

Visible Language

the journal of
visual communication
research

Peterson et al.

06 – 51

funded research developing a systematic framework called the *Pictorial Trapezoid*, which offers greater control in producing new pictures with generative AI, and describing how an AI might be trained for semiotic precision in distinct research contexts

implications for providing greater control over image-to-image generation

Trischler

52 – 79

funded research exploring designers' current and desired uses of existing design knowledge

implications for the dissemination of design knowledge

Zender

80 – 83

announcement of a new publishing model for *Visible Language*

Visible
Language

57
• 3

the journal of visual communication research

december, 2023

ISSN 0022-2224

Published continuously
since 1967.

2023

december

Naomi Baron	<i>The American University, Washington, D.C.</i>
Michael Bierut	<i>Pentagram, New York, NY</i>
Ann Bessemans	<i>Hasselt University & PXL-MAD (Media, Arts, Design, School of Arts, Hasselt, Belgium)</i>
Charles Bigelow	<i>Type designer</i>
Matthew Carter	<i>Carter & Cone Type, Cambridge, MA</i>
Keith Crutcher	<i>Cincinnati, OH</i>
Meredith Davis	<i>Emerita Alumni Distinguished Graduate Professor, Department of Graphic and Industrial Design, College of Design / NC State University</i>
Mary Dyson	<i>University of Reading, UK</i>
Jorge Frascara	<i>University of Alberta, Canada</i>
Ken Friedman	<i>Chair Professor of Design Innovation Studies, College of Design and Innovation, Tongji University</i>
Michael Golec	<i>School of the Art Institute of Chicago, Chicago, IL</i>
Judith Gregory	<i>University of California-Irvine, Irvine, CA</i>
Dori Griffin	<i>University of Florida School of Art + Art History</i>
Kevin Larson	<i>Microsoft Advanced Reading Technologies</i>
Aaron Marcus	<i>Aaron Marcus & Associates, Berkeley, CA</i>
Tom Ockerse	<i>Rhode Island School of Design, Providence, RI</i>
Sharon Poggenpohl	<i>Estes Park, CO</i>
Michael Renner	<i>The Basel School of Design, Visual Communication Institute, Academy of Art and Design, HGK FHNW</i>
Stan Ruecker	<i>IIT, Chicago, IL</i>
Katie Salen	<i>DePaul University, Chicago, IL</i>
Peter Storkerson	<i>Champaign, IL</i>
Karl van der Waarde	<i>Swinburne University of Technology, Melbourne</i>
Mike Zender	<i>University of Cincinnati, Cincinnati, OH</i>

Contents

december

Matthew Peterson
Ashley L. Anderson
Kayla Rondinelli
Helen Armstrong

The Pictorial Trapezoid:
Adapting McCloud's Big Triangle for
Creative Semiotic Precision in Generative
Text-to-Image AI

o6 —

5¹

D.J. Trischler

Visible Language Online Archive
Improvement Study:
Envisioning Visible Language's Next
Generation Online Archive

5² —

79

Mike Zender

A New Model

8o —

83

The Pictorial Trapezoid:

Adapting McCloud's Big Triangle for Creative Semiotic Precision in Generative Text-to-Image AI

Matthew Peterson
Ashley L. Anderson
Kayla Rondinelli
Helen Armstrong

Abstract

Generative artificial intelligence (AI) is rapidly being adopted in diverse research contexts that, given the specificity of theoretical frameworks and research objectives, require a high degree of semiotic precision in AI output. With text-to-image generative models, the selection of subject matter and subsequent stylistic variation both have the potential to influence measurable desired outcomes. A major challenge in using generative models in design research is achieving a form of fidelity between a visual representation and a corresponding concept that must be conveyed. Scott McCloud's Big Triangle categorizes a broad range of visual representational stylistic variation, largely based on comic art. We extend the Big Triangle with a more systematically described framework called the Pictorial Trapezoid, which offers greater control in producing new pictures with generative AI. We provide a case study of the process by which we developed the Pictorial Trapezoid, and demonstrate its efficacy for an additional two research use cases. In each case we differentiate project-specific criteria for selecting what is being represented and visualizing that selection. Finally, we describe how an AI might be trained for semiotic precision in distinct research contexts using the Pictorial Trapezoid.

Keywords

Artificial intelligence
Generative AI
Semiotics
Text-to-image
Visual representation

Introduction

Generative artificial intelligence (AI) is actively and publicly transforming designers' capabilities in professional creative work (Liu et al., 2023; Turchi et al., 2023), and design research will similarly be impacted. As design researchers seek novel ways to engage with text-to-image AI technology, they will need to exert control over visualization within the confines of rigorous investigations. Such control must address visual style, a major and well-documented component of generative imagery prompting (e.g., "colorful risograph," Midjourney, n.d.[a]), but little attention has been paid to how to make informed stylistic decisions in highly demanding research contexts where stylistic variation has the potential to influence measurable desired outcomes.

Even if generative AI can produce almost anything, it is unprepared to produce the right thing without proper guidance. We are concerned with the form that guidance can take in research contexts, and to that end we utilize the case of a research project that, due to its nuance, exposes an important distinction between representational fidelity and conceptual fidelity. The former concerns the degree to which a picture resembles the entity in the real world to which it refers, while the latter concerns the degree to which a picture clearly calls to mind an abstract concept. Our objective is to highlight the potential of text-to-image AI for achieving appropriate visual representations by coordinating representational and conceptual fidelity. Scott McCloud's (1993) Big Triangle is a theoretical framework and picture plotting mechanism that categorizes a broad range of visual representational variation, largely based on comic art. In developing the Big Triangle, McCloud "was looking for a way to put all of comics' visual vocabulary (e.g., pictures, words, specialized symbols) into some kind of easily understood map of possibilities" (McCloud, n.d.[a]). We adapt the Big Triangle for formative use with generative AI. Adaptation is necessary because, in McCloud's own words, the Big Triangle "isn't a particularly scientific or exact tool" that "glosses over important aspects of artist intent and viewer interpretation" (McCloud, n.d.[b]). We are concerned with viewer interpretation rather than artist intent, and we are particularly concerned with precision in semiotic description, as precision is necessary in guiding generative AI.

The following section of this paper describes the early stages of the federally funded research project "Discovering Design Innovations for the Promotion of Insights by Intelligence Analysts," or Case 1 (see Acknowledgments for attribution). The process for Case 1 is itself a secondary contribution of this paper. We used exploratory visual analysis and image generation to derive criteria for continued design investigation, as guided by the Big Triangle.

The third section fulfills the adaptation of the Big Triangle suggested by Case 1, resulting in what we call the Pictorial Trapezoid, our primary contribution. The Pictorial Trapezoid permits the manipulation of style — readily available in text-to-image AI — in an exacting fashion that is conducive to the goal-based nature of research requiring customized visualization. We demonstrate its applicability beyond the initial research use case in the fourth section, using it to isolate stylistic characteristics relevant to two other research projects. Case 2 is the second author's ongoing social-belonging intervention for Black undergraduate students in engineering, in which participants will generate visual metaphorical representations derived from their experiences of belonging. Case 3 is drawn from Waisberg et al.'s (2023) work (as well as Balas & Micieli, 2022), which is centered on visualizing personal neuro-ophthalmic visual phenomena as experienced by patients. Finally, as a tertiary contribution and in the fifth section, we address the potential training of generative models with the Pictorial Trapezoid, which could facilitate human-AI teaming for research projects including and beyond the cases described here.

Case 1: Problematizing Concept Visualization

Intelligence analysts active in the intelligence community described for us conditions of their profession that precipitated the investigation of Case 1, which explores a novel use of imagery in interface design that could assist language analysts in their work. An intelligence analyst engages with an extensive collection of communication events and data from a considerable number of sources in a process of sensemaking. Analysts face significant difficulty in reconnecting with relationships found among data points from distinct sources when the data is multitudinous, and when they step away at the end of the workday. Case 1 is ongoing and ultimately explores how AI might promote analysts' serendipitous recognition of past data connections, and how AI may also proactively find connections between data sources that can be presented to analysts.

Due to the nature of intelligence analysis and an already extensive workflow, we envisioned a system that reduces cognitive burden and supports specific analytic tasks (we are limited here in our ability to disclose details of intelligence analysis workflows). Thus our investigation takes into account the obstacles impeding efficient memorization and cognitive processing and examines how a new interface could best contribute to analysts' recognition and recall. The human mind is well suited to relating pictures and visual structures to verbal representations, forming

strong referential and associative connections in long-term memory (Sadoski & Paivio, 2001). Relating language units to pictures provides more opportunities for representing and re-representing information. Furthermore, information need not be represented in the format in which it was initially experienced — e.g., a picture can call to mind a word.

Case 1 is based on the conceptual peg hypothesis, which asserts that concrete imagery can serve as a “conceptual peg” upon which to hang more abstract — and we presume more extensive — information (Anderson et al., 1977; Spoehr & Lehmkuhle, 1982). For this project conceptual pegs take the form of small pictographs, which can be mentally associated with a source of intelligence information (e.g., a recorded conversation, a summary report) and then incorporated into an interface for recalling and revisiting that information. Such pictographs thus have redintegrative value, meaning they promote “arousal of any response by a part of the complex of stimuli that originally aroused that response” (Merriam-Webster Dictionary) or the mental “restoration (of something) to a state of wholeness” (that “something” here being intelligence information; Oxford Languages Dictionary). In the simplest form of interface for Case 1, the pictographs that function as conceptual pegs could take the place of regular file icons, resulting in unique instead of categorical icons. In a more novel form, multiple pictographs could represent what would otherwise be a single file, or they could be positioned and rearranged in ways that do not respect normal hierarchical file organization structures. But these interface considerations are being addressed in a later stage of Case 1 that is beyond the scope of this paper.

AI can create the kinds of pictorial representations necessary for a conceptual peg interface, and it can present potential connections between different representations for additional insight, but it is not immediately apparent how it can do so well, and what limitations must be enforced. To that end, we devised three tasks to guide our exploration of the picturable aspects of intelligence information sources, and the nature of any resultant pictures themselves.

First Task:

Develop a semiotic catalog of representative intelligence information sources’ concepts, and analyze the concepts’ redintegrative value both to intelligence analysis and as determinants of pictographs.

The conceptual peg hypothesis suggests that pictographs in our imagined interface emphasize concrete aspects of information. Pictographs with this function are essentially metonymic or synecdochic, with only a portion of the information visualized despite the

intent that they represent, by association, the respective information in full. Creation of conceptual pegs requires two phases:

- 1.
Selection of a target concept from the information.
- 2.
Visualization of that target concept.

Given the information workflow in intelligence analysis, an analyst would likely indicate an area of interest, from which an AI could identify one of the few concepts internal or adjacent to that area, and then visualize the selected concept. Because the AI is involved in selection, we need to understand the range of possible concepts available. To that end, we identified two representative intelligence sources, an upstream example of a recorded conversation (a publicly available surreptitious recording of Billy Graham with and by Richard Nixon) and a downstream example of a report (a Bellingcat report on Russia’s bombing of Mykolaiv; Sheldon, 2023). “Upstream” and “downstream” refer to where in an extended workflow these kinds of sources would be encountered by analysts. Collaborating members of the intelligence community were limited in what they could share, but they confirmed that these sources were reasonably authentic representatives. In order to characterize these representative sources of intelligence, two coders identified embedded concepts according to the following categories: people, places, things, actions, adjectives, metaphors, and weapons (the latter of which is relevant only to the downstream report, as an example of ad hoc categories that are specific to intelligence customer interests). Furthermore, ChatGPT was utilized to identify idioms, to ensure that we were accounting for figurative language that is culturally bound, which would be encountered frequently in language translation activities. Documentation of concepts permitted a characterization of the sources’ profiles, as detailed in Table 1.

A sample visual display from the Bellingcat report is shown in Figure 1, and Figure 2 demonstrates the parallel and nested structure of nonverbal representations, in contrast to the serial structure of verbal representations. Sample armatures are visible in Figure 3, in another case of nesting in pictures. Armatures are visual marks that are non-objective indices (after Peirce; Burks, 1949), not depicting anything real in the world as much as guiding interpretation of pictures and words, as adjuncts to those semiotic elements (Peterson et al., 2021).

While the upstream conversation and the downstream report are only two examples of intelligence sources, the difference in their profiles suggests a substantial range of conceptual material, and we anticipate that they adequately cover the kinds of language and visuals that an AI will have at its disposal in selecting concepts from which to create

Elements and Concept Categories	Upstream Conversation	Downstream Report
Language units (primarily sentences)	113	60
People	18	51
Places	8	97
Things	19	383
Actions	80	233
Adjectives	14	127
Metaphors	24	4
Idioms ¹	6	5
Weapons ²	0	22
Visual displays	0	10
Top-level pictures	0	7
Enclosed pictures (within visual displays)	0	6
Armatures (non-depictive visual adjuncts)	0	97

Table 1.

Upstream and downstream intelligence information source profiles. (1) Idioms were determined by ChatGPT. (2) Weaponry is a theme particular to the downstream report (Sheldon, 2023), and thus it represents topical content.

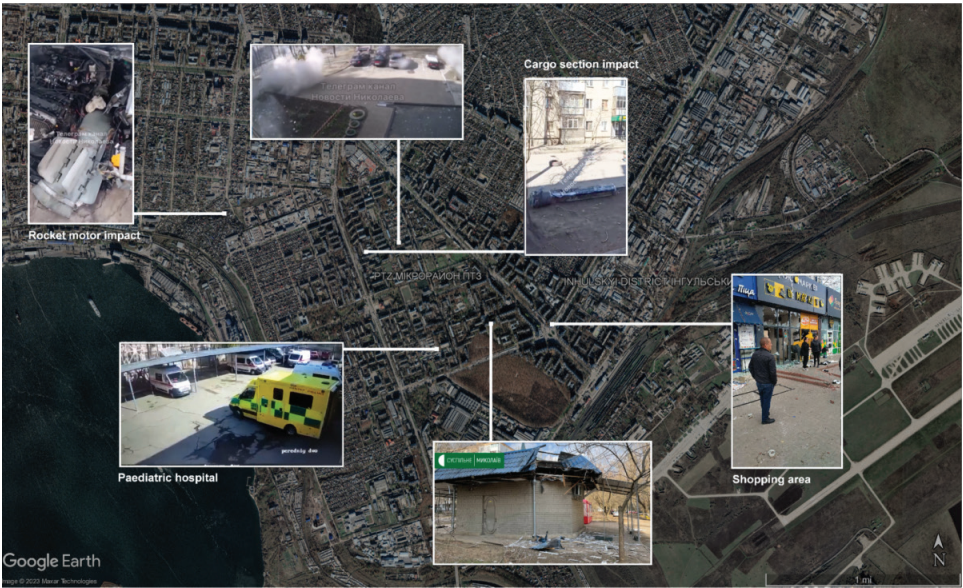


Figure 1.

Visual display highlighting the locations of pictures taken by media sources. This and subsequent visuals are taken from Sheldon (2023).

conceptual pegs. However, in practice, other data that is not initially conceptual may be of interest to analysts. For instance, language analysts may want to capture non-speech sounds such as a dog barking, music, or construction noises. Furthermore, metaphors that are not explicitly articulated in a source may be suggested by context. For instance, Nixon and Graham discuss how Nixon “threw away the text at the last [minute]” during a speech when a boy saluted him, and operator notes (i.e., contextual notes written by an analyst) reveal both that Nixon cried at that moment, and that the seemingly impromptu event was actually planned and performed. The metaphor *crocodile tears* accurately characterizes this situation, and an AI might recognize the relevance of that metaphor, thus generating a concept not directly evidenced in data. (Whether this is desirable is another matter.)

The upstream conversation includes a number of metaphors. Some overt metaphors are identifiable as figurative rhetorical speech — such as “I have to fight against the tide” — and some subtler metaphors are identifiable because of conceptual metaphor theory (Lakoff & Johnson, 1980) — such as “I got all that in there” (when Graham describes points he made in a New York Times editorial). The downstream report,

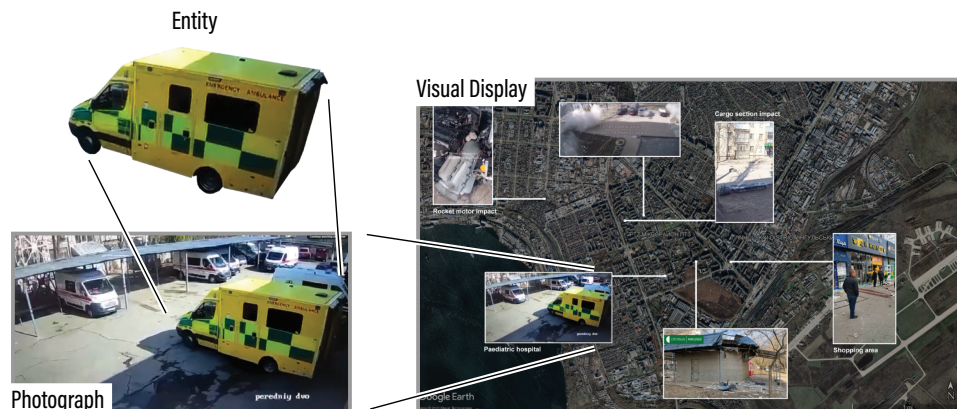


Figure 2.

Nesting structure of visual displays:
a visual display contains multiple
photographs, which may in turn
contain discriminable entities (from
Sheldon, 2023).

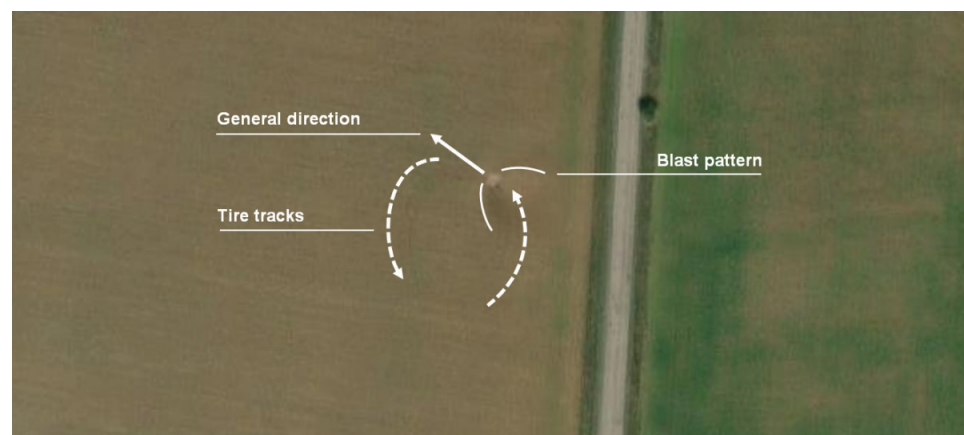


Figure 3.

Armatures over an aerial
photograph to scaffold
interpretation of a blast site (from
Sheldon, 2023).

in contrast, has minimal metaphor usage — it is a dense accounting of factual information.

Most of the concepts we have cataloged are unlikely to serve as useful material for conceptual pegs. We propose criteria for the concepts and pictographs embodied in conceptual pegs, two of which are relevant to the selection phase: essential and distinctive.

1.

Essential criterion for selecting concepts: the concept from which a pictograph will be generated should be crucial to the intelligence source in the context of an intelligence storyline, whenever possible.

2.

