

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Schema Drift: Relational Concept Stability Across Repeated Comparison

Permalink

<https://escholarship.org/uc/item/9sv0h8x9>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 46(0)

Authors

Vagnino, Richard

Walker, Caren M.

Publication Date

2024

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Schema Drift: Relational Concept Stability Across Repeated Comparison

Richard Vagnino (rvagnino@ucsd.edu)

Department of Philosophy, University of California, San Diego

Caren M. Walker (carenwalker@ucsd.edu)

Department of Psychology, University of California, San Diego

Abstract

Analogical reasoning is one of the most common ways individuals bring previous experience to bear on unfamiliar situations. Most theories describe this process as a structured comparison that involves mapping the relational properties between a familiar source and unfamiliar target. This both allows the transfer of useful inferences from the source to the target and highlights the common structure shared by both analogs, represented by an abstract schema. This schema can help with identifying and reasoning about structurally similar situations in the future. While researchers have studied how representations of source and target analogs undergo alterations as a result of this mapping process, little attention has been paid to how the abstract schemas thought to guide future analogical reasoning might similarly change with use. We explore this question in two experiments and present evidence that suggests abstract schemas do indeed drift under certain conditions.

Keywords: analogical reasoning; abstraction; schema; conceptual change; rerepresentation

Introduction

“Nature never repeats herself,” or so the saying goes. And yet, the more we experience and learn about the world, the better we tend to be at navigating this fundamental uncertainty. How, despite the uniqueness of the present, are we able to leverage our knowledge of the past to help us understand new situations and overcome unfamiliar problems? One of the most common and powerful ways to bring the familiar to bear on the novel is through analogical reasoning.

In the broadest sense, analogical reasoning can be thought of as the “structured comparison between two situations” where similarities between relational properties in both domains are of special importance (Vendetti et al. 2014, p. 1172). Relational properties, as the name suggests, express relations between two or more targets. For instance, the property of being a predator is a relational property, because it consists in the relationship between at least two things (namely a predator and some other entity, its prey). Relational properties are often contrasted with what are sometimes called *attributes* — properties that can be expressed by predicates which take a single argument, such as “being the color red” (Holyoak and Thagard, 1995, p. 25).

While some relational properties involve fairly simple relations (e.g., being larger than or to the left of), many are abstract (e.g., being the cause of). This emphasis on relational properties highlights an important feature of analogical cognition, namely that analogs (i.e., the situations or entities

that are the subject of comparison) can differ, sometimes radically, in terms of their attributes as long as they share a sufficient number of relational properties; in other words, analogies depend on similarities between relational structures (although surface similarity may play an important role in the detection and spontaneous generation of analogies) (Holyoak and Koh, 1987).

While there are numerous computational models of analogical reasoning on offer, most divide the process into a number of distinct phases, which include the selection of some relevant source analog (which may be preceded by a phase where cues guiding this process are detected), followed by a mapping process in which correspondences between the relational structure of analogs are recognized and connected, and ending with the transfer or projection of inferences from the more familiar *source* analog to the novel *target* in order to guide reasoning, action, or learning (Gentner, 1983; Holyoak and Thagard, 1995).

Some theories suggest (or even require) that mappings between analogs be one-to-one (meaning every element in the source analog maps onto a distinct element in the target analog) and structurally consistent (meaning that elements between analogs are mapped onto their counterparts). However, in everyday episodes of analogical cognition, it is rarely the case that source and target analogs happen to be structured this way purely by accident. One solution to this mismatch suggested by both theoretical models and experimental evidence is that both the analogical source and target can undergo *rerepresentation*, such that analogs which, strictly speaking, do not share a common (i.e., identical or synonymous) relational structure can be transformed to make mapping possible. For instance, the scenarios “the tiger broke out of its enclosure” and “the criminal broke out of jail” can both be understood in terms of the same relations (i.e., both involve entities breaking free from physical captivity); however, “the tiger broke out of its enclosure” and “Mary got out of a possessive relationship” cannot be understood in this way (Day & Asmuth, 2017). Instead, these situations need to be represented in more abstract terms — a process known as *semantic ascent* — in order to be mapped (Oberholzer et al., 2018). In this case, both might be represented to reflect the relation of entities escaping a confining situation or becoming free. In other cases, rerepresentation may create rather than merely reveal similarities between situations, when both are judged to be members of the same schema-governed category. In such instances, the schema-governed category may be projected from the more familiar or *canonical*

exemplar, to the less obvious or accepted one (Oberholzer et al., 2018).

