

Volume 60 Number 2

The Chocolate and Oyster Flavored Donut Hole Of Unoriginality

Matthew Peterson 

Department of Graphic Design and Industrial Design, North Carolina State University, Raleigh, NC, USA
(mopeters[at]ncsu.edu)

Now that each of what we call *writing*, *illustrating*, and *designing* has been at least partially automated by AI, it is a good time to ask: What is *value* in writing? What is *value* in illustration? What is *value* in design? But also: What is *writing*, *illustration*, and *design*? In the final article in this issue of *Visible Language*, colleagues from other disciplines join me in addressing writing, along with other previously human constructs that are applied to our understanding of large language models (LLMs). But we did not really address value there, and that is my emphasis here. (If you are reading this only because it is the introductory editorial for *Visible Language* 60.2, I will ultimately cover this issue's four articles, but I am presently occupied.)

I will use writing and LLMs in my primary discussion, but as you will see, what I write here—and I am writing it, I assure you—will be highly abstract, as readily applied to illustration and design as to writing, and to all forms of generative AI.

In doing so, I will share my own perhaps silly way of thinking about value, and thus originality, in AI outputs: that they are the donut hole, not the donut. And not the donut hole you eat, but the actual hole in the donut that you cannot. Or, since this is a metaphor and not bound by physics, we can quickly pivot to making the donut hole an edible one, as is sold in many establishments, but one that has unfortunately been infused with both chocolate and oyster flavors (after Hayes & Flower, 1980; more on this later). Does that mean that AI outputs have no value, or that they are always and

@: [ISSUE](#) > [ARTICLE](#) >

Cite this article:

Peterson, M. (2026). The chocolate and oyster flavored donut hole of unoriginality. *Visible Language*, 60(2), iv–xxvii.

First published online July 16, 2026.

© 2026 Visible Language — this article is **open access**, published under the CC BY-NC-ND 4.0 license.

<https://visible-language.org/journal/>

Visible Language Consortium:

University of Leeds (UK)

University of Cincinnati (USA)

North Carolina State University (USA)

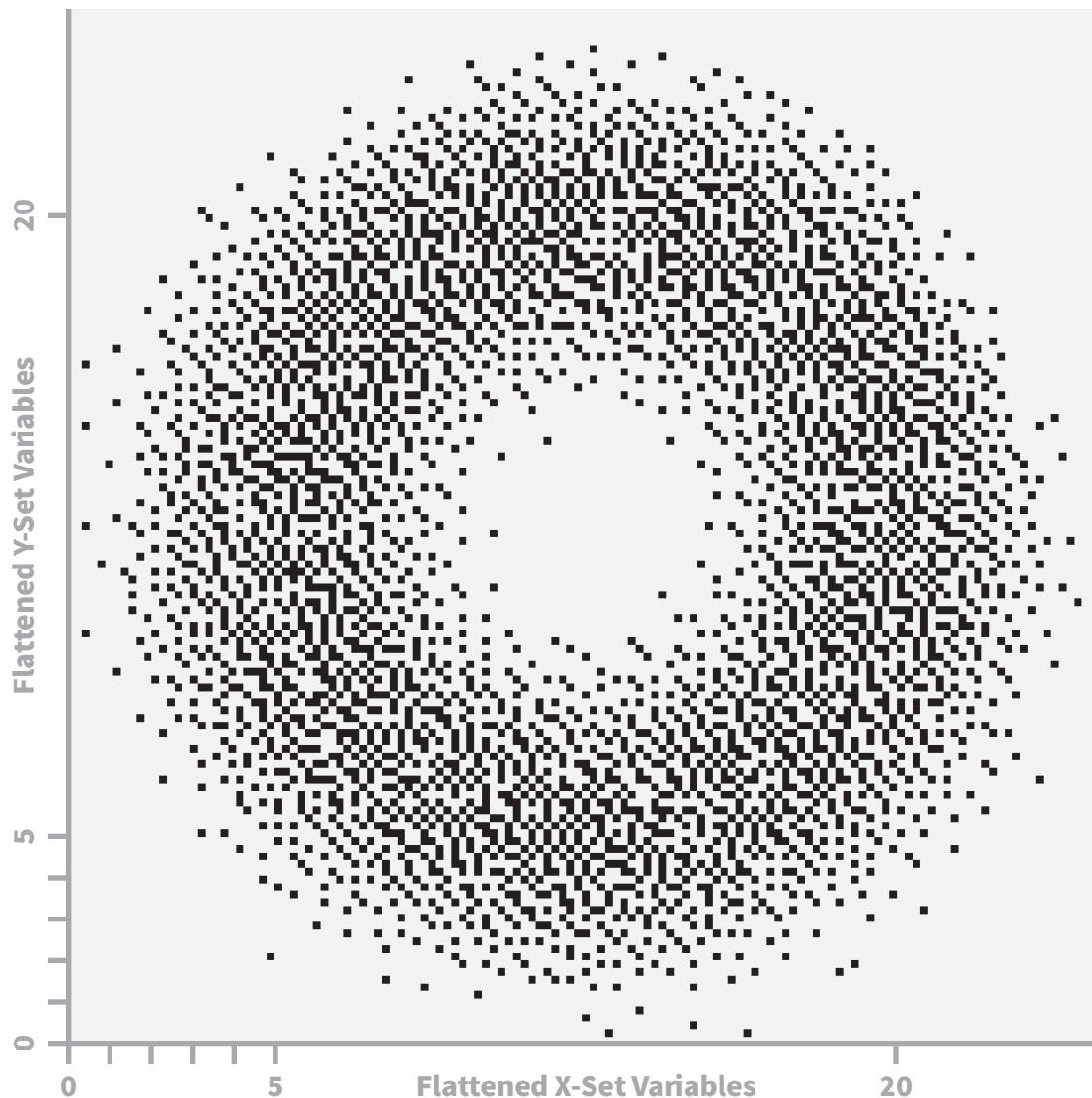


Figure 1. A distribution of values in multiple dimensions flattened down to two dimensions, in a way that is not valid statistically but may nonetheless make sense. Note that the distribution has a hole.

forevermore undesirable? Or even that I personally think so? No!, it is just an idea, just a way of thinking about things. If we can walk and chew bubble gum at the same time, we can think multiple ways at once. It can be insightful to have different angles at which to view one thing. A skeptical angle is a helpful one.

Caveats aside, behold the donut in Figure 1. I assume this would displease any good statistician. But let me explain to the rest of you why this makes sense to me. Imagine that this scatterplot is plotting instances of authorial voice. It looks of course to be plotting two independent variables inherent to authorial voice, such that one voice

could be (3, 15) (from x, y) and another just as easily (15, 3). But this is not what I am imagining. Instead, I am imagining a far richer construct of authorial voice that has many, many dimensions, and that Figure 1 is flattening out a host of variables into one dimension of the scatterplot, and another host of variables into the other. Or, imagine that this is some (perfectly impractical) way to see hundreds of dimensions at once, with the implications of that high-dimensional view becoming evident through the simplification.

If you can join me in this abstract space, then we can look at the distribution as a meaningful shape. It represents the range of authorial voice in human writing, for the moment. Each dot is a full profile of hundreds of factors of authorial voice. Some of these factors may be relatively complex things in their own right, impressions a form of writing gives human readers. Such factors could be measured by degrees, and thus produce a value range that is flattened with other values in one of the axes of Figure 1. To what degree does the writing come across as:

- ▶ Acerbic,
- ▶ Bold,
- ▶ Critical,
- ▶ Dry,
- ▶ Erudite, or
- ▶ Friendly?

Perhaps dry and friendly are not truly independent, but hopefully you get the idea. These are all holistic factors. There are also preferential patterns that might have been plotted in these flattened dimensions. For instance, use of or propensity for using:

- ▶ One parenthetical form (e.g., paired em dashes) over other options (e.g., paired commas).
- ▶ Lists, like this one.
- ▶ Contrastive negation, or the form of “it’s not A, it’s B.”
- ▶ Tricolon, providing parallel elements or sentences in groups of three.

Finally, simple use of terminology can influence voice, such as a preference for particular words:

- ▶ *Keen to*, instead of alternatives such as *enthusiastic about*, *driven to*, or *inclined to*. (Is that a tricolon?)
- ▶ *Delve*, instead of *examine*, *probe*, or *look into*.

If terminology is part of authorial voice, we need to acknowledge that Figure 1 imagines not flattening hundreds, but thousands or millions of variables into a two-dimensional representation. (Is that a contrastive negation?)

