

UC Santa Barbara

UC Santa Barbara Previously Published Works

Title

Tracking the topology of neural manifolds across populations

Permalink

<https://escholarship.org/uc/item/9rh8p4tj>

Journal

Proceedings of the National Academy of Sciences of the United States of America,
121(46)

ISSN

0027-8424

Authors

Yoon, Iris HR
Henselman-Petrusek, Gregory
Yu, Yiyi
[et al.](#)

Publication Date

2024-11-12

DOI

10.1073/pnas.2407997121

Copyright Information

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

Peer reviewed



Tracking the topology of neural manifolds across populations

Iris H. R. Yoon^{a,1} , Gregory Henselman-Petrusek^b, Yiyi Yu^c, Robert Ghrist^{d,e}, Spencer LaVere Smith^c, and Chad Giusti^{f,1}

Affiliations are included on p. 8.

Edited by Tatyana O. Sharpee, The Salk Institute and University of California, San Diego, La Jolla, CA; received April 22, 2024; accepted October 9, 2024, by Editorial Board Member Thomas D. Albright

Neural manifolds summarize the intrinsic structure of the information encoded by a population of neurons. Advances in experimental techniques have made simultaneous recordings from multiple brain regions increasingly commonplace, raising the possibility of studying how these manifolds relate across populations. However, when the manifolds are nonlinear and possibly code for multiple unknown variables, it is challenging to extract robust and falsifiable information about their relationships. We introduce a framework, called the method of analogous cycles, for matching topological features of neural manifolds using only observed dissimilarity matrices within and between neural populations. We demonstrate via analysis of simulations and in vivo experimental data that this method can be used to correctly identify multiple shared circular coordinate systems across both stimuli and inferred neural manifolds. Conversely, the method rejects matching features that are not intrinsic to one of the systems. Further, as this method is deterministic and does not rely on dimensionality reduction or optimization methods, it is amenable to direct mathematical investigation and interpretation in terms of the underlying neural activity. We thus propose the method of analogous cycles as a suitable foundation for a theory of cross-population analysis via neural manifolds.

circular coordinates | neural encoding | neural coding propagation | analogous cycles | persistent homology

Among the most successful models for coding in populations of biological or artificial neurons is the neural manifold model. In this model, individual neurons correspond to receptive fields, spatially localized regions in some metric space or manifold (Fig. 1*A*). The aggregate input to and state of the population corresponds to a point in the neural manifold, and a neuron is active precisely when that point is within its corresponding receptive field; a common model for the firing rate of each neuron is the tuning curve, given by a bump function supported within the receptive field (1–3). Historically, neural manifolds were discovered by a direct comparison of the activity of individual neurons to known stimuli or behaviors which carried known geometric structure (1, 4–6). Recent efforts have developed tools for detecting and studying the structure of neural manifolds intrinsically, without reference to external correlates (7–9). As we discover and characterize more intricate neural manifolds, the question now becomes understanding what they represent and how they are related to one another.

When a neural manifold is linear (Fig. 1*A*, *Left*), there are many robust methods available for extracting its structure from observations of the system (7, 10), and basic correlation analysis can be applied to compare such spaces. However, nonlinear neural manifolds are common (8, 11–14), and the most familiar among these are circular coordinates (Fig. 1*A*, *Right*). For example, simple cells in the primary visual cortex and head direction cells in the hippocampus encode a notion of angle using a circular coordinate system, while the quasi-periodic activity of pacemaker circuits encodes a circular coordinate temporally. Machine learning approaches have been employed to detect these structures in vivo and to compare them across populations (15–20). However, such approaches often require a good initial dimensionality estimate and reduction, and they aim to align local geometry (15–18). For systems such as grid and conjunctive cells in the entorhinal cortex, where different modules have incompatible geometries, such methods can be difficult to apply (6, 21). The authors posit that the development of a mathematically well-founded theory of neural computation in and among neural manifolds requires a bottom-up approach, in which quantitative tools for detecting and interacting with these systems derive from the structure of neural manifolds themselves. Recent work has demonstrated that tools from applied algebraic topology (22–24) provide an effective mathematical and computational framework for studying the intrinsic shape of neural population activity (Fig. 1*B*) (8, 25–27).

Significance

Neural manifolds provide a quantitative, geometric language for describing how the activity profiles of individual neurons combine to encode complex information. However, neurons in distinct neural populations may exhibit very different statistics in response to similar stimuli. As a result, the geometry of the corresponding neural manifolds may look quite different. We introduce the method of analogous cycles, which allows for robust and falsifiable inference of matchings between features of neural manifolds that code, at the population level, for related nonlinear topological features. This method is deterministic, does not involve dimensionality reduction or estimation, and requires only observations of intrinsic activity, without reference to external stimuli. We demonstrate its effectiveness on both simulations and in vivo experimental data.

Author contributions: I.H.R.Y., R.G., and C.G. designed research; I.H.R.Y. and C.G. performed research; I.H.R.Y. analyzed data; G.H.-P. developed software; Y.Y. and S.L.S. collected experimental data and provided neuroscientific expertise; and I.H.R.Y., G.H.-P., Y.Y., R.G., S.L.S., and C.G. wrote the paper.

The authors declare no competing interest.

This article is a PNAS Direct Submission. T.O.S. is a guest editor invited by the Editorial Board.

Copyright © 2024 the Author(s). Published by PNAS. This open access article is distributed under Creative Commons Attribution-NonCommercial-NoDerivatives License 4.0 (CC BY-NC-ND).

¹To whom correspondence may be addressed. Email: hyoon@wesleyan.edu or chad.giusti@oregonstate.edu.

This article contains supporting information online at <https://www.pnas.org/lookup/suppl/doi:10.1073/pnas.2407997121/-DCSupplemental>.

Published November 8, 2024.

Topological methods describe the structure of a neural manifold in terms of the intersections of receptive fields, providing a coarse summary of the shape of the space (*SI Appendix, section 1A2*). Such an encoding is natural in the context of population activity, where pairwise similarity of spike trains provides a proxy for overlap of receptive fields (Fig. 1*B* and *SI Appendix, Fig. S1*). This discards fine information about geometry in favor of robust representation of mesoscale, nonlocal information like connectivity. Finding precisely two distinct paths from any point to any other in the neural manifold, for example, indicates the presence of a circular coordinate (Fig. 1*B* and *SI Appendix, Fig. S1*). Encoding this structure as linear algebra, we obtain a tool, persistent homology (22–24, 28), for studying shape that detects and characterizes nonlinear *cycles* in neural manifolds using only observations of population activity (8, 9, 12, 25, 27, 29). These are represented as a persistence diagram, denoted PD, which summarizes the multiscale structure of the cycles and their relative significance* (Fig. 1*B, Right*).

Once we have detected a nonlinear neural manifold, we are left with the problem of assigning semantics. Without knowledge of the function of the population, how can we determine which, if any, external variables—stimuli, behavior, or structured activity of other neurons—are faithfully encoded by this structure? Our goal is not to assign semantics to individual neurons, but rather to use observations of population activity to determine which, if any, variables the population encodes. As the structure may be inherently high-dimensional, we cannot a priori isolate features to study. Thus, the methods must work simultaneously with all significant structure in both neural manifolds.

