



Facultad de Ciencias

TRABAJO FIN DE GRADO
Grado en Matemáticas

El teorema de Borsuk-Ulam y algunas aplicaciones

Autora: Irene Rodríguez Hernández

Tutores: Adrián Fidalgo Díaz y Beatriz Molina Samper

Curso 2024-2025

Índice

Resumen	3
Introducción	4
1. Preliminares	6
1.1. Grupo fundamental y revestimientos	6
1.2. Levantado de caminos y homotopías	8
1.3. Símplices	10
2. El Teorema de Borsuk-Ulam. Versiones equivalentes	15
2.1. Las formulaciones clásicas	17
2.2. El Teorema de Lusternik-Schnirelmann	21
2.3. Lema $N + 1$ de Fan	22
2.4. Lema de Tucker	28
3. Demostración del Teorema de Borsuk-Ulam	30
3.1. Caso $n = 1$	30
3.2. Caso $n = 2$	30
3.3. Caso general	33
4. Consecuencias del Teorema de Borsuk-Ulam	42
4.1. Teorema del punto fijo de Brouwer	42
4.2. Lema de Sperner	48
4.3. Juego de Hex	52
4.3.1. Historia y funcionamiento del juego	52
4.3.2. Hex implica Brouwer	58
4.3.3. Brouwer implica Hex	59
4.3.4. El juego de Hex n -dimensional	62
4.4. Problema del collar	63
Referencias	67

Resumen

El Teorema de Borsuk-Ulam establece que, dada una aplicación continua de la n -esfera sobre $\mathbb{R}^n$, existen dos puntos antipodales para los que la aplicación toma el mismo valor.

En este trabajo se demostrará la equivalencia entre el Teorema de Borsuk-Ulam y varios enunciados, entre ellos el Teorema de Lusternik-Schnirelmann, el Lema $N + 1$ de Fan y el Lema de Tucker. A continuación, se procederá a la demostración en el caso $n = 1$, $n = 2$ y el caso general.

Finalmente, se explorarán algunas de sus aplicaciones: el Teorema del punto fijo de Brouwer, el Lema de Sperner, el Teorema de Hex y el problema del collar.

Introducción

En todo momento existen dos puntos antipodales sobre la superficie terrestre que comparten la misma presión y temperatura. Este curioso hecho es una consecuencia del Teorema de Borsuk-Ulam: dada una aplicación continua de la n -esfera en el espacio euclídeo n -dimensional, existen dos puntos antipodales para los que la aplicación toma el mismo valor. El teorema recibe su nombre de los matemáticos polacos Karol Borsuk y Stanisław Ulam: este último lo propuso como conjetura, y fue Borsuk quien lo demostró por primera vez en 1933.

Aparentemente sencillo, este teorema adquiere una mayor relevancia al observar los múltiples campos en los que podemos encontrar aplicaciones: el Teorema del punto fijo de Brouwer y el Teorema del sándwich de jamón tal vez sean sus consecuencias más conocidas, pero sus aplicaciones van más allá de la topología: aparece en la geometría (lema de Radon, [Gui]), en el procesamiento de imágenes digitales [Pet16], en teoría de grafos (número cromático de grafos de Kneser, [Lov78]), o en programación no lineal [Kaw25], entre otros.

Parte del motivo por el que se pueden encontrar aplicaciones tan extensas está, sin duda, relacionado a la extensa red de posibles formulaciones equivalentes que permiten adaptarlo a contextos topológicos, combinatorios o conjuntistas. En este trabajo exploraremos algunas de ellas. Por supuesto, tal cantidad de enunciados equivalentes da lugar a todo un abanico de formas de abordar la demostración. Se puede elegir demostrar el lema de Tucker, el Teorema de Lusternik-Schnirelmann [LS35], o se puede probar una de las formulaciones clásicas mediante grupos de homología o mediante herramientas topológicas. Surgen también diversas generalizaciones del teorema.

Como ya hemos dicho, el Teorema de Borsuk-Ulam puede enunciarse mediante diversas formulaciones equivalentes, y es en esto en lo que se centrará el capítulo 2. Comenzaremos por las formulaciones clásicas, cinco enunciados que tratan directamente con aplicaciones antipodales y continuas. A continuación, se presenta el Teorema de Lusternik-Schnirelmann: un enunciado que trata con recubrimientos de la esfera por conjuntos cerrados, y otro análogo que considera recubrimientos por abiertos. Los dos últimos, conocidos como el Lema $N + 1$ de Fan y el Lema de Tucker, son enunciados de carácter combinatorio relacionados a triangulaciones y etiquetados.

Tras probar las equivalencias descritas, procederemos a la demostración del teorema en el capítulo 3. La prueba es relativamente sencilla para los casos $n = 1$ y $n = 2$, a partir de conocimientos básicos sobre el grupo fundamental. El caso general se trata mediante una demostración geométrica que utiliza herramientas estudiadas en la asignatura de Topología Algebraica.

La última sección del trabajo se dedica, como promete el título, a la exploración de algunas de las múltiples aplicaciones del teorema. El Teorema del punto fijo de Brouwer, que establece que toda aplicación continua de una bola cerrada en sí misma tiene un punto fijo, se puede probar directamente mediante el Teo-

rema de Borsuk-Ulam, y existe además una construcción directa en la que los puntos antipodales que tienen la misma imagen son los que nos dan el punto fijo. Daremos, además, una versión más débil del teorema a la que es equivalente. Le sigue el Lema de Sperner, un resultado combinatorio relacionado al etiquetado de símlices, en cuya demostración utilizaremos el Teorema de Borsuk-Ulam. A continuación, hablaremos del juego de Hex, un juego de estrategia para dos jugadores que se turnan seleccionando casillas hexagonales del tablero. El Teorema de Hex afirma que el juego de Hex tiene siempre un ganador, un hecho equivalente al Teorema del punto fijo de Brouwer. Finalmente se trata el problema del collar, que consiste en el reparto equitativo de las joyas de un collar entre dos ladrones matemáticos, realizando la menor cantidad posible de cortes. Este reparto siempre se puede hacer con, como mucho, tantos cortes como tipos de joya haya a repartir, un hecho que, de nuevo, se prueba mediante el Teorema de Borsuk-Ulam.

1. Preliminares

En este primer apartado introduciremos algunas definiciones con las que vamos a trabajar, así como algunos resultados vistos en el grado que serán necesarios en las secciones posteriores.

En primer lugar, vamos a fijar algunas notaciones que emplearemos a lo largo de la memoria.

- Denotaremos por $\mathbb{S}^n$ a la esfera $\mathbb{S}^n = \{\mathbf{x} \in \mathbb{R}^{n+1} : \|\mathbf{x}\|_2 = 1\}$, y por E^n al conjunto $\{\mathbf{x} \in \mathbb{R}^n : \|\mathbf{x}\|_2 \leq 1\}$. Denotaremos $\hat{\mathbb{S}}^n$ a la esfera en norma 1 y $\mathbb{S}_\infty^n$ a la esfera en norma infinito:

$$\hat{\mathbb{S}}^n = \{\mathbf{x} \in \mathbb{R}^{n+1} : \sum_{i=1}^{n+1} |x_i| = 1\},$$

$$\mathbb{S}_\infty^n = \{\mathbf{x} \in \mathbb{R}^{n+1} : \max_{i=1}^{n+1} |x_i| = 1\}.$$

- El conjunto $\{\mathbf{x} \in \mathbb{R}^{n+1} : \|\mathbf{x}\|_2 = 1, x_{n+1} \geq 0\}$ es el casquete superior de $\mathbb{S}^n$, al que llamaremos U . Llamaremos L al casquete inferior definido de forma análoga.
- En todo momento denotaremos al intervalo $[0, 1]$ como I .

1.1. Grupo fundamental y revestimientos

Las siguientes definiciones son conocidas, pero merecen ser repasadas dado que las usaremos con frecuencia.

Definición 1. *Un camino en un espacio topológico X es una aplicación continua $\alpha : I \rightarrow X$. Si $\alpha(0) = \alpha(1) = x_0$ para cierto $x_0 \in X$, decimos que α es un lazo en X con punto base x_0 .*

Dados dos caminos α y β tales que $\alpha(1) = \beta(0)$, definimos la operación concatenación de caminos como

$$\alpha\beta(t) = \begin{cases} \alpha(2t) & \text{si } t \in [0, \frac{1}{2}] \\ \beta(2t - 1) & \text{si } t \in [\frac{1}{2}, 1]. \end{cases}$$

Definición 2. *Dados X e Y espacios topológicos y $f, g : X \rightarrow Y$ aplicaciones continuas, decimos que f y g son homótopas y escribimos $f \simeq g$ si existe una aplicación $F : X \times I \rightarrow Y$ continua tal que $F(x, 0) = f(x)$ y $F(x, 1) = g(x)$ para todo $x \in X$. A dicha aplicación F se le llama homotopía entre f y g .*

Además, si existe un subconjunto $A \subseteq X$ tal que $f|_A = g|_A$ y existe una homotopía $F : X \times I \rightarrow Y$ entre f y g tal que $F(a, t) = f(a) = g(a)$ para todo $a \in A$, decimos que f y g son homótopas relativamente a A , escrito $f \simeq g \text{ rel } A$.

En particular, si $X = [0, 1]$ y $A = \{0, 1\}$, diremos que se trata de una homotopía de caminos.

Ser homótopos, ser homótopos relativamente a un conjunto y ser caminos homótopos son relaciones de equivalencia.

Definición 3. El grupo fundamental de X relativo al punto base x_0 , escrito $\pi_1(X, x_0)$, es el conjunto de las clases de homotopía relativa a $\{0, 1\}$ de los lazos en X con punto base x_0 . Es un grupo con la operación de concatenación de caminos al tomar clases, es decir, $[\alpha][\beta] = [\alpha\beta]$ para $[\alpha], [\beta] \in \pi_1(X, x_0)$ cualesquiera.

Denotaremos por $f : (X, x_0) \rightarrow (Y, y_0)$ a una aplicación $f : X \rightarrow Y$ tal que $f(x_0) = y_0$.

Definición 4. Si $f : (X, x_0) \rightarrow (Y, y_0)$ es una aplicación continua, definimos el homomorfismo inducido por f como la aplicación

$$f_* : \pi_1(X, x_0) \rightarrow \pi_1(Y, y_0)$$

$$f_*([\alpha]) = [f \circ \alpha].$$

En efecto, es un homomorfismo, puesto que

$$f_*([\alpha][\beta]) = [f \circ (\alpha\beta)] = [(f \circ \alpha)(f \circ \beta)] = [(f \circ \alpha)][(f \circ \beta)].$$

Recordemos que $f_* \circ g_* = (f \circ g)_*$, y $id_* = id_{\pi_1(X, x_0)}$.

El ejemplo básico de grupo fundamental no trivial es $\pi_1(\mathbb{S}^1, (1, 0)) \cong \mathbb{Z}$. Para probar que esto es así, pasamos por la definición de revestimiento y el levantamiento único de caminos y homotopías, que recordaremos a continuación.

No entraremos en más detalles sobre el grupo fundamental, puesto que es un concepto básico de la topología algebraica que ya hemos estudiado en el grado. Se puede encontrar más información consultando [Mun07, Capítulo 9].

Definición 5. Dada una aplicación continua y sobreyectiva $p : E \rightarrow X$, decimos que un abierto $V \subset X$ está bien recubierto si $p^{-1}(V) = \bigcup_{i \in I} S_i$, siendo esta una unión disjunta de abiertos, y donde $p|_{S_i}$ establece un homeomorfismo entre S_i y V . A los S_i los llamaremos hojas de V por p . Aunque es un abuso de notación, puesto que $p|_{S_i}$ tiene llegada en X , escribiremos $p|_{S_i} : S_i \rightarrow V$.

Decimos que una aplicación $p : E \rightarrow X$ es una proyección recubridora (o aplicación recubridora) si es continua, sobreyectiva y para todo $x \in X$ existe un entorno de x que sea un abierto bien recubierto por p . Decimos que (E, p) es un revestimiento de X .

Definición 6. Sea $p : (E, e_0) \rightarrow (X, x_0)$ una proyección recubridora, y consideramos una aplicación continua $f : (Y, y_0) \rightarrow (X, x_0)$. Llamamos elevación

(también llamadas *levantado* o *levantamiento*) de f a una aplicación continua $f' : (Y, y_0) \rightarrow (E, e_0)$ tal que $p \circ f' = f$.

$$\begin{array}{ccc} & & (E, e_0) \\ & \nearrow f' & \downarrow p \\ (Y, y_0) & \xrightarrow{f} & (X, x_0) \end{array}$$

1.2. Levantado de caminos y homotopías

Veamos algunos resultados de cara a las homotopías, a las proyecciones recubridoras y a las elevaciones.

Proposición 1. Sean $p : (E, e_0) \rightarrow (X, x_0)$ una proyección recubridora y $f : (Y, y_0) \rightarrow (X, x_0)$ una aplicación continua. Si Y es conexo, entonces las elevaciones de f son únicas.

Demostración. [Hat02, Proposición 1.34] Supongamos que f'_1 y f'_2 son dos elevaciones de la aplicación $f : (Y, y_0) \rightarrow (X, x_0)$. Veamos que el conjunto A de puntos donde coinciden es abierto, cerrado y no vacío, por lo que necesariamente ha de ser el total (dado que Y es conexo).

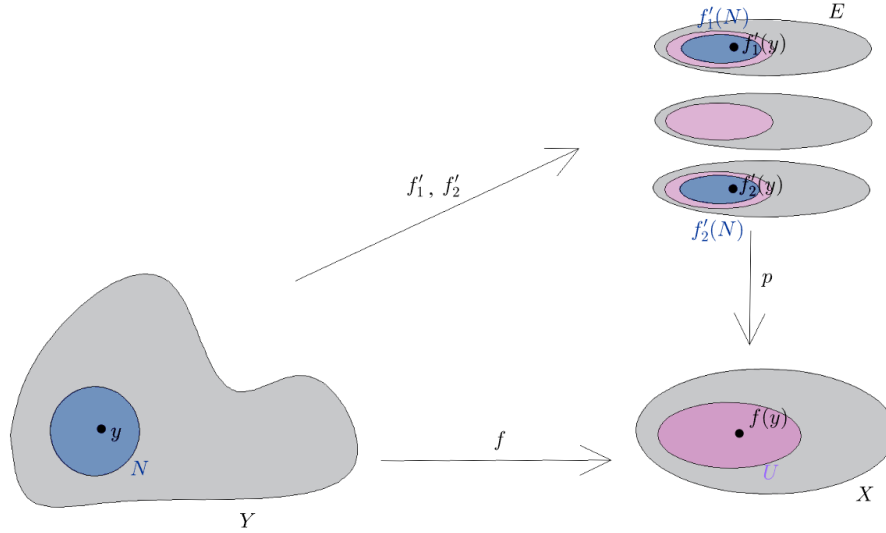


Figura 1: Dos elevaciones distintas.

Sea $y \in Y$, y sea U un abierto bien recubierto de $f(y)$ en X . Sean S_1 y S_2 las hojas de U por p a las que pertenecen $f'_1(y)$ y $f'_2(y)$, respectivamente. S_1 y

S_2 son conjuntos abiertos. Como f'_1 y f'_2 son continuas por definición, existe un entorno N de y tal que $f'_1(N) \subset S_1$ y $f'_2(N) \subset S_2$.

Si $S_1 = S_2$, entonces $p|_{S_1} = p|_{S_2}$, luego

$$f'_1(y) = p|_{S_1}^{-1}(f(y)) = p|_{S_2}^{-1}(f(y)) = f'_2(y).$$

Además, como las hojas son disjuntas, si $S_1 \neq S_2$ se tiene que $f'_1(y) \neq f'_2(y)$. Es decir, $S_1 = S_2$ si y solo si $f'_1(y) = f'_2(y)$. Distinguiamos ambos casos:

Supongamos que $f'_1(y) \neq f'_2(y)$. Hemos visto que, entonces, $S_1 \neq S_2$. Como $f'_1(N) \subset S_1$ y $f'_2(N) \subset S_2$ y como S_1 y S_2 son disjuntas, esto implica que $f'_1(N)$ y $f'_2(N)$ son disjuntos, de modo que $f'_1(z) \neq f'_2(z)$ para todo $z \in N$. Deducimos que el conjunto de puntos donde no coinciden es abierto, luego su complementario, A , es cerrado.

Supongamos ahora que $f'_1(y) = f'_2(y)$. Entonces, $S_1 = S_2$. Como $p \circ f'_1 = p \circ f'_2$ y p es inyectiva en $S_1 = S_2$, entonces $f'_1(z) = f'_2(z)$ para todo $z \in N$, de modo que A es abierto.

Finalmente, como $f'_1(y_0) = f'_2(y_0) = e_0$, se tiene que A no es vacío. $\square$

Proposición 2. *Toda proyección recubridora es abierta.*

Demostración. Sea $p : E \rightarrow X$ una proyección recubridora y sea A un subconjunto abierto de E . Vamos a probar que $p(A)$ es abierto, probando para ello que es entorno de todos sus puntos.

Tomemos un punto $p(e)$ con $e \in A$. Sea V un abierto bien recubierto de $p(e)$ y S la hoja de V a la que pertenece e . $A \cap S$ es abierto de S , y como $p|_S$ es homeomorfismo, $p|_S(A \cap S) = p(A \cap S)$ es abierto de V , y por tanto de X . Como $p(e) \in p(A \cap S) \subset p(A)$, concluimos que $p(A)$ es entorno de $p(e)$, como queríamos probar. $\square$

El siguiente resultado es conocido como el Teorema de elevación de caminos.

Proposición 3. *Si $p : (E, e_0) \rightarrow (X, x_0)$ es una proyección recubridora y α es un camino en X con punto inicial x_0 , existe un único camino α'_{e_0} en E con punto inicial e_0 que sea elevación de α .*

Demostración. Tomamos R un recubrimiento de X por conjuntos abiertos bien recubiertos. Consideramos una partición de I , $0 = t_0 < t_1 < \dots < t_n = 1$, de forma que α lleve cada $[t_i, t_{i+1}]$, $0 \leq i \leq n-1$, a un abierto bien recubierto en R . Sabemos que esto lo podemos hacer por el Lema del número de Lebesgue [Mun07, Lema 27.5], que establece que como I es un espacio métrico compacto existe $\delta > 0$ tal que para cada subconjunto con diámetro menor que δ existe un elemento del recubrimiento que lo contiene.

Definimos $\alpha'_{e_0}(0) = e_0$. Supongamos que $\alpha'_{e_0}(t)$ está definida para $0 < t \leq t_i$. Sea $U \in R$ un abierto bien recubierto conteniendo $\alpha([t_i, t_{i+1}])$. Definimos ahora

α'_{e_0} en el intervalo $[t_i, t_{i+1}]$ como $\alpha'_{e_0}(t) = (p|_S)^{-1}(\alpha(t))$, donde S es la hoja de U por p en la que está $\alpha'_{e_0}(t_i)$. Como $p|_S$ es homeomorfismo, entonces α'_{e_0} es continua en $[t_i, t_{i+1}]$. Continuamos por recurrencia finita hasta que esté definida en todo I .

Con esto hemos probado la existencia y con la Proposición 1 queda probada la unicidad, pues el intervalo $[0, 1]$ es conexo. $\square$

Utilizaremos sin demostrarlo el Teorema de elevación de homotopías, del cual podemos encontrar una demostración en [GH18, Capítulo 5].

Proposición 4. *Sea $p : (E, e_0) \rightarrow (X, x_0)$ una proyección recubridora, y sea $f : (Y, y_0) \rightarrow (X, x_0)$ una aplicación con una elevación $f' : (Y, y_0) \rightarrow (E, e_0)$. Entonces, toda homotopía $F : Y \times I \rightarrow X$ tal que $F(y, 0) = f(y)$ para todo $y \in Y$ admite una elevación a la homotopía $F' : Y \times I \rightarrow E$ con $F'(y, 0) = f'(y)$ para todo $y \in Y$. Además, si $g(y) = F(y, 1) = (p \circ F')(y, 1)$ entonces $g'(y) = F'(y, 1)$ es una elevación de g .*

Proposición 5. *Sea α y γ caminos en X tales que $\alpha \simeq \gamma$ rel $\{0, 1\}$, y sea $p : (E, e_0) \rightarrow (X, x_0)$ una proyección recubridora. Entonces, sus elevaciones α'_{e_0} y γ'_{e_0} son homótopas relativamente a $\{0, 1\}$. En particular, tienen el mismo punto final.*

Demostración. Sea $F : I \times I \rightarrow X$ una homotopía entre α y γ relativa a $\{0, 1\}$, y denotemos $x_0 = \alpha(0) = \gamma(0)$, $x_1 = \alpha(1) = \gamma(1)$. Por el Teorema de elevación de homotopías, F tiene una elevación $F' : I \times I \rightarrow E$ tal que $F'(s, 0) = \alpha'_{e_0}(s)$ para cada $s \in I$. Además, como $\gamma_{e_0}(s) = F(s, 1) = p \circ F'(s, 1)$, el mismo teorema nos dice que $\gamma'_{e_0}(s) = F'(s, 1)$ es una elevación de $\gamma_{e_0}(s)$. Luego F' es una homotopía entre α'_{e_0} y γ'_{e_0} . Falta ver que la homotopía es relativa a $\{0, 1\}$:

Para todo $t \in I$ se tiene que $p \circ F'(0, t) = F(0, t) = x_0$. Por lo tanto, $F'(\{0\} \times I) \subset p^{-1}(x_0)$, que es un subconjunto discreto de E . Como $F'(\{0\} \times I)$ es un conjunto conexo (puesto que es la imagen de un conjunto conexo por una aplicación continua) contenido en un conjunto discreto, ha de ser unipuntual, y como $F'(0, 0) = \alpha'_{e_0}(0) = e_0$, solo puede ser $\{e_0\}$.

Como $p \circ F'(1, t) = F(1, t) = x_1$ para todo t , entonces $F'(\{1\} \times I) \subset p^{-1}(x_1)$. Por el mismo razonamiento que antes, $F'(\{1\} \times I) = F'(1, 0) = \alpha'_{e_0}(1)$. $\square$

1.3. Símplices

Repasaremos varias definiciones y resultados relativos a complejos simpliciales vistas en las asignaturas del grado y en [Mun18, Capítulo 1].

Definición 7. *Un n -símplice $\sigma \subset \mathbb{R}^m$ es la envolvente convexa de los puntos afínmente independientes $\{a_0, \dots, a_n\} \subset \mathbb{R}^m$. Escribimos $\sigma = \langle a_0, \dots, a_n \rangle$, y decimos que n es la dimensión del símplice.*

Los *símplices* engendrados por subconjuntos de $\{a_0, \dots, a_n\}$ se llaman *caras* de σ . Las caras de dimensión 0 se llaman *vértices*. La *frontera* (geométrica) de σ es la unión de las caras propias de σ , es decir, la unión de las caras distintas del propio σ . El *interior* (geométrico) de σ son los puntos de σ que no están en la frontera.

Proposición 6. Si x es un punto de un *símplice* σ , existe una única cara que lo contenga en su interior. Llamamos a esta cara el *soporte* de x .

Definición 8. Un *complejo simplicial* en $\mathbb{R}^m$ es un conjunto finito K de *símplices* de $\mathbb{R}^m$ tales que se cumplen las siguientes propiedades:

1. Cada cara de un *símplice* de K es un *símplice* de K .
2. La intersección de dos *símplices* de K es vacía o es una cara de cada uno de ellos.

La *dimensión* de K ($\dim(K)$) es r si K contiene algún r -*símplice* pero no contiene ningún $(r+1)$ -*símplice*. Un subconjunto L de K que sea un *complejo simplicial* se dice que es un *subcomplejo simplicial* de K .

El *poliedro* $|K|$ de un *complejo simplicial* K es el subespacio de $\mathbb{R}^n$ formado por la unión de los *símplices* de K .

Definición 9. Sea K un *complejo simplicial*. Decimos que un *complejo* K' es una *subdivisión* de K si cada *símplice* de K' está contenido en un *símplice* de K , y cada *símplice* de K es unión finita de *símplices* de K' .

Definición 10. Sea $\sigma = \langle a_0, \dots, a_n \rangle$ un *símplice* en $\mathbb{R}^m$. Su *baricentro* es el punto

$$\hat{\sigma} = \sum_{i=0}^n \frac{1}{n+1} a_i.$$

Definimos la *subdivisión baricéntrica* de K como el *complejo* K' formado por los *símplices* de la forma $\langle \hat{\sigma}_1, \dots, \hat{\sigma}_k \rangle$ donde $\sigma_1, \dots, \sigma_k$ son *símplices* de K tales que σ_i es una cara propia de σ_{i+1} para todo $i \in \{0, \dots, k-1\}$.

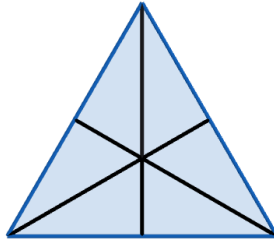


Figura 2: Subdivisión baricéntrica de un 2-símplice.

Definición 11. Una *triangulación* de un espacio topológico X es un *complejo simplicial* K y un *homeomorfismo* $h : |K| \rightarrow X$.

