

Intuitive Network Topology

Sami R. Yousif and Elizabeth M. Brannon
Department of Psychology, University of Pennsylvania

Topology is the branch of mathematics that seeks to understand and describe spatial relations. A number of studies have examined the human perception of topology—in particular, whether adults and young children perceive and differentiate objects based on features like closure, boundedness, and emptiness. Topology is about more than “wholes and holes,” however; it also offers an efficient language for representing network structure. Topological maps, common for subway systems across the world, are an example of how effective this language can be. Inspired by this idea, here we examine “intuitive network topology.” We first show that people readily differentiate objects based on several different features of topological networks. We then show that people both remember and match objects in accordance with their topology, over and above substantial variation in their surface features. These results demonstrate that humans possess an intuitive understanding for the basic topological features of networks, and hint at the possibility that topology may serve as a format for representing relations in the mind.

Public Significance Statement

Topology is infamously unintuitive. Objects like Klein bottles, mobius strips, and the Alexander horned sphere all have surprising, unusual properties. Yet some instances of topology are striking for how intuitive they are: Topological subway maps, for instance, seem somehow more natural than veridical, Euclidean representations of space. Here, we explore this latter sort of topology. We show that people are surprisingly sensitive to topological network structure, such that they will readily discriminate, match, and remember items on its basis. We argue that this sensitivity to topology may be indicative of an intuitive “language” for representing spatial relations in general, even beyond the domain of spatial cognition.

Keywords: topology, geometry, spatial perception, spatial cognition, relations

What do a coffee mug and a donut have in common? Topologically speaking, everything. The two are homeomorphic: The two can be continuously transformed into each other without any cutting or gluing. In the world of topology, these two very different forms are exactly the same.

Topology—the branch of mathematics that concerns itself with spatial relations—is often profoundly unintuitive in exactly this way. That a donut shape can be transformed into a coffee mug is surprising to many (see Figure 1A). That an object (i.e., a mobius strip) can have only one surface and one edge is difficult to comprehend, even when looking directly at such a figure (see Figure 1C). Yet there are cases for which thinking topologically is quite natural—more so, even, than canonical representations of space. Consider Figure 1D, for instance. This is a topological map of the London underground,

first drawn this way in 1931 by English designer Harry Beck. Despite the complexity of the map, we have no trouble understanding it; with just a glance, it is easy to see how to get from one place to another. This near-unnoticeable simplicity hints at the possibility that at least certain forms of topological representation (i.e., those which describe networks) are highly intuitive—that we may be endowed with a basic capacity to reason in topological terms.

What differentiates the topological map of the London underground (and other subway maps) from most canonical forms of maps is that it is not a veridical rendering of Euclidean space. What the map aims to capture is not exact geometric relations (e.g., precise spatial information like distances or angles), but, instead, the relations between various stops and lines. Angles and distances—key properties of Euclidean geometry—are purposefully distorted in this map.

Karen Schloss served as action editor.

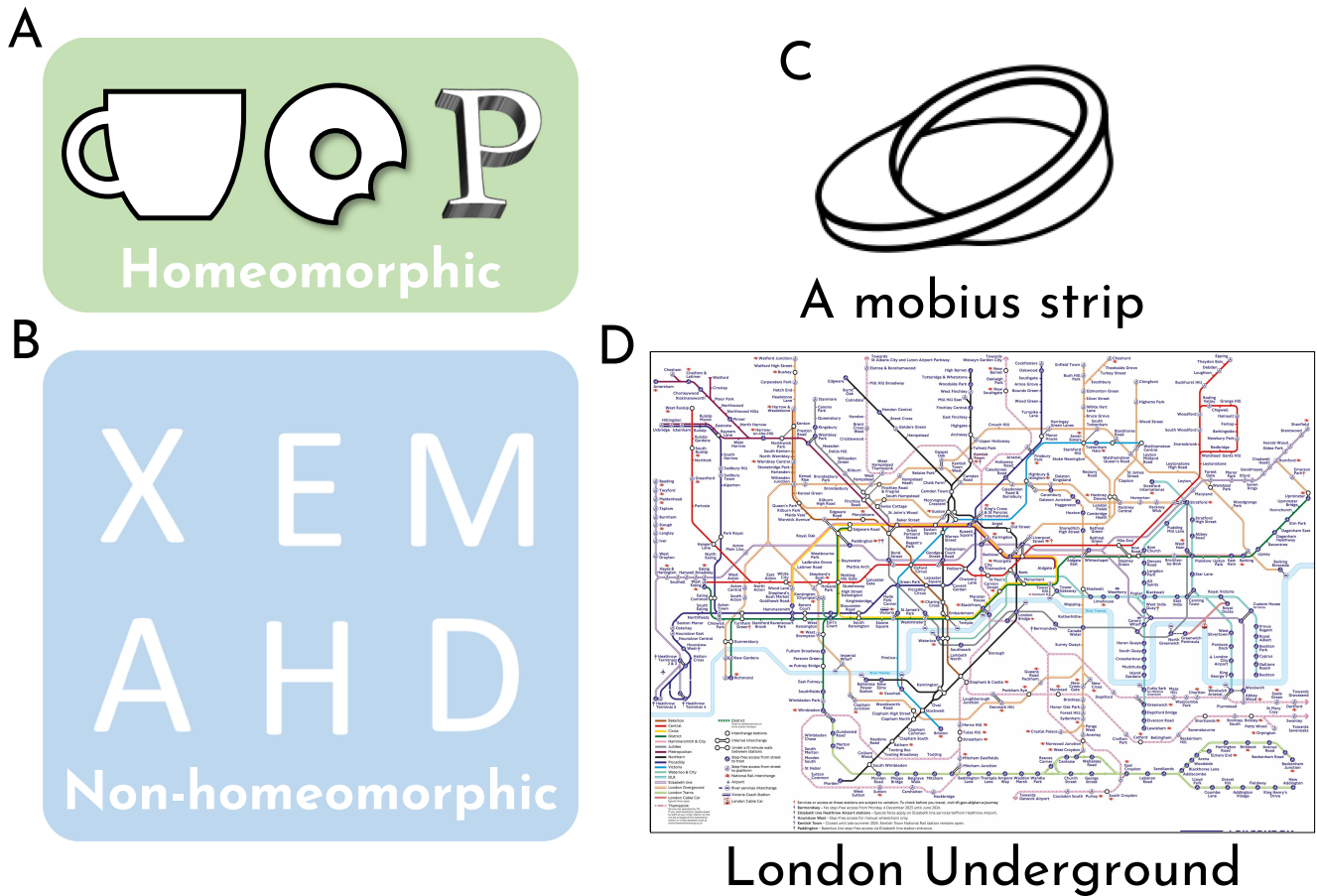
The findings of this study have not been presented at any other venues, nor published anywhere online before the publication of this article. We thank Richard Aslin, Tal Boger, Sam Clarke, Russell Epstein, and Frank Keil for feedback on earlier iterations of these ideas and/or feedback on early drafts of the article. We thank Melissa Kibbe for detailed, thoughtful feedback during the review process. For helping with stimulus creation, we thank Yijin Hu. Finally, we wish to thank all the members of the Developing Minds lab at the University of Pennsylvania for their feedback on this work. The authors have no conflicts of interest. Preregistrations, stimuli, raw data, and other supplemental materials are available on our

Open Science Framework page: <https://osf.io/w8p6s/>.

Sami R. Yousif served as lead for conceptualization, data curation, formal analysis, investigation, methodology, project administration, resources, software, validation, visualization, and writing—original draft. Elizabeth M. Brannon served as lead for funding acquisition and supervision and served in a supporting role for conceptualization. Sami R. Yousif and Elizabeth M. Brannon contributed equally to writing—review and editing.

Correspondence concerning this article should be addressed to Sami R. Yousif, Department of Psychology, University of Pennsylvania, 425 South University Avenue, Stephen A. Levin Building, Philadelphia, PA 19104-6241, United States. Email: sryousif@sas.upenn.edu

Figure 1
Examples of Topological Forms



Note. (A) An example of three homeomorphic shapes. Famously, a coffee mug and a donut have the same topology. The letter “P,” if printed in a three-dimensional form, also shares this topology. (B) An example of English letters which have a unique topology; none of these is like the other. (C) A depiction of a mobius strip, an object known for its unique, unintuitive topology, as it has only one side and one edge. (D) A modern rendition of the original “topological map” designed for the London underground, from *Transport for London*, by Transport for London, 2024 (<https://tfl.gov.uk>). Reprinted with permission. See the online article for the color version of this figure.

Notice, for example, how all angles come in 45° increments, or how most stops along a given line are equidistant. In reality, the tracks are not so neatly laid, and the stops are not so evenly distributed. From the perspective of the passengers on the underground, however, this map captures exactly what matters most: where they are in relation to the rest of the system, and what path they should take to reach their destination most efficiently. The geometric details do not matter, but the spatial relations do.

Topology concerns itself with properties that are preserved following deformations such as twisting and stretching. Network topology, therefore, concerns itself with the shapes of networks independent of such continuous deformations. A given line could be twisted or stretched such that a dozen stops were removed or added, but it would still be (functionally) one direct path. In the case of subway maps, what matters functionally is only whether or where you have to switch lines; stops that contain line switches are functional nodes in a way that other stops are not.

