



Learning orientations: a discrete geometry model

Y. Dabaghian¹ 

Received: 25 July 2020 / Revised: 17 October 2021 / Accepted: 25 October 2021
© The Author(s), under exclusive licence to Springer Nature Switzerland AG 2021

Abstract

In the mammalian brain, many neuronal ensembles are involved in representing spatial structure of the environment. In particular, there exist cells that encode the animal's location and cells that encode head direction. A number of studies have addressed properties of the spatial maps produced by these two populations of neurons, mainly by establishing correlations between their spiking parameters and geometric characteristics of the animal's environments. The question remains however, how the brain may intrinsically combine the direction and the location information into a unified spatial framework that enables animals' orientation. Below we propose a model of such a framework, using ideas and constructs from algebraic topology and synthetic affine geometry.

Keywords Topological data analysis · Finite synthetic geometry · Neural decoding · Spatial representation · Head direction cells · Hippocampus

Mathematics Subject Classification 51E15 · 52C99 · 54J05

1 Introduction and background

Spatial cognition in mammals is based on an internalized representation of space—a *cognitive map*¹ that emerges from neuronal activity in several regions of the brain (O'Keefe and Nadel 1978; Derdikman and Moser 2011; Grieves and Jeffery 2017; Tolman 1948; McNaughton 1996; Moser et al. 2008; Lisman et al. 2017). The type of information encoded by a specific neuronal population is discovered by establishing correspondences between its spiking parameters and spatial characteristics of the environment. For example, ascribing the xy -coordinates to every spike produced by

¹ Throughout the text, terminological definitions are given in *italics*.

✉ Y. Dabaghian
Yuri.A.Dabaghian@uth.tmc.edu

¹ Department of Neurology, The University of Texas McGovern Medical School, 6431 Fannin St, Houston, TX 77030, USA

the hippocampal principal neurons according to the animal's (in the experiments, typically rat's) position at the moment of spiking, produces distinct clusters, indicating that these neurons, the so-called *place cells*, fire only within specific locations—their respective *place fields* (O'Keefe and Dostrovsky 1971; Best and White 1998). The layout of the place fields in a spatial domain $\mathcal{E}$ —the *place field map* $M_{\mathcal{E}}$ (Fig. 1A)—thus defines the temporal order of the place cells' spiking activity during the animal's navigation, which is a key determinant of the cognitive map's structure. Hence, tagging the spikes with the location information can be viewed as a mapping from a cognitive map $\mathcal{C}$ into the navigated space,

$$f_{\sigma} : \mathcal{C} \rightarrow \mathcal{E}, \quad (1\sigma)$$

referred to as *spatial mapping* in Babichev et al. (2016). Similarly, tagging the spikes produced by certain neurons in the postsubiculum (and in few other brain regions (Taube et al. 1990; Wiener and Taube 2005)) with the rat's head direction angle φ produces clusters in the space of planar directions—the circle S^1 , thus defining a mapping

$$f_{\eta} : \mathcal{C} \rightarrow S^1. \quad (1\eta)$$

The angular domains in which specific *head direction cells* become active can be viewed as *head direction fields* in S^1 , similar to the hippocampal place fields in the navigated space. The corresponding *head direction map*, M_{S^1} , determines the order in which the head direction cells spike during the rat's movements (Fig. 1B, Taube et al. (1996); Muller et al. (1996)).

The preferred angular domains depend weakly, if at all, on the rat's position, just as place fields are overall decoupled from the head or body orientation (see however Jercog et al. 2019; Rubin et al. 2014). Thus, the following discussion will be based on the assumption that both cell populations contribute to an *allocentric* representation of the ambient space: the place cells encode a topological map of locations (Gothard et al. 1996; Alvernhe et al. 2008, 2011, 2012; Dabaghian et al. 2014; Wu and Foster 2014), whereas head direction cells augment it with angular information (Taube 1998; Valerio and Taube 2012; McNaughton et al. 2006; Savelli and Knierim 2019).

Topological model. The physiological and the computational mechanisms by which a cognitive map comes into existence remain vague (McNaughton 1996; Moser et al. 2008; Lisman et al. 2017). However, certain insights into its structure can be obtained through geometric and topological constructions. For example, a place field map $M_{\mathcal{E}}$ can be viewed as a *cover* of the navigated environment $\mathcal{E}$ by the place fields v_i ,

$$\mathcal{E} = \cup_i v_i, \quad (2)$$

and used to link the topology of $\mathcal{E}$ to the topological structure of the cognitive map $\mathcal{C}$. Indeed, according to the Alexandrov-Čech theorem, if every nonempty set of overlapping place fields, $v_{i_0, i_1, \dots, i_n} \equiv v_{i_0} \cap v_{i_1} \cap \dots \cap v_{i_n} = v_{i_0, i_1, \dots, i_n} \neq \emptyset$, is represented by an abstract simplex, $v_{i_0, i_1, \dots, i_n} = [v_{i_0}, v_{i_1}, \dots, v_{i_n}]$, then the homologies of the resulting simplicial complex $\mathcal{N}_{\sigma}$ —the *nerve* of the map $M_{\mathcal{E}}$ —match the homologies of the underlying space $H_*(\mathcal{N}_{\sigma}) = H_*(\mathcal{E})$, provided that all the overlaps $v_{i_0, i_1, \dots, i_n}$

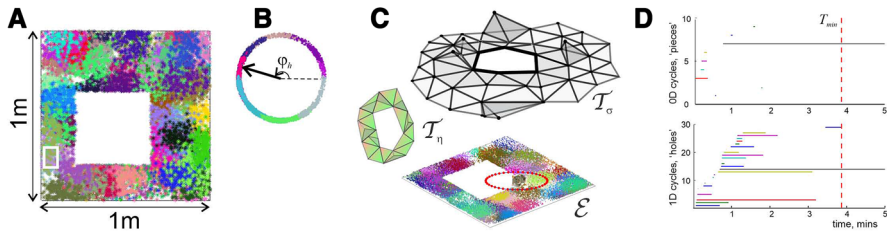


Fig. 1 Basic topological constructions. **A.** Simulated place field map M_E with place fields scattered randomly in a 1×1 m square environment $\mathcal{E}$ with a square hole in the middle. Clusters of dots of a particular color represent individual place fields. **B.** A head direction field map M_{S^1} covers the space of directions, S^1 . Clusters of colored dots mark specific head direction fields v_h , centered each at its preferred angle φ_h . **C.** The net pool of place cell coactivities is represented by the coactivity complex $\mathcal{T}_\sigma(t)$ (top right), which provides a developing topological representation of the environment $\mathcal{E}$ (bottom). The head direction cells map a circular space of directions S^1 (shown as a ring around the rat). The net pool of head direction cell activities is schematically represented the coactivity complex $\mathcal{T}_\eta(t)$. **D.** The timelines of the separate pieces (top panel) and holes (bottom panel) in the complex $\mathcal{T}_\sigma(t)$ are shown as horizontal bars. At the onset of the navigation, $\mathcal{T}_\sigma(t)$ contains many spurious topological defects that disappear after a certain “learning period” $T_{\min}^\sigma$, leaving behind a few persistent loops that define the topological shape of $\mathcal{T}_\sigma(t)$ (Dabaghian et al. (2012)). Similar behavior is exhibited by the head direction coactivity complex $\mathcal{T}_\eta(t)$

are contractible. This implies that $\mathcal{N}_\sigma$ and $\mathcal{E}$ have the same *topological shape*—same number of connectivity components, holes, cavities, tunnels, etc. Hatcher (2002). The same line of arguments allows relating the head direction map with the topology of the space of directions S^1 (Fig. 1C).

It must be emphasized however, that reasoning in terms of place and head direction fields may not capture the brain’s intrinsic principles of processing spiking information, e.g., explain how either the location or the direction signals contribute to animal’s spatial awareness, because the experimentally constructed firing fields are nothing but artificial constructions used to interpret and visualize spiking data (Hargreaves et al. 2007; Sargolini et al. 2006). Addressing the brain’s intrinsic space representation mechanisms requires carrying the analyses directly in terms of spike times, without invoking auxiliary correlates between neuronal activity and the observed environmental features.

Fortunately, the approach motivated by the nerve theorem can be easily transferred into a “spiking” format. Indeed, one can view a combination of the coactive place cells—a *cell assembly* (Hebb 1949; Harris 2005; Buzsaki 2010)—as an abstract *coactivity simplex*,

$$\sigma_i = [c_{i_0}, c_{i_1}, \dots, c_{i_n}] \quad (3)$$

that activates when the rat crosses its *simplex field* v_{σ_i} —a domain where all the cells $c_i \in \sigma_i$ are coactive (Curto and Itskov 2008). By construction, this domain is defined by the overlap of the corresponding place fields $v_{\sigma_i} = v_{i_0, i_1, \dots, i_n}$, and may hence be viewed as the the projection of σ_i into $\mathcal{E}$ under the mapping (1 σ). Note that if two coactivity simplexes overlap, their respective fields also overlap, $\sigma_i \cap \sigma_j \neq \emptyset \Leftrightarrow v_{\sigma_i} \cap v_{\sigma_j} \neq \emptyset$. Thus, if a cell c_i is shared by a set U_i of simplexes, $U_i = \{\sigma : \sigma \cap c_i \neq \emptyset\}$, then its place field is formed by the union of the corresponding σ -fields,

$$v_i = \cup_{\sigma \in U_i} v_\sigma.$$

If a simplex σ_i first appears at the moment t_i , then the net pool of neuronal activities produced by the time t gives rise to a time-developing simplicial *coactivity complex*

$$\mathcal{T}_\sigma(t) = \cup_{t_i < t} \sigma_i$$

that inflates ($\mathcal{T}_\sigma(t) \subseteq \mathcal{T}_\sigma(t')$ for $t < t'$), and eventually *saturates*, converging to the nerve complex's structure, i.e., $\mathcal{T}_\sigma(t) \approx \mathcal{N}_\sigma$, for $t \gtrsim T_\sigma^*$. Analyses based on simulating rat's moving through randomly scattered place fields show that, e.g., for a small environment $\mathcal{E}$ illustrated on Fig. 1A, the rate of new simplexes' appearance slacks in about $T_\sigma^* \approx 6$ minutes (Babichev et al. 2016), which provides an estimate for the time required to map $\mathcal{E}$.

The topological dynamics of $\mathcal{T}_\sigma(t)$ can be described using Persistent Homology theory (Zomorodian and Carlsson 2005; Edelsbrunner and Harer 2010; Kang et al. 2021), which allows identifying the ongoing shape of $\mathcal{T}_\sigma(t)$ based on the times of its simplexes' first appearance. Typically, $\mathcal{T}_\sigma(t)$ starts off with numerous topological defects that tend to disappear as the information provided by the spiking place cells accumulates (see Dabaghian et al. 2012; Arai et al. 2014; Hoffman et al. 2016; Basso et al. 2016 and Fig. 1D). Hence the minimal period $T_{\min}^\sigma$ required to recover the "physical" homologies $H_*(\mathcal{E})$ provides an estimate for the time necessary to learn topological connectivity of the environment, which, for the case illustrated on Fig. 1A, is about $T_{\min}^\sigma \approx 4 - 5$ minutes (Dabaghian et al. 2012; Arai et al. 2014; Basso et al. 2016; Hoffman et al. 2016; Alvernhe et al. 2012; Piet et al. 2018).

Importantly, the coactivity complex may be used not only as a tool for estimating learning timescales, but also as a schematic representation of the cognitive map's developing structure, providing a context for interpreting the ongoing neuronal activity. Indeed, a consecutive sequence of σ -fields visited by the rat,

$$\Upsilon = \{v_{i_0}, v_{i_1}, \dots, v_{i_n}, \dots\}, \quad (4)$$

captures the shape of the underlying physical trajectory $s \subset \Upsilon$ (Guger et al. 2011; Jensen and Lisman 2000; Frank et al. 2000; Brown et al. 1998; Zhang et al. 1998; Huang et al. 2009). The corresponding chain of the place cell assemblies ignited in the hippocampal network is represented by the *simplicial path*

$$\tilde{\sigma} = \{\sigma_1, \sigma_2, \dots, \sigma_n, \dots\}. \quad (5\sigma)$$

The fact that this information allows interpreting certain cognitive phenomena (Pfeiffer and Foster 2013; Johnson and Redish 2007; Dragoi and Tonegawa 2011) suggests that the animal's movements are faithfully monitored by neuronal activity, i.e., that in sufficiently well-developed complexes (e.g., for $t > T_{\min}^\sigma$) simplicial paths capture the shapes of the underlying trajectories (Guger et al. 2011; Jensen and Lisman 2000; Frank et al. 2000; Brown et al. 1998; Zhang et al. 1998; Huang et al. 2009). For the referencing convenience, this assumption is formulated as two model requirements:

R1. Actuality. At any moment of time, there exists an active assembly σ that represents the animal's current location.

