

TEMat

Topología y neurociencia: la percepción del espacio

✉ Erika Magdalena
Herrera Machado^a
Universidad de La Laguna
erikaherreramachado@gmail.com

Resumen: La introducción de herramientas matemáticas en la neurociencia ha contribuido al rápido progreso y expansión experimental de esta disciplina. En este trabajo se pretende dar una visión de algunas aplicaciones de la topología en neurociencia. Comenzaremos con la interpretación matemática de ciertos fenómenos experimentales relacionados con las neuronas de posicionamiento. Seguidamente se definirá el código neuronal, para continuar con una pequeña introducción a los complejos simpliciales, elementos matemáticos que nos ayudarán a abordar la cuestión de interpretar estos fenómenos. Tras presentar la noción de nervio, una conocida herramienta de la topología algebraica, presentaremos un importante resultado: el lema del nervio, clave para entender el funcionamiento de las células de posicionamiento. Finalizaremos el trabajo estudiando los llamados códigos convexos y algunos resultados sobre obstrucciones locales a la convexidad de los códigos.

Abstract: The introduction of mathematical tools in neuroscience has contributed to the rapid progress and experimental expansion of this discipline. This paper aims to give a vision of some applications of topology in neuroscience. We will begin with the mathematical interpretation of certain experimental phenomena related to place neurons. The neural code will then be defined, to continue with a short introduction to simplicial complexes, the mathematical elements that will help us to interpret these phenomena. After presenting the notion of nerve, a known tool of algebraic topology, we will present an important result: the nerve lemma, key to understanding the functioning of place cells. We will finish the work by studying the so called convex codes and some results on local obstructions to the convexity of the codes.

Palabras clave: complejo simplicial, código neuronal, códigos convexos, células de posicionamiento, lema del nervio.

MSC2010: 55U05, 55U10.

Recibido: 13 de octubre de 2019.

Aceptado: 3 de febrero de 2021.

Agradecimientos: Quisiera dar las gracias a mis tutores, Francisco Javier Díaz Díaz y Edith Padrón Fernández, por su tiempo, dedicación e inestimables consejos, y a las Matemáticas, por su incommensurable aportación a mi vida.

Referencia: HERRERA MACHADO, Erika Magdalena. «Topología y neurociencia: la percepción del espacio». En: *TEMat*, 5 (2021), págs. 43-55. ISSN: 2530-9633. URL: <https://temat.es/articulo/2021-p43>.

^aEl actual trabajo se desarrolló como parte de la asignatura Trabajo de Fin de Grado del Grado en Matemáticas, Universidad de La Laguna.

1. Introducción

Sin haberse percatado, cuando el lector se ha movido por el ambiente en el que se encuentra leyendo este artículo, su cerebro, en un proceso natural, ha formado una representación de tal experiencia, dotándolo de contexto espacial para los recuerdos y experiencias pasadas. En este trabajo se pretende estudiar mediante técnicas matemáticas la manera en la que la actividad del cerebro representa la información espacial del entorno. Las referencias esenciales han sido los artículos de Curto [3] y Curto *et al.* [4].

Lo que acabamos de comentar es solo un ejemplo de las muchas aplicaciones de la neurociencia. Esta ciencia multidisciplinar estudia la estructura, el desarrollo y el funcionamiento del sistema nervioso, centrándose en el cerebro y su impacto en el comportamiento y funciones cognitivas del individuo.

Los primeros estudios se atribuyen a los egipcios, autores del texto médico más antiguo que conoce la historia, denominado «papiro Edwin Smith». Este data de 1700 a. C. y en él se analizan el cerebro, las meninges, la médula espinal y el líquido cefalorraquídeo, aunque tanto los egipcios como el griego Aristóteles adoptaron la creencia de que nuestra conciencia, imaginación y memoria estaban enraizadas en el corazón. Fue aproximadamente en el año 170 a. C. cuando el trabajo del médico griego Galeno cuestionó esta opinión.

Importantes descubrimientos en esta ciencia han acontecido hasta nuestros días, destacando nombres como Descartes, Galvani, Golgi, Ramón y Cajal o Rabi, entre otros, por sus aportaciones. Sin embargo, no es hasta el siglo XX cuando la neurociencia pasa a ser reconocida como una disciplina académica en sí misma, experimentando su más rápido progreso a mitad de siglo. A ello ha contribuido en gran medida la introducción de técnicas matemáticas y, en particular, topológicas.

Para poner al lector en contexto, imagínese explorando un entorno. A medida que avanza, una *neurona de posicionamiento* se activa cuando se sitúa en una zona del espacio bien delimitada, llamada *campo de posicionamiento*. De esta manera, el entorno completo puede quedar determinado por la actividad de un conjunto de estas neuronas. Considerando los campos de posicionamiento como una colección de abiertos en un espacio topológico, es posible definir un código neuronal a partir de dicha colección. Es más, ciertos estudios revelan que los campos de posicionamiento son esencialmente convexos, lo cual motiva que en este trabajo pongamos énfasis en examinar los llamados *códigos neuronales convexos*.

Para estudiar satisfactoriamente lo mencionado anteriormente, serán claves los complejos simpliciales, utilizados históricamente en la triangulación de espacios topológicos y que no son otra cosa que conjuntos de estructuras geométricas elementales: segmentos lineales, triángulos, tetraedros y sus análogos en mayor dimensión. Sin embargo, para nuestras aplicaciones en neurociencia utilizaremos su versión puramente combinatoria: el complejo simplicial abstracto.

2. Topología en neurociencia

Nuestro cerebro cuenta con más de 86 000 millones de neuronas, células eléctricamente activas conectadas entre sí mediante complejas redes. Para transmitir información, una neurona se comunica con su vecina mediante impulsos eléctricos, que se generan como respuesta a estímulos o de manera espontánea. La neurociencia estudia el sistema nervioso, su estructura y funcionamiento, centrándose en el cerebro y su impacto en el comportamiento y las funciones cognitivas.

La pregunta a responder sería: ¿cómo la actividad conjunta de las neuronas representa la información del mundo exterior?

Los primeros descubrimientos relevantes en esta temática tuvieron lugar en la década de los 50, cuando Hodgkin y Huxley estudiaron la dinámica de los impulsos eléctricos para una neurona aislada [10]. En ese momento, incluso sin tener en cuenta la red neuronal que la rodeaba, ya parecía una quimera predecir cuándo una neurona iba a activarse.

Sin embargo, Hubel y Wiesel, ambos premio Nobel en 1981 por sus aportaciones en el estudio del área visual de la corteza cerebral, realizaron en los años 60 un experimento consistente en mostrarle a un gato despierto patrones de barras negras sobre un fondo blanco. Estudiando algunas neuronas por separado, encontraron que cada una respondía a un ángulo diferente de dicha barra, como se muestra en la figura 1, donde la frecuencia de impulsos eléctricos de una neurona se dispara para un ángulo de 45° .