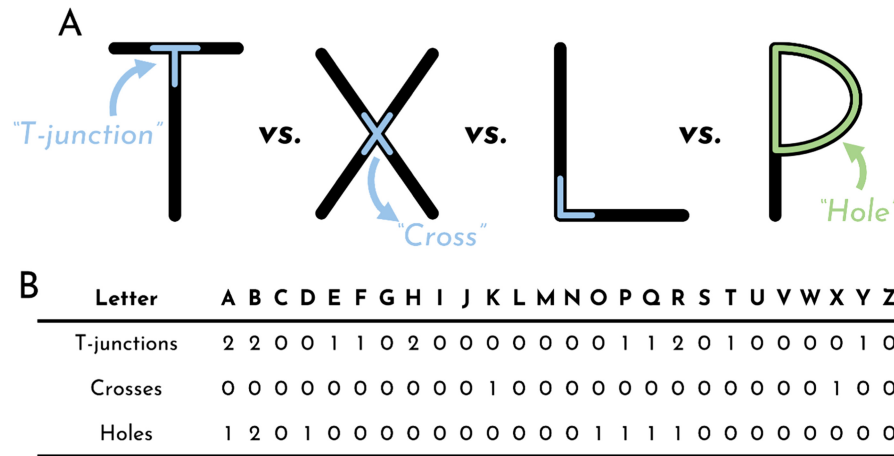
Topological networks can be described by a sort of “language” that consists of a few discrete “parts.” Consider, for instance,

different tracks in a network that are either in the shape of a “T” or an “L.” A “T” shape and an “L” shape have much in common (see Figure 2A). They are both made up of two adjacent lines. They both have one vertex. (They are both contained in the world “topology.”) But, if we think about these shapes as networks, the “T” has one thing that an “L” does not. If the answer is not obvious, it may help to consider a different formulation of the question: What prevents the “T” from being continuously transformed into an “L”?

The answer is that a “T” shape has a vertex at which three lines meet (see Figure 2A). The “L” shape, in contrast, has only a 2-point vertex. A 2-point vertex, in topological terms, is irrelevant. One can easily imagine how the “L” could be twisted at that vertex to form a straight line (like a straw with a built-in bend). This is to say that the “L” is homeomorphic with a straight line, but not with a “T.” No matter how you twist or stretch or bend the “T,” it is 3-point vertex prevents it from becoming a straight line. This 3-point vertex (or “T-junction”) is a functional “building block” of a larger structure.

It is also possible to have 4-point vertices (a “cross”; see Figure 2A), or 5- or 6-point vertices (as in an asterisk), or any n -point vertex. Any

Figure 2
Network Topology



Note. (A) An example of the relative topological “parts” of the letters “T,” “X,” “L,” and “P.” (B) The number of 3-point vertices (T-junctions), 4-point vertices (crosses), and holes of all 26 letters of the English alphabet. Letters like “C,” “M,” and “Z” are the same on all the dimensions, and, in this case, are homeomorphic. (The fact that two shapes have the same topological features does not necessarily mean that they are homeomorphic; just like the order of words in a sentence affects its meaning, how topological features are combined affects their overall structure.) See the online article for the color version of this figure.

such n -point vertex (greater than $n=2$) is a distinct topological form; none can be transformed into any other. From a topological perspective, these basic building blocks are the only units that matter. It does not matter, for instance, how many stops are placed along a line in the subway system, or how many times that line bends. It only matters whether and how two lines intersect with one another. Those intersections reflect the choice points relevant to passengers.

In addition to the various forms of vertices, there is one other essential building block of network topologies. This one is easier to spot. Consider “T” versus “P.” What is the difference between them? Both have a single T-junction. Both are made up of two lines. (Both are contained in the word “topology”!). But only the latter has a hole. Holes, like T-junctions and crosses, are a distinct topological feature of networks.¹

Although we do not typically think or talk about objects in topological terms, it is easy to see how these building blocks are intuitive: The difference between “T” and “P” feels almost so obvious that it need not be stated; the difference between “T” and “L” might be slightly less obvious, but the difference between them is nevertheless clearly appreciable. A question arises, then, about whether (or to what extent) people naturally represent objects and spaces with respect to their network topology.

Prior Work on Topology

Sensitivity to certain kinds of topological structure has been examined extensively for over four decades (see, e.g., Chen, 1982, 2005; Zhou et al., 2010). This research program began with a claim that the visual system represents topological structure (Chen, 1982; but see Rubin & Kanwisher, 1985, as well as Chen, 1990), but has been extended to show how other processes, like apparent motion (Chen, 1985), working memory (Wei et al., 2019), and number perception (He et al., 2015) are influenced by topology. This sensitivity to topology appears to be deep-seated: Human infants (Chien et al., 2012;

Kibbe & Leslie, 2016) and such distant species as bees (Chen et al., 2003) exhibit sensitivity to basic topological properties like closure.

Most of this prior work focuses on the topology of objects, rather than the sort of network topology that is represented in a topological subway map. Though related, network topology and object topology differ in important ways. Consider the letters “T” and “L,” for instance. From the perspective of object topology—if we imagine these letters as three-dimensional objects that we can touch—they are homeomorphic. They are both just wholes without holes, no different from a cube, a sphere, or a pyramid. If one imagines these shapes not as pieces of paper, but as segments in a maze, however, the difference between them becomes clearer. If you were navigating a maze and you reached the 2-point vertex of the “L,” you would have no choice to make; there is only one path to follow. If, instead, you found yourself at the 3-point

¹ We offer the following note about the terminology throughout our article: We refer primarily to “holes,” “crosses,” and “T-junctions.” However, these are not the only terms used to describe these structures. What we’ve referred to as a “hole” here is sometimes referred to as a “face.” This would be especially true when thinking about surfaces rather than nodes occupying empty space. For instance, each of the four visible surfaces of a canonical square pyramid may be considered a “face” with three vertices. Similarly, “T-junctions” and “crosses” are really just vertices connecting different numbers of edges. The vertex at the apex of a square pyramid is a 4-point vertex (“cross”) and the vertices at the base are 3-point vertices (“T-junctions”). Networks can sometimes be described in terms of their “degree sequence,” which is just an ordered list of the number of degrees (edges) at each vertex. A square pyramid has a degree sequence of 4-3-3-3-3, as there are five vertices, one connecting three edges and four connecting three. “Graph invariants,” that is, networks that are topologically identical, will always have the same degree sequence; however, two networks with identical degree sequence will not necessarily be graph invariant. Also, T-junctions are sometimes broken down into more detailed shapes like Y-junctions, but this difference is not topologically relevant, and therefore we refer to all possible 3-point vertices as T-junctions.

vertex of the “T,” you would have to pick a side. No matter which direction you came from, there would be two branching paths from that point forward. Likewise, at a 4-point vertex (a cross) there would be three options; at a 5-point vertex there would be four; and so on. Holes, too, would affect how you navigate. The simplest topological form containing a hole is just a circle, or an “O.” If you were trapped inside a maze of that form, you would never have any decisions to make, but you would never escape, either.

The difference between object topology and network topology matters in practice. It may alter the interpretation of some prior work. For example, consider Experiments 1 and 2 of Wei et al. (2019). They compare three key conditions: A no-change condition, a nontopological change condition, and a topological change condition (see their Figure 1). Through the lens of network topology, the nontopological change condition actually contains a meaningful topological change. It compares an “E” and an “H.” The former has a single T-junction, whereas the latter has two.

While features of object topology like closure, overlap, and embedding have been widely studied (Wei et al., 2019; but see also Kenderla et al., 2023; Kibbe & Leslie, 2016; Renz et al., 2000; Rips, 2020), including in work which has utilized graph-like stimuli rather than filled-in shapes (see Kanbe, 2013), features of network topology like T-junctions and crosses have not been explicitly studied in this way (but see, e.g., Lowet et al., 2018). It remains unclear whether to what extent these different kinds of topology are related.

The kind of topology studied here, in contrast with prior work on object topology, offers a “language” for describing relations, whether those be the relations of social groups, institutional structures, or logistical networks (for a review of the study of relations in cognitive science, see Hafri & Firestone, 2021). In this way, network topology, by virtue of capturing said relations, may lie at the heart of human cognition.

Current Study

Here, we investigate sensitivity to some of the “building blocks” of relational, topological reasoning (e.g., crosses, T-junctions, and holes). These parts are meaningful because they are functional: This simple “language,” consisting of two types of functional “units,” is sufficient not only to describe the entire English alphabet (see Figure 2B) but all conceivable network structures. Therefore, we ask: To what extent do people intuitively “speak” the language of topological representation? To what extent do they readily represent and differentiate the functional building blocks of topological networks? We answer these questions using three different experimental paradigms. Experiments 1 and 2 use an odd-one-out paradigm to ask whether participants discriminate based on topological relations. Experiment 3 uses a memory paradigm to ask whether participants are more likely to false alarm when lures share topological features from previously encoded items. Experiment 4 requires participants to make explicit similarity judgments and asks whether changes in topology alter perceived similarity more than changes in other surface features.

Experiment 1a: Odd One Out (Three Segments)

In a first experiment, we tested sensitivity to topological structure in the most straightforward way possible. Participants completed a simple “odd-one-out” game modeled after the “intruder” paradigm used by Dehaene et al., (2006). Participants were tasked with

selecting which of six fictional “letters” was not like the others in cases where five of the six forms had the same topology and one differed slightly. If people are sensitive to topological structure, they should be able to successfully identify the odd one out. At first blush, this may not seem like a bold prediction: It should be easy to pick a single nontriangle out of a set of triangles. However, some of the comparisons are much less obvious. See the example trial in Figure 3A. Can you tell which item is the odd one out?

One strength of this design is that we can ask not only whether people distinguish items based on their topology, but also how they distinguish items based on their topology. Are there certain topological features (e.g., holes, T-junctions) that are more salient than others? Does the likelihood of selecting the odd-one-out scale with the number of topological differences of the distinct item? In other words, our stimulus set was designed to allow us to assess whether people are broadly sensitive to topological differences, or whether they are only sensitive to certain ones, sometimes.