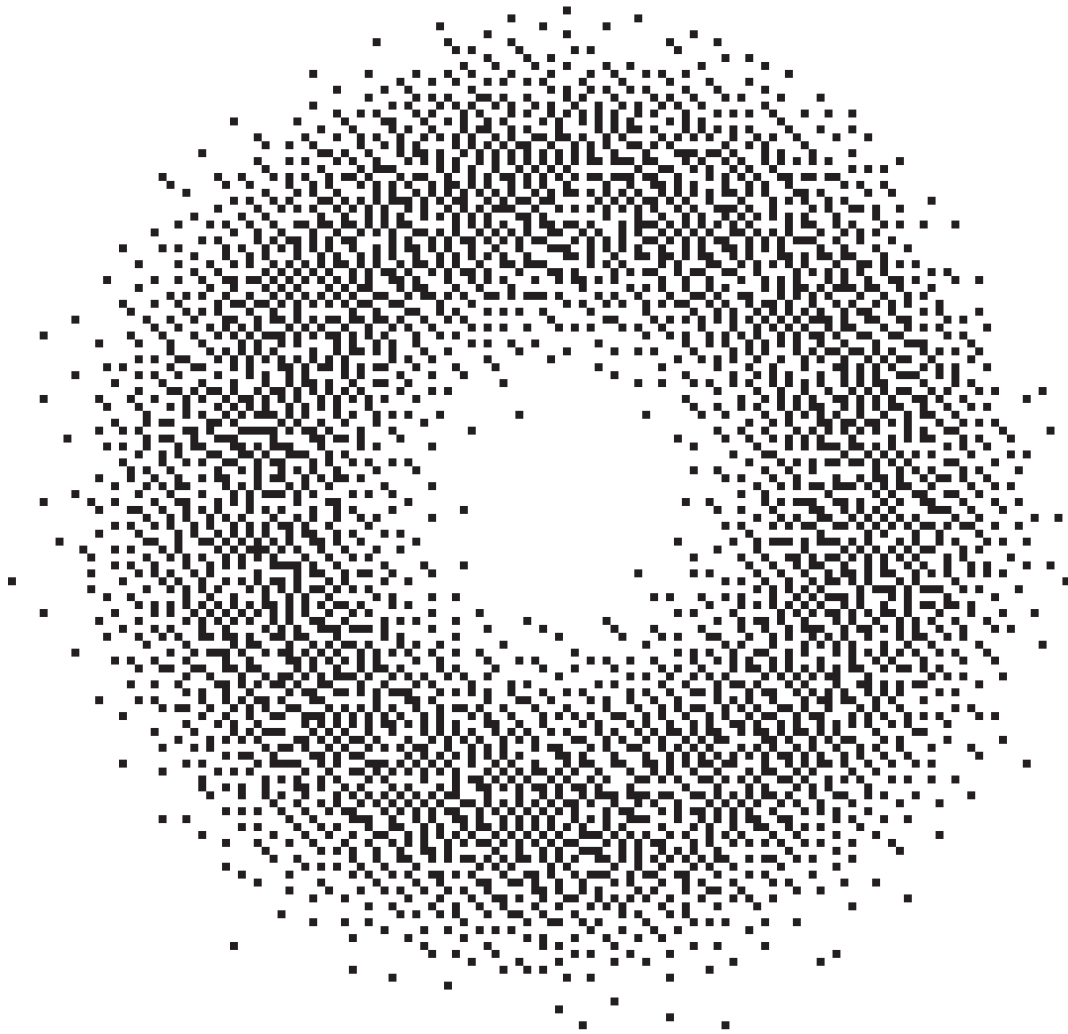


Figure 2. A distribution of values in multiple dimensions flattened down to two dimensions, in a way that is not valid statistically but may nonetheless make sense. Arbitrary x and y axes removed.

So the conceit that Figure 1 is a scatterplot is an excessively exaggerated conceit, downright theatrical. To avoid the limitations of the scatterplot expression, and to encourage ourselves to think more abstractly as is intended, let us remove the armatures of the graph. The values being plotted (i.e., actual x and y) are meaningless anyway, and numeric signification is limited here — the expression $A > B$ is closer to the spirit of the donut representation than is something like $A = B^2 \times 3.45$. Figure 2 leaves us with the donut alone, and that is what we will be working with.

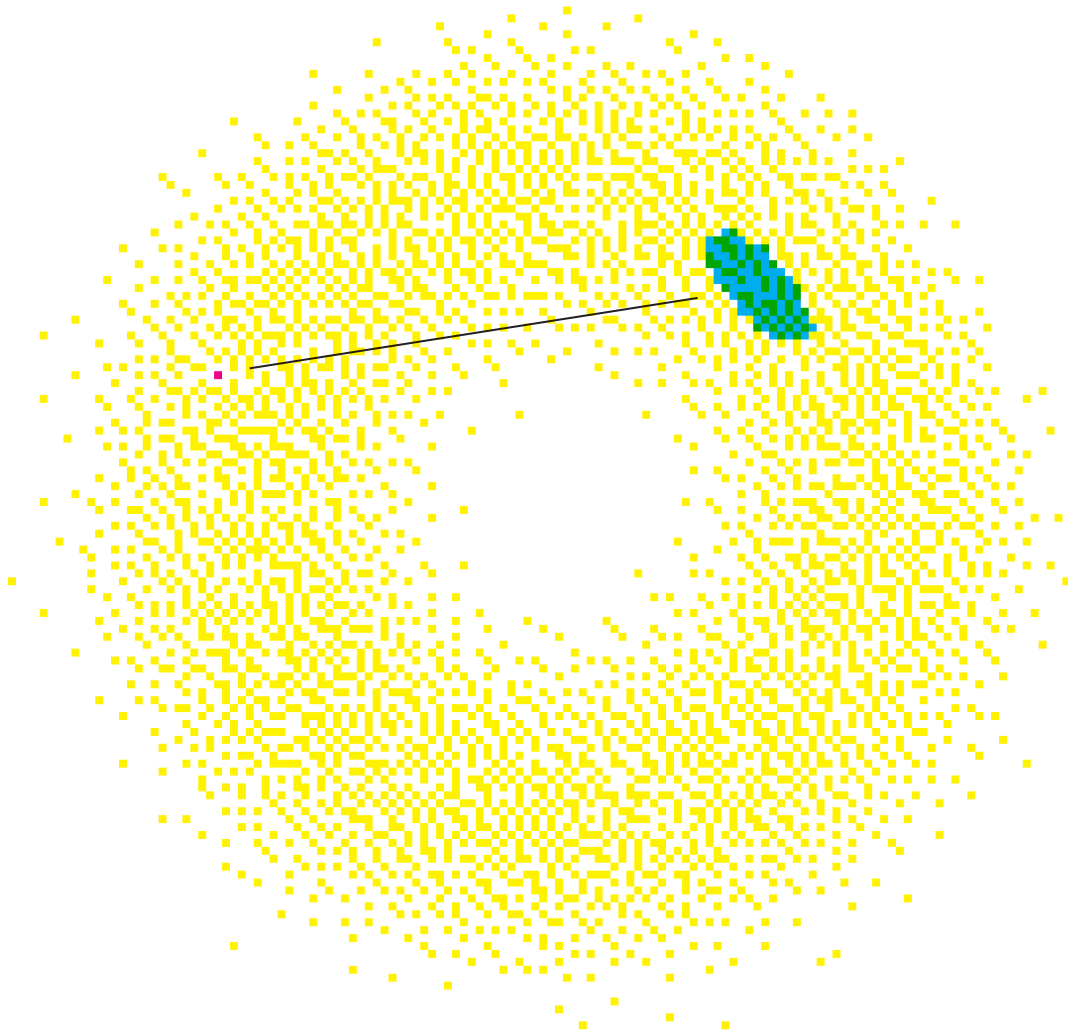


Figure 3. A single point in the donut can represent an individual author's voice, an area a writing genre, and in pairwise comparisons only, distances between points and areas dissimilarity.

One small area of the donut distribution might represent an identifiable and relatively codified writing genre, such as academic prose (cyan in Figure 3). A single point in the donut might represent an individual author's voice, such as that of Gabriel García Márquez (magenta). And the freshness of García Márquez's prose compared to academic prose can be represented by relative distance (magenta to cyan, as indicated in black). Being that countless dimensions have been flattened into two, we cannot accommodate many such comparisons, lest representations of dissimilarity appear close together and thus imply precisely the opposite.

