

Theory is Shapes

Matthew Varona, Maryam Hedayati, Matthew Kay, and Carolina Nobre

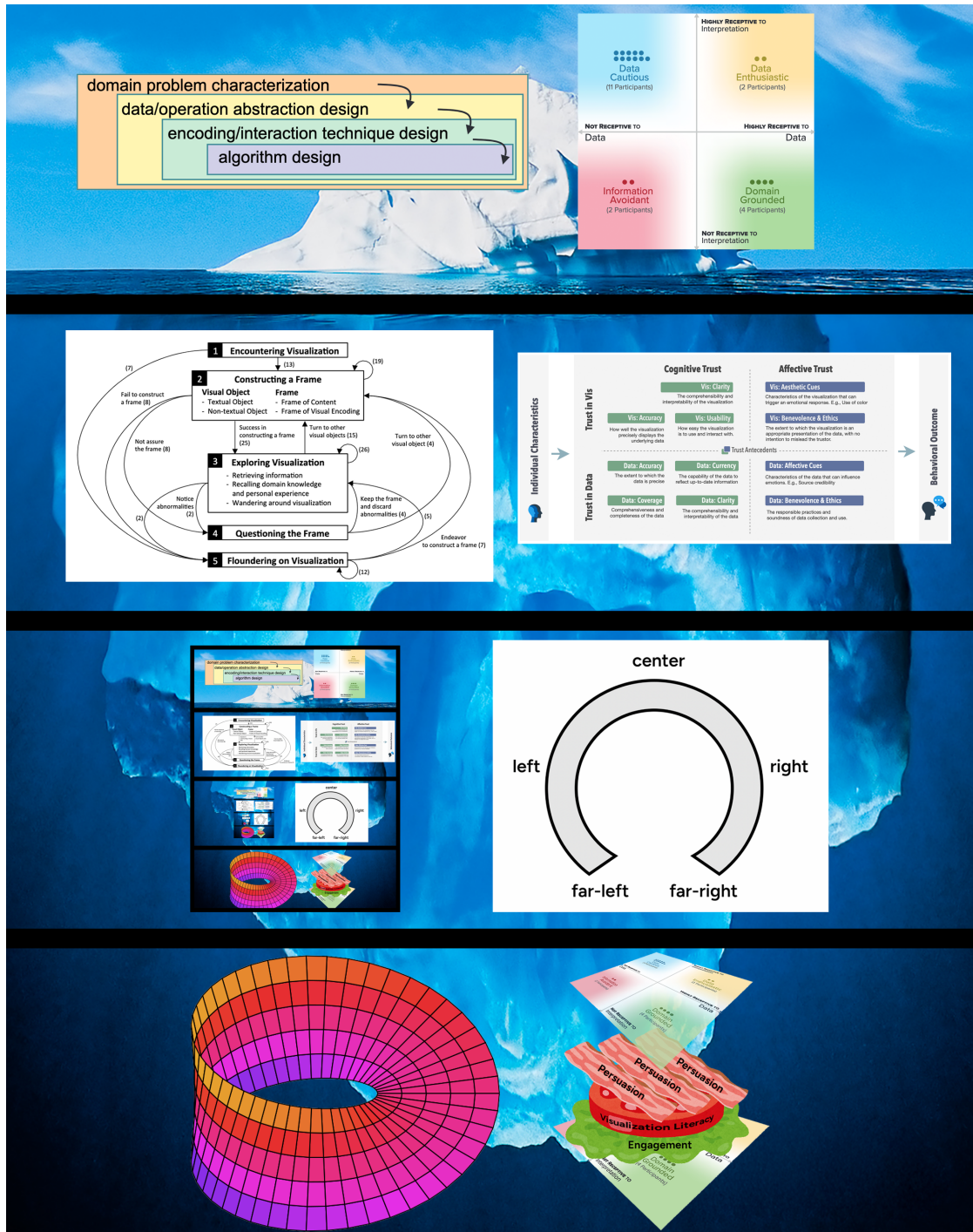


Fig. 1: The iceberg of visualization theory figures.

Abstract—"Theory figures" are a staple of theoretical visualization research. Common shapes such as Cartesian planes and flowcharts can be used not only to explain conceptual contributions, but to think through and refine the contribution itself. Yet, theory figures tend to be limited to a set of standard shapes, limiting the creative and expressive potential of visualization theory. In this work, we explore how the shapes used in theory figures afford different understandings and explanations of their underlying phenomena. We speculate on the value of visualizing theories using more expressive configurations, such as icebergs, horseshoes, Möbius strips, and BLT sandwiches. By reflecting on figure-making's generative role in the practice of theorizing, we conclude that theory is, in fact, shapes.

Index Terms—Theory, shapes, sandwiches

1 INTRODUCTION

Theory gives shape to scholarly work. It can condense messy phenomena into structured explanations, abstracting away details in order to highlight specific patterns or relationships. In visualization research, conceptual contributions like models, frameworks, and taxonomies underpin how we design systems and study users. Theory is not just textual; it is often visual, too. Figures that depict a theory’s components and relationships—what we call “theory figures”—aid researchers in both communicating and constructing theoretical contributions. Such figures are common in systematic literature reviews, qualitative research, or any other work that presents conceptual contributions.

