



Long-horizon associative learning as a unifying framework for statistical learning across scales

Lucas Benjamin^{a,1} , Ana Fló^a , Fosca Al Roumi^a , and Ghislaine Dehaene-Lambertz^{a,1}

Contributed by Ghislaine Dehaene-Lambertz; received May 28, 2025; accepted June 28, 2026; reviewed by Richard N. Aslin, Luca Bonatti, and Nicholas B. Turk-Browne

Sensory inputs are rich with temporal patterns that unfold across multiple timescales. Uncovering these regularities is essential for anticipating future events and navigating the environment efficiently. Numerous models have been proposed to account for learning at specific temporal scales; however, they are often designed in isolation and rely on narrowly tuned statistical measures, limiting their generalizability to other paradigms. In contrast, humans typically learn without prior knowledge of the underlying structure or the relevant timescale at which regularities occur. Here, we present a unifying account of statistical learning that spans a wide range of temporal dependencies, from adjacent and nonadjacent transitions to complex network structures. This model, long-horizon associative learning, offers a biologically grounded implementation of the successor representation, or equivalently, the free energy minimization model. Reanalyzing data from 11 previously published studies, we show that a single neural mechanism captures both local statistical regularities and higher-order structural properties. This mechanism rests on graded temporal overlap of associative traces and is governed by a single free parameter (β). This initial domain-general associative learning process, emerging from the graded structure of associations, may later scaffold to higher-level operations such as grouping, categorization, rule abstraction, and memory formation. Overall, this framework offers a conceptual synthesis that bridges disparate strands of the statistical learning literature and reframes apparent paradigm-specific effects as different expressions of a common underlying computation.

learning | infant | language | EEG | modeling

Speech, music, movements, or actions: many everyday events unfold in time, reaching us as continuous streams of sensory information. These streams are rich, diverse, and temporally structured, embedding complex dependencies across multiple timescales: from immediate transitions between consecutive events (1) to more distant regularities linking elements separated by long-range intervals (2, 3). Humans are remarkably skilled at discovering these statistical structures. For instance, listeners exposed to continuous speech are sensitive to the likelihood that one syllable follows another, a mechanism proposed to support early word learning in infants (1, 4). They also detect nonadjacent linguistic relationships, such as “is”...“ing” regardless of the number of intervening syllables (5, 6), and even infer higher-order structure in network-like sequences (7–9): when stimuli are drawn from distinct graph-based clusters, participants detect transitions between clusters even when adjacent transition probabilities remain constant.

Although some attempts have been made to aggregate adjacent and nonadjacent dependency learning under a common associative framework (10), much of the statistical and network learning literature still relies on distinct models tailored to specific experimental designs. These models can be broadly categorized by the temporal scale they target: local (transitions between adjacent elements), intermediate (nonadjacent dependencies), and high-order (network-level structure) (Fig. 1). This multiplicity implicitly assumes that participants engage in numerous parallel, scale-specific computations. Such fragmentation poses serious memory and efficiency challenges, especially for infants, whose cognitive resources, even if limited, support impressive learning capacities (1, 5, 11, 12).

To reconcile learning across timescales within a unified framework, we suggest that both short- and long-range associations arise from a single neural mechanism grounded in Hebbian principles, whereby learning results from overlapping neural representations of successive elements (49). Because neural traces decay gradually, successive elements overlap in time, enabling associations not only between consecutive events but also between temporally distant events, creating a “long-horizon” associative field. This mechanism is both parsimonious and biologically plausible, relying on sustained patterns of neural

Significance

Humans readily extract structure from sequences, detecting regularities that span from local transitions to large-scale network organization. Yet these phenomena are typically explained by separate models tailored to specific tasks. We discuss a unifying framework that shows this diversity can arise from a single mechanism: long-horizon associative learning. This biologically plausible process integrates information across time via graded associative traces. By replacing a fragmented set of statistical measures with a common computational principle, this framework reveals a shared basis for diverse statistical learning effects. It provides a simple and quantitative way to compare learning across paradigms, developmental stages, and species, opening avenues for linking behavioral findings with neural mechanisms.

Author affiliations: ^aCognitive Neuroimaging Unit, CNRS Equipe de recherche labellisée 9003, INSERM U1354, Commissariat à l’Energie Atomique et aux énergies alternatives, Université Paris-Saclay, NeuroSpin Center, Gif/Yvette 91191, France

Author contributions: L.B. and G.D.-L. designed research; L.B. performed research; L.B. analyzed data; G.D.-L. funding acquisition; and L.B., A.F., F.A.R., and G.D.-L. wrote the paper.

Reviewers: R.N.A., Haskins Laboratories Inc; L.B., Université Pompeu Fabra; and N.B.T.-B., Yale University.

Competing interest statement: One of the authors (A.F.) co-authored in 2022 a review paper with one of the reviewers of this study (Richard Aslin).

Copyright © 2026 the Author(s). Published by PNAS. This article is distributed under [Creative Commons Attribution-NonCommercial-NoDerivatives License 4.0 \(CC BY-NC-ND\)](https://creativecommons.org/licenses/by-nc-nd/4.0/).

¹To whom correspondence may be addressed. Email: lucas.benjamin78@gmail.com or gdehaene@gmail.com.

This article contains supporting information online at <https://www.pnas.org/lookup/suppl/doi:10.1073/pnas.2513423123/-/DCSupplemental>.

Published July 28, 2026.