R2. Specificity. Different place cell assemblies represent different domains in $\mathcal{E}$, i.e., σ -simplexes serve as unique indexes of the animal's location in a given map $\mathcal{C}$.

An implication of these requirements is that the simplex fields cover the explored surfaces (2) and that if the consecutive simplexes in (5 σ) are *adjacent*, i.e., no simplexes ignite between σ_i and σ_{i+1} (schematically denoted below as $\sigma_i \bullet \rightarrow \sigma_{i+1}$) then the corresponding σ -fields are adjacent or overlap.

Head orientation map. Using the same line of arguments, one can deduce the topology of the space of directions by building a dynamic *head direction coactivity complex* $\mathcal{T}_\eta(t)$ from the simplexes

$$\eta_j = [h_{j1}, \dots, h_{jn}],$$

which designate the assemblies of head direction cells $h_{j1}, h_{j2}, \dots, h_{jn}$. If a simplex η_j first activates at the moment t_j , then

$$\mathcal{T}_\eta(t) = \cup_{t_j \leq t} \eta_j.$$

As the complex $\mathcal{T}_\eta(t)$ develops, it forms a stage for representing the head direction cell spiking structure: in full analogy with (5 σ), traversing a physical trajectory $s(t)$ induces a sequence of active η -simplexes, or a head direction simplicial path

$$\tilde{\eta} = \{\eta_1, \eta_2, \dots, \eta_n, \dots\}, \quad (5\eta)$$

in which different η -simplexes represent distinct directions, at all locations. As spiking information accumulates, the topological structure of $\mathcal{T}_\eta(t)$ converges to the structure of nerve complex $\mathcal{N}_\eta$ induced by the head direction fields' cover of S^1 —every η -simplex projects into its respective head direction field v_η under the mapping (1 η). Simulations demonstrate that in the environment shown on Fig. 1A, a typical coactivity complex $\mathcal{T}_\eta(t)$ saturates in about $T_*^\eta \approx 2$ minutes, while the persistent homologies of $\mathcal{T}_\eta(t)$ filtered according to the times of η -simplexes' first appearances reveal the circular topology of the space of directions in about $T_{\min}^\eta \approx 1.5$ minutes.

From the biological point of view however, these results do not provide an estimate for *orientation learning time*: by itself, $T_{\min}^\eta$ may be viewed as the time required to learn head directions at a particular location, in every environment, whereas learning to orient in $\mathcal{E}$ implies knowing directions at every location and an ability to link orientations across locations. The latter is a much more extensive task, which, as it will be argued below, requires additional specifications and interpretations.

The following discussion is dedicated to constructing phenomenological models of orientation learning using algebraic topology and synthetic geometry approaches. In Sect. 2, we construct and test a direct generalization of the topological model, similar to the one used in Dabaghian et al. (2012); Arai et al. (2014); Basso et al. (2016); Hoffman et al. (2016) and demonstrate that it fails to produce biologically viable predictions for the learning period. In Sect. 3, the topological approach is qualitatively generalized using an alternative scope of ideas inspired by synthetic geometry. In Sect. 4, it is demonstrated that the resulting framework allows incorporating additional neurophysiological mechanisms and acquiring the topological connectivity of the environment

in a biologically viable time, thus revealing a new level of organization of the cognitive map, as discussed in Sect. 5.

2 Topological model of orientation learning

Orientation coactivity complex. The model requirements **R1** and **R2**, applied to both hippocampal and head direction activity, imply that the animal's location and orientation are represented, at any moment of time, by an active σ -simplex and an active η -simplex. Thus, the net pattern of activity in the hippocampal and in the head direction networks defines a (σ, η) pair—a single *oriented*, or *pose* simplex

$$\zeta = [\sigma, \eta],$$

(the latter term is borrowed from robotics (Thrun et al. 2005; Heinze et al. 2018; Savelli and Knierim 2019)). Restricting a ζ -simplex to its maximal subsimplexes spanned, respectively, by the place- or the head direction cells defines the projections into its positional and directional components,

$$\pi_\sigma : \zeta \rightarrow \sigma, \quad (6\sigma)$$

$$\pi_\eta : \zeta \rightarrow \eta, \quad (6\eta)$$

which permits terminology such as “ ζ is located at σ ,” “ ζ is directed toward η ,” “a location σ is directed by η ,” “ η is applied at σ ,” etc. Thus, one may refer to the σ -simplexes as to *locations* and to the η -simplexes as to *directions*, implying, depending on the context, either the items encoded in the cognitive map, or the σ/η -fields, or both.

As in the previously discussed cases, the collection of pose simplexes produced up to a moment t forms an *orientation coactivity complex* $\mathcal{T}_\zeta(t)$ that schematically represents the net pool of conjunctive patterns generated by the place- and the head direction cells accumulated since the onset of the navigation. In particular, the combinations of cells ignited along a physical path s induces an *oriented simplicial path*

$$\tilde{\zeta} = \{\zeta_1, \zeta_2, \dots, \zeta_n, \dots\}, \quad (7)$$

which runs through $\mathcal{T}_\zeta(t)$. The transitions from a given active pose simplex, ζ_i , to the next, ζ_{i+1} , occur at discrete moments $t_1, t_2, \dots, t_n, \dots$, when either the σ - or the η -component of ζ_i deactivates and the corresponding component of ζ_{i+1} ignites. Thus, the simplicial paths $\tilde{\sigma}$ and $\tilde{\eta}$ can be produced from the oriented path (7) using (5). In contrast with (5 σ) and (5 η), the σ - and η -simplexes in such paths are indexed uniformly, according to the indexes of (7),

$$\begin{aligned}\tilde{\sigma} &= \{\sigma_1, \sigma_2, \dots, \sigma_n\}, & (8\sigma) \\ \tilde{\eta} &= \{\eta_1, \eta_2, \dots, \eta_n\}. & (8\eta)\end{aligned}$$

(8σ) or in (8η) (but not in both of them simultaneously) may coincide, e.g., the location σ_i may remain the same during several timesteps, while the η -activity changes, or vice versa.

Since the rat can potentially run in any direction at any location (unless stopped by an obstacle), there are no *a priori* restrictions on the order of the place cell and the head direction cells spiking activity. This observation is formalized by another model requirement:

R3. Independence. A given head direction cell assembly η may become coactive with any place cell assembly σ and vice versa, with independent σ - and η -spiking parameters.

In model's terms, this implies that the development of the coactivity complex $\mathcal{T}_\eta(t)$ and its ultimate saturated structure $\mathcal{N}_\eta$ is the same at any location σ , and vice versa, the *saturated* structure of $\mathcal{T}_\sigma(t)$ is independent from η -activity.

However, since the activities in the hippocampal and in the head direction cell networks represent complementary aspects of the same movements, certain characteristics of the simplicial paths (8) are coupled. Specifically, in light of **R1–R2**, a connected physical trajectory s should induce a connected σ -path in the place cell complex $\mathcal{T}_\sigma$, together with a connected η -path in the head direction complex $\mathcal{T}_\eta$. Similarly, a *looping* trajectory should induce periodic sequences of simplexes,

$$\begin{aligned}\tilde{\sigma}_o &= \{\sigma_1, \sigma_2, \dots, \sigma_n, \sigma_1, \sigma_2, \dots\}, \\ \tilde{\eta}_o &= \{\eta_1, \eta_2, \dots, \eta_n, \eta_1, \eta_2, \dots\}.\end{aligned}$$

In other words, making a loop in physical space $\mathcal{E}$ should induced σ - and η -loops. Thus, without referencing the physical trajectory, the model requires

R4. Topological consistency. The simplicial paths (8) should be connected and a simple periodic σ -path should induce a simple periodic η -path and vice versa.

Orientation learning. As discussed above, getting rid of the topological defects in $\mathcal{T}_\sigma(t)$ and in $\mathcal{T}_\eta(t)$ allows faithful topological classification of physical routes in terms of the neuronal (co)activity. Thus, the “topological maturation” of these complexes can be viewed as a schematic representation of the learning process. The concept of oriented simplicial paths embedded into the orientation coactivity complex $\mathcal{T}_\zeta(t)$ allows a similar interpretation of the spatial orientation learning—as acquiring an ability to distinguish between qualitatively disparate moving sequences. Indeed, knowing how to orient in a given space, viewed as a cognitive ability to reach desired places from different directions via a suitable selection of intermediate locations and turns, may be interpreted mathematically as an ability to classify trajectories using topologically inequivalent classes of oriented simplicial paths (8). From an algebraic-topological perspective, this may be possible after the orientation complex $\mathcal{T}_\zeta(t)$ acquires its correct topology.

To establish the latter, note that the complex $\mathcal{T}_\zeta(t)$ has the same nature as $\mathcal{T}_\sigma(t)$ and $\mathcal{T}_\eta(t)$ —it is an emerging temporal representation of a nerve complex, induced from a cover of a certain *orientation space* $\mathcal{O}$ that combines $\mathcal{E}$ and S^1 . Since the preferred angles of the head direction cells remain the same at all locations Wiener and Taube (2005), the space of directions S^1 represented by these cells does not “twist” as the rat moves across $\mathcal{E}$, which implies that the orientation space has a direct product structure $\mathcal{O} = \mathcal{E} \times S^1$ (Hatcher 2002). One may thus combine (1 σ) and (1 η) to construct a *joint spatial mapping*, $f_\zeta = (f_\sigma, f_\eta)$, that associates instances of simultaneous activity of place- and head direction cell groups with domains in the orientation space,

$$f_\zeta : \mathcal{T}_\zeta \rightarrow \mathcal{O}.$$

For example, if a given place cell c maps into a field $v = f_\sigma(c)$, then the coactivity of a pair $\zeta = [c, h]$ (the smallest possible combined coactivity) can be mapped into $\mathcal{O}$ by shifting v along the corresponding fiber S^1 according to the angle field of the head direction component of ζ , $\varphi = f_\eta(h)$ (Fig. 2A). The resulting *orientation fields*, $v_{\zeta_i} = f_\zeta(\zeta_i)$, form a cover of the orientation space,

$$\mathcal{O} = \cup_i v_{\zeta_i},$$

whose nerve $\mathcal{N}_\zeta$ is reproduced by the temporal orientation complex $\mathcal{T}_\zeta(t)$.

Just as the σ - and η -simplexes, each pose simplex ζ_k has a well-defined appearance time, t_k , due to which the orientation complex is *time-filtered*, $\mathcal{T}_\zeta(t) \subseteq \mathcal{T}_\zeta(t')$ for $t < t'$. Applying Persistent Homology techniques (Zomorodian and Carlsson 2005; Edelsbrunner and Harer 2010), one can compare topological shape defined by the time-dependent Betti numbers of $\mathcal{T}_\zeta(t)$ with the shape of the orientation space $\mathcal{O} = \mathcal{E} \times S^1$, and thus quantify the orientation learning process.

As an illustration of this approach, we simulated the rat’s movements on a circular runway, for which the total representing space for the combined place cell and head direction cell coactivity forms a 2D torus (Fig. 2A). To simplify modeling, we used movement direction as a proxy for the head direction, although physiologically these parameters not identical (Cei et al. 2014; Raudies et al. 2015; Laurens and Angelaki 2018; Shinder and Taube 2011, 2014). Computations show that the correct topological shapes of the place- and the head direction complexes emerge in about 1 minute (Fig. 2B,C), while the transient topological defects in $\mathcal{T}_\zeta(t)$ disappear in about 80 minutes (Fig. 2D)—a surprisingly large value that exceeds behavioral outcomes by an order of magnitude (Alvernhe et al. 2012; Piet et al. 2018). Even assuming that biological learning may involve only a partial reconstruction of the orientation space’s topology, the persistence diagrams shown on Fig. 2D indicate that $\mathcal{T}_\zeta(t)$ is not only riddled with holes for over 40 minutes, but also that it remains disconnected ($b_0(\mathcal{T}_\zeta) > 1$) for up until 17 minutes, i.e., according to the model, the animal should not be able to acquire a connected map of a simple annulus after completing multiple lapses across it.