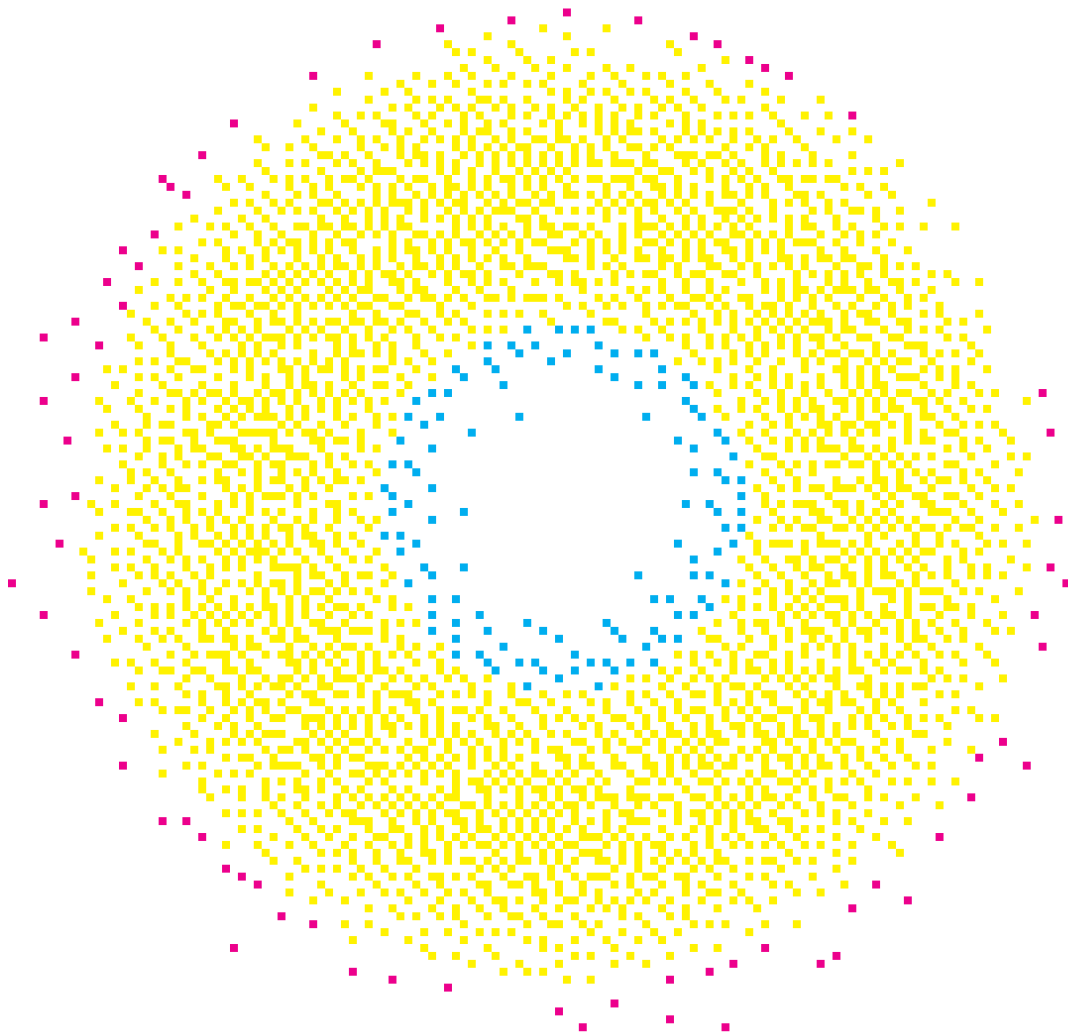


Figure 4. The donut's periphery could represent outliers in the conventional sense, as unique or innovative, while the inner points could represent the mundane. But the donut is a flexible concept.

The relativity of this absurd expression of variance means that we can rapidly pivot to other uses of the imaginary donut distribution. For instance, let us consider periphery in the donut. Points on the outside of the donut (e.g., magenta in Figure 4) could represent outliers in the regular sense. An iconoclastic or downright bizarre writing style may be on the outskirts. And maybe that style proves influential and other dots begin to cluster around it over time. But the abstract flexibility of the donut means that points on the inner edge (e.g., cyan) could just as easily not be conventional outliers, even in the same plotting of values. Dots approaching the center of the donut could represent the mundane by definition, if we wish them to do so.

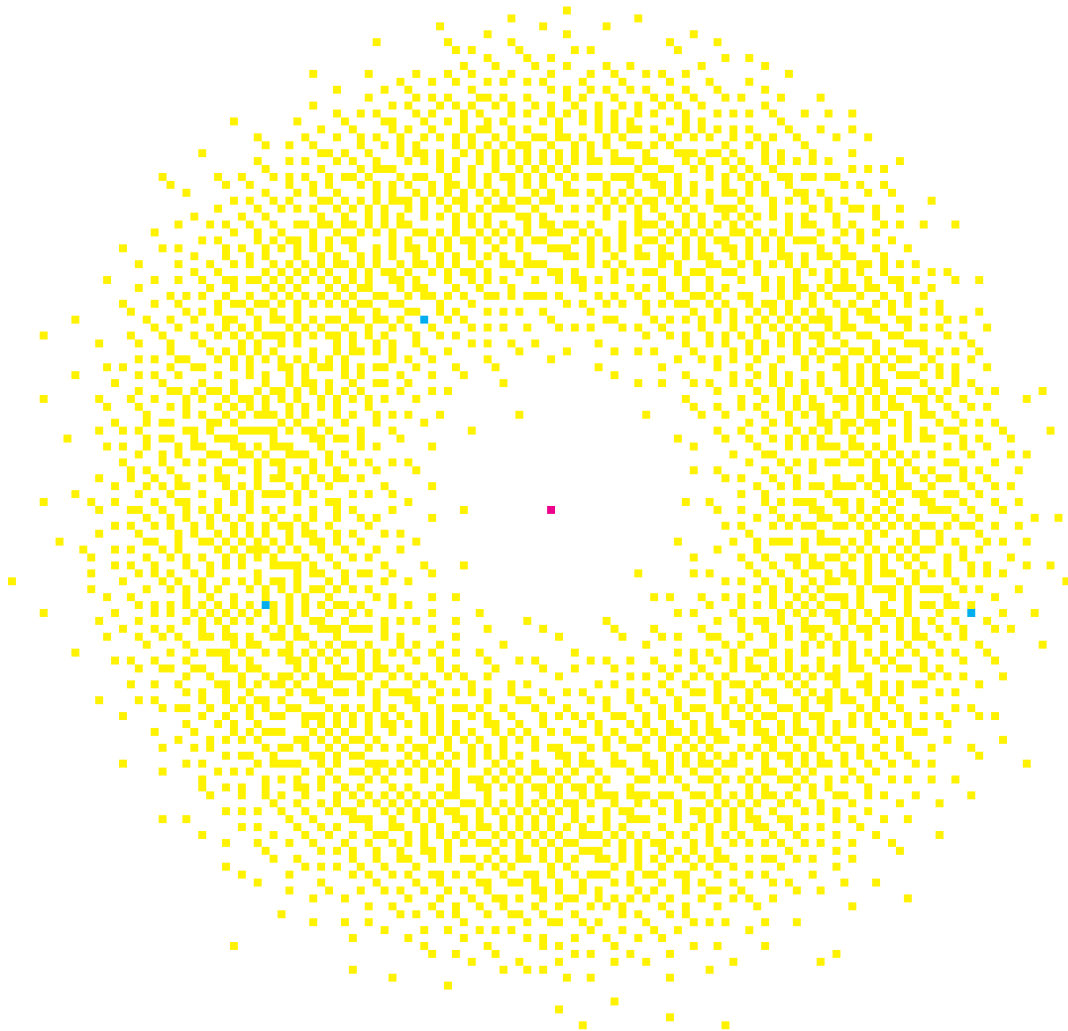


Figure 5. As a vague representation of a distribution, multiple points can be averaged, with that average plotted on the donut (or in its hole). No attempt is made to make this mathematically sound.

However, I would like to remain flexible as to what the center of the donut represents. So let us forget how the center edge represented mundanity as described with Figure 4, and instead look at a new dot appearing in magenta in Figure 5, relatively near the absolute center. This will be an imaginary dot—a second order of the imaginary, because the entire donut is itself imaginary. Call it a theoretical dot, one that would appear if you were able to average together the three cyan dots within the existing donut. Since the conceit is that the donut is a pluridimensional scatterplot, the magenta dot is simply an averaging of each of the collapsed-dimensional values of the three cyan dots. If the

flattened values for the cyan dots are (5.7, 10.1), (9.3, 16.8), and (22.1, 9.9), then the magenta value is (12.4, 12.3).

