

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

The Role of Holes in Whack-a-Mole: Investigating Procedural Learning with Graded Statistical Structures Across State Space Sizes

Permalink

<https://escholarship.org/uc/item/31w516nk>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Chang, Cherrie

Li, Tianyi

Hartshorne, Joshua K

Publication Date

2026

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

The Role of Holes in Whack-a-Mole: Investigating Procedural Learning with Graded Statistical Structures Across State Space Sizes

Cherrie Chang (cherriechang15@gmail.com)¹ & Tianyi Li² & Joshua Hartshorne¹

¹Department of Communication Sciences & Disorders, MGH Institute of Health Professions

²Department of Psychology & Neuroscience, Boston College

Abstract

Procedural learning—the implicit acquisition of motor and cognitive skills through repeated practice—has been proposed to support abilities from typing to language acquisition. Traditional serial reaction time paradigms study procedural learning using stimulus sequences with simple underlying statistical structures: deterministic or with binary probability levels, involving only 4 possible states. We introduce a novel paradigm where participants learn sequences with graded probability levels, generated from transition matrices with a continuous gradient of transition probabilities. We found robust effects of both surprisal (transition probability from previous state) and entropy (uncertainty of next state) on reaction times across 5 state space sizes (4–8), with no interaction between surprisal and entropy, indicating learners tracked specific transition probabilities even at high-uncertainty states. These findings demonstrate that procedural learning is sensitive to fine-grained statistical structure and scales beyond traditional 4-state paradigms, supporting theoretical proposals linking procedural learning mechanisms to skill acquisition in complex domains like language.

Keywords: procedural learning; implicit learning; statistical learning; procedural memory; surprisal; entropy

Introductory

How often have you typed “form” instead of “from”? Or ended a word with “-tino” instead of “-tion”, perhaps even in an embarrassingly glaring context like a conference paper? The typos we make are not random; they reflect the statistical patterns in the language we use. Experienced typists learn that some keystroke sequences are highly probable, causing their fingers to automatically perform ones like “for” mistakenly, or perform them so quickly they slip up, like in the case of “-tino” (Grudin, 1983; Pinet et al., 2016; Rumelhart & Norman, 1982). Yet typists cannot consciously report these probabilities or articulate their finger movements at the speed of execution—this knowledge is expressed in the automatic motor performance of typing, not in conscious awareness (Crump & Logan, 2010). This implicit form of learning, where statistical patterns are unconsciously acquired and expressed as behavior rather than declared knowledge, is called procedural learning (Nissen & Bullemer, 1987; Squire, 2004).

Besides typing, procedural learning underlies skill acquisition across many domains, from proprioceptive skills like riding a bike to perceptual-cognitive skills like reading mirrored text (Kassubek et al., 2001). More recently, procedural learning has been proposed to support increasingly diverse and sophisticated abilities, notably grammar acquisition in

language (Kidd & Arciuli, 2016; Lammertink et al., 2017; Lum et al., 2014; Ullman & Pierpont, 2005). These proposals invite closer scrutiny of its scalability: can the procedural learning system handle the large-scale, complex statistical structures that characterize domains like language? Consider a suggestive historical case: while QWERTY keyboards (26 letter keys) enabled fluent touch typing through procedural learning, Chinese typewriters (2,000+ character keys) never achieved comparable automaticity and were eventually abandoned (for QWERTY keyboards!) (Mullaney, 2017). If the Chinese typewriter wasn’t adopted because it required typists to learn complex motor sequences involving many keys, then this failure raises a fundamental question: does procedural learning fail beyond a certain level of statistical complexity?

Traditional paradigms for studying procedural learning are ill-suited to answering this question. The primary method for measuring procedural learning has been the Serial Reaction Time Task (SRTT) (Nissen & Bullemer, 1987). In the SRTT, participants respond to stimuli appearing in one of four positions on screen using corresponding keystrokes. Unbeknownst to them, the sequence of stimulus positions follows a deterministic pattern. Over time, reaction times decrease, demonstrating procedural learning of the keystroke sequence in correspondence to the stimulus positions’ deterministic pattern. While the SRTT provided early evidence for procedural learning, it suffers from a critical confound: participants often develop explicit awareness of the sequence, making it difficult to isolate implicit procedural learning effects (Lustig, 2022; Song et al., 2008), resulting in low test-retest reliability of the task (Oliveira et al., 2023; West et al., 2018).