Method

Transparency and Openness

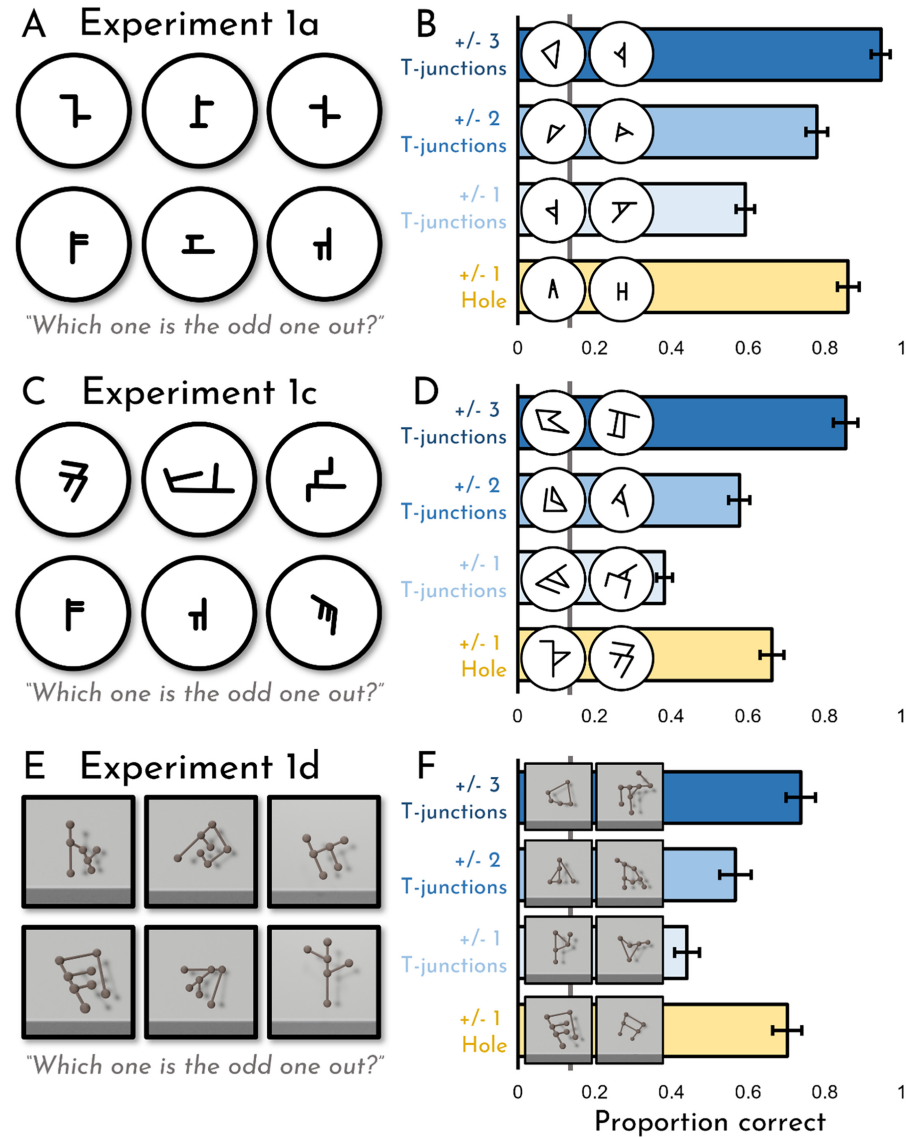
Here, and for all subsequent experiments in this article, the sample sizes, primary dependent variables, and key statistical tests were chosen in advance and were preregistered. This information, as well as raw data and materials, are available on a public Open Science Framework (OSF) page (<https://osf.io/w8p6s/>; Yousif & Brannon, 2024).

Participants

Fifty participants completed the task online via Prolific. All participants were adults currently residing in the United States. The exclusion criteria we preregistered were highly conservative, and thus no participants were excluded. No demographic data were collected. This study was approved by the relevant Institutional Review Board.

Stimuli

There were 42 distinct stimuli, six exemplars for each of seven distinct topologies. In terms of [T-junctions-holes], the seven unique topologies were: [0-0], [0-1], [1-0], [1-1], [2-0], [2-1], [3-1]. ([0-0] refers to an item with zero T-junctions and zero holes, like an “L,” whereas [1-0] refers to an item with one T-junction and zero holes, like a “T.”) All stimuli are available on our OSF page (and are described using a similar notation). Each of the six stimuli on each trial was presented inside of a circle, which was designed to be approximately 1 in. in diameter on a typical computer display. Stimuli were designed to be as heterogeneous as possible while consisting of exactly three line segments. This means that each stimulus could be drawn using three straight lines, ignoring any vertices. For instance, a “T” could be thought of as one line resting on top of another line, or it could be thought of as three lines joining at a central point. For our purposes here, this would count as only two line segments. We held the number of line segments constant to minimize variation across stimuli and avoid potential confounds. Additionally, the use of three line segments ensured that stimuli were as simple and straightforward as possible.

Figure 3*Experiment 1*

Note. Design (A) and results (B) of Experiments 1a. Design (C) and results (D) of Experiments 1c. Design (E) and results (F) of Experiments 1d. The images superimposed on the bars represent a single contrast with that number/kind of topological differences. Error bars represent $\pm 1 SE$. As there are six options, chance performance is equal to 16.67%, indicated by the gray line. The correct answer in (A) is the top-left item; the correct answer in (C) is the top-right item; the correct item in (E) is the top-middle item; in each of these cases, the correct answer has one fewer T-junction than the other items. See the online article for the color version of this figure.

Procedure

Participants were introduced to the “odd-one-out” task in which they would see six “letters” and were asked to identify which one was not like the others. Participants were told that these letters were from a hypothetical alphabet that did not necessarily follow the rules of English letters. They were further told that the letters were presented in a random orientation. They were given no other instruction about how they were meant to complete the task. Each trial consisted of a 3×2 grid of stimuli (see Figure 3A), from which observers were

prompted to select the odd one out. There was no time limit on their responses, and they were not instructed to hurry. Participants completed two representative practice trials, the data from which were not recorded, before beginning the task.

Not all topologies were compared against all others. Of the 21 possible comparisons, we selected 12. The 12 comparisons were as follows: [0-0 vs. 0-1], [0-0 vs. 1-0], [0-1 vs. 1-1], [1-0 vs. 1-1], [1-1 vs. 2-1], [2-0 vs. 1-0], [2-0 vs. 2-1], [2-1 vs. 3-1], [2-0 vs. 0-0], [1-1 vs. 3-1], [2-1 vs. 0-1], [3-1 vs. 0-1]. These were selected so that (a) the odd one out always differed either in the number of

holes or in the number of T-junctions, but not both, and (b) the odd one out would differ from the rest by different degrees (e.g., on some trials, there may be a difference of one T-junction such that five items have two T-junctions and one has only one T-junctions; on others, there may be a difference of two or three T-junctions). Each of the 12 comparisons was presented a total of four times: Twice with one of the topologies in the majority and twice with the other topology in the majority (i.e., for the comparison [0-0 vs. 0-1], there were two trials for which 0-0 was the odd one out and two trials for which 0-1 was the odd one out). The specific exemplars that were chosen for each topology, as well as the locations in which they appeared, were fully randomized with no constraints. This meant that, for a given comparison [such as 0-0 vs. 0-1], participants might see different exemplars.

Results and Discussion

Results from this experiment can be seen in Figure 3B (for more detailed results, see Figure S1 in the additional online materials on the OSF page). Participants were well above-chance discriminating topological differences. This was true for all comparison types, ± 1 hole, $t(49) = 24.49$, $p < .001$, $d = 3.46$; ± 1 T-junction, $t(49) = 17.44$, $p < .001$, $d = 2.47$; ± 2 T-junctions, $t(49) = 22.06$, $p < .001$, $d = 3.12$; ± 3 T-junctions, $t(49) = 31.08$, $p < .001$, $d = 4.40$. Moreover, this was true for each of the 12 unique comparisons ($ps < .001$). As is evident from the figure, differences in holes were more salient than T-junctions; participants were better at detecting a difference of one hole than differences of one, $t(49) = 12.85$, $p < .001$, $d = 1.82$ or even two, $t(49) = 4.83$, $p < .001$, $d = .68$, T-junctions, but worse than differences of three T-junctions $t(49) = 4.93$, $p < .001$, $d = .70$.

While some of these differences may seem obvious once pointed out, keep in mind that participants were given no instructions whatsoever about how they were meant to differentiate the items. They could have attempted to differentiate based on the angle of the lines, the length of the lines, the length of whole item, and so on. As the example in Figure 3A illustrates, some of these comparisons can be quite challenging even with knowledge of the topological distinctions. Yet, even for this comparison—the most difficult one in the set—participants were well above chance.

Experiment 1b: Odd One Out (Four Segments)

Experiment 1a relied exclusively on “letters” made up of exactly three line segments. This eliminated any confound with the number of line segments, but it also limited the available number of topology forms: There are only so many ways to create one hole and one T-junction with three line segments. So, to make the task slightly more complex, we constructed a new stimulus set comprised of letters made up of four line segments and replicated the study.

Method

Everything about the design and procedure was identical to Experiment 1a, except that all stimuli consisted of four line segments instead of three. Fifty new participants completed the task via Prolific. There were no exclusions.

Results and Discussion

Participants were again well above-chance discriminating based on topological differences. This was true for all comparison types, ± 1

hole, $t(49) = 13.21$, $p < .001$, $d = 1.87$; ± 1 T-junction, $t(49) = 8.32$, $p < .001$, $d = 1.18$; ± 2 T-junctions, $t(49) = 12.51$, $p < .001$, $d = 1.77$; ± 3 T-junctions, $t(49) = 17.37$, $p < .001$, $d = 2.46$, and for all but two of the 12 unique comparisons ($ps < .05$; eight of the 12 were $p < .001$). Differences in holes were again relatively more salient than T-junctions; participants were better at detecting a difference of one hole than differences of one, $t(49) = 9.39$, $p < .001$, $d = 1.33$, or even two, $t(49) = 3.42$, $p = .001$, $d = .48$, T-junctions. Thus, these results replicate and extend the results of Experiment 1a (see Figure S2 in the additional online materials at the OSF page for a summary figure; see Figure S3 in the additional online materials at the OSF page for complete results).

Experiment 1c: Odd One Out (Three to Five Segments)

A potential concern is that exemplars of a given topology may be said to resemble each other to a large degree. Even if one struggles to articulate or quantify what exactly, other than their topology, makes two items more similar, it is hard to deny that this is a reasonable argument. To address this concern, we replicated Experiments 1a and 1b, except that each trial contained stimuli with three, four, and five line segments, all intermixed (see Figure 3C). Still, just as before, only one of the items differed topologically. In this design, the surface features of the objects varied much more on each trial. We reasoned that, as the stimuli were significantly more heterogeneous, the differences in topology should be less salient. Thus, if participants still successfully identified the odd one out, it would provide stronger evidence of genuine sensitivity to topological structure (rather than some other low-level property).

