

SEQL: Category learning as progressive abstraction using structure mapping

Sven E. Kuehne (skuehne@ils.nwu.edu)

Kenneth D. Forbus (forbus@ils.nwu.edu)

Department of Computer Science, Northwestern University
1890 Maple Avenue, Evanston, IL 60201 USA

Dedre Gentner (gentner@nwu.edu)

Bryan Quinn (bquinn@nwu.edu)

Department of Psychology, Northwestern University
2029 Sheridan Rd., Evanston, IL 60201 USA

Abstract

The nature of categories and their acquisition is one of the central open questions in Cognitive Science. We suggest that categories are represented via structured descriptions and formed by a process of progressive abstraction, through successive comparison with incoming exemplars. This paper describes how SEQL (Skorstad, Gentner, & Medin, 1988), a computer model for category learning, which is based on SME (Falkenhainer *et al* 1986, 1989; Forbus *et al* 1994) can be used to simulate a recent categorization experiment (Ramscar & Pain, 1996), using a new algorithm, *Generalization and Exemplar Learning* (GEL). We demonstrate that SEQL produces behavior consistent with human subjects.

Introduction

Similarity is often viewed as central to categorization. For instance, prototype theories of categorization posit that categorization decisions are made on the basis of the similarity of an entity to the prototypical member of that category (Rosch 1975). However, similarity-based accounts have been criticized recently on the grounds that they fail to capture the role of theory-based knowledge in category formation (Murphy & Medin, 1985; Murphy & Allopenna 1994; Wisniewski, 1995). Indeed, subjective similarity can sometimes be disassociated from the probability of category membership (Keil, 1989; Rips, 1989; Gelman & Wellman, 1991). For example, bats share many perceptual and behavioral characteristics with birds (e.g., they both fly). Yet despite these similarities, we do not categorize bats as birds. Rather, bats are categorized as mammals because of nonobvious but theoretically central properties such as giving birth to live young.

This research attempts to bridge the gap between similarity and categorization. We suggest that many of the problems with similarity-based accounts stem from viewing similarity in terms of featural commonalities (see Goldstone, 1994). We agree that theory-based knowledge is important to categorization, but we do not see that as inimical to a role for similarity, provided similarity is modeled appropriately. Recent studies have provided evidence that the

process of structural alignment (Markman & Gentner, 1993; Medin, Goldstone & Gentner, 1993) is a part of the category learning process (Lassaline & Murphy, 1998; Ramscar & Pain, 1996). These studies have shown that similarity, when understood as the result of an alignment process, is capable of incorporating theory-based knowledge and higher order relations into the process of category learning (Gentner & Medina, 1998).

This paper describes how a computer model for category learning, SEQL (Skorstad, Gentner & Medin, 1988), which uses structural alignment to incrementally compute abstractions, can be used to simulate recent results in category learning. We begin by outlining structure-mapping theory and evidence that structural alignment plays an important role in categorization. We then describe a *progressive abstraction* model of category learning, and the GEL (*Generalization and Exemplar Learning*) algorithm (see also Blok & Gentner, 2000) which implements it in SEQL. Finally, we show how the results of a recent study exploring the role of structural alignment in category learning (Ramscar & Pain, 1996; Darrington, Lingstadt, & Ramscar, 1998) can be simulated using SEQL.

Structural alignment and category-based inference

Structure-mapping theory (Gentner, 1983, 1989) provides an account of analogy and similarity based on comparisons of structured representations. According to structure-mapping theory, structural alignment takes as input two structured representations (*base* and *target*) and produces as output a set of mappings. Each mapping consists of a set of *correspondences* that align items in the base with items in the target and a set of *candidate inferences*, which are surmises about the target made on the basis of the base representation plus the correspondences. The constraints on the correspondences include *structural consistency*, i.e., that each item in the base maps to at most one item in the target and vice-versa (the *1:1 constraint*) and that if a correspondence between two statements is included in a mapping, then so must correspondences between its arguments (the *parallel connectivity* constraint). Which mapping is chosen

