

Caricaturing Shapes in Visual Memory



Zekun Sun, Subin Han, and Chaz Firestone^{id}

Department of Psychological and Brain Sciences, Johns Hopkins University

Psychological Science
1–14

© The Author(s) 2024

Article reuse guidelines:

sagepub.com/journals-permissions

DOI: 10.1177/09567976231225091

www.psychologicalscience.org/PS



Abstract

When representing high-level stimuli, such as faces and animals, we tend to emphasize salient features—such as a face's prominent cheekbones or a bird's pointed beak. Such *mental caricaturing* leaves traces in memory, which exaggerates these distinctive qualities. How broadly does this phenomenon extend? Here, in six experiments ($N = 700$ adults), we explored how memory automatically caricatures basic units of visual processing—simple geometric shapes—even without task-related demands to do so. Participants saw a novel shape and then immediately adjusted a copy of that shape to match what they had seen. Surprisingly, participants reconstructed shapes in exaggerated form, amplifying curvature, enlarging salient parts, and so on. Follow-up experiments generalized this bias to new parameters, ruled out strategic responding, and amplified the effects in serial transmission. Thus, even the most basic stimuli we encounter are remembered as caricatures of themselves.

Keywords

caricature, shape, memory, complexity, open data, open materials

Caricature is the practice of exaggerating distinctive features when representing some stimulus or concept (Fig. 1a). Though often used intentionally for political or comedic effect (Perkins, 1975), human cognition sometimes engages in a caricaturing process of its own, encoding and even misremembering stimuli in exaggerated form. For example, caricatures are often judged as the best likeness of familiar faces (more so than the actual faces themselves; Rhodes et al., 1987), and participants are slower to differentiate a face from its caricature when first seeing the face and then the caricature than vice versa (suggesting that the first stimulus was encoded in exaggerated form; Mauro & Kubovy, 1992). Such exaggeration is thought to aid related processes, such as recognition and categorization (Benson & Perrett, 1991; Chang et al., 2002; Mauro & Kubovy, 1992; O'Toole et al., 1997; Rhodes et al., 1987), and may be related to neural pattern separation that protects mnemonic representations from interference by keeping them distinct (Yassa & Stark, 2011; for exaggeration effects explained by different neural mechanisms, see Zhao et al., 2021, as well as Chanales et al., 2021, and Drascher & Kuhl, 2022).

Mental Caricatures: From Faces to Shapes

How general is the phenomenon of mental caricature? Does it apply to any stimulus, regardless of its category, class, or context, and regardless of the demands of a given task? Indeed, it has been speculated that caricatured representation may be appropriate not just for faces, but also for other visual stimuli such as animals and objects. For example, Rhodes et al. (1987) suggest that upon encountering a new species of bird, we might encode the bird with its distinctive features exaggerated, so as to better distinguish it from other birds. Early studies provided similar evidence for simple visual patterns; for example, when a shape has a gap differentiating it from other shapes, the gap is misremembered as wider (Garner, 1962; see also Gibson, 1929). But unlike caricatured encoding in face memory, this result was more localized rather than holistic, and it was likely

Corresponding Author:

Chaz Firestone, Johns Hopkins University, Department of Psychological and Brain Sciences
Email: chaz@jhu.edu

induced by explicit task demands (which involved differentiating a target shape from distractors). More generally, whereas caricature-like processing can be adaptive when one is explicitly required to distinguish multiple items held in memory at one time (see, e.g., Chunharas et al., 2022, on repulsion effects), it is unknown whether memories of simple, single visual forms show caricature biases in the absence of considerable external pressure to distinguish them from one another (Donderi, 2006).

Here, we explore caricatured representation for what are among the simplest visual stimuli we encounter: basic geometric shapes (Fig. 1b). Shape stimuli have a long history in research on visual perception and memory (Attneave & Arnoult, 1956), and they are interesting here too for several reasons. First, compared to faces, familiar objects, or even novel categories (e.g., Davis & Love, 2010; Gauthier & Tarr, 1997), random geometric shapes have little social significance, few preexisting associations, and few salient characteristics by which to organize them, making them well-suited to avoid contamination by high-level goals, knowledge, and other such factors. Second, despite their simplicity, geometric shapes can vary along several well-defined parameters, affording control over their features in ways that are not always possible with more naturalistic stimuli like faces and animals. Third and finally, advances in computational geometry allow for succinct and standardized measures of shapes' information density (Feldman & Singh, 2005; Siddiqi et al., 1999),

Statement of Relevance

We often portray images in caricatured form, exaggerating their distinctive qualities for political or comedic effect. Do our memories also engage in a caricaturelike process when encoding and remembering what we see? The present work explored a caricature bias for the most basic visual stimuli we encounter: simple geometric shapes. When observers saw a shape and had to reproduce it from memory, they tended to create *shape caricatures* that exaggerated each shape's qualities, even though the task itself gave them no obvious reason to do so. This work suggests that our minds encode the world around us in ways that emphasize what is distinctive, even when the relevant stimuli have no particular significance and even when there are no task demands that require it.

making it possible to quantify memory distortions in precise and illuminating ways.

At the same time, exploring caricatures for such stimuli invites a further question: What does it mean to caricature a shape? How does one exaggerate a meaningless, contextless, blob? For face caricatures, it is relatively clear which features to emphasize and how to do so. For example, a caricature of Albert Einstein

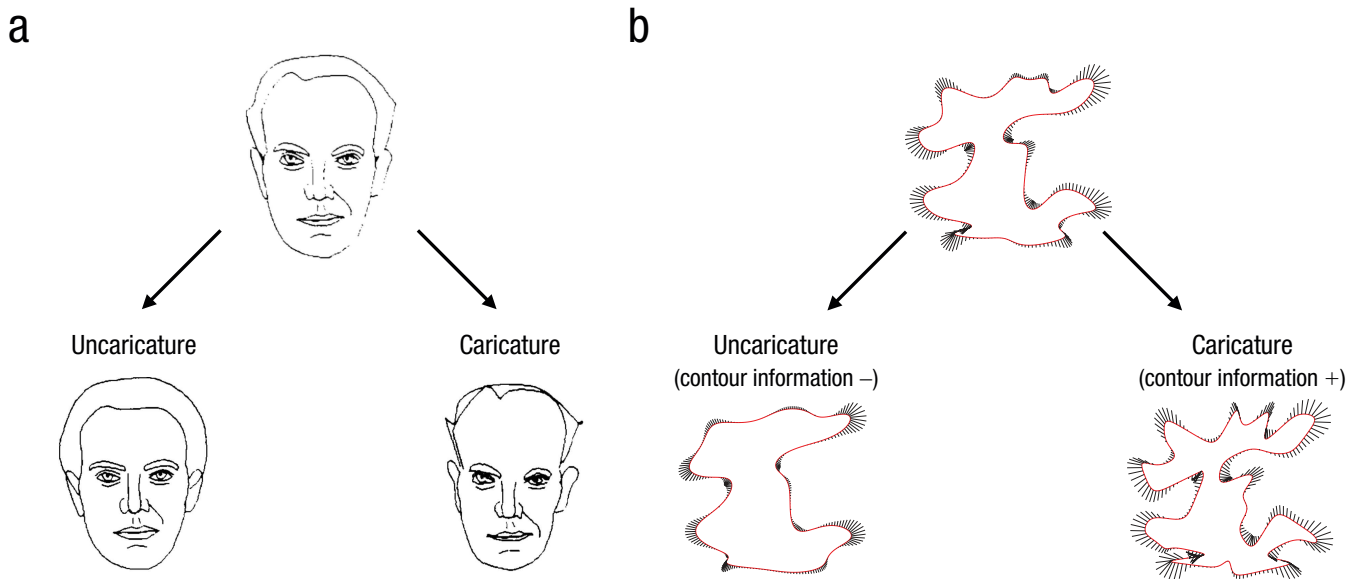


Fig. 1. (a) Illustration of facial caricature, adapted from Rhodes et al. (1987). (b) The present work explores caricatures for basic visual shapes by manipulating their features to normalize or exaggerate them. The normalization process decreases the magnitude of the turning angles along a shape's contour; in contrast, exaggeration increases these magnitudes. The radiating bars indicate the surprisal values of points along the shape's contour, which increase as the shape becomes caricatured.

might make his hair frizzier, his cheeks rounder, and so on. More formally, work in computer graphics has automated caricature generation by computing the distance between a given face and a norm comprising many faces, and then adjusting the face's features away from the norm (Brennan, 1982; see also Lee et al., 2000; Rhodes et al., 1987). But it is not obvious that such an approach even applies to geometric shapes, if only because there is no straightforward candidate for a "normal" or "average" shape (or, perhaps, there are many such candidates, and it would be unclear which to choose; see, e.g., Amalric et al., 2017).

