

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Visuo-Spatial Memory Processing and the Visual Impedance Effect

Permalink

<https://escholarship.org/uc/item/1q66c1v7>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 37(0)

Authors

Albrecht, Rebecca

Schultheis, Holger

Fu, Wai-Tat

Publication Date

2015

Peer reviewed

Visuo-Spatial Memory Processing and the Visual Impedance Effect

Rebecca Albrecht (albrechr@cs.uni-freiburg.de)

Center for Cognitive Science, University of Freiburg

Holger Schultheis (schulth@informatik.uni-bremen.de)

Cognitive Systems, University of Bremen

Wai-Tat Fu (wfu@illinois.edu)

Department of Computer Science, University of Illinois at Urbana-Champaign

Abstract

Models of spatial reasoning often assume distinct visual and spatial representations. In particular, the visual impedance effect – slower response time when more visual details are represented in three-term series spatial reasoning tasks – has been taken as evidence for the distinctive roles of visual and spatial representations. In this paper, we show that a memory model of spreading activation based on the ACT-R architecture can explain the visual impedance effect without the assumption of distinct visual and spatial representations. Using the same memory representation, varying levels of visual features associated with an object are represented in the model. The visual impedance effect is explained by the spreading activation mechanism of ACT-R. The model not only provides a more parsimonious explanation to the visual impedance effect, but also leads to testable predictions of a wide range of memory effects in spatial reasoning.

Keywords: Visual impedance, memory processing, scalable representation, spreading activation, ACT-R, relational reasoning, mental model theory.

Introduction

Processing visual and spatial information is among the most crucial human abilities, because it permeates virtually everything we do (imagine moving in / through the environment without being able to process the visual and spatial information available from your surroundings).

In a seminal paper Ungerleider and Mishkin (1982) argue that in the primate brain two separate pathways are responsible for processing visuo-spatial information: The *what* pathway and the *where* pathway which are associated with the temporal and parietal lobe, respectively. The *what* pathway mainly processes information related to object identification and recognition (e.g., color), while the *where* pathway mainly processes spatial information (e.g., object location or movement). This distinction has subsequently received additional support from many behavioural and neuroscientific studies (e.g., Milner & Goodale, 2008; Klauer & Zhao, 2004).

The existence of these two distinct neural pathways has given rise to the assumption that visuo-spatial information processing in humans draws on two distinct types of mental representations: Visual and spatial representations. Although this assumption is shared by the two main theories of visuo-spatial information processing, the mental model theory (P. Johnson-Laird, 1998) and the theory of mental imagery (Kosslyn, Thompson, & Ganis, 2006), the nature of the representations and, in particular, the relation between the two types of representations (see Sima, Schultheis, &

Barkowsky, 2013, for an in-depth discussion of this point) remains largely to be determined.

In this paper we argue that previous research has not sufficiently considered the role of memory when studying and comparing visual and spatial representations. We demonstrate our case by presenting a memory model that explains the visual impedance effect without the assumption that visual and spatial representations have distinctive functional roles in spatial reasoning. Instead, the model assumes that visual and spatial information can *both* be represented similarly as memory items. However, the model predicts that sometimes additional visual details may slow down the *maintenance* of the memory representations of the information. This is different from the argument by Knauff and Johnson-Laird (2002), who argued that visual representations of information may slow down the *reasoning* process. As a result, contrary to their arguments Knauff and Johnson-Laird (2002), the visual impedance effect does not provide any support to the claim that visual and spatial relations are represented distinctively, nor does it imply that an abstract spatial mental model can lead to a faster reasoning process. Our model not only provides a more parsimonious explanation to the visual impedance effect, but it also has the advantages of having more generality and continuity with other theories in cognitive sciences.

Visual Impedance

Three-term series problems (P. N. Johnson-Laird, 1972) have played a prominent role in investigating spatial and visual representations (e.g., Shaver, Pierson, & Lang, 1975; Knauff & Johnson-Laird, 2002; Rauh, Hagen, Kuss, Schlieder, & Strube, 2005; Schultheis, Bertel, & Barkowsky, 2014; Schultheis & Barkowsky, 2013; Sima et al., 2013). A three-term series problem constitutes a deductive relational reasoning problem in which the relation between two objects, *A* and *C*, has to be inferred given the relations between objects *A* and *B* as well as the relation between *B* and *C*. For example, given the information that (a) the dog is left of the cat and (b) the mouse is left of the dog, participants may be asked to verify the statement that the mouse is left of the cat. Similarly, knowing that (a) the dog is dirtier than the cat and (b) the mouse is dirtier than the dog, one can verify the statement that the mouse is dirtier than the cat.

