

The Structure-Mapping Engine: A Multidecade Interaction Between Psychology and Artificial Intelligence

Dedre Gentner¹  and Kenneth Forbus² 

¹Department of Psychology, Northwestern University, and ²Department of Computer Science, Northwestern University

Abstract

This article describes the structure-mapping engine (SME) and its relation to psychological theory and research. SME was created in 1986 as a simulation of structure-mapping theory (SMT) and is still in use, both on its own and as part of larger scale simulations such as CogSketch and Companion that capture analogy's roles in other cognitive processing. Over the 4 decades since artificial intelligence (AI) first appeared, there has been continual interaction between AI research and human research. We begin by briefly reviewing SMT and the basic construction of SME. After comparing SME with other simulations, we then describe some specific contributions of SME to our understanding of human analogical processing. We close by proposing that these psychological models can become a new technology for AI.

Keywords

analogy, structure mapping, similarity, cognitive simulation, symbolic AI

Analogical ability—the ability to recognize and reason about common relational patterns across different contexts—is a core mechanism in human cognition. It is central in learning and discovery and plays a large role in everyday reasoning. For example, when Ruth Bader Ginsburg argued against the U.S. Supreme Court's decision in July 2013 to invalidate parts of the Voting Rights Act, she used the following analogy: “Throwing out [parts of the Voting Rights Act] when it has worked and is continuing to work to stop discriminatory changes is like throwing away your umbrella in a rainstorm because you are not getting wet.” (Katz, 2015). Although this analogy compares two vastly different situations, it shares a common pattern of relations, leading to an inference: Because it would be a mistake to throw away your umbrella in a rainstorm (because you are not getting wet), it is likewise a mistake to discard key parts of the Voting Rights Act (because discrimination seems to be under control).

Carrying out analogies such as this one is a sophisticated ability. At the concrete level, throwing away an umbrella has nothing in common with removing part of a congressional act. Yet human adults can readily

match these two situations and perceive common relational patterns. How does this happen? This article describes a psychological theory (structure mapping) and a computational model (the structure-mapping engine, or SME) that aim to capture the process of analogical thinking.

Structure-Mapping Theory

According to structure-mapping theory (Gentner, 1983, 2010), understanding an analogy involves finding a matching relational structure between the two items being compared. More precisely, an analogical comparison involves establishing a *structural alignment* between two represented cases (the *base* and *target*). In structural alignments, like relations are aligned, and the entities in the two cases are placed in correspondence on the basis of the relational alignment. Mappings must satisfy *structural consistency*, meaning

Corresponding Author:

Dedre Gentner, Department of Psychology, Northwestern University

Email: gentner@northwestern.edu

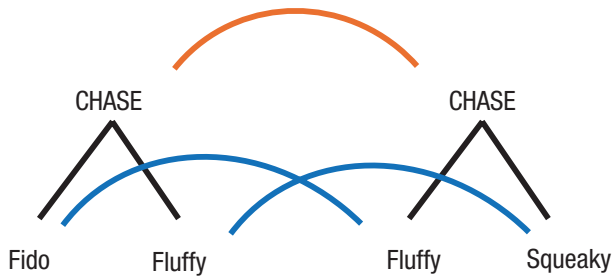


Fig. 1. Simple analogy illustrating structural alignment and the role of parallel connectivity and one-to-one correspondence in structural consistency. Having aligned the relation CHASE(Fido, Fluffy) with the relation CHASE(Fluffy, Squeaky; red arc), we must now align the arguments according to their roles in the two relations (by parallel connectivity). Thus, Fido corresponds to Fluffy, and Fluffy corresponds to Squeaky (blue arcs). Fluffy cannot also match the second instances of Fluffy because that would violate one-to-one correspondence.

that there must be a *one-to-one correspondence* between the elements of the two analogues, and *parallel connectivity*—that is, if two propositions correspond, then their arguments must correspond as well. To illustrate, consider a very simple analogy: “Fido chased Fluffy. Likewise, Fluffy chased Squeaky.”

In structural alignment, the critical step is aligning the common relation—in this case “chase” (see Fig. 1). Then, by parallel connectivity, the entities are placed into correspondence according to like roles in the relation. Thus, Fido corresponds to Fluffy, and Fluffy corresponds to Squeaky. The requirement of one-to-one correspondence means that we cannot also put Fluffy into correspondence with Fluffy, even though they match perfectly at the entity level. Psychological research has borne out the claim that people seek structural consistency in analogies (Krawczyk et al., 2004; Markman & Gentner, 1993).