Distinctive criterion for selecting concepts: the concept should not be common across intelligence sources, or else it cannot serve as a strong memory trace back to a particular source. (While it is useful to track shared characteristics across data, we consider this to be covered by the tag, a conventional interface element that could coexist with conceptual pegs.)

Entire categories documented in Table 1 are likely to be useful only in rare cases. For instance, people who are essential to a storyline are likely to appear repeatedly in traffic, in which case their presence in one intelligence source will not distinguish that source from others. Adjectives are not likely to be essential or distinctive in many cases, as they are more frequently repeated in language.

The conceptual peg hypothesis emphasizes the concreteness of pegs, and readily visualized concepts, such as *dove* and *salute*, will likely make the most effective pegs from language. Pictures, by their nature, will tend to be distinctive, and we anticipate that when pictures are included in traffic, they will often make useful sources for pegs, should analysts determine them to be essential. For instance, the ambulance in Figure 2 is a reasonable conceptual peg not only for the visual display in which it is nested, but for the entire downstream report in which that visual display is nested in turn.

Second Task:

Analyze variation in AI-generated pictographs according to McCloud's Big Triangle, and evaluate the pictographs' redintegrative value to intelligence analysts.

Following selection, the visualization phase of conceptual peg generation requires an accounting of representational granularity in signification, or a measure of meaningful visual complexity.

For the second task, we identified a parsimonious subset of concepts from the upstream and downstream sources that meet the selection criteria as conceivably essential and distinctive concepts, and offer the opportunity to explore a wide range of AI-generated pictographs. Selected concepts include: *Richard Nixon*, *dove*, *crocodile tears*, and *(rocket artillery) attack*. We chose this set of concepts because it covers a range of challenges for visualization, from specific to general concrete nouns (Richard Nixon to dove), to metaphorical nouns (crocodile tears), to actions (attack). We then generated copious pictographs for those concepts using McCloud's (1993) Big Triangle, to ensure that our variation was extensive and systematic. We found the Big Triangle to be unique among the copious descriptions of semiotic distinctions in the literature, because it is systematic, with all considered variation anchored to only three qualities, and because it has comparative granularity, with McCloud providing 116 visual examples related to one another in a single chart (pp. 52–53). McCloud's anchoring of semiotic variation with three qualities is vaguely reminiscent of Charles Sanders Peirce's distinction of icon, index, and symbol (Burks, 1949), but Peirce did not provide visual evidence akin to McCloud's. Rune Pettersson's career-summarizing *Information Design* book series (Pettersson, 2015a–f), in its collective 2,056 pages, demonstrates a feature of the literature on visual semiotics, that there are extensive means by which visual characteristics and meaning have been elucidated, but without the combination of systematicity and granularity available in the Big Triangle (cf. Pettersson's, 2015c, pp. 327–329, Picture Circle, which covers much the same ground conceptually but does not give researchers the same wealth of examples to compare against). The Big Triangle is a uniquely rich starting point among the options in the literature.

The Big Triangle is shown in Figure 4, and has been used extensively in art and design education (*Understanding Comics* has 9,924 Google Scholar citations as of late 2023, and the Big Triangle is one of two main theoretical contributions in the book, along with a typology of narrative transitions, p. 74).

Though McCloud was concerned with illustration for comics, his framework can describe a full semiotic range of pictorial representation. As pictures approach the reality corner, they look increasingly like the real-world entities they signify. As they approach the meaning corner, they gain a degree of conceptual purity, stripped of specific details. And as they approach the picture plane corner, their formal characteristics are emphasized at the expense of semiotic clarity.

Note that in Figure 4 we excluded an additional right-hand portion of the Big Triangle that positions words in relation to pictures. Dual coding theory (Sadoski & Paivio, 2001) and working memory theory (Baddeley, 2018) both postulate distinct cognitive subsystems for

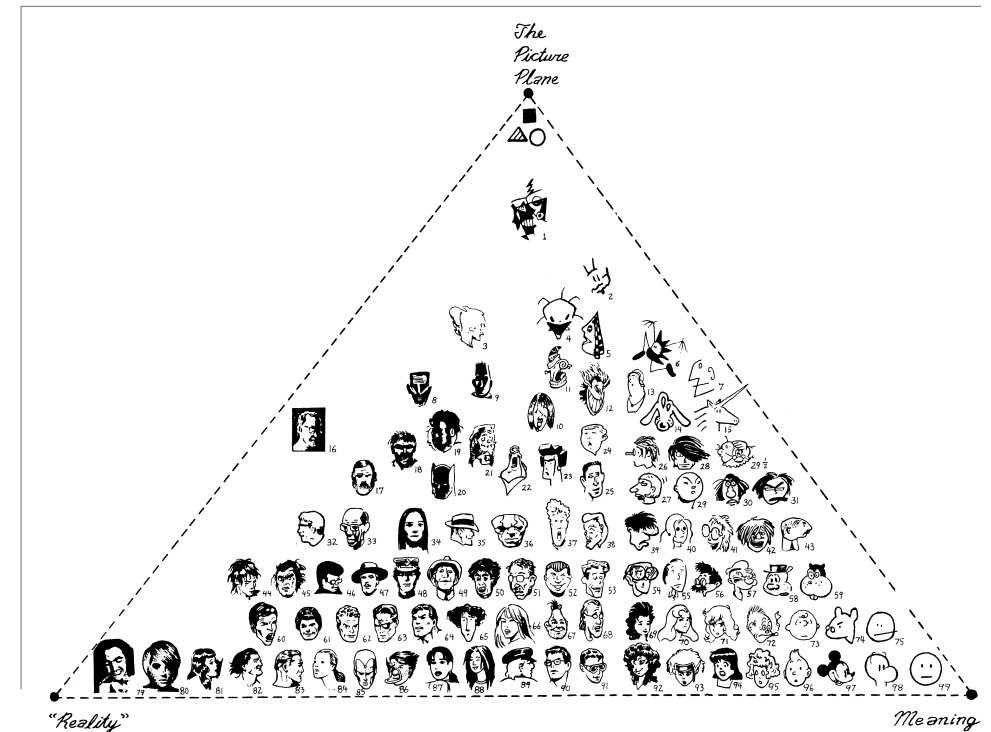


Figure 4.

Modified version of McCloud's
(1993) Big Triangle (pp. 52–53),
excluding a right-hand portion of
words and their visual expression.

processing verbal and nonverbal representations, with distinct formats for those representations, and thus we feel that integrating words with pictures in the Big Triangle suggests a greater degree of correspondence than exists.

We generated hundreds of pictographs from the subset of concepts, continuing until we had sufficient variety to plot the pictographs in a range of representational granularity (and ultimately through the Big Triangle). We chose 10 levels of granularity because that was the point at which we began to find it difficult to differentiate levels when plotting pictographs. Beyond this temporary judgment of convenience, an ideal level of granularity would need to be determined in future work, which could be accomplished in an experimental study by asking participants to sort a large number of pictographs in order between two given anchors (at the extremes, and as predetermined). Akin to intercoder reliability, consistency of ordering across participants could be calculated to determine, per a standard of agreement, which pictographs are consistently placed above and below one another. An operational granularity would

be the number of levels that can be ordered readily and consistently. In the short term, we believe our method produces adequate variety to make reasonable judgments.

Two of our proposed criteria for conceptual pegs are relevant to the visualization phase: distinctive and compact (as noted in the first task, the *distinctive* criterion is also relevant to the selection phase).

1.

Distinctive criterion for visualizing concepts: the pictograph generated from a concept should perceptibly differ from other generated pictographs, or else it cannot serve as a strong memory trace back to a particular source.

2.

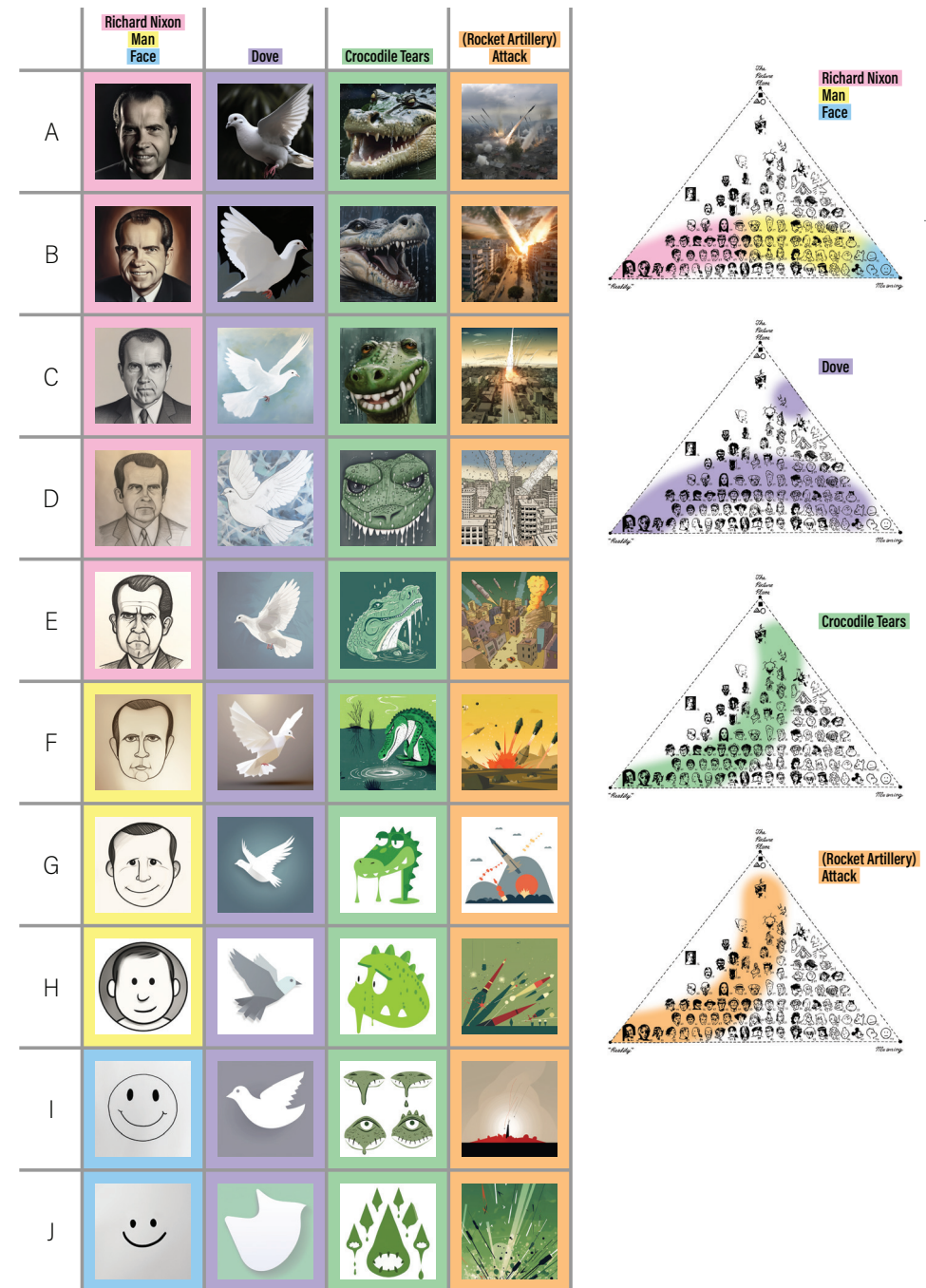
Compact criterion for visualizing concepts: the pictograph should have limited complexity such that it is discernable at a small size, as the project premise is based on a multitude of conceptual pegs being visually accessible at any given time in a dynamic interface, so that serendipitous connections can occur.

When we attempted to visualize Richard Nixon in a range of representational granularity, we realized that increasing abstraction reduces the specificity of signification. Figure 5 plots the following: as a concept, *Richard Nixon* becomes simply *man* in row F when Nixon's distinguishing characteristics are lost through simplification; and *man* is likewise reduced to *face* in row I when gendered signifiers are lost. Thus, there is a region of the Big Triangle that the specific concept *Richard Nixon* can occupy, and another region where a picture can suggest a face but not a man. (We acknowledge that this is an oversimplification of gender that is not reflected in many personal identities.) When the other concepts' pictographs in rows A–J are plotted in the Big Triangle, *dove* approaches the meaning corner (at I) but does not reach it, while the final pictograph (at J) is somewhat discontinuous, when approaching the picture plane corner. But *crocodile tears* and *(rocket artillery) attack* do not approach the meaning corner. We were unable to create equivalent dynamics for the disparate concepts, and we believe this is due to the nature of the concepts themselves. *Face* and *dove* are what we consider to be basic concepts, nearly universal things in the world (Rosch, 1978). *Face* is truly universal, and it is no surprise that we humans readily recognize abstract marks as faces, and that simple representations have emerged that are easy to replicate. *Dove*, in contrast, is culturally dependent, but also often culturally important, and it likewise has common representations. It cannot be reduced as much as *face* can and remain recognizable. *Crocodile tears* and *(rocket artillery) attack* are multi-component concepts, and the latter is additionally an action. These are not basic concepts.

Figure 5.

opposite page

Variable signification matrix of pictographs generated for the concepts selected from the upstream and downstream intelligence sources, with conceptual ranges plotted into the Big Triangle (McCloud, 1993). All pictographs were generated with Midjourney.



Particular concept classifications may best be visualized in particular areas of the triangle — for instance, specific places (e.g., New Orleans) versus kinds of places (e.g., a city). Photographic depictions of a hospital may be ideal for representing a particular hospital, whereas more iconic visualizations may be useful in more generally representing the concept *hospital*.

Ultimately, the plotting exercise evident in Figure 5 suggests a relative “sweet spot” with an ideal intersection of McCloud’s (1993) reality, meaning, and picture plane qualities. The highlighted region of Figure 6 approximates the sweet spot for Case 1, and the gray regions are likely to fail to meet one of the selection or visualization criteria. The region at A is near the reality corner, and pictographs in this area may be too complex to be compact. The region at B is near the meaning corner, and pictographs in this area may be too generic to be distinctive. And while the *essential* criterion concerns concept selection, pictographs in the region at C — near the picture plane corner — may be so abstract as to be unidentifiable, thus recursively violating that criterion. That is, such pictographs may interrupt the concept-pictograph signification chain. We consider these only to be rules of thumb, with exceptions anticipated.

Figure 6.

A possible sweet spot in the Big Triangle for conceptual peg pictographs in Case 1 (modified from McCloud, 1993, pp. 52–53).

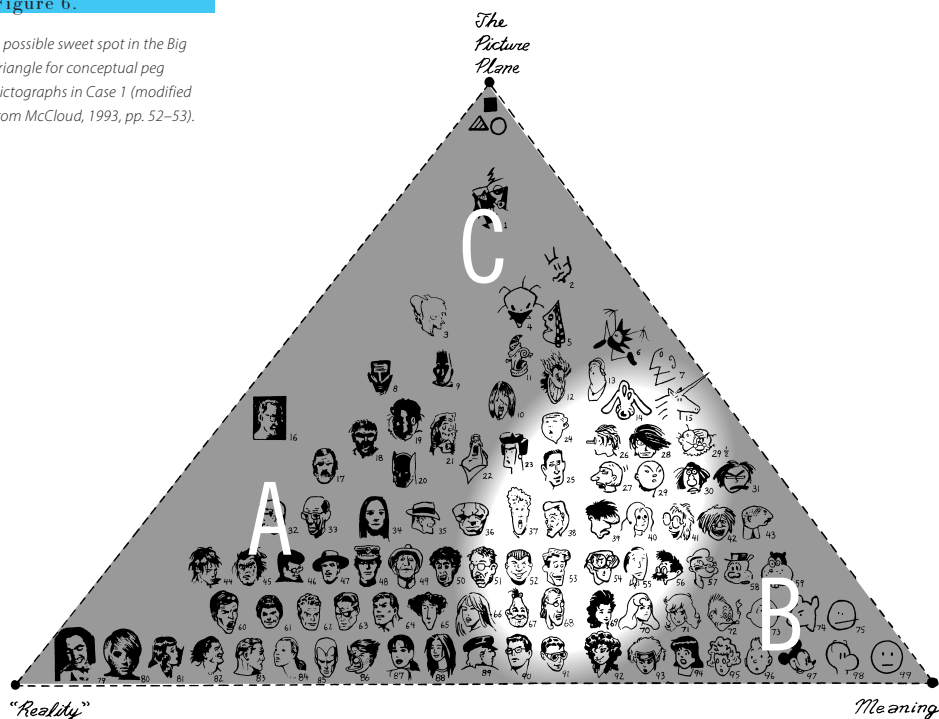


Figure 7.

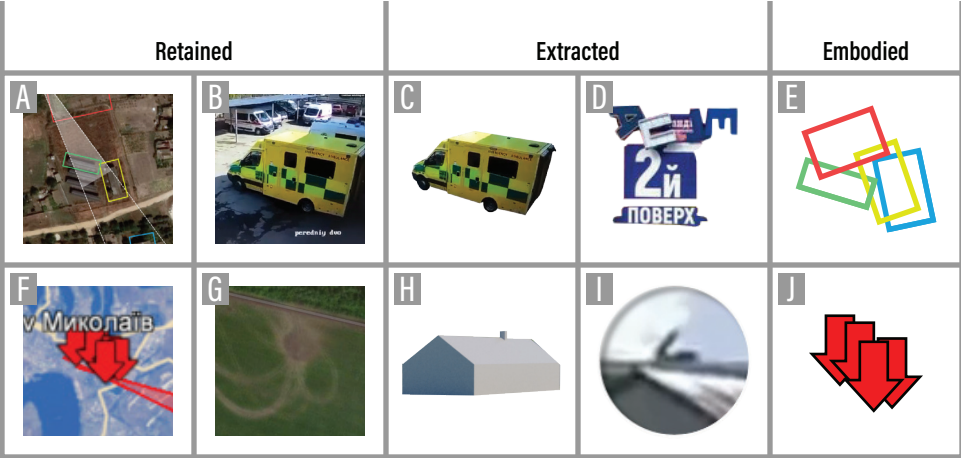
Conceptual-blended pictographs for the upstream and downstream sources. Each combines two concepts. All pictographs were generated with Midjourney.

	Dove & Senate	Russian Army & River	Cluster Munitions & Sundial
A			
B			
C			

Due to the conceptual richness of the upstream and downstream intelligence sources, the selection and visualization criteria may be served by conceptual blending, the representation of two or more concepts in a single hybrid pictograph. A single conceptual peg reference to two essential concepts should better characterize that source, or be highly distinctive. It would also increase the chance that a concept in a given source that subsequently becomes contextually important is directly visible as a pictograph in a conceptual peg interface. Cunha (2022) demonstrated considerable variation of conceptual blending of small pictographs through the use of generative models. We have utilized conceptual blending for the upstream and downstream sources, as shown in Figure 7. The results vary in the degree to which entities are blended. *Dove & Senate* at A and B in Figure 7 are fanciful hybrid entities, while *Dove & Senate* at C retains both dove and Senate entities, though the size differential is surprising, with the dove towering over the US Capitol. *Russian army & river* at A in Figure 7 places a commander in the river, a realistic but silly arrangement. The *cluster munitions & sundial* examples in Figure 7 appear to be munitions-flavored sundials; that is, the sundial is the primary motif. Though they combine two concepts apiece, these pictographs are only moderately more complex than other generations.