To address this need, in ref. 30, the authors introduced the mathematics underlying the method of analogous cycles (Fig. 1*C* and *SI Appendix, section 1A6*). Given the persistent homology of two related systems and a measure of cross-dissimilarity between the two systems, the method of analogous cycles encodes how families of receptive fields in each system witness those in the other (Fig. 1*C* and *SI Appendix, section 1A4*). Leveraging this encoding, it enumerates all possible relations between their topological features consistent with the cross-dissimilarity measure. This method thus describes how representations of complex information in neural population activity evolve across brain regions and allows us to assign semantics from observations of stimuli and behaviors. These results are amenable to theoretical mathematical analysis using existing theory from topology and geometry.

In this paper, we apply the method of analogous cycles to the study of coding in neural systems. We validate the tool in simulation studies against ground truth, establish statistical tests for significance of feature identification, and demonstrate how the method can be applied in both simulated and experimentally observed neural systems.

Results

Assigning Semantics to Topology of Neural Manifolds in Simulated Visual Systems. Our fundamental question is whether the mathematical framework in ref. 30 can identify shared circular features in neural systems. To test this, we first consider a sequence of simulated neural populations for which we have access to ground truth. We investigate how topological structure in feed forward networks designed to represent or discard features

of presented stimuli is matched across layers. We consider both networks that encode individual features of stimuli and those which synthesize inputs to describe more sophisticated structure. **Identifying circular features in neural activity induced by structured stimuli.** Can the analogous cycles method identify the features of a stimulus that are encoded by a population of neurons? To address this question, we developed a modified version of a standard grating stimulus with interesting topology and a computational model of simple cells in V1 (Fig. 2*A1*). Our stimuli are videos of a masked grating on a square torus,[†] (Fig. 2*B*). The location of the grating and the orientation of the grating vary continuously (*SI Appendix, section 2A1*), creating three intrinsic circular coordinates. We simulated simple cells via samples from an inhomogeneous Poisson process with rate given by a rectified dot product between a modified Gabor filter and the stimulus video (Fig. 2*C* and *SI Appendix, section 2A2*). This population of simple cells, by design, encodes all of the significant topological features of the stimulus videos, providing a collection of expected ground-truth matches among encoded features against which we can compare the output of the analogous cycles method.

For the video stimuli, we randomly sampled 400 out of 40,000 stimulus images, and we computed a pairwise dissimilarity matrix D_{stim} via the L_2 distance on the orientation and location of the circular mask (*SI Appendix, section 1B3*). For the simple cells, we computed a dissimilarity matrix D_{simp} via a windowed cross-correlation dissimilarity (*Materials and Methods*). From these two matrices, we computed persistence diagrams PD(stim) and PD(simp) (Fig. 2*E, Left, Center*, without colors). The persistence diagrams indicated that both the stimuli and the neural manifold for the simulated simple cell population have multiple circular features. In both cases, three of the features are significant, quantified using a quantile test on distance from the birth equals death line (*Materials and Methods*). However, simply observing that both populations carry circular features does not provide us with a method for comparing them. Are the circular structures observed in the activity of the simulated simple cells driven by the stimulus? If so, which feature of the stimulus is encoded by which circular feature of the simple cells?

The method of analogous cycles requires a measure of dissimilarity between the stimuli and the activity of each simple cell. To each video, we associated a binary vector with one entry per frame of the video, with value 1 at frame t if the presented stimulus video is displaying an image within a threshold distance in the space of images (*SI Appendix, section 1B4*). To compute the dissimilarity matrix $D_{\text{stim, simp}}$ for each image and simulated cell, we applied the windowed cross-correlation dissimilarity (*Materials and Methods*) to the corresponding binary vector and spike train. The method of analogous cycles computed with D_{stim} , D_{simp} , and $D_{\text{stim, simp}}$ as inputs matched three pairs of points between PD(stim) and PD(simp) (Fig. 2*E, Left, Center*, with colors). To validate, we compared the results to a geometric null model (*Materials and Methods*, Fig. 2*E, Left matrix*) and computed the likelihood of observing the teal, pink, and yellow pairs ($P < 0.005$, $N = 200$ trials).

We can verify the results of our computations by visualizing the matches. Visualizing the cycles detected in PD(stim) indicates that the points represented by the teal square, yellow pentagon, and pink star, respectively, represent orientation, x -coordinates, and y -coordinates in the stimulus video (*SI Appendix, Fig. S25*). Comparing the matched cycles of PD(simp) against the geometry

*While persistent homology can be used to study geometry of any dimension, in this paper, we focus on the case of 1-dimensional persistence diagrams, which reflect circular structure.

[†]A square region becomes a square torus by identifying the left and right edges and by identifying the top and bottom edges, as in the video game Asteroids.