Biological implications of mismatches between the experimental and the modeled estimates of learning timescales are intriguing: since the model’s quantifications are based on “topological accounting” of place- and head direction cell coactivities

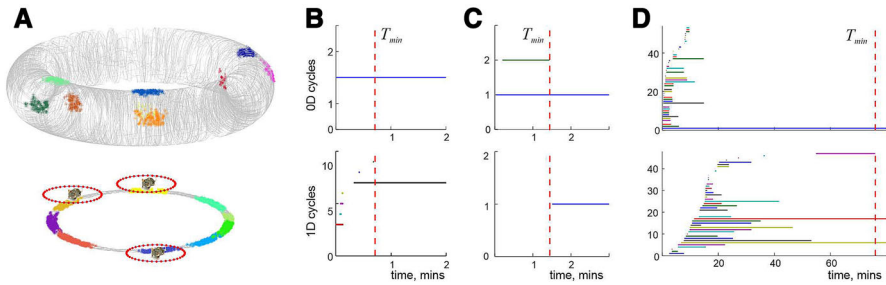


Fig. 2 Topological leaning dynamics. **A.** The orientation space for a rat moving on a circular runway (bottom) is a topological torus (top), $\mathcal{O} = T^2$. Clusters of colored dots show examples of the simplex fields $v_{\xi_{ij}}$ of the basic coactivity combinations $\xi_{ij} = [c_i, h_j]$. The timelines of the topological loops in the place cell complex $\mathcal{T}_\sigma(t)$ (panel **B**, horizontal bars) and in the head direction cell complex $\mathcal{T}_\eta(t)$ (panel **C**) disappear in under a minute; a stable loop in $0D$ and a stable loop in $1D$ in each case indicate that both the runway and the direction space are topological circles. **D.** The orientation coactivity complex $\mathcal{T}_\zeta(t)$ contains many topological defects that take over an hour to disappear: the spurious $0D$ loops contract in about 17 minutes (top panel), and the spurious $1D$ loops persist for 80 minutes. For an “open field” environment (Fig. 1A) a similar topological learning process takes hours. These estimates exceed the experimental learning timescales, suggesting that the physiological orientation learning involves additional mechanisms

induced by the animal’s movements, the root of the problem seems to lay not in the mathematical side of the model, but in the biological assumptions that underlie the computations. Specifically, the overly long learning periods suggest that building a spatial map from the movement-triggered neuronal coactivity alone does not take into account certain principal components of the learning mechanism. In other words, the fact that the animal seems to produce correct representations of the environment much faster than it would be possible from the influx of navigation-triggered data, implies that the brain can bypass the necessity to discover every bit of information empirically, i.e., that building a cognitive map may be accelerated by “generating information from within,” via autonomous network dynamics.

Physiologically, this conclusion is not surprising: the phenomena associated with spatial information processing through endogenous hippocampal activity are well known. Many experiments have demonstrated that the animal can *replay* place cells during the quiescent states (Wu and Foster 2014; Karlsson and Frank 2009; Ólafsdóttir et al. 2018) or sleep (Ji and Wilson 2007; Louie and Wilson 2001), in the order in which they have fired during preceding active exploration of the environment, or *preplay* place cells in sequences that represent future trajectories (Pfeiffer and Foster 2013; Johnson and Redish 2007; Dragoi and Tonegawa 2011). These phenomena are commonly viewed as manifestations of the animal’s “mental explorations” of its cognitive map, which help acquiring, sustaining and retrieving memories (Hopfield 2010; Zeithamova et al. 2012).

However, one would expect a different functional impact of replays and preplays on spatial learning. Since replays represent past experiences, they cannot accelerate acquisition of new spatial information—in the model’s terms, reactivation of simplexes that are already included into the coactivity complex cannot alter its shape (in absence of synaptic and structural instabilities (Roux et al. 2017; van de Ven et al. 2016; Babichev et al. 2018, 2019)). In contrast, preplaying place cell combinations that have

not yet been triggered by previous physical moves may speed up the learning process. Yet, from the modeling perspective, it is *a priori* unclear which specific trajectories may be preplayed by the brain, or, in computational terms, which specific simplicial trajectories should be “injected” into the simulated cognitive map $\mathcal{T}_c(t)$ to simulate the preplays that may accelerate learning. Experiments suggest that “natural” connections between locations are the straight runs (Pfeiffer and Foster 2013; Valerio and Taube 2012; McNaughton et al. 2006); however, implementing such runs would require a certain “geometrization” of the topological model. In the following, we propose a geometric implementation of preplays that help to expedite learning process and open new perspectives modeling spatial representations.

3 Geometric model of orientation learning

The motivation for an alternative approach comes from the observation that combining inputs from the place and the head direction cells offers a possibility of establishing different arrangements of the locations in the hippocampal map. Indeed, common interpretations of the head direction cells’ functions suggest that the rat’s movements guided by a fixed head direction activity trace approximately straight paths, whereas shifts in η -activity indicate curved segments of the trajectory, turns, etc. Chen et al. (1994); Wiener and Taube (2005); Savelli and Knierim (2019). This implies that the brain may use the head direction cells’ outputs to represent shapes of the paths encoded by the place cells, to align their segments, identify collinearities, their incidences, parallelness, etc.

To address these structures and their properties we will use the following definitions:

D1. Simplexes σ_1 and σ_2 are *aligned in η -direction*, if they may ignite during an uninterrupted activity of a fixed η -simplex. In formal notations, $(\sigma_1, \sigma_2) \triangleleft \eta$.

D2. Two η -aligned simplexes are *η -adjacent*, $\{\sigma_1 \rightarrow \sigma_2 | \eta\}$, if the ignition of σ_2 follows immediately the ignition of σ_1 , with no other cell groups igniting in-between.

D3. An ordered sequence of σ -simplexes forms an *η -oriented alignment* if each pair of consecutive simplexes, (σ_i, σ_{i+1}) in (8 σ), is η -adjacent, i.e., if the oriented path,

$$\ell = \{[\sigma_1, \eta], [\sigma_2, \eta], \dots, [\sigma_n, \eta]\},$$

never changes direction. The notation

$$\ell = \{\sigma_1, \sigma_2, \dots, \sigma_n | \eta\} \equiv \{\bar{\sigma} | \eta\} \quad (9)$$

highlights the set of *collinear* locations $\bar{\sigma} = (\sigma_1, \sigma_2, \dots, \sigma_n)$ and the η -simplex that orients it. The bar in $\bar{\sigma}$ is used to distinguish an alignment from a generic simplicial path $\bar{\sigma}$.

D4. An alignment ℓ_1 *augments* an alignment ℓ_2 , if both ℓ_1 and ℓ_2 can be guided by an uninterrupted η -activity ($\bar{\sigma}_1 \bowtie \bar{\sigma}_2 \Leftrightarrow (\bar{\sigma}_1 \cup \bar{\sigma}_2) \triangleleft \eta$). Conversely, a proper subset $\bar{\sigma}'$ of an aligned set $\bar{\sigma}$ forms its proper *subalignment* ($\bar{\sigma}' \subset \bar{\sigma} \Leftrightarrow \{\bar{\sigma}' | \eta\} \times \{\bar{\sigma} | \eta\}$).

D5. Two alignments ℓ_1 and ℓ_2 *overlap*, if they share a location σ ($\ell_1 \cap \ell_2 = \sigma \Leftrightarrow \sigma \in \bar{\sigma}_1 \cap \bar{\sigma}_2$).

D6. A location σ' lays *outside* of an η -alignment ℓ_η , if it aligns with any σ from ℓ_η along a direction different from η ($\sigma' \notin \ell_\eta \Leftrightarrow \exists \sigma \in \ell_\eta, \eta' \neq \eta : (\sigma, \sigma') \triangleleft \eta'$).

D7. Two alignments are *parallel*, if they are directed by the same or opposite η -activity, without augmenting each other, i.e., if one $\pm\eta$ -alignment contains a location outside of the other one, ($\ell_1 \parallel \ell_2 \Leftrightarrow \eta_1 = \pm\eta_2$, and $\exists \sigma : \sigma \in \bar{\sigma}_1, \sigma \notin \bar{\sigma}_2$, where $f_\eta(-\eta) \cong \pi + f_\eta(\eta)$).

D8. A *yaw* is an oriented path in which a sequence of η -simplexes ignites at a fixed location σ ,

$$y = \{[\sigma, \eta_1], [\sigma, \eta_2], \dots, [\sigma, \eta_n]\}.$$

Thus, yaws may be viewed as structural opposites of the alignments, which is emphasized by the notation

$$y = \{\eta_1, \eta_2, \dots, \eta_n | \sigma\} \equiv \{\hat{\eta} | \sigma\}$$

that highlights the range of η -simplexes, $\hat{\eta}$, ignited at the *axis of the yaw*, σ .

D9. A *clockwise turn* is an oriented path $\tilde{\zeta}_+ = \{[\sigma_1, \eta_1], [\sigma_2, \eta_2], \dots, [\sigma_n, \eta_n]\}$ with a growing angular sequence, i.e., the angle φ_i that represents the element η_i is not greater than the next one, $\varphi_{i+1} \geq \varphi_i$. A *counterclockwise turn* $\tilde{\zeta}_-$ is an oriented path with a decreasing angular sequence, $\varphi_{i+1} \leq \varphi_i$.

The latter definition is due to the observation that η -simplexes can be ordered according to the angles they represent, i.e., $\eta_i < \eta_{i+1}$ iff $\varphi_i < \varphi_{i+1}$, which also allows defining the angle between alignments,

$$\angle(\ell_{\eta_i}, \ell_{\eta_j}) \equiv \angle(\eta_i, \eta_j) \equiv |\varphi_i - \varphi_j| = \Delta\varphi_{ij}.$$

In light of the definitions **D1–D9**, the model requirements **R1–R4** imply that the entire ensembles of the active place and head direction cells are involved into geometric arrangements, e.g., every location σ belongs to an alignment directed by a η -simplex, and conversely, every η -simplex directs a nonempty alignment through the σ -map. The question is, whether this collection of alignments is sufficiently complete to allow self-contained geometric reasoning in terms of collinearities, incidences, parallelisms, etc., i.e., does it form a self-contained geometry?

The standard approach to answering this question is based on verifying a set of axioms, in this case—the axioms of affine geometry, applied to the elements of a suitable set $A = \{x, y, z, \dots\}$ and its select subsets $\ell_1, \ell_2, \dots \in A$:

A1. Any pair of distinct elements of $x \neq y$ is included into a unique subset ℓ .

A2. There exists an element x outside of any given subset ℓ , $x \notin \ell$.

A3. For any subset ℓ and an element $x \notin \ell$, there exists a unique subset ℓ_x that includes x , but does not overlap with ℓ .

If these axioms (referred to as *A-axioms* below) are satisfied, then the subsets $\ell_1, \ell_2, \dots$, can be viewed as lines because interrelationships among them and with

other elements of A reproduce the familiar geometric incidences between lines and points in the Euclidean plane (Hilbert 1992; Hilbert and Cohn-Vossen 1999; Batten 1997; Kerteszi 1976). However, the set of geometries established via the $\mathbf{A}$ -axioms is much broader than its main “motivating example:” the standard planar affine geometry $\mathcal{A}_E$ is but a specific model implementing the $\mathbf{A}$ -axioms using the infinite set of infinitesimal points and infinite lines (Hilbert 1992; Hilbert and Cohn-Vossen 1999; Batten 1997; Kerteszi 1976). In fact, it is also possible to use the $\mathbf{A}$ -axioms to establish geometry on finite sets, thus producing finite affine planes. This is important for modeling cognitive maps encoded by the physiological networks that contain finite numbers of neurons and thus may represent finite sets of locations and alignments. Specifically, a possible adaptation of the $\mathbf{A}$ -axioms using spiking semantics, is the following:

A1_n. Any pair of distinct locations $\sigma_1 \neq \sigma_2$ belongs to a unique alignment, i.e., σ_1 and σ_2 may ignite in sequence during the activity of a single η -simplex ($\forall \sigma_1, \sigma_2, \exists! \eta, \bar{\sigma} : (\sigma_1, \sigma_2) \times \bar{\sigma} \triangleleft \eta$).

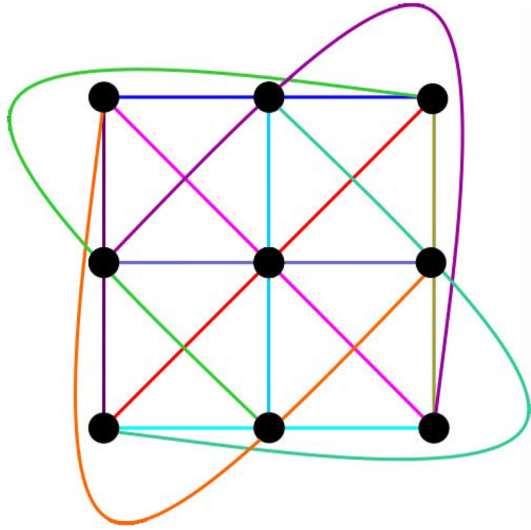
A2_n. The location-encoding network can activate a group of cells σ to represent a location outside of any given alignment ($\forall \ell, \exists \sigma \notin \ell$).

A3_n. For any alignment ℓ and a location $\sigma \notin \ell$, there exists a unique alignment ℓ_σ parallel to ℓ that passes through σ ($\forall \ell, \sigma \notin \ell, \exists! \ell_\sigma : \sigma \in \ell_\sigma, \ell \parallel \ell_\sigma$).