Still, rerepresentation appears to have its limits. In one study, Day and Asmuth (2017) had participants in two conditions — *Same* and *Different* — rate two pairs of situations — a Base comparison and Test comparison — in terms of similarity on a 15 point scale from *Very dissimilar* to *Very similar*. Similarity was chosen as the variable of interest based on previous work in which it was observed to be a reliable indicator of representational change (Boroditsky, 2017; Goldstone, Lippa & Shiffrin, 2001). In the *Same* condition, the first comparison consisted of a *Standard* sentence and a sentence related to the *Standard* in a particular way, which they called *Relation A* (e.g., a sentence about a student headed to college and a baby bird learning to flying might both express the relation LEAVING_HOME). The second comparison also contained the *Standard* sentence and a new sentence that also expressed *Relation A* (e.g., a sentence about a baseball team travelling out of state for a game). The first comparison in the *Different* condition likewise consisted of the *Standard* sentence and *Relation A*; however, the second comparison brought together the *Standard* sentence and a sentence that expressed a different relation, which they called *Relation B* (e.g., the previous sentence about the college student and a sentence about an expedition to a distant land expressing the relation EXPLORING_UNKNOWN). Participants in the *Same* condition reported significantly higher similarity ratings for the second comparison than their counterparts in the *Different* condition, suggesting that once a situation or entity has undergone rerepresentation during the structure mapping process, it appears less likely to easily undergo rerepresentation again.

Both inference projection resulting from a successful analogical mapping and rerepresentation to promote analogical coherence have been shown to lead to long lasting effects in the representations of target and source domains. For instance, Blanchette and Dunbar (2002) had participants read an informational passage on an unfamiliar target. In the experimental condition, the final paragraph contained a description of an analogous situation (i.e., it served as a source analog). The control condition featured the same informational passage first; however, the final paragraph simply included more details about the target. After a brief distractor, both groups were given a recognition test containing a combination of statements either present in the text, not present in the text, or not present but *inferable by analogy* from the source passage. They found that participants in the experimental condition were nearly twice as likely as those in the control to falsely recognize analogical inferences as part of the original text.

Likewise, Vendetti et al. (2014) found that non-alignable differences in the source analog (e.g., properties of the source that differ from the target and cannot be mapped between the two) were assimilated — rather than temporarily suppressed during mapping — permanently altering source analog representations. Vendetti et al. hypothesized that, since non-

alignable differences weaken the coherence of an analogy, those differences would either be temporarily suppressed and made less salient (but still remain a part of the source analog) or permanently altered (i.e., assimilated) in order to better subserve coherence. They found evidence for the latter.

Because the structure mapping process highlights the system of shared relations common to both analogs, analogical reasoning has been suggested as an important mechanism for the acquisition of relational concepts during development (Gentner & Kurtz, 2005). Similarly, reasoners who encounter multiple instances of analogies which share the same relational structure will eventually derive an abstract representation of that structure in the form of a schema (Gick & Holyoak, 1980; Gentner & Hoyos, 2017). These schemas help individuals identify structurally similar situations in the future and reason about them more efficiently. Despite interest in the effects of rerepresentation on source and target analogs, little work, to the authors' knowledge, has considered how abstract schemas may likewise change with repeated activation. The present study aims to address this lacuna.

The Present Study

Our study examines the question of whether relational schemas derived from structured comparisons can undergo change or “drift” after repeated activations involving semantically disparate domains. For instance, will the relational schema for *robbery* represented as (THIEF, GOODS, VICTIM) (Gentner & Kurtz, 2005, p. 153) change if used repeatedly in subsequent comparisons from domains such as chemistry (e.g., one atom stealing an electron from another), entertainment (e.g., an entertainer stealing the show), and romance (e.g., stealing your lovers heart)? Do only source and target analogs undergo rerepresentation to promote analogical coherence, or are schemas likewise subject to alteration in order to accommodate a wider range of figurative comparisons?