Well-chosen figures can do the work of an entire page of exposition [5]. A Cartesian plane invites the reader to imagine a two-dimensional space of variation, with distinct quadrants and axes. Similarly, a flowchart can depict complex processes or dependencies at a glance. In doing the heavy lifting to explain a theory, shapes can also become synonymous with the theory itself, helping it stick in the minds of readers. When designing a figure, we are not merely illustrating a theory; we are, in a very real sense, thinking through shapes.

Yet, figures in visualization research rely heavily on a narrow canon of shapes—grids, flowcharts, matrices, and the occasional Euler diagram.¹ Reliance on these familiar shapes can limit the creative potential of visualization theory. What might we learn if we embraced the value of more expressive shapes, such as Möbius strips, icebergs, and BLT sandwiches? Could these shapes afford new ways of constructing, organizing, and sharing visualization theory?

In this work, we explore the landscape of theory figures in visualization research. We begin by discussing the diagrammatic affordances of four conventional shapes in theory figures—Cartesians, flowcharts, matrices, and set diagrams. We then speculate on the value of using more unique shapes, considering examples like “The BLT Theory of Visualization Consumption”. We conclude by reflecting on how theory-figure-making can be not just a matter of communication, but a core practice of theorization.

2 BACKGROUND

Scholars have long reflected on the visual and rhetorical forms of research itself. Prior work includes automatic generation of IEEE VIS paper titles based on common naming conventions [20], as well as using computer vision to analyze the quality of a paper solely based on its layout [28].² While these approaches are mainly quantitative and focus on different paper elements than ours, they share with our work a curiosity about how form shapes meaning in research.

Our interest in theory figures also draws on the area of visual semiotics, most notably Jacques Bertin’s Semiology of Graphics [8]. Bertin demonstrated that graphical forms—points, lines, areas, and their arrangements—can encode information and meaning. We extend this perspective to contemporary theory figures, themselves structured graphics whose shapes influence interpretation.

-
- Matthew Varona and Carolina Nobre are with the University of Toronto.
E-mails: [varona, cnobre]@cs.toronto.edu.
 - Maryam Hedayati is with Princeton University.
E-mail: mhedayati@princeton.edu.
 - Matthew Kay is with Northwestern University.
E-mail: mjskay@northwestern.edu.

¹In this paper, we liberally use the term “shapes” to refer to the different visual forms theory can take. Strictly speaking, some of these aren’t exactly shapes, but you could write an entire alt.vis paper in the time it takes to say “different visual forms” 59 times, which is the amount of times we say “shape” in this paper.

²The resulting system seemed to score papers higher if they had lots of large, colorful figures. Thankfully, the system doesn’t look at actual content, so we suspect our paper would score fairly well based on that metric alone.

3 CONVENTIONAL FIGURE SHAPES AND THEIR AFFORDANCES

The shape of a theory figure is not neutral. A figure’s geometry guides how a model is interpreted—for example, a flowchart encourages the reader to follow along possible paths, while a table invites comparison across different categories. We describe these qualities as a figure’s *affordances*, adapting the term from human-computer interaction (HCI) research [25]. While affordances in HCI are about the actions possible with an interface, we use the term more loosely, extending it into the domain of conceptual and interpretive possibilities. In the context of theory figures, affordances capture how a shape invites particular ways of thinking—both for the theorist constructing the model and for the reader interpreting it.

We begin by looking at four of the most common shapes seen in theory figures: Cartesian planes, matrices, networks, and set diagrams. Each has distinct affordances that influence the structure and emphasis of the theories they depict. These affordances serve dual roles: they invite certain interpretations from a reader, and by extension, influence outcomes from the practice of theorizing. By first understanding how these familiar shapes guide our thinking, we can later begin to imagine more adventurous shapes for future theory figures.

3.1 Cartesian planes

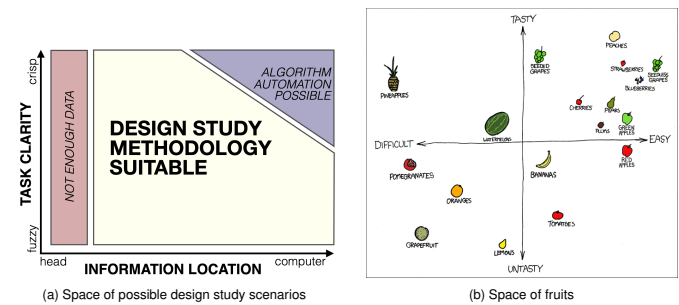


Fig. 2: Cartesian planes afford continuous positioning as well as discrete segmentation of their elements. In (a), Sedlmair et al. [27] provide guidance on when it is suitable to conduct a design study, based on task clarity and information location. The xkcd comic in (b) organizes fruits along axes of tastiness and ease of eating [30].

Cartesian planes consist of two axes that represent continuous variables. They are commonly used to map a conceptual space, situating ideas, phenomena, or actors based on their positions along two meaningful dimensions. This makes the Cartesian plane especially well-suited to theories that involve trade-offs or correlations between variables.

