

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Local Sampling with Momentum Accounts for Human Random Sequence Generation

Permalink

<https://escholarship.org/uc/item/3gz154kg>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 43(43)

Authors

Castillo, Lucas

Leon Villagra, Pablo

Chater, Nicholas

et al.

Publication Date

2021

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Local Sampling with Momentum Accounts for Human Random Sequence Generation

Lucas Castillo¹

Pablo León-Villagrà¹

Nick Chater²

Adam N. Sanborn¹

¹University of Warwick, United Kingdom

²Warwick Business School, United Kingdom

Abstract

Many models of cognition assume that people can generate independent samples, yet people fail to do so in random generation tasks. One prominent explanation for this behavior is that people use learned schemas. Instead, we propose that deviations from randomness arise from people sampling locally rather than independently. To test these explanations, we teach people one- and two-dimensional arrangements of syllables and ask them to generate random sequences from them. Although our results reproduce characteristic features of human random generation, such as a preference for adjacent items and an avoidance of repetitions, we also find an effect of dimensionality on the patterns people produce. Furthermore, model comparisons revealed that local sampling accounted better for participants' sequences than a schema account. Finally, evaluating the importance of each models' constituents, we show that the local sampling model proposed new states based on its current trajectory, rather than an inhibition-of-return-like principle.

Keywords: sampling for inference; random generation; representation learning

Introduction

People's behavior is often inconsistent, and at times irrational. Consider *probability matching*, the phenomenon by which if someone believes a particular option to be successful 90% of the time, they pick that option 90% of the time, despite rationality dictating they choose it always (Vulkan, 2000). This behavior has been explained by sampling: that people mentally sample one option according to the probability that it will be successful (Vul, Goodman, Griffiths, & Tenenbaum, 2014). This assumption that responses are independently sampled is pervasive across models of cognition (e.g. Nosofsky, 1984) and independent and identically distributed (*iid*) samples are accumulated in sequential sampling models of decision-making (e.g. Ratcliff, 1978).

While many models assume that what people see and how they respond is driven by samples that are independent of one another, studies of random generation repeatedly show that people produce highly predictable sequences. When asked to generate random numbers, for example, people produce adjacent numbers disproportionately often, change direction (from ascending to descending, or vice versa) too rarely, and seldom repeat the number they have just produced (Wagenaar, 1972). These effects do not seem to be the result of thinking too hard about generating random numbers: the predictability of random sequences is exacerbated by asking participants to perform concurrent tasks or increasing the rates at which sequences have to be generated (Cooper, 2016;

Towse, 1998). These two perspectives stand in apparent contradiction. Why can people produce *iid* samples when categorizing or making decisions, but struggle when explicitly asked to do so?

Schemas Versus Local Sampling

One prominent explanation of people's systematic deviations from randomness is that they engage in a different process altogether: Instead of sampling, they generate sequences based on learned schemas, such as counting up or down. According to this account, these sequences are monitored by the central executive, which switches schemas when randomness is perceived to decline (Baddeley, Emslie, Kolodny, & Duncan, 1998; Cooper, 2016).

Here we propose an alternative explanation to account for these stereotypical patterns: when people generate random sequences, they are tapping into a general cognitive ability to produce samples for inference – a sophisticated mental algorithm that does not sample independently, but locally (Sanborn & Chater, 2016; Chater et al., 2020). Local sampling has been used to explain other deviations from normativity, such as the anchoring effect (Lieder, Griffiths, Huys, & Goodman, 2018) or the 'unpacking' effect (Dasgupta, Schulz, & Gershman, 2017); and other human deviations from *iid* sampling, such as the long-ranging autocorrelations in the temporal structure of many cognitive activities (Zhu, Sanborn, & Chater, 2018).

A popular family of local sampling algorithms is Markov Chain Monte Carlo (MCMC). MCMC algorithms produce a series of states by randomly proposing changes to the current state, and then transitioning depending on the relative probability of the current state versus the proposed new state. Random-Walk Metropolis-Hastings (RW-MCMC), an algorithm from this family, already presents some of the features found in human random sequences: When sampling from a uniform distribution, such as a limited range of numbers, it can favor close items, since proposals are commonly small, local, perturbations of the most recent sample. Furthermore, most RW-MCMC algorithms will rarely repeat the current state when sampling from a uniform distribution, since the proposed state, by definition, will be as likely as the last-visited state. A preference to maintain direction between samples could also be incorporated in a sampling model, for example, in a Langevin-MCMC sampler where momentum is maintained between samples (Horowitz, 1991).