Para referirnos a las triangulaciones, nos permitiremos hacer abuso del lenguaje y referirnos al complejo simplicial en lugar de al homeomorfismo. Veamos algunos ejemplos:

1. Sean $e_1 = (1, 0, 0)$, $e_2 = (0, 1, 0)$ y $e_3 = (0, 0, 1)$. Consideramos el complejo simplicial en $\mathbb{R}^3$ cuyos símlices son aquellos que tienen como vértices uno, dos o tres elementos del conjunto $\{e_1, e_2, e_3, -e_1, -e_2, -e_3\}$ que estén en el mismo octante. Esta es una triangulación de $\hat{\mathbb{S}}^2$. Si en cada símlice hacemos una subdivisión baricéntrica, el nuevo complejo también triangula $\hat{\mathbb{S}}^2$. Observemos que el máximo diámetro de los símlices de la nueva triangulación es menor.

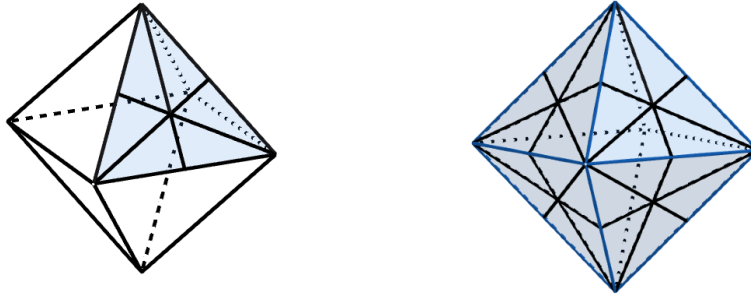


Figura 3: Dividiendo cada 2-símlice como en la figura de la izquierda (que no es un complejo simplicial) obtenemos la triangulación de la derecha.

2. Sea $\sigma = \langle a_0, \dots, a_d \rangle$ un símplex en $\mathbb{R}$. Llamamos a $\sigma \times [0, 1]$ su *prisma simplicial*, y denotamos $v_i = (a_i, 0)$ y $v'_i = (a_i, 1)$, $i \in \{0, \dots, d\}$. El complejo dado por los símlices $\sigma_j = \langle v_0, v_1, \dots, v_j, v'_j, v'_{j+1}, \dots, v'_d \rangle$ y todas sus caras nos da una triangulación del prisma simplicial (Figuras 4 y 5).



Figura 4: Triangulación de un prisma simplicial si $n = 1$.

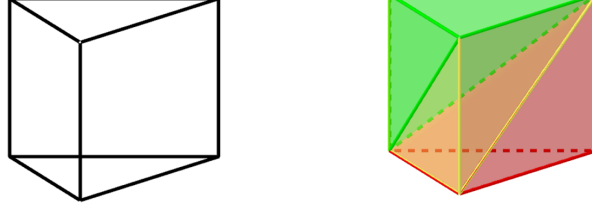


Figura 5: Triangulación de un prisma simplicial si $n = 2$.

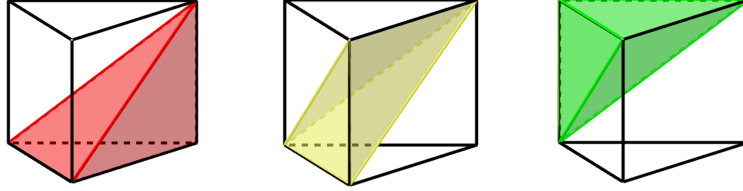


Figura 6: Símplices para el caso $n = 2$.

Proposición 7. [Arm83, Lema 6.4] Sea K un complejo simplicial de dimensión n , y sea K^1 su primera subdivisión baricéntrica. Denotamos por $\mu(K)$ el mayor diámetro de los símplices de K , y lo análogo para K^1 . Entonces,

$$\mu(K^1) \leq \frac{n}{n+1} \mu(K).$$

Demostración. El diámetro de un símplice es el máximo de las longitudes de sus caras de dimensión 1. Sea $\sigma = \langle \hat{\sigma}_1, \hat{\sigma}_2 \rangle \in K^1$, donde $\hat{\sigma}_1$ y $\hat{\sigma}_2$ son los baricentros de dos símplices σ_1 y σ_2 de K , siendo σ_2 una cara de σ_1 . Entonces $\sigma \subset \sigma_1$. Si σ_1 es de dimensión k , se tiene que $\mu(\sigma) \leq \frac{k}{k+1} \mu(\sigma_1) \leq \frac{n}{n+1} \mu(K)$. $\square$

Corolario 1. Sea K una triangulación de dimensión n y $K^1, K^2, \dots, K^m, \dots$ subdivisiones baricéntricas sucesivas. Entonces, $\mu(K^m) \rightarrow 0$ cuando $m \rightarrow \infty$.

Demostración. Por la Proposición 7, el diámetro de la subdivisión baricéntrica de un n -símplice es como mucho $\frac{n}{n+1}$ veces el diámetro del símplice, luego si el diámetro de K es d , entonces el diámetro de K^m es $(\frac{n}{n+1})^m d$, que tiende a 0 cuando m tiende a infinito. $\square$

Proposición 8. Dada una aplicación $f : V(K) \rightarrow X$, donde $X \subset \mathbb{R}^n$ y $V(K)$ son los vértices de un complejo simplicial K , podemos extenderla a una aplicación continua $f' : |K| \rightarrow \mathbb{R}^n$ afín en cada símplice de K .

Demostración. Si $\mathbf{x} \in |K|$, existe $\sigma \in K$ tal que σ es el soporte de $\mathbf{x}$, con vértices $v_0, \dots, v_k$. Entonces, $\mathbf{x} = \sum_{i=0}^k \alpha_i v_i$ para ciertos α_i , con $\alpha_i \geq 0$ y

$$\sum_{i=0}^k \alpha_i = 1.$$

Definimos $f'(\mathbf{x}) = \sum_{i=0}^k \alpha_i f(v_i)$. Observemos que f' coincide con f en los vértices de K y es afín en cada símple. Está bien definida, porque cada punto x tiene un único símple que sea su soporte, y es continua en $|K|$. $\square$

Definición 12. Dada una aplicación $f : V(K) \rightarrow X$, llamaremos a la aplicación f' de la proposición anterior la extensión afín de f a K .

2. El Teorema de Borsuk-Ulam. Versiones equivalentes

El Teorema de Borsuk-Ulam recibe su nombre de los matemáticos polacos Stanisław Ulam y Karol Borsuk. El primero de ellos lo propuso como una conjetura, y fue el segundo quien lo probó en 1933. Enunciaremos primero la formulación original, y probaremos la equivalencia con sus otras versiones. En el siguiente capítulo daremos la demostración del teorema. Durante la mayor parte de esta sección nos basaremos en [Mat03, Capítulo 2.1].

En primer lugar damos dos definiciones muy sencillas que son conceptos elementales para hablar del Teorema de Borsuk-Ulam:

Definición 13. Decimos que dos puntos $x, y \in \mathbb{R}^n$ son antipodales si $x = -y$.

Durante esta memoria, será frecuente restringir esta definición a puntos de $\mathbb{S}^{n-1}$, de $\hat{\mathbb{S}}^{n-1}$ o de $\mathbb{S}_{\infty}^{n-1}$.

Definición 14. Sea $A \subseteq B \subseteq \mathbb{R}^n$ y $X \subseteq \mathbb{R}^m$. Decimos que una aplicación $f : B \rightarrow X$ es antipodal en A si $f(-x) = -f(x)$ para todo $x \in A$.

Decimos que una aplicación $f : B \rightarrow X$ es antipodal si es antipodal en todo B , es decir, si para todo $x \in B$ se tiene que $f(-x) = -f(x)$.

Si $B = \mathbb{S}^n$ y $X = \mathbb{S}^m$, a las aplicaciones antipodales también se las conoce como aplicaciones que conservan antípodas, por motivos obvios: si dos puntos son antipodales, sus imágenes por una de estas aplicaciones también lo son.

Veamos algunos ejemplos de aplicaciones antipodales:

- La aplicación $f : \mathbb{S}^n \rightarrow \mathbb{S}^n$ que lleva cada punto a su antípoda (ver Figura 7).

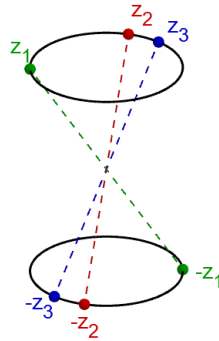


Figura 7: La aplicación $f : \mathbb{S}^1 \rightarrow \mathbb{S}^1$ dada por $f(z) = -z$ es antipodal.

- Cualquier rotación de $\mathbb{S}^n$. Por ejemplo, una rotación en $\mathbb{S}^1$ con un ángulo φ viene dada por $f(x, y) = (x\cos(\varphi) - y\sin(\varphi), x\sin(\varphi) + y\cos(\varphi))$. Luego $f(-x, -y) = (-x\cos(\varphi) + y\sin(\varphi), -x\sin(\varphi) - y\cos(\varphi)) = -f(x, y)$.
- La aplicación $f : \mathbb{S}^1 \rightarrow \mathbb{S}^1$ dada por $f(z) = z^n$, con n impar. Dado $\mathbf{z} = e^{i\theta}$, se tiene que $f(\mathbf{z}) = e^{ni\theta}$, y

$$f(-\mathbf{z}) = f(e^{i(\theta+\pi)}) = e^{ni(\theta+\pi)} = e^{ni\theta} e^{ni\pi} = -e^{ni\theta} = -f(\mathbf{z}).$$

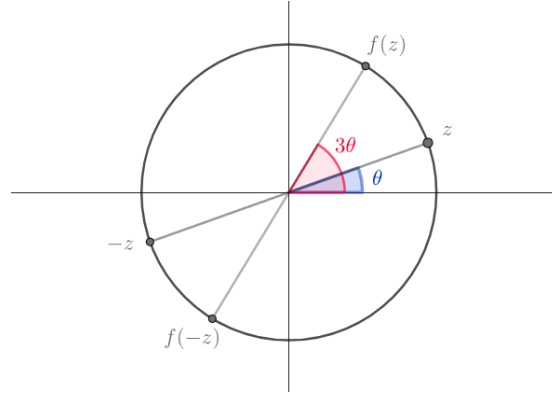


Figura 8: La aplicación $f(z) = z^3$ lleva puntos antipodales a puntos antipodales.

- Sea $w \in \mathbb{R}^{n+1}$. La proyección ortogonal $P : \mathbb{R}^{n+1} \rightarrow H$ de un vector v sobre el hiperplano $H = w^\perp$ puede escribirse como $P(v) = v - \frac{v \cdot w}{w \cdot w} w$, donde $\cdot$ representa el producto escalar. Entonces,

$$P(-v) = -v - \frac{-v \cdot w}{w \cdot w} w = -(v - \frac{v \cdot w}{w \cdot w} w) = -P(v),$$

luego P es antipodal.

Aunque los resultados que enunciaremos a continuación tratan con aplicaciones definidas en $\mathbb{S}^n$ donde se considera la norma 2, también se cumplen para $\hat{\mathbb{S}}^n$ y $\mathbb{S}_\infty^n$, dado que se puede tomar un homeomorfismo entre ellas que mantiene la antipodalidad. Por ejemplo, podemos considerar el homeomorfismo que lleva cada punto $\mathbf{x}$ de $\mathbb{S}^n$ al punto de corte de $\hat{\mathbb{S}}^n$ con la semirrecta que comienza en el origen y pasa por $\mathbf{x}$, representado en la Figura 9. Esto será útil, por ejemplo, a la hora de tomar triangulaciones en algunas demostraciones.

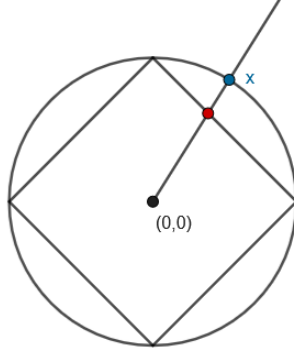


Figura 9: Homeomorfismo entre $\mathbb{S}^n$ y $\hat{\mathbb{S}}^n$.

Enunciamos a continuación el Teorema de Borsuk-Ulam.

Teorema 1. *Para toda aplicación continua $f : \mathbb{S}^n \rightarrow \mathbb{R}^n$ existe un punto $\mathbf{x} \in \mathbb{S}^n$ tal que $f(\mathbf{x}) = f(-\mathbf{x})$.*

Es frecuente mencionar este enunciado acompañado del siguiente, a modo de curiosidad: en todo momento hay dos puntos antipodales de la Tierra que tienen la misma temperatura y presión atmosférica. Es un claro caso particular de lo anterior: si la temperatura y la presión atmosférica son continuas, y se considera la superficie terrestre como una esfera, se puede llevar cada punto de la Tierra a un punto del plano tomando como coordenadas su presión y su temperatura. Aplicando el teorema, habrá dos puntos antipodales que lleguen al mismo punto de este plano.

2.1. Las formulaciones clásicas

A continuación, veremos que los tres siguientes enunciados son equivalentes al Teorema de Borsuk-Ulam:

2. Para toda aplicación antipodal y continua $f : \mathbb{S}^n \rightarrow \mathbb{R}^n$ existe un punto $\mathbf{x} \in \mathbb{S}^n$ tal que $f(\mathbf{x}) = 0$.
3. No existe ninguna aplicación antipodal y continua $f : \mathbb{S}^n \rightarrow \mathbb{S}^{n-1}$.
4. No existe ninguna aplicación continua $f : E^n \rightarrow \mathbb{S}^{n-1}$ que sea antipodal en $\partial E^n = \mathbb{S}^{n-1}$.

Demostración.

(1 $\Rightarrow$ 2) Sea $f : \mathbb{S}^n \rightarrow \mathbb{R}^n$ una aplicación antipodal y continua. Aplicando (1), existe $\mathbf{x} \in \mathbb{S}^n$ tal que $f(\mathbf{x}) = f(-\mathbf{x})$. Además, por ser antipodal, $f(\mathbf{x}) = -f(-\mathbf{x})$. Necesariamente ha de ser $f(\mathbf{x}) = 0$.

(2 $\Rightarrow$ 1) Sea $f : \mathbb{S}^n \rightarrow \mathbb{R}^n$ continua. La aplicación $g : \mathbb{S}^n \rightarrow \mathbb{R}^n$, definida como $g(\mathbf{x}) = f(\mathbf{x}) - f(-\mathbf{x})$ es antipodal, puesto que

$$g(-\mathbf{x}) = f(-\mathbf{x}) - f(\mathbf{x}) = -(f(\mathbf{x}) - f(-\mathbf{x})) = -g(\mathbf{x}),$$

y es continua por serlo f . Por (2), existe $\mathbf{x}_0$ en $\mathbb{S}^n$ tal que $g(\mathbf{x}_0) = 0$, es decir, tal que $f(\mathbf{x}_0) = f(-\mathbf{x}_0)$.

(2 $\Rightarrow$ 3) Supongamos que $f : \mathbb{S}^n \rightarrow \mathbb{S}^{n-1}$ es una aplicación antipodal y continua. Consideramos la inclusión $i : \mathbb{S}^{n-1} \rightarrow \mathbb{R}^n$. La aplicación $i \circ f : \mathbb{S}^n \rightarrow \mathbb{R}^n$ es antipodal. Por (2), existe $\mathbf{x} \in \mathbb{S}^n$ tal que $i \circ f(\mathbf{x}) = 0$, luego $f(\mathbf{x}) = 0$. Pero, como $0 \notin \mathbb{S}^{n-1}$, hemos llegado a un absurdo.

(3 $\Rightarrow$ 4) Supongamos que $f : E^n \rightarrow \mathbb{S}^{n-1}$ es una aplicación continua y antipodal en $\mathbb{S}^{n-1}$. Sea $\pi : \mathbb{R}^{n+1} \rightarrow \mathbb{R}^n$ la proyección en las n primeras coordenadas, es decir, $\pi(x_1, \dots, x_{n+1}) = (x_1, \dots, x_n)$. Observemos que $\pi(U) = E^n$. Consideremos la aplicación $g : \mathbb{S}^n \rightarrow \mathbb{S}^{n-1}$ definida de la siguiente forma:

$$g(\mathbf{x}) = \begin{cases} (f \circ \pi)(\mathbf{x}) & \text{si } \mathbf{x} \text{ está en el casquete superior } U \\ -(f \circ \pi)(-\mathbf{x}) & \text{si } \mathbf{x} \text{ está en el casquete inferior } L \end{cases}$$

Está bien definida, porque en $U \cap L$ se tiene que, por ser f antipodal,

$$(f \circ \pi)(x_1, \dots, x_n, 0) = f(x_1, \dots, x_n) = -f(-x_1, \dots, -x_n) = -(f \circ \pi)(-x_1, \dots, -x_n, 0).$$

Si $\mathbf{x}$ está en U , entonces $-g(\mathbf{x}) = -(f \circ \pi)(\mathbf{x})$. Además, $-\mathbf{x}$ está en L y $g(-\mathbf{x}) = -(f \circ \pi)(\mathbf{x})$, luego $g(-\mathbf{x}) = -g(\mathbf{x})$. Puesto que g es continua y antipodal, se contradice (3).

(4 $\Rightarrow$ 2) Sea $f : \mathbb{S}^n \rightarrow \mathbb{R}^n$ una aplicación antipodal tal que $f(\mathbf{x}) \neq 0$ para todo $\mathbf{x} \in \mathbb{S}^n$. Sea $h : f(\mathbb{S}^n) \rightarrow \mathbb{S}^{n-1}$ que lleva $\mathbf{x}$ a $\frac{\mathbf{x}}{\|\mathbf{x}\|}$ y consideremos la aplicación $p : \mathbb{S}^n \cap \{\mathbf{x} \in \mathbb{R}^{n+1} : x_{n+1} \geq 0\} \rightarrow E^n$ la proyección en las n primeras coordenadas. Definamos $g : E^n \rightarrow \mathbb{S}^{n-1}$ como

$$g(\mathbf{x}) = (h \circ f \circ p^{-1})(\mathbf{x}),$$

donde p^{-1} lleva cada punto al casquete superior de $\mathbb{S}^n$. Sabemos que g es continua por ser la composición de aplicaciones continuas, y como para todo $\mathbf{x} \in \mathbb{S}^{n-1}$ se tiene que

$$\begin{aligned} g(-\mathbf{x}) &= (h \circ f \circ p^{-1})(-\mathbf{x}) = (h \circ f)(-x_1, \dots, -x_n, \sqrt{1 - \|\mathbf{x}\|^2}) \\ &= (h \circ (-f))(x_1, \dots, x_n, -\sqrt{1 - \|\mathbf{x}\|^2}) = (-h \circ f)(x_1, \dots, x_n, -\sqrt{1 - \|\mathbf{x}\|^2}) \\ &= -(h \circ f \circ p^{-1})(\mathbf{x}) = -g(\mathbf{x}). \end{aligned}$$

Luego g es antipodal en $\mathbb{S}^{n-1}$, contradiciendo (4). $\square$

Otro resultado al que el Teorema de Borsuk-Ulam es equivalente es el siguiente [Gra03, Capítulo 5, Teorema 5.2]:

5. Una aplicación antipodal y continua $f : \mathbb{S}^{n-1} \rightarrow \mathbb{S}^{n-1}$ no es homótopa a una aplicación constante.

Antes de demostrarlo, probaremos la siguiente proposición:

Proposición 9. *Sea $f : \mathbb{S}^{n-1} \rightarrow X$ una aplicación continua. Son equivalentes:*

1. *f es homótopa a una aplicación constante.*
2. *Existe una extensión continua de f , $\hat{f} : E^n \rightarrow X$, tal que $\hat{f}|_{\mathbb{S}^{n-1}} = f$.*

Además, si se cumple lo anterior, también es cierto que:

3. *f_* es el homomorfismo trivial, es decir, el que lleva todos los elementos al neutro.*

Si $n = 2$, los tres enunciados son equivalentes.

Demostración.

(1 $\Rightarrow$ 2) Sea $\varphi : \mathbb{S}^{n-1} \times [0, 1] \rightarrow E^n$ la aplicación dada por $\varphi(\mathbf{x}, t) = t\mathbf{x}$. Es una aplicación cociente que relaciona los puntos de la forma $(\mathbf{x}, 0)$ con $\mathbf{x} \in \mathbb{S}^{n-1}$. Además, $\varphi|_{\mathbb{S}^{n-1} \times (0, 1]}$ define un homeomorfismo entre $\mathbb{S}^{n-1} \times (0, 1]$ y $E^n \setminus \{0\}$.

Sea $f : \mathbb{S}^{n-1} \rightarrow X$ una aplicación homótopa a la aplicación constante igual a p , que denotaremos por $c_p : \mathbb{S}^{n-1} \rightarrow X$, y sea $F : \mathbb{S}^{n-1} \times [0, 1] \rightarrow X$ una homotopía de c_p a f . Definimos $\hat{f} : E^n \rightarrow X$ como

$$\hat{f}(\mathbf{y}) = \begin{cases} F(\varphi^{-1}(\mathbf{y})) & \text{si } \mathbf{y} \neq 0 \\ p & \text{si } \mathbf{y} = 0. \end{cases}$$

Tenemos el siguiente diagrama conmutativo:

$$\begin{array}{ccc} \mathbb{S}^{n-1} \times [0, 1] & \xrightarrow{\varphi} & E^n \\ F \downarrow & \nearrow \hat{f} & \\ X & & \end{array}$$

Por la propiedad universal del cociente, $\hat{f}$ es continua si y solo si $F = \hat{f} \circ \varphi$ es continua. Pero F es continua por ser una homotopía, luego también $\hat{f}$ es continua.

Como $\hat{f}|_{\mathbb{S}^{n-1}} = F \circ \varphi^{-1}|_{\mathbb{S}^{n-1}} = F|_{\mathbb{S}^{n-1} \times \{1\}} = f$, se tiene que $\hat{f}$ es una extensión continua de f .

(2 $\Rightarrow$ 1) Sea $\hat{f} : E^n \rightarrow X$ una aplicación continua tal que $\hat{f}|_{\mathbb{S}^{n-1}} = f$. Entonces, $F = \hat{f} \circ \varphi$, con φ como en la implicación anterior, es continua y cumple que

$$F(\mathbf{x}, 0) = (\hat{f} \circ \varphi)(\mathbf{x}, 0) = \hat{f}(\mathbf{0}),$$

$$F(\mathbf{x}, 1) = (\hat{f} \circ \varphi)(\mathbf{x}, 1) = \hat{f}(\mathbf{x}) = f(\mathbf{x}),$$

luego F es una homotopía entre la aplicación constantemente igual a $\hat{f}(\mathbf{0})$ y f .

(2 $\Rightarrow$ 3) Consideramos la inclusión $i : \mathbb{S}^{n-1} \rightarrow E^n$. Entonces, $\hat{f} \circ i = f$, por lo que $\hat{f}_* \circ i_* = f_*$. Pero i_* es el homomorfismo trivial, ya que tiene llegada en $\pi_1(E^n)$, que es el grupo trivial.

(3 $\Rightarrow$ 1 cuando $n = 2$) Sea $p : \mathbb{R} \rightarrow \mathbb{S}^1$ la proyección recubridora dada por $p(x) = (\cos(2\pi x), \sin(2\pi x))$, y consideramos su restricción $p_0 : I \rightarrow \mathbb{S}^1$. $[p_0]$ es un lazo en $\mathbb{S}^1$ que genera $\pi_1(\mathbb{S}^1, x_0)$.

La aplicación $p_0 \times \text{Id} : I \times I \rightarrow \mathbb{S}^1 \times I$ es una aplicación cociente por ser continua, abierta y sobreyectiva (por ser p_0 una proyección recubridora), lleva los puntos $(0, t)$ y $(1, t)$ a los puntos (x_0, t) y es inyectiva en $(0, 1) \times I$.

Por otro lado, como f_* es el homomorfismo trivial, el lazo $f \circ p_0$ es el elemento neutro de $\pi_1(X, f(x_0))$, luego existe una homotopía $F : I \times I \rightarrow X$ relativa a $\{0, 1\}$ entre $f \circ p_0$ y el lazo constantemente igual a $f(x_0)$. Además, F lleva los puntos de la forma $(0, t)$, $(1, t)$ y $(x, 1)$ al $f(x_0)$.

F induce una aplicación continua $H : \mathbb{S}^1 \times I \rightarrow X$ que es una homotopía entre f y una aplicación constante, con $H \circ (p_0 \times \text{Id}) = F$. $\square$

Ahora demostraremos que el enunciado (5) es equivalente al Teorema de Borsuk-Ulam. En concreto, utilizaremos la Proposición 9 para probar que es equivalente al enunciado (4).

Demostración.

(4 $\Rightarrow$ 5) Supongamos que existe una aplicación antipodal $f : \mathbb{S}^{n-1} \rightarrow \mathbb{S}^{n-1}$ homótopa a una aplicación constante. Por la Proposición 9, existe una extensión continua $\hat{f} : E^n \rightarrow \mathbb{S}^{n-1}$. Como $\hat{f}|_{\mathbb{S}^{n-1}} = f$, y f es antipodal, entonces $\hat{f}$ es antipodal en $\mathbb{S}^{n-1}$, contradiciendo (4).