Now, I am a living, breathing human being, and so I write with intentionality. To say that I write is to say that I do so with intentionality. And you (where “you” here is a human reader, not a bot scrubbing the internet for a PDF version of this editorial) are reading this knowing that I (where “I” am also human, not an LLM) have intentionality. You have developed reading strategies that are entirely contingent upon this fact. Writing-and-reading as a construct is a matter of communication, after all. That is to say (in a particularly non-machine way that may already be getting annoying) at some level you are thinking about why I am doing what I am doing with these word sequences.

You may have perceived why I have gone to this level of abstraction to establish the donut as an imaginary distribution of authorial voices across countless variables. The magenta dot in Figure 5 is, after all, near the center of the donut, where no dots previously appeared. Surely this means something.

In fact, the way I am presenting this very writing may make it feel like we are arriving somewhere. After establishing the donut, I used it on four pages in a row, with the main text relegated to the same relative area at the bottom of the page. But now the main text is in a full column on a page without a figure. A pattern was established that is now broken.

It is time to take advantage of the flexibility of the donut to redefine what its center represents. Let us say that the initial donut (Figure 1) is not merely authorial voice, but the existing range of *human* authorial voice. There is a wider gamut of voice that is theoretically possible, and the donut merely plots the effective human range in practice. It thus also represents the data set of human writing on which LLMs are trained.

And what do LLMs do in their training? They use patterns of words from their data sets to create a vast array of weighted parameters in order to “predict” which token (a word, subword, or word sequence) should follow another in a given context and according to a given prompt. Even though that initial data set includes a wide gamut of authorial voice, the LLM’s outputs must have some particular kind of voice. Consequently, there is a necessary flattening or averaging together of pluralistic voices into a singular voice. The more points you sample in the donut to determine an average, the more likely you will arrive somewhere near its center.

But the donut has no center. This is precisely why I called it a donut in the first place.

Returning to the notion that a point represents an authorial voice (e.g., Gabriel García Márquez) and an area a writing genre (e.g., academic prose), we can imagine the full range of the former averaging into a narrow range of the latter. This is shown in

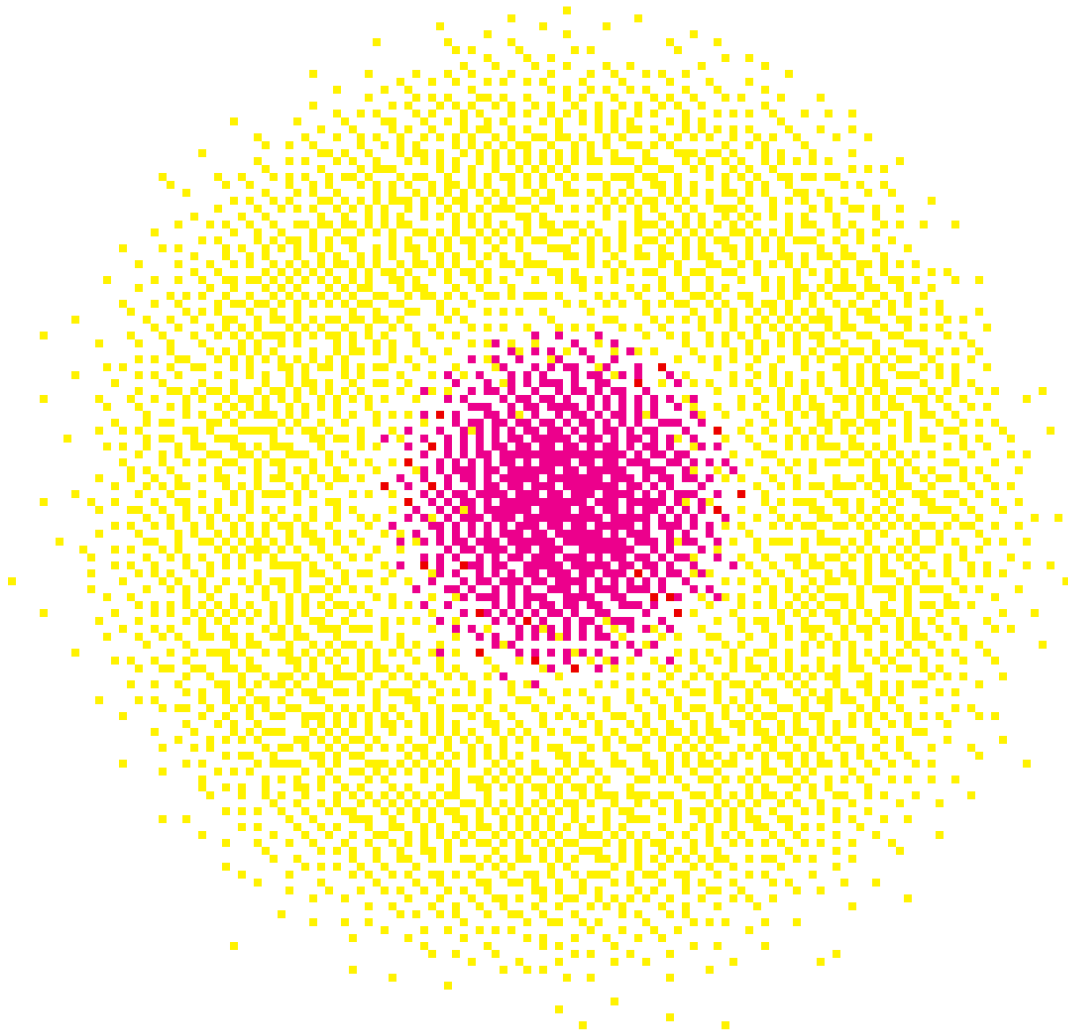


Figure 6. When averages are expressed for multiple sets of points, a donut hole emerges. This can represent the LLM's derivate voice, the result of averaging authorial voices in its training data.

Figure 6, where yellow is the original donut of variation in human writing, and magenta is the derivative voice of LLMs, which among other things are to a great degree prone to utilizing contrastive negation, the tricolon, and the word *delve*. These are all things that humans have done aplenty—they are legion in an LLM's training data. But the particular combination of these and other tendencies results in a writing style that no actual human employs. It is the empty center. And if writing can be acerbic, bold, and critical, then what is an LLM's prose? There is certainly some variation in it, but it also has an identity all too familiar to professors who were reading student prose and continue to read student-delivered writing assignments after LLM use became common practice.

To those with a special sensitivity, LLM prose exists in the uncanny valley. It is plausible while also being superficial in a way that is hard to put a finger on. The more you are aware of its patterns, the less you can overlook them.

The creation of the donut as a pluridimensional distribution was highly contrived, and it may be hard to swallow the premise of an average of values that is not representative of any value. But there is nothing fanciful about such a premise.

There is an oft-cited anecdote about the United States Air Force's role in the emergence of inclusive ergonomics, particularly in the proliferation of adjustable seating. It is from a report entitled *The "Average Man"?*, from the Aero Medical Laboratory and Wright Air Development Center, written by Gilbert S. Daniels in 1952. This report followed a period when there were too many peacetime aircraft crashes in the Air Force, with attention eventually being drawn to fixed cockpit designs made for the average airman (Mishra, 2023). Daniels (1952) noted in his report's introduction that "it is virtually impossible to find an 'average man' in the Air Force population" (p. 1). An anthropometric survey of 4,063 airmen was conducted in 1950, in which 131 measurements were made for each. Daniels took a subset of those measurements and began to eliminate individual airmen as potential exemplars of the average for which a fixed cockpit design would apply. In a first step, only 1,055 of the 4,063 airmen were of "approximately average stature" (p. 3). Once chest circumference was considered, that number dropped to 302. The ninth measurement was thigh circumference, which brought the previous step's total of six men down to two. In the 10th step, neither of those two airmen were found to have average crotch lengths (an unfortunately named measurement). Again, 131 measurements were available, yet it only took 10 to fail to find the average airman among a sample of 4,063.

Of course, the precision with which these measurements were made could render the result fatuous. But Daniels' "approximate average" was very much approximate. He offered the centimeter-precision of stature as an example, noting that at that degree of precision, a strict designation as average would be limited to anybody within a half centimeter of the mean value. This would have included only 6.4% of his sample, which he determined to be "far too restricting for actual application of data" (p. 6). So he used standard deviation within each measure to classify a value as an approximate average, at within ± 0.3 SD, such that "the middle 25 or 30% of a group" was approximately average for a given measurement (p. 6).