Method

Everything about the design and procedure was identical to Experiments 1a and 1b, except that the stimuli were made up of three, four, and five line segments. Fifty new participants completed the task via Prolific. There were no exclusions.

This experiment included the 42 stimuli used in Experiment 1a, the 42 stimuli used in Experiment 1b, and 42 new stimuli (comprised of five line segments) designed specifically for this experiment. Thus, there was a much larger pool of stimuli which could be selected on a given trial. The items were selected randomly with the only constraint being that each trial must consist of two items with three segments, two items with four segments, and two items with five segments. The odd one out was nevertheless randomly determined and could have had any number of line segments (i.e., the number of line segments of the odd one out was not counterbalanced across trials). As is evident from Figure 3C, this design results in stimuli that were considerably more heterogeneous.

Results and Discussion

Participants were again well above-chance discriminating based on topological differences. This was true for all comparison types, ± 1 hole, $t(49) = 15.89$, $p < .001$, $d = 2.25$; ± 1 T-junction, $t(49) = 10.26$, $p < .001$, $d = 1.45$; ± 2 T-junctions, $t(49) = 15.02$, $p < .001$, $d = 2.12$; ± 3 T-junctions, $t(49) = 21.24$, $p < .001$, $d = 3.00$, and for all 12 unique comparisons ($ps < .05$; 11 of the 12 were $p < .001$). Differences in holes were again relatively more salient than T-junctions; participants were better at detecting a difference of 1-hole than differences of one, $t(49) = 12.10$, $p < .001$, $d = 1.71$,

and two, $t(49) = 4.65$, $p < .001$, $d = .66$, T-junctions. Thus, these results replicate and extend the results of Experiment 1a (see Figure S1 in the additional online materials at the OSF page).

Thus, these results replicate and extend the results of Experiments 1a and 1b, showing that the previously observed differences in choice are unlikely to be explained by any low-level confound in the stimulus set.

Experiment 1d: 3D Stimuli (Three to Five Segments)

The stimuli tested so far occupy a strange middle-ground: They are not quite like normal objects, but they are not quite networks either. Is there something about these letter-like stimuli that is influencing how participants respond? To address this possibility, we created a new stimulus set. Unlike the previous stimuli, these new stimuli were 3D rendered to have depth and shading like ordinary objects. Additionally, they were constructed to look more like networks in that they were given visible nodes and edges (see Figure 3E). The goal of this experiment is simply to determine if the previous patterns of results generalize to a novel stimulus set—to ensure that the results of the previous experiments are not due to any idiosyncrasies in the original stimuli.

Method

Everything about the design and procedure was identical to Experiments 1c, except that the stimuli were 3D renderings constructed in blender (see Figure 3E). The full stimulus set is available on our OSF page. Fifty new participants completed the task via Prolific. There were no exclusions.

Results and Discussion

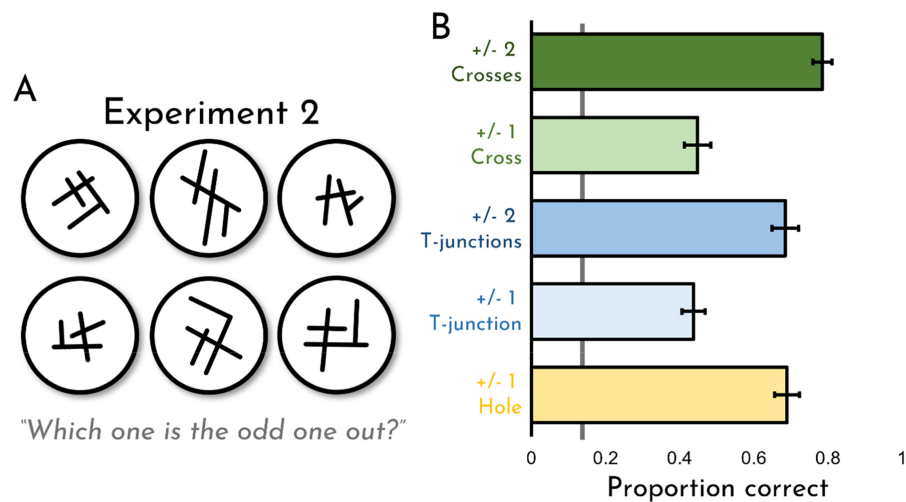
Participants were again well above-chance discriminating based on topological differences. This was true for all comparison types, ± 1

hole, $t(49) = 14.15$, $p < .001$, $d = 2.00$; ± 1 T-junction, $t(49) = 8.35$, $p < .001$, $d = 1.18$; ± 2 T-junctions, $t(49) = 9.68$, $p < .001$, $d = 1.37$; ± 3 T-junctions, $t(49) = 15.09$, $p < .001$, $d = 2.13$, and for all 12 unique comparisons ($p < .001$). These results are almost identical to the results of Experiment 1c. Therefore, behavior is unlikely to be explained by any idiosyncrasies of either stimulus set. This conclusion is not obvious. The stimuli used in this experiment differ in at least one crucial way from those used in the previous experiments. Because “nodes” in the networks are represented by spheres, T-junctions are a more visible (in that they are always marked by a separate sphere). While this should not influence how people interpret the topology of the networks, it may well have led participants to use various heuristics to identify the odd one out. The fact that behavior is largely consistent across the two stimulus sets indicates that participants are likely solving the problems in the same way in both cases—a further indication that participants are truly sensitive to the topological structure per se.

Experiment 2: More Complex Topologies

Experiments 1a, 1b, and 1c demonstrated a clear sensitivity to certain topological structures (T-junctions, holes). However, T-junctions and holes are not representative of all possible topologies. By constraining the stimulus sets to only these features, we limited the set of possible network topologies tested. Thus, the goal of Experiment 2 was to expand the range of topologies that we tested and to further increase the complexity of the stimuli. Rather than manipulating just T-junctions and holes, we added a third topological “part”: 4-point vertices, or “crosses” (see Figure 2A). An example trial can be seen in Figure 4A (and for more examples, all stimuli are available on our OSF page). Successful odd-one-out detection in this task would provide even stronger evidence for genuine sensitivity to topology.

Figure 4
Experiment 2



Note. Design (A) and results (B) of Experiment 2. Error bars represent ± 1 SE. As there are six options, chance performance is equal to 16.67%, indicated by the gray line. The correct answer in (A) is the bottom middle shape, as it has two crosses (like the other items) but no T-junctions (whereas the other items have two crosses and one T-junction). See the online article for the color version of this figure.

Method

Everything about the design and procedure was identical to the previous three experiments, except that (a) the stimuli were always made up of exactly four line segments, (b) they could consist of T-junctions, holes, and crosses, and (c) we used a new stimulus set that expanded on the stimuli used in Experiment 1b. Fifty new participants completed the task via Prolific. There were no exclusions.

For this task, we generated a new stimulus set from scratch. This consisted of 11 unique topologies, which were in terms of [crosses-T-junctions-holes]: [0-0-0], [0-0-1], [0-1-0], [0-1-1], [0-2-0], [0-2-1], [1-1-0], [1-0-1], [2-0-0], [2-0-1], and [2-1-0]. In this notation, the network 0-0-1 would be one with zero crosses, zero T-junctions, and a single hole (i.e., a quadrilateral). Given that there were six unique exemplars of each topology, this resulted in 66 unique items.

Again, not all possible topologies were compared against all others. Instead, we preselected 15 comparisons (six of which varied in their number of crosses, seven in their number of T-junctions, and three in their number of holes). These selections were made so that there were some instances in which there was only a single difference (e.g., a difference between 1 and 2 T-junctions), and some instances in which there were two differences (e.g., the comparison [2-0-1] vs. [0-0-1] would mean that the odd one out had two more or two fewer crosses than the other items). The trials in which holes differed were the only exception; there was never a case where the odd one out differed by more than one hole. This was done because (a) it is already clear from the prior data that people are highly sensitive to the presence of absence of holes, and (b) allowing for multiple holes needlessly complicates the stimuli.

The topological comparisons were the same for all participants, however the individual exemplars were selected randomly and so differed across participants. Each comparison was presented a total of 4 times, 2 times with one of the topologies in the majority, and 2 times with the other topology in the majority. Given that there were 15 possible comparisons, this resulted in a total of 60 total trials.

Results and Discussion

The results of this experiment can be seen in Figure 4B. Participants were again well above-chance discriminating based on topological differences. This was true for all comparison types, ± 1 hole, $t(49) = 15.78$, $p < .001$, $d = 2.23$; ± 1 T-junction, $t(49) = 8.73$, $p < .001$, $d = 1.23$; ± 2 T-junctions, $t(49) = 14.57$, $p < .001$, $d = 2.06$; ± 1 cross, $t(49) = 7.85$, $p < .001$, $d = 1.11$; ± 2 crosses, $t(49) = 23.80$, $p < .001$, $d = 3.37$. As before, differences in holes were relatively salient; participants were better at detecting differences in holes than differences in T-junctions, $t(49) = 10.69$, $p < .001$, $d = 1.51$, or crosses, $t(49) = 7.85$, $p < .001$, $d = 1.10$. These results replicate and extend the results of Experiment 1a–1d. They demonstrate that people are highly sensitive to several distinct topological features (i.e., holes, T-junctions, and crosses), even when the stimuli vary along multiple dimensions across the stimulus set.

Experiment 3: Memory

The experiments thus far have demonstrated that people are sensitive to topology, at least when they are explicitly asked to reason about objects' structural similarity. However, that adults can sometimes reason about topology may not be too surprising. After all,

we have established that topology is a kind of relational reasoning, and we know that adult participants have the capacity to reason about relations. So, here, we examined not just whether adults can discriminate based on topology, but whether they spontaneously represent objects based on their topology. Using the same sorts of stimuli we used in the previous task, we designed a memory task in which participants were shown six different items with two distinct topologies (three of each) and were given several seconds to encode them as well as possible. After a delay, participants saw a single test item and were asked to indicate whether that item had been present in the previous set of six. In one third of trials, the test item was one of the six items presented in the encoding phase. In the other two thirds of trials, the test item was new. Crucially, however, half of the time that the test item was new, it matched one of the two topologies present in the encoding phase; on the other half of trials, it was of a third distinct topology (see Figure 5A). We reasoned that if people represent objects in terms of their topology, they should be more likely to falsely indicate that they had seen an item that had a familiar topology, even if it had not been seen previously.