The Alternating Serial Reaction Time Task (ASRT) addresses these issues by introducing a probabilistic design. In the ASRT, deterministic (d) and random (r) trials alternate (e.g., d-r-d-r), creating sequences with binary probability levels for transitioning from one position to the next: high (e.g. $p_{d \rightarrow r \rightarrow d}$) vs. low (e.g. $p_{r \rightarrow d \rightarrow r}$) (Howard Jr. & Howard, 1997). This design minimizes explicit awareness and dramatically improves reliability (Farkas et al., 2024; Oliveira et al., 2023). However, a key limitation of ASRT is that its bimodal distribution of transition probabilities represents an overtly simplistic statistical structure, lacking the continuous probability distributions found in real-world domains like language. Finally, in both SRTT and ASRT, experiment designs have mostly tested sequences with 4 positions (Hedenius et al., 2011; Janacek et al., 2012), creating an sandbox domain with only 4 states, despite

real-world domains where procedural learning is implicated often involving far larger state spaces. Whether procedural learning can track continuous statistical structures, and whether it scales to larger state spaces, remains untested.

Here, we introduce a new paradigm that retains ASRT’s probabilistic design but introduces a graded statistical structure by having a smooth continuum of transition probabilities between positions. We construct transition matrices to generate stimulus position sequences where positions differ in their transition probability distributions, ranging from deterministic to uniformly distributed. We also vary the number of positions (4, 5, 6, 7, 8) to test whether procedural learning scales with state space size. This allows us to ask: (1) Are learners sensitive to fine-grained differences in transition probabilities, beyond simple high vs. low probability contrasts? (2) Do learning trajectories change as the number of states increases?

Hypotheses

We make the following predictions: **(H1) Surprisal Effect:** If participants are learning the statistical structure underlying the sequence and not merely practicing motor responses, trial-level reaction times (RTs) will be predicted by surprisal—the negative log probability of the observed transition—such that higher-surprisal (less predictable) transitions elicit slower responses. **(H2) Previous Position Entropy Effect:** Previous position entropy—the Shannon entropy (Shannon, 1948) of the transition probability distribution from the previous position—will modulate performance similarly, with higher-entropy (more uncertain) previous positions leading to slower RTs. **(H3) Surprisal $\times$ Entropy Interaction:** We further predict that surprisal effects will be stronger following low-entropy positions, where predictions are more confident and prediction errors more costly (surprisal $\times$ entropy interaction). **(H4) State Space Size Effect:** These effects will be attenuated as the number of positions increases, reflecting greater difficulty in learning more complex statistical structures over more states.

Method

Participants

A simulation-based power analysis on pilot data ($N=11$) indicated ~80 total participants would achieve 80% power to detect a learning effect. We recruited 251 (133 female, 118 male) total participants as a conservative estimate¹. Participants were recruited via Prolific, screened for >95% approval rating from previous studies, and limited to using desktops/laptops. Participants were instructed to leave the experiment if not using a physical QWERTY keyboard, though this was not software-enforced. Participants ranged in age from 18 to 75 years ($\mu=35.6$, $\sigma=11.9$), representing 32 nationalities and 20 primary languages. Participants were compensated at \$12.00/hr (median completion time = ~20-40 mins). All participants provided informed consent prior to the experiment. The study was approved by the Institutional Review Board at the Mass General Hospital Institute of Health Professions.

¹We accidentally recruited one extra participant.

Experiment Design

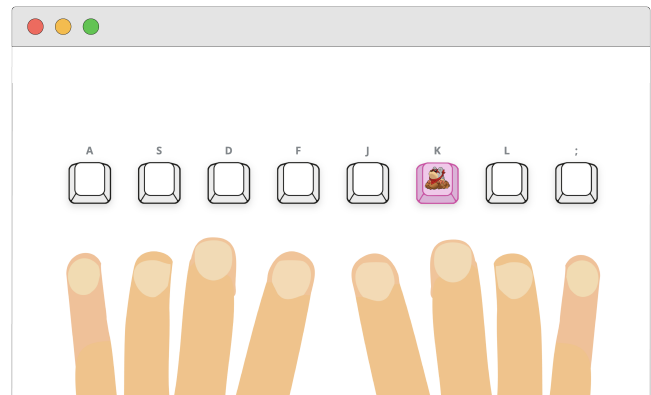


Figure 1: An example trial screen in the 8-position condition.

Like other SRTT designs, our experiment tasks participants to respond to a stimulus appearing in one of several evenly-spaced positions on screen as quickly and accurately as possible by pressing a corresponding key. In our implementation, this stimulus is a mole garbed in a bright red bib and matching sunglasses (Figure 1). Participants are evenly divided into 5 groups of 50 per condition, differing in the number of positions (n ; 4, 5, 6, 7, 8) the mole can appear in².

A participant is first given a text-based tutorial accompanied by an animated demonstration on which key to press in response to each position, then instructed to work through a set of $2n$ practice trials, such that each position is visited twice in random order. After the practice phase, the participant moves on to complete 20 blocks of trials, each separated with a self-paced break. Each block consists of 10 trials per position ($10n$). This design choice results in longer blocks for conditions with more positions, but ensures each position is visited an equal number of times on average across conditions. In both practice and main trials, participants are given feedback on correctness via a pop-up short message (a checkmark vs. “Try again!” + an error tone), and allowed to retry each failed trial until they respond correctly. Following standard SRTT protocol, there is also a 120ms response-stimulus interval (RSI) between trials Howard et al. (2004). We found during piloting that allowing retries, combined with the 120ms RSI, prevented participants from making compensatory errors, where they unintentionally press the wrong key in the next trial in an attempt to correct their current trial (Rabbitt, 1966). We record response time, keyboard response and correctness for each trial. After each block, participants are given adaptive feedback based on their accuracy and speed. At the end of the experiment, participants complete a brief questionnaire probing their explicit awareness of any patterns in the mole’s appearances (Table 1).