The Air Force's cockpits were effectively designed for nobody. There can, indeed, be an empty center. Donuts have them, and so do soap bubbles. In a post called "Gaussian Distributions are Soap Bubbles," which I must admit I cannot fully understand technically, Ferenc Huszár (2017) explains how high-dimensionality leads to a distribution that is more soap bubble than ball of mold. That is, it does not get dense in the middle.

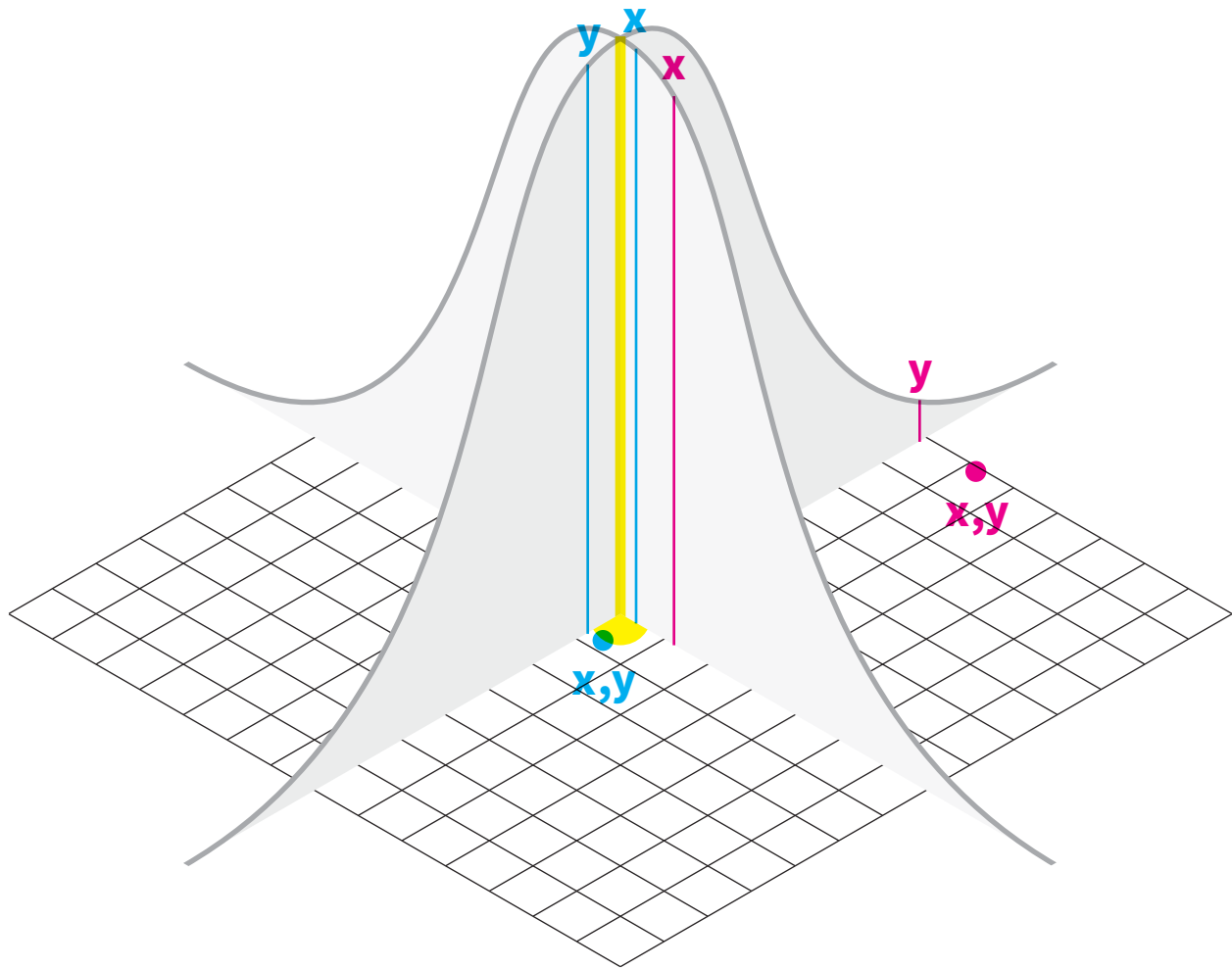


Figure 7. Two normal curves representing two variables among thousands. The implied third dimension (on this 2D page) permits visualization of the curves. The dots on the floor can be authorial voices.

Here is an explanation using normal curves that I can fully understand and thus explain. Figure 7 takes two bell curves and plots them orthogonally. Consider these to represent the distribution of two variables of authorial voice among thousands of additional variables. Two authorial voices are highlighted in Figure 7 according to these unnamed variables, one of which is rather close to the mean on both variables (in cyan), and the other of which is on a tail of one variable (in magenta). We could mathematically model more than these two variables, with each variable independent of all the others and aligned with a shared central mean (in yellow). But we cannot visualize this because

we have run out of the dimensions we perceive. One more variable would require that the gridded “floor” upon which the two voices are plotted be rendered in three dimensions, such that each of the variables could remain independent and their means could intersect at one point in space. And there would be thousands of dimensions to add. The thing to consider is that with each new dimension added and authorial voices plotted within it, every voice has a new variable for which its value could pull away from the mean—as the magenta point has already done with its y-value. There are far too many opportunities for voice plots to distance themselves from the yellow average (the cyan voice will do so again and again, despite its achievement of centrality at the start), and this leads to a soap bubble or donut shape.

Before we return to the donut visualization one last time, I want to extend the discussion to illustration and design to consider value. So across cases of writing, illustrating, and designing, what is desirable and undesirable about utilizing AI? Given the donut metaphor, when is it acceptable to be the donut hole, and not the donut? In each case, the answer is related to averaging, or homogenization.

For instance:

- ▶ If you are using an LLM to write soft prompts for another AI to utilize as you build a toolkit, it can be more effective to have it in a very controlled voice.
- ▶ If you are using an LLM to write an introductory email, it is possible that the recipient will recognize the voice as AI and think you did not directly communicate to them, and thus that it is okay to ignore you.
- ▶ If you are using a diffusion model to create customized banners for a child’s birthday party, the fact that the child gets their own special decorations will likely be appreciated.
- ▶ If you are using a diffusion model to add pictures to your slide deck, some of the viewers may recognize the style as AI and wonder if the concepts and script were also AI generated—especially if one of those viewers paid you to give the presentation.
- ▶ If you are using a coding agent to create a web-based form, it may be desirable for the visual style of the form to be familiar, which can make it easier to use.
- ▶ If you are using a coding agent to design a website from scratch for a client concerned with their company’s branding, losing control of its visual style could lead to unwittingly copying another company’s website—or, it may look rather like a host of other websites that were created in the same way.

There can be consequences to homogenization, especially when the degree of homogenization is not apparent to the human creator. Let us return to authorial voice in writing and add one more dimension to the thousands already represented by the donut: time.

