

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Do We Really Reason about a Picture as the Referent?

Permalink

<https://escholarship.org/uc/item/5hn6q98j>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 25(25)

ISSN

1069-7977

Authors

Shimojima, Atsushi
Fukaya, Takugo

Publication Date

2003

Peer reviewed

Do We Really Reason about a Picture as the Referent?

Atsushi Shimojima^{1,2} and Takugo Fukaya

¹School of Knowledge Science, Japan Advanced Institute of Science and Technology
1-1 Asahi-dai, Tatsunokuchi, Nomi-gun, Ishikawa 923-1292, Japan

²ATR Media Information Science Laboratories
2-2-2 Hikari-dai, Keihanna Science City, Kyoto 619-0288, Japan

ashimoji@jaist.ac.jp tfukaya@atr.co.jp

Abstract

A significant portion of the previous accounts of inferential utilities of graphical representations (e.g., Sloman, 1971; Larkin & Simon, 1987) implicitly relies on the existence of what may be called *inferences through hypothetical drawing*. However, conclusive detections of them by means of standard performance measures have turned out to be difficult (Schwartz, 1995). This paper attempts to fill the gap and provide positive evidence to their existence on the basis of eye-tracking data of subjects who worked with external diagrams in transitive inferential tasks.

Schwartz (1995) cites an intriguing example to distinguish what he calls “reasoning about a picture as the referent” from “reasoning about the picture’s referent.” Suppose you are given the picture of a hinge in Figure 1 and asked to tell whether the two marks on the legs will meet if the hinge closes. If, in solving this problem, (a) you assume that the upper line of the hinge picture swings down to the lower line, (b) infer that the mark on the upper line will meet the bottom mark, and then (c) conclude that the mark on the upper leg of the hinge will meet the bottom mark, you are *reasoning about the picture as the referent*. This is so because (a) you consider a hypothetical (not actual) operation on the given diagram, namely, the downward swing of the upper line, (b) you make an inference about the results of this hypothetical operation, and (c) you project this inference to a inference about the physical hinge.

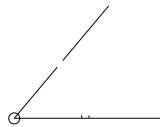


Figure 1: A schematic picture of a hinge (Schwartz, 1995).

The alternative way to solve this problem is to use the picture only as an information source or a memory-aid. You (a) interpret the given picture to obtain information about the initial state of the hinge, (b) assume that the upper leg of the hinge swings down to the lower leg, and (c) infer that the mark on the upper leg will meet the mark on the lower leg. Here again you consider a hypothetical operation and make an inference about it. Unlike

the previous case, however, the hypothesized operation is *not* an operation on the hinge picture, but an operation on the physical hinge the picture depicts. You may refer back to the hinge picture now and then, but what your inference is about is the movement of the upper leg of the hinge, not the movement of the upper line of the hinge picture. You are *reasoning about the picture’s referent*, rather than the picture itself.

Well, the concept of “reasoning about the picture itself” is thus clear, but is it real? Are we really engaged in that type of inferences in some cases? If so, how do we tell when we are?

Schwartz (1995) himself *assumed* the existence of that type of inferences, and went on to investigate what features of pictorial representations encourage it. As it turned out, however, it was not easy to determine, from response time and solution performance alone, when a subject was engaged in the second type of inferences. The precise nature of this type of inferences is simply unknown, and we have little clue to its “cash value,” namely, its impact on subject performance in inferential tasks. This led Schwartz to coin an operational definition of this type of inferences, identifying it with what he calls “feature-based reasoning.”

In this paper, we will back up a little and test the existence assumption of the first type of inferences itself. Our experiment consisted in observing eye-movements of subjects working with diagrams in transitive inference tasks, and it was specifically designed to verify whether the subjects were engaged in reasoning about diagrams themselves. The issue is important since, as we will see shortly, a significant portion of the previous accounts of the inferential utility of graphical representations (Sloman, 1971; Larkin & Simon, 1987; Narayanan, Suwa, & Motoda, 1995) relies on the existence of such inferential procedures. And for this precise reason, positive evidence for their existence would amount to the discovery of a new form inferences that generalizes the seemingly diverse phenomena reported in the diagrammatic reasoning literature. To reason about a picture itself, if it ever occurs, is to exploit a matching between geometrical or topological constraints on a graphical representation and constraints on its referent (Shimojima, 1995), and in that sense, its discovery would also be the demonstration of one precise way a cognitive agent interacts with graphical representations to unburden inferential loads.

