecting the River Highway to the mountains, is Mountain Rd. Most of Etna lies west of Mountain Rd. just north of its intersection with the River Highway. Etna is built around four streets that surround the Town Park. On the eastern edge of the park, there is a white Gazebo. The Gazebo is used to house the town band during afternoon concerts. Along the eastern edge of the Town Park runs Mountain Rd. The other three streets in Etna are each only a block long. Along the southern border of the park runs Maple St. Maple St. is lined with large maple trees. These maples, when they come alive with color in the fall are an attraction for many tourists. Across the street from the park, on separate sides, lie three of the town's main buildings--the Town Hall, the Store, and the School. Across the street from the east side of the park is the Town Hall. The Town Hall is the oldest structure in the town and one of the buildings around which the town was built. Across the street from the north side of the park is the Store. People often gather at the Store to find out the latest town news. Across the street from the west side of the park is the School. The little red, one-roomed schoolhouse is the original school built when the town was founded. At the northwest corner of River Highway and Mountain Rd. is the Gas Station. One of the mechanics from the Gas Station sits in front of the station office and waves to all the cars that drive past.

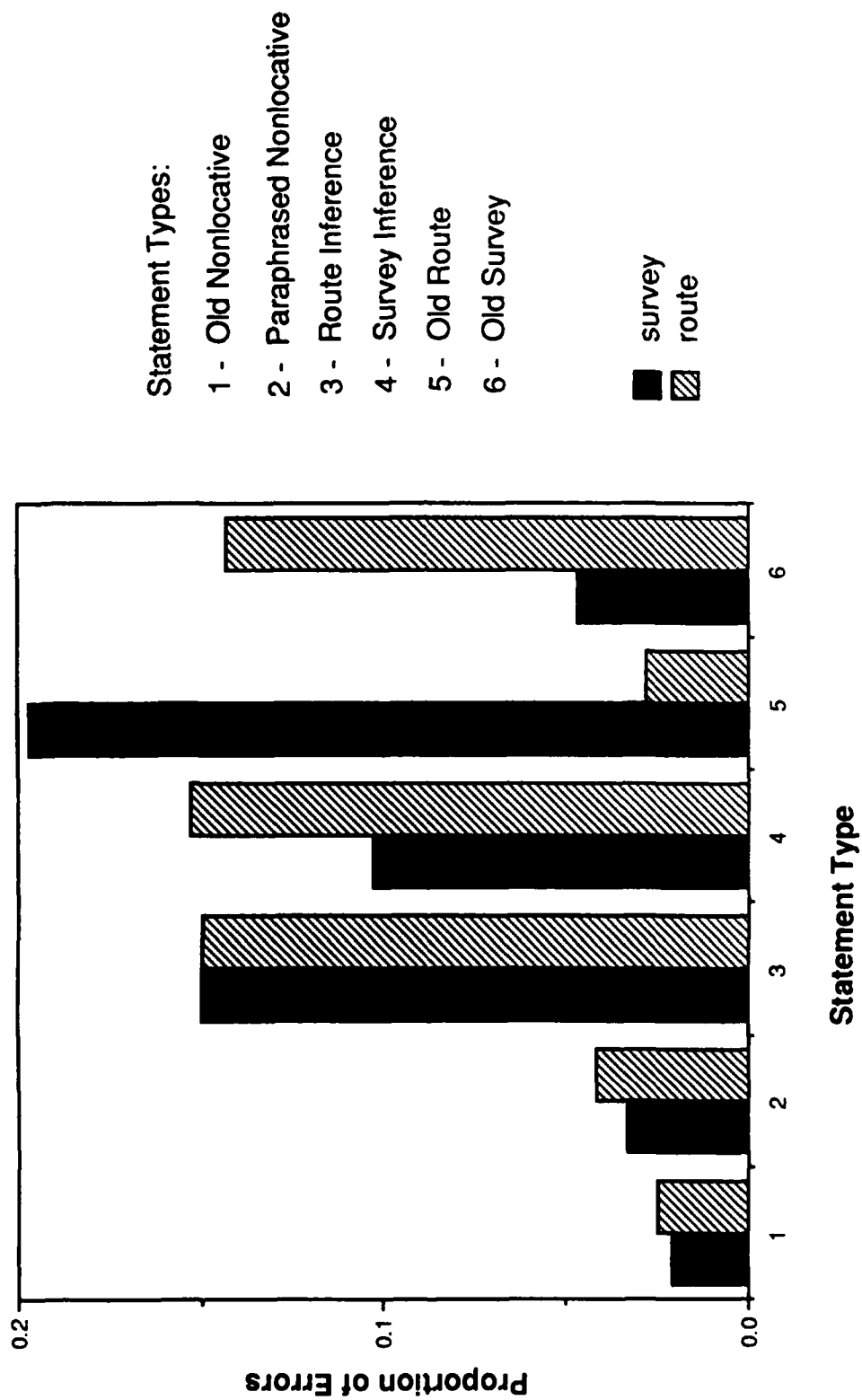
Study Times: Experiment 1



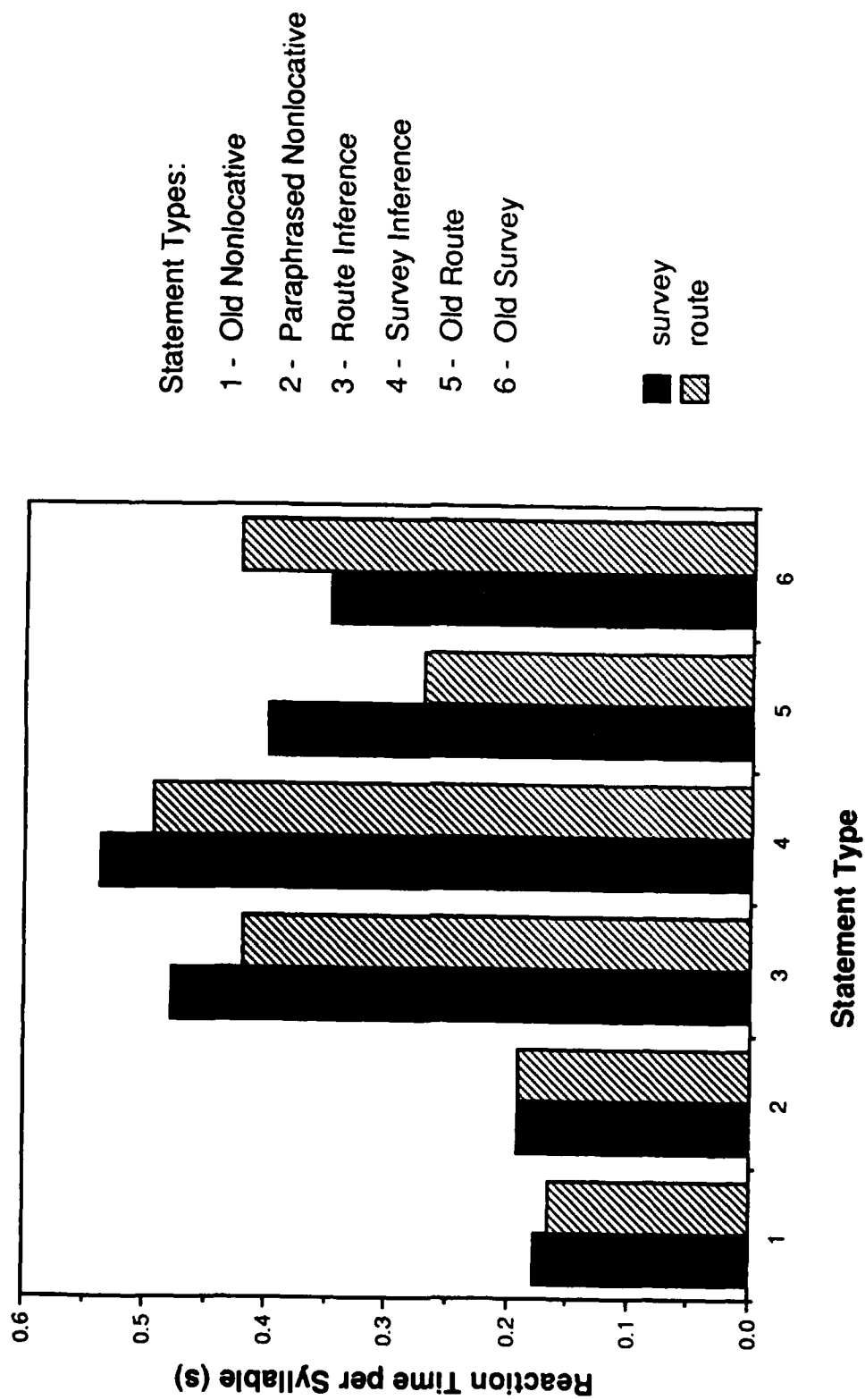
Reaction Time per Syllable: Experiment 1