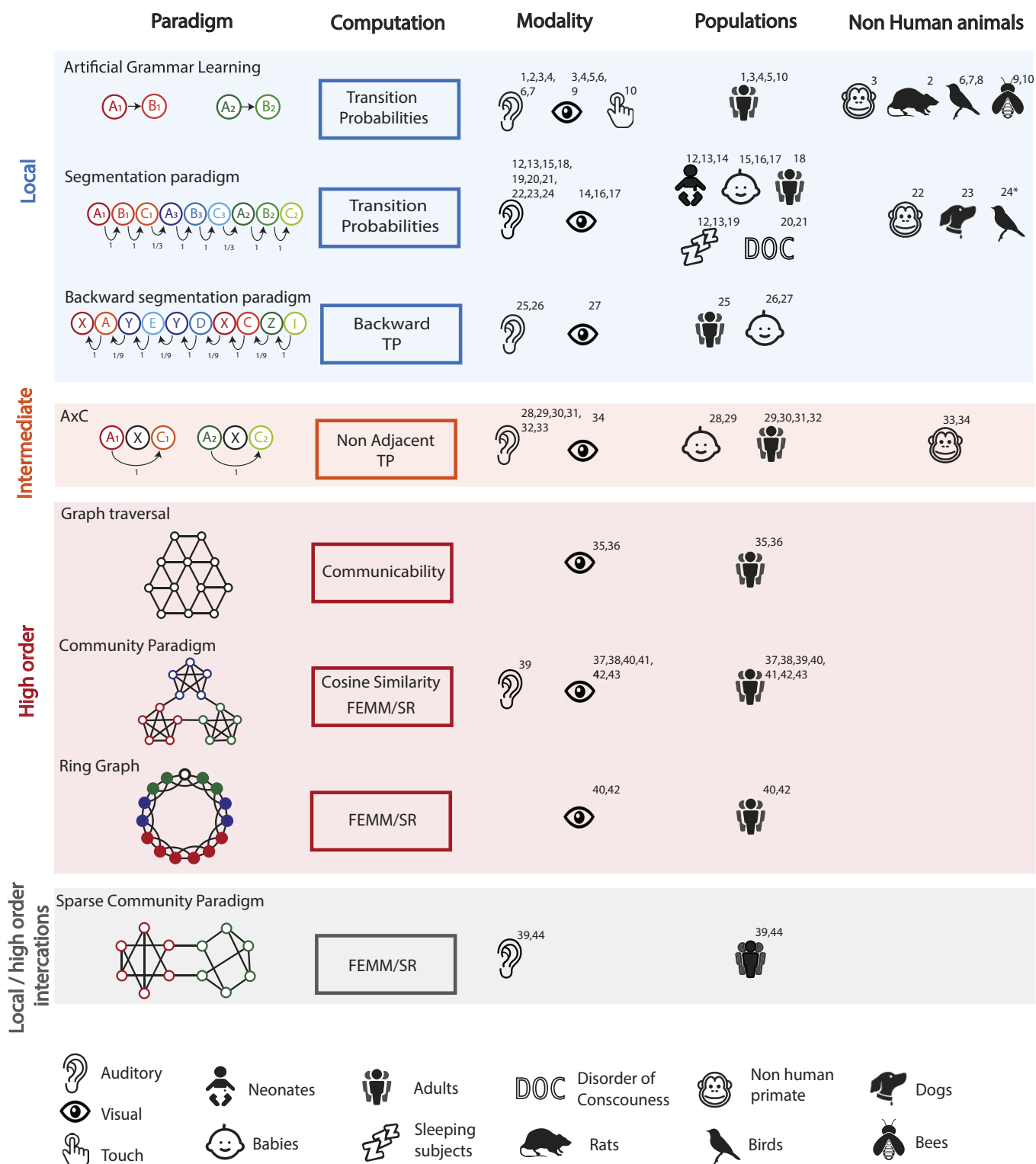


Fig. 1. Overview of statistical learning paradigms categorized by the scale of temporal dependencies: local, intermediate, and high-order. The figure highlights some studies demonstrating learning at different scales of statistical structure, across various sensory modalities and participating populations. It should be noted that this does not necessarily represent the ability of populations to learn or not learn each paradigm, but rather the current state of our knowledge on this subject and on what has been formally tested. 1. Harris (4), 2. Toro and Trobalon (13), 3. Milne et al. (14), 4. Conway and Christiansen (15), 5. Reber (16), 6. Santolin et al. (17), 7. James et al. (18), 8. Menhart et al. (19), 9. Avarguès-Weber et al. (20), 10. Conway and Christiansen (15), 11. Nicholls et al. (21), 12. Fló et al. (22, 23), 13. Benjamin et al. (11), 14. Bulf et al. (24), 15. Saffran et al. (1), 16. Fiser and Aslin (25), 17. Kirkham et al. (26), 18. Saffran et al. (1, 27), 19. Batterink and Zhang (28), 20. Xu et al. (29), 21. Benjamin et al. (30), 22. Hauser et al. (31), 23. Boros et al. (32), 24. Takahasi (33), 25. Perruchet and Desauty (34), 26. Pelucchi et al. (35), 27. Tummelsthammer et al. (36), 28. Buiatti et al. (37), 29. Gómez et al. (6), 30. Peña et al. (5), 31. Endress (38), 32. Newport and Aslin (12), 33. Newport et al. (39), 34. Sonnweber et al. (40), 35. Garvert et al. (7), 36. Mark et al. (41), 37. Schapiro et al. (3, 42), 38. Karuza et al. (43, 44), 39. Benjamin et al. (8), 40. Lynn et al. (45), 41. Ren et al. (46), 42. Stiso et al. (9), 43. Kakaei et al. (47), 44. Benjamin et al. (48).

activity that could, in principle, be implemented with a limited number of neurons. It can also be formalized mathematically, yielding parameters that can be directly linked to measurable neural properties.

The free energy minimization model [FEMM (45)] and the Successor Representation [SR (50)] model provide formal instantiations of long-horizon associative learning (L-HAL). Initially proposed to explain network learning in humans (FEMM) and cognitive map formation for spatial navigation (SR), the two models are mathematically equivalent despite their conceptually distinct origins. Both characterize associative strength as a weighted combination of near and distant transitions: a linear mixture of the transition probability matrix for adjacent elements (A_1) and higher temporal orders ($A_{\Delta t}$), with exponential decay as a function of Δt .^{*} A single parameter, β , governs this decay: when $\beta \rightarrow \infty$, only adjacent transitions contribute; when $\beta \rightarrow 0$, all elements become equally associated regardless of distance. Thus, estimating β directly specifies the associative timescale.

The model is agnostic to the type of primitive (syllables, voices, tones, or any other units) and can operate over any representational stream. However, it requires that each stream consist of discrete, comparable elements forming a coherent sequence. For example, statistics can be computed on structured streams of syllables with random voices producing them or on structured streams of voices producing random syllables (22). Within any given representational layer, a single β captures the associative timescale; distinct layers may operate concurrently with different decay constants and developmental trajectories. The model thus describes an early, domain-general associative process operating over any segmentable sequence, whose outputs scaffold higher-level operations such as grouping, categorization, rule abstraction, and memory formation.

We previously tested the plausibility of such long-horizon associations in adults using an auditory “sparse community” paradigm in which tone sequences were governed by both local and high-order dependencies. Tones were drawn from a network organized in two clusters, so that transitions occurred either within or between clusters while adjacent transition probabilities were held constant. Participants later judged within-cluster transitions as more familiar, even when they had never been presented; whereas between-cluster transitions were judged less familiar, despite being presented as often as the within-cluster ones (8) (see Fig. 1, Local/High-order interaction). Participants were thus sensitive to the higher-order structure of the network, implicitly completing missing transitions consistent with that structure while rejecting those that violated it. A subsequent MEG study revealed neural signatures consistent with decaying-overlap dynamics (48), supporting the idea that even high-order structure learning and generalization can emerge from a very simple associative rule.

Given its general formulation, the L-HAL framework should capture temporal associations at all scales, from adjacent to long-range dependencies. We therefore challenged the model with a broad set of published datasets spanning local, nonadjacent, and network-level learning. The β parameter may vary across domains, developmental stages, or cognitive states. However, we might expect it to fall within a restricted range that captures the main signatures of human learning, thereby offering a unified, parsimonious account of statistical learning across scales. Thus, we provide an explicit comparison of the data with model predictions using a unique β value previously optimized (8, 45), and evaluate

domain validity from predictions obtained across a broad range of β values.

Model Description

Our analyses drew on publicly available datasets encompassing a wide range of behavioral and neurophysiological measures, including familiarity rankings, reaction times in healthy adults, looking preferences in infants, ERP and fMRI. Collectively, these datasets cover the full range of temporal scales depicted in Fig. 1. We qualitatively compared the experimental results with the model's predictions and quantitatively assessed the goodness of fit when sufficient data were available.

Formal Description of the Model. Given the formal equivalence between FEMM and SR (*Discussion*) and for the sake of consistency with our previous work, we anchored the description in the FEMM formalism. This model captures a trade-off between short and long-distance associations: instead of linking only consecutive elements (i and $i - 1$), it also binds each element i to earlier ones ($i - 2$, $i - 3$, etc.) with a probability that decreases exponentially with temporal distance, following a Boltzmann distribution (*SI Appendix*).

The estimated model can then be written as

$$\hat{A} = \sum_{\Delta t=0}^{+\infty} Q(\Delta t) A_{\Delta t+1} \quad \text{With } Q(\Delta t) = \frac{1}{Z} e^{-\beta \Delta t}, \quad [1]$$

$A_{\Delta t}$: Nonadjacent transition probability matrix of order Δt between elements of the sequence,

β : Generalization factor varying from 0 (no generalization) to 1 (uniform generalization),

Z : Constant value $Z = 1/(1 - e^{-\beta})$,

Δt : Distance between items of the sequence.