Validating these axioms over the net pool of spiking activities produced by the hippocampal and head direction cells would establish a discrete-geometric structure encoded by (σ, η) -neuronal activity. However, the requirements imposed by the $\mathbf{A}_n$ -axioms may not be compatible with physiological mechanisms that operate the corresponding networks, as well as with these networks’ functions. Indeed, the configurations formed by the connections in finite planes are typically non-planar (Fig. 3), whereas physiological computations combining place and head direction cells’ activities appear to enable geometric planning in planar environments (Valerio and Taube 2012; McNaughton et al. 2006). Second, the combinations of locations that form “relational lines” according to the $\mathbf{A}_n$ -axioms may not have the standard properties of their Euclidean counterparts, e.g., they may include sequences of locations that cannot be consistently mapped into straight Euclidean paths (Fig. 3). In contrast, experiments show that neuronal activity during animals’ movements along straight arrangements of σ -fields, as well as their offline preplays/replays (Matar and Daw 2018; Byrne and Becker 2004), dovetail with the definitions **D1–D7**. Third, given a large number of the encoded locations (in rats, about 3×10^4 of active place cells in small environments Ziv et al. (2013)) and a very large set of possible co-active cell combinations (Buzsaki 2010; Babichev et al. 2016), a finite set of η -simplexes may not suffice to align all pairs of locations in the sense of the definition **D7**.

On the other hand, certain key features of finite affine planes, e.g., the necessity of having k parallel lines in every direction and a fixed number, $k + 1$, of lines passing through each location (Hilbert 1992; Hilbert and Cohn-Vossen 1999; Batten 1997; Kerteszi 1976) are reflected in the network. Indeed, the existence of a fixed population of head direction assemblies results in a fixed number $k \cong N_\eta/2$ of distinct alignments passing through any location and the same number of directions running across the cognitive map.

Fig. 3 Finite affine plane of the third order, $\mathcal{A}_3$, with 9 points (black dots) aligned according to the **A**-axioms in 12 lines (colors) forms a non-planar configuration



Together, these observations suggest that the brain may combine certain aspects of finite geometries and Euclidean plane discretizations. Rather than trying to recognize the net geometry of the resulting representations from the onset, one may adopt a constructive “bottom up” approach: it may be possible to interpret certain *local* properties of spiking activity as basic geometric relations and then follow how such relations accumulate at larger scales, yielding *global* geometric frameworks. For example, it can be argued that place cell activities can be aligned locally, i.e., that ignitions of a specific head direction assembly can accompany transitions of activity from one place cell assembly to an adjoining one. It is also plausible that, in stable network configurations, the selection of cell groups involved in such transitions is limited or even unambiguous. Also, given the number of place cells ($N_c \approx 3 \times 10^5$) and typical assembly sizes (60 – 300 cells) (Buzsaki 2010), there should be enough place cell combinations to represent a sufficient set of alignments, overlaps between them, etc. (Babichev et al. 2016; Perin et al. 2011; Reimann et al. 2017).

Thus, in addition to (mostly topological) requirements **R1–R4**, one can consider the following neuro-geometric rules:

G1. Any two adjacent locations align in a unique direction ($\forall \sigma_1 \rightarrow \sigma_2, \exists! \eta : (\sigma_1, \sigma_2) \triangleleft \eta$).

G2. A location adjacent to a given one may be recruited in any direction ($\forall \eta, \sigma_1, \exists! \sigma_2 : (\sigma_1 \rightarrow \sigma_2) \triangleleft \eta$).

G3. The location-encoding network can explicitly represent the overlap between any two non-parallel alignments ($\forall \ell_1, \ell_2, \eta_1 \neq \eta_2, \exists! \sigma : \ell_1 \cap \ell_2 = \sigma \in \ell_1, \ell_2$).

In contrast with the **A**_n-axioms that aim to establish large-scale properties of a “cognitive” affine plane as a whole, the **G**-rules define local geometric relationships induced by local mechanisms controlling neuronal activity. In particular, **G1** ascertains a possibility of aligning any two adjacent locations, rather than any two locations as required by the **A1**_n. The rule **G2** is complementary: it posits that if an active η -

combination is selected, then the activity can propagate from a given σ_1 to a specific adjacent σ_2 . Lastly, the rule **G3** allows reasoning about the locations, alignments, incidences, etc., assuming that all these elements can be physiologically actualized.

As a first application of the **G**-rules, notice that a σ -path, viewed as a sequence of adjacent σ s, induces a unique ordered η -sequence, i.e., a $\tilde{\eta}$ -path in $\mathcal{T}_\eta(t)$. Together, these paths define an oriented trajectory $\tilde{\zeta}'$, formed by uniquely directed straight links between adjacent locations, which can be graphically represented by a *directed polygonal chain* of σ -locations (Fig. 4A). Conversely, the fact that a generic $\tilde{\zeta}$ projects into an ordered sequence of adjacent σ -simplexes, $\tilde{\sigma} = \pi_\sigma(\tilde{\zeta})$, implies that oriented paths can be aligned into the polygonal chains, $\tilde{\zeta} \rightarrow \tilde{\zeta}'$ (Fig. 4B). From the perspective of the topological model discussed in Section 2, this means that each path $\tilde{\sigma}$ can be “lifted” from $\mathcal{T}_\sigma(t)$ into a unique oriented path $\tilde{\zeta}' \in \mathcal{T}_\zeta(t)$ by a back projection, $\tilde{\zeta}' = \pi_\sigma^{-1}(\tilde{\sigma})$. In the following, the term “simplicial path” will refer to the polygonal chains only, unless explicitly stated otherwise, and the “prime” notation will be suppressed.

Reversing the order of simplexes in a chain and inverting the corresponding η -sequence,

$$\tilde{\sigma}_+ = \{\sigma_{i_1}, \sigma_{i_2}, \dots, \sigma_{i_n}\} \rightarrow \tilde{\sigma}_- = \{\sigma_{i_n}, \sigma_{i_{n-1}}, \dots, \sigma_{i_1}\}, \quad (10\sigma)$$

$$\tilde{\eta}_+ = \{\eta_{i_1}, \eta_{i_2}, \dots, \eta_{i_n}\} \rightarrow \tilde{\eta}_- = \{\pm\eta_{i_n}, \pm\eta_{i_{n-1}}, \dots, \pm\eta_{i_1}\}. \quad (10\eta)$$

where the angle $-\eta$ is diametrically opposite to η , $f_\eta(-\eta_i) \approx \pi + f_\eta(\eta_i)$. The “+” sign in (10 η) corresponds to “backing up” along the path $\tilde{\sigma}_+$ and the “−” sign to reversing the moving direction, either by implementing the required physical steps or by flipping the order of the replayed or preplayed sequences (Foster and Wilson 2006; Ambrose et al. 2016) (in open fields, place cell spiking is omnidirectional Chen et al. (2018)). The Since move reversal does not affect the σ -paths’ geometries,

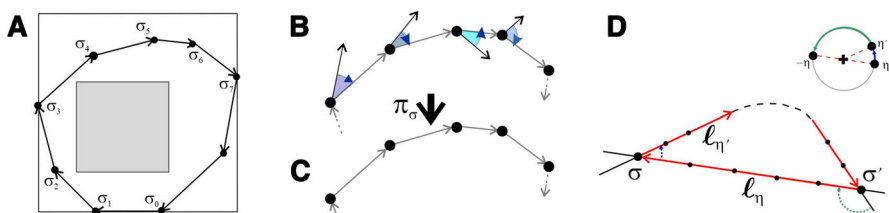


Fig. 4 G-rules based geometric constructions. **A.** A directed polygonal chain connecting pairs of adjacent locations along the navigated trajectory in the environment shown on Fig. 1A. **B.** A schematic representation of a discrete homotopy from a generic ζ -path to a polygonal chain: the η -components shift towards the unique directions that align the adjacent locations (gray arrows). **C.** The resulting polygonal chain (a combination of yaws and straight runs) projects by π_σ into an undirected chain connecting the adjacent locations—a fragment of the chain shown on the panel A. **D.** Lemma 2: If two nonparallel alignments, ℓ_η and $\ell_{\eta'}$, produce two intersections σ and σ' , then there exists a $\tilde{\sigma}$ -trajectory that forms a noncontractible simple loop, while the corresponding $\tilde{\eta}$ -trajectory forms a contractible segment (top left corner), in contradiction with **R4**

the transformations (10) can be regarded as equivalence relationships, which do not reference physical trajectory:

R5. Reversibility. *Simplicial σ -paths related via (10) are geometrically identical.*

In accordance with **R5**, a given η -oriented alignment, $\ell_+ = \{\bar{\sigma}|\eta\}$, and its inverse, $\ell_- = \{\bar{\sigma}|- \eta\}$, define the same collinear sequence, i.e., $\pi_\sigma(\ell_+) = \pi_\sigma(\ell_-) = \bar{\sigma}$ —a natural observation that motivates the definition **D7**. In particular, a pair of $\pm\eta$ adjacent locations is also geometrically adjacent ($\sigma_i \bullet \sigma_j \Leftrightarrow \{\sigma_i \rightarrow \sigma_j|\eta\}$ or $\{\sigma_i \leftarrow \sigma_j|- \eta\}$), which allows representing trajectories by *undirected polygonal chains* connecting adjacent σ -fields (Fig. 4C). For example, a bending chain corresponds to a clockwise turn ζ_+ as well as to its counterclockwise counterpart $\tilde{\zeta}_-$ (both project to the same σ -path, $\pi_\sigma(\zeta_+) = \pi_\sigma(\tilde{\zeta}_-) = \tilde{\sigma}$); a closed chain—to a loop that can be traversed in clockwise or in counterclockwise direction ($\tilde{\sigma}_o = \pi_\sigma(\zeta_{o-}) = \pi_\sigma(\tilde{\zeta}_{o+})$), etc.

Returning to the link between the **G**-rules and the remaining two **A**- or **A_n**-axioms, it can be observed that the **A2_n**-axiom is an immediate consequence of **G2**: if the network is capable of actualizing up to N_η simplexes adjacent to a given one along N_η available directions, then $N_\eta - 2$ of them will necessarily lay outside of a given alignment. The argument for the existence and uniqueness of parallel lines (axiom **A3_n**) can be organized into the following two lemmas:

Lemma 1 *If σ is a location outside of an alignment ℓ_η , $\sigma \notin \ell_\eta$, and ℓ'_η is an alignment directed by η at σ , $\ell_\eta \neq \ell'_\eta$, then ℓ_η and ℓ'_η do not overlap.*

Proof . Assume that the overlap exists, $\ell_\eta \cap \ell'_\eta = \sigma' \neq \emptyset$. Since η is a unique index of directions, the location σ' is η -aligned with its adjacent locations both in ℓ_η and in ℓ'_η . Thus, ℓ_η and ℓ'_η augment each other ($\bar{\sigma}_1 \bowtie \bar{\sigma}_2$), forming a single joint η -alignment that passes through σ , in contradiction with the original assumption $\sigma \notin \ell_\eta$. $\square$

Lemma 2 *Two non-parallel lines cannot intersect more than once.*

Proof . Consider two alignments ℓ_η and $\ell_{\eta'}$, $\eta \neq \eta'$, with $\ell_\eta \cap \ell_{\eta'} = \sigma$. Without loss of generality (change $\ell_\eta \rightarrow \ell_{-\eta}$ if necessary), we may assume that the angle between them is sharp, $\angle(\ell_\eta, \ell_{\eta'}) < \pi/2$ (Fig. 4C). Consider an oriented path ζ that starts at σ in η' -direction, i.e., along $\ell_{\eta'}$. If $\ell_{\eta'}$ crosses ℓ_η again at a location σ' , then the path ζ may turn back at σ' and continue along $\ell_{-\eta}$ towards σ , then continue along $\ell_{\eta'}$ again, etc., yielding a single closed σ -path. On the other hand, the corresponding η -path links η and η' at the first turn and then goes back from η' to $-\eta$ at the second turn, forming a contractible segment, in contradiction with **R4**. $\square$

In effect, these two lemmas validate the constructive definition of parallelness **D7** and point at an alternative form of the axiom **A3_n**: *If two locations σ_1 and σ_2 are aligned along η , then any two alignments ℓ_1 and ℓ_2 directed through σ_1 and σ_2 by a $\eta' \neq \eta$ are parallel.*

The foregoing discussion suggests that the spatial framework represented by the place cells and the head direction cells forms neither a naïve discretization of the Euclidean plane nor a conventional finite geometry, as defined by the standard **A**-axioms. Rather, combining the location and the direction information can be used to

capture a certain sub-collection of geometric arrangements, e.g., a particular set of alignments, which may then be used for navigation and geometric planning (Valerio and Taube 2012; McNaughton et al. 2006). Correspondingly, orientation learning can be interpreted as a process of establishing and expanding such arrangements (e.g., prolonging shorter alignments, completing partial ones, etc.) and accumulating them in the cognitive map.