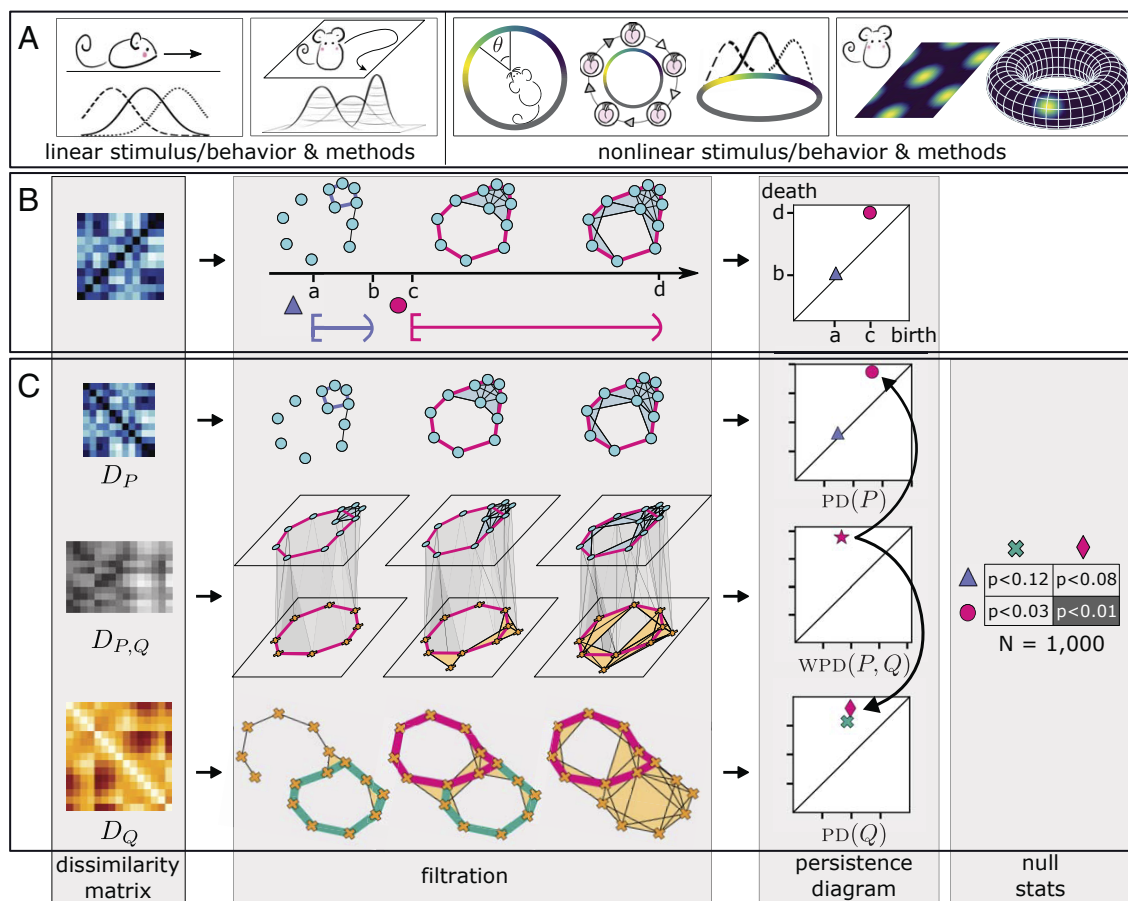


Fig. 1. Topological comparison of neural manifolds across populations. (A) Tuning curves assemble to represent linear (Left) and nonlinear (Right) neural manifolds. (B) Detection of circular features from a dissimilarity matrix via persistent homology. Given a dissimilarity matrix (Left), a filtration of simplicial complexes (Middle) encodes the system at varying thresholds. The birth and death thresholds for topological features (Middle) are summarized in a persistence diagram (Right). (C) The method of analogous cycles matches features. Given systems P , Q , with dissimilarity matrices D_P , D_Q , compute corresponding persistence diagrams $PD(P)$, $PD(Q)$ (Top, Bottom). For a cross-system dissimilarity matrix $D_{P,Q}$ (Middle row, Left), create a witness filtration (Middle row, Center), with features summarized in the witness persistence diagram (Middle row, Right). The point (star) indicates one shared feature between the two systems. The analogous cycles method matches any representations of this feature in $PD(P)$ (pink circle) and $PD(Q)$ (pink diamond) consistent with $D_{P,Q}$. The null model matching matrix (Far Right) gives the probability that the method returns a match between each pair of points in $PD(P)$ and $PD(Q)$ under the geometric null model (Materials and Methods). Shading at the (i, j) entry indicates that the corresponding pair is matched through the computed $WPD(P, Q)$.

induced by the stimuli provides verification of their semantics (SI Appendix, Fig. S26).

Tracking and falsifying feature propagation across populations. The previous experiment validates the method of analogous cycles for populations which faithfully encode the topology of the input space. However, neural computation across populations involves selection and synthesis of features. In real neural systems, we therefore expect that only a subset of features would propagate between populations, while new features that are not intrinsic to neural manifolds of upstream populations would be generated. This led us to ask whether this method correctly identifies which, if any, topological features are shared with other populations and which are novel.

To address this question, we extended our simulated visual system to include two additional simulated neural populations, each receiving input from the population of simulated simple cells (Fig. 2A). The neurons in these populations were trained to carry unimodal tuning curves for grating orientation and motion direction in the stimulus movies, respectively (Fig. 2D and SI Appendix, sections 2A3 and 2A4). In both cases, the coding properties of the population are well described by circular neural manifolds. The orientation-sensitive population selects a

single feature from those represented by the upstream simple cell population, while the motion-sensitive population synthesizes position information across a short time window, and thus does not reflect any existing feature described by the simple cell neural manifold. Persistence diagrams computed from the windowed cross-correlation dissimilarity matrices D_{ori} and D_{dir} for both populations consisted of a single point (Fig. 2E), confirming that the simulated population activity was well constrained to the expected neural manifolds.

We applied the method of analogous cycles to compare the neural manifold for the simulated simple cell population and both of these new circular neural manifolds. The cross-system dissimilarity matrices $D_{simp, ori}$, $D_{simp, dir}$ were computed by the windowed cross-correlation dissimilarity of spike trains (Materials and Methods). The method identified one pair of analogous points between $PD(simp)$ and $PD(ori)$ (Fig. 2E, Top row, teal), correctly identifying the teal diamond point in $PD(simp)$ as the feature of the simple cell neural manifold that corresponds to the circular feature in neural manifold of the new orientation-selective population. Comparison with the geometric null model (Fig. 2E, Top Right) indicates that this match is statistically significant ($P < 0.005$, $N = 200$ trials). On the