Knauff and Johnson-Laird (2002) conducted experiments

in which they compare participants performance in solving such three-term series for different types of relations. Specifically, Knauff and Johnson-Laird (2002) distinguish between *visual*, *visuo-spatial*, and *control* relations: Visual relations are relations that are easy to envisage visually (e.g., dirtier); visuo-spatial relations are relations that are easy to envisage spatially and visually (e.g., to the left of); control relations are relations that are hard to envisage both spatially and visually (e.g., better). The main finding reported by Knauff and Johnson-Laird (2002) is that reasoning about visual relations takes significantly more time than reasoning about either visuo-spatial or control relations. This comparatively poor performance of reasoning with visual representations has been termed the *visual impedance effect*.

The explanation provided for the visual impedance effect assumes that for all types of relations the actual reasoning process involves (spatial) mental models (Knauff, Fangmeier, Ruff, & Johnson-Laird, 2003). For visuo-spatial relations and control relations the given information is directly represented in such a spatial mental model and, thus, can immediately be used for reasoning. For visual relations, however, the given information is initially represented by a visual representation (e.g., a visual mental image) that does not support reasoning. To solve the reasoning problem, an additional step for building a spatial mental model is required (Knauff, 2009). Note that this explanation of the visual impedance effect assumes that the comparatively poor performance with visual relations is due to problems associated with the *reasoning* process. Against this background it seems remarkable that all available computational models that formalise reasoning with spatial mental models do not explain the visual impedance effect (Krumnack, Bucher, Nejasmic, Nebel, & Knauff, 2011; Ragni & Knauff, 2013; Khemlani, Trafton, Lotstein, & Johnson-Laird, 2012).

We propose that the visual impedance effect is not a reasoning effect, but a memory effect. Following Schultheis and Barkowsky (2011), we assume that spatial and visual representations are not two distinct, qualitatively different types of representations, but that visuo-spatial representations lie along a continuum that characterizes how specifically amodal / spatial or modality-specific / visual a representation is. Depending on the current task context, representations can flexibly scale along this continuum (i.e., become more or less visual) without a need to modify the reasoning processes working on the representations. Given such scalable representations, we assume that visual relations give rise to more complex representations, because more visual details are represented. Specifically, due to spreading activation of memory items, the additional visual details slower down the access to the information that needs to be processed during the later reasoning process; and, thus, slows down the overall reaction times when solving the three-term series problem.

In the following we first introduce a scalable, abstract representation for relational information that supports reasoning. We then present our model that combines this repre-

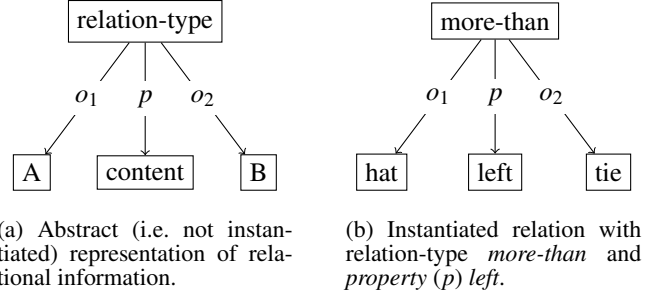


Figure 1: Uniform representation of relational information. A relation consists of a *relation-type*, two objects (o_1, o_2), and a property (p).

sentation with ACT-R’s spreading activation mechanism. We illustrate how our model explains the visual impedance effect and present an ACT-R implementation and simulation of our model. We conclude in highlighting the contribution of our modeling work as well as interesting questions for future work.

Memory Representation

In order to understand the dependency between a representation of relational information in memory on the one hand and a reasoning process on the other we will first introduce a scalable, abstract representation for relational information. In particular, our abstract representation distinguishes between the relation in a mathematical sense, i.e., as it is deemed suitable for reasoning, and the meaning of a relation.

To abstractly represent relational information of the type employed in three-term series problems, we consider *more-than* relation types. A *more-than* relation type consists of three different pieces of information, two objects (o_1, o_2) and a property (p), i.e. *more-than*(o_1, p, o_2). An example is depicted in Figure 1). Intuitively, our representation can be understood as “object o_1 has more of a property p than object o_2 ”. Or more concretely, “the hat is more *left* than the tie” for the visuo-spatial relation “left” and “the hat is more *dirty* than the tie” for the visual relation “dirty” (cf. Figure 1).