4 Synthesizing cognitive geometry

As a basic example of a geometric map learning, consider an oriented trajectory $\zeta(t)$ that starts with an alignment $\ell_1 = \{\sigma_0, \sigma_1 | \eta_1\}$, followed by a yaw at σ_1 , and continues along $\ell_2 = \{\sigma_1, \sigma_2 | \eta_2\}$, reaching σ_2 at the moment t_2 (Fig. 5A). As in Sec. 2, the head and the motion directions are identified for simplicity. If σ_0 , σ_1 and σ_2 are the only locations in the emerging affine map $\mathcal{A}(t_2)$, then σ_2 is adjacent to σ_0 and hence it must align with σ_0 along a certain η -direction η_{20} (assuming a generic case, in which ℓ_1 and ℓ_2 are nonparallel, $\eta_1 \neq \pm \eta_2$). Representing this alignment in the parahippocampal network, i.e., producing the corresponding imprints in the synaptic architecture via plasticity mechanisms (Leuner and Gould 2010; Caroni et al. 2012; Brown and O. 2020), requires actualization by igniting σ_2 and σ_0 consecutively during the activity of a particular η_{20} . This can be achieved either by navigating between the corresponding σ -fields or off-line, via autonomous network activity. In the former case, the connection $\ell_{20} = \{\sigma_2, \sigma_0 | \eta_{20}\}$ is incorporated into the map after the animal arrives to σ_0 from σ_2 , i.e., at the “empirical learning” timescale discussed in Section 2. In the latter case, ℓ_{20} may form at the spontaneous spiking activity timescale (milliseconds (Wu and Foster 2014; Karlsson and Frank 2009; Ólafsdóttir et al. 2018; Ji and Wilson 2007; Louie and Wilson 2001; Pfeiffer and Foster 2013; Johnson and Redish 2007; Dragoi and Tonegawa 2011)), as soon as the animal reaches σ_2 , which clearly accelerates the formation of the cognitive map.

This illustrates the model’s general approach: although the geometric constructions were discussed in Section 3 in reference to spiking produced during the animal’s movements, they also apply to endogenous spiking activity. In other words, the **G**-rules can be used “imperatively,” for producing geometric structures in the cognitive map autonomously, based on available information rather than physical navigation. In particular, preplays can be used for aligning locations with specific η -assemblies by preplaying straight “home runs,” as soon as the physical trajectory assumes a suitable configuration.

The physiological processes that enforce transitions of activity between cell assemblies are currently studied both experimentally and theoretically (Harris 2005; Buzsaki 2010; Laurens and Angelaki 2018); in case of the head direction and place cells, the corresponding network computations may be guided by sensory (e.g., visual) and idiothetic (proprioceptive, vestibular and motor) inputs and involve a variety neurophysiological mechanisms (Haggerty and Ji 2015; Knierim et al. 1998; Chen et al. 2013; Zhang et al. 2014; Laurens and Angelaki 2018). However, the principles of utilizing such mechanisms for acquiring a map of orientations can be illustrated using basic, self-contained algorithms that rely on the information provided by the hippocam-

pal and head direction spiking. Specifically, the history of the η -assemblies' ignitions allows estimating the direction between the loci of a polygonal chain according to

$$\Delta\varphi_{ij} = \frac{1}{N} \sum_{k=i}^j \varphi_k n_k,$$

where n_k is the number of spikes produced by the k^{th} head direction assembly η_k , φ_k is the corresponding angle (i.e., $\eta_k = \pi_\eta(\zeta_k)$ and $\varphi_k = f_\eta(\eta_k)$), N is the total number of spikes. If η_k is characterized by a Poisson firing rate μ_k , then the number of spikes that it produces over an ignition period Δt_k can be estimated as $n_k = \mu_k \Delta t_k$. Assuming for simplicity that all rates are the same $\mu_k = \mu$, the angular shifts can be estimated from the individual ignitions' duration and the total navigation time T ,

$$\Delta\varphi_{ij} = \frac{1}{T} \sum_{k=i}^j \varphi_k \Delta t_k. \quad (11)$$

The η -simplex required to perform a home run from σ_j to σ_i can then be selected as the one whose discrete angle is closest to $\varphi_j = \varphi_i + \Delta\varphi_{ij}$, i.e.,

$$\eta_j = \min_{\eta} (\varphi_j - f_\eta(\eta)). \quad (12)$$

In particular, (11) and (12) allow estimating the required direction from σ_2 to σ_0 and thus identifying the simplex η_{20} that needs to direct the corresponding home run preplay ℓ_{20} . Other models can be built by modifying or altering these rules.

The next move continues along $\ell_3 = \{\sigma_2, \sigma_3 | \eta_3\}$, arriving to σ_3 at the moment t_3 , which allows preplaying connections to previously visited locations along ℓ_{31} and ℓ_{30} in the map $\mathcal{A}_C(t_3)$ (Fig. 5B). If $\zeta(t)$ is a right turn ($\eta_1 < \eta_2 < \eta_3$), then the line ℓ_{31} lays between ℓ_{30} and ℓ_3 , and, according to **G3**, overlaps with ℓ_{20} at $\bar{\sigma}_1$, which will thus lay between σ_0 and σ_2 , $\ell_{31} = \{\sigma_3, \bar{\sigma}_1, \sigma_1 | \eta_{31}\}$. Also, since ℓ_3 and ℓ_1 are non-parallel, they produce an overlap at $\bar{\sigma}_2$ that extends the "seed alignments" $\ell_1(t_2)$ and $\ell_3(t_2)$ to $\ell_1(t_3) = \{\sigma_0, \sigma_1, \bar{\sigma}_2 | \eta_1\}$ and $\ell_3(t_3) = \{\bar{\sigma}_2, \sigma_2, \sigma_3 | \eta_3\}$. The locations within the alignments ℓ_1 and ℓ_3 are ordered correspondingly, e.g., σ_1 falls between σ_0 and $\bar{\sigma}_2$, and σ_2 falls between $\bar{\sigma}_2$ and σ_3 .

Note that since $\bar{\sigma}_1$ and $\bar{\sigma}_2$ can be viewed as adjacent, it is also possible to form an additional alignment $\ell_x = \{\bar{\sigma}_1, \bar{\sigma}_2 | \eta_x\}$, which induces two additional locations by intersecting ℓ_{30} and ℓ_{12} (Fig. 5C). However, the orientations of the existing segments of the trajectory do not determine the direction η_x , and ℓ_x can therefore be viewed as a "provisional" alignment that may be actualized once the explicit information specifying its orientation emerges. Nevertheless, such alignments and the incidences that they induce may be incorporated into the hippocampal map, to accelerate its topological dynamics.

As the turn continues, the next segment connects to σ_4 along $\ell_4 = \{\sigma_3, \sigma_4 | \eta_4\}$, allowing home run preplays ℓ_{40} , ℓ_{41} , ℓ_{42} , which produce additional intersections, augmenting the lines ℓ_{30} , ℓ_{31} and ℓ_{43} in specific order (Fig. 5D). Subsequent segments

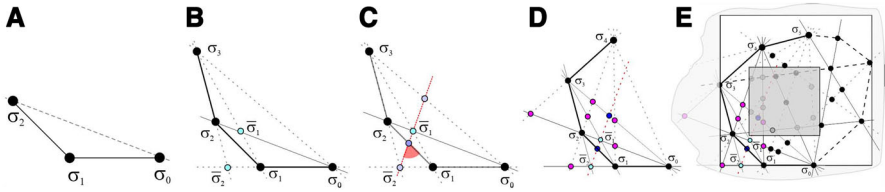


Fig. 5 Aligning the cognitive map. **A** The endpoints of the initial two segments of the trajectory are connected by a home run preplay ℓ_{20} (dashed line). **B** Reaching the next location, σ_3 , allows preplaying home runs to σ_0 and σ_1 and introducing the intersections $\bar{\sigma}_1$ and $\bar{\sigma}_2$ into the map. **C** The direction of the alignment connecting the new points $\bar{\sigma}_1$ and $\bar{\sigma}_2$ remain undefined, so $\ell_x = \{\bar{\sigma}_1, \bar{\sigma}_2 | \eta_x\}$ and may thus be viewed as provisional (topological) alignment in $\mathcal{A}_C(t_3)$. **D** The location σ_4 induces at several additional preplays to previously visited locations. **E** At each step t_n , the adjacent segments of the trajectory, $\sigma_i \bullet \sigma_j$ induce a connectivity graph $\mathcal{G}_\sigma(t_n)$, and the corresponding clique complex $\mathcal{T}_\sigma^\ell(\mathcal{G}_\sigma)$ schematically represents the topological structure of the aligned map $\mathcal{A}_C(t_n)$. **F** The decays of the cell assemblies representing the unvisited locations (shaded areas) induces the required topological dynamics

of the trajectory can generate ever larger sets of locations and alignments but the map learning process can be terminated when the map $\mathcal{A}_C(t_n)$ stabilizes topologically (see below). At each step, the acquired collection of alignments embedded into the unfolding cognitive map sustains its ongoing geometric structure.

Other alignments may be produced by more complex relationships within the existing configurations and their maps, as suggested, e.g., by Desargues or Pappus theorems. However, in finite planes these relationships may not be necessitated by the incidence axioms and require additional properties and implementing mechanisms (Hilbert 1992; Hilbert and Cohn-Vossen 1999; Batten 1997; Kerteszi 1976). This scope of questions falls beyond this discussion and will be addressed elsewhere.

Topological quantification of geometric learning. The assumption of the model is that the influx of endogenously generated σ - and ζ -simplexes accelerates the emergence of an aligned cognitive map with the correct topological shape. Testing this hypothesis requires computing the persistent homologies of the corresponding *aligned complexes* $\mathcal{T}_\sigma^\ell(t)$ and $\mathcal{T}_\zeta^\ell(t)$; however, the algorithm described above produces the locations σ_i and their appearance times t_i without specifying cells that comprise a given simplex or detailing how these cells are shared between simplexes, which is required for the homological computations.

To extract the needed information, consider a graph $\mathcal{G}_\sigma$ whose links correspond to the adjacent simplexes, i.e., vertexes $v_i, v_j \in \mathcal{G}_\sigma$ are connected if $\sigma_i \bullet \sigma_j$. If each σ_i acts as an assembly, i.e., ignites when all of its vertex-cells (3) activate and if the adjacent simplexes share vertexes, i.e., $\sigma_i \bullet \sigma_j \Leftrightarrow \sigma_i \cap \sigma_j \neq \emptyset$ (required for spatiotemporal contiguity, see Babichev and Dabaghian (2018)), then each $\mathcal{G}_\sigma$ -link marks at least one putative cell c_k shared by σ_i and σ_j . In a conservative estimate (assuming, e.g., no “redundant” cells that manifest themselves within just one assembly), the set of $\mathcal{G}_\sigma$ -links terminating at a given vertex σ thus defines the neuronal decomposition (3) of the corresponding simplex. Same analyses allow restoring neuronal decompositions for η -simplexes and constructing the ζ -simplexes, thus producing cliques simplicial complexes $\mathcal{T}_\sigma^\ell(t)$, $\mathcal{T}_\eta^\ell(t)$ and $\mathcal{T}_\zeta^\ell(t)$.

If, according to the requirement **R1**, the resulting σ - and η -fields cover their respective representing spaces $\mathcal{E}$ and S^1 , then the ζ -fields cover the orientation space $\mathcal{O} = \mathcal{E} \times S^1$, and the nerves associated with these covers, along with their temporal representations, $\mathcal{T}_\sigma^\ell(t)$ and $\mathcal{T}_\zeta^\ell(t)$, should have the required topological properties. However, this argument has a principal caveat: some locations induced through endogenous network activity may correspond to physically inaccessible domains in $\mathcal{E}$, which may divert the evolution of the resulting coactivity complex from the topology of the place field nerve of the navigated environment. Simulations show that indeed, the “autonomously constructed” complex $\mathcal{T}_\sigma^\ell(t)$ tends to acquire a trivial shape ($b_{n>0}(\mathcal{T}_\sigma^\ell) = 0$) irrespective of the shape of the underlying $\mathcal{E}$.

A solution to this problem may be based on exploring functional differences between place cell combinations $\hat{\sigma}_i$ that represent “physically allowed” locations and the combinations $\hat{\sigma}_k$ that represent “physically prohibited” regions. One would expect that in a confined environment, the former kind of cell groups should reactivate regularly due to animal’s (re)visits, whereas the latter kind is never “validated” through actual exploration— $\hat{\sigma}_k$ s may activate only during the occasional preplays or replays. Taking advantage of this difference, let us assume that cell assemblies have a finite lifetimes (Buzsaki 2010; Harris 2005), i.e., that 1) the probability of an assembly’s disappearance after an inactivity period t_σ is

$$p_\sigma \simeq e^{-t_\sigma/\tau_\sigma},$$

where τ_σ is σ ’s mean decay period, and 2) that the decay process resets ($t_\sigma = 0$) after each reactivation of σ (for some physiological motivations and references see Babichev et al. (2018, 2019)).