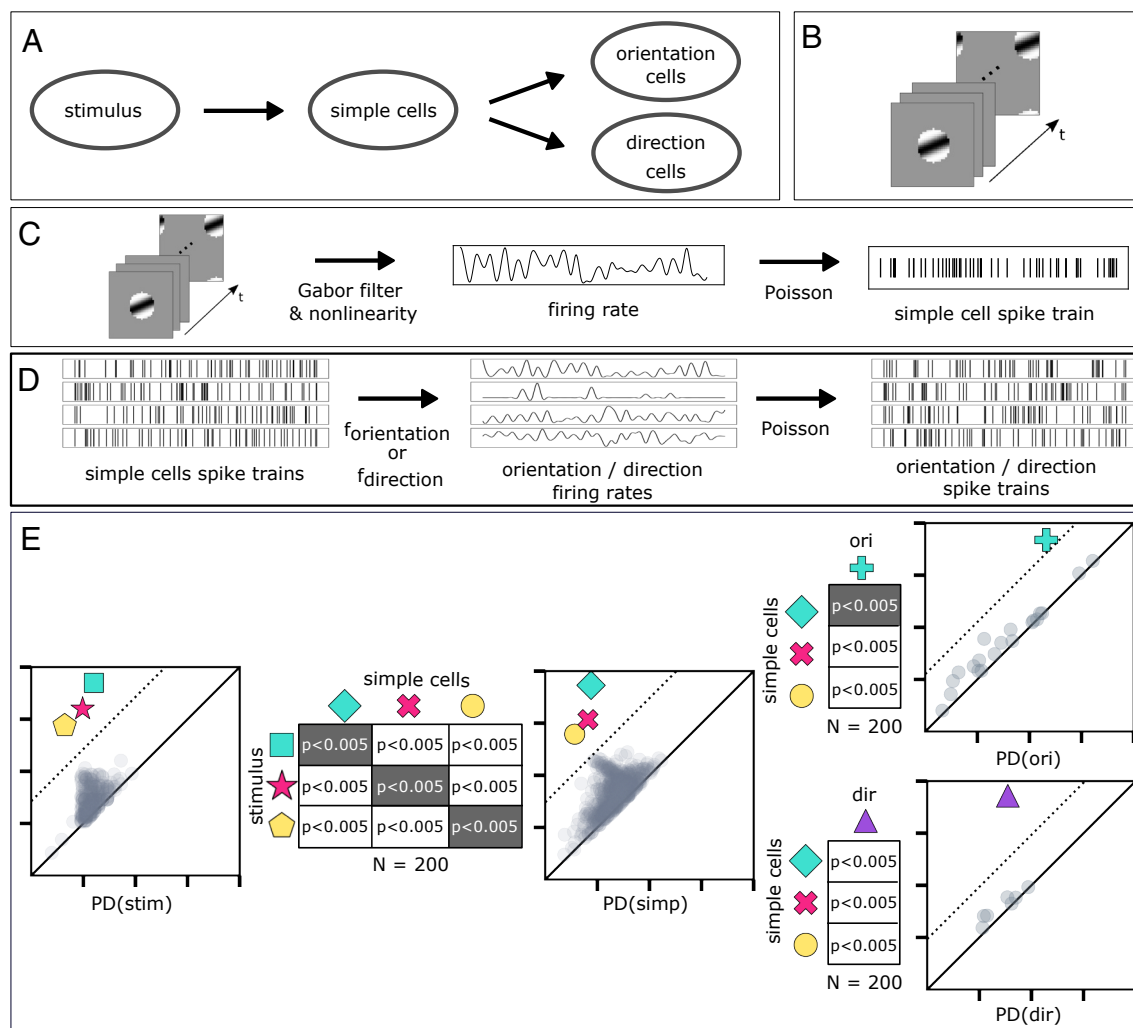


Fig. 2. Analogous cycles correctly match features across a simulated visual system. (A) Architecture of the simulation. Stimulus videos are presented to a population of simulated spiking simple cells. Spikes from these cells are presented to populations tuned to orientation and direction features of the stimulus. (B) Stimuli are 40,000-frame videos of a circular grating with fixed orientation, with the center moving in a fixed direction with fixed speed, looping to the opposite edge at boundaries. Orientations, starting locations, and movement directions are sampled uniformly. (C) Simulation of spiking simple cells. Firing rate is computed as the rectified dot product between a modified Gabor filter and frames of the stimulus video. Spikes are sampled from the resulting inhomogeneous Poisson process. (D) Firing rates of orientation-sensitive cells are given by $f_{\text{orientation}}$, which is described by a feed-forward neural network trained to approximate a unimodal tuning curve on orientations using the simple cell spike trains as input (SI Appendix, section 2A3). Output is obtained by sampling from the resulting firing rates using an inhomogeneous Poisson process. Direction cells are simulated similarly. (E) Analogous cycles match cycles across systems. (Left) All three significant circular features of the stimulus are matched to those of the simple cells. Matches are indicated by color. All matches are statistically significant. (Top Right) One circular feature (teal diamond) in the simple cell neural manifold is matched to the unique significant feature (teal cross) of orientation-sensitive population, implying a projection of encoded information from the simple cells. (Bottom Right) No matches are found between the simple cells and direction cells, indicating that the direction cells encode a novel feature.

other hand, the method of analogous cycles correctly identified no matches between PD(simp) and PD(dir) (Fig. 2 E, Bottom Right), indicating that the circular feature encoded by the direction-sensitive population must either be synthesized from upstream information but not intrinsically described by the earlier neural manifold, or projected from some unobserved population.

Disentangling Conjunctive Cell Coding in the Simulated Entorhinal Cortex. While the method of analogous cycles was effective in feature identification and selection in the simulated visual system, it is possible that there is some aspect of those simulations that was particularly well suited to topological analysis. As the method is intended to be a general tool, we next asked whether the choice of the model or the architecture of the simulation impacts the performance of the method. In order to address this concern, we applied the method to simulated

populations of entorhinal cortex neurons (Fig. 3A). Using experimental recordings of location and head directions from rats engaging in foraging behavior taken from ref. 8, we simulated populations of grid cells, head-direction cells, and conjunctive cells. Grid cells, located in the dorsocaudal medial entorhinal cortex, are activated when an animal's position coincides with a regular grid that spans the environment (6). We simulated grid cell firing rates using the continuous attractor network model, as in ref. 31. Head direction (HD) cells, which fire when an animal's head is aligned with a specific direction, were simulated using selected tuning curves (SI Appendix, section 2B). Conjunctive grid cells, a population of grid cells whose firing is modulated by a preferred head direction, can be found in the deep layers of the medial entorhinal cortex (32, 33). We computed the firing rate of each such conjunctive cell as the minimum firing rate of a model grid cell and a model HD cell to a given stimulus (Fig. 4B).

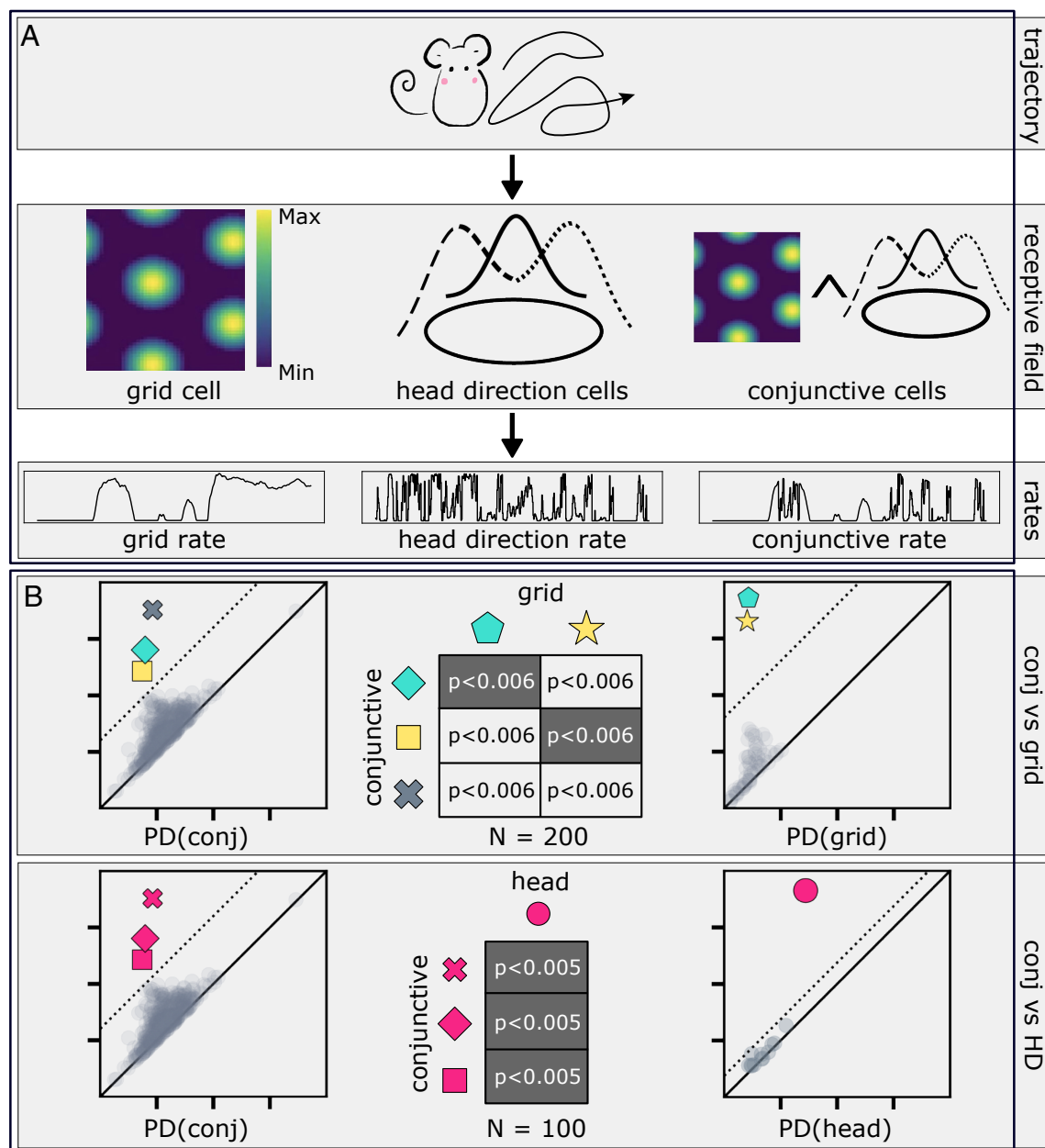


Fig. 3. Analogous cycles assign semantics to simulated navigational system. (A) Firing rate simulation. Given a rat trajectory, the grid cells are simulated using a continuous attractor networks model. The head-direction (HD) cells are simulated by tuning curves. The conjunctive cells firing rates are computed by the minimum firing rates of the grid cells and HD cells. (B) (Top) Analogous cycles between the grid and conjunctive cells indicated by the colors. The teal and yellow points in the conjunctive PD encode the torus arising from the grid cell organization. (Bottom) Analogous cycles between the HD and conjunctive cells. The single significant feature in PD(head) is analogous to a combination of the significant features in PD(conj). Since the diamond and square points of the conjunctive PD are analogous to the two points in the grid PD, we conclude that the cross point of the conjunctive PD must encode the cyclicity of HD cells.