Contrasting Schema and Sampling Accounts

In the standard random number generation task, disentangling schema accounts from local sampling is difficult for two reasons: First, participants' associations between numbers exist prior to the task and are unknown to the experimenter. Second, the items that participants have to randomize already encode a one-dimensional spatial arrangement (the number line) and learned sequences (counting up or down).

In the current study, we introduce a novel experimental design that allows us to adjudicate between the predictions of these two accounts. We ask participants to generate random sequences after learning a display of syllables, thus weakening prior preferences for counting up or down. We experimentally manipulate the dimension of the display, with participants learning either one-dimensional or two-dimensional arrangements. Finally, we control the domains' local structure by randomizing the displays' adjacencies.

This design allows us to investigate the effect of the dimensionality of the domain on deviations from randomness often found in human random generation. We examine three common features of human random generation: a tendency not to repeat items, a tendency to transition to proximate locations in the representation, and a tendency to make fewer changes of direction. Given that random generation rests on a general cognitive ability, we expect participants to deviate from randomness in similar ways regardless of the dimension of the domain.

We are also interested in evaluating the ability of schema and sampling models to account for participants' sequences directly. Manipulating the global and local structure of the domains allows us to dissociate the predictions of the two model classes: If the stereotypical patterns of human random generation are due to learned schemas, then the display arrangement will not affect the produced sequences. Instead, participants should be biased by syllable frequencies in English (i.e. their associations prior to the task), possibly weighted by the frequency of syllables in the experimental training block.

Our manipulation can also inform details of the sampling model: A preference for maintaining direction may result from either a proposal mechanism whereby consecutive samples share a common direction (we call this momentum) or from one where the last-visited sample is avoided (inhibition of return, IOR; see Johnson et al., 2013). While in the one-dimensional case both these proposal mechanisms suggest the same result, the two-dimensional structure allows us to distinguish between the two: while IOR disfavors visiting the previous location, momentum would avoid locations near the previous one as well. As an illustration, consider a traveller who first visits Paris and then London. If they follow IOR they might go to Berlin next, since they only avoid Paris. Instead, if they follow their current direction, they will be more likely to choose Dublin as their next destination.

We first evaluate the effect of our manipulation on indicators of the randomness of a sequence that are common in

the psychological literature. We then contrast schema and sampling models and explore the models' ability to capture qualitative features of human random samples.

Experiment

In our pre-registered experiment¹, we first asked participants to learn a one- or two-dimensional display of syllables. Then participants had to produce random sequences of those syllables in two blocks.

Participants

42 participants took part in the experiment ($M_{\text{age}} = 24.95$, $SD = 9.67$; 27 Female, 14 Male, 1 Non-Binary), and two were excluded following pre-registered criteria (see Procedure and Design). Participants were recruited from the university participant pool and were required to have English as their first language. They received a flat fee of £3.5, plus a bonus of up to £1.8 depending on their performance in the learning stage. The average payment was £4.64, and the experiment took about 30 minutes to complete.

Materials

We generated seven 2-letter syllables ending in *a* to avoid varying ease in transitions due to rhyming. In choosing syllables, we considered both the frequencies of syllables and the frequencies of syllable pairs in the Brown corpus, aiming for a homogeneous set. The selected syllables were: *ca*, *ha*, *la*, *ma*, *na*, *pa*, *ta*.

Procedure and Design

The experiment was conducted online using Microsoft Teams, which was also used to record participants. In the experiment, participants first learned either a one- or two-dimensional display of syllables (Learning stage), then verbally produced these randomly at a constant pace (Random Generation stage). Participants learned only one arrangement of syllables, but produced random sequences in two blocks. Participants were randomly allocated to one of five possible random arrangements, which determined the adjacency of syllables in the display.