Here, we approach shape caricaturing from the perspective of information theory. It has long been proposed that the information content of a shape is naturally characterized in terms of changes to its bounding contour, such as vertices, protuberances, curves, and other deviations from smoothness (Attneave & Arnoult, 1956; Norman et al., 2001). Along these lines, we conceive shape caricaturing as a process that increases these changes and further exaggerates this contour information so that the turns of a shape's visible boundary appear even more salient and distinctive (Fig. 1b). With this approach in hand (see the Method section for more detail), we ask whether the mind engages in a caricaturing process even for stimuli with no obvious norm and even when there is no task-related pressure to do so (such as long retention periods, active interference from other stimuli, and so on).

The Present Experiments: Reproduce the Shape

Do people spontaneously misremember even basic visual stimuli in exaggerated form? To test this question, we asked whether participants who must reproduce a recently seen geometric shape tend to create a caricatured version of that shape. We generated a library of novel shapes whose contour information could be manipulated; then we showed these shapes to participants, one at a time, and simply asked them to reproduce what they had just seen a moment earlier. We then analyzed memory distortions in terms of changes in contour information. To foreshadow the key results, shapes were consistently remembered as exaggerated versions of themselves.

Open Practices Statement

Demos of these experiments can be viewed at <https://perceptionresearch.org/caricature/>. Data, experiment code, stimuli, and other relevant materials for all studies are available at <https://osf.io/7grk8/>. These studies were not preregistered.

Experiments 1 and 2: Shape Caricatures in Visual Memory

As an initial test of shape caricatures in visual working memory, Experiments 1 and 2 asked participants to briefly view a shape and then adjust a copy of that shape until it looked like the one they had just viewed.

Method

Participants. We recruited 50 participants for each experiment from the online platform Prolific (<https://www.prolific.co/>). A power analysis based on a pilot suggested that this sample would have power above 99% to reveal effects of the sort explored here. This experiment and all others reported here were approved by the Johns Hopkins University Institutional Review Board.

Stimuli. Our task requires that a shape be adjustable, such that its features can change in a continuous fashion while still preserving, in some meaningful way, a sense of its being the same object. To achieve this, we first created 30 "parent" shapes, which were highly irregular polygons that served as the maximally exaggerated shape within its shape family (i.e., the shape with the greatest amount of information along its contour). The shape-generation process started with selecting a set of randomly located points to be vertices of the shape's edges. We then connected these points using the method of Delaunay triangulation. Facets along the boundary of this triangle mesh were removed until the resulting polygon had a predetermined number of sides. Finally, the boundary of the polygon was resampled to 1,000 sequential points and smoothed to appear natural or organic. Additional constraints included a minimum angle of at least 10° and a maximum angle of 170°, so that the turns on the contours of shapes would be discernible.

For each of these shapes, we gradually smoothed its contours to create a spectrum of similar shapes (using ShapeToolBox 1.0; Feldman & Singh, 2006). A box mask was applied to an increasing number of consecutive points on the contour of the parent shape so that the curve consisting of the points in the mask was flattened to the averaged value along each axis. We started by setting up the mask size as 4 points and then smoothed the shape with that mask twice to create the next shape; then we increased the size of the mask by one unit to create the next 25 shapes, eventually creating a series of 51 structurally similar but gradually smoothed shapes (one parent shape and its 50 children). All the images were shown at an approximate size of 200 × 150 pixels on the participant's display, slightly differing across images.

We computed the contour surprisal (i.e., shape information) for all shapes (including parent shapes and

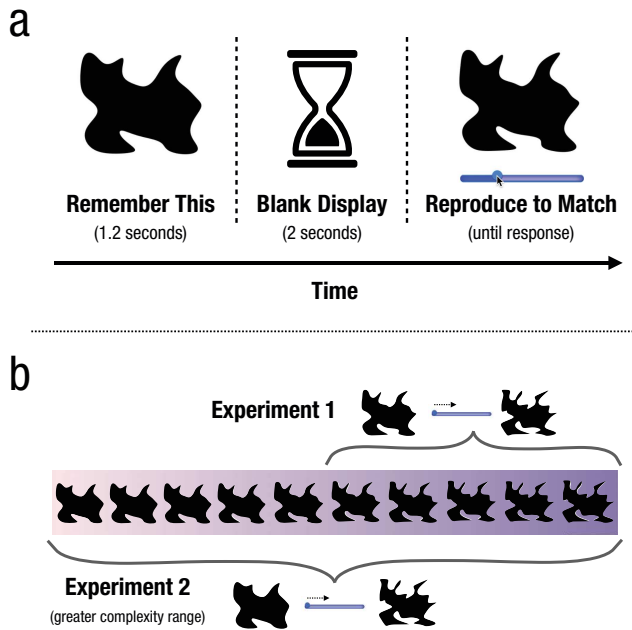


Fig. 2. Illustration (a) of our shape-reproduction task, seen in Experiment 1. Each trial begins with a random-looking shape appearing for 1.2 s. After a 2-s delay, participants see a copy of that shape (appearing at a different contour-information level) and are asked to reproduce the shape they just saw by moving a slider to exaggerate or normalize the test shape. The shapes in Experiment 1 are drawn from a range of entropy values within a single shape family (b); there are 30 shape families in all (only one of which is shown here). Experiment 2 doubled the range by including even more simple shapes; thus, the target shapes were sampled from a wider degree of exaggeration, and moving the slider caused more dramatic changes of shapes' appearance. In the actual experiments, the sign of the slider was randomly assigned between participants, and the starting position of the slider was randomized on each trial. (Four additional experiments tested other ranges of information density. See more details in endnote 2.)²

derivative shapes)—a measure based on the magnitude of point-to-point turns along the contour of a shape (Fig. 1b). An intuitive way to capture this measure might be to imagine a person walking along the contour of a shape; the more often and dramatically this person changes direction (so that their next step was not easily predictable from their previous step), the less predictable and the higher the surprisal of that step. The cumulative surprisal is the sum of the surprisal of the set of points along a shape's contour.¹ This quantification derives from information-theoretic approaches to visual perception (Attneave, 1954; Feldman & Singh, 2005), which are complemented by experimental evidence that the visual system uses such contour information to guide object recognition, feature detection, and other processes (e.g., Baker et al., 2021; Barenholtz et al., 2003; Norman et al., 2001). This information-theoretic measure is especially appropriate for our purposes, because the exaggeration of features in caricaturing a shape would amplify the overall information of the shape's contour. As contour information increases, the

shapes appear to have exaggerated features, such as amplified curvature, enlarged salient parts, and so on.

As can be seen in Figure 2b, the shapes yielded by this procedure are relatively complex, with the center of the slider representing a fairly angular or pointed shape (rather than an overly smooth one). Thus, Experiment 2 used the very same design as Experiment 1 but with twice the range of contour information by including more rounded and blunt simple shapes (so that the center of the range corresponded to a much simpler and normalized shape). To achieve this, we used the same procedure as before to further normalize the shapes used in Experiment 1, and we derived another 50 shapes continuing into the simple end of the spectrum. The new spectrum that resulted included 101 shapes from a given family—the 51 shapes used in Experiment 1 and the 50 newly derived shapes. (Note that there were still 51 steps on the responding slider, each step corresponding to one of every pair of shapes in a sequence of 101 shapes, resulting in greater change between each step of the slider).