A core affordance of the Cartesian plane is the ability to partition. Axes can divide the plane into four quadrants, creating distinct regions based on the variables of the design space. This quadrant logic lends itself well to typologies: e.g., “high-high,” “low-high,” “high-low,” and “low-low” patterns, such as in Fig. 2b [30]. This type of theory figure has seen popular use in the form of the “political compass”, which claims to organize political views along economic (left-right) and social (authoritarian-libertarian) axes [1].

In addition to quadrants, Cartesian planes afford more granular “sub-areas” that can represent meaningful subsets or scenarios. Sedlmair et al. [27] use this technique to help researchers decide when a visualization design study is appropriate, as seen in Fig. 2a. These two types of grouping can complement or complicate each other; for example, a four-quadrant space could contain sub-areas that cross axis boundaries, showing edge cases or commonalities between quadrants. Through these different types of segmentation, Cartesian planes are able to represent both continuous and discrete elements of a theory.

While coordinate systems in figures are generally limited to 2 dimensions, we acknowledge the potential of 3-or-more-dimensional conceptual spaces as well. However, such figures may suffer from poor readability until academia breaks free from the tyranny of static paper formats [13, 19]. In lieu of interactivity or animation, Cartesian planes

can represent additional variables through the use of color, size, or shape, similar to scatter plots.

3.2 Matrices and tables

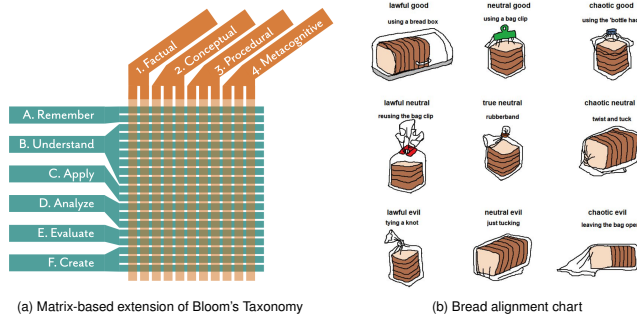


Fig. 3: Matrices organize knowledge into a neat grid, allowing rapid comparison and scanning. In (a), Adar and Lee use a matrix to modify an existing educational taxonomy [6]. The meme in (b) organizes different ways of storing bread based on character alignments from the game Dungeons & Dragons [2].

Matrices, in their simplest form, are two-dimensional grids that organize information into rows and columns. Each cell represents the intersection of a specific row and column, allowing for categorical or ordinal comparisons along two dimensions. Unlike Cartesian planes, which are primarily continuous, matrices favor discrete, combinatorial structure. They invite viewers to scan, compare, and categorize, rather than interpolate.

A popular example is the alignment chart from the role-playing game Dungeons & Dragons. Alignment charts use two axes—Lawful/Neutral/Chaotic and Good/Neutral/Evil—to explain the values and behaviors of in-game characters. The format has become ubiquitous in online culture, used to sort everything from fictional characters to bread storage methods, as seen in Fig. 3b [2]. Keen readers may notice that matrices with ordered axes (such as alignment charts) are just subdivided Cartesian planes. While this may be true for D&D charts, matrices generally ignore position within their cells. Subdivided Cartesian planes, however, still treat position within each region as meaningful.

One can extend the logic of matrices into hierarchical tables, where one or both axes have nested categories instead of a flat list. These are often used to show hierarchy in taxonomies, as in Adar and Lee's cognitive taxonomy for communicative visualizations Fig. 3a [6]. Depending on how hierarchy is introduced, it is possible for the table to lose some of its matrix properties (for example, if the uniform grid is interrupted or cells are grouped).

3.3 Networks

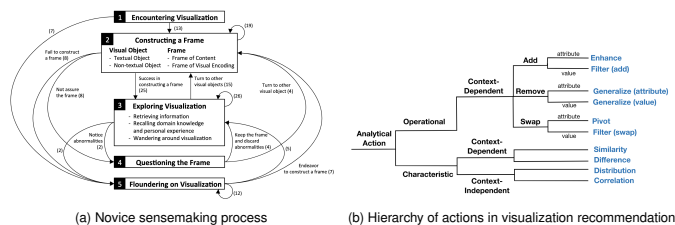


Fig. 4: Networks are suitable for representing processes, connections, or hierarchies. In (a), Lee et al. [22] depict the sensemaking process as a directed graph containing cycles and branches. In (b), Lee et al. show the hierarchy of actions in visualization recommendation systems.

Networks consist of nodes and edges, and can surface relationships, dependencies, and flows between entities. They often occur as

flowcharts or directed graphs that depict processes. Unlike Cartesian planes or matrices, which afford positioning and categorization, networks excel at showing connections and sequences. They can branch, recurse, and converge, making them an ideal choice for visualizing elements with complex relationships.

A key affordance of networks is path-following: arrows can trace progression from one node to the next. Networks also afford recursion and feedback in a way that other diagram types do not—cycles and loops can be drawn directly into the structure, making it easy to represent messy sequences. Lee et al. [22] use these techniques to depict how novices make sense of visualizations, which is an iterative and non-linear process (Fig. 4a).