Learning stage Participants were presented with a display composed of seven hexagons arranged in either a two-dimensional grid or a single row, depending on the experimental condition (see Figure 1 for an example). Each hexagon contained a syllable, and participants had to learn the location of the syllables. Because only one syllable was visible at a time, participants had to select which syllable they wanted to see next, and they could only choose among hexagons that were adjacent to the currently visible hexagon. There was a one-second delay between the presentation of syllables during which participants were instructed to anticipate the next syllable.

Once participants expressed that they had learned the display, they were tested on all seven cells in random order. To

¹<https://osf.io/q3yrj>

continue to the main experiment, they had to answer two consecutive tests without errors. In case of error, participants continued to learn the display and could get re-tested. If participants failed the test four times or exceeded the maximum learning time of 10 minutes, the experiment was terminated, and the participant was excluded. Participants spent an average of 5.71 minutes ($SD = 2.38$) learning the syllables, with no participant exceeding the ten-minute limit. Two participants were excluded for failing the test four times (average failed attempts = 0.65, $SD = 0.86$; excluded participants' average duration = 9m 54s).

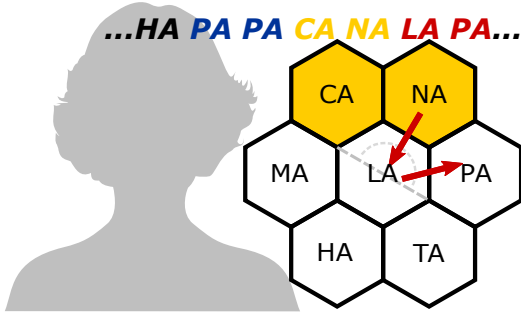


Figure 1: Syllable repetition (blue), adjacency (yellow) and turning point (red) for the example sequence *ha pa pa ca na la pa*. Turning points occur whenever a change in direction between two successive moves exceeds 90° .

Random generation stage In the random generation stage, participants had to produce the learned syllables as unpredictably as possible for five minutes ($M = 303$ s, $SD = 4.45$). They were instructed to do so as if “drawing a syllable out from a hat, saying it out loud, putting it back, shuffling, then repeating the process”, following the instructions in Baddeley (1966). Participants did not see the display of syllables during the random generation stage and instead saw a flashing dot on their screen, appearing at a pace of 80 times per minute. Participants were instructed to produce syllables every time the dot appeared, but produced a slightly lower rate than targeted (400), with an average of 356.04 syllables ($SD = 71.71$) after five minutes. The median gap between each syllable a participant uttered was 843ms, and the median standard deviation in these utterances was 190ms. Despite not seeing the display, all participants named each syllable at least once in each sequence.

At the end of this stage, there was a small break to allow participants to rest. Then, both learning and generation stages were repeated: participants saw the same display of syllables and were tested, and once two blocks were recalled correctly, they produced random syllables for another five minutes ($M = 303$ s, $SD = 3.06$). Participants spent an average of one minute and 44 seconds re-learning and re-testing the display ($SD = 47$ s), and the average number of failed attempts in this second test was low ($M = 0.18$, $SD = .38$).

	M	$F(1, 39)$	p	BF_{10}
Repeats	-1.76	30.63	<.001	744
Adjacency	+.30	38.58	<.001	> 1000
Turning Point	-.08	9.91	.003	6
Phase Length ₁	-.07	10.47	.002	9
Phase Length ₂	-.04	.59	.45	1/13
Phase Length ₃	-.12	2.71	.11	1/7

Table 1: Mean values of the difference score d for six descriptive measures of randomness. Values other than 0 represent deviations from the theoretical value expected from *iid*.

Results

Descriptive Measures

We calculated the following measures from the transcribed sequences, adapting previous indices (see Towse & Neil, 1998) to two-dimensional domains where applicable:

- The proportion of repeated syllables (R).
- The proportion of adjacent syllables in the display (A).
- The proportion of turning points (TP). A TP is a transition for which the absolute difference between the current and previous direction is larger than 90° .
- Phase lengths (PL): A phase is the number of syllables before each turning point. We obtained the proportions of phases of lengths 1, 2, and 3 relative to all phases.