Prior research suggests that relational concepts (e.g., *predator*, which is defined by its extrinsic relations to other concepts) are more malleable than entity concepts (e.g., *banana*, which is defined by certain intrinsic properties) across a wider range of contexts (Asmuth & Gentner, 2016). Two possible hypotheses could explain the purported mutability of relational concepts (which we take to be akin to our notion of drift). Either a) the representational structure of relational concepts (expressed by abstract schemas) is fixed, and their contextual flexibility is entirely due to the rerepresentation of the concepts that act like arguments, occupying the open slots in their structure (e.g., rerepresenting the concept “electron” to fit the GOODS role in the *robbery* schema); or b) the representational structure of relational concepts can — under the right conditions — undergo change as well. The present study comprises two experiments designed to test these competing hypotheses. Experiment 1 aims to replicate and extend the study conducted by Day and Asmuth (2017), further investigating

Table 1: Order of Stimuli Presentation

Comparison	Same	Different
Base	Standard + A	Standard + B
Test 1	Standard + A	Standard + C
Test 2	Standard + A	Standard + D
Test 3	Standard + A	Standard + E
Test 4	Standard + A	Standard + B

the rerepresentation effect they reported. Experiment 2 directly investigates whether repeated activations of the same relational concept across different semantic domains can impact the representational structure of that concept.

Experiment 1

Our goal in this first study was to replicate Day and Asmuth's (2017) findings and investigate the rerepresentational effect they describe more deeply. In order to see what impact repetition would have on this effect, we modified their design to include four test comparisons instead of one. We also altered the final comparison in the *Different* condition to express the same relation as the initial base comparison to assess whether participants would still rate a sentence pair expressing the base relation highly after being exposed to numerous non-matching comparisons.

Participants

30 undergraduates at a large, public West coast university participated in the study in exchange for credit. This sample size was based on Day and Asmuth's original study.

Materials and Procedure

The study was administered online using the Qualtrics platform. After reading the instructions, participants were shown two example comparisons to familiarize them with the task. Subsequently, participants were shown 10 pairs of sentences, one at a time, split into two conditions (*Same* and *Different*). For each pair of sentences, participants were asked to rate how similar they perceived the two situations by moving a slider along a 15-point scale, ranging from *Very dissimilar* to *Very similar*.

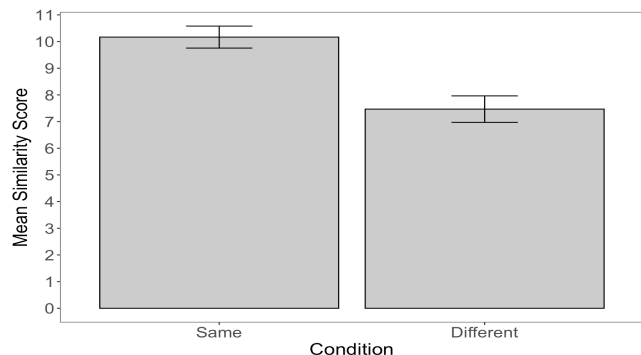
Stimuli consisted of two *Standard* sentences (one for each condition) and ten test sentences. For the *Same* condition, the *Standard* sentence was written to be paired with five test sentences all expressing the same relation (see Table 2). For instance, the first sentence in Table 2 about the tiger with a penchant for freedom, and the second sentence about the careless computer scientist, are both examples of "something dangerous being accidentally released." For the *Different* condition, the *Standard* sentence was designed to be compatible with five test sentences expressing four different relations in total (see Table 3). The sentences about studious Alice and awkward Joel in Table 3, for example, can both be represented as examples of individuals "avoiding social interactions with others." All sentences pairs were piloted to ensure high similarity ratings prior to any effects generated by the study design.