Procedure

In the task, participants briefly saw a single shape, which then disappeared; a copy of it then returned, and participants had to adjust the copy to match what they had seen.

On each trial, a shape from one of the 30 families appeared on the screen at a random level of information density between 20% and 80% of the full range. The shape remained on the display for 1.2 s. Next it disappeared, and was replaced by a blank display that lasted for 2 s. Finally, the shape reappeared in the same location with a slider located below it, this time at a different random level of information (sampled from the full range). Participants were instructed to move the slider to adjust the second shape until they thought it was identical to the shape that appeared at first (which was indeed an option available to them), with no time pressure to respond (Fig. 2a).

By moving the slider from one end to the other, participants were iterating the sequence of shapes derived from the parent shape. Though there were 51 discrete steps on the slider corresponding to 51 different shapes, these shapes varied smoothly enough that the adjustment process felt much like a continuous animation of a shape evolving back and forth until a satisfactory frame was found.

Overall, the experiment consisted of 34 trials, including one practice trial at the beginning, 30 test trials corresponding to the 30 unique shape families, and three catch trials appearing in the early, middle, and late stages of the session. (In catch trials, certain locations along the slider completely transformed the shape

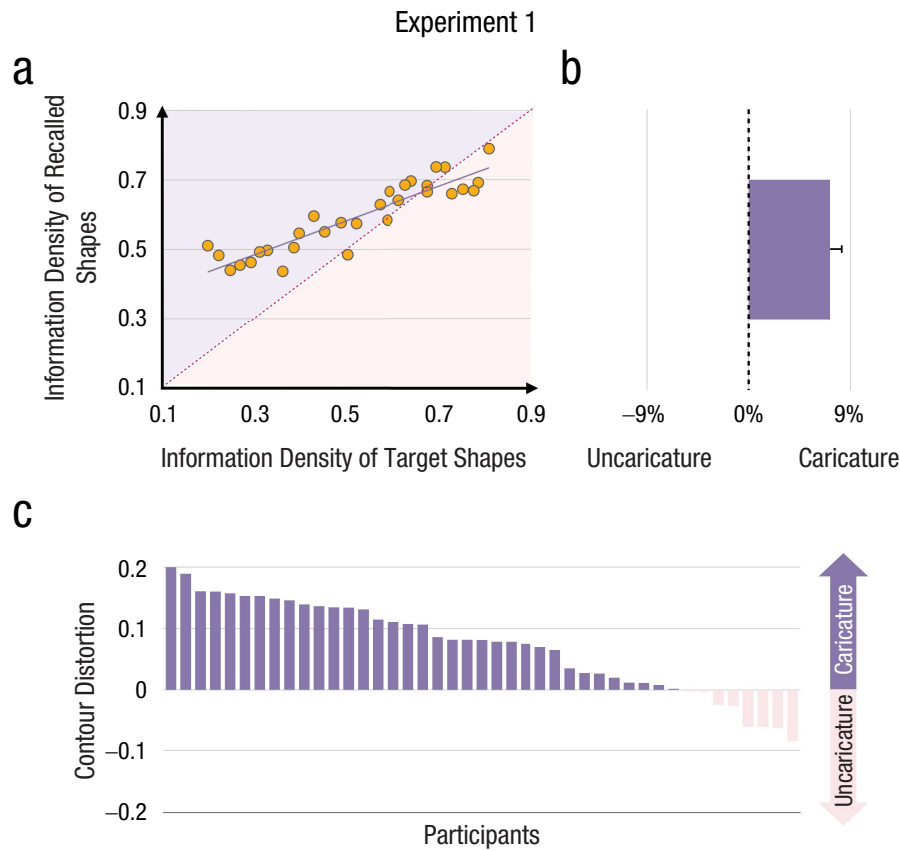


Fig. 3. Results of Experiment 1. In (a) is illustrated the recall bias at different levels of complexity, with presented-shape complexity on the *x*-axis and recalled shape complexity on the *y*-axis. The majority of points fall into the purple space, meaning that the information density of recalled shapes was higher than the information density of the presented shapes. In (b) we illustrate the recall bias collapsed across all complexity levels, showing that recalled shapes were significantly caricatured by 7.19% on the entire scale. Error bars depict $+1$ *SEM* of the difference between the mean of recalled values and presented values. In (c) we show how a strong majority of participants reproduced exaggerated versions of target shapes.

into a member of a different family that did not appear elsewhere in the experiment; any participant who ever answered in this region of slider space was determined not to have been paying attention.) The sign of the slider was randomly assigned across participants so that moving the slider rightward either exaggerated or normalized a shape.

Readers can experience this task for themselves at <https://perceptionresearch.org/caricature/>.

Results

In Experiment 1, 7 participants were excluded for failing to submit a complete data set, leaving 43 participants with analyzable data. For ease of presentation, we normalized the range of contour information from 0 to 1 (0 represents the maximally normalized shape and 1 the maximally exaggerated shape).

As predicted by an encoding-by-caricature approach, participants tended to overestimate the information density of the remembered shapes, increasing their contour information by an average of 13.76% relative to the actually presented shapes' contour information (Fig. 3). This pattern can be seen (and quantified) in several ways. One straightforward way is to consider the average slider position of recalled shapes: Out of 51 total steps, where the average shape was presented at step 26, the bias corresponded to approximately 4 steps. Another way is to consider the average normalized information density of the remembered shapes (0.59) compared to the average normalized information density of the presented shapes (0.52). A paired *t* test confirmed that the information density of reproduced shapes was significantly greater than the true average information density of the target shapes, $t(42) = 6.25$, $p = 1.72 \times 10^{-7}$, $d = 0.95$, $SE = 0.011$, $CI_{bias} = 0.072 [0.049, 0.095]$ (Fig. 3b).

This pattern of results was also fairly consistent across participants, with 81% of participants trending in the expected direction (Fig. 3c). In other words, participants adjusted the test shapes to be exaggerated versions of the original shapes they had just seen seconds earlier. Figure 3a shows these results for each of the information-density levels used as targets in the experiment; most of the recalled values are greater than the true values (see Experiment 3 for evidence that even these very complex shapes are actually misremembered in exaggerated form as well).

In Experiment 2, 4 participants were excluded for failing to provide a complete data set, and 6 participants failed to pass at least one catch trial, leaving 40 participants with analyzable data. Again, participants tended to reproduce caricatured versions of the original shapes. The averaged contour information of recalled shapes was 0.56, significantly higher than the true average information of 0.52, $t(39) = 4.60$, $p = 4.36 \times 10^{-5}$, $d = 0.73$, $SE = 0.0088$, $CI_{\text{bias}} = 0.041 [0.024, 0.058]$. Seventy-eight percent of participants trended in this direction (i.e., recalling shapes as having exaggerated contours). Thus, the memory distortion observed in Experiment 1 generalizes to a wider range of shapes that vary more considerably in information density, with participants again caricaturing shapes in memory.

These results thus provide initial evidence for the hypothesis that the mind encodes and stores even novel contextless shapes as more informationally dense than they really are, caricaturing even some of the most basic stimuli we encounter.

Experiment 3: Forced Choice

Experiments 1 and 2 revealed a bias wherein novel shapes were remembered in exaggerated form. However, these results could be specific to certain methodological choices, especially the response method of continuously adjusting a shape using a slider. For example, because the slider cycles through the full range of available shapes, participants were able to see many candidate answers in ways that may have biased their final choice. This response modality also likely introduced a kind of regression effect to the center of the responding slider: Unless the target shape corresponds to the very middle of the slider, the uncaricatured options and caricatured options are unbalanced in terms of available responses. For example, recalling a shape that appeared at the 80th percentile of contour information gives participants only 20% of the space to caricature it but 80% to uncaricature it, which may bias responses in that direction (and vice versa for simpler shapes).