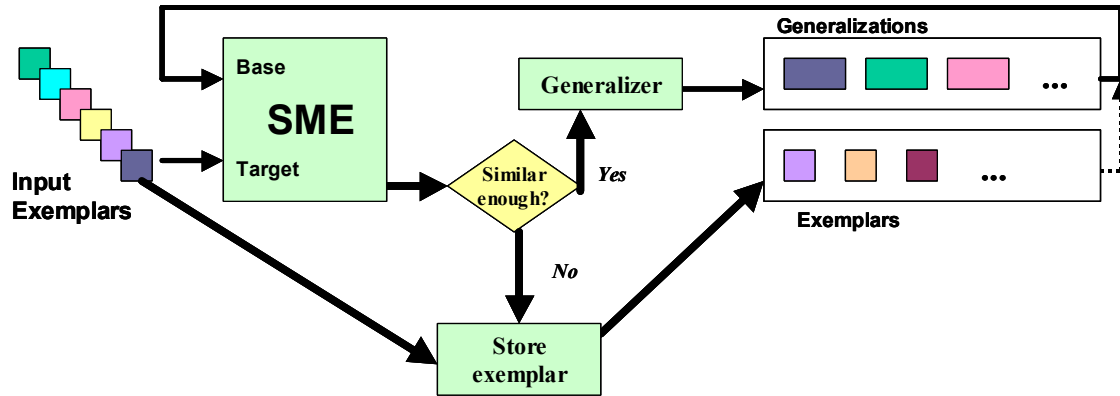


Figure 1: The GEL strategy as implemented in SEQL

is governed by the *systematicity* constraint: Preference is given to mappings that match systems of relations in the base and target. Each of these theoretical constraints is motivated by the role analogy plays in cognitive processing. The 1:1 and parallel connectivity constraints ensure that the candidate inferences of a mapping are well-defined. The systematicity constraint reflects a (tacit) preference for inferential power in analogical arguments.

The structure-mapping view offers new explanations relevant to categorization. For example, there is evidence that humans prefer matches that share common higher-order relations over those that share only object-attributes (Gentner, Rattermann & Forbus, 1993) and seek systems of commonalities rather than over isolated commonalities when carrying out comparisons (Clement & Gentner, 1991). These tacit preferences would contribute to theory-based biases in categorization. Ramsar and Pain (1996) tested this possibility by asking subjects to sort the Gentner, Rattermann & Forbus stories (which varied systematically in their similarity relations) into categories of their own choosing. As discussed below, subjects classified the stories by common causal structure rather than by common object features.

Progressive abstraction via GEL

We propose a model of category learning by progressive abstraction. On this account, the representation of a category is a structured description generated by successive comparisons with incoming exemplars. The idea is that a set of *generalizations* and a set of *exemplars* are maintained. Each generalization is regarded as a separate category. The exemplars are items that are too dissimilar to the existing categories to be assimilated into any of them. As each new exemplar E arrives, it is first compared with the existing generalization(s) G_i , using structural alignment. If the match is good enough (as compared to a preset threshold) then E is assimilated into G_i by replacing G_i with the structural overlap. If no generalization is sufficiently similar, then E is compared with each stored exemplar E_i . If one of those matches is over threshold, then their overlap becomes a new generalization, and E_i is removed from the stored exemplars. Otherwise, E is added to the set of stored exem-

plars. This, intuitively, is the *Generalization and Exemplar Learning* algorithm.

Thus GEL works by constantly folding new exemplars into an evolving category representation. It is intended to capture the incremental, conservative nature of concept learning. Early learning is typically highly contextualized, with many specific, concrete details (Gentner, 1989; Medin & Ross, 1989). With experience successive comparisons among exemplars weather away the concrete details and eventually reveal the abstractions beneath.

SEQL (Skorstad, Gentner, & Medin, 1988) provides a platform for implementing this algorithm. SEQL was designed as a cognitive simulation toolkit that could be used to implement a variety of schema abstraction strategies, ranging from pure exemplar to pure prototype strategies, as well as mixed strategies. The 1988 paper demonstrated that a model involving abstraction based on successive comparisons could model effects due to presentation order in concept learning. Figure 1 shows the architecture of SEQL using the GEL strategy.