For sequences following a Markovian rule—the current element is only dependent on the previous one— $A_{\Delta t} = A^{\Delta t}$. Eq. 1 can be rewritten:

$$\hat{A} = \sum_{\Delta t=0}^{+\infty} Q(\Delta t) A^{\Delta t+1} \quad \text{With } Q(\Delta t) = \frac{1}{Z} e^{-\beta \Delta t}, \quad [2]$$

Which simplifies as (*SI Appendix*):

$$\hat{A} = (1 - e^{-\beta}) A (I - e^{-\beta} A)^{-1}, \quad [3]$$

I : Identity matrix (same size as A)

Computation of the Model for Each Paradigm. We carefully curated landmark studies spanning multiple scales of statistical learning. They mostly shared a common structure: participants were exposed to sequences of stimuli during training, then tested on congruent and incongruent items in a postlearning phase. Learning was inferred from differences in familiarity ratings or deviance responses between these item types.

We simulated the associative strength between elements from the training material, computed the predicted familiarity for each test condition, and compared these predictions with published empirical results. To maintain parsimony and avoid paradigm-specific overfitting, we first fixed β to be 0.06, the value proposed by Lynn et al. (45) and that we also found in our previous work (8). As a second step, for each study, we identified the range of β values yielding the

^{*} Δt corresponds to the distance between items of the sequence and not time per se.

best fit, exploring values from 10^{-4} to 10^2 to determine the optimal range explaining performance across paradigms.

The following approach was used:

- For Markovian sequences, we computed the analytical association matrix using Eq. 3.
- For AXC and artificial grammar paradigms, where A is insufficient to fully describe the training material, we estimated FEMM associations directly from the corpus of training sequences. This was done by computing and summing (with an exponential decay factor) all TP orders within each training sequence (Eq. 1). The final $\hat{A}$ matrix was obtained by averaging the estimated association matrix across all training sequences experienced by participants.

This method estimated the optimal learning after convergence, ignoring possible local imbalance or curriculum effects (see later in the discussion for a model including learning dynamics). Once the participants' theoretical optimal mental model was derived from the training material, we computed the predicted familiarity for each test condition:

- For a pair (ab), the estimated familiarity is $\hat{A}(ab)$. This is the case for all network paradigms (Fig. 2 E–G and I)
- For longer test sequences, we computed a weighted sum of all pairwise associations between elements. For instance, the familiarity of a triplet abc was estimated as $\hat{A}(ab) + \hat{A}(bc) + e^{-\beta} \hat{A}(ac)$. This procedure was applied to all paradigms involving test sequences with more than two elements (Fig. 2 A–D). In the artificial grammar paradigm, the state “end” was treated as an item, following the original study.
- In the specific case of Lynn et al. (45), we followed their original method to compare two types of networks: Community and Lattice structures, by computing the average FEMM familiarity across all transitions within each network.

We obtained quantitative predictions of model familiarity for each experimental condition and compared them with published experimental measures.

Results

All the results are presented in Fig. 2.

Local Scale.

Segmentation paradigms. Sensitivity to adjacent transition probabilities has long been recognized as a core mechanism of statistical learning, most notably demonstrated in the seminal speech segmentation study in 8-mo-old infants (1). In such paradigms (1, 11, 30, 51), participants hear continuous streams of syllables (or other auditory, visual, or haptic stimuli) containing pseudowords (ABC). Within words, transition probabilities are high ($TP = 1$), but drop to $1/3$ across word boundaries. After brief exposure, participants successfully extract words ($A_i B_i C_i$) as they distinguish them from part-words ($B_i C_i A_i$) straddling a word boundary, through differences in looking time, familiarity rating, chunk recall accuracy, or EEG responses (11, 22, 23, 27, 29, 30, 52).

Comparison with the model. FEMM simulations reproduced these effects: for $\beta = 0.06$, word familiarity exceeded part-word familiarity by 19% (Fig. 2A). As adjacent transition probability is the first order and dominant component of the FEMM, it is not surprising that a TP difference induced a FEMM difference and that larger β values further amplify the word vs part-word difference (Fig. 2A).

Grammar learning paradigms. Milne et al.'s grammar-learning study (14) offered an ideal test case, combining simplicity with fully reported training and test sequences. Human participants and macaques were exposed to 8 different training sequences, then judged whether new five-item sequences followed the same grammar. Their judgments correlated with local transition probabilities in both visual and auditory modalities for both humans and macaques.

Comparison with the model. FEMM predictions, computed from the training set, closely matched participants' behavior (Visual $R = -0.69$ $P < 0.05$, Auditory $R = -0.74$ $p < 0.01$ for $\beta = 0.06$ – Fig. 2B), indicating that the model accurately captured observed learning patterns. This was not the case for macaques: only higher β values provided a good fit to their data, suggesting that monkeys primarily encoded first-order transition probabilities, neglecting longer-distance associations (SI Appendix, Fig. S2).

Intermediate Scale.

AXC paradigm. At the intermediate level, sensitivity to nonadjacent dependencies has been proposed to underlie grammatical learning, such as tense agreement in English (“is ... ing,” “has ... ed”). Peña et al. (5) showed that adults can learn nonadjacent transition probabilities of the form $P(C|XA)$, where the syllable A predicts C despite an intervening X syllable when AXC triplets were separated by subliminal pauses. In that experiment, participants later preferred test items preserving the A–C dependency (Rule Word $AX'C$ where X' is novel) over Part-Words ($XC_i A_i$). Gómez (6) further found that learning improved as the variability of possible Xs increased (2, 6, 12, or 24 syllables), a result that cannot be explained by learning the constant A–C dependency alone.

Comparison with the model. We estimated FEMM familiarity scores for Words, Part-Words, and Rule-Words in the classical case with three possible intervening Xs (37) and training on isolated words. Following the authors' interpretation that pauses segmented the stream, we computed the model on each triplet (figure 1 in (5), Fig. 2C). For $\beta = 0.06$, the model predicted Words familiarity ($A_i X C_i$) > Rule-Words familiarity ($A_i X' C_i$) > Part-Words familiarity ($XC_i A_i$), closely matching behavioral data (5, 12, 37). Consistent with empirical findings (5), removing pauses during training (i.e., using a continuous stream) abolished the familiarity advantage for Rule-Words over Part-Words. Varying the number of possible Xs (2, 12, 22, 30, 48, 52) replicated Gómez's result: the Word > non-Word familiarity difference increased with X variability (8 – Fig. 2D). This pattern reflects FEMM's trade-off between local and distant associations: when X variability grows, adjacent transitions (AX, XC) weaken while the A–C link remains constant, biasing predictions toward the nonadjacent dependency. Finally, unlike in local paradigms, the best-fitting β range for intermediate-scale learning shifts toward lower values, but still displays a good overlap with the local paradigm range of validity (Fig. 2 C and D).

Higher-Order Scale.

Garvert et al.'s network paradigm. At a higher-order scale, sequence structure is often described using network formalisms. Garvert et al. (7) showed that hippocampal–entorhinal fMRI activity during navigation of a 12-node network correlated with Communicability, a graph-based measure of distance between nodes. Although conceptually similar to FEMM or SR, Communicability penalizes long-range dependencies more strongly.