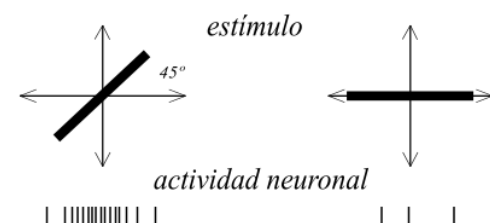


Figura 1: Respuesta neuronal a estímulo visual.

Hubel y Wiesel habían descubierto las *neuronas de orientación* [6], cuya actividad, por tanto, se podría predecir mirando únicamente los estímulos de la pantalla. Se pensó entonces que cada neurona actuaba como un sensor para una característica particular de la escena visual.

Una década más tarde tiene lugar un descubrimiento similar, llevado a cabo por el neurocientífico John O'Keefe, premio Nobel de Medicina en 2014 por sus descubrimientos de células asociadas al posicionamiento espacial del individuo. Mediante experimentos consistentes en estudiar un roedor moviéndose por un entorno acotado, O'Keefe se percató de que ciertas neuronas del *hipocampo* respondían de forma selectiva a diferentes localizaciones de su entorno físico [9]. Estas células, que sirven como sensores de posición en el espacio, se conocen como *neuronas de posicionamiento*. Profundizaremos sobre este tema en la siguiente sección.

Nota 1. El hipocampo es un pequeño órgano situado en lo que se conoce como el *sistema límbico* (conjunto de estructuras interconectadas cuya función se relaciona con los estados emocionales). Se le asocian procesos ligados a la memoria, las emociones y la navegación espacial. ◀

Sistema Límbico

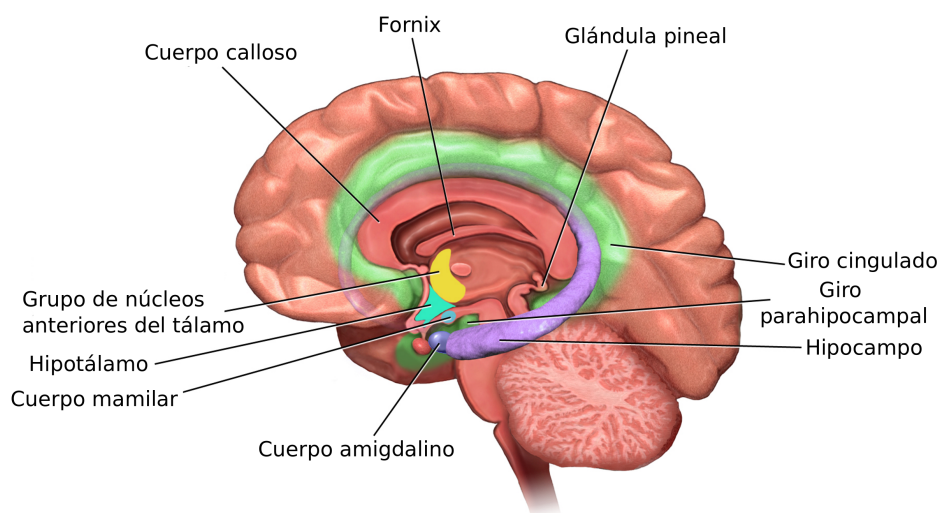


Figura 2: Sistema límbico. Imagen modificada a partir de un original de Blausen.com [1] con licencia © CC BY 3.0 Unported.

3. Neuronas de posicionamiento

Como hemos mencionado previamente, existe un tipo de células cerebrales conocidas como neuronas de posicionamiento cuyo ratio de impulsos eléctricos aumenta cuando el individuo se sitúa en una zona de su entorno físico denominada *campo de posicionamiento* asociado a la neurona. En nuestro estudio consideraremos un solo campo de posicionamiento asociado a cada neurona, aunque en entornos más grandes que los reproducidos en un laboratorio, cada neurona puede llegar a tener varios de estos campos asociados. Los campos de posicionamiento describen circuitos con importantes implicaciones para la memoria, dado que proveen el contexto espacial para los recuerdos y experiencias pasadas.

Durante un largo tiempo solo se podía monitorizar una neurona. Sin embargo, cuando la monitorización simultánea de varias células fue posible, se demostró mediante inferencia estadística que la posición del individuo podía deducirse de la actividad colectiva de las células de posicionamiento [2].

La figura 3 esquematiza el experimento que pasamos a detallar.

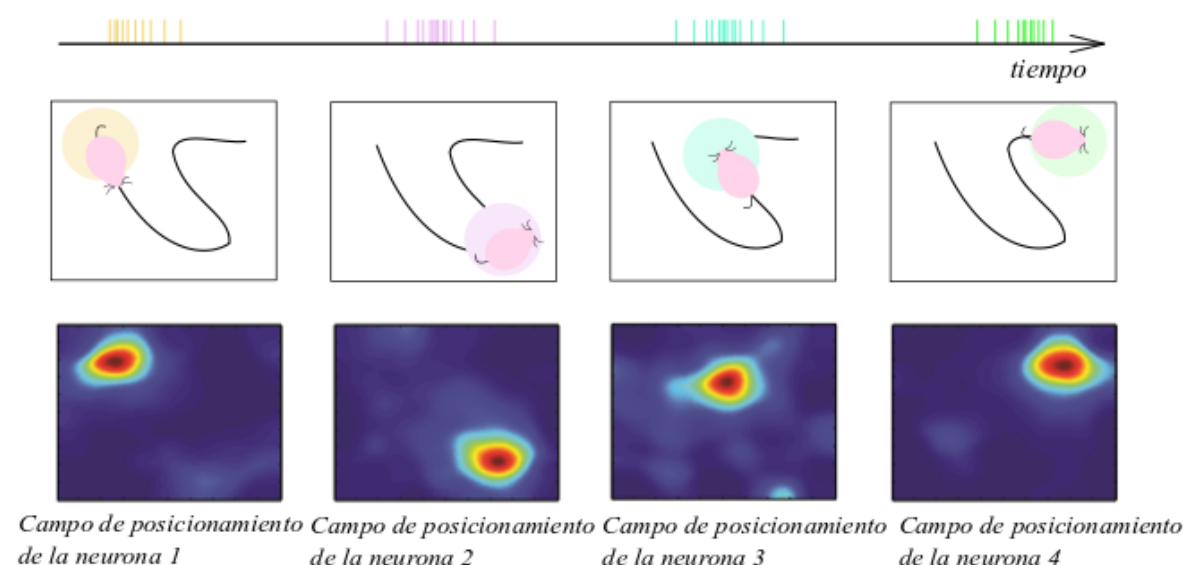


Figura 3: Mapa de calor. Republicado con autorización de la American Mathematical Society a partir de un trabajo de Curto [3, figure 2]; permiso obtenido a través de Copyright Clearance Center, Inc. Figura no sujeta a licencia Creative Commons.