Like other common shapes, networks can depict hierarchy in the form of trees. Lee et al. [21] use a tree structure to taxonomize the types of actions used in visualization recommendation systems (Fig. 4b). An advantage of trees is that they can efficiently show many layers of hierarchy, while also letting readers trace a path from the root category to a leaf. Aside from hierarchies and processes, networks could be used to show many-to-many relationships or dependencies, although care should be taken to minimize edge crossings for readability.³

3.4 Set diagrams

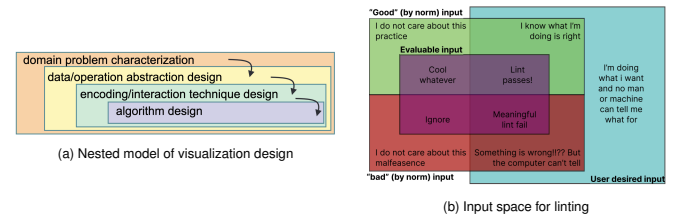


Fig. 5: Set diagrams can show intersection, membership, and dependency. Munzner's model in (a) depicts visualization design layers as a series of nested subsets [24]. Meanwhile in (b), Crisan and McNutt use an Euler diagram to show the relationships between good, bad, and user-desired inputs for linting [11].

Set diagrams show how concepts relate to each other through overlap and nesting. The primary types of set diagrams are Euler diagrams and Venn diagrams—the main difference between the two is that Venn diagrams depict all possible relationships, while Euler diagrams can depict a relevant subset of relationships. These figures are best suited for theories about membership and shared properties. Like networks, they can also be used to convey dependency.

A well-known example in visualization research is Munzner's nested model of visualization design (Fig. 5a) [24], which uses a layered, set-like figure to show how lower-level design decisions (such as algorithms) are contained within higher-level concerns (like tasks and domain problems). While the nested model shows subset relationships, set diagrams are also capable of showing intersection or disjointness. In Fig. 5b, Crisan and McNutt characterize the space of inputs for linters using an Euler diagram [11].

Set diagrams afford an intuitive way to communicate inclusion and overlap: they make it easy to see which concepts belong together, where categories intersect, and which areas remain distinct. They can also hint at hierarchical relationships through nesting, while still accommodating cross-cutting categories. However, their focus on membership means they are less effective for expressing dimensionality or continuous variation.

3.5 Hybrid figures

Some theory figures combine multiple shapes into hybrids, capturing complexity that a single form might lack. For instance, consider El-hamdadi et al.'s framework of trust in visualizations (Fig. 6) [14]. This figure nests a matrix inside a flowchart, showing that trust is influenced

³Or, if you want worse readability for some reason, you can deliberately maximize edge crossings [12].

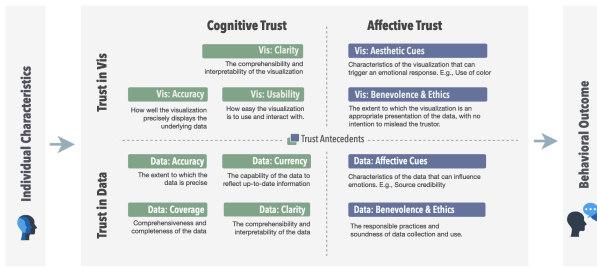


Fig. 6: Elhamdadi et al.'s [14] framework of trust in visualizations contains both matrix and flowchart elements.

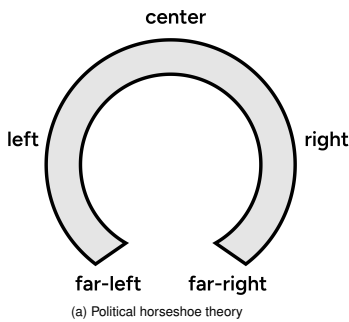
by individual traits and affects behavioral outcomes. Hybrid figures can enable greater expressiveness, but they come with trade-offs: complexity might make them harder to read, and an overloaded design risks diluting the clarity that simpler shapes provide. While we assert that many theories can be communicated with a single shape, hybrid figures remain a solid option when a theory spans multiple structural logics.

4 TOWARDS MORE EXPRESSIVE THEORY FIGURES

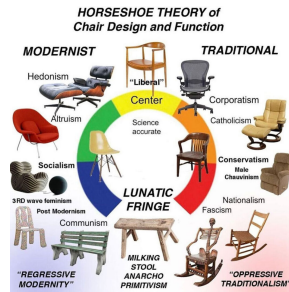
Theory figures hold great communicative power, as demonstrated by the variety of affordances offered by conventional shapes. It follows from this logic, however, that the narrow set of currently-used theory shapes constrains the expressiveness of theory work. Notice that common shapes tend to be based on mathematical concepts; this reflects the computer science background of many visualization researchers. Mathematical precision, however, is not the only virtue of theory. How might we use more *metaphorical* figures to connect with readers on a deeper level—for example, by capitalizing on our primal fear of maritime hazards, or by tantalizing hungry readers who haven't had lunch yet?

In this section, we discuss four examples of expressive shapes with the potential to describe, organize, and communicate visualization theory in novel ways. In some cases, we demonstrate the utility of these shapes by using them to propose example theories. These examples are necessarily half-formed: the intellectual labor of applying these shapes to valid, useful theories is beyond the scope of this provocation. Instead, our aim is to reveal how quickly a shape can suggest a narrative, regardless of whether or not that narrative withstands scrutiny.

