

J. Albella Martínez
A. Fernández Tojo
A. M. González Rueda
A. Mascato García

EDITORES

As matemáticas do veciño

Actas do Seminario de Iniciación á Investigación

2014
2015

10º Aniversario



INSTITUTO DE MATEMÁTICAS

ACTAS DO SEMINARIO
DE
INICIACIÓN Á INVESTIGACIÓN

CURSO 2014 – 2015

Editores:

Jorge Albella Martínez

Adrián Fernández Tojo

Ángel M. González Rueda

Ana Mascato García

© 2013 Seminario de Iniciación á Investigación.
Instituto de Matemáticas da Universidade de Santiago de Compostela

Coordina:

Seminario de Iniciación á Investigación (SII)
seminarios3c@gmail.com

Edita:

Instituto de Matemáticas da Universidade de Santiago de Compostela

Imprime:

Campus na nube
C/ Casa Gradín (Campus Vida), baixo
15782 Santiago de Compostela
A Coruña

ISSN: 2171-6536

Depósito Legal: C 1873-2015

Mathematics is the most beautiful and most powerful creation of the human spirit.

Stefan Banach (1892-1945),

Today's scientists have substituted Mathematics for experiments, and they wander off through equation after equation, and eventually build a structure which has no relation to reality.

Nikola Tesla (1856-1943).

Prefacio

Cando o actual comité do SII me pediu que escribise un prólogo a esta nova edición das Matemáticas do veciño, só podía pensar: “Dez anos, xa pasaron dez anos”. Non me resulta difícil recordar como naceu o Seminario. Nas nosas estadías todos asistimos a seminarios máis ou menos formais, todos demos charlas máis ou menos técnicas, sabíamos (ou non), grazas a eles, que facía a veciña de despacho. Sen embargo, pouco coñecíamos do que investigaban os nosos amigos aquí. Non tiñamos, polo tanto, máis intención que a de encher un baleiro. E pasalo ben, claro, sempre pasalo ben. Non queríamos durar tanto. Perdón, non crimos que esto fose durar tanto. Fixemos unha declaración de intencións: charlas sen profesores, sen presentacións sofisticadas, para un público heteroxéneo e comprensibles para todos. Sen máis normas, e sen complexos, botamos a andar. E tanto andamos, que dez anos e toda unha vida despois me vexo ocupando o lugar que tivo a profesora Elena Vázquez Abal nas nosas primeiras Matemáticas do veciño. Se mo permitides, vou citar unhas palabras que escribiu ela naquel primeiro prefacio: *Levaron á práctica [...] o poñeren en común o seu traballo, as súas dúbidas, os seus obxectivos, a súa (in)experiencia oradora, facendo do descubrimento científico un labor máis solidario e menos solitario*. Non atopo palabras mellores para describir o noso Seminario, o voso Seminario. Grazas, de corazón, a todos os que manteñen vivo este proxecto, esta realidade que fundamos. Só me resta pedirllas a tódolos que colaboraron que disfruten del e da marabillosa etapa que están a vivir. Dez anos son un mundo, pero pasan nun suspiro.

Pontedeume, 17 de xuño de 2015.

Ana Belén Rodríguez Raposo

Índice xeral

Introdución	1
I Resumos das charlas	3
Cristina Vidal Castiñeira “Acciones polares”	5
Lorena Saavedra López “Disconxugación: Funcións de Green de signo constante”	11
Víctor Sanmartín López “¿Podemos confiar en las sucesiones?”	17
Jorge Rodríguez Veiga “Selección y asignación de recursos para la contención de un incendio forestal”	23
Jorge Losada Rodríguez “Macedonia fraccionaria”	29
Saray Busto Ulloa “Diseñando generadores”	35
Gonzalo Castiñeira Veiga “Un problema de narices”	41
Jesús Conde Lago “O último teorema de Fermat”	47
Lucía López Somoza “La ecuación de Hill”	53
Marcos Matabuena Rodríguez “El camino para ganar la NBA: técnicas de predicción matemática”	59

Néstor León Delgado	
“Geometría simpléctica superior motivada desde la teoría de campos lagrangiana”	65
M^a Isabel Borrajo García	
“Incendios forestales y estadística”	71
Pedro Pablo Campo Díaz	
“Dinámica de exoplanetas y exosatélites”	77
Jose Ameijeiras Alonso	
“Dando una vuelta a la estadística: los datos circulares”	83
II 10º Aniversario	89
J. Carlos Díaz Ramos	
“Simetría e forma”	91
Julio González Díaz	
“Votando secuencialmente”	97
Ariadna Arias	
“Coloquio: Indicadores de producción científica”	103

Introdución

É imposible que unha área das matemáticas progrese se as achegas que se fan nela quedan esquecidas nos caixóns dos seus descubridores. A comunicación entre os investigadores é fundamental, pero tamén o é a interdisciplinaridade, sobre todo entre os que comezan as súas carreiras investigadoras.

Con este espírito nace o Seminario de Iniciación a Investigación (SII), unha entidade encadrada dentro do Instituto de Matemáticas. O SII ten por finalidade que aqueles que se están a dar os seus primeiros pasos como investigadores teñan a oportunidade de escoitar aos seus compañeiros que traballan noutros departamentos e de expoñer as súas ideas.

As actividades do SII consisten nun conxunto de charlas que teñen lugar durante todo o curso académico na Facultade de Matemáticas da Universidade de Santiago de Compostela. Abertas a todo o mundo, estas reunións, en xeral quincenais, son un lugar para a discusión, o afloramento de ideas e a vida social na Facultade alén da rutina investigadora ou docente. Nelas, profesores, alumnos e investigadores teñen a oportunidade de coñecerse, emprender proxectos comúns e descubrir novos intereses. Ademais, para os poñentes supón unha oportunidade única de desenvolver competencias transversais fundamentais para as súas carreiras como son falar en público, a capacidade argumentativa e a adecuación á audiencia, pois cabe salientar que as charlas están destinadas a xente que non é especialista no tema do que tratan.

E imprescindible destacar ademais a riqueza da procedencia dos poñentes dos SII. Moi a miúdo temos o pracer de poder escoitar a xente chegada doutras facultades ou incluso doutras universidades, o cal da idea da capacidade de convocatoria e o alcance transversal das actividades do SII.

O Comité Organizador do SII, encargado de organizar as actividades do SII, facelas públicas e atender as necesidades loxísticas das mesmas, é tamén o responsable de elaborar estas actas que reflicten o enorme esforzo que, entre poñentes, oíntes e organizadores, estamos a realizar para que este proxecto sexa posible. Os propios poñentes foron os encargados de revisar os resumos das charlas, de xeito que cada quen tivo que corrixir un correspondente a unha área distinta á da súa especialidade, asegurando así que estes sexan comprensibles para todos.

Por último engadir que, como non podería ser doutra maneira, o curso que vén haberá cambios destinados a mellorar as actividades que o SII leva a cabo. Quizais o máis importante é a renovación da metade do Comité Organizador, de forma que os membros máis experimentados darán paso as novas xeracións de investigadores

para así manter vivo o espírito iniciador ao que previamente facíamos alusión.

Agradecementos

Non podemos esquecernos de todos aqueles aos que tanto lles debemos e que fan posible o SII. Aos anteriores organizadores do SII polo seu esforzo e consellos e por suposto aos poñentes: Jesús R. Aboal, Jose Ameijeiras, Ariadna Arias, M^a Isabel Borrajo, Saray Busto, Pedro Pablo Campo, Gonzalo Castiñeira, Jesús Conde, J. Carlos Díaz, Julio González, Néstor León, Lucía López, Jorge Losada, Marcos Matabuena, Salvador Naya, Juan José Nieto, Jorge Rodríguez, Lorena Saavedra, Víctor Sanmartín, Javier Tarrío-Saavedra e Cristina Vidal.

Tamén agradecemos sinceramente a elaboración do prefacio destas actas a Ana Belén Rodríguez Raposo, quen hai dez anos comezou, xunto cos seus compañeiros do primeiro comité organizador, este proxecto.

Santiago de Compostela, 15 de outubro do 2015.

O Comité Editorial.

Parte I

Resumos das charlas

Acciones polares

Cristina Vidal Castiñeira

Departamento de Geometría y Topología

1 de Octubre de 2014

Introducción

Según Felix Klein, la geometría es el estudio de aquellas propiedades de un espacio que son invariantes bajo la acción de un grupo de transformaciones. Un área importante dentro de la geometría diferencial es el estudio de las subvariedades de una variedad de Riemann dada. En particular estamos interesados en subvariedades con un alto grado de simetría, esto es, aquellas sobre las que actúa un grupo de isometrías suficientemente grande. Este interés nos lleva al concepto de acción isométrica y, en particular, al de acción polar, el cual definiremos en este artículo.

Las nociones y resultados citados en este artículo se pueden encontrar en [1] y [3].

Acciones isométricas

Para poder definir acciones isométricas, necesitamos previamente de dos conceptos básicos, como son grupo de Lie y la operación diferenciable de un grupo sobre una variedad. Por lo que procedamos a definirlos.

Definición 1. *Un grupo de Lie es un grupo topológico Hausdorff G dotado de una estructura de variedad diferenciable C^k compatible con la estructura de grupo, es decir, tal que se verifica que la aplicación*

$$\begin{aligned}\varphi : G \times G &\rightarrow G \\ (x, y) &\mapsto xy^{-1}\end{aligned}$$

es diferenciable C^k .

Definición 2. *Se dice que un grupo de Lie G opera diferenciablemente a la izquierda sobre una variedad diferenciable X , si existe una aplicación diferenciable*

$$\begin{aligned}\psi : G \times X &\rightarrow X \\ (a, x) &\mapsto \psi(a, x) = ax,\end{aligned}$$

tal que:

1. $(ab)x = a(bx)$, es decir, $\psi(ab, x) = \psi(a, \psi(b, x))$, $\forall x \in X, \forall a, b \in G$.

2. $\forall a \in G$, la aplicación

$$\begin{aligned}\bar{a} : X &\rightarrow X \\ x &\mapsto ax\end{aligned}$$

es un difeomorfismo.

Observación 3. Análogamente definimos la operación diferenciable a la derecha de un grupo sobre una variedad. Con ambas definiciones tenemos la operación diferenciable de un grupo sobre una variedad.

Ahora ya estamos en las condiciones necesarias para definir acción isométrica.

Definición 4. Sea $\bar{M}$ una variedad de Riemann y G un grupo de Lie. Decimos que G actúa isométricamente sobre $\bar{M}$ si existe una aplicación diferenciable

$$\begin{aligned}\varphi : G \times \bar{M} &\rightarrow \bar{M} \\ (g, p) &\mapsto gp\end{aligned}$$

que satisface $(gg')p = g(g'p)$ para todo $g, g' \in G$ y $p \in \bar{M}$, y tal que la aplicación

$$\begin{aligned}\varphi_g : \bar{M} &\rightarrow \bar{M} \\ p &\mapsto gp\end{aligned}$$

es una isometría de $\bar{M}$ para todo $g \in G$. A tal ϕ se la conoce como acción isométrica en $\bar{M}$. Se denota por $I(\bar{M})$ al grupo de isometrías de $\bar{M}$, el cual se sabe que es un grupo de Lie.

Por cada punto $p \in \bar{M}$, la órbita de la acción de G a través de p se define como el conjunto

$$G \cdot p = \{gp : g \in G\}$$

y el grupo de isotropía o estabilizador en p como

$$G_p = \{g \in G : gp = p\}.$$

Veamos a continuación algunos ejemplos de acciones isométricas.

Ejemplo 5.

1. Sea S^1 la esfera usual en $\mathbb{R}^2$ (la circunferencia). La acción del grupo $SO(2) = SO(2, \mathbb{R})$ en S^1 es una acción isométrica. Veámoslo:

Toda acción isométrica de $SO(2)$ en S^1 es de la forma

$$\begin{aligned}\varphi : SO(2) \times S^1 &\rightarrow S^1 \\ (A, x) &\mapsto A \cdot x = Ax,\end{aligned}$$

donde

$$A = \begin{pmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{pmatrix}.$$

a) Veamos primero que $SO(2)$ es un grupo de Lie:

- $\forall A, B \in SO(2)$

$$\begin{aligned} A \cdot B &= \begin{pmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{pmatrix} \cdot \begin{pmatrix} \cos(\phi) & -\sin(\phi) \\ \sin(\phi) & \cos(\phi) \end{pmatrix} \\ &= \begin{pmatrix} \cos(\theta + \phi) & -\sin(\theta + \phi) \\ \sin(\theta + \phi) & \cos(\theta + \phi) \end{pmatrix} \in SO(2). \end{aligned}$$

- $\forall A, B, C \in SO(2)$

$$\begin{aligned} (AB)C &= \begin{pmatrix} \cos(\theta + \phi) & -\sin(\theta + \phi) \\ \sin(\theta + \phi) & \cos(\theta + \phi) \end{pmatrix} \cdot \begin{pmatrix} \cos(\xi) & -\sin(\xi) \\ \sin(\xi) & \cos(\xi) \end{pmatrix} \\ &= \begin{pmatrix} \cos(\theta + \phi + \xi) & -\sin(\theta + \phi + \xi) \\ \sin(\theta + \phi + \xi) & \cos(\theta + \phi + \xi) \end{pmatrix} \\ &= \begin{pmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{pmatrix} \cdot \begin{pmatrix} \cos(\phi + \xi) & -\sin(\phi + \xi) \\ \sin(\phi + \xi) & \cos(\phi + \xi) \end{pmatrix} = A(BC). \end{aligned}$$

- Basta tomar $\theta = 0$ y tenemos el elemento neutro:

$$A = \begin{pmatrix} \cos(0) & -\sin(0) \\ \sin(0) & \cos(0) \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = I \in SO(2).$$

- $\forall A \in SO(2) \exists A^{-1}$,

$$A^{-1} = \begin{pmatrix} \cos(\theta) & \sin(\theta) \\ -\sin(\theta) & \cos(\theta) \end{pmatrix} = \begin{pmatrix} \cos(-\theta) & -\sin(-\theta) \\ \sin(-\theta) & \cos(-\theta) \end{pmatrix} \in SO(2),$$

tal que

$$AA^{-1} = A^{-1}A = I.$$

Por tanto, $SO(2)$ es un grupo. Y tomando la aplicación

$$\begin{aligned} \psi : SO(2) \times SO(2) &\rightarrow SO(2), \\ (A, B) &\mapsto AB^{-1} \end{aligned}$$

que es una aplicación diferenciable (por ser composición de aplicaciones diferenciables), tenemos que $SO(2)$ es un grupo de Lie.

b) Veamos ahora que la acción φ es una acción isométrica. Dado que

$$Ax = \begin{pmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} \cos(\theta)x_1 - \sin(\theta)x_2 \\ \sin(\theta)x_1 + \cos(\theta)x_2 \end{pmatrix},$$

se obtiene que

$$\begin{aligned}
 (AB)x &= \begin{pmatrix} \cos(\theta + \phi) & -\sin(\theta + \phi) \\ \sin(\theta + \phi) & \cos(\theta + \phi) \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \\
 &= \begin{pmatrix} \cos(\theta + \phi)x_1 - \sin(\theta + \phi)x_2 \\ \sin(\theta + \phi)x_1 + \cos(\theta + \phi)x_2 \end{pmatrix} \\
 &= \begin{pmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{pmatrix} \begin{pmatrix} \cos(\phi)x_1 - \sin(\phi)x_2 \\ \sin(\phi)x_1 + \cos(\phi)x_2 \end{pmatrix} = A(Bx).
 \end{aligned}$$

Además, como la aplicación

$$\begin{aligned}
 \varphi_A : S^1 &\rightarrow S^1 \\
 x &\mapsto Ax
 \end{aligned}$$

tiene como matriz asociada A , que es la matriz de una rotación, tenemos que φ_A es una isometría.

Por tanto, la acción de $SO(2)$ en S^1 es una acción isométrica con órbita S^1 .

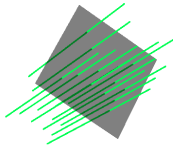
2. La acción de $SO(2)$ en $\mathbb{R}^2$ es una acción isométrica, y la prueba es análoga al caso anterior. Las órbitas en este caso serían todas las circunferencias concéntricas en el origen y el propio origen.
3. Sea S^2 la esfera usual en $\mathbb{R}^3$. El grupo $SO(3) = SO(3, \mathbb{R})$ consiste en las isometrías lineales que preservan la orientación, es decir, aplicaciones lineales que preservan la distancia y de determinante $+1$. Como toda aplicación lineal deja al origen como punto fijo, vemos que cualquier rotación lleva S^2 en sí misma. Por lo que la acción de $SO(3)$ en $\mathbb{R}^3$ sería otro ejemplo de acción isométrica.

Acciones polares

Definición 6. Una acción isométrica de un grupo G sobre una variedad de Riemann $\bar{M}$ se dice polar si su foliación de órbitas es polar, es decir, si existe una subvariedad inmersa Σ de $\bar{M}$ que interseca a todas las órbitas de la G -acción, y para cada $p \in \Sigma$, el espacio tangente de Σ en p , $T_p\Sigma$, y el espacio tangente de la órbita a través de p en p , $T_p(G \cdot p)$, son ortogonales. En tal caso, la subvariedad Σ se llama sección de la G -acción. Cualquier acción polar admite secciones a través de un punto dado.