It has recently been experimentally verified that grid cells carry a toroidal neural manifold (8, 34), well-parameterized by two independent circular coordinates, while HD cells have a circular neural manifold (35). We thus expect conjunctive cells to have a neural manifold parameterized by three independent circular features. We computed internal dissimilarity matrices D_{grid} , D_{HD} , D_{conj} , and cross-system dissimilarity matrices $D_{\text{grid, conj}}$, $D_{\text{HD, conj}}$ using the windowed cross-correlation dissimilarity (Materials and Methods), and obtained persistence diagrams PD(grid), PD(head), and PD(conj) that verify our expected count of circular features in the three simulated populations (Fig. 3B).

We applied the analogous cycles method to the populations of grid and conjunctive cells to test whether the method correctly matches the two circular features of the conjunctive cells that arise from the grid cells. The method identifies two pairs of analogous cycles between the grid and conjunctive persistence diagrams, indicated by the pairs of teal and yellow points in Fig. 3B. Comparison with the geometric null model (Fig. 3B) indicates that a randomly arising match between either the teal points or the yellow points are both highly unlikely ($P < 0.006$, $N = 200$ trials). We thus conclude that these two features of the conjunctive cell stimulus space represent the toroidal coordinates described by the constituent grid cells.

We then applied the analogous cycles method between the conjunctive and the HD cells. As illustrated in Fig. 3B, the method matches the significant feature (pink circle) in PD(head) to the three significant features in the PD(conj). Such a many-to-one match involves a linear combination of the corresponding cycles, indicating that there is some interaction among the three circular coordinates in the chosen cycle basis for the conjunctive cells (SI Appendix, Figs. S27 and S28). Here, we conclude that the chosen cycle basis for PD(conj) correlated the head direction and direction of motion. One can verify this deduction by constructing a change of basis for the cycles in PD(conj) which isolates the match to the cycle which is unmatched to points in PD(grid). As before, comparing with the geometric null model (Fig. 3B) indicates that these matches are statistically significant ($P < 0.005$, $N = 200$ trials).

Matching In Vivo Neural Manifolds for the Primary Visual Cortex and Anterolateral Region. Finally, we asked whether the analogous cycles method is robust enough to identify shared features of neural manifolds in experimental data. To address this question, we considered in vivo recordings from two mouse neural systems that are known to share coding properties: the primary visual cortex (V1) and the anterolateral (AL) region (Fig. 4A) (36–38).

We extracted single-cell spike trains from two-photon calcium imaging traces in both regions from a head-fixed mouse (Fig. 4A). Recordings were made while presenting black–white full contrast drifting gratings. Four different orientations were presented twice with opposite drifting directions for twenty trials (Fig. 4B, Materials and Methods). We aggregated the spike trains across

the trials (Fig. 4B) and performed preprocessing steps to remove neurons that fire uniformly and unreliably (SI Appendix, section 2C1). All internal and cross-system dissimilarity matrices were computed using the windowed cross-correlation dissimilarity (Materials and Methods). We computed persistence diagrams PD(V1) and PD(AL) for the two regions, both of which contained only two points (Fig. 4C). For such small persistence diagrams, our usual method of selecting a significance threshold for points in a persistence diagram cannot be applied. In this case, we use an alternative significance threshold given by the largest lifetime from a sample of persistence diagrams computed from synthetic spike trains that match the observed population statistics (Materials and Methods).

We applied the analogous cycles method to these two diagrams, and identified a matched pair of points as shown in Fig. 4C. Comparison to the geometric null model (Fig. 4C) indicates that the identified analogous pair is statistically significant ($P < 0.005$, $N = 200$ trials). The fact that V1 and AL neurons encode a matched feature is consistent with prior observations that neurons in V1 and AL share coding properties for orientation of visual stimuli (39, 40). Visualizing the spike trains for neurons in V1 and AL that constitute the matched cycles shows these subpopulations of neurons in V1 and AL have similar spiking profiles, without reference to stimuli (SI Appendix, Fig. S29).

Discussion

In this study, we introduced and validated the method of analogous cycles, a method for comparing topological structure in neural manifolds across populations. We demonstrated that the method reliably matches related circular coordinates between neural manifolds and rejects matching cycles that are not related. The method utilizes only the data of a cross-dissimilarity matrix to deterministically identify matches. Furthermore, the method works in the context of multiple simultaneously represented features, providing tools for comparing complex neural representations without descending to the level of activity of individual neurons. While we did not investigate such data in this study, these methods can be applied to other high-dimensional neural manifolds with nontrivial topology, such as spheres. Together, these features provide a robust tool set for assigning semantics to nonlinear structure within neural manifolds, and understanding how structures interact across populations.

In contrast to some prior studies that compare neural coding across populations, the method does not require any a priori knowledge of encoded variables or the use of optimization methods (15–17, 41). While some existing studies on information transmission across neural populations focus on anatomical connections (40, 42, 43) and spike train propagation (44–47), this method provides mathematically robust tools for studying and falsifying the propagation of nonlinear structure in codes at the population-level. Unsupervised approaches such as manifold learning have been successful in comparing neural manifolds across populations, but such analyses are hindered when alignment of local geometry is difficult or such an alignment does not exist (6, 15–18, 21). Prior studies using topological methods (8, 9, 12, 25, 27, 29) have been effective at describing structure of neural manifolds, but these did not provide tools for studying multisystem structure.

Looking forward, many existing methods perform dimensionality reduction on population vectors to mitigate computations and the effects of noise. Our method deals with dissimilarity between neural units, providing interpretability but leaving in

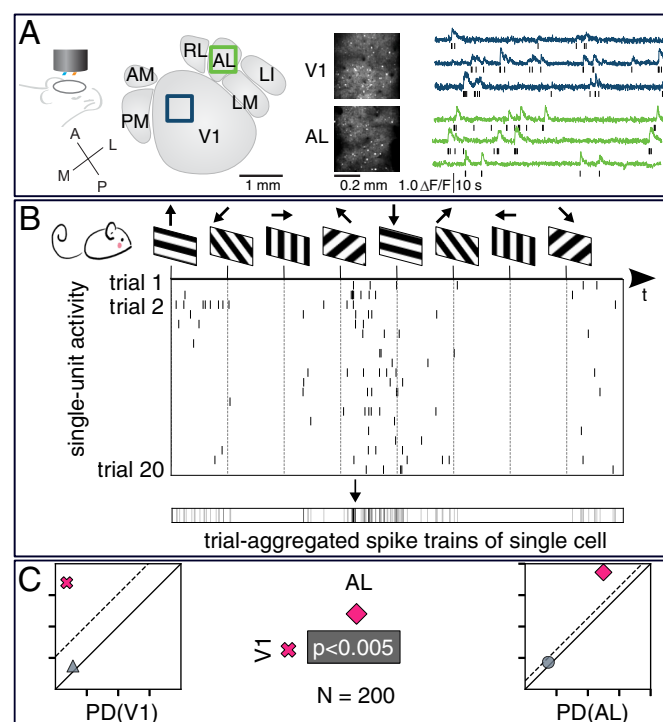


Fig. 4. Analogous cycles for neural coding propagation on experimental data. (A) Schematic (Top) of dual region two-photon calcium imaging in primary (V1) and AL visual areas (Left). Spike trains were inferred from the calcium sensor dynamics (Right). (B) The mouse was presented with a video consisting of drifting gratings of four orientations. Each stimulus was repeated 20 times. For each cell, the trial-aggregated spike train was computed. (C) The analogous cycles method identifies a pair of points in the V1 and AL persistence diagram that encodes the same circular information.

