

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

What can Information Extraction from Scenes and Causal Systems Tell us about Learning from Text and Pictures?

Permalink

<https://escholarship.org/uc/item/6188h50d>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 32(32)

ISSN

1069-7977

Authors

Eitel, Alexander
Scheiter, Katharina
Schuler, Anne

Publication Date

2010

Peer reviewed

What can Information Extraction from Scenes and Causal Systems Tell us about Learning from Text and Pictures?

Alexander Eitel (a.eitel@iwm-kmrc.de)

Knowledge Media Research Center, Konrad-Adenauer Str. 40,
72072 Tübingen, Germany

Katharina Scheiter (k.scheiter@iwm-kmrc.de)

Knowledge Media Research Center, Konrad-Adenauer Str. 40,
72072 Tübingen, Germany

Anne Schöler (a.schoeler@iwm-kmrc.de)

Knowledge Media Research Center, Konrad-Adenauer Str. 40,
72072 Tübingen, Germany

Abstract

Numerous studies have shown that the gist in photorealistic pictures of scenes is extracted after very short presentation times. So far, the investigation of gist extraction has been limited to pictures of scenes. The present study investigated whether the gist in pictures of causal systems, which are typically used as instructional material, is extracted as fast as the gist in pictures of scenes, and whether more than just the gist is rapidly extracted from a causal system (i.e., information concerning its details and functioning). Schematic and photorealistic pictures of scenes and causal systems were presented to subjects ($N = 24$) at different presentation times. Results showed that the gist in causal systems is extracted as fast as in scenes, and that an initial understanding of the functioning of schematic causal systems is also rapidly acquired. Results are discussed in the light of their implications for learning from text and pictures.

Keywords: Gist; Scene Perception; Causal Systems; Learning from Text and Pictures

Learning from Text and Pictures

In multimedia research it is a well known finding that learning from text and pictures leads to better retention and recall than learning from text alone (Levie & Lentz, 1982). Moreover, when students have to learn about causal systems, they are better able to apply their knowledge to produce creative solutions to problem-solving questions after learning from text and pictures than after learning from text alone (Mayer, 1989). Accordingly, learning from both text and pictures leads to higher comprehension than learning from text alone. Despite the fact that the beneficial effects of pictures for learning are well established in the research literature, far less is known concerning how the pictorial information is processed during learning.

An exception is an early study by Hegarty and Just (1993), in which comprehension was assessed via questions about the kinematics of different pulley systems. Students were better able to infer motion in the pulley system when previously learning from text and pictures than when learning from text alone or picture alone. Eye tracking data

furthermore revealed that during learning from text and pictures of pulley systems, subjects first processed text information, and then switched to the picture in order to integrate the information from both sources.

Unlike Hegarty and Just (1993), various studies showed that when subjects were confronted with information from text and pictures, they often initially looked at the picture for a short time before they started to read the text. This pattern of processing has been shown for advertisements (Rayner et al., 2001), comics (Carroll, Young, & Guertin, 1991), real-world scenes (Underwood, Jebbett, & Roberts, 2004), and biology schoolbooks (Mak, 2008). In a study by Stone and Glock (1981), in which subjects had to learn how to build a cardboard loading cart from text and schematic pictures, subjects first looked at the picture for 1000 to 2000 ms, before they started to read the text. According to the authors, subjects initially looked at the picture in order to get a first impression (i.e., gist) of what the material was about. However, it is yet unclear what role looking briefly at a picture prior to reading a text may play for understanding the presented content. At least this phenomenon has not been directly addressed in research on learning from text and pictures. However, there has been ample research in basic cognitive psychology about the extraction of information from briefly looking at pictures of scenes.

Extraction of Information from Scenes

In an early study of Biederman and colleagues (1974), subjects had to select one out of two labels, which they judged to better describe a picture of a jumbled vs. coherent scene. In coherent scenes, when the two labels were similar (e.g., “shopping plaza” vs. “busy road and stores”), accuracy of selecting the right label was at 100% for the majority of subjects after 300 ms of presentation. When the two labels were dissimilar, a ceiling effect in accuracy of selecting the right label occurred after only 100 ms. The authors concluded that information about the gist of a scene is already extracted after a single fixation, which enabled subjects to perform the task correctly. This is in line with

the findings from Henderson and Hollingworth (1999), who state that the average fixation duration during scene viewing is about 330 ms.