Tipos de acciones polares en $\mathbb{R}^3$

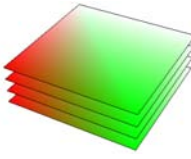
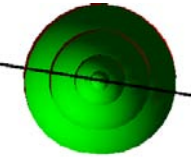
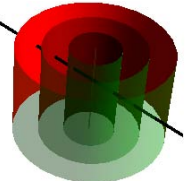
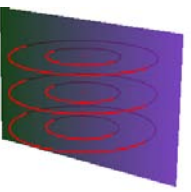
1.



Acción de $\mathbb{R}$ en $\mathbb{R}^3$.

Órbitas: rectas paralelas.

Sección: cualquier plano ortogonal a una de las rectas será ortogonal a todas.

2.  Accion de $\mathbb{R}^2$ en $\mathbb{R}^3$. (Tambien es la accion de $O(2)$ en $\mathbb{R}^3$).
Orbitas: planos paralelos.
Seccion: cualquier linea recta que interseque a uno de los planos ortogonalmente, los intersecara a todos.
3.  Accion de $SO(3)$ en $\mathbb{R}^3$.
Orbitas: un punto y todas las esferas concentricas en torno a dicho punto.
Seccion: cualquier linea recta a traves de dicho punto.
4.  Accion de $SO(2) \times \mathbb{R}$ en $\mathbb{R}^3$.
Orbitas: una linea recta y cilindros coaxiales, alrededor de dicha recta.
Seccion: cualquier recta que interseque al eje ortogonalmente.
5.  Accion de $SO(2)$ en $\mathbb{R}^3$.
Orbitas: un eje de puntos fijos (la linea recta) y circunferencias concentricas en torno a ese eje de puntos fijos.
Seccion: cualquier plano que contenga al eje de puntos fijos.
6. LA ACCION TRIVIAL Accion del grupo de la matriz identidad, I , en $\mathbb{R}^3$.
Orbitas: todos los puntos de $\mathbb{R}^3$.
Seccion: todo $\mathbb{R}^3$.
7. LA ACCION TRANSITIVA Accion de $\mathbb{R}^3$ en $\mathbb{R}^3$.
Orbitas: todo $\mathbb{R}^3$.
Seccion: cualquier punto de $\mathbb{R}^3$ (ya que el vector tangente asociado es el vector 0, y por tanto siempre es ortogonal).

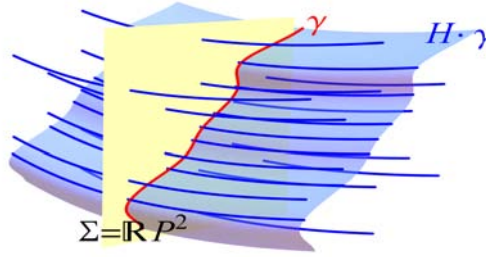
Aplicacion de las acciones polares

Una de las aplicaciones de las acciones polares la podemos ver en el articulo pendiente de publicacion [2], donde se estudian las hipersuperficies reales con dos curvaturas principales en los planos proyectivo e hiperbolico complejos.

Gran parte de los resultados este trabajo se obtuvieron gracias a las acciones polares, ya que la idea de la construccion de tales hipersuperficies, por ejemplo en el plano proyectivo complejo, $\mathbb{C}P^2$, es la que mostraremos a continuacion.

Tomamos la acción del grupo $H = U(1) \times U(1) \times U(1)$ en $\mathbb{C}^3$ y la cocientamos por la fibración de Hopf ($z \sim \lambda z$ con $\lambda \in S^1, z \in \mathbb{C}^3$). Lo que obtenemos es que H actúa sobre $\mathbb{C}P^2$ polarmente con sección $\Sigma = \mathbb{R}P^2$ y con órbitas principales $S^1 \times S^1$.

Lo que principalmente se hace en dicho artículo es tomar dichas órbitas de la acción polar (las cuales son toros bidimensionales) y buscar una curva γ en la sección Σ de manera que la unión de dichas órbitas a lo largo de la curva γ nos proporcione la hipersuperficie con dos curvaturas principales, $H \cdot \gamma$, que estábamos buscando.



El resultado es mucho más complejo que este esquema, pero esto nos muestra una breve idea de la utilidad de tales acciones.

Bibliografía

- [1] Berndt, J., Console, S. y Olmos, C. (2003). *Submanifolds and holonomy*, Chapman & Hall/CRC Research Notes in Mathematics, **434**, Chapman and Hall/CRC, Boca Raton, FL.
- [2] Díaz-Ramos, J. C., Domínguez-Vázquez, M. y Vidal-Castiñeira, C. *Real hypersurfaces with two principal curvatures in complex projective and hyperbolic planes*, arXiv:1310.0357v1 [math.DG].
- [3] Domínguez-Vázquez, M. (2013). *Isoparametric foliations and polar actions on complex space forma*, Tesis doctoral, Universidad de Santiago de Compostela.

Disconjugación: Funcións de Green de signo constante

Lorena Saavedra López

Departamento de Análise Matemática

15 de outubro de 2014

Introdución

Centrándonos no concepto de disconjugación para ecuacións diferenciais lineais ordinarias de orde n [Co], que caracteriza o máximo número de ceros que pode posuír unha solución dedúcese certas propiedades cualitativas das solucións de determinadas ecuacións diferenciais non homoxéneas.

Disconjugación

Definición 1. Sexan $p_1(t), \dots, p_n(t) \in L^2(I)$.

Unha ecuación diferencial de orde n

$$y^{(n)}(t) + p_1(t)y^{(n-1)}(t) + \dots + p_n(t)y(t) = 0, \quad (1)$$

dise disconjugada nun intervalo I se cada solución non trivial ten menos de n ceros en I , contados con multiplicidade.

Visto que os coeficientes da ecuación diferencial tal como foi definida pertencen ó espazo $L^2(I)$, as nosas posibles solucións pertencerán ó seguinte espazo:

$$\begin{aligned} H^n(I) &= \left\{ u: I \rightarrow \mathbb{R} \mid u^{(k)} \in L^2(I), k = 0, \dots, n \right\} \\ &= \left\{ u \in C^{n-1}(I) \mid u^{(n-1)} \in AC(I), u^{(n)} \in L^2 \right\}. \end{aligned}$$

De seguido, danse unha serie de resultados sobre este concepto.

Proposición 2. Sexa I un intervalo, entón existe $\delta > 0$ tal que (1) é disconjugada en calquera subintervalo de lonxitude menor ca δ .

Proposición 3. Se a ecuación (1) é disconjugada en I existirá unha única solución non trivial satisfacendo as seguintes condicións:

$$y^{(\nu_i)}(a_i) = \beta_{i\nu_i}, \quad \nu_i = 0, \dots, r_i - 1, \quad i = 1, \dots, m,$$

con $r_1 + \dots + r_m = n$ e $a_i \in I$ para todo $i = 1, \dots, m$, sempre e cando algún dos $\beta_{i\nu_i} \neq 0$.

PALABRAS CLAVE: Disconjugación, sistemas de Markov e Descartes, función de Green.

Proposición 4. *Se o intervalo $I = (a, b)$ é aberto e cada solución non trivial de (1) ten menos de n ceros distintos, entón cada solución ten menos de n ceros contados coa súa multiplicidade.*

Vense, a continuación, uns tipos de sistemas que permitirán obter distintas condicións sobre as ecuacións diferenciais.

Definición 5. *Sexan $y_1, \dots, y_n$ funcións de clase H^n nun intervalo I . Dise que forman un sistema de Čebyšev se calquera combinación non trivial de $y_1(t), \dots, y_n(t)$ ten menos de n ceros.*

É evidente que a ecuación (1) é disconjugada en I se, e só se, algún sistema fundamental de solucións e, polo tanto, cada sistema fundamental é de Čebyšev.

Definición 6. *Sexan $y_1, \dots, y_n$ funcións de H^n nun intervalo I . Dise que forman un sistema de Markov se os n Wronskianos*

$$W(y_1, \dots, y_k) = \begin{vmatrix} y_1 & \dots & y_k \\ \vdots & \ddots & \vdots \\ y_1^{(k-1)} & \dots & y_k^{(k-1)} \end{vmatrix}, \quad k = 1, \dots, n,$$

son positivos en I .

Definición 7. *Sexan $y_1, \dots, y_n$ funcións de H^n nun intervalo I . Dise que forman un sistema de Descartes se tódolos Wronskianos ordeados*

$$W(y_{i_1}, \dots, y_{i_k}) = \begin{vmatrix} y_{i_1} & \dots & y_{i_k} \\ \vdots & \ddots & \vdots \\ y_{i_1}^{(k-1)} & \dots & y_{i_k}^{(k-1)} \end{vmatrix}, \quad 1 \leq i_1 < \dots < i_k \leq n, \quad k = 1, \dots, n,$$

son positivos en I .

Teorema 8. *A ecuación diferencial lineal (1) ten un sistema fundamental de solucións de Markov se, e só se, o operador L ten unha representación*

$$Ly \equiv v_1 v_2 \dots v_n \frac{d}{dt} \left(\frac{1}{v_n} \frac{d}{dt} \left(\dots \frac{d}{dt} \left(\frac{1}{v_2} \frac{d}{dt} \left(\frac{1}{v_1} y \right) \right) \right) \right),$$

onde $v_k > 0$ son funcións de clase H^{n-k+1} para $k = 1, \dots, n$.

Vexamos un exemplo de segunda orde no que aparece reflectido o anterior resultado.

Exemplo 9. *Sexa $m > 0$, tal que $m \in (0, \pi)$ consideramos o problema*

$$y''(t) + m^2 y(t) = 0, \quad t \in (0, 1].$$

Obtemos un sistema fundamental de solucións da forma

$$\begin{aligned} y_1(t) &= \sin(mt), \\ y_2(t) &= -\cos(mt). \end{aligned}$$

Podemos verificar que é un sistema de Markov:

- $W_1 = \sin(mt) > 0$, xa que $mt \in (0, \pi)$.
- *Calculemos W_2 :*

$$W_2 = \begin{vmatrix} \sin(mt) & -\cos(mt) \\ m \cos(mt) & m \sin(mt) \end{vmatrix} = m > 0.$$

Polo tanto podemos construír v_1 e v_2 .

$$\begin{aligned} v_1 &= \sin(mt), \\ v_2 &= \frac{W_2}{\sin^2(mt)} = \frac{m}{\sin^2(mt)}. \end{aligned}$$

A construción do operador L sería agora

$$\begin{aligned} Ly &\equiv \sin(mt) \frac{m}{\sin^2(mt)} \frac{d}{dt} \left(\frac{\sin^2(mt)}{m} \frac{d}{dt} \left(\frac{y}{\sin(mt)} \right) \right) \\ &= \frac{m}{\sin(mt)} \frac{d}{dt} \left(\frac{y'(t) \sin(mt) - m \cos(mt) y(t)}{m} \right) \\ &= \frac{1}{\sin(mt)} (y''(t) \sin(mt) + m y'(t) \cos(mt) \\ &\quad + m^2 \sin(mt) y(t) - m \cos(mt) y(t)) \\ &= y''(t) + m^2 y(t). \end{aligned}$$

Apoiándonos nesta caracterización obtéñense unha serie de resultados:

Teorema 10. *A ecuación (1) ten un sistema fundamental de solucións de Markov no intervalo compacto $I = [a, b]$ se, e só se, é disconxugada en I .*

Proposición 11. *O conxunto de tódalas ecuacións (1) disconxugadas nun intervalo compacto I é conexo e aberto.*

Teorema 12. *Se a ecuación (1) ten un sistema fundamental de solucións de Markov nun intervalo $I = [a, b]$ ou $I = [a, b)$ entón ten un sistema fundamental de solucións de Descartes en I .*

Teorema 13. *A ecuación (1) ten un sistema fundamental de solucións de Descartes nun intervalo compacto $I = [a, b]$ se, e só se, é disconxugada en I .*

Función de Green

Construcción da función de Green

Sexa f unha función de $L^2(I)$, $a_1 < \dots < a_m$ puntos de I e $r_1, \dots, r_m$ enteiros positivos cuxa suma é n .

O problema multipunto

$$\begin{aligned} Ly(t) &= f(t), \quad t \in I, \\ y^{(\nu_i)}(a_i) &= 0 \quad \nu_i = 0, \dots, r_i - 1; \quad i = 1, \dots, m, \end{aligned} \quad (2)$$

ten unha solución única, porque para $f \equiv 0$, a ecuación $Ly = 0$ soamente ten a solución $y = 0$ con n ceros, por ser disconxugada.

A única solución pódese representar da seguinte forma

$$y(t) = \int_a^b G(t, s) f(s) ds,$$

onde a función de Green, $G(t, s)$, está definida polas seguintes propiedades:

- Como función de t , $G(t, s)$ é solución da ecuación (1) nos intervalos $[a, s)$ e $(s, b]$, satisfacendo as condicións de fronteira

$$G^{(\nu_i)}(a_i, s) = 0 \quad \nu_i = 0, \dots, r_i - 1; \quad i = 1, \dots, m, \quad s \in I.$$

- Como función de t , $G(t, s)$ e as súas primeiras $n - 2$ derivadas son continuas en $t = s$, mentres que

$$\frac{\partial^{n-1}}{\partial t^{n-1}} G(s^+, s) - \frac{\partial^{n-1}}{\partial t^{n-1}} G(s^-, s) = 1, \quad s \in I.$$

Baixo estas condicións, non é difícil verificar que a función y será a única solución do problema (2).

Signo da función de Green

Estudiaremos a partir de agora os ceros da función de Green, para un $s \in I$ fixado. Sexa $P(t) = (t - a_1)^{r_1} \dots (t - a_m)^{r_m}$. Tense o seguinte resultado.

Lema 14. *Se a ecuación (1) é disconxugada e $G(t, s)$ é a función de Green asociada ó problema (2). Entón no cadrado $[a, b] \times [a, b]$ cúmprese*

$$G(t, s) P(t) \geq 0.$$

Teorema 15. *Se a ecuación (1) é disconxugada en $[a_1, a_m]$ e $G(t, s)$ é a función de Green asociada ó problema (2). Cúmprese que $G(t, s)/P(t) > 0$ para $a_1 < s < a_m$ e $a_1 \leq t \leq a_m$.*

Lema 16. *Se temos unha ecuación do tipo (1) con coeficientes constantes*

$$y^{(n)}(t) + c_1 y^{(n-1)}(t) + \dots + c_n y^{(n)}(t) = 0, \quad (3)$$

que é disconxugada nun intervalo $[a, b]$. Entón (3) é disconxugada en calquera intervalo de lonxitude menor ou igual a $b - a$.

Proposición 17. *Consideramos o problema (2) asociado ó operador con coeficientes constantes*

$$Ly \equiv y^{(n)}(t) + a_1 y^{(n-1)}(t) + \dots + a_n y(t) = 0.$$

Se esta ecuación é disconjugada en $[a_1, a_m]$ podemos afirmar que a función de Green satisfai a tese do Lema 14, é dicir,

$$G(t, s) P(t) \geq 0,$$

nos seguintes conxuntos:

$$\begin{aligned} A &= [a_1, a_m] \times [a_1, a_m], \\ B &= [2a_1 - a_m, 2a_m - a_1] \times [a_1 - 2a_m, a_1], \\ C &= [2a_1 - a_m, 2a_m - a_1] \times [a_m, 2a_m - a_1]. \end{aligned}$$

Vexamos un exemplo no que ademais se cumpre $G(t, s) P(t) \geq 0$ nos conxuntos

$$\begin{aligned} D &= [2a_1 - a_m, a_1] \times [a_1, a_m], \\ E &= [a_m, 2a_m - a_1] \times [a_1, a_m]. \end{aligned}$$

Exemplo 18. *Consideramos o problema de segunda orde no intervalo $[-1, 2]$*

$$\begin{aligned} u''(t) + m^2 u(t) &= 0, \quad t \in [-1, 2], \\ u(0) = u(1) &= 0. \end{aligned} \tag{4}$$

Obtense, tendo en conta as condicións impostas na definición de función de Green, a seguinte función de Green para o problema anterior:

$$G(t, s) = \begin{cases} \frac{-\operatorname{sen}(ms) \operatorname{sen}(m(1-t))}{m \operatorname{sen}(m)}, & 0 \leq s \leq t \leq 1, \\ \frac{-\operatorname{sen}(m(1-s)) \operatorname{sen}(mt)}{m \operatorname{sen}(m)}, & t < s \leq 1, \quad t \leq 0, \\ \frac{\operatorname{sen}(m(t-s))}{m}, & 1 < s \leq t, \\ -\frac{\operatorname{sen}(m(t-s))}{m}, & t < s < 0, \\ 0, & \text{noutro caso.} \end{cases}$$

Neste caso $P(t) = t(t-1)$.