To address this, Experiment 3 used a forced-choice paradigm, giving only two options for all target shapes

across different levels of complexity. This experiment thus asks whether the caricature effect goes beyond any one response modality and also whether it is more uniform over different complexity ranges than Experiments 1 and 2 may have seemed to suggest (in particular, whether it also arises for shapes on the complex end of the spectrum).

Method

One hundred participants were recruited for this experiment from Prolific. This sample size was double that of Experiments 1 and 2 because the data collected by the two-alternative forced-choice (2AFC) paradigm are sparser than the data collected by the previous response modality.

The critical difference between this experiment and the previous two experiments is that two candidate shapes (instead of an adjustable slider) served as response options during the recall phase, and participants simply selected one of them. Both candidate shapes deviated from the true target by two steps in opposite directions on the shape spectrum used in Experiment 2. In other words, one option was a caricatured version of the target shape, and the other was an uncaricatured version of the same shape (see Fig. 4a). This paradigm precluded participants from viewing the entire shape spectrum during the recall phase and also eliminated or attenuated any independent bias toward the center of the responding range. If memories of a shape are genuinely biased toward its caricature, then participants should tend to choose the caricatured version over the uncaricatured version.

Results

Six participants were excluded for failing at least one catch trial, leaving 94 participants for analyses.

Once again, the results supported a caricature bias in visual memory. When recalling the target shape from two alternatives (a caricatured shape and an uncaricatured shape), participants were significantly more likely to select the caricatured option (62% for the caricatured option; $t(93) = 8.35$, $p = 6.26 \times 10^{-13}$, $d = 0.86$, $SE = 0.015$, $CI_{\text{difference}} = 0.12 [0.092, 0.15]$). Seventy-three out of 94 participants trended in this direction.

Importantly, without a slider to draw responses toward its center, it became especially clear that the caricature bias occurred across the whole range of shape complexities. Figure 4b shows results for each of the information-density levels used as targets in the experiment: Even the most information-dense shapes tended to be caricatured in recall (with no tendency for this effect to drop off with complexity), suggesting that the

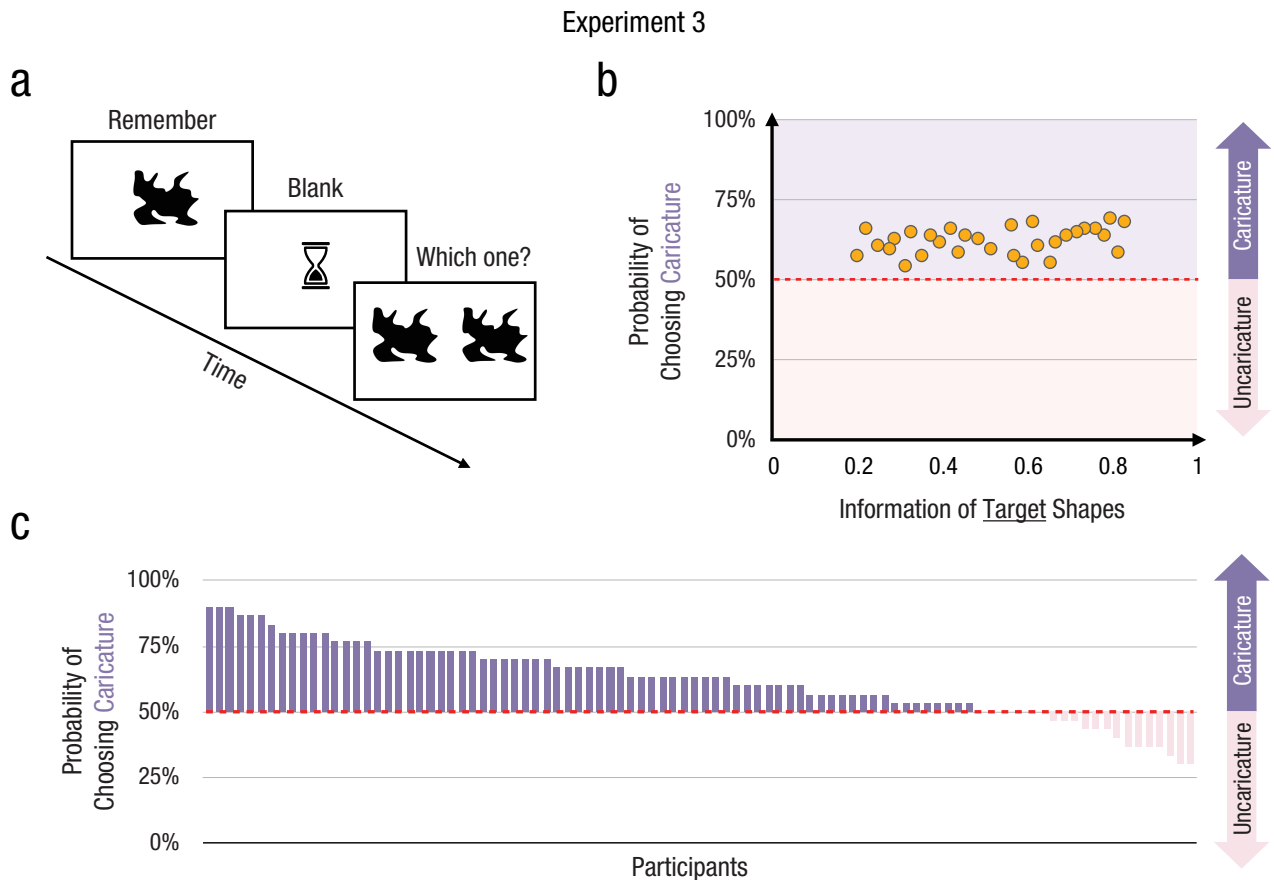


Fig. 4. Results of Experiment 3. Participants were asked to view a shape and then select from two options (one more caricatured and one more uncaricatured) within the same shape family as the target shape (a). Participants were significantly more likely to choose the caricatured shape (b), extending our findings to another response method. Moreover, this recall bias (*y*-axis) emerged across the full range of complexity levels (*x*-axis); all points (averaged responses) fell into the “caricature” space (purple), indicating that shapes at each complexity level were more likely to be recalled as their caricatured versions. As shown in (c), a strong majority of participants more frequently chose exaggerated versions of the target shapes.

caricature effect does not depend on how complex the target is and indeed persists even for very complex shapes. This result thus (a) demonstrates the consistency of the caricature effect across further variations in experimental design; (b) reveals that the response modality (a bounded slider) may have masked the effect's consistency in previous experiments; and (c) helps to rule out the possibility that regression to the center of the slider (and the range of sampled shapes) was responsible for the effects observed earlier.³

Together with results from Experiments 1 and 2, this 2AFC task provides converging evidence for the caricature bias in visual memory.

Experiments 4 and 5: Ruling out Response Bias

We have described the present results as a bias to remember shapes in exaggerated form. However, there

may be alternative explanations. For example, perhaps participants simply enjoyed looking at more complex images, and so adjusted their responses in the direction of exaggeration to match this preference. Or perhaps participants responded strategically: If they expected their memories to lose detail over time, for example, they might compensate for such expected losses by intentionally choosing a more informationally dense shape. Experiments 4 and 5 ruled out these alternatives.

Method

As in Experiment 3, 100 participants were recruited from Prolific for Experiments 4 and 5.

Experiment 4 proceeded in the same way as Experiment 2 except that the 30 shapes were divided into two testing blocks: a *memory* block and a *perception* block. In the memory block, participants were asked to reproduce shapes exactly as in Experiment 2, by adjusting a

shape after the target shape had disappeared. In the perception block, the target shape remained on the screen the entire time, and participants simply adjusted another shape to match the target shape, which was still visible right in front of them (see Fig. 5a). The 30 shapes from Experiment 2 were divided into two groups of 15 shapes (one group for each block), randomly chosen anew for each participant. Block order was counterbalanced across participants. Each block started with a practice trial and included two catch trials. This design aimed to tease apart memory distortions and response biases favoring exaggerated shapes. If the caricature bias found in previous experiments merely arose from response biases, then the bias should also apply to the online viewing cases (as in studies of the “El Greco fallacy”; Firestone, 2013; Firestone & Scholl, 2014; Valenti & Firestone, 2019). However, if the perception block fails to produce a caricature bias similar to the bias found in the memory block, then the bias we observed earlier is unlikely to be explained by general response biases of this sort.

