

Random Behavior Is Stable Across Tasks and Time

Tal Boger¹, Sami R. Yousif², Samuel D. McDougale^{3, 4}, and Robb B. Rutledge^{3, 4}

¹Department of Psychological and Brain Sciences, Johns Hopkins University

²Department of Psychology and Neuroscience, University of North Carolina

³Department of Psychology, Yale University

⁴Wu Tsai Institute, Yale University

Whether it's choosing a tennis serve or escaping a predator, the ability to behave randomly provides a range of adaptive benefits. Decades of work explore how people both produce and detect randomness, revealing profound nonrandom biases and heuristics in our mental representations of randomness. But how is randomness realized in the mind? Do individuals have a "one-size-fits-all" conception of randomness that they employ across different tasks and time points? Or do they instead use simple context-specific strategies? Here, we develop a model that reveals individual differences in how humans attempt to generate random sequences. Then, in three experiments, we reveal that random behavior is stable across both tasks and time. In Experiment 1, participants generated sequences of random numbers and one-dimensional random locations. Behavior was remarkably consistent across the two tasks. In Experiment 2, we gave participants both a random-number-generation and a two-dimensional random-location-generation task, such that the tasks diverged in structure. We again observed stable individual differences across tasks. Finally, in Experiment 3, we collected data from the same participants as in Experiment 2, but 1 year later; we found stable individual differences across that span. Across all experiments, we find idiosyncratic behaviors that are stable across tasks and time. Thus, we suggest that a trait-like randomness generator exists in the mind.

Public Significance Statement

Understanding random behavior is crucial to understanding the mind. The presence—or absence—of randomness separates signal from noise, order from disorder, and meaningful patterns from mere coincidences. But how is randomness realized in the mind and implemented in behavior? One idea is that the mind attempts to generate randomness differently in different settings—that is, that our random behavior may differ from one task to another or from one day to the next. Another possibility is that random behavior is more trait-like—that is, that random behavior is stable across tasks and even across time. In three experiments, we provide evidence for the latter. We suggest that random behaviors may be realized by a stable individual randomness generator.

Keywords: randomness, computational modeling, random generation, statistical biases

Supplemental materials: <https://doi.org/10.1037/xge0001755.supp>

Among the foremost achievements of any cognitive system is the ability to act in seemingly random ways. A rabbit might randomly weave through a field to evade a fox; a tennis player might randomly choose their serves to keep their opponent guessing; and a competitor in a game may choose to behave erratically to deceive their opponent. Acting randomly aids in survival, learning, play, and

more. It is perhaps surprising, then, that humans exhibit many biases in tasks that require thinking about or generating randomness. For example, people judge sequences (e.g., a series of coin flips) with more "alternations" in states as more random (Kahneman & Tversky, 1972; Figure 1A), people think the past history of a random event affects its future probability (the "gambler's fallacy"; Clotfelter &

This article was published Online First April 10, 2025.

Wilma Bainbridge served as action editor.

Tal Boger  <https://orcid.org/0000-0002-5476-8416>

Tal Boger and Sami R. Yousif are equal contributors. These data are available at <https://osf.io/ebycj/>. Robb B. Rutledge holds equity in Maia. The other authors declare no competing interests. Robb B. Rutledge is supported by the National Institute of Mental Health (Grant R01MH124110). The authors thank Chaz Firestone for useful feedback and discussion.

Tal Boger played a lead role in data curation and formal analysis and an equal role in conceptualization, investigation, methodology, visualization, writing—original draft, and writing—review and editing. Sami R. Yousif played an equal role in conceptualization, data curation, formal analysis,

investigation, methodology, visualization, writing—original draft, and writing—review and editing. Samuel D. McDougale played a supporting role in conceptualization, methodology, supervision, and writing—review and editing. Robb B. Rutledge played a lead role in funding acquisition and a supporting role in conceptualization, project administration, supervision, and writing—review and editing.

Correspondence concerning this article should be addressed to Tal Boger, Department of Psychological and Brain Sciences, Johns Hopkins University, 3400 North Charles Street, Baltimore, MD 21218, United States, or Sami R. Yousif, Department of Psychology and Neuroscience, University of North Carolina, 235 East Cameron Avenue, Chapel Hill, NC 27599, United States. Email: tboger1@jhu.edu or sami.yousif@unc.edu

Cook, 1993), and so on (Arkes et al., 1988; Gilovich et al., 1985; Roese & Vohs, 2012). These biases in judgments of randomness seem to persist across many different tasks and modalities (Sherman et al., 2022; Yu et al., 2018).

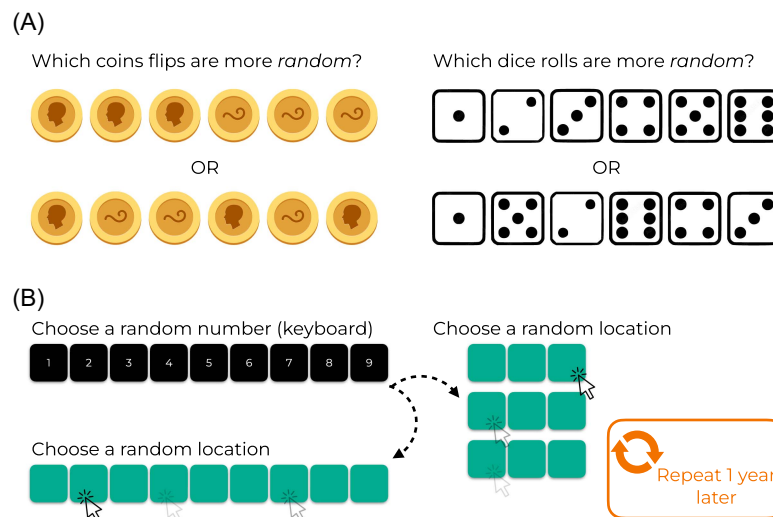
Biases of randomization are not just a quirk of human cognition. Rather, they may reveal something fundamental about how humans interact with their environments. Randomness is important for how we perceive: Detecting the movement and appearance of a snake in a pattern of foliage may reduce our chance of injury. Randomness is important for how we learn: Perceiving randomness modulates learning via simple conditioning (Wagner & Rescorla, 1972), statistical learning (Saffran et al., 1996, 1999; Turk-Browne et al., 2005; for review, see Sherman et al., 2020), exploration (Gershman, 2018, 2019; Tomov et al., 2020), and even higher level language acquisition (Kelly & Martin, 1994). Randomness is important for how we remember: Statistical regularity improves our memory (Sherman & Turk-Browne, 2020). And randomness drives how we act: The ability to act randomly can help prey animals evade predators (Moore et al., 2017; Szopa-Comley & Ioannou, 2022), tennis players optimize their choice of serve (Walker & Wooders, 2001), and soccer players make decisions in penalty shootouts (Misirlisoy & Haggard, 2014). In this way, unpredictable (i.e., random) behavior is advantageous and important to any biological system (see, e.g., “protean behavior”; Driver & Humphries, 1988).

For this reason, a large body of work has advanced our understanding of how—and why—our representations of randomness

may be biased or nonrandom (see Bar-Hillel & Wagenaar, 1991; Falk & Konold, 1997; Griffiths & Tenenbaum, 2001, 2003; Griffiths et al., 2018; Kareev, 1992; Rapoport & Budescu, 1992; Warren et al., 2018). For example, some suggest that biases in randomness may be a result of a Bayesian inference process over observed data (Griffiths & Tenenbaum, 2001, 2003); others suggest that such biases in fact reflect unbiased estimates of our statistical environments (Warren et al., 2018). Here, we add to this discussion by asking how randomness is realized in the mind. Specifically, we ask: Are random behaviors the product of separate, domain-specific heuristics that change over time? Or instead, might randomness be realized in a trait-like fashion, such that each person has a unique “randomness generator” that persists both across tasks and across time?

We measure human random behavior using randomness generation tasks in different domains and across different time points (Figure 1B). Asking people to generate—rather than judge—random sequences has provided rich insights into how the mind represents randomness (see A. D. Baddeley, 1966; Treisman & Faulkner, 1987; Wagenaar, 1972). Indeed, in much the same way that a rich literature on the *perception* of randomness has illuminated several aspects of human randomness (as above), so too has a rich literature on the *generation* of randomness illuminated similar and other questions (see Nickerson, 2002, for a discussion of the two together). For example, an overalternation bias pervades not only the perception of randomness but also the generation of randomness,

Figure 1
Studying Human Randomness



Note. (A) Which set of coin flips appears more random? Although both sequences are equally random, people consistently say that the second is more random than the first. This is one example of the “over-alternation bias,” whereby people conflate randomness with changes in state. This can also be expressed in higher level forms than mere repetitions, such as “jumps” in number, changes in monotonicity, and more. For example, the second sequence of dice rolls appears more random than the first. (B) We asked participants to generate sequences of random numbers (using their keyboard) and random locations (either in one or two dimensions by clicking on boxes with their mouse). The same participants completed these tasks 1 year later. Thus, our approach leverages large-scale online data collection to find stable random behavior across tasks and time. See the online article for the color version of this figure.

insofar as people have a propensity to not repeat choices (and thus alternate between options more than would otherwise be expected; A. Baddeley et al., 1998; Castillo et al., 2024; Cooper, 2016; Feng & Rutledge, 2024; Rutledge et al., 2009). Along similar lines, when generating random sequences, people tend to exhaust all available choices more quickly than would be expected by chance (Ginsburg & Karpiuk, 1994; Towse & Neil, 1998). Additionally, research on the generation of randomness has revealed that people can learn to generate sequences nearly indistinguishable from true randomness (Neuringer, 1986) and that the ability to generate random sequences relates to many other cognitive functions, such as working memory and executive control more broadly (A. Baddeley et al., 1998; Biesaga & Nowak, 2024; Jahanshahi et al., 2006). Thus, studying randomness generation both provides supporting evidence for findings in randomness perception and opens the door to new sets of questions about the function and structure of randomness in the mind.