Podemos ver que no intervalo $[0, 1]$ a ecuación é disconjugada para $m \in (0, \pi)$.

As solucións deste problema son da forma

$$u(t) = \alpha_1 \operatorname{sen}(mt) + \alpha_2 \operatorname{sen}(m(1-t)).$$

Podemos suponer que $u(0) = 0$, senon faise unha traslacion da primeira raiz a cero, e teriamos as

$$u(0) = \alpha_2 \sin(m)$$

o que implica que $\alpha_2 = 0$, xa que estamos a suponer que $m = k\phi$ para $k \in \mathbb{Z}$. Polo tanto a solucion sera da forma

$$u(t) = \alpha_1 \sin(mt)$$

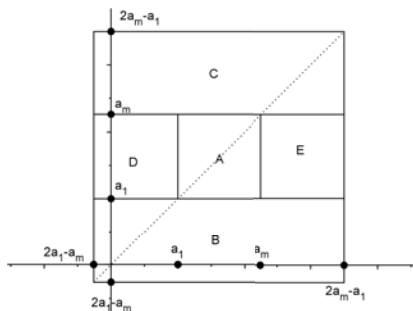
Así, $u(t) = 0$ se, e so se, $\sin(mt) = 0$, xa que $\alpha_1 = 0$ para que a solucion sexa non trivial. Isto cumprese se, e so se, $mt = k\phi$ sendo $k = 0, 1, \dots$, as que o segundo cero atopariámolo nun punto $t = \phi/m$.

Este punto $\phi/m \in [0, 1]$ se, e so se, $m \in [0, \phi]$. Polo tanto, a ecuacion (4) sera disconjugada se, e so se, $m \in [0, \phi/2]$.

Polo resultado anterior sabemos que $P(t)G(t+s) \geq 0$ nos conxuntos A, B e C . Ademais, para este problema, tamen se cumpre a desigualdade nos conxuntos D e E . Dado que en ambos subconxuntos $P(t) = t(t-1) \geq 0$, tense que cumprir que $G(t+s) \geq 0$.

Efectivamente, conxunto D , tense que $m(1-s) \in (0, \phi)$ e $mt \in (\phi, 0)$ polo que $G(t+s) \geq 0$.

Analogamente no conxunto E , $m(1-t) \in (\phi, 0)$ e $ms \in (0, \phi)$, e así $G(t+s) \geq 0$.



Nun futuro ser a interesante poder proporcionar algun resultado xeral que nos indicase que sucede nos conxuntos E e D de forma xeral.

Bibliografía

[Co] Coppel, W.A. (1971). *Disconjugacy*. Lecture Notes in Mathematics. Springer-Verlag.

¿Podemos confiar en las sucesiones?

Víctor Sanmartín López

Departamento de Geometría y Topología

29 de octubre de 2014

Introducción

El desarrollo de la topología general estuvo ligado desde sus comienzos a la búsqueda de una buena noción de convergencia y, tradicionalmente, convergencia significaba convergencia de sucesiones. Por ejemplo, los primeros intentos de Fréchet de encontrar una clase general de espacios adecuada para el estudio de nociones topológicas se centraron en el uso de la convergencia secuencial y así, en 1906, al mismo tiempo que introducía la noción de espacio métrico, fue de los primeros en plantear el siguiente problema:

Caracterizar la clase de los espacios que pueden ser completamente determinados por sus sucesiones convergentes.

La propuesta de Fréchet consistió en introducir axiomáticamente la noción de convergencia en un conjunto. Una convergencia en X sería una colección de pares $((x_n); x)$ sujeta a unos axiomas, siendo (x_n) una sucesión en X y x un punto de X . En 1923 Aleksandrov y Urysohn modificaron los axiomas originales de Fréchet dando lugar a los más tarde denominados espacios $\mathcal{L}^*$ que, junto con diversas variantes, se conocen actualmente como *espacios de convergencia*.

Sin embargo, pronto estuvieron claras las limitaciones de este enfoque y, del mismo modo que la teoría de integración había revelado la insuficiencia de la convergencia secuencial para los propósitos del análisis, se hizo evidente que las sucesiones no eran adecuadas para los propósitos de la topología general, dando como resultado que el desarrollo posterior de la topología se fundamentase en otros conceptos tales como sistemas de entornos, operadores de clausura o conjuntos abiertos. Sin embargo, aunque las sucesiones no permiten describir la topología de espacios topológicos generales, sí lo hacen en muchos espacios importantes y son más simples e intuitivas que otras nociones de convergencia posteriores tales como las redes o los filtros. Por ello, se consideró de interés la determinación del verdadero alcance de su utilidad y así, en 1962, Venkataraman planteó de nuevo el problema de Fréchet en términos más concretos:

Caracterizar la clase de los espacios topológicos que pueden ser completamente determinados por sus sucesiones convergentes.

La respuesta más cumplida fue la dada por Franklin en 1965 y 1967 cuando, en dos artículos [1, 2] con el significativo título de “*Spaces in which sequences suffice*”, introdujo los denominados espacios secuenciales.

Espacios métricos y primero numerables

La topología de un espacio métrico puede ser caracterizada mediante la convergencia de sucesiones y muchos conceptos topológicos admiten una expresión simple con dicha convergencia. Por ejemplo:

- *Un subconjunto A de un espacio métrico es abierto si, y solo si, toda sucesión convergente a un punto de A finaliza en A .*
- *La adherencia de un conjunto está constituida por los límites de las sucesiones de puntos del conjunto.*
- *La continuidad equivale a la continuidad secuencial.*
- *La compacidad equivale a la compacidad secuencial.*

La convergencia de sucesiones en espacios topológicos se define de manera natural y verifica muchas de las propiedades elementales (convergencia de sucesiones estacionarias, convergencia de subsucesiones de una convergente, etc.); pero no todas, por ejemplo, la unicidad del límite: para topologías gruesas es habitual que una sucesión tenga varios límites, pudiendo llegar al extremo, para la topología trivial, de que toda sucesión converja a todo punto.

Las propiedades arriba citadas de la convergencia secuencial en espacios métricos tampoco se verifican en espacios topológicos generales. Es fácil encontrar contraejemplos. De nuevo, es preciso imponer condiciones adicionales y lo habitual es usar el primer axioma de numerabilidad debido a que permite replicar muchos de los argumentos de espacios métricos, sin más que substituir la colección de las bolas de radio $1/n$ centradas en un punto por una base local numerable $\{V_n \mid n \in \mathbb{N}\}$ que verifique $V_1 \supset V_2 \supset \dots$ (tal base local se puede construir a partir de una base local numerable arbitraria $\{U_n \mid n \in \mathbb{N}\}$ definiendo, para cada $n \in \mathbb{N}$, $V_n = U_1 \cap U_2 \cap \dots \cap U_n$). De esta manera se obtiene que, en espacios que satisfacen el primer axioma de numerabilidad, se verifican muchas de las propiedades citadas para espacios métricos. Entre ellas las siguientes:

- *Un subconjunto A es abierto si, y solo si, toda sucesión convergente a un punto de A finaliza en A .*
- *La adherencia de un conjunto está constituida por los límites de las sucesiones de puntos del conjunto.*
- *El carácter Hausdorff del espacio equivale a que todas las sucesiones convergentes tengan límite único.*

- *La continuidad equivale a la continuidad secuencial.*
- *La compacidad numerable equivale a la compacidad secuencial.*

Como se observa, ha sido preciso substituir la compacidad por la compacidad numerable (en espacios métricos ambos conceptos son equivalentes) lo que podría inducir a pensar que, con una adecuada reformulación, las propiedades anteriores también se verificarían en espacios topológicos generales. Sin embargo, tomando por ejemplo la recta real $\mathbb{R}$, si consideramos sobre ella las topologías discreta y conumerable, es sencillo comprobar que con ambas topologías las sucesiones convergentes son las mismas, sin embargo, se trata por definición de espacios topológicos diferentes.

Este ejemplo pone de manifiesto que la convergencia secuencial no determina la topología del espacio sino una clase de topologías, aunque en dicha clase será posible distinguir una de ellas susceptible de ser caracterizada mediante la convergencia de sucesiones; tales topologías serán justamente las de los espacios secuenciales. Se introduce a continuación la terminología y las notaciones que serán empleadas en lo sucesivo.

Si (x_n) es una sucesión en un conjunto X y $A \subset X$, se dirá que (x_n) *finaliza en* A cuando exista $\nu \in \mathbb{N}$ tal que $x_n \in A$ para todo $n \geq \nu$.

Definición 1. Sean (X, τ) un espacio topológico, (x_n) una sucesión en X y $x \in X$. Se dice que (x_n) *converge a* x en (X, τ) , y se representa por $(x_n) \rightarrow x$, si (x_n) *finaliza en todo entorno de* x .

Es suficiente comprobar que la sucesión finaliza en todos los entornos de alguna base local de x , por ejemplo en los entornos abiertos.

Entre las consecuencias directas de la definición se tiene que, si (x_n) es una sucesión convergente a x , toda subsucesión (x_{n_k}) converge también a x , y que la continuidad implica la continuidad secuencial.

Espacios secuenciales

En lo que sigue (X, τ) será un espacio topológico arbitrario.

Definición 2. Un subconjunto $A \subset X$ es secuencialmente abierto en (X, τ) si toda sucesión en X convergente en (X, τ) a un punto de A finaliza en A .

Dado que un conjunto abierto es entorno de cada uno de sus puntos, se tiene:

Proposición 3. Todo subconjunto abierto de X es secuencialmente abierto.

El complementario de un conjunto secuencialmente abierto en (X, τ) se denominará secuencialmente cerrado en (X, τ) . La proposición anterior asegura que todo subconjunto cerrado de X es secuencialmente cerrado.

Puede comprobarse mediante argumentos sencillos que si (X, τ) es un espacio topológico arbitrario, los conjuntos secuencialmente abiertos forman una topología que denotaremos por τ_s , que además es más fina que la topología de partida, es decir, $\tau \subset \tau_s$. Además, puede establecerse dentro del retículo de topologías del conjunto X , la relación de equivalencia *ser secuencialmente equivalentes*, que relaciona aquellas topologías con las mismas sucesiones convergentes. De esta manera, es posible argumentar que τ y τ_s son topologías sobre X secuencialmente equivalentes. Es más, τ_s es la topología más fina de entre las secuencialmente equivalentes a τ . Hechas estas consideraciones, pasamos a definir los espacios secuenciales, que, como veremos, responden a la cuestión que se planteaba en la introducción.

Definición 4. *Un espacio topológico (X, τ) se dice secuencial si todo subconjunto secuencialmente abierto es abierto.*

Se comprueba de manera inmediata que:

Proposición 5. *Para un espacio topológico (X, τ) equivalen:*

1. (X, τ) es secuencial.
2. Todo subconjunto secuencialmente cerrado es cerrado.
3. $\tau = \tau_s$.

De este modo, puede decirse que la convergencia de sucesiones determina la topología de un espacio secuencial. En este punto, utilizando simplemente la definición, es posible comprobar que los espacios secuenciales constituyen una clase de espacios mayor que la de los primero numerables. En términos más sencillos, todo espacio primero numerable (y por tanto métrico) es secuencial.

Un primer análisis de gran interés sobre estos nuevos espacios que acabamos de definir consiste en comprobar como se comportan ante las principales operaciones topológicas. Así, el producto de espacios secuenciales no es, en general, secuencial. No obstante, puede imponerse alguna condición sobre uno de los factores, por ejemplo la compacidad secuencial local, para obtener un producto secuencial.

Por otra parte, los espacios secuenciales se conservan por paso al cociente y mediante la suma topológica, en el sentido que precisan las dos siguientes proposiciones:

Proposición 6. *La suma topológica de cualquier familia de espacios secuenciales es un espacio secuencial.*

Proposición 7. *Sea X un espacio secuencial y $f: X \rightarrow Y$ una aplicación cociente. Entonces Y es un espacio secuencial.*

Por ́ltimo, conviene mencionar que la secuencialidad no es una propiedad estrictamente hereditaria, aunque ańadiendo la hiṕtesis de abierto o cerrado sobre el subespacio obtenemos un subespacio secuencial.

Tras todas estas consideraciones, sería conveniente saber cuáles son exactamente los espacios secuenciales. Es sabido que constituyen una clase de espacios más amplia que la de los espacios primero numerables, sin embargo, puede decirse mucho más sobre ellos. En este sentido, se enuncia a continuación la caracterización de los espacios secuenciales dada por Franklin [1].

Teorema 8. *Para cualquier espacio topoĺgico (X, τ) equivalen:*

1. (X, τ) es secuencial,
2. (X, τ) es cociente de un espacio métrico,
3. (X, τ) es cociente de un espacio que verifica el primer axioma de numerabilidad.

Una vez establecido el hecho de que, para espacios secuenciales, la convergencia de sucesiones determina la topología, se comprueba a continuación hasta que punto permite describirla en los mismos o parecidos términos en los que describe la de los espacios métricos o los que verifican el primer axioma de numerabilidad. Para ello se examinará la validez, en espacios secuenciales, de alguno de los resultados enunciados al inicio de la presente acta.

En primer lugar, se tiene la caracterización de los conjuntos abiertos como secuencialmente abiertos puesto que así se han definido los espacios secuenciales. Sin embargo, un punto adherente de un subconjunto no es necesariamente límite de una sucesión de puntos del subconjunto, como sí sucedía en los espacios métricos y primero numerables. Los espacios que verifican esta ́ltima equivalencia son los llamados espacios de Fréchet-Urysohn y tienen interés por sí mismos. Se trata de una clase de espacios situados entre los primeros numerables y los espacios secuenciales.

En cuanto a la equivalencia entre continuidad y continuidad secuencial, se tiene lo siguiente:

Proposición 9. *Sean X un espacio secuencial, Y un espacio topoĺgico arbitrario y $f: X \rightarrow Y$ una aplicación. Entonces, f es secuencialmente continua si, y solo si, f es continua.*

Sin embargo, en las condiciones de la proposición anterior, sucede algo sorprendente. No es cierto que si f es secuencialmente continua en un punto $x \in X$, f sea continua en x . Pero, ¿cómo es posible que en cierto sentido la continuidad equivalga a la continuidad secuencial pero una aplicación secuencialmente continua en un punto pueda ser no continua en el mismo? Esta discrepancia no debería extrańar

puesto que, como es sabido, incluso para funciones reales de variable real es preciso utilizar una versión numerable del axioma de elección para demostrar que secuencialmente continua en un punto implica continua en el punto.

En cuanto a la unicidad del límite y el carácter Hausdorff, es conocido el hecho que a continuación se expone.

Proposición 10. *Todo espacio que verifique el primer axioma de numerabilidad y en el cual cada sucesión convergente tenga un único límite, es de Hausdorff.*

Resulta natural, por lo tanto, preguntarse si la propiedad de ser secuencial puede substituir al primer axioma de numerabilidad en la proposición anterior. La respuesta en general es no. Sin embargo, añadiendo la condición de ser localmente numerablemente compacto sobre el espacio de partida, si se podría afirmar el carácter Hausdorff del mismo.

Para finalizar, en lo que concierne a la equivalencia entre compacidad numerable y secuencial, se tiene lo siguiente.

Proposición 11. *Para un espacio X secuencial y paracompacto equivalen:*

1. X es compacto.
2. X es numerablemente compacto.
3. X es secuencialmente compacto.

Aunque el tema expuesto lleva siendo tratado mucho tiempo (desde los años cincuenta), todavía siguen apareciendo nuevas publicaciones hoy en día. La mayoría de ellas se centran en estudiar bajo qué condiciones de los factores, el producto de espacios secuenciales es secuencial. Con todo y como conclusión, se observa que pese a que las sucesiones no permiten caracterizar la topología un espacio topológico arbitrario, resultan de gran utilidad en una amplia clase de espacios, como es la clase de los espacios secuenciales.

Bibliografía

- [1] Franklin, S. P. (1965). *Spaces in which sequences suffice*. Fund. Math. **57** , pp. 107-115.
- [2] Franklin, S. P. (1967). *Spaces in which sequences suffice II*. Fund. Math. **61** , pp. 51-56.

Selección y asignación de recursos para la contención de un incendio forestal

Jorge Rodríguez Veiga

Departamento de Estadística e Investigación Operativa

12 de noviembre de 2014

Introducción

El diseño de los sistemas de apoyo a la decisión en logística es un área muy activa de investigación y aplicaciones en la moderna investigación operativa. En este marco, para el control de incendios forestales es esencial tomar decisiones eficientes debido a que los recursos son limitados. En este sentido, la teoría económica juega un papel central en la gestión de los incendios forestales. Precursores en el estudio económico de los incendios forestales fueron Headley [3] y Sparhawk [6], que describen como establecer una gestión óptima de manejo de un incendio.

El marco teórico utilizado para identificar la forma más eficiente de administrar los costes de un incendio forestal fue el Cost Plus Net Value Change ($C + NVC$) (Gorte y Gorte [2]). En él, se pretende minimizar el coste para el uso de los recursos en la lucha contra un incendio, más un coste producido por las hectáreas de terreno quemadas, donde no sólo se debe tener en cuenta las pérdidas materiales producidas por el incendio (árboles, bienes urbanos, etc.), sino también la repoblación o la reconstrucción de estas zonas.