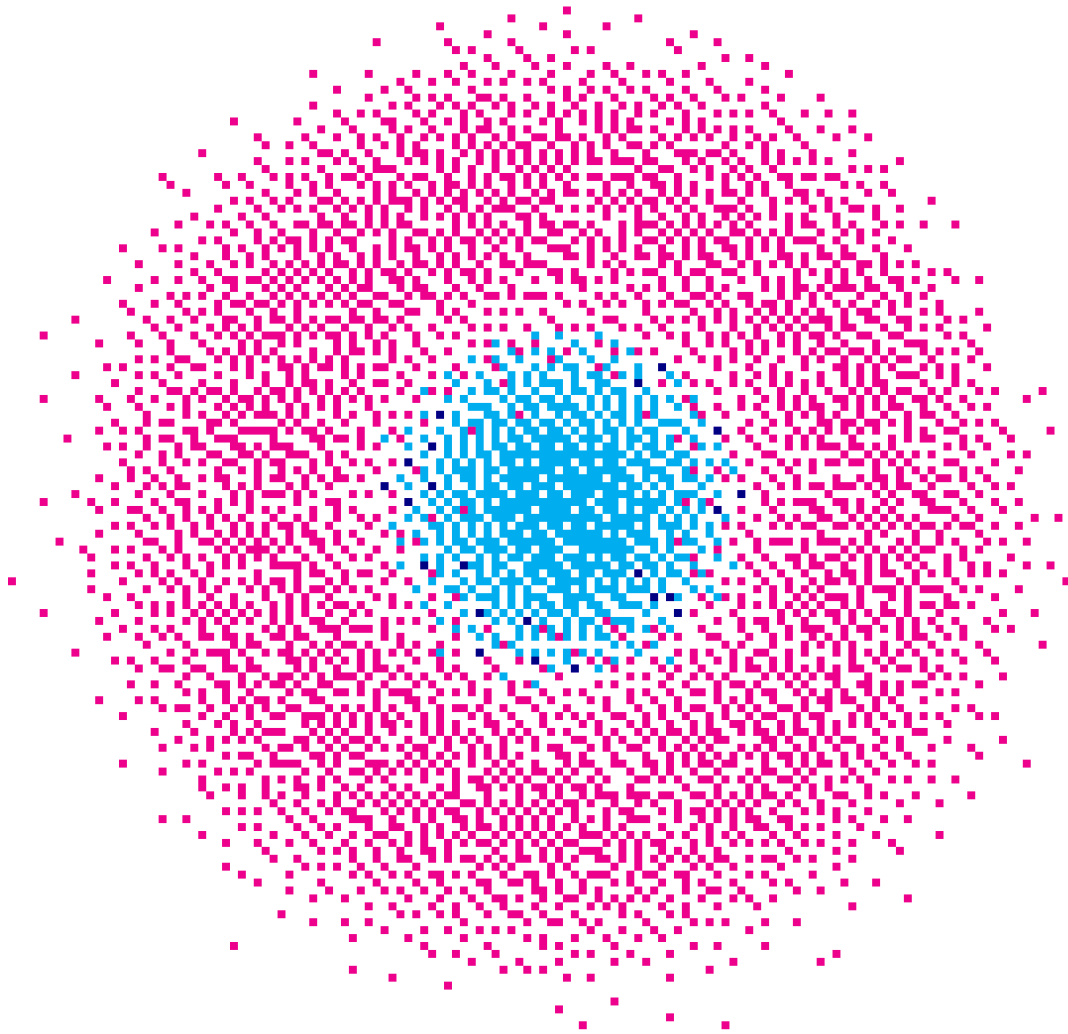


Figure 8. The donut of human variance and the donut hole of LLM variance. In an early time period ($t=1$), LLMs are a novelty and human variance is relatively unaffected.

With apologies for the repetition, Figure 8 (Figure 6 with different coloring) represents the range of human authorial voice in writing in magenta, and the corresponding range of LLMs in cyan. Unconventional forms of writing are on the periphery—or the authorial voices that in part define those forms—and the totality of the magenta donut thus represents the diversity of human authorial voice. I have adjusted the coloring for Figure 8 to make overlaps more apparent. Note that there are some points in this rendering where cyan and magenta create blue through overprinting. There are around 20 blue dots, though that number does not mean anything particularly.

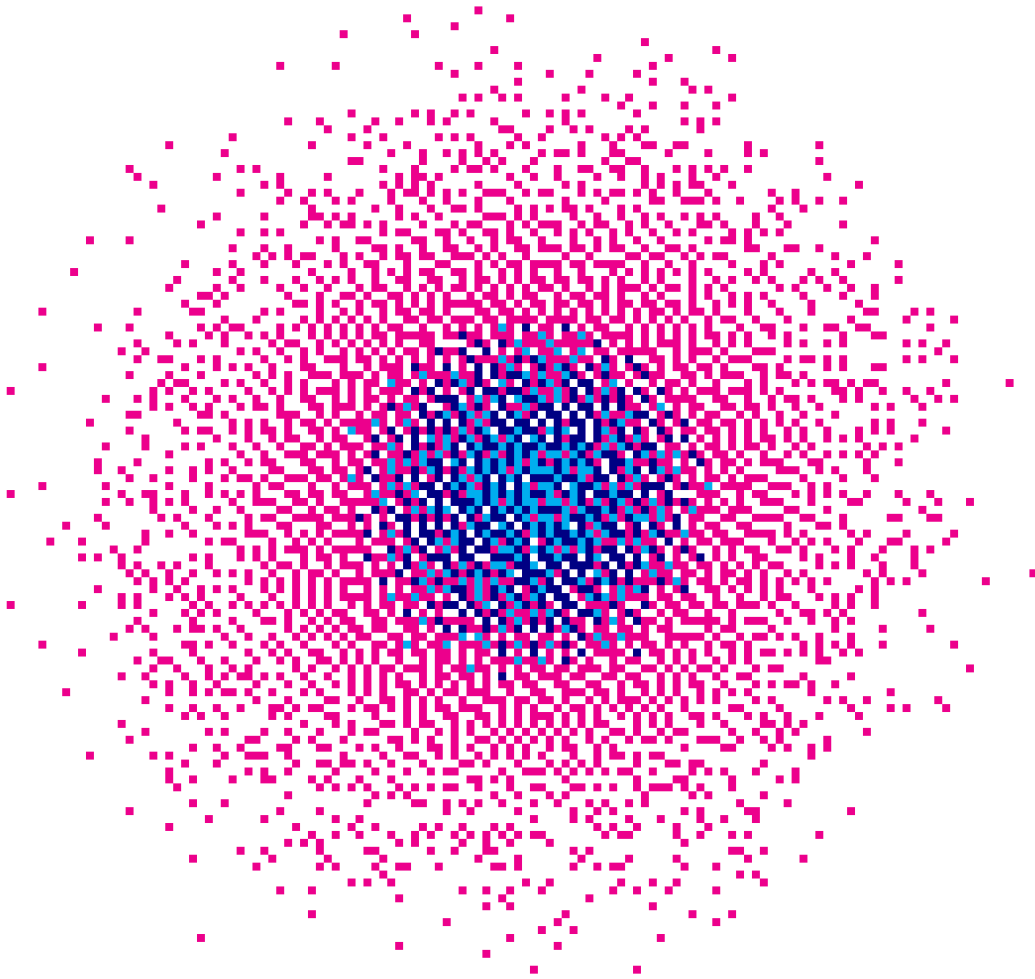


Figure 9. The donut of human variance and the donut hole of LLM variance. In a formative time period ($t=2$), LLM use has drawn the human range toward the center, but the gamut remains wide.

Figure 9 is a later moment in time, after LLMs hit the market and became familiar to millions of users. At this point, LLMs are being used to produce massive volumes of texts. Kids are growing up and have to produce texts themselves, as students and as professionals. And plenty of people, kids and adults alike, are using LLMs when they do not have to, but simply want to—for blog posts, for emails, and for anything else. Instead of these new writers developing their own voices, much of the language they produce in fact comes through an LLM, and it takes on that voice. Maybe not entirely, but it pushes what would otherwise be their unique human authorial voice a little toward the center.



Figure 10. The donut of human variance and the donut hole of LLM variance. In a mature time period ($t=3$), LLM use has drawn the human range toward the center. Outliers become more rare.

The longer this happens, the more real the LLM voice becomes, because it increasingly dominates the corpus that the next generation of LLMs is trained on, yes, but also because the human beings who are reading an increasing proportion of writing in the LLM voice begin to emulate it, even when they are not actively using AI. Figure 10 captures this period, where the original donut hole is now part donut. If you look at the outliers, they remain, and you might think we retain a good deal of variety in writing. Certainly we do, and texts produced years and centuries before LLM technology was invented remain available somewhere. Maybe those could represent the outliers.

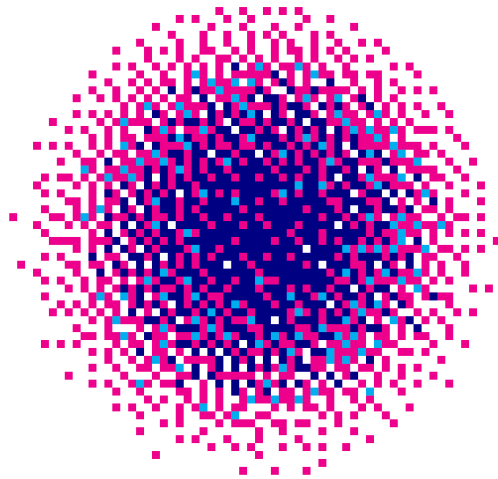


Figure 11. The donut of human variance and the donut hole of LLM variance. In a depressing time period ($t=4$), LLM use and model collapse has resulted in a form of human collapse—LLM emulation.

But I think, in practice, that what is being read most is what has been written most recently—not only on social media and in blog posts, but in academic writing and many places beyond. And so the outliers in Figure 10 may instead represent practicing writers who did not take to the LLM, and who continue writing without it. But as of Figure 11, as the center increasingly dominates, the outliers disappear—not because those authors died out, but because they are crowded out in the cacophony of texts. Their unique forms of writing are diminished in the overall writing market share.