More immediately, studying the generation of randomness (rather than the perception of randomness) allows us to not only infer general heuristics of human randomization but also extract precise patterns on an individual level (see, e.g., Schulz et al., 2012). By asking people to generate sequences in multiple domains, we test whether human randomization biases persist across *tasks* (i.e., whether a person's behavior in a random-number-generation task may look similar to their behavior in a random-location-generation task). By asking people to generate sequences across lengthy delays, we test whether these biases persist across *time* (i.e., whether a person's behavior in a random task today will resemble their behavior in that task—or a different one!—1 year later).

Our computational modeling approach allows us to characterize domain-general, time-invariant random behaviors on an individual level. In much the same way that people have consistent extraverted, risk-averse, or spontaneous behaviors (Poropat, 2009), we suggest that people have consistent random behaviors.

Method

Participants and Procedure

Participants in all experiments were recruited from the online platform Prolific (for a discussion of the reliability of this subject pool, see Peer et al., 2017). Subjects received compensation upon completing the experiment. The experiments were approved by the Yale University Institutional Review Board. Experiments 1 and 2 recruited 200 unique subjects each. Experiment 3 invited back all participants from Experiment 2 1 year later. Eighty-four people completed Experiment 3.

In Experiment 1, participants completed a random-number-generation task and a one-dimensional random-location-generation task back-to-back. The order of these two tasks was randomly counterbalanced across participants; 100 participants were randomly selected to complete the number-generation task first, and the other 100 participants completed the location-generation task first. In each task, subjects generated a sequence of 250 random choices (in the number-generation task, these were made with their keyboard; in the location-generation task, these were made by clicking on a box). On each “trial,” participants saw either an empty screen except for instructions (in the case of the number-generation task) or a row of nine black boxes (in the case of the location-generation

task) and generated a random choice. The choice appeared on screen for 750 ms—either via the chosen number flashing on screen or by the chosen box turning teal—before the choice disappeared. In this way, participants were generating random sequences with no look back, as they could not see any previous choices.

Experiment 2 was the same as Experiment 1, except that now participants completed a two-dimensional location-generation task instead of a one-dimensional location-generation task. So, instead of the boxes appearing in a 1-by-9 grid, they now appeared in a 3-by-3 grid. This two-dimensional task differs in representational format from the number task, allowing us to test the limits of this domain-general. Note that participants could not input responses to the number-generation task with the number pad (and could instead only do it via the numbers on the top of the keyboard).

Experiment 3 used the exact same design as Experiment 2, even down to the order of tasks (i.e., if a subject completed the two-dimensional task in Experiment 2, they completed it first in Experiment 3, too). Data for Experiment 3 were collected about 1 year after Experiment 2 (an average of about 11 months and 2 weeks later).

Exclusions

Prior to data analysis, we excluded participants who behaved in clearly nonrandom ways. This is because highly nonrandom participants are significantly *easier* to capture with a model. For instance, a participant who selects each number successively (e.g., 1–2–3–4–5) will be easily modeled, but such behavior reveals nothing meaningful about how randomness works in the mind. Thus, we are excluding participants only to be conservative; we want to know if we can capture systematicity in the behavior of even the most random participants in our sample. Note that all our results hold (and are in fact stronger) if we perform no exclusions (as shown in our [Supplemental Material](#)).

We used two metrics to determine nonrandom behaviors. First, we excluded any participant who selected any of the nine options fewer than 10 total times out of 250 trials (i.e., less than 4% of the time). Second, we computed an average numerical distance for each participant—the mean absolute difference between consecutive choices—and excluded participants whose value was below a certain threshold (2.4 for the number-generation and one-dimensional location-generation tasks, 1.27 in the two-dimensional location-generation tasks). Both these exclusion criteria were chosen based on the fact that discrete uniform distributions do not contain these properties approximately 99.99% of the time. In other words, 99.99% of truly random sequences (generated by 250 discrete uniform selections from the numbers 1 to 9) meet these criteria. Note that the average numerical distance threshold in the number-generation and one-dimensional location-generation tasks differs from that of the two-dimensional location-generation tasks due to a difference in format. In the number and one-dimensional tasks, average numerical distance is calculated via differences between consecutive choices—that is, subtracting one choice from the next. However, in the two-dimensional tasks, average numerical distance is calculated via Euclidean distance between consecutive choices. So, for example, the distance from the top-left box (i.e., [1, 1]) to the bottom-right box ([3, 3]) is 2.83.

We calculated these exclusion criteria for each task in each experiment. Within an experiment, a participant's data were

excluded from *both* tasks if they failed to meet either criterion in either task. This is because, in order to test cross-task or cross-time performance, we wanted equal numbers of participants in both tasks of each experiment. These exclusion criteria left us with 141 participants in Experiment 1 and 142 participants in Experiment 2 (out of 200 recruited in each). In Experiment 3, we excluded participants who failed to meet either criterion in either task or in either time of measurement (i.e., either Experiment 3 or Experiment 2, such that, again, we had equal numbers of participants for each comparison). This excluded 31 of 84 total participants, leaving 53 total.

Emergent Properties

Throughout our analyses, we rely on a set of emergent properties (what some might call “sequence features”) that capture high-level features of random sequences. The first emergent property we examined was the number of repeats—the number of times a subject contributes the same choice two trials in a row. The second was average numerical distance; as above, this is defined as the mean of the difference between consecutive choices in the number-generation and one-dimensional location-generation tasks and as the mean of the Euclidean distance between consecutive choices in the two-dimensional location-generation task. The final property was direction switches; this is defined as the number of times the sequence changes in monotonicity. For example, the sequence “1-5-7-2-4” contains two direction switches: one from 7 to 2 (as, prior to choosing 2, the sequence was increasing from 5 to 7) and one from 2 to 4 (as, prior to choosing 4, the sequence was decreasing from 7 to 2). However, this instantiation of direction switches holds only for the number-generation and one-dimensional location-generation tasks, where sequences vary only in one dimension.

In the two-dimensional location-generation task, we defined direction switches in an angular way that generalizes across dimensions. To account for the angular movement in a two-dimensional grid while keeping the same principles, we defined a direction switch as a movement to any box that lies on the same side of the normal vector formed by the previous two choices. In other words, we first calculated the angle formed by c^{t-1} and c^t (i.e., $\text{atan}(c^{t-1} - c^t)$), where c^t is the choice made on trial t (and c^{t-1} is the choice made on trial $t-1$). Then, we considered each possible next-choice c^{t+1} as being in the same “direction” if it was on the opposite side of the normal vector formed by these two choices; so, c^{t+1} is not a direction switch if $-90 \leq \text{atan}(c^{t-1} - c^t) - \text{atan}(c^t - c^{t+1}) \leq 90$. This definition maps onto the number and one-dimensional tasks; consider the points in these tasks as changing only in x , with y being constant. Choices increasing in x (i.e., increasing from 1 to 9) all have an “angle” of 0, and thus, monotonic increases have a difference of 0 between consecutive choices (implying no direction switch). However, if a choice flips a sequence from increasing to decreasing, then the “angle” is 180 degrees, resulting in a direction switch using the above definition. Additionally, note that choices *on* the normal vector itself are not direction switches in this definition; for example, the sequence “1-5-7” on a keypad (i.e., two-dimensional grid) does not count as a direction switch.

Modeling

We fit our models with a nonlinear program solver from the MATLAB optimization toolbox. We fit a model to each participant

in each task, allowing us to compare parameters and predictions across tasks (and across time).

In each task, we “fit” a Markov chain, in other words extracting participant transition probabilities from each individual sequence, to use as a baseline model. Let S be the space of possible choices in the task, c^t be the chosen value on trial t , and $P(S)^{t+1}$ be the probability distribution over the sample space for the next trial. Thus, $P(c^t, S)$ is the transition probability from choice c^t to any possible choice in S . The Markov chain is exactly this; in other words, the Markov chain’s prediction for the next choice is defined as $M : P(S)^{t+1} = P(c^t, S)$.