During the experiment, participants were first shown a *Base Comparison* (terminology adopted from Day & Asmuth, 2017) consisting of the *Standard* sentence and a test sentence and presented with the prompt "How similar are these two situations." They rated this sentence pair on the 15-point scale described above. In the *Same* condition, participants were subsequently shown four more test comparisons, each composed of the *Standard* sentence and a test sentence expressing the same relation as the *Base Comparison*. The *Different* condition followed the same format, except that first, second, and third comparisons expressed different relations than the *Base Comparison*. In both conditions, we randomized the order of test comparisons 1-3 to mitigate the impact of order effects.

Results and Discussion

Since trials in Day and Asmuth (2017) consisted of one Base Comparison and a single Test Comparison, we began our analysis with Test Comparison 1 (Test 1 in Table 1). Using a paired-samples t-test, we found a significant difference between conditions (see Figure 1; $t(29) = 5.34, p < .001, d = .086$), with similarity ratings for Same condition trials ($M = 10.17, SD = 2.26$) significantly higher than the Different condition trials ($M = 7.47, SD = 2.72$). These results closely resemble the findings of the original study.

Figure 1: Test 1 Mean Similarity Score by Condition



Since we were also interested in the effect of repetition on both groups, we analyzed Test Comparison 4 as well. Recall that for the *Different* condition, participants once again saw the *Standard* sentence plus a test sentence reflecting the same relation as the *Base Comparison* for that condition. In this instance, a paired-samples t-test did not reveal a significant difference between conditions ($t(29) = -1.62, p = .117$), with similarity ratings for Same condition ($M = 11.1, SD = 2.56$) and Different condition ($M = 12.17, SD = 2.56$) trials remaining strong.

Table 4: Mean Similarity Scores for All Trials

Condition	Test 1	Test 2	Test 3	Test 4
Same	10.17	10.9	12.2	11.1
Different	7.47	7.7	6.67	12.17

Table 2. *Same Condition Sentences*

Sentence Type	Event Description	Relation Description
Standard	As the zoo keeper was busy cleaning its habitat, the Siberian tiger was able to escape its open cage into the nearby city.	
Relation A	While testing a network security system, the computer scientist inadvertently released a destructive virus onto the internet.	RELEASE
Relation A	While working on a cure for disease, the virologist accidentally carried a sample out of the lab on their shoe.	RELEASE

Table 4 shows the means for both conditions for all four test trials and illustrates the general pattern we found in our analysis: a characteristic drop in similarity ratings for comparisons not matching the *Base Comparison* relation that persists over numerous repetitions. Further, the rerepresentation effect persists even after participants have read and rated numerous comparisons that do not reflect the *Base Comparison* relation (as was the case in the *Different* condition). This suggests the rerepresentation effect described by Day and Asmuth (2017) is quite robust.

Experiment 2

Our aim in the second study was to directly investigate whether repeated activations of the same relational concept across a wide range of different contexts could induce a change to its representational structure in order to promote coherence during its use in a variety of comparisons. We further modified Day and Asmuth's (2017) design, eliminating the *Standard* sentence, but retaining one sentence between subsequent comparisons.

Participants

60 undergraduates at the same university participated in the study in exchange for credit. Participants were randomly assigned to one of two conditions, with 30 per condition. An additional three undergraduates were tested but excluded, because the difference between their initial and final similarity ratings fell more than two standard deviations outside the mean and deemed outliers. Again, the sample size was based on Day and Asmuth's original study.

Materials and Procedure

The study was administered online using the Qualtrics platform. Participants were divided into two groups: experimental and control. Both groups received the same instructions and example comparisons as participants in Experiment 1. Both groups were shown five pairs of sentences, one at a time, followed by a brief distractor task, and then one final comparison. For each pair of sentences, participants were asked to rate how similar they perceived the two situations to be to one another by moving a slider along a 15-point scale, ranging from *Very dissimilar* to *Very*

similar. Participants were also prompted to provide one sentence explanations for their ratings.

Stimuli consisted of 21 sentences. Sentences designed for the experimental group were further divided into two versions each composed of eight sentences. Sentences in the first version, when compared, expressed the RELEASE schema (i.e., something harmful being unintentionally released). Sentences in the second version expressed the AVOID schema (i.e., individuals avoiding contact with someone or something). The control group stimuli contained some sentences recycled from the experimental group trials, as well as five additional sentences not used in the experimental group trials.