Method

Many of the display and design details are the same as in Experiment 1c, except as noted below. Overall, 50 new participants completed the task via Prolific. There were no exclusions.

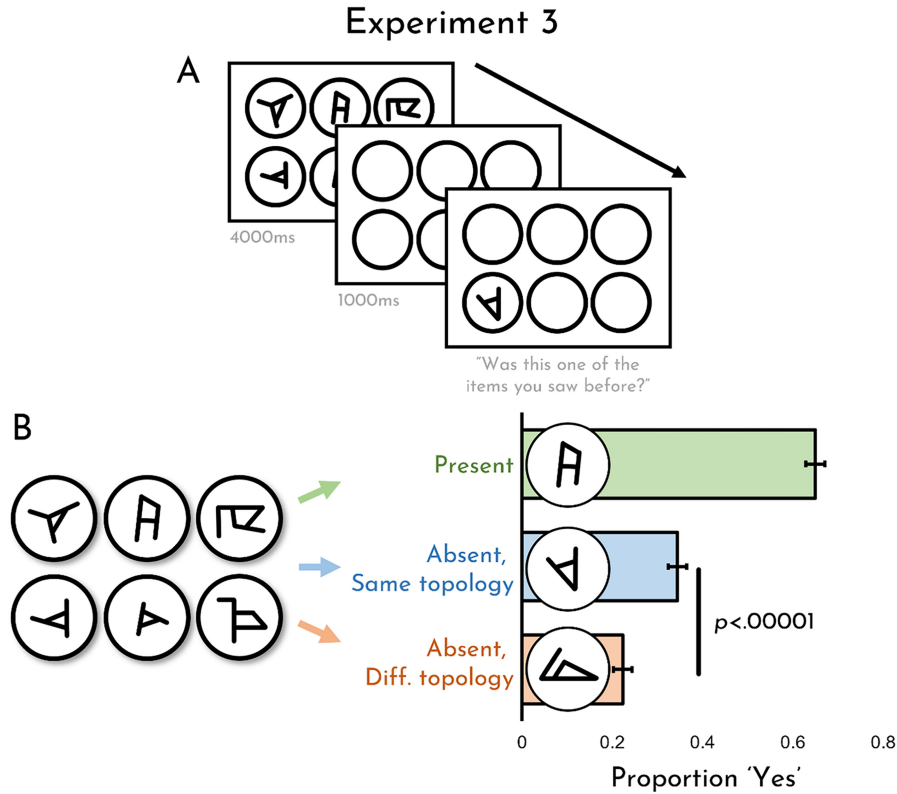
We used the stimuli from Experiment 1c (i.e., the full set of 126 stimuli comprises six exemplars for each of seven topologies with three, four, and five line segments). On each trial, participants were shown six random items, selected with the constraint that there were three items with each of two different topologies (randomly chosen) and an equal number of items with three, four, and five line segments (i.e., two of each). The only additional constraint was that for each topology there were two items with the same number of line segments and one with a different number. There were no constraints on which topologies were presented together. The six items were shown for 4 s for encoding, followed by a 1 s delay before a single test stimulus was presented. Participants were prompted to indicate whether the test item had been present in the encoding phase pressing "Y" for yes or "N" for no.

There were three possible trial types (see Figure 5B): "present" trials in which the test item had indeed been present in the encoding phase of the trial; "absent—same topology" trials in which the item was not present but matched one of the two topologies present in the encoding phase; and "absent—different topology" trials in which the item had not been presented and was of a third distinct topology not present in the encoding phase. There were 40 of each trial type, for a total of 120 trials.

Results and Discussion

The results of this experiment can be seen in Figure 5B. Unsurprisingly, participants were much more likely to say that they had seen an item before when they truly had, compared to both the absent—same topology, $t(49) = 12.09$, $p < .001$, $d = 1.71$, and absent—different topology, $t(49) = 14.84$, $p < .001$, $d = 2.10$, trials. Critically, however, people were also more likely to false alarm to new same-topology test items compared to different-topology test items, $t(49) = 7.20$, $p < .001$, $d = 1.02$. One possible deflationary explanation of the observed results is that they are caused by a specific subset

Figure 5
Experiment 3



Note. Design (A) and results (B) of Experiment 3. Error bars represent $\pm 1 SE$. See the online article for the color version of this figure.

of the trials (for instance, the ones in which there is a difference in holes). Intuitively, one could imagine how participants might false alarm to a stimulus with a hole if they had just seen six items with holes (and less likely to false alarm if there was no hole after they had just seen items with holes). To address this possibility, we separately analyzed only those trials in which there was no difference in holes (i.e., wherein all stimuli, at encoding and test, have the same number of holes). Still, participants were more likely to false alarms to same-topology test items, $t(49) = 5.34$, $p = .014$, $d = .76$.

This tendency to false alarm to items of the same topology suggests that people are not merely sensitive to topological structure when they are asked to be. Instead, people may be encoding these objects, at least to some extent, with respect to their topology—such that they are more likely to confuse two distinct items with the same topology. These findings therefore raise the intriguing possibility that topology is part of the “language” of spatial representation used to represent objects and spaces.

Experiment 4: Match-to-Sample Task

The previous two tasks demonstrated that people generally make discriminations based on topological structure and that they may remember objects partly with respect to their topology. But what if participants are asked to make explicit similarity judgments? Here, we had participants judge which of the two items was more like a sample item. One of the items was a topology-match: It had the

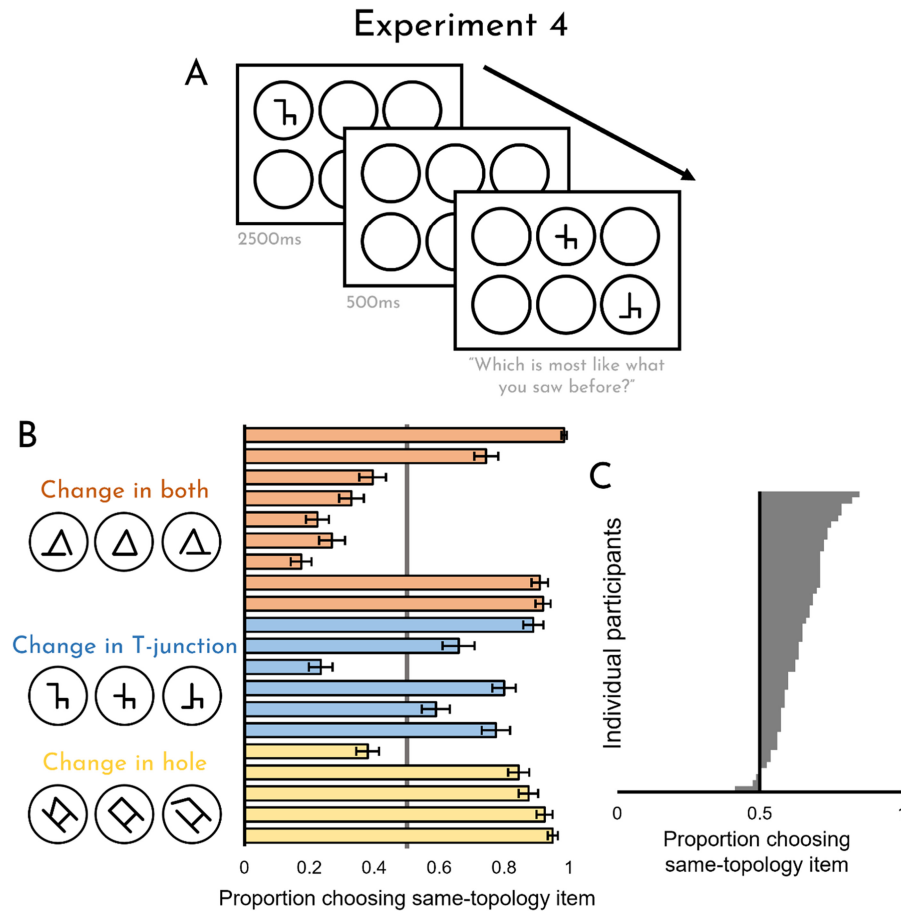
exact same topology as the sample item, but was different in some other specific way (e.g., via the shifting or rotation of an individual line segment). The other item was a topology-mismatch: It had a different topology from the sample item but was objectively more like the sample item than the topology-match was (for a more detailed explanation, see Method). In other words, we created a tension between surface-level similarity and topological similarity: One item was objectively more similar to the sample (e.g., in terms of the physical distance that one segment moved), but the other item was matched in terms of its topology (the topological matches were always twice as different from the sample in terms of objective change; see Figure 6A and B). Simply, we asked whether people might more readily match items based on topology rather than surface-level similarity.

Method

Many of the display and design details are the same as in Experiment 3, except as noted below. Fifty new participants completed the task via Prolific. There were no exclusions.

We created a new stimulus set where each trial consisted of (a) a sample item, (b) a topology-mismatch (a variation of the sample item with a fixed change, via translation or rotation, that disrupted topology), and (c) a topology-match (a variation of the sample item with twice as much physical change as the topology-mismatch but which maintained the sample topology as the sample). In other words, the

Figure 6
Experiment 4



Note. Design (A) and results (B–C) of Experiment 4. Displayed here is the proportion of participants matching on topology, broken down by item (B) and participant (C). Error bars represent ± 1 SE. See the online article for the color version of this figure.

topology-mismatch was more physically similar to the sample than the topology-match was, but the topology-match retained the sample's topological form. For instance, in Figure 6A, the sample item is shown in the top-left panel. The two variations are shown in the bottom-right panel. The topology-mismatch was created by moving the leftward horizontal line halfway down the vertical line. This topology-match was created by moving that same leftward horizontal line fully down the vertical line. This topology-match involved that same line moving by twice as much, but the result was an item that shared the topology of the sample. For all 20 unique items that we created, the topology-match always differed from the sample by twice as much as the topology-mismatch in exactly this way. The stimuli were designed by a research assistant (with instructions from the authors) who had no knowledge of the hypotheses.