We compared the Markov chain to our own model. Our model for the number-generation and one-dimensional location-generation tasks had three parameters: a stay parameter, a side-switch parameter, and a direction-switch parameter. Each parameter was bounded $(-1, 1)$. These parameters are all built on top of a naive random model with no prior assumptions about transition probabilities; in other words, these models adjust choice likelihood from a uniform distribution over possible choices. Thus, each model starts with $M : P(S)^{t+1} = U(S) = \frac{1}{9}$, as all nine choices are equally likely. After all parameters were added, we applied a softmax function to the result, giving us probabilities summing to 1.

The stay parameter, α , adjusts the probability of the number that was just chosen (i.e., it makes a repeat either more or less likely). So, this means that $P(S = c^t)^{t+1} = \frac{1}{9} + \alpha$. For most participants, the stay parameter was at its lower bound, meaning that most participants very rarely repeated choices.

The side-switch parameter, β , adjusts the probability of the set of numbers on the opposite “side” of the sequence from what was just chosen. The set of numbers 1–9 has two sides: 1–4 and 6–9 (as does a row of nine boxes). This parameter was subtracted from the probability of choosing a number on the same side as the previous choice and added to the probability of choosing a number on the opposite side. Note, however, that the parameter $\beta \in [-1, 1]$, meaning that the parameter could actually *increase* the probability of staying on the same side (and reduce the probability of switching sides) if it was negative. A mathematical operationalization of the side-switch parameter is $P(S < 5 | c^t < 5)^{t+1} = P(S < 5 | c^t < 5) - \beta$ and $P(S > 5 | c^t < 5)^{t+1} = P(S > 5 | c^t < 5) + \beta$ (and vice versa for choices on the other side, i.e., $c^t > 5$).

In the two-dimensional case, the side-switch parameter is defined in a way that accounts for the structure of the task: A side switch in our 3-by-3 grid is a movement to any choice that is not adjacent to the previously chosen box. This preserves all the features of side switches in the number and one-dimensional case. Specifically, from the center point of the grid, *nothing* is a side switch; similarly, from the number 5 (or the fifth box in the row), *nothing* is a side switch. Additionally, this definition ensures that a side switch occurs in at least one dimension of the grid. For example, suppose we are considering what is a side switch from the top-left box. Across the rows of the grid, it is certainly a side switch to jump from the first row to the third, regardless of the column chosen, but so too across the columns of the grid it is a side switch to go from the first column to the third, regardless of the row chosen. This operationalization considers both cases as side switches.

These parameters were added into the model in the order that they are presented here; first, we adjusted the probability of the previous choices; then, we adjusted the probability of both sides of the sequence before adding a final direction-switch parameter. This

parameter, γ , works in much the same way as the side-switch parameter, but accounts for direction instead. It is added to choices that imply a direction switch from the previous choice and subtracted from choices that do not imply a direction switch from the previous choice. So, $P(S > c^t | c^t < c^{t-1})^{t+1} = P(S > c^t) + \gamma$ and $P(S < c^t | c^t < c^{t-1})^{t+1} = P(S < c^t) - \gamma$ (and vice versa for the opposite direction, where $c^t > c^{t-1}$).

As mentioned above, operationalizing direction switches in two dimensions is computationally complex. Adding it to our model squashed the other parameters and affected the results negatively. Thus, our model in the two-dimensional tasks does not include this parameter.

Putting together all our definitions in the order that our model computes them, we have the following:

$$M: P(S)^{t+1} = U(S) = \frac{1}{9}, \quad (1)$$

$$P(S = c^t)^{t+1} = \frac{1}{9} + \alpha, \quad (2)$$

$$P(S < 5 | c^t < 5)^{t+1} = P(S < 5) - \beta \text{ (and vice versa)}, \quad (3)$$

$$P(S > c^t | c^t < c^{t-1})^{t+1} = P(S > c^t) + \gamma \text{ (and vice versa)}. \quad (4)$$

After these computations, we transformed the outputs into a probability distribution that summed to 1 using softmax. All model simulation results reported in the article are performed by generating these probability distributions on each “trial” and then sampling the distribution several times, resulting in several different model-generated sequences for each participant.

Transparency and Openness

Readers can experience all of our experiments—exactly as they were presented to participants—at <https://perceptionstudies.github.io/randomness/>. All data, experiment scripts, models, and more are publicly available at <https://osf.io/ebycj/>.

Note that our modeling choices and experiments were not pre-registered. This is because we were open to a variety of possibilities and outcomes before collecting data and thus did not want to commit to a specific model or approach. We are indeed open to the possibility that other models may outperform ours (e.g., Angelike & Musch, 2024; Castillo et al., 2024); our model aims to demonstrate stability across tasks and time. Furthermore, we take considerable steps to check that our model is not overfitting and ensure that our results hold regardless of exclusion criteria (see Supplemental Material).

Results

Experiment 1: Consistent Behavior in Random-Number-Generation and Random-Location-Generation Tasks

In Experiment 1, subjects ($N = 142$ after exclusions, see the Method section) generated a sequence of 250 random numbers (using their keyboard) and 250 random locations (using their mouse to click on one box in a row of nine; Figure 2A). The order of these two tasks was randomly counterbalanced across participants. Although such a task may seem artificial (as people are rarely asked to generate sequences of random numbers or locations in their everyday lives), this approach has been useful in

revealing the basic tendencies of human random behavior (see, e.g., Kahneman & Tversky, 1972). This approach is fruitful in part because people behave systematically even when they *think* they are behaving unpredictably. Our paradigm exploits that fact. We collected a rich, cross-domain data set of random sequences. Each subject contributed 250-item-long sequences of both random numbers and random locations. We asked: Do subjects’ sequences in these two tasks share behavioral signatures?

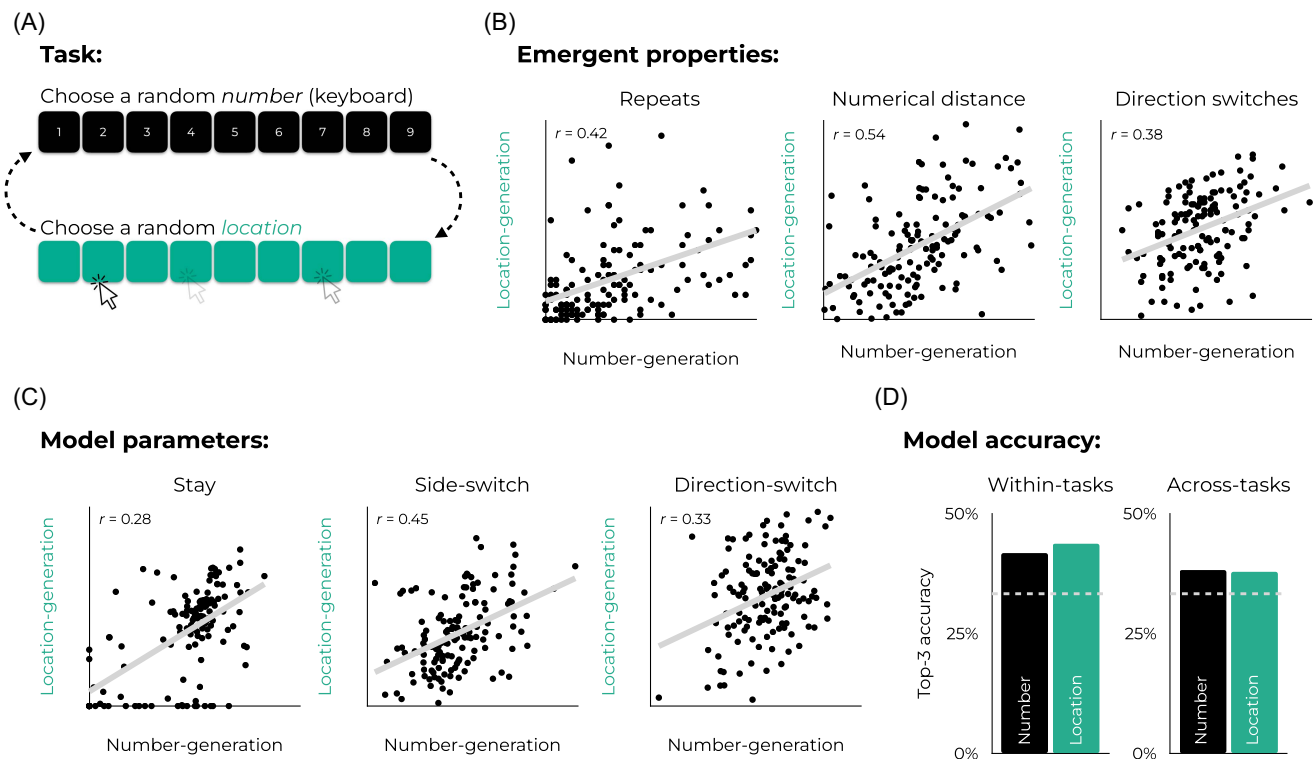
We excluded subjects who contributed substantially nonrandom sequences (e.g., a sequence of the same number each time; criteria are described in detail in our Method section). Note that this stacks the deck *against* our predictions, as people who behave non-randomly are easier to predict and thus may display even more systematic behavior across tasks. Our conclusions do not change when running all our analyses on data without any exclusions; see Supplemental Material. Then, we examined several emergent properties (i.e., what some might refer to as “sequence features”) of the sequences generated by subjects. Specifically, we computed the number of repeats (i.e., the total number of times that a participant selected the same option multiple times successively), the average numerical distance (i.e., the average distance between consecutive choices), and the number of direction switches (i.e., the total number of times the sequence switches from consecutively ascending numbers to descending numbers, or vice versa) for each sequence. Similar properties have been used as metrics of success in previous works examining randomness generation (see Castillo et al., 2024; Towse & Neil, 1998).