In this simple analogy, we began with identical input relations. More often, the relations are semantically similar but not identical. In this case, people may arrive at a slightly more abstract relation that captures their common core (*minimal ascension*). For example, in the following analogy, one might recode both events as *pursuing*: “Fido raced around after Fluffy. Likewise, Mary repeatedly texted Fred.” As long as there is a core of identical meaning, the relations can be abstracted to match. This allows for *tiered identity*.

A further claim of structure mapping is the *systematicity principle*—a preference for interpretations in which the lower order relational matches are connected by higher order constraining relations, such as causal, mathematical, or spatial relations (see Fig. 2). Systematicity guides the selection of which commonalities enter into the mapping and which inferences to project. It

reflects a tacit preference for coherence and predictive power in comparisons.

Achieving a structural alignment leads to further processes that benefit learning. First, inferences are often projected from one case to the other. For example, in the analogy above, suppose you were told “Fido raced after Fluffy, so Fluffy hid.” Assuming you believe these two events are analogous, you would draw the inference that Squeaky might also have hidden. We refer to analogical inferences as “candidate inferences” because they are not guaranteed to be true. But if the analogy is a good one, we may expend considerable effort to test the inference. Many scientific experiments are generated in this way (Dunbar, 1993). A second outcome of structural alignment is that the common relational structure becomes more salient and may form the seed of a new abstraction. Third, *alignable differences* may spontaneously emerge from the mapping. Alignable differences are differences that play corresponding roles in the two cases (see Fig. 3). Last, relations that have been provisionally aligned may be replaced by a common abstraction, resulting in a more abstract representation of the case (Gentner, 2010).

The structure-mapping engine

SME (Forbus, Ferguson, et al., 2017) aims to simulate the processes that humans use during analogical matching. It takes structured representations as its inputs. In earlier research using SME, these representations were mostly hand-coded; however, in the past 2 decades, SME has mostly operated over representations that were computed independently by other systems, such as language, vision, or artificial intelligence (AI) reasoning systems (Forbus & Hinrichs, 2017). SME’s use of structured representations makes it a symbolic model, although hybrid systems are sometimes used to provide its inputs (e.g., Forbus et al., 2025), and versions of SME have been built that incorporate distributed representations (e.g., Forbus, Liang, & Rabkina, 2017).

SME operates in a three-stage process from local matches to global interpretations. Figures 2 and 3 show this process, using an analogy that explains heat flow (the target) by comparing it to water flow (the base; adapted from Buckley, 1979). Each case consists of knowledge represented in structured networks of concepts and relations. This analogy assumes that the reader knows that water flows from a higher pressure to a lower pressure and that the pressure is higher in the beaker than in the vial. The goal is to communicate that heat flows from a higher temperature to a lower temperature, just as water flows from a higher pressure to a lower pressure.¹