En el año 2013 nace el proyecto LUMES, cuyo objetivo principal es el desarrollo de nuevas tecnologías avanzadas para la lucha integral contra los grandes incendios forestales, lo que permite reducir la superficie quemada, el número de incendios y la generación de un recinto de seguridad en las operaciones que reducen significativamente la tasa de accidentes de los participantes (técnicos, brigadistas y pilotos).

Una de las tareas de este proyecto es determinar una selección óptima de recursos aéreos, que incluya el número y tipo de recursos necesarios para la contención de un incendio forestal. Ésta no es una tarea sencilla, pues se cuenta con un gran número de posibilidades. La dificultad se incrementa si la tarea no consiste únicamente en seleccionar el conjunto de recursos, sino que también estos se han de asignar a los diferentes periodos de tiempo en el transcurso del incendio.

Se expondrán dos modelos diferentes. En el primero, y como introducción al planteamiento del problema, se estudia y trabaja con un modelo tomado del trabajo realizado por Donovan y Rideout [1]. Este modelo tiene como objetivo determinar el número de recursos necesarios para contener un incendio desde un instante inicial

PALABRAS CLAVE: programación lineal entera; gestión de recursos aéreos; incendios forestales.

(cuando aún no se ha realizado el despliegue de ningún tipo de recurso), minimizando los costes y daños relacionados con el fuego. Al identificar la combinación óptima de los recursos de extinción de incendios, los modelos aplican la teoría C+NVC, mencionada anteriormente, al incendio.

El segundo modelo se realiza tras estudiar las necesidades demandadas por la empresa de servicios aéreos del proyecto. En éste no sólo debe obtenerse el número necesario de aeronaves para contener el incendio, sino que se deben asignar con precisión a los distintos periodos de tiempo. Además, a petición de la empresa demandante, deberán tenerse en consideración ciertas restricciones sobre el tiempo de empleo de los recursos.

Modelo de selección de recursos

El problema de determinar la selección de recursos que consigan contener el incendio desde un instante inicial (que se traduce en construir una línea alrededor del incendio) a mínimo coste (C+NVC) se plantea como un problema de programación lineal entera, pues los medios de lucha son unidades indivisibles.

Comenzamos introduciendo la notación necesaria, para posteriormente formular el problema.

Los parámetros del modelo se podrán dividir en referentes a los recursos y referentes al incendio. Respecto a los recursos, tendremos para cada uno de ellos, el coste por hora (C_i), coste fijo (P_i), rendimiento (PR_i) y tiempo que tarda en llegar al incendio (A_i). Donde i hace referencia al índice del recurso, $i \in \mathcal{I} = \{1, \dots, n\}$, siendo n el número de recursos disponibles.

Los parámetros referentes al incendio serán cada uno de los periodos en los que se tiene una estimación en la evolución del mismo (H_j) y para cada uno de esos periodos, el incremento del perímetro (PER_j), incremento del NVC (NVC_j) y perímetro (SP_j). Donde j es el índice de cada periodo de tiempo, $j \in \mathcal{J} = \{1, \dots, m\}$, siendo m el número de periodos.

Como variables de decisión binarias, tendremos una variable Z_i que nos indicará con 1 si la aeronave i es seleccionada, D_{ij} que tomará el valor 1 si el recurso i se emplea hasta el periodo j e, Y_j que toma el valor 1 cuando el incendio no está contenido en el instante j . Como variables reales, tenemos L_j que será el perímetro construido por los recursos hasta el periodo j (km).

La formulación matemática siguiente modela la función objetivo (suma de los costes) y las limitaciones que se imponen para identificar la asignación óptima en el problema de la selección de recursos para la contención de un incendio forestal. Resumidamente, estas limitaciones/restricciones serán que el incendio se ha de acotar en un periodo de tiempo concreto (restricción (2)), y tres restricciones para establecer cuando las aeronaves son seleccionadas o no (restricción (3)) y cuando el incendio está o no contenido (restricciones (4) y (5)).

$$\min \quad \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} C_i H_j D_{ij} + \sum_{i \in \mathcal{I}} P_i Z_i + \sum_{j \in \mathcal{J}} NVC_j Y_{j-1} \quad (1)$$

Sujeto a:

$$\sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} (H_j - A_i) PR_i D_{ij} \geq \sum_{j \in \mathcal{J}} PER_j Y_{j-1} \quad (2)$$

$$\forall i \in \mathcal{I}, \sum_{j \in \mathcal{J}} D_{ij} \leq Z_i \quad (3)$$

$$\forall j \in \mathcal{J}, SP_j N_j - L_j \leq M Y_j \quad (4)$$

$$\forall j \in \mathcal{J}, \sum_{i \in \mathcal{I}} (H_j - A_i) PR_i D_{ij} = L_j \quad (5)$$

Modelo de selección y asignación de recursos

Con el fin de atender posibles requisitos alternativos que atañen a las aeronaves encargadas de la extinción del incendio, se modifica el modelo anterior con el fin de que éste tenga en consideración aspectos tales como la asignación de los recursos a lo largo de diferentes periodos de tiempo, de forma que no todos los recursos seleccionados han de operar desde el instante inicial, sino que pueden incorporarse al incendio en cualquier instante en el cual se precisen; que el algoritmo debe de poder ejecutarse en cualquier instante de tiempo, no únicamente en el momento de detección del incendio; si el incendio no está controlado, debe de haber al menos una aeronave en cada frente, y se han de incorporar restricciones en el tiempo de uso de las aeronaves (*Circular Operativa 16-B* [5]).

Se harán unas observaciones iniciales, pues los periodos de tiempo se construirán tomando intervalos de 10 en 10 minutos (para poder ajustar los descansos y tiempos de vuelo con cada intervalo). A partir de estos datos, se obtienen los periodos de tiempo H_j con los que se va a trabajar, que irán desde un periodo temporal anterior al actual (H_0) hasta el último momento del que se tiene una estimación de la evolución del incendio (H_m). En caso de que no se tenga estimación de la evolución del incendio en algún periodo H_j , se tomarán para ese periodo, el perímetro del siguiente periodo temporal del cual se tenga una estimación, y para el incremento del área, la parte proporcional al incremento entre la anterior estimación que se tenga y la siguiente.

En la notación necesaria para formular el problema obviaremos los parámetros y variables descritos anteriormente. De este modo, tendremos como parámetros referentes a los recursos, los tiempos de utilización de las aeronaves: tiempo de vuelo sin descansos (TV_i), tiempo de descanso (TD_i) y tiempo de vuelo diario ($TVDi$). Luego, para establecer la situación actual de la aeronave en el momento en el que se ejecuta el modelo, tendremos un parámetro para establecer el tiempo que lleva volado la aeronave desde su último descanso (TA_i), el número periodos de descanso que lleva realizados ($TDMI_i$) y el tiempo de vuelo diario realizado (TVT_i). Finalmente, también se tiene en cuenta como parámetro el número de periodos de tiempo que comprende un descanso (TDM_i).

Referente a los parámetros del incendio, tendremos a mayores: el número de frentes del incendio (nF), que servirá para determinar el número mínimo de aero-

naves que han de estar presentes en cada periodo de tiempo, y el número máximo de aeronaves que se pueden gestionar en éste ($nMax$).

Por último, las variables de decisión que tendremos en cuenta en el modelo (sin contar las ya descritas en el modelo anterior), harán referencia a la situación de las aeronaves: como variables binarias tendremos, B_{ij} que toma el valor 1 cuando el recurso i se selecciona para entrar en el incendio en el periodo j ; VU_{ij} que indica con 1 si la aeronave i se encuentra volando al incendio en el periodo j ; DES_{ij} para determinar los periodos j en los que la aeronave i ha de realizar el descanso, y FD_{ij} que toma el valor 1 cuando el descanso de la aeronave i finaliza en el periodo j . Por último, tendremos una variable entera, μ_j , que toma el número de recursos faltantes hasta alcanzar el mínimo de recursos que tiene que haber en el incendio en el periodo j .

A continuación se describe la formulación matemática que modela la función objetivo y las restricciones para cumplir la demanda de la empresa:

$$\begin{aligned} \text{mín} \quad & \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} C_i H_{j+1} D_{ij} - \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} H_j C_i B_{ij} - \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} (H_{j+1} - H_j) C_i DES_{ij} + \sum_{i \in \mathcal{I}} P_i Z_i + \\ & \sum_{j \in \mathcal{J}} NVC_j Y_{j-1} + Y_m + \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} VU_{ij} + \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} 0,1 H_j B_{ij} + \sum_{j \in \mathcal{J}} M' \mu_j \end{aligned} \quad (6)$$

Sujeto a :

$$\begin{aligned} \sum_{j \in \mathcal{J}} PER_j Y_{j-1} \leq & \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} (H_{m+1} - H_j - A_i) PR_i B_{ij} - \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} (H_{m+1} - H_{j+1}) PR_i D_{ij} - \\ & \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} (H_{j+1} - H_j) PR_i DES_{ij} \end{aligned} \quad (7)$$

$$\forall i \in \mathcal{I}, \forall t \in \mathcal{T}, \quad \sum_{k=1}^t B_{ik} - DES_{it} - \sum_{k=1}^{t-l_i} B_{ik} \leq VU_{it} \quad (8)$$

$$\forall t \in \mathcal{T}, \quad \sum_{i \in \mathcal{I}} \sum_{k=1}^t B_{ik} - \sum_{i \in \mathcal{I}} VU_{it} - \sum_{i \in \mathcal{I}} DES_{it} - \sum_{i \in \mathcal{I}} \sum_{k=1}^{t-1} D_{ik} + \mu_t \geq nF Y_{t-1} \quad (9)$$

$$\forall i \in \mathcal{I}, \quad \sum_{t \in \mathcal{T}} D_{it} \leq Z_i \quad (10)$$

$$\forall i \in \mathcal{I}, \quad \begin{cases} B_{i1} + \sum_{t=2}^m (m+1) B_{it} \leq m Z_i & , \text{ si } A_i = 0 \\ \sum_{t \in \mathcal{T}} B_{it} \leq Z_i & , \text{ en otro caso} \end{cases} \quad (11)$$

$$\forall i \in \mathcal{I}, \forall t \in \mathcal{T}, \quad \begin{cases} TV_i \geq \sum_{k=1}^t (H_{t+1} - H_k + TA_i - TDMI_i (H_{k+1} - H_k)) B_{ik} - \sum_{k=1}^t (H_t - H_k) D_{ik} - \sum_{k=1}^t \frac{TD_i}{TDM_i} DES_{ik} - \sum_{k=1}^t TV_i FD_{ik} & , TA_i = 0 \text{ ó } TA_i \geq 2 \\ TV_i \geq (H_{t+1} - H_1 + TA_i - TDMI_i (H_{k+1} - H_k)) B_{i1} + \sum_{k=2}^t (TV_i + H_{t+1} - H_t) B_{ik} - \sum_{k=1}^t (H_t - H_k) D_{ik} - \sum_{k=1}^t \frac{TD_i}{TDM_i} DES_{ik} - \sum_{k=1}^t TV_i FD_{ik} & , \text{ en otro caso} \end{cases} \quad (12)$$

$$\forall i \in \mathcal{I}, \forall t \in \mathcal{T}, \quad \left\{ \begin{array}{l} 0 \leq \sum_{k=1}^t (H_{t+1} - H_k + TA_i - TDMI_i(H_{k+1} - H_k)) B_{ik} - \\ \sum_{k=1}^t (H_t - H_k) D_{ik} - \sum_{k=1}^t \frac{TD_i}{TDM_i} DES_{ik} - \sum_{k=1}^t TV_i FD_{ik} \quad , TA_i = 0 \text{ ó } \\ TA_i \geq 2 \\ \\ 0 \leq (H_{t+1} - H_1 + TA_i - TDMI_i(H_{k+1} - H_k)) B_{i1} + \\ \sum_{k=2}^t (TV_i + H_{t+1} - H_t) B_{ik} - \sum_{k=1}^t (H_t - H_k) D_{ik} - \\ \sum_{k=1}^t \frac{TD_i}{TDM_i} DES_{ik} - \sum_{k=1}^t TV_i FD_{ik} \quad , \text{ en otro caso} \end{array} \right. \quad (13)$$

$$\forall i \in \mathcal{I}, \forall t \in \mathcal{T}, \quad \left\{ \begin{array}{l} \sum_{k=t-TDM_i+1}^t DES_{ik} \geq TDM_i FD_{it} \quad , \text{ si } t - TDM_i \geq 0 \\ \\ \sum_{k=1}^t DES_{ik} \geq (TDM_i - TDMI_i) FD_{it} \quad , \text{ en otro caso} \end{array} \right. \quad (14)$$

$$\forall i \in \mathcal{I}, \forall t \in \mathcal{T}, \quad DES_{it} \leq \sum_{k=t}^{t+TDM_i} FD_{ik} \quad (15)$$

$$\forall t \in \mathcal{T}, \quad \begin{aligned} M Y_t \geq & - \sum_{i \in \mathcal{I}} \sum_{k=1}^t (H_{t+1} - H_k - A_i) PR_i B_{ik} + \sum_{i \in \mathcal{I}} \sum_{k=1}^t (H_{t+1} - H_{k+1}) PR_i D_{ik} + \\ & \sum_{i \in \mathcal{I}} \sum_{k=1}^t (H_{k+1} - H_k) PR_i DES_{ik} + SP_t Y_{t-1} \end{aligned} \quad (16)$$

$$\forall i \in \mathcal{I}, \quad \sum_{t \in \mathcal{T}} H_{t+1} D_{it} \geq \sum_{t \in \mathcal{T}} H_{t+1} B_{it} \quad (17)$$

$$\forall i \in \mathcal{I}, \quad \sum_{t \in \mathcal{T}} (H_{m+1} - H_t) B_{it} - \sum_{t \in \mathcal{T}} (H_{m+1} - H_{t+1}) D_{it} \leq TV D_i - TV T_i \quad (18)$$

$$\forall i \in \mathcal{I}, \quad \begin{aligned} A_i Z_i \leq & \sum_{t \in \mathcal{T}} (H_{m+1} - H_t) B_{it} - \sum_{t \in \mathcal{T}} (H_{m+1} - H_{t+1}) D_{it} - \\ & \sum_{t \in \mathcal{T}} (H_{t+1} - H_t) DES_{it} \end{aligned} \quad (19)$$

$$\forall i \in \mathcal{I}, \forall t \in \mathcal{T}, \quad \sum_{k=1}^t B_{ik} - \sum_{k=1}^t D_{ik} \leq DES_{it} \quad (20)$$

$$\forall t \in \mathcal{T}, \quad \sum_{i \in \mathcal{I}} \sum_{k=1}^t B_{ik} - \sum_{i \in \mathcal{I}} \sum_{k=1}^{t-1} D_{ik} \leq nMax Y_{t-1} \quad (21)$$

siendo $l_i = \{t \in \mathcal{T} : H_{t+1} = H_2 + A_i\} + N^\circ$ periodos descanso durante viaje en la ecuación (9); y M, M' valores suficientemente grandes, como por ejemplo:

$$M = \max_{t \in \mathcal{T}} SP_t + \sum_{i \in \mathcal{I}} (H_{m+1} - H_2) PR_i,$$

$$M' = 100 \left((H_{m+1} - H_1) \sum_{i \in \mathcal{I}} C_i + \sum_{t \in \mathcal{T}} NV C_t \right).$$

La función objetivo (6) establecerá algo análogo a la función objetivo (1), salvo que además se incluirán ciertos términos para tener control sobre las variables Y_m , VU_{ij} y B_{ij} , y otro término (el último), que penalizará la función objetivo en caso de que no se alcance el número mínimo de aeronaves en el incendio en un periodo concreto de tiempo.

La restricción (7) sería análoga a la (2); la (10), (11) y (17) a la (3), y la (16) a la (5). Para establecer el número máximo y mínimo de aeronaves se establecen las restricciones (21) y (9) respectivamente. Y finalmente, las restricciones restantes, (8), (12), (13), (14), (15), (18), (19), (20), se incorporan para obligar a fijar los periodos de vuelo máximo diario, en qué periodos de tiempo las aeronaves estarán volando hacia el incendio, en cuáles estarán trabajando sobre éste y los periodos en los que las aeronaves deberán realizar los descansos oportunos.

Líneas futuras

Una línea de desarrollo, es la incorporación de estocasticidad en el modelo. Mientras que en un problema determinístico de programación matemática, todos los parámetros que aparecen en su formulación son números conocidos, en programación estocástica dichos parámetros (o por lo menos algunos de estos) son desconocidos, aunque para ellos se conoce o se puede estimar su distribución de probabilidad.

En este sentido estuvo estudiándose el trabajo de tesis de Lee [4], basado en la misma idea que el primer modelo determinista (Donovan y Rideout [1]), pero aumentando su alcance al considerar las características de crecimiento del incendio, tales como perímetro y área quemada, estocásticas.