Supongamos que un ratón comienza a explorar su entorno desde la esquina superior izquierda. En ese momento se activa la neurona de posicionamiento 1, cuya actividad se refleja como en la barra superior de la imagen. A medida que el animal avanza se activará el resto de neuronas monitorizadas. La parte inferior de la imagen representa un mapa de calor de la concentración de actividad neuronal en cada campo de posicionamiento. Las áreas rojas reflejan un alto ratio de impulsos eléctricos, mientras que las azules denotan actividad nula (véase el artículo de Giusti *et al.* [5]).

Los estudios sobre las neuronas de posicionamiento se abordan fundamentalmente desde dos perspectivas. Por un lado, la *teoría de la red neuronal* se centra en entender cómo la actividad de las neuronas depende de las propiedades de la red. Sin embargo, en este trabajo nos centraremos en la *teoría del código neuronal*, que presta atención a las relaciones entre la actividad neuronal y los estímulos externos.

Como nos muestra la figura 3, una vez calculados los campos de posicionamiento de las neuronas de posicionamiento del animal, basta ver cuál de estas células se activa para conocer su posición en su entorno físico. Por consiguiente, la idea principal de la investigación que se aborda desde esta perspectiva es que las neuronas de posicionamiento disparan impulsos eléctricos en diferentes lugares del ambiente de tal forma que el entorno entero se representa por la actividad del conjunto de neuronas.

Partiendo de los conceptos de neuronas de posicionamiento y campos de posicionamiento, veremos que se puede inferir información sobre la topología del espacio donde se mueve el ratón.

4. Código neuronal

Los campos de posicionamiento pueden interpretarse como un recubrimiento abierto del entorno físico considerado, donde cada abierto corresponde al lugar donde se activa una neurona.

Ejemplo 2. Consideremos cuatro neuronas monitorizadas cuyos campos de posicionamiento son los de la figura 4. $U_1 \cap U_2 \cap U_3 = U_2$ representa el lugar donde se activan la neurona 1, 2 y 3 a la vez, hecho que se puede representar con la cadena 1110, conocida como palabra de un código neuronal (colección de cadenas formadas por ceros y unos). En este caso, las palabras están formadas por cuatro dígitos, dado que estamos considerando cuatro neuronas, donde el 1 significa que la neurona está activa y el 0 que está en estado de reposo.

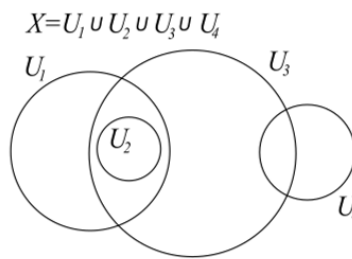


Figura 4: Campos de posicionamiento de cuatro neuronas.

Para simplificar en ciertas ocasiones la notación, es interesante considerar una representación alternativa para las palabras de un código $\mathcal{C}$. Así, definimos el **soporte** de una palabra $c = a_1 \cdots a_n \in \mathcal{C}$ como

$$\text{supp}(c) = \{i \in \{1, \dots, n\} \mid a_i = 1\} \subset \{1, \dots, n\},$$

donde n representa el número de neuronas monitorizadas.

Para la palabra 1110 tenemos que $\text{supp}(1110) = \{1, 2, 3\}$. Obsérvese que $\text{supp}(00 \cdots 0) = \emptyset$.

Por tanto, para referirnos a las palabras de un código neuronal, podemos utilizar indistintamente cadenas de unos y ceros o subconjuntos de $\{1, \dots, n\}$.

Sin embargo, para formar un código neuronal a partir de una colección de abiertos, no consideraremos todas las posibles palabras, sino aquellas que representen una intersección de abiertos no cubierta por el resto, como muestra la siguiente definición.

Definición 3. Sea X un espacio topológico y consideremos una colección de abiertos $\mathcal{U} = \{U_1, \dots, U_n\}$ de X . El **código de $\mathcal{U}$** se define por

$$\mathcal{C}(\mathcal{U}) = \left\{ \sigma \subset \{1, \dots, n\} \mid \bigcap_{i \in \sigma} U_i - \left(\bigcup_{j \in \{1, \dots, n\} - \sigma} U_j \right) \neq \emptyset \right\}.$$

Nota 4. Para facilitar la lectura, utilizaremos la notación $U_\sigma = \bigcap_{i \in \sigma} U_i$.

Cada palabra $\sigma \in \mathcal{C}(\mathcal{U})$ corresponde a una intersección de abiertos no cubierta por el resto. Nótese que $U_\emptyset = \bigcap_{i \in \emptyset} U_i = X$, por lo que $\emptyset$ estará en el código cuando la colección $\mathcal{U}$ no recubra el espacio en su totalidad. Por otro lado, como el número de neuronas a estudiar es finito, las colecciones de abiertos que consideraremos también lo serán.

Observemos en la figura 5 que la palabra 0100 no está incluida en el código. Esto es debido a que el abierto U_2 está contenido en U_1 y U_3 , y esta palabra no estaría reflejando que en esa porción del espacio también se activan las neuronas 1 y 3.

En definitiva, las intersecciones cubiertas por uniones de otros abiertos no reflejan toda la información sobre la relación de activación entre las neuronas.

Por otro lado, podríamos pensar en un código, no necesariamente asociado a un recubrimiento, de la siguiente manera.

$$\mathcal{C}(\mathcal{U})$$

$\{1\}$	$\{1,3\}$	1000	1010
$\{3\}$	$\{3,4\}$	0010	0011
$\{4\}$	$\{1,2,3\}$	0001	1110

Figura 5: Las dos representaciones del código de la colección de abiertos de la figura 4.

Definición 5. Se conoce como **código neuronal**, o simplemente **código**, a toda colección de cadenas binarias $\mathcal{C} \subset \{0, 1\}^n$. Cada elemento de $\mathcal{C}$ se denomina **palabra**. ◀

Hemos visto que toda colección de abiertos de un espacio topológico define un código neuronal. Es natural plantearse si todo código neuronal puede ser asociado a una colección de abiertos en algún espacio real d -dimensional. El siguiente resultado da respuesta a esta pregunta.

Proposición 6. Sea $\mathcal{C} \subset \{0, 1\}^n$ un código neuronal. Para todo $d \geq 1$ existirá una colección $\mathcal{U}$ de abiertos en $\mathbb{R}^d$ tal que $\mathcal{C} = \mathcal{C}(\mathcal{U})$.