4.1 Horseshoe



(a) Political horseshoe theory



(b) Horseshoe theory of chairs

Fig. 7: Horseshoes evoke parallelism between seemingly opposite points. They have seen use in politics as well as furniture design. [26]

A horseshoe curves an existing 1-dimensional continuum unto itself, bringing the ends of the spectrum closer together. The most popular use of this shape is in political discourse: the idea that the extreme left and right are more similar to each other than to their moderate counterparts (Fig. 7a). The political horseshoe theory is widely criticized by political scientists, who have pointed out that the far left and right differ vastly in their beliefs and motivations [10, 16]. Nevertheless, horseshoe

theory continues to be referenced in popular discourse—proof that a compelling shape can outlive a weak theory.

Even though the political horseshoe theory is horsecrap, the shape itself has great communicative potential. It quickly conveys parallelism between two sides of a spectrum or two related processes. The ends of a horseshoe reach towards each other but do not meet⁴, showing convergence while keeping end points distinct. Furthermore, the U-shape introduces a vertical axis along which another variable can be encoded.

One could imagine folding other linear concepts—like the processes discussed in Sec. 3.3—into horseshoe shapes to highlight unexpected points of similarity. Whether that is useful or misleading depends, as always, on what the shape is being asked to represent.

4.2 Icebergs

The iceberg is a shape with strong cultural currency, particularly in online subcultures, where it is used to rank layers of knowledge [3]. In these diagrams, the visible tip of the iceberg represents common understanding, while the submerged layers represent increasingly niche or insider knowledge. Its main affordance is communicating obscurity. The iceberg implies not just hierarchy or ranking, but concealment. It suggests that what is visible is incomplete, and that meaningful or important information lies beyond the view of the general public, as we have done in Fig. 1 through our Iceberg of Theory Figures.

Icebergs can also afford a sense of snobbery. To descend into the depths of the iceberg is to claim superior understanding, making it a solid choice for researchers hoping to signal that they are cooler than other people.⁵

Still, the iceberg is a powerful shape for theories that depend on surface/depth metaphors: symptoms vs. root causes, visible outcomes vs. invisible mechanisms. Its visual form encodes a clear conceptual split between what is seen and what is submerged, inviting the reader to descend into the underexplored depths.

A close relative of the iceberg is the tier list. Originating from video game communities, tier lists started out as a way to rank the competitive viability of game characters. They are now used to rank basically anything. While tier lists lack the visible-invisible dichotomy of icebergs, they come with their own unique affordances. Among players of fighting games, competing and winning using low-tier characters is a source of pride, showing one's drive to succeed against the odds [29]. Thus, tier lists could describe situations where there is tension between easy options and “noble” options. They are also more flexible than icebergs; tier lists can use non-competitive criteria, like which characters are most likely to believe in Santa Claus [4].

4.3 Möbius Strip

A Möbius strip is a mathematical object made by taking a strip of paper, giving it a half-twist, then joining the ends together. The resulting shape is one-sided, despite originating from a two-sided object (Fig. 8a). This non-orientable surface has various fascinating applications. Consider a “crab canon”: a musical piece where a melody is simultaneously played forwards and backwards. By printing two halves of the melody back-to-back on a piece of paper, then turning it into a Möbius strip, one can simultaneously traverse the original and reversed melodies.⁶⁷

Other applications include conveyor belts (ensuring that the surface receives an even distribution of wear) and computer print cartridges [7].

When used for theory, the Möbius strip affords a sense of infinity and return. Tracing the surface of the strip will eventually bring you back to your starting point, but with a subtle shift—what feels like the “other side” is actually the same surface, just further along. It is

⁴That would make it a circle.

⁵Of course, the iceberg in our teaser figure is not used this way at all. We mainly wanted to be able to make the recursion joke in the third layer.

⁶For a video demonstration of Bach's crab canon as a Möbius strip, see [here](#).

⁷Technically, this depends on how the notes are being read. A Möbius strip introduces a half-twist, meaning the score “flips” vertically at the seam where the start and end are joined. However, if you assume that whoever is reading the music can tell up from down, it still works. For a version where vertical flipping does not matter, see [here](#).

well-suited for theories that show how a process repeats with variation, or how two seemingly distinct states are actually paired aspects of the same phenomenon. Whereas the horseshoe highlights parallelism between extremes, the Möbius strip erases the boundary between them. Apparent opposites become one and the same side of a loop.

In a similar class of shapes to the Möbius loop is the ouroboros. It is an ancient symbol of a snake devouring its own tail. While the ouroboros also represents the infinite, it frames this concept with more futility than a Möbius strip—a snake eating its own tail is doomed to repeat its destiny. As such, the ouroboros tends to be used in a more negative sense.