Similarly, Loftus, Nelson, and Kallman (1983) conducted a study in which subjects were asked to decide whether the picture of a scene had already been presented or not. Subjects were told to base their decision either on general properties of the picture or on detail information. When the decision was based on general properties of the picture, performance increased much less between 250, 500, and 1000 ms presentation time than when the decision was based on detail information. The authors concluded that most holistic information in scenes is extracted from the first fixation (about 330 ms; Henderson & Hollingworth, 1999) and subsequent fixations have the primary purpose of identifying relevant details.

Castelhano and Henderson (2008) also provided evidence for a rapid extraction of holistic information from pictures of scenes by presenting photos of scenes to subjects for a short time (25 – 250 ms) and later asking them whether a specific detail had been depicted in the scene. The detail in question was either consistent (e.g., fire hydrant) or inconsistent (e.g., tea set) with the gist of the scene (e.g., street scene) but was never actually present. Between 42 and 250 ms presentation time, subjects more often affirmed that the detail in question was present in the scene when the detail was consistent than when it was inconsistent with the gist of the scene. The authors concluded that a rapidly acquired (42 – 250 ms) scene gist was responsible for more affirmative responses to details consistent with scene gist by activating information about the scene's content and basic-level category.

To conclude, studies in basic cognitive psychology consistently demonstrate that information about the gist (e.g., general topic) in photos of scenes is extracted within the first fixation. Later fixations are presumably made to scan the scene for details.

Aims of the Current Study

The aim of the current study was to apply and compare findings from basic cognitive research on gist extraction from scenes to learning from text and pictures to better understand the role that pictures might play during the latter.

Unlike with scenes, there has yet not been much research about the extraction of information from instructional material. As mentioned before, in the study from Stone and Glock (1981), looking at the picture for 1000 to 2000 ms was interpreted as the time it took subjects to extract the gist. This is much longer than the time it takes subjects to extract the gist from scenes (< 250 ms). However, Stone and Glock interpreted the time subjects initially looked at the picture before reading the text as the time required to extract its gist. Subjects could also have extracted the gist within the first fixation (about 330 ms) as in scenes and looked at the picture up to 2000 ms only in order to scan it for details. Thus, it is still unclear when information about the gist and details is extracted in pictures of instructional material.

Information extraction in pictures of causal systems was investigated, since they are often used as instructional material in studies on learning from text and pictures (e.g., Hegarty & Just, 1993; Mayer, 1989). It was expected that once the gist of a causal system has been extracted, subjects would use the remaining time to understand the functioning of the depicted system. Hence, with longer presentation times knowledge about the functioning of the system should improve. In the aforementioned studies on learning about causal systems (e.g., Mayer, 1989), mostly schematic pictures of causal systems have been used (e.g., line drawings). On the other hand, gist extraction from scenes has been investigated by presenting photorealistic pictures to subjects (e.g., Castelhano & Henderson, 2008). In the present study, it was investigated whether these findings on gist extraction could be extended to schematic pictures of causal systems from studies on learning from text and pictures (e.g., Hegarty & Just, 1993). To overcome the confound that in previous research mostly photorealistic pictures of scenes and schematic pictures of causal systems were used, schematic and photorealistic depictions of both, scenes and causal systems were directly compared to each other in the current study. The degree of realism is considered to be a continuum with schematic line drawings on the one end and photos of natural objects on the other end. The less similar an illustration is to its real-world referent with respect to shape, details, color, and texture the more schematic it is. It can be expected that in general information will be extracted more easily from schematic depictions than from realistic ones, because the prior do contain fewer elements which can be recognized more easily due to better contrasts etc. However, effects of realism were not the focus of the study; rather this variable was solely introduced to bridge the gap between prototypical materials used in scene perception research (photorealistic scenes) and research on learning from text and pictures (schematic depictions of causal systems).