Demostración. En primer lugar, ordenamos los elementos de $\mathcal{C}$ como $\{c_1, \dots, c_m\}$. Como $\mathbb{R}^d$ es un espacio de Hausdorff¹, para cada $c_k \in \mathcal{C}$ existe un punto distinto $x_k \in \mathbb{R}^d$ y un entorno abierto N_k de x_k , de forma que ningún par de conjuntos N_k se intersequen. Definimos, para cada $j \in \{1, \dots, n\}$,

$$U_j = \bigcup_{k \in A_j} N_k, \text{ donde } A_j = \{k \in \{1, \dots, m\} \mid j \in \text{supp}(c_k)\},$$

y tomamos $\mathcal{U} = \{U_1, \dots, U_n\}$. Por construcción, $\mathcal{C} = \mathcal{C}(\mathcal{U})$. ■

Visualicemos la anterior demostración con un ejemplo en $\mathbb{R}$.

Ejemplo 7. Tomamos el código $\mathcal{C} = \{000, 010, 110, 101\} = \{c_1, c_2, c_3, c_4\}$, asociado a tres neuronas. Elegimos el abierto $N_i = (i - 3/2, i - 1/2)$ para la palabra c_i , como en la figura 6, con $i = 1, \dots, 4$.

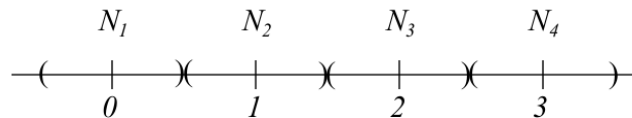


Figura 6: Abiertos en la recta real.

Tenemos que $\text{supp}(c_1) = \emptyset$, $\text{supp}(c_2) = \{2\}$, $\text{supp}(c_3) = \{1, 2\}$ y $\text{supp}(c_4) = \{1, 3\}$, luego $\mathcal{U} = \{U_1, U_2, U_3\}$, con $U_1 = N_3 \cup N_4$, $U_2 = N_2 \cup N_3$ y $U_3 = N_4$, como se muestra en la figura 7.

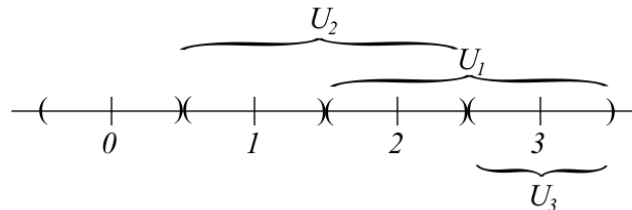


Figura 7: Colección de abiertos resultado de uniones de los abiertos anteriores.

¹Un espacio de Hausdorff es un espacio topológico en el que puntos distintos tienen entornos disjuntos.

Calculamos su código:

$$\begin{array}{ll}
 \sigma = \{1\} & \rightarrow U_1 - (U_2 \cup U_3) = \emptyset, \\
 \sigma = \{2\} & \rightarrow U_2 - (U_1 \cup U_3) \neq \emptyset, \\
 \sigma = \{3\} & \rightarrow U_3 - (U_1 \cup U_2) = \emptyset, \\
 \sigma = \{1, 2\} & \rightarrow U_{\{1,2\}} - U_3 \neq \emptyset, \\
 \sigma = \{1, 3\} & \rightarrow U_{\{1,3\}} - U_2 \neq \emptyset, \\
 \sigma = \{2, 3\} & \rightarrow U_{\{2,3\}} - U_1 = \emptyset, \\
 \sigma = \{1, 2, 3\} & \rightarrow U_{\{1,2,3\}} = \emptyset, \\
 \sigma = \emptyset & \rightarrow U_{\emptyset} - U_{\{1,2,3\}} \neq \emptyset.
 \end{array}$$

Y obtenemos que $\mathcal{C}(\mathcal{U}) = \{\emptyset, \{2\}, \{1, 2\}, \{1, 3\}\} = \{000, 010, 110, 101\} = \mathcal{C}$.

Cabe puntualizar que hemos obtenido una colección de abiertos en $\mathbb{R}$, pero puede ser generalizado para dimensiones superiores. ◀

5. Complejos simpliciales asociados a un código

Comenzaremos la sección introduciendo las nociones de complejo simplicial y complejo simplicial abstracto, y veremos cómo asociarlas a los códigos neuronales (se recomienda consultar el libro de Munkres [8]). Consideraremos únicamente complejos simpliciales finitos, dado que, como comentamos previamente, el número de neuronas monitorizadas es siempre finito.

Los complejos simpliciales finitos son espacios topológicos construidos utilizando estructuras más sencillas, los *simplices*.

Será necesario introducir previamente la idea de conjunto de puntos geoméricamente independiente.

Definición 8. Diremos que un conjunto de puntos $\{a_0, \dots, a_n\}$ de $\mathbb{R}^N$ es **geoméricamente independiente** si los vectores $\{\overrightarrow{a_0 a_1}, \dots, \overrightarrow{a_0 a_n}\}$ son linealmente independientes. ◀

Así, se define la noción de *simplex* como sigue.

Definición 9. Sea $\{a_0, \dots, a_n\}$ un conjunto de puntos geoméricamente independiente de $\mathbb{R}^N$. Un **n -simplex** σ generado por $a_0, \dots, a_n$, que denotaremos por $\sigma = \langle a_0, \dots, a_n \rangle$, consiste en el conjunto de todos los puntos x de $\mathbb{R}^N$ de la forma

$$x = \sum_{i=0}^n t_i a_i, \quad \text{donde } \sum_{i=0}^n t_i = 1 \text{ y } t_i \geq 0 \text{ para todo } i. \quad \blacktriangleleft$$

Como podemos apreciar en la figura 8, un 0-simplex es un punto; un 1-simplex, un segmento; un 2-simplex, un triángulo, y un 3-simplex, un tetraedro en $\mathbb{R}^3$.

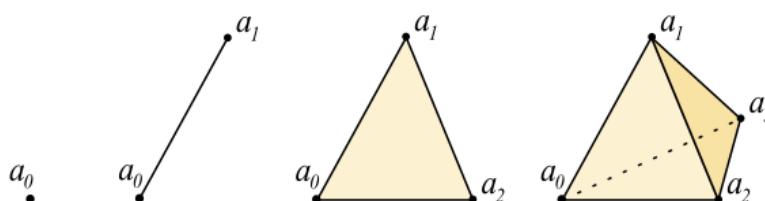


Figura 8: Simplexes con hasta cuatro vértices.

Los $n + 1$ puntos $a_0, \dots, a_n$ que generan el *simplex* σ se denominan *vértices*, y n es la *dimensión* de σ . Cualquier *simplex* generado por un subconjunto de $\{a_0, \dots, a_n\}$ se denomina *cara* de σ .