On each trial, the sample (one of 20 items) would appear in one of the six cells of a 3×2 grid (we recycled the grid display from the previous experiments to make the task slightly more interesting; see Figure 6A). It would appear for 2.5 s. After a 500 ms delay, two test items were presented in two other random cells of the 3×2

grid. Participants were prompted to indicate by mouse click "which of the images [was] most like the one [they] saw before." Response was registered via brief visual feedback regardless of their choice.

Because the items were mirrored (the topology-mismatch was always the midpoint between the other two items), the sample item that was presented could be either the default item that we generated or the corresponding, same-topology item. The order was counterbalanced. As there were 20 unique items, two orders in which the items could appear, and two repetitions of every trial type, there were a total of 80 trials.

Results and Discussion

The results of this experiment can be seen in Figure 6, broken down by item (Figure 6B) and participant (Figure 6C). Participants overwhelmingly matched based on object topology at the item level, about 64% of the time; $t(19) = 2.23$, $p = .038$, $d = .50$, as well as the participant level, $t(49) = 10.84$, $p < .001$, $d = 1.53$. This was independently true for those trials in which there was a change in the presence of holes, $t(49) = 18.31$,

$p < .001$, $d = 2.59$; T-junctions, $t(49) = 6.28$, $p < .001$, $d = .89$; and both, $t(49) = 3.15$, $p = .003$, $d = .45$. There were however some items for which participants seemed to have preferred the more physically similar object to the topological match. These were primarily the items that involved rotation. As one can see in the items for themselves, the degree of rotation was (intentionally) large, and this created situations where, for instance, a line would rotate through another line in the image to the other side (see the additional online materials at the OSF page).

That people sometimes matched based on surface features does not undermine the overall finding that on average people match based on topology. After all, the topology-matched items were always objectively twice as different as the topology-mismatched items. Despite this, over 90% of participants were selected based on topology more often than not (see Figure 6C). In any case, topology is surely not the only way that people categorize or represent objects; it is one of many. Curvature, for instance, is topologically irrelevant, yet nobody would argue that people are not sensitive to curvature (or that, in a different task, if people occasionally categorized based on curvature, this would provide evidence against the relevance of topology). It is telling that (the vast majority of) participants consistently categorized based on topology rather than surface features, even if there are some exceptions.

General Discussion

Here, we have explored whether and how people are sensitive to topological structures. In the first set of experiments, we showed that participants can discriminate between objects based on their topology, even when given no instructions about how they were meant to discriminate (Experiments 1a–1d). Next, we showed that this ability to discriminate items based on their topologies extends to stimuli with more complex topologies (Experiment 2). In all four experiments, we found not only that people on average discriminated based on topology, but that the likelihood of identifying the odd one out scaled with the number of topological changes. In a fifth experiment (Experiment 3), we showed that people's memories are influenced by topology: People are more likely to falsely report having seen a novel item that matches the topology of a familiar item than a novel item with a novel topology. Finally, we showed that people are more likely to explicitly match objects based on their topology rather than their surface features, to the extent that people will select based on topology even when the other option is objectively twice as similar to the sample (Experiment 4). Collectively, these results suggest that, despite the fact that human adults rarely reason explicitly in topological terms, they nevertheless act on an intuitive notion of network topology.

Perceiving Relations

Topology is, ultimately, about relations—how the pieces of space connect to one another. This is why topological maps are intuitive: They capture the relevant relations without the distraction of spatial detail. Reasoning about relations has a long history in cognitive science, in perception (Franconeri et al., 2012; Kim & Biederman, 2012; Kranjec et al., 2014; Roth & Franconeri, 2012; for review, see Hafri & Firestone, 2021), language (Bowerman & Choi, 2003; Landau, 2017, 2020; Landau & Jackendoff, 1993), and social cognition (Ding et al., 2017; Papeo & Abassi, 2019; Papeo et al., 2017). In

fact, virtually all the work cited here cuts across at least two of those disciplines, speaking to the centrality of relations in perception and cognition.

Typically, the sorts of relations that are studied in cognitive science are ones with corresponding words (e.g., above, next to, push, etc.). The topological relations studied here, in contrast, are ones that are less well represented in language. Of course, topological relations can be captured by language. We do have a word for “hole,” and here we have assigned the terms “T-junctions” and “crosses” to refer to 3-point vertices and 4-point vertices, respectively. Yet these latter terms are not ones that we would readily use to describe these structures. Most people encountering a three-way versus four-way intersection would likely describe it as (just) an intersection—even though we clearly interact with these two types of intersections quite differently. Even the ordinary concept of a “hole” is tricky. Is the “hole” in the letter “R” the same as a bullet hole? Not quite: The latter we think of as an intrusion on an otherwise complete object, whereas the former is a part of the object. Thus, while the topological features we have studied here are clearly intuitive on some level (people readily discriminate on their basis; people's memories are influenced by them; etc.), they are not straightforwardly captured by ordinary language. This further supports the possibility that, like “core” geometric properties such as symmetry, centrality, and curvature, the topological knowledge studied here may develop independently from language.

What we have shown here is that people are sensitive to these topological features in some form. This is not to say that people are treating the letter-like stimuli used here as networks themselves. Though our stimuli bear some resemblance to the structure of subway maps—as they were designed with such maps in mind—it seems more likely that people are viewing them as objects rather than network structures per se. An open question remains about whether there is in fact a common topological representation underlying the representation of certain objects as well as certain large-scale spaces.

Alternative Explanations

We have shown that adults are generally sensitive to topological structure—but why? One obvious possibility is that adults only appear to be sensitive to topology because topological changes result in other relevant changes. Differences in topology may influence perceived complexity, for instance. Prior work has shown that the complexity of figures influences attention (Sun & Firestone, 2021), encoding in language (Sun & Firestone, 2022b), and even aesthetic preferences (Sun & Firestone, 2022a). Just glancing at the stimuli used here, there is no doubt that complexity varies across different topological forms to some extent.

However, it is not clear how the complexity of topological forms should be measured. There are numerous ways of characterizing psychological and visual complexity (for review, see Donderi, 2006), none of which seems apt to measure the complexity of figures like these. Consider one metric of complexity used for describing visual forms: “structural surprisal” (see Feldman & Singh, 2006; Sun & Firestone, 2021). This metric is meant to capture the cumulative “surprisal” of each branch of a skeletal figure; figures with more branches with more twists and turns would be more “surprising” and therefore more complex. This measure alone, however, cannot explain how a lower “t” differs from a capital “T,” or how a capital

“T” differs from a capital “L.” Structural surprisal is designed for application to skeletal figures for which the curvature of every line matters. Of course, curvature is irrelevant in the world of topology.

Nevertheless, it should not come as a surprise that topology and complexity are related to some extent. Topology is meant to capture relations, after all, and relations can vary in their complexity. What we mean to say is not that topology is unrelated to complexity, but rather that the ways we would typically operationalize complexity are unlikely to explain the results observed here (both in principle and in practice). For instance, the stimuli used in Experiments 1c, 1d, 2, and 3 vary dramatically in their surface complexity, despite only subtle changes in topology. The stimuli used in Experiment 4 were designed to equate most surface features (i.e., they use the same number of lines, the lines are equated for length, etc.), and thus should have comparable complexity as well. To reduce these results to “mere” differences in complexity requires formulations of complexity that largely recapitulate the topological structure (see Bonchev & Buck, 2005).

Topology has also been considered one possible way of capturing the large-scale spatial representations that support human navigation (see, e.g., Coutrot et al., 2022; R. A. Epstein et al., 2017). As of yet, however, this has been little more than speculation. There have been no detailed investigations into whether or how topological relations for large-scale spaces may be represented. Yet there is some prior work in the general field of spatial navigation that provides hints at the use of topological representation. For instance, when people are asked to recreate mental maps of familiar areas, they are consistently biased to draw turns as if they were 90° angles (even those that differed substantially from 90°; see Byrne, 1979; Moar & Bower, 1983). This suggests people represent the presence of an intersection, but not the magnitude of the turn itself—a tendency that makes sense through the lens of network topology.

The present work may then serve as a blueprint, of sorts, to investigate topological representations on larger scales. Of particular interest may be how topological forms of representation are related to other forms of representation, like coordinate-based formats (see, e.g., Yousif, 2022; Yousif & Keil, 2021; Yousif et al., 2023).

Topology, “Core Knowledge,” and a “Language of Thought”

Core knowledge theory posits that human cognition is founded on few specific, discrete systems (e.g., objects, actions, numbers, and space) that have deep phylogenetic roots (Carey, 2009; E. S. Spelke, 2000; E. S. Spelke & Kinzler, 2007). Topology itself may not be a kind of core knowledge, but it interestingly cuts across several possible domains of core knowledge, like space/geometry (see Dehaene et al., 2006; Izard et al., 2011; E. Spelke et al., 2010), objects (see Baillargeon, 1995; Baillargeon et al., 1985), and even number (see, e.g., He et al., 2015). As topology offers a discrete, low-dimensional means of representing how things relate to one another, it may be at the heart of how we represent not only spaces and objects, but also social networks, institutional structures, and more.

Perhaps topology offers a sort of “language of thought” (see Fodor, 1975; for a recent review, see Quilty-Dunn et al., 2022) for representing relations. While the topological structure described here lacks some of the properties we might demand of a true “language of thought” (e.g., a subject–predicate structure; for discussion, see Camp, 2009), it has some others. For instance, it has

compositional semantics: Topological networks are described by (a) what parts they contain and (b) how those parts are combined (see Clarke, 2023). Therefore, from a few building blocks (holes, T-junctions, crosses, etc.) arise an infinite number of geometric/relational forms. Thus, “intuitive topology” may be about more than topology itself; it may provide a window into the structure of thought more generally. Insofar as we have shown here that people are surprisingly sensitive to topological structure (and that people remember and match items based on their topology), there is reason to further examine whether this sort of thinking is related to relational representation more broadly.