In addition to our dapper mole, several design choices are made to facilitate participant understanding and engagement. During the practice phase, each on-screen position is labelled with its corresponding key character (e.g. “A”, “S”, “D”) to help

²One extra participant in the 5-position condition.

Table 1: Post-task questionnaire assessing explicit awareness.

Item	Question
Q1	Did you notice anything special about the task? <i>Choices: Yes, No</i>
Q1b	If yes: What did you notice? (<i>free text</i>)
Q2	Did you notice any regularity in where the mole appeared? <i>Choices: Yes, No</i>
Q2b	If yes: Can you describe the regularity? (<i>free text</i>)
Q3	Did you use any strategy to help you respond faster? <i>Choices: Yes, No</i>
Q3b	If yes: Can you describe the strategy? (<i>free text</i>)
Q4	There WAS a regularity in the sequence. What do you think it was? (<i>free text</i>)
Q4b	How confident are you in your answer? <i>Scale: 1 (Not at all) to 5 (Very confident)</i>

familiarize participants with the keyboard mappings. These labels disappear in the main trials to prevent explicit sequence encoding via visual labels, but reappear during retry of failed trials. Key mappings followed conventional QWERTY finger layouts to minimize cognitive load (Table 2). However, piloting revealed that the wider gap between keys “F” and “J” introduced a discrepancy between the evenly-spaced on-screen positions and the unevenly-spaced keyboard keys, causing more errors in the middle positions, particularly at larger n . Shifting either hand’s placement to close the extra gap did not resolve this, as the positions were still corresponding half to one hand and half to the other. After testing different visual aids, we settled on adding two hand diagrams to the bottom of the screen, each finger aligned with its target position (Figure 1). We received feedback that this helped participants better map the spatial layout of the on-screen positions to their corresponding keys. Finally, on-screen positions were rendered as 3D keyboard keys that light up and visibly depress on keypress, reinforcing the diegetic link between screen and keyboard to encourage motor rather than visual-spatial encoding.

Transition Matrices

Our experiment deviates from other SRTT designs in how the position sequence and its probabilistic structure is constructed. In each experiment run, we use a first-order Markov process to generate the sequence of positions the mole appears in. This process is defined by a transition matrix, where each entry in the matrix defines the probability of transitioning from one position to another. Following the five conditions in our experiment, we designed five “original” transition matrices of sizes 4x4, 5x5, 6x6, 7x7, and 8x8 to describe the transition probabilities for each number of positions. These matrices were constructed to have graded and increasing levels of entropy across rows, such that position 1 has lowest entropy—having the most predictable next-position-transitions—followed by the second; with the last position, position n , having highest entropy. For example, in the 4x4 matrix (Figure 2), position 1 has a high

Table 2: Keyboard keys mapped to each on-screen position from left to right for different number of positions.

n positions	keys (left to right)
4	{D, F, J, K}
5	{S, D, F, J, K}
6	{S, D, F, J, K, L}
7	{A, S, D, F, J, K, L}
8	{A, S, D, F, J, K, L, ;}

probability of transitioning to position 2 ($p_{1 \rightarrow 2} = 0.97$), a low probability of transitioning to position 1 ($p_{1 \rightarrow 1} = 0.02$), and never transitions to positions 3 or 4; while position 4 has a more uniform distribution of transition probabilities to all other positions ($p_{4 \rightarrow 1} = 0.38$, $p_{4 \rightarrow 2} = 0.02$, $p_{4 \rightarrow 3} = 0.23$, $p_{4 \rightarrow 4} = 0.34$). This gives rise to a gradient of entropy values across the positions, resulting in a more continuously distributed statistical structure underlying the sequence of positions the stimulus mole appears in, allowing us to investigate how the procedural learning system differentially picks up varying levels of probability in a given statistical structure.