Bibliografía

- [1] Donovan, G. and Rideout, D. (2003). *An integer programming model to optimize resource allocation for wildfire containment*, Forest Science, **49**(2), pp. 331–335.
- [2] Gorte, J. K. and Gorte, R. W. (1979). *Application of Economic Techniques to Fire Management-A status Review and Evaluation*, USDA Forest Service General Technical Report, INT-53.
- [3] Headley, R. (1916). *Fire Suppression, District 5*, United States Department of Agriculture, Forest Service, Washington.
- [4] Lee, W. J. (2006). *A Stochastic Mixed Integer Programming Approach to Wild-fire Management Systems*, Tesis doctoral, Texas A&M University.
- [5] Ministerio de Fomento, Dirección General de Aviación Civil, España (2001). *Anexo nº 1 a Circular Operativa 16-B*, Madrid.
- [6] Sparhawk, W. N. (1925). *The use of liability ratings in planning forest fire protection*, Journal of Agricultural Research, **30** (8), pp 693–792.

Macedonia fraccionaria

Jorge Losada Rodríguez

Departamento de Análise Matemática

26 de novembro de 2014

O cálculo fraccionario é unha activa rama da análise matemática na que podemos falar de integrais e derivadas de orde non enteira. O nome é enganoso, pois a xeneralización non se limita unicamente a ordes racionais (ou fraccionarias). Tamén é válida para integrais e derivadas de orde un número real calquera.

Durante un longo tempo, esta disciplina foi considerada esotérica e sen aplicacións. Sen embargo, nas últimas décadas, o interese polos operadores diferenciais e integrais de orde non enteira aumentou de xeito notable, probándose ademais a súa utilidade en multitude de situacións. A pesares disto, o cálculo fraccionario segue a ser un descoñecido para parte da comunidade matemática e científica.

No fondo, as ideas esenciais do cálculo fraccionario son comúns a moitas áreas da matemática e daquela, os resultados deberían estar enlazados entre si.

Introdución

Semella que a nosa maneira de entender e analizar o mundo que nos rodea evoluciona do *enteiro*, ou algo aínda máis particular, cara o *fraccionario* ou incluso cara un concepto máis xeral.

Lembremos, cando cativos, dánsenos de inicio os números naturais, mais deseguido xorden –por estrita necesidade– os números enteiros. Anos máis tarde, cando precisamos da división, preséntansenos os números racionais ou fraccionarios, que xunto cos misteriosos irracionais, forman o corpo dos números reais. Conxunto que será logo estendido aos números complexos.

Desde un punto de vista xeométrico, ocorre esencialmente o mesmo. Inicialmente, aparecen os conceptos de punto, curva, superficie e volume; que están asociados ás dimensións enteiras: nula, unha, dúas e tres, respectivamente. Non obstante, os relativamente recentes traballos de B. Mandelbrot (1924–2010), poñen de manifesto a importancia de conceptos xa introducidos antes historicamente, coma o de dimensión de Hausdorff, dimensión que poder tomar como valor calquera número real positivo. Así, a xeometría amplía os seus horizontes e o concepto de dimensión acada un novo significado, que non deixa nunca de ser coherente cos resultados do pasado.

Na análise matemática, onde as ben coñecidas derivadas e integrais xogan un papel fundamental, debemos ser tamén conscientes de que estes operadores tan só

son axeitados para un conxunto moi selecto de funcións. De feito, esta observación é a motivación fundamental das dúas grandes revolucións no cálculo do século XX: a teoría da medida por un lado que, con H. Lebesgue como protagonista, xeneraliza a idea de integral e doutra banda, a teoría das distribucións de L. Schwartz, que estende a idea orixinal de derivada.

Non obstante, tamén poderíamos seguir o sendeiro antes indicado, e intentar definir derivadas e integrais de orde non enteira que, por suposto, xeneralicen os ben coñecidos operadores diferenciais e integrais de orde enteira. Este é o camiño do cálculo fraccionario!

Se cadra, despois poderemos falar de funcións continuas sen derivada en ningún punto (pénsese na función de Weierstrass), pero con derivada de certa orde α , con $0 < \alpha < 1$.

Como definir unha integral de orde non enteira

Como xa dixemos, o obxectivo do cálculo fraccionario é interpolar aos operadores diferenciais e integrais de orde enteira. Así, seguindo as ideas de B. Riemann e J. Liouville, se comezamos analizando o que ocorre ao iterar o operador integral sobre unha función f axeitada, observamos que, integrando por partes e procedendo despois por indución:

$$\begin{aligned} I_a^1 f(t) &= \int_a^t f(s_1) ds_1, \\ I_a^2 f(t) &= \int_a^t \int_a^{s_1} f(s_2) ds_2 ds_1 = \int_a^t (t-s)f(s) ds, \\ &\vdots \\ I_a^n f(t) &= \int_a^t \int_a^{s_1} \cdots \int_a^{s_{n-1}} f(s_n) ds_n \cdots ds_1 = \frac{1}{(n-1)!} \int_a^t (t-s)^{n-1} f(s) ds; \end{aligned}$$

onde $a \in \mathbb{R}$ e $n \in \mathbb{N}$. Daquela, empregando a función gamma para xeneralizar ao factorial da derradeira expresión, obtemos a seguinte definición.

Definición 1. Sexan $\alpha > 0$, $a \in \mathbb{R}$ e $f: [a, +\infty) \rightarrow \mathbb{R}$ unha función. A integral, segundo Riemann-Liouville, de orde α da función f con punto base a ven dada por,

$${}^{\text{RL}}I_a^\alpha f(t) = \frac{1}{\Gamma(\alpha)} \int_a^t (t-s)^{\alpha-1} f(s) ds, \quad t > a.$$

Observación 2. Obviamente, se $\alpha = n \in \mathbb{N}$, entón ${}^{\text{RL}}I_a^\alpha f$ coincide co operador integral iterado n veces sobre a función f .

A composición de operadores integrais de orde fraccionaria é unha cuestión importante. Pode probarse que se $\alpha, \beta > 0$, entón ${}^{\text{RL}}I_a^\alpha {}^{\text{RL}}I_a^\beta f = {}^{\text{RL}}I_a^{\alpha+\beta} f$.

Consideremos agora o seguinte espazo de funcións.

Definición 3. *Sexa $[a, b]$ un intervalo de $\mathbb{R}$. Diremos que $f: [a, b] \rightarrow \mathbb{R}$ é unha función de Hölder con expoñente $0 < \lambda < 1$ se*

$$|f(x_1) - f(x_2)| < A |x_1 - x_2|^\lambda \quad \text{para todo } x_1, x_2 \in [a, b],$$

onde A é unha constante real. Denotaremos a tal conxunto por $H^\lambda[a, b]$.

Observación 4. *Se $f \in H^\lambda[a, b]$ para algún $0 < \lambda < 1$, entón f é unha función continua. Tense tamén, que certas función fractais, coma a de Weierstrass, son funcións de Hölder para un valor do expoñente λ moi relacionado coa súa dimensión de Hausdorff, que lembremos, é un número entre 0 e 1.*

Pois ben, dados $0 < \lambda < 1$ e $f \in H^\lambda[a, b]$, pode probarse que se $\alpha + \lambda < 1$, entón ${}^{\text{RL}}I_a^\alpha f \in H^{\lambda+\alpha}$. É dicir, as integrais fraccionarias de Riemann-Liouville aumentan a regularidade das funcións nos espazos de Hölder.

Derivadas de orde non enteira

Agora, despois da noción de integral *fraccionaria* de orde $\alpha > 0$, a idea de derivada de orde α é imprescindible.

Asociando a derivada de orde $\alpha > 0$ coa integral de orde $-\alpha$, un síntese tentado a substituír α por $-\alpha$ na Definición 1. A pouco que nos fixemos, darémonos conta de que iso non conduce a bo destino, pois obteremos unha integral diverxente. Así, a definición de derivada de orde $\alpha > 0$ debe ser máis coidadosa.

Definición 5. *Sexan $\alpha > 0$, $a \in \mathbb{R}$ e $f: [a, +\infty) \rightarrow \mathbb{R}$ unha función. A derivada, segundo Riemann-Liouville, de orde α da función f con punto base a ven dada por,*

$$\begin{aligned} {}^{\text{RL}}D_a^\alpha f(t) &\equiv {}^{\text{RL}}I_a^{-\alpha} f(t) = D^n {}^{\text{RL}}I_a^{n-\alpha} f(t) \\ &= \frac{1}{\Gamma(n-\alpha)} \frac{d^n}{dt^n} \int_a^t (t-s)^{n-\alpha-1} f(s) ds, \quad t > a; \end{aligned}$$

onde $n = \lceil \alpha \rceil$, sendo $\lceil \alpha \rceil$ o menor número natural maior ou igual que α .

Observación 6. *Dado $n \in \mathbb{N}$, para obter a derivada de orde $n-1 < \alpha \leq n$ dunha función, calculamos primeiro a integral fraccionaria de orde $n-\alpha$ e despois, derivamos n veces. Polo tanto, segundo Riemann-Liouville, a derivada de orde $\alpha > 0$ dunha función é unha derivada de orde enteira dunha integral de orde fraccionaria.*

Nótese que de acordo coa Definición 5, a derivada dunha función constante de orde $\alpha > 0$, $\alpha \notin \mathbb{N}$, non é a función identicamente nula.

Doutra banda, agora que xa sabemos calcular integrais e derivadas de orde arbitraria, podemos plantexar ecuacións diferenciais de orde fraccionaria coas que intentar perfeccionar certos modelos físicos. A maior dificultade aparece á hora de impor as condicións iniciais axeitadas que garantan a unicidade da solución, pois deben ser dadas de xeito non fisicamente interpretable a día de hoxe.

Por riba disto, a resolución de ecuacións diferenciais de orde fraccionaria é realmente complexa. Pois, por exemplo, a transformada de Laplace (moi adecuada para a resolución de ecuacións diferenciais de orde enteira) perde parte da súa utilidade.

As trabas comentadas anteriormente, levaron ao físico e matemático italiano M. Caputo a propor unha nova definición de derivada de orde fraccionaria $\alpha > 0$. A súa idea é moi sinxela: se Riemann e Liouville primeiro integran para logo derivar, Caputo opera en orde inversa.

Definición 7. Sexan $\alpha > 0$, $a \in \mathbb{R}$ e $f: [a, +\infty) \rightarrow \mathbb{R}$ unha función. A derivada, segundo Caputo, de orde $\alpha > 0$ da función f con punto base a defínese como,

$$\begin{aligned} {}^c D_a^\alpha f(t) &\equiv {}^{\text{RL}} I_a^{-\alpha} f(t) = {}^{\text{RL}} I_a^{n-\alpha} D^n f(t) \\ &= \frac{1}{\Gamma(n-\alpha)} \int_a^t (t-s)^{n-\alpha-1} f^{(n)}(s) ds, \quad t > a; \end{aligned}$$

onde, novamente, $n = \lceil \alpha \rceil$.

Observación 8. Para calcular agora a derivada de orde $\alpha > 0$ dunha función, é preciso que exista a súa derivada de orde $\lceil \alpha \rceil \in \mathbb{N}$. Nótese que $\lceil \alpha \rceil \geq \alpha$. Así, por exemplo, se queremos calcular unha semiderivada (derivada de orde $1/2$) de f , é preciso que exista f' e esta existencia non parece, en principio, moi coherente.

Notemos tamén que a derivada, segundo Caputo, dunha función constante é, para calquera orde $\alpha > 0$, a función idénticamente nula.

A derivada fraccionaria de Caputo, ao igual que a integral e a derivada fraccionaria de Riemann-Liouville, carece hoxe en día dunha interpretación física ou xeométrica clara e aceptada pola maioría da comunidade matemática. Sen embargo, a súa principal ventaxa respecto da proposta de Riemann-Liouville, radica nas aplicacións.

Cando consideramos ecuacións diferenciais de orde fraccionaria coa derivada de Caputo, as condicións iniciais poden ser expresadas empregando derivadas de orde enteira e ademais, a transformada de Laplace recobra, nesta situación, a súa grande utilidade.

A característica esencial tanto da derivada fraccionaria de Riemann-Liouville, coma da de Caputo, é que son operadores globais. Lembremos que as derivadas usuais de orde enteira son operadores locais, isto é, a derivada n -ésima dunha función nun punto, depende dos valores de tal función nunha veciñanza arbitrariamente pequena de tal punto. Sen embargo, as derivadas de Riemann-Liouville e Caputo teñen en conta a historia da función, pois o seu valor nun punto t depende dos valores da función no intervalo $[a, t]$. É por isto que se di que *as derivadas de orde fraccionaria teñen memoria*.

A propiedade que vimos de comentar, fai que os modelos con ecuacións diferenciais de orde fraccionaria sexan especialmente útiles na análise de procesos afectados polo pasado, así coma no estudo de propiedades de certos materiais viscoelásticos ou viscoplásticos.

É tamén natural e útil preguntarse pola relación entre as derivadas e as integrais de orde fraccionaria, serán operadores inversos? Así, coma no caso de orde enteira sábese que para $\alpha > 0$

$${}^{\text{RL}}D_a^\alpha [{}^{\text{RL}}I_a^\alpha f](t) = f(t) \quad \text{e} \quad {}^{\text{C}}D_a^\alpha [{}^{\text{RL}}I_a^\alpha f](t) = f(t), \quad \text{para } t > a;$$

mais, se por exemplo, $0 < \alpha < 1$ entón,

$${}^{\text{RL}}I_a^\alpha [{}^{\text{RL}}D_a^\alpha f](t) = f(t) + c_1 (t - a)^{\alpha-1} \quad \text{e} \quad {}^{\text{RL}}I_a^\alpha [{}^{\text{C}}D_a^\alpha f](t) = f(t) + c_2,$$

onde $c_1, c_2 \in \mathbb{R}$ son constantes que dependen da función f considerada.

Nas referencias [3, 4] pode consultarse unha introdución amena e completa do cálculo fraccionario.

Sistemas de numeración anormais

O sistema de numeración decimal que empregamos a diario é ben coñecido. De igual xeito, os sistemas binario ou ternario, non presentan ningunha dificultade e así, podemos traballar, sen maior complicación, en sistemas de numeración posicionais nos que a base sexa un número natural calquera.

Pero... por que non considerar un sistema de numeración posicional no que a base sexa un número real $\alpha > 1$ calquera? Tal cuestión está, en certo sentido, moi relacionada con todo o anterior. Explícome; unha vez familiarizados coas dimensións enteiras, preguntámonos pola interpretación dunha dimensión fraccionaria, e aparece así a xeometría fractal e os obxectos fractais, entre os que se atopa o grafo da xa citada función de Weierstrass. Desde o punto de vista da análise, o cálculo fraccionario intenta dar resposta á curiosidade polo significado dunha derivada (ou integral) de orde non enteira. Logo a pregunta anterior podémola pensar como o xerme da álgebra (ou teoría de números) fraccionaria unha vez superado os conceptos de número fraccionario e número irracional, orixe verdaeiro de todas as cuestións anteriores. Voltemos logo á álgebra!

Comecemos reproducindo ideas clásicas. Dado un número $c > 0$, teríamos que

$$c = (a_n a_{n-1} \cdots a_0 \cdot a_{-1} a_{-2} \cdots)_\alpha,$$

se, e só se,

$$c = a_n \alpha^n + a_{n-1} \alpha^{n-1} + \cdots + a_0 + a_{-1} \alpha^{-1} + a_{-2} \alpha^{-2} + \cdots,$$

onde para cada $i = n, n-1, \dots, 1, 0, -1, \dots$ esiximos $a_i \in \{0, 1, \dots, [\alpha]\}$.

Un caso especialmente curioso ocorre, como se pode consultar en [1], cando consideramos como base do sistema de numeración o número de ouro. Isto é, cando

$$\alpha = \tau = \frac{1 + \sqrt{5}}{2}.$$

Trátase do coñecido como τ -system, que foi introducido por G. Bergman cando apenas contaba doce anos de idade.

A propiedade característica do número de ouro

$$\tau^n = \tau^{n-1} + \tau^{n-2}, \quad \text{para cada } n \in \mathbb{Z}, \quad (1)$$

dota a este sistema dunhas regras de cálculo que, en ocasións, son extremadamente sinxelas. Convén notar tamén, que cando empregamos como base un número irracional, entón certos números irracionais teñen unha representación finita.

Tendo en conta que

$$\frac{1 + \sqrt{5}}{2} + \left(\frac{1 + \sqrt{5}}{2} \right)^{-2} = 2,$$

deducimos que $2 = (10,01)_\tau$. Pero, empregando (1), tamén podemos expresar o 2 como $(1,11)_\tau$, $(10,0011)_\tau$, $(10,001011)_\tau$, $(1,101011)_\tau, \dots$ Sen embargo, $2 = (10,01)_\tau$ é a expresión máis cómoda e sinxela; pois non ten dous uns consecutivos e daquela, non é posible aplicar a regra $100 = 011$ deducida de (1).