For an illustration of these indices, see Figure 1. Because these measures are dependent on the display’s dimension, we normalized them by obtaining a difference score d for each descriptive measure, which equalled the log of the observed proportion minus the log of the proportion expected from independent and identically distributed random samples (*iid* hereafter). We tested people’s random generation in two Bayesian linear mixed-effects models that included a random intercept u_0 per participant, using the *rstanarm* R package with default priors (Goodrich, Gabry, Ali, & Brilleman, 2020)². The first:

$$\hat{d} = \beta_0 + u_0 \times \text{participant} \quad (1)$$

tested whether people’s descriptive measures significantly deviated from *iid* (i.e. whether $\beta_0 \neq 0$). We expected the difference scores for Repeats and Turning Points to be lower than 0, and the one for Adjacency to be higher than 0, consistent with previous literature. The second model:

$$\hat{d} = \beta_0 + \beta_1 \times \text{Dimension} + \beta_2 \times \text{Block} + \beta_3 \times \text{Dimension} \times \text{Block} + u_0 \times \text{participant} \quad (2)$$

²For all models we fit both Bayesian and frequentist equivalents. The Bayes Factors (BF_{10} : evidence for the alternative hypothesis relative to the null) are interpreted as suggested by Jeffreys (1961).

tested whether these deviations were influenced by the global arrangement of the display participants learned (i.e. whether $\beta_1 \neq 0$). Consistent with the idea that random generation rests on a general cognitive ability, we expected no effects on these scores due to *Dimension*.

Because each participant produced two sequences, we included *Block* as a predictor as well as its interaction with *Dimension*, to check that any possible effects were not driven by fatigue. Because there were no credible differences due to *Block* or to the *Block* $\times$ *Dimension* interaction, we relegate those statistical analyses to Table 2, and only discuss the effects of *Dimension* in the main text.

	Factor	Δ	$F(1, 38)$	p	BF_{10}
Rep	Dim	+.13	.04	.85	1/19
	Blk	+.67	1.67	.20	1/11
	Dim $\times$ Blk	+.32	0.09	.76	1/23
Adj	Dim	+.37	22.49	<.001	261
	Blk	+.01	.16	.69	1/50
	Dim $\times$ Blk	-.06	1.23	.27	1/30
Turn Pt	Dim	-.13	8.63	.006	3
	Blk	-.004	0.06	.81	1/39
	Dim $\times$ Blk	+.05	1.27	.27	1/22
Ph Len ₁	Dim	-.18	38.80	<.001	>1000
	Blk	-.01	0.35	.55	1/34
	Dim $\times$ Blk	+.07	3.82	.06	1/7
Ph Len ₂	Dim	+.13	1.92	.17	1/7
	Blk	+.02	.10	.75	1/33
	Dim $\times$ Blk	-.10	1.08	.31	1/19
Ph Len ₃	Dim	+.56	18.95	<.001	22
	Blk	+.25	3.60	.07	1/4
	Dim $\times$ Blk	+.04	0.2	.88	1/26

Table 2: Effects of Dimension and Block order on d . Differences expressed as Δ . $\Delta Dim = 1D - 2D$; $\Delta Blk = 1st - 2nd$; $\Delta Interaction = \Delta Blk_{Dim=1} - \Delta Blk_{Dim=2}$.

The local arrangement of the syllables had no significant effect on any of the four measures of randomness (all $ps > .50$; all $BF_{10} < 1/5$), and so we aggregate data across arrangements in our subsequent analysis. Consistent with our hypotheses and previous literature, participants deviated from *iid* randomness in multiple ways (see Table 1). They repeated (R) syllables far less than expected from *iid*, and transitioned to adjacent syllables (A) more often than *iid*. They also made significantly fewer changes in direction than *iid* (TP), with fewer runs of one item (PL1) than expected but no credible difference from randomness in runs of two items (PL2) or three items (PL3).

These behaviors, however, depended on the display participants learned (see Table 2 and Figure 2), which we did not expect. Participants in the one-dimensional condition transitioned to adjacent syllables more often, making changes in direction less frequently, and having fewer runs of length one and more runs of length three. Interestingly, an ex-

ploratory post-hoc analysis revealed that participants in the two-dimensional condition did not make fewer changes of direction than expected by chance ($F(1, 19) = .22$, $p = .64$, $BF_{10} = 1/11$).