We further assume that for a concrete relational statement each of the three arguments is associated with features in memory that represent the arguments’ meaning. We define the *content* of an object (or property) as the tuple of features necessary to represent the object o , i.e. $content(o) = (f_1^o, \dots, f_n^o)$, and property p , i.e. $content(p) = (f_1^p, \dots, f_n^p)$, in memory. Figure 2 shows graphically how relational statements and their contents are represented. This representation includes an abstract representation suitable for reasoning (Figure 2, above the red line) and the memory representation defined by the content of the relational information (Figure 2, below the red line).

Scalability of this representation is defined in terms of the number of features involved in representing the relational statement. In particular, the representation can scale to become more or less visual depending on how much visual fea-

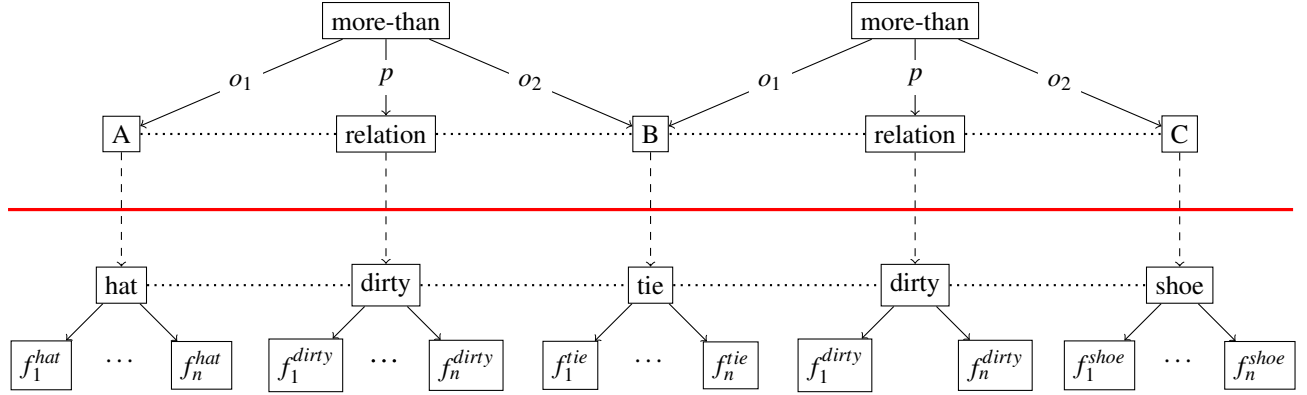


Figure 2: Representation of relational information ‘the hat is dirtier than the tie’ and ‘the tie is dirtier than the shoe’ in memory assuming a hierarchical representation of information. Above the red line is a representation which supports reasoning with relations. Below the red line is the memory representation of the objects and property used to elaborate the content of relational information. Scalability is defined in terms of the number of features necessary to represent relational information.

tures are associated with the relational statement. For example, the relation “dirtier” may be associated with features like “dirt”, “mud”, “black dots”, etc. and, thus yield a more visual representation than the relation “to the left of”, which may only be associated with a single feature “position”.

Cognitive Model

In this section we describe an ACT-R model that explains the visual impedance effect. Employing the above described representation, the model explains the effect as a memory phenomenon arising from spreading activation.

ACT-R Spreading Activation

ACT-R realises working memory as a structured set of buffers (Anderson, 2007). Buffers hold declarative information, so-called *chunks*. A chunk is a set of key-value (or slot-value) pairs. For example, a chunk representation of the introduced ‘more-than’ relation has three slots (o_1, p, o_2) to which values (often also chunks) can be assigned. Behaviour in ACT-R is produced by the repeated application of production rules that fit a current working memory state and change the working memory state according to their definition. Changes to the working memory come about by requests to modules that are associated with buffers. Modules process requests by updating the chunks contained in their buffers. This processing is associated with a time cost and, in some cases, has an uncertain outcome.

The ACT-R declarative module (sometimes called declarative memory) holds all declarative information known to a model such as, for example, the complete representation depicted in Figure 2. The time it takes the declarative module to process requests depends on the activation values assigned to candidate chunks. While a number of mechanisms can influence the chunks’ activations, we focus our analysis to spreading activation.