Several design criteria/constraints were imposed on the construction of these original transition matrices. First, these matrices had to be doubly stochastic to ensure uniform visitation to each position over time, preventing confounding effects from position frequency on learning. Second, self-transitions (e.g. $1 \rightarrow 1$) should have near-zero probabilities to reduce data loss, as data from repeated trials have been historically thrown out from analysis in SRTT tasks to prevent confounding motor effects from repeating the same keystroke (Janacek et al., 2012; Oliveira et al., 2023). Third, the graded entropy levels across the positions should translate to graded probabilities across all pair-wise (bigram) transitions as well, as having more fine-grained, continuous levels of transition probabilities is

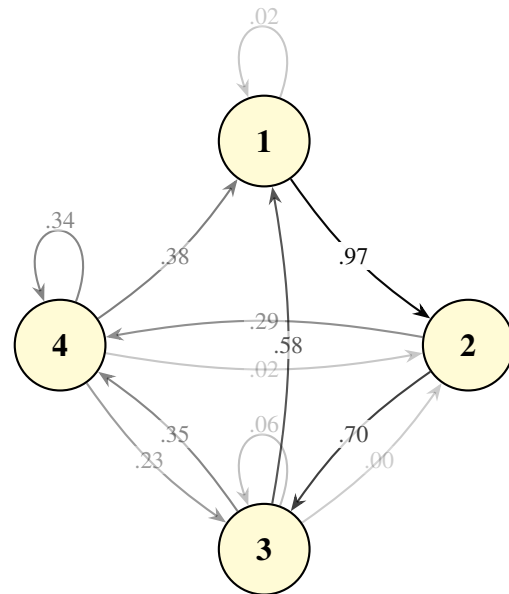


Figure 2: State transition graph of the 4x4 transition matrix. Nodes = positions; edges = transition probabilities.

our experiment's key distinction from other SRTT tasks, such as the ASRT paradigm, which only has 2 probability levels³. Fourth, the increase in entropy across positions (i.e. rows) should be significant between one another and as linear as possible to ensure a smooth gradient with meaningfully distinct levels in predictability across positions.

Strictly satisfying all these constraints simultaneously proved mathematically impossible, so we constructed each matrix via numerical optimization. The search space consisted of all $n \times n$ non-negative matrices for each condition. Row sums equal to 1 were enforced as hard constraints; column sums equal to 1 (double stochasticity) were enforced as approximate equality constraints with tolerance 10^{-4} , since exact double stochasticity proved incompatible with the entropy gradient objective. The remaining design goals were encoded as soft penalty terms in an objective function: (1) squared distance between the sorted per-row conditional entropies and a linearly-spaced target profile spanning 0 to $\log_2(n)$ bits; (2) a spread penalty encouraging the entropy range (max – min) to approach its theoretical maximum; (3) variance of successive entropy gaps, penalizing non-linearity; and (4) a distinctness penalty for entropy pairs within 0.1 bits of each other. The weighted sum of these terms was minimized using Sequential Least Squares Programming (`scipy.optimize.minimize`) with 3–5 random restarts per matrix size, selecting the solution with the lowest final objective value.

The five original matrices we constructed using this algorithm has these properties: 1) near fully doubly stochastic—all row and column sums to 1 (± 0.0001); 2) approach a uniform stationary distribution after 7 state transitions, where the stationary visitation for each state (position) are within 0.001 of each other; 3) had a graded entropy structure across states—each matrix had a position entropy range of 0.061–0.181 bits for position 1 ($\mu=0.107$, $\sigma=0.048$), and 1.695–2.835 bits for position n ($\mu=2.319$, $\sigma=0.451$). Within a matrix each position increases in entropy from the previous by 0.295–0.699 bits ($\mu=0.442$, $\sigma=0.085$); and 4) avoided self loops as much as possible—all of the matrices had <0.08 probabilities in the left diagonal up until the last two positions. Within each matrix size condition, all participants experience the same underlying entropy gradient across positions, but the mapping between screen positions and entropy levels is randomized through a double permutation procedure that shuffles the original matrix (of that size) using the same random permutation to produce a new matrix of the same size. This preserves the statistical properties of the original matrix—transition probabilities and entropy values—but decouples entropy levels from specific screen positions to prevent confounding from position-specific effects. The position sequence for each participant is subsequently generated using their individual shuffled transition matrix, so the sequence each participant experiences is different both within and between matrix size conditions.

³Transition probabilities are controlled at the trigram level (e.g. $1 \rightarrow 2 \rightarrow 3$) in the ASRT.

Statistical Analysis

Variables & Exclusion Criteria

The main independent variables are surprisal and previous position entropy. Surprisal is defined as the negative log probability of the actual next-position given the current position, derived from the transition matrix used to generate the sequence for that participant. Previous positional entropy is defined as the Shannon entropy of the transition probabilities from the previous position, also derived from the transition matrix. Both surprisal and previous positional entropy are mean-centered within each matrix size condition to facilitate interpretation of main effects and interactions in subsequent models. We also create a binary variable indicating whether the current position is a repetition of the previous position (i.e. self-transition). This variable is included as a covariate in subsequent models to control for potential motor effects from repeating the same keystroke.