SEQL uses a simulation of structural alignment as a component both for simulating category-based inference and for the comparison process used in category learning. The Structure-Mapping Engine (SME) (Falkenhainer *et al.* 1986, 1989; Forbus *et al.* 1994) is a cognitive simulation of analogical matching. Given base and target descriptions, SME finds globally consistent interpretations via a local-to-global match process. The base description is either an exemplar from the list of stored exemplars (or the first exemplar of the input sequence) or a stored generalized description. The target description is always a new exemplar from the input sequence. SME begins by proposing correspondences, called *match hypotheses*, in parallel between statements in the base and target. Then, SME filters out structurally inconsistent match hypotheses. Mutually consistent collections of match hypotheses are gathered into global mappings using a greedy merge algorithm. An evaluation procedure based on the systematicity principle is used to compute the *structural evaluation* for each mapping.

We will need a few conventions to describe the algorithm underlying SEQL and GEL. The *overlap* found when comparing two descriptions is defined as the set of statements in the base which have correspondences in the best mapping

SME produces for those two descriptions. To determine if an exemplar and a concept description are “sufficiently similar”, we need to use the notion of structural evaluation to provide a numerical summary of similarity that is independent of the size of the descriptions. Previous cognitive simulation studies using SME (c.f. Falkenhainer, Forbus, & Gentner, 1989) have demonstrated the structural evaluation of the strongest mapping to be correlated with such judgments. Consequently, let $SE(b, t)$ be the structural evaluation score of the best mapping SME finds for the comparison between descriptions b and t . Our numerical summary of similarity is defined as follows:

$$NSIM(d, e) = SE(d, e) / SE(d, d)$$

that is, the structural evaluation score of the concept description d compared with e , normalized by the score of the concept description compared to itself. This limits $NSIM$ to being between 0 and 1, which simplifies the comparison between the structural evaluation score and a threshold value. We use a threshold T to decide whether or not a match between an exemplar and a concept description is considered “good”, i.e. sufficiently similar:

$$NSIM(d, e) \geq T$$

The threshold T determines how conservative the system will be: If $T = 1.0$, then no abstraction will occur, since only perfect matches would be grouped. If $T = 0.0$, then any two descriptions could match, leading quickly to an empty description as the concept representation.

The Generalization Algorithm

If the base and target are judged sufficiently similar by the process just described, a generalization process will transform the most promising¹ mapping between the two descriptions into a new generalized description. The generalization algorithm creates a new description by copying the predicate structure in the overlap of the mapping, substituting generic labels for particular entity names but preserving the predicate structure, including the specific predicate names used².

The GEL algorithm

The GEL strategy is based on two observations: (1) Natural sequences of input stimuli may consist of examples from several different categories. Furthermore, some natural concepts involve disjunctive descriptions. This suggests maintaining multiple generalizations. In the case where all input stimuli are assumed to be from the same category (e.g., supervised learning), the set of generalizations serves as a disjunctive concept definition. Otherwise, each generalization represents a distinct category. (2) Exemplars that are initially singular may only be the first of many similar ex-

emplars to come. This suggests maintaining a list of *singular exemplars* which currently do not fit any existing generalization well. Singular exemplars are transformed into generalizations if a later exemplar matches them best.

Generalization and Exemplar Learning (GEL): A concept description consists of G , a list of generalizations, and S , a set of exemplars. Elements of S are structured descriptions, and elements of G are tuples of a structured description and an integer N indicating the number of exemplars assimilated into that generalization. If nothing is known about a concept, G and S are initially empty. When a new exemplar E_i arrives, the following occurs:

Procedure GEL (E_i, G, S, T)

1. For each $\langle G_i, N_i \rangle \in G$,
 - a. If $NSIM(G_i, E_i) > T$ then $\langle G_i, N_i \rangle \leftarrow \langle \text{Generalization}(G_i, E_i), N_i + 1 \rangle$; go to 4.
2. For each $S_i \in S$,
 - a. If $NSIM(S_i, E_i) > T$ then let $G_{\text{new}} = \text{Generalization}(S_i, E_i)$;
 $G \leftarrow G \cup \{\langle G_{\text{new}}, 2 \rangle\}$; $S \leftarrow S - \{S_i\}$; go to 4.
3. $S \leftarrow S \cup \{E_i\}$.
4. Sort G by N_i 's so that it will be searched from maximum to minimum number of exemplars involved.