Each participant contributed two sequences—one from each task—allowing us to conduct model-free analyses asking whether subjects behaved similarly across the two tasks (Figure 2B). We observed strong cross-task correlations for each of our three emergent properties: On the subject level, the number of repeats, $r(140) = 0.42$, $p < .001$; average numerical distance, $r(140) = 0.54$, $p < .001$; and the number of direction switches, $r(140) = 0.38$, $p < .001$, were all correlated across tasks.

We also analyzed other features of these sequences, like their “redundancy” (i.e., the uniformity of each response alternative, as in Towse & Neil, 1998). Interestingly, we found no correlation in redundancy across these two tasks, $r(140) = 0.10$, $p = .23$. The absence of a significant correlation here suggests that the consistency we observe for the above properties reflects something deep about the sequence-generation process and is unlikely to be explained by some other domain-general factor such as effort.

To further probe subtle individual differences in random generation, we developed a descriptive model for each participant’s sequence in each task. This model was fit with maximum likelihood estimation (implemented via a nonlinear program solver from the MATLAB optimization toolbox; see the Method Section). The model contains only three parameters: (a) a stay parameter (which captures the participant’s propensity to repeat choices), (b) a side-switch parameter (which captures the participant’s propensity to switch “sides” in the response space), and (c) a direction-switch parameter (which captures the participant’s propensity to switch “directions” in the response space). Note that these parameters are quite general and can be applied to nearly any one-dimensional space (and we later show how this is true in two dimensions too). We compare our model to a Markov chain that predicts a participant’s next choice using their transition probabilities between choices in each sequence (see Supplemental Material for discussion of an

Figure 2
Stable Random Behavior Across a Number-Generation and One-Dimensional Location-Generation Task



Note. (A) Subjects generated two random sequences: one of random numbers (by pressing number on their keyboard) and one of random locations (by clicking on boxes a row of boxes with their mouse; the order in which subjects completed the two tasks was randomly counterbalanced across subjects). (B) Emergent properties of the sequences for each subject. All three properties—number of repeats, average numerical distance, and number of direction switches—were significantly correlated across tasks. (C) Model parameters for each subject. Note that each parameter was significantly correlated across tasks, suggesting domain-general random behavior. (D) Model accuracy for predicting the next item in the sequence. Our models successfully predicted subject behavior both within tasks (i.e., using the model fit to number-generation to predict data in the number-generation task; 41.8% in the number-generation task and 43.7% in the location-generation task) and across tasks (i.e., when using the model fit to location-generation to predict data in the number-generation task, we observe 38.2% Top 3 accuracy, and the opposite yields 37.8% accuracy). Note that bars are labeled by the model used; so in the across-tasks graph, the “number” bar denotes the performance of the number-generation model in predicting location-generation data. See the online article for the color version of this figure.

additional, second-order Markov baseline). Our model is described in full detail—including parameter definitions and derivations—in the Method section, and all data and code are available in our Open Science Framework repository (<https://osf.io/ebycj/>).

Our model captured individual differences in randomness generation (Figure 2C): Each of the three parameters was significantly correlated across tasks, suggesting consistent domain-general behavior on the individual level, stay parameter: $r(140) = 0.28$, $p < .001$; side-switch parameter: $r(140) = 0.45$, $p < .001$; direction-switch parameter: $r(140) = 0.33$, $p < .001$. Note that our model builds on (and replicates) classic biases in human randomness, such as the overalternation bias (A. Baddeley et al., 1998; Castillo et al., 2024; Cooper, 2016; Kahneman & Tversky, 1972; Walker & Wooders, 2001). Specifically, the stay parameter—which adjusts the likelihood of repeating the previous choice—was at (or near) its lower bound for almost every subject (median stay parameter was below -0.9 , with the lower bound being -1), indicating that subjects rarely repeated choices.

To further evaluate the performance of our model, we used each participant’s parameters to simulate 100 new sequences. Then, we

also simulated 100 sequences for each participant using just their transition probabilities (i.e., the Markov chain baseline). For each of these 100 simulations, we computed the average numerical distance and number of direction switches for each subject’s simulated sequence and asked whether our model had a higher correlation with the true average numerical distance and direction switches (i.e., from the participant’s real sequence) than the Markov chain baseline. In other words, for each simulation, we computed two correlations across participants—one between our model-generated sequence and the data and one between the Markov chain baseline-generated sequence and the data—and asked which correlation was higher. We expected that our model would outperform the Markov chain baseline in most simulations if it is in fact capturing meaningful variation in random behavior. And that is what we observed: Our model outperformed the Markov chain in terms of both average numerical distances and direction switches in every simulation we conducted.

We next measured performance more directly by asking: How often does the model-generated sequence correctly predict an individual’s actual choices? In other words, given a participant’s last two responses (a window long enough to instantiate our parameters),

does our model successfully predict the next response? In principle, it is possible that our models successfully describe emergent properties of the data, but fail to predict specific values.

When comparing model-generated sequences from the random-number-generation model to the true random-number-generation sequences, we found that the true choice was among the Top 3 most-likely model predictions 41.8% of the time ($p < .001$ in the Wilcoxon signed-rank test compared to fits to 100 shuffles of the participant's data, which yielded an average accuracy of 33.8%; Figure 2C). In the one-dimensional random-location-generation task, the model achieved 43.7% Top 3 accuracy ($p < .001$ compared to 33.2% average accuracy for fits to 100 shuffles). These results suggest that the model can predict individual response choices beyond just summary biases.

Most importantly, the models also made accurate response-specific predictions *across* tasks. We asked whether these results reflected patterns that are not just applicable to *any* participant, but rather specific and unique to each individual by comparing the model's accuracy for each participant to what would be predicted by using the parameters from every other participant—in other words, a leave-one-out prediction accuracy. The model fit to the location-generation task predicted a participant's data in the number-generation task with 38.2% Top 3 accuracy ($p < .001$ compared to 34.0% in the leave-one-out baseline). Further, the number-generation model predicted the location-generation data with 37.8% accuracy ($p < .05$ compared to 36.4% based on leave one out). Thus, our model successfully predicts variance in specific choices both within and across tasks on an individual level.

Collectively, these results suggest that (a) aspects of human randomization are correlated across tasks, (b) our model captures meaningful variation in human randomization, and (c) cross-task similarities are explained by individual idiosyncrasies and not just group-level heuristics.

Experiment 2: Consistent Behavior in Tasks With Different Surface-Level Features

Although Experiment 1 revealed cross-task similarities, one might argue that the two tasks share an underlying representational format (i.e., a straight line in space), providing a simple explanation for the observed similarities (i.e., without invoking a stable randomness generation). After all, human adults are known to represent numbers spatially via a mental number line (Dehaene et al., 1993). In Experiment 2 ($N = 143$ after exclusions), we tested this domain-generalizability by replacing the one-dimensional random-location-generation task with a two-dimensional random-location-generation task.

The benefit of the two-dimensional approach is that, although it is still fundamentally a spatial randomization task, the format of this task differs from that of the number task. Rather than being arranged in a straight line, the cells in this experiment were organized in a 3-by-3 grid (Figure 3A). Mappings of the numbers 1–9 onto a 3-by-3 grid are substantially weaker than a typical mental number line (Darling et al., 2017; Darling & Havelka, 2010). Note that one can map the numbers 1–9 onto a 3-by-3 grid via the number pad of a keyboard. However, in our experiments, we disabled the number pad, such that participants could generate random numbers only by using the top row of their keyboards, where the numbers are arranged in a

straight line. Thus, observing similarities across the two tasks would further support the idea that random behavior is domain-general.

As in Experiment 1, we first examined model-free correlations in emergent properties across the two tasks. Again, we found strong cross-task correlations in each of the properties: number of repeats, $r(141) = 0.51$, $p < .001$, Figure 3B; average numerical distance, $r(141) = 0.44$, $p < .001$; and number of direction switches, $r(141) = 0.32$, $p < .001$. The cross-task correlation for side switches and direction switches is especially striking given that they are instantiated for one dimension (i.e., for the number-generation task) differently from how they are instantiated for two dimensions (i.e., in polar space for the two-dimensional location-generation task; both definitions are described in the Method section). Note that because of this difference in instantiation, our final model of the two-dimensional data does not include a direction-switch parameter.