The main dependent variables are reaction time (RT)—log-transformed to reduce skewness—and accuracy (proportion of correct responses) on non-practice/main trials. Incorrect trials, trials with RT >1000 ms, and trials with RTs greater than 3 median absolute deviations above the participant's median RT (outliers) are excluded. On the participant level, participants with overall accuracy below 90%, median RT >1000 ms, or with outliers making up more than 20% of total trials are excluded from analysis. We also screened for explicit awareness of the underlying structure by filtering for participants who reported noticing patterns in their questionnaire responses, but found no participant who could describe the structure accurately. After these exclusions, 203 participants remained (80% of original sample): 44 in the 4-position condition, 44 in 5-position, 40 in 6-position, 43 in 7-position, and 32 in 8-position. Lastly, we discarded data from the first trial of each block to account for blocks being implicitly treated as “restarts” and not following an existing probability distribution by participants.

Modeling Approach

We use linear mixed effects regression (LMER) to model the log-transformed RTs. Models are fit using the `lmerTest` package in R. P-values are obtained using the Satterthwaite approximation for degrees of freedom, as implemented in the `lmerTest` package (Kuznetsova et al., 2017).

We first determine the optimal random and fixed effects structure using the final block of trials in each condition, fitting candidate models to each of the five matrix size conditions separately and selecting the structure with the lowest summed Akaike information criterion (AIC; a metric of model fit balancing goodness-of-fit against model complexity) across all conditions (Akaike, 1998). Models are initially fit with the maximal random effects structure supported by the design (Barr et al., 2013): random intercepts for participant and position, and random slopes by participant for surprisal, previous positional entropy, their interaction, and whether the current position is a repetition of the previous position. Where maximal models fail to converge, we use bounded linear mixed models

(blmer/bglmer) as a first remedy. If singularity persists, we iteratively simplify the random effects structure, removing random slopes one at a time and retaining the most complex structure that converges without singularity. Among these converging models, we select the simplest structure without increasing AIC. For fixed effects, starting with the full model including surprisal, previous positional entropy, their interaction, and repetition, we use AIC-based stepwise backward selection. An effect is retained if removing it increases AIC.

Using this selected structure, we then fit separate models to each individual block within each matrix size condition to examine learning trajectories. Finally, we extract the effect size estimates from each model and test whether they are moderated by matrix size and block using linear regression without random effects, as the data are aggregated coefficients.

Results

As a preliminary sanity check, we confirmed that RTs decreased across blocks ($\beta = -0.01$, $t = -11.72$, $p < .001$, 95% CI [-0.007, -0.005]), establishing that performance improvement occurred.

Surprisal & Previous Position Entropy (H1, H2, H3)

Following the model selection procedure described above, the final converged model included by-participant random intercepts and slopes for previous position entropy, random intercepts for position, and fixed effects for surprisal, previous position entropy, their interaction, and repetition. In the final blocks, all predicted effects were statistically reliable across conditions, with significance in at least 4 out of 5 conditions for all fixed effects (Table 3). Higher-surprisal transitions elicited slower responses (H1), as did transitions following higher-entropy positions (H2). The negative interaction indicated that surprisal effects were stronger following low-entropy (i.e., more predictable) positions, consistent with prediction-based learning (H3). Repetition trials showed faster responses, likely reflecting motor facilitation. This model structure converged without singular fits across all 100 models (20 blocks $\times$ 5 conditions) and so was used for all subsequent analyses.

State Space Size (H4)

To examine learning trajectories, we fit this selected model structure separately to each block within each matrix size condition, yielding 100 models total (20 blocks $\times$ 5 conditions). To test whether the effects of surprisal, previous position entropy, and their interaction were moderated by matrix size and block, we first extracted the fixed effect estimates from the 100 models, then fit linear models predicting each effect size as a function of matrix size, block, and their interaction. This approach allows us to assess how sensitivity to surprisal and entropy emerges and changes over the course of learning, and whether these trajectories differ across state space sizes.

The results (Figure 3) shows an increasing trend for surprisal effect on learning over blocks ($\beta = 0.00$, $p < 0.001$), with smaller matrix sizes having a significant larger effect compared to larger matrix sizes ($\beta = -0.01$, $p < 0.001$; H4).

Table 3: Fixed effects from final-block models across all matrix size conditions

Predictor	Mean β	Range	Prop. Sig.
Surprisal	0.034	[0.018, 0.046]	5/5***
Previous Entropy	0.056	[0.027, 0.103]	4/5
Surprisal $\times$ Entropy	-0.016	[-0.033, -0.001]	2/5
Repetition	-0.179	[-0.278, -0.106]	5/5***

Note. Mean β averaged across 5 matrix size conditions.

Prop. Sig. = conditions with $p < .05$; *** = all conditions $p < .001$.