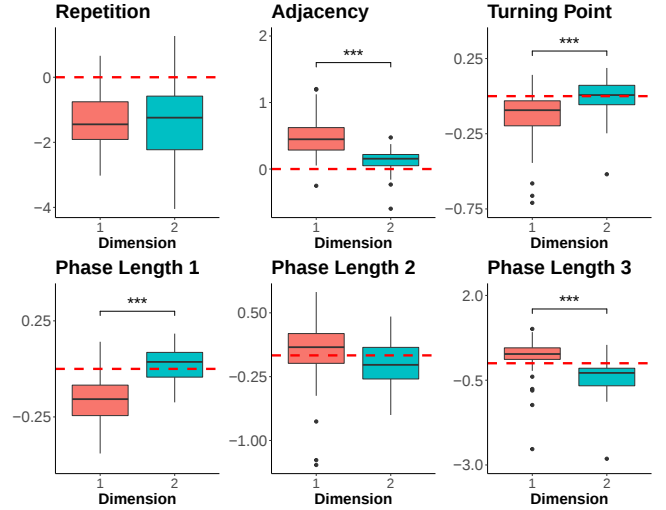


Figure 2: Deviations from the expected value (red dashed line) for six descriptive measures of random generation, according to the dimensionality of the learned representation. Participants in the 1D condition transitioned to adjacent syllables more frequently, and engaged in longer runs without turning, than participants who learned a two-dimensional representation

Model Fits

To test whether a sampling account could predict participants' sequences better than a schema account, we compared two models embodying the principles that each account would use. Each model was the weighted sum of four 'model parts' (described below). The likelihood of a syllable i was calculated as:

$$L_i = \sum_{n=1}^4 w_n \times \text{softmax}(\beta_n \times mp_n) \quad (3)$$

where each model part mp was the vector of weights for each syllable, scaled by β (bounded to be ≥ 0 to avoid the model parts to express the inverse prediction), and w was the model part weight ($0 \leq w \leq 1$ and $\sum w = 1$). The values of β and w were fit to each random sequence (i.e. two per participant).

Because a schema account would expect participants to utter syllables proportional to their relative frequency in the English language, possibly weighted by their frequency during the learning stage, we composed the schema model of two *Corpus* and two *Learning Sequence* parts. These parts weighted each syllable by considering the relative frequencies of syllables (*Unigram*) or pairs of syllables (*Bigram*), in the Brown English Corpus and in each participants' learning stage, respectively.

Instead, a sampling account would expect participants to explore locally and avoid repetitions. Thus, the sampling

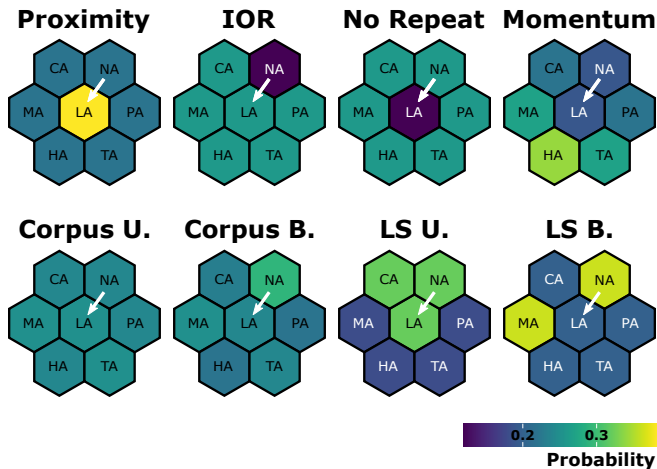


Figure 3: Probabilities assigned by each model part for the example participant in Figure 1. The last uttered syllables were *na*, *la*, and the learning sequence was: *ca*, *na*, *pa*, *la*, *ma*, *ha*, *ta*, *la*, *na*, *ca*. The figure presents the base probabilities of each model, which were scaled by β (see text).

model consisted of a *Proximity* part that weighted each syllable based on the Euclidean distance from the last uttered syllable and a *No Repeat* part that weighted syllables based on whether a syllable was repeated. To account for people's tendency not to make turns, we postulated two additional model parts: a tendency to maintain direction (*Momentum*), and an avoidance for the last-visited syllable (inhibition of return, *IOR*; see Johnson et al., 2013). *IOR* weighted syllables based on whether a syllable was the same as the penultimate uttered one; and *Momentum* weighted based on whether a syllable would deviate from the trajectory drawn in the display between the penultimate and the last syllable uttered (see Figure 3 for example probabilities). This 'model part' approach allowed us to measure the contribution of each principle in predicting participants' sequences.