The sort step improves the likelihood that good generalizations will be found, if they exist, by preferentially building up strong existing generalizations. The assimilation of singular examples in step 2 serves a similar purpose.

A Case Study

Ramscar and Pain (1996) investigated the role of structural similarity in categorization by having subjects sort stories taken from the “Karla the Hawk” similarity experiments (Gentner, Rattermann, & Forbus, 1993). Each set had six stories: a base plus five variants. All variants except OO have first-order relational commonalities with the base (i.e., common events: e.g., *chase/pursue*). In addition, some variants share object attributes (i.e., similar entities: e.g., *hawk/eagle*). In addition, some share higher-order relations such as common causal structure. The following relationships hold:

B, the *base* story. The properties of the other stories in the set are defined with respect to this story.

LS, the *literally similar* story, shares all levels with the base: object attributes, first-order events and higher-order relational structure..

TA, the *analogy* story, shares first-order and higher-order relations with the base, but not object attributes.

MA, the *mere-appearance* story, shares first-order relations and object attributes with the base, but not higher-order relational structure.

FA, the *false analogy* story, shares only first-order relations with the base. (It is an analog of the MA story)

¹ Generally SME produces more than one mapping for a comparison. SEQL only considers the mapping with the highest structural evaluation score.

² In the case of matches involving non-identical functions, the function in the base description is used.

OO, the *objects-only* story, shares only object attributes with the base.

Ramscar and Pain gave subjects each 10 sets of six stories, one set at a time. Stories within a set were presented in randomized order, and the order of sets of stories was also randomized. For each set, subject read every story in the set to familiarize themselves with the stories, and then were asked to “group the stories into the categories that seem most natural and appropriate to you. These groups can range from putting every member of the story set into the same group, to putting each story into a group on its own.”

Ramscar & Pain (1996) identified ten types of groupings, or which only five exceeded 3% responding. The two most common groupings were based on a shared relational structure, with the base either included (B-LS-TA, FA-MA, OO -- 79.5%) or separately classified (B, LS-TA, FA-MA, OO -- 8%). The next two groupings were based on shared first-order relations (B-LS-TA-FA-MA, OO -- 4%) or on object similarities (5%) -- either B-LS-MA-OO, FA-TA or B-OO, LS-MA, FA-TA.

Overall, subjects showed a strong preference to group stories that shared systematic relational structure, consistent with the predictions of structure-mapping theory; 87.5% of the groupings were based on common first-order and higher-order relations.

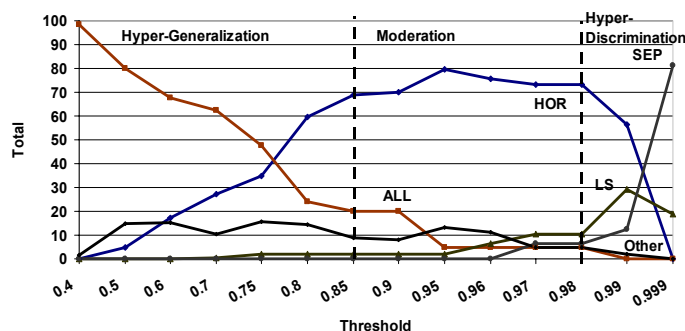


Figure 2: Simulation Results. Threshold values (x axis) are expanded towards the right to focus on key behavior

Simulation

Can GEL and SEQL capture this categorization behavior? To test this, we used representations of eight of the “Karla the Hawk” story sets that had been generated for an earlier study (Forbus, Gentner, & Law, 1995). These representations were generated before the Ramscar and Pain experiment, so that the choices made in creating the representations were made independently. Each set consists of five stories, a base (B), a literal similarity (LS), a true analogy (TA), a false analogy (FA), and a mere appearance (MA) story. (The object-only (OO) stories had not been represented in the previous experiment.) For each of the eight sets of stories, we ran SEQL on every possible sequences (i.e., $5! = 120$). We identified five types of groupings:

Higher-order Relational Structure (HOR): All groupings based on shared structural relations. This includes Ramscar and Pain’s top two groupings (B-LS-

TA, MA-FA and B, LS-TA, MA-FA), as well as groupings in which MA and FA are separated, e.g., B-LS-TA, FA, MA and B, LS-TA, MA, FA.