As noted in the first task, when pictures are included in traffic (i.e., the flow of intelligence information in an intelligence community), they will likely prove useful as sources for conceptual pegs. While words, being purely symbolic and conventional, are repeatedly experienced across contexts, pictures are most often uniquely encountered in particular contexts. Thus they can form a more direct memory trace to concepts already in a distinctive form. Furthermore, pictures have long been known to create stronger memory traces than words, with multimodal memory traces stronger yet. We have focused on text-to-image generations of conceptual pegs not because textually-constituted concepts should be favored, but because these generations are translational across codes, and are thus especially problematic. Figure 8 provides examples of pictographs generated from pictures in the downstream intelligence report. Four examples are retained from their sources, merely cropped and sized to fit a pictographic scheme (A,B,F,G in Figure 8). Four are extracted from pictures, additionally silhouetted (C,D,H,I). Finally, two are embodied, or redrawn from key elements (E,J). We assume that it is desirable to retain the visual characteristics of the source pictures in the generated conceptual pegs, in order to reinforce the modal memory of engaging with the original source, even when the resulting pictograph falls outside of the sweet spot.

Figure 8.
Potential conceptual pegs drawn from pictorially-constituted concepts in the downstream Bellingcat report (Sheldon, 2023).



Third Task:
Analyze variation in AI-generated pictographs according to visual style as manipulable through prompting.

Because of the inherent representational complexity of pictures, every picture has a distinguishable visual style — or a degree of conformance to an aesthetic tradition — though styles are multifaceted and have fuzzy boundaries, and cannot thus be delineated with great precision. In the practice of guiding a text-to-image generative model, textual prompts determine stylistic characteristics as deviations from a possible “house style.” We assume that an AI that automatically generates visual conceptual pegs would have a sensitivity to visual style comparable to what is achievable through prompting. Furthermore, we assume that such an AI would be trained to generate conceptual pegs with a predetermined degree of stylistic variation, and that this training would derive from best practices of aesthetic prompting.

Prompt engineering for generative AI is an emerging discipline, particularly for text-to-image generative models that allow users to create pictures from natural language text descriptions. It is the iterative process of generating and refining pictures within text-to-image AI through trial and error. Research suggests that prompting is a learned rather than an intuitive skill (Openlaender et al., 2023). As interest grows, resources to improve prompting are being developed to support a growing community of practitioners. Through practice, users have identified specific keywords and phrases that serve as modifiers and can be included in various permutations to influence output — by increasing the quality of a picture or rendering it in a particular visual style (Liu & Chilton, 2022). Prompts usually contain several terms corresponding to a subject, style, medium, and additional details (Martins et al., 2023). *Subject* refers to the main visual element of the picture, while *style* refers to distinct visual characteristics. *Medium* is used to specify a picture’s apparent material means of creation — for instance, rendered as a photograph, painting, or drawing. Figure 9 shows a series of mountains rendered in five prompted styles: photograph, painting, sketch, cartoon, and computer icon. It is immediately apparent, as reflected in empty cells in Figure 9, that visual style places limitations on complexity. For instance, cartoons are by definition not photorealistic, and thus they will never occupy row A in Figure 9. Conversely, to abstract a photograph below row B would require a stylistic deviation that would fundamentally diminish its photographic quality, shifting the resulting pictograph into another column of Figure 9 altogether; thus photographs

Figure 9.

Variable signification matrix of pictographs generated for the mountain concept. Stylistic variation delimits representational granularity, as is evident in empty cells. All pictographs were generated with Midjourney.

	Photograph	Painting	Sketch	Cartoon	Computer Icon
A					
B					
C					
D					
E					
F					
G					
H					
I					
J					

24

are not plotted below row B.

Prompts were strategically written to achieve the range of pictographs needed to stylistically differentiate mountains. Plain text descriptions indicating subject (i.e., "a mountain") and medium (e.g., "a pencil sketch") were used with selected parameters to alter settings in Midjourney (n.d.[b]). This platform was selected because it provided easy access to granular control of stylistic parameters compared to other text-to-image generators. The *style* parameter was used to reduce the degree to which Midjourney's stylistic bias impacted outcomes. Higher values impose a stronger house style on outcomes, leading to pictures that may be less reflective of the prompt, while lower values result in less stylized pictures that are more closely related to the prompt. This parameter is set to 100 by default, but can be adjusted in the range of 0 to 1000. We used

Table 2.

Sample prompts used to control variation of mountain generations with Midjourney. "Source" refers to cells in Figure 9; "--v 5" is notation for Midjourney's version five mode; "--s" is notation for the style parameter; "--q" is notation for the quality parameter.

Photograph – A



high resolution photograph of a typical mountain, unobstructed view --v 5

Painting – E



a stylized painting of a mountain inspired by van gogh but still as realistic --v 5 --s 50 --q 2

Sketch – H



a realistic sketch of a mountain --v 5 --s 50 --q 2

Painting – I



a mountain, painted outline only, simple, minimal, no detail, painted by a high school student --q 2 --v 5

Cartoon – I



a triangle that barely looks like a mountain, simple, cartoon --q 2 --v 5

Computer Icon – J



a very simple diagram which indicates a mountain using just one line which looks like an icon --v 5 --s 50 --q 2

Midjourney's *chaos* parameter to increase variability in outcomes. Higher values yield more unusual or less expected results, while lower values yield a reliable similarity across results. This parameter is set to 0 by default and can be increased up to 100 (Midjourney, n.d.[b]). Achieving outcomes that were increasingly abstract, given our interest in exploring the full range represented in the Big Triangle, required alternative methods. For the lower rows of Figure 9, the general shape of a mountain was used to guide the system — for instance, prompting Midjourney to generate “a triangle that barely looks like a mountain.” To achieve a specific visual style or degree of abstraction, we prompted Midjourney to generate results matched to particular skill levels — for instance, a high school student or Vincent Van Gogh. Similarly, pictographs were generated to represent extremes from most photorealistic to most abstract in order to delineate a middle ground. Finally, previous generations were added as visual prompts either with written prompts or with other pictures to generate middle-ground outcomes. Sample prompts are provided in Table 2.

Implications

The above exploration elucidated underlying issues that subsequently impacted interface prototypes we are not presenting here. We identified two phases of visual conceptual peg creation: selection of a target concept and its subsequent visualization as a pictograph. We proposed three criteria for conceptual pegs: essential, distinctive, and compact. As shown in Table 3, isolated concepts should be both essential and distinctive, while the pictographs that represent those concepts should be both distinctive and compact. We also added nuance to McCloud's (1993) Big Triangle in separating selection and visualization, culminating in our proposal that concepts themselves must shift as representational granularity is modulated. This

Table 3.

Proposed concept selection
and visualization criteria for
pictographs that function as
conceptual pegs.

Criterion	Concept Selection	Concept Visualization
Essential	Concept is crucial in the context of a storyline	—
Distinctive	Concept is not too common across cuts in the context of a storyline	Pictograph is perceptibly different from other pictographs in the context of a storyline
Compact	—	Pictograph is discernable at a small size

is best explained in the necessary shifting of *Richard Nixon to man to face* in Figure 5 as visual complexity is reduced. More specifically for the visual conceptual peg premise, regions of the Big Triangle have characteristics that interact with our criteria, and that might suggest a sweet spot at some distance from McCloud's reality, meaning, and picture plane corners, as seen in Figure 6.

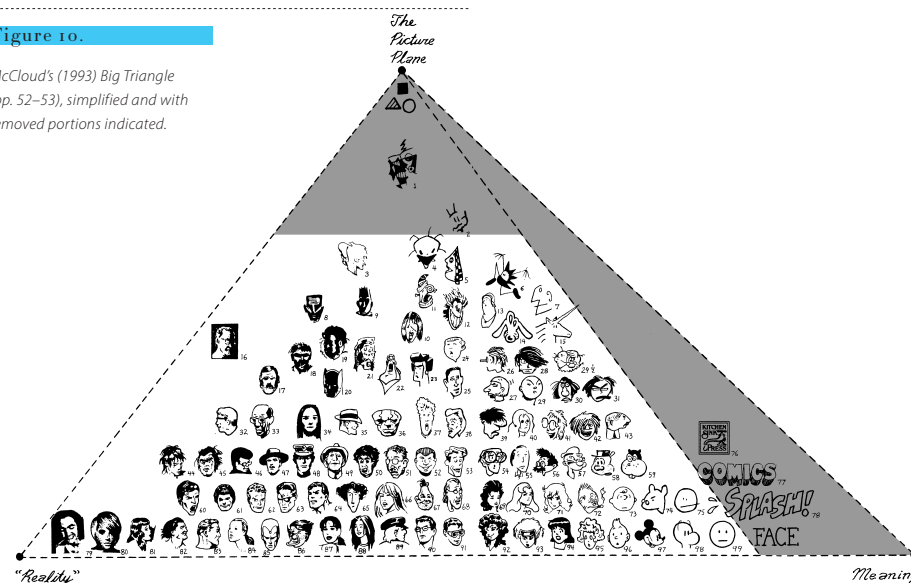
The Pictorial Trapezoid

The exploration described above was conducted in the service of a speculative interface design project that ultimately led to visual prototyping. But the adaptation of McCloud's (1993) Big Triangle precipitated by this exploration suggests a revised framework with a systematic description. This requires a further step in adaptation. In review, the Big Triangle as presented by McCloud includes a central triangular portion with corners that represent extreme qualities achievable in pictures (Figure 10):

- Near the reality corner, pictures can be “received” in a more straightforward fashion (p. 49) because they visually resemble real entities out in the world.
- Near the meaning corner, pictures are “more abstracted” and “require greater levels of perception, more like words” (p. 49).
- Near the picture plane corner, pictures are “non-iconic” and “no

Figure 10.

McCloud's (1993) Big Triangle
(pp. 52–53), simplified and with
removed portions indicated.



attempt is made to cling to resemblance or meaning" (p. 50).

The Big Triangle has a wedge off of the meaning-picture plane axis that plots text, typographic flourishes, and logotypes. McCloud demonstrates efficacy for understanding comic art in seeing continuity between pictures and words, particularly in terms of comics production (pp. 138–161). But we feel that moving beyond comics, such explicit continuity from pictures to text and typography is not supported by the nature of human working memory (Baddeley, 2018) and mental representations (Sadoski & Paivio, 2001). The visuospatial sketchpad of working memory processes depictive representations, while the phonological loop processes linguistic representations (Baddeley, 2018). These systems and their codes are fundamentally distinct from one another, suggesting separation in a corresponding framework, not continuity. We thus ignored this portion of the Big Triangle in our exploration, focusing exclusively on pictorial signification.

The underlying agenda for our exploration especially problematized representational fidelity and conceptual fidelity. Representational fidelity refers to the degree of accuracy to which a picture resembles its referent. To achieve high representational fidelity, there must be an entity in the real world to compare the picture against (in contrast to concepts like *love* and *war*). Conceptual fidelity concerns the connection of pictures to concepts that can be articulated more abstractly (e.g., a picture of a chair reading as "chair"), and refers to the degree to which a picture represents a given concept over alternatives. High representational fidelity is achieved in the Big Triangle's left-hand side, while high conceptual fidelity is achieved in its right-hand side. We found an interaction between representational fidelity and conceptual fidelity. While McCloud plots a range of what are largely faces in his rendering of the Big Triangle, we noted that the more realistic a picture, the less it can represent a generic concept. As seen in Figure 5, a highly articulated face has too much detail to read simply as "face" — it may instead become a white man with moderate representational fidelity, and Richard Nixon with high representational fidelity.

Furthermore, in as much as McCloud primarily plots faces in his Big Triangle, representations approaching the picture plane corner lose any resemblance, and thus cannot be faces. As we excised the verbal language portion of the Big Triangle, our emphasis on pictorial signification suggested trimming off the top as well (Figure 10, highlighted). The kinds of marks McCloud places near this corner do not cease to exist; but they are fundamentally different from actual pictures, which we consider to suggest depiction. This separation of representation types aligns with Peirce's distinction of icon, index, and symbol (Burks, 1949). The icon is exclusive to what we retain of the Big Triangle, as depictions of real-world entities. The index will "refer to or call attention to" something else (p. 677),

and we follow Peterson et al. (2021) in classifying "armatures" or "adjuncts" — marks such as lines, arrows, and boxes — as indices that are subservient to depictive entities and written language (though these are not the only kinds of indices). The symbol is entirely arbitrary and must be learned, like the written letters and words of natural language.

Figure 11 presents our adaptation of McCloud's (1993) Big Triangle, which we call the Pictorial Trapezoid. It exclusively plots variations of Peirce's icon (Burks, 1949). McCloud intended the Big Triangle to represent "the total pictorial vocabulary of comics or of any of the visual arts" (p. 51). The Pictorial Trapezoid is more aligned with design than with art. While the Big Triangle supports appreciation, and what Rosenblatt (1978) calls aesthetic reading, the Pictorial Trapezoid is intended to support generating pictures for what Rosenblatt (1978) calls efferent reading, or a person's engagement with media as a means to an end, a way to "carry away" some benefit from the experience (p. 24). For Case 1, we sought to facilitate intelligence analysts' use of visual conceptual pegs to establish stronger memory traces for intelligence information in the service of sensemaking. To establish the Pictorial Trapezoid as a separable contribution that extends McCloud's work, we describe it more systematically here. (While McCloud discusses the relationship between reality and meaning corners, and reality and picture plane corners, he does not directly address the relationship between meaning and picture plane corners; this is an example of a gap in systematicity.)

Figure 11 plots the *mountain* concept through the Pictorial Trapezoid because mountains can be rendered to varying degrees of fidelity while still retaining the connection to the basic concept (cf. *face to man* to *Richard Nixon* in Figure 5). The labeled corners include (A) the depictive vertex, (B) the conceptual vertex, and (C) the non-objective vertex, though the latter is placed outside of the trapezoid because visuals near that corner would not be recognizable enough to signify anything in particular. Instead, corners D and E approximate a limitation of identifiable signification.

Figure 12 introduces a numeric scheme for plotting pictures in the Pictorial Trapezoid. Each of three digits indicates relative proximity to one of the vertices, in the order of ABC. Thus Position 199 (P-199) is nearest to the depictive vertex (the first digit, in the A slot in Figure 12, is "1"). Consistent with the Big Triangle, as pictures approach a corner of the Pictorial Trapezoid, they are constrained with respect to the others (e.g., P-111 is not possible). Along the outer edges of the trapezoid, pictures vary in particular ways. The A-B axis maximizes differentiation according to complexity (or degree of detail); the B-C axis maximizes differentiation according to clarity (or how readily the picture communicates a distinct concept); and the C-A axis maximizes differentiation of saliency (for instance,

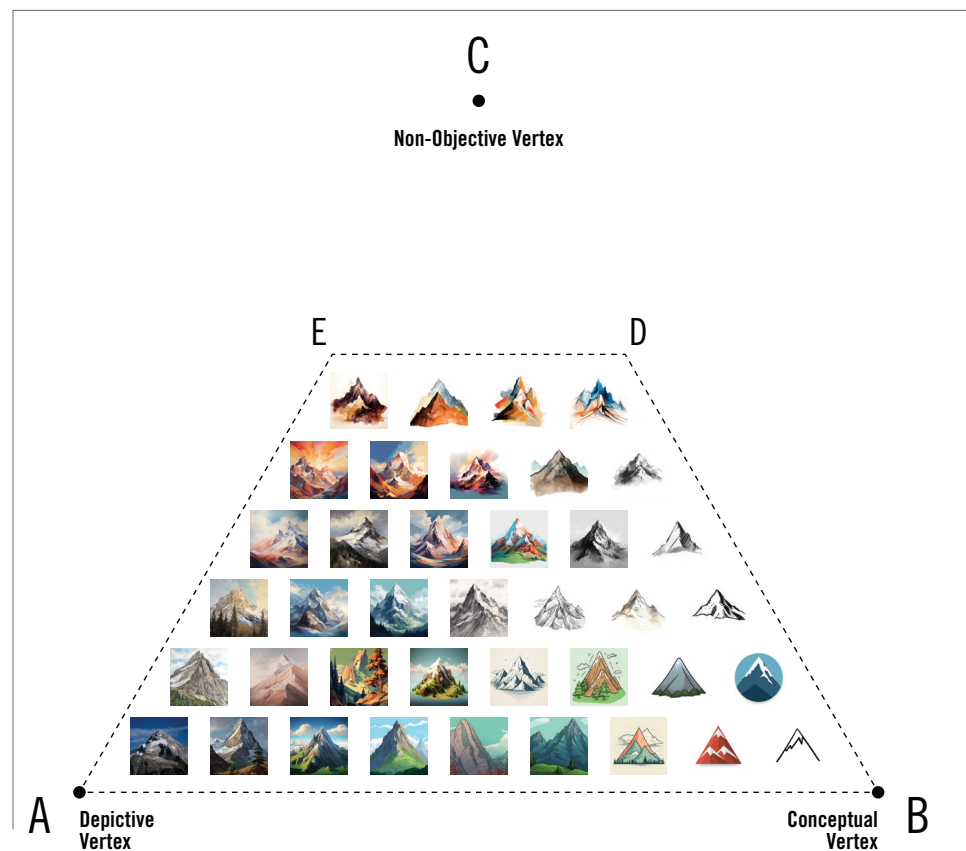


Figure 11.

The Pictorial Trapezoid, a revision of McCloud's (1993) Big Triangle, with pictographs of the mountain concept. All pictographs were generated with Midjourney.

at P-199, the thing being represented is salient; at an imaginary P-991, stylistic qualities such as linework would themselves be salient, with no remaining sense of any thing being represented).

Figure 13 summarizes trends within the Pictorial Trapezoid, and Figure 14 plots these trends as continua separate from the trapezoid schematic. Moving outward from the depictive vertex at equal distances from corners B and C, pictures change from more realistic to more abstract. Deviating from this line toward B results in conceptual abstraction, where simple concepts like *face* need little differentiation, while deviating toward C results in figurative abstraction. Moving outward from the conceptual vertex at equal distances from corners A and C, pictures change from more generic to more specific (in the sense that Richard Nixon is a specific man). And moving outward from the non-objective vertex at equal

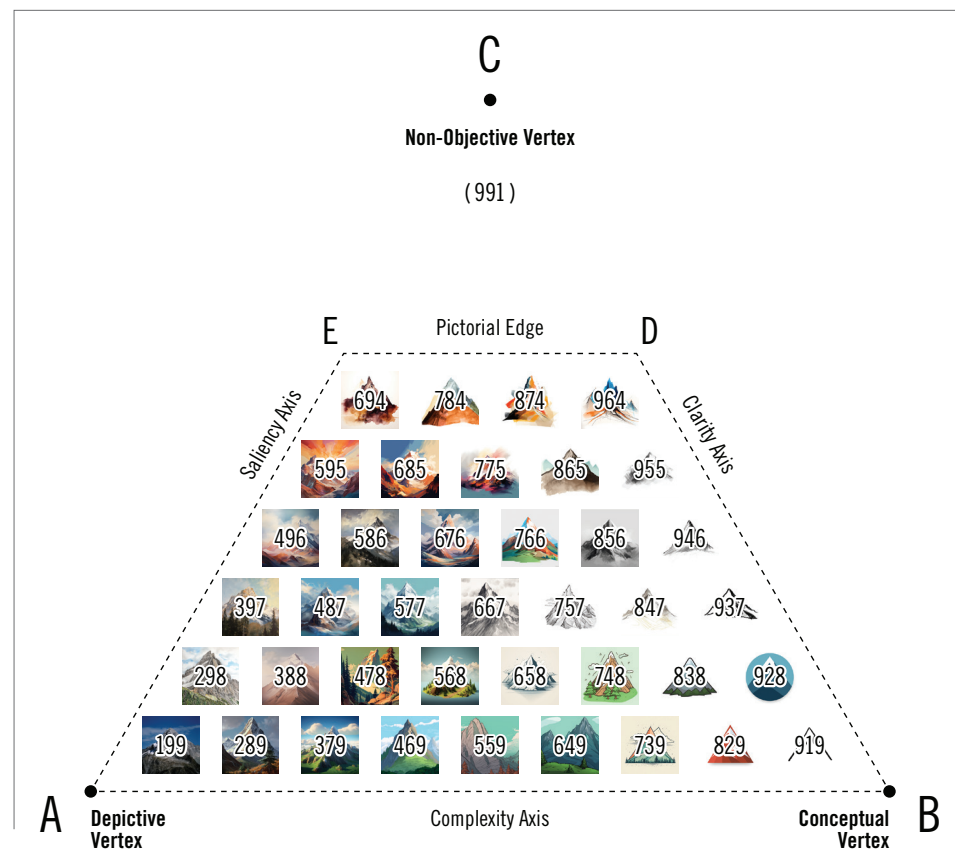


Figure 12.