The spreading activation of a chunk c in ACT-R is defined in terms of a signal strength S between chunk c and all chunks

d which are a part of the current working memory state (i.e., assigned to a buffer). Formally, the signal strength between chunk c and a chunk d is computed as $S_{d,c} = S - \ln(fan_d)$. The signal strength depends on the number of outgoing connections of chunk d , a concept which has been termed *fan* of chunk d . The signal strength additionally depends on a global constant S which has been interpreted psychologically as an approximation of the declarative memory size. The complete spreading activation of a chunk c is calculated by ¹

$$sa(c) = \sum_{d \text{ in working memory}} S_{d,c}$$

The time to retrieve a chunk c from declarative memory is defined with respect to the activation of chunk c , in our case $RT(c) = a \cdot e^{-sa(c)}$, where a is a constant. The higher the activation of a chunk the lower the response time. For spreading activation a greater fan implies a lower activation and, thus, a higher response time.

Example. It may be helpful to more closely consider how the spreading activation mechanism explains the visual impedance effect. The visual impedance effect is measured in the time to verify a given conclusion. When a conclusion needs to be verified, information from the first and second relational statement have already been integrated in a mental representation. Depending on assumptions stated in a reasoning theory this mental representation may, for example, be a mental model (Ragni & Knauff, 2013) or a relational inference, i.e., a relational statement, (Braine & O’Brien, 1998). In either case, this mental representation needs to be retrieved from declarative memory in order to verify the conclusion. In

¹For the sake of representation simplicity we assume that spreading activation is enabled for every buffer and that all buffers are assigned the same weight, which sum up to a total of 1. Additionally, in our scenario the working memory holds the same number of chunks in every request. Thus, we leave out the weight of the ACT-R spreading activation equation.

the following, we will use a mental model as the mental representation. The argument for a relational inference is analogue and the simulation results for the retrieval time only differ by a constant factor due to the different representation.

Consider our scalable representation defined for relational information, i.e., $more-than(o_1, p, o_2)$. When a mental model is requested from the declarative module the conclusion is represented in the model's working memory (e.g., the goal buffer). Thus, o_1 , o_2 and p are potential sources for spreading activation (cp. Figure 3). The signal strength between the mental model chunk (*model*) and the content of a relation p is then $S_{p,model} = S - \ln(fan(p)) + r$. The fan is influenced by the number of features associated with the relation, i.e. $fan(p) = content(p)$, and the constant r approximating the fan associated with reasoning representation (e.g. a second mental model in declarative memory). The signal strength can thus be calculated as $S_{p,model} = S - \ln(content(p))$. Consequently, the more features are necessary to represent the content of a relation, the higher the retrieval time of a target chunk (due to the higher fan).

Now consider the concrete relations introduced for the visual impedance effect, that is, visual relations like "dirtier", visuo-spatial relations like "to the left of" and control relations like "better". If we assume that the visuo-spatial relation "to the left of" can be represented using one feature (e.g., the position), the fan is $fan(left\ of) = content(left\ of) = 1$ (Figure 3b). For visual relations like "dirtier", on the other hand, more features need to be represented (e.g., dirt, mud, etc.) Therefore, the fan associated with "dirtier" is higher than the fan associated with "to the left-of", i.e. $fan(dirtier) > fan(left\ of)$, and the signal strength between chunks "dirtier" and *model* is lower than between "to the left of" and *model*, i.e. $S_{dirtier,model} < S_{left\ of,model}$ (Figure 3a). Assuming that the features representing objects o_1 and o_2 are the same in both cases², the spreading activation that the mental model chunk receives from the visual relation "dirtier" is given by

$$sa(model) = \underbrace{S - \ln(k)}_{S_{hat,model}} + \underbrace{S - \ln(n)}_{S_{dirtier,model}} + \underbrace{S - \ln(l)}_{S_{shoe,model}}$$