Subject-specific parameters were correlated across the two tasks: stay parameter, $r(141) = 0.48$, $p < .001$; side-switch parameter, $r(141) = 0.41$, $p < .001$; Figure 3C. Furthermore, even with only two parameters, our model again outperformed the Markov chain in generating simulated data for each participant: As in Experiment 1, our model outperformed the Markov chain baseline in both emergent properties and both tasks for all 100 simulations. Despite the fundamental differences between a random-number-generation and a two-dimensional random-location-generator task, participants behaved similarly across the two tasks.

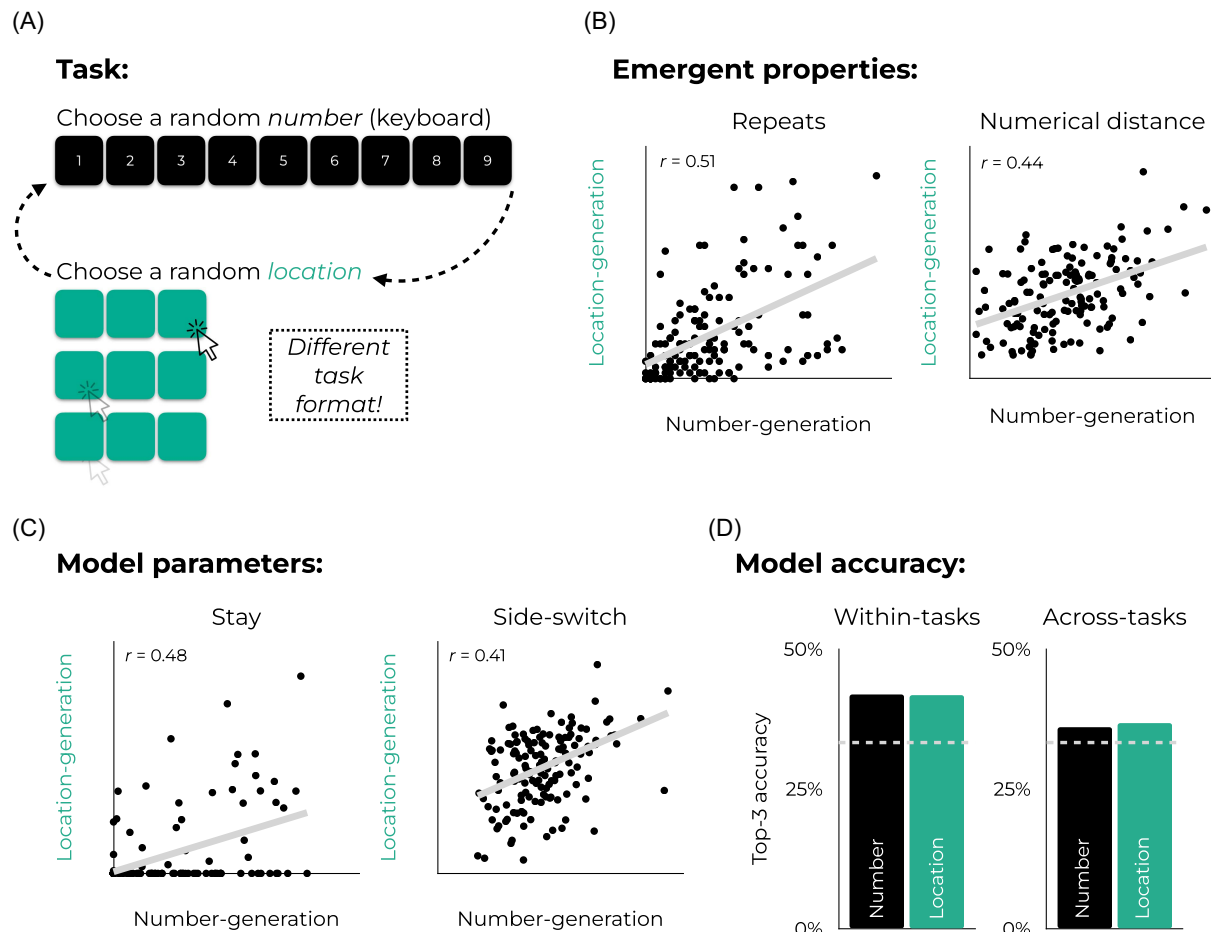
As before, we tested the accuracy of our models both within and across tasks. Observing cross-task predictiveness would be especially surprising, as the parameters are instantiated differently across the two tasks. However, we again observed consistently strong cross-task model accuracy (Figure 3D): The number-generation model predicted the correct value in its Top 3 most-probable values 41.9% of the time ($p < .001$ compared to 33.6% average accuracy for fits to 100 shuffles), and the two-dimensional location-generation model predicted the Top 3 value 41.7% of the time ($p < .001$ compared to 33.2% average accuracy for fits to 100 shuffles). Furthermore, the model fit to the two-dimensional location data predicted the number-generation data with 36.8% Top 3 accuracy ($p < .001$ compared to 33.2% leave-one-out accuracy), and the model fit to the number-generation data predicted the two-dimensional location data with 36.0% Top 3 accuracy ($p < .05$ compared to 38.3% leave-one-out accuracy). The high baseline accuracy in this latter case suggests that 2D location data are well-described by group-level heuristics.

Together, these results suggest that individual patterns of randomization are stable across tasks even when the tasks differ substantially in their representational format and features.

Experiment 3: Consistent Behavior Over Time

If random choice behavior truly is trait-like, perhaps the most straightforward feature we should expect to find is stability over time—much like how some personality traits are stable over time. Here, we asked whether we could predict a subject's performance on our tasks 1 year after they completed the original task. This experiment also controls for potential confounds related to collecting data from multiple tasks on the same day; if a stable randomness generator exists

Figure 3
Stable Randomness in Tasks With Different Surface-Level Features



Note. (A) Subjects generated two sequences of random numbers: As before, one was generated by pressing random numbers on the keyboard. The other was generated by clicking on random locations with the mouse in a two-dimensional grid. The task was designed such that the representational format differs between the two tasks (unlike in Experiment 1, where one-dimensional grids could be represented as number lines). (B) Emergent properties were correlated on the individual level. (C) Model parameters for each subject. Unlike in Experiment 1, there was no direction-switch parameter here due to its different instantiation in two dimensions. The other two parameters were significantly correlated across tasks. (D) Model accuracy for predicting the next item in the sequence. As before, the model successfully predicted subject behavior at rates above chance within tasks (41.9% for number-generation, 41.7% for two-dimensional location-generation). We also found evidence for cross-task predictive power (36.8% for two-dimensional location-generation model predicting number-generation data, and 36.0% for the opposite). See the online article for the color version of this figure.

on an individual level, then a subject's random behavior should contain similar patterns at any time at which it is measured.

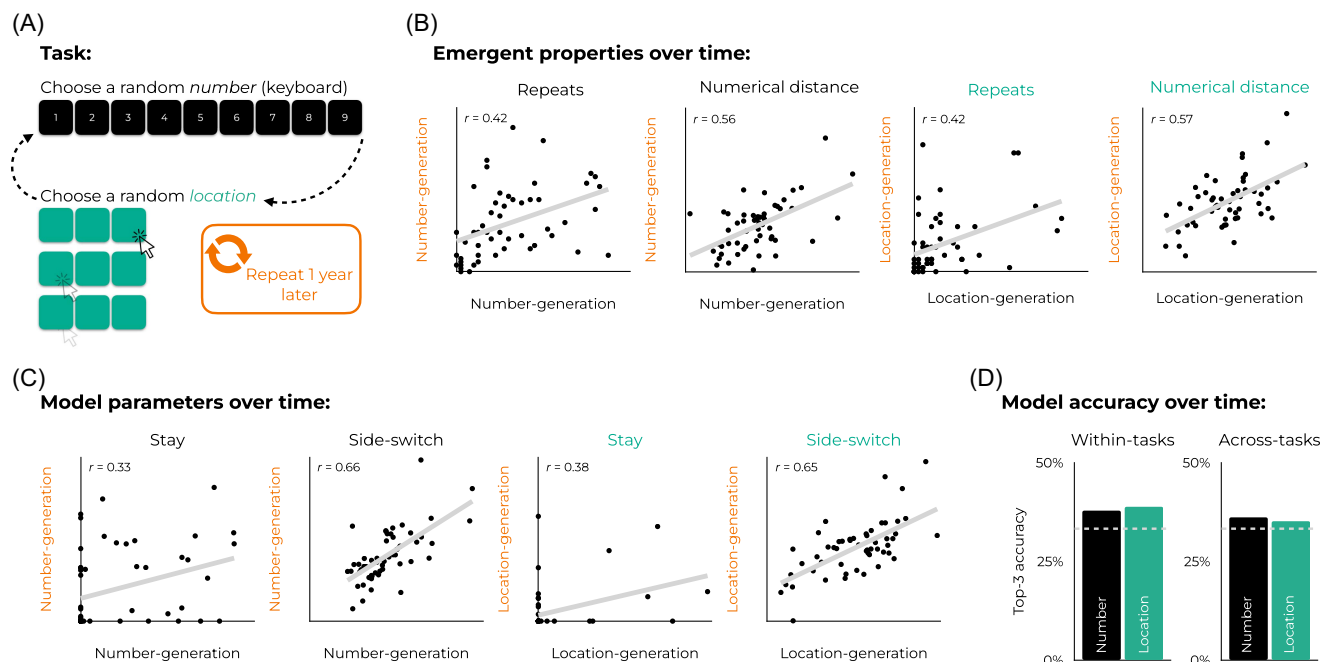
To test this, in Experiment 3, we collected longitudinal data from the same participants that completed Experiment 2; these participants ($N = 53$ after exclusions) generated new sequences of random numbers and two-dimensional random locations approximately 1 year after the initial study (mean of 50 weeks later; Figure 4A).