Así, é fácil ver que para obter a expresión máis sinxela dun número dado no τ -system, o que debemos facer é eliminar todos os uns consecutivos. Empregamos para iso a fórmula indicada anteriormente. Nótese logo, que na expresión máis sinxela en base τ dun número real, nunca poderá haber dous uns consecutivos. Tampouco é difícil probar (ver [2]) que para calquera base $\alpha > 1$, a expresión máis sinxela dun número dado en dita base, sempre existe e é única.

É dicir, cando consideramos sistemas de numeración posicionais con base non enteira, os díxitos da expresión dun número non son independentes entre si, independencia que si é certa se a base é un número enteiro. Parece como se tales sistemas, ao igual que as derivadas de orde non enteira, *tivesen memoria*.

Conclusión

A curiosidade por considerar números reais arbitrarios en ámbitos nos que os enteiros son os usuais, lévanos desde orixes distantes a características e situacións comúns que deben ser investigadas.

Bibliografía

- [1] Bergman, G. (1957). *A number system with an irrational base*, Mathematics Magazine, **31**(2), pp. 98–110.
- [2] Kempner, A. J. (1936). *Anormal systems of numeration*, The American Mathematical Monthly, **43**(10), pp. 610–617.
- [3] Miller, K. S. e Ross, B. (1993). *An Introduction to the Fractional Calculus and Fractional Differential Equations*, J. Wiley & Sons.
- [4] Oldham, K. B. e Spanier, J. (1974). *The Fractional Calculus*, Acad. Press.

Diseñando generadores

Saray Busto Ulloa

Departamento de Matemática Aplicada

10 de diciembre de 2014

La demanda energética en constante crecimiento y la rápida disminución de las fuentes convencionales han inducido el desarrollo de nuevos dispositivos de generación de energía que aprovechan los recursos renovables disponibles en el planeta. En este ámbito del desarrollo de nuevas tecnologías se encuentran los dispositivos de generación de energía de corrientes de eje horizontal (HATT).

El estudio del comportamiento del fluido en torno a estos dispositivos puede ser realizado mediante la dinámica de fluidos computacional (CFD). Para llevar a cabo estas simulaciones, introduciremos las ecuaciones de Navier-Stokes incompresibles, seleccionaremos un modelo de turbulencia e indicaremos los principales algoritmos para el tratamiento del movimiento rotatorio de la geometría.

Ecuaciones de Navier-Stokes

En esta sección, se introducirán las ecuaciones generales de la dinámica de fluidos. A partir de ellas, se deducirán las ecuaciones de Navier-Stokes incompresibles que permitirán modelar el comportamiento del fluido en el entorno del dispositivo de captación de energía de corrientes.

Consideremos un sistema de referencia euleriano y una región Ω en el espacio $\mathbb{R}^3$. Denotemos por ρ la densidad del fluido, $\mathbf{v}$ la velocidad lineal y $\mathbf{n}$ el vector normal a la frontera. Entonces el principio de conservación de la masa afirma que la tasa de variación de la masa en un recinto $\mathcal{W} \subset \Omega$ coincide con el flujo de materia que entra a través de su frontera, es decir,

$$\frac{d}{dt} \int_{\mathcal{W}} \rho dV = - \int_{\partial \mathcal{W}} \rho \mathbf{v} \cdot \mathbf{n} dS.$$

Suponiendo que los campos son suficientemente regulares, aplicando el teorema de Gauss y teniendo en cuenta que la igualdad anterior ha de ser cierta para todo recinto $\mathcal{W}$, se puede concluir que

$$\frac{\partial \rho}{\partial t} - \operatorname{div} \rho \mathbf{v} = 0 \tag{1}$$

en el dominio $\mathcal{W}$.

Por otro lado, la ley de conservación del momento lineal se formula diciendo que la tasa de variación del momento en un recinto coincide con la suma de la resultante de las fuerzas distribuidas y de las fuerzas superficiales actuando sobre él, menos el flujo de momentos que entra a través de la frontera, es decir,

$$\frac{d}{dt} \int_{\mathcal{W}} \rho \mathbf{v} dV = - \int_{\partial \mathcal{W}} \rho \mathbf{v} (\mathbf{v} \cdot \mathbf{n}) dS + \int_{\partial \mathcal{W}} \mathbf{T} \mathbf{n} dS + \int_{\mathcal{W}} \rho \mathbf{f} dV,$$

siendo $\mathbf{T}$ el tensor de tensiones de Cauchy y $\mathbf{f}$ las fuerzas distribuidas. Teniendo en cuenta el teorema de la divergencia y la ecuación de conservación de la masa, se obtiene

$$\frac{\partial \rho \mathbf{v}}{\partial t} + \operatorname{div}(\rho \mathbf{v} \otimes \mathbf{v}) - \operatorname{div} \mathbf{T} = \rho \mathbf{f}.$$

Además, en el caso de fluidos de Coleman-Noll, el tensor de tensiones es función afín del tensor de velocidades de deformación:

$$\mathbf{T} = -\pi \mathbf{I} + \mathbf{T}_v,$$

donde $\mathbf{T}_v$ representa la parte viscosa y π es la presión. Finalmente, asumiendo un fluido newtoniano bajo la hipótesis de Stokes, se llega a la ley constitutiva

$$\mathbf{T} = -\pi \mathbf{I} - \frac{2}{3} \mu \operatorname{div} \mathbf{v} \mathbf{I} + \mu (\operatorname{grad} \mathbf{v} + \operatorname{grad} \mathbf{v}^t),$$

donde μ es la viscosidad dinámica o molecular.

Con el objetivo de completar el sistema anterior sería necesario definir la ecuación de conservación de la energía. Sin embargo, el modelo puede ser simplificado gracias a la consideración de un flujo adiabático y a través del estudio de las ecuaciones adimensionalizadas. Todo ello, permite resolver de manera desacoplada la temperatura y los campos de velocidades y presiones.

Por otro lado, no existe una compresión apreciable del fluido de modo que la ecuación (1) se reduce a

$$\operatorname{div} \mathbf{v} = 0.$$

Además, se observa que el número de Reynolds es del orden de 10^6 , lo que aparentemente, permitiría despreciar los efectos viscosos (llegando a las ecuaciones de Euler para flujos incompresibles). Sin embargo, debemos tener en cuenta la presencia de la capa límite sobre la superficie de las palas, lo que obliga a considerar tanto los términos de inercia como los viscosos.

Por lo tanto, bajo las hipótesis de un fluido con comportamiento newtoniano, en régimen adiabático, a densidad constante, con velocidades subsónicas y en el que las únicas fuerzas distribuidas son las gravitatorias, se obtienen las siguientes ecuaciones del modelo límite adimensionalizadas:

$$\operatorname{div}^* \mathbf{v}^* = 0,$$

$$\frac{\partial \mathbf{v}^*}{\partial t^*} + \operatorname{div}^* (\mathbf{v}^* \otimes \mathbf{v}^*) + \operatorname{grad}^* \pi^* - \frac{1}{Re} \Delta^* \mathbf{v}^* = -\frac{1}{Fr^2} \mathbf{e}_3, \quad (2)$$

denominadas ecuaciones de Navier-Stokes incompresibles adimensionalizadas. Puesto que el problema que queremos resolver no presenta una superficie libre, no será necesario estudiar la componente hidrostática del campo de presiones. De este modo, la ecuación (2) puede ser reescrita como

$$\frac{\partial \mathbf{v}^*}{\partial t^*} + \text{div}^* (\mathbf{v}^* \otimes \mathbf{v}^*) + \text{grad}^* \pi'^* - \frac{1}{Re} \Delta^* \mathbf{v}^* = 0,$$

con π'^* la presión reducida, y sólo se tendrá en cuenta el término gravitatorio si se desea calcular la presión total.

Finalmente, teniendo en cuenta las simplificaciones consideradas, llegamos al modelo límite de las ecuaciones de Navier-Stokes para flujos incompresibles:

$$\text{div} \mathbf{v} = 0,$$

$$\rho \frac{\partial \mathbf{v}}{\partial t} + \rho \text{div} (\mathbf{v} \otimes \mathbf{v}) + \text{grad} \pi - \mu \Delta \mathbf{v} = 0. \quad (3)$$

Modelo de turbulencia

El elevado valor del número de Reynolds nos indica que el flujo presenta un carácter turbulento. Esto obliga a considerar dos escalas de simulación: la del flujo macroscópico, cuyas propiedades queremos determinar, y la de las fluctuaciones, que se producen en una escala correspondiente al espesor de la capa límite. Aunque, en un principio, no nos interesa detallar los fenómenos producidos en la escala reducida, su efecto sobre el flujo medio sí debe ser tenido en cuenta. La pequeña escala a la que se producen las fluctuaciones de los campos de velocidades y de presión imposibilita la simulación directa de la turbulencia, que exigiría unos tamaños de malla demasiado reducidos. Por ello, se hace necesario considerar modelos de turbulencia que permiten determinar la viscosidad turbulenta, es decir, los efectos viscosos producidos por las escalas más pequeñas del fluido sobre el flujo medio, mediante la resolución de un sistema de ecuaciones. Existen distintos tipos de modelos de turbulencia que pueden ser clasificados en RANS (*Reynolds Averaged Navier-Stokes*), LES (*Large Eddy Simulation*) y DES (*Detached Eddy Simulation*).

La potencia computacional de la que se dispone en la actualidad hace que los modelos de turbulencia LES y DES resulten, todavía, demasiado costosos. Por ello, en el desarrollo de este trabajo, se ha decidido trabajar con modelos RANS.

Los modelos de turbulencia tipo RANS se basan en la descomposición del flujo en un promedio estadístico más una fluctuación:

$$\pi = \Pi + \pi', \quad \mathbf{v} = \mathbf{V} + \mathbf{v}'.$$

Sustituyendo en (3) y promediando la ecuación resultante, se obtiene que

$$\rho \frac{\partial \mathbf{V}}{\partial t} + \rho \text{div} (\mathbf{V} \otimes \mathbf{V}) + \text{grad} \Pi - \mu \Delta \mathbf{V} - \text{div} \tau^R = 0,$$

donde

$$\tau^R = -\rho \langle \mathbf{v}' \otimes \mathbf{v}' \rangle$$

se denomina tensor de tensiones de Reynolds.

Con el objetivo de obtener un modelo cerrado, se establece la hipótesis de Bousinesq que asume que, dada la viscosidad dinámica turbulenta μ_t , el tensor de tensiones de Reynolds es de la forma

$$\tau^R = \frac{1}{3} \text{tr}(\tau) \mathbf{I} + 2\mu_t D(\mathbf{V})$$

con

$$\tau = \frac{2}{3} \rho k + 2\mu_t \text{grad} \mathbf{V}, \quad D(\mathbf{V}) = \frac{1}{2} (\text{grad} \mathbf{V} + \text{grad} \mathbf{V}^t)$$

y k la energía cinética turbulenta. El término de la traza del tensor τ , que aparece en la ley anterior, puede ser absorbido por el término de la presión media

$$\Pi^* = \Pi - \frac{1}{3} \text{tr}(\tau).$$

Por lo tanto, para formular un modelo cerrado nos resta proporcionar una expresión para el cálculo de μ_t . En este trabajo, nos centraremos en el modelo de dos ecuaciones $k - \omega$ SST que asume una expresión para la viscosidad turbulenta de la forma

$$\mu_t = \frac{\rho k}{\omega} \frac{1}{\max \left[\frac{1}{\alpha^*}, \frac{S F_2}{a_1 \omega} \right]}$$

con ω la tasa de disipación específica, S la velocidad de deformación y a_1 una constante de cierre. Las ecuaciones de transporte para este modelo, despreciando los términos fuente y los efectos de compresibilidad, se escriben:

$$\frac{\partial \rho k}{\partial t} + \text{div}(\rho k \mathbf{v}) = \text{div}(\Gamma_k \text{grad} k) + \overline{G_k},$$

$$\frac{\partial \rho \omega}{\partial t} + \text{div}(\rho \omega \mathbf{v}) = \text{div}(\Gamma_\omega \text{grad} \omega) + G_\omega + D_\omega.$$

donde Γ_k y Γ_ω son las difusividades efectivas de k y ω , $\overline{G_k}$ la producción de energía cinética turbulenta debida a los gradientes de la velocidad media, G_ω la producción de ω , D_ω el término de difusión cruzada resultante y F_2 una función de mezcla (ver [2] para más detalles del modelo).

Modelización de la rotación

Para finalizar la descripción del modelo de resolución, debemos tener en cuenta las deformaciones y movimientos que puede sufrir la geometría, y con ello la malla. Diferenciaremos dos técnicas distintas que permiten simular el movimiento dinámico de una malla sin necesidad de proceder a su reconstrucción en cada paso de tiempo:

- *Moving Reference Frame* (MRF). Realiza una aproximación en estado estacionario que permite la asignación de velocidades angulares a cada región. Esta técnica puede resultar efectiva como primera aproximación del comportamiento del fluido. Sin embargo, se basa en la consideración de una posición prefijada de la geometría lo que le impide capturar efectos de carácter evolutivo, como serían por ejemplo, el paso de una pala de un rotor por delante de la torre o la influencia de dos rotores próximos girando a distintas velocidades.
- *Sliding Mesh*. Permite la simulación transitoria de un flujo bajo movimientos de traslación y rotación. Se centra en el desplazamiento de una de las regiones frente a la otra y para ello exige la definición de una zona *interface* entre ambas. Un buen dato inicial que proporcionar a esta técnica sería la solución obtenida mediante MRF.

El movimiento de tipo rotatorio de las hélices del dispositivo a modelar, puede entonces ser simulado de forma efectiva mediante *Sliding Mesh* inicializado con MRF.

Resultados numéricos

Con el objetivo de simular el flujo en el entorno de un dispositivo de generación de energía de corrientes (ver Figura 1), empleamos el software comercial Ansys Fluent.

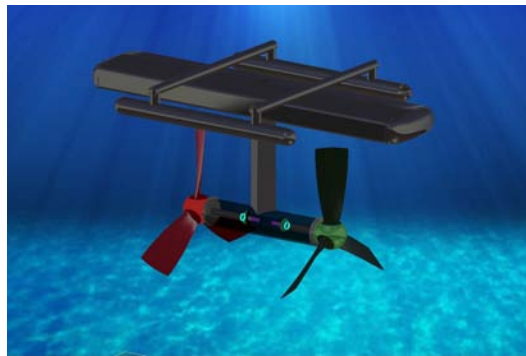


Figura 1: HATT diseñado por la empresa Magallanes Renovables S.L.

A continuación incluimos algunos de los resultados obtenidos al considerar una velocidad de corriente de 1,5 m/s y una velocidad de rotación de 10 rpm (ver [1] para más detalles de las simulaciones realizadas y de los resultados obtenidos).

Estudiando la Figura 2, vemos que la región afectada por el flujo directamente perturbado por el rotor emplazado aguas arriba se sitúa entre las palas del rotor situado aguas abajo y no sobre ellas. De este modo, dichas palas perciben un flujo limpio, lo que justificaría la consideración de unas velocidades de rotación iguales pero de signos opuestos para ambos rotores.

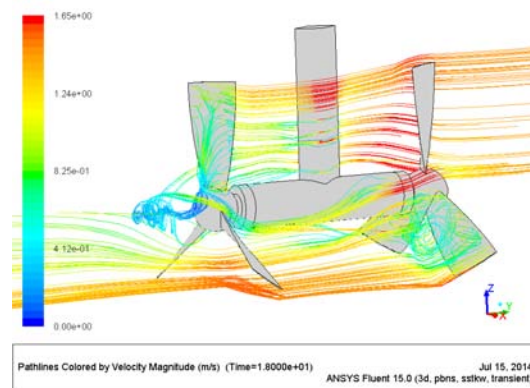


Figura 2: Líneas de corriente coloreadas por la magnitud de la velocidad.

Los momentos ejercidos sobre las palas proporcionan información acerca de la energía extraída por el dispositivo de captación de energía de corrientes. El movimiento rotacional de las turbinas da lugar a distintas configuraciones geométricas, pudiéndose observar que los momentos obtenidos dependen de la posición relativa ocupada por las palas (ver Figura 3). Este hecho destaca la importancia de simular el movimiento de los rotores mediante el modelo *Sliding Mesh*.

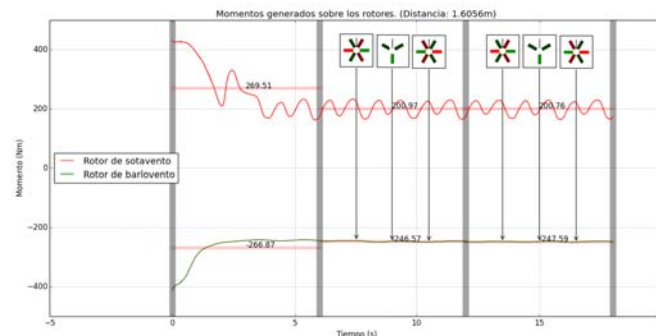


Figura 3: Momentos obtenidos sobre los rotores a lo largo de tres giros completos.

Bibliografía

- [1] Busto, S. (2014). *Análisis mediante CFD de un dispositivo de generación de energía de corrientes*, Trabajo de master, Universidad de Santiago de Compostela.
- [2] Wilcox, D. (1988). *Turbulence modelling for CFD*, DCW Industries.