Incremental positions within the Pictorial Trapezoid. All pictographs were generated with Midjourney.

distances from corners A and B, and starting at the pictorial edge, pictures change from being stylized and nearly indeterminate in terms of what they represent to being readily identifiable through a straightforward presentation. These trends interact and place limits on the represented concepts themselves. For instance, a common line drawing of a chair at P-919 is available to convey the pure *chair* concept. But *chair* cannot be as directly conveyed by a photograph at P-199 that is apparently *the ratty recliner on the corner in front of the frat house I saw when I drove into work yesterday*, nor by a Cubist painting above P-775.

The systematic description embodied in Figures 11–14 operationalizes the Pictorial Trapezoid. Though our plotting of mountain pictographs is rough, the vertices and axes as described permit

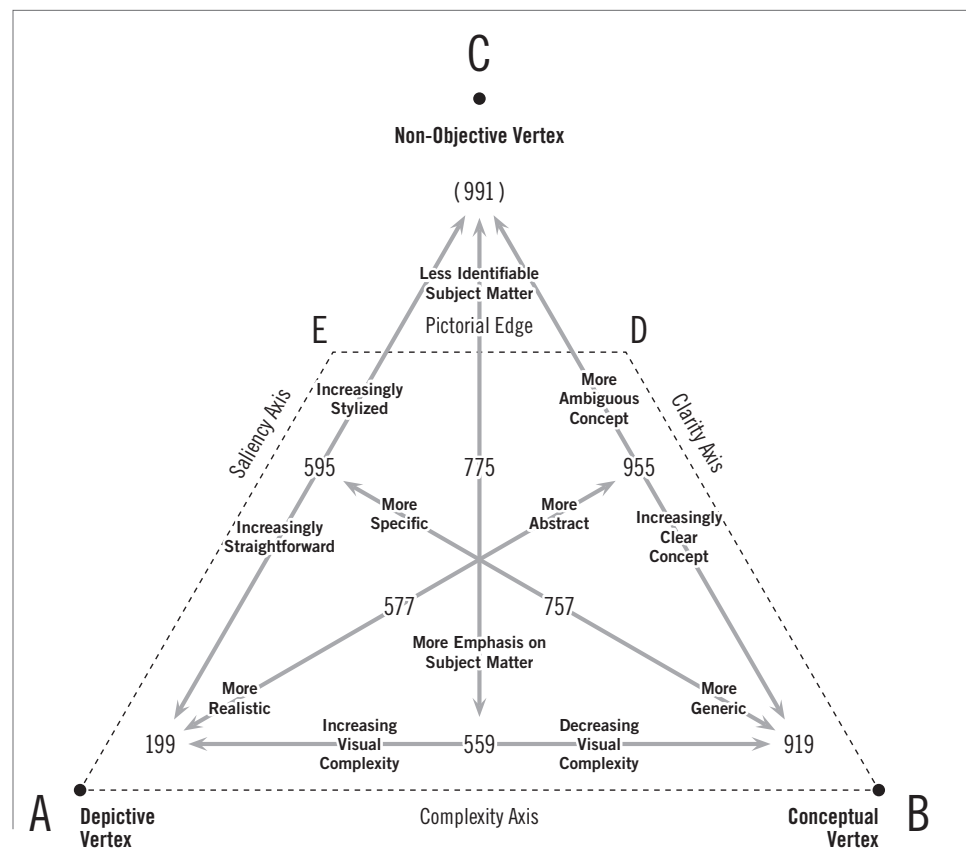


Figure 13.

Semiotic trends in the Pictorial Trapezoid.

an evidence-based plot. Independent Likert scales for the three axes, and for vectors drawn from the three vertices, could be utilized to develop a visual corpus with a variety of pictures at each position, either in a controlled study with a sufficient sample of users, or through extensive work by multiple researchers achieving intercoder reliability. While pilot testing may suggest revised terminology, Table 4 shows how individual dimensions could be used to characterize individual pictures, with only two items apiece necessary for plotting pictures into trapezoid positions. For instance, the pictograph at P-478 in Figure 11 (and unobscured in Figure 10) can be characterized as a “slightly realistic” (A=4), “moderately specific example of” (B=7), and “relatively

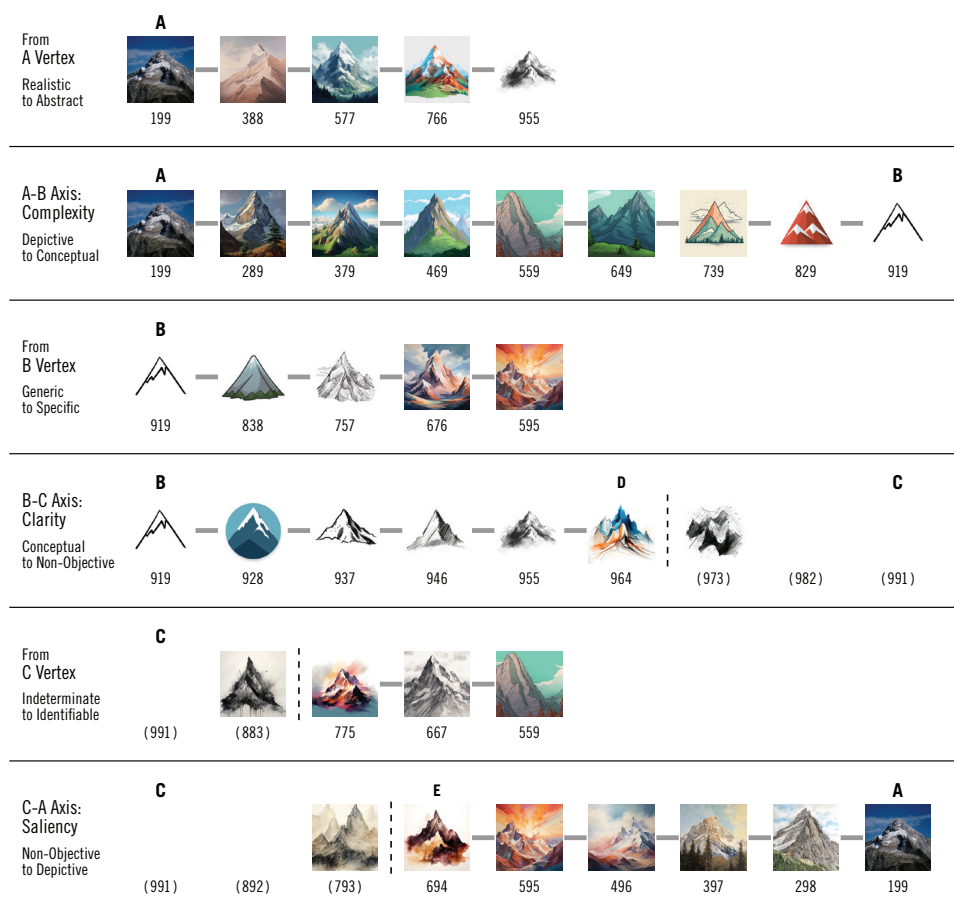


Figure 14.

Axes and vertex vectors derived from the Pictorial Trapezoid for the mountain concept. Dashed lines represent the pictorial edge, beyond which visuals are semiotically unidentifiable or non-objective. Examples of visuals that seemingly pass this threshold and no longer represent mountains are included as reference points (P-793, P-883, and P-973). All pictographs were generated with Midjourney.

	Complexity Axis (A-B)	Conceptual Vertex (B)	Non-Objective Vertex (C)
1	Extremely visually complex	An extremely generic example	Extremely stylized and impossible to determine what the picture represents
2	Very visually complex	A very generic example	
3	Moderately visually complex	A moderately generic example	
4	Slightly visually complex	A slightly generic example	Extremely stylized but identifiable
5	Between complex and simple	Between generic and specific	Very stylized
6	Slightly visually simple	A slightly specific example	Moderately stylized
7	Moderately visually simple	A moderately specific example	Slightly stylized
8	Very visually simple	A very specific example	Relatively straightforward
9	Extremely visually simple	An extremely specific example	Very straightforward

Table 4.

Potential Likert-style scales for coding pictures and, given level, plotting them on the Pictorial Trapezoid. Any two axes and/or vertex vectors, measured independently, can in tandem plot a picture. Prompts would differ; e.g., conceptual vertex: "How generic or specific is the thing this picture represents? For example, Beyoncé and Dolly Parton are specific, while 'person' is generic."

straightforward" (C=8) mountain. As we demonstrated with Figure 9, stylistic distinctions such as "painting" and "cartoon" are expected to skew pictures to a narrower range of values than are represented in the trapezoid overall. A visual corpus of plotted pictures would reveal value ranges according to such stylistic distinctions, which are routinely utilized in text-to-image AI.

The Pictorial Trapezoid was suggested by the exploration of Case 1 as detailed in the previous section, and we have further refined it here to incrementally improve its efficacy in a formative capacity. It emerged from an idiosyncratic investigation, so we next consider its relevance to other investigations that were not responsible for its development. Because the Pictorial Trapezoid effectuates semiotic control, applicability to other investigations is enhanced if it can enforce a theoretically grounded sweet spot of signification, as we found in Case 1 (Figure 6).

Applicability of the Pictorial Trapezoid

To further illustrate the potential of the Pictorial Trapezoid for formative use, we highlight two use cases of text-to-image AI in research contexts. Case 2 is an ongoing doctoral-qualifying investigation to develop and evaluate the efficacy of mediated rescripting, a social-belonging intervention for

Black undergraduate students in engineering. In this project, AI-generated pictures will be created to convey visual metaphors based on students' experiences with stereotype threat. Much like Case 1, Case 2 is explicitly positioned as design research; it considers how visual representational style might be leveraged with established psychological theory to mitigate threats to belonging. Case 3 is a recently published patient diagnostic care intervention for individuals experiencing rare visual conditions (Waisberg et al., 2023). Researchers used AI-generated pictures to convey patient-reported descriptions of the perceptual experiences associated with such phenomena to promote empathy among clinicians. Visual representations of optical phenomena may be especially helpful in understanding unique visual conditions not represented adequately through standard visual test results. Increased clinician empathy has consistently been linked to positive outcomes for patients and healthcare professionals (Moudatsou et al., 2020).

In both additional cases, coordinating representational and conceptual fidelity is vital to creating effective visualizations that provide benefits to users. For the initial Case 1, we divided the process of creating visual conceptual pegs into two phases: choosing a specific concept (selection), and depicting that concept (visualization). We also outlined three criteria for conceptual pegs: essential, distinctive, and compact. We revisit that structure here. Finally, we discuss where sweet spots, or ideal ranges of visual representational style for trained AI models, might fall within the Pictorial Trapezoid for the additional cases.

Case 2: Visualizing Belonging for Black Undergraduate Engineering Students

Despite substantial growth in science, technology, engineering, and math (STEM) fields at collegiate and professional levels, Black students remain underrepresented among engineering bachelor degree recipients in comparison with their peers (NCSES, 2023). Some researchers have attributed this pattern to negative experiences threatening students' sense of belonging within engineering programs (Lee et al., 2020; Strayhorn, 2018). These threats influence the degree to which students believe their racial identity is compatible with the engineering discipline. Researchers suggest developing interventions that reduce students' doubts about social belonging to mitigate the negative effects of stereotype threat (Totonchi, 2021; Walton & Cohen, 2007). In this context, sense of belonging refers to the degree to which individuals see themselves as "socially connected" in their academic environment (Walton & Cohen, 2007, p. 82).

The second author is presently engaged in Case 2, a study that establishes and tests mediated rescripting as a belonging intervention for Black undergraduate engineering students. (Participant recruitment is still underway.) Mediated rescripting is based on imagery

rescripting, a therapeutic method that involves patients “visually recalling and reexperiencing” mental images and thoughts related to distressing events, and then with the help of a therapist changing the imagery “to produce a more favorable outcome” (Rusch et al., 2000, p. 9). Mediated rescripting seeks to increase access to imagery rescripting by utilizing text-to-image AI technology to generate pictures capturing personally relevant mental imagery formalized as visual metaphors; which could in the future be operationalized in an app with an embedded and trained AI. Case 2 requires participants to create pictures based on metaphors derived from their personal experiences of stereotype threat, with the goal of achieving visual-conceptual correspondence. The premise is that having Black engineering students engage strategically with metaphorical imagery can strengthen their sense of belonging and engineering identity, thus decreasing anxiety and stress.

The Case 2 study consists of two distinct phases followed by a period of ongoing engagement. Participants are first asked to reflect on challenges faced in their engineering program, particularly those related to their sense of belonging. Next, participants are asked to write corresponding metaphors using the prompt: “x is like y, because...” or “x is y, because...” In a subsequent stage, participants collaborate with the session facilitator to write prompts that adequately depict their metaphors.

Examples of pictures likely to be created are shown in Table 5. These pictures were generated by the second author using DALL-E 2, as proof of concept explorations and not outcomes from the actual study — the completed study may result in pictures featuring a wider stylistic range. The pictures are based on metaphors formulated from themes of belonging expressed by Black undergraduate STEM students in a study by Strayhorn (2011), which sought to reveal how Black students describe their sense of belonging, including how it is “inhibited or inspired” (p. 220). Each set of pictures captures two versions of the same metaphor: the negative picture is meant to illustrate the “as is” experience of threatened belonging, while the alternative picture is meant to illustrate a reimagined, more positive sense of belonging. Prompting participants to identify challenges to their sense of belonging will help evoke the initial negative imagery. In contrast, prompts that ask them to think about times their sense of belonging has been affirmed or to consider what that affirmation might look like will help evoke the alternative positive imagery. Picture 1a in Table 5 captures the negative experience of feeling ignored and unseen by instructors and peers in class, with the student’s face obscured with a white, hazy cloud, while picture 1b portrays the alternative experience of feeling confident and visible, with the student standing confidently in front of a crowd. While the context and style of the two pictures differ, they work together to communicate alternative realities for the same fundamental

Experience	Metaphor	(a) Negative	(b) Alternative
(1) Feeling ignored or unseen by instructors and peers	Invisibility		
(2) Being inundated with constant reminders of negative perceptions and stereotypes about Black students	Stuck in a constant stream or barrage of attacking orbs		

Table 5.

Sets of metaphorical imagery generated using DALL-E 2, representing experiences of being ignored (1a & 1b) and being inundated with negative stereotypes (2a & 2b), based on themes identified by Strayhorn (2011).

concept. Picture 2a in Table 5 captures a student’s experience of being attacked with negative stereotypes about Black students, as represented by orbs, while picture 2b portrays a student controlling or harnessing power over the orbs. The alternative picture is achieved by altering the subject’s facial expression and hands. A concept primarily links the first pair (1a and 1b), while both concept and direct visualization link the second pair (2a and 2b).

While research suggests that effective imagery rescripting relies heavily on the “goodness of simulation” or the fidelity of the simulated experience, few picture-based iterations of imagery rescripting problematize conceptual and representational fidelity or formalize criteria for picture-making (Looney et al., 2021). In Case 1, we identified two phases of conceptual peg creation: concept selection and concept visualization. For Case 2 the structure holds, but we use different qualifiers: experience selection, metaphor selection, and metaphor visualization. The proposed selection and visualization criteria for metaphors in mediated rescripting are provided in Table 6. The criteria for selection require that the experience be personally relevant, emotionally salient, and can be expressed metaphorically. During the intervention, participants are asked to reflect on personal experiences related to belonging in engineering. While

Criterion	Experience Selection	Metaphor Selection	Metaphor Visualization
Personally Relevant	Experience must be derived directly from the individual's thoughts, feelings, and behaviors	—	—
Metaphorical	Experience must be able to be expressed metaphorically	Metaphor can be mapped back onto the experience through entailments	Picture contains key visual elements referencing source and target domains
Emotionally Salient	Experience should evoke emotion and be memorable	Metaphor evokes emotions similar to those elicited by the actual experience	Picture features intentional stylistic choices that convey emotions; feels right

Table 6.

Proposed concept selection and visualization criteria for visualized metaphors used for mediated rescripting.

others may share similar experiences, they must be directly related to the participant's personal thoughts, feelings, and behaviors. Once formulated through the elicited metaphor prompt (Low, 2017), the experience should adequately map onto components of a conceptual metaphor (Lakoff & Johnson, 1980), which includes a target domain, source domain, and entailments. The target domain is the experience being described or the general sense of belonging; the source domain is a scenario that can be used to provide insight into the target domain; and *entailment* refers to characteristics of the source domain that can be logically mapped onto the target domain, along with the implications of these mappings. The selected experience and metaphor should be memorable and evoke emotions similar to those elicited by the actual experience. Corresponding visual criteria include key visual elements that reference source and target domains and stylistic choices that convey corresponding emotions.

There is thus considerable interaction between the conceptual (i.e., metaphorical) and the depictive (i.e., personal experiential) in the visualized products of mediated rescripting. In both examples in Table 5, the student is themselves represented, along with an environment (1a and 1b) or antagonists (2a and 2b). After plotting a few visualized metaphor examples, we suspect that a reasonable sweet spot might fall in the Pictorial Trapezoid's lower left as shown in gray in Figure 15; other areas (in white) may fail to meet one of the selection or visualization criteria. We do not extend the gray area fully to the depictive vertex (lower left corner), because pictures in this area may be too realistic and straightforward to be read as metaphorical. Additionally, current text-to-image generators may struggle to adequately render pictures that must depict reality at too high a degree of representational fidelity, as the systems are notorious for distorting body

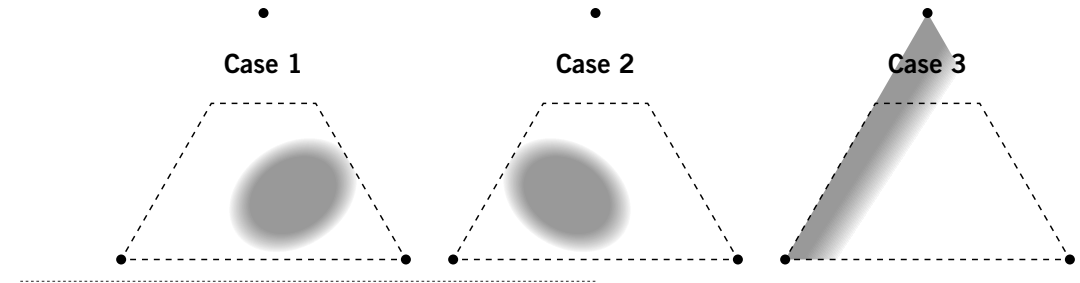


Figure 15.