Experiment 5 aimed to rule out strategic responding as an explanation of our findings. In this experiment, we surveyed participants for their intuition about the expected results of our memory task. The procedure began with a brief introduction to the background and the research question from the previous studies, including a short video demonstrating the experimental trials in Experiments 1 and 2. After participants indicated that they understood the experiment, they were presented with two predictions about the outcome of the experiment, with example shapes depicting the possible results; they were then asked to select which prediction they thought was right (Fig. 5b). Two options were given to participants: *more simple* (shown by example shapes being normalized) or *more complex* (shown by example shapes being exaggerated); participants indicated their prediction by clicking on the corresponding button. If participants have a shared assumption about the biases that might arise in recalling these shapes, then such beliefs could interact with the effect found in previous experiments. However, if they do not, then such beliefs are perhaps unlikely to explain our results.

Results and discussion

Both experiments supported a genuine memory distortion, rather than other forms of bias.

In Experiment 4, 5 participants were excluded for not passing at least one catch trial, leaving 95 participants for analysis. As in earlier experiments, participants in the memory condition reproduced shapes in caricatured form, $t(94) = 5.57$, $p = 2.39 \times 10^{-7}$, $d = 0.62$, $SE = 0$, $CI_{\text{bias}} = 0.036$ [0.024, 0.048]. There was also a very

small bias in the caricatured direction for perception trials, where the target shape remained on the screen, $t(94) = 3.15$, $p = .0022$, $d = 0.32$, $SE = 0.0023$, $CI_{\text{bias}} = 0.007$ [0.0027, 0.012]. However, this bias was several times smaller than the memory bias (0.7% vs. 3.6%), and significantly so, $t(94) = 4.44$, $p = 2.45 \times 10^{-5}$, $d = 0.46$, $SE = 0.0065$, $CI_{\text{difference}} = 0.029$ [0.016, 0.041]. If participants simply had a preference to set the slider to the complex end (because, e.g., they find complex shapes visually appealing and preferred to look at them), then the same effects should have arisen in perception trials as in memory trials. This result thus suggests that a response bias favoring exaggerated shapes cannot fully explain the caricature distortions we have observed.

In Experiment 5, we found that there were no consensus assumptions about the expected results of this kind of task, in ways that are encouraging for our memory interpretation. The two options (*more simple* and *more complex*) were chosen at almost identical rates: 51 participants chose “more complex,” and 49 chose “more simple” (binomial probability test, $p = .92$). Evidently, there was no clear preference between these two predictions, suggesting the lack of strong or consistent predictions about this sort of design. This pattern of results is ideal for our interpretation, because a strong bias in either direction might have been problematic for our account. If most participants thought there would be a complexity bias, then perhaps the participants in our task were acting to fulfill that bias. If most participants thought there would be a simplicity bias, then perhaps the participants in our task were acting to overcome that bias. But a near 50/50 split suggests that there was no consistent expectation among participants that it could produce our results. Recall, for example, that our earlier biases appeared in well over half of participants (81% in Experiment 1 and 78% in Experiment 2). Given the results of the present survey experiment, the participants who drove our earlier results must have carried both simplicity and complexity expectations in ways that seemingly rule out such expectations as the (sole) cause of our effects.

Experiment 6: Accumulating Biases in Serial Reproduction

Although the effects in our earlier experiments were reliable and robust, their magnitude was often subtle, resulting in a caricature bias of 4% to 6% when considered in terms of the full range of information densities on the scale. One consequence of these effect sizes is that any one instance of this bias is likely to be small in practice. However, in the real world, we often recall the same objects and events multiple times, and even relay them to others who in turn recall what we have

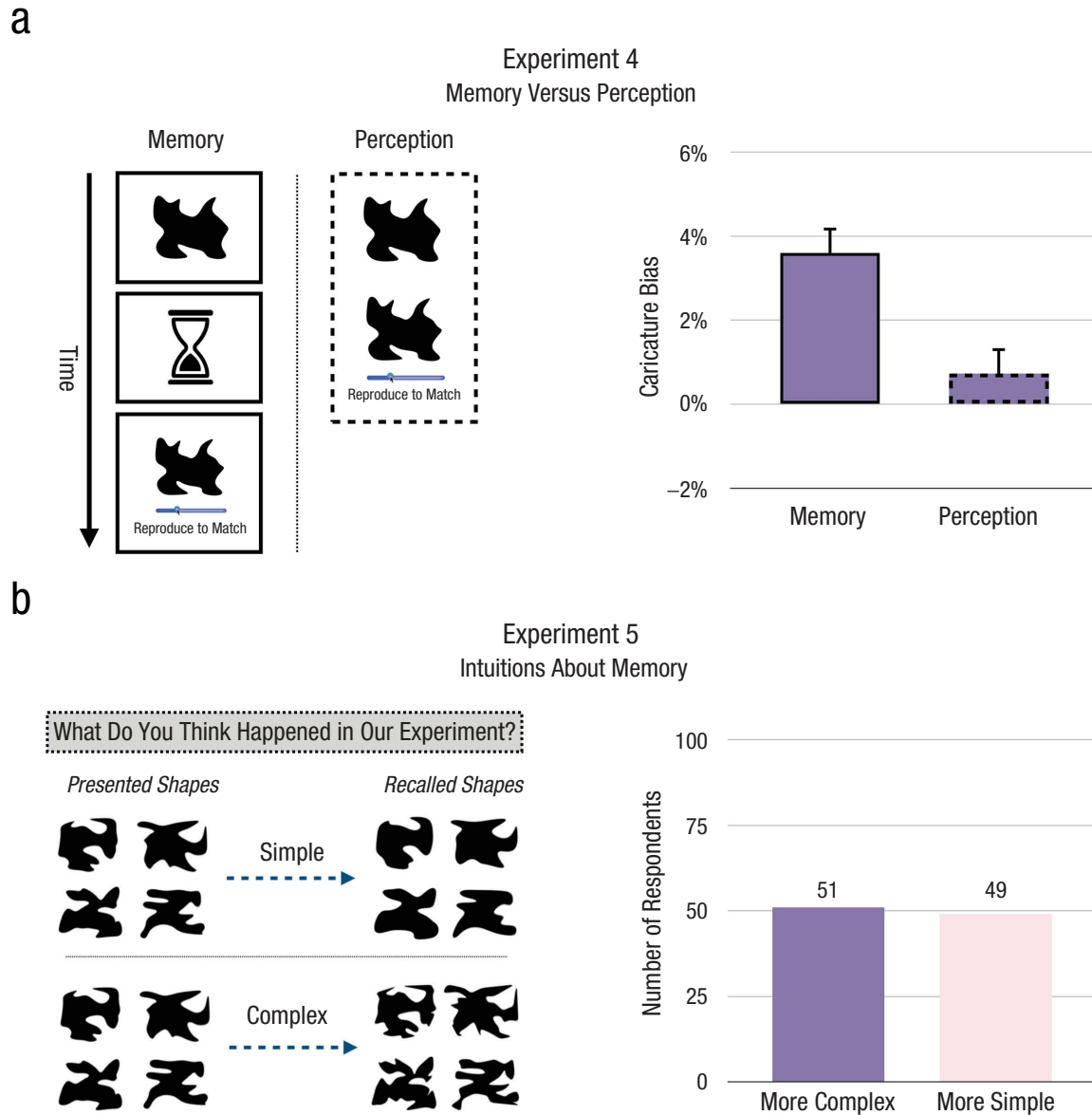


Fig. 5. Results of Experiments 4 and 5. Experiment 4 (a) replicated our memory effect, but also included a perception condition in which the adjustable shape and the target shape remained on the screen at the same time. Under these conditions, there was only a very minimal caricature bias (0.7%), and it was many times smaller than the memory bias (3.6%), consistent with a genuine memory distortion in our previous studies. Experiment 5 asked participants to predict the results of our prior experiments. Participants were shown a toy version of the task (b) and were asked what kind of bias they thought would occur. They were almost perfectly split in their predictions, suggesting no strong or pervasive predicted bias toward either exaggerations or simplification. Error bars represent 1 SEM of the difference between conditions.

told them. Might iterated recall of this sort amplify the biases observed here?