A little more precisely, the main question of this paper is whether the following type of inferences really exist:

In the presence of a visually accessible graphical representation, an agent assumes a hypothetical transformation of the graphic that adds a premise to it, infers about the result of the transformation, and translates an obtained conclusion to a conclusion about the referent of the graphic.

In the following, we will call this type of inferences *inferences through hypothetical drawing* (“HD inference” for short). We will first clarify the breadth of this concept by reviewing some representative studies of diagrammatic reasoning and showing the inferential processes described in them are special cases of HD inferences. We will then report our main experiment designed to verify the existence of HD inferences.

Examples

Sloman (1971) presented one of the earliest, yet clearest discussions on the phenomenon of hypothetical drawing. He observed that we can solve certain inferential problems more easily when we use diagrammatic representations than when we use only symbolic representations. For example, if AB in Figure 2 represents a ditch, and CD represents a movable plank, how should we move the plank until it lies across the ditch?

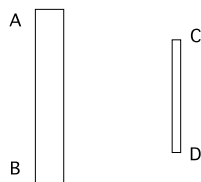


Figure 2: The Ditch-Plank Problem (Sloman 1971)

According to him, our ability to solve these problems efficiently “seems to depend on the availability of a battery of ‘subroutines’ which we can bring to bear on parts of spatial configurations, transforming them in specific ways representing changes of certain sorts in *other* configurations.” In many cases, these subroutines are performed internally: we only “imagine or envisage rotations, stretches and translations of parts” of a diagram. Thus, for the ditch-plank problem, we can “envisage” moving the smaller rectangular leftward with a little rotation. This is a hypothetical operation on parts of the relevant diagram, an instance of hypothetical drawing in our sense.

Beside their well-known analysis of information-indexing functions of diagrams, Larkin and Simon (1987) discussed instances of graphical representations that afford hypothetical drawing. Among them is the type of graphs commonly used in macro economics (Figure 3), where the abscissa represents the quantity of a commodity produced or demanded, and the ordinate represents the price at which that quantity will be supplied

or purchased. The line D is a demand curve, and the lines S and S' are supply curves in different conditions, and E and E' are equilibriums in different supply conditions.

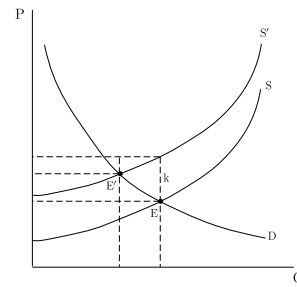


Figure 3: A Supply and Demand Graph from Macro-Economics (Larkin and Simon 1987)

Larkin and Simon asked what the effect on the equilibrium will be of imposing a manufacturer’s tax of k dollars on a unit of the commodity if the initial curve is S and the initial equilibrium E . Expressing that change in the graph amounts to moving up the curve labeled as “ S ” vertically to the position of the “ S' ” curve, and it would move the intersection with the “ D ” curve to the position labeled as “ E' .” This shift of the intersection in turn means a price increase and a quantity decrease of the commodity, with the price increase less than the amount of the tax. Here, a simple operation on an element (the “ S ” curve) of the graph lets us infer effects of the corresponding change in a particular market situation. Although Figure 3 actually shows the result of that operation as the “ S' ” curve and therefore the operation is not strictly hypothetical, we might well perform the same inference just by assuming the operation. According to Larkin and Simon, the “great utility of the diagram” arises from “the fact that it makes explicit the relative positions of the equilibrium points, so that the conclusions can be read off with the help of simple, direct perceptual operations.” “Perceptual operations” in this passage are clearly cases of hypothetical drawing in our sense.