In place of a *Standard* sentence, we structured the comparisons to always retain one sentence from the previous pair. For instance, the first comparison in an experimental group trial might consist of this pair:

As the zoo keeper was busy cleaning its habitat, the Siberian tiger was able to escape its open cage into the nearby city.

While testing a network security system, the computer scientist inadvertently released a destructive virus onto the internet.

The following comparison would then begin with the second sentence from the previous pair, coupled with a new sentence expressing the same relation:

While testing a network security system, the computer scientist inadvertently released a destructive virus onto the internet.

When the instructor turned around to write something on the board, Eric slipped out of the boring lecture to vandalize the school's restrooms.

In the experimental group trials, every comparison expressed the same relation. In the control group, no common relation was preserved across all comparisons. In both groups, after the distractor task, participants saw one final comparison, which was nearly identical to the first comparison except for two minor differences. First, a small number of arbitrary details unrelated to the relational

Table 3. *Different Condition Sentences*

Sentence Type	Event Description	Relation Description
Standard	Despite invitations from other students, Alice spent her entire Friday night in the school library.	
Relation B	Andrew was very serious about his musical ambitions, and he made a point of avoiding any distractions that would keep him from his practicing.	FOCUSED_ON
Relation C	Joel had always been somewhat socially awkward, and when he started his new job he avoided eating lunch in the breakroom with the other employees.	AVOID

structure of the sentences were altered in order to lessen the likelihood that participants would immediately recognize the sentence and simply recall their previous rating. For instance, “Siberian tiger” was changed to “Burmese python” in one final comparison. Otherwise, the two sentences were identical. Research suggests that details like these are easily abstracted away during structured comparisons and therefore should have no impact on similarity ratings. Second, the order of the two sentences was flipped in order to preserve the overall structure described above.

Table 5 lays out the order of comparisons used in our design. We predicted that after reading numerous comparisons across a wide range of different semantic contexts that participants in the experimental condition would rate the final comparison significantly lower than they rated the first comparison, despite the two pairs of sentences being virtually identical.

Table 5: Comparison Order for Experiment 2

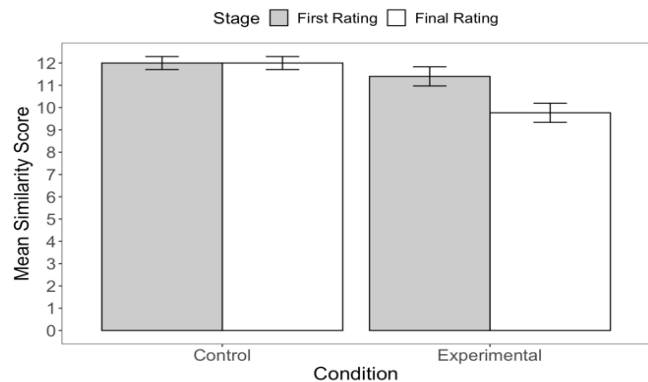
Phase	Comparison
Test Comparison 1	Sentence A + Sentence B
Test Comparison 2	Sentence B + Sentence C
Test Comparison 3	Sentence C + Sentence D
Test Comparison 4	Sentence D + Sentence E
Test Comparison 5	Sentence E + Sentence B*
Distractor	
Test Comparison 6	Sentence B* + Sentence A*

Results and Discussion

A paired-samples *t*-test revealed a significant difference between the final similarity judgements between groups (see Figure 2; $t(29) = -3.36, p = .002$), with similarity ratings for experimental group trials ($M = 9.77, SD = 2.35$) being significantly lower than control condition trials ($M = 12, SD = 1.6$). We also ran a paired-samples *t*-test on trial 1 for both groups $t(29) = -1.07, p = .29$ which revealed no significant difference between the experimental group ($M = 11.4, SD = 2.35$) and control ($M = 12, SD = 1.6$) confirming the effect was due to our manipulation and not a pre-existing difference

between groups. Additionally, we ran a paired-samples *t*-test within each condition to analyze the effect of repeated comparison on initial vs. final similarity judgements. Again, we found a significant difference between initial and final similarity judgements in the experimental group (see Figure 2; $t(29) = -2.69, p = .011$), with similarity ratings for the final comparison ($M = 9.77, SD = 2.35$) significantly lower than the initial comparison ($M = 11.4, SD = 2.35$).