Possible sweet spots for the three case studies: Case 1, a speculative interface project for intelligence analysis, through which the Pictorial Trapezoid was suggested; Case 2, an in-progress intervention on sense of belonging; and Case 3, a study concerning patients' visual phenomena (Waisberg et al., 2023).

and facial features (Strickland, 2022). McCloud (1993) hypothesized that the comic art strategy of visually simplifying protagonists in contrast to more detailed environments and antagonists allows viewers to better project themselves into the character of a protagonist. If McCloud's premise holds, this would complicate plotting a mediated rescripting picture into a single position in the Pictorial Trapezoid. Thus, while Case 2 is elucidated by the Pictorial Trapezoid, it also may suggest that the dual nature of source and target in visual metaphor can complicate plotting and that positioning within the trapezoid will benefit from further consideration. (We consider this merely a possibility, and have not plotted separate areas in Figure 15.)

Case 3: Visualizing Unique Neuro-Ophthalmic Visual Phenomena

Case 3 concerns the visualization of neuro-ophthalmic visual phenomena experienced by patients with oscillopsia stemming from multiple sclerosis, Charles Bonnet syndrome, and the Pulfrich phenomenon (Waisberg et al., 2023). Oscillopsia is characterized by perceptual illusions of movement and elaborate hallucinations. Empathy, or the ability to relate and understand what others are experiencing, is a valuable skill for health care providers. Prior research shows that increased empathy positively influences patient outcomes, improves patient compliance, and minimizes claims of malpractice (Hickson et al., 2002; Kelley et al., 2014; Riess, 2015). For healthcare providers who work with patients experiencing rare and unique visual conditions, understanding the full patient experience is especially challenging; after all, providers cannot see through anyone else's eyes, only their own. To help bridge this gap, Waisberg et al. (2023) propose using text-to-image AI generators to visualize patient-reported descriptions of unique neuro-ophthalmic visual phenomena.

In an earlier exploration, Balas and Micieli (2022) used text-to-image AI to illustrate the experience of visual snow syndrome based on a verbal description provided by a 17-year-old patient. Visual snow

syndrome (VSS) is a neurological condition characterized by a persistent flickering of small dots throughout the visual field, like static on a TV. VSS is an ideal candidate for visualization because diagnoses of the condition cannot be made based on exams. Clinicians instead rely on patient descriptions. Further, Balas and Micieli suggest that objectively visualizing the experience of visual conditions can be difficult because different patients are likely to describe them using unique language. To address this, they used popular text-to-image generators, including DALL-E, Midjourney, and Stable Diffusion, to generate representative pictures for educational use. Prompts were written based on a patient description and narrowed to a selection of two pictures they thought conformed to the descriptions and established visualizations of the phenomenon.

Similar to Case 2, ensuring that AI-generated pictures appropriately represent patient descriptions was a primary concern for Case 3. Medical researchers do not problematize conceptual and representational fidelity or formalize criteria beyond suggesting that descriptions be picturable and consistent with existing understandings of the phenomenon (Balas & Micieli, 2022; Zhu et al., 2007). For Case 3, we keep with the two-phase structure but append new qualifiers, as description selection and perception visualization, with three new criteria for pictures conveying

Table 7.

*Proposed concept selection and
visualization criteria for pictures
conveying visual conditions.*

Criterion	Description Selection	Perception Visualization
Unique	Description represents personal concerns	—
Specific	Description is both thorough and precise	—
Rich	Description includes text that can be matched with corresponding visuals	Picture is a complete and detailed representation of the patient's visual field

visual conditions: unique, specific, and rich. Phases and criteria are provided in Table 7.

After plotting sample patient pictures, we believe that a reasonable sweet spot might track along the saliency axis (Figure 13) of the Pictorial Trapezoid, as shown in Figure 15. Because patients in Case 3 are conveying optical phenomena, depictions should be rich, and thus AI should operate within the indicated portion of the trapezoid. However, in reflecting upon the nature of neuro-ophthalmic visual phenomena, we believe there may be extreme cases in which a patient may see artifacts of light more than distinct features of a world around them. In this case visualizations may pass over the pictorial edge and, from others' perspectives, be non-objective — that is, such visualizations may not strictly be “pictures” according to our use of the term. As with the potential dual-plotting we noted with Case 2, this is a complicating factor for the Pictorial Trapezoid that possibly challenges its fundamental trapezoidal form (reactivating the portion of the Big Triangle we removed). In this case we have reflected the complication in Figure 15 with a sweet spot that extends outward.

Despite these complications, Cases 2 and 3 are elucidated by the Pictorial Trapezoid, and the distinction of selection and visualization we established in Case 1 describes criteria for generative AI when context-specific qualifiers are utilized — e.g., concept selection in Case 1, and description selection in Case 3. We now address technological dynamics upon which use of the Pictorial Trapezoid in research contexts like these cases depends.

Training AI with the Pictorial Trapezoid

In our three research use cases, humans utilize text-to-image generative models to produce unique visualizations that represent either information (Case 1) or experience (Cases 2 and 3). In Case 1, per subsequent prototyping studies not described here, the prompt consists of textual information highlighted by the user, from which an AI must derive a discrete concept. For Case 2, students describe their personal metaphorical conceptualizations through targeted prompting, and for Case 3 the patient's verbal description of their optical phenomena is converted into a text-based prompt. In none of the three cases is the user asked to be a prompt engineer — they are assumed to have very little, if any, expertise around prompting.

Popular text-to-image generators like Midjourney, DALL-E, and Stable Diffusion combine large language models (LLMs) with diffusion models. An LLM is a type of deep learning model that is trained to detect and learn patterns between sequences of words. Using these learned

patterns, the LLM can analyze a text prompt and convert the text's meaning into machine readable output. This output then guides a diffusion model. Diffusion models are trained through a process in which noise has been added to a labeled picture to break it down. The model then reverts the process, removing the visual noise until a picture is produced that corresponds to the initial text prompt (Lyu et al., 2022). In this way, text-to-image generators produce unique pictures in response to text input.

To put into practice Cases 1–3, developers could build on an existing foundation model like Midjourney. Foundation models are trained on huge quantities of unlabeled data at scale, typically through self-supervised learning. As a result these generalized models can be adapted to a variety of downstream tasks (CRFM, 2021). If the specified downstream task requires more customized results, developers can build on top of existing foundation models to refine output. Each of Cases 1–3 as described would require this type of customization to constrain visual stylistic range.

Such customization could be achieved through fine-tuning and prompt-tuning. Both techniques, working together, can improve the results of a foundation model for specific applications. To fine-tune a model, developers prepare a new model to work with the foundation model, and this fine-tuned model can be trained on a new dataset (OpenAI, n.d.). When applied in Cases 1–3, the model could be trained using labeled pictures that fall within the respective sweet spot of the Pictorial Trapezoid. The resulting text-to-image generative model would produce new unique pictures that fit desirable parameters. The system could continue to learn and refine output once it was put into use, even identifying visual preferences of individuals via input through the interface.

In concert with fine-tuning, prompt-tuning could afford another efficient method of narrowing image output to a specified stylistic range. Lester et al. (2021) characterize prompt-tuning as “a simple yet effective mechanism for learning ‘soft prompts’ to condition frozen language models to perform specific downstream tasks” (p. 3045); the mechanism automates the textual prompting process, enabling human users to input fairly simple or straightforward prompts. Unbeknownst to these users, automated prompts on the backend narrow the range of possible image output, thus negating the need to enter complex paragraphs of text to control stylistic variation. Prompt-tuning may be particularly appropriate for research use cases such as Cases 2 and 3, in which a student or patient who is not particularly skilled at prompting needs to communicate a rich experience. This can be a back-and-forth process, with the AI producing an initial visualization within the desired stylistic range, and the user offering additional text prompts to refine it (or for Case 3, a health care provider

describing a generated picture back to a patient in their own words, when the patient cannot see the picture clearly, in order for the patient to clarify).

Building on existing large foundation models and relying on fine-tuning and prompt-tuning to narrow the range of visual output would be a relatively inexpensive process that would not require retraining the foundation models. However, this approach does present security concerns when foundation systems like ChatGPT or Midjourney are run by third-party vendors. In such a scenario, users share data that is then stored on servers owned by the companies that operate them (Ray, 2023). Issues of privacy and security that arise in both intelligence (Case 1) and healthcare (Cases 2 and 3) contexts would require developers instead to build off of open source foundation models to produce systems that could be deployed and controlled locally. Once such models are tuned, text-to-image generators provide a clear avenue for operationalizing use.

Conclusion

We have provided a revision of McCloud's (1993) influential Big Triangle in the form of the Pictorial Trapezoid. We detailed the design research context from which the Pictorial Trapezoid emerged as a first case study, and applied it to an additional two research use cases. While these cases provide some breadth to our investigation, they nonetheless represent a limitation. The value of the Pictorial Trapezoid currently rests on argumentation, and future work is necessary to empirically validate its positional scheme beyond impressions of the trapezoid's apparent uses as demonstrated in Cases 1–3.

We have described the Pictorial Trapezoid systematically, permitting its use not only in gauging semiotic factors such as representational fidelity and conceptual fidelity in pictures (i.e., summative analysis), but in teaming with generative AI to facilitate principled visualization in exacting research contexts (i.e., formative use). Fulfilling the potential of the latter would require training an AI and learning from that process. To that end, we have provided copious descriptive figures and outlined options and challenges for operationalizing the Pictorial Trapezoid for text-to-image generative models. We hope this helps designers cross the threshold from a sophisticated understanding of visualization to engagement in an exciting and impactful technological revolution, in collaboration with computer scientists and other experts. Generative AI is and will continue to be utilized in research contexts across disciplines, and where visualization is concerned, a high degree of semiotic precision will be necessary to maximize the impact of interventions. The Pictorial Trapezoid can be one tool in achieving that precision.

Acknowledgments

This material is based upon work done, in whole or in part, in coordination with the Department of Defense (DoD). Any opinions, findings, conclusions, or recommendations expressed in this material are those of the authors and do not necessarily reflect the views of the DoD and/or any agency or entity of the United States Government. In addition to the authors, Isha Parate was involved in the Case 1 investigation.

References

- Anderson, R. C., Goetz, E. T., Pichert, J. W., & Halff, H. M. (1977). Two faces of the conceptual peg hypothesis. *Journal of Experimental Psychology: Human Learning and Memory*, 3(2), 142–149. <https://psycnet.apa.org/doi/10.1037/0278-7393.3.2.142>
- Baddeley, A. D. (2018). *Exploring working memory: Selected works of Alan Baddeley*. Routledge.
- Balas, M., & Micieli, J. A. (2022). Visual snow syndrome: Use of text-to-image artificial intelligence models to improve the patient perspective. *Canadian Journal of Neurological Sciences*, 1–2. <https://doi.org/10.1017/cjn.2022.317>
- Burks, A. W. (1949). Icon, index, and symbol. *Philosophy and Phenomenological Research*, 9(4), 673–689. <https://doi.org/10.2307/2103298>
- Center for Research on Foundation Models (CRFM) (2022). *On the opportunities and risks of foundation models: Special report*. Stanford Institute for Human-Centered Artificial Intelligence. <https://arxiv.org/abs/2108.07258>
- Cunha, J. M. (2022). *Generation of concept representative symbols: Towards visual conceptual blending* [dissertation]. Universidade de

Coimbra, Portugal.

- Hickson, G. B., Federspiel, C. F., Pichert, J. W., Miller, C. S., Gauld-Jaeger, J., & Bost, P. (2002). Patient complaints and malpractice risk. *Journal of the American Medical Association*, 287(22), 2951–2957. <https://doi.org/10.1001/jama.287.22.2951>
- Kelley, J. M., Kraft-Todd, G., Schapira, L., Kossowsky, J., & Riess, H. (2014). The influence of the patient-clinician relationship on healthcare outcomes: A systematic review and meta-analysis of randomized controlled trials. *PLOS ONE*, 9(4), e94207. <https://doi.org/10.1371/journal.pone.0094207>
- Lakoff, G., & Johnson, M. (1980/2003). *Metaphors we live by*. University of Chicago Press.
- Lee, M. J., Collins, J. D., Harwood, S. A., Mendenhall, R., & Hunt, M. B. (2020). If you aren't White, Asian or Indian, you aren't an engineer: Racial microaggressions in STEM education. *International Journal of STEM Education*, 7(1), 1–16. <https://doi.org/10.1186/s40594-020-00241-4>
- Lester, B., Al-Rfou, R., Constant, N. (2021). The power of scale for parameter-efficient prompt tuning. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing* (pp. 3045–3059). <https://aclanthology.org/2021.emnlp-main.243.pdf>
- Liu, V., & Chilton, L. B. (2022). Design guidelines for prompt engineering text-to-image generative models. *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems* (pp. 1–23). <https://doi.org/10.1145/3491102.3501825>
- Liu, V., Vermeulen, J., Fitzmaurice, G., & Matejka, J. (2023). 3DALL-E: Integrating text-to-image AI in 3D design workflows. *Proceedings of the 2023 ACM Designing Interactive Systems Conference*, 1955–1977. <https://doi.org/10.1145/3563657.3596098>
- Looney, K., El-Leithy, S., & Brown, G. (2021). The role of simulation in imagery rescripting for post-traumatic stress disorder: A single case series. *Behavioural and Cognitive Psychotherapy*, 49(3), 257–271. <https://doi.org/10.1017/S1352465820000806>
- Low, G. (2017). Eliciting metaphor in education research: Is it really worth the effort? In F. Ervas, E. Gola, & M. G. Rossi (Eds.), *Metaphor in communication, science and education* (pp. 249–266). Walter de

Gruyter.

Lyu, Y., Wang, X., Lin, R., & Wu, J. (2022). Communication in human–AI co-creation: Perceptual analysis of paintings generated by text-to-image system. *Applied Sciences*, 12(22), 11312. <https://doi.org/10.3390/app122211312>

Martins, T., Cunha, J. M., Correia, J., Machado, P. (2023). Towards the evolution of prompts with MetaPrompter. In C. Johnson, N. Rodríguez-Fernández, & S. M. Rebelo (Eds.), *Artificial intelligence in music, sound, art and design*. Springer. https://doi.org/10.1007/978-3-031-29956-8_12

Midjourney (n.d.[a]). Midjourney stylize parameter (webpage). <https://docs.midjourney.com/docs/stylize>

Midjourney (n.d.[b]). Midjourney parameter list (webpage). <https://docs.midjourney.com/docs/parameter-list>

McCloud, S. (1993). *Understanding comics: The invisible art*. Harper Perennial.

McCloud, S. (n.d.[a]). The Big Triangle (webpage). <https://www.scottmccloud.com/4-inventions/triangle/index.html>

McCloud, S. (n.d.[b]). The Big Triangle (cont) [webpage]. <https://www.scottmccloud.com/4-inventions/triangle/12.html>

Moudatsou, M., Stavropoulou, A., Philalithis, A., & Koukouli, S. (2020). The role of empathy in health and social care professionals. *Healthcare*, 8(1), 26. <https://doi.org/10.3390/healthcare8010026>

National Center for Science and Engineering Statistics (NCSES) (2023). Diversity and STEM: Women, minorities, and persons with disabilities 2023. *Special Report NSF 23-315*. Alexandria, VA: National Science Foundation. <https://ncses.nsf.gov/wmpd>

OpenAI (n.d.). Fine tuning (webpage). <https://platform.openai.com/docs/guides/fine-tuning>

Oppenlaender, J., Linder, R., & Silvennoinen, J. (2023). Prompting AI art: An investigation into the creative skill of prompt engineering. *arXiv*. <https://doi.org/10.48550/arXiv.2303.13534>

Peterson, M., Delgado, C., Tang, K.-S., Bordas, C., & Norville, K. (2021). A taxonomy of cognitive image functions for science curriculum materials: Identifying and creating 'performative' visual displays. *International Journal of Science Education*, 43(2), 314–343. <https://doi.org/10.1080/09500693.2020.1868609>

doi.org/10.1080/09500693.2020.1868609

Pettersson, R. (2015a). *Information design 1: Message design*. International Institute for Information Design. <https://www.iiid.net/PublicLibrary/Pettersson-Rune-ID1-Message-Design.pdf>

Pettersson, R. (2015b). *Information design 2: Text design*. International Institute for Information Design. <https://www.iiid.net/PublicLibrary/Pettersson-Rune-ID2-Text-Design.pdf>

Pettersson, R. (2015c). *Information design 3: Image design*. International Institute for Information Design. <https://www.iiid.net/PublicLibrary/Pettersson-Rune-ID3-Image-Design.pdf>

Pettersson, R. (2015d). *Information design 4: Graphic design*. International Institute for Information Design. <https://www.iiid.net/PublicLibrary/Pettersson-Rune-ID4-Graphic-Design.pdf>

Pettersson, R. (2015e). *Information design 5: Cognition*. International Institute for Information Design. <https://www.iiid.net/PublicLibrary/Pettersson-Rune-ID5-Cognition.pdf>

Pettersson, R. (2015f). *Information design 6: Predecessors & pioneers*. International Institute for Information Design. <https://www.iiid.net/PublicLibrary/Pettersson-Rune-ID6-Predecessors.pdf>

Ray, S. (2023). Samsung bans ChatGPT among employees after sensitive code leak. *Forbes*, May 2, 2023. <https://www.forbes.com/sites/siladityaray/2023/05/02/samsung-bans-chatgpt-and-other-chatbots-for-employees-after-sensitive-code-leak/?sh=1686e72c6078>

Riess, H. (2015). The impact of clinical empathy on patients and clinicians: Understanding empathy's side effects. *AJOB Neuroscience*, 6(3), 51–53. <https://doi.org/10.1080/21507740.2015.1052591>

Rosch, E. R. (1978). Principles of categorization. In E. R. Rosch & B. B. Lloyd (Eds.), *Cognition and categorization* (pp. 27–48). John Wiley & Sons.

Rosenblatt, L. M. (1978). *The reader, the text, the poem: The transactional theory of the literary work*. Sage.

Rusch, M. D., Grunert, B. K., Mendelsohn, R. A., & Smucker, M. R. (2000). Imagery rescripting for recurrent, distressing images. *Cognitive and Behavioral Practice*, 7(2), 173–182. [https://doi.org/10.1016/S1077-7229\(00\)80028-0](https://doi.org/10.1016/S1077-7229(00)80028-0)

Sadoski, M., & Paivio, A. (2001). *Imagery and text: A dual coding theory of*

reading and writing. Lawrence Erlbaum.

Sheldon, M. (2023). Anatomy of a shelling: How Russian rocket artillery struck Mykolaiv. *Bellingcat* (January 27, 2023). <https://www.bellingcat.com/news/2023/01/27/anatomy-of-a-shelling-how-russian-rocket-artillery-struck-mykolaiv/>

Spoehr, K. T., & Lehmkuhle, S. W. (1982). *Visual information processing*. W. H. Freeman and Company.