Comparison with the model. We computed FEMM familiarity estimates for each network transition and reanalyzed the average activation within the hippocampal–entorhinal cluster identified by Garvert et al. Mean activation per transition was significantly correlated with FEMM predictions ($R = -0.29$, $P < 0.05$ for $\beta = 0.06$)

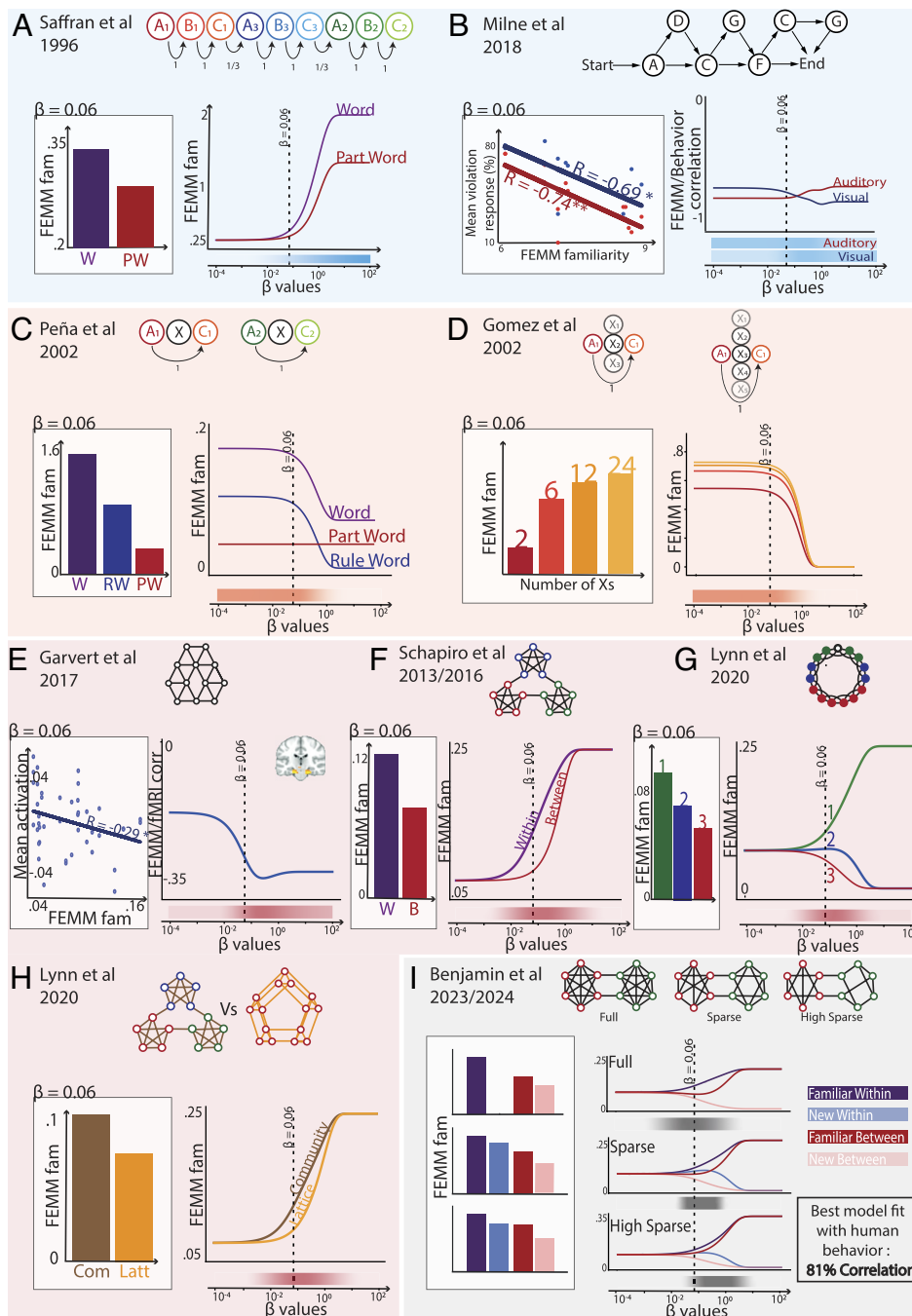


Fig. 2. Results from the analysis of the statistical learning paradigms across different temporal scales using the Free-Energy Minimization Model (FEMM). In this figure, the *Left* panels in white display the model predictions for a previously optimized β value ($\beta = 0.06$), and the *Right* panels the model predictions for β values between 10^{-4} and 10^2 . Dotted lines show $\beta = 0.06$, and shaded bars below illustrate the β range for which the model accurately predicts the available human data. (A) In word segmentation tasks using familiarization with continuous speech streams, participants are subsequently more familiar with embedded “Words” than with “Part-Words” that span boundaries. For this local paradigm, higher β values better predict the Word vs. Part-Word differences. (B) Artificial grammar learning paradigms target local transition probability learning. FEMM familiarity values significantly correlate with mean violation responses reported by Milne et al. (14) in both auditory (red) and visual (blue) modalities. Macaques’ performance modeling (SOM) revealed higher β values than in humans. (C) At the intermediate scale, sensitivity to nonadjacent dependencies has been demonstrated using AXC paradigms. The reported preference for RuleWord over PartWord is also correctly accounted for by FEMM. (D) Increasing the number of X elements in the AXC paradigm affects participants’ performances in a Word vs NonWord discrimination task, despite the nonadjacent transition probability $P(C|XA) = 1$ remaining constant across conditions. FEMM estimates (*Left* bar-plot) for varying numbers of intervening X elements (ranging from 2 to 29) mirror the behavioral results reported by Gómez et al. (6) (*Right* bar-plot), showing that word preference increases with the number of possible intervening syllables. (E) Using fMRI, Garvert et al. (7) showed that when exposed to a sequence deriving from a network structure, a cluster in the entorhinal cortex significantly correlated with Communicability. FEMM also shows a significant correlation with the average activation within this cluster. (F) In the community paradigm (3), participants distinguished transitions within vs. between clusters despite identical local transition probabilities. FEMM reproduces this distinction as shown in the bar plot comparing the familiarity of within and between transitions. (G) In a ring graph paradigm, Lynn et al. (45) showed graded familiarity judgments based on graph distance (2 vs. 3 steps apart). FEMM accurately predicts these gradations. Distance 1 corresponds to the connection between white and green nodes, Distance 2 to white and blue nodes, and Distance 3 to white and red nodes. (H) Lynn et al. (45) also showed that learning a community structure was easier (faster reaction time) than a lattice structure. FEMM predicts this effect, with higher average familiarity for community (brown) than lattice (orange) graphs. (I) In a modified version of the community paradigm (sparse and highly sparse designs), we removed some within-community transitions to be able to test learning of local and high-order structures. FEMM predicted generalization within-community, assigning higher familiarity to nonexistent within-community transitions than to familiar between-community ones. Behavioral and MEG results confirmed this prediction. We showed that FEMM significantly outperformed the other statistical models proposed in the literature.

(Fig. 2E), supporting the model's relevance for explaining brain responses to structured sequences.

Community network. Schapiro et al. (3) introduced the community network paradigm to demonstrate humans' sensitivity to network structure, even in the absence of informative local statistics. The network consists of three communities—groups of nodes highly interconnected internally but weakly connected to other groups—where each node has exactly four neighbors, and thus a uniform probability of 0.25. Despite this flat local structure, participants reliably distinguished transitions within communities from those between communities (3, 8, 42, 44, 45, 48, 53, 54). Lynn et al. (45) further contrasted this community network with a lattice network of identical size and average connectivity and found that participants trained on a lattice structure exhibited lower familiarity (slower reaction times) compared to those trained on a community structure.

Comparison with the model. Originally proposed to explain these effects, FEMM accurately predicted the observed drop in familiarity at community boundaries (45) (Fig. 2F) and reproduced the higher overall familiarity for community vs. lattice networks. Indeed, average FEMM familiarity across all transitions was higher for the community network than for the lattice network, consistent with participants' faster reaction times when learning the community structure (Fig. 2H).