Hence, the current study addressed the question whether the gist in causal systems would be extracted as fast as the gist in scenes. Further, details were assumed to be better extracted at longer presentation times compared to shorter ones. Finally, it was investigated whether the functioning of causal systems would be understood and whether this depended on presentation time.

Method

Participants and Design

Twenty-four students (15 female, 9 male, average age: $M = 23.83$ years, $SD = 3.50$) from the University of Tuebingen, Germany, took part in the experiment for either payment or course credit. The experiment followed a $2 \times 2 \times 4$ design, with *Type* (scene vs. causal system), *Realism* (schematic vs. realistic) and *Presentation Time* (150 vs. 600 vs. 2000 vs. 6000 ms) serving as within-subjects factors.

Materials and Procedure

The materials in the experiment comprised 80 pictures of scenes and 80 pictures of causal systems. In a pilot study, subjects had to rate the number of objects in each picture, and to categorize the pictures with respect to their degree of realism and type (scene vs. causal system); the rated number of objects was the same for realistic and schematic pictures, and only unambiguous illustrations were used in the study. A scene depicted an everyday situation. A causal system always had a certain purpose (e.g., pulling weight). It consisted of multiple components, where at least one component was influenced by another – hence, removing one component would have changed the functioning of the system. In the experiment, for both scenes and causal systems, half of the pictures were schematic, the other half realistic. This led to four different categories of pictures in the experiment (see Figure 1). Each picture appeared in the center of the computer screen and covered nearly the whole screen size. An experimental session consisted of 8 training trials and 160 experimental trials.

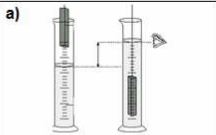
	schematic	realistic
Causal systems	a) 	b) 
Scenes	c) 	d) 

Figure 1: Categorization of pictures used in the experiment. Pictures could either depict a scene or a causal system and could be either schematic or realistic.

Each experimental trial started with the presentation of the word “ready?”, which remained on the screen until a key was pressed. After pressing a key, the word “ready?” was replaced by the fixation cross, which was displayed for 800 ms. Then a picture (scene vs. causal system, schematic vs. realistic) appeared for either 150, 600, 2000 or 6000 ms, respectively, and was immediately masked afterwards. Both pictures and presentation times were presented in a randomized order. After each picture, a statement about the gist, then about details, and then about the functioning of the picture was presented and students were asked to respond to these statements (see Measures section for details). The statement concerning the functioning was presented only after pictures of causal systems. After responding to the last statement (detail or functioning), the trial was over and the word “ready?” reappeared, which marked the beginning of a new trial. An experimental trial for a single picture lasted about 15 seconds. The whole experimental session lasted approximately 45 minutes.

Measures

After viewing each picture, participants had to respond to either two or three statements about the picture depending on the experimental condition. All statements were in a two-alternative-forced-choice format, where students had to choose between a “yes” and a “no” response by pressing one of two keys on a keyboard. In half of the trials, “yes” was the correct response, in the other half of the trials “no” was correct. The first statement was about the gist of the picture. For instance, students were asked to decide whether a scene could be identified as “happy people” (see Figure 1d) or whether a causal system could be identified as “electric circuit” (see Figure 1b). Statements about the gist always consisted of only one to three words. In the second statement, participants had to judge whether specific details had been present in the scene (e.g., “presents are lying under the tree”; see Figure 1c) or in the causal system (e.g., “an eye is depicted”; see Figure 1a) just seen. Details were not relevant to either the meaning of the scene or the functioning of the causal system. Moreover, details were depicted in the periphery rather than in the center of the picture so that they were less likely to be seen within the first fixation. The third statement was presented only after pictures of causal systems, and was about the functioning of the depicted system. In order to be able to answer statements about the functioning correctly, inferences were required (e.g., “If the block is pulled out of the test tube, then liquids are at the same level in both test tubes”; see Figure 1a). It is important to note that statements concerning the functioning could be answered correctly only by relating multiple objects from the picture to each other; they could not be answered correctly solely based on prior knowledge that might have been activated once the causal system had been recognized correctly. The detail and functioning statements consisted of one sentence each.