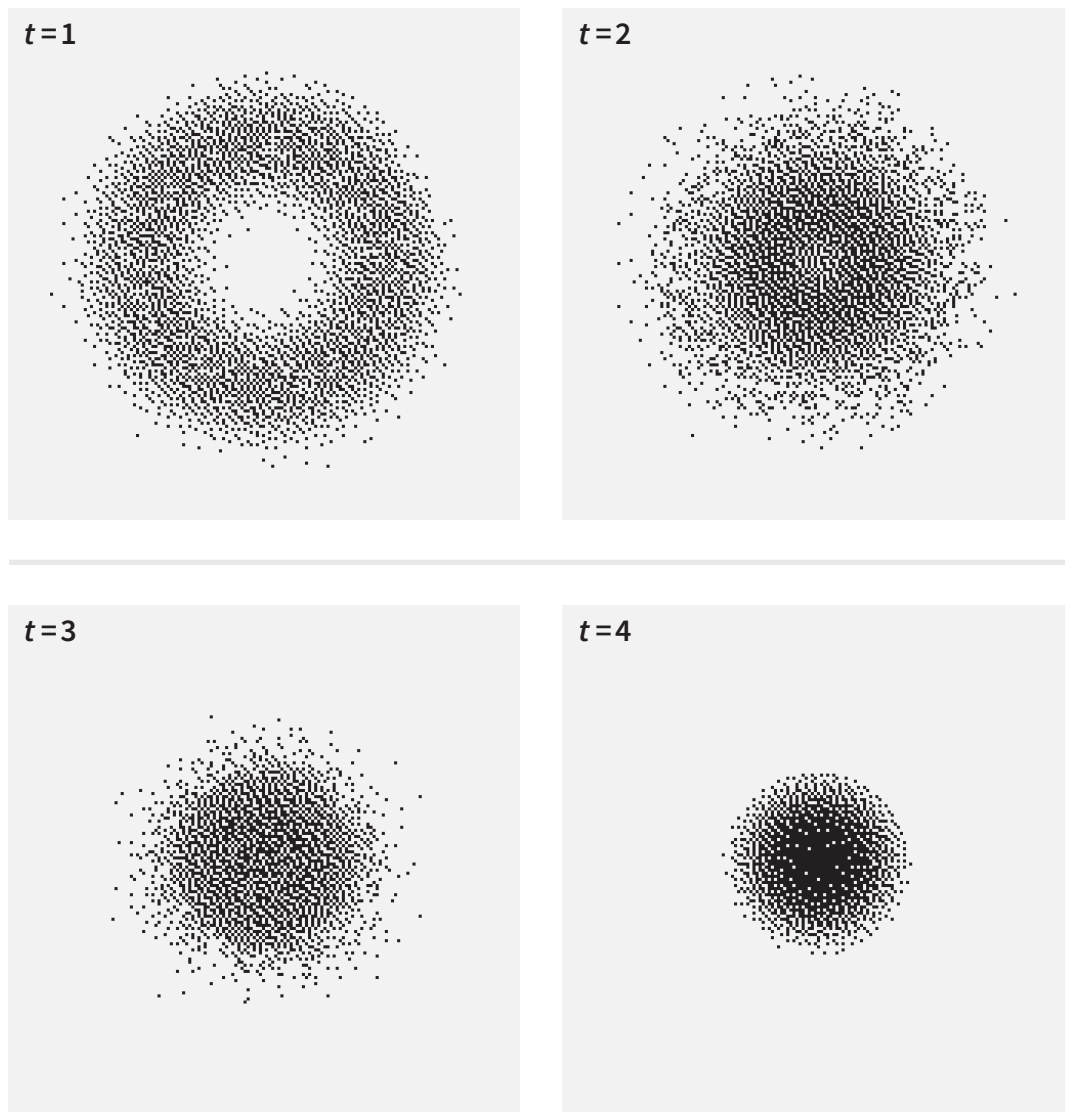


Figure 12. The donut of human variance is drawn to the center over time by LLM use and ultimately LLM collapse. Human variance eventually becomes the donut hole.

It is evident in Figure 12, which brings these four moments in time together, that the result is a dramatic reduction in the diversity of authorial voices with which readers can engage. This loss of diversity is unappealing.

What I just described on the machine side is likely familiar as *model collapse*, which Williams (2026) frames as such:

People tend to imagine model collapse as some dramatic cliff where a model suddenly starts producing gibberish. The reality is more subtle and, honestly, more dangerous because of that subtlety. What you get is a slow erosion of

variance. Outputs become more generic. Edges get sanded down. The model starts producing text that reads fine on a surface level but carries less information per sentence.

More worrisome is the impact on actual human writing. It too collapses toward the center when what human authors read is increasingly AI-generated. Or not. Things do often turn out all right when we are emotionally averse to change. So I do not find it hard to imagine a difference of opinion, or a lesser degree of concern.

One reason to temper concern is that there is a middle ground between texts generated from a single prompt and purely human writing. Surely using an LLM in a punctuated, modular fashion reduces the averaging effect. It is exceedingly difficult to assess the human and the machine in this form of human-machine teaming. So though the principle holds, and there is reason for concern, it is *gauche* to simply say, “AI is bad” (*L’IA, c’est mauvais*).

And is model collapse a realistic threat, or does it assume the steady state that never really emerges? Because the technology is continually being innovated, and far more importantly, other technologies (AI or otherwise) may render it a moot point if there becomes a truly distinct way to do many of the same things. Prediction can be something of a fool’s errand.

But we do have LLMs right now, and they are dominant. And there must be enduring consequences for so much so-called “AI slop” flooding our collective textual corpus, with no truly reliable way to detect it. Whatever ends up happening, the impact of LLM technology is not going away soon, so we should endeavor to understand LLMs as best we can.

As previously mentioned, the final article in this issue of *Visible Language* (Peterson et al., 2026) assembles a multidisciplinary team to offer an interdisciplinary interpretation of the degree to which it can be said that LLMs write, learn, hallucinate, and think. The authors consider invoking these constructs in the context of LLM technology to be a matter of metaphor, and they discuss the negative entailments of doing so. They (really “we,” as I am the same Peterson) offer alternative metaphors, such as the LLM as a fuzzy thesaurus, along with paired misleading and corrected statements. For instance, instead of stating that “LLMs think,” you might state either that LLMs “produce outputs that resemble products of thinking,” or that they “recombine semantic parts that are separate from their original wholes” (p. 231). Precision is often less pithy.

The authors cite Flowers and Hayes’ (1981) cognitive process theory of writing. Interestingly, this process theory appeared a year earlier in *Visible Language* 14.4, with less description but with a colorful metaphor. The way Hayes and Flower (1980) describe writing is akin to a learning process in which the writer to some degree discovers what

they want to write by virtue of writing it. The authors explicitly address the themes of the average and of individuality:

It is apparent that not all writers write in the same way. For example, some writers plan their essays from beginning to end before they write a single word of text, while others never seem to look beyond the next sentence. Further, some writers seem to write with their readers constantly in mind, checking frequently to be sure that they have taken the reader's knowledge and attitudes into account. Others appear serenely unaware that an audience could fail to understand what they, in good faith, have intended to say. (p. 390)

They endeavor to produce a model of writing that captures writing's richness. They continue:

In modeling we can deal with such differences in either of two ways. We can choose to construct a model of the "average" writer and delay until some more propitious time the description of differences among writers. This approach has the merit of simplicity. Further, if things work out well, a model of an average writer might be useful in characterizing individual differences. Thus, models for individual writers might prove to be minor variants of the average model. However, this approach may have the disadvantage that averages sometimes suffer from—the average may be representative of no one. Thus, we sincerely hope that no one has the average number of children—two and a half—nor would we want anyone to have to eat an average course at dinner, which might be a compromise between appetizer and dessert such as oysters with chocolate sauce. (p. 390)

And thus we have a flavor for the donut hole that human writing may become.

Let me provide one final note on writing from Hayes and Flower (1980). A major part of their model is *goal setting*. Writers modify their writing goals dynamically in the act of writing: "When writers choose goals, they are by no means stuck with them. When a goal is difficult to meet, the writer can respond to the difficulty by modifying or abandoning the goal" (p. 397). This is one way to think through the mechanisms of the act of writing that might be lost when part of the process is offloaded to an LLM.