In fact, Larkin and Simon also considered situations where the demand is more elastic. On the surface of the graph, this amounts to considering cases where the curve labeled “ D ” are flatter. On these assumptions, the vertical distance of “ E' ” and “ E ” would be less while their horizontal distance would be greater. This in turn means that on such “elastic” demands, the price increase is less while the quantity decrease is greater as the result of taxing of k . This time, the operations on the “ D ” curve are purely hypothetical, and their results are not physically reflected in the graph. Yet we can “see directly” the effects of such hypothetical operations on the graph, and infer about different demand situations of the market.

Narayanan et al. (1995) analyzed verbal protocols of subjects working on “behavior hypothesis problems,” where they used a schematic diagram of a mechanical device to predict its behaviors on a given input opera-

tion. For example, subjects were given the diagram in Figure 4 and asked to predict how the state of the entire device will change when high-pressure gas is pumped in through hole *A* in the direction of arrow.

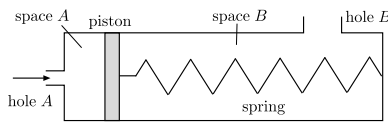


Figure 4: Diagram for a Behavior Hypothesis Problem (Narayanan et al. 1995)

The process model proposed by Narayanan and his colleagues about the inferential procedure in this task explicitly contains a sub-process which we would call “hypothetical drawing.” In their model, subjects process three types of internal representations when solving a behavior hypothesis problem: “conceptual frames” that encode conceptual knowledge about the device, and two types of “visual representations” that encode information about the given diagram itself, rather than the device depicted by it. In particular, in the process called “visualization,” subjects simulate spatial behaviors by incrementally modifying these internal representations of the diagram. They “transform the represented diagram by manipulating diagram elements that represent individual components”, where manipulations contains such general operations as *move*, *rotate*, *copy*, and *delete* as well as component specific ones such as *compress-spring*.

The influential work on “mental animation” by Hegarty (1992) might be also taken to refer to a form of HD-inferences on pulley diagrams, but her own characterization of the process seems neutral about whether the mental animation is about the movements of the graphical elements of a pulley diagram itself or about the movements of the physical parts of the actual pulley system.

Trafton and Trickett (2001) present interesting research on “spatial transformations”, defined as cognitive operations on “internal (i.e., mental) image” or “external image (i.e., a scientific visualization on a computer-generated image).” Thus, spatial transformations seem to comprise hypothetical drawing, except that the latter are confined to hypothesized operations on external image as opposed to mental operations on internal image.

Experiment

Assumption Our experiment was based on the methodological assumption that a mental operation (such as hypothetical drawing) on a visual scene is correlated with eye-fixation patterns on the scene, when the mental operation involves a elongated scan of the scene. For example, when we imagine a person in our visual scene to walk from the scene’s left end to the right end, we assume that the eye fixation is significantly more likely to move from the left end to the right end than when no such imagination takes place.

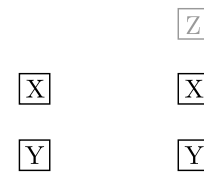


Figure 5: A vertical position diagram expressing the premise that *X* is bigger than *Y* (left). The result of updating the left picture by adding the second premise that *Z* is bigger than *X* (right).

Of course, there has been significant controversy as to whether correlations really exist between the content of mental imaging and eye movements, but recent studies (Brandt & Stark, 1997; Demarais & Cohen, 1998) started to reaccumulate positive evidence for correlation at least in the case of elongated scanning or image movements. In particular, the experiments by Brand and Stark (1997) indicate the existence of such correlation when the subjects are engaged in imaging tasks on the base of a given visual scene (a grid picture), rather than pure imaging tasks in closed-eye conditions. This provides a strong support to our methodological assumption.

Main Design Suppose, in a deductive inference task, you were shown the leftmost picture in Figure 5 in a display, while hearing the premise that *X* is bigger than *Y*. The natural interpretation of the picture is then that the “aboveness” relation holding between the square labeled “*X*” and the square labeled “*Y*” means the “bigger-than” relation holding between the objects *X* and *Y*. Now suppose you hear the second premise that *Z* is bigger than *X*. You are then asked whether *Z* is bigger than *Y*. Since the aboveness relation means the bigger-than relation in the current semantic convention, if this second premise is to be added to the picture, a square labeled “*Z*” should be placed above the square labeled “*X*,” resulting in the updated picture in Figure 5. It would then be easy to answer the question, since the “*Z*” square is clearly above the “*Y*” square, meaning that *Z* is bigger than *Y*.