Literal Similarity and Base (LS) (Conservative Similarity): All groupings in which only the literal similarity story is grouped with the base story and the true analogy story is kept separate, i.e. B-LS, TA, MA-FA and B-LS, TA, MA, FA.

ALL: The grouping that includes all stories, i.e. B-LS-TA-MA-FA.

SEP: No groupings at all; each story separate.

Other: All other unclassifiable groupings.

The simulation results are shown in Figure 2. The x axis shows the threshold value and the y axis the percentage of groupings found at that value. Each data point represents the average of 960 simulation runs (120 sequences * 8 sets). For the whole graph (14 threshold values) we ran a total of 13,440 sequences, which required more than 60,000 SME matches.

The first thing to notice is the range of patterns that results as the threshold is varied. As we move from low to high thresholds the pattern changes from indiscriminate lumping of all stories together on the left (*hyper-generalization*) to finicky separation of each story into each own category at the extreme right (*hyper-discrimination*). Below $T = 0.85$, the majority of groupings produced are of the ALL type (hyper-generalization), and when $T > 0.99$, the stories are grouped into individual categories.

In the threshold range between 0.85 to 0.98 we see a pattern that matches well with human results: There is a strong preference for groupings based on common higher-order relational structure. At 0.85 the number of groupings based on higher-order relations reaches about 70 percent. At the same time, the number of unclassifiable Other groupings drops below 10 per cent. Higher-order groupings remain dominant until the threshold becomes extremely high. As we move to the upper bound of the humanlike range, there is first a rise in the more conservative LS groupings.³ Finally, as the threshold approaches 1.0, almost no stories will be grouped together, leading to a total separation – each story will form its own category

A comparison of our results with the study by Ramscar and Pain would suggest that their data correspond to a threshold value around 0.95 in our simulation. The only major difference is that their human subjects did not produce as many ALL groupings (comparable to Ramscar and Pain’s Type 3 groupings) as SEQL does it in our experiment.

Order Effects

According to the progressive alignment model, there should be effects of the order in which the items are received. Assuming a moderate threshold, SEQL’s new

³ The rising number of LS groupings for threshold values beyond 0.95 stems entirely from differentiating previous HOR groupings. (The groupings for each threshold value are mutually exclusive.)

categories will initially be highly conservative -- closely identified with the first few instances that comprise them. Subsequent comparisons will lead to further abstractions and 'wear away' specific features. Thus the generalization will gradually come to match a wider range of new items. This means that the overlap among the first few exemplars determines the initial encoding of the category. Thus we should be able to manipulate the likelihood that SEQL (or human subjects) arrives at the HOR grouping by varying the initial exemplars.

In further studies we tested the progressive alignment prediction (Quinn et al., in preparation). We ran a simulation contrasting *structure-promoting* orders with *surface-promoting* orders. A *structure-promoting* order is one in which stories that overlap in relational structure: e.g., (B, LS, TA) occur first, and a *surface-promoting* order is one in which stories with common attribute matches (LS, MA, FA) occur first. The results of the simulation showed that on structure-promoting orders, SEQL produced structure-based groupings 90% of the time. In contrast, for surface-promoting orders, SEQL produced structure-based groupings only 70% of the time.

Thus early presentation of structural or surface commonalities will bias the kinds of categories formed. This result is consistent with the prior findings of order effects in human category learning (Elio & Anderson 1984; Wattenmaker 1993; Medin & Bettger 1994). However, to achieve a closer test, Quinn, Kuehne, Gentner and Forbus (in preparation) ran human subjects on a serial categorization study contrasting structure-promoting and surface-promoting orders.⁴ Subjects received the stories one at a time on a workstation, and were allowed to see each story only once. After reading the five stories in each set, participants were asked to group them into categories that seemed "most natural and appropriate." Half the subjects received the stories in *structure-promoting orderings* -- e.g., {B, LS, TA, FA, MA}. The other half received them in *surface-promoting* orders -- e.g., {MA, LS, TA, FA, B}. (Four variants were used in each condition)