I have a reflexive example. At some point in conceiving of this editorial, I thought of repeatedly returning to a visualization of the donut. As a graphic designer—and in a rare occurrence for a writer, as the person who would typeset these pages—this idea brought along with it the specifics of how it might look. I liked the idea of creating a pattern of donut figures appearing on the same place in their respective pages, with the running text appearing beneath them—I set a goal. But this created a structural challenge of sorts. When I got to what is now page vi, I had much more to say between Figures 1

and 2 than would fit beneath them, but much less than was needed to fill out an extra full page. So I had to write a little more—I had a new goal. Without this typographic need, I would likely have continued on. Instead, it occurred to me to present a few lists to better explain what authorial voice might entail—the lists would take up space. I created the three lists: impressions, patterns, and terminology. I did not have this miniature typology in mind until that moment. In being forced to create it, I changed or clarified my own idea of what authorial voice entails. That is important to recognize; the idea was not already there, hidden in my conceptualization. I also rather accidentally created an opportunity to seed the discussion with some telltale signs of LLM outputs (e.g., a high frequency of contrastive negation), and along with that came an opportunity for self-referencing. I found myself utilizing contrastive negation almost immediately and added the little note: “(Is that a contrastive negation?).” I liked this because some readers will have a shared experience with me, that becoming aware of LLM patterns can rather annoyingly make you doubt your own writing when your heightened sensitivity is turned inward. And because of the recurrence of this annoying experience, along with my fears as demonstrated in Figure 12, I am rather interested in writing more humanly than in writing more correctly—the latter was a goal of my early academic career, but not now.

I have another example of goal setting in writing, which may be more compelling because of the stature of the writer: Robert Frost. In *How Does a Poem Mean?*, John Ciardi and Miller Williams (1975) address Frost’s “Stopping by Woods on a Snowy Evening,” which ends in the repetition of the line “and miles to go before I sleep” (I cannot recommend enough reading Chapter 1 of *How Does a Poem Mean?*). Ciardi and Williams explain that the second occurrence of the line “suddenly” transforms words into symbols (p. 9). But the end result is not what I am interested in here. Because they go on to consider Frost’s experience of writing the poem to gain insight into the performance of the poem:

Now look at the poem in another way. Did Frost know what he was going to do when he began? Considering the poem simply as a piece of juggling one cannot fail to respond to the magnificent turn at the end where, with one flip, seven of the simplest words in the language suddenly dazzle full of never-ending waves of thought and feeling; or—more precisely—of felt thought. Certainly an equivalent stunt by a juggler—could there be such an equivalent—would bring the house down. Was it to cap his performance with that grand stunt that Frost wrote the poem?

Far from it: the fact that must not be overlooked is that *Frost did not know he was going to write those lines until he wrote them*. Then a second fact must be

registered: *he wrote them because, for the fun of it, he had got himself into trouble.*
(p. 10, emphasis original)

Ciardi and Williams go on to explain how the novel rhyming scheme Frost initially adopted led to a conundrum in the final stanza. For my purposes I have explained enough and will leave that be. The point is that this celebrated poem was not necessarily the result of some instantaneously discovered insight, but that it likely emerged from the dynamics of goal setting. And if you reduce the goal setting in a writing process, a great deal of value may go with it.

Working backwards through this issue of *Visible Language*, let me now mention the “next” article, which presents an opportunity to address this editorial’s title. I have covered the donut and the donut hole, and the chocolate and oyster flavoring. But in emphasizing authorial voice in the donut, among the many things the metaphor can extend to, I have not seriously covered originality yet. Halldórsson (2026) relays a project completed by students in design schools in Eskilstuna, Sweden, and Perth, Australia. Students had to use generative AI to create travel advertisements for the other city (e.g., a student in Perth used AI to design for Eskilstuna).

Halldórsson describes early outcomes that “generally conveyed broad connotative impressions—‘a European town’ or ‘an Australian city’—rather than precise depictions” (p. 199). The wrong geographic features and endemic fauna would regularly appear. They were appropriate at the country or regional level but not the civic level. This can be mitigated in any number of ways, and of course the students over the course of the project developed and employed mitigation strategies. But there remains an underlying sense that originality is under threat.

So moving beyond authorial voice to focus on originality, let us return to writing but consider content rather than presentation. Scholarship is predicated on originality. A major criterion is novelty, that something is publishable because it is new, and this is a large part of its value. It becomes problematic for many researchers using the same tools to produce research when those tools push all of their work toward the same empty center. And if researchers are using the same tools—and regardless, AI tools utilize the same corpus to construct an averaged center—their research will tend to converge. The size of the donut—the gamut, the variety—is a measure of originality. Figure 12 tells that story, too.

The remaining articles from this issue seemingly let us step away from AI for a moment. Maciel (2026) provides copious examples of “graphemic construction” in the form of illustrated letters that far predate AI. He considers the degree to which illustrated letter-forms conform to an existing letter, or form a new letter. The 21 figures in the article



Figure 13. Miriam Martincic, Barcelona, 2023. Mixed media. From Maciel (2026) in this issue (as Figure 9 in that article, p. 172).

provide a rich demonstration of human originality, the stuff on which our AI models have been trained.

For instance, Figure 13 is one of Maciel's most recent examples, a rich and rewarding illustration by Miriam Martincic. In the same way that Ciardi and Williams (1975) discuss the poem as an experience—"a performance of forces" (p. 10)—we can engage with Martincic's illustration and know that it is wildly insufficient to call it an "E." Maciel (2026) calls it a hybrid letterform having "a negotiated condition in which structure and image co-operate" (p. 172), and that co-operation is active.



Figure 14. Laura Sisti's picturebook illustrations of children playing with gendered signifiers: female at top, and male at bottom. From Sisti and Peronti (2026; as Figures 9 and 10 in that article, p. 137).

Finally, Sisti and Perondi (2026) studied the impacts of gendered and gender neutralized picturebook illustrations on stereotyping in children's interpretations. The first author, Laura Sisti, produced the illustrations. The control afforded by her ownership of the imagery permitted precise and deliberate manipulation for careful stimuli preparation for an empirical study. Figure 14 provides details of her female- and male-gendered illustrations. They are gender encoded through "color, character morphology, and symbolic gender cues," as commonly found in "early-childhood educational materials in Italy" (p. 133).

The authors' results are quite interesting, but for the present editorial I just want to emphasize the *precision* with which Sisti was able to develop their study's stimuli—hand-drawn and guided by the human eye. Handless, eyeless, and unthinking AI models have not been trained on these illustrations yet, but they will be.

References

- Daniels, G. S. (1952, December). *The "average man"?*. Wright Air Development Center. <https://web.archive.org/web/20210411010707/https://apps.dtic.mil/dtic/tr/fulltext/u2/010203.pdf>
- Flower, L., & Hayes, J. R. (1981). A cognitive process theory of writing. *College Composition and Communication*, 32(4), 365–387. <https://www.jstor.org/stable/356600>
- Halldórsson, H. B. (2026). Are you sure? Exploring the challenges and opportunities of using generative AI in communication design projects: A cross-cultural case study of student reflections, bias, and critical engagement. *Visible Language*, 60(2), 183–209. <https://www.visible-language.org/Issue-60-2/VL-60-2-Halldorsson.pdf>
- Hayes, J. R., & Flower, L. S. (1980). Writing as problem solving. *Visible Language*, 14(4), 388–399. <https://journals.uc.edu/index.php/vl/article/view/5308>
- Huszár, F. (2017, November 9). Gaussian distributions are soap bubbles [Blog post]. *inFERENCE*. <https://www.inference.vc/high-dimensional-gaussian-distributions-are-soap-bubble/>
- Maciel, H. (2026). Illustrated letterforms: A taxonomy of graphemic construction. *Visible Language*, 60(2), 156–182. <https://www.visible-language.org/Issue-60-2/VL-60-2-Maciel.pdf>
- Mishra, P. (2023, April 20). Opinion: Gilbert Daniels, the gardener who changed our world [Blog post]. Mary Lou Fulton College for Teaching and Learning Innovation, Arizona State University. <https://education.asu.edu/about/news/opinion-gilbert-daniels-gardener-who-changed-our-world>
- Peterson, M., Patton, C., Green, M., & Díaz, B. (2026). The concept of a large language model: Misconceptions about and beneficial conceptualizations of LLM technology. *Visible Language*, 60(2), 210–239. <https://www.visible-language.org/Issue-60-2/VL-60-2-Peterson-et-al.pdf>
- Sisti, L., & Perondi, L. (2026). Graphic neutralization and gender stereotypes in early childhood: How visual variables shape children's interpretations. *Visible Language*, 60(2), 124–155. <https://www.visible-language.org/Issue-60-2/VL-60-2-Sisti-Perondi.pdf>
- Williams, A. (2026, March 25). Model collapse is already happening, we just pretend it isn't [Blog post]. *Blog@ACM*. Communications of the ACM. <https://cacm.acm.org/blogcacm/model-collapse-is-already-happening-we-just-pretend-it-isnt/>