In our experiment, however, neither the experimenter nor the subject added this second premise to the picture. Instead, we observed how the subject’s eyes move when the second premise was given and when the question was asked. If a subject is engaged in an HD-inference, then he or she must assume a hypothetical operation of drawing a “*Z*” square above the “*X*” square and then evaluate how this hypothetical operation would change the configuration of the picture. Thus, on the basis of the methodological assumption described above, we expect that the subject’s fixation points would move between the blank area above the “*X*” square (the “hypothetical drawing position”) and the “*Y*” square. If the subject uses the picture only as a memory-aid for the first premise or uses it for no specific purpose, then the fixation points would not necessarily move in this way. Thus, we expect that the subject’s eye-fixation pattern should reveal the use of

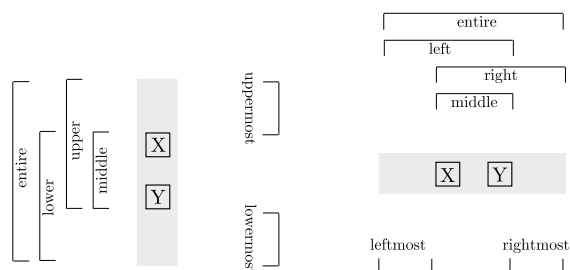


Figure 6: Forms of vertical and horizontal position diagrams used in the experiment. The adjectives (“upper,” “lower,” “middle,” etc.) expressing different areas of diagrams are also defined.

hypothetical drawing, if it ever occurs.

Preparation of Stimuli We created a total of 48 transitive inference problems as stimuli. Each problem consists of two premises and a yes-no question, where the first premise relates one object X to the other Y with a transitive relation R , the second premise relates Y to the third object Z with the same relation R or its inverse $\bar{R}$, and the question is whether these premises imply that Z has the relation R or $\bar{R}$ to the first object X .

A total of 6 pairs of a transitive relation and its inverse were mentioned in the problems: *in front of* and *behind*, *to the east of* and *to the west of*, *bigger than* and *smaller than*, *brighter than* and *darker than*, *heavier than* and *lighter than*. We asked 7 independent subjects to make visualizability and spatializability assessments about each pair of relations (Knauff & Johnson-Laird, 2000). The subjects used the scale from 1 (very difficult) to 7 (very easy). As the result, the mean ratings ranged from a moderately visualizable pair (5.63 for *bigger than* and *smaller than*) to a moderately unvisualizable pair (2.75 for *brighter than* and *darker than*) and from the highly spatializable pair (6.38 for *in front of* and *behind*) to a moderately unsatializable pair (2.38 for *brighter than* and *darker than*).

We designed the problems so that they may systematically vary in two respects. First, they varied in the syntactic structure of the picture displayed with the first premise. That is, a half of the problems come with vertical position diagrams such as the left diagram in Figure 6, while the other half come with horizontal position diagrams such as the right diagram in Figure 6.

Secondly, the problems varied in the semantic convention associated with the displayed picture. That is, for each problem where a syntactic relation R' indicates the target relation R , there was a corresponding problem where the same syntactic relation R' indicates the inverse target relation $\bar{R}$. For example, the problem cited in the last section has a counter-part problem where the aboveness relation in the displayed diagram indicates the *smaller than* relation rather than the *bigger than* relation.

Thus, the problems were divided into eight types

Table 1: Hypothetical drawing positions (H.D.P.) of each type of problems used in the experiment.