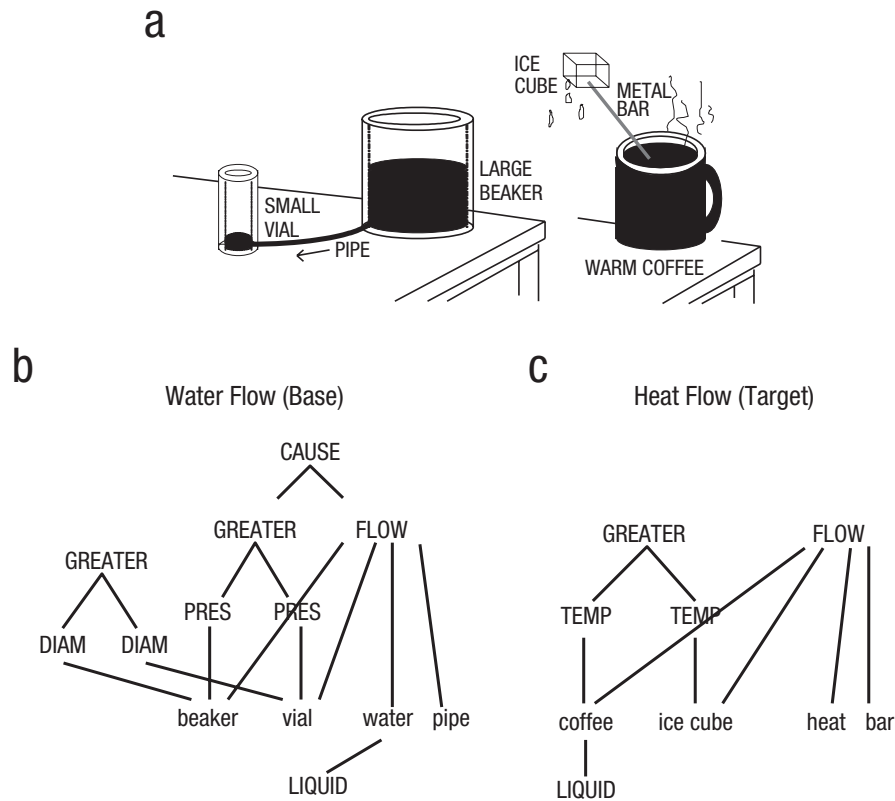


Fig. 2. Analogy between water flow and heat flow. This (a) analogy between water flow and heat flow shows (b, c) structured representations of two situations. These representations aim to capture a person's knowledge about a domain. Here we assume greater knowledge about water flow than about heat flow. The representations include entities (e.g., "beaker") and attributes (i.e., properties of entities; e.g., LIQUID(coffee)), as well as functions that denote parts or dimensions of entities, for example, PRES(x , y). These representations also include relations (i.e., predicates that connect two or more entities), for example, FLOW(water, beaker, vial, pipe), meaning that water flows from the beaker to the vial via the pipe. Higher order constraining relations serve to connect lower order relations into a system, for example, CAUSE(GREATER PRES (beaker)) and CAUSE(FLOW(water, beaker, vial, pipe)).

Stage 1 is a local, structurally blind free-for-all in which all possible identity matches between individual propositions are made. In addition, for each match, parallel connectivity introduces matches between their arguments, regardless of structural consistency (Fig. 2a). In Stage 2, the match hypotheses are combined into larger clusters (or "kernels") that must be structurally consistent (Fig. 2b). The systematicity bias is used to rate kernels, with deeper kernels preferred over shallower ones, even if they have the same number of statements. In Stage 3, structurally consistent kernels are combined into one or more final mappings (Fig. 2c). Each mapping consists of the set of correspondences between the propositions and entities in the base and target, candidate inferences that project information from the base to the target, and a score indicating SME's estimate of similarity.²

Analogical mapping is an example of a graph-matching problem, a well-known NP-complete problem for which, in general, the computational cost increases exponentially with the size of the input representations.³ Analogy researchers have dealt with this problem in two ways. One way to limit the size of the problem is to first select a structure in one analogue and then project it to the other (e.g., DORA: Doumas et al., 2008; LISA: Hummel & Holyoak, 2003; IAM: Keane, 1990). The other way is to begin with semantic alignment, as in SME and SIAM (Goldstone, 1994). SME's initial tiered identity stage creates a limited starting set of identical or near-identical matches that is then refined by applying structural consistency and systematicity. SME then uses an efficient approximate merging algorithm,⁴ enabling it to scale to realistic representations.