Experiment 6 explored this question using *serial reproduction*, in which one participant's responses serve as the stimuli for another participant (Bartlett, 1932). Analogous to the game of “telephone,” a message or stimulus is shown to an initial observer, who then transmits it to another observer—but not before modifying the stimulus (often unintentionally) according to their own priors or biases. By repeating this procedure

several times, the transmission chain accumulates the collective bias of the group and converges on a well or attractor that serves as a kind of equilibrium point.

This method has been used to study a variety of biases (e.g., Kalish et al., 2007; Langlois et al., 2021; Uddenberg & Scholl, 2018), though often for more complex tasks or with high-level stimuli. Here, we applied this method to our very simple shape-memory task. This allowed us to explore (a) how such distortions accumulate over time and (b) whether there exists a

boundary for caricature distortion over which the shapes may not be continuously exaggerated.

Method

Participants. We aimed to have 30 transmission chains for each of 30 shapes, with each chain comprising 10 steps (i.e., nine observer-to-observer transmissions). We also allowed each participant to make 30 judgments (one for each of the 30 shapes). In total, this led to a target sample of 300 participants with admissible data, which required 316 participants total (see below for more information). All participants were recruited from Prolific.

Stimuli and procedure. From the participant's point of view, the experiment was quite similar to Experiments 1 and 2: A shape was presented on a given trial, and the participants reconstructed it after a short delay. However, in almost all cases, the stimuli they were viewing had come from a previous participant who had completed the same procedure; an exception occurred when a participant was at the beginning of a transmission chain.

From the experimenter's point of view, there were 10 rounds of shape judgments for each transmission chain, 30 transmission chains per shape, and 30 shapes total (i.e., 900 chains of 10 steps in all). The 30 shapes to be viewed in the first round initially appeared at one of three levels of contour information: 10 shapes at a relatively low level of contour information (approximately one quarter of the way up the scale), 10 at a moderate level (middle point of the scale), and the other 10 at a relatively high level (roughly three quarters of the way up the scale). Every participant occupied the same position in every chain they participated in: For example, if a participant was in the fourth position of a given chain, he or she was in the fourth position of each of the 30 shape chains that were part of the session.

Results

The results again revealed a caricature bias, but this time with a much greater final magnitude (Fig. 6a). A sample chain from this experiment appears in Figure 6b; the shape starts off at a fairly low level of complexity, and by the end it has grown several new jagged appendages. This pattern pervaded the stimulus set, resulting in a subjectively appreciable degree of caricaturing. Compared to the original shapes, the shapes reproduced in the 10th round had much greater contour information (0.72 vs. 0.52), $t(29) = 14.97$, $p = 3.54 \times 10^{-15}$, $d = 2.73$, $SE = 0.013$, $CI_{\text{bias}} = 0.19$ [1.64, 0.21], so that the shapes were eventually biased approximately 20% away from their true averaged values in terms of the entire scale and nearly 40% away when considered in terms

of their initial values. Average contour information increased in each round without tending to pause or reverse.

Figure 6a shows the temporal evolution of the shapes, where each line represents the average value across all reproduction chains at the three given starting levels of shape information. As shown, these chains gradually converge into the exaggeration region of space as the experiment advances. Indeed, even for chains starting at a high level of contour information (indicated by the green line), the effect of regression to the center of the slider did not drive the reproduction chain back to a more moderate level. This pattern of convergence thus demonstrated a robust bias to remember shapes as being increasingly information dense or more exaggerated than they really were, and in ways that can be easily visualized and appreciated.

Even though the caricature effects observed here may cause only a subtle distortion for any one shape on any one trial (which may be hard to recognize at the level of each caricaturing step), the present results suggest that the accumulation of such biases over time can be significant indeed. Considering that daily life often involves repeatedly remembering and recalling various objects, the biases we observed here may well accumulate in real-world contexts as well.

General Discussion

Memory rarely replicates what we see; it reconstructs past experiences with biases and distortions. What kinds of biases arise for the most basic stimuli we encounter? Here, we explored how memory engages in a process of caricature, even for simple geometric shapes, and even without explicit demand to remember or distinguish multiple objects at the same time, preexisting schematic associations or knowledge, or other contextual factors that might tend to encourage such biases.

Adding Information to Memories

It was not a foregone conclusion that the present experiments would turn out this way (i.e., with a bias toward exaggeration, or indeed with any directional bias at all). If anything, memories tend to lose detail over time—and so a natural prediction might have been that shapes like those used here would be remembered as simpler or less detailed versions of themselves (Cooper et al., 2019). At the same time, memories often add information that wasn't present in the encoded stimulus or event, and so in certain situations they can end up being richer or more detailed than what was actually shown to participants. For example, memory may generate events that never occurred (Kominsky et al., 2021; Loftus & Palmer, 1974), add objects that were not originally in

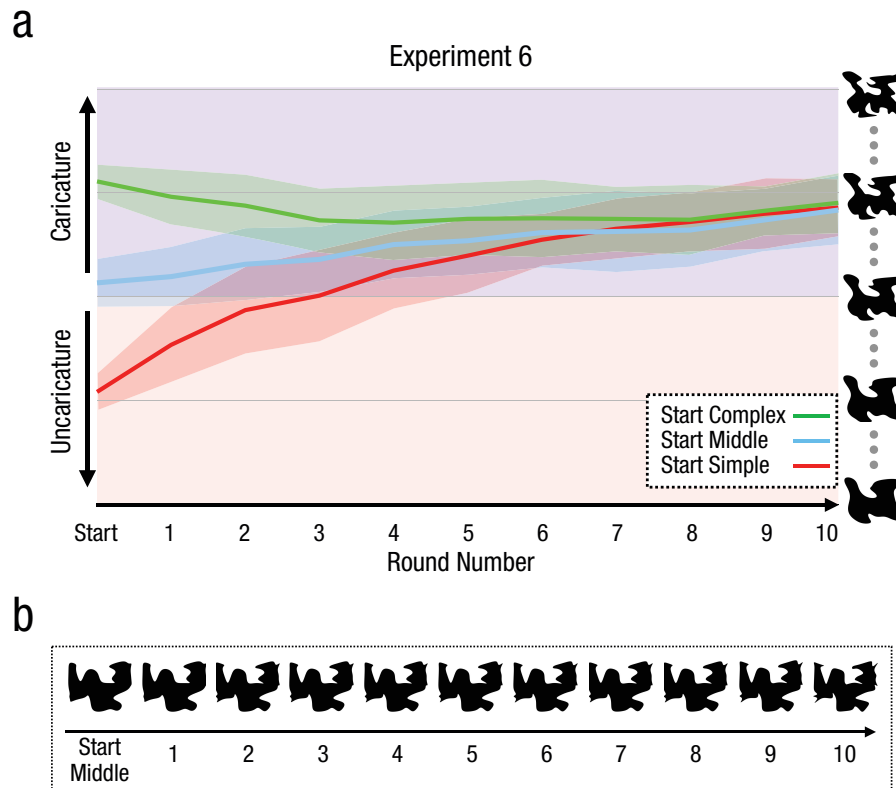


Fig. 6. Performance on the serial reproduction task of Experiment 6, in which the recalled shapes of one participant were used as stimuli for the next participant. The three colored lines in (a) represent the average complexity values of the first step in the transmission chains, from highly normalized (red), medially normalized (blue), and highly exaggerated (green). The three lines gradually converge as the transmission chains proceed, ending with a bias toward the information-dense side of the scale. The collective error is about 20% of the entire range of information densities used in the experiment. The shaded band around each line indicates a 95% confidence interval. In (b) we show the evolution of one of the shapes in the experiment as it passed from participant to participant.

visual scenes (Bainbridge & Baker, 2020; Intraub & Richardson, 1989), complete ambiguous visual patterns (Bartlett, 1932; Brewer & Treyns, 1981), run forward events seen previously (Freyd & Finke, 1984; Hafri et al., 2022), and even enhance the quality of images (Rivera-Aparicio et al., 2021).