There has been recent interest in the notion of a “language of thought” for geometric representation (see Dehaene et al., 2022; Sablé-Meyer et al., 2021, 2022; see also Al Roumi et al., 2021; Amalric et al., 2017). One hypothetical “language” uses repetition, concatenation, and embedding to generate the sorts of basic shapes that have been observed in human cultures for tens of thousands of years (Sablé-Meyer et al., 2022). The promise of this research program is obvious: It offers to describe the fundamental units of spatial representation.

A related controversy is whether there is a pictorial syntax or “grammar” that can be used to describe visual images without language (Lande, 2024). This relates to fundamental issues in vision science, such as how shapes are perceived constantly across viewpoints and how geometric forms are represented (see, e.g., W. Epstein & Park, 1963; Green, 2023; Slater & Morison, 1985). For example, how do we know that a hand continues to be a hand, even when we view it from many different perspectives? This is a question about the format of the underlying representation (see Yousif, 2022). One proposal is that the visual system uses “shape skeletons” to represent the internal structure of shapes. A substantial body of work demonstrates that the visual system is indeed sensitive to the skeletal structure of shapes, in particular the medial axis (see, e.g., Ayzenberg & Lourenco, 2019, 2022; Ayzenberg et al., 2019, 2022; Firestone & Scholl, 2014). Shape skeletons stop short of providing a complete language of forms, however. There exists a deeper question: How is the structure of a shape skeleton itself represented? Topology—or graph theory more generally—may offer some answers.

There are, in other words, many pervasive questions about the “language” (or languages; see Green, 2023) used to efficiently represent visual information. We suggest that while the network topological “language” we target here is unlikely to be sufficient to describe all visual information, it may be a part of a suite of syntaxes used by the visual system to represent basic visual relations.

Topology abstracts away from spatial detail like curvature, or jagged lines, and so on (structures that feature prominently in the language proposed by Sablé-Meyer and colleagues, for instance). Thus, topology is not a replacement for other “languages” that have been proposed. Yet the existence of topological maps (see Figure 1D) is a testament to the value of a simplified spatial language. Extracting away from metric detail and focusing instead on relations ironically improves our ability to process the relevant information. Perhaps topology acts as a sort of “assembly language” on which other languages build. One could imagine a primitive topological representation of a city map, for instance, on which precise details are gradually superimposed.

As an “assembly language,” topology has many virtues: (a) its parts (holes, T-junctions, crosses) are few; (b) its parts are discrete and clearly defined (i.e., there is no such thing as a large or a

small hole; there is just “hole”); (c) it has compositional semantics (such that complex entities arise out of simple parts); (d) it is infinitely generative; and (e) it is maximally domain-general (i.e., it is not just a description of space, it is just as well as description of any relational system).

This is not to say that topological representation in the mind perfectly mirrors topological representation in the external world. Just because networks can be described in a topological language does not mean that the mind necessarily represents them in this way. What we mean is that network topology, construed this way, has features that make it efficient for representing complex structures and that, for this reason, the mind may rely on a representational language that is similar in nature. Given the way that relational thinking may pervade human cognition—the way that topology may act as a sort of relational “language” which is key to several domains of “core knowledge”—these questions are worth exploring further.

If these topological building blocks are a part of a functional language of thought in the mind, a natural question arises about the development of this language: To what extent is the sort of basic topological thinking intuitive to young children—or even infants? A few studies have explored sensitivity to topological properties of objects (e.g., closure, overlap) in infancy (Chien et al., 2012; Kibbe & Leslie, 2016) in early childhood (Kenderla et al., 2023) and in animal models (Chen et al., 2003), and there is consensus that at least some aspects of topology are a vital part of object representation. A question remains about whether and/or how the network-like topological relations studied here are related to object topology. Investigations of both object and network topology—in adults as well as children—may similarly provide clues to the foundations of mental representation.

Constraints on Generality

These studies were conducted on a diverse sample of adults from the United States. Thus, we are not assuming that these findings generalize beyond this group. An open question remains about whether these findings would generalize to individuals without formal training in mathematics and geometry (see, e.g., Dehaene et al., 2006).

Conclusion

Here, we have shown that people intuitively understand the language of topological networks: They readily identify items that differ in their topology from other topology-matched items; they match items based on topology rather than surface-level similarity; and they even remember, in part, with respect to topological structure. These results speak to a means of relational representation that may be foundational not only to spatial representation, but mental representation more broadly. This intuitive sense of network topology may be a window to understanding the mind’s natural language for representing relations.

References

- Al Roumi, F., Marti, S., Wang, L., Amalric, M., & Dehaene, S. (2021). Mental compression of spatial sequences in human working memory using numerical and geometrical primitives. *Neuron*, 109(16), 2627–2639.e4. <https://doi.org/10.1016/j.neuron.2021.06.009>
- Amalric, M., Wang, L., Pica, P., Figueira, S., Sigman, M., & Dehaene, S. (2017). The language of geometry: Fast comprehension of geometrical primitives and rules in human adults and preschoolers. *PLoS Computational Biology*, 13(1), Article e1005273. <https://doi.org/10.1371/journal.pcbi.1005273>
- Ayzenberg, V., Chen, Y., Yousif, S. R., & Lourenco, S. F. (2019). Skeletal representations of shape in human vision: Evidence for a pruned medial axis model. *Journal of Vision*, 19(6), Article 6. <https://doi.org/10.1167/19.6.6>
- Ayzenberg, V., Kamps, F. S., Dilks, D. D., & Lourenco, S. F. (2022). Skeletal representations of shape in the human visual cortex. *Neuropsychologia*, 164, Article 108092. <https://doi.org/10.1016/j.neuropsychologia.2021.108092>
- Ayzenberg, V., & Lourenco, S. F. (2019). Skeletal descriptions of shape provide unique perceptual information for object recognition. *Scientific Reports*, 9(1), Article 9359. <https://doi.org/10.1038/s41598-019-45268-y>
- Ayzenberg, V., & Lourenco, S. F. (2022). Perception of an object’s global shape is best described by a model of skeletal structure in human infants. *eLife*, 11, Article e74943. <https://doi.org/10.7554/eLife.74943>
- Baillargeon, R. (1995). Physical reasoning in infancy. In M. S. Gazzaniga (Ed.), *The cognitive neurosciences* (pp. 181–204). MIT Press.
- Baillargeon, R., Spelke, E. S., & Wasserman, S. (1985). Object permanence in five-month-old infants. *Cognition*, 20(3), 191–208. [https://doi.org/10.1016/0010-0277\(85\)90008-3](https://doi.org/10.1016/0010-0277(85)90008-3)
- Bonchev, D., & Buck, G. A. (2005). Quantitative measures of network complexity. In D. Bonchev & D. H. Rouvray (Eds.), *Complexity in chemistry, biology, and ecology* (pp. 191–235). Springer US.
- Bowerman, M., & Choi, S. (2003). Space under construction: Language-specific spatial categorization in first language acquisition. In D. Gentner & S. Goldin-Meadow (Eds.), *Language in mind: Advances in the study of language and thought* (pp. 387–427). MIT Press.
- Byrne, R. W. (1979). Memory for urban geography. *The Quarterly Journal of Experimental Psychology*, 31(1), 147–154. <https://doi.org/10.1080/14640747908400714>
- Camp, E. (2009). A language of baboon thought? In R. Lurz (Ed.), *The philosophy of animal minds* (pp. 108–127). Cambridge University Press.
- Carey, S. (2009). *The origin of concepts*. Oxford University Press.
- Chen, L. (1982). Topological structure in visual perception. *Science*, 218(4573), 699–700. <https://doi.org/10.1126/science.7134969>
- Chen, L. (1985). Topological structure in the perception of apparent motion. *Perception*, 14(2), 197–208. <https://doi.org/10.1068/p140197>
- Chen, L. (1990). Holes and wholes: A reply to Rubin and Kanwisher. *Perception & Psychophysics*, 47(1), 47–53. <https://doi.org/10.3758/BF03208163>
- Chen, L. (2005). The topological approach to perceptual organization. *Visual Cognition*, 12(4), 553–637. <https://doi.org/10.1080/13506280444000256>
- Chen, L., Zhang, S., & Srinivasan, M. V. (2003). Global perception in small brains: Topological pattern recognition in honey bees. *Proceedings of the National Academy of Sciences*, 100(11), 6884–6889. <https://doi.org/10.1073/pnas.0732090100>
- Chien, S. H. L., Lin, Y. L., Qian, W., Zhou, K., Lin, M. K., & Hsu, H. Y. (2012). With or without a hole: Young infants’ sensitivity for topological versus geometric property. *Perception*, 41(3), 305–318. <https://doi.org/10.1068/p7031>
- Clarke, S. (2023). Compositionality and constituent structure in the analogue mind. *Philosophical Perspectives*, 37(1), 90–118. <https://doi.org/10.1111/phpe.12182>
- Coutrot, A., Manley, E., Goodroe, S., Gahnstrom, C., Filomena, G., Yesiltepe, D., Dalton, R. C., Wiener, J. M., Hölscher, C., Hornberger, M., & Spiers, H. J. (2022). Entropy of city street networks linked to future spatial navigation ability. *Nature*, 604(7904), 104–110. <https://doi.org/10.1038/s41586-022-04486-7>
- Dehaene, S., Al Roumi, F., Lakretz, Y., Planton, S., & Sablé-Meyer, M. (2022). Symbols and mental programs: A hypothesis about human singularity. *Trends in Cognitive Sciences*, 26(9), 751–766. <https://doi.org/10.1016/j.tics.2022.06.010>
- Dehaene, S., Izard, V., Pica, P., & Spelke, E. (2006). Core knowledge of geometry in an Amazonian indigene group. *Science*, 311(5759), 381–384. <https://doi.org/10.1126/science.1121739>