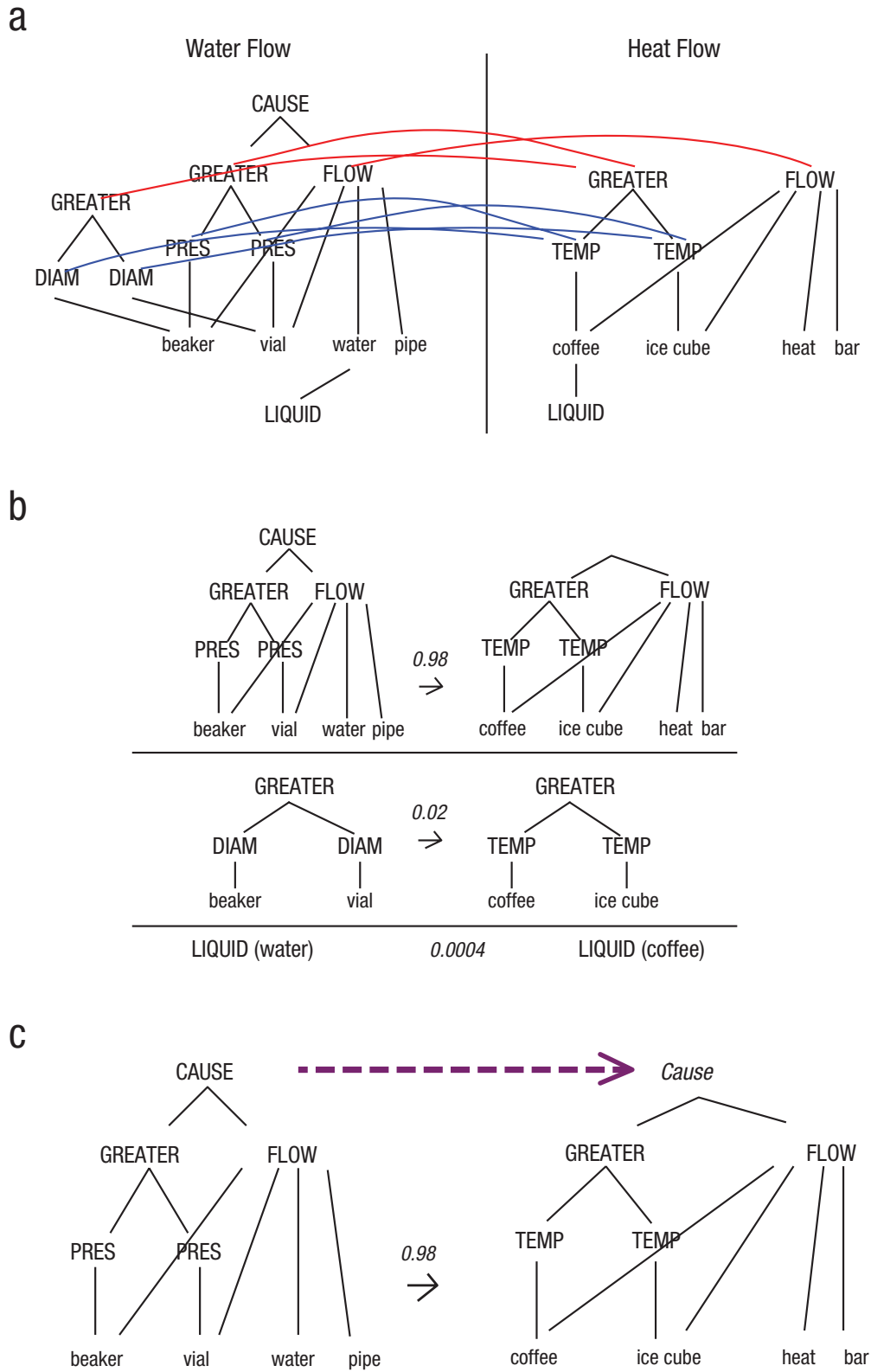


Fig. 3. (continued on next page)

Fig. 3. An example illustrating SME's three-stage process. The red and blue arcs in (a) Stage 1 represent match hypotheses based on identical predicates and further matches required by parallel connectivity, respectively. Entity correspondences are not shown for clarity. Note that the requirement of one-to-one correspondences does not hold at this stage. For example, both GREATER PRES(beaker), PRES(vial) and GREATER DIAM(beaker), DIAMETER(vial) are matched to GREATER TEMP(coffee), TEMP(ice cube). In (b) Stage 2, structurally consistent substructures (called "kernels") are extracted from the Stage 1 forest. Each kernel is given a structural evaluation score on the basis of the number of predicates and the depth of its structure. Here, the top kernel has the highest structural evaluation score. In (c) Stage 3, the final interpretation, the kernels are merged, beginning with the one with the highest structural evaluation and adding more kernels if possible while maintaining structural consistency. (In this analogy, no other kernels can be added to the winning kernel from Stage 2.) In addition, candidate inferences can be drawn by projecting predicates from the base to the target. (Candidate inferences must belong to the common system in the base.) Here, the candidate inference is that heat flow is caused by the difference in temperature (just as water flow is caused by a difference in pressure).