place these concerns. Extending the method to work with dimensionality-reduced population vectors, as used in ref. 8, would allow us to integrate strengths of existing methods with analogous cycles. For example, many studies identify network states in aggregate population vectors, assign meanings to each state using external stimuli/behavior, and then interpret neural systems via network state. A suitably augmented analogous cycles method could identify families of network states more intricate than clusters. Furthermore, while the method presented in this paper identifies geometrically independent features across neural manifolds, in some cases where those features have complex interactions, the algorithmically selected match may not align with intuition, and applying significance thresholding may return false negatives for matching. We give an example of such a pair of neural manifolds in *SI Appendix, section 3E* and describe an ad hoc method for correcting such misalignments. Developing formally justified methods for ameliorating these issues will require further research.

In summary, the analogous cycles method allows us to study how nonlinear coordinate systems encoded by distinct populations of neurons are related, without estimating tuning curves or relying on known stimuli or behavior. Complementing recent developments in multiregion imaging techniques (48, 49), the method provides avenues for the study of information flow and computation across brain regions and for comparison of coding properties between homologous regions in distinct organisms.

Materials and Methods

Dissimilarity between Spike Trains and Firing Rates. Let $\vec{x}$ and $\vec{y}$ be vectors of the same dimensionality that each represent either a spike train or a firing rate. We define the similarity between $\vec{x}$ and $\vec{y}$ to be the normalized sum of shifted convolutions, for a limited range of displacement values specified by the shift parameter ℓ :

$$\text{Sim}_\ell(\vec{x}, \vec{y}) = \frac{\sum_{n=-\ell}^{\ell} \sum_m \vec{x}_m \vec{y}_{m+n}}{\|\vec{x}\|_2 \|\vec{y}\|_2}.$$

Given a collection of time series $P = \{\vec{x}_i\}_{i=1}^K$, we compute pairwise dissimilarity as windowed cross-correlation dissimilarity

$$\text{Dis}_\ell(\vec{x}, \vec{y}) = \begin{cases} 1 - \frac{\text{Sim}_\ell(\vec{x}, \vec{y})}{M} & \text{if } \vec{x} \neq \vec{y} \\ 0 & \text{if } \vec{x} = \vec{y} \end{cases}$$

where $M = \max_{\vec{x}, \vec{y} \in P} \text{Sim}_\ell(\vec{x}, \vec{y})$. See *SI Appendix, section 1B2* for a thorough discussion on dissimilarity. A comparison of this dissimilarity measure for varying ℓ and a comparison to existing spike train dissimilarities can be found in *SI Appendix, section 3C*.

Persistent Homology. For a neural population or system P with N elements, we perform topological data analysis using an $N \times N$ pairwise dissimilarity matrix D_P . Let $\{\epsilon_i\}_{i=1}^m$ be a collection of parameters. We construct a filtration of simplicial complexes

$$X_{\epsilon_1} \subseteq X_{\epsilon_2} \subseteq \dots \subseteq X_{\epsilon_m},$$

where X_{ϵ_k} consists of N vertices and has n -simplex $[v_0, \dots, v_n]$ precisely when all pairwise dissimilarity among the listed elements is at most ϵ_k . Computing homology in dimension 1 with $\mathbb{Z}_2$ coefficients, we obtain a sequence of vector spaces summarizing the circular features in each simplicial complex at each parameter. We track the birth and death of such circular features via a persistence diagram, denoted $\text{PD}(P)$, where a feature that is born at parameter b and dies at d is represented by a point in the plane with coordinates (b, d) . See *SI Appendix, section 1A3* for a gentler exposition and *SI Appendix, section 3D* for experiments establishing robustness of persistence diagrams against various noise in simulations.

Significant Points on Persistence Diagrams. The points on a persistence diagram that are far from the diagonal line represent mesoscale features of the system, while the remaining points are often considered noise. To identify important features, we compute the empirical distribution of the lifetimes $d - b$ of all points, and threshold at $Q_3 + 3(Q_3 - Q_1)$, where Q_i is the i th quartile. We indicate this threshold value as a dotted line on the persistence diagram, taking any point above the threshold to be significant. A comparison of methods for identifying significant points in persistence diagrams is given in *SI Appendix, section 3A*.

If a persistence diagram contains insufficiently many points to apply the nonparametric statistical test to the distributions of lifetimes, we can apply an alternative test for significance developed in ref. 8. We sample a population of shuffled spike trains by binning our data spike trains into 0.0751 s bins and randomly permuting the entries of the resulting vector. We compute a persistence diagram for this population as for the data. Applying this process $N = 10,000$ times, we obtain a distribution of lifetimes, and take our significance threshold to be the largest observed lifetime. In *SI Appendix, section 3B*, we detail alternative tests for significance one can employ when dealing with small persistence diagrams.

Analogous Cycles. Let P, Q be two systems. The inputs to the analogous cycles method are the internal dissimilarity matrices D_P, D_Q , and the cross-system dissimilarity matrix $D_{P,Q}$. Let $\text{PD}(P)$ and $\text{PD}(Q)$ denote the dimension-1 persistence diagrams of D_P and D_Q . From $D_{P,Q}$, we construct a sequence of witness complexes

$$W_{P,Q}^{\psi_1} \subseteq W_{P,Q}^{\psi_2} \subseteq \dots \subseteq W_{P,Q}^{\psi_t}.$$

The witness complex $W_{P,Q}^{\psi_k}$ is a simplicial complex with vertices P and an n -simplex $[p_0, \dots, p_n]$ precisely when there exists a point $q \in Q$ such that $D_{P,Q}[p_j, q] \leq \psi_k$ for all $j \in \{0, \dots, n\}$. We compute the homology in dimension 1 with $\mathbb{Z}_2$ coefficients and summarize the results as the witness persistence diagram, denoted $\text{WPD}(P, Q)$ (50). By the Dowker duality theorem, $\text{WPD}(P, Q)$ and $\text{WPD}(Q, P)$ are identical (51, 52). See *SI Appendix, section 1A4* for details.

For each significant point $w \in \text{WPD}(P, Q)$, we construct auxiliary complexes which admit inclusions to the filtration of complexes on P (resp. Q) and the witness complex at a maximal parameter $\psi(w)$. We leverage induced maps on homology to construct a linear system whose solutions enumerate matches $p_w \subseteq \text{PD}(P)$ and $q_w \subseteq \text{PD}(Q)$ through w . For details, see *SI Appendix, section 1A6* and ref. 30. These (possibly empty) collections of significant points in the persistence diagrams are referred to as matched, or analogous, cycles. The output is illustrated via color-coordination of points in $\text{PD}(P)$ and $\text{PD}(Q)$. For a visualization of the output including the witness persistence diagrams, see *SI Appendix, Fig. S24*.