Definición 10. Llamaremos **complejo simplicial** K en $\mathbb{R}^N$ a una colección finita de símlices en $\mathbb{R}^N$ tal que:

1. Cada cara de un símplex de K es también un símplex de K .
2. La intersección de cualquier par de símlices de K da como resultado una cara de ambos o el conjunto vacío.

Ejemplo 11. Para la figura 9 tenemos el complejo simplicial K formado por el 2-símplex $\langle a_0, a_1, a_2 \rangle$, el 1-símplex $\langle a_2, a_3 \rangle$, el 2-símplex $\langle a_3, a_4, a_5 \rangle$ y todas sus caras.

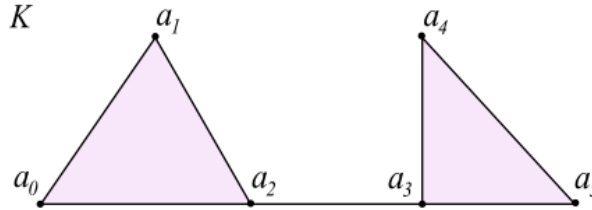


Figura 9: Complejo simplicial K .

Definición 12. Se llama **espacio subyacente** o **poliedro** de K , y se denota por $|K|$, a la unión de los símlices de K como subconjunto de $\mathbb{R}^N$, con N suficientemente grande, dotado de la topología inducida por la usual en $\mathbb{R}^N$.

El concepto de complejo simplicial abstracto que introduciremos a continuación será el principal nexo de unión entre las matemáticas y los estudios en el campo de la neurociencia que tratamos en este trabajo.

Definición 13. Llamamos **complejo simplicial abstracto** a una colección $\mathcal{S}$ de conjuntos finitos no vacíos verificando que, si A es un elemento de $\mathcal{S}$, entonces también lo es cada subconjunto no vacío de A .

Cada elemento A de $\mathcal{S}$ se llamará símplex de $\mathcal{S}$, y su dimensión es su cardinal menos uno. Cada subconjunto no vacío de A se denominará cara de A . La dimensión de $\mathcal{S}$ es la mayor de las dimensiones de sus símlices, y el conjunto de vértices V de $\mathcal{S}$ es la unión de todos los 0-símlices de $\mathcal{S}$.

Definición 14. Sea K un complejo simplicial y $V = \{a_0, \dots, a_n\}$ su conjunto de vértices. Se llama **esquema de vértices** de K a la colección $\mathcal{K}$ de todos los subconjuntos $\{a_{i_0}, \dots, a_{i_k}\}$ de V tales que los vértices $a_{i_0}, \dots, a_{i_k}$ generan un símplex de K .

Definición 15. Llamamos **realización geométrica** de un complejo simplicial abstracto $\mathcal{S}$ a un complejo simplicial K satisfaciendo que existe una biyección entre el esquema de vértices de K y $\mathcal{S}$.

A continuación, veamos cómo asociar un complejo simplicial abstracto a un código neuronal.

Definición 16. Sea $\mathcal{C} \subset \{0, 1\}^n$ un código neuronal. Se llama **complejo simplicial del código** $\mathcal{C}$, denotado por $\Delta(\mathcal{C})$, al complejo simplicial abstracto más pequeño con conjunto de vértices $\{1, \dots, n\}$ que contiene al soporte de todas las palabras de $\mathcal{C}$. Las palabras correspondiendo a las caras de $\Delta(\mathcal{C})$ de dimensión máxima se dirán **maximales**.

Por la definición de complejo simplicial abstracto, los subconjuntos de los soportes de todas las palabras de $\mathcal{C}$ deben pertenecer a $\Delta(\mathcal{C})$. Por tanto,

$$\Delta(\mathcal{C}) = \{\sigma \subset \{1, \dots, n\} \mid \sigma \neq \emptyset \text{ y } \sigma \subset \text{supp}(c) \text{ para algún } c \in \mathcal{C}\}.$$

Nota 17. Obsérvese que de ahora en adelante se denotarán de la misma manera los complejos simpliciales abstractos y sus realizaciones geométricas. Con el fin de visualizar mejor lo que se expondrá en las siguientes secciones, el lector puede pensar los complejos simpliciales abstractos como sus realizaciones geométricas.

Ejemplo 18. A continuación, retomando el código de la figura 5, veamos una realización geométrica del complejo simplicial abstracto asociado al mismo.

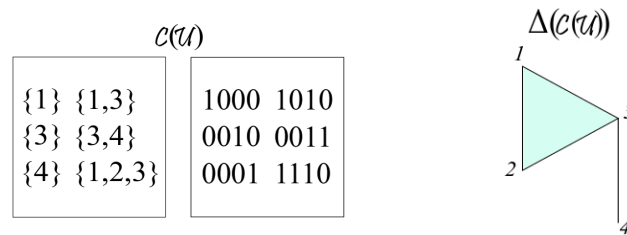


Figura 10: Código y realización geométrica del complejo simplicial abstracto que genera.

De la palabra 1110, cuyo soporte es $\{1, 2, 3\}$, se obtiene el 2-símplice de vértices 1, 2 y 3, y al tratarse del menor complejo simplicial abstracto que lo contiene, consideramos también todas sus caras. Esto explica que en $\Delta(\mathcal{C}(\mathcal{U}))$ consideremos el símplice $\{1, 2\}$ aunque la palabra 1100 no esté en el código. ◀

6. Nervio y lema del nervio

Comenzaremos esta sección introduciendo la noción de *nervio de un recubrimiento abierto*, otra herramienta para modelizar la información espacial que codifican las neuronas de posicionamiento. Seguidamente, veremos la relación entre el nervio de un recubrimiento finito $\mathcal{U}$ y el complejo simplicial del código del recubrimiento. El resultado principal es el *lema del nervio*, que relaciona la estructura del nervio de un recubrimiento finito con la estructura topológica del espacio recubierto, mediante lo que definiremos como equivalencia homotópica.

Definición 19. Llamamos **nervio de un recubrimiento finito de abiertos** $\mathcal{U} = \{U_1, \dots, U_n\}$ de un espacio topológico X al conjunto

$$\mathcal{N}(\mathcal{U}) = \left\{ \sigma \subset \{1, \dots, n\} \mid \sigma \neq \emptyset \text{ y } \bigcap_{i \in \sigma} U_i \neq \emptyset \right\}.$$

El nervio de un recubrimiento finito es un complejo simplicial abstracto, cuyos símlices están formados por los índices de los abiertos con intersección no vacía.

Ejemplo 20. Si retomamos el ejemplo que se ha ido presentando a lo largo del trabajo, se tiene que una realización geométrica del nervio del recubrimiento finito que mostramos de nuevo es la que se aprecia en la figura 11.