Ring graph. Lynn et al. also investigated a different network topology: a 15-node ring-graph in which each node connects to four neighbors (+1/+2/-1/-2). Thus, nodes are separated by one, two, or three transitions. Measuring the familiarity of each pair of this network using reaction times revealed a graded preference for directly connected pairs (distance 1), followed by unconnected pairs at distance 2, and then distance 3.

Comparison with the model. FEMM reproduced this hierarchy of familiarity, as shown by Lynn and colleagues (55) (Fig. 2G), with predicted familiarity decreasing systematically with increasing graph distance.

Local and High-Order Interaction.

Sparse and high sparse community networks. In previous studies, we introduced the sparse community design (8, 48), in which one (sparse) or two (high sparse) possible transitions per node inside each community were removed. It created a 2-by-2 design used during test where a transition between two sounds could be Familiar or New (heard vs. unheard during training) and occurred Within or Between communities (the two sounds were from the same community or not). Participants systematically "completed" the missing within-community transitions—judging them as familiar despite never hearing them—indicating they reconstructed the latent structure. Conversely, between-community transitions were judged less familiar than within-community ones, even though both occurred equally often.

To explain this behavior, we systematically compared leading models from the statistical and network learning literature, including transition probabilities, nonadjacent transition probabilities, n-grams, communicability... FEMM provided the best fit (8, 48) (Fig. 2I). Remarkably, across the full, sparse, and high-sparse community designs, only a narrow range of β values fitted the data (Fig. 2I), making these paradigms powerful probes for pinpointing the associative timescale.

Discussion

How can we efficiently extract multiple scales of structure from temporal input? This question has long engaged researchers both in cognitive neuroscience and machine learning, driving computational

innovations in artificial learning agents, including the emergence of recurrent connections in artificial networks (56), or memory traces learning methods for spiking networks (57). The framework for statistical learning presented in this study seeks to connect formal mathematical and biological descriptions, explaining how humans learn conditional probabilities across different temporal structures. Indeed, its dynamics can be instantiated without recurrent connections, by simple physiologically plausible mechanisms such as sustained neural activity. When items' representations persist over time, their temporal overlap lets the brain form associations across various timescales. This mechanism enables the simultaneous learning of a wide range of conditional probabilities within a sequence. These persistent neural traces are assumed by the models to decay exponentially (Fig. 3B), a hypothesis supported by recent empirical evidence (28, 30).

Our goal was to show that this single associative principle can account for the diversity of findings reported in the statistical learning literature, across multiple temporal scales. While the Free-Energy Minimization Model (FEMM) provided the original mathematical formalization, it was originally developed in the predictive coding framework. Here, we adopt the more neutral term L-HAL to emphasize that the same formal mechanism can arise passively from overlapping neural traces, without invoking explicit prediction or error correction. The model accurately reproduces experimental data while generalizing across paradigms that have traditionally required distinct, domain-specific models. More broadly, effects interpreted as abstract rule learning (43, 44)—such as nonadjacent (AXC) dependencies or network completion—may instead emerge naturally from long associative dynamics. It also bridges two domains often studied independently: statistical learning and network learning. While network learning is traditionally viewed as abstract structure learning requiring explicit mental maps (46, 58), our results suggest that these high-order effects can emerge naturally from the same implicit associative dynamics seen in statistical learning. Finally, the model's minimal architecture makes it biologically plausible: learning across timescales reduces to a single Hebbian process governed by gradually decaying memory traces.

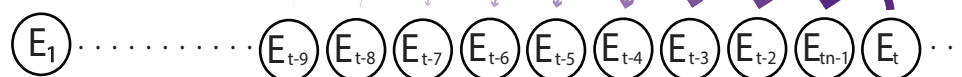
A Single Free Parameter to Capture Integration Across Statistical Scales. L-HAL explains how associations extend across time through a single parameter, β , which tunes the balance between precision and generalization. High β values restrict learning to local transitions, while low values promote integration across distant elements. Estimating β across species or ages could thus reveal differences in associative timescales and generalization capacity. Cross-species data support this interpretation: in Milne et al. (14), macaques' behavior was best captured by high β values, indicating reliance on first-order associations (SI Appendix, Fig. S2). In human infants, β may start high—reflecting immature neurons' limited capacity for sustained firing—favoring short-range associations, consistent with evidence of local transition learning in sleeping neonates. Implicit learning paradigms combined with EEG recordings could be used to track how β evolves with age. While variation in integration timescale as measured with β is likely to reflect only a part of the developmental maturation of statistical learning abilities, measuring β may nonetheless provide a testable functional link between regional brain maturation and changes in performance.

Note, however, that the measured β is often not derived directly from the output of the LHAL mechanism, but rather from the behavioral outcomes of the decision and response processes. For instance, in an active and explicit behavioral version of the sparse

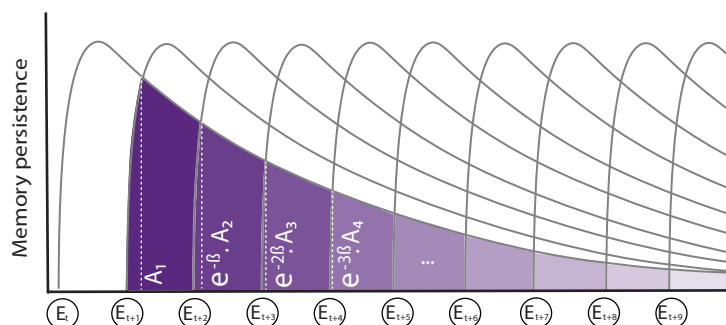
A Theoretical associative strength :

$$\hat{A} = \frac{1}{Z} \sum_{\Delta t=0}^{+\infty} e^{-\beta \Delta t} \cdot A_{\Delta t+1} \quad \text{With } A_n = P(E_t | X, \dots, XE_{t-n})$$

Input Sequence :



B Long Horizon associative learning implementation



C

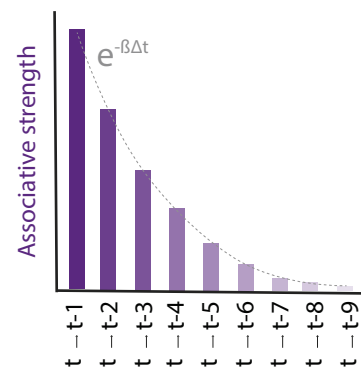


Fig. 3. Description of the FEMM and its L-HAL implementation. (A) Formal equation of the FEMM, used to compute the association matrix ($\hat{A}$) between all pairs of elements in a sequence. $\hat{A}$ reflects the cumulative contribution of adjacent and nonadjacent transition probabilities across all temporal distances, weighted by an exponential decreasing function. The diagram on the right illustrates the associative strength between a current element E_t and its predecessors. It shows strong association with the immediately preceding element (E_{t-1} , adjacent TP) and weaker associations with earlier elements ($E_{t-2}, E_{t-3}, \dots$; nonadjacent TP) consistent with the exponential decay. (B) Graphical representation of the FEMM's implementation through a L-HAL strategy. If the mental representation of each element decays exponentially over time, the overlap between successive representations allows for Hebbian associations to form between both adjacent and nonadjacent elements, following a penalty computed with FEMM. (C) Schematic depiction of the exponentially decaying associative weights between the current element and its predecessors, illustrating connections to both consecutive ($t \rightarrow t-1$) and nonconsecutive ($t \rightarrow t-2$; $t \rightarrow t-3$; ...) elements of the sequence.

community paradigm, we estimated an average β of 0.06. In contrast, using the same paradigm in an implicit and passive MEG study yielded a higher estimate of $\beta = 0.5$. While this difference may partly result from variations in measurement sensitivity, it may also reflect the impact of active engagement. The effort to memorize, by recruiting stronger top-down processes, could contribute to higher maintenance of remote elements (longer traces = lower β – Fig. 3 A and B). Given ongoing debates about the role of attention in statistical learning (29, 30, 59–64), careful estimation of β across different paradigms might help clarify the exact role of attention in learning.