Figure 2: Initial vs. Final Similarity Scores



These findings confirmed our prediction that repeated comparisons involving the same relational schema across different contexts would lead to lower similarity ratings on the final comparison. This finding is especially striking for two reasons. First, as noted previously, the first and final sentence pairs are virtually identical. This means participants in the experimental condition routinely rated the final pair significantly less similar despite having seen, and rated highly, essentially the same pair only minutes prior. Further, participants in the control condition showed no difference on average in their similarity judgements between the first and final comparisons. This suggests that under certain conditions, the representational structure of relational concepts — which underlie the similarity judgements observed in Experiments 1 and 2 — can be altered.

General Discussion

Relational concepts (represented as abstract schemas) play a key role in our ability to reason by analogy. When we recognize that an unfamiliar problem or situation is structurally similar to one we've faced in before, we're able to transfer insights from previous experience, allowing what we've learned in the past to help guide us in the present. Most models of analogical reasoning require that the relational structures of two scenarios be mapped before this kind of knowledge transfer can occur; however, it is rarely the case that two different situations or objects just happen to share identical relational structures. To address this, researchers hypothesized that mismatching relational structures could undergo rerepresentation, resulting in changes that would promote analogical coherence and satisfy the constraints of models like Structure Mapping Theory (Kurtz, 2005; Oberholzer et al., 2018).

This study investigates whether relational schemas can undergo change when repeatedly deployed across a wide range of different contexts. We were motivated, in part, by prior findings suggesting that relational concepts are more semantically flexible than entity concepts (Asmuth & Gentner, 2016), and that once a representation has been altered (i.e., rerepresented) for the sake of comparison, subsequent attempts to alter it further prove difficult (Day & Asmuth, 2017). We took as our starting point two possible hypotheses: Either a) the representational structure of relational concepts is fixed (and their observed mutability is entirely due to the rerepresentation of the concepts that act like arguments, occupying the open slots in their structure); or b) the representational structure of relational concepts can — under the right conditions — undergo change.

In Experiment 1, we successfully replicated Day and Asmuth's (2017) findings, and provided further evidence for rerepresentation during instances of structured comparison. While both findings provide evidence that rerepresentation does, in fact, occur, it also casts doubt on the notion that rerepresentation alone is the panacea it has often been made out to be. In short, once a representation has been altered to fit a given comparison, it becomes, in a sense, stubborn. It could be, as Day and Asmuth suggest, that unless strongly motivated to do so, individuals simply do not exert the additional processing required to rerepresent the same situation more than once. Or perhaps rerepresented structures are like overworked dough, stiff and resistant to further change. Whatever the case may be, these findings are notable given the importance of rerepresentation to structure mapping theory, which has been hypothesized to play a central role in analogical reasoning, similarity judgements, and classification, among other cognitive processes (Markman & Gentner, 2001).

In Experiment 2, we provide evidence that repeated activations of the same relational schema across disparate semantic domains can induce change in that schema. Unlike in Experiment 1, we do not present participants with the same *Standard* sentence in every comparison. Instead, Experiment 2 only retains the second sentence of each comparison,

shifting it to the first position for the subsequent pair. Given that both Experiments involve repeated comparisons using the same relation, what might explain the observed difference in results? One possibility is that the presence of the *Standard* sentence in Experiment 1 acted as a cue to the *Base Comparison*, further reinforcing the effect observed by Day and Asmuth (2017). This may explain why, in the absence of such a reminder, repeated comparisons across semantically disparate domains are more likely to induce schema drift.