Proportion of Errors: Experiment 1



Reaction Time per Syllable: Experiment 2



Proportion of Errors: Experiment 2

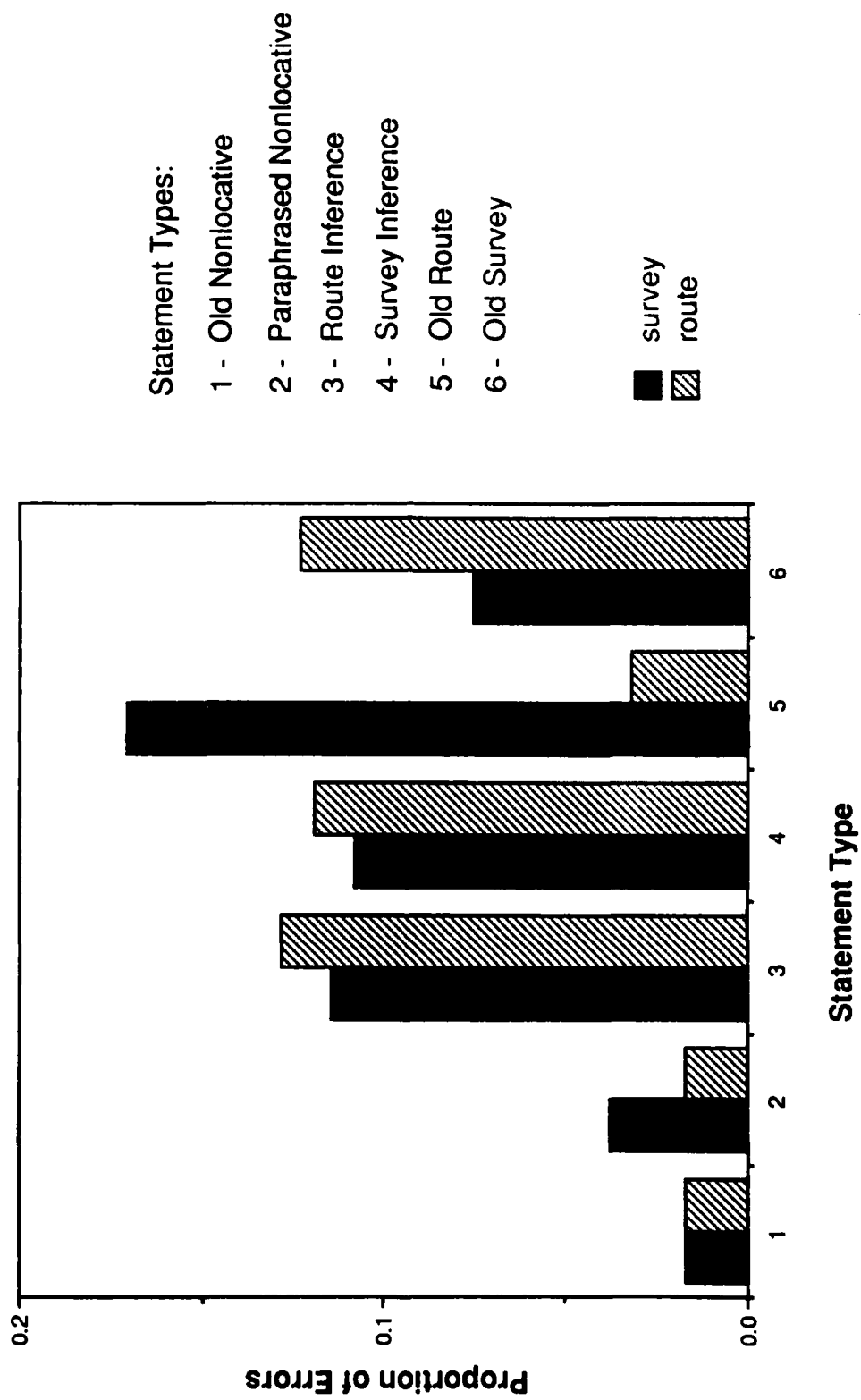


Table 1. Mean response times (in seconds) for each dimension for all experiments. Means are given separately for reclining and upright observers.

		Experiment				
		1	2	3	4	5
d i m e n s i o n		Upright				
	head/feet	1.57	1.36	1.59	--	1.50
	front/back	1.84	1.58	1.81	--	1.72
	left/right	2.21	2.02	2.26	--	2.07
d i m e n s i o n		Reclining				
	head/feet	--	--	--	2.42	2.14
	front/back	--	--	--	2.26	1.82
	left/right	--	--	--	3.25	2.59

Table 2

Mean response times (in seconds) for third person narratives

Posture	Direction					
	Head	Feet	Front	Back	Left	Right
Upright	1.52	1.34	1.43	1.62	2.07	2.16
Mean	1.43		1.53		2.12	
Reclining	2.06	1.90	1.58	1.77	2.69	2.30
Mean	1.98		1.67		2.50	

Table 3

Mean response times (in seconds) for object-centered narratives

Posture	Direction					
	Top	Bottom	Front	Back	Left	Right
Upright	1.59	1.66	1.99	2.51	2.99	2.58
Mean	1.62		2.25		2.78	
Reclining	3.14	3.85	2.51	3.15	3.51	2.93
Mean	3.49		2.83		3.22	

Table 4.

Mean response times (in seconds) to object questions
for third person narratives

Posture	Direction					
	Head	Feet	Front	Back	Left	Right
Upright	1.91	1.93	2.01	2.22	2.51	2.15
Mean	1.92		2.11		2.33	
Reclining	2.81	2.47	2.01	2.15	3.42	3.36
Mean	2.64		2.08		3.39	

Table 5.

Mean response times (in seconds) to direction and object questions

Posture	Direction					
	Head	Feet	Front	Back	Left	Right
DIRECTION PROBES						
Upright	1.57	1.55	1.69	1.93	2.57	2.00
Mean	1.56		1.81		2.29	
Reclining	2.50	2.49	2.09	1.89	2.69	2.80
Mean	2.49		1.99		2.74	
OBJECT PROBES						
Upright	1.65	1.83	1.87	2.07	2.52	1.97
Mean	1.74		1.97		2.25	
Reclining	2.50	2.30	1.81	2.28	2.60	2.59
Mean	2.40		2.05		2.59	