As the main dependent variable, the percent correct was computed. Each correct response (both hits and correct rejections) was coded with 1, each incorrect response with 0. Multiplied by 100, percent correct was 100% at maximum and 50% at chance level and was computed separately for the three types of statements (gist, details, and functioning). Mean reaction times (RT) for responses to the different statements served as a second dependent variable in the experiment. Eye tracking data were assessed as well, but will not be reported here for space reasons.

Results

Overall, results revealed that there was no speed-accuracy trade-off, since there was no significant negative correlation between accuracy and RT ($r = .24, p = .26$). Thus, only accuracy to statements about the gist, details, and the functioning will be analyzed here.

Gist

T-tests revealed that both in scenes ($t(23) = 31.32, p < .001$) and in causal systems ($t(23) = 11.42, p < .001$), accuracy to

gist statements was above chance level at the shortest presentation time (150 ms), which speaks in favor of an early extraction of gist from both, scenes and causal systems (see Figure 2).

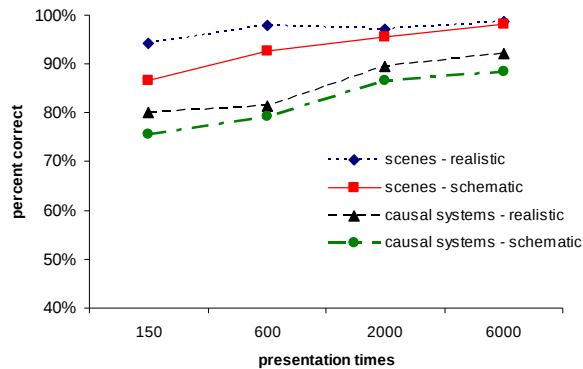


Figure 2: Accuracy to statements about the gist in schematic and realistic pictures of scenes and in schematic and realistic pictures of causal systems.

A 2 (*Type*: scenes vs. causal systems) $\times$ 2 (*Realism*: realistic vs. schematic) $\times$ 4 (*Presentation Time*: 150 vs. 600 vs. 2000 vs. 6000 ms) repeated-measures ANOVA was conducted to analyze accuracy for statements about gist. There was a significant main effect of *Type*, indicating that statements about the gist were answered more accurately ($F(1, 23) = 96.68, p < .001$) in scenes than in causal systems (see Figure 2), which is probably due to a higher difficulty in recognizing the general topic of causal systems. There were also significant main effects of *Realism* ($F(1, 23) = 8.41, p = .01$) and *Presentation Time* ($F(3, 69) = 17.51, p < .001$) meaning that gist extraction was better in realistic than in schematic pictures, and improved with longer presentation times. There were no interactions (all $p_s > .05$).

Details

T-tests revealed that both in scenes ($t(23) = 5.04, p < .001$) and in causal systems ($t(23) = 2.31, p = .03$), accuracy to statements about details was above chance level at 150 ms (see Figure 3).

A 2 (*Type*) $\times$ 2 (*Realism*) $\times$ 4 (*Presentation Time*) repeated-measures ANOVA revealed significant main effects of *Type* ($F(1, 23) = 25.63, p < .001$) and *Presentation Time* ($F(3, 69) = 12.46, p < .001$) on accuracy to detail statements. As expected, details were recognized more accurately at longer presentation times both in scenes and in causal systems. While there was no main effect of *Realism* ($F(1, 23) = 1.79, p = .19$) it interacted significantly with *Type* ($F(1, 23) = 13.08, p < .001$). Bonferroni tests showed that detail extraction was better in realistic than in schematic pictures of scenes ($p = .03$), whereas it tended to be worse in realistic than in schematic pictures of causal systems ($p = .065$). There were no further interactions (all $F_s < 1$).