Type	Diagram	H.D.P.	Type	Diagram	H.D.P.
a	vertical	uppermost	b	vertical	lowermost
c	horizontal	rightmost	d	horizontal	leftmost
e	horizontal	leftmost	f	horizontal	rightmost
g	vertical	lowermost	h	vertical	uppermost

$a, \dots, h$, according to the combinations of the premise types, the syntactic types of accompanying diagrams, and the semantic rules associated with them. Note that different types of problems would have different *hypothetical drawing positions*, namely, places in diagrams where new squares would be drawn if second premises were added. For example, the problem cited in the beginning of this section has the hypothetical drawing position in the uppermost area of the diagram. If, however, the second premise were that the bear is smaller than the cat, the hypothetical drawing position would be the lowermost position. Table 1 shows the hypothetical drawing position for each problem type.

Procedures We asked 10 subjects (2 females and 8 males, ages: 22–28) to solve the 48 problems thus produced. In each case, the subject wore a cap equipped with NAC eye-tracker EMR-8, facing to a 90-inch projection screen 10 feet away. With the jaw placed on a fixed support, the subject first solved six practice problems. After the calibration of eye-marks, the subject then solved 24 problems with horizontal position diagrams and 24 more problems with vertical position diagrams. The premises and the question was given in tape-recorded voice, with the first premise accompanied by a diagram displayed in the computer display. In each trial, the subject pushed the space bar of a keyboard to proceed to the next premise, to the question, or to the next problem. The subject pushed the “4” key to answer “yes” to the question and the “6” key to answer “no.” There was no key for going back to the previous stage or for correcting the answer. The movement of the subject’s right eye was tracked and recorded throughout the session.

Due to a calibration failure, we had to exclude the entire data of one subject from our analysis. Of the remaining 432 trials, the subjects gave correct answers in 362 trials (84 %), and we analyzed only these trials. We avoided trials with incorrect answers in order to exclude, as much as possible, accidental eye-movements caused by simple mishearing of premises or questions.

Results

Overall tendency Figures 7 and 8 show typical patterns of eye fixations during individual trials. Each individual figure shows the fixation pattern in the period after the first premise was presented with the diagram until the subject pushed an answering key. The sizes of circles express the lengths of fixation durations.

The left figure in Figure 7 shows eye fixations during a solution of a type- a problem, and it is visually clear that eye fixations concentrated in the upper area of the

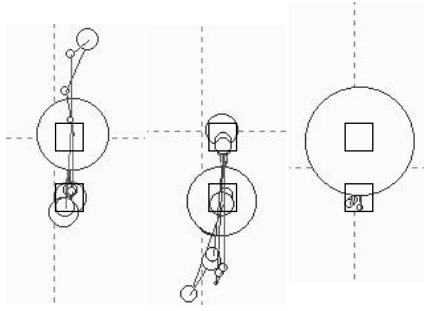


Figure 7: Eye fixations in a trial on a type-*a* problem (left), on a type-*b* problem (middle), and on a type-*a* problem (right); the sizes of circles express durations of fixations; approximate positions of two squares in the actual diagram are also indicated.

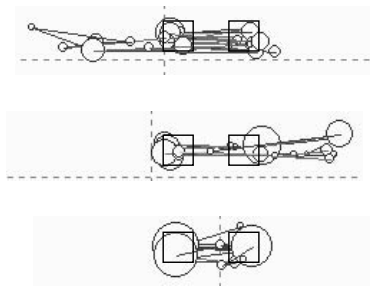


Figure 8: Eye fixations in a trial on a type-*e* problem (top), on a type-*c* problem (middle), and on a type-*c* problem (bottom).

diagram. The middle figure shows eye fixations during a solution of a type-*b* problem, and this time, they gathered in the lower area. Thus, in both cases, the subject eyes moved between the existing squares and the hypothetical drawing position. In fact, the large-sample Wilcoxon signed-rank test shows that the net fixation time is significantly longer in the uppermost part than in the lowermost part in type-*a* and type-*h* problems ($n = 92, z = 3.78, p < 0.001$, two tails) while it is significantly longer in the lowermost part in type-*b* and type-*g* problems ($n = 92, z = 3.70, p < 0.001$, two tails).

On the other hand, there are a significant number of cases where eyes did not move to hypothetical drawing positions, but concentrated on the small middle part of the diagram. The right figure shows a typical example of that fixation pattern.