Strayhorn, T. L. (2011). Sense of belonging and African-American student success in STEM: Comparative insights between men and women. In H. T. Frierson & W. F. Tate (Eds.), *Beyond stock stories and folktales: African Americans' paths to STEM fields* (Vol. 11, pp. 213–226). Emerald Group Publishing Limited. [https://doi.org/10.1108/S1479-3644\(2011\)0000011014](https://doi.org/10.1108/S1479-3644(2011)0000011014)

Strayhorn, T. L. (2018). *College students' sense of belonging: A key to educational success for all students* (2nd ed.). Routledge. <https://doi.org/10.4324/9781315297293>

Strickland, E. (2022). DALL-E 2's failures are the most interesting thing about it. *IEEE Spectrum*. <https://spectrum.ieee.org/openai-dall-e-2>

Totonchi, D. A., Perez, T., Lee, Y., Robinson, K. A., & Linnenbrink-Garcia, L. (2021). The role of stereotype threat in ethnically minoritized students' science motivation: A four-year longitudinal study of achievement and persistence in STEM. *Contemporary Educational Psychology*, 67, 102015. <https://doi.org/10.1016/j.cedpsych.2021.102015>

Turchi, T., Carta, S., Ambrosini, L., & Malizia, A. (2023). Human-AI co-creation: Evaluating the impact of large-scale text-to-image generative models on the creative process. In L. D. Spano, A. Schmidt, C. Santoro, & S. Stumpf (Eds.), *End-user development* (pp. 35–51). Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-34433-6_3

Waisberg, E., Ong, J., Masalkhi, M., Zaman, N., Kamran, S. A., Sarker, P., Lee, A. G., & Tavakkoli, A. (2023). Text-to-image artificial intelligence to aid clinicians in perceiving unique neuro-ophthalmic visual phenomena. *Irish Journal of Medical Science*. <https://doi.org/10.1007/s11845-023-03315-8>

[org/10.1007/s11845-023-03315-8](https://doi.org/10.1007/s11845-023-03315-8)

Walton, G. M., & Cohen, G. L. (2007). A question of belonging: Race, social fit, and achievement. *Journal of Personality and Social Psychology*, 92(1), 82–96. <https://doi.org/10.1037/0022-3514.92.1.82>

Zhu, X., Goldberg, A. B., Eldawy, M., Dyer, C. R., & Strock, B. (2007). A text-to-picture synthesis system for augmenting communication. In *Proceedings of the 22nd national conference on artificial intelligence* (pp. 1590–1595). AAAI Press. <https://cdn.aaai.org/AAAI/2007/AAAI07-252.pdf>

Authors

Matthew Peterson

Matthew Peterson, PhD, is Associate Professor of Graphic & Experience Design at North Carolina State University. His research focuses on visual representation, especially in the development of novel interfaces and environments. He integrates design with other disciplines, with publications, projects, and proposals in collaboration with experts in STEM education, engineering, psychology, advertising, biology, physics, and data science. Peterson is also engaged in advocating for and describing a more rigorous design discipline through publications and presentations.

Ashley L. Anderson

Ashley L. Anderson is a PhD in Design candidate at North Carolina State University, focusing on issues of visual representation, visual metaphor, generative AI, and design for social psychological intervention. Prior to entering the PhD program, Anderson earned her Master of Graphic Design at NC State. Her current dissertation work evaluates the efficacy of mediated rescripting, a picture-based intervention designed to improve belonging for Black undergraduate engineering students.

Kayla Rondinelli

Kayla Rondinelli is a UX designer, graphics illustrator, and artist based in Raleigh, NC. She is a Master of Graphic & Experience Design student at North Carolina State University, and works as a graduate research assistant in the College of Design and as a freelance graphic designer. Her interest in reducing the carbon impact of the built environment has pushed her to focus on sustainable design practices. Rondinelli hopes her work provides a platform for conversations concerning environmental conservation, equitable distribution of the burdens of climate change, and preservation of the natural world.

Helen Armstrong

Helen Armstrong is Professor of Graphic & Experience Design at North Carolina State University, where she is director of the MGXD program. Her research focuses on accessible design, digital rights, and machine learning. Her books include *Graphic Design Theory*, *Digital Design Theory*, *Participate*, and *Big Data, Big Design: Why Designers Should Care About Artificial*

Intelligence. Armstrong is a past member of the AIGA National Board of Directors, the editorial board of *Design and Culture*, and a former co-chair of the AIGA Design Educators Community. Her work has been recognized by *Print* and *HOW*, and included in numerous publications in the U.S. and the U.K. Armstrong is the proud mom of a kid with disabilities and is a fierce advocate for designing interfaces and experiences that are inclusive and intelligent.

Visible Language

Online Archive Improvement Study:

Envisioning
Visible Language's
Next Generation
Online Archive.

D.J. Trischler

Abstract

This study aimed to identify strategies to improve *Visible Language's* online archive design for a broader readership. The author used a combination of exploratory methods with current and potential *Visible Language* readers to identify designers' current and desired uses of design knowledge. The findings led to three key design strategies: article curation, advanced search filters, and a dynamic reading experience. The strategies extend beyond *Visible Language* to design archives that wish to improve and broaden engagement with their content.

Keywords

Academic Journals
Academic Publishing
Reading Experience
Design Knowledge
Accessible Scholarship
Online Archives

Introduction

Overview

Visible Language is the longest-standing scholarly journal about graphic and communication design. As such, the journal has accumulated an online archive that houses decades of valuable design knowledge defining and reflecting an ever-evolving discipline.

Subscribers to *Visible Language* are mostly university libraries and individuals within academia. The author speculates that most online readers are academics and access articles through databases instead of directly through the *Visible Language* online archive.

This study sought to identify strategies to improve *Visible Language's* online archive design for a broader readership. The questions in Table 1 structured this two-phase study that explored designers' current uses of design knowledge (Phase One) and their desired uses (Phase Two). The author used a combination of exploratory methods with current and potential *Visible Language* readers to answer these questions. Their responses laid the foundations for an improved *Visible Language* online archive via three design strategies with diagrams visualizing their meanings. The strategies extend beyond *Visible Language* to design archives that wish to improve and broaden engagement with their content.

Table 1

The study's research questions with
their corresponding research phase.

#	Questions	Phase
1	How do designers currently use design knowledge published in academic journals and other publications about graphic and communication design?	One
2	Which journals and publications do they go to regularly?	One
3	What kinds of articles are they looking for?	One/Two
4	How are the articles read/used?	One/Two
5	How could designers/would designers like to use design knowledge?	Two
6	What features would make the VL online archive more useful?	Two

Study Definitions

Visible Language

Visible Language (VL) aims to help inform authors, editors, and subjects as they collaborate to make design knowledge. Though it has evolved, VL has held a core purpose to increase designers' knowledge. VL's founding purpose in 1967, when it was named *The Journal of Typographic Research*, was "to report and to encourage scientific investigation of our alphabetic and related symbols (Wrolstad, 1967)." The journal renamed itself to *Visible Language* four years later because of the constraints of the previous name. The new name opened the content boundaries, and the new tagline reflected its expansion — "The Journal for Research on the Visual Media of Language Expression" (Wrolstad, 1971). The editor wrote at that time, "This Journal represents what could be the first concerted effort to organize our investigation of every respect of this visual medium of language expression." VL sought to present "a range of interests" regarding the conception, formation, reception, and interpretation of writing, typography, signs, and more (Wrolstad, 1971).

In 1987, Sharon Helmer Poggenpohl, VL's new editor, changed the tagline to "the quarterly concerned with all that is involved with our being literate" (1987). The "range of interests" expanded again, which may reflect significant technological shifts. VL covers during this era reflect the tension and possibilities caused by desktop computers. The articles leaned theoretical as designers sought to orient themselves within an ever-changing discipline and world.

Currently — under VL's third editor, Mike Zender — the tagline reads "the journal of visual communication research" (2022). The online archive states, "*Visible Language* supports the community of communication design scholars, researchers, and practitioners through the publication of rigorous, relevant communication design research" (Zender, n.d.). This description is more targeted than previously. The change reflects broader influences like more scholarly design journals and a shifting discipline, including the rise of UX/UI research. Zender and the editorial team, of which I, the author, am a part, encourage design researchers toward empirical, evidence-based design studies reminiscent of the original mandate toward "scientific investigation."

Design

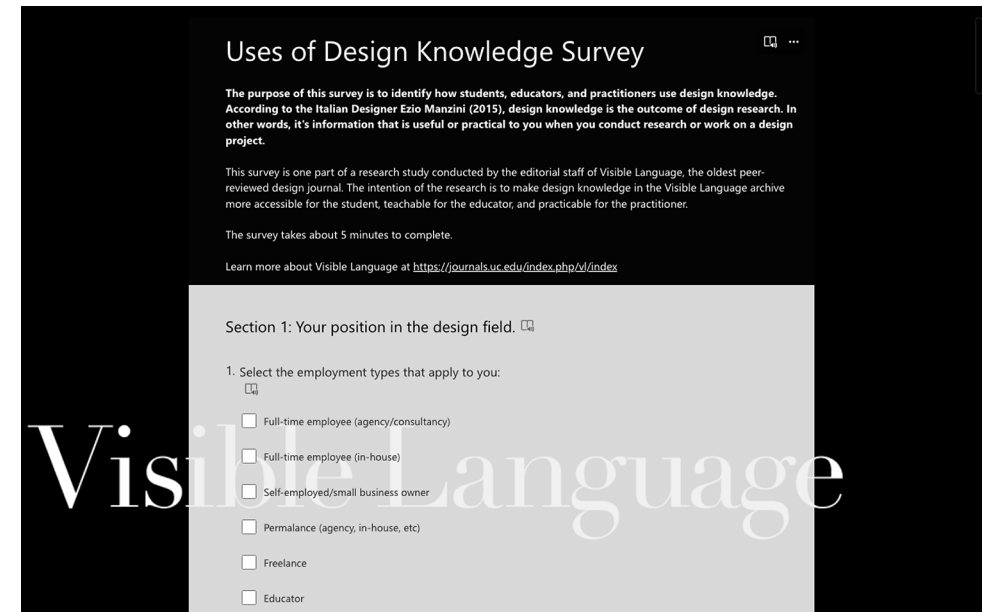
Graphic and Communication Design (Design) often looks like arranging and rearranging letterforms, colors, and shapes to achieve better graphic and communication design processes and outcomes. The ever-change list of what graphic and communication design encompasses includes user experience design, user interface design, design research, type design, visual identity design, information design, service design, design educators, and more.

Designers

By designers, this study refers to people who create and study design. Being that designers are often one of many interdisciplinary collaborators working on a given project or study, this study extended its audience to include readers and authors beyond traditional design, including, but not limited to, visual anthropologists or neuroscientists interested in design knowledge.

Design Knowledge

Design knowledge is the outcome of design research (Manzini, 2015). It is information gathered through studies for, into, and through design (Frayling, 1994). Developing design knowledge usually involves systematized processes but is not exclusive to scholarly sources like VL. Design knowledge fills a gap for designers, whether it is scholarly or not.



Uses of Design Knowledge Survey

The purpose of this survey is to identify how students, educators, and practitioners use design knowledge. According to the Italian Designer Ezio Manzini (2015), design knowledge is the outcome of design research. In other words, it's information that is useful or practical to you when you conduct research or work on a design project.

This survey is one part of a research study conducted by the editorial staff of Visible Language, the oldest peer-reviewed design journal. The intention of the research is to make design knowledge in the Visible Language archive more accessible for the student, teachable for the educator, and practicable for the practitioner.

The survey takes about 5 minutes to complete.

Learn more about Visible Language at <https://journals.uc.edu/index.php/vl/index>

Section 1: Your position in the design field.

1. Select the employment types that apply to you:

- ☐ Full-time employee (agency/consultancy)
- ☐ Full-time employee (in-house)
- ☐ Self-employed/small business owner
- ☐ Permalence (agency, in-house, etc)
- ☐ Freelance
- ☐ Educator

Methods

Phase One: Study Design and Participation

Figure 1*Phase One Survey*

Phase One — current uses of design knowledge — involved a literature review, survey, and semi-structured interviews. For the literature review, the author mainly studied AIGA — the Professional Association of Design in the U.S. — materials to understand designers' current needs and uses of design knowledge. These resources included the 2019 Design Census (AIGA, 2019), Davis' pulse on GCD's current and future state (Davis, 2021), and *Design POV: An In-Depth Look at the Design Industry Now* (AIGA, 2021). While they included international perspectives on design, the content reflected a mostly American point of view.

The author recruited 41 survey participants via personal emails and LinkedIn posts. The "Uses of Design Knowledge" survey included 15 Questions. The first set asked about the participant's role in design, the second set asked about their uses of design knowledge. Aspects of the survey language and descriptions relied on the 2019 AIGA Design Census and a research study about journal reading methods (Subramanyam, 2013).

The author recruited two designers for semi-structured interviews through the survey. The participants — design practitioners — responded to seven questions about design knowledge acquisition. The participants offered valuable insights into reading and presenting design research.

Phase One: Analysis

Literature Review

The literature review revealed rapid technological advancements, sprawling globalization, and calls for social impact. AIGA's studies are directed toward creating design knowledge for designers working within a world and discipline marked by constant change. Undoubtedly, VL's archive reflects decades of change in design — a search of the word "change" produced fifty-six articles. This problem space — designers' responding to constant change — aligns well with VL's core purpose and correlates with what designers shared in this study.

Survey and Semi-Structured Interviews

The author analyzed the survey and semi-structured interviews — along with the literature review findings — to form the seven Design Criteria (Table 2). The following paragraphs describe the criteria with evidence found through Phase One's methods.

Design Criterion One:

*Curate and organize articles
for better searchability*

There was a desire to make VL articles easier to find. One survey participant suggested "better systems of organizing and cataloging." Another said they find "curation" of design knowledge helpful, which the author takes to mean that information is organized and presented by themes or subjects. Others implied that they wished they could search by specific methods (like usability tests) or standard questions (like why you should not use a logo in an email footer). The current VL online archive presents articles chronologically by issues and volumes. Existing search options include filters for dates and authors for readers who already have these specifics in mind. The VL online archive can better organize or curate articles to give readers more avenues into design knowledge.

Table 2

Phase One Design Criteria

#	Design Criteria	Category
1	Curate and organize articles for better searchability	Organization
2	Guide readers through an evidence-based design process	Grouping
3	Connect readers to the various origin stories within the field	Grouping
4	Elevate content that addresses a rapidly changing world	Grouping
5	Bridge the gap between the scholarship and practice	Presentation
6	Summarize article at various scales	Presentation
7	Increase connection between readers and writers	Presentation

Design Criterion Two:

*Guide readers through
evidence-based processes*

This criterion reflects the desire to use design research processes within design practice. Only 14 of 41 (34%) of the participants identified as a design researcher. Yet many of the participants rely on design knowledge. 25 (60%) said they seek design knowledge to identify methods; 19 (46%) use design knowledge to back up their views and decisions; and 16 (39%) use design knowledge to prove the effectiveness of their designs. Therefore, a curated archive should include groupings of articles highlighting design methods and processes that others, especially practitioners, can easily find and utilize.

Design Criterion Three:

*Connect readers to the
various origin stories
within the field*

The VL online archive holds articles about the origins of designed things. An interview participant said they stumbled upon a VL article about ISO icons while looking to "understand why certain icons were systematized and drawn with a particular perspective — why do women have triangular dresses?" They wanted to "understand the origin of these things" to inform decisions on icon systems for a client project. These types of articles — the origin stories of design and designed things — could be curated for better searchability.

.....
Design Criterion Four:

*Elevate content that
addresses a rapidly
changing world*

One of the interview participants wished for more information on how to “build an interface that allows the user to compose their [own dashboard]? ... That tool has to allow that user to build [charts] and allow all of our other customers to build similar things.” In other words, designers are asking questions there may not be answers to yet. Grappling with changing technologies is familiar to VL. For instance, an article from 1998 titled “Writing in the Age of Email: The Impact of Ideology versus Technology” grapples with evolving language and “computer mediated communication” (Baron). Grouping articles like this — where designers in the past grappled with change — could offer the contemporary designer helpful insights in the face of yet-to-be-answered questions about changing technologies.

.....
Design Criterion Five:

*Bridge the gap between
scholarship and practice*

The VL editorial team wanted to know where most designers go for design knowledge. Five of the 41 (12%) participants had read VL. Eighteen (43%) go to “scholarly resources” for design knowledge. Otherwise, they use non-scholarly resources like podcasts, films/video, books, industry publications, and other popular resources. Bridging the gap between scholarship and practice might mean using language — both visual and verbal — that welcomes design practitioners or students new to scholarly archives.

.....
Design Criterion Six:

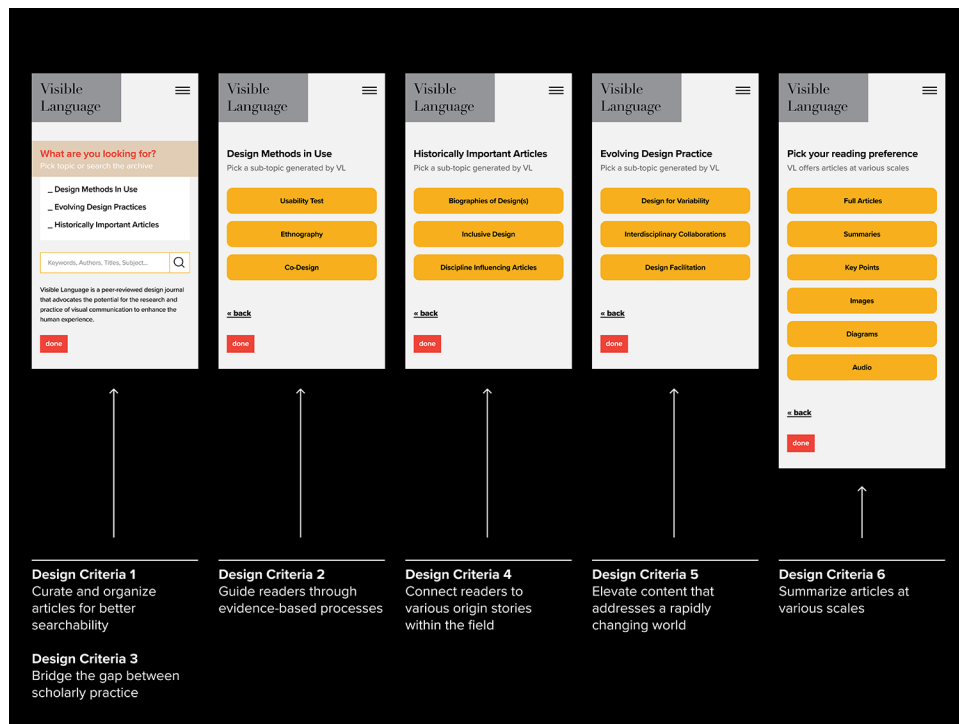
*Summarize articles at
various scales.*

One interview participant confessed, “I stopped for like three hours today and didn’t get any work done because I was skimming the internet or reading... I feel bad stopping and doing research.” The other said, “What are we trying to solve for? How did we do it? What did we learn? What’s next? Get me to the bullet points — include a slide that says TL; DR (Too Long; Didn’t Read).” Thirty-six survey participants (87%) said they skimmed articles — just over half said they read from start to finish. Many readers may prefer summaries or takeaways from an article — *Shi Ji* starts articles with highlights, and *The International Journal of Design* uses symbols to communicate article content. There is room to explore dynamic and abstracted expressions of VL online archive content that take less time to understand than reading full scholarly articles.