SME has some advantages as a model of human analogical processing. First, because it uses the same process for overall similarity as for analogy, SME can capture the effect of progressive alignment, in which carrying out close comparisons potentiates carrying out purely relational comparisons (Haryu et al. 2011; Kotovsky & Gentner, 1996). Second, SME's systematicity bias allows it to capture the finding that people choose which predicates to include in mappings on the basis of whether they belong to a shared systematic set of relations (Clement & Gentner, 1991) and that people prefer to map inferences from a more systematic case to a less systematic one rather than the reverse (Bowdle & Gentner, 1997). Last, SME's alignment-first process means that it is not necessary to choose which base structure to project in advance. Rather, the best match can emerge from the comparison process itself. This allows SME to capture analogy's role in spontaneous discovery and in early learning. For example, research shows that comparing two things can help children perceive a common relational pattern—even when that pattern was not perceived in either of the separate items (Christie & Gentner, 2010).

The claim that analogies are initially processed via symmetric alignment may seem implausible, given that many analogies and metaphors seem to be strongly directional. For example, the statement "Some salesmen are bulldozers" makes sense, whereas its reverse "Some bulldozers are salesmen" does not. This might lead one to believe that we understand "Some salesmen are bulldozers" by extracting a relational pattern from the base ("bulldozers") and projecting it to the target ("salesmen"). But this intuition is deceptive. Wolff and Gentner (2011) showed people strongly directional metaphors briefly and asked them to rate their comprehensibility. When people were required to answer quickly (by 500 ms), they rated forward and reversed metaphors as equally comprehensible

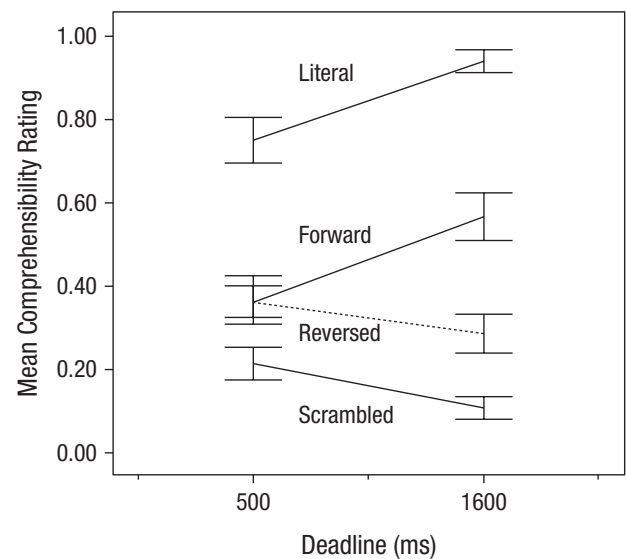


Fig. 4. Proportion of statements judged comprehensible in Wolff and Gentner's (2011) Experiment 3. When people had to respond by 500 ms, they were unable to distinguish forward from reversed comparison statements, but they rated scrambled (nonsense) statements significantly low in comprehensibility and literally true statements far higher in comprehensibility, indicating that they were making meaningful judgments. This suggests that very early processing is done by symmetric alignment. Later in processing (at 1,600 ms), people strongly preferred the forward direction over the reverse direction, suggesting that the projection of inferences is a later directional process. This pattern suggests that, by 500 ms, people have a sense of comprehension taking place, and this sense is independent of direction—consistent with an initial alignment process rather than with initial directional projection.

(Fig. 4). This was not due to failure to process meaning—even at 500 ms, people rejected nonsense (scrambled) sentences and accepted literally true statements. Later in processing—by 1,600 ms—people strongly preferred the forward direction. This is consistent with an early symmetric alignment process

followed by a later process of projecting inferences, as in SME.

Large-scale simulations with SME

SME has been used to model similarity-based retrieval using a two-phase process called “many are called but few are chosen” (MAC/FAC; Forbus et al., 1995). Its initial phase uses a massively parallel, inexpensive process that treats structures like vectors. This produces a small number of structured descriptions that are then filtered via SME. MAC/FAC has successfully simulated psychological results, including the finding that surface matches are commonly retrieved (e.g., Gentner et al. 1993; Trench & Minervino 2015). SME has also been used to model generalization over examples, as in the sequential analogical generalization engine (SAGE; Kandaswamy & Forbus, 2012). In SAGE, SME is used to compare incrementally added examples to produce probabilistic schemas a series of similar examples. An early version of SAGE successfully modeled infants’ analogical generalization across simple grammatical patterns (Kuehne et al., 2000).