A similar tendency was found in trials with horizontal diagrams. On the one hand, a majority cases had eye fixations between the existing squares and the hypothetical drawing positions. See the top and the middle figure in Figure 8. In fact, the net fixation time is significantly longer in the rightmost part than in the leftmost part in type-*c* and type-*f* problems ($n = 91, z =$

Table 2: Counts of trials classified in terms of the areas covered by eye fixations and the problem types, where problems were presented with vertical position diagrams.

	lower	upper	entire	middle	total
type- <i>a</i> & <i>h</i>	5 (-3.4)	31 (2.5)	8 (0.6)	45 (0.4)	89
type- <i>b</i> & <i>g</i>	35 (3.4)	9 (-2.5)	5 (-0.6)	40 (-0.4)	89
total	40	40	13	85	178

Table 3: Counts of trials classified in terms of the areas covered by eye fixations and the problem types, where problems were presented with horizontal position diagrams.

	right	left	entire	middle	total
type- <i>d</i> & <i>e</i>	3 (-3.6)	28 (2.1)	11 (1.2)	52 (0.4)	94
type- <i>c</i> & <i>f</i>	33 (3.7)	9 (-2.1)	4 (-1.2)	44 (-0.4)	90
total	36	37	15	96	184

3.47, $p < 0.001$, two tails), while it is significantly longer in the leftmost part in type-*d* and type-*e* problems ($n = 92, z = 3.06, p < 0.002$, two tails). On the other hand, another majority of trials showed concentrations of fixation points in the small middle part of the diagram (as shown in the bottom figure).

Fixation ranges What range of the diagram did eyes cover during each trial? Tables 2 and 3 show the counts of trials classified in terms of the areas of diagrams covered by fixations and the types of problems being solved. Table 2 shows the case with vertical diagrams, while Table 3 shows the case with horizontal diagrams. As both tables show, about a half of the trails have eye fixations concentrated on the middle range of diagrams, consisting of the two squares and the blank space in between. The other half of the trails, however, had wider fixation ranges, spreading over either up, down, to the left, to the right, or to the entire diagram. And in these cases, fixation ranges have a very significant dependency on types of problems both for the vertical cases ($\chi^2 = 35.59, p < .001$) and the horizontal cases ($\chi^2 = 38.62, p < .001$). That is, when a subject was engaged in a type-*b* or type-*g* problem, eye fixations tended to spread over the lower range of the diagram, while with a type-*a* or type-*h* problem, they tended to cover the upper range. With a type-*c* or type-*f* problem, eye fixations tended to cover the right-side range of the diagram, while with a type-*d* or type-*e* problem, they tended to cover the left-side range.

Discussions

Overall, the cases in our data can be divided into two broad classes: the class of cases where fixation points concentrated on the middle part of the diagram (181 cases in total), and the class of cases where fixation points spread beyond the middle part (181 cases in total). In the second class of cases, the subject exhibited a

very strong tendency to move their eyes back and forth between the actual graphical elements (i.e. two squares) and the hypothetical drawing position.

Now, it is not our concern to investigate what inferential strategy or strategies the subjects were engaged in the first class of cases, where they fixed their eyes to the small area consisting of the actually drawn squares. Whatever the explanation may be, whatever inferential strategies this class of cases may represent, our data clearly indicate that, in the second class of cases, where a subject was *not* engaged in those inferential strategies, he or she seems to be engaged in an inference through hypothetical drawing. And this type of cases accounts for a half of the trials that we observed. So, our results indicate not only the existence of HD-inferences, but their generality or pervasiveness in this type of inferential tasks.

Moreover, our data indicate that the strong dependency was preserved over a certain syntactic variance of the displayed diagram: it was observed no matter the displayed diagram was vertical one or horizontal one. Perhaps more importantly, the tendency was preserved over semantic variances of the displayed diagram too. For example, it was observed both in case the aboveness relation in diagrams means the smaller than relation and in case the same aboveness relation means the bigger than relation. This indicates that the observed fixation patterns were controlled by the particular syntactic and semantic rules associated with the diagram itself.