(5 $\Rightarrow$ 4) Supongamos que existe una aplicación continua $f : E^n \rightarrow \mathbb{S}^{n-1}$ antipodal en $\mathbb{S}^{n-1}$. Sea $g : \mathbb{S}^{n-1} \rightarrow \mathbb{S}^{n-1}$ definida como $g = f|_{\mathbb{S}^{n-1}}$. Como f es antipodal en $\mathbb{S}^{n-1}$, entonces g es antipodal. Observemos que f es una extensión continua de g a E^n , luego por la Proposición 9 sabemos que f es homótopa a una aplicación constante. Esto se contradice con (5). $\square$

2.2. El Teorema de Lusternik-Schnirelmann

Los enunciados que se presentan a continuación, conocidos como el Teorema de Lusternik-Schnirelmann, también son equivalentes al Teorema de Borsuk-Ulam:

6. Para todo recubrimiento de $\mathbb{S}^n$ por $n + 1$ conjuntos cerrados $F_1, \dots, F_{n+1}$ existe $i \in \{1, \dots, n + 1\}$ tal que F_i contiene un par de puntos antipodales. Es decir, tal que $F_i \cap (-F_i) \neq \emptyset$.
7. Para todo recubrimiento de $\mathbb{S}^n$ por $n + 1$ conjuntos abiertos $U_1, \dots, U_{n+1}$ existe $i \in \{1, \dots, n + 1\}$ tal que U_i contiene un par de puntos antipodales. Es decir, tal que $U_i \cap (-U_i) \neq \emptyset$.

Demostraremos a continuación que, efectivamente, son equivalentes a los enunciados anteriores.

Demostración.

(1 $\Rightarrow$ 6) Sea $F_1, \dots, F_{n+1}$ un recubrimiento por cerrados de $\mathbb{S}^n$, y considere-mos la aplicación $f : \mathbb{S}^n \rightarrow \mathbb{R}^n$ dada por $f(\mathbf{x}) = (d(\mathbf{x}, F_1), \dots, d(\mathbf{x}, F_n))$, donde $d(\mathbf{x}, F_i) = \inf\{d(x, y) : y \in F_i\}$ denota la distancia de x a F_i . Es una aplicación continua de $\mathbb{S}^n$ en $\mathbb{R}^n$, luego existe $\mathbf{x}_0 \in \mathbb{S}^n$ tal que $f(\mathbf{x}_0) = f(-\mathbf{x}_0)$, es decir,

$$(d(\mathbf{x}_0, F_1), \dots, d(\mathbf{x}_0, F_n)) = (d(-\mathbf{x}_0, F_1), \dots, d(-\mathbf{x}_0, F_n)).$$

Como F_i es un cerrado, sabemos que $x \in F_i$ si y solo si $d(\mathbf{x}, F_i) = 0$. Tenemos dos casos posibles:

1. Si $\mathbf{x}_0 \in F_i$ para algún $1 \leq i \leq n$, entonces $d(\mathbf{x}_0, F_i) = 0 = d(-\mathbf{x}_0, F_i)$, y por tanto $-\mathbf{x}_0 \in F_i$.
2. Si no, necesariamente $\mathbf{x}_0 \in F_{n+1}$. Entonces $d(\mathbf{x}_0, F_i) \neq 0$, $1 \leq i \leq n$, y por tanto $d(-\mathbf{x}_0, F_i) \neq 0$, $1 \leq i \leq n$. Como consecuencia, $-\mathbf{x}_0 \notin F_i$ para ningún $1 \leq i \leq n$. Como $F_1, \dots, F_{n+1}$ es un recubrimiento de $\mathbb{S}^n$, entonces ha de ser $-\mathbf{x}_0 \in F_{n+1}$.

(6 $\Rightarrow$ 7) Sea $U_1, \dots, U_{n+1}$ un recubrimiento por abiertos de $\mathbb{S}^n$. Tomamos para cada $\mathbf{x} \in \mathbb{S}^n$ un entorno abierto $V_{\mathbf{x}}$ tal que $\overline{V_{\mathbf{x}}} \subseteq U_i$ para algún $i \in \{1, \dots, n + 1\}$. Como $\{V_{\mathbf{x}} : \mathbf{x} \in \mathbb{S}^n\}$ es un recubrimiento por abiertos del compacto $\mathbb{S}^n$, entonces existe un subrecubrimiento finito $V_{\mathbf{x}_1}, \dots, V_{\mathbf{x}_k}$. Definamos $F_j = \bigcup_{\overline{V_{\mathbf{x}_i}} \subseteq U_j} \overline{V_{\mathbf{x}_i}}$. Observemos que $F_j \subseteq U_j$ para todo $j \in \{1, \dots, n + 1\}$.

Tenemos ahora un recubrimiento por cerrados $F_1, \dots, F_{n+1}$ tal que $F_i \subseteq U_i$ para todo $i \in \{1, \dots, n + 1\}$. Por (6), existe j tal que F_j contiene una pareja de puntos antipodales, y como $F_j \subseteq U_j$, también U_j los contiene.

(7 $\Rightarrow$ 2) Supongamos por reducción al absurdo que $f : \mathbb{S}^n \rightarrow \mathbb{R}^n$ es una aplicación antipodal tal que $f(\mathbf{y}) \neq 0$ para todo $\mathbf{y} \in \mathbb{S}^n$. Consideremos los abiertos $V_1, \dots, V_{n+1}$ de $\mathbb{R}^n$ dados por

$$V_i = \{\mathbf{x} \in \mathbb{R}^n : x_i > 0\}, \quad i = 1, \dots, n,$$

$$V_{n+1} = \{\mathbf{x} \in \mathbb{R}^n : x_1 + \dots + x_n < 0\}.$$

Se tiene que $\{V_1, V_2, \dots, V_{n+1}\}$ es un recubrimiento por abiertos de $\mathbb{R}^n \setminus \{0\}$, porque o bien al menos una coordenada es mayor que cero o bien la suma de todas es menor que cero (no puede darse el caso de que todas sean cero).

Como f es continua podemos considerar $U_1 = f^{-1}(V_1), \dots, U_{n+1} = f^{-1}(V_{n+1})$, que es recubrimiento por abiertos de $\mathbb{S}^n$ porque $f(\mathbf{y}) \neq 0$ para todo $\mathbf{y} \in \mathbb{S}^n$. Veamos que $\{U_1, U_2, \dots, U_{n+1}\}$ es un recubrimiento por abiertos de $\mathbb{S}^n$ que no cumple (7).

Si $\mathbf{x}$ está en cierto U_i con $i \leq 1 \leq n$, entonces $f(\mathbf{x}) \in V_i$. Luego

$$-f(\mathbf{x}) = f(-\mathbf{x}) \notin V_i,$$

porque cambia el signo de x_i . Por tanto, $-\mathbf{x} \notin U_i$. Por otro lado, si $f(\mathbf{x}) \in V_{n+1}$, entonces $x_1 + \dots + x_n < 0$. Se tiene entonces que $-x_1 - \dots - x_n > 0$, y por tanto $f(-\mathbf{x}) \notin V_{n+1}$.

Concluimos que ningún U_i contiene a la vez dos puntos antipodales, en contradicción con (7), como queríamos probar. $\square$

2.3. Lema $N + 1$ de Fan

Las dos últimas formulaciones del Teorema de Borsuk-Ulam son resultados combinatorios. Es necesario que introduzcamos previamente algunas definiciones. Recordemos que denotamos por $\hat{\mathbb{S}}^n = \{\mathbf{x} \in \mathbb{R}^{n+1} : \sum_{i=1}^{n+1} |x_i| = 1\}$ a la esfera en norma 1.

Definición 15. *Vamos a decir que una triangulación T es simétrica si para todo símple σ de la triangulación se cumple que $-\sigma \in T$.*

Por ejemplo, las sucesivas subdivisiones baricéntricas de una triangulación simétrica son, de nuevo, triangulaciones simétricas.

Definición 16. *Dado $m \in \mathbb{N}$, llamaremos m -etiquetado a una aplicación de la forma*

$$L : V(T) \rightarrow \{1, -1, 2, -2, \dots, m, -m\},$$

donde $V(T)$ son los vértices de la triangulación.

Decimos que el etiquetado de una triangulación simétrica es antipodal si $L(-v) = -L(v)$ para todo $v \in V(T)$.

Un etiquetado tiene un borde complementario si tiene dos vértices adyacentes (es decir, vértices de un 1-símplice de la triangulación) cuyas etiquetas son i y $-i$.

Un símple es alternado positivo si sus etiquetas son

$$\{k_1, -k_2, k_3, \dots, (-1)^{i+1} k_i, \dots, (-1)^{l+1} k_l\},$$

con $1 \leq k_1 < k_2 < \dots < k_l$, y es alternado negativo si sus etiquetas son

$$\{-k_1, k_2, -k_3, \dots, (-1)^i k_i, \dots, (-1)^l k_l\},$$

con $1 \leq k_1 < k_2 < \dots < k_l$.

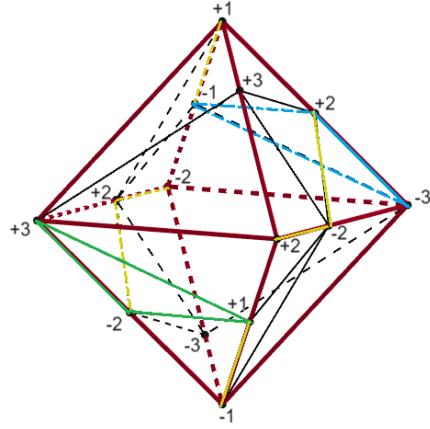


Figura 10: 3-etiquetado antipodal de una triangulación de $\mathbb{S}^2$. En amarillo se indican los bordes complementarios, en azul los símlices alternados negativos y en verde los alternados positivos.

Con estas definiciones, podemos dar otro enunciado equivalente al Teorema de Borsuk-Ulam, conocido como el Lema $N + 1$ de Fan [NS13].

8. Toda triangulación simétrica de $\hat{\mathbb{S}}^n$ con un $(n + 1)$ -etiquetado antipodal tiene o bien un borde complementario o bien un n -símple alternado positivo.

Observemos que, en caso de tener un símple alternado positivo, el hecho de ser un etiquetado antipodal implica que también tiene un símple alternado negativo.

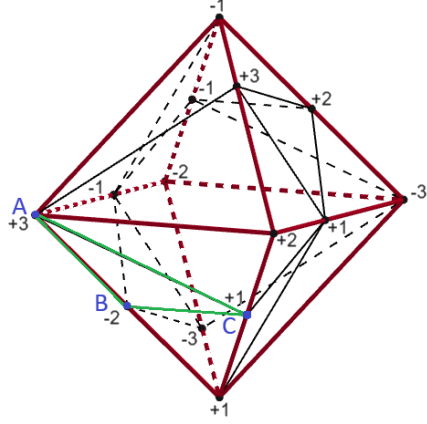


Figura 11: 3-etiquetado antipodal de una triangulación de $\mathbb{S}^2$ sin bordes complementarios.

Para entender informalmente este enunciado, fijémonos en el símplexe alternado positivo de la Figura 11. Si cambiamos la etiqueta del vértice A (y también la de $-A$, para mantener la antipodalidad), obtendremos un borde complementario. Si cambiamos la etiqueta de B a $+2$, formamos un símplexe alternado negativo, y si la cambiamos por cualquier otra tenemos de nuevo un borde complementario. Si cambiamos la etiqueta de C a $+3$, obtenemos un símplexe alternado positivo, y cambiándola por cualquier otro valor se obtiene un borde complementario.

Demostración. ($2 \Rightarrow 8$) Sea T una triangulación simétrica de $\hat{\mathbb{S}}^n$ con un $(n+1)$ -etiquetado antipodal L que no tenga bordes complementarios. Sea w_i el punto de $\mathbb{R}^{n+1}$ con la i -ésima coordenada igual a n y las demás iguales a -1 , esto es, $w_i = (-1, \dots, -1, n, -1, \dots, -1)$. Definimos $w_{-i} = -w_i$ para $1 \leq i \leq n+1$

Sean $W_+ = \{w_1, \dots, w_{n+1}\}$ y $W_- = \{-w_1, \dots, -w_{n+1}\}$. Sea

$$W = W_+ \cup W_-. \quad (1)$$

W está contenido en el hiperplano $H = \{\mathbf{x} \in \mathbb{R}^{n+1} : \sum_{i=1}^{n+1} x_i = 0\}$. Observamos que la aplicación $\pi|_H : H \rightarrow \mathbb{R}^n$ dada por

$$\pi|_H(x_1, \dots, x_n, x_{n+1}) = (x_1, \dots, x_n)$$

es un homeomorfismo antipodal entre H y $\mathbb{R}^n$, y que además $\pi|_H(\mathbf{x}) = \mathbf{0}$ si y solo si $\mathbf{x} = \mathbf{0}$.

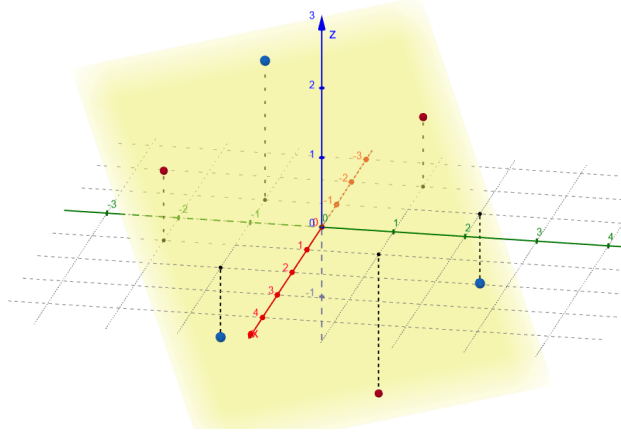


Figura 12: Para $n = 2$: W_+ , en azul; W_- , en rojo; y el hiperplano H .

Sea $\hat{h} : V(T) \rightarrow W$ la aplicación dada por

$$\hat{h}(v) = \begin{cases} w_{L(v)} & \text{si } L(v) \text{ impar,} \\ -w_{L(v)} & \text{si } L(v) \text{ par.} \end{cases}$$

La aplicación $\hat{h}$ es antipodal por ser L antisimétrico: si $L(v)$ impar,

$$\hat{h}(v) = w_{L(v)} = w_{-L(-v)} = -w_{L(-v)} = -\hat{h}(-v),$$

mientras que si es par,

$$\hat{h}(v) = -w_{L(v)} = -w_{-L(-v)} = w_{L(-v)} = -\hat{h}(-v).$$

Sea $h : \hat{\mathbb{S}}^n \rightarrow H$ su extensión afín en cada símplice de T (véase Proposición 8). La aplicación h es continua. Veamos que es antipodal.

Si $\mathbf{x} \in \sigma = \langle v_0, \dots, v_k \rangle$ con coordenadas $\alpha_0, \dots, \alpha_k$, por ser T simétrica se tiene que $-\mathbf{x} \in -\sigma = \langle -v_0, \dots, -v_k \rangle$ con las mismas coordenadas $\alpha_0, \dots, \alpha_k$. Entonces,

$$h(-\mathbf{x}) = \sum_{i=0}^k \alpha_i \hat{h}(-v_i) = \sum_{i=0}^k \alpha_i (-\hat{h}(v_i)) = - \sum_{i=0}^k \alpha_i \hat{h}(v_i) = -h(\mathbf{x}).$$

Entonces, $\pi|_H \circ h : \hat{\mathbb{S}}^n \rightarrow \mathbb{R}^n$ es una aplicación antipodal y continua. Por el Teorema de Borsuk-Ulam, existe $\mathbf{x} \in \hat{\mathbb{S}}^n$ tal que $(\pi|_H \circ h)(\mathbf{x}) = \mathbf{0}$. Como $\pi|_H(\mathbf{x}) = \mathbf{0}$ si y solo si $\mathbf{x} = \mathbf{0}$, entonces $h(\mathbf{x}) = \mathbf{0}$.

Sea σ el soporte de $\mathbf{x}$ en T , y veamos que o bien σ o bien $-\sigma$ es alternado positivo. Sabemos que $\mathbf{0} \in h(\sigma)$. Sea $V = \{\hat{h}(v) : v \in V(\sigma)\}$. V tiene como

mucho $n+1$ elementos y $V \subset W$. Como hemos elegido L de forma que no tenga bordes complementarios, no existen v_1, v_2 vértices de σ tales que

$$L(v_1) = -L(v_2),$$

luego V no contiene ningún conjunto de la forma $\{w_i, -w_i\}$. Podemos entonces escribir V como $V = \{w_j\}_{j \in B_+} \cup \{-w_j\}_{j \in B_-}$ para dos subconjuntos disjuntos B_+, B_- de $\{1, \dots, n+1\}$. Observemos que

$$\{w_j\}_{j \in B_+} \subset W_+ \text{ y } \{-w_j\}_{j \in B_-} \subset W_-.$$

Notemos que, considerando el producto escalar usual,

$$w_i \cdot w_i = (-1)^2 n + n^2 = n(n+1)$$

para todo $1 \leq i \leq n$, y que, si $i \neq j$, entonces

$$w_i \cdot w_j = (-1)^2(n-1) - 2n = -n-1.$$

Sea $v = \sum_{j \in B_+} w_j - \sum_{j \in B_-} w_j$ la suma de los vectores de V . Si $i \in B_+$, entonces

$$w_i \cdot v = w_i \cdot \sum_{j \in B_+} w_j - w_i \cdot \sum_{j \in B_-} w_j =$$

$$= n(n+1) + (|B_+| - 1)(-n-1) - |B_-|(-n-1) = (n+1)(n+1 - |B_+| + |B_-|).$$

Observemos que $w_i \cdot v = 0$ si y solo si $|B_+| = n+1$ y $|B_-| = 0$. Esto ocurre si y solo si $V = W_+$. En caso contrario, $w_i \cdot v > 0$.

Ahora, si $i \in B_-$,

$$-w_i \cdot v = w_i \cdot \sum_{j \in B_+} w_j - w_i \cdot \sum_{j \in B_-} w_j =$$

$$= (n+1)|B_+| + n(n+1) + (|B_-| - 1)(-n-1) = (n+1)(n+1 + |B_+| - |B_-|).$$

Igual que antes, $-w_i \cdot v = 0$ si y solo si $|B_+| = 0$ y $|B_-| = n+1$, es decir, si y solo si $V = W_-$. En cualquier otro caso $-w_i \cdot v > 0$.

Como $\mathbf{0} \in h(\sigma)$, entonces la envolvente convexa de V contiene al origen. Es decir, existen $\alpha_1, \dots, \alpha_m \in \mathbb{R}$ tales que $\alpha_j \geq 0$ para todo j , $\sum_{j=1}^m \alpha_j = 1$ y $\alpha_1 w_{i_1} + \dots + \alpha_m w_{i_m} = \mathbf{0}$, donde $V = \{w_{i_1}, \dots, w_{i_m}\}$. Entonces,

$$0 = \mathbf{0} \cdot v = (\alpha_1 w_{i_1} + \dots + \alpha_m w_{i_m}) \cdot v,$$

por lo que no puede ser que $w \cdot v > 0$ para todo $w \in V$. Por lo anterior, solo es posible que $V = W_+$ o $V = W_-$. Supongamos que es lo primero, es decir,

$$V = \{\hat{h}(v) : v \in V(\sigma)\} = \{w_1, \dots, w_{n+1}\}.$$

Hemos definido $\hat{h}(v)$ como $w_{L(v)}$ o $-w_{L(v)}$ según si $L(v)$ es par o impar, respectivamente, pero en V solo tenemos vectores de la forma w_i , luego todo $L(v)$ par ha de ser negativo, y todo $L(v)$ impar positivo. Luego σ tiene etiquetas $\{1, -2, 3, \dots, (-1)^n(n+1)\}$. Hemos encontrado un símplice alternado positivo.

Supongamos ahora que $V = W_-$. Por el mismo razonamiento, σ tiene etiquetas $\{-1, 2, -3, \dots, (-1)^{n+1}(n+1)\}$. Como el etiquetado L es antipodal, entonces $-\sigma$ es un símplice con etiquetado $\{1, -2, 3, \dots, (-1)^n(n+1)\}$, luego también tenemos en este caso un símplice alternado positivo.

(8 $\Rightarrow$ 2) Supongamos que $h : \hat{S}^n \rightarrow \mathbb{R}^n$ es una aplicación continua y antipodal, $h(\mathbf{x}) = (h_1(\mathbf{x}), \dots, h_n(\mathbf{x}))$, tal que no existe $\mathbf{x} \in \hat{S}^n$ con $h(\mathbf{x}) = \mathbf{0}$. Sea $f : \hat{S}^n \rightarrow \mathbb{R}^{n+1}$ dada por $f(\mathbf{x}) = (h_1(\mathbf{x}), \dots, h_n(\mathbf{x}), -\sum_{i=1}^n h_i(\mathbf{x}))$.

La aplicación f es continua y lleva $\hat{S}^n$ al hiperplano

$$H = \left\{ \mathbf{x} \in \mathbb{R}^{n+1} : \sum_{i=1}^{n+1} x_i = 0 \right\} \subset \mathbb{R}^{n+1}.$$

Además, es antipodal por serlo h . Como h no se anula en ningún punto, tampoco lo hace f .

Sea T una triangulación simétrica de $\hat{S}^n$ y sea W como como en la ecuación (1). Definimos un etiquetado $L : V(T) \rightarrow \{1, -1, 2, -2, \dots, n+1, -(n+1)\}$ como $L(v) = i$, donde i es tal que w_i sea el elemento de W más cercano a $f(v)$ en $\mathbb{R}^{n+1}$. Si hubiera dos a la misma distancia, elegimos el que haga que $L(v)$ tenga el menor valor absoluto. Observemos que esta definición daría problemas si w_i y $-w_i$ fueran los elementos más cercanos a $f(v)$, pero para que eso ocurriese tendría que ocurrir que $f(v) = 0$. Como $f(v) \neq 0$ para todo vértice v , el etiquetado está bien definido.

Como f es antipodal, entonces $f(v)$ está a la misma distancia de cada w_i de la que está $f(-v)$ a $-w_i$, luego el etiquetado L es antipodal. Por el Lema $N+1$ de Fan, o bien existe un símplice de T con etiquetas $\{1, -2, 3, \dots, (-1)^n(n+1)\}$ o bien existe un borde complementario con etiquetado $\{i, -i\}$.

Consideramos ahora subdivisiones baricéntricas sucesivas T^m de T para cada $m \in \mathbb{N}$. Observemos que los diámetros de estas triangulaciones tienden a 0, por el corolario 1. Repitiendo el razonamiento, para cada T^m tomamos el etiquetado L^m análogo al establecido anteriormente, y resulta que o bien existe un símplice alternado positivo o bien existe un borde complementario. En particular, hay infinitas triangulaciones cumpliendo lo primero, o hay infinitas triangulaciones cumpliendo lo segundo.

Supongamos que estamos en el primer caso. Considerando las triangulaciones para las que existe un símplice alternado positivo podemos tomar una sucesión de los baricentros de estos símplices alternados positivos. Por construcción de los etiquetados, los elementos más cercanos a los vértices son, respectivamente, $\{w_1, -w_2, w_3, \dots, (-1)^n w_{n+1}\}$. Como los diámetros de los símplices alternados positivos convergen a 0, la sucesión de baricentros $\{\hat{z}_m\}_{m=1}^\infty$ converge a un punto

límite z y, por ser f continua, $f(\hat{z}_m)$ converge a $f(z)$, que es equidistante de todos los puntos $\{w_1, -w_2, w_3, \dots, (-1)^n w_{n+1}\}$. El único punto con esta propiedad es el cero, luego z ha de cumplir que $f(z) = \mathbf{0}$, y por tanto que $h(z) = \mathbf{0}$, en contradicción con que no exista $\mathbf{x} \in \hat{\mathbb{S}}^n$ con $h(\mathbf{x}) = \mathbf{0}$.

Supongamos ahora que estamos en el segundo caso, es decir, existen infinitos bordes con etiquetado $\{i, -i\}$. De nuevo, tomando la sucesión de los baricentros, obtenemos un punto límite z tal que $f(z)$ está a la misma distancia de w_i y $-w_i$. Este punto solo puede ser el cero, luego repitiendo el razonamiento anterior llegamos a una contradicción. $\square$

2.4. Lema de Tucker

El último enunciado equivalente al Teorema de Borsuk-Ulam que veremos es el Lema de Tucker.

9. Sea T una triangulación de $\hat{E}^n = \{\mathbf{x} \in \mathbb{R}^n : \sum_{i=1}^n |x_i| \leq 1\}$ que sea simétrica en la frontera, $\hat{\mathbb{S}}^{n-1}$. Sea L un n -etiquetado de T que sea antipodal en $V(T) \cap \hat{\mathbb{S}}^{n-1}$. Entonces, T tiene un borde complementario.

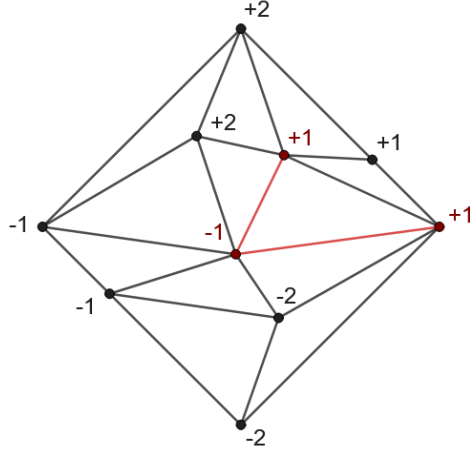


Figura 13: Triangulación de $\hat{E}^2$ con un etiquetado antipodal en la frontera. En rojo se indican los bordes complementarios.

Demostración.