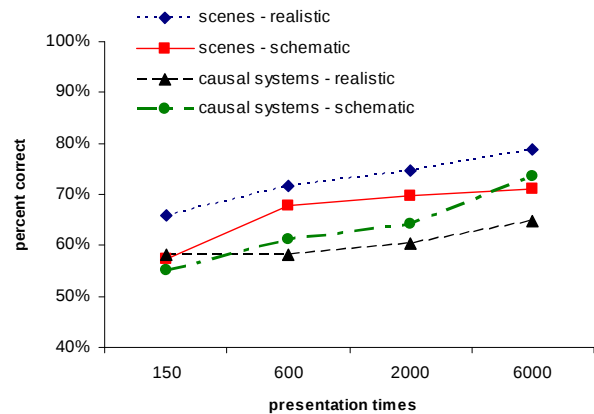


Figure 3: Accuracy to statements about details in schematic and realistic pictures of scenes and in schematic and realistic pictures of causal systems.

Functioning

T-tests revealed that accuracy to statements about the functioning of realistic pictures of causal systems was at chance level for 150, 600 and 2000 ms (all $p_s > .05$). Only at the longest presentation time of 6000 ms, accuracy was above chance level ($t(23) = 5.29, p < .001$). On the other hand, accuracy to statements about the functioning of schematic causal systems was already above chance level ($t(23) = 3.86, p = .001$) at 600 ms presentation time (see Figure 4). Only at the shortest presentation time of 150 ms, accuracy was at chance level ($p > .05$).

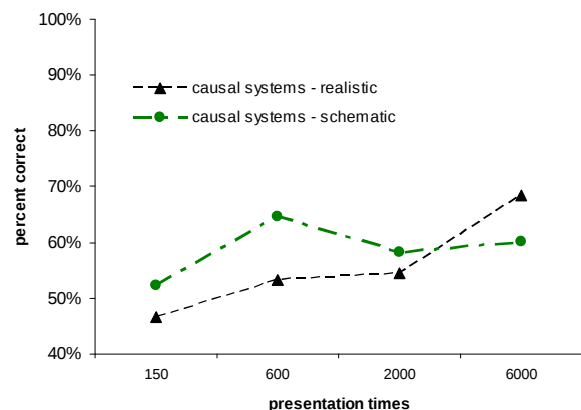


Figure 4: Accuracy to statements about the functioning in schematic and realistic pictures of causal systems.

A 2 (*Realism*) $\times$ 4 (*Presentation Time*) repeated-measures ANOVA revealed a significant main effect of *Presentation Time* ($F(3, 69) = 9.20, p < .001$) and a significant interaction *Realism*Presentation Time* ($F(3, 69) = 3.13, p = .03$), meaning that for realistic pictures of causal systems the functioning was understood better at longer presentation times, which was not the case for schematic ones.

Bonferroni comparisons confirmed that in schematic causal systems, longer presentation times (2000, 6000 ms)

did not lead to further improvements in understanding of the functioning (both $p_s > .05$) compared to 600 ms presentation time. Thus, it can be concluded that in schematic causal systems, an initial understanding was rapidly acquired (at 600 ms), and at longer presentation times schematic causal systems might have solely been scanned for details. In realistic pictures of causal systems there was no understanding of the functioning but at the longest presentation time of 6000 ms. Bonferroni tests showed that understanding of the functioning still improved between 2000 and 6000 ms ($p = .02$). Thus, it took subjects longer to understand the functioning of realistic pictures of causal systems (6000 ms) than to understand the functioning of schematic ones (600 ms). Subjects probably still attended to realistic pictures of causal systems at 6000 ms in order to extract their functioning. This possibly led to less attention to details at 6000 ms, which could have resulted in the marginally lower performance in detail extraction ($p = .054$) for realistic compared to schematic pictures of causal systems after 6000 ms presentation time.

Discussion

The present study aimed at investigating the extraction of different information (gist, details, and the functioning) from briefly attending to schematic and realistic pictures of scenes and causal systems.