It's also worth noting the difference between the dips in similarity ratings between Experiments 1 and 2. In Experiment 1, similarity ratings in the *Different* condition for test comparisons 1-3 drop below the halfway point on the scale, suggesting that participants judged the compared situations to be at least somewhat dissimilar. However, in Experiment 2, while the drop is significant, the mean similarity ratings in the final judgement of the experimental condition remain above the halfway mark. This might suggest that changes in relational concept structure of the kind we observed are subtle: enough to allow for the concept to accommodate a wide range of applications, but not so much as to significantly alter its character. It also hints at a potentially novel method for studying the structure of representations. The representations used in Experiment 1 (e.g., "As the zoo keeper was busy cleaning its habitat, the Siberian tiger was able to escape its open cage into the nearby city") can be described as relatively *rich* or *dense*; they contain numerous entities and objects, which in turn possess a variety of features (Gentner & Kurtz, 2005). By comparison, the structure of a relational concept like AVOID is relatively *sparse*, possessing relatively few such properties or connections. If denser representations undergo more dramatic rerepresentations by virtue of having more features to map or assimilate, the magnitude of the drop in similarity ratings resulting from these changes could serve as an indirect measure of representational complexity.

Taken together, we believe the findings from Experiments 1 and 2 militate in favor of the hypothesis that relational concepts can "drift" under the right circumstances. In particular, it appears the relational concepts may drift when used in repeated comparisons across different contexts over a relatively brief span of time. While this kind of scenario is relatively uncommon in everyday experience (we are rarely accosted on the street by strangers demanding we entertain a variety of abstractly related scenarios), it seems plausible that our relational schemas do gradually change over time to accommodate new, structurally similar experiences.

Acknowledgements

This work was partly motivated by discussions with the UC San Diego graduate students in PSYC 247. We also thank the members of the Early Learning and Cognition Lab who provided critical feedback throughout the development of the project. Finally, we are grateful to Samuel Day and Jennifer Asmuth for providing details about the work that inspired the

current experiments. This work was partly funded by NSF grant SBE-2047581 to C. Walker.

References

- Asmuth, J., & Gentner, D. (2017). Relational categories are more mutable than entity categories. *Quarterly Journal of Experimental Psychology*, 70(10), 2007-2025.
- Blanchette, I., & Dunbar, K. (2002). Representational change and analogy: How analogical inferences alter target representations. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 28(4), 672.
- Boroditsky, L. (2007). Comparison and the development of knowledge. *Cognition*, 102(1), 118-128.
- Day, S. B., & Asmuth, J. (2017). Re-representation in comparison and similarity. In *CogSci*.
- Gentner, D. (1983). Structure-mapping: A theoretical framework for analogy. *Cognitive science*, 7(2), 155-170.
- Gentner, D., & Hoyos, C. (2017). Analogy and abstraction. *Topics in cognitive science*, 9(3), 672-693.
- Gentner, D., & Kurtz, K. J. (2005). Relational categories. In *Categorization inside and outside the laboratory: Essays in honor of Douglas L. Medin*. (pp. 151-175). American Psychological Association.
- Gick, M. L., & Holyoak, K. J. (1980). Analogical problem solving. *Cognitive psychology*, 12(3), 306-355.
- Goldstone, R. L., Lippa, Y., & Shiffrin, R. M. (2001). Altering object representations through category learning. *Cognition*, 78(1), 27-43.
- Holyoak, K. J., & Koh, K. (1987). Surface and structural similarity in analogical transfer. *Memory & cognition*, 15(4), 332-340.
- Holyoak, K. J., & Thagard, P. (1995). *Mental leaps: Analogy in creative thought*. MIT press.
- Kurtz, K. J. (2005). Re-representation in comparison: Building an empirical case. *Journal of Experimental & Theoretical Artificial Intelligence*, 17(4), 447-459.
- Oberholzer, N., Trench, M., Kurtz, K. J., & Minervino, R. A. (2018). Analogies without commonalities? Evidence of re-representation via relational category activation. *Frontiers in psychology*, 9, 2441.
- Vendetti, M. S., Wu, A., Rowshanshad, E., Knowlton, B. J., & Holyoak, K. J. (2014). When reasoning modifies memory: Schematic assimilation triggered by analogical mapping. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 40(4), 1172.