.....
Design Criterion Seven:

*Increase connection
between readers
and writers*

The final design criterion encourages connections among current and potential VL community members. “It would be great to hear from the authors. I sometimes go on YouTube and listen to a book tour lecture before I buy the book and read it. It is nice to hear the author directly...” a survey participant wrote. 17 (41%) engage design knowledge in “dialogue with others.” Additional responses referred to ideas like organizing symposia, workshops, and discussions. The online archive could catalyze connections across articles, themes, and backgrounds of authors. Readers and authors alike could contribute and socialize in the archive through a comment section or wiki-like pages.



Phase Two: Study Design and Participation

Figure 2

The prototype with Design Criterion from Phase One. The prototype did not include Design Criterion Seven; however, it will appear in the final proposal.

Phase Two demonstrated six of the seven the design criteria from Phase One through the low-fidelity archive interface prototype in Figure 2. The prototype offered curated content that reflected the noteworthy topic spaces articulated in Phase 1. It also focused on approachable language and the option to engage scholarly articles at various scales for designers unfamiliar with scholarly spaces.

The author used Maze, an online usability testing platform, to share the prototype and a card sorting exercise with participants. Each of these activities involved qualitative questions before or after the activity. There was a pre- and post-question with the prototype, a question to follow up the card sort, a general question for anything else on a participant's mind, and a request for contact information — a total of seven prompts.

The purpose of Phase Two was not to evaluate the interface usability but to explore designers' desired uses of design knowledge. What do they wish they knew, and how do they prefer to engage design knowledge?

The author recruited participants through personal emails, sharing on LinkedIn and AIGA chapter Facebook groups, and the University of Cincinnati's Ullman School of Design newsletter. Of the 64 participants who arrived at the study's home page, 36 (56%) responded to the pre-prototype question, 39 (60%) interacted with the prototype, 26 (40%) responded to the prototype follow-up question, 28 (43%) participated in the card sort, 24 (37%) responded to the card sort follow-up question, 24 (37%) responded to the question about anything else participants would like to add, and 24 (37%) responded to the last question about sharing their email. Responses included a participant's decision to "skip" a question.

The first question activity asked participants what they wished they knew about visual communication design education, research, and practice. This question produced 79 coded snippets of text across 35 participants.

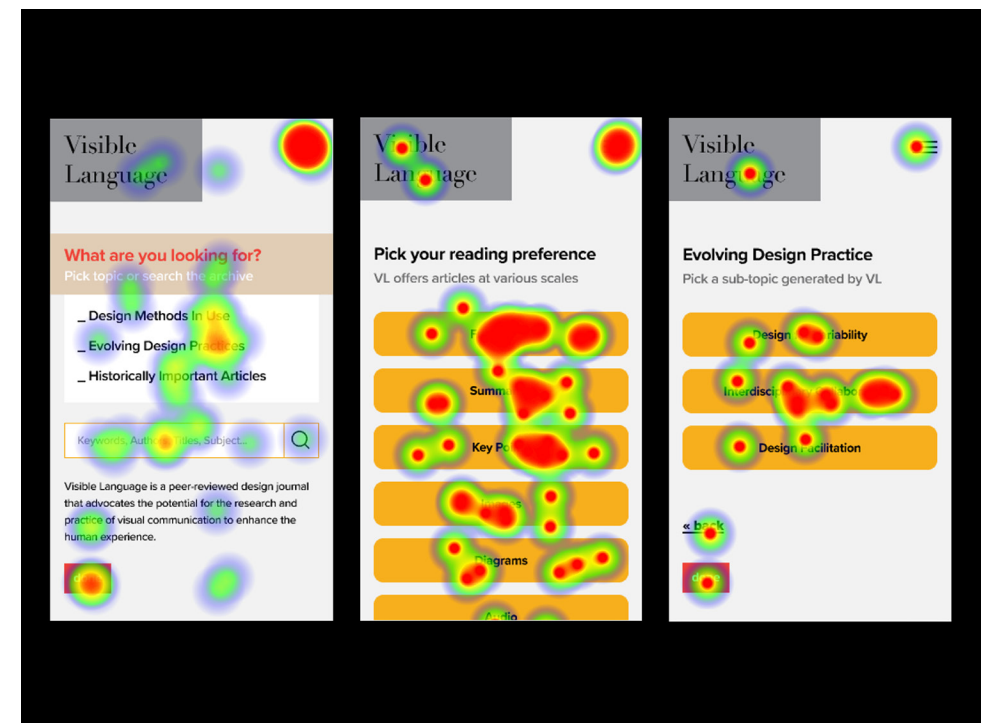


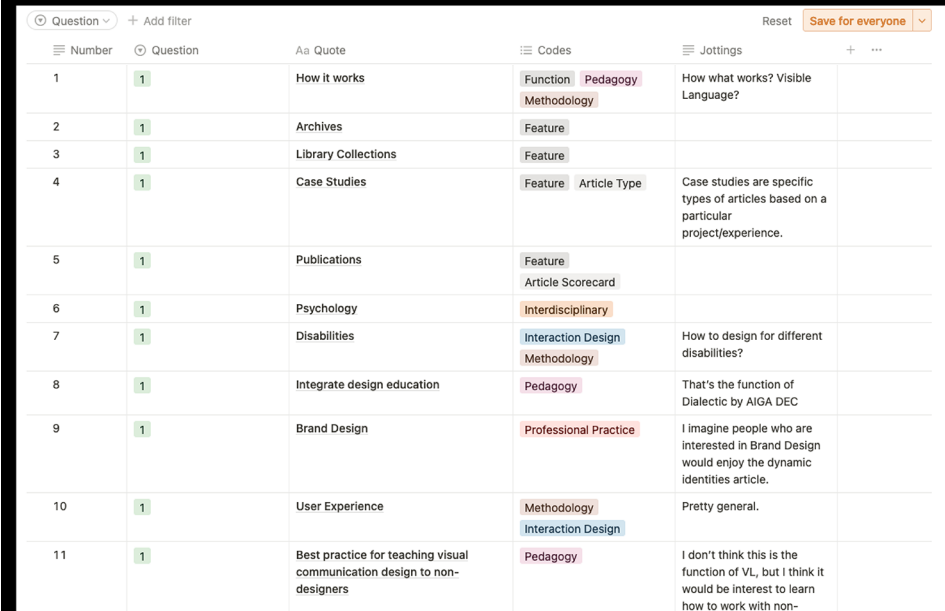
Figure 3

Examples of Heat Maps from the Phase Two Study.

The second part of the first activity invited participants to navigate the low-fidelity prototype (Figure 3). The prototype's primary purpose was to facilitate more responses from the participants regarding what they wished for in the VL online archive in question two. After browsing the prototype, participants were asked what other topics or features they would want from the VL online archive. Playing with the prototype generated 31 more coded snippets.

The second activity — a card sort — asked participants to group a pre-defined list of article reading preferences (based on Design Criterion Six). The list included “Summaries,” “Key Points,” “Images,” “Full Articles,” “Diagrams,” “Audio,” and “Lecture/Presentation.” Participants moved the preferences into the following groups, “Would Use,” “May Use,” and “Would Not Use.” Immediately after the card sorting, the participants could share what they would expect from their most preferred reading preferences. This question produced 34 coded snippets.

The last activity asked the participants to add any additional comments related to the study. Five participants shared 13 additional snippets. The activity ended with a request to stay in contact about future studies and ongoing updates about VL — 13 participants gave their email addresses of the 59 who began the Maze study (22%).



Question	+	Add filter	Reset	Save for everyone
Number	Question	As Quote	Codes	Jottings
1	1	How it works	Function Pedagogy Methodology	How what works? Visible Language?
2	1	Archives	Feature	
3	1	Library Collections	Feature	
4	1	Case Studies	Feature Article Type	Case studies are specific types of articles based on a particular project/experience.
5	1	Publications	Feature Article Scorecard	
6	1	Psychology	Interdisciplinary	
7	1	Disabilities	Interaction Design Methodology	How to design for different disabilities?
8	1	Integrate design education	Pedagogy	That's the function of Dialectic by AIGA DEC
9	1	Brand Design	Professional Practice	I imagine people who are interested in Brand Design would enjoy the dynamic identities article.
10	1	User Experience	Methodology Interaction Design	Pretty general.
11	1	Best practice for teaching visual communication design to non-designers	Pedagogy	I don't think this is the function of VL, but I think it would be interest to learn how to work with non-

Phase Two: Analysis

Coding Qualitative Responses

Figure 4

Coding Phase Two Snippets

The next paragraphs describe the codes that result from analyzing participant responses in Phase Two and how they began to inform the article curation suggested by Design Criterion One in Phase Two. Each grouping includes a title that reflects the design knowledge participants seek, a brief description, and quotes from the participants.

Code/Curated Group	Instances
Alternative Design Practices	18
Emerging Techniques	10
Historiography	24
Interdisciplinary Models	18
Accessible Experiences	17
Methodology	26
Pedagogy	27
Professional Practice	35

Table 3

Phase Two Qualitative Codes/
Curation System (Alphabetical
order).

.....
*Designers seek “Critical
or Alternative Design Practices.”*

This code marked instances affiliated with inclusive design or alternatives to “business-as-usual design” (Wizinsky, 2022). This might mean search filters related to an author’s identity and a specific collection of articles tagged by “Alternative Design Practices.”

- “Search by various tags or categories, i.e., female authors, queer authors, ethnic or religious perspectives, particular topics like design curriculum.”
- “Criticism that positions Design as a Postcapitalist Field of Knowledge.”
- “Transitional Design.”

.....
*Designers seek
“Emerging Techniques”*

This code related to Design Criteria Five in Phase One is about elevating content that addresses a rapidly changing world. Similarly, it reflects the participants’ wish to stay on the forefront of, or at the very least, keep up with change. VL’s future online archive should include articles curated by “Emerging Techniques,” an iteration of “Evolving Design Practices,” as portrayed in the Phase Two prototype.

- “Artificial Intelligence.”

— “Another trend is the rapid development of technology and the materials that can be used in design practice today.”

— “Impacts of technology and design on sociopolitical worlds.”

.....
*Designers seek
“Historiography.”*

This code was associated with the origins and meaning of visual things and design history (like Design Criterion Three). The code has a higher code count (24 snippets) and was present across the activities in Phase Two. The code represents a consolidation of the “Historically Important Articles” collection in the prototype and emphasizes lineages in design. VL’s future online archive should include articles curated with “Historiography” that speak to the origin of designed things and the design discipline.

— “I mentioned other sub-fields of design earlier, such as service design, transitional design, and innovation design—how did these modes of design begin and evolve? Can we see the history using images?”

— “A high-level overview of a topic with an opportunity to dig deeper into the topic.”

— “Track the generation and interrogation of knowledge streams.”

.....
*Designers seek “
Interdisciplinary Models”*

Quotes attached to this code cited design plus another discipline or work outside traditional design training (like “community cultivation” on design projects). The code evolved from the “Interdisciplinary Collaborations” button on the prototype. The heat map of the button’s screen revealed 10 engagements, making it the most popular topic on its page for “Evolving Design Practices.”

— “Interdisciplinary models.”

— “Connections with other fields of research.”

— “Somatic science of communication.”

.....
*Designers seek examples
of “Accessible Experiences.”*

This code reflected instances about designing better human-centered, “Accessible Experiences” across varied formats, contexts, and processes. Participants were concerned with improving the experiences of subjects in design research or the intended audiences of design projects. VL’s archive should curate “Accessible Examples” from their collection of articles.

— “Translating words into other languages from accessibility.”

— “Why, in the majority of design programs, is there a wasteland of ‘communication’ courses: the “how to design” courses never have anything to do with the experience of the “the user/audience?”

— “Better outreach to vision impaired.”

Designers Seek “Methodologies”

This code described instances that point to a desire for reliable research and design methods. The author iterated on the “Design Methods in Use” section title — the most visited in the prototype (18 participants) — to create “Methodologies.” It is the broadest of the codes, appearing in 26 instances, almost always combined with other codes.

— “Finding subjects for research interviews.”

— “More ways to create low fidelity prototypes with users during workshops (in-person and virtual).”

— “Find frameworks.”

Designers Seek “Pedagogical Frameworks”

This code characterized instances focused on learning and teaching design. While the prototype did not explicitly reference pedagogy, it was among the highest-level codes with 27 snippets about the classroom, curriculum, and theory. As a result, the code should be a curated group in VL’s future online archive.

— “Integrate design education.”

— “How curriculum is created for a class from scratch.”

— “I just started teaching as an adjunct, so anything about visual communication design education!”

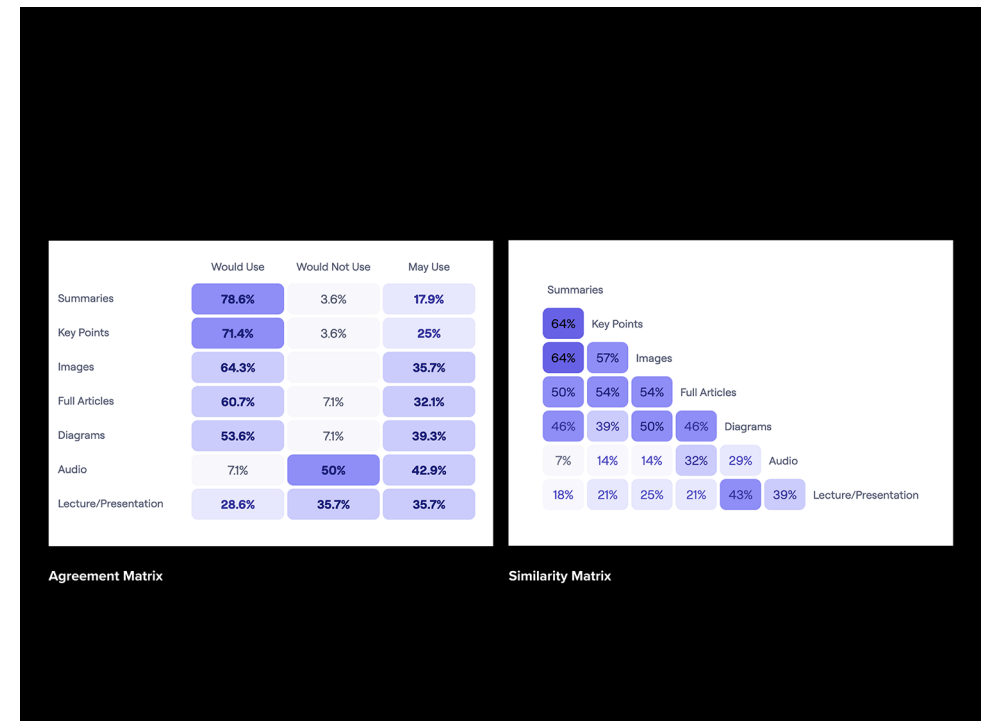
Practitioners seek “Design Knowledge.”

This code referenced descriptions of designers’ needs in commercial settings like studios or in-house design jobs. It was the most frequently used code (35 snippets). The prototype did not include specific connections to articles tagged “Professional Practice,” but its inclusion is critical in VL’s future online archive. The VL online archive must facilitate access to professional practitioners who may need help finding articles through scholarly databases.

— “How to design things well but not follow the usual trends.”

— “I wish there were more robust sustainability resources and certifications for our industry.”

— “Innovative practice and organizational models.”



Analyzing Reading Preferences

Figure 5

Reading Preference Results

The card sort showed (Figure 5) that reading preferences like summaries, key points, and images are more useful to the participants. However, all the options had some attractiveness. Written responses offered valuable perspectives behind the participants’ expectations, which indicated a more dynamic reading experience — in addition to traditional articles. A dynamic reading experience could include article summaries, links to related content, and visual representations (for example, a system map that depicts the interconnectivity of themes or authors). An extension like this could bridge the gap between scholarship and practice.

— “Everyone’s time is limited, so having a solid idea of what is going on via key points or a summary before reading a whole article is useful.”

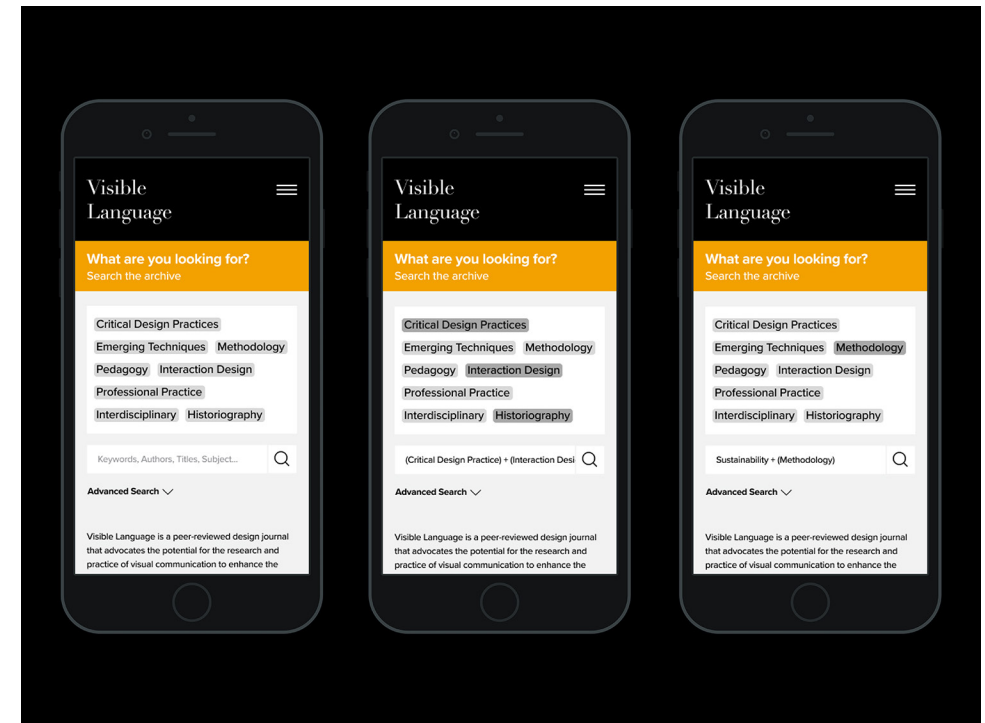
— “Lectures/presentations would ideally include suggested talking points or any additional context (i.e., reading list) from the author that would help discuss the topic in the classroom.”

— “It would be nice to add some sort of ‘related articles,’ ‘related topics,’ or ‘other articles by this author’ with links to previously published works.”

Results

During the analysis of Phase One of this study, the author produced seven Design Criteria that directed the design of a low-fidelity prototype in Phase Two. That prototype helped the author identify expanded topics for curation and a dynamic reading experience.

As a result, the author consolidated Phase 1 and Phase 2 into a set of Design Strategies that emphasize the study’s most important takeaways. The Design Strategies are focused priorities that will help VL improve and broaden engagement in the “Next Generation” VL online archive. Design Strategy One and Two emphasize curation. Design Strategy Three focuses on article summarization and building connections across authors, editors, and readers through a dynamic reading experience.