As a baseline, we also compared the models with an *iid* sampling model, which predicted the same probability for each syllable. We compared model fits by calculating BIC scores and relative BIC weights.

Most participants were consistent across the two random generation stages: 27/40 participants had the same model achieve the lowest BIC scores in their two sequences (see Figure 4). The local sampling model best explained the performance of participants in 49/80 sequences (significantly more than expected by chance, one-tailed binomial test $p < .001$), while the schema and *iid* models best explained 16 ($p = .99$) and 10 ($p = 1$) sequences, respectively. The schema model performed best in the one-dimensional domain (12 sequences; 4 in the two-dimensional domain), whereas the local sampling (23; 26) and the *iid* models (4; 6) had similar performance in both domains.

To compare the relative contribution of the models' constituent parts, we fitted every possible model within the

schema and local sampling families consisting of 1-4 parts and computed their BIC weights. We used the model BIC weights as an approximation to the probability of the model given the data (see Neath & Cavanaugh, 2012), which allowed us to compute the probability of each model part within each class of models (as the sum of the BIC weights of the models that included it). This resulted in the *Learning Sequence (Bigram)* model achieving the best fit for the family of schema models, with an advantage over the next best model of 374 in BIC and a BIC weight of 1. For local sampling models, the best model was the *No Repeat + Momentum* model, with an advantage over the next best model of 306 in BIC and a BIC weight of 1.

Finally, we compared the full seven-parameter local sampling model to a model with only *Learning Sequence (Bigram)* as a predictor. We expected that the learning sequence model would perform better since it can encompass both how people learn the task and the spatial structure in the display (since participants had to learn the display by spatially navigating it) despite having only one parameter. However, the performance of both models was roughly equal, with the learning sequence best predicting 31 sequences, and 34 sequences best predicted by the local sampling model.

Discussion

Using a novel experimental design, we investigated how people generated random sequences after learning one- or two-dimensional displays of syllables. Overall, participants avoided repetitions, preferred adjacent items, and turned fewer times than expected, consistent with previous experiments in random generation (Cooper, 2016; Towse, 1998). However, contrary to our original hypothesis, we also found a significant effect of dimensionality on adjacency, turning point indices, and phase lengths: While participants across both one- and two-dimensional conditions preferred adjacent syllables, this preference was much higher for one-dimensional displays. Moreover, in one-dimensional structures participants produced ascending or descending runs frequently, resulting in lower than chance turning points and phase lengths. In contrast, in two-dimensional displays, participants produced turning points and phase lengths as expected by *iid* samples.

These results are important for understanding human random generation, as they highlight that the sample domain can influence the ability to produce random sequences. One possible explanation for this difference is that people adopt different cognitive mechanisms when producing samples from one- and two-dimensional domains. Alternatively, both domains might use the same mechanism but differ in the features of these mechanisms, or the same mechanism could produce different behavior by adapting to the sampling environment. For example, from a sampling perspective, it may be that the sampling adopts different parameters in one- and two-dimensional settings, or that more efficient algorithms are adopted in higher-dimensional domains. From a schema