Structure-mapping theory holds that the same analogical processes operate throughout human cognition in both perceptual and conceptual domains. If so, SME should be able to simulate that same range of phenomena. The evidence so far is that it does. In visual problem-solving, for example, the CogSketch system (Forbus et al., 2011) provides a model of high-level human vision. Stimuli are provided either via hand-drawn sketches, via copy/paste from slides, or via deep learning computer vision systems operating over images or Kinect video (Forbus et al., 2025). CogSketch automatically computes visual qualitative representations such as positional relations (e.g., “above,” “left”) and topological relations (e.g., “touching,” “inside”), decomposes shapes into edges, and recognizes larger groups. The same representations are used for geometric analogies,⁵ for visual oddity tasks, and for Raven’s progressive matrices (RPM), among others (Forbus & Lovett, 2021). The RPM task is viewed as the best predictor of human fluid intelligence, and the CogSketch model scores in the 61st percentile⁶ (Lovett & Forbus, 2017)—evidence that SME performs in a humanlike way on a key test of human intelligence.

Modeling large-scale tasks involves embedding SME into other systems such as CogSketch. Cognitive architectures are models that seek to capture many different cognitive phenomena, closer to what minds can do (Newell 1994). SME, MAC/FAC, and SAGE play central roles in a cognitive architecture called Companion (Forbus & Hinrichs, 2017). For example, combining

analogy with the natural language capabilities of Companion has enabled it to learn by reading and, with CogSketch, by multimodal reading. Analogical learning within Companion has also been used to learn games and textbook problem-solving, as well as to model aspects of human conceptual change (Forbus & Hinrichs, 2017).

SME has been used in deployed systems. For example, in a Companion-based information kiosk that answers questions about Northwestern University’s computer science department, analogy is used in interpreting questions (Wilson et al. 2019). The most widespread deployment of SME has been in *sketch worksheets* (Forbus et al., 2020), which rely on CogSketch’s visual representations. Instructors in geoscience and other topics use CogSketch to construct a problem for students and to annotate the facts CogSketch automatically produces with advice.⁷ When a student attempts a problem, SME is used to compare the student and instructor sketches. It identifies alignable differences between the two sketches and uses these differences to give instructor-provided advice. Thus, the students get targeted feedback and instructors get detailed student data (Forbus et al., 2020). Sketch worksheets have been used in classes at multiple age levels. Sketch worksheets utilize SME’s ability to do a humanlike similarity comparison over visual materials. These deployments, along with the breadth of phenomena SME has been used to model, demonstrate SME’s robustness.

SME is a symbolic model. A natural question is how it relates to deep learning and/or generative AI models. Deep learning of visual components has been found to be useful in a hybrid system that produces inputs for CogSketch (Forbus et al., 2025). Large language models have also been used in conjunction with a symbolic language system to improve learning by reading (Nakos & Forbus, 2023). Some tighter integrations have been explored (e.g., using vector embeddings as a component in SME’s structural evaluation calculations; Liang et al., 2016). However, whether generative AI systems are capable of doing analogical processing without hybridization with a symbol system remains an open question at this point.

SME has been a focus of fruitful interaction between AI research and human research for 4 decades. We suggest that attention to analogical processing could be useful for AI more broadly. Analogical learning is incremental, unlike deep learning, which operates offline. Analogical learning is also more data-efficient than deep learning (e.g., Forbus et al., 2025). Last, analogical learning produces inspectable models with structured representations that are easily understood, in contrast to the opaque distributed representations built by deep

learning systems. Thus, we suggest that human models of analogical processing can lead to more humanlike AI systems.

Recommended Reading

- Forbus, K. D., Ferguson, R. W., Lovett, A., & Gentner, D. (2017). (See References). Describes how the structure-mapping engine works.
- Forbus, K. D., & Hinrichs, T. (2017). (See References). Summarizes cognitive simulation and AI experiments using the Companion architecture.
- Gentner, D. (2010). (See References). Describes how the structure-mapping engine affects human learning.
- Gentner, D., & Smith, L. A. (2013). Analogical learning and reasoning. In D. Reisberg (Ed.), *Oxford library of psychology* (pp. 668–681). Oxford University Press. Summarizes research on analogical comparison.