People produced more structural grouping in the structure-promoting orders (88.9%) than in the surface-promoting orders (71.4). The corresponding proportions for SEQL were 85.7 and 71.4. Interestingly, when subjects were required to write out justifications for their groupings, their percentages of structural groupings increased (to 92.9 and 79.8 in structure-promoting and surface-promoting orders, respectively). Although these results bear out our predictions concerning order, there was also a discrepancy: unlike human subjects, SEQL did not produce surface groupings when given the surface-promoting orders; instead it produced various 'other' groupings. This suggests that the object representations we used in our simulation were not as rich as human representations. Forbus and Gentner (1989) *specificity conjecture* suggests that more low-order informa-

tion would lead to SEQL more closely modeling the human data.

Discussion

A growing number of studies suggest that structural alignment plays a central role in categorization. This paper shows that a simulation based on progressive abstraction via successive comparisons can model some of the recent human results in this area. The Generalization and Exemplar Learning strategy for SEQL presented here provides a computational model for progressive abstraction. In our experiments we have shown that SEQL using GEL can produce results that are consistent with human categorization behavior. By using SME, a simulation of structure-mapping theory that already has considerable support in terms of psychological plausibility (c.f. Gentner & Markman, 1997), SEQL extends the computational account that mirrors the psychological evidence that implicates structural alignment in categorization.

SEQL's behavior in the 'moderate threshold' range fits fairly well with that of Ramsar and Pain's human subjects. What about its behavior when given more extreme thresholds? We speculate that these kinds of threshold effects occur in human similarity as well. For example, a person identifying her own door key sets a higher match threshold than a person who is simply asked to identify an object as a key. In the other direction, a person who simply needed a small metal object (say, to weight down a transparency sheet) would apply a more liberal match criterion -- a key or a coin would be equally fitting.

Although SEQL provides an interesting model for some aspects of categorization, it has several important limitations. There is no recovery or limiting mechanism that prevents successive generalizations from becoming so abstract that they have no inferential power. The use of a high similarity threshold makes such vacuous abstractions less likely, but what is needed is a better understanding of when and how such outcomes are (or are not) avoided in learning. SEQL also does not model the effects of learning contrasting categories. Nevertheless, SEQL does suggest even in its current state that some of the mysteries of categorization may be solved by further explorations of progressive abstraction via structural alignment.

Acknowledgements

We thank Tom Mostek for technical assistance and Michael Ramsar for providing stimuli, and both of them, as well as Ken Kurtz and Jeff Loewenstein, for valuable discussions. This research was supported by the Cognitive Science and Artificial Intelligence Divisions of the Office of Naval Research.

References

Blok, S. V., & Gentner, D. (2000). Reasoning from shared structure. *Proceedings of the 22nd Meeting of the Cognitive Science Society*.

⁴ Seven of the Karla the Hawk story sets were used (the same sets used by Ramsar and Pain (1996)). As in the above simulations, each story set contained a base (B) story, a literal similarity (LS) story, a true analogy (TA) story, a mere appearance (MA) story, and a false analogy (FA).