Finally, significant individual differences might also modulate the β value. Although challenging to measure at the individual level, Lynn and colleagues (45) showed that an individually fitted β parameter correlated with participants' memory capacity in an n-back task, providing initial evidence that individuals may differ in their generalization/accuracy trade-off, or possibly reflecting attentional engagement in both tasks. We still lack intrasubject test–retest reliability of this parameter to investigate this question. Given the focus on measurement reliability in the recent literature (52), a specific study on the consistency of individual β parameters across statistical learning tasks is an important next step.

Biological Implementation. We proposed L-HAL as a biologically grounded implementation of FEMM or SR. Such an implementation predicts traces of past elements slowly decaying with time. In a previous MEG experiment, we found evidence consistent with this: neural representations of tones decayed exponentially over time (or number of elements—see limitation paragraph), with overlapping activity across up to eight successive tones

(Fig. 3 B and C). Moreover, MEG signals reflected the probability of upcoming tones based on this decay pattern, capturing associations at multiple orders of transition probabilities (48). Recently, Kahn and colleagues (54) found learning properties compatible with a trace-like implementation of SR, rather than an alternative encoding with element reactivations. This further supports L-HAL as a good candidate to implement FEMM or SR in implicit sequence learning tasks.

Model Extension to Spatial Associations. The principles underlying associations between temporally ordered events (a single dimension) described here can, theoretically, be extended to any space, regardless of the axes or the number of dimensions. In particular, for spatial patterns, the association strength (Fig. 3B) can be defined by replacing Δt by $\Delta \text{distance}$, with $\Delta \text{distance} = \max(|\Delta x|, |\Delta y|)$, or any other 2D distance metric (SI Appendix). While this formulation provides a computational analogue of FEMM for spatial patterns, the equivalence does not hold at the level of neural implementation. Implementing such two-dimensional long-distance associations would likely require mechanisms beyond simple sustained neural activity (e.g., the serial allocation of attention to limited spatial regions of a scene, noisy grid-cell, etc). Characterizing these mechanisms, however, lies beyond the scope of the present manuscript.

Related ideas have emerged in the reinforcement learning and spatial navigation literature. In these domains, hippocampal place cells have been proposed to represent maps of probabilistic future states and rewards by encoding the successor representation rather than simple positional cognitive maps (50, 65, 66). Because future states may encompass multiple steps ahead, the underlying

mathematical problem is analogous to integrating transition probabilities across different temporal orders. The successor representation (SR) provides a formal solution to this problem. SR is defined as a weighted sum of probabilistic future states: $SR = \sum_{\Delta t} \gamma^{\Delta t} A^{\Delta t}$, where A is the transition probability matrix between successive states and $\gamma^{\Delta t}$ (with $0 < \gamma < 1$) is a discount factor. Although FEMM and SR emerged from different research contexts (statistical learning vs. cognitive map), they are mathematically equivalent under the transformation $\gamma = e^{-\beta}$. This convergence suggests that similar computational principles may underlie sequential statistical learning, spatial pattern recognition, and spatial navigation, albeit with distinct neural implementations.

Model Limitations and Further Directions. Our goal here was to highlight the power of low-level associative mechanisms and their potential role in supporting a wide range of statistical learning phenomena. However, such mechanisms are unlikely to explain all aspects of human sequence learning. In many cases, additional processes may build upon this initial layer, whose output might be further integrated into higher-level structures (67), such as syllables into words, on which L-HAL may operate again. In the speech domain particularly, this basic process may be superseded by other constraints: prosodic boundaries are known to block learning across the boundary (11, 68, 69), suggesting that associative learning is constrained by the structural domain in which it occurs. The ability to extract AxC structure is similarly modulated by developmental stage and input format [continuous vs. segmented streams; native vs. foreign languages (70–72)], implying that prior knowledge or architectural constraints either directly modulate β values or act downstream on the model's output. Dissociating these possibilities and controlling for the influence of input structure on learning will require further investigation.

Because L-HAL formalizes a graded associative process based on overlapping neural traces, it inherently supports probabilistic generalizations that extend beyond the observed input. It enables extrapolation to unobserved but statistically compatible combinations, such as “phantom words” (73) or graph completion in the sparse community paradigm (8, 48). However, it is radically different from rule-based learning (74–79), which involves explicit abstraction: the input is no longer encoded as sensory exemplars but instead represented in terms of labels. For instance, two alternating tones of 800 and 1,000 Hz may be abstracted as an ABAB pattern, and generalized to any alternating stimuli, tones, or visual stimuli (79). This symbolic form of generalization is not accessible to L-HAL and likely emerges at a later processing stage.

Further research is needed to clarify how these forms of learning interact and whether rule-based learning might emerge as a strategy to reduce the cognitive cost of maintaining long-range associations. It remains also to be determined whether the decay in associative strength is governed primarily by time or by the number of intervening elements (as implemented here). Fig. 2 showed that for most paradigms, the model predictions are not drastically different with a 10-fold β factor. Within a range of medium-paced sequences, considering time or number of elements would yield similar results for short elements like syllables or tones.

In this paper, the model is initially described using a theoretical transition matrix to compute the familiarity estimate. However, this approach overlooks the effects of incomplete learning during sequence exposure, as well as curriculum effects such as recency or primacy bias for instance (80–82). The model can be readily extended to address these limitations by replacing the absolute familiarity estimate with a dynamic one, updated after the presentation of each element i . Formally, Eq. 1 is replaced by

$$\hat{A}(i) = \sum_{\Delta t=0}^{+\infty} Q(\Delta t) \tilde{A}(i)_{\Delta t+1} \text{ With } Q(\Delta t) = \frac{1}{Z} e^{-\beta \Delta t},$$

where $\hat{A}(i)$ is the model estimation after the presentation of element i in the sequence and $\tilde{A}(i)$ the estimated transition probability conditioned on the sequence up to element i . This matrix reflects the actual sequence experienced by subjects and can incorporate biases (for instance, recency bias can be implemented by restricting $\tilde{A}(i)$ to recently observed items). In this formulation, L-HAL becomes a dynamic framework that captures both online learning during sequence presentation and the amount of exposure required to learn each sequence, as well as curriculum-related effects. Such a dynamic formulation could help account for individual learning trajectories rather than just the final result of efficient learning as in the current study.

Data, Materials, and Software Availability. All codes for this study are available at <https://osf.io/g65d7> (83).

ACKNOWLEDGMENTS. This work received funding from the European Research Council under the European Union's Horizon 2020 research and innovation program (Grant Agreement No. 695710 and 101142651) and from the French government managed by the Agence Nationale de la Recherche as part of the France 2030 program, under reference ANR-23-IAHU-0010. L.B. thanks Les Treilles Foundation for its support. The authors used an AI-assisted language tool to improve the clarity and grammar of the manuscript. All scientific content and interpretations are the sole responsibility of the authors.