The results demonstrate that the gist was rapidly extracted (< 150 ms) in both scenes and causal systems, confirming prior research from gist extraction in scenes (e.g., Castelano & Henderson, 2008) and expanding it to instructional material. Moreover, details were recognized more accurately at longer presentation times, which is in line with prior research from detail extraction in scenes (Loftus et al., 1983). Comprehension of the functioning quickly reached an asymptote in schematic pictures of causal systems (at 600 ms). In realistic pictures of causal systems, however, subjects needed more time to understand the functioning, which might have impaired detail extraction at longer presentation times because subjects might have split their attention between details and objects that they assumed to be relevant for understanding the functioning of the system. The analysis of the eye tracking data will reveal whether these assumptions hold true.

Influences on Comprehension of the Functioning of Causal Systems

More familiarity can possibly account for the faster comprehension of the functioning in schematic than in realistic pictures of causal systems. Schematic causal systems often appear in textbooks, but students are seldom faced with and almost never learn from photorealistic pictures of causal systems. Hence, a lack of familiarity with realistic pictures of causal systems could explain why understanding of schematic causal systems reached an asymptote very quickly (600 ms), whereas understanding of realistic pictures of causal systems was still at chance level at 600 ms and at 2000 ms presentation time.

Moreover, there might be an influence of domain-specific knowledge on comprehension of the functioning of causal systems. Unfortunately, in the present study no prior knowledge test could be administered, because causal systems were from many different domains (biology, chemistry, physics, engineering, and mechanics) and thus a prior knowledge test for each domain would have been too long. However, a demographic questionnaire was presented to participants that assessed their prior knowledge with regard to their last school grades in the respective school subjects and their general interest in the different domains. No participant had both very good school grades and a high interest in each of the aforementioned domains. Thus, it is highly unlikely that a participant could answer to all statements about the functioning of causal systems solely by relying on high prior knowledge. To test the influence of prior knowledge on the comprehension of causal systems in the respective domain on a more fine-grained level, further studies will be conducted.

Does the Gist of Causal Systems Help in Learning from Subsequent Text?

Studies that experimentally varied the sequence of presenting a text and a corresponding picture (Kulhavy et al., 1993; Ullrich & Schnitz, 2008) have shown that processing of a picture before the corresponding learning text can foster learning. Kulhavy and colleagues (1993) obtained better learning outcomes when a map was presented before a text. According to the authors, the structure of the map helped subjects in learning from subsequent text. However, in these studies, a picture was presented for either three to five minutes, or even without time constraints, which presumably led to a detailed mental model of the picture that was later integrated with the text and thus resulted in higher learning outcomes.

Results of the present study suggest that the gist in causal systems is extracted after very short presentation times (< 150 ms). It is unlikely that the short presentation (150 ms) of a picture already leads to a detailed mental model of the picture. Presumably, it rather acts as a scaffold (Friedman, 1979). Friedman (1979) assumed that subsequent information can then be added to that scaffold, thereby facilitating incremental mental model construction from pictures (and text, cf. Hegarty & Just, 1993). Moreover, Castelano and Henderson (2007) suggested that gist extraction leads to priming of the spatial structure of a picture, and this spatial structure “lingers in memory and can facilitate later perceptual and cognitive operations and behavior” (p. 760). If the gist already provided a scaffold of a picture that can be held in memory for some time, then later information from the text could be added to that scaffold, which could result in better learning. To test this assumption, we further plan to conduct studies, in which the picture of a causal system is presented for a short time (e.g., 150 ms) before a subsequent learning text.

Does Comprehension of the Functioning of Causal Systems Help in Learning from Subsequent Text?

The current results suggest that 600 ms can be enough to gain a preliminary comprehension of the functioning of schematic causal systems. As mentioned before, Stone and Glock (1981) showed that when subjects learned from text and pictures, they first attended to the picture for 1000 to 2000 ms before they started to read the text. Thus, this initial attention on the picture was probably long enough not only to extract the gist but also to gain a preliminary comprehension of the pictures' functioning, which in turn could have led to better learning from subsequent text.

However, it is not yet clear whether subjects in the study from Stone and Glock (1981) actually gained a preliminary comprehension of the picture after initially attending to. Subjects could also have attended to the picture in the first place because it merely was more visually appealing than the text. In this case, the initial attention to the picture possibly might not have been helpful for learning. Thus, further studies will have to investigate whether attending to a causal system for the time necessary to understand its functioning (i.e., 600 ms) fosters learning from subsequent text when subjects are instructed to attend to the system to understand its functioning versus when they are not.