Un problema de narices

Gonzalo Castiñeira Veiga

Departamento de Matemática Aplicada

28 de enero de 2015

Introducción

En el presente problema se considera un estudio del flujo aéreo nasal para una cavidad de un individuo en particular, a petición de la empresa NASAL S. L. (Laboratorio de Sistemas Avanzados de Flujo Aéreo Nasal). Los principales objetivos eran la comparación de los datos obtenidos por la empresa con software libre, con el código comercial ANSYS y por otro lado, desarrollar una metodología que permitiese realizar de forma sistemática simulaciones con distintas condiciones de respiración.

Para el estudio de este problema destacan tres dificultades: la obtención de la geometría y el mallado de la misma, la modelización correcta del problema (cuáles son las hipótesis a considerar) así como la complejidad de su resolución y la interpretación de resultados debido a la falta de experimentación.

Expondremos en primer lugar una descripción del problema, a continuación expondremos las ecuaciones matemáticas que gobiernan el flujo y finalmente abordaremos los resultados obtenidos.

Descripción del problema

La cavidad nasal tiene como objetivo principal filtrar, calentar y humidificar el aire inspirado antes de que llegue a los pulmones a la vez que facilita la sensación de olor.

Anatomía y Fisiología

La primera característica importante que llama nuestra atención es la geometría de la cavidad nasal, la cual podemos dividir en 4 partes:

- Vestíbulo: En esta parte el aire entra por los orificios anteriores de las fosas nasales, también conocidos como narinas, de forma prácticamente vertical a la superficie que definen los orificios nasales.

- Válvula nasal: En esta región el flujo adquiere características notables, como picos de velocidad, zonas con mayor esfuerzo cortante o caídas de presión, mencionados en la bibliografía.
- Cornetes: Estas estructuras óseas van a ser en gran medida las responsables de caracterizar el tipo de flujo dentro de la cavidad nasal.
- Nasofaringe: Las cavidades izquierda y derecha se encuentran, produciendo la mezcla del aire correspondiente a cada lado, lo que produce la aparición de remolinos y comportamiento caótico antes de abandonar la cavidad nasal por la faringe hacia los pulmones.

La geometría presenta diferentes factores que condicionan las características del flujo para que se produzcan apropiadamente las funciones nasales que son el olfato, la filtración, humidificación y acondicionamiento del aire antes de llegar a los pulmones, evitando daños alveolares.

Como consecuencia a que no todas las condiciones climáticas confieren al aire características adecuadas, podemos distinguir las narices leptorrinas, mesorrinas y platirrinas (ver Figura 1), correspondientes desde los climas más fríos y secos a los más calientes y húmedos. Por tanto, los primeros tendrán un flujo más turbulento (mayor tiempo de residencia) y los últimos más laminar.



Figura 1: Nariz leptorrina, mesorrina y platirrina, respectivamente.

Una vez el aire entra por las narinas, es necesario también un proceso de filtración y defensa, que nos proteja ante agentes externos dañinos.

Motivación para el uso del CFD

La función de filtración y defensa es también la que impide un eficiente consumo de medicamentos por vía nasal. Muchos fármacos tomados por vía tópica oral, son digeridos en el estómago antes de que el organismo haya absorbido sus componentes eficientemente, sin embargo, los fármacos consumidos por vía nasal que se depositan en la superficie de la cavidad, logran optimizar el objetivo final del fármaco. En este punto, es necesario el estudio de deposición de partículas así como los tiempos de residencia en las cavidades nasales.

La anatomía nasal impide la toma directa de patrones del flujo respiratorio. Con el desarrollo exponencial tecnológico, ha cobrado gran auge la simulación fluidodinámica computacional (CFD) para este tipo de problemas.

Modelización matemática

Para comenzar con la descripción del modelo, partiremos de las ecuaciones que describen de forma general el comportamiento de un fluido Newtoniano con transferencia de calor. Estas son las ecuaciones de cantidad de movimiento, conservación de la masa, conservación de la energía y la ecuación de estado:

$$\begin{aligned}\frac{\partial \rho}{\partial t} + \operatorname{div}(\rho \mathbf{v}) &= 0, \\ \frac{\partial(\rho \mathbf{v})}{\partial t} + \operatorname{div}(\rho \mathbf{v} \otimes \mathbf{v}) - \operatorname{div} \mathbb{T}_\nu + \operatorname{grad} \pi &= \rho \mathbf{g}, \\ \frac{\partial(\rho h)}{\partial t} + \operatorname{div}(\rho h \mathbf{v}) - \mathbb{T}_\nu \cdot \mathbb{D} - \dot{\pi} &= -\operatorname{div} \mathbf{q}_c + f, \\ \pi &= \rho R \theta,\end{aligned}$$

siendo ρ la densidad, $\mathbf{v}$ la velocidad, $\mathbb{T}_\nu$ la parte viscosa del tensor de tensiones de Cauchy, π la presión del fluido, $\mathbf{g}$ la gravedad atmosférica, h la entalpía específica, $\mathbb{D}$ la parte simétrica del gradiente del campo de velocidades, $\mathbf{q}_c$ el flujo de calor por conducción, f la fuente de calor, R la constante universal de los gases y θ la temperatura del fluido.

A continuación realizaremos una adimensionalización, es decir, substituiremos nuestras variables por las correspondientes reescaladas por las magnitudes características del problema (denotadas con subíndice cero). Esto nos permite despreciar aquellos términos no influyentes y obtener un modelo simplificado para cada caso.

En el proceso del análisis dimensional introducimos números adimensionales, de los cuales destacamos el número de Reynolds, que compara los términos de inercia frente a los viscosos. Suele utilizarse como indicador para saber si estamos ante un flujo laminar o turbulento, sin embargo, debido a la complejidad de la geometría, no resulta sencillo realizar esta decisión debido al comportamiento caótico incluso para caudales bajos, asociados a valores de Reynolds inferiores al establecido como umbral. Por tanto, los autores han elegido como convenio considerar flujo laminar aquellos en el rango de respiración tranquila (hasta los 15L/min) y turbulento para respiración forzada. En la literatura se justifica la hipótesis de flujo estacionario, la cual imitaremos con el objetivo de comparar nuestros resultados obtenidos.

- Respiración tranquila: Para caudales entre 7.5L/min y 15L/min tenemos un flujo estacionario, laminar e incompresible con transferencia de calor, modelado por las siguientes ecuaciones:

$$\begin{aligned}\operatorname{div} \mathbf{v} &= 0, \\ \operatorname{div}(\rho \mathbf{v} \otimes \mathbf{v}) - \frac{1}{Re} \operatorname{div}(\mathbb{T}_\nu) + \operatorname{grad} \pi' &= 0, \\ \operatorname{div}(\rho h \mathbf{v}) - \frac{1}{Pe} \operatorname{div}(k \operatorname{grad} \theta) &= 0, \\ \pi_0 &= \rho R \theta,\end{aligned}$$

siendo Pe el número adimensional de Peclet y k la conductividad térmica.

- Respiración forzada: Para caudales a partir de 15L/min tenemos un flujo estacionario, turbulento e incompresible, con transferencia de calor. Las ecuaciones son análogas al anterior reemplazando $\frac{1}{Re}div(\mathbb{T}_\nu)$ por

$$div((\mu + \mu_t)(grad\mathbf{v} + (grad\mathbf{v})^T)),$$

con μ_t la viscosidad turbulenta, obtenida mediante un modelo de turbulencia.

Por completitud, mencionamos a continuación los tres tipos de tratamiento para la modelar las escalas más pequeñas en un flujo turbulento.

- **DNS**, *Direct Numerical Simulation*, consiste en la realización numérica directa del problema. Por tanto, se hace imprescindible una malla suficientemente fina que sea capaz de captar dichas estructuras, lo que conlleva un elevado coste computacional.
- **LES**, *Large eddy simulation*, tratan de resolver los torbellinos de mayor tamaño y modelizan la influencia de los de menor tamaño. La resolución de la malla será fina especialmente cerca de las paredes, sin embargo, su coste computacional es admisible.
- **RANS**, *Reynolds Average Navier Stokes*, utiliza los promedios temporales de las variables. El coste computacional es relativamente bajo.

En nuestro trabajo optaremos por utilizar el modelo RANS, $k-\omega$ SST por ser, de los modelos RANS probados, el que mejor resultados parece dar en comparación con los experimentos.

Condiciones de contorno

Para las condiciones de contorno nos encontramos con dos posibilidades.

- Condición de presiones: En la entrada consideraremos temperatura y presión ambiental, mientras que en la salida impondremos una presión negativa responsable de la entrada de aire a la cavidad y una temperatura de $\theta = 304K$. En las paredes imponemos velocidad nula y una temperatura constante de $\theta = 306K$. Esta sería la opción más realista si tenemos en cuenta como se produce la respiración.
- Condición de velocidades: En la entrada impondremos un caudal específico y temperatura ambiental. En la salida imponemos condición de salida libre y en las paredes nuevamente velocidad nula y una temperatura constante de $\theta=306$ K. Esta condición es menos realista, pero menos costosa computacionalmente. Será la condición que utilizaremos en nuestras simulaciones.

Resultados

Una vez modelizado el problema realizamos mallas adecuadas para ejecutar nuestra simulación. El mallado en un problema en CFD es una etapa esencial de la cual dependerán los resultados que obtendremos en nuestras simulaciones.

En primer lugar, comparemos dos simulaciones analógicas realizadas con las diferentes condiciones de contorno expuestas (Figura 2). Si bien a medida que nos alejamos de la entrada las diferencias disminuyen, en el caso de los caudales, la presión de un orificio es mucho mayor que el otro mientras que en el caso de las presiones uno de los orificios tendrá mayor caudal. Este último hecho se corresponde con la realidad, pues podemos llegar a respirar hasta cuatro veces más por un orificio que por el otro.

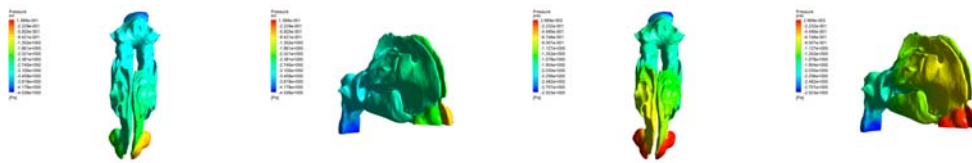


Figura 2: Condición caudales (izquierda) y presiones (derecha).

En la Figura 3 mostramos los resultados obtenidos mediante la simulación de nuestros modelos considerando los caudales de $7.5L/min$, $10L/min$ y $15L/min$. Podemos observar como el aire llega a la región olfativa con velocidades muy bajas para evitar daños en los nervios olfativos. También podemos apreciar que debido a que nos encontramos en la parte posterior de los cornetes, los campos de velocidad apuntan hacia la nasofaringe por donde el aire abandona la cavidad nasal.



Figura 3: Líneas de corriente y contornos de velocidad local. Caudal $7.5L/min$.

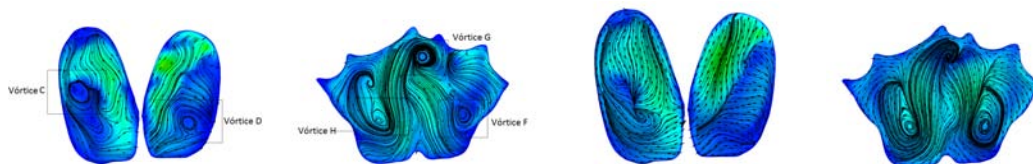


Figura 4: Vórtices para los caudales $15L/min$ y $40L/min$.

En las dos primeras secciones de la Figura 4 podemos observar para $15L/min$ la aparición de numerosos vórtices donde el aire de ambas cavidades se mezcla. En

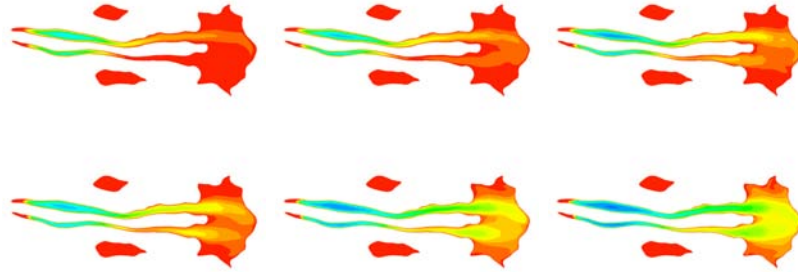


Figura 5: Temperaturas en función del aumento de caudal.

particular encontramos las conocidas como bifurcaciones de Hopf, que nos indican cierto comportamiento transitorio del flujo. Por otro lado, tenemos los resultados obtenidos para $40L/min$ mediante un modelo RANS, donde podemos observar que, al contrario de lo que se esperaba, hay menos vórtices, lo que nos hace reconsiderar el problema mediante otro método para el tratamiento de la turbulencia.

El acondicionamiento del aire es esencial para una función respiratoria adecuada. Observamos en la Figura 5 como el aire llega a la nasofaringe con menor temperatura a medida que aumenta la velocidad de inspiración.

Trabajo futuro

Como trabajo futuro es necesaria la búsqueda de simplificaciones en nuestro dominio así como un estudio de las condiciones de contorno expuestas. También es necesaria la comparación experimental para la verificación de resultados. Finalmente, el objetivo será desarrollar herramientas útiles para los especialistas, que les permitan la realización de diagnósticos de una forma no invasiva, eficaz y veraz.

Bibliografía

- [1] Quadrio, M., Pipolo, C., Corti, S., Lenzi, R., Messina, F., Pesci, C. y Felisati G. (2013). *Review of computational fluid dynamics in the assessment of nasal air flow and analysis of its limitations*, Eur Arch Otorhinolaryngol, **271**, pp. 2349–54.
- [2] Wen, J., Inthavong, K., Tu, J. y Wang, S. (2008). *Numerical simulations for detailed airflow dynamics in a human nasal cavity*, Respiratory Physiology & Neurobiology, **161**, pp. 125–135.
- [3] Weinhold, I. (2004). *Numerical simulation of airflow in the human nose*, Eur Arch Otorhinolaryngol, **261**, pp. 452–455.

O último teorema de Fermat

Jesús Conde Lago

Departamento de Álgebra

11 de febreiro de 2015

Introdución

O último teorema de Fermat afirma que a ecuación $x^n + y^n = z^n$ non ten solucións enteiras non triviais sendo $n > 2$ un número enteiro. É un dos problemas matemáticos máis coñecidos e un dos principais impulsos do desenvolvemento da teoría alxébrica de números durante o século XIX. Foi considerado por Pierre de Fermat en 1637 e non resolto ata 1995 por Andrew Wiles, quen demostrou un caso particular da conxectura de Taniyama-Shimura que implicaba o último teorema de Fermat.

O noso obxectivo será facer un percorrido dende a formulación do enunciado en 1637 por Pierre de Fermat ata a súa demostración en 1995 por Andrew Wiles. O que nos permitirá establecer importantes resultados que permitiron a Ernst Kummer probar o último teorema de Fermat para o caso de expoñentes primos regulares e que asentaron as bases da teoría alxébrica de números.

O último teorema de Fermat

Enunciado e primeiros resultados (ata 1839)

O matemático francés Pierre de Fermat (1601-1665) escribiu en 1637 na marxe do seu exemplar da *Arithmetica* de Diofanto, xunto ao problema da búsqueda das ternas pitagóricas, a seguinte nota:

“É imposible descompoñer un cubo en dous cubos, un bicadrado en dous bicadrados, e en xeral, unha potencia calquera, máis alá do cadrado, en dúas potencias do mesmo expoñente. Encontrei unha demostración realmente marabillosa, pero a marxe deste libro é moi pequena para poñela.”

Pierre de Fermat

O enunciado foi coñecido posteriormente como o último teorema de Fermat:

Teorema 1 (O último teorema de Fermat). *A ecuación*

$$x^n + y^n = z^n$$

non ten solucións enteiras non triviais ($xyz \neq 0$) para todo enteiro $n > 2$.

PALABRAS CLAVE: Último teorema de Fermat; número complexo ideal; ideal; primo regular.

Descoñécese se Fermat tiña tal demostración, pero si se sabe que tiña unha proba para o caso $n = 4$ empregando o método de descenso infinito que el mesmo introduciu. Isto permitiu unha importante simplificación do problema. Pois todo enteiro $n > 2$ é múltiplo de 4 ou dun primo $p \neq 2$ e a ecuación $x^{ab} + y^{ab} = z^{ab}$ pódese escribir como $(x^a)^b + (y^a)^b = (z^a)^b$, entón tra-la proba do caso $n = 4$ o problema reducíase a probar que $x^p + y^p = z^p$ non ten solucións enteiras non triviais para todo primo $p \neq 2$.

Nas seguintes décadas demostrouse o enunciado para certos expoñentes p primos. Por exemplo, para $p = 3$ foi parcialmente resolto por Euler en 1770 e de forma completa por Gauss, ambos traballando con números alxébricos da forma $a + b\sqrt{-3}$ con a e b enteiros. Posteriormente, para $p = 5$ foi probado por Dirichlet e Legendre de forma independente entre os anos 1825 e 1828, para $p = 7$ por Lamé en 1839, ...