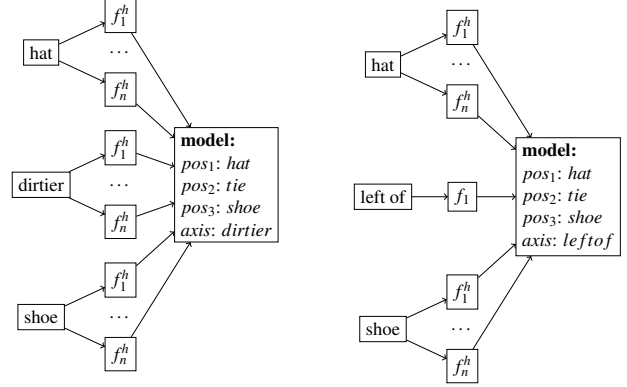
and from the spatial relation "to the left of" is given by

$$sa(model) = \underbrace{S - \ln(k)}_{S_{hat,model}} + \underbrace{S - \ln(1)}_{S_{left\ of,model}} + \underbrace{S - \ln(l)}_{S_{shoe,model}}$$

Obviously, the mental model chunk receives more spreading activation for visuo-spatial relations. Therefore, the retrieval time is lower (see Figure 4).

ACT-R implementation. Our ACT-R model is based on the ACT-R implementation of PRISM (Ragni, Fangmeier, &

²Instead of keeping the features of the objects fixed and varying the features of the relation, it would also be possible to represent the objects by more of less features depending on the relation type. This would not impact the explanatory power of our model w.r.t the visual impedance effect.



(a) Fan for chunks in working memory for conclusion relation "dirtier". The content of relation "dirtier" needs to be represented by more than one feature (e.g., dirty, mud, black spots, etc).

(b) Fan for chunks in working memory for conclusion relation "to the left of". The content of relation "to the left of" can be represented by a single feature (e.g., the position).

Figure 3: Example illustration of memory representations for visual (a) and visual-spatial relations. Chunks active in the working memory are "hat", "shoe", and "dirtier" (3a), and "to the left of" (3b). The target chunk is a mental model that needs to be retrieved in order to verify the conclusion. Due to the higher fan the target chunk (here a mental model) receives more spreading activation for the conclusion relation "to the left of" than for the conclusion relation "dirtier". Thus, the retrieval time is higher for the conclusion relation "dirtier".

Brüssow, 2010), which assumes that exactly one retrieval is necessary to verify a given conclusion.

For a prototypical task such as "the hat is dirtier than the tie", the "tie is dirtier than the shoe", "is the hat dirtier than the shoe?" we define the mental representation as either a mental model chunk or a relational inference chunk which is stored in declarative memory, e.g.,

- (r1 ISA model pos1 hat pos2 tie pos3 shoe rel dirtier)
- (r2 ISA inference o1 hat o2 shoe rel dirtier)

Additionally we represent the features associated with objects and the property of a relation as content chunks, e.g.,

- (l1 ISA content id left-of feature l1)
- (d1 ISA content id dirtier feature dd)

Source of spreading activation is a representation of the conclusion in the goal buffer, e.g.,

(rell ISA more-than o1 hat p dirtier o2 shoe)

We define one production rule which requests a mental representation chunk (model or inference) from declarative memory. This request does not specify any restrictions on slot values other than the type being either a mental model or an inference. The time it takes the declarative module to answer this request depends on the fan associated with the objects and the property of the relation, that is, the number of content chunks associated with the objects and the property.

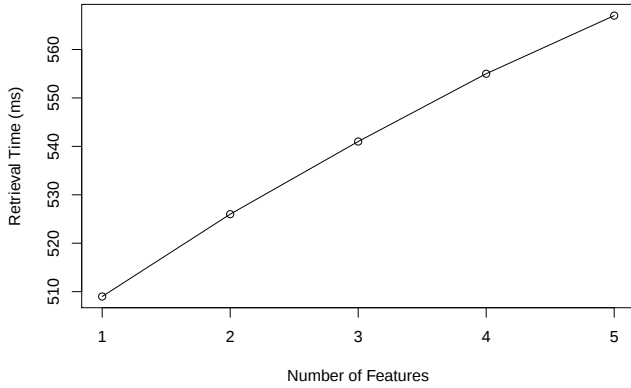


Figure 4: Retrieval time for a mental representation chunk (here mental model chunk) when a conclusion needs to be verified. The response time increases independently of a concrete task or reasoning theory with the number of visual features associated with the objects and the relation.

All ACT-R parameters are set to their default values. We approximate parameter S by the logarithm of the average size of the model’s declarative memory (i.e. $S = 3$). Figure 4 shows how the retrieval time increases with the number of features associated with the content of a relation.