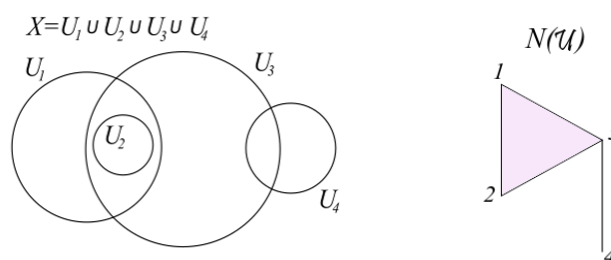


Figura 11: Recubrimiento finito de abiertos de un espacio topológico X y nervio asociado.

Una intersección no vacía de dos abiertos, como $U_3 \cap U_4$, se traduce en el símplice $\{3, 4\}$, mientras que una intersección no vacía de tres abiertos, como $U_1 \cap U_2 \cap U_3$ genera el 2-símplice $\{1, 2, 3\}$. U_2 y U_4 no tienen intersección, por lo que no hay ningún símplice de vértices $\{2, 4\}$. ◀

El lector, al observar las figuras 10 y 11, se habrá podido percatar de que el nervio $\mathcal{N}(\mathcal{U})$ del recubrimiento y el complejo simplicial $\Delta(\mathcal{C}(\mathcal{U}))$ asociado al código coinciden. En efecto, esto es cierto en general.

Proposición 21. Sea $\mathcal{U}$ un recubrimiento finito de un espacio topológico. Entonces, $\mathcal{N}(\mathcal{U}) = \Delta(\mathcal{C}(\mathcal{U}))$.

Demostración. Por definición, está claro que toda palabra del código pertenece al nervio, por lo que deducimos que $\mathcal{C}(\mathcal{U}) \subset \mathcal{N}(\mathcal{U})$. Entonces, $\Delta(\mathcal{C}(\mathcal{U})) \subset \mathcal{N}(\mathcal{U})$, ya que vimos que el nervio es un complejo simplicial abstracto y $\Delta(\mathcal{C}(\mathcal{U}))$ se define como el complejo simplicial abstracto más pequeño que contiene al código.

Por otro lado, toda palabra $\sigma \in \mathcal{N}(\mathcal{U})$ tiene una palabra maximal $\omega \in \mathcal{N}(\mathcal{U})$ tal que $\sigma \subset \omega$. Entonces, $\bigcap_{i \in \omega} U_i \neq \emptyset$ y, para todo $\tau \not\subset \omega$, $\bigcap_{j \in \tau \cup \omega} U_j = \emptyset$ (en otro caso, ω no sería maximal). Por tanto, deducimos que $\bigcap_{i \in \omega} U_i - \bigcup_{j \notin \omega} U_j = \bigcap_{i \in \omega} U_i \neq \emptyset$, es decir, $\omega \in \mathcal{C}(\mathcal{U})$. Así, concluimos $\sigma \in \Delta(\mathcal{C}(\mathcal{U}))$. ■

Por otra parte, cabe añadir que el nervio, bajo determinadas condiciones, refleja la topología del espacio. Veamos previamente un ejemplo que ilustra esta situación.

Ejemplo 22. La figura 12 representa dos recubrimientos finitos diferentes de un anillo. En el primero tenemos tres abiertos cuyas intersecciones son contráctiles. En el segundo caso, tenemos un recubrimiento formado por dos abiertos cuya intersección no es contráctil. Asimismo, en el primer caso, el nervio del recubrimiento refleja la topología del espacio recubierto, lo que no ocurre con el nervio del segundo recubrimiento del anillo. ◀

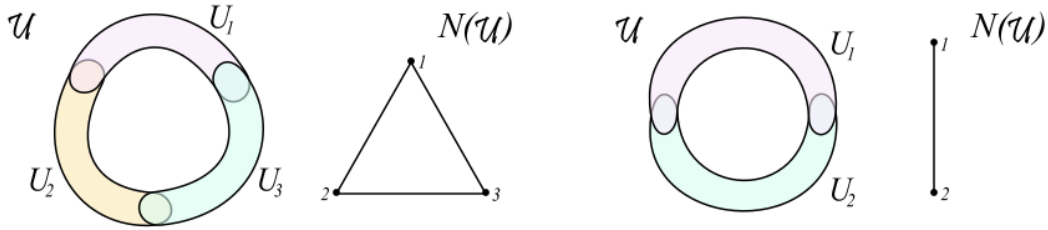


Figura 12: Buenos y malos recubrimientos.

Definición 23. Para un espacio topológico X , un recubrimiento finito abierto $\mathcal{U} = \{U_1, \dots, U_n\}$ es un **buen recubrimiento** si toda intersección no vacía de abiertos de $\mathcal{U}$ es contráctil. ◀

Nota 24. Intuitivamente, un espacio contráctil es aquel que puede ser reducido de forma continua a un punto dentro de ese mismo espacio. Tras estudiar los conceptos que siguen, conoceremos la definición matemática de esta idea, que se recoge en la definición 27. ◀

El siguiente concepto topológico nos permite introducir cuándo dos espacios se pueden deformar continuamente uno en el otro, en cuyo caso diremos que los espacios son *homotópicamente equivalentes*.

Definición 25. Dos aplicaciones continuas $f, g: X \rightarrow Y$ se dicen **homótopas** si existe una aplicación continua $F: X \times [0, 1] \rightarrow Y$ tal que $F(x, 0) = f(x)$ y $F(x, 1) = g(x)$, para todo $x \in X$. Se denota por $F: f \simeq g$, y la aplicación F se denomina **homotopía** entre f y g . ◀

Definición 26. Se dice que dos espacios X e Y son **homotópicamente equivalentes**, y lo denotaremos $X \simeq Y$, si existen aplicaciones continuas $f: X \rightarrow Y$ y $g: Y \rightarrow X$ tales que $g \circ f \simeq \text{id}_X: X \rightarrow X$ y $f \circ g \simeq \text{id}_Y: Y \rightarrow Y$. ◀

Definición 27. Un espacio topológico X se dice **contráctil** si la aplicación identidad sobre X es homótopa a una aplicación constante. ◀

A continuación, presentaremos el resultado que muestra lo que ya anunciábamos previamente: para buenos recubrimientos finitos, el nervio proporciona información de la topología del espacio.

Teorema 28 (lema del nervio). Sea $\mathcal{U}$ un buen recubrimiento finito de un espacio topológico X . Entonces, el poliedro $|\mathcal{N}(\mathcal{U})|$ asociado al nervio de $\mathcal{U}$ es homotópicamente equivalente a X .

Se puede consultar la demostración de este resultado en el libro de Kozlov [7].