Significance of Analogous Cycles from Geometric Null Models. Random or shuffled dissimilarity matrices provide a null model inconsistent with any geometry and empirically always produce no significant matches. A better null model would assume an underlying shared geometry consistent with the number of cycles observed in each system, in which features are randomly aligned.

To construct such a null model in the relevant 1-dimensional case, let P, Q be observed systems with n_P and n_Q elements, and let D_P, D_Q , and $D_{P,Q}$ be the dissimilarity matrices. Let m_P and m_Q denote the number of significant points in $\text{PD}(P)$ and $\text{PD}(Q)$, and assume $m_P \geq m_Q$. Construct a m_P -dimensional torus $T_P^{\frac{1}{2}}$ and let P' be a uniform sample of n_P points from that torus. Select an m_Q -dimensional torus embedded within T_P , and let Q' denote a sample of n_Q points from this subspace. Take $D_{P',Q'}$ to be the cross-distance matrix for these points. Let $\{(p_w, q_w) | w \in \text{WPD}(P', Q') \text{ significant}\}$ be the output of the method of analogous cycles applied to D_P, D_Q and $D_{P',Q'}$ for a sample from this null model. For each $p_i \in \text{PD}(P)$ and $q_j \in \text{PD}(Q)$, let $F_{p,q}$ be the number of significant $w \in \text{WPD}(P', Q')$ such that $p_i \in p_w$ and $q_j \in q_w$, and summarize

[‡] An m -dimensional torus has m independent circular coordinates and can be obtained by taking an m -dimensional cube and identifying its opposite faces.

the outputs in a frequency matrix F whose (i, j) -entry is F_{p_i, q_j} . We repeat this sampling procedure $N \geq 100$ times and report the empirical probability matrix $\frac{F}{N}$ to provide nonparametric P -values for matches.

Experimental Data. All animal procedures and experiments were approved by the Institutional Animal Care and Use Committee of the University of North Carolina at Chapel Hill or the University of California Santa Barbara and performed in accordance with the regulation of the US Department of Health and Human Services. Cranial windows were surgically implanted in mice expressing the fluorescent calcium sensor protein GCaMP6s (53) in layer 2/3 pyramidal neurons in the primary visual cortex (V1) and higher visual area AL (54, 55). Large field-of-view two-photon calcium imaging (56, 57) measured neuronal activity responses to visual stimuli. Fluorescence dynamics were converted to inferred spike trains (58). Visual stimuli were displayed on a 60 Hz LCD monitor (9.2 cm $\times$ 15 cm). All stimuli were displayed in full contrast. Stimuli consisted of drifting square-wave gratings. We repeated the stimulus presentation 20 times. The experimental data are presented in ref. 55. See *SI Appendix, section 2C* for data preprocessing steps.

Computations. Code and experimental data can be found at the following Github repository https://github.com/irishryoon/analogous_neural. All persistent homology computations are performed via Eirene (59).

Data, Materials, and Software Availability. Code and experimental data have been deposited in https://github.com/irishryoon/analogous_neural (60).

ACKNOWLEDGMENTS. This work was supported by NSF DMS 1854683 (I.H.R.Y. and C.G.), NSF DMS 1854748 (G.H.-P.), AFOSR FA9550-21-1-0334 (R.G.), and AFOSR FA9550-21-1-0266 (C.G.). We thank Carina Curto and Vladimir Itskov for helpful conversations.

Author affiliations: ^aDepartment of Mathematics and Computer Science, Wesleyan University, Middletown, CT 06459; ^bPacific Northwest National Laboratory, Richland, WA 99352; ^cDepartment of Electrical and Computer Engineering, University of California, Santa Barbara, CA 93106; ^dDepartment of Mathematics, University of Pennsylvania, Philadelphia, PA 19104; ^eDepartment of Electrical & Systems Engineering, University of Pennsylvania, Philadelphia, PA 19104; and ^fDepartment of Mathematics, Oregon State University, Corvallis, OR 97331