To emphasize the contribution of the locations imprinted into the network structure due to physical activity over the computationally induced locations, the latter may be attributed with a shorter decay period, $\tau_{\hat{\sigma}} = \tau \ll T_{\min}^\sigma$, whereas the former may be treated as semi-stable $\tau_{\hat{\sigma}} \gg \tau_{\hat{\sigma}}$, e.g., for basic estimates, one can use $\tau_{\hat{\sigma}} = \infty$. Lastly, the transition between $\hat{\sigma}$ s and $\hat{\sigma}$ s is modeled by stabilizing the decaying assemblies upon validation, i.e.,

$$\tau_\sigma = \begin{cases} \tau_{\hat{\sigma}} = \tau & \text{before animal visits } \nu_\sigma, \\ \tau_{\hat{\sigma}} = \infty & \text{after animal visits } \nu_\sigma. \end{cases}$$

With this plasticity rule, physically permitted locations $\hat{\sigma}$ should maintain their presence in the map, whereas the prohibited locations $\hat{\sigma}$ should decay, revealing the physical shape of the environment (Fig. 5E).

To verify this approach, the semi-random foraging trajectory simulated in Section 2 was replaced with a polygonal chain trajectory consisting of straight moves and random yaws. The preplays were then modeled by injecting straight alignments into the coactivity graph as soon as the required information became available (for details see Babichev et al. (2019)). Based on the results of Babichev et al. (2018, 2019), the decay rate $\tau = 0.5$ secs was selected to model the dynamics of the unstable locations.

For these parameters, the homological characteristics of the resulting “flickering” coactivity complex evaluated using ZigZag Homology techniques (Carlsson and Silva 2010; Carlsson et al. 2009; Edelsbrunner et al. 2002), quickly became stable: the Betti numbers stabilized at $b_{n \leq 1}(T_{\zeta}^{\ell}) = 1$, $b_{n > 1}(T_{\zeta}^{\ell}) = 0$, in $T_{\min}^{\zeta} \approx 5$ minutes, which approximately matches the hippocampal learning time $T_{\min}^{\sigma}$ and demonstrates that geometric organization of the cognitive map brings learning dynamics to the biologically viable timescale.

5 Discussion

The proposed models of orientation learning are built by combining inputs from the hippocampal place cells and the head directions cells. Experiments demonstrate that these two populations of neurons are coupled: in slowly deforming environments, their spiking activities remain highly correlated, pointing at a unified cognitive spatial framework that involves both locations and orientations (Knierim et al. 1995; Yoganarasimha and Knierim 2005; Hargreaves et al. 2007; Sargolini et al. 2006). The goal of this study is to combine topological and geometric approaches to model such a framework, and to evaluate the corresponding learning dynamics.

The first model (Sect. 2) is based on the observation that both the hippocampal and the head direction maps are of a topological nature: while the place cells encode a qualitative, elastic map of the navigated environment $\mathcal{E}$ (Gothard et al. 1996; Alvernhe et al. 2008, 2011, 2012; Dabaghian et al. 2014; Wu and Foster 2014), the head direction cells map the space of directions, S^1 (Taube 1998). A combination of place and head direction cells’ inputs can hence be used to construct an extended topological map of oriented locations $\mathcal{O}$, which has a structure of a direct product $\mathcal{E} \times S^1$ —a natural framework for describing the kinematics of rats’ movements.

The second model (Sect. 3) is structurally similar (a discrete map of directions is associated with each location), but involves constructions that define an additional, geometric layer of the cognitive map’s architecture. In particular, this model allows viewing spatial orientation learning from a geometric perspective—not only as a process of discovering connections between locations, but also establishing shapes of location arrangements, e.g., straight or turning paths, their incidences, intersections, junctions, etc. In neuroscience literature, such references are commonly made in relation to the physical geometry of the representing spaces $\mathcal{E}$ and S^1 , e.g., the “straightness” of a σ -field arrangement implies simply that it can be matched by a Euclidean line in the environment where the rat is observed (Valerio and Taube 2012; McNaughton et al. 2006). However, understanding the geometric structure of the cognitive map requires interpreting neuronal activity in systems’ own terms, rather than through the parameters of exterior geometry.

The key observation underlying the geometric model is that the activity of head direction cells “tags” the activity of place cells in a way that allows an *intrinsic* geometric interpretation of the combined spiking patterns, i.e., defining alignments, turns, yaws, etc., in terms of neuronal spiking parameters. A famous quote attributed to D. Hilbert proclaims that “*the axioms of geometry would be just as valid if one replaced*

the undefined terms ‘point, line, and plane’ with ‘table, chair, and beer mug’...” (Blumenthal 1935). From such perspective, this model aims at constructing a synthetic “location and compass” neuro-geometry in terms of the temporal relationships between the spike trains, without using extrinsic references or *ad hoc* measures, which may be a general principle for how space and geometry emerge from neuronal activity.

In order to emphasize connections with conventional geometries, the model is formulated in a semi-axiomatic form. However, in contrast with the standard affine A -axioms or their direct analogues, the A_n -axioms, the rules **G1–G3** serve not just as formal assertions that lay logical foundations for geometric deductions, but also as reflections of physiological properties of the networks that implement the computations. First, since the networks contain a finite number of neurons and may actualize a finite set of locations and alignments, the emergent geometry is finitary. Second, certain notions that in standard discrete affine plane $\mathcal{A}$ are introduced indirectly, *relationally*, become *constructive* in the “cognitive” affine plane. For example, directions defined through equivalence classes of parallel lines in $\mathcal{A}$ (Hilbert 1992; Hilbert and Cohn-Vossen 1999; Batten 1997; Karteszi 1976), are defined explicitly in $\mathcal{A}_C$ using the directing η -activities. Third, certain elements of the geometric structure are actualized explicitly through the network’s architecture, e.g., a fixed number of alignments passing through every location is implemented by cell assemblies wired into the head direction network (Bassett et al. 2018; Redish et al. 1996). Other properties are not prewired but *acquired* during a particular learning experience and reflect both the physical structure of a specific environment and intrinsic mechanisms of spatial information processing.

6 Appendix: Spatial learning

Spatial learning in mammals is a complex, distributed process, sustained by coherent effort of many strongly interconnected brain structures, which jointly yield an internalized, coherent, cognitively accessible representation of a physical environment—a cognitive map (Grieves and Jeffery 2017; Moser et al. 2008; Derdikman and Moser 2011; McNaughton 1996; Lisman et al. 2017). A key role in this process is played by hippocampus—evolutionarily one of the oldest brain parts, that is critical for rapid pattern encoding and erasure, storing conjunctive associations between stimuli, cortical representations, etc. (Kim and Frank 2009; Madroñal et al. 2016; Kragel et al. 2020[103] Cheng 2013; Zhang and Manahan-Vaughan 2015; Poulter et al. 2018).

A key property that links hippocampal principal neurons to spatial learning is spatial selectivity of their spiking: each cell fires in its preferred location that is largely independent of the animal’s behavior, body or head orientation, movement direction, etc. (O’Keefe and Nadel 1978; O’Keefe and Dostrovsky 1971; Derdikman and Moser 2011; Grieves and Jeffery 2017; Best and White 1998; Moser et al. 2008; McNaughton 1996; Lisman et al. 2017). Spatial layout of the firing fields hence controls the order of the hippocampal neurons’ activity and is therefore commonly viewed as the main correlate of the cognitive map (Frank et al. 2000; Guger et al. 2011; Brown et al. 1998; Jensen and Lisman 2000; Zhang et al. 1998; Huang et al. 2009), although the exact link between the two remains opaque Vorhees and Williams (2014). A common

(and historically oldest) approach views spatial learning as tuning the individual cells to their respective firing fields. Once such tuning is complete for most cells in the ensemble (for rats' place fields this takes about four minutes Frank et al. 2000), it is presumed that a large-scale hippocampal map has emerged (Frank et al. 2000; Wilson and McNaughton 1993; Wills et al. 2010). Further learning is tied to slower information flow from hippocampus to cortex and back—a multifaceted process happening at several timescales, that depends on specific experiences, number of exposures, behavioral control, modality, synaptic mechanisms, etc. (Rothschild et al. 2017; Weber and Sprekeler 2018; Rolls 2018; Michelmann et al. 2021).

Analyses of these phenomena are based on interpreting population activity (Wilson and McNaughton 1993; Mau et al. 2020; Zhang et al. 1998; Pouget et al. 2000), and building specific spatial representations from spiking data (Savelli and Knierim 2019; Dabaghian 2021; Kang et al. 2021; Gluck et al. 2003). However, the principles that could guide or constrain such constructions remain controversial. For example, the shapes and the locations of place fields can be affected by visual, olfactory, vestibular, kinesthetic, and other cues as well as more subtle cognitive aspects of navigation, e.g., the animal's goals and discrepancies between the expected and actual location of navigational targets (Wood et al. 2000; Butler et al. 2019; Chen et al. 2013; Erdem and Hasselmo 2012; Kragel et al. 2020; Ji and Wilson 2007; Rothschild et al. 2017). It is hence unclear which part of this information is embedded in the place cell spiking and transmitted to downstream neurons, along with the consequences for the type of spatial properties that might form the basis for the cognitive map.

The dominant assumption within the field has been that the hippocampal map incorporates detailed geometric information arising through proprioceptive and sensory cues (Solomon et al. 2019; Stella et al. 2013; Terrazas et al. 2005; Moser and Moser 2008). However, a growing volume of experimental results suggest that hippocampal map is topological in nature, providing a schematic representation of space, akin to a subway-map (Alvernhe et al. 2008, 2011, 2012; Gothard et al. 1996; Dabaghian et al. 2014; Place and Nitz 2020; Krupic et al. 2018; Chen et al. 2014, 2012). From a biological standpoint, such maps may be viewed as rough-and-ready frameworks into which geometric details could be situated over time (Chen et al. 2013; Erdem and Hasselmo 2012; Yoder et al. 2011; Knierim et al. 1995, 1998; Jercog et al. 2019). Hippocampal ability to acquire and store conjunctive associations allows coupling the outputs of the animal's directional system, notably head direction cells, with hippocampal connectivity map (Zhang 1996; Zugaro et al. 2003; Hargreaves et al. 2007).

Since many specifics of this process remain unknown, modeling may be based on describing the structure of the information flow itself, without detailing the mechanisms that may produce or process the contributing spiking activity. The topological nature of the cognitive map suggests using topological analyses, notably persistent and zigzag homology methods that allow assessing the map's topological dynamics (Chen et al. 2014, 2012; Dabaghian et al. 2012; Kang et al. 2021), taking into account a wide scope of physiological phenomena—brain waves (Arai et al. 2014; Basso et al. 2016; Hoffman et al. 2016), synaptic imperfections Dabaghian (2018), plasticity (Babichev et al. 2018, 2019), internal dynamics (replays and preplays) (Babichev et al. 2019; Wu and Foster 2014; Ólafsdóttir et al. 2018), etc. (Dabaghian 2021). The proposed model aims to extend this approach to address a long standing problem: how complementary

inputs (topological, geometrical or non-spatial) provided by different brain parts may combine into a full representation of the environment (Wood et al. 2000; van der Veldt et al. 2021; Sargolini et al. 2006; O'Reilly and McClelland 1994; Komorowski et al. 2009; Eichenbaum 2004).

7 Methods

The computational algorithms used in this study were described in (Dabaghian et al. 2012; Arai et al. 2014; Basso et al. 2016; Hoffman et al. 2016; Babichev et al. 2016; Babichev and Dabaghian 2018; Babichev et al. 2016, 2018, 2019).

The environment shown on Fig. 1A is similar to the arenas used in electrophysiological experiments (Hafting et al. 2005)[137]. The simulated trajectory represents exploratory spatial behavior that does not favor one segment of the environment over another.

Place cell spiking probability was modeled as a Poisson process with the rate

$$\lambda(r) = f e^{-\frac{(r-r_0)^2}{2s^2}},$$

where f is the maximal rate and s defines the size of the firing field centered at $r_0 = (x_0, y_0)$ Barbieri et al. (2004). In addition, spiking probability was modulated by the θ -waves, which also define the temporal window $w \approx 250$ ms (about two θ -periods) for detecting the place cell spiking coactivity (Arai et al. 2014; Mizuseki et al. 2009). The place field centers r_0 for each computed place field map were randomly and uniformly scattered over the environment.

Persistent Homology Theory computations were performed using Javaplex computational software (Adams et al. 2014) as described in Dabaghian et al. (2012); Arai et al. (2014); Basso et al. (2016); Hoffman et al. (2016); Babichev et al. (2016). Usage of zigzag persistent homology methods is described in Babichev et al. (2018, 2019).

Acknowledgements The author is grateful to D. Morozov for providing Zigzag persistent homology simulating software. The work was supported by the NSF grant 1901338.

Declarations

Conflict of interest The author states that there is no conflict of interest.