Nestes anos destacou a figura da francesa Sophie Germain, a matemática máis relevante do século XIX e cuxos traballos abriron en 1823 unha nova vía na búsqueda dunha demostración do último teorema de Fermat. Se anteriormente o problema reducírase ao estudo para os expoñentes primos $p \neq 2$, Germain dividiu esta simplificación do último teorema de Fermat en dous casos complementarios:

- *Primer caso do último teorema de Fermat:* non existen solucións enteiras x, y, z non triviais para a ecuación $x^p + y^p = z^p$ con $p \nmid xyz$.
- *Segundo caso do último teorema de Fermat:* non existen solucións enteiras x, y, z non triviais para a ecuación $x^p + y^p = z^p$ con $p \mid xyz$.

Esta división non é casual, pois está motivada pola proba que a propia Germain fixo do primeiro caso do último teorema de Fermat para os expoñentes primos p tales que $2p + 1$ tamén é primo (3, 5, 11, 23, 29, 41, ...). Nese mesmo ano Legendre estendeu en 1823 este criterio aos primos p tales que $4p + 1$, $8p + 1$, $10p + 1$, $14p + 1$ ou $16p + 1$ é primo (ver [4] para consultar o artigo orixinal de Legendre, grazas ao cal se conservan os traballos de Germain).

Moitos dos avances obtidos na búsqueda da demostración do último teorema de Fermat nos últimos 150 anos estudan por separado ambos casos, sendo en xeral o segundo caso tecnicamente máis complicado ca o primeiro. Este é o caso dos traballos de Kummer, os cales lle permitirían demostrar o último teorema de Fermat para os expoñentes primos regulares.

Ernst Kummer e a teoría alxébrica de números

Ernst Kummer, na metade do século XIX, foi o primeiro matemático en abordar seriamente o problema para todos os expoñentes. Comezou traballando no anel

$$\mathbb{Z}[\zeta_p] = \{a_0 + a_1\zeta_p + \dots + a_m\zeta_p^m \mid 0 \leq m \in \mathbb{Z}, a_i \in \mathbb{Z} \forall i = 0, \dots, m\}$$

onde $\zeta_p \neq 1$ é unha raíz complexa p -ésima primitiva da unidade, e observando que neste anel a ecuación $x^p + y^p = z^p$ pódese escribir como

$$z^p = x^p + y^p = (x + y)(x + \zeta_p y) \cdots (x + \zeta_p^{p-1} y). \quad (1)$$

Lamé, quen realizou as mesmas consideracións que Kummer, afirmou en 1847 que tiña unha demostración do último teorema de Fermat empregando que $\mathbb{Z}[\zeta_p]$ é un dominio de factorización única para todo primo p . Isto non é certo, pois tan só o será para os primos $p \leq 19$, tal como conxeturou o propio Kummer e se demostrou en 1950.

Lamé, considerando a factorización anterior (1), probou que os factores $x + \zeta_p^m y$ e $x + \zeta_p^n y$ son coprimos dous a dous en $\mathbb{Z}[\zeta_p]$ sempre que $m \neq n$. A partir deste resultado e da factorización xa mencionada, concluiu que

$$\exists a_i \in \mathbb{Z}[\zeta_p] \text{ tal que } x + \zeta_p^i y = a_i^p, \quad \forall i = 1, \dots, p-1.$$

De forma implícita, estaba a afirmar erroneamente que $\mathbb{Z}[\zeta_p]$ é un dominio de factorización única. Partindo agora das igualdades $x + \zeta_p^i y = a_i^p$ obtidas, Lamé chegou a unha contradicción, de forma que non podían existir enteiros non triviais verificando a ecuación $x^p + y^p = z^p$ para $p \neq 2$. Obtendo así unha demostración do último teorema de Fermat.

Como xa se dixo, a demostración de Lamé tiña un erro. Sen embargo a idea non estaría desencamiñada, de feito sería a mesma estratexia que logo empregaría Kummer na súa demostración para os expoñentes primos regulares.

Co fin de solventar o problema da factorización única nestes aneis, Kummer observou que en $\mathbb{Z}[\zeta_p]$ sempre existe unha factorización única, non en termos de elementos do anel senón en termos de ideais do anel, aos que el denominou “números ideais complexos”. Estes elementos son os coñecidos na actualidade, coa notación de Dedekind, como ideais dun anel.

Considerando agora na factorización anterior (1) da ecuación cada factor como un ideal, temos unha factorización da p -ésima potencia do ideal (z) como produto de ideais:

$$(z)^p = (x + y)(x + \zeta_p y) \cdots (x + \zeta_p^{p-1} y) \subseteq \mathbb{Z}[\zeta_p].$$

Da mesma forma que fixo Lamé, buscaremos unha factorización única, pero neste caso será unha factorización única en ideais primos. Para isto, se empregarán estes dous importantes resultados: a factorización única en ideais primos no anel de enteiros alxébricos $\mathcal{O}_L$ para toda extensión de corpos $L|\mathbb{Q}$ e que $\mathbb{Z}[\zeta_p]$ é o anel de enteiros de $\mathbb{Q}(\zeta_p)$.

Teorema 2. *Sexa $L|\mathbb{Q}$ unha extensión de corpos. Cada ideal $\mathfrak{a}$ do anel de enteiros alxébricos $\mathcal{O}_L$ diferente de 0 e 1 admite unha factorización*

$$\mathfrak{a} = \mathfrak{p}_1 \cdots \mathfrak{p}_r$$

en ideais primos non nulos $\mathfrak{p}_i$ de $\mathcal{O}_L$, a cal é única salvo a orde dos factores.

Teorema 3. *Sexa ζ_p unha raíz p -ésima primitiva da unidade. Dado o corpo ciclotómico $\mathbb{Q}(\zeta_p)$, unha $\mathbb{Z}$ -base do anel de enteiros de $\mathbb{Q}(\zeta_p)$ está dada por $\{1, \zeta_p, \dots, \zeta_p^{p-2}\}$, é dicir,*

$$\mathcal{O}_{\mathbb{Q}(\zeta_p)} = \mathbb{Z} + \mathbb{Z}\zeta_p + \cdots + \mathbb{Z}\zeta_p^{p-2} = \mathbb{Z}[\zeta_p].$$

Agora, de forma análoga á demostración fallida de Lamé, os ideais $(x + \zeta_p^m y)$ e $(x + \zeta_p^n y)$ son coprimos dous a dous en $\mathbb{Z}[\zeta_p]$ sempre que $m \neq n$. Entón pola factorización única en ideais primos en $\mathbb{Z}[\zeta_p]$ temos que

$$(x + \zeta_p^i y) = \mathfrak{a}_i^p \quad \forall i = 0, \dots, p-1. \quad (2)$$

Se cada $\mathfrak{a}_i$ fose un ideal principal, é dicir,

$$(x + \zeta_p^i y) = \mathfrak{a}_i^p = (\alpha_i)^p \quad \forall i = 0, \dots, p-1,$$

entón obteríamos

$$x + \zeta_p^i y = \varepsilon_i \alpha_i^p, \quad \varepsilon_i \in \mathbb{Z}[\zeta_p]^*, \quad \forall i = 0, \dots, p-1.$$

Esta igualdade é moi similar á empregada por Lamé. De feito, Kummer tamén chegou a unha contradición a partir desta igualdade, demostrando así o último teorema de Fermat para certos casos, aqueles nos que os ideais $\mathfrak{a}_i$ sexan principais.

O que nos estamos preguntando indirectamente é se $\mathbb{Z}[\zeta_p]$ é un dominio de ideais principais. O propio Kummer indicou que a resposta é negativa, de feito é inmediato se recordamos que todo dominio de ideais principais é dominio de factorización única e que dito anel só terá factorización única cando $p \leq 19$. Consciente de que a condición de ser un dominio de ideais principais era demasiado forte, Kummer propúxose intentar evitar esta hipótese tan restrictiva.

A idea será buscar as condicións nas cales todos os ideais $\mathfrak{a}$ de $\mathbb{Z}[\zeta_p]$ verificando $(\beta) = \mathfrak{a}^p \subseteq \mathbb{Z}[\zeta_p]$, son ideais principais. En esencia, o que estamos a facer é medir canto de lonxe se atopa un anel de ser un dominio de ideais principais. Para isto Kummer definiu o “grupo de clases de ideais” dun corpo K , de forma intuitiva e sen precisar moito, como o grupo cociente dos ideais do anel $\mathcal{O}_K$ sobre aqueles que son principais. Denotarase este grupo abeliano por $\text{Cl}(K)$.

Ademais, definirase o “número de clases” de K , h_K , como o cardinal grupo de clases de ideais de K . Este número será finito, o coñecido como teorema da finitude do número de clase, o cal é un dos teoremas máis importantes da teoría alxébrica de números. Isto, nos permite caracterizar, por exemplo, os dominios de ideais principais como aqueles que teñen número de clases 1.

No noso caso particular, ao traballar cos aneis de enteiros alxébricos $\mathbb{Z}[\zeta_p]$, denotaremos os corpos ciclotómicos por $K_p := \mathbb{Q}(\zeta_p)$ e o número de clases correspondente por $h_p := h_{K_p}$. Polo tanto, o que faremos será asignar a cada primo p un número enteiro h_p .

Volvendo á demostración de Kummer, necesitábase que se $(\beta) = \mathfrak{a}^p \subseteq \mathbb{Z}[\zeta_p]$ entón existese $\alpha \in \mathbb{Z}[\zeta_p]$ tal que $\mathfrak{a} = (\alpha)$ fose un ideal principal en $\mathbb{Z}[\zeta_p]$. Vexamos cando isto é certo. Supoñamos que $(\beta) = \mathfrak{a}^p$, así $\mathfrak{a}^p$ é un ideal principal. Polo tanto a p -ésima potencia de $\mathfrak{a}$ é o elemento identidade en $\text{Cl}(K_p)$, e entón a orde de $\mathfrak{a} \in \text{Cl}(K_p)$ é 1 ou p . Se $\text{Cl}(K_p)$ non tivese elementos de orde p , $\mathfrak{a}$ tería orde 1 e sería principal, que é o que nos interesa. Polo teorema de Lagrange, sabemos que $\text{Cl}(K_p)$ ten un elemento de orde p se e só se p é divisor da orde de $\text{Cl}(K_p)$, é dicir, se e só se o número de clases h_p é divisible por p .

Polo tanto, interésannos os primos p tales que p non divide ao número de clases h_p . Chamaremos primos regulares a este tipo de números primos.

Definición 4. *Un primo p é regular se non divide ao seu número de clases h_p .*

Desta forma temos que se p é un primo regular, entón cada ideal $\mathfrak{a}_i$ de (2) é un ideal principal. Obtendo así

$$x + \zeta_p^i y = \varepsilon_i \alpha_i^p, \quad \varepsilon_i \in \mathbb{Z}[\zeta_p]^*, \quad \forall i = 0, \dots, p-1.$$

A partir de aquí, como xa se dixo, Kummer logrou chegar á mesma contradicción á que chegaba Lamé na súa demostración, probando así o último teorema de Fermat para todos os expoñentes primos regulares.

Teorema 5 (O último teorema de Fermat para expoñentes primos regulares). *A ecuación*

$$x^p + y^p = z^p$$

non ten solucións enteiras non triviais para todo primo regular $p \neq 2$.

En realidade, Kummer probou por separado os dous casos que estableceu Sophie Germain, sendo ambos moi similares pero o segundo máis complexo tecnicamente (ver [1] e [6]).

Mencionar que, se ben os primos p para os cales $\mathbb{Z}[\zeta_p]$ é un dominio de factorización única (os de número de clases 1) son menores ou iguais que 19, polo menos entre os primeiros primos hai moitos máis que son regulares (non son regulares, por orde: 37, 59, 67, 101, ... ; pero segue sendo un problema aberto se o conxunto de primos regulares é finito). Polo tanto, o resultado de Kummer foi un avance notable na búsqueda da demostración do último teorema de Fermat, ademais da importancia que tivo no desenvolvemento da teoría alxébrica de números (ver [2] e [5]).

Andrew Wiles e a súa demostración

Xa na metade do século XIX, momento de aparición de novos métodos na teoría de números e na álgebra en xeral, o matemático Yutaka Taniyama propúxose averiguar se as curvas elípticas definidas sobre un corpo de números son parametrizables mediante funcións automorfas. Búsqueda que foi acotada posteriormente por Goro Shimura en 1964 e André Weil en 1967, resultando así a conxectura de Taniyama-Shimura: toda curva elíptica racional é unha función modular.

A priori, esta conxectura proporcionaría un camiño para probar a conxectura de Hasse-Weil. Sen embargo, Gunther Frey observou en 1985 que dada unha hipotética solución enteira (a, b, c) da ecuación $x^p + y^p = z^p$ para un expoñente primo $p \neq 2$, entón a curva elíptica definida sobre $\mathbb{Q}$

$$y^2 = x(x - a^p)(x + b^p),$$

denominada curva de Frey, non sería modular. Isto relacionaba o último teorema de Fermat coa conxectura de Taniyama-Shimura (ver [3]).

A afirmación de Frey foi demostrada por Kenneth Ribet en 1986, mediante a proba da conxectura épsilon (caso particular da conxectura da modularidade de Jean-Pierre Serre), quen demostrou que a curva de Frey non é modular.

Teorema 6 (Conxectura épsilon ou teorema de Ribet). *Toda curva de Frey non é modular.*

Isto permitiu afirmar que a conxectura de Taniyama-Shimura implicaba o último teorema de Fermat, o que levou ao matemático británico Andrew Wiles a buscar unha demostración do último teorema de Fermat, a cal se reduciría a probar que toda curva elíptica “semiestable” é modular (versión semiestable da conxectura de Taniyama-Shimura-Weil), dado que toda curva de Frey é “semiestable”.

Tras anos de traballo, Wiles presentou a súa demostración nunha conferencia de Cambridge en 1993, pero esta contiña un erro. Non foi ata 1995 cando Andrew Wiles chegou á demostración do último teorema de Fermat mediante a proba do caso “semiestable” da conxectura de Taniyama-Shimura.

Teorema 7 (Caso “semiestable” da conxectura de Taniyama-Shimura). *Toda curva elíptica racional semiestable é modular.*

Desta forma, Andrew Wiles deu resposta en 1995 a un dos maiores desafíos matemáticos da historia tras máis de 350 anos dende o seu enunciado en 1637 por Pierre de Fermat.

Bibliografía

- [1] Borevich, Z. I. e Chafarevich I. R. (1967). *Théorie de nombres*, Ed. facsímil, Gauthier-Villars.
- [2] Ireland, K. F. e Rosen, M. I. (1982). *A Classical Introduction to Modern Number Theory*, Springer-Verlag.
- [3] Kato, K. & Kurokawa, N. e Saito, T. (2000). *Number Theory 1: Fermat’s dream*, Providence, Rhode Island: American Mathematical Society.
- [4] Legendre M. (1823). *Recherches sur quelques objets d’analyse indéterminée et particulièrement sur le théorème de Fermat*, Mémoires de l’Académie Royale des Sciences de l’Institut de France.
- [5] Neukirch, J. (1999). *Algebraic Number Theory*, Springer-Verlag.
- [6] Ribenboim P. (1999). *Fermat’s Last Theorem*, Springer-Verlag.

La ecuación de Hill

Lucía López Somoza

Departamento de Análisis Matemático

25 de Febrero de 2015

Resumen

La ecuación de Hill es una ecuación diferencial con innumerables aplicaciones en ingeniería y física, especialmente en problemas relacionados con fenómenos oscilatorios.

Comentaremos las propiedades de oscilación de las soluciones del problema homogéneo, así como la estabilidad de las mismas. También hablaremos de la existencia de una sucesión de autovalores asociados a los problemas de Dirichlet y periódicos.

En cuanto al caso no homogéneo, analizaremos la existencia de solución única para el problema periódico, relacionándola con el signo de la función de Green asociada.

Caso homogéneo

Consideremos la ecuación diferencial lineal homogénea de segundo orden

$$y'' + p(t)y' + q(t)y = 0, \quad (1)$$

con $p, q \in \mathcal{L}_{loc}^\infty(\mathbb{R})$ (es decir, funciones medibles y acotadas en cualquier intervalo compacto salvo en un conjunto de medida nula). Buscaremos soluciones de (1) que verifiquen que $y \in \mathcal{C}^1(\mathbb{R})$, $y' \in \mathcal{AC}(\mathbb{R})$ e $y'' \in \mathcal{L}_{loc}^\infty(\mathbb{R})$, donde $\mathcal{AC}(\mathbb{R})$ denota al conjunto de funciones absolutamente continuas en $\mathbb{R}$.

Nos interesarán los fenómenos de oscilación de dichas soluciones. El Teorema de separación de Sturm proporciona un primer resultado al respecto:

Teorema 1 (de separación de Sturm). *Si y_1 e y_2 son dos soluciones linealmente independientes de $y'' + p(t)y' + q(t)y = 0$, entonces los ceros de estas funciones son distintos y se verifica que y_1