1. D. H. Hubel, T. N. Wiesel, Receptive fields of single neurones in the cat's striate cortex. *J. Physiol.* **148**, 572–591 (1959).
2. J. O'Keefe, N. Burgess, Geometric determinants of the place fields of hippocampal neurons. *Nature* **381**, 425–428 (1996).
3. C. M. Niell, M. P. Stryker, Highly selective receptive fields in mouse visual cortex. *J. Neurosci.* **28**, 7520–7536 (2008).
4. R. Ben-Yishai, R. L. Bar-Or, H. Sompolinsky, Theory of orientation tuning in visual cortex. *Proc. Natl. Acad. Sci. U.S.A.* **92**, 3844–3848 (1995).
5. J. O'Keefe, J. O. Dostrovsky, The hippocampus as a spatial map. Preliminary evidence from unit activity in the freely-moving rat. *Brain Res.* **34**, 171–175 (1971).
6. T. Hafting, M. Fyhn, S. Molden, M. B. Moser, E. I. Moser, Microstructure of a spatial map in the entorhinal cortex. *Nature* **436**, 801–806 (2005).
7. A. Rubin *et al.*, Revealing neural correlates of behavior without behavioral measurements. *Nat. Commun.* **10**, 4745 (2019).
8. R. J. Gardner *et al.*, Toroidal topology of population activity in grid cells. *Nature* **602**, 123–128 (2022).
9. E. Rybakken, N. Baas, B. Dunn, Decoding of neural data using cohomological feature extraction. *Neural Comput.* **31**, 68–93 (2019).
10. V. Villette, A. Malvache, T. Tressard, N. Dupuy, R. Cossart, Internally recurring hippocampal sequences as a population template of spatiotemporal information. *Neuron* **88**, 357–366 (2015).
11. A. De, R. Chaudhuri, Common population codes produce extremely nonlinear neural manifolds. *Proc. Natl. Acad. Sci. U.S.A.* **120**, e2305853120 (2023).
12. R. Chaudhuri, B. Gerçek, B. Pandey, A. Peyrache, I. Fiete, The intrinsic attractor manifold and population dynamics of a canonical cognitive circuit across waking and sleep. *Nat. Neurosci.* **22**, 1512–1520 (2019).
13. H. Zhang, P. D. Rich, A. K. Lee, T. O. Sharpee, Hippocampal spatial representations exhibit hyperbolic geometry that expands with experience. *Nat. Neurosci.* **26**, 131–139 (2023).
14. Y. Zhou, B. H. Smith, T. O. Sharpee, Hyperbolic geometry of the olfactory space. *Sci. Adv.* **4**, eaq1458 (2018).
15. A. D. Degenhart *et al.*, Stabilization of a brain-computer interface via the alignment of low-dimensional spaces of neural activity. *Nat. Biomed. Eng.* **4**, 672 (2020).
16. M. Ganjali, A. Mehridehnavi, S. Rakhshani, A. Khorasani, Unsupervised neural manifold alignment for stable decoding of movement from cortical signals. *Int. J. Neural Syst.* **34**, 2450006 (2024).
17. J. A. Gallego *et al.*, Cortical population activity within a preserved neural manifold underlies multiple motor behaviors. *Nat. Commun.* **9**, 4233 (2018).
18. J. A. Gallego, M. G. Perich, R. H. Chowdhury, S. A. Solla, L. E. Miller, Long-term stability of cortical population dynamics underlying consistent behavior. *Nat. Neurosci.* **23**, 260–270 (2020).
19. C. Pandarinath *et al.*, Inferring single-trial neural population dynamics using sequential auto-encoders. *Nat. Methods* **15**, 805–815 (2018).
20. C. Pandarinath *et al.*, Latent factors and dynamics in motor cortex and their application to brain-machine interfaces. *J. Neurosci.* **38**, 9390–9401 (2018).
21. H. Stensola *et al.*, The entorhinal grid map is discretized. *Nature* **492**, 72–78 (2012).
22. R. Ghrist, Barcodes: The persistent topology of data. *Bull. Am. Math. Soc.* **45**, 61–75 (2008).
23. G. E. Carlsson, Topology and data. *Bull. Am. Math. Soc.* **46**, 255–308 (2009).
24. H. Edelsbrunner, J. Harer, Persistent homology—A survey. *Discrete Comput. Geom.* **453**, 1–26 (2008).
25. C. Curto, V. Itskov, Cell groups reveal structure of stimulus space. *PLoS Comput. Biol.* **4**, e1000205 (2008).
26. Y. Dabaghian, F. Mémoli, L. Frank, G. Carlsson, A topological paradigm for hippocampal spatial map formation using persistent homology. *PLoS Comput. Biol.* **8**, 1–14 (2012).
27. C. Giusti, E. Pastalkova, C. Curto, V. Itskov, Clique topology reveals intrinsic geometric structure in neural correlations. *Proc. Natl. Acad. Sci. U.S.A.* **112**, 13455–13460 (2015).
28. H. Edelsbrunner, D. Letscher, A. Zomorodian, "Topological persistence and simplification" in *Proceedings 41st Annual Symposium on Foundations of Computer Science* (2000), pp. 454–463.
29. G. Singh *et al.*, Topological analysis of population activity in visual cortex. *J. Vis.* **8**, 11.1–11.18 (2008).
30. H. R. Yoon, R. Ghrist, C. Giusti, Persistent extension and analogous bars: Data-induced relations between persistence barcodes. *J. Appl. Comput. Topol.* **7**, 571–617 (2023).
31. J. Covey *et al.*, Recurrent inhibitory circuitry as a mechanism for grid formation. *Nat. Neurosci.* **16**, 318–324 (2013).
32. K. Gerle *et al.*, Grid cells are modulated by local head direction. *Nat. Commun.* **11**, 4228 (2020).
33. F. Sargolini *et al.*, Conjunctive representation of position, direction, and velocity in entorhinal cortex. *Science (New York, N.Y.)* **312**, 758–762 (2006).
34. A. Guanella, D. Kiper, P. Verschure, A model of grid cells based on a twisted torus topology. *Int. J. Neural Syst.* **17**, 231–240 (2007).
35. L. Giocomo *et al.*, Topography of head direction cells in medial entorhinal cortex. *Curr. Biol.* **24**, 252–262 (2014).
36. Y. Yu, J. N. Stirman, C. R. Dorsett, S. L. Smith, Selective representations of texture and motion in mouse higher visual areas. *Curr. Biol.* **32**, 2810–2820 (2022).
37. M. L. Andermann, A. M. Kerlin, D. K. Roumis, L. L. Glickfeld, R. C. Reid, Functional specialization of mouse higher visual cortical areas. *Neuron* **72**, 1025–1039 (2011).
38. Q. Wang, A. Burkhalter, Area map of mouse visual cortex. *J. Comp. Neurol.* **502**, 339–357 (2007).
39. J. H. Marshel, M. Garrett, I. Nauhaus, E. M. Callaway, Functional specialization of seven mouse visual cortical areas. *Neuron* **72**, 1040–1054 (2011).
40. L. Glickfeld, M. Andermann, V. Bonin, R. C. Reid, Cortico-cortical projections in mouse visual cortex are functionally target specific. *Nat. Neurosci.* **16**, 219–226 (2013).
41. M. Dabagia, K. P. Kording, E. L. Dyer, Aligning latent representations of neural activity. *Nat. Biomed. Eng.* **7**, 337–343 (2023).
42. J. Lanciego, F. Wouterlood, A half century of experimental neuroanatomical tracing. *J. Chem. Neuroanat.* **42**, 157–183 (2011).
43. L. Sinich, G. Blasdel, Oriented axon projections in primary visual cortex of the monkey. *J. Neurosci. Soc. Neurosci.* **21**, 4416–4426 (2001).
44. T. Vogels, L. Abbott, Signal propagation and logic gating in networks of integrate-and-fire neurons. *J. Neurosci. Soc. Neurosci.* **25**, 10786–10795 (2005).
45. M. Diesmann, M. O. Gewaltig, A. Aertsen, Stable propagation of spikes in cortical neural networks. *Nature* **402**, 529–533 (2000).
46. M. O. Gewaltig, M. Diesmann, A. Aertsen, Propagation of cortical synfire activity: Survival probability in single trials and stability in the mean. *Neural Networks* **14**, 657–673 (2001).
47. S. Moldakarimov, M. Bazhenov, T. J. Sejnowski, Feedback stabilizes propagation of synchronous spiking in cortical neural networks. *Proc. Natl. Acad. Sci. U.S.A.* **112**, 2545–2550 (2015).
48. J. J. Jun *et al.*, Fully integrated silicon probes for high-density recording of neural activity. *Nature* **551**, 232–236 (2017).
49. C. K. Kim, A. Adhikari, K. Deisseroth, Integration of optogenetics with complementary methodologies in systems neuroscience. *Nat. Rev. Neurosci.* **18**, 222–235 (2017).
50. V. de Silva, G. Carlsson, "Topological estimation using witness complexes" in *SPBG'04 Symposium on Point Based Graphics 2004* (2004).
51. C. H. Dowker, Homology groups of relations. *Ann. Math.* **56**, 84–95 (1952).
52. S. Chowdhury, F. Mémoli, A functorial Dowker theorem and persistent homology of asymmetric networks. *J. Appl. Comput. Topol.* **2**, 115–175 (2018).
53. T. W. Chen *et al.*, Ultrasensitive fluorescent proteins for imaging neuronal activity. *Nature* **499**, 295–300 (2013).
54. Y. Yu, J. N. Stirman, C. R. Dorsett, S. L. Smith, Selective representations of texture and motion in mouse higher visual areas. *Curr. Biol.* **32**, 2810–2820 (2022).
55. Y. Yu, J. N. Stirman, C. R. Dorsett, S. L. Smith, Visual information is broadcast among cortical areas in discrete channels. *eLife* **13**, RP97848 (2024).
56. C. H. Yu, J. N. Stirman, Y. Yu, R. Hira, S. L. Smith, Diesel2p mesoscope with dual independent scan engines for flexible capture of dynamics in distributed neural circuitry. *Nat. Commun.* **12**, 6639 (2021).
57. J. N. Stirman, I. T. Smith, M. W. Kudenov, S. L. Smith, Wide field-of-view, multi-region, two-photon imaging of neuronal activity in the mammalian brain. *Nat. Biotechnol.* **34**, 857–862 (2016).
58. E. A. Pnevmatikakis *et al.*, Simultaneous denoising, deconvolution, and demixing of calcium imaging data. *Neuron* **89**, 285–299 (2016).
59. G. Henselman, R. Ghrist, Matroid filtrations and computational persistent homology. arXiv [Preprint] (2016). <https://doi.org/10.48550/arXiv.1606.00199> (Accessed 8 January 2023).
60. I. Yoon, Analogous cycles. GitHub. https://github.com/irishryoon/analogous_neural. Deposited 8 April 2024.