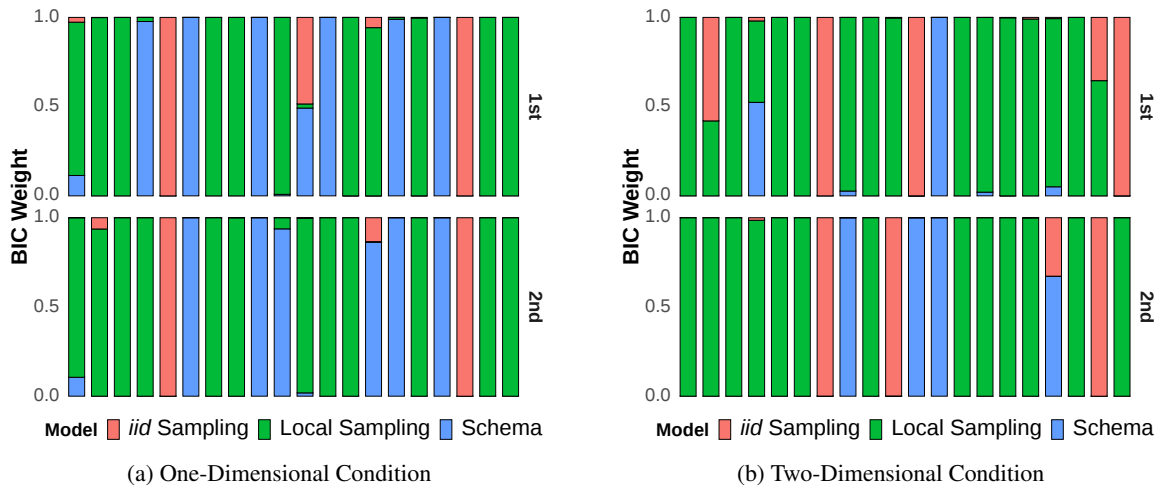


Figure 4: Distribution of BIC weights across the *iid*, local sampling and schema models. Each column represents one participant. Top and bottom rows are the first and second sequences a participant produced.

perspective, it is plausible that schemas are domain specific, or that monitoring processes are more constrained in higher dimensions.

In addition to finding an effect of global structure on peoples' random sequences, our model comparisons highlighted that, especially in two-dimensional domains, frequency-based models performed much worse than local sampling models. This result was especially striking given that the training participants received captured aspects of local sampling: Participants had to learn syllables adjacent to one another, and they naturally followed the previous direction when learning them, preferring to continue forwards rather than backward. A comparison of the contributions of each model part suggested that avoiding repetition and momentum were key features of the random generation process. Both features are important for sampling efficiently, local samplers without them are bound to produce frequent turns, only rarely reaching the edges of the domain.

Our results are consistent with recent models of human causal learning (Bramley, Dayan, Griffiths, & Lagnado, 2017), human exploration (Collignon & Lucas, 2019), category learning (Markant, Settles, & Gureckis, 2016), and even perception of ambiguous images (Gershman, Vul, & Tenenbaum, 2012), characterizing human cognition as an incremental process, anchored on the last, local, state.

Future Directions

Our experiment instructed participants to produce syllables at a constant rate, but participants did not always follow that rate exactly, exhibiting moderate deviations throughout the experiment. Following previous work that highlighted the close relationship between random generation and cognitive load or time constraints (Cooper, 2016; Towse, 1998), this suggests that these delays can be predictive of the participants' behavior. Future work should examine if the temporal structure of participants' sequences can inform patterns in their

generation of random items. Furthermore, these changes in production rates might also reflect fatigue, with participants re-using samples to reduce cognitive load, making fatigue a potentially useful predictor of increased stereotypical behavior.

Previous accounts of random generation emphasized the role of the central executive in monitoring the randomness of sequences (Baddeley et al., 1998; Cooper, 2016). Future work should explore extending our local sampling approach to include similar mechanisms, for example, by adapting the sampling process and monitoring measures of randomness of the sequences (for a review on such approaches, see Andrieu & Thoms, 2008). Combining these ideas into a unified theoretical model would give a new perspective on the role of randomness in human inference, and insight into human cognition more generally.

Acknowledgments

We thank Victoria Eshelby for helpful discussion at the early stages of this project. We thank the three anonymous reviewers for their feedback and suggestions. This work was supported by a European Research Council grant (817492-SAMPLING).

References

- Andrieu, C., & Thoms, J. (2008). A tutorial on adaptive MCMC. *Statistics and Computing*, 18(4), 343–373.
- Baddeley, A. (1966). Short-term memory for word sequences as a function of acoustic, semantic and formal similarity. *Quarterly Journal of Experimental Psychology*, 18(4), 362–365.
- Baddeley, A., Emslie, H., Kolodny, J., & Duncan, J. (1998). Random generation and the executive control of working memory. *The Quarterly Journal of Experimental Psychology: Section A*, 51(4), 819–852.