(8 $\Rightarrow$ 9) Sea T una triangulación de $\hat{E}^n$ simétrica en $\hat{\mathbb{S}}^{n-1}$ y sea L un n -etiquetado antipodal en la frontera. Como todo n -etiquetado es también un $(n+1)$ -etiquetado, por (8) sabemos que si no hay un n -símplice alternado positivo en T entonces existe un borde complementario. Sin embargo, no pue-

de existir este n -símplice alternado positivo, porque solo tenemos las etiquetas $\{1, -1, 2, -2, \dots, n, -n\}$, cuando necesitaríamos tener etiquetas con al menos $n + 1$ valores absolutos diferentes.

(9 $\Rightarrow$ 4) Por reducción al absurdo, sea $f : \hat{E}^n \rightarrow \hat{S}^{n-1}$ una aplicación continua antipodal en $\hat{S}^{n-1}$. Sea $\varepsilon = \frac{1}{\sqrt{n}}$. Como f es una aplicación continua definida en un espacio compacto, entonces es uniformemente continua, luego para nuestro ε existe $\delta > 0$ tal que si $\|\mathbf{x} - \mathbf{y}\|_\infty < \delta$ entonces $\|f(\mathbf{x}) - f(\mathbf{y})\|_\infty < 2\varepsilon$. Tomamos una triangulación T de $\hat{E}^n$ tal que cada símplice tenga diámetro (considerando la norma del supremo) menor que δ .

Observemos que hemos elegido ε de tal manera que para cada $\mathbf{x} \in \hat{S}^{n-1}$ al menos una de sus componentes x_j es mayor o igual que ε (en caso contrario, $\sum_{i=1}^n x_i^2 < \sum_{i=1}^n \varepsilon^2 = n \frac{1}{n} = 1$, contradiciendo que $\mathbf{x} \in \hat{S}^{n-1}$). Gracias a esto, la aplicación $k : V(T) \rightarrow \{1, \dots, n\}$ dada por

$$k(v) = \min\{i : |f_i(v)| \geq \varepsilon\}$$

está bien definida, donde $f_i(v)$ es la i -ésima componente de $f(v)$. Además, como f es antipodal en la frontera,

$$k(-v) = \min\{i : |f_i(-v)| \geq \varepsilon\} = \min\{i : |-f_i(v)| \geq \varepsilon\} = k(v)$$

para cada vértice v en $\hat{S}^{n-1}$.

Ahora, sea $\lambda : V(T) \rightarrow \{1, -1, 2, -2, \dots, n, -n\}$ la aplicación dada por

$$\lambda(v) = \begin{cases} +k(v) & \text{si } f_{k(v)}(v) > 0 \\ -k(v) & \text{si } f_{k(v)}(v) < 0 \end{cases}$$

Veamos que λ es antipodal en la frontera. Supongamos sin pérdida de generalidad que $f_{k(v)}(v) > 0$ para cierto $v \in \hat{S}^{n-1}$. Entonces,

$$\lambda(-v) = -k(-v) = -k(v) = -\lambda(v).$$

Estamos en condiciones de aplicar el lema de Tucker, que nos dice que existe un borde complementario. Sean v_1 y v_2 los vértices de este borde (cuya distancia es menor que δ , por el diámetro de la triangulación elegida) con etiquetas $i > 0$ y $-i$ respectivamente. Es decir,

$$i = \lambda(v_1) = k(v_1) = \min\{i : |f_i(v_1)| \geq \varepsilon\}$$

y, como $f_i(v_1) > 0$, se cumple que $f_i(v_1) \geq \varepsilon$. De forma análoga se deduce que $f_i(v_2) \leq -\varepsilon$. Luego $\|f(v_1) - f(v_2)\|_\infty \geq 2\varepsilon$, en contradicción con que, si $\|\mathbf{x} - \mathbf{y}\|_\infty < \delta$, entonces $\|f(\mathbf{x}) - f(\mathbf{y})\|_\infty < 2\varepsilon$. $\square$

3. Demostración del Teorema de Borsuk-Ulam

Demostraremos primero los casos $n = 1$ y $n = 2$ antes de proceder a la demostración en el caso general. El primero utiliza el Teorema de los valores intermedios [Mun07, Teorema 24.3], que requiere la conexión de $\mathbb{S}^1$. Para el caso $n = 2$ utilizamos nuestro conocimiento del grupo fundamental de la circunferencia. El caso general es más complicado y se puede tratar de distintas maneras. Una demostración de corte combinatorio se puede encontrar en [Mat03, Sección 2.3], mientras que [Mas91, Capítulo 15, Teorema 2.4] la aborda mediante grupos de homología. La demostración que presentaremos en esta memoria es la demostración geométrica de [Mat03, Sección 2.2].

3.1. Caso $n = 1$

Vamos a demostrar la versión (2) del Teorema de Borsuk-Ulam, es decir, probaremos que para toda aplicación antipodal y continua $f : \mathbb{S}^1 \rightarrow \mathbb{R}$ existe un punto $\mathbf{x} \in \mathbb{S}^1$ tal que $f(\mathbf{x}) = 0$.

Sea $f : \mathbb{S}^1 \rightarrow \mathbb{R}$ antipodal y continua. Si $f(\mathbf{x}) = 0$ para todo $\mathbf{x} \in \mathbb{S}^1$, no hay nada que probar.

Supongamos ahora que existe un punto $\mathbf{y} \in \mathbb{S}^1$ tal que $f(\mathbf{y}) > 0$. Entonces, por ser antipodal, $f(-\mathbf{y}) < 0$. Por el Teorema de los valores intermedios, existe un punto $\mathbf{x}_0 \in \mathbb{S}^1$ tal que $f(\mathbf{x}_0) = 0$.

3.2. Caso $n = 2$

Nos basaremos en [Mun07, Sección 57]. Probaremos la versión (5) de la Sección 2.1 (una aplicación antipodal y continua $f : \mathbb{S}^{n-1} \rightarrow \mathbb{S}^{n-1}$ no es homótopa a una aplicación constante).

Sea $b_0 = (1, 0)$, y sea $h : \mathbb{S}^1 \rightarrow \mathbb{S}^1$ una aplicación antipodal. Podemos pedir que $h(b_0) = b_0$, porque en caso contrario bastaría con hacer una rotación r de manera que $r(h(b_0)) = b_0$. Las rotaciones son antipodales, por lo que $r \circ h$ sigue siendo antipodal.

Identificando $\mathbb{R}^2$ con el plano complejo, vamos a considerar la aplicación $\rho : \mathbb{S}^1 \rightarrow \mathbb{S}^1$ que lleva z a z^2 . Como estamos trabajando en $\mathbb{S}^1$, podemos escribir $z = e^{i\alpha}$, luego $\rho(z) = e^{2i\alpha}$.

Veamos que es ρ una proyección recubridora. Es continua, y si $x \in \mathbb{S}^1$ podemos escribirlo como $x = e^{i\alpha}$ para cierto α , y se tiene que

$$\rho(e^{i\frac{\alpha}{2}}) = x,$$

por lo que también es sobreyectiva.

Tenemos que ver también que para cada punto de $\mathbb{S}^1$ existe un entorno que sea un abierto bien recubierto. Sean $V_1 = \mathbb{S}^1 \setminus \{(1, 0)\}$ y $V_2 = \mathbb{S}^1 \setminus \{(-1, 0)\}$.

Observemos que $\mathbb{S}^1 = V_1 \cup V_2$. Veamos que son abiertos bien recubiertos. Se tiene que $\rho^{-1}(V_1) = S_1 \cup S_2$, donde

$$S_1 = \{e^{i\delta} : \delta \in (0, \pi)\} = \{(x, y) \in \mathbb{S}^1 : y > 0\},$$

$$S_2 = \{e^{i\delta} : \delta \in (\pi, 2\pi)\} = \{(x, y) \in \mathbb{S}^1 : y < 0\},$$

y $\rho^{-1}(V_2) = S_3 \cup S_4$, donde

$$S_3 = \{e^{i\delta} : \delta \in \left(-\frac{\pi}{2}, \frac{\pi}{2}\right)\} = \{(x, y) \in \mathbb{S}^1 : x > 0\},$$

$$S_4 = \left\{e^{i\delta} : \delta \in \left(\frac{\pi}{2}, \frac{3\pi}{2}\right)\right\} = \{(x, y) \in \mathbb{S}^1 : x < 0\}.$$

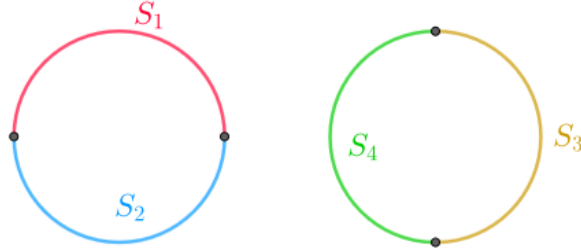


Figura 14: S_1 , S_2 , S_3 y S_4 .

Las aplicaciones $\rho|_{S_1}$ y $\rho|_{S_2}$ son homeomorfismos de S_1 y S_2 en V_1 , respectivamente, y ocurre lo mismo con $\rho|_{S_3}$, $\rho|_{S_4}$ y V_2 . Es decir, hemos probado que ρ es una proyección recubridora. En particular, por la Proposición 2, es abierta.

Por ser sobreyectiva, continua y abierta ρ es una aplicación cociente que relaciona los puntos z y $-z$ de $\mathbb{S}^1$. Por esto y porque $h(-z) = -h(z)$ se tiene que $\rho(h(z)) = \rho(h(-z))$.

Consideramos la aplicación continua $k : \mathbb{S}^1 \rightarrow \mathbb{S}^1$ dada por

$$k(e^{i\alpha}) = (h(e^{i\frac{\alpha}{2}}))^2.$$

Observemos que podemos escribir cada punto de $\mathbb{S}^1$ como $e^{i\alpha}$ o como $e^{i(\alpha+2m\pi)}$ con $m \in \mathbb{Z}$, y sus respectivas imágenes por k son $h(e^{i\frac{\alpha}{2}})^2$ y $h(e^{i(\frac{\alpha}{2}+m\pi)})^2$. Si m es par, $e^{i\frac{\alpha}{2}} = e^{i(\frac{\alpha}{2}+m\pi)}$, luego las imágenes coinciden. Si es impar, $e^{i\frac{\alpha}{2}}$ y $e^{i(\frac{\alpha}{2}+m\pi)}$ son puntos antipodales. Pero como h es antipodal, al tomar su imagen

por h solo cambia el signo, por lo que al tomar el cuadrado ambos coinciden. Es decir, k está bien definida.

Si $z = e^{i\alpha} \in \mathbb{S}^1$, entonces

$$(k \circ \rho)(z) = (k \circ \rho)(e^{i\alpha}) = k(e^{2i\alpha}) = (h(e^{i\alpha}))^2 = (\rho \circ h)(e^{i\alpha}) = (\rho \circ h)(z).$$

Luego $k \circ \rho = \rho \circ h$, es decir, k hace conmutativo el siguiente diagrama:

$$\begin{array}{ccc} \mathbb{S}^1 & \xrightarrow{h} & \mathbb{S}^1 \\ \rho \downarrow & & \downarrow \rho \\ \mathbb{S}^1 & \xrightarrow{k} & \mathbb{S}^1 \end{array}$$

Observemos que $\rho(b_0) = b_0^2 = b_0$, luego $k(b_0) = (k \circ \rho)(b_0) = (\rho \circ h)(b_0) = b_0$.

Sea ahora γ un camino en $\mathbb{S}^1$ de b_0 en $-b_0$. Su proyección $\alpha = \rho \circ \gamma$ es un lazo, luego por definición γ es la elevación de α a b_0 .

$$\begin{array}{ccc} & & \mathbb{S}^1 \\ & \nearrow \gamma & \downarrow \rho \\ I & \xrightarrow{\alpha} & \mathbb{S}^1 \end{array}$$

Veamos que α no es homótopo al lazo constante e_{b_0} . Supongamos que $\alpha \simeq e_{b_0}$ rel $\{0, 1\}$ para llegar a un absurdo. Ambos son caminos que empiezan en b_0 , luego por la Proposición 5 sus elevaciones son homótopas rel $\{0, 1\}$. En particular, tienen el mismo punto final. Pero la elevación de e_{b_0} acaba en b_0 , mientras que γ acaba en $-b_0$, y llegaríamos a un absurdo. Por lo tanto, $[\alpha]$ no es el neutro.

Veamos ahora que k_* no es el homomorfismo trivial. Como $h \circ \gamma$ es un camino en $\mathbb{S}^1$ de b_0 en $-b_0$, por lo anterior sabemos que $[\rho \circ h \circ \gamma]$ no es el neutro, y como $k \circ \rho = \rho \circ h$ esto quiere decir que $[k \circ \rho \circ \gamma] = k_*([\rho \circ \gamma])$ no es el elemento neutro. Entonces el homomorfismo $k_* : \pi_1(\mathbb{S}^1, b_0) \rightarrow \pi_1(\mathbb{S}^1, b_0)$ no es el trivial.

Nos falta ver que h_* tampoco es el homomorfismo trivial.

Probemos que si un homomorfismo de grupos $\varphi : \mathbb{Z} \rightarrow \mathbb{Z}$ no es trivial, entonces φ es inyectivo. Supongamos por reducción al absurdo que $\varphi : \mathbb{Z} \rightarrow \mathbb{Z}$ es un homomorfismo de grupos no inyectivo. Entonces, existe $a \in \mathbb{Z}$, $a \neq 0$, tal que $\varphi(a) = 0$. Pero $\varphi(a) = a\varphi(1) = 0$. Como $\mathbb{Z}$ es dominio de integridad, no tiene divisores de cero, luego $\varphi(1) = 0$. Por lo tanto, $\varphi(b) = b\varphi(1) = 0$ para todo $b \in \mathbb{Z}$, y se tiene que φ es el homomorfismo trivial.

Utilizando lo anterior, como k_* no es el homomorfismo trivial, entonces k_* es inyectivo. También lo es ρ_* , puesto que es equivalente a multiplicar por 2 en $\mathbb{Z}$. Por lo tanto, también $k_* \circ \rho_* = (k \circ \rho)_*$ es inyectivo. Como

$$\rho_* \circ h_* = (\rho \circ h)_* = (k \circ \rho)_*,$$

el homomorfismo h_* ha de ser inyectivo, y por tanto no es trivial.

Finalmente, si h fuera homótopa a una aplicación constante, por la Proposición 9 tendríamos que h_* sería el homomorfismo trivial, contradiciendo lo anterior. Concluimos que h_* no es homótopa a una constante, probando así el Teorema de Borsuk-Ulam para $n = 2$.

3.3. Caso general

Durante la demostración, recurriremos a un tipo de aplicaciones a las que llamaremos *genéricas*, que nos asegurarán que el razonamiento que seguiremos sea válido. Para ello, necesitaremos ver que, en efecto, es posible encontrar “fácilmente” estas aplicaciones genéricas. Es con este propósito que introducimos el concepto de *nunca denso*, así como los tres lemas posteriores.

Definición 17. Decimos que un subconjunto A de un espacio topológico X es *nunca denso* si $\text{Int}(\overline{A}) = \emptyset$.

Lema 1. Las tres propiedades siguientes son equivalentes:

1. A es nunca denso en X .
2. $\text{Int}(X \setminus A)$ es denso en X .
3. Para todo abierto U no vacío existe otro abierto no vacío $U' \subset U$ tal que $U' \cap A = \emptyset$.

Demostración.

(1 $\Rightarrow$ 2) Supongamos que A es un conjunto nunca denso, y sea U un abierto no vacío. El conjunto $(X \setminus \overline{A}) \cap U$ es un abierto de U por ser $X \setminus \overline{A}$ un abierto de X , y no es vacío, porque, si lo fuera, se tendría que $U \subset \overline{A}$, y por tanto $U = \text{Int}(U) \subset \text{Int}(\overline{A}) = \emptyset$. Luego $\text{Int}(X \setminus A) = X \setminus \overline{A}$ interseca a todo abierto no vacío de X , por lo que es denso en X .

(2 $\Rightarrow$ 3) Si $\text{Int}(X \setminus A)$ es denso en X , para todo U abierto no vacío se tiene que $(X \setminus \overline{A}) \cap U$ es un abierto de U no vacío, luego existe un abierto U' contenido en $(X \setminus \overline{A}) \cap U$. Entonces, $U' \cap A \subset (X \setminus \overline{A}) \cap U \cap A = \emptyset$.

(3 $\Rightarrow$ 1) Utilizaremos reducción al absurdo, suponiendo que $\text{Int}(\overline{A}) \neq \emptyset$. Como $\text{Int}(\overline{A})$ es abierto, existe un abierto $U \subset \text{Int}(\overline{A}) \subset \overline{A}$. Por hipótesis, existe $U' \subset U$ tal que $U' \cap A = \emptyset$. Entonces, $U' \subset \overline{A} \setminus A$. Esto es absurdo. $\square$

Lema 2. Dado un polinomio $p(x_1, \dots, x_k)$ con coeficientes en $\mathbb{R}$ que no sea idénticamente nulo, el conjunto $Z(p) = \{\mathbf{x} \in \mathbb{R}^k : p(\mathbf{x}) = 0\}$ de puntos donde se anula es nunca denso.

Demostración. Lo demostraremos por inducción sobre el número de variables del polinomio. Sea $k = 1$ y sea $p(x)$ un polinomio no idénticamente nulo. El

cardinal del conjunto $\{x \in \mathbb{R} : p(x) = 0\}$ es, como mucho, el grado de p . Es un conjunto cerrado (puesto que es unión finita de conjuntos cerrados), y su interior es vacío, luego es nunca denso.

Supongamos que se cumple para $k-1$, y veamos qué ocurre para un polinomio $p(x_1, \dots, x_k)$ en k variables que no sea el nulo. Supongamos que existe una bola abierta $B \subset Z(p)$, y tomamos un hiperplano H dado por

$$x_k - a_1x_1 - a_2x_2 - \dots - a_{k-1}x_{k-1} - b = 0.$$

Escribimos $h = a_1x_1 + \dots + a_{k-1}x_{k-1} + b$. Sea $\tilde{p}(x_1, \dots, x_{k-1}) = p(x_1, \dots, x_{k-1}, h)$. Podemos elegir H de forma que $\tilde{p}$ no sea idénticamente nulo, veámoslo por reducción al absurdo. Supongamos que $\tilde{p}$ es idénticamente nulo en todo hiperplano, y elegimos un H_1 con su respectivo h_1 . Como $p(x_1, \dots, x_{k-1}, x_k) = 0$ en todo H_1 , entonces p es divisible por $x_k - h_1$. Repitiendo el razonamiento para otros hiperplanos $H_2, \dots, H_r$ distintos con sus respectivos $h_2, \dots, h_r$ se deduce que p es divisible por $x_k - h_2, \dots, x_k - h_r$, luego p tiene grado mayor o igual que r en x_k . Podemos concluir que el grado de p en x_k es arbitrariamente grande, lo que es absurdo. Luego podemos elegir H que haga que el polinomio $\tilde{p}$ no sea siempre nulo. Es un polinomio en $k-1$ variables, luego el conjunto de puntos donde se anula es nunca denso.

Observemos que $B \cap H$ es un abierto de H (que se puede identificar con $\mathbb{R}^{k-1}$), puesto que B es un abierto de $\mathbb{R}^k$. Por otro lado,

$$\begin{aligned} B \cap H &\subset \{\mathbf{x} \in \mathbb{R}^k : x_k - a_1x_1 - \dots - a_{k-1}x_{k-1} - b = 0 \text{ y } p(\mathbf{x}) = 0\} \\ &= \{(x_1, \dots, x_{k-1}) \in \mathbb{R}^{k-1} : p(x_1, \dots, x_{k-1}, h) = 0\} = Z(\tilde{p}). \end{aligned}$$

Entonces, $B \cap H$ no contiene ninguna bola abierta cuya intersección con $Z(\tilde{p})$ sea vacía, luego $Z(\tilde{p})$ no es nunca denso, en contra de la hipótesis de inducción. Esto implica que no existe una bola abierta $B \subset Z(p)$. Como $Z(p)$ es cerrado, entonces $\text{Int}(\overline{Z(p)}) = \text{Int}(Z(p)) = \emptyset$, luego $Z(p)$ es nunca denso. $\square$

Lema 3. *La unión finita de conjuntos nunca densos es un conjunto nunca denso.*

Demostración. Basta probarlo para $n = 2$, dado que el caso general se sigue por inducción. Sean A_1 y A_2 dos conjuntos nunca densos, y sea B una bola abierta. Como A_1 es nunca denso, existe $B'_1 \subset B$ tal que $B'_1 \cap A_1 = \emptyset$. Como A_2 es nunca denso y B'_1 es una bola abierta, existe $B'_2 \subset B'_1$ tal que $B'_2 \cap A_2 = \emptyset$. Entonces, para toda bola abierta B existe una bola abierta $B'_2 \subset B$ tal que $B'_2 \cap (A_1 \cup A_2) = \emptyset$. $\square$

Introducimos a continuación un ejemplo de conjunto nunca denso que utilizaremos en la demostración.

Vamos a decir que una aplicación afín $h : \sigma \rightarrow \mathbb{R}^n$, donde σ es un símplex de dimensión $n+1$, es genérica si $h^{-1}(\mathbf{0})$ no interseca caras de σ de dimensión

menor que n . Es decir, si h es genérica, o bien $h^{-1}(\mathbf{0})$ es el vacío o bien es un segmento cuyo interior está en el interior de σ y sus extremos están en dos caras distintas de σ .

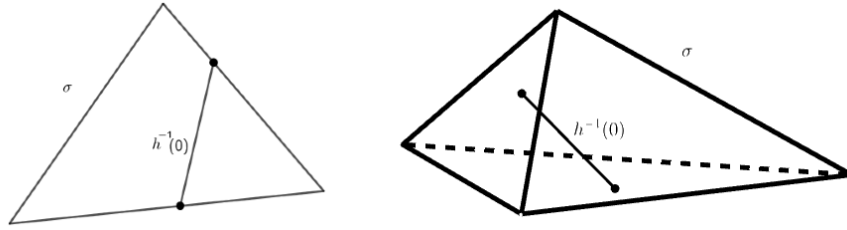


Figura 15: $h^{-1}(\mathbf{0})$ para una aplicación genérica si $n = 1$ o $n = 2$.

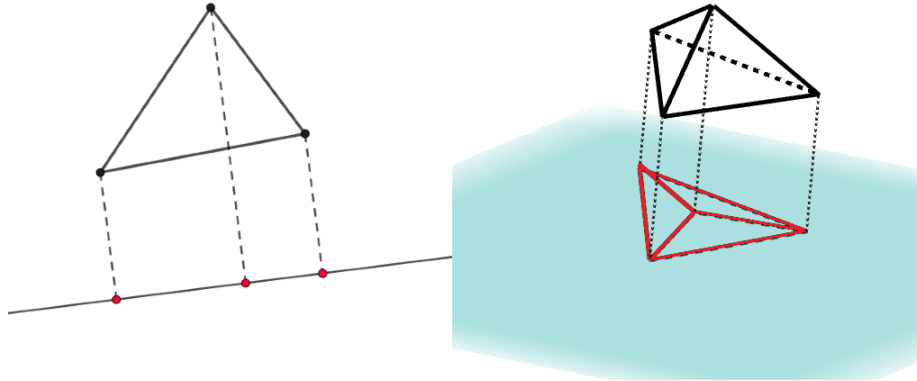


Figura 16: Si $\mathbf{0}$ es uno de los puntos marcados en rojo, entonces h no es genérica. En el primer caso, $h^{-1}(\mathbf{0})$ interseca un símple de dimensión 0, y en el segundo un símple de dimensión 0 o 1.

Cada aplicación afín $h : \mathbb{R}^{n+1} \rightarrow \mathbb{R}^n$ dada por $h(\mathbf{x}) = A\mathbf{x} + \mathbf{b}$ se puede representar como una matriz $(A|\mathbf{b})$ de dimensión $n \times (n+2)$ con entradas en $\mathbb{R}$. Vamos a probar que las aplicaciones no genéricas son los ceros de un polinomio.