El lema del nervio establece, en otras palabras, que el nervio codifica la percepción del espacio. Para comprender por qué nos resulta realmente práctico este resultado, supongamos que hemos monitorizado las neuronas de posicionamiento de un roedor que se mueve por un espacio X cuya forma no conocemos. Según el lema del nervio, no necesitamos mirar directamente al espacio X para saber su forma (por ejemplo, para ver si tiene agujeros), pues el poliedro del nervio (o, equivalentemente, el poliedro del complejo simplicial del código) del recubrimiento finito que definen las neuronas con sus campos de posicionamiento (siempre y cuando sea un buen recubrimiento) ya proporciona esa información.

$$\begin{array}{ccc} \mathcal{U} & \searrow & \mathcal{C}(\mathcal{U}) \\ \downarrow & & \downarrow \\ \mathcal{N}(\mathcal{U}) & = & \Delta(\mathcal{C}(\mathcal{U})) \end{array}$$

Por otro lado, se nos podría plantear lo siguiente: si los campos de posicionamiento de nuestras neuronas monitorizadas definen un buen recubrimiento $\mathcal{U}$ del espacio y tenemos un camino directo para estudiar la forma topológica del lugar (que es considerar directamente $\mathcal{N}(\mathcal{U})$ y aplicar el lema del nervio), ¿por qué resulta de interés estudiar el código $\mathcal{C}(\mathcal{U})$? La respuesta se apoya en que determinadas investigaciones demuestran que el código captura información que el nervio pierde, como, por ejemplo, la relación de contenido entre los abiertos del recubrimiento [4]. Veamos un ejemplo.

Ejemplo 29. La figura 13a muestra una colección $\mathcal{U} = \{U_1, U_2, U_3, U_4\}$ de conjuntos convexos abiertos en $\mathbb{R}^N$.

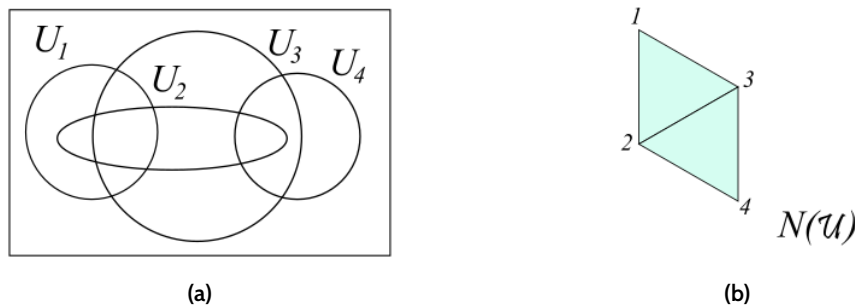


Figura 13: Colección de conjuntos convexos abiertos en $\mathbb{R}^N$ y realización geométrica del nervio de la misma.

Se tiene que

$$\mathcal{C} = \mathcal{C}(\mathcal{U}) = \{\emptyset, \{1\}, \{3\}, \{4\}, \{1, 2\}, \{1, 3\}, \{2, 3\}, \{3, 4\}, \{1, 2, 3\}, \{2, 3, 4\}\}.$$

Y, además,

$$\mathcal{N}(\mathcal{U}) = \{\{1\}, \{2\}, \{3\}, \{4\}, \{1, 2\}, \{1, 3\}, \{2, 3\}, \{2, 4\}, \{3, 4\}, \{1, 2, 3\}, \{2, 3, 4\}\}.$$

$\{1, 2, 3\}$ y $\{2, 3, 4\}$ corresponden a las palabras maximales. El nervio (realización geométrica en la figura 13b), como comentábamos, pierde información acerca de las relaciones de contenidos de los abiertos U_i de $\mathcal{U}$, pero esta información sí la recoge el código. Por ejemplo, el hecho de que $U_2 \subset U_1 \cup U_3$ se refleja en $\mathcal{C}$, dado que cualquier palabra que contenga un 2 contiene también un 1 o un 3. ◀

7. Códigos convexos

Se ha observado experimentalmente que los campos de posicionamiento son esencialmente convexos [5] (véase la figura 3), entendiéndolos como campos de posicionamiento en $\mathbb{R}^N$. Es decir, que dado cualquier par de puntos del conjunto, el segmento que los une está incluido en él. Esto motiva que pongamos énfasis

en estudiar los llamados códigos convexos, que no son otra cosa que códigos asociados a una colección de abiertos convexos. En caso de que esta familia de abiertos determine un recubrimiento, está claro que sería un buen recubrimiento, dado que toda intersección de convexos es convexa y, por tanto, contráctil.

Para facilitar la labor de determinar si un código es o no convexo, utilizaremos como herramienta las relaciones, que se definen en el artículo de Curto *et al.* [4] como sigue.

Definición 30. Dado $\mathcal{C}(\mathcal{U})$, se define una **relación** como el par (σ, τ) , con $\sigma, \tau \subset \{1, \dots, n\}$, tal que $\bigcap_{i \in \sigma} U_i \subset \bigcup_{j \in \tau} U_j$, donde $\sigma \neq \emptyset$, $\sigma \cap \tau = \emptyset$ y $\bigcap_{i \in \sigma \cup \{j\}} U_i \neq \emptyset$ para todo $j \in \tau$. ◀

Ejemplo 31. Para nuestro recubrimiento, el conjunto de relaciones es

$$\{(\{1, 4\}, \emptyset), (\{1, 2, 4\}, \emptyset), (\{1, 3, 4\}, \emptyset), (\{2, 3, 4\}, \emptyset), (\{1, 2, 3, 4\}, \emptyset), \\ (\{2\}, \{1\}), (\{2\}, \{3\}), (\{2\}, \{1, 3\}), (\{1, 2\}, \{3\}), (\{2, 3\}, \{1\})\}.$$

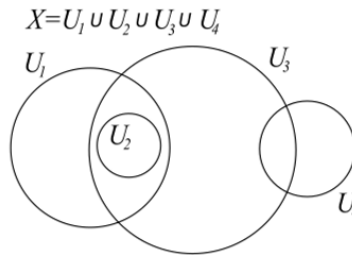


Figura 14: Recubrimiento finito de abiertos de un espacio topológico X .

A modo de ejemplo, el par $(\{1, 2\}, \{3\})$ es una relación dado que $U_1 \cap U_2 = U_2$ está contenido en U_3 y $U_1 \cap U_2 \cap U_3 \neq \emptyset$. ◀

Las obstrucciones locales a la convexidad son la condición que impide a un código ser convexo. Para caracterizarlas, usamos los pares definidos anteriormente y razonamos de la siguiente manera.

Supongamos que partimos de un recubrimiento finito donde todos los abiertos U_j son convexos. Si (σ, τ) es una relación asociada al recubrimiento, por el lema del nervio, el nervio asociado al recubrimiento de $\bigcap_{i \in \sigma} U_i$, que se denota $\mathcal{N}(\{\bigcap_{i \in \sigma \cup \{j\}} U_i\}_{j \in \tau})$, debe tener el mismo tipo de homotopía que $\bigcap_{i \in \sigma} U_i$, luego debe ser contráctil. Así, si $\mathcal{N}(\{\bigcap_{i \in \sigma \cup \{j\}} U_i\}_{j \in \tau})$ resulta no ser contráctil, concluimos que algunos U_j no son convexos.