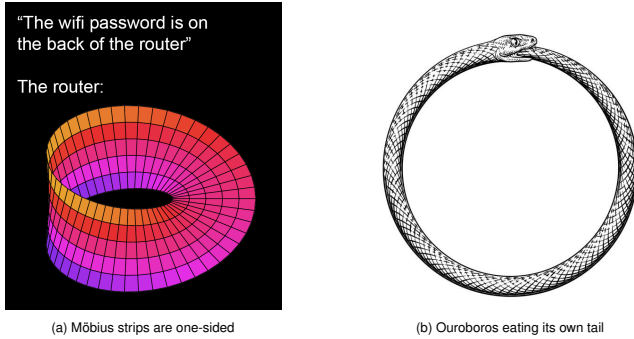


Fig. 8: A Möbius strip affords infinity and looping; the mythical ouroboros can express similar ideas, but with an air of self-defeatism.

4.4 BLT

A bacon, lettuce, and tomato (BLT) sandwich, as its name would suggest, consists of (at minimum) bacon strips⁸, lettuce, and tomato slices sandwiched between two slices of bread. The multi-ingredient, layered structure of a BLT makes it well-suited to describe theoretical models where there is complex interplay between elements. In a BLT, each ingredient maintains its distinct identity while contributing to the whole. The bread provides structural foundation, while inner ingredients can represent related concepts of varying importance and function. Just as the ingredients of a BLT come together to afford deliciousness, a BLT theory is worth more than the sum of its theoretical components.

The main diagrammatic benefit of the BLT is that its ingredients can be overlaid onto any “bread” (i.e. an existing theory figure). This “bread” serves as a base layer that can contain any foundational theory, while the inner ingredients represent adaptations or contextual factors that modify the base theory. Consider the example in Fig. 9. Through an in-depth interview study, He et al. characterized the space of information receptivity: an individual’s openness to receiving data and interpretations of data [17]. In their work, they posit that information receptivity underpins other factors in the interplay between visualization and reader, such as literacy, engagement, and persuasion. While the original authors of the paper advocated for this in a compelling manner, we believe the point would have been driven even further home if it were expressed as the **BLT Theory of Visualization Consumption**.

As the 6th most popular sandwich in the United States [15], BLTs hold immense cultural relevance. Consumers of BLTs often hold strong opinions about the “correct” way to make the sandwich. For instance, food critic J. Kenji López-Alt asserts that “[a] BLT is a tomato sandwich, seasoned with bacon. From this basic premise, all else follows.” [23] BLTs thus afford debates about theoretical primacy—which component is the true “tomato” of your theory, the essential element whose freshness is paramount and around which everything else is organized? Similarly, a BLT theory could be remixed with mayonnaise, avocado, or any number of auxiliary ingredients.⁹

We urge the community to look into the vast possibilities of food-based theories. Perhaps there could exist the complementary Grilled

⁸We note that our BLT sandwich uses vegan bacon.

⁹The authors of this paper claim no official stance on what constitutes an authentic BLT. Add ingredients at your own peril.

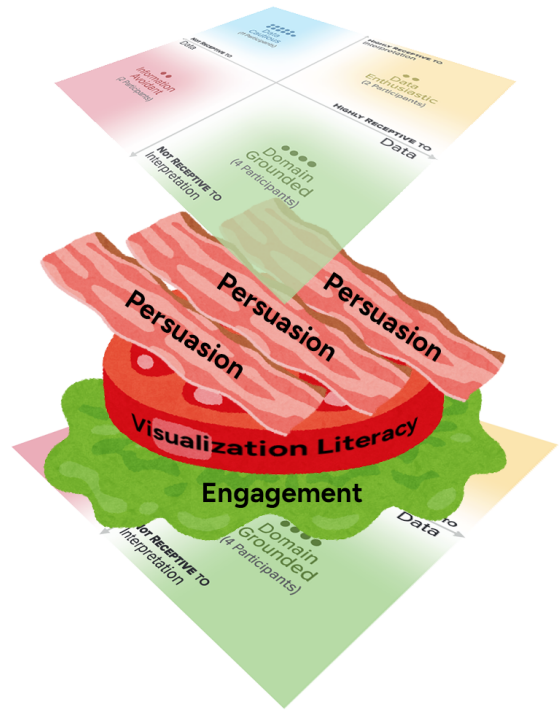


Fig. 9: The BLT Theory of Visualization Consumption, depicting the interplay between visualization and reader. The sandwich consists of a visualization literacy tomato, engagement lettuce, and persuasion bacon, sandwiched between the foundational concept of information receptivity.

Cheese Theory and Tomato Soup Theory, or a time-consuming and laborious process could be captured by a Risotto Framework. The apex of theory figures might come in the form of a Möbius strip-shaped pastry—even better if it has some kind of filling.¹⁰

5 FREQUENTLY ASKED QUESTIONS

5.1 Your new shapes suck!

First of all, that’s not a question. Second of all, no, they do not. The shapes presented here offer unique affordances that traditional figures like Cartesian planes and matrices simply cannot. A BLT captures distinct ingredients of a unified whole; a Möbius strip elegantly embodies infinity. While these shapes may not yet be common in academic publications, their expressive and analogous potential is undeniable. Our aim is to demonstrate that embracing new shapes can expand the ways we develop and communicate theory. In time, we hope to see the IEEE VIS proceedings filled with fantastical figures.