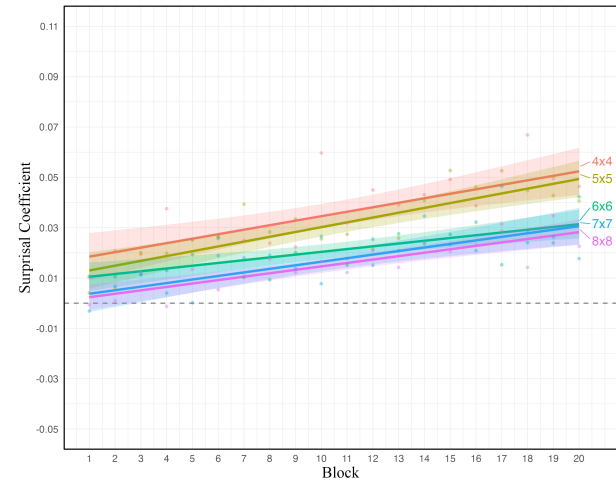


Figure 3: Surprisal effect across learning blocks per condition.

Similarly, as (Figure 4) shows, previous position entropy has an increasing trend over blocks ($\beta = 0.00$, $p < 0.001$), with smaller matrix sizes showing a larger effect ($\beta = -0.01$, $p < 0.001$; H4). The interaction between matrix size and block is also significant for previous position entropy effects, indicating that the increase in previous position entropy effect over blocks is more pronounced for smaller matrix sizes ($\beta = -0.00$, $p < 0.01$).

The interaction effect between surprisal and previous position entropy (Figure 5) shows a decreasing trend over blocks ($\beta = -0.00$, $p < 0.01$), indicating that the modulation of surprisal

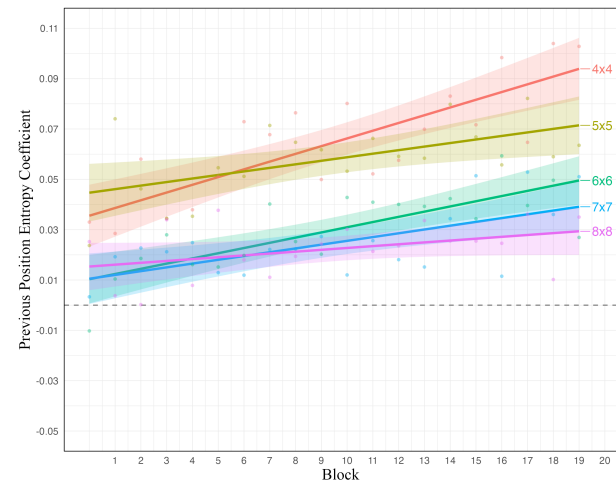


Figure 4: Previous position entropy effect across learning blocks per condition.

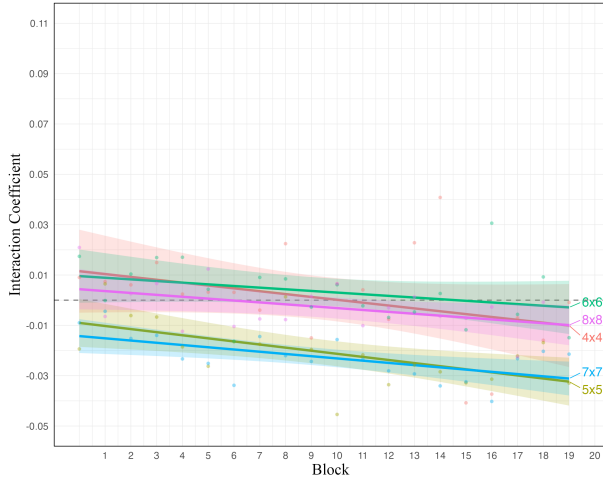


Figure 5: Surprisal x previous position entropy effect across learning blocks per condition.

effects by previous position entropy becomes less pronounced with learning. However, neither matrix size nor the interaction between matrix size and block significantly predicted the interaction effect, suggesting that this trend is consistent across different state space sizes.

To examine whether the effect of repetition is also moderated by matrix size and learning blocks, we also ran a model predicting repetition effects from matrix size, block, and their interaction. The results indicate that repetition effects did not significantly change over blocks ($\beta = 0.00$), but differed by matrix size ($\beta = -0.05$), and the interactions of them ($\beta = 0.00$), suggesting that the effect of repetition was stable over learning but varied across different state space sizes.

Discussion

Our study provides clear evidence that procedural learning is sensitive to graded statistical structure in motor sequences. Across all five state space sizes tested (4, 5, 6, 7, 8), we found robust effects of both surprisal and previous position entropy on reaction times (H1, H2). These effects demonstrate that participants tracked the fine-grained transition probabilities embedded in the sequence, rather than simply practicing motor responses. Our findings are particularly striking given the brevity of training: participants completed only 20 blocks over ~20–40 minutes, receiving just 10 exposures per position per block. Despite this limited exposure, the trajectory analyses revealed that sensitivity to both surprisal and previous position entropy emerged early and continued to develop throughout the experiment. This rapid learning suggests that the procedural learning system is capable of representing and responding to distributional information more nuanced and complex than simple binary patterns from limited exposure to a sequence.