Remarkably, a participant's sequence generated at the first time point shared features with the sequence they generated 1 year later. This was true of all three emergent properties in both the number-generation task: number of repeats: $r(51) = 0.42$, $p < .01$; average numerical distance: $r(51) = 0.56$, $p < .001$; number of direction switches: $r(51) = 0.51$, $p < .001$; and the two-dimensional location-generation

task: number of repeats: $r(51) = 0.42$, $p < .01$; average numerical distance: $r(51) = 0.57$, $p < .001$; number of direction switches: $r(51) = 0.61$, $p < .001$; Figure 4B.

Subject-specific parameters in our model were also correlated over time, both in the number-generation task: stay parameter: $r(51) = 0.33$, $p = .01$; side-switch parameter: $r(51) = 0.66$, $p < .001$; direction-switch parameter: $r(51) = 0.48$, $p < .001$; and in the location-generation task: stay parameter: $r(51) = 0.38$, $p < .01$; side-switch parameter: $r(51) = 0.65$, $p < .001$; Figure 4C.

Next, as a stronger test of stability across time, we compared our model's simulated data for each task based on model parameters from 1 year earlier to a Markov chain's simulated data based on the more recently collected data. Participant data matched our model's predictions based on parameters from year-old data more closely

Figure 4*Predicting Random Behavior, 1 Year Later*

Note. (A) One year after finishing data collection for Experiment 2, we invited the same subjects to generate two new random sequences (one of numbers and one of two-dimensional locations). This allowed us to compare each subject's sequence to themselves 1 year later in Experiment 3. (B) All emergent properties were highly correlated when comparing a subject's data to their new sequence, collected 1 year later. (C) Model parameters for each subject over time. Notice how highly correlated the side-switch parameter is. Note that the modest correlation in the stay parameter is partly due to the parameter values being *so* consistent that many participants had the minimum parameter value (-1) when we fit the model to their data from both Experiment 2 and Experiment 3. (D) Model accuracy over tasks and time. The model accurately predicted choices within a given task over time (left; "within-tasks"; 37.9% for number-generation, 38.8% for location-generation). The model even predicted choices across both task and time together (right, "across-tasks"; 36.1% for old number-generation data predicting new two-dimensional location-generation data, and 36.6% for the opposite), although it does not significantly outperform the leave-one-out baseline. See the online article for the color version of this figure.

than the current Markov chain in both the number-generation task (where our model outperformed the Markov chain in 66/100 simulations in terms of average numerical distance and in 99/100 simulations in terms of direction switches) and the two-dimensional location-generation task (where our model outperformed the Markov chain in 99/100 simulations in terms of average numerical distance and 93/100 simulations in terms of direction switches).

Despite the above correlations, a subject's current data may contain subtle, choice-level differences that significantly deviate from their initial data. However, this is not what we observed: Our model, when fit to each subject's initial time point data, successfully predicted variance in their next choice 1 year later in both the number-generation task (37.9%, $p < .05$ compared to 36.4% leave-one-out accuracy) and the location-generation task (38.8%, $p < .05$ compared to 35.2% leave-one-out accuracy; Figure 4D). We also examined performance across both tasks and across time, simultaneously. Using subject-specific parameters from the initial number-generation task to predict data in the location-generation task 1 year later yielded 36.1% accuracy; and using the initial location-generation model to predict current number-generation data resulted in 36.6% accuracy. However, neither significantly outperformed the leave-one-out baseline ($p = .18$ using number-generation parameters for predicting location-generation data 1 year

later, $p = .77$ using location-generation parameters for predicting number-generation data 1 year later).

Recall that this task involves generating sequences that are as random as possible; the "right" way to do these tasks is to be unpredictable. Yet, not only can we reliably predict an individual's behavior; we can do so across tasks and over the span of one full year. This suggests that random behavior is highly idiosyncratic and perhaps trait-like.

Discussion

The ability to perceive and generate randomness is a core part of cognition. How we perceive randomness influences how we read financial information, how we understand weather forecasts, and more. Similarly, our ability to act randomly may be vital in a range of game-theoretic scenarios (e.g., if you have an upper hand on your opponent, you may still wish to behave in such a way that does not reveal your enhanced knowledge; see Mazor et al., 2024). Here, we show not only that humans exhibit the same characteristic biases when trying to act randomly in a variety of settings, but also that an individual's specific random behaviors are consistent across tasks and across time. We demonstrate this through a model that captures subject-specific behavioral patterns in different domains (i.e.,

number vs. one- and two-dimensional space) and at different time points (i.e., the span of a year). This focus on *individual* characteristics of random generation presents a shift from how random behavior is classically studied (i.e., on the group level).

In light of our findings, we speculate that there exists a trait-like “randomness generator” in the mind. Even though each person’s so-called generator may exhibit the same overall tendencies (e.g., well-known biases like an overemphasis on state alternations; Kahneman & Tversky, 1972), each one may be tuned in slightly different ways (for instance, because of a unique “diet” of statistical information or a different internal model of what “random” means). Along these lines, we suggest that many classic works on randomness that demonstrate group-level departures from true random behavior may have meaningful individual-level counterparts. For example, Bar-Hillel and Wagenaar (1991), Griffiths and Tenenbaum (2003), and Griffiths et al. (2018) each analyzed features of “subjective” randomness in the mind. While these works provide much insight into human randomness, they analyze group-level heuristics; thus, “subjective” is taken to mean “non-objective.” Here, our approach advances the notion of “subjective” proposed in these works by demonstrating that randomness may be instantiated in the mind in a truly unique fashion for each person (i.e., “subjective” as “individual”). In practice, perhaps this means that we should think about one’s ability to generate randomness in the way that we think about other cognitive or emotional traits—that is, as individual differences that persist across time in tractable ways.

This perspective opens the door to many novel questions. Insofar as random behavior is trait-like, what other trait-like capacities might it be related to? Might random behavior predict more general cognitive capacities—like how executive function or numerical acuity (Halberda et al., 2008), for instance, do? Although random behaviors have been shown to relate to executive function (A. Baddeley et al., 1998; Jahanshahi et al., 2006), we suggest that this relationship may extend beyond mere “performance”: Two individuals who “perform” equally well in our task may still have distinct patterns of responses that our model can detect and differentiate. Thus, our model is most likely capturing more than general differences in executive function (though this is nevertheless an interesting avenue for future work). Additionally, given the importance of randomness to behavior, one might wonder how early in development this putative randomness generator arises (see Towse & McLachlan, 1999), how its development tracks (if at all) with the development of other cognitive abilities (like number acuity), and whether it remains stable throughout the lifespan (Gauvrit et al., 2017). Furthermore, might nonhuman animal minds contain a similar sort of generator (Lee et al., 2004), and, if so, can that behavior be modeled using the same sorts of parameters used here? What are the neural correlates of these processes, and might they reveal further features of an individual’s randomness generator (Jahanshahi & Dimberger, 1998; Jahanshahi et al., 1998)? Finally, randomness generation and randomness perception often show similar cognitive signatures (Nickerson, 2002). But do they in fact rely on the same underlying process? If so, might a trait-like “randomness perceiver” exist in ways similar to the trait-like “randomness generator” we find here?

The simple model presented here is not perfect; we are open to the idea that other variants of our model or other kinds of models altogether may better predict random behavior, perhaps by including additional (or different) parameters. What we wish to

emphasize here is that a simple model combining insights from past research (e.g., tendencies to repeat, exhaust choices, etc., as in Castillo et al., 2024; Towse & Neil, 1998) is able to predict behavior across tasks and time. It is striking, in our view, that our ability to predict behavior across tasks is almost as good as our ability to predict behavior within tasks. This is a strong indication that random behavior can be explained by an individual trait-like generator, rather than task-, context-, or domain-specific generators. Still, any future improvements to our model would enhance our ability to understand subtle differences in individual behavior and the psychological mechanisms that give rise to these differences.

What Does It Mean That Randomness Is Stable Over (Long Periods of) Time?