Consideremos el $(n+1)$ -símple $\sigma = \langle \mathbf{0}, e_1, \dots, e_{n+1} \rangle$, donde

$$e_i = (0, \dots, 0, 1, 0, \dots, 0)$$

con 1 en la posición i . Si $h^{-1}(\mathbf{0})$ corta a σ en una de las caras de dimensión $n-1$ que contienen al cero, entonces

$$\begin{pmatrix} a_{1,1} & a_{1,2} & \cdots & a_{1,n+1} \\ a_{2,1} & a_{2,2} & \cdots & a_{2,n+1} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n,1} & a_{n,2} & \cdots & a_{n,n+1} \end{pmatrix} \begin{pmatrix} x_1 \\ \vdots \\ x_{i-1} \\ 0 \\ x_{i+1} \\ \vdots \\ x_{j-1} \\ 0 \\ x_{j+1} \\ \vdots \\ x_{n+1} \end{pmatrix} + \begin{pmatrix} b_1 \\ b_2 \\ \vdots \\ b_n \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}.$$

Es decir, tenemos un sistema de n ecuaciones con $n - 1$ incógnitas

$$\begin{cases} a_{1,1}x_1 + \cdots + a_{1,i-1}x_{i-1} + a_{1,i+1}x_{i+1} + \cdots + a_{1,j-1}x_{j-1} + a_{1,j+1}x_{j+1} + \cdots + a_{1,n+1}x_{n+1} = 0 \\ a_{2,1}x_1 + \cdots + a_{2,i-1}x_{i-1} + a_{2,i+1}x_{i+1} + \cdots + a_{2,j-1}x_{j-1} + a_{2,j+1}x_{j+1} + \cdots + a_{2,n+1}x_{n+1} = 0 \\ \vdots \\ a_{n,1}x_1 + \cdots + a_{n,i-1}x_{i-1} + a_{n,i+1}x_{i+1} + \cdots + a_{n,j-1}x_{j-1} + a_{n,j+1}x_{j+1} + \cdots + a_{n,n+1}x_{n+1} = 0. \end{cases}$$

Tiene solución si y solo si el rango de $(A|\mathbf{b})$ es menor o igual que $n - 1$, es decir, si y solo si

$$\begin{vmatrix} a_{1,1} & \cdots & a_{1,i-1} & a_{1,i+1} & \cdots & a_{1,j-1} & a_{1,j+1} & \cdots & a_{1,n+1} & b_1 \\ a_{2,1} & \cdots & a_{2,i-1} & a_{2,i+1} & \cdots & a_{2,j-1} & a_{2,j+1} & \cdots & a_{2,n+1} & b_2 \\ \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ a_{n,1} & \cdots & a_{n,i-1} & a_{n,i+1} & \cdots & a_{n,j-1} & a_{n,j+1} & \cdots & a_{n,n+1} & b_n \end{vmatrix} = 0.$$

Este determinante es un polinomio evaluado en

$$(a_{k,l})_{l \in \{1, \dots, n+1\} \setminus \{i,j\}}, (b_k)_{k \in \{1, \dots, n\}}.$$

En resumen, para que $h^{-1}(\mathbf{0})$ no corte a esa cara, el polinomio dado por el determinante evaluado en esos puntos tiene que valer 0. Por el Lema 2, las aplicaciones h tales que $h^{-1}(\mathbf{0})$ corta a esa cara son un conjunto nunca denso.

Haciendo lo mismo para el resto de caras, por el Lema 3 se tiene que el conjunto de aplicaciones no genéricas es nunca denso.

Con esto, tenemos las herramientas necesarias para demostrar el Teorema de Borsuk-Ulam en el caso general. Probaremos por reducción al absurdo la versión (2) del teorema considerando la norma 1, que dice que para toda aplicación antipodal y continua $f : \hat{\mathbb{S}}^n \rightarrow \mathbb{R}^n$ existe un punto $\mathbf{x} \in \hat{\mathbb{S}}^n$ tal que $f(\mathbf{x}) = 0$.

Demostración. Supongamos que $f(\mathbf{x}) \neq 0$ para todo $\mathbf{x} \in \hat{\mathbb{S}}^n$. Como f no tiene ceros y está definida en el compacto $\hat{\mathbb{S}}^n$, se tiene que $\|f(\mathbf{x})\| > 0$, así que existe $\varepsilon > 0$, que no depende de $\mathbf{x}$, tal que $\|f(\mathbf{x})\| \geq \varepsilon$ para cada $\mathbf{x} \in \hat{\mathbb{S}}^n$.

La aplicación f es continua en un conjunto compacto, luego es uniformemente continua. Entonces, para este ε existe un $\delta > 0$ tal que $\|f(\mathbf{x}) - f(\mathbf{y})\| < \frac{\varepsilon}{2}$ si $\|\mathbf{x} - \mathbf{y}\| < \delta$.

Paso 1: Triangulamos $\hat{\mathbb{S}}^n \times I$.

Dado un símple σ en $\mathbb{R}^{n+1}$, recordemos que $\sigma \times I$ es su prisma simplicial. Observemos que $\hat{\mathbb{S}}^n$ es unión de una cantidad finita de símlices, luego $\hat{\mathbb{S}}^n \times I$ es unión de sus respectivos prismas simpliciales. Vamos a buscar una triangulación de $\hat{\mathbb{S}}^n \times I$ cuyos símlices tengan diámetro menor que δ .

En primer lugar, tomamos la triangulación de $\hat{\mathbb{S}}^n \times \{0\}$ dada por los símlices de la forma

$$\Delta^n = \{(x_1, \dots, x_{n+1}, 0) \in \mathbb{R}^{n+2} : x_i \geq 0, \sum_{i=1}^{n+1} x_i = 1\}$$

y sus reflexiones por $G = \{-1, 1\}^{n+1} \times \{0\}$, y la análoga para $\hat{\mathbb{S}}^n \times \{1\}$. En el resto de $\hat{\mathbb{S}}^n \times I$, triangulamos cada prisma simplicial $\sigma \times I$ con una triangulación T_σ como hemos visto en los preliminares, de forma que la unión de estas T_σ , a la que denotaremos T^1 , contenga a las triangulaciones de $\hat{\mathbb{S}}^n \times \{0\}$ y $\hat{\mathbb{S}}^n \times \{1\}$. Ahora, tomamos T^2 la subdivisión baricéntrica de T^1 , y continuamos haciendo subdivisiones baricéntricas sucesivas hasta obtener una triangulación $T = T^r$ de $\hat{\mathbb{S}}^n \times I$ tal que el diámetro de sus símlices sea menor que δ . Esto se puede hacer por la Proposición 7. Llamaremos T_0 a la triangulación dada por el subcomplejo contenido en $\hat{\mathbb{S}}^n \times \{0\}$, y T_1 a la dada por el subcomplejo contenido en $\hat{\mathbb{S}}^n \times \{1\}$.

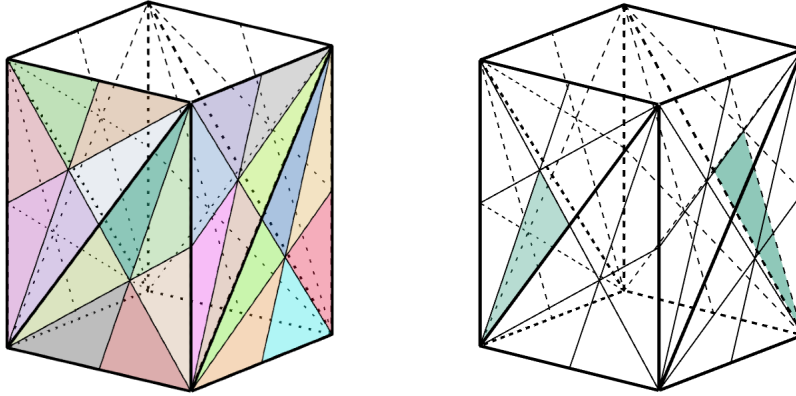


Figura 17: Una posible triangulación de $\hat{\mathbb{S}}^2 \times I$. Un símple y su imagen por ν .

Sea $\nu : \hat{\mathbb{S}}^n \times I \rightarrow \hat{\mathbb{S}}^n \times I$ la aplicación que lleva $(\mathbf{x}, t)$ en $(-\mathbf{x}, t)$. Observemos que cada símple $\sigma \in T$ va a su “opuesto” $\nu(\sigma) \in T$, y al revés. Además, si $(\mathbf{x}, t) \in \sigma \cap \nu(\sigma)$, entonces tanto $(\mathbf{x}, t)$ como $(-\mathbf{x}, t)$ están en σ , luego $(\mathbf{0}, t) \in \sigma$.

Necesariamente $(\mathbf{0}, t) \in \hat{\mathbb{S}}^n \times I$, pero $(0, \dots, 0) \notin \hat{\mathbb{S}}^n$. Como esto no puede ocurrir, $\sigma \cap \nu(\sigma) = \emptyset$.

Paso 2: Definimos g y una homotopía F entre g y f .

Tomamos una aplicación antipodal $g : \hat{\mathbb{S}}^n \cong \hat{\mathbb{S}}^n \times \{0\} \rightarrow \mathbb{R}^n$ dada por una proyección ortogonal de forma que tenga solo dos ceros, y que ambos ceros de g queden en el interior de dos símlices de T_0 de dimensión n , como muestran las figuras 18 y 19.

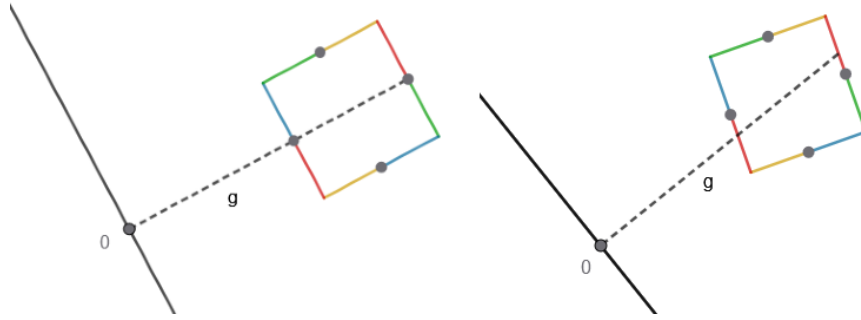


Figura 18: Caso $n = 1$. A la izquierda, una proyección incorrecta. A la derecha, una correcta.

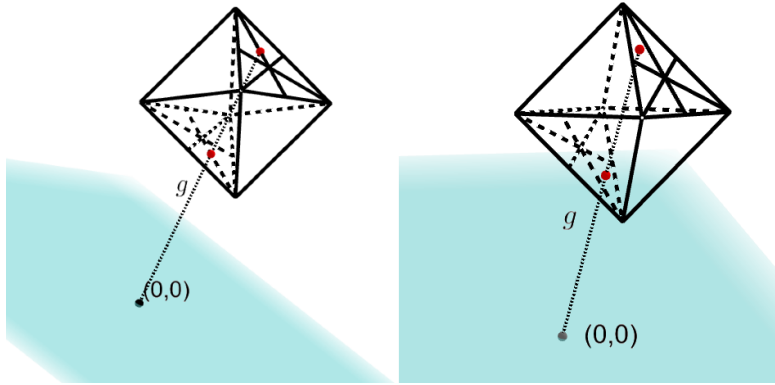


Figura 19: Caso $n = 2$. A la izquierda, una proyección incorrecta. A la derecha, una correcta.

Sea $F : \hat{\mathbb{S}}^n \times I \rightarrow \mathbb{R}^n$ la homotopía entre f y g dada por

$$F(\mathbf{x}, t) = (1 - t)g(\mathbf{x}) + tf(\mathbf{x}).$$

Observemos que

$$F(\nu(\mathbf{x}, t)) = F(-\mathbf{x}, t) = (1-t)g(-\mathbf{x}) + tf(-\mathbf{x}) = -(1-t)g(\mathbf{x}) - tf(\mathbf{x}) = -F(\mathbf{x}, t).$$

Paso 3: Extendemos $F|_{V(T)}$ de manera afín.

Consideramos $\bar{F} : \hat{\mathbb{S}}^n \times I \rightarrow \mathbb{R}^n$ la extensión afín de $F|_{V(T)}$ como en la Proposición 8.

Veamos que $\bar{F}$ no se anula en $\hat{\mathbb{S}}^n \times \{1\}$. Notemos que $\bar{F}|_{\hat{\mathbb{S}}^n \times \{1\}}$ es la extensión afín de $F|_{V(T) \cap (\hat{\mathbb{S}}^n \times \{1\})} = f|_{V(T_0)}$, y denotemos $\bar{f} = \bar{F}|_{\hat{\mathbb{S}}^n \times \{1\}}$. Ahora, para todo $\mathbf{y} \in \hat{\mathbb{S}}^n \times \{1\}$ existe un símple $\sigma \in T_1$ tal que $\mathbf{y} \in \sigma$.

Los valores de $\bar{f}$ en cada símple están comprendidos entre los valores de f en sus vértices, luego

$$\|f(\mathbf{y}) - \bar{f}(\mathbf{p})\| \leq \max_{\mathbf{x} \in \sigma} \|f(\mathbf{y}) - f(\mathbf{x})\|$$

para cada $\mathbf{p} \in \sigma$. Hemos elegido T de forma que $\|\mathbf{y} - \mathbf{x}\| < \delta$. En particular, $\|f(\mathbf{y}) - \bar{f}(\mathbf{y})\| \leq \frac{\varepsilon}{2}$. Ahora,

$$\|\bar{f}(\mathbf{y})\| \geq \|f(\mathbf{y})\| - \|f(\mathbf{y}) - \bar{f}(\mathbf{y})\| \geq \varepsilon - \frac{\varepsilon}{2} > 0.$$

Por otro lado, como g es afín (por ser una proyección) y $\bar{F}$ también, y coinciden en los vértices de T , ambas aplicaciones coinciden en $\hat{\mathbb{S}}^n \times \{0\}$. Recordemos que g tiene exactamente dos ceros y que están en el interior de símlices de dimensión n en T_0 , luego esto también se cumple para $\bar{F}$.

Paso 4: Tomamos una perturbación pequeña de $\bar{F}$ que sea genérica.

Diremos que una aplicación $G : \hat{\mathbb{S}}^n \times I \rightarrow \mathbb{R}^n$ es genérica si lo es en cada símple de dimensión $n+1$ de T .

Vamos a buscar una aplicación $P : \hat{\mathbb{S}}^n \times I \rightarrow \mathbb{R}^n$ afín en cada símple tal que $P(\nu(\mathbf{x}, t)) = -P(\mathbf{x}, t)$ y tal que $\hat{F} = \bar{F} + P$ satisfaga:

- Si $\bar{F}$ no se anula en un símple de $T_1 \cup T_0$, entonces $\hat{F}$ tampoco se anula.
- Si σ es un símple de T_0 que contiene a uno de los dos ceros de g , entonces $(\bar{F} + P)|_{\sigma}$ es una biyección entre σ y un símple τ de dimensión n en $\mathbb{R}^n$ que contiene al origen en su interior, es decir, $\mathbf{0} \in \text{int}(\tau)$.

Es decir, si todos los valores de P están suficientemente cerca de 0, entonces $\hat{F}$ no se anulará en $\hat{\mathbb{S}}^n \times \{1\}$ y tendrá dos ceros en $\hat{\mathbb{S}}^n \times \{0\}$, que están en el interior de símlices de dimensión n .

Antes de comenzar la demostración hemos visto que el conjunto de aplicaciones no genéricas en un símple es nunca denso, y el conjunto de símlices

de T es finito. Con esto, se puede probar que podemos elegir una aplicación P que toma valores suficientemente pequeños como para cumplir lo anterior y que haga que $\hat{F} = \bar{F} + P$ sea genérica.

Paso 5: La contraimagen del cero por $\hat{F}$ nos da un camino.

Para cada s3mplice $\tau \in T$ de dimensi3n n hay dos posibilidades: o bien es una cara de dos s3mplices de dimensi3n $n + 1$ de T (si $\tau \notin T_0 \cup T_1$), o bien es la cara de un 3nico s3mplice de dimensi3n $n + 1$ de T (si $\tau \in T_0 \cup T_1$).

Como $\hat{F}^{-1}(\mathbf{0}) = \cup_{\sigma_i \in T} \hat{F}|_{\sigma_i}^{-1}(\mathbf{0})$, y cada $\hat{F}|_{\sigma_i}^{-1}(\mathbf{0})$ o es vac3o o es un segmento como en la Figura 15, se tiene que:

- En $T_1 \cup T_0$ solo hay dos puntos que sean ceros de $\hat{F}$. Estos ceros est3n en s3mplices de dimensi3n n en T_0 , luego son cara de un 3nico s3mplice de dimensi3n $n + 1$. Un segmento que tenga por extremo uno de estos ceros no puede “continuar” en otro segmento contenido en otro s3mplice de dimensi3n $n + 1$, porque este s3mplice no existe.
- En $T \setminus T_1 \cup T_0$, cada segmento termina en un s3mplice de dimensi3n n que es cara de dos s3mplices de dimensi3n $n + 1$. No puede ser que en uno de los extremos no comience otro segmento, porque como $\hat{F}$ se anula en un punto del $(n + 1)$ -s3mplice tambi3n ha de hacerlo en todo un segmento. (Ver Figura 20).

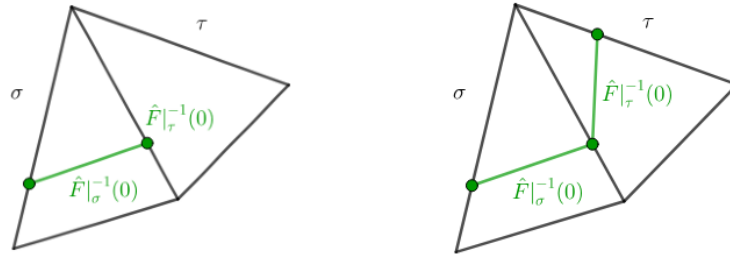


Figura 20: El caso de la izquierda no puede darse, porque $\hat{F}|_{\tau}^{-1}(\mathbf{0})$ ser3a un punto. Solo puede ocurrir lo que se indica en la imagen de la derecha.

Entonces, $\hat{F}^{-1}(\mathbf{0})$ contiene al camino formado por segmentos que hemos descrito. Llamaremos γ a ese camino.

Paso 6: γ es sim3trico bajo ν .

Si γ tiene longitud l , parametrizamos $\gamma : [0, 1] \rightarrow \hat{S}^n \times I$ de forma que $\gamma(z)$ es el punto de γ que est3 a distancia zl de uno de los ceros de $\hat{F}$ en $\hat{S}^n \times \{0\}$, con $z \in [0, 1]$. Veamos que $\gamma(z) = \nu(\gamma(z))$.

Como $\gamma(z) \in \hat{F}^{-1}(\mathbf{0})$, entonces $\hat{F}(\gamma(z)) = \bar{F}(\gamma(z)) + P(\gamma(z)) = 0$. Pero hemos definido P de forma que $P(\nu(\mathbf{x}, t)) = -P(\mathbf{x}, t)$ para todo $(\mathbf{x}, t) \in \hat{\mathbb{S}}^n \times I$. Además, hemos visto antes que

$$F(\mathbf{x}, t) = -F(\nu(\mathbf{x}, t)).$$

Luego $\bar{F}$ cumple que $-\bar{F}(\mathbf{x}, t) = \bar{F}(\nu(\mathbf{x}, t))$ en los vértices por coincidir con F y en el resto de puntos por ser afín en cada símplice, por ser simétricos los símplices. Es decir, $\bar{F}(\nu(\gamma(z))) + P(\nu(\gamma(z))) = 0$, por lo que $\nu(\gamma(z)) \in \hat{F}^{-1}(\mathbf{0})$.

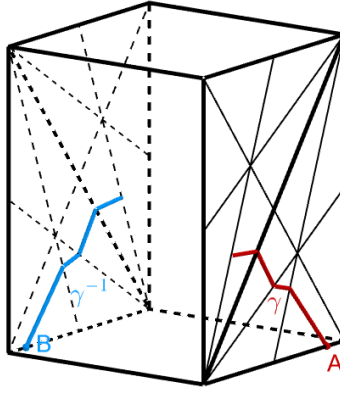


Figura 21: El camino simétrico a recorrer γ desde A es recorrer el camino inverso, γ^{-1} , desde B .

Llamemos A y B a los ceros de $\hat{F}$ en $\hat{\mathbb{S}}^n \times \{0\}$, con $\gamma(0) = A$ y $\gamma(1) = B$. Si nos movemos por γ desde A , sabemos que los puntos simétricos a los que recorramos también están en $\hat{F}^{-1}(\mathbf{0})$ por el párrafo anterior. Como comienzan en B , han de estar en la imagen de γ , luego corresponden a recorrer el camino inverso γ^{-1} . Es decir, $\nu(\gamma(z)) = \gamma(1 - z)$.

Paso 7: ν tiene un punto fijo.

Como $\nu(\gamma(z)) = \gamma(1 - z)$, en particular se tiene que $\nu(\gamma(\frac{1}{2})) = \gamma(\frac{1}{2})$, luego ν tiene un punto fijo. Un punto fijo para ν ha de cumplir que $(\mathbf{x}, t) = (-\mathbf{x}, t)$, es decir, $\mathbf{x} = -\mathbf{x}$. Solo puede ser el $\mathbf{0}$. Pero $\mathbf{0} \notin \hat{\mathbb{S}}^n$, luego hemos llegado a una contradicción. $\square$

4. Consecuencias del Teorema de Borsuk-Ulam

4.1. Teorema del punto fijo de Brouwer

El Teorema del punto fijo de Brouwer, demostrado para el caso $n = 3$ en 1904 por Piers Bohl y en 1909 por Brouwer, que también demostró en 1910 el caso general, es de gran importancia en la topología. En este apartado probaremos que se puede deducir a partir del Teorema de Borsuk-Ulam. Después veremos que el Teorema del punto fijo de Brouwer es equivalente a probar que la esfera no es un espacio contráctil.

Dada una aplicación f de un espacio X en sí mismo, un punto fijo es un elemento $\mathbf{x} \in X$ tal que $f(\mathbf{x}) = \mathbf{x}$. El enunciado del Teorema del punto fijo de Brouwer es el siguiente:

Teorema 2. *Toda aplicación continua $f : E^n \rightarrow E^n$ tiene un punto fijo.*

Daremos dos demostraciones de este hecho, ambas obtenidas del Teorema de Borsuk-Ulam. La primera será más breve, pero la segunda cuenta con la ventaja de mostrar con mayor claridad la relación entre ambos teoremas, utilizando una construcción directa en la que los puntos antipodales que tienen la misma imagen son los que nos dan el punto fijo. Comencemos por la demostración más breve, de [Wei21]. Recordamos la definición de retracto:

Definición 18. *Dado A subespacio de X , decimos que una aplicación continua $r : X \rightarrow A$ es una retracción si $r(a) = a$ para todo $a \in A$. Se dice en este caso que A es un retracto de X .*

Demostración. Supongamos que no es cierto el Teorema 2, es decir, tenemos una aplicación continua $f : E^n \rightarrow E^n$ que no tiene ningún punto fijo. Vamos a buscar una contradicción con la versión (4) del Teorema de Borsuk-Ulam, que afirma que no existe ninguna aplicación continua de E^n en $\mathbb{S}^{n-1}$ antipodal en $\mathbb{S}^{n-1}$.

Como $\mathbf{x} \neq f(\mathbf{x})$ para todo $\mathbf{x} \in E^n$, podemos considerar para cada punto la semirrecta

$$\mathbf{x} + t(f(\mathbf{x}) - \mathbf{x}), \text{ con } t \in [0, +\infty),$$

que parte de $\mathbf{x}$ y pasa por $f(\mathbf{x})$. Sea $h : E^n \rightarrow \mathbb{S}^{n-1}$ la aplicación que envía $\mathbf{x}$ al punto donde dicha semirrecta corta a $\mathbb{S}^{n-1}$.

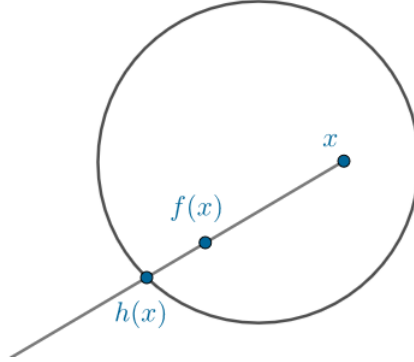


Figura 22: $h(\mathbf{x})$ para cierto $\mathbf{x} \in E^n$.

Veamos que h es continua. Si consideramos el punto $(x_1, \dots, x_n) \in E^n$, su imagen por h será $\mathbf{x} + t(\mathbf{x})(f(\mathbf{x}) - \mathbf{x})$, donde $t(\mathbf{x})$ se elige de forma que lo anterior esté en $\mathbb{S}^{n-1}$. Para ello, es necesario que

$$(x_1 + t(\mathbf{x})y_1)^2 + \dots + (x_n + t(\mathbf{x})y_n)^2 = 1,$$

donde $y_i = f_i(\mathbf{x}) - x_i$. Operando, se tiene que

$$\sum_{i=1}^n x_i^2 + t^2 \left(\sum_{i=1}^n y_i^2 \right) + 2t \left(\sum_{i=1}^n x_i y_i \right) = 1,$$

de donde

$$t(\mathbf{x}) = \frac{(\sum_{i=1}^n -x_i y_i) + \sqrt{(\sum_{i=1}^n x_i y_i)^2 - (\sum_{i=1}^n y_i^2)((\sum_{i=1}^n x_i^2) - 1)}}{\sum_{i=1}^n y_i^2}.$$

El radicando es mayor o igual que cero, porque $(\sum_{i=1}^n x_i^2) - 1 \leq 0$ para cada punto de E^n , luego $-(\sum_{i=1}^n y_i^2)((\sum_{i=1}^n x_i^2) - 1) \geq 0$. El denominador no se anula, porque solo lo haría si $f_i(\mathbf{x}) = x_i$ para todo i , pero entonces se tendría que $f(\mathbf{x}) = \mathbf{x}$, en contra de que f no tiene puntos fijos. Luego t es una función continua. Entonces, h también lo es.