The weak, partially significant interaction between surprisal and previous position entropy (H3) offers converging evidence: even at high-entropy positions where the next position is less predictable, participants still exhibited sensitivity to the surprisal of specific transitions. This suggests that learners

are not merely detecting a binary “some positions are harder” distinction, but are truly tracking the different probabilities of individual transitions regardless of overall uncertainty.

The persistence of surprisal and previous position entropy effects across matrix sizes offers compelling evidence for the robustness of procedural learning across state space sizes (H4). While effect sizes were larger for smaller matrices, sensitivity to surprisal and previous position entropy remained statistically reliable even in the 8x8 condition. This scaling beyond the typical 4-position SRTT paradigm is encouraging for propositions that extend procedural learning as a supporting mechanism to skill acquisition in complex, naturalistic domains like language (Ullman, 2004).

Limitations

Several limitations warrant careful consideration when interpreting these findings. First, our transition matrices remain relatively small compared to the state spaces encountered in natural domains like language. Second, the short duration of our task means it does not fully represent the extended timescales typical of real-world procedural learning. Third, we made methodological decisions such as how to handle random effects structures when fitting separate models to individual blocks that may affect the quantitative estimates of effect sizes, though we believe these are unlikely to alter the qualitative pattern of results. These modeling choices were necessary to maintain convergence and interpretability across our per-block analyses, but future work should explore the robustness of specific effect size estimates to different modeling approaches. Finally, we ran a post-hoc power analysis on the pilot data on each condition separately, and found that ~80 participants *per condition* was required to achieve 80% power for detecting learning effects. While this is probably due to the small pilot sample (~2 participants per condition), it suggests that our study may be underpowered to detect smaller effects, particularly in the larger matrix size conditions. Future work with larger samples will be important for confirming the robustness of these effects and further exploring how they scale with state space size.

Implications for Understanding Procedural Learning

At face value, these results have important implications for understanding the relationship between procedural learning and language acquisition. Our task represents a substantially more complex design than traditionally used in procedural learning research, yet we still observe fine-grained sensitivity to surprisal and entropy that converges with recent quantitative analyses of language processing, where both surprisal (Levy, 2008) and entropy (Frank, 2013) independently predict reading times. This convergence increases our confidence that serial reaction time tasks can serve as valid models for language learning processes. More broadly, it increases the plausibility that procedural learning mechanisms and language learning rely on shared underlying cognitive systems, consistent with theoretical proposals linking the two (Ullman, 2004).