Accordingly, our model accounts for the visual impedance effect by predicting that verification of conclusions for three-term series problems involving visual relations such as “dirtier” take more time than for problems involving visuo-spatial relations. Interestingly, our model predicts that the visual impedance effect should not be restricted to the use of visual relations, but should arise whenever the reasoner is inclined to associate multiple (visual) features with a relational statement. As discussed below, existing evidence supports this prediction.

Conclusion

Knauff and Johnson-Laird argued that the reason why items that could easily be envisaged would lead to slower response times was that visual representations of irrelevant features slowed down the reasoning process. We provided an alternative explanation: the easily envisaged items took longer to be accessed in memory because they were associated with more visual features, which slowed down their access time as predicted by the spreading activation mechanism. Contrary to the argument by Knauff and Johnson-Laird, we did not find that the visual impedance effect provided any support to the claim that easily envisaged items were represented by a visual representation that was functional different from a (spatial) mental model, nor did the results support the claim that “visual imagery as the medium for reasoning would be implausible” (Knauff & Johnson-Laird, 2002).

Our work shows that combining the concept of scalable representation structures with spreading activation provides a more parsimonious explanation to the visual impedance effect, as the proposed model does not assume distinct visual and spatial representations or a specific reasoning process. The current model uses memory representation of objects and memory processes that have been used to explain a wide range of memory effects (e.g., in previous ACT-R models of memory tasks). The current model therefore has the potential to lead to a wide range of testable predictions on the effects of memory in spatial reasoning, such as effects of individual differences in working memory capacity, interference effects, or effects of memory decay. In addition to the original visual impedance effect (Knauff & Johnson-Laird, 2002), our modeling work also explains moderations of the effect that have been reported. If the visual impedance effect is due to memory processing as assumed in the proposed model, it should scale with the model’s ability and necessity to represent specific features in order to maintain a representation suitable for reasoning. Consistent with our model, research shows that blind people show no visual impedance effect (Knauff, 2009)—perhaps because they are less inclined to represent objects with visual features, or the number of visual features tend to be lower for blind people. Furthermore, people who have a higher tendency to visualize object details show a stronger visual impedance effect (Castañeda & Knauff, 2013), because they tend to represent more visual features as other groups.

The proposed cognitive model investigates the impact of the memory representations of visual features on spatial reasoning. The model, however, does not make any assumption on the reasoning process, as the reasoning tasks are the same across the different conditions in the studies on visual impedance (Knauff & Johnson-Laird, 2002). In other words, the explanation of the visual impedance effect by our model is independent of the reasoning process. For example, if we apply the reasoning process in the PRISM model (Ragni et al., 2010), in which the main difference in level of difficulty in spatial reasoning tasks is characterized by the number of focus operations on the represented objects, we will have the same number of focus operations in each condition, and the only difference is how quickly the model can assess the objects represented in memory as the focus operations are applied. However, we should point out that the PRISM model by itself does not seem to be able to explain the visual impedance effect. On the other hand, our model can be used with other reasoning theories (e.g., (Krumnack et al., 2011; Braine & O’Brien, 1998)) to explain the visual impedance effects. In other words, our model suggests that the visual impedance effect can be explained by memory processes rather than reasoning processes.

The goal of this paper is to show that the visual impedance effect can be explained without committing to a unique spatial representation that is distinct from visual representation. This is consistent with the idea that the long debate about the

role of visual imagery in spatial reasoning can be resolved by considering the visuo-spatial representation as a continuum (Schultheis and Barkowsky (2011)), with varying levels of visual (and spatial) features represented in memory. Another advantage of this approach is that by utilizing memory representations and mechanisms, the model is more readily compared and tested against a wide range of cognitive phenomena beyond spatial reasoning. We believe that our modeling approach and results constitute an important first step towards studying the impact of memory processing in human reasoning and the nature of spatial and visual representations. While previous theories and studies mostly restricted considerations to the reasoning process or the representation, we define a link between these concepts. As a result, our approach also highlights promising avenues for future work, both empirical and computational, to shed more light on aspects of reasoning processes and representations.

Empirically, we propose experiments that explicitly control the number of represented features, both of objects and relations. Such an experiment would yield valuable results on the effect size with respect to the number of features necessary to represent concepts. A possible approach is using a high and low similarity conditions similar to Folk and Luce (1987), that is, where more or less features need to be represented in order to draw conclusions (e.g., “the red hat is dirtier than the tie” vs. “the red hat is dirtier than the red tie” vs. “the red hat is dirtier than the blue tie”).