Previous work has hinted at the idea that random behavior might be stable across tasks (Castillo et al., 2024; Yu et al., 2018). Recent work has even extended this to show that it applies to nonuniform or recently learned distributions (Castillo et al., 2024). More generally, one may expect random behavior to be consistent across tasks insofar as it relates to statistical understanding—which children exhibit early in development (Xu & Garcia, 2008). However, such works have emphasized stability on the group level (e.g., that repetitions are consistent in different modalities), rather than on the individual level. Thus, our approach here adds critical evidence to the idea that random behaviors may be realized by an individual random generator. Furthermore, previous work has not, to our knowledge, tested stability in random behavior over time, as we have here. Thus, our work builds on previous work that hints at such possibilities, and explores these new questions with a novel computational model fit to a unique data set.

But what are the implications of this stability in random behavior across time? Many different explanations can be provided as to how and why random behavior turns out to be stable. One possibility is that human randomness is in fact realized like a biased version of a computer’s pseudorandom generator, with slightly different algorithms for each individual. Another possibility is that random behaviors draw on mechanisms for other stable processes in the mind that are not unique random behaviors. Under this view, random behaviors would be stable, though not necessarily because of a random generator. Rather, they would come to be stable because another part of the mind contributes to random behavior (e.g., intuitions about probability or numerical cognition, executive function ability; A. D. Baddeley, 1966; A. Baddeley et al., 1998). Showing that random behavior indeed stays consistent across time opens many questions along these lines, and answering these questions will ultimately help illuminate how randomness is realized in the mind.

Constraints on Generality

Participants for this study were recruited via the online platform Prolific; thus, we obtained a diverse sample of U.S. adults. We do not assume these findings generalize beyond this group.

One additional consideration is whether (or to what extent) the patterns of randomness observed here reflect biases of random behavior itself as opposed to misconceptions of what it means to be random. For instance, we take “random” to mean “unpredictable,” but perhaps participants take “random” to mean “uniform.” Had we run a slightly different version of the task in which participants were

incentivized to be unpredictable, we may have observed different biases. Note that similar biases seem to arise in strategic games that do not rely on a definition of randomness, e.g., rock-paper-scissors; Batzilis et al., 2019. However, various strategic considerations might lead to more or less random behavior, in the case of mimicry, e.g., Baker & Rachlin, 2001; Belot et al., 2013. This strikes us as a valuable question for future work. The modeling approach taken here offers a method for comparing subtle individual differences in random behavior and allows such work to test whether changes in information that individuals have might lead to systematic changes in model parameters.

Additionally, the model here could surely be improved. For example, models with a “cycling parameter,” which captures the increasing likelihood of selecting an option based on how long ago it was last selected, may prove more accurate (and more generalizable) for these tasks (see Angelike & Musch, 2024). As it stands, the current work serves as an existence proof that there are stable individual differences in random behavior that can be modeled across tasks and time.

Conclusion

People’s ability to perceive and generate random sequences reveals deep insights into human cognition. Here, using a simple computational model, we demonstrate that highly idiosyncratic aspects of human randomization are consistent across tasks and across time, suggesting that a stable, trait-like random generator might exist in the mind. This work opens the door to a wide range of new questions about the nature of this random generator in the mind.

References

- Angelike, T., & Musch, J. (2024). *An improved modeling approach to investigate biases in human random number generation*. Open Science Framework.
- Arkes, H. R., Faust, D., Guilmette, T. J., & Hart, K. (1988). Eliminating the hindsight bias. *Journal of Applied Psychology*, 73(2), 305–307. <https://doi.org/10.1037/0021-9010.73.2.305>
- 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. <https://doi.org/10.1080/713755788>
- Baddeley, A. D. (1966). The capacity for generating information by randomization. *The Quarterly Journal of Experimental Psychology*, 18(2), 119–129. <https://doi.org/10.1080/14640746608400019>
- Baker, F., & Rachlin, H. (2001). Probability of reciprocation in repeated prisoner’s dilemma games. *Journal of Behavioral Decision Making*, 14(1), 51–67. [https://doi.org/10.1002/1099-0771\(200101\)14:1<51::AID-BDM365>3.0.CO;2-K](https://doi.org/10.1002/1099-0771(200101)14:1<51::AID-BDM365>3.0.CO;2-K)
- Bar-Hillel, M., & Wagenaar, W. A. (1991). The perception of randomness. *Advances in Applied Mathematics*, 12(4), 428–454. [https://doi.org/10.1016/0196-8858\(91\)90029-I](https://doi.org/10.1016/0196-8858(91)90029-I)
- Batzilis, D., Jaffe, S., Levitt, S., List, J. A., & Picel, J. (2019). Behavior in strategic settings: Evidence from a million rock-paper-scissors games. *Games*, 10(2), Article 18. <https://doi.org/10.3390/g10020018>
- Belot, M., Crawford, V. P., & Heyes, C. (2013). Players of Matching Pennies automatically imitate opponents’ gestures against strong incentives. *Proceedings of the National Academy of Sciences of the United States of America*, 110(8), 2763–2768. <https://doi.org/10.1073/pnas.1209981110>
- Biesaga, M., & Nowak, A. (2024). The role of the working memory storage component in a random-like series generation. *PLOS ONE*, 19(1), Article e0296731. <https://doi.org/10.1371/journal.pone.0296731>
- Castillo, L., León-Villagrà, P., Chater, N., & Sanborn, A. (2024). Explaining the flaws in human random generation as local sampling with momentum. *PLOS Computational Biology*, 20(1), Article e1011739. <https://doi.org/10.1371/journal.pcbi.1011739>
- Clotfelter, C. T., & Cook, P. J. (1993). The “gambler’s fallacy” in lottery play. *Management Science*, 39(12), 1521–1525. <https://doi.org/10.1287/mnsc.39.12.1521>
- Cooper, R. P. (2016). Executive functions and the generation of “random” sequential responses: A computational account. *Journal of Mathematical Psychology*, 73, 153–168. <https://doi.org/10.1016/j.jmp.2016.06.002>
- Darling, S., Allen, R. J., & Havelka, J. (2017). Visuospatial bootstrapping: When visuospatial and verbal memory work together. *Current Directions in Psychological Science*, 26(1), 3–9. <https://doi.org/10.1177/0963721416665342>
- Darling, S., & Havelka, J. (2010). Visuospatial bootstrapping: Evidence for binding of verbal and spatial information in working memory. *Quarterly Journal of Experimental Psychology*, 63(2), 239–245. <https://doi.org/10.1080/17470210903348605>
- Dehaene, S., Bossini, S., & Giraux, P. (1993). The mental representation of parity and number magnitude. *Journal of Experimental Psychology: General*, 122(3), 371–396. <https://doi.org/10.1037/0096-3445.122.3.371>
- Driver, P. M., & Humphries, D. A. (1988). *Protean behaviour: The biology of unpredictability*. Clarendon Press.
- Falk, R., & Konold, C. (1997). Making sense of randomness: Implicit encoding as a basis for judgment. *Psychological Review*, 104(2), 301–318. <https://doi.org/10.1037/0033-295X.104.2.301>
- Feng, G. W., & Rutledge, R. B. (2024). Surprising sounds influence risky decision making. *Nature Communications*, 15(1), Article 8027. <https://doi.org/10.1038/s41467-024-51729-4>
- Gauvrit, N., Zenil, H., Soler-Toscano, F., Delahaye, J. P., & Brugger, P. (2017). Human behavioral complexity peaks at age 25. *PLOS Computational Biology*, 13(4), Article e1005408. <https://doi.org/10.1371/journal.pcbi.1005408>
- Gershman, S. J. (2018). Deconstructing the human algorithms for exploration. *Cognition*, 173, 34–42. <https://doi.org/10.1016/j.cognition.2017.12.014>
- Gershman, S. J. (2019). Uncertainty and exploration. *Decision*, 6(3), 277–286. <https://doi.org/10.1037/dec0000101>
- Gilovich, T., Vallone, R., & Tversky, A. (1985). The hot hand in basketball: On the misperception of random sequences. *Cognitive Psychology*, 17(3), 295–314. [https://doi.org/10.1016/0010-0285\(85\)90010-6](https://doi.org/10.1016/0010-0285(85)90010-6)
- Ginsburg, N., & Karpiuk, P. (1994). Random generation: Analysis of the responses. *Perceptual and Motor Skills*, 79(3), 1059–1067. <https://doi.org/10.2466/pms.1994.79.3.1059>
- Griffiths, T. L., Daniels, D., Austerweil, J. L., & Tenenbaum, J. B. (2018). Subjective randomness as statistical inference. *Cognitive Psychology*, 103, 85–109. <https://doi.org/10.1016/j.cogpsych.2018.02.003>
- Griffiths, T. L., & Tenenbaum, J. B. (2001). Randomness and coincidences: Reconciling intuition and probability theory. *Proceedings of the 23rd annual conference of the cognitive science society*, 370–375.
- Griffiths, T. L., & Tenenbaum, J. B. (2003). Probability, algorithmic complexity, and subjective randomness. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 25(25), 480–485.
- Halberda, J., Mazocco, M. M., & Feigenson, L. (2008). Individual differences in non-verbal number acuity correlate with maths achievement. *Nature*, 455(7213), 665–668. <https://doi.org/10.1038/nature07246>
- Jahanshahi, M., & Dirnberger, G. (1998). The left dorsolateral prefrontal cortex and random generation of responses: Studies with transcranial magnetic stimulation. *Neuropsychologia*, 37(2), 181–190. [https://doi.org/10.1016/S0028-3932\(98\)00092-X](https://doi.org/10.1016/S0028-3932(98)00092-X)
- Jahanshahi, M., Profice, P., Brown, R. G., Ridding, M. C., Dirnberger, G., & Rothwell, J. C. (1998). The effects of transcranial magnetic stimulation over the dorsolateral prefrontal cortex on suppression of habitual counting during random number generation. *Brain: A Journal of Neurology*, 121(8), 1533–1544. <https://doi.org/10.1093/brain/121.8.1533>