This in turn ensures that the observed eye movements were indicative of mental operations on the present diagram scene (hypothetical drawing), rather than mental operations on some other independent image constructed for the objects or situation represented by the diagram. For such mental imaging, if it ever existed, would not change its syntax and semantics as those of the displayed diagram do, and hence would be reflected in a unified fixation pattern independent of the diagram's syntactic and semantic variation. Our data showed that on the contrary, eye fixation patterns tracked the diagram's syntactic and semantic variation.

Conclusions

Thus, our experiment gave clear evidence for the existence of HD-inferences, reasoning about a picture as the referent. Intuitively, this might come with no surprise, but its theoretical implications are quite significant.

First, it provides empirical support to the processes postulated in previous studies of diagrammatic reasoning. The concepts of envisaging operation (Sloman, 1971), perceptual operation (Larkin & Simon, 1987), and visualizing operation (Narayanan et al., 1995) now have significant empirical muscle and do their explanatory jobs more convincingly.

Moreover, the concept of HD-inference identifies the common structure of the seemingly diverse processes thus postulated, and hence has a potential to make the existing accounts more systematic. Specifically, to infer through hypothetical drawing is, after all, to exploit geometrical or topological constraints on graphics to rea-

son about the distal object. Those structural constraints on graphics serve as scaffolding to one's deductive endeavor, even if it may be eventually thrown away. Of course, different types of graphics (e.g., maps, graphs, line drawings) have different types of structural constraints, and hence the kinds and qualities of supplied scaffolding must vary accordingly. In particular, the match and mismatch between structural constraints on graphics and constraints on their referents will be an important determinant of the quality of the scaffolding. Thus, our finding suggests, this factor is one of the dominant ones that regulate the way cognitive agents interact with external graphics in inferential tasks.

Reference

- Brandt, S. A., & Stark, L. W. (1997). Spontaneous Eye Movements during Visual Imagery Reflect the Content of the Visual Scene. *Journal of Cognitive Neuroscience*, 9(1), 27–38.
- Demarais, A. M., & Cohen, B. H. (1998). Evidence for image-scanning eye movements during transitive inference. *Biological Psychology*, 49, 229–247.
- Hegarty, M. (1992). Mental Animation. In Glasgow, J., Narayanan, N. H., & Chanrasekaran, B. (Eds.), *Diagrammatic Reasoning: Cognitive and Computational Perspectives*, pp. 535–575. AAAI Press, Menlo Park, CA.
- Knauff, M., & Johnson-Laird, P. N. (2000). Visual and Spatial Representations in Relational Reasoning. In *Proceedings of the Twenty-Second Annual Conference of the Cognitive Science Society*, pp. 759–763.
- Larkin, J. H., & Simon, H. A. (1987). Why a Diagram Is (Sometimes) Worth Ten Thousand Words. In Glasgow, J., Narayanan, N. H., & Chanrasekaran, B. (Eds.), *Diagrammatic Reasoning: Cognitive and Computational Perspectives*, pp. 69–109. AAAI Press, Menlo Park, CA.
- Narayanan, N. H., Suwa, M., & Motoda, H. (1995). Hypothesizing Behaviors from Device Diagrams. In Glasgow, J., Narayanan, N. H., & Chanrasekaran, B. (Eds.), *Diagrammatic Reasoning: Cognitive and Computational Perspectives*, pp. 501–534. AAAI Press, Menlo Park, CA.
- Schwartz, D. L. (1995). Reasoning about the Referent of a Picture versus Reasoning about the Picture as the Referent: an Effect of Visual Realism. *Memory and Cognition*, 23(6), 709–722.
- Shimojima, A. (1995). Operational Constraints in Diagrammatic Reasoning. In Barwise, J., & Allwein, G. (Eds.), *Logical Reasoning with Diagrams*, pp. 27–48. Oxford University Press, Oxford.
- Sloman, A. (1971). Interactions between Philosophy and AI: the Role of Intuition and Non-Logical Reasoning in Intelligence. *Artificial Intelligence*, 2, 209–225.
- Trafton, J. G., & Trickett, S. B. (2001). A New Model of Graph and Visualization Usage. In *Proceedings of the Twenty-Third Annual Conference of the Cognitive Science Society*, pp. 1048–1053.