Transparency

Action Editor: Robert L. Goldstone

Editor: Robert L. Goldstone

Declaration of Conflicting Interests

The authors declared no potential conflicts of interest with respect to the research, authorship, and/or publication of this article.

ORCID iDs

Dedre Gentner  <https://orcid.org/0000-0002-5120-6688>

Kenneth Forbus  <https://orcid.org/0000-0003-2067-5227>

Notes

1. These representations are meant to capture psychological construals of a situation. Construals of the same external information may vary across and within people in different contexts. We do not claim that structured representations are the only kind of representation that humans use; however, we maintain that this is one important kind of available representation.
2. Other evaluation processes happen outside of the structure-mapping engine, such as checking whether the inferences are valid in the target and (for goal-directed analogies) noting whether the inferences are relevant to the goal.
3. There are special cases in which graph matching is more tractable (e.g., van Rooij et al., 2008), but we do not limit the structure-mapping engine to those cases, preferring instead an approximate algorithm.
4. The structure-mapping engine uses what in computer science is called a “greedy algorithm” that does not guarantee optimal answers but is very efficient. The complexity of the structure-mapping engine is $O(n^2 \log(n))$.
5. To download the geometric analogy model plus stimuli, see the CogSketch distribution at <https://www.qrg.northwestern.edu/software/cogsketch/index.html>.
6. The structure-mapping engine's success on the Raven's task is not susceptible to the complaint that the answers were already in its knowledge base. Unlike large language models, it is not

trained on data from the web, and its knowledge and reasoning are inspectable.

7. More than 30 sketch worksheets that cover introductory geoscience classes are available online through the Science Education Resource Center at Carleton College.

References

- Bowlde, B., & Gentner, D. (1997). Informativity and asymmetry in comparisons. *Cognitive Psychology*, 34, 244–286.
- Buckley, S. (1979). *Sun up to sun down*. McGraw-Hill.
- Christie, S., & Gentner, D. (2010). Where hypotheses come from: Learning new relations by structural alignment. *Journal of Cognition and Development*, 11(3), 356–373.
- Clement, C., & Gentner, D. (1991). Systematicity as a selection constraint in analogical mapping. *Cognitive Science*, 15, 89–132.
- Doumas, L. A. A., Hummel, J. E., & Sandhofer, C. M. (2008). A theory of the discovery and predication of relational concepts. *Psychological Review*, 115, 1–43.
- Dunbar, K. (1993). Concept discovery in a scientific domain. *Cognitive Science*, 17, 391–434.
- Forbus, K., Usher, J., Lovett, A., Lockwood, K., & Wetzel, J. (2011). CogSketch: Sketch understanding for cognitive science research and for education. *Topics in Cognitive Science*, 3(4), 648–666.
- Forbus, K. D., Chen, K., Xu, W., & Usher, M. (2025). Hybrid primal sketch: Combining analogy, qualitative representations, and computer vision for scene understanding. *Cognitive Systems Research*, 93, Article 101390. <https://doi.org/10.1016/j.cogsys.2025.101390>
- Forbus, K. D., Ferguson, R. W., Lovett, A., & Gentner, D. (2017). Extending SME to handle large-scale cognitive modeling. *Cognitive Science*, 41, 1152–1201. <https://doi.org/10.1111/cogs.12377>
- Forbus, K. D., Garnier, B., Tikoff, B., Marko, W., Usher, M., & Mclure, M. (2020). Sketch worksheets in STEM classrooms: Two deployments. *AI Magazine*, 41(1), 19–32. <https://doi.org/10.1609/aimag.v41i1.5189>
- Forbus, K. D., Gentner, D., & Law, K. (1995). MAC/FAC: A model of similarity-based retrieval. *Cognitive Science*, 19, 141–205.
- Forbus, K. D., & Hinrichs, T. (2017). Analogy and qualitative representations in the Companion cognitive architecture. *AI Magazine*, 38(4), 34–42. <https://doi.org/10.1609/aimag.v38i4.2743>
- Forbus, K. D., Liang, C., & Rabkina, I. (2017). Representation and computation in cognitive models. *Topics in Cognitive Science*, 9, 694–718. <https://doi.org/10.1111/tops.12277>
- Forbus, K. D., & Lovett, A. (2021). Same/different in visual reasoning. *Current Opinion in Behavioral Sciences*, 37, 63–68. <https://doi.org/10.1016/j.cobeha.2020.09.008>
- Gentner, D. (1983). Structure-mapping: A theoretical framework for analogy. *Cognitive Science*, 7, 155–170.
- Gentner, D. (2010). Bootstrapping the mind: Analogical processes and symbol systems. *Cognitive Science*, 34(5), 752–775.
- Gentner, D., Rattermann, M. J., & Forbus, K. D. (1993). The roles of similarity in transfer: Separating retrieval from inferential soundness. *Cognitive Psychology*, 25, 524–575.