From Basic Cognitive Research to Educational Settings

The study demonstrated that a well established effect in basic cognitive research (rapid gist extraction in scenes) can be found with instructional material as well, thereby providing further insights into the roles that pictures may play in learning from text and pictures. As such, this study can be considered as a starting point of an interdisciplinary approach that tries to better understand the processes that take place during learning from pictures and text through systematically applying findings from basic cognitive psychology to educational scenarios. Besides leading to a better understanding of the learning process, in the long run this approach may also provide recommendations for efficient instructional designs in educational settings.

References

- Biederman, I., Rabinowitz, J. C., Glass, A. L., & Stacy, E. W. (1974). On the information extracted from a glance at a scene. *Journal of Experimental Psychology: General*, 103, 597-600.
- Carroll, P. J., Young, J. R., & Guertin, M. S. (1991). Visual analysis of cartoons: A view from the far side. In K. Rayner (Ed.), *Eye movements and Visual Cognition: Scene Perception and Reading*. New York: Springer.
- Castelhano, M. S., & Henderson, J. M. (2007). Initial scene representations facilitate eye movement guidance in visual search. *Journal of Experimental Psychology: Human Perception and Performance*, 33, 753-763.
- Castelhano, M. S., & Henderson, J. M. (2008). The influence of color on the perception of scene gist. *Journal of Experimental Psychology: Human Perception and Performance*, 34, 660-675.
- Friedman, A. (1979). Framing pictures: The role of knowledge in automatized encoding and memory for gist. *Journal of Experimental Psychology: General*, 108, 316-355.
- Hegarty, M., & Just, M. A. (1993). Constructing mental models of machines from text and diagrams. *Journal of Memory and Language*, 32, 717-742.
- Henderson, J. M., & Hollingworth, A. (1999). High-level scene perception. *Annual Review of Psychology*, 50, 243-271.
- Kulhavy, R. W., Stock, W. A., Verdi, M. P., Rittschoff, K. A., & Savenye, W. (1993). Why maps improve memory for text: The influence of structural information on working memory operations. *European Journal of Cognitive Psychology*, 5, 375-392.
- Levie, W. H., & Lentz, R. (1982). Effects of text illustrations: A review of research. *Educational Communication and Technology*, 30, 195-232.
- Loftus, G. R., Nelson, W. W., & Kallman, H. J. (1983). Differential acquisition rates for different types of information from pictures. *Quarterly Journal of Experimental Psychology-A*, 35, 187-198.
- Mak, P. (2008). Effects of references from text to picture on the processing of school texts: Evidence from eye tracking. In A. Maes & S. Ainsworth (Eds.), *Proceedings EARLI Special Interest Group Text and Graphics: Exploiting the opportunities - Learning with textual, graphical, and multimodal representations*. Tilburg: Tilburg University.
- Mayer, R. E. (1989). Systematic thinking fostered by illustrations in scientific text. *Journal of Educational Psychology*, 81, 240-246.
- Rayner, K., Rotello, C. M., Steward, A. J., Keir, J., & Duffy, S. A. (2001). Integrating text and pictorial information: Eye movements when looking at print advertisements. *Journal of Experimental Psychology: Applied*, 7, 219-226.
- Stone, D. E., & Glock, M. E. (1981). How do young adults read directions with and without pictures? *Journal of Educational Psychology*, 73, 419-426.
- Ullrich, M., & Schnotz, W. (2008). Integration of picture and text: Effects of sequencing and redundancy on learning outcomes. In A. Maes & S. Ainsworth (Eds.), *Proceedings EARLI Special Interest Group Text and Graphics: Exploiting the opportunities - Learning with textual, graphical, and multimodal representations*. Tilburg: Tilburg University.
- Underwood, G., Jebbett, L., & Roberts, K. (2004). Inspecting pictures for information to verify a sentence: Eye movements in general encoding and in focused search. *Quarterly Journal of Experimental Psychology-A*, 57, 165-182.