Restringiendo h a $\mathbb{S}^{n-1}$, obtenemos la aplicación identidad $\text{Id}_{\mathbb{S}^{n-1}}$: en efecto, si $\mathbf{x} \in \mathbb{S}^{n-1}$, el punto de corte con $\mathbb{S}^{n-1}$ de cualquier semirrecta que pase por $\mathbf{x}$ es el propio $\mathbf{x}$. Dicho de otro modo, h es una retracción de E^n en $\mathbb{S}^{n-1}$. En particular, esto implica que h es antipodal en $\mathbb{S}^{n-1}$, puesto que la identidad lo es. Esto se contradice con la versión (4) del Teorema de Borsuk-Ulam. Luego f ha de tener un punto fijo. $\square$

La idea intuitiva de la segunda demostración, [Su97], es la siguiente: en lugar de trabajar en E^n , lo haremos en el n -cubo $[-1, 1]^n$ (la bola cerrada en norma

infinito). Observemos que $\partial[-1, 1]^{n+1} = \mathbb{S}_\infty^n$. Tomemos las caras

$$S_+ = \{\mathbf{x} \in \mathbb{S}_\infty^n : \mathbf{x} = (x_1, \dots, x_n, 1)\},$$

$$S_- = \{\mathbf{x} \in \mathbb{S}_\infty^n : \mathbf{x} = (x_1, \dots, x_n, -1)\}.$$

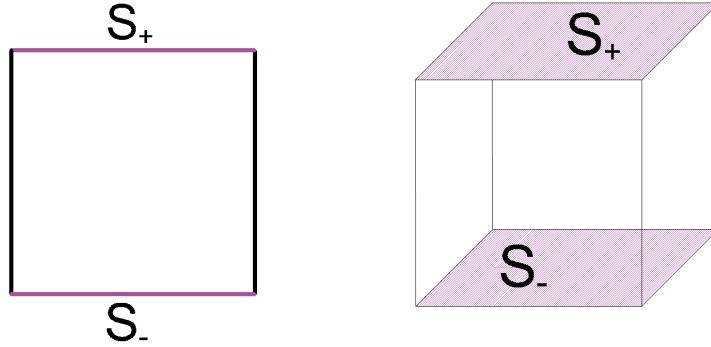


Figura 23: S_+ y S_- para los casos $n = 1$ y $n = 2$.

Dada una aplicación continua $f : [-1, 1]^n \rightarrow [-1, 1]^n$ vamos a definir una aplicación antipodal continua $g : \mathbb{S}_\infty^n \rightarrow \mathbb{R}^n$ de forma que:

1. Dado $x \in S_+$, si $g(\mathbf{x}) = \mathbf{0}$ entonces $\mathbf{x}$ es un punto fijo de f .
2. Dado $x \in S_-$, $g(-\mathbf{x}) = -g(\mathbf{x})$.
3. En $\mathbb{S}_\infty^n \setminus (S_+ \cup S_-)$, g es continua, antipodal y no nula.

Si lo conseguimos, podemos aplicar la versión (2) del Teorema de Borsuk-Ulam y obtener $\mathbf{x} \in \mathbb{S}_\infty^n$ tal que $g(\mathbf{x}) = \mathbf{0}$. Por las propiedades que cumple g , podemos asumir que $\mathbf{x} \in S_+$, que es el punto fijo del Teorema de Brouwer. Vista la idea, procedamos a la demostración del Teorema 2.

Demostración. Sea $f : [-1, 1]^n \rightarrow [-1, 1]^n$ una aplicación continua. Denotemos por $\pi : \mathbb{R}^{n+1} \rightarrow \mathbb{R}^n$ la proyección $\pi(x_1, \dots, x_{n+1}) = (x_1, \dots, x_n)$.

Paso 1: Definimos g en $S_+ \cup S_-$.

En S_+ , definimos g como $g_+(\mathbf{x}) = \pi(\mathbf{x}) - f(\pi(\mathbf{x}))$, y en S_- la definimos como $g_-(\mathbf{x}) = \pi(\mathbf{x}) + f(-\pi(\mathbf{x}))$. Observemos que si consideramos un $\mathbf{x} \in S_+$ tal que $g_+(\mathbf{x}) = \mathbf{0}$, entonces $\pi(\mathbf{x}) - f(\pi(\mathbf{x})) = \mathbf{0}$, y tendríamos que $\pi(\mathbf{x})$ sería un punto fijo de f . Esta definición asegura que g será antipodal en $S_+ \cup S_-$, ya que si $\mathbf{x} \in S_+$, entonces

$$g_-(-\mathbf{x}) = \pi(-\mathbf{x}) + f(-\pi(-\mathbf{x})) = -\pi(\mathbf{x}) + f(\pi(\mathbf{x})) = -g_+(\mathbf{x}).$$

Paso 2: Definimos g en el ecuador.

Sea $S_0 = \{\mathbf{x} \in \mathbb{S}_\infty^n : \mathbf{x} = (x_1, \dots, x_n, 0)\}$. Definimos g en S_0 como

$$g_0(\mathbf{x}) = \pi(\mathbf{x}) + \frac{g(x_1, \dots, x_n, -1) + g(x_1, \dots, x_n, 1)}{2}.$$

Así definida, g_0 es antipodal en S_0 :

$$\begin{aligned} g_0(-\mathbf{x}) &= \pi(-\mathbf{x}) + \frac{g_-(-x_1, \dots, -x_n, -1) + g_+(-x_1, \dots, -x_n, 1)}{2} \\ &= -\pi(\mathbf{x}) + \frac{-g_+(x_1, \dots, x_n, 1) - g_-(x_1, \dots, x_n, -1)}{2} = -g_0(\mathbf{x}). \end{aligned}$$

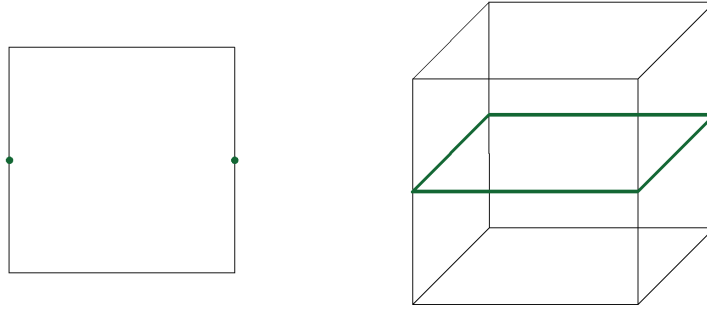


Figura 24: En color verde, S_0 para los casos $n = 1$ y $n = 2$.

Paso 3: Estudiamos g en el corte de una cara con S_0 , S_+ o S_- .

Sea F una de las otras caras del n -cubo, es decir, F es de la forma

$$F_i^+ = \{\mathbf{x} \in \mathbb{S}_\infty^n : \mathbf{x} = (x_1, \dots, x_{i-1}, 1, x_{i+1}, \dots, x_{n+1})\}$$

o

$$F_i^- = \{\mathbf{x} \in \mathbb{S}_\infty^n : \mathbf{x} = (x_1, \dots, x_{i-1}, -1, x_{i+1}, \dots, x_{n+1})\}$$

para algún $1 \leq i \leq n$. Por la definición de g , tenemos que

$$g_i(\mathbf{x}) = \begin{cases} x_i - f_i(\pi(\mathbf{x})) & \text{en } F \cap S_+ \\ x_i + f_i(-\pi(\mathbf{x})) & \text{en } F \cap S_-, \end{cases}$$

donde por g_i y f_i nos referimos a la proyección i -ésima de g y f , respectivamente.

Se tiene que $|f_i| \leq 1$, porque f_i tiene llegada en $[-1, 1]$. Tenemos que, si $x_i = 1$, entonces $g_i(\mathbf{x}) \geq 0$ en $(F \cap S_+) \cup (F \cap S_-) = F \cap (S_+ \cup S_-)$. En el caso $x_i = -1$, entonces $g_i(\mathbf{x}) \leq 0$ en $F \cap (S_+ \cup S_-)$. Luego para todo $\mathbf{x} \in F \cap (S_+ \cup S_-)$ se cumple que $g_i(\mathbf{x})$ o bien es nula o tiene el mismo signo que x_i .

Sea ahora $\mathbf{x} \in S_0$ tal que $\mathbf{x} \in F_i^+$ o $\mathbf{x} \in F_i^-$, es decir, tal que la coordenada i -ésima x_i cumpla que $|x_i| = 1$. Como $\mathbf{x} \in S_0$, por la definición de g en S_0 se tiene que

$$g_i(\mathbf{x}) = \pi_i(\mathbf{x}) + \frac{g_i(x_1, \dots, x_n, -1) + g_i(x_1, \dots, x_n, 1)}{2},$$

que o bien es 1 más una cantidad mayor o igual que 0, o bien es -1 más una cantidad menor o igual que 0 (porque, como ya hemos visto, como $(x_1, \dots, x_n, \pm 1) \in F \cap (S_+ \cup S_-)$ entonces $g_i(\mathbf{x})$ es cero o tiene el mismo signo que x_i). En ambos casos, tiene estrictamente el signo de x_i .

Paso 4: Extendemos g linealmente.

Vamos a extender g linealmente al resto de puntos de $\mathbb{S}_\infty^n$: para cada $\mathbf{x}$ tal que $0 \leq x_{n+1} \leq 1$, definimos

$$g(\mathbf{x}) = x_{n+1}g_+(x_1, \dots, x_n, 1) + (1 - x_{n+1})g_0(x_1, \dots, x_n, 0). \quad (2)$$

Y análogamente, para cada $\mathbf{x}$ tal que $-1 \leq x_{n+1} \leq 0$,

$$g(\mathbf{x}) = -x_{n+1}g_-(x_1, \dots, x_n, -1) + (1 + x_{n+1})g_0(x_1, \dots, x_n, 0). \quad (3)$$

La función g está definida y es continua en cada uno de los siguientes 5 subconjuntos que recubren $\mathbb{S}_\infty^n$: S_0 , S_+ , S_- , los puntos tales que $0 \leq x_{n+1} \leq 1$ y los puntos tales que $-1 \leq x_{n+1} \leq 0$. Recordemos cada una de estas definiciones:

$$g(\mathbf{x}) = \begin{cases} \pi(\mathbf{x}) - f(\pi(\mathbf{x})) & \text{si } \mathbf{x} \in S_+ \\ \pi(\mathbf{x}) + f(-\pi(\mathbf{x})) & \text{si } \mathbf{x} \in S_- \\ \pi(\mathbf{x}) + \frac{g_-(x_1, \dots, x_n, -1) + g_+(x_1, \dots, x_n, 1)}{2} & \text{si } \mathbf{x} \in S_0 \\ x_{n+1}g_+(x_1, \dots, x_n, 1) + (1 - x_{n+1})g_0(x_1, \dots, x_n, 0) & \text{si } 0 \leq x_{n+1} \leq 1 \\ -x_{n+1}g_-(x_1, \dots, x_n, -1) + (1 + x_{n+1})g_0(x_1, \dots, x_n, 0) & \text{si } -1 \leq x_{n+1} \leq 0. \end{cases}$$

Paso 5: Comprobamos que g cumple las propiedades deseadas.

Comprobemos que g está bien definida, es continua y es antipodal. Para empezar, veamos que g coincide en S_0 con sus definiciones en (2) y en (3); en S_+ con (2), y en S_- con (3).

- En S_0 , la coordenada x_{n+1} es 0, luego

$$\begin{aligned} g(\mathbf{x}) &= x_{n+1}g(x_1, \dots, x_n, 1) + (1 - x_{n+1})g(x_1, \dots, x_n, 0) \\ &= g(x_1, \dots, x_n, 0) = g_0(\mathbf{x}) \end{aligned}$$

y

$$\begin{aligned} g(\mathbf{x}) &= -x_{n+1}g(x_1, \dots, x_n, -1) + (1 + x_{n+1})g(x_1, \dots, x_n, 0) \\ &= g(x_1, \dots, x_n, 0) = g_0(\mathbf{x}). \end{aligned}$$

- En S_+ , la coordenada x_{n+1} es 1.

$$\begin{aligned} g(\mathbf{x}) &= x_{n+1}g(x_1, \dots, x_n, 1) + (1 - x_{n+1})g(x_1, \dots, x_n, 0) \\ &= g(x_1, \dots, x_n, 1) = g_+(\mathbf{x}). \end{aligned}$$

- En S_- , $x_{n+1} = -1$.

$$\begin{aligned} g(\mathbf{x}) &= -x_{n+1}g(x_1, \dots, x_n, -1) + (1 + x_{n+1})g(x_1, \dots, x_n, 0) \\ &= g(x_1, \dots, x_n, -1) = g_-(\mathbf{x}). \end{aligned}$$

Luego g está bien definida y es continua en $\mathbb{S}_\infty^n$.

Falta ver que es antipodal. Recordemos que sí que lo es en S_0 , S_+ y S_- . Dado $\mathbf{x}$ tal que $0 \leq x_{n+1} \leq 1$, entonces $-1 \leq -x_{n+1} \leq 0$, y

$$\begin{aligned} g(-\mathbf{x}) &= -(-x_{n+1})g_-(x_1, \dots, x_n, -1) + (1 + (-x_{n+1}))g_0(x_1, \dots, x_n, 0) \\ &= x_{n+1}(-g_+(x_1, \dots, x_n, 1)) + (1 - x_{n+1})(-g_0(x_1, \dots, x_n, 0)) \\ &= -x_{n+1}g_+(x_1, \dots, x_n, 1) + (1 - x_{n+1})g_0(x_1, \dots, x_n, 0) = -g(\mathbf{x}), \end{aligned}$$

luego es antipodal, como queríamos.

Veamos finalmente que si un punto no está en S_+ ni S_- , entonces su imagen por g no puede ser $\mathbf{0}$. Si $\mathbf{x} \notin S_+$, $\mathbf{x} \notin S_-$, ha de existir cierto k , con $1 \leq k \leq n$, tal que $|x_k| = 1$. Vamos a ver que $g_k(\mathbf{x}) \neq 0$. Sabemos que $g_k(x_1, \dots, x_n, 1)$ y $g_k(x_1, \dots, x_n, -1)$ son cero o tienen el mismo signo que x_k , y también que $g_k(x_1, \dots, x_n, 0)$ tiene estrictamente el mismo signo que x_k . Entonces, para el caso $0 \leq x_{n+1} \leq 1$,

$$g_k(\mathbf{x}) = x_{n+1}g_k(x_1, \dots, x_n, 1) + (1 - x_{n+1})g_k(x_1, \dots, x_n, 0)$$

donde $x_{n+1} \geq 0$ y $1 - x_{n+1} > 0$ (recordemos que $\mathbf{x} \notin S_+$), luego $g_k(\mathbf{x})$ tiene el mismo signo que x_k .

Como g es antipodal, de lo anterior se deduce que si para $-1 \leq x_{n+1} \leq 0$ entonces $g_k(\mathbf{x})$ tiene el mismo signo que x_k . En ambos casos, como $g_k(\mathbf{x}) \neq 0$, se tiene que $g(\mathbf{x}) \neq \mathbf{0}$.

Como g es antipodal, aplicando la versión (2) del Teorema de Borsuk-Ulam sabemos que existe un punto $\mathbf{x}_0$ tal que $g(\mathbf{x}_0) = \mathbf{0}$. Esto no puede ocurrir en ninguna cara que no sea S_+ o S_- . Podemos suponer sin pérdida de la generalidad que $\mathbf{x}_0$ está en S_+ (en caso contrario bastaría con considerar $-\mathbf{x}_0$, cuya imagen también es $\mathbf{0}$ por ser g antipodal). Entonces, $g(\mathbf{x}_0) = \mathbf{0} = \pi(\mathbf{x}_0) - f(\pi(\mathbf{x}_0))$, por lo que $\pi(\mathbf{x}_0)$ es un punto fijo para f . $\square$

Recordamos la siguiente definición.

Definición 19. Decimos que un espacio X es contráctil si la aplicación identidad Id_X es homótopa a una aplicación constante.

Sabemos que todo espacio contráctil es simplemente conexo, pero lo contrario no es cierto. En efecto $\mathbb{S}^n$ no es contráctil pese a ser simplemente conexo para $n \geq 2$.

Proposición 10. *Son equivalentes:*

1. $\mathbb{S}^n$ no es contráctil.
2. El Teorema del punto fijo de Brouwer.

Demostración. ($2 \Rightarrow 1$) Supongamos por reducción al absurdo que $\text{Id}_{\mathbb{S}^n}$ sí que es homótopa a una aplicación c_{x_0} constantemente igual a $\mathbf{x}_0$. Por la Proposición 9, existe una extensión continua $r : E^{n+1} \rightarrow \mathbb{S}^n$ tal que $r|_{\mathbb{S}^n} = \text{Id}_{\mathbb{S}^n}$. Se tiene que r es una retracción de E^{n+1} en $\mathbb{S}^n$.

Consideremos la aplicación $g : E^{n+1} \rightarrow E^{n+1}$ definida como $g(\mathbf{x}) = -r(\mathbf{x})$. Veamos que g no tiene puntos fijos, llegando a una contradicción con el Teorema del punto fijo de Brouwer. Si tenemos un punto fijo, es decir, si

$$\mathbf{x} = g(\mathbf{x}) = -r(\mathbf{x}),$$

necesariamente este punto $\mathbf{x}$ ha de estar en $\mathbb{S}^n$, por lo que

$$\mathbf{x} = -r|_{\mathbb{S}^n}(\mathbf{x}) = -\text{Id}_{\mathbb{S}^n}(\mathbf{x}) = -\mathbf{x}.$$

Esto solo es posible si $\mathbf{x} = \mathbf{0}$, pero $\mathbf{0} \notin \mathbb{S}^n$. Luego g es una aplicación continua de E^n en sí mismo que no tiene ningún punto fijo.

($1 \Rightarrow 2$) De nuevo, utilizaremos reducción al absurdo, suponiendo que tenemos una aplicación continua $f : E^{n+1} \rightarrow E^{n+1}$ que no tiene puntos fijos. Podemos construir la misma h que describimos anteriormente en la Figura 22. Observemos que $h|_{\mathbb{S}^n} = \text{Id}_{\mathbb{S}^n}$, ya que si $\mathbf{x} \in \mathbb{S}^n$ entonces el punto de corte de la semirrecta con $\mathbb{S}^n$ es el propio $\mathbf{x}$. Se tiene entonces que h es una extensión continua de $\text{Id}_{\mathbb{S}^n}$. Por la Proposición 9 se concluye que $\text{Id}_{\mathbb{S}^n}$ es homótopa a una aplicación constante, contradiciendo (1). $\square$

El primer enunciado es una consecuencia directa de la versión (5) del Teorema de Borsuk-Ulam, que nos decía que ninguna aplicación antipodal y continua de $\mathbb{S}^n$ en $\mathbb{S}^n$ es homótopa a una aplicación constante.

4.2. Lema de Sperner

El lema de Sperner fue demostrado por el matemático alemán Emanuel Sperner en 1928 [Spe28]. Vimos en la asignatura de Topología Algebraica que este lema implica el Teorema del punto fijo de Brouwer. En esta sección veremos que el Teorema de Borsuk-Ulam (en particular, la versión 8) implica directamente el lema de Sperner [NS13].

En primer lugar, es necesario que demos varias definiciones.

Definición 20. Denotamos el *símplice*

$$\Delta^n = \{(x_1, \dots, x_{n+1}) \in \mathbb{R}^{n+1} : x_i \geq 0, \sum_{i=1}^{n+1} x_i = 1\},$$

y sea T una triangulación de Δ^n . Un *etiquetado de Sperner* es una aplicación

$$l : V(T) \rightarrow \{1, 2, \dots, n+1\}$$

de forma que para todo $v = (v_1, \dots, v_{n+1}) \in V(T)$ entonces $l(v) \in \{i : v_i \neq 0\}$. Decimos que un *símplice* $\sigma \in T$ etiquetado de Sperner tiene *etiquetado pleno* si $l|_{V(\sigma)}$ es sobreyectiva.

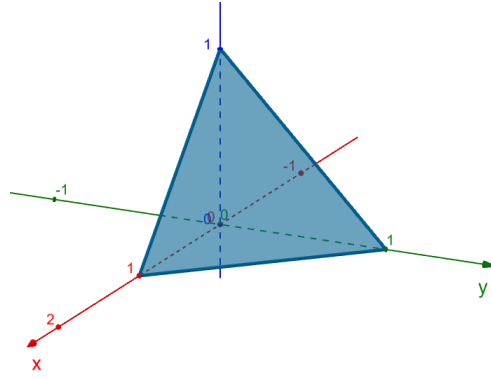


Figura 25: Δ^2 .

Una definición alternativa del etiquetado de Sperner es, dada una subdivisión de $\Delta^n = \langle v^1, \dots, v^{n+1} \rangle$ con una triangulación T , una aplicación

$$l : V(T) \rightarrow \{1, 2, \dots, n+1\}$$

tal que $l(v) \in \{i_1, \dots, i_k\}$ si $v \in \langle v^{i_1}, \dots, v^{i_k} \rangle$. Ambas definiciones son equivalentes: $v_i = 0$ si y solo si v está en el interior de una cara propia de Δ^n que no contiene a v^i .

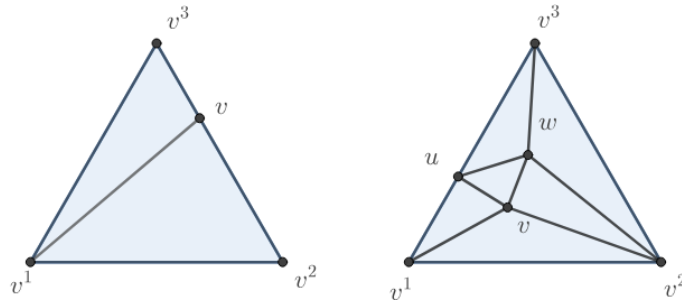


Figura 26: Ejemplos de subdivisiones de Δ^2 .

Fijémonos en la triangulación de la izquierda en la Figura 26. Para que la figura muestre un etiquetado de Sperner, cada v^i ha de estar etiquetado con i , y la etiqueta de v ha de ser 2 o 3. En ambos casos, uno de los dos 2-símplices tiene etiquetado pleno. En la triangulación de la derecha, como $u \in \langle v^1, v^3 \rangle$, entonces u ha de tener la etiqueta 1 o 3, mientras que las etiquetas de v y de w pueden ser cualesquiera en $\{1, 2, 3\}$.

El lema de Sperner es el siguiente:

Proposición 11. *Toda triangulación de Δ^n con un etiquetado de Sperner tiene un n -símplice σ tal que $l|_\sigma$ tiene etiquetado pleno.*

Demostración. Sea S una triangulación de Δ^n con un etiquetado de Sperner l , y sea $G = \{-1, 1\}^{n+1}$. Observemos que, si $g = (g_1, \dots, g_{n+1}) \in G$ y $v \in \Delta^n$ entonces $gv = (g_1v_1, \dots, g_{n+1}v_{n+1}) \in \hat{\mathbb{S}}^n$, donde g refleja v en las coordenadas tales que $g_i = -1$. Ocurre lo mismo para todo símplice $\sigma \in S$: si $\sigma = \langle a^1, \dots, a^m \rangle$, entonces $g\sigma = \langle ga^1, \dots, ga^m \rangle$ es la reflexión de σ en las coordenadas con $g_i = -1$. Se puede comprobar que los símplices $\{g\sigma : \sigma \in S, g \in G\}$ y sus caras son un complejo simplicial que nos da una triangulación T de $\hat{\mathbb{S}}^n$. Esta triangulación es simétrica por construcción: dado $\sigma \in T$, existe $g \in G$ tal que $\sigma = g\tau$ para cierto $\tau \in S$. Entonces, $-g \in G$ y $-\sigma = -g\tau \in T$.

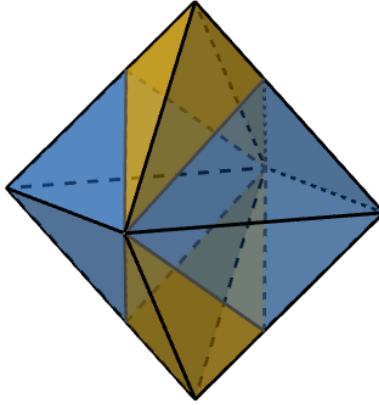


Figura 27: Una triangulación de $\hat{\mathbb{S}}^2$ construida como se ha descrito.

Extendamos el etiquetado l a un etiquetado L de T dado por

$$L(gv) = g_{l(v)}(-1)^{l(v)+1}l(v),$$

con $v \in S$. Observemos que $|L(gv)| = |g_{l(v)}| |(-1)^{l(v)+1}| |l(v)| = |l(v)|$.

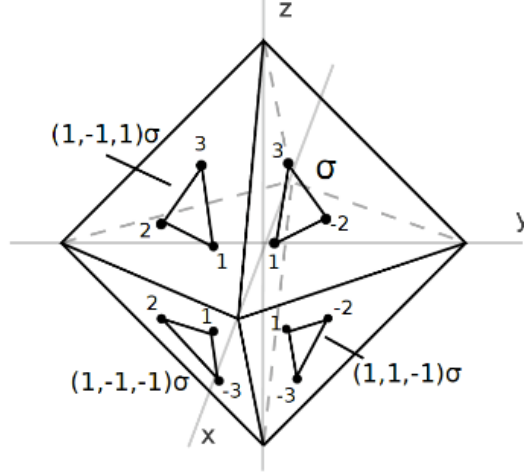


Figura 28: [NS13]: Un símplece alternado positivo $\sigma \in T$ que proviene de un símplece con etiquetado pleno en S , y los símpleces $g\sigma$ para $g = (1, 1, -1)$, $g = (1, -1, 1)$ y $g = (1, -1, -1)$, con sus etiquetados por L correspondientes.