- Clement, C. A., & Gentner, D. (1991). Systematicity as a selection constraint in analogical mapping. *Cognitive Science*, 15, 89-132.
- Darrington, S., Lingstadt, T., & Ramscar, M. (1998). Analogy as a subprocess of categorization. *Proceedings of the 20th Annual Conference of the Cognitive Science Society* (pp. 279-284). Erlbaum.
- Elio, R. & Anderson, J.R. (1984) The effects of information order and learning mode on schema abstraction. *Memory & Cognition* 12(1), 20-30.
- Falkenhainer, B., Forbus, K. D., & Gentner, D. (1986). The Structure-mapping Engine, *Proceedings of AAAI-86*, Philadelphia, PA.
- Falkenhainer, B., Forbus, K. D., & Gentner, D. (1989). The Structure-mapping Engine: An algorithm and examples. *Artificial Intelligence*, 41: 1-63
- Forbus, K. D., Ferguson, R. W., & Gentner, D. (1994). Incremental Structure Mapping, *Proceedings of the 16th Annual Conference of the Cognitive Science Society*, Erlbaum.
- Forbus, K. D., & Gentner, D. (1989). Structural evaluation of analogies: What counts? *Proceedings of the Eleventh Annual Conference of the Cognitive Science Society*, 34, 341-348.
- Forbus, K. D., Gentner, D., Everett, J. O., & Wu, M. (1997). Towards a computational model of evaluating and using analogical inferences, *Proceedings of the 19th Annual Conference of the Cognitive Science Society* (pp. 229-234). Erlbaum.
- Forbus, K. D., Gentner, D. & Law, K. (1995). MAC/FAC: a model of similarity based retrieval. *Cognitive Science*, 19(2): 141-205.
- Gelman, S.A., & Wellman, H.M. (1991). Insides and essences: Early understandings of the non-obvious. *Cognition*, 38, 213-244.
- Gentner, D. (1983) Structure-mapping: A theoretical framework for analogy. *Cognitive Science*, 7: 155-170
- Gentner, D. (1989) The mechanisms of analogical learning. In S. Vosniadou & A. Ortony (Eds.), *Similarity and analogical reasoning* (pp 199-241). London: Cambridge University Press. (Reprinted in *Knowledge Acquisition and Learning*, 1993, 673-694.)
- Gentner, D., & Markman, A. B. (1997). Structure mapping in analogy and similarity. *American Psychologist*, 52, 45-56.
- Gentner, D., & Medina, J. (1998). Similarity and the development of rules. *Cognition*, 65, 263-297.
- Gentner, D., Rattermann, M. J. & Forbus, K. (1993). The roles of similarity in transfer. *Cognitive Psychology*, 25: 524-575.
- Goldstone, R. L. (1994). The role of similarity in categorization: Providing a groundwork. *Cognition* 52(2), 125-157.
- Goldstone, R. L., Medin, D. L., & Gentner, D. (1991). Relational similarity and the non-independence of features in similarity judgments. *Cognitive Psychology*, 23, 22-264.
- Keil, F.C. (1989). *Concepts, kinds, and cognitive development*. Cambridge, MA: MIT Press.
- Kotovskiy, L., & Gentner, D. (1996). Comparison and categorization in the development of relational similarity. *Child Development*, 67, 2797-
- Lassaline, M.E., & Murphy, G.L. (1998) Alignment and category learning. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 24 (1), 144-160.
- Markman, A. B., & Gentner, D. (1993). Structural alignment during similarity comparisons. *Cognitive Psychology*, 25, 431-467.
- Medin, D.L. & Bettger, J.G. (1994). Presentation order and recognition of categorically related examples. *Psychonomic Bulletin & Review*, 1(2), 250-254.
- Medin, D. L., Goldstone, R., and Gentner, D. (1993). Respects for similarity. *Psychological Review*, 100(2), 254-278.
- Murphy, G.L., & Medin, D.L. (1985). The role of theories in conceptual coherence. *Psychological Review*, 92, 289-315.
- Murphy, G.L. & Allopenna, P.D. (1994) The locus of knowledge effects in concept learning. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 20, 904-919.
- Ramscar, M., & Pain, H. (1996). Can a real distinction be made between cognitive theories of analogy and categorization?, *Proceedings of the 18th Annual Conference of the Cognitive Science Society* (pp. 346-351), Erlbaum
- Rips, L.J. (1989) Similarity, typicality, and categorization. In S. Vosniadou & A Ortony (Eds.), *Similarity and analogical reasoning* (pp. 21-59). New York: Cambridge University Press.
- Rosch, E. (1975). Cognitive representations of semantic categories. *Journal of Experimental Psychology: General*, 104, 192-233.
- Skorstad, J., Gentner, D., & Medin, D. (1988). Abstraction processes during concept learning: a structural view. *Proceedings of the 10th Annual Conference of the Cognitive Science Society*. Erlbaum.
- Wisniewski, E.J. (1995) Prior knowledge and functionally relevant features in concept learning. *Journal of Experimental Psychology: Learning Memory and Cognition*, 21, 449-468.