However, unlike memory distortions that recruit high-level knowledge or schemas (Bae et al., 2015; Hemmer & Steyvers, 2009), the caricature bias we explore here adds information to single, novel objects that have never been seen before, making it less likely that this bias relies on long-term memory or preexisting associations.

Caricatures and Memory Repulsion

Why, then, does this happen? One reason this bias may occur is to enhance distinctiveness and aid later recognition, as in studies of face caricatures (e.g., Lee et al., 2000). Recent work at the intersection of machine

learning and visual perception has also shown that a kind of caricatured encoding of videos can improve human performance in detecting visual misinformation (Fosco et al., 2022).

These results may thus be related to the hypothesis that remembered objects are often repulsed from one another, so that similar items (such as colored circles) are encoded in ways that amplify the differences between them (e.g., Bae & Luck, 2017; Chunharas et al., 2022; Scotti et al., 2021). The current work could be seen as either (a) extending this phenomenon even further, by showing that even memorizing a single item can engage this process of repulsion from its possible “peers” (even when those other objects need not be remembered or recalled on a given trial), or (b) revealing a more general mechanism whereby salient object information is automatically remembered in exaggerated form. Though the present work does not fully unravel the relationship between caricature biases and repulsion effects, future work may investigate

whether caricaturing makes shapes more distinctive than their counterparts and whether the present caricature bias assists later processes, such as recognition and categorization.

Caricatures Beyond Faces

Although most of the scientific (and popular) attention paid to caricatures has focused on faces and bodies, caricaturing may also arise when visually representing information in different domains. For example, adults and children who are given the task of conveying concepts through drawings tend to exaggerate those aspects of the object that best distinguish it from neighboring concepts, and they may distort the same object in different ways depending on the level of abstraction being considered (Fan et al., 2020). The present work adds to this literature as well. It may not be so surprising for the mind to amplify the distance between a new object and a norm composed of similar objects; indeed, the primary approach used in early studies of caricatures relies on norms of this sort, such as a face norm averaged from many faces, or a prototype of a species of animal (e.g., Corneille et al., 2004). What is distinctive about the present results, however, is the absence of an obvious norm or prototype for novel shapes, at least of the sort studied here (though see Sablé-Meyer et al., 2022, for a discussion of the notion of geometric primitives). Thus, the effects reported here may reflect a more general process: Even when there is no explicit pressure to do so, visual memory emphasizes and exaggerates what is salient about the world around us.

Transparency

Action Editor: Rachael Jack

Editor: Patricia J. Bauer

Author Contributions

Zekun Sun: Conceptualization; Data curation; Formal analysis; Investigation; Methodology; Resources; Software; Visualization; Writing – original draft; Writing – review & editing.

Subin Han: Conceptualization; Data curation; Methodology; Validation; Writing – review & editing.

Chaz Firestone: Conceptualization; Funding acquisition; Project administration; Supervision; Writing – original draft; Writing – review & editing.

Declaration of Conflicting Interests

The author(s) declared that there were no conflicts of interest with respect to the authorship or the publication of this article.

Open Practices

Demos of these experiments are available at <https://perceptionresearch.org/caricature/>, so readers can experience these tasks as the participants did. The data, experiment code, stimuli, and other relevant materials for all studies are available at <https://osf.io/7grk8/>. These studies

were not preregistered. This article has received the badges for Open Data and Open Materials. More information about the Open Practices badges can be found at <http://www.psychologicalscience.org/publications/badges>.



ORCID iD

Chaz Firestone  <https://orcid.org/0000-0002-1247-2422>

Notes

1. In accordance with Feldman and Singh (2005), this measure treats convex and concave regions differently in their information content: The convex direction (an outward turning) is assumed to be slightly more likely than the concave direction (an inward turning). In other words, points of convex curvature are more surprising (and thus convey more information) than otherwise equivalent points of concave curvature. This parameter of asymmetry reflects psychological findings of the special status of negative curvature in perception (e.g., Barenholtz et al., 2003).

2. Indeed, to further explore the role of baseline complexity in mental caricaturing, we conducted four additional variations of these experiments, testing this effect over other complexity ranges: (a) 0.25 to 0.75 of the scale used in Experiment 1 (i.e., a symmetric contraction of the range); (b) 0.5 to 1 of the scale used in Experiment 1 (i.e., a sample biased toward the more complex end of the spectrum); (c) 0 to 0.5 of the scale used in Experiment 2 (i.e., a sample biased toward the simpler end of the spectrum); and (d) 0.25 to 0.75 of the scale used in Experiment 2 (equivalent to -0.5 to $+0.5$ of the range of Experiment 1). All four experiments showed a significant caricature effect: (a) $t(36) = 5.10$, $p = 1.10 \times 10^{-5}$; (b) $t(39) = 8.83$, $p = 7.62 \times 10^{-11}$; (c) $t(48) = 4.63$, $p = 2.77 \times 10^{-5}$; and (d) $t(39) = 5.01$, $p = 1.81 \times 10^{-5}$. This suggests a highly reliable and consistent caricature bias no matter how the shapes are sampled. We thank a reviewer for comments that inspired this approach.

3. Thoughtful readers may be curious whether this bias holds true for even extremely complex shapes. Though generating more and more complex shapes tends to produce odd stimuli that often give the appearance of having some other recognizable identity, thoughtful comments by a reviewer led us to run another experiment following the design for Experiment 3, in which the contour information spectrum of the experiment was expanded by 50% at the complex end of the spectrum. This experiment still showed a significant caricature effect, both for the full 0 to 1.5 range (57% choosing more complex; $t(89) = 4.83$, $p = 5.59 \times 10^{-6}$) and also just for the new shapes tested (57% choosing more complex; $t(89) = 3.45$, $p < .001$), suggesting that the caricaturing effect observed here really does apply across a wide range of complexity values.

References

- Amalric, M., Wang, L., Pica, P., Figueira, S., Sigman, M., & Dehaene, S. (2017). The language of geometry: Fast comprehension of geometrical primitives and rules in human adults and preschoolers. *PLOS Computational Biology*,