Veamos que L está bien definido. Supongamos que un vértice de T se puede escribir como gv y como $\hat{g}\hat{v}$, con $g, \hat{g} \in G$ y $v, \hat{v} \in S$. Luego $g_i v_i = \hat{g}_i \hat{v}_i$ para todo i . Como g_i y $\hat{g}_i$ solo pueden valer 1 o -1 , entonces $|v_i| = |\hat{v}_i|$ para todo i , y como $v \in S$ solo puede ser que $v_i = -\hat{v}_i$ si $v_i = \hat{v}_i = 0$. En cualquier caso, $v_i = \hat{v}_i$ para todo i , y $v = \hat{v}$. Si $v_i \neq 0$, esto significa que $g_i = \hat{g}_i$. Como l es un etiquetado de Sperner, por definición se tiene que $l(v) \in \{i : v_i \neq 0\}$. Luego

$$L(gv) = g_{l(v)}(-1)^{l(v)+1}l(v) = \hat{g}_{l(v)}(-1)^{l(v)+1}l(v) = L(\hat{g}v) = L(\hat{g}\hat{v}).$$

Hemos construido T de forma que es una triangulación simétrica. El etiquetado L es antipodal:

$$L(-gv) = L((-g)v) = (-g_{l(v)})(-1)^{l(v)+1}l(v) = -(g_{l(v)}(-1)^{l(v)+1}l(v)) = -L(gv)$$

para todo $gv \in T$. Además, como l es un etiquetado de Sperner, todas sus etiquetas $\{1, 2, \dots, n+1\}$ son positivas. Esto significa que l no tiene bordes complementarios (ver Definición 16). Si $v, w \in S$, entonces $l(v) = l(w)$ si y solo si

$$L(gv) = g_{l(v)}(-1)^{l(v)+1}l(v) = g_{l(w)}(-1)^{l(w)+1}l(w) = L(gw).$$

Es decir, si hubiera dos vértices adyacentes gv y gw (con $g \in G$, $v, w \in S$) etiquetados por $L(gv) = j$ y $L(gw) = -j$, entonces $l(v) \neq l(w)$. Vimos antes que $|L(gv)| = |l(v)|$. Entonces,

$$j = |L(gv)| = l(v) \neq l(w) = |L(gw)| = j,$$

absurdo. Luego necesariamente no puede haber bordes complementarios en L . Estamos en condiciones de aplicar la versión 8 del Teorema de Borsuk-Ulam: existe un n -símplice $g\sigma$ de T alternado positivo.

Cada vértice v de σ cumple que $l(v) = |l(v)| = |L(gv)|$. Como $g\sigma \in T$ es alternado positivo por L , entonces σ tiene etiquetado pleno por l , como queríamos probar. $\square$

Una versión más fuerte del lema dice que, de hecho, el número de símlices con etiquetado pleno es impar [Spe28, Sección 3].

4.3. Juego de Hex

4.3.1. Historia y funcionamiento del juego

El juego de Hex recibe su nombre del tablero en el que se juega, formado por casillas hexagonales. Fue inventado por el danés Piet Hein en 1942, mientras buscaba una demostración para el teorema de los cuatro colores. Poco después, lo comercializó bajo el nombre de Con-Tac-Tix. Una idea similar se le ocurrió en 1948 a John Nash, que consideró en su lugar un tablero formado por cuadrados que se consideraban adyacentes si compartían un lado o una de las esquinas. Más información sobre sus orígenes se puede encontrar en [HR06].

Las reglas son simples. En un tablero de $n \times m$ hexágonos como el de la Figura 29, dos jugadores se turnan para colorear uno de ellos en cada turno, uno de ellos en rojo, y el otro en azul. El jugador rojo gana si consigue conectar mediante hexágonos rojos los lados superior e inferior. El jugador azul gana si consigue conectar los lados izquierdo y derecho con hexágonos azules.

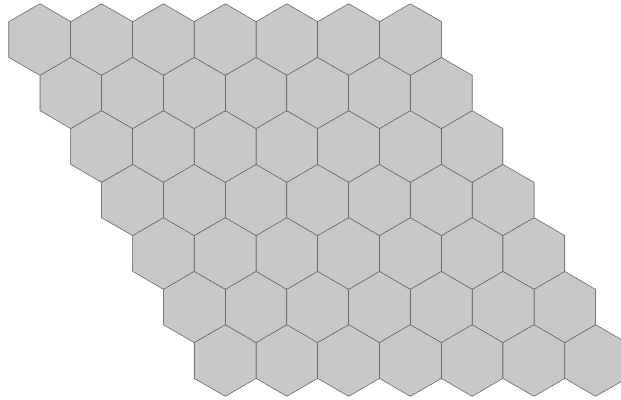


Figura 29: Tablero de Hex de dimensión 7x7

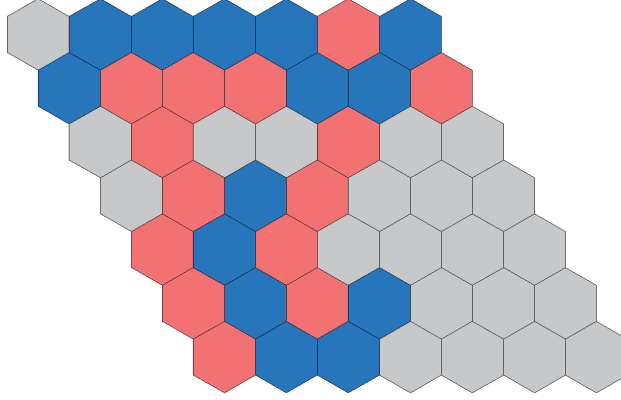


Figura 30: Juego de Hex en el que ha ganado el jugador azul.

Una forma equivalente de entender este tablero es el grafo de los puntos $(z_1, z_2) \in \mathbb{N}^2$ tales que $1 \leq z_1 \leq n$, $1 \leq z_2 \leq n$, donde dos v rtices $z = (z_1, z_2)$ y $z' = (z'_1, z'_2)$ est n conectados si se cumplen las condiciones:

1. $\|z - z'\|_\infty = 1$.
2. O bien $z_1 \leq z'_1$ y $z_2 \leq z'_2$, o bien $z'_1 \leq z_1$ y $z'_2 \leq z_2$.

A los v rtices del grafo los llamaremos casillas, y diremos que dos casillas son adyacentes si est n unidas por una arista.

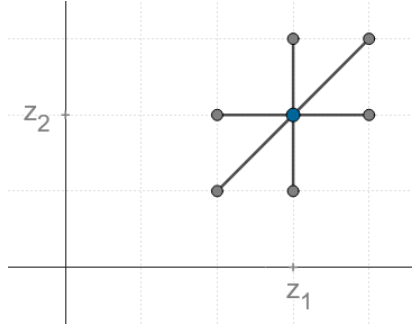


Figura 31: El v rtice (z_1, z_2) y sus v rtices adyacentes.

Llamamos $B_{n,m} = ([1, n] \times [1, m]) \cap \mathbb{N}^2$ al conjunto de v rtices del grafo, y decimos que $B_{n,m}$ es el tablero de Hex de dimensi n $n \times m$.

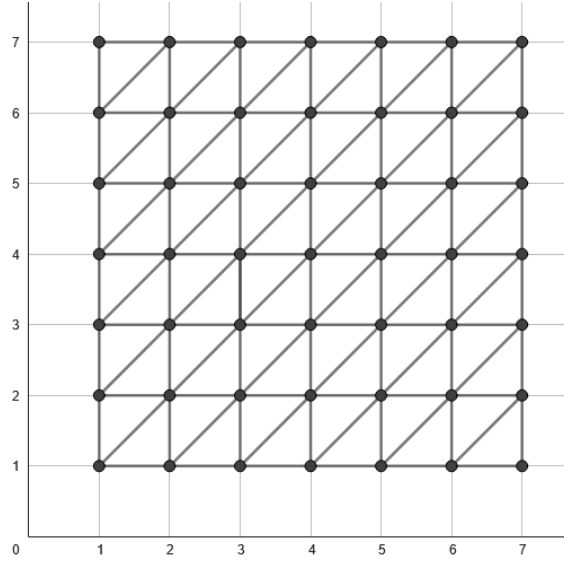


Figura 32: Tablero de Hex de dimensión 7x7 visto como un grafo.

Definición 21. Denotaremos por O al conjunto de vértices con primera coordenada igual a 1, E al conjunto de los vértices con primera coordenada igual a n , S al conjunto de vértices que tengan la segunda coordenada igual a 1 y N al conjunto de los que tengan la segunda coordenada igual a m .

Diremos que un conjunto de casillas $\{z^1, \dots, z^r\}$ es un camino si todas pertenecen al mismo jugador y z^i y z^{i+1} son adyacentes para $i = 1, \dots, r-1$. Diremos que un camino es ganador para R si contiene por lo menos un vértice de N y uno de S . Análogamente, un camino es ganador para A si contiene por lo menos un vértice de E y uno de O .

Un conjunto $B \subset \mathbb{N}^2$ decimos que es conexo si no existen subconjuntos B_1 y B_2 tales que $B = B_1 \cup B_2$, $B_1 \neq \emptyset$, $B_2 \neq \emptyset$ y, para todos $z \in B_1$, $z' \in B_2$, z y z' no sean adyacentes.

Teorema 3. (Teorema de Hex) Consideramos el tablero de Hex de dimensión $n \times n$, y supongamos que está recubierto por dos conjuntos disjuntos A y R . Entonces, o bien A contiene un camino ganador para el jugador A , o bien R contiene un camino ganador para el jugador R .

Es fácil entender, intuitivamente, que este hecho es cierto, y que el juego no puede acabar en empate (esto también se cumple en un tablero de dimensión $n \times m$). Volvamos a pensar en el tablero de casillas hexagonales y fijémonos en la Figura 33: la existencia de un camino ganador para A impide que se pueda crear uno ganador para R , incluso si todas las demás casillas pertenecieran a R .

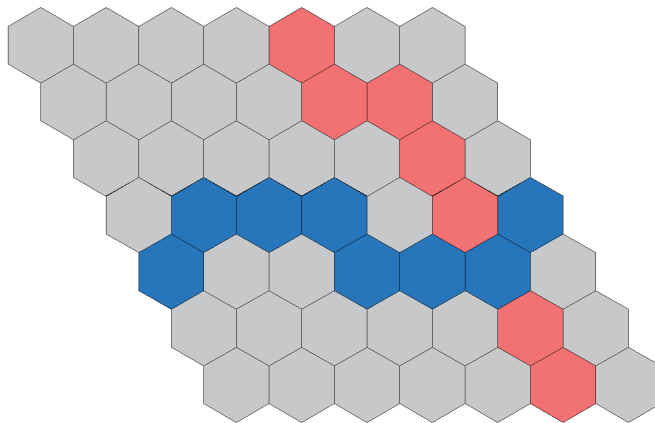


Figura 33: El camino ganador para A “impide” que haya un camino ganador para R .

Como no puede acabar en empate, podemos deducir que en un tablero $n \times n$ existe una estrategia ganadora para el primer jugador [HJ63], [Bec08]. Supongamos, por reducción al absurdo, que el primer jugador no tiene estrategia ganadora. Por el Teorema de Zermelo, ha de tenerla el segundo. Sin embargo, si esto ocurre, el primer jugador también la tiene: basta con que haga un primer movimiento al azar (convirtiéndose, en cierto modo, en el segundo jugador) y después continúe con la estrategia ganadora del segundo jugador. Esto se contradice con la hipótesis de que el segundo jugador tiene estrategia ganadora, concluyéndose que es el primero el que la tiene. A esto se le conoce como *robo de estrategia*.

Es importante señalar que conocer la existencia de una estrategia ganadora no equivale a conocer la estrategia en sí. De hecho, solo se conoce la estrategia ganadora en tableros de dimensión menor o igual que 9×9 . En la Figura 34 se representa el árbol de juego para un tablero 2×2 en el que el primer jugador es el azul. Este árbol nos da una estrategia ganadora para A , cuyo primer movimiento consiste en empezar por la casilla superior derecha o la inferior izquierda.

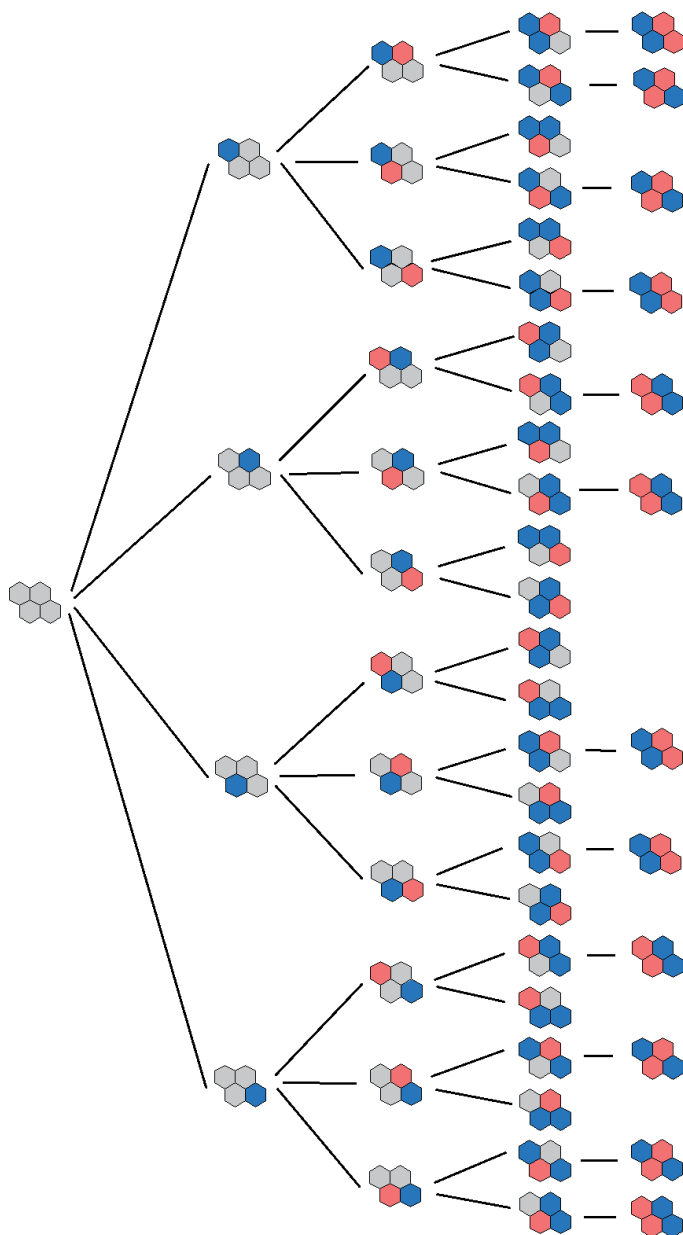


Figura 34: Árbol de juego en un tablero 2×2 .

Veamos algunos puzles sencillos basados en el juego de Hex para familiarizarnos con el juego [Hex].

1. Figura 35. Juega *A* y gana en un movimiento. Este puzle es una modificación que hace más sencillo uno propuesto por Piet Hein.

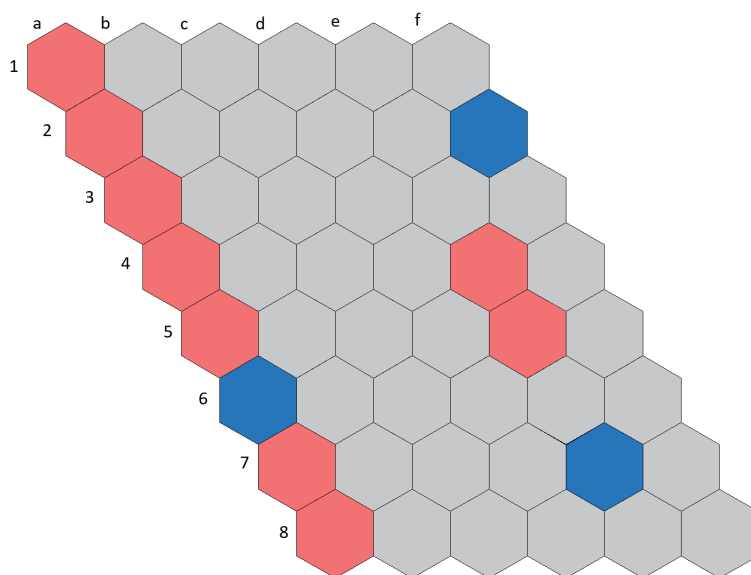


Figura 35: Modificación de un puzle de Piet Hein.

Solución: *A* tiene que elegir la casilla *c5*. *R* no podrá evitar que la conecte con *a6*, porque si *R* elige la casilla *b5* entonces *A* elegirá la *b6* y viceversa. Si *R* intenta evitar que se conecte a *f2*, entonces *A* elegirá *d6*, donde cada movimiento que pueda hacer *R* para evitar que conecte *c5* y *e7* puede ser “contrarrestado” por *A*. Ocurre lo mismo si *R* intenta evitar que se conecte a *e7*, en este caso tomando *A* la casilla *d3*.

2. Figura 36. Juega *A* y gana en un movimiento.

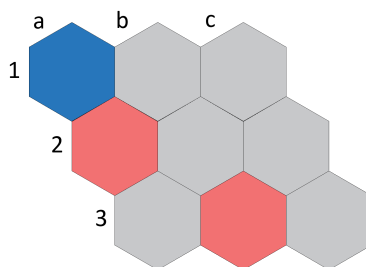


Figura 36: Puzle de Piet Hein.

Solución: A perdería si eligiera las casillas b1, a3 o c3, porque R elegiría c1 y en su siguiente turno podría ganar eligiendo b2 o c2. Del mismo modo, A pierde eligiendo c2, porque R escogería b2 y en su siguiente turno conseguiría un camino con b1 o c1; y A pierde si elige c1 porque R elegiría b1, y en su siguiente turno b2 o a3. Solo es razonable que A elija b2. Desde esta posición se puede ver que para ganar basta con “interrumpir” los intentos de R de formar un camino ganador.

En las siguientes dos secciones demostraremos siguiendo [Gal79] la equivalencia entre el Teorema del punto fijo de Brouwer con el Teorema de Hex.

4.3.2. Hex implica Brouwer

Veamos que el Teorema de Hex implica el Teorema del punto fijo de Brouwer para E^2 .

Demostración. Consideramos $\mathbb{R}^2$ dotado de la norma del supremo $\|\cdot\|_\infty$. Sea $f : I \times I \rightarrow I \times I$ una función continua, con $f = (f_1, f_2)$. Como $I \times I$ es compacto, sabemos que f es uniformemente continua. Es decir, dado $\varepsilon > 0$ existe $\delta_\varepsilon > 0$ tal que $\|f(\mathbf{x}) - f(\mathbf{y})\|_\infty < \varepsilon$ para todo $\mathbf{x}, \mathbf{y} \in I \times I$ tales que $\|\mathbf{x} - \mathbf{y}\|_\infty < \delta_\varepsilon$, y además podemos tomar este δ_ε de forma que $\delta_\varepsilon < \varepsilon$. Consideraremos un tablero de Hex (como en la Figura 32) de dimensión $k \times k$ tal que $\frac{1}{k} < \delta_\varepsilon$. Vamos a tomar los conjuntos A^+, A^-, R^+, R^- de forma que un vértice $z = (z_1, z_2) \in B_{k,k}$ pertenezca a uno de ellos si $\frac{z}{k} = (\frac{z_1}{k}, \frac{z_2}{k})$ es trasladado por f al menos una distancia ε hacia la derecha, izquierda, arriba o abajo, respectivamente. Esto es:

$$\begin{aligned} A^+ &= \{z \in B_{k,k} : f_1(\frac{z}{k}) - \frac{z_1}{k} > \varepsilon\} & R^+ &= \{z \in B_{k,k} : f_2(\frac{z}{k}) - \frac{z_2}{k} > \varepsilon\} \\ A^- &= \{z \in B_{k,k} : \frac{z_1}{k} - f_1(\frac{z}{k}) > \varepsilon\} & R^- &= \{z \in B_{k,k} : \frac{z_2}{k} - f_2(\frac{z}{k}) > \varepsilon\} \end{aligned}$$

Lo que estamos haciendo es colocar un tablero de Hex de dimensión suficientemente grande dentro de $I \times I$.

Sean $R = R^+ \cup R^-$ y $A = A^+ \cup A^-$. Observemos que un vértice z no pertenece a $A \cup R$ si y solo si $\|f(\frac{z}{k}) - \frac{z}{k}\|_\infty \leq \varepsilon$, luego no necesariamente recubren todo el tablero. Si conseguimos probar que, de hecho, $A \cup R$ no recubre $B_{k,k}$ hemos terminado: en este caso, para todo $\varepsilon > 0$ existe un vértice z_ε tal que $\|f(\frac{z_\varepsilon}{k}) - \frac{z_\varepsilon}{k}\|_\infty \leq \varepsilon$. Sea $x_\varepsilon = \frac{z_\varepsilon}{k}$. La sucesión $\{x_{\frac{1}{n}}\}_{n=1}^\infty$ en el compacto $I \times I$ tiene una subsucesión convergente hacia un punto x^* . Podemos suponer directamente que

$$x_{\frac{1}{n}} \xrightarrow{n \rightarrow \infty} x^*.$$

Por la continuidad de f y de la norma,

$$\|f(x^*) - x^*\|_\infty = \lim_{n \rightarrow \infty} \|f(x_{1/n}) - x_{1/n}\|_\infty \leq \lim_{n \rightarrow \infty} \frac{1}{n} = 0,$$

luego $f(x^*) = x^*$, con lo que x^* es un punto fijo.

Queremos probar que $A \cup R$ no recubre $B_{k,k}$. Veamos primero que si $z \in A^+$ y $z' \in A^-$ entonces z y z' no son adyacentes. Supongamos que lo son. Entonces, $f_1(\frac{z}{k}) - \frac{z_1}{k} > \varepsilon$ y $\frac{z'_1}{k} - f_1(\frac{z'}{k}) > \varepsilon$. Sumando ambas expresiones, se tiene que

$$f_1\left(\frac{z}{k}\right) - f_1\left(\frac{z'}{k}\right) + \frac{z'_1}{k} - \frac{z_1}{k} > 2\varepsilon$$

Por otro lado, como hemos escogido z y z' adyacentes,

$$\frac{z'_1}{k} - \frac{z_1}{k} \leq \left\| \frac{z'}{k} - \frac{z}{k} \right\|_\infty = \frac{1}{k} < \delta_\varepsilon < \varepsilon,$$

luego $\frac{z_1}{k} - \frac{z'_1}{k} > -\varepsilon$. Sumando esta expresión con la obtenida en el párrafo anterior, llegamos a que $f_1(\frac{z}{k}) - f_1(\frac{z'}{k}) > \varepsilon$.

Como z y z' son adyacentes, entonces $\|\frac{z}{k} - \frac{z'}{k}\|_\infty = \frac{1}{k} < \delta_\varepsilon$. Pero

$$\left\| f\left(\frac{z}{k}\right) - f\left(\frac{z'}{k}\right) \right\|_\infty \geq f_j\left(\frac{z}{k}\right) - f_j\left(\frac{z'}{k}\right) > \varepsilon,$$

contradiendo que f sea uniformemente continua. Luego z y z' no son adyacentes. Análogamente, si $z \in R^+$ y $z' \in R^-$ entonces z y z' no son adyacentes. Esto se cumple para cualesquiera $z \in A^+$, $z' \in A^-$ y para cualesquiera $z \in R^+$, $z' \in R^-$, de lo que se concluye que los conjuntos R^+ y R^- no son contiguos, y que A^+ y A^- tampoco (donde decimos que dos subconjuntos de vértices B y C de un grafo son contiguos si existen $b \in B$, $c \in C$ tales que b y c sean adyacentes).

Veamos ahora que $A^+ \cap E = \emptyset$. Sea Q un conjunto conexo en A . Como A^+ y A^- no son contiguos, para que sea conexo es necesario que Q esté contenido en uno de ellos. Si A^+ tuviera un punto de E , es decir, de la forma (a, k) , entonces f estaría moviendo $(a, k)/k = (a/k, 1)$ al menos una distancia ε hacia la derecha, de modo que la segunda coordenada sería mayor que 1. Esto es imposible, porque f tiene llegada en $I \times I$. Luego A^+ no contiene ningún punto de E . Ocurre lo mismo para A^- , que no contiene ningún punto de O . Es decir, ningún camino contenido en A puede contener a la vez puntos de E y de O . Análogamente se concluye que ningún camino contenido en R contiene a la vez puntos de N y S . Por tanto, no hay caminos ganadores, y por el Teorema de Hex no puede ser que A y R recubran $B_{k,k}$. $\square$

4.3.3. Brouwer implica Hex

Veamos ahora que el Teorema del punto fijo de Brouwer implica el Teorema de Hex.

Demostración. Supongamos, por reducción al absurdo, que el tablero de Hex de dimensión $k \times k$ (de nuevo, visto como en la Figura 32) tiene una partición

en dos conjuntos de vértices A y R de forma que no haya un camino en A que una O y E ni un camino en R que una N y S . Sean $\tilde{O}$ el conjunto de vértices de A tales que existe un camino en A que los conecta a O , $\tilde{E}$ los demás puntos de A (es decir, $\tilde{E} = A \setminus \tilde{O}$), $\tilde{S}$ los vértices de R tales que existe un camino en R que los conecta a S y $\tilde{N}$ los demás elementos de R (es decir, $\tilde{N} = R \setminus \tilde{S}$).