- Bramley, N. R., Dayan, P., Griffiths, T. L., & Lagnado, D. A. (2017). Formalizing Neurath's ship: Approximate algorithms for online causal learning. *Psychological Review*, 124(3), 301.
- Chater, N., Zhu, J.-Q., Spicer, J., Sundh, J., León-Villagrà, P., & Sanborn, A. (2020). Probabilistic biases meet the Bayesian brain. *Current Directions in Psychological Science*, 29(5), 506–512.
- Collignon, N., & Lucas, C. (2019). Epistemic drive and memory manipulations in explore-exploit problems. In *Proceedings of the 41st Annual Meeting of the Cognitive Science Society* (pp. 1540–1546).
- Cooper, R. P. (2016). Executive functions and the generation of “random” sequential responses: A computational account. *Journal of Mathematical Psychology*, 73, 153–168.
- Dasgupta, I., Schulz, E., & Gershman, S. J. (2017). Where do hypotheses come from? *Cognitive Psychology*, 96, 1–25.
- Gershman, S. J., Vul, E., & Tenenbaum, J. B. (2012). Multistability and perceptual inference. *Neural Computation*, 24(1), 1–24.
- Goodrich, B., Gabry, J., Ali, I., & Brilleman, S. (2020). *rstanarm: Bayesian applied regression modeling via Stan*. Retrieved from <https://mc-stan.org/rstanarm> (R package version 2.21.1)
- Horowitz, A. M. (1991). A generalized guided Monte Carlo algorithm. *Physics Letters B*, 268(2), 247–252.
- Jeffreys, H. (1961). *Theory of probability* (3rd ed ed.). Oxford: Oxford University Press.
- Johnson, M. R., Higgins, J. A., Norman, K. A., Sederberg, P. B., Smith, T. A., & Johnson, M. K. (2013). Foraging for Thought: An Inhibition-of-Return-Like Effect Resulting From Directing Attention Within Working Memory. *Psychological Science*, 24(7), 1104–1112.
- Lieder, F., Griffiths, T. L., Huys, Q. J., & Goodman, N. D. (2018). The anchoring bias reflects rational use of cognitive resources. *Psychonomic Bulletin & Review*, 25(1), 322–349.
- Markant, D. B., Settles, B., & Gureckis, T. M. (2016). Self-directed learning favors local, rather than global, uncertainty. *Cognitive Science*, 40(1), 100–120.
- Neath, A. A., & Cavanaugh, J. E. (2012). The Bayesian information criterion: Background, derivation, and applications: The Bayesian information criterion. *Wiley Interdisciplinary Reviews: Computational Statistics*, 4(2), 199–203. doi: 10.1002/wics.199
- Nosofsky, R. M. (1984). Choice, similarity, and the context theory of classification. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 10(1), 104.
- Ratcliff, R. (1978). A theory of memory retrieval. *Psychological Review*, 85(2), 59.
- Sanborn, A. N., & Chater, N. (2016). Bayesian brains without probabilities. *Trends in Cognitive Sciences*, 20(12), 883–893.
- Towse, J. N. (1998). On random generation and the central executive of working memory. *British Journal of Psychology*, 89(1), 77–101.
- Towse, J. N., & Neil, D. (1998). Analyzing human random generation behavior: A review of methods used and a computer program for describing performance. *Behavior Research Methods, Instruments, & Computers*, 30(4), 583–591.
- Vul, E., Goodman, N., Griffiths, T. L., & Tenenbaum, J. B. (2014). One and done? Optimal decisions from very few samples. *Cognitive Science*, 38(4), 599–637.
- Vulkan, N. (2000). An Economist's Perspective on Probability Matching. *Journal of Economic Surveys*, 14(1), 101–118.
- Wagenaar, W. A. (1972). Generation of random sequences by human subjects: A critical survey of literature. *Psychological Bulletin*, 77(1), 65–72.
- Zhu, J.-Q., Sanborn, A. N., & Chater, N. (2018). Mental Sampling in Multimodal Representations. In *Proceedings of the 32nd International Conference on Neural Information Processing Systems* (pp. 5753–5764).