5.2 Aren’t horseshoe theory and the political compass massive oversimplifications?

Yes! Theory is useful precisely because it abstracts detail away, and yet, poorly conceived abstraction can be misleading. Theory figures lie at the core of this tension: they wrangle some real-world concept into neat, diagrammable dimensions, in a delicate balancing act between simplicity and nuance. We maintain that political horseshoe theory’s failings have little to do with the horseshoe itself compared to its uncritical development and adoption.

A tempting response to these concerns would be to add an ever-increasing number of caveats or expansions to a theory, in the hopes of defending against nitpicks from reviewers and skeptics. Sociologist Kieran Healy calls these “nuance traps”, arguing that an obsession with nuance detracts from the abstraction necessary for theories to be useful,

¹⁰I am getting hungry. Are you getting hungry?

compelling, and even correct [18]. As the saying goes, all models are wrong, but some are useful.¹¹

5.3 How is this related to visualization?

Visualization theory is fundamental to how we conduct research. Our conceptual understandings of visualization processes, users, and phenomena cascade into how we conduct studies and build systems. By diving into figure design, we hope to expand the theoretical playground that visualization researchers occupy. That said, we believe the concept of diagrammatic affordances is relevant to anybody concerned with communicating and constructing theory, even beyond visualization.

5.4 Should we *always* use shapes to theorize?

Not necessarily! Shapes are great, but there are certainly cases where they aren't the best way to think through theory. For example, some theories might consist of mathematical relationships that are difficult to map into a coherent visual, and are better expressed in formal notation.

That said, we maintain that thinking about shapes is more often than not a useful exercise. Theorycrafting with figures could serve a similar function as the practice of *memoing* in qualitative methods. In memoing, the researcher notes down observations and ideas about low-level codes and themes, crystallizing their theoretical arguments through the process of writing [9]. Even if a theory ultimately resists being drawn, the act of trying to give it a visual form can clarify relationships, reveal gaps, or surface hidden assumptions.

6 CONCLUSION

This work proposes the idea that **theory is shapes**. We discuss the affordances offered by common shapes, demonstrating how shapes influence theory. We then describe a new set of exciting and expressive shapes with which to do theory work. However, these shapes are only the beginning. Having drawn attention to the value of shapes, we hope the community continues to sink down into the lowest layers of the theory figure iceberg.

ACKNOWLEDGMENTS

We wish to thank Naaz Sibia and Patrick Lee for their valuable feedback on this work. We also thank the website irasutoya.com for providing the free clip art used to construct the BLT.