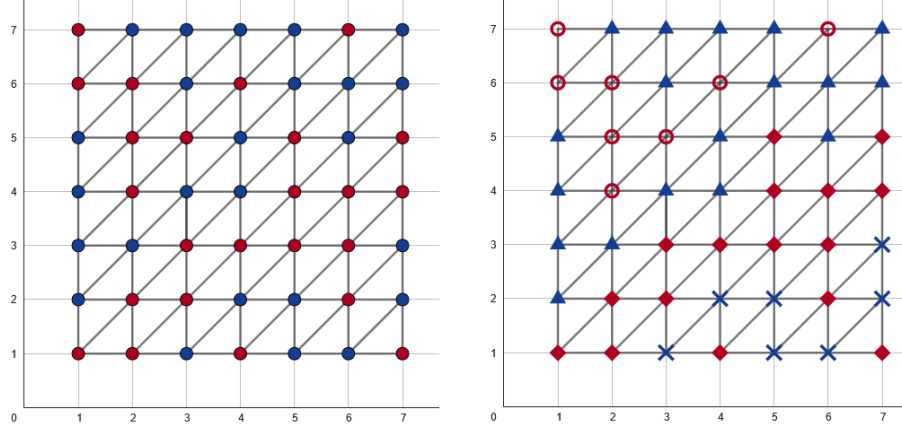


Figura 37: Dada la partición de la izquierda, en la imagen de la derecha se han representado los conjuntos $\tilde{O}$ con triángulos azules, $\tilde{E}$ con cruces azules, $\tilde{S}$ con cuadrados rojos y $\tilde{N}$ con círculos rojos.

Sea $f : B_{k,k} \rightarrow B_{k,k}$, dada por

$$f(z) = \begin{cases} z + (1, 0) & \text{si } z \in \tilde{O} \\ z - (1, 0) & \text{si } z \in \tilde{E} \\ z + (0, 1) & \text{si } z \in \tilde{S} \\ z - (0, 1) & \text{si } z \in \tilde{N} \end{cases}$$

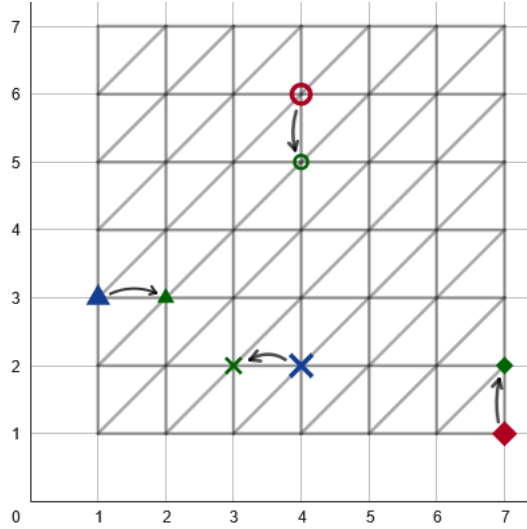


Figura 38: En verde, las imágenes por f de cuatro vértices de la partición anterior.

Veamos que está bien definida, es decir, que $f(z) \in B_{k,k}$. Sea $z \in \tilde{S}$. Estamos suponiendo que no existe un camino ganador en R , y sabemos que z forma parte de un camino que lo une con S , luego en ese camino no puede haber ningún vértice de N . En particular, z no puede estar en N , luego al sumar 1 a su segunda coordenada seguimos teniendo un vértice de $B_{k,k}$ (no nos salimos del tablero). El razonamiento para un vértice de $\tilde{O}$ es análogo.

Si $z \in \tilde{N}$ sabemos que $z \notin S$, porque si no tendríamos un camino ganador para R , en contra de nuestra hipótesis. Como $z \notin S$, la segunda coordenada es un entero estrictamente mayor que 1, luego al restar $z - (0, 1)$ tendremos un vértice que sigue estando en $B_{k,k}$. De nuevo, el razonamiento para $\tilde{E}$ es análogo.

La representación del tablero que estamos considerando es una triangulación de $[1, k]^2$. Podemos extender la aplicación $i \circ f : B_{k,k} \rightarrow [1, k]^2$ (donde $i : B_{k,k} \rightarrow [1, k]^2$ es la inclusión) definida en los vértices a una aplicación afín

$$\hat{f} : [1, k]^2 \rightarrow [1, k]^2$$

como en la Proposición 8. Veamos que $\hat{f}$ no tiene puntos fijos, contradiciendo el Teorema del punto fijo de Brouwer.

La aplicación $\hat{f}$ tiene un punto fijo si y solo si existe $x = \lambda_1 v_1 + \lambda_2 v_2 + \lambda_3 v_3$ (siendo v_1, v_2 y v_3 los vértices de un símple de dimensión 2 de la triangulación) tal que

$$\hat{f}(x) = \lambda_1 f(v_1) + \lambda_2 f(v_2) + \lambda_3 f(v_3) = \lambda_1 v_1 + \lambda_2 v_2 + \lambda_3 v_3 = x.$$

En los vértices, $\hat{f}$ coincide con f , luego $\hat{f}(v_i) = v_i + w_i$, donde w_i es $(1, 0)$, $(-1, 0)$, $(0, 1)$ o $(0, -1)$, según si v_i pertenece a $\tilde{O}$, $\tilde{E}$, $\tilde{S}$ o $\tilde{N}$ respectivamente. Es decir, $\hat{f}$ tiene un punto fijo si y solo si

$$\lambda_1(v_1 + w_1) + \lambda_2(v_2 + w_2) + \lambda_3(v_3 + w_3) = \lambda_1 v_1 + \lambda_2 v_2 + \lambda_3 v_3.$$

En este caso, $\lambda_1 w_1 + \lambda_2 w_2 + \lambda_3 w_3 = 0$, o lo que es equivalente, 0 pertenece a $\langle w_1, w_2, w_3 \rangle$.

Veamos finalmente que 0 no está en la envolvente convexa de los vectores que desplazan los vértices de cualquiera de los triángulos. Nótese que $\tilde{O}$ y $\tilde{E}$ no tienen puntos adyacentes, es decir, ningún triángulo contiene a la vez un vértice de $\tilde{O}$ y uno de $\tilde{E}$. O, equivalentemente, ningún triángulo tiene un vértice que se mueva $(1, 0)$ y otro que se mueva $(-1, 0)$ al tomar sus imágenes por f . Lo mismo ocurre con $\tilde{S}$ y $\tilde{N}$ en las direcciones $(0, 1)$ y $(0, -1)$. Por el principio del palomar, los tres vértices son desplazados por, como mucho, dos vectores distintos de $\{(1, 0), (-1, 0), (0, 1), (0, -1)\}$, y esos vectores no son opuestos, de modo que han de estar en el mismo cuadrante. Esto implica que 0 no está en la envolvente convexa de estos vectores.

Luego $\hat{f}$ no tiene puntos fijos. Es decir, hemos encontrado una aplicación continua de $I \times I$ en sí mismo que no tiene puntos fijos, contradiciendo el Teorema del punto fijo de Brouwer. $\square$

4.3.4. El juego de Hex n -dimensional

El juego de Hex se puede generalizar a n dimensiones y n jugadores. Para ello, consideramos $B_{k_1, \dots, k_n}$ el grafo de los puntos $(z_1, \dots, z_n) \in \mathbb{N}^n$ tales que $1 \leq i \leq k_i$ para todo $i \in \{1, \dots, n\}$, donde los vértices $z = (z_1, \dots, z_n)$ y $z' = (z'_1, \dots, z'_n)$ están conectados si se cumplen las condiciones:

1. $\|z - z'\|_\infty = 1$.
2. O bien $z_i \leq z'_i$ para todo i , o bien $z'_i \leq z_i$ para todo i .

Consideremos un tablero n -dimensional $B = B_{k_1, \dots, k_n}$. Ahora, el Teorema de Hex nos dice:

Teorema 4. (Teorema de Hex n -dimensional) Para cada i , definimos

$$H_i^- = \{z \in B : z_i = 1\},$$

$$H_i^+ = \{z \in B : z_i = k_i\}.$$

Para toda aplicación $L : B \rightarrow \{1, \dots, n\}$ existe $i \in \{1, \dots, n\}$ tal que $L^{-1}(i)$ contiene un conjunto conexo que une H_i^- y H_i^+ .

El Teorema de Hex que habíamos visto anteriormente era equivalente al Teorema del punto fijo de Brouwer para E^2 . También es cierto que el Teorema

de Hex n -dimensional es equivalente al Teorema del punto fijo de Brouwer en E^n . La demostración que hemos dado para el caso $n = 2$ puede adaptarse sin demasiada dificultad al caso general, y no la desarrollaremos en este trabajo por brevedad. La prueba puede encontrarse en [Gal79].

4.4. Problema del collar

Imaginemos que queremos repartir un collar con k tipos distintos de joyas entre m ladrones, haciendo el mínimo número de cortes posible (sin incluir el cierre del collar), de forma que cada ladrón acabe con el mismo número de joyas de cada tipo, asumiendo que el número de joyas de cada tipo es par. Este problema lo resolvieron para $m = 2$ C. H. Goldberg y D. B. West en [GW85], usando inducción sobre el número de joyas. Un año más tarde, N. Alon y D. B. West dieron en [AW86] una demostración (también para el caso $m = 2$) más breve utilizando el Teorema de Borsuk-Ulam, y generalizaron el problema para todo $m > 2$. En esta sección veremos la demostración del problema para $m = 2$ utilizando el Teorema de Borsuk-Ulam.

Para la demostración, trataremos con una versión continua del problema.

Definición 22. Llamamos k -coloraciones a las aplicaciones $c : I \rightarrow \{1, \dots, k\}$ que asignan a cada punto $x \in I$ un único color i , con $1 \leq i \leq k$, de forma que la unión de puntos de color i es medible para todo i .

Dada una k -coloración, llamaremos bisección de tamaño r a una secuencia de números $0 = y_0 < y_1 < \dots < y_r < y_{r+1} = 1$ tal que exista un subconjunto $J \subset \{0, 1, \dots, r\}$ que haga que $\cup\{[y_{i-1}, y_i] : i \in J\}$ contenga exactamente la mitad de cada color.

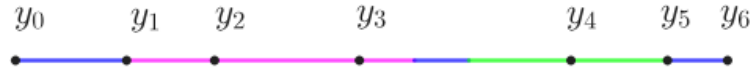


Figura 39: Bisección de tamaño 5 de una 3-coloración.

Vamos que dos ladrones pueden repartirse de forma equitativa los k colores de un intervalo realizando como mucho k cortes. Más precisamente, tenemos el siguiente resultado:

Teorema 5. Toda k -coloración tiene una bisección de tamaño como mucho k .

Demostración. Dada una k -coloración, queremos construir una función continua y antipodal $f : \mathbb{S}^k \rightarrow \mathbb{R}^k$.

Sea $\mathbf{x} = (x_1, \dots, x_{k+1}) \in \hat{\mathbb{S}}^k$. Sea $z_0(\mathbf{x}) = 0$, y definimos

$$z_i(\mathbf{x}) = \sum_{j=1}^i |x_j|,$$

con $1 \leq i \leq k$. Observemos que $z_k(\mathbf{x}) = 1$. De esta forma, un punto de $\hat{\mathbb{S}}^k$ representa una posible división del intervalo $[0, 1]$ en $k+1$ segmentos (algunos podrían ser vacíos). Sea $m_j(i, \mathbf{x})$ la medida del color j en el segmento $[z_{i-1}(\mathbf{x}), z_i(\mathbf{x})]$, con $j = 1, \dots, k$. Definimos $f_j : \mathbb{S}^k \rightarrow \mathbb{R}$ por

$$f_j(\mathbf{x}) = \sum_{i=1}^{k+1} \text{sign}(x_i) m_j(i, \mathbf{x}).$$

Cada f_j indica la cantidad del color j de diferencia entre los ladrones si repartimos a uno los intervalos $[z_{i-1}(\mathbf{x}), z_i(\mathbf{x})]$ tales que $x_i < 0$ y al otro los intervalos $[z_{i-1}(\mathbf{x}), z_i(\mathbf{x})]$ tales que $x_i > 0$ (en caso de que $x_i = 0$, el intervalo es vacío). Finalmente, podemos definir f como

$$f(\mathbf{x}) = (f_1(\mathbf{x}), f_2(\mathbf{x}), \dots, f_k(\mathbf{x})).$$

De esta forma, si encontramos un punto $\mathbf{x} \in \hat{\mathbb{S}}^k$ tal que $f(\mathbf{x}) = \mathbf{0}$, habremos encontrado una partición de I tal que, al repartir los intervalos $[z_{i-1}(\mathbf{x}), z_i(\mathbf{x})]$ según el signo de cada x_i , ambos ladrones tengan la misma cantidad de todos los colores.

Veamos que, fijado i , la aplicación $m_j(i, \mathbf{x})$ es continua:

Sea $\varepsilon > 0$. Sabemos que z_i y z_{i-1} son continuas, luego existe $\delta > 0$ tal que, si $|x - x_0| < \delta$ entonces $|z_i(x) - z_i(x_0)| < \varepsilon$, $|z_{i-1}(x) - z_{i-1}(x_0)| < \varepsilon$. Denotemos por $\tilde{m}_j(A)$ la medida del color j en un conjunto A . Para ε suficientemente pequeño y $|x - x_0| < \delta$, hay dos posibilidades: que uno de los intervalos

$$[z_{i-1}(x), z_i(x)], [z_{i-1}(x_0), z_i(x_0)]$$

esté contenido en el otro, o que se solapen sin contenerse. En el primer caso, podemos suponer sin pérdida de generalidad que $[z_{i-1}(x_0), z_i(x_0)] \subset [z_{i-1}(x), z_i(x)]$, y tenemos que

$$\begin{aligned} |m_j(i, x) - m_j(i, x_0)| &= |\tilde{m}_j([z_{i-1}(x), z_{i-1}(x_0)) \cup (z_i(x_0), z_i(x)])| \\ &\leq |z_i(x) - z_i(x_0)| + |z_{i-1}(x) - z_{i-1}(x_0)| < 2\varepsilon. \end{aligned}$$

En el segundo caso, suponiendo que $z_{i-1}(x_0) \leq z_{i-1}(x)$,

$$\begin{aligned} |m_j(i, x) - m_j(i, x_0)| &= |\tilde{m}_j([z_{i-1}(x_0), z_{i-1}(x)]) - \tilde{m}_j([z_i(x_0), z_i(x)])| \\ &\leq |z_i(x) - z_i(x_0)| + |z_{i-1}(x) - z_{i-1}(x_0)| < 2\varepsilon. \end{aligned}$$

Luego $m_j(i, \mathbf{x})$ es continua. La función $sign$ no es continua en 0, pero observemos que si $x_i = 0$ entonces

$$z_i(\mathbf{x}) = z_{i-1}(\mathbf{x}) + 0^2 = z_{i-1}(\mathbf{x}),$$

luego $[z_{i-1}(\mathbf{x}), z_i(\mathbf{x})] = \emptyset$. Se tiene entonces que $m_j(i, \mathbf{x}) = 0$, y por tanto $sign(x_i)m_j(i, \mathbf{x}) = 0$. Luego $sign(x_i)m_j(i, \mathbf{x})$ es continua. La aplicación f es continua si y solo si lo es cada f_j , y f_j es una suma finita de funciones continuas.

Veamos que f es antipodal. Como

$$z_i(\mathbf{x}) = \sum_{j=1}^i x_j^2 = \sum_{j=1}^i (-x_j)^2 = z_i(-\mathbf{x}),$$

deducimos que $[z_{i-1}(\mathbf{x}), z_i(\mathbf{x})] = [z_{i-1}(-\mathbf{x}), z_i(-\mathbf{x})]$, luego

$$m_j(i, \mathbf{x}) = m_j(i, -\mathbf{x}).$$

En consecuencia,

$$f_j(\mathbf{x}) = \sum_{i=1}^{k+1} sign(x_i)m_j(i, \mathbf{x}) = - \sum_{i=1}^{k+1} sign(-x_i)m_j(i, \mathbf{x}) = -f_j(-\mathbf{x}),$$

y $f(\mathbf{x}) = -f(-\mathbf{x})$.

Aplicando la versión 2 del Teorema de Borsuk-Ulam, existe $\hat{\mathbf{x}} \in \hat{\mathbb{S}}^k$ tal que $f(\hat{\mathbf{x}}) = \mathbf{0}$. Se tiene que para cada j se cumple $f_j(\hat{\mathbf{x}}) = 0$, es decir,

$$\sum_{\substack{sign(\hat{x}_i)=+1 \\ i=1}}^{k+1} m_j(i, \hat{\mathbf{x}}) = \sum_{\substack{sign(\hat{x}_l)=-1 \\ l=1}}^{k+1} m_j(l, \hat{\mathbf{x}}).$$

Esto significa que, si damos al primer ladrón los intervalos $[z_{i-1}, z_i]$ tales que $sign(x_i) = +1$, y al segundo los $[z_{i-1}, z_i]$ tales que $sign(x_i) = -1$, ambos tendrán la misma cantidad del color j para todo j . Es decir, cortando I en los puntos $z_1, \dots, z_k$, obtenemos una bisección de tamaño k . $\square$

Sabiendo esto, resolver el problema discreto es sencillo. Si tenemos un total de $2n$ joyas, con $2a_i$ de cada color i , $1 \leq i \leq k$, basta con colorear el segmento $[\frac{j-1}{2n}, \frac{j}{2n}]$ del color de la j -ésima joya para poder aplicar el resultado anterior. Es posible que la bisección que encontremos no corte al intervalo en los puntos $\frac{j}{2n}$, sino que divida alguna joya en varios trozos. Sin embargo, si tenemos un corte de esta forma, como el número de joyas que recibe cada ladrón ha de ser $a_i \in \mathbb{Z}$, ha de haber algún otro corte que divida una joya. Desplazamos estos dos cortes de forma que no cambie el total de cada ladrón, hasta que uno de ellos alcance uno de los extremos del segmento al que pertenece.

Haciendo inducción en el número de cortes que dividen alguna joya de color i , y repitiendo el proceso para cada color, podemos conseguir una bisección de

tamaño como mucho k cuyos cortes estén únicamente en puntos de la forma $\frac{j}{2n}$. Esto resuelve el problema del collar.

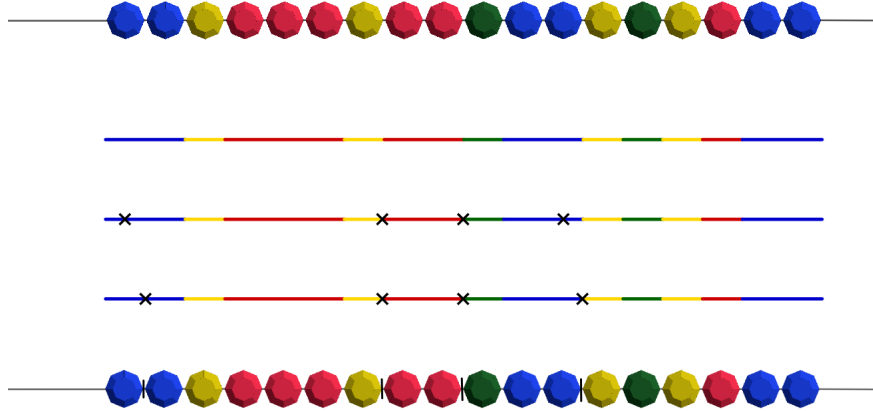


Figura 40: Dado un collar con joyas de 4 colores, tomamos la 4-coloración correspondiente y encontramos una bisección. Ajustamos los extremos de los intervalos hasta que todos ellos sean de la forma $\frac{j}{2n}$, y repartimos el collar con esos cortes.

Referencias

- [Arm83] M. A. Armstrong. “Basic Topology”. En: (1983). DOI: 10.1007/978-1-4757-1793-8. URL: <http://link.springer.com/10.1007/978-1-4757-1793-8>.
- [AW86] Noga Alon y Douglas B. West. “The Borsuk-Ulam Theorem and Bisection of Necklaces”. En: *Proceedings of the American Mathematical Society* 98 (4 1986). ISSN: 00029939. DOI: 10.2307/2045739.
- [Bec08] József Beck. “Combinatorial Games: Tic-Tac-Toe Theory”. En: (ene. de 2008). DOI: 10.1017/CB09780511735202.
- [Gal79] David Gale. “The Game of Hex and the Brouwer Fixed-Point Theorem”. En: *The American Mathematical Monthly* 86 (10 1979). ISSN: 00029890. DOI: 10.2307/2320146.
- [GH18] Marvin J. Greenberg y John R. Harper. *Algebraic topology: A first course*. 2018. DOI: 10.1201/9780429502408.
- [Gra03] Andrzej Granas. *Fixed point theory*. Ed. por James Dugundji. Springer, 2003. ISBN: 0-387-21593-X. DOI: 10.1007/978-0-387-21593-8.
- [Gui] Craig R Guilbault. “AN ELEMENTARY DEDUCTION OF THE TOPOLOGICAL RADON THEOREM FROM BORSUK-ULAM”. En: ().
- [GW85] Charles H Goldberg y Douglas B West. “Bisection of circle colorings*”. En: *SIAM J. ALG. DISC. METH* 6 (1 1985). URL: <http://www.siam.org/journals/ojsa.php>.
- [Hat02] Allen Hatcher. *Algebraic topology*. Cambridge: Cambridge University Press, 2002, págs. xii+544. ISBN: 0-521-79160-X; 0-521-79540-0.
- [Hex] HexWiki. *Piet Hein’s puzzles*. URL: https://www.hexwiki.net/index.php/Piet_Hein%27s_puzzles#Puzzles.
- [HJ63] A. W. Hales y R. I. Jewett. “Regularity and Positional Games”. En: *Transactions of the American Mathematical Society* 106 (2 feb. de 1963), pág. 222. ISSN: 00029947. DOI: 10.2307/1993764.
- [HR06] Ryan B. Hayward y Jack van Rijswijk. “Hex and combinatorics”. En: *Discrete Mathematics* 306 (19-20 oct. de 2006), págs. 2515-2528. ISSN: 0012-365X. DOI: 10.1016/J.DISC.2006.01.029.
- [Kaw25] Hidefumi Kawasaki. “AN APPLICATION OF BORSUK-ULAM’S THEOREM TO NONLINEAR PROGRAMMING”. En: *Journal of the Operations Research Society of Japan* 68 (ene. de 2025), págs. 21-29. DOI: 10.15807/jorsj.68.21.
- [Lov78] L. Lovász. “Kneser’s conjecture, chromatic number, and homotopy”. En: *Journal of Combinatorial Theory, Series A* 25.3 (nov. de 1978), págs. 319-324. ISSN: 0097-3165. URL: <https://www.sciencedirect.com/science/article/pii/0097316578900225>.

- [LS35] L. Lusternik y L. Schnirelmann. “Méthodes Topologiques dans les problèmes variationnels. I. Espaces à un nombre fini de dimensions”. En: *The Mathematical Gazette* 19.234 (1935), págs. 235-236. ISSN: 0025-5572. DOI: 10.2307/3605899. URL: <https://doi.org/10.2307/3605899>.
- [Mas91] William S. Massey. “A Basic Course in Algebraic Topology”. En: 127 (1991). DOI: 10.1007/978-1-4939-9063-4. URL: <http://link.springer.com/10.1007/978-1-4939-9063-4>.
- [Mat03] Jiri. Matousek. *Using the Borsuk-Ulam Theorem : Lectures on Topological Methods in Combinatorics and Geometry*. 1st ed. 2003. Springer Berlin Heidelberg, 2003. ISBN: 3-540-76649-9. DOI: 10.1007/978-3-540-76649-0.
- [Mun07] James R. Munkres. *Topología*. Ed. por Ángel (Ferrández Izquierdo) Ferrández. 2^a ed. Prentice-Hall, 2007. ISBN: 9788420531809.
- [Mun18] James R. Munkres. *Elements of algebraic topology*. 2018. DOI: 10.1201/9780429493911.
- [NS13] Kathryn L. Nyman y Francis Edward Su. “A Borsuk-Ulam equivalent that directly implies Sperner’s Lemma”. En: *American Mathematical Monthly* 120 (4 2013). ISSN: 00029890. DOI: 10.4169/amer.math.monthly.120.04.346.
- [Pet16] James F. Peters. *Computational Proximity : Excursions in the Topology of Digital Images*. 1st ed. 2016. Intelligent Systems Reference Library, 102. Cham: Springer International Publishing, 2016. ISBN: 3-319-30262-0. DOI: 10.1007/978-3-319-30262-1. URL: <https://doi.org/10.1007/978-3-319-30262-1>.
- [Spe28] E. Sperner. “Neuer beweis für die invarianz der dimensionszahl und des gebietes”. En: *Abhandlungen aus dem Mathematischen Seminar der Universität Hamburg* 6.1 (dic. de 1928), págs. 265-272. ISSN: 1865-8784. DOI: 10.1007/BF02940617. URL: <https://doi.org/10.1007/BF02940617>.
- [Su97] Francis Edward Su. “Borsuk-Ulam implies brouwer: A direct construction”. En: *American Mathematical Monthly* 104 (9 1997). ISSN: 00029890. DOI: 10.2307/2975293.
- [Wei21] Max Weinstein. “Combinatorial perspectives on Borsuk-Ulam and Brouwer”. En: (2021).