References

- Adams H., Tausz A., Vejdemo-Johansson M.: javaPlex: A Research Software Package for Persistent (Co)Homology. In: Hong H., Yap C. (eds) *Mathematical Software – ICMS 2014*. ICMS 2014. Lecture Notes in Computer Science, vol 8592. Springer, Berlin (2014). <http://appliedtopology.github.io/javaplex/>
- Alvernhe, A., Van Cauter, T., Save, E., Poucet, B.: Different CA1 and CA3 representations of novel routes in a shortcut situation. *J. Neurosci.* **28**(29), 7324–33 (2008)
- Alvernhe, A., Save, E., Poucet, B.: Local remapping of place cell firing in the Tolman detour task. *Eur. J. Neurosci.* **33**(9), 1696–705 (2011)
- Alvernhe, A., Sargolini, F., Poucet, B.: Rats build and update topological representations through exploration. *Animal Cogn.* **15**, 359–368 (2012)
- Ambrose, R., Pfeiffer, B., Foster, D.: Reverse replay of hippocampal place cells is uniquely modulated by changing reward. *Neuron* **91**(5), 1124–36 (2016)
- Arai, M., Brandt, V., Dabaghian, Y.: The effects of theta precession on spatial learning and simplicial complex dynamics in a topological model of the hippocampal spatial map. *PLoS Comput. Biol.* **10**, 1003651 (2014)
- Babichev, A., Dabaghian, Y.: Topological schemas of memory spaces. *Front. Comput. Neurosci.* (2018). <https://doi.org/10.3389/fncom.2018.00027>
- Babichev, A., Cheng, S., Dabaghian, Y.: Topological schemas of cognitive maps and spatial learning. *Front. Comput. Neurosci.* **10**, 18 (2016)
- Babichev, A., Mémoli, F., Ji, D., Dabaghian, Y.: A topological model of the hippocampal cell simplex network. *Front. Comput. Neurosci.* **10**, 50 (2016)
- Babichev, A., Morozov, D., Dabaghian, Y.: Robust spatial memory maps encoded by networks with transient connections. *PLoS Comput. Bio.* **14**(9), 1006433 (2018)
- Babichev, A., Morozov, D., Dabaghian, Y.: Replays of spatial memories suppress topological fluctuations in cognitive map. *Netw. Neurosci.* **3**(2), 1–18 (2019)
- Barbieri, R., Frank, L., Nguyen, D., Quirk, M., Solo, V., Wilson, M., Brown, E.: Dynamic analyses of information encoding in neural ensembles. *Neural Comput.* **16**, 277–307 (2004)
- Bassett, J., Wills, T., Cacucci, F.: Self-organized attractor dynamics in the developing head direction circuit. *Current Biol.* **28**(4), 609–15.e3 (2018)
- Basso, E., Arai, M., Dabaghian, Y.: The effects of γ -synchronization on spatial learning in a topological model of the hippocampal spatial map. *PLoS Comput. Biol.* **12**, 9 (2016)
- Bast T., Wilson, I., Witter, M., Morris, R.: From Rapid Place Learning to Behavioral Performance: A Key Role for the Intermediate Hippocampus. *PLoS Biol.* **7**(4): e1000089
- Batten, L.: *Combinatorics of Finite Geometries*. Cambridge University Press, Cambridge (1997)
- Best, P., White, A.: Hippocampal cellular activity: a brief history of space. *Proc. Natl. Acad. Sci.* **95**, 2717–19 (1998)
- Blumenthal, O.: *Lebensgeschichte*. In: *David Hilbert Gesammelte Abhandlungen*, **3**: 388–429 Springer-Verlag, Berlin (1935)
- Brown, R., Donald, O.: Hebb and the Organization of Behavior: 17 years in the writing. *Mol. Brain* **13**, 55 (2020)
- Brown, E., Frank, L., Tang, D., Quirk, M., Wilson, M.: A statistical paradigm for neural spike train decoding applied to position prediction from ensemble firing patterns of rat hippocampal place cells. *J. Neurosci.* **18**, 7411–7425 (1998)
- Brun, V., Solstad, T., Kjelstrup, K., Fyhn, M., Witter, M., Moser, E., Moser, M-B.: Progressive increase in grid scale from dorsal to ventral medial entorhinal cortex
- Butler, W., Hardcastle, K., Giocomo, L.: Remembered reward locations restructure entorhinal spatial maps. *Science* **363**(6434), 1447–1452 (2019)
- Buzsaki, G.: Neural syntax: cell assemblies, synapsembles, and readers. *Neuron* **68**, 362–385 (2010)
- Byrne, P., Becker, S.: Modeling mental navigation in scenes with multiple objects. *Neural Comput.* **16**(9), 1851–72 (2004)
- Carlsson, G., Silva, Vd., Morozov, D.: Zigzag persistent homology and real-valued functions. *Proceedings of the 25th annual symposium on Computational geometry*. Aarhus, Denmark: ACM: 247–256 (2009)
- Carlsson, G., Silva, Vd.: Zigzag Persistence. *Found. Comput. Math.* **10**, 367–405 (2010)
- Caroni, P., Donato, F., Muller, D.: Structural plasticity upon learning: regulation and functions. *Nat. Rev. Neurosci.* **13**, 478–490 (2012)

- Cei, A., Girardeau, G., Drieu, C., El Kanbi, K., Zugaro, M.: Reversed theta sequences of hippocampal cell assemblies during backward travel. *Nat. Neurosci.* **17**(5), 719–724 (2014)
- Chen, G., King, J., Burgess, N., O'Keefe, J.: How vision and movement combine in the hippocampal place code. *Proc. Natl. Acad. Sci.* **110**(1), 378–83 (2013)
- Chen, Z., Kloosterman, F., Brown, E., Wilson, M.: Uncovering spatial topology represented by rat hippocampal population neuronal codes. *J. Comput. Neurosci.* **33**, 227–255 (2012)
- Chen, L., Lin, L., Green, E., Barnes, C., McNaughton, B.: Head-direction cells in the rat posterior cortex. *I. Exp. Brain Res.* **101**(1), 8–23 (1994)
- Chen, Z., Gomperts, S., Yamamoto, J., Wilson, M.: Neural representation of spatial topology in the rodent hippocampus. *Neural Comput.* **26**(1), 1–39 (2014)
- Chen, G., King, J., Lu, Y., Cacucci, F., Burgess, N.: Spatial cell firing during virtual navigation of open arenas by head-restrained mice. *Elife* **7**, 34789 (2018)
- Cheng, S.: The CRISP theory of hippocampal function in episodic memory. *Front. Neural Circuits* **7**(88), 1 (2013)
- Curto, C., Itskov, V.: Cell groups reveal structure of stimulus space. *PLoS Comput. Biol.* **4**, 1000205 (2008)
- Dabaghian, Y., Brandt, V., Frank, L.: Reconciving the hippocampal map as a topological template. *eLife* (2014). <https://doi.org/10.7554/eLife.03476>
- Dabaghian, Y.: From Topological Analyses to Functional Modeling: The Case of Hippocampus. *Front. Comput. Neurosci.* **14**, (2021)
- Dabaghian, Y.: Through synapses to spatial memory maps: a topological model. *Sci. Reps.* **9**, 572 (2018)
- Dabaghian, Y., Mémoli, F., Frank, L., Carlsson, G.: A topological paradigm for hippocampal spatial map formation using persistent homology. *PLoS Comput. Biol.* **8**, 1002581 (2012)
- Derdikman, D., Moser, E.: A manifold of spatial maps in the brain. In *Space, Time and Number in the Brain*, pp. 41–57. Academic Press (2011)
- Dragoi, G., Tonegawa, S.: Preplay of future place cell sequences by hippocampal cellular assemblies. *Nature* **469**, 397–401 (2011)
- Edelsbrunner, H., Letscher, D., Zomorodian, A.: Topological persistence and simplification. *Disc. Comput. Geom.* **28**, 511–533 (2002)
- Edelsbrunner, H., Harer, J.: Computational topology: an introduction. *Amer. Math. Soc* (2010)
- Eichenbaum, H.: Hippocampus: cognitive processes and neural representations that underlie declarative memory. *Neuron* **44**, 109–120 (2004)
- Erdem, U., Hasselmo, M.: A goal-directed spatial navigation model using forward trajectory planning based on grid cells. *Eur. J. Neurosci.* **35**(6), 916–931 (2012)
- Foster, D., Wilson, M.: Reverse replay of behavioural sequences in hippocampal place cells during the awake state. *Nature* **440**, 680–683 (2006)
- Frank, L., Brown, E., Wilson, M.: Trajectory encoding in the hippocampus and entorhinal cortex. *Neuron* **27**(1), 169–178 (2000)
- Gluck, A., Meeter, M., Myers, C.: Computational models of the hippocampal region: linking incremental learning and episodic memory. *Trends Cogn. Sci.* **7**(6), 269–276 (2003)
- Gothard, K., Skaggs, W., McNaughton, B.: Dynamics of mismatch correction in the hippocampal ensemble code for space: interaction between path integration and environmental cues. *J. Neurosci.* **16**(24), 8027–8040 (1996)
- Grievens, R., Jeffery, K.: The representation of space in the brain. *Behav. Process.* **135**, 113–31 (2017)
- Guger, C., Gener, T., Pennartz, C., Brotons-Mas, J., Edlinger, G., Bermúdez, I., Badia, S., Verschure, P., Schaffelhofer, S., Sanchez-Vives, M.: Real-time position reconstruction with hippocampal place cells. *Front. Neurosci.* **5**, 85 (2011)
- Hafting, T., Fyhn, M., Molden, S., Moser, M.-B., Moser, E.I.: Microstructure of a spatial map in the entorhinal cortex. *Nature* **436**, 801–806 (2005)
- Haggerty, D., Ji, D.: Activities of visual cortical and hippocampal neurons co-fluctuate in freely moving rats during spatial behavior. *eLife* **4**, e08902 (2015)
- Hargreaves, E., Yoganarasimha, D., Knierim, J.: Cohesiveness of spatial and directional representations recorded from neural ensembles in the anterior thalamus, parasubiculum, medial entorhinal cortex, and hippocampus. *Hippocampus* **17**, 826–841 (2007)
- Harris, K.: Neural signatures of cell assembly organization. *Nat. Rev. Neurosci.* **6**, 399–407 (2005)
- Hatcher, A.: Algebraic Topology. Cambridge University Press, Cambridge; New York (2002)
- Hebb, D.: The organization of behavior: A neuropsychological theory. Wiley; Chapman & Hall, Hoboken (1949)