- Jahanshahi, M., Saleem, T., Ho, A. K., Dirnberger, G., & Fuller, R. (2006). Random number generation as an index of controlled processing. *Neuropsychology*, 20(4), 391–399. <https://doi.org/10.1037/0894-4105.20.4.391>
- Kahneman, D., & Tversky, A. (1972). Subjective probability: A judgment of representativeness. *Cognitive Psychology*, 3(3), 430–454. [https://doi.org/10.1016/0010-0285\(72\)90016-3](https://doi.org/10.1016/0010-0285(72)90016-3)
- Kareev, Y. (1992). Not that bad after all: Generation of random sequences. *Journal of Experimental Psychology: Human Perception and Performance*, 18(4), 1189–1194. <https://doi.org/10.1037/0096-1523.18.4.1189>
- Kelly, M. H., & Martin, S. (1994). Domain-general abilities applied to domain-specific tasks: Sensitivity to probabilities in perception, cognition, and language. *Lingua*, 92, 105–140. [https://doi.org/10.1016/0024-3841\(94\)90339-5](https://doi.org/10.1016/0024-3841(94)90339-5)
- Lee, D., Conroy, M. L., McGreevy, B. P., & Barraclough, D. J. (2004). Reinforcement learning and decision making in monkeys during a competitive game. *Cognitive Brain Research*, 22(1), 45–58. <https://doi.org/10.1016/j.cogbrainres.2004.07.007>
- Mazor, M., Firestone, C., & Phillips, I. (2024). *Pretending not to know reveals a powerful capacity for self-simulation*. PsyArXiv. https://osf.io/preprints/psyarxiv/pgxrz_v1
- Misirilsoy, E., & Haggard, P. (2014). Asymmetric predictability and cognitive competition in football penalty shootouts. *Current Biology*, 24(16), 1918–1922. <https://doi.org/10.1016/j.cub.2014.07.013>
- Moore, T. Y., Cooper, K. L., Biewener, A. A., & Vasudevan, R. (2017). Unpredictability of escape trajectory explains predator evasion ability and microhabitat preference of desert rodents. *Nature Communications*, 8(1), Article 440. <https://doi.org/10.1038/s41467-017-00373-2>
- Neuringer, A. (1986). Can people behave “randomly”? The role of feedback. *Journal of Experimental Psychology: General*, 115(1), 62–75. <https://doi.org/10.1037/0096-3445.115.1.62>
- Nickerson, R. S. (2002). The production and perception of randomness. *Psychological Review*, 109(2), 330–357. <https://doi.org/10.1037/0033-295X.109.2.330>
- Peer, E., Brandimarte, L., Samat, S., & Acquisti, A. (2017). Beyond the turk: Alternative platforms for crowdsourcing behavioral research. *Journal of Experimental Social Psychology*, 70, 153–163. <https://doi.org/10.1016/j.jesp.2017.01.006>
- Poropat, A. E. (2009). A meta-analysis of the five-factor model of personality and academic performance. *Psychological Bulletin*, 135(2), 322–338. <https://doi.org/10.1037/a0014996>
- Rapoport, A., & Budescu, D. V. (1992). Generation of random series in two-person strictly competitive games. *Journal of Experimental Psychology: General*, 121(3), 352–363. <https://doi.org/10.1037/0096-3445.121.3.352>
- Roese, N. J., & Vohs, K. D. (2012). Hindsight bias. *Perspectives on Psychological Science*, 7(5), 411–426. <https://doi.org/10.1177/1745691612454303>
- Rutledge, R. B., Lazzaro, S. C., Lau, B., Myers, C. E., Gluck, M. A., & Glimcher, P. W. (2009). Dopaminergic drugs modulate learning rates and perseveration in Parkinson’s patients in a dynamic foraging task. *The Journal of Neuroscience*, 29(48), 15104–15114. <https://doi.org/10.1523/JNEUROSCI.3524-09.2009>
- Saffran, J. R., Aslin, R. N., & Newport, E. L. (1996). Statistical learning by 8-month-old infants. *Science*, 274(5294), 1926–1928. <https://doi.org/10.1126/science.274.5294.1926>
- Saffran, J. R., Johnson, E. K., Aslin, R. N., & Newport, E. L. (1999). Statistical learning of tone sequences by human infants and adults. *Cognition*, 70(1), 27–52. [https://doi.org/10.1016/S0010-0277\(98\)00075-4](https://doi.org/10.1016/S0010-0277(98)00075-4)
- Schulz, M. A., Schmalbach, B., Brugger, P., & Witt, K. (2012). Analysing humanly generated random number sequences: A pattern-based approach. *PLOS ONE*, 7(7), Article e41531. <https://doi.org/10.1371/journal.pone.0041531>
- Sherman, B. E., Graves, K. N., & Turk-Browne, N. B. (2020). The prevalence and importance of statistical learning in human cognition and behavior. *Current Opinion in Behavioral Sciences*, 32, 15–20. <https://doi.org/10.1016/j.cobeha.2020.01.015>
- Sherman, B. E., & Turk-Browne, N. B. (2020). Statistical prediction of the future impairs episodic encoding of the present. *Proceedings of the National Academy of Sciences of the United States of America*, 117(37), 22760–22770. <https://doi.org/10.1073/pnas.2013291117>
- Sherman, B. E., Yousif, S. R., Reiner, C., & Keil, F. (2022). The speed of statistical perception. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 44(44), 2738–2744.
- Szopa-Comley, A. W., & Ioannou, C. C. (2022). Responsive robotic prey reveal how predators adapt to predictability in escape tactics. *Proceedings of the National Academy of Sciences of the United States of America*, 119(23), Article e2117858119. <https://doi.org/10.1073/pnas.2117858119>
- Tomov, M. S., Truong, V. Q., Hundia, R. A., & Gershman, S. J. (2020). Dissociable neural correlates of uncertainty underlie different exploration strategies. *Nature Communications*, 11(1), Article 2371. <https://doi.org/10.1038/s41467-020-15766-z>
- Towse, J. N., & McLachlan, A. (1999). An exploration of random generation among children. *British Journal of Developmental Psychology*, 17(3), 363–380. <https://doi.org/10.1348/026151099165348>
- 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. <https://doi.org/10.3758/BF03209475>
- Treisman, M., & Faulkner, A. (1987). Generation of random sequences by human subjects: Cognitive operations or psychological process? *Journal of Experimental Psychology: General*, 116(4), 337–355. <https://doi.org/10.1037/0096-3445.116.4.337>
- Turk-Browne, N. B., Jungé, J., & Scholl, B. J. (2005). The automaticity of visual statistical learning. *Journal of Experimental Psychology: General*, 134(4), 552–564. <https://doi.org/10.1037/0096-3445.134.4.552>
- Wagenaar, W. A. (1972). Generation of random sequences by human subjects: A critical survey of literature. *Psychological Bulletin*, 77(1), 65–72. <https://doi.org/10.1037/h0032060>
- Wagner, A. R., & Rescorla, R. A. (1972). *Inhibition in Pavlovian conditioning: Application of a theory*. Inhibition and Learning.
- Walker, M., & Wooders, J. (2001). Minimax play at Wimbledon. *The American Economic Review*, 91(5), 1521–1538. <https://doi.org/10.1257/aer.91.5.1521>
- Warren, P. A., Gostoli, U., Farmer, G. D., El-Dereby, W., & Hahn, U. (2018). A re-examination of “bias” in human randomness perception. *Journal of Experimental Psychology: Human Perception and Performance*, 44(5), 663–680. <https://doi.org/10.1037/xhp0000462>
- Xu, F., & Garcia, V. (2008). Intuitive statistics by 8-month-old infants. *Proceedings of the National Academy of Sciences of the United States of America*, 105(13), 5012–5015. <https://doi.org/10.1073/pnas.0704450105>
- Yu, R. Q., Gunn, J., Osherson, D., & Zhao, J. (2018). The consistency of the subjective concept of randomness. *Quarterly Journal of Experimental Psychology: Human Experimental Psychology*, 71(4), 906–916. <https://doi.org/10.1080/17470218.2017.1307428>

Received July 24, 2024

Revision received February 4, 2025

Accepted February 14, 2025 ■