1. J. R. Saffran, R. N. Aslin, E. L. Newport, Statistical learning by 8-month-old infants. *Science* **274**, 1926–1928 (1996).
2. A. Basirat, S. Dehaene, G. Dehaene-Lambertz, A hierarchy of cortical responses to sequence violations in three-month-old infants. *Cognition* **132**, 137–150 (2014).
3. A. C. Schapiro, T. T. Rogers, N. I. Cordova, N. B. Turk-Browne, M. M. Botvinick, Neural representations of events arise from temporal community structure. *Nat. Neurosci.* **16**, 486–492 (2013).
4. Z. S. Harris, From phoneme to morpheme. *Language* **31**, 190 (1955).
5. M. Peña, L. L. Bonatti, M. Nespor, J. Mehler, Signal-driven computations in speech processing. *Science* **298**, 604–607 (2002).
6. R. L. Gómez, Variability and detection of invariant structure. *Psychol. Sci.* **13**, 431–436 (2002).
7. M. M. Garvert, R. J. Dolan, T. E. Behrens, A map of abstract relational knowledge in the human hippocampal-entorhinal cortex. *eLife* **6**, 1–20 (2017).
8. L. Benjamin, A. Fló, F. Al Roumi, G. Dehaene-Lambertz, Humans parsimoniously represent auditory sequences by pruning and completing the underlying network structure. *eLife* **12**, e86430 (2023).
9. J. Stiso *et al.*, Neurophysiological evidence for cognitive map formation during sequence learning. *eNeuro* **9**, ENEURO.0361-21.2022 (2022).
10. A. D. Endress, S. P. Johnson, When forgetting fosters learning: A neural network model for statistical learning. *Cognition* **213**, 104621 (2021).
11. L. Benjamin *et al.*, Tracking transitional probabilities and segmenting auditory sequences are dissociable processes in adults and neonates. *Dev. Sci.* **26**, e13300 (2022).
12. E. L. Newport, R. N. Aslin, Learning at a distance I. Statistical learning of non-adjacent dependencies. *Cogn. Psychol.* **48**, 127–162 (2004).
13. J. M. Toro, J. B. Trobalón, Statistical computations over a speech stream in a rodent. *Percept. Psychophys.* **67**, 867–875 (2005).
14. A. E. Milne, C. I. Petkov, B. Wilson, Auditory and visual sequence learning in humans and monkeys using an artificial grammar learning paradigm. *Neuroscience* **389**, 104–117 (2018).
15. C. M. Conway, M. H. Christiansen, Modality-constrained statistical learning of tactile, visual, and auditory sequences. *J. Exp. Psychol. Learn. Mem. Cogn.* **31**, 24–39 (2005).
16. A. S. Reber, Implicit learning of artificial grammars. *J. Verbal Learning Verbal Behav.* **6**, 855–863 (1967).
17. C. Santolin, O. Rosa-Salva, L. Regolin, G. Vallortigara, Generalization of visual regularities in newly hatched chicks (*Gallus gallus*). *Anim. Cogn.* **19**, 1007–1017 (2016).
18. L. S. James, H. Sun, K. Wada, J. T. Sakata, Statistical learning for vocal sequence acquisition in a songbird. *Sci. Rep.* **10**, 1–18 (2020).
19. O. Menyhart, O. Kolodny, M. H. Goldstein, T. J. DeVogd, S. Edelman, Juvenile zebra finches learn the underlying structural regularities of their fathers' song. *Front. Psychol.* **6**, 571 (2015).
20. A. Avargués-Weber *et al.*, Different mechanisms underlie implicit visual statistical learning in honey bees and humans. *Proc. Natl. Acad. Sci. U.S.A.* **117**, 25923–25934 (2020).
21. E. K. Nicholls, N. Hempel de Ibarra, Bees associate colour cues with differences in pollen rewards. *J. Exp. Biol.* **217**, 2783–2788 (2014).
22. A. Fló, L. Benjamin, M. Palu, G. Dehaene-Lambertz, Statistical learning beyond words in human neonates. *eLife* **13**, RP101802 (2025).
23. A. Fló, L. Benjamin, M. Palu, G. Dehaene-Lambertz, Sleeping neonates track transitional probabilities in speech but only retain the first syllable of words. *Sci. Rep.* **12**, 4391 (2022).