Esto da pie a definir el concepto de obstrucción local a la convexidad de un código como sigue.

Definición 32. Una relación (σ, τ) es una obstrucción local de $\mathcal{C}$ si $\tau \neq \emptyset$ y $\mathcal{N}(\{\bigcap_{i \in \sigma \cup \{j\}} U_i\}_{j \in \tau})$ no es contráctil. ◀

Por tanto, si $\mathcal{C}$ tiene obstrucciones locales, entonces $\mathcal{C}$ no es un código convexo.

Proposición 33. Si para cada relación (σ, τ) se tiene que $\bigcap_{i \in \sigma \cup \tau} U_i \neq \emptyset$, entonces $\mathcal{C}$ no tiene obstrucciones locales.

Demostración. La condición $\bigcap_{i \in \sigma \cup \tau} U_i \neq \emptyset$ implica que $\mathcal{N}(\{\bigcap_{i \in \sigma \cup \{j\}} U_i\}_{j \in \tau})$ es el símplex completo con conjunto de vértices τ , que es contráctil. Si esto se cumple para toda relación del código, entonces ninguna puede dar lugar a una obstrucción local. ■

Cabe comentar que la ausencia de obstrucciones locales a la convexidad del código no garantiza que el código sea convexo. Sin embargo, conocer de antemano que no se pueden encontrar obstrucciones locales a la convexidad es de utilidad para evitar buscar este tipo de obstrucciones.

8. Conclusión

En la actualidad, la experimentación en neurociencia está presenciando un periodo de rápido progreso y expansión, generando un gran volumen de nuevos datos que pueden comprenderse mejor con la introducción de técnicas y herramientas matemáticas. En particular, a lo largo de la última década, se han estado implementando cada vez más los métodos topológicos.

Con este trabajo se ha pretendido mostrar que la topología puede desempeñar un papel relevante en el desarrollo de la neurociencia.

Hemos comprobado que, a medida que un individuo camina a través de un entorno, la actividad de sus neuronas define los campos de posicionamiento. Estos, que pueden entenderse como un recubrimiento finito $\mathcal{U}$ de un espacio topológico X , además, generan un código neuronal $\mathcal{C}(\mathcal{U})$. A su vez, hemos visto que a todo código $\mathcal{C}$ se le puede asociar una colección de abiertos en $\mathbb{R}^d$, con $d \geq 1$, de forma que $\mathcal{C} = \mathcal{C}(\mathcal{U})$. A partir de estos códigos, se puede definir un complejo simplicial abstracto, $\Delta(\mathcal{C}(\mathcal{U}))$, que coincidirá con el nervio asociado al recubrimiento finito $\mathcal{U}$, $\mathcal{N}(\mathcal{U})$. El resultado principal del artículo, el lema del nervio, afirma que si $\mathcal{U}$ constituye lo que llamamos un buen recubrimiento finito de X , la forma del espacio se puede conocer estudiando las propiedades del elemento matemático $\mathcal{N}(\mathcal{U})$. Además de ello, hemos prestado atención a estudiar si los códigos son códigos convexos, afirmación que se puede descartar si se identifica la presencia de obstrucciones locales a la convexidad.

Referencias

- [1] BLAUSEN.COM. «Medical gallery of Blausen Medical 2014». En: *WikiJournal of Medicine* 1.2 (2014). ISSN: 2002-4436. <https://doi.org/10.15347/WJM/2014.010>.
- [2] BROWN, Emery N.; FRANK, Loren M.; TANG, Dengda; QUIRK, Michael C., y WILSON, Matthew A. «A Statistical Paradigm for Neural Spike Train Decoding Applied to Position Prediction from Ensemble Firing Patterns of Rat Hippocampal Place Cells». En: *Journal of Neuroscience* 18.18 (1998), págs. 7411-7425. ISSN: 0270-6474. <https://doi.org/10.1523/JNEUROSCI.18-18-07411.1998>.
- [3] CURTO, Carina. «What can topology tell us about the neural code?» En: *American Mathematical Society. Bulletin. New Series* 54.1 (2017), págs. 63-78. ISSN: 0273-0979. <https://doi.org/10.1090/bull/1554>.
- [4] CURTO, Carina; GROSS, Elizabeth; JEFFRIES, Jack; MORRISON, Katherine; OMAR, Mohamed; ROSEN, Zvi; SHIU, Anne, y YOUNGS, Nora. «What makes a neural code convex?» En: *SIAM Journal on Applied Algebra and Geometry* 1.1 (2017), págs. 222-238. <https://doi.org/10.1137/16M1073170>.
- [5] GIUSTI, Chad; PASTALKOVA, Eva; CURTO, Carina, y ITSKOV, Vladimir. «Clique topology reveals intrinsic geometric structure in neural correlations». En: *Proceedings of the National Academy of Sciences of the United States of America* 112.44 (2015), págs. 13455-13460. ISSN: 0027-8424. <https://doi.org/10.1073/pnas.1506407112>.
- [6] HUBEL, David H. y WIESEL, Torsten N. «Receptive fields of single neurones in the cat's striate cortex». En: *The Journal of Physiology* 148.3 (1959), págs. 574-591. <https://doi.org/10.1113/jphysiol.1959.sp006308>.
- [7] KOZLOV, Dmitry. *Combinatorial algebraic topology*. Vol. 21. Algorithms and Computation in Mathematics. Berlín: Springer, 2008. <https://doi.org/10.1007/978-3-540-71962-5>.
- [8] MUNKRES, James R. *Elements of algebraic topology*. Menlo Park: Addison-Wesley Publishing Company, 1984. ISBN: 978-0-201-04586-4.
- [9] O'KEEFE, John y DOSTROVSKY, Jonathan. «The hippocampus as a spatial map. Preliminary evidence from unit activity in the freely-moving rat». En: *Brain Research* 34.1 (1971), págs. 171-175. ISSN: 0006-8993. [https://doi.org/10.1016/0006-8993\(71\)90358-1](https://doi.org/10.1016/0006-8993(71)90358-1).
- [10] RINZEL, John. «Discussion: Electrical excitability of cells, theory and experiment: Review of the Hodgkin-Huxley foundation and an update». En: *Bulletin of Mathematical Biology* 52.1 (1990), págs. 5-23. ISSN: 0092-8240. [https://doi.org/10.1016/S0092-8240\(05\)80003-5](https://doi.org/10.1016/S0092-8240(05)80003-5).