- 13(1), Article e1005273. <https://doi.org/10.1371/journal.pcbi.1005273>
- Attneave, F. (1954). Some informational aspects of visual perception. *Psychological Review*, 61(3), 183–193.
- Attneave, F., & Arnoult, M. D. (1956). The quantitative study of shape and pattern perception. *Psychological Bulletin*, 53(6), 452–471.
- Bae, G.-Y., & Luck, S. J. (2017). Interactions between visual working memory representations. *Attention, Perception, & Psychophysics*, 79, 2376–2395.
- Bae, G.-Y., Olkkonen, M., Allred, S. R., & Flombaum, J. I. (2015). Why some colors appear more memorable than others: A model combining categories and particulars in color working memory. *Journal of Experimental Psychology: General*, 144(4), 744–763.
- Bainbridge, W. A., & Baker, C. I. (2020). Boundaries extend and contract in scene memory depending on image properties. *Current Biology*, 30(3), 537–543.
- Baker, N., Garrigan, P., & Kellman, P. J. (2021). Constant curvature segments as building blocks of 2D shape representation. *Journal of Experimental Psychology: General*, 150(8), 1556–1580.
- Barenholtz, E., Cohen, E. H., Feldman, J., & Singh, M. (2003). Detection of change in shape: An advantage for concavities. *Cognition*, 89(1), 1–9.
- Bartlett, F. C. (1932). *Remembering: A study in experimental and social psychology*. Cambridge University Press.
- Benson, P. J., & Perrett, D. I. (1991). Perception and recognition of photographic quality facial caricatures: Implications for the recognition of natural images. *European Journal of Cognitive Psychology*, 3(1), 105–135.
- Brennan, S. E. (1982). *Caricature generator* [Doctoral dissertation]. Massachusetts Institute of Technology.
- Brewer, W. F., & Treyens, J. C. (1981). Role of schemata in memory for places. *Cognitive Psychology*, 13(2), 207–230.
- Chanales, A. J., Tremblay-McGaw, A. G., Drascher, M. L., & Kuhl, B. A. (2021). Adaptive repulsion of long-term memory representations is triggered by event similarity. *Psychological Science*, 32(5), 705–720.
- Chang, P. P., Levine, S. C., & Benson, P. J. (2002). Children's recognition of caricatures. *Developmental Psychology*, 38(6), 1038–1051.
- Chunharas, C., Rademaker, R. L., Brady, T. F., & Serences, J. T. (2022). An adaptive perspective on visual working memory distortions. *Journal of Experimental Psychology: General*, 151(10), 2300–2323. <https://doi.org/10.1037/xge0001191>
- Cooper, R. A., Kensinger, E. A., & Ritchey, M. (2019). Memories fade: The relationship between memory vividness and remembered visual salience. *Psychological Science*, 30(5), 657–668.
- Corneille, O., Huart, J., Becquart, E., & Brédart, S. (2004). When memory shifts toward more typical category exemplars: Accentuation effects in the recollection of ethnically ambiguous faces. *Journal of Personality and Social Psychology*, 86(2), 236–250.
- Davis, T., & Love, B. C. (2010). Memory for category information is idealized through contrast with competing options. *Psychological Science*, 21(2), 234–242.
- Donderi, D. C. (2006). Visual complexity: A review. *Psychological Bulletin*, 132(1), 73–97.
- Drascher, M. L., & Kuhl, B. A. (2022). Long-term memory interference is resolved via repulsion and precision along diagnostic memory dimensions. *Psychonomic Bulletin & Review*, 29(5), 1898–1912. <https://doi.org/10.3758/s13423-022-02082-4>.
- Fan, J. E., Hawkins, R. D., Wu, M., & Goodman, N. D. (2020). Pragmatic inference and visual abstraction enable contextual flexibility during visual communication. *Computational Brain & Behavior*, 3(1), 86–101.
- Feldman, J., & Singh, M. (2005). Information along contours and object boundaries. *Psychological Review*, 112(1), 243–252.
- Feldman, J., & Singh, M. (2006). Bayesian estimation of the shape skeleton. *Proceedings of the National Academy of Sciences, USA*, 103(47), 18014–18019.
- Firestone, C. (2013). On the origin and status of the “El Greco fallacy.” *Perception*, 42(6), 672–674.
- Firestone, C., & Scholl, B. J. (2014). “Top-down” effects where none should be found: The El Greco fallacy in perception research. *Psychological Science*, 25(1), 38–46.
- Fosco, C., Josephs, E., Andonian, A., Lee, A., Wang, X., & Oliva, A. (2022). *Deepfake caricatures: Amplifying attention to artifacts increases deepfake detection by humans and machines* [arXiv preprint]. arXiv:2206.00535.
- Freyd, J. J., & Finke, R. A. (1984). Representational momentum. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 10(1), 126–132.
- Garner, W. R. (1962). *Uncertainty and structure as psychological concepts*. Wiley.
- Gauthier, I., & Tarr, M. J. (1997). Becoming a “greeble” expert: Exploring mechanisms for face recognition. *Vision Research*, 37(12), 1673–1682.
- Gibson, J. J. (1929). The reproduction of visually perceived forms. *Journal of Experimental Psychology*, 12(1), 1–39.
- Hafri, A., Boger, T., & Firestone, C. (2022). Melting ice with your mind: Representational momentum for physical states. *Psychological Science*, 33(5), 725–735.
- Hemmer, P., & Steyvers, M. (2009). A Bayesian account of reconstructive memory. *Topics in Cognitive Science*, 1(1), 189–202.
- Intraub, H., & Richardson, M. (1989). Wide-angle memories of close-up scenes. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 15(2), 179–187.
- Kalish, M. L., Griffiths, T. L., & Lewandowsky, S. (2007). Iterated learning: Inter-generational knowledge transmission reveals inductive biases. *Psychonomic Bulletin & Review*, 14(2), 288–294.
- Kominsky, J. F., Baker, L., Keil, F. C., & Strickland, B. (2021). Causality and continuity close the gaps in event representations. *Memory & Cognition*, 49(3), 518–531.
- Langlois, T. A., Jacoby, N., Suchow, J. W., & Griffiths, T. L. (2021). Serial reproduction reveals the geometry of visuospatial representations. *Proceedings of the National Academy of Sciences, USA*, 118(13), Article e2012938118.
- Lee, K., Byatt, G., & Rhodes, G. (2000). Caricature effects, distinctiveness, and identification: Testing the face-space framework. *Psychological Science*, 11(5), 379–385.

- Loftus, E. F., & Palmer, J. C. (1974). Reconstruction of automobile destruction: An example of the interaction between language and memory. *Journal of Verbal Learning and Verbal Behavior*, 13(5), 585–589.
- Mauro, R., & Kubovy, M. (1992). Caricature and face recognition. *Memory & Cognition*, 20(4), 433–440.
- Norman, J. F., Phillips, F., & Ross, H. E. (2001). Information concentration along the boundary contours of naturally shaped solid objects. *Perception*, 30(11), 1285–1294.
- O'Toole, A. J., Vetter, T., Volz, H., & Salter, E. M. (1997). Three-dimensional caricatures of human heads: Distinctiveness and the perception of facial age. *Perception*, 26(6), 719–732.
- Perkins, D. (1975). A definition of caricature and caricature and recognition. *Studies in Visual Communication*, 2(1), 1–24.
- Rhodes, G., Brennan, S., & Carey, S. (1987). Identification and ratings of caricatures: Implications for mental representations of faces. *Cognitive Psychology*, 19(4), 473–497.
- Rivera-Aparicio, J., Yu, Q., & Firestone, C. (2021). Hi-def memories of lo-def scenes. *Psychonomic Bulletin & Review*, 28, 928–936. <https://doi.org/10.3758/s13423-020-01829-1>
- Sablé-Meyer, M., Ellis, K., Tenenbaum, J., & Dehaene, S. (2022). A language of thought for the mental representation of geometric shapes. *Cognitive Psychology*, 139, Article 101527.
- Scotti, P. S., Hong, Y., Leber, A. B., & Golomb, J. D. (2021). Visual working memory items drift apart due to active, not passive, maintenance. *Journal of Experimental Psychology: General*, 150(12), 2506–2524. <https://doi.org/10.1037/xge0000890>
- Siddiqi, K., Shokoufandeh, A., Dickinson, S. J., & Zucker, S. W. (1999). Shock graphs and shape matching. *International Journal of Computer Vision*, 35(1), 13–32.
- Uddenberg, S., & Scholl, B. J. (2018). Teleface: Serial reproduction of faces reveals a whiteward bias in race memory. *Journal of Experimental Psychology: General*, 147(10), 1466–1487. <https://doi.org/10.1037/xge0000446>.
- Valenti, J., & Firestone, C. (2019). Finding the “odd one out”: Memory color effects and the logic of appearance. *Cognition*, 191, Article 103934.
- Yassa, M. A., & Stark, C. E. (2011). Pattern separation in the hippocampus. *Trends in Neurosciences*, 34(10), 515–525.
- Zhao, Y., Chanals, A. J., & Kuhl, B. A. (2021). Adaptive memory distortions are predicted by feature representations in parietal cortex. *Journal of Neuroscience*, 41(13), 3014–3024.