24. H. Bulf, S. P. Johnson, E. Valenza, Visual statistical learning in the newborn infant. *Cognition* **121**, 127–132 (2011).
25. J. Fiser, R. N. Aslin, Statistical learning of new visual feature combinations by infants. *Proc. Natl. Acad. Sci. U.S.A.* **99**, 15822–15826 (2002).
26. N. Z. Kirkham, J. A. Slemmer, S. P. Johnson, Visual statistical learning in infancy: Evidence for a domain general learning mechanism. *Cognition* **83**, B35–B42 (2002).
27. J. R. Saffran, E. K. Johnson, R. N. Aslin, E. L. Newport, Statistical learning of tone sequences by human infants and adults. *Cognition* **70**, 27–52 (1999).
28. L. J. Batterink, S. Zhang, Simple statistical regularities presented during sleep are detected but not retained. *Neuropsychologia* **164**, 108106 (2022).
29. C. Xu *et al.*, Statistical learning in patients in the minimally conscious state. *Cereb. Cortex* **33**, 2507–2516 (2022).
30. L. Benjamin *et al.*, The role of conscious attention in auditory statistical learning: Evidence from patients with impaired consciousness. *iScience* **28**, 111591 (2025).
31. M. D. Hauser, E. L. Newport, R. N. Aslin, Segmentation of the speech stream in a non-human primate: Statistical learning in cotton-top tamarins. *Cognition* **78**, B53–B64 (2001).
32. M. Boros *et al.*, Neural processes underlying statistical learning for speech segmentation in dogs. *Curr. Biol.* **31**, 5512–5521.e5 (2021).
33. M. Takahashi, H. Yamada, K. Okanoya, Statistical and prosodic cues for song segmentation learning by Bengalese finches (*Lonchura striata* var. *domestica*). *Ethology* **116**, 481–489 (2010).
34. P. Perruchet, S. Desauty, A role for backward transitional probabilities in word segmentation? *Mem. Cognit.* **36**, 1299–1305 (2008).
35. B. Pelucchi, J. F. Hay, J. R. Saffran, Learning in reverse: Eight-month-old infants track backward transitional probabilities. *Cognition* **113**, 244–247 (2009).
36. K. Tummetshammer, D. Amso, R. M. French, N. Z. Kirkham, Across space and time: Infants learn from backward and forward visual statistics. *Dev. Sci.* **20**, e12474 (2016).
37. M. Buiatti, M. Peña, G. Dehaene-Lambertz, Investigating the neural correlates of continuous speech computation with frequency-tagged neuroelectric responses. *NeuroImage* **44**, 509–519 (2009).
38. A. D. Endress, Learning melodies from non-adjacent tones. *Acta Psychol. (Amst.)* **135**, 182–190 (2010).
39. E. L. Newport, M. D. Hauser, G. Spaepen, R. N. Aslin, Learning at a distance II. Statistical learning of non-adjacent dependencies in a non-human primate. *Cogn. Psychol.* **49**, 85–117 (2004).
40. R. Sonnweber, A. Ravnani, W. T. Fitch, Non-adjacent visual dependency learning in chimpanzees. *Anim. Cogn.* **18**, 733–745 (2015).
41. S. Mark, R. Moran, T. Parr, S. W. Kennerley, T. E. J. Behrens, Transferring structural knowledge across cognitive maps in humans and models. *Nat. Commun.* **11**, 1–12 (2020).
42. A. C. Schapiro, N. B. Turk-Browne, K. A. Norman, M. M. Botvinick, Statistical learning of temporal community structure in the hippocampus. *Hippocampus* **26**, 3–8 (2016).
43. E. A. Karuza, A. E. Kahn, S. L. Thompson-Schill, D. S. Bassett, Process reveals structure: How a network is traversed mediates expectations about its architecture. *Sci. Rep.* **7**, 1–9 (2017).
44. E. A. Karuza, A. E. Kahn, D. S. Bassett, Human sensitivity to community structure is robust to topological variation. *Complexity* **2019**, 8379321 (2019).
45. C. W. Lynn, A. E. Kahn, N. Nyema, D. S. Bassett, Abstract representations of events arise from mental errors in learning and memory. *Nat. Commun.* **11**, 2313 (2020).
46. X. Ren, H. Zhang, H. Luo, Dynamic emergence of relational structure network in human brains. *Prog. Neurobiol.* **219**, 102373 (2022).
47. E. Kakaei, S. Aleshin, J. Braun, Visual object recognition is facilitated by temporal community structure. *Learn. Amp Mem.* **28**, 148–152 (2021).
48. L. Benjamin, M. Sablé-Meyer, A. Fló, G. Dehaene-Lambertz, F. Al Roumi, Long-horizon associative learning explains human sensitivity to statistical and network structures in auditory sequences. *J. Neurosci.* **44**, e1369232024 (2024).
49. D. O. Hebb, *The Organization of Behavior* (Lawrence Erlbaum, 2002).
50. P. Dayan, Improving generalization for temporal difference learning: The successor representation. *Neural Comput.* **5**, 613–624 (1993).
51. A. Fló *et al.*, Newborns are sensitive to multiple cues for word segmentation in continuous speech. *Dev. Sci.* **22**, e12802 (2019).
52. E. S. Isbilen, S. M. McCauley, E. Kidd, M. H. Christiansen, Statistically induced chunking recall: A memory-based approach to statistical learning. *Cogn. Sci.* **44**, e12848 (2020).
53. A. E. Kahn, E. A. Karuza, J. M. Vettel, D. S. Bassett, Network constraints on learnability of probabilistic motor sequences. *Nat. Hum. Behav.* **2**, 936–947 (2018).
54. A. E. Kahn, D. S. Bassett, N. D. Daw, Trial-by-trial learning of successor representations in human behavior. *PLoS Comput. Biol.* **21**, e1013696 (2025).
55. C. W. Lynn, D. S. Bassett, How humans learn and represent networks. *Proc. Natl. Acad. Sci. U.S.A.* **117**, 29407–29415 (2020).
56. J. L. Elman, Finding structure in time. *Cogn. Sci.* **14**, 179–211 (1990).
57. G. Bellec *et al.*, A solution to the learning dilemma for recurrent networks of spiking neurons. *Nat. Commun.* **11**, 3625 (2020).
58. J. C. R. Whittington *et al.*, The Tolman-Eichenbaum machine: Unifying space and relational memory through generalization in the hippocampal formation. *Cell* **183**, 1249–1263.e23 (2020).
59. G. G. Ambrus *et al.*, When less is more: Enhanced statistical learning of non-adjacent dependencies after disruption of bilateral DLPFC. *J. Mem. Lang.* **114**, 104144 (2020).
60. L. J. Batterink, D. Choi, Optimizing steady-state responses to index statistical learning: Response to Benjamin and colleagues. *Cortex* **142**, 379–388 (2021).
61. L. J. Batterink, K. A. Paller, Statistical learning of speech regularities can occur outside the focus of attention. *Cortex* **115**, 56–71 (2019).
62. L. Benjamin, G. Dehaene-Lambertz, A. Fló, Remarks on the analysis of steady-state responses: Spurious artifacts introduced by overlapping epochs. *Cortex* **142**, 370–378 (2021).
63. H. Liu, T. A. Forest, K. Duncan, A. S. Finn, What sticks after statistical learning: The persistence of implicit versus explicit memory traces. *Cognition* **236**, 105439 (2023).
64. J. M. Toro, S. Sinnett, S. Soto-Faraco, Speech segmentation by statistical learning depends on attention. *Cognition* **97**, B25–B34 (2005).
65. I. Momennejad, Learning structures: Predictive representations, replay, and generalization. *Curr. Opin. Behav. Sci.* **32**, 155–166 (2020).
66. K. L. Stachenfeld, M. M. Botvinick, S. J. Gershman, The hippocampus as a predictive map. *Nat. Neurosci.* **20**, 1643–1653 (2017).
67. E. D. Thiessen, A. T. Kronstein, D. G. Hufnagle, The extraction and integration framework: A two-process account of statistical learning. *Psychol. Bull.* **139**, 792–814 (2013).
68. A. Christophe, E. Dupoux, J. Bertoncini, J. Mehler, Do infants perceive word boundaries? An empirical study of the bootstrapping of lexical acquisition. *J. Acoust. Soc. Am.* **95**, 1570–1580 (1994).
69. M. Shukla, M. Nespor, J. Mehler, An interaction between prosody and statistics in the segmentation of fluent speech. *Cogn. Psychol.* **54**, 1–32 (2007).
70. C. Kabdebon, M. Pena, M. Buiatti, G. Dehaene-Lambertz, Electrophysiological evidence of statistical learning of long-distance dependencies in 8-month-old preterm and full-term infants. *Brain Lang.* **148**, 25–36 (2015).
71. E. Marchetto, L. L. Bonatti, Words and possible words in early language acquisition. *Cogn. Psychol.* **67**, 130–150 (2013).
72. E. Marchetto, L. L. Bonatti, Finding words and word structure in artificial speech: The development of infants' sensitivity to morphosyntactic regularities. *J. Child Lang.* **42**, 873–902 (2015).
73. A. D. Endress, J. Mehler, The surprising power of statistical learning: When fragment knowledge leads to false memories of unheard words. *J. Mem. Lang.* **60**, 351–367 (2009).
74. F. Al Roumi, S. Marti, L. Wang, M. Amalric, S. Dehaene, Mental compression of spatial sequences in human working memory using numerical and geometrical primitives. *Neuron* **109**, 2627–2639.e4 (2021).
75. F. Al Roumi, S. Planton, L. Wang, S. Dehaene, Brain-imaging evidence for compression of binary sound sequences in human memory. *eLife* **12**, e84376 (2023).
76. S. Dehaene, F. Al Roumi, Y. Lakretz, S. Planton, M. Sablé-Meyer, Symbols and mental programs: A hypothesis about human singularity. *Trends Cogn. Sci.* **26**, 751–766 (2022).
77. S. Planton *et al.*, A theory of memory for binary sequences: Evidence for a mental compression algorithm in humans. *PLoS Comput. Biol.* **17**, e1008598 (2021).
78. J. Quilty-Dunn, N. Porot, E. Mandelbaum, The best game in town: The reemergence of the language-of-thought hypothesis across the cognitive sciences. *Behav. Brain Sci.* **46**, e261 (2022).
79. G. F. Marcus, K. J. Fernandes, S. P. Johnson, Infant rule learning facilitated by speech. *Psychol. Sci.* **18**, 387–391 (2007).
80. F. Bulgarelli, D. J. Weiss, Anchors aweigh: The impact of overlearning on entrenchment effects in statistical learning. *J. Exp. Psychol. Learn. Mem. Cogn.* **42**, 1621–1631 (2016).
81. A. L. Gebhart, R. N. Aslin, E. L. Newport, Changing structures in midstream: Learning along the statistical garden path. *Cogn. Sci.* **33**, 1087–1116 (2009).
82. J. A. Jungé, B. J. Scholl, M. M. Chun, How is spatial context learning integrated over signal versus noise? A primacy effect in contextual cueing. *Vis. Cogn.* **15**, 1–11 (2007).
83. L. Benjamin, M. Author, Data from "Long-horizon associative learning as a unifying framework for statistical learning across scales." Open Science Framework. <https://osf.io/g65d7/>. Deposited 13 March 2026.