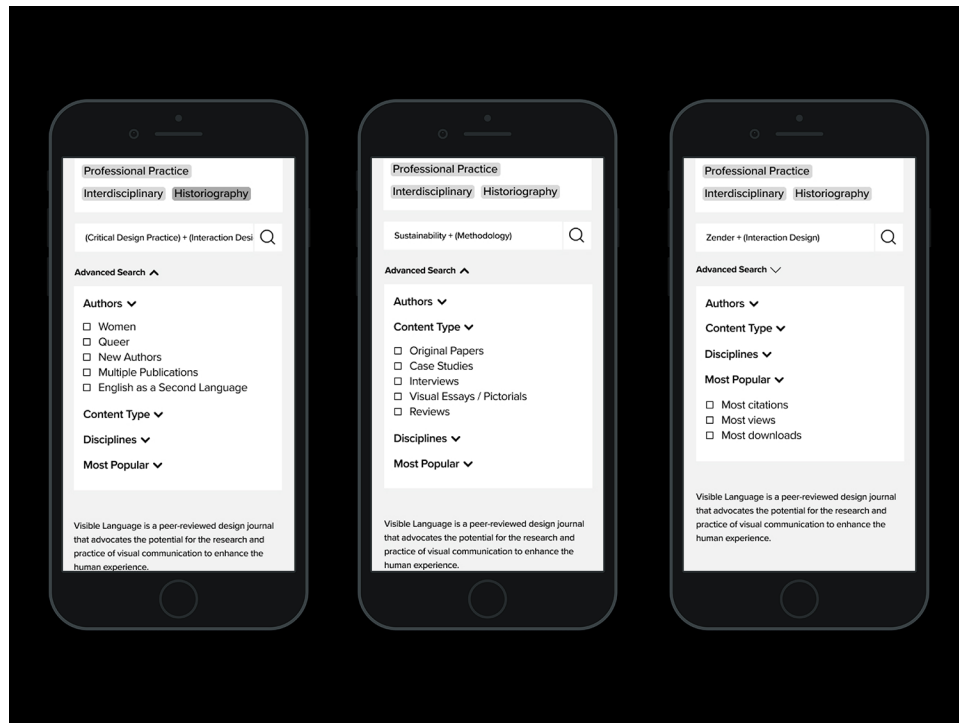


Design Strategy One: Simplify curation groupings and maximize variations so readers can quickly retrieve articles related to the most relevant topics

Figure 6

*Specific Tags and Open Search in
VL’s proposed “Next Generation”
online archive.*

The prototype in Phase Two offered three groups of specific articles without the possibility of combinations. VL’s proposed next-generation online archive presents eight combinable themes plus an open search. For instance, a reader could look for articles that combine Critical Design Practice, Interaction Design, and Historiography or look for articles about Methodology with keywords related to sustainability (Figure 6). The new system enables a user-friendly resource for readers — academic and non-academic alike — to quickly find articles based on the topics designers in this study shared that were most relevant.

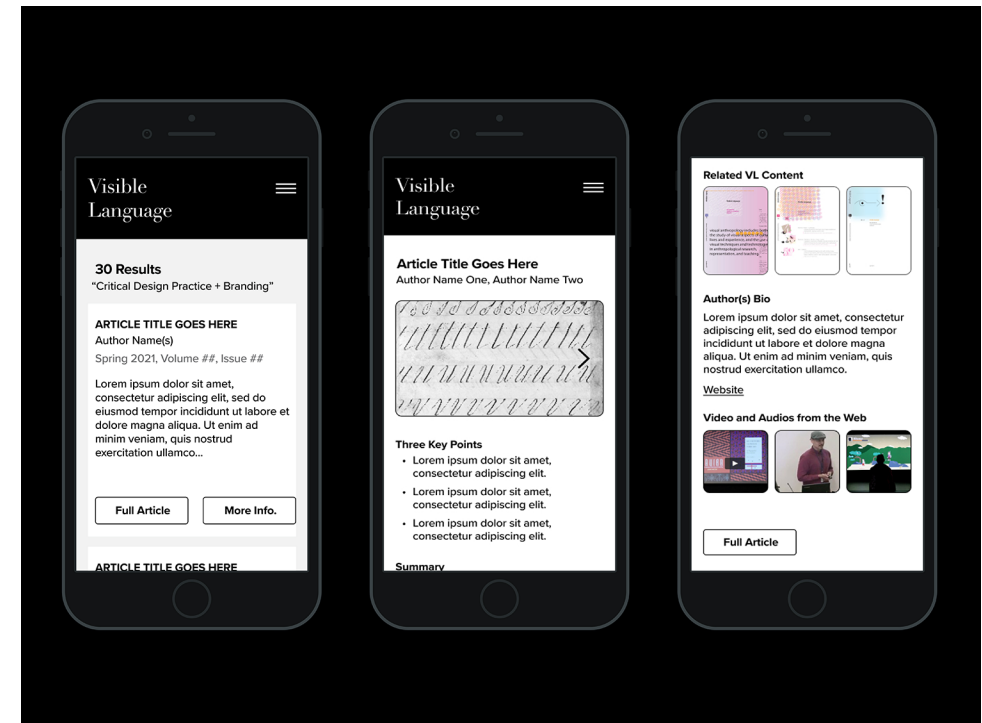


Design Strategy Two: Offer advanced search filters, like author background or specific methods

Figure 7

Advanced Search Filters in VL's proposed "Next Generation" online archive.

Beyond curating topics and areas of interest, VL's next-generation online archive should give the readers more detailed search options related to author identity and experience, related disciplines, specific methods, content types, and popularity (Figure 7). A detailed sub-search in VL's future online archive opens additional entry pathways into the archive's design knowledge.



Design Strategy Three: Create a dynamic reading experience that enables readers to discover more about the author, topic, and key points before deciding to read entire articles.

Figure 8

Dynamic Reading Experience in VL's proposed "Next Generation" online archive.

VL's future archive should include possibilities for readers to familiarize themselves with an article, its authors, and general topic space before – or instead of – reading an entire article. The dynamic reading experience requires more definition in future studies. However, it could look like articles – or curated groups – having a summary, key points, images, diagrams, and tools that deliver, on the one hand, vital content for skimming and, on the other hand, opportunities for deep dives (Figure 8). Either way, the dynamic reading experience should present content in different modalities. Readers could partially generate the pre-reading experience content, wiki-style. There may also be opportunities for AI-compiled summaries and organization.

Discussion

Overview

This study started with the belief that the *VL* digital archive holds a vital collection of design knowledge that is meaningful and valuable for contemporary and future designers as they grapple with an unpredictable discipline and world. However, the archive needs improvements to reach more designers, especially those who do not work in academia.

The author uncovered seven Design Criteria during Phase One to direct the prototype design in Phase Two. The Design Criteria suggested how *VL* might organize and present archive content into relevant categories and potential ways to connect readers to content, other readers, and writers. The prototype in Phase Two led the author to identify curated groups and a dynamic reading experience. The author then combined the findings from Phase One and Two into three key Design Strategies for *VL*'s future online archive. These strategies shaped the proposed archive design with flexible curation, filtering, and a dynamic reading experience. Innovations like these could unlock the deep repositories of design knowledge hidden behind both an inferior interface and, as one Phase Two participant described, "dry" articles that, if presented in more dynamic ways, may attract more readers, especially those less likely to visit a scholarly archive.

Limitations

This researcher conducted this study between fall 2021 and summer 2022 as an editorial assistant to *VL* and a graduate student at the University of Cincinnati's Ullman School of Design at the College of Design, Architecture, Art, and Planning. Funding for the editorial assistantship came from an internal grant directed toward editorial assistants. *VL*'s editors — professors at the Ullman School of Design — wrote the budget with this study in mind. The grant covered the author's time but did not include funding for promotion or recruitment. As such, the study consists of a modest — but insight-filled — sample size of participants mostly from, but not limited to, the United States.

The study's results should not be confused with generalizable knowledge. It did not fall under traditional human-subject research and did not go through the University of Cincinnati's IRB approval process. With these points in mind, the reader should

consider this content as a case study representing a product-improvement project — utilizing an evidence-based process — with localized results. The insights — or design knowledge — above is speculative, not prescriptive.

Finally, multiple people — including the *VL* editorial team and an additional editorial assistant — graciously contributed feedback and assistance during this study. Bias is inevitable within an analysis of a product or service by the people leading the product or service and this article did not go through an outside peer-review process. The *VL* editorial team encourages increased outside participation and external studies of the archive, which will be elaborated on in the next section.

Next steps

This study, and subsequent studies, of the *VL* online archive falls into the camp of "investigations into academic journal design and reading experiences." (Barness & Papaelias, 2021) *VL* is a design journal led by designers, allowing it to act as a testing ground for micro and macro journal design improvements that potentially generate knowledge — for other journals. "Such changes would be welcomed by journal readers" (ibid. p. 561).

An immediate next step for the *VL* online archive is the "evaluation, refinement, and production" of the design strategies above (Hanington & Martin, 2012). A future study could include methods like think-aloud protocol, usability tests, and more research through design with current and potential *VL* readers using further developed prototypes of *VL*'s next-generation archive.

Another study could examine how Artificial Intelligence might identify, organize, and classify the archive using the tags in this study. A future *VL* editorial assistant could archive and begin curating articles if budget limitations prevent using machine learning.

Additionally, this study instigated further research into the *VL* online archive dynamic reading experience, which could have implications beyond the specific journal. A future study could dive further into this area, highlighting the best modalities for demonstrating the interconnectedness of a theme, article, authors, methods, and more.

Dear VL Reader

If you shared feedback and perspectives in this study, thank you. There will be future opportunities for participation as the editorial team works on the “Next Generation” VL online archive. Please email Editor, Mike Zender, mike.zender@uc.edu, if you would like to participate.

References

ALGA. (2019). *Design Census 2019*.

ALGA. (2021). *Design POV: An In-Depth Look at the Design Industry Now, Executive Summary*. <https://www.aiga.org/aiga-design-pov-reports>

Baron, N. S. (1998). Writing in the age of email: The impact of ideology versus technology. *Visible Language*, 32 (1), 35.

Barness, J., & Papaelias, A. (2021). Readable, Serious, Traditional: Investigating Scholarly Perceptions of the Visual Design and Reading Experiences of Academic Journals. *She Ji: The Journal of Design, Economics, and Innovation*, 7(4), 540-564. <https://doi.org/10.1016/j.sheji.2021.10.005>

Davis, M. (2021). *Design Futures*. <https://www.aiga.org/resources/design-futures-research>

(1987). The Avant-Garde and The Text: A Special Issue of *Visible Language*. *Visible Language*. <https://journals.uc.edu/index.php/VL/issue/view/379>

Frayling, C. (1994). *Research in art and design (Royal College of Art Research Papers, vol 1, no 1, 1993/4)*. <https://researchonline.rca.ac.uk/id/eprint/384>

Hanington, B., & Martin, B. (2012). *Universal Methods of Design: 100 Ways to Research Complex Problems, Develop Innovative Ideas, and Design Effective Solutions*. Rockport Publishers. https://play.google.com/store/books/details?id=pCVATlpzYfUC&source=gbs_api

Manzini, E. (2015). *Design, when everybody designs: An introduction to design for social innovation*. MIT press. <https://books.google.com/books?hl=en&lr=&id=BDnqBgAAQBAJ&oi=fnd&pg=PR5&dq=Design+when+everyone+designs&ots=l7n3oMFeEi&sig=gMMtdDE89K7f mNW-ISRQOPFIwA>

Subramanyam, R. (2013). Art of reading a journal article: Methodically and effectively. *J Oral Maxillofac Pathol*, 17(1), 65-70. <https://doi.org/10.4103/0973-029X.110733>

Wizinsky, M. (2022). *Design after Capitalism*. MIT Press. http://books.google.com/books?id=nMgyEAAAQBAJ&hl=&source=gbs_api

Wrolstad, M. E. (1967). A Prefatory Note to the First Number. *Visible Language*, 1. <https://journals.uc.edu/index.php/VL/issue/view/295>

Wrolstad, M. E. (1971). Visible Language: The Journal for Research on the Visual Media of Language Expression. *Visible Language*, 5 - 12. <https://journals.uc.edu/index.php/VL/issue/view/314>

(2022). December Issue. *Visible Language*, 56.3. <https://journals.uc.edu/index.php/VL/issue/view/488/225>

Zender, M. (n.d.). *Visible Language*. <https://journals.uc.edu/index.php/VL/index>

Author

D.J. Trischler

D.J. Trischler is an Assistant Professor of Communication Design at The University of Cincinnati's School of Design at DAAP. His design research and creative interests focus on reflective visual communication design methods and artifacts that foster inclusivity. His Master of Design thesis research investigated the possibilities of impacting neighborhood quality of life through bottom-up, neighborhood-centered design methodologies. Price Hill in Cincinnati, Ohio, was the site for his thesis and is also the location for NODES (Neighborhoods of Designed Engagement Systems), an interdisciplinary team project that D.J. currently co-leads. NODES aims to increase digital equity through free public Wi-Fi, provide access to digital public library content, and encourage connection between neighbors. Before entering academia, D.J. had over ten years of professional experience, including independent design practice. His portfolio includes civic-centered place branding projects like the official Typeface for Chattanooga, Tennessee, the visual identity design and consultation for the Chattanooga Public Library, and community flags for several neighborhoods in Cincinnati. His work is featured in *The Atlantic*, *Nieman Labs*, and the *Designing Brand Identities* textbook.

A New Model

Mike Zender

Change

Many of you may have followed the dust-up surrounding *Design Studies* this past summer. Just as that unfortunate disagreement happened, a mature author, researcher, and contributor to *Visible Language* asked about getting immediate open access for a forthcoming manuscript. His grant agency required open access and the research subject was timely so we were sensitive to his request. However, *Visible Language* has operated under a hybrid model where subscribers get access to articles immediately and everyone else has access after a one-year embargo, so this request would require an exception. We agreed on a temporary measure to meet his need but the two simultaneous events launched us on a reevaluation of our publishing model in light of a thorough examination prevalent publishing models. We did not like the picture.

The publishing world is moving to fully Open Access. (OA) Open Access proponents seem to have honorable aims including making knowledge freely available to everyone. Forget for a moment that many OA proponents work in Universities where students pay for knowledge acquisition or in government agencies that are supported by laws. We at *Visible Language* share OA proponents' honorable aims. We certainly think knowledge is a great benefit and our hearts are committed to its widest possible dissemination. That's why we publish.

If OA proponents are pushing for an eventual future where nearly everyone publishes everything for free that system exists already and it's called the internet. We are all familiar with the generous diversity of quality there. Perhaps someday in an AI mediated world a search and rating algorithm will facilitate finding quality work in the free-for-all. I'm not sure who will pay for that system, perhaps the equivalent of Google ads. Then we'll need funding to buy the ads to push our work to the top of the AI results. I for one will be happy to miss that world.

For now, we must rely on the quality and authority that comes from publishers. Lots of recent articles bring that system into question. Just this week the journal *Nature* had to retract another article. Something about cold fusion (again). Our world seems to be devolving in the trustworthiness dimension.

Journals used to be housed mostly in trusted institutions like Universities and Societies whose role was both validating quality and ensuring dissemination. These direct sources of validation published articles and bore the costs. Today we mostly have publishers as third parties who may publish a stable of journals covering an immense variety of topics. I'm going to steer clear of blaming the money-making associated with this to avoid devilish divisions on the evils of money. Suffice to say, publishing is work and work is seldom free, especially very good work.

Regarding payment, the trend has become for authors to pay publication fees to publishers to cover costs. This we oppose. It limits publication to those who can afford it. This seems completely opposed to aim of Open Access but seen the other way around, that is, the only ones with access to distribution of their work are those with resources. What about the poor or unfunded researcher with interesting but novel findings? Pay-to-publish stratifies the publishing world into the rich and powerful and everyone else. Sorry, that's not a model for us.

Disliking the current picture, we imagined alternatives. We applied for and received funding to conduct a research study to look at what designers want from *Visible Language*. That study is reported in this issue by D. J. Trischler's article immediately preceding this one. In response to our preliminary research findings and current events, this summer we sketched several different approaches to an improved publishing model. We reimagined a centralized peer-review system much like the Associated Press, which would vet manuscripts and offer them to publishers to pick up and disseminate for a fee. We imagined authors being paid and supported by on-line adds and/or "likes" similar to YouTube channels. We dreamed of a user-centered system where users are pinged about a new manuscript of potential interest and if, after reading the Abstract, they read the full article either the individual or their institution is charged a small fee. We conceived alternatives, and more, we shared our ideas with peers. After that and considerable internal deliberation, we winnowed to new model for *Visible Language* based on the past but fit for the future.

Visible Language staff has decided to move back to the direct model where the experts do both vetting and dissemination. We plan to form a consortium of four or five peer institutions to jointly publish *Visual Language*. Thanks to publishing systems such as OJS, our costs including printing are low enough that if each institution contributes a very few thousand dollars annually the journal will be sustained. Each institution would have a seat on the journal governing Editorial Board as an Associate Editor, and equal say in everything from the direction of the journal to articles published. Specifically, we propose that:

- each supporting institution will have a functional seat for one of their faculty/leaders on the *Visible Language* Editorial Board as an Associate Editor;
- the *Visible Language* Editorial Board will maintain complete editorial control of the journal: guide the journal's direction, review articles, plan issues, solicit contributions, etc.;
- each institution will be recognized as a member of the consortium and Editorial Board which directs *Visible Language* with their logo/name on the journal back cover;

- the Editor in Chief will come from one of the consortium institutions, perhaps on a rotating basis, perhaps with 3–5-year renewable terms, elected by the Board;
- UC and UCPress, not-for-profit institutions, will still "own" and publish *Visible Language* as the host institution and professional publication administrators;
- during formation taking about one year, the Consortium will re-write the journal's charter to reflect the new model;
- the University of Cincinnati will enter into collaborative agreements (or other suitable legal documents) with each institution to formalize the consortium.
- once the Consortium is formed, the journal's direction will be determined by the new *Visible Language* Editorial Board.

We believe this model, if it works, has several advantages. It makes the journal sustainable and independent, not beholden to any individual funder or single special interest. It ensures diversity of perspective within unity of purpose, and the potential for diversity of region, language, and culture for a truly global perspective. It provides ample energy with different Associate Editors representing different interests. It broadens access to alternate networks to enrich intellectual quality and rigor. It provides a continual pool of successors to become the next Editor in Chief. Consortium institutions may include research institutes and professional firms of high stature potentially strengthening the integration of theoretical and professional Design knowledge.

There are potential downsides inherent in any diverse body which will require wisdom, tact, some skill to manage. Fortunately, we have independent publishing professionals at UCPress to help us navigate and operate. Our current contract with UCPress (an in-house contract) stipulates that VL is totally controlled by the Editor and Editorial Staff at the UC School of Design. UCPress provides expert advice and publishing/hosting services on contract, nothing more. *Visible Language* has its own financial accounts. When the new Consortium is formed, these lines of authority and accountability will continue to be clearly stipulated.

We're excited at this evolution of *Visible Language*. We think it could become a model that others can adapt that might help to solve some of the problems that pervade scholarly publishing today. We look forward to having the consortium in place by the first issue of 2025.

Stay tuned.

Mike Zender, Editor, *Visible Language*