References

- Akaike, H. (1998). Information theory and an extension of the maximum likelihood principle. In *Selected papers of hirotugu akaike* (pp. 199–213). Springer.
- Barr, D. J., Levy, R., Scheepers, C., & Tily, H. J. (2013). Random effects structure for confirmatory hypothesis testing: Keep it maximal. *Journal of Memory and Language*, 68(3), 255–278. <https://doi.org/10.1016/j.jml.2012.11.001>
- Crump, M. J., & Logan, G. D. (2010). Episodic contributions to sequential control: Learning from a typist's touch. *Journal of Experimental Psychology: Human Perception and Performance*, 36(3), 662.
- Farkas, B. C., Krajcsi, A., Janacsek, K., & Nemeth, D. (2024). The complexity of measuring reliability in learning tasks: An illustration using the Alternating Serial Reaction Time Task. *Behavior Research Methods*, 56(1), 301–317. <https://doi.org/10.3758/s13428-022-02038-5>
- Frank, S. L. (2013). Uncertainty reduction as a measure of cognitive load in sentence comprehension. *Topics in Cognitive Science*, 5(3), 475–494. <https://doi.org/https://doi.org/10.1111/tops.12025>
- Grudin, J. T. (1983). Error patterns in novice and skilled transcription typing. In *Cognitive aspects of skilled typewriting* (pp. 121–143). Springer.
- Hedenius, M., Persson, J., Tremblay, A., Adi-Japha, E., Veríssimo, J., Dye, C. D., Alm, P., Jennische, M., Bruce Tomblin, J., & Ullman, M. T. (2011). Grammar predicts procedural learning and consolidation deficits in children with Specific Language Impairment. *Research in Developmental Disabilities*, 32(6), 2362–2375. <https://doi.org/10.1016/j.ridd.2011.07.026>
- Howard, D. V., Howard, J. H., Japikse, K., DiYanni, C., Thompson, A., & Somberg, R. (2004). Implicit Sequence Learning: Effects of Level of Structure, Adult Age, and Extended Practice. *Psychology and Aging*, 19(1), 79–92. <https://doi.org/10.1037/0882-7974.19.1.79>
- Howard Jr., J. H., & Howard, D. V. (1997). Age differences in implicit learning of higher order dependencies in serial patterns. *Psychology and Aging*, 12(4), 634–656. <https://doi.org/10.1037/0882-7974.12.4.634>
- Janacsek, K., Fiser, J., & Nemeth, D. (2012). The best time to acquire new skills: Age-related differences in implicit sequence learning across the human lifespan. *Developmental Science*, 15(4), 496–505. <https://doi.org/10.1111/j.1467-7687.2012.01150.x>
- Kassubek, J., Schmidtke, K., Kimmig, H., Lücking, C. H., & Greenlee, M. W. (2001). Changes in cortical activation during mirror reading before and after training: An fMRI study of procedural learning. *Cognitive Brain Research*, 10(3), 207–217. [https://doi.org/https://doi.org/10.1016/S0926-6410\(00\)00037-9](https://doi.org/https://doi.org/10.1016/S0926-6410(00)00037-9)
- Kidd, E., & Arciuli, J. (2016). Individual Differences in Statistical Learning Predict Children's Comprehension of Syntax. *Child Development*, 87(1), 184–193. <https://doi.org/10.1111/cdev.12461>
- Kuznetsova, A., Brockhoff, P. B., & Christensen, R. H. B. (2017). lmerTest Package: Tests in Linear Mixed Effects Models. *Journal of Statistical Software*, 82(13), 1–26. <https://doi.org/10.18637/jss.v082.i13>
- Lammertink, I., Boersma, P., Wijnen, F., & Rispens, J. (2017). Statistical Learning in Specific Language Impairment: A Meta-Analysis. *Journal of Speech, Language, and Hearing Research*, 60(12), 3474–3486. https://doi.org/10.1044/2017_JSLHR-L-16-0439
- Levy, R. (2008). Expectation-based syntactic comprehension. *Cognition*, 106(3), 1126–1177.
- Lum, J. A. G., Conti-Ramsden, G., Morgan, A. T., & Ullman, M. T. (2014). Procedural learning deficits in specific language impairment (SLI): A meta-analysis of serial reaction time task performance. *Cortex*, 51, 1–10. <https://doi.org/10.1016/j.cortex.2013.10.011>
- Lustig, C. (2022). *The transition from implicit to explicit knowledge representations in implicit sequence learning* [Text.thesis.doctoral, Universität zu Köln]. <http://www.uni-koeln.de/>
- Mullaney, T. S. (2017). *The Chinese Typewriter: A History*. MIT Press.
- Nissen, M. J., & Bullemer, P. (1987). Attentional requirements of learning: Evidence from performance measures. *Cognitive Psychology*, 19(1), 1–32. [https://doi.org/10.1016/0010-0285\(87\)90002-8](https://doi.org/10.1016/0010-0285(87)90002-8)
- Oliveira, C. M., Hayiou-Thomas, M. E., & Henderson, L. M. (2023). The reliability of the serial reaction time task: Meta-analysis of test–retest correlations. *Royal Society Open Science*, 10(7), 221542. <https://doi.org/10.1098/rsos.221542>
- Pinet, S., Ziegler, J. C., & Alario, F.-X. (2016). Typing is writing: Linguistic properties modulate typing execution. *Psychonomic Bulletin & Review*, 23(6), 1898–1906.
- Rabbitt, P. A. (1966). Errors and error correction in choice-response tasks. *Journal of Experimental Psychology*, 71(2), 264.
- Rumelhart, D. E., & Norman, D. A. (1982). Simulating a skilled typist: A study of skilled cognitive-motor performance. *Cognitive Science*, 6(1), 1–36. <https://api.semanticscholar.org/CorpusID:18037595>
- Shannon, C. E. (1948). *A mathematical theory of communication* (Vol. 27, pp. 379–423). The Bell System Technical Journal. <https://doi.org/10.1002/j.1538-7305.1948.tb01338.x>
- Song, S., Howard, J. H., & Howard, D. V. (2008). Perceptual sequence learning in a serial reaction time task. *Experimental Brain Research*, 189(2), 145–158. <https://doi.org/10.1007/s00221-008-1411-z>
- Squire, L. R. (2004). Memory systems of the brain: A brief history and current perspective. *Neurobiology of Learning and Memory*, 82(3), 171–177. <https://doi.org/https://doi.org/10.1016/j.nlm.2004.06.005>
- Ullman, M. T. (2004). Contributions of memory circuits to language: The declarative/procedural model. *Cognition*,

- Towards a New Functional Anatomy of Language*, 92(1), 231–270. <https://doi.org/10.1016/j.cognition.2003.10.008>
- Ullman, M. T., & Pierpont, E. I. (2005). Specific Language Impairment is not Specific to Language: The Procedural Deficit Hypothesis. *Cortex*, 41(3), 399–433. [https://doi.org/10.1016/S0010-9452\(08\)70276-4](https://doi.org/10.1016/S0010-9452(08)70276-4)
- West, G., Vadillo, M. A., Shanks, D. R., & Hulme, C. (2018). The procedural learning deficit hypothesis of language learning disorders: We see some problems. *Developmental Science*, 21(2), e12552. <https://doi.org/10.1111/desc.12552>