Computationally, assuming visual impedance is in fact an effect in memory processing our results can be used to further examine reasoning theories. The ACT-R implementation of the PRISM model represents the complete mental model in one chunk and approximates focus operations as a constant factor. However, according to the ACT-R theory information is usually represented as linked lists. Thus, if a mental model was represented as a linked list a focus operation would in fact be a request to declarative memory. In this case, our memory model would predict a linear increase in the response time for visuo-spatial relations with the number of focus operations.

References

- Anderson, J. R. (2007). *How can the human mind occur in the physical universe?* Oxford University Press.
- Braine, M. D., & O'Brien, D. P. (1998). *Mental logic*. Psychology Press.
- Castañeda, L. E. G., & Knauff, M. (2013). Individual differences, imagery and the visual impedance effect. In M. Knauff, M. Pauen, N. Sebanz, & I. Wachsmuth (Eds.), *Proceedings of the 35rd annual conference of the cognitive science society*. Austin, TX: Cognitive Science Society.
- Folk, M. D., & Luce, R. D. (1987). Effects of stimulus complexity on mental rotation rate of polygons. *Journal of Experimental Psychology: Human Perception & Performance*, 13, 395 - 404.
- Johnson-Laird, P. (1998). Imagery, visualization, and thinking. *Perception and cognition at century's end*, 441-467.
- Johnson-Laird, P. N. (1972). The three-term series problem. *Cognition*, 1, 57-82.
- Khemlani, S., Trafton, J. G., Lotstein, M., & Johnson-Laird, P. (2012). A process model of immediate inferences. In *Proceedings of the 11th international conference on cognitive modeling* (p. 151).
- Klauser, K. C., & Zhao, Z. (2004). Double dissociations in visual and spatial short-term memory. *Journal of Experimental Psychology: General*, 133(3), 355-381.
- Knauff, M. (2009). A neuro-cognitive theory of deductive relational reasoning with mental models and visual images. *Spatial Cognition & Computation*, 9(2), 109-137.
- Knauff, M., Fangmeier, T., Ruff, C. C., & Johnson-Laird, P. N. (2003). Reasoning, models, and images: Behavioral measures and cortical activity. *Journal of Cognitive Neuroscience*, 15(4), 559-573.
- Knauff, M., & Johnson-Laird, P. (2002). Visual imagery can impede reasoning. *Memory & Cognition*, 30(3), 363-371.
- Kosslyn, S. M., Thompson, W. L., & Ganis, G. (2006). *The case for mental imagery*. New York: OUP.
- Krumnack, A., Bucher, L., Nejasmic, J., Nebel, B., & Knauff, M. (2011). A model for relational reasoning as verbal reasoning. *Cognitive Systems Research*, 11, 377-392.
- Milner, A., & Goodale, M. (2008). Two visual systems reviewed. *Neuropsychologia*, 46(3), 774 - 785.
- Ragni, M., Fangmeier, T., & Brüssow, S. (2010). Deductive spatial reasoning: From neurological evidence to a cognitive model. In *Proceedings of the 10th international conference on cognitive modeling* (pp. 193-198).
- Ragni, M., & Knauff, M. (2013). A theory and a computational model of spatial reasoning with preferred mental models. *Psychological review*, 120(3), 561.
- Rauh, R., Hagen, C., Kuss, T., Schlieder, C., & Strube, G. (2005). Preferred and alternative mental models in spatial reasoning. *Spatial Cognition & Computation*, 5(2), 239-269.
- Schultheis, H., & Barkowsky, T. (2011). Casimir: an architecture for mental spatial knowledge processing. *Topics in Cognitive Science*, 3(4), 778-795.
- Schultheis, H., & Barkowsky, T. (2013). Variable stability of preferences in spatial reasoning. *Cognitive Processing*, 14(2), 209 - 211.
- Schultheis, H., Bertel, S., & Barkowsky, T. (2014). Modeling mental spatial reasoning about cardinal directions. *Cognitive Science*, 38(8), 1521-1561.
- Shaver, P., Pierson, L., & Lang, S. (1975). Converging evidence for the functional significance of imagery in problem solving. *Cognition*, 3(4), 359 - 375.
- Sima, J. F., Schultheis, H., & Barkowsky, T. (2013). Differences between spatial and visual mental representations. *Frontiers in psychology*, 4.
- Ungerleider, L., & Mishkin, M. (1982). Two cortical systems. *Analysis of Visual Behavior*, 549-586.