REFERENCES

- [1] The Political Compass. <https://www.politicalcompass.org/>. 2
- [2] Alignment Charts. *Know Your Meme*, Sept. 2009. <https://knowyourmeme.com/memes/alignment-charts>. 3
- [3] Iceberg Charts. *Know Your Meme*, Feb. 2016. <https://knowyourmeme.com/memes/iceberg-charts>. 4
- [4] Tier Lists. *Know Your Meme*, July 2016. <https://knowyourmeme.com/memes/tier-lists>. 4
- [5] "A picture is worth a thousand words." <https://www.merriam-webster.com/dictionary/a+picture+is+worth+a+thousand+words>, July 2025. 2
- [6] E. Adar and E. Lee. Communicative Visualizations as a Learning Problem. *IEEE Transactions on Visualization and Computer Graphics*, 27(2):946–956, Feb. 2021. doi: 10.1109/TVCG.2020.3030375 3
- [7] S. Alagappan. The Timeless Journey of the Möbius Strip. *Scientific American*. <https://www.scientificamerican.com/article/the-timeless-journey-of-the-moebius-strip/>. 4
- [8] J. Bertin. *Semiology of Graphics: Diagrams, Networks, Maps*. ESRI Press, 2011. 2
- [9] B. Birks, Y. Chapman, and K. Francis. Memoing in qualitative research: Probing data and processes. *Journal of Research in Nursing*, 13(1):68–75, Jan. 2008. <https://doi.org/10.1177/1744987107081254>. doi: 10.1177/1744987107081254 6
- [10] S. Choat. 'Horseshoe theory' is nonsense – the far right and far left have little in common. *The Conversation*, May 2017. <http://theconversation.com/horseshoe-theory-is-nonsense-the-far-right-and-far-left-have-little-in-common-77588>. 4
- [11] A. Crisan and A. M. McNutt. Linting is People! Exploring the Potential of Human Computation as a Sociotechnical Linter of Data Visualizations. In *Proceedings of the Extended Abstracts of the CHI Conference on Human Factors in Computing Systems*, pp. 1–10, Apr. 2025. <http://arxiv.org/abs/2502.07649>. doi: 10.1145/3706599.3716230 3
- [12] S. Di Bartolomeo, M. Lang, and C. Dunne. The worst graph layout algorithm ever. <https://osf.io/4hfy9>, Aug. 2022. doi: 10.31219/osf.io/4hfy9 3
- [13] P. Dragicevic, Y. Jansen, A. Sarma, M. Kay, and F. Chevalier. Increasing the Transparency of Research Papers with Explorable Multiverse Analyses. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*, pp. 1–15. ACM, Glasgow Scotland Uk, May 2019. doi: 10.1145/3290605.3300295 2
- [14] H. Elhamedi, A. Stefkovics, J. Beyer, E. Moerth, H. Pfister, C. X. Bearfield, and C. Nobre. Vistrust: A Multidimensional Framework and Empirical Study of Trust in Data Visualizations. *IEEE Transactions on Visualization and Computer Graphics*, 30(01):348–358, Jan. 2024. <https://www.computer.org/csdl/journal/tg/2024/01/10308716/1RMfAeipy4>. doi: 10.1109/TVCG.2023.3326579 3, 4
- [15] A. Greiner. What's America's Favorite Sandwich? | YouGov. <https://today.yougov.com/consumer/articles/24609-whats-americas-favorite-sandwich>. 5
- [16] P. H. P. Hanel, N. Zarzeczna, and G. Haddock. Sharing the Same Political Ideology Yet Endorsing Different Values: Left- and Right-Wing Political Supporters Are More Heterogeneous Than Moderates. *Social Psychological and Personality Science*, 10(7):874–882, Sept. 2019. <https://doi.org/10.1177/1948550618803348>. doi: 10.1177/1948550618803348 4
- [17] H. A. He, J. Walny, S. Thoma, S. Carpendale, and W. Willett. Enthusiastic and Grounded, Avoidant and Cautious: Understanding Public Receptivity to Data and Visualizations. *IEEE Transactions on Visualization and Computer Graphics*, 30(1):1435–1445, Jan. 2024. <https://ieeexplore.ieee.org/abstract/document/10292909>. doi: 10.1109/TVCG.2023.3326917 5
- [18] K. Healy. Fuck Nuance. *Sociological Theory*, 35(2):118–127, June 2017. <https://doi.org/10.1177/0735275117709046>. doi: 10.1177/0735275117709046 6
- [19] J. Heer, M. Conlen, V. Devireddy, T. Nguyen, and J. Horowitz. Living Papers: A Language Toolkit for Augmented Scholarly Communication. In *ACM User Interface Software & Technology (UIST)*. 2023. <https://idl.uw.edu/papers/living-papers>. doi: 10.1145/3586183.3606791 2
- [20] M. Lang, F. K. Opálený, P. Ulbrich, and L. Nikolajević. * (This Name Can Be Automatically Generated). In *Alt.Vis 2022*, Aug. 2022. <https://openreview.net/forum?id=CktC0F4a5DK>. 2
- [21] D. J.-L. Lee, V. Setlur, M. Tory, K. Karahalios, and A. Parameswaran. Deconstructing Categorization in Visualization Recommendation: A Taxonomy and Comparative Study. *IEEE Transactions on Visualization and Computer Graphics*, 28(12):4225–4239, Dec. 2022. doi: 10.1109/TVCG.2021.3085751 3
- [22] S. Lee, S.-H. Kim, Y.-H. Hung, H. Lam, Y.-A. Kang, and J. S. Yi. How do People Make Sense of Unfamiliar Visualizations?: A Grounded Model of Novice's Information Visualization Sensemaking. *IEEE Transactions on Visualization and Computer Graphics*, 22(1):499–508, Jan. 2016. doi: 10.1109/TVCG.2015.2467195 3
- [23] J. K. López-Alt. The Best BLT Sandwich. <https://www.seriousseats.com/ultimate-blt-sandwich-bacon-lettuce-tomato-recipe>. 5
- [24] T. Munzner. A Nested Model for Visualization Design and Validation. *IEEE Transactions on Visualization and Computer Graphics*, 15(6):921–928, Nov. 2009. <http://ieeexplore.ieee.org/document/5290695/>. doi: 10.1109/TVCG.2009.111 3
- [25] D. Norman. *The Design of Everyday Things: Revised and Expanded Edition*. Basic Books, Nov. 2013. 2
- [26] T. Rosen and L. Mörke. BOMB Magazine | Henrike Naumann by Tobias Rosen and Luise Mörke. *BOMB Magazine*, 2023. <https://bombmagazine.org/articles/2023/02/01/armchair-ideology-henrike-naumann-interview/>. 4
- [27] M. Sedlmair, M. Meyer, and T. Munzner. Design Study Methodology: Reflections from the Trenches and the Stacks. *IEEE Transactions on Visualization and Computer Graphics*, 18(12):2431–2440, Dec. 2012. <https://ieeexplore.ieee.org/>

¹¹This common aphorism points out that all models are inherently wrong because they simplify complex realities. Arguably, this statement itself glosses over the range of abstraction possible within theory, although it is still a nice phrase to keep in mind. One might even say that “all aphorisms are wrong, but some are useful”.

abstract / document / 6327248?casa_token = JMwC_nn7YMYAAAAA :
6XcnDiNcHnKT3p1yzHcQU9VaH6ScRaGCYlp_Z -
phiR6NijaSU6V_cicef18T6xJijHqV6-iJBQ. doi: 10.1109/TVCG.2012.213 2

- [28] C. von Bearnensquash. Paper Gestalt. <https://vision.ucsd.edu/sites/default/files/publications/pdfs/Paper%20Gestalt.pdf>. 2
- [29] C. Wagar. How to Pick a Fighting Game Character. <https://critpoints.net/2023/06/03/how-to-pick-a-fighting-game-character/>, June 2023. 4
- [30] xkcd. Fuck Grapefruit. <https://xkcd.com/388/>. 2