- Heinze, S., Narendra, A., Cheung, A.: Principles of insect path integration. *Current Biol.* **28**(17), R1043–R1058 (2018)
- Hilbert, D.: *The Foundations of Geometry*. Open Court, La Salle, IL (1992)
- Hilbert, D., Cohn-Vossen, S.: *Geometry and the imagination*. AMS Chelsea Pub, Providence (1999)
- Hoffman, K., Babichev, A., Dabaghian, Y.: A model of topological mapping of space in bat hippocampus. *Hippocampus* **26**, 1345–1353 (2016)
- Hopfield, J.: Neurodynamics of mental exploration. *Proc. Natl. Acad. Sci.* **107**, 1648–1653 (2010)
- Huang, Y., Brandon, M., Griffin, A., Hasselmo, M., Eden, U.: Decoding movement trajectories through a T-maze using point process filters applied to place field data from rat hippocampal region CA1. *Neural Comput.* **21**(12), 3305–3334 (2009)
- Jensen, O., Lisman, J.E.: Position reconstruction from an ensemble of hippocampal place cells: contribution of theta phase coding. *J. Neurophysiol.* **83**, 2602–2609 (2000)
- Jercog, P., Ahmadian, Y., Woodruff, C., Deb-Sen, R., Abbott, L., Kandel, E.: Heading direction with respect to a reference point modulates place-cell activity. *Nat. Commun.* **10**(1), 1–8 (2019)
- Ji, D., Wilson, M.: Coordinated memory replay in the visual cortex and hippocampus during sleep. *Nat. Neurosci.* **10**, 100–7 (2007)
- Johnson, A., Redish, A.: Neural Ensembles in CA3 transiently encode paths forward of the animal at a decision point. *J. Neurosci.* **27**, 12176–89 (2007)
- Kang, L., Xu, B., Morozov, D.: Evaluating state space discovery by persistent cohomology in the spatial representation system. *Front. Comput. Neurosci.* **15**(28), 616748 (2021)
- Karlsson, M., Frank, L.: Awake replay of remote experiences in the hippocampus. *Nat. Neurosci.* **12**, 913–918 (2009)
- Kartesz, F. *Introduction to Finite Geometries*, Disquisitiones Math. Hung., 7, Akademiai Kiado, Budapest (1976)
- Kim, S., Frank, L.: Hippocampal lesions impair rapid learning of a continuous spatial alternation task. *PLoS One* **4**(5), 5494 (2009)
- Knierim, J., Kudrimoti, H., McNaughton, B.: Place cells, head direction cells, and the learning of landmark stability. *J. Neurosci.* **15**(3), 1648–59 (1995)
- Knierim, J., Kudrimoti, H., McNaughton, B.: Interactions between idiothetic cues and external landmarks in the control of place cells and head direction cells. *J. Neurophysiol.* **80**(1), 425–46 (1998)
- Komorowski, R., Manns, J., Eichenbaum, H.: Robust conjunctive item-place coding by hippocampal neurons parallels learning what happens where. *J. Neurosci.* **29**(31), 9918–29 (2009)
- Kragel, J., VanHaerents, S., Templer, J., Schuele, S., Rosenow, J., Nilakantan, A., Bridge, D.: Hippocampal theta coordinates memory processing during visual exploration. *eLife* **9**, e52108 (2020)
- Krupic, J., Bauza, M., Burton, S., O'Keefe, J.: Local transformations of the hippocampal cognitive map. *Science* **359**(6380), 1143–1146 (2018)
- Laurens, J., Angelaki, D.: The brain compass: a perspective on how self-motion updates the head direction cell attractor. *Neuron* **97**(2), 275–89 (2018)
- Leuner, B., Gould, E.: Structural plasticity and hippocampal function. *Annu. Rev. Psychol.* **61**, 111–140 (2010)
- Lisman, J., Buzsáki, G., Eichenbaum, H., Nadel, L., Ranganath, C., Redish, A.: Viewpoints: how the hippocampus contributes to memory, navigation and cognition. *Nat. Neurosci.* **20**, 1434–48 (2017)
- Louie, K., Wilson, M.: Temporally structured replay of awake hippocampal ensemble activity during rapid eye movement sleep. *Neuron* **29**, 145–156 (2001)
- Madroñal, N., Delgado-García, J., Fernández-Guizán, A., Chatterjee, J., Köhn, M., Mattucci, C., Jain, A., Tsetsenis, T., Illarionova, A., Grinevich, V., Gross, C., Gruart, A.: Rapid erasure of hippocampal memory following inhibition of dentate gyrus granule cells. *Nat. Commun.* **18**(7), 10923 (2016)
- Mattar, M., Daw, N.: Prioritized memory access explains planning and hippocampal replay. *Nat. Neurosci.* **21**(11), 1609–1617 (2018)
- Mau, W., Hasselmo, M., Cai, D.: The brain in motion: How ensemble fluidity drives memory-updating and flexibility. *eLife* **9**, e63550 (2020)
- McNaughton, B., Battaglia, F., Jensen, O., Moser, E., Moser, M.: Path integration and the neural basis of the “cognitive map”. *Nat. Rev. Neurosci.* **7**(8), 663–78 (2006)
- McNaughton, B.: Cognitive cartography. *Nature* **381**(6581), 368–9 (1996)
- Michelmann, S., Price, A., Aubrey, B., Strauss, C., Doyle, W., Friedman, D., Dugan, P., Devinsky, O., Devore, S., Flinker, A., Hasson, U., Norman, K.: Moment-by-moment tracking of naturalistic learning and its underlying hippocampo-cortical interactions. *Nat. Comm.* **12**(1), 5394 (2021)

- Mizuseki, K., Sirota, A., Pastalkova, E., Buzsaki, G.: Theta oscillations provide temporal windows for local circuit computation in the entorhinal-hippocampal loop. *Neuron* **64**, 267–280 (2009)
- Moser, E., Kropff, E., Moser, M.-B.: Place cells, grid cells, and the brain's spatial representation system. *Ann. Rev. Neurosci.* **31**(1), 69–89 (2008)
- Moser, E., Moser, M.-B.: A metric for space. *Hippocampus* **18**, 1142–1156 (2008)
- Muller, R., Ranck, J., Jr., Taube, J.: Head direction cells: properties and functional significance. *Curr. Opin. Neurobiol.* **6**, 196–206 (1996)
- O'Keefe, J., Nadel, L.: *The Hippocampus as a cognitive map*. Oxford, London (1978)
- O'Keefe, J., Dostrovsky, J.: The hippocampus as a spatial map. Preliminary evidence from unit activity in the freely-moving rat. *Brain Res.* **34**(1), 171–175 (1971)
- Ólafsdóttir, H., Bush, D., Barry, C.: The Role of Hippocampal Replay in Memory and Planning. *Current Biol.* **28**(1), R37–R50 (2018)
- O'Reilly, R., McClelland, J.: Hippocampal conjunctive encoding, storage, and recall: avoiding a trade-off. *Hippocampus* **4**(6), 661–682 (1994)
- Perin, R., Berger, T., Markram, H.: A synaptic organizing principle for cortical neuronal groups. *Proc. Natl. Acad. Sci.* **108**(13), 5419–24 (2011)
- Pfeiffer, B., Foster, D.: Hippocampal place-cell sequences depict future paths to remembered goals. *Nature* **497**, 74–9 (2013)
- Piet, A., El Hady, A., Brody, C.: Rats adopt the optimal timescale for evidence integration in a dynamic environment. *Nat. Commun.* **9**(1), 4265–77 (2018)
- Place, R., Nitz, D.: Cognitive maps: distortions of the hippocampal space map define neighborhoods. *Current Biol.* **30**(8), R340–R342 (2020)
- Pouget, A., Dayan, P., Zemel, R.: Information processing with population codes. *Nat. Rev. Neurosci.* **1**, 125–132 (2000)
- Poulter, S., Hartley, T., Lever, C.: The neurobiology of mammalian navigation. *Current Biol.* **28**(17), R1023–R1042 (2018)
- Raudies, F., Brandon, M.P., Chapman, G.W., Hasselmo, M.: Head direction is coded more strongly than movement direction in a population of entorhinal neurons. *Brain Res.* **1621**, 355–367 (2015)
- Redish, A., Elga, A., Touretzky, D.: A coupled attractor model of the rodent Head Direction system. *Netw. Comput. Neural Syst.* **7**(4), 671–685 (1996)
- Reimann, M., Nolte, M., Scolamiero, M., Turner, K., Perin, R., Chindemi, G., Dlotko, P., Levi, R., Hess, K., Markram, H.: Cliques of Neurons Bound into Cavities Provide a Missing Link between Structure and Function. *Front. Comput. Neurosci.* **11**(48) (2017)
- Rolls, E.: The storage and recall of memories in the hippocampo-cortical system. *Cell Tissue Res.* **373**(3), 577–604 (2018)
- Rothschild, G., Eban, E., Frank, L.: A cortical-hippocampal-cortical loop of information processing during memory consolidation. *Nat. Neurosci.* **20**(2), 251–259 (2017)
- Roux, L., Hu, B., Eichler, R., Stark, E., Buzsaki, G.: Sharp wave ripples during learning stabilize the hippocampal spatial map. *Nat. Neurosci.* **20**, 845–853 (2017)
- Rubin, A., Yartsev, M., Ulanovsky, N.: Encoding of head direction by hippocampal place cells in bats. *J. Neurosci.* **34**(3), 1067–1080 (2014)
- Sargolini, F., Fyhn, M., Hafting, T., McNaughton, B., Witter, M., Moser, M., Moser, E.: Conjunctive representation of position, direction, and velocity in entorhinal cortex. *Science* **312**(5774), 758–62 (2006)
- Savelli, F., Knierim, J.: Origin and role of path integration in the cognitive representations of the hippocampus: Computational insights into open questions. *J. Exp. Biology*, **222**:jeb188912 (2019)
- Shinder, M., Taube, J.: Active and passive movement are encoded equally by head direction cells in the anterodorsal thalamus. *J. Neurophysiol.* **106**(2), 788–800 (2011)
- Shinder, M.E., Taube, J.S.: Self-motion improves head direction cell tuning. *J. Neurophysiol.* **111**, 2479–2492 (2014)
- Solomon, E., Lega, B., Sperling, M., Kahana, M.: Hippocampal theta codes for distances in semantic and temporal spaces. *Proc. Natl. Acad. Sci.* **116**(48), 24343–24352 (2019)
- Stella, F., Cerasti, E., Treves, A.: Unveiling the metric structure of internal representations of space. *Front. Neural Circuits.* **7**, (2013)
- Taube, J.: Head direction cells and the neurophysiological basis for a sense of direction. *Prog. Neurobiol.* **55**, 225–256 (1998)

- Taube, J., Muller, R., Ranck, J.: Head-direction cells recorded from the postsubiculum in freely moving rats. *J. Neurosci.* **10**(420–435), 436–447 (1990)
- Taube, J., Goodridge, J., Golob, E., Dudchenko, P., Stackman, R.: Processing the head direction cell signal: a review and commentary. *Brain Res. Bull.* **40**, 477–484 (1996)
- Terrazas, A., Krause, M., Lipa, P., Gothard, K., Barnes, C., McNaughton, B.: Self-motion and the hippocampal spatial metric. *J. Neurosci.* **25**(35), 8085–96 (2005)
- Thrun, S., Burgard, W., Fox, D.: Probabilistic Robotics. MIT Press, Cambridge (2005)
- Tolman, E.: Cognitive maps in rats and men. *Psychol. Rev.* **55**, 189–208 (1948)
- Valerio, S., Taube, J.: Path integration: How the head direction signal maintains and corrects spatial orientation. *Nat. Neurosci.* **15**, 1445–1453 (2012)
- van de Ven, G., Trouche, S., McNamara, C., Allen, K., Dupret, D.: Hippocampal offline reactivation consolidates recently formed cell assembly patterns during sharp wave-ripples. *Neuron* **92**(5), 968–74 (2016)
- van der Veldt, S., Etter, G., Mosser, C., Manseau, F., Williams, S.: Conjunctive spatial and self-motion codes are topographically organized in the GABAergic cells of the lateral septum. *PLOS Biol.* **19**(8), 3001383 (2021)
- Vorhees, C., Williams, M.: Assessing spatial learning and memory in rodents. *ILAR J.* **55**(2), 310–332 (2014)
- Weber, S., Sprekeler, H.: Learning place cells, grid cells and invariances with excitatory and inhibitory plasticity. *eLife* **7**, e34560 (2018)
- Wiener, S., Taube, J. (eds.): Head direction cells and the neural mechanisms of spatial orientation. MIT Press (2005)
- Wills, T., Cacucci, F., Burgess, N., O'Keefe, J.: Development of the hippocampal cognitive map in preweanling rats. *Science* **328**(5985), 1573–1576 (2010)
- Wilson, M., McNaughton, B.: Dynamics of the hippocampal ensemble code for space. *Science* **261**(5124), 1055–8 (1993)
- Wood, E., Dudchenko, P., Robitsek, R., Eichenbaum, H.: Hippocampal neurons encode information about different types of memory episodes occurring in the same location. *Neuron* **27**(3), 623–33 (2000)
- Wu, X., Foster, D.: Hippocampal replay captures the unique topological structure of a novel environment. *J. Neurosci.* **34**, 6459–6469 (2014)
- Yoder, R., Clark, B., Taube, J.: Origins of landmark encoding in the brain. *Trends Neurosci.* **34**(11), 561–571 (2011)
- Yoganarasimha, D., Knierim, J.: Coupling between place cells and head direction cells during relative translations and rotations of distal landmarks. *Exper. Brain Res.* **160**, 344–359 (2005)
- Zeithamova, D., Schlichting, M., Preston, A.: The hippocampus and inferential reasoning: Building memories to navigate future decisions. *Front. Hum. Neurosci.* **266**: 70 (2012)
- Zhang, K.: Representation of spatial orientation by the intrinsic dynamics of the head-direction cell ensemble: a theory. *J. Neurosci.* **16**, 2112–2126 (1996)
- Zhang, S., Manahan-Vaughan, D.: Spatial olfactory learning contributes to place field formation in the hippocampus. *Cerebral Cortex* **25**(2), 423–432 (2015)
- Zhang, K., Ginzburg, I., McNaughton, B., Sejnowski, T.: Interpreting neuronal population activity by reconstruction: unified framework with application to hippocampal place cells. *J. Neurophysiol.* **79**(2), 1017–44 (1998)
- Zhang, S., Schönfeld, F., Wiskott, L., Manahan-Vaughan, D.: Spatial representations of place cells in darkness are supported by path integration and border information. *Front. Behav. Neurosci.* **8**, 222 (2014)
- Ziv, Y., Burns, L., Cocker, E., Hamel, E., Ghosh, K., Kitch, L., El Gamal, A., Schnitzer, M.: Long-term dynamics of CA1 hippocampal place codes. *Nat. Neurosci.* **16**(3), 264–6 (2013)
- Zomorodian, A., Carlsson, G.: Computing persistent homology. *Discrete Comput. Geom.* **33**, 249–274 (2005)
- Zugaro, M., Arleo, A., Berthoz, A., Wiener, S.: Rapid spatial reorientation and head direction cells. *J. Neurosci.* **23**(8), 3478–82 (2003)