- Ding, X., Gao, Z., & Shen, M. (2017). Two equals one: Two human actions during social interaction are grouped as one unit in working memory. *Psychological Science*, 28(9), 1311–1320. <https://doi.org/10.1177/0956797617707318>
- Donderi, D. C. (2006). Visual complexity: A review. *Psychological Bulletin*, 132(1), 73–97. <https://doi.org/10.1037/0033-2909.132.1.73>
- Epstein, R. A., Patai, E. Z., Julian, J. B., & Spiers, H. J. (2017). The cognitive map in humans: Spatial navigation and beyond. *Nature Neuroscience*, 20(11), 1504–1513. <https://doi.org/10.1038/nn.4656>
- Epstein, W., & Park, J. N. (1963). Shape constancy: Functional relationships and theoretical formulations. *Psychological Bulletin*, 60(3), 265–288. <https://doi.org/10.1037/h0040875>
- Feldman, J., & Singh, M. (2006). Bayesian estimation of the shape skeleton. *Proceedings of the National Academy of Sciences*, 103(47), 18014–18019. <https://doi.org/10.1073/pnas.0608811103>
- Firestone, C., & Scholl, B. J. (2014). “Please tap the shape, anywhere you like” shape skeletons in human vision revealed by an exceedingly simple measure. *Psychological Science*, 25(2), 377–386. <https://doi.org/10.1177/0956797613507584>
- Fodor, J. A. (1975). *The language of thought* (Vol. 5). Harvard University Press.
- Franconeri, S. L., Scimeca, J. M., Roth, J. C., Helseth, S. A., & Kahn, L. E. (2012). Flexible visual processing of spatial relationships. *Cognition*, 122(2), 210–227. <https://doi.org/10.1016/j.cognition.2011.11.002>
- Green, E. J. (2023). A pluralist perspective on shape constancy. *The British Journal for the Philosophy of Science*, 58(1), <https://doi.org/10.1086/727427>
- Hafri, A., & Firestone, C. (2021). The perception of relations. *Trends in Cognitive Sciences*, 25(6), 475–492. <https://doi.org/10.1016/j.tics.2021.01.006>
- He, L., Zhou, K., Zhou, T., He, S., & Chen, L. (2015). Topology-defined units in numerosity perception. *Proceedings of the National Academy of Sciences*, 112(41), E5647–E5655. <https://doi.org/10.1073/pnas.1512408112>
- Izard, V., Pica, P., Spelke, E., & Dehaene, S. (2011). Flexible intuitions of Euclidean geometry in an Amazonian indigene group. *Proceedings of the National Academy of Sciences*, 108(24), 9782–9787. <https://doi.org/10.1073/pnas.1016686108>
- Kanbe, F. (2013). On the generality of the topological theory of visual shape perception. *Perception*, 42(8), 849–872. <https://doi.org/10.1068/p7497>
- Kenderla, P., Kim, S. H., & Kibbe, M. M. (2023). Competition between object topology and surface features in children’s extension of novel nouns. *Open Mind*, 7, 93–110. https://doi.org/10.1162/opmi_a_00073
- Kibbe, M. M., & Leslie, A. M. (2016). The ring that does not bind: Topological class in infants’ working memory for objects. *Cognitive Development*, 38, 1–9. <https://doi.org/10.1016/j.cogdev.2015.12.001>
- Kim, J. G., & Biederman, I. (2012). Greater sensitivity to nonaccidental than metric changes in the relations between simple shapes in the lateral occipital cortex. *NeuroImage*, 63(4), 1818–1826. <https://doi.org/10.1016/j.neuroimage.2012.08.066>
- Kranjec, A., Lupyan, G., & Chatterjee, A. (2014). Categorical biases in perceiving spatial relations. *PLoS ONE*, 9(5), Article e98604. <https://doi.org/10.1371/journal.pone.0098604>
- Landau, B. (2017). Update on “what” and “where” in spatial language: A new division of labor for spatial terms. *Cognitive Science*, 41(S2), 321–350. <https://doi.org/10.1111/cogs.12410>
- Landau, B. (2020). Learning simple spatial terms: Core and more. *Topics in Cognitive Science*, 12(1), 91–114. <https://doi.org/10.1111/tops.12394>
- Landau, B., & Jackendoff, R. (1993). Whence and whither in spatial language and spatial cognition? *Behavioral and Brain Sciences*, 16(2), 255–265. <https://doi.org/10.1017/S0140525X00029927>
- Lande, K. J. (2024). Pictorial syntax. *Mind & Language*. Advance online publication. <https://doi.org/10.1111/mila.12497>
- Lowet, A. S., Firestone, C., & Scholl, B. J. (2018). Seeing structure: Shape skeletons modulate perceived similarity. *Attention, Perception, & Psychophysics*, 80(5), 1278–1289. <https://doi.org/10.3758/s13414-017-1457-8>
- Moar, I., & Bower, G. H. (1983). Inconsistency in spatial knowledge. *Memory & Cognition*, 11(2), 107–113. <https://doi.org/10.3758/BF03213464>
- Papeo, L., & Abassi, E. (2019). Seeing social events: The visual specialization for dyadic human–human interactions. *Journal of Experimental Psychology: Human Perception and Performance*, 45(7), 877–888. <https://doi.org/10.1037/xhp0000646>
- Papeo, L., Stein, T., & Soto-Faraco, S. (2017). The two-body inversion effect. *Psychological Science*, 28(3), 369–379. <https://doi.org/10.1177/0956797616685769>
- Quilty-Dunn, J., Porot, N., & Mandelbaum, E. (2022). The best game in town: The re-emergence of the language of thought hypothesis across the cognitive sciences. *Behavioral and Brain Sciences*, 46, Article e261. <https://doi.org/10.1017/S0140525X22002849>
- Renz, J., Rauh, R., & Knauff, M. (2000, November). Towards cognitive adequacy of topological spatial relations. In C. Freksa, W. Brauer, C. Habel, & K. F. Wender (Eds.), *Spatial cognition II: Integrating abstract theories, empirical studies, formal methods, and practical applications* (pp. 184–197). Springer.
- Rips, L. J. (2020). Possible objects: Topological approaches to individuation. *Cognitive Science*, 44(11), Article e12916. <https://doi.org/10.1111/cogs.12916>
- Roth, J. C., & Franconeri, S. L. (2012). Asymmetric coding of categorical spatial relations in both language and vision. *Frontiers in Psychology*, 3, Article 464. <https://doi.org/10.3389/fpsyg.2012.00464>
- Rubin, J. M., & Kanwisher, N. (1985). Topological perception: Holes in an experiment. *Perception & Psychophysics*, 37(2), 179–180. <https://doi.org/10.3758/BF03202856>
- Sablé-Meyer, M., Ellis, K., Tenenbaum, J., & Dehaene, S. (2022). A language of thought for the mental representation of geometric shapes. *Cognitive Psychology*, 139, Article 101527. <https://doi.org/10.1016/j.cogpsych.2022.101527>
- Sablé-Meyer, M., Fagot, J., Caparos, S., van Kerkhove, T., Amalric, M., & Dehaene, S. (2021). Sensitivity to geometric shape regularity in humans and baboons: A putative signature of human singularity. *Proceedings of the National Academy of Sciences*, 118(16), Article e2023123118. <https://doi.org/10.1073/pnas.2023123118>
- Slater, A., & Morison, V. (1985). Shape constancy and slant perception at birth. *Perception*, 14(3), 337–344. <https://doi.org/10.1068/p140337>
- Spelke, E., Lee, S. A., & Izard, V. (2010). Beyond core knowledge: Natural geometry. *Cognitive Science*, 34(5), 863–884. <https://doi.org/10.1111/j.1551-6709.2010.01110.x>
- Spelke, E. S. (2000). Core knowledge. *The American Psychologist*, 55(11), 1233–1243. <https://doi.org/10.1037/0003-066X.55.11.1233>
- Spelke, E. S., & Kinzler, K. D. (2007). Core knowledge. *Developmental Science*, 10(1), 89–96. <https://doi.org/10.1111/j.1467-7687.2007.00569.x>
- Sun, Z., & Firestone, C. (2021). Curious objects: How visual complexity guides attention and engagement. *Cognitive Science*, 45(4), Article e12933. <https://doi.org/10.1111/cogs.12933>
- Sun, Z., & Firestone, C. (2022a). Beautiful on the inside: Aesthetic preferences and the skeletal complexity of shapes. *Perception*, 51(12), 904–918. <https://doi.org/10.1177/03010066221124872>
- Sun, Z., & Firestone, C. (2022b). Seeing and speaking: How verbal “description length” encodes visual complexity. *Journal of Experimental Psychology: General*, 151(1), 82–96. <https://doi.org/10.1037/xge0001076>
- Wei, N., Zhou, T., Zhang, Z., Zhuo, Y., & Chen, L. (2019). Visual working memory representation as a topological defined perceptual object. *Journal of Vision*, 19(7), Article 12. <https://doi.org/10.1167/19.7.12>
- Yousif, S. R. (2022). Redundancy and reducibility in the formats of spatial representations. *Perspectives on Psychological Science*, 17(6), 1778–1793. <https://doi.org/10.1177/17456916221077115>
- Yousif, S. R., & Brannon, E. M. (2024, March 6). *Intuitive network topology*. <https://osf.io/w8p6s>

- Yousif, S. R., Forrence, A. D., & McDougle, S. D. (2023). A common format for representing spatial location in visual and motor working memory. *Psychonomic Bulletin & Review*. <https://doi.org/10.3758/s13423-023-02366-3>
- Yousif, S. R., & Keil, F. C. (2021). The shape of space: Evidence for spontaneous but flexible use of polar coordinates in visuospatial representations. *Psychological Science*, 32(4), 573–586. <https://doi.org/10.1177/0956797620972373>
- Zhou, K., Luo, H., Zhou, T., Zhuo, Y., & Chen, L. (2010). Topological change disturbs object continuity in attentive tracking. *Proceedings of the National Academy of Sciences*, 107(50), 21920–21924. <https://doi.org/10.1073/pnas.1010919108>

Received September 1, 2023

Revision received March 19, 2024

Accepted March 29, 2024 ■