- Goldstone, R. L. (1994). Similarity, interactive activation, and mapping. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 20(1), 3–28.
- Haryu, E., Imai, M., & Okada, H. (2011). Object similarity bootstraps young children to action-based verb extensions. *Child Development*, 82(2), 674–686.
- Hummel, J. E., & Holyoak, K. J. (2003). A symbolic-connectionist theory of relational inference and generalization. *Psychological Review*, 110(3), 220–264.
- Kandaswamy, S., & Forbus, K. (2012). Modeling learning of relational abstractions via structural alignment. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 34, 545–550.
- Katz, E. D. (2015). Justice Ginsburg's umbrella. In S. R. Bagenstos & E. D. Katz (Eds.), *A nation of widening opportunities? The Civil Rights Act at fifty*. Michigan Publishing.
- Keane, M. T. G. (1990). Incremental analogizing: Theory and model. In K. J. Gilhooly, M. T. G. Keane, R. H. Logie, & G. Erdos (Eds.), *Lines of thinking: Reflections on the psychology of thought* (Vol. 1). John Wiley & Sons.
- Kotovsky, L., & Gentner, D. (1996). Comparison and categorization in the development of relational similarity. *Child Development*, 67, 2797–2822.
- Krawczyk, D., Holyoak, K., & Hummel, J. (2004). Structural constraints and object similarity in analogical mapping and inference. *Thinking and Reasoning*, 10, 85–104.
- Kuehne, S. E., Gentner, D., & Forbus, K. D. (2000). Modeling infant learning via symbolic structural alignment. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 22, 286–291.
- Liang, C., Paritosh, P., Rajendran, V., & Forbus, K. D. (2016). Learning paraphrase identification with structural alignment. In G. Brewka (Ed.), *IJCAI'16: Proceedings of the Twenty-Fifth International Joint Conference on Artificial Intelligence* (pp. 2859–2865). AAAI Press.
- Lovett, A., & Forbus, K. D. (2017). Modeling visual problem solving as analogical reasoning. *Psychological Review*, 124(1), 60–90. <https://doi.org/10.1037/rev0000039>
- Markman, A. B., & Gentner, D. (1993). Structural alignment during similarity comparisons. *Cognitive Psychology*, 25, 431–467.
- Nakos, C., & Forbus, K. D. (2023). Using large language models in the companion cognitive architecture: A case study and future prospects. *Proceedings of the AAAI Symposium Series*, 2(1), 356–359. <https://doi.org/10.1609/aaais.v2i1.27700>
- Newell, A. (1994). *Unified theories of cognition*. Harvard University Press.
- Trench, M., & Minervino, R. A. (2015). The role of surface similarity in analogical retrieval: Bridging the gap between the naturalistic and the experimental traditions. *Cognitive Science*, 39(6), 1292–1319.
- van Rooij, I., Evans, P., Muller, M., Gedge, J., & Wareham, T. (2008). Identifying sources of intractability in cognitive models: An illustration using analogical structure mapping. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 30, 915–920.
- Wilson, J., Chen, K., Crouse, M. C., Nakos, C., Ribeiro, D., Rabkina, I., & Forbus, K. D. (2019, August 2–5). *Analogical question answering in a multimodal information kiosk* [Paper presentation]. Proceedings of the Seventh Annual Conference on Advances in Cognitive Systems, Cambridge, MA, United States.
- Wolff, P., & Gentner, D. (2011). Structure-mapping in metaphor comprehension. *Cognitive Science*, 35(8), 1456–1488. <https://doi.org/10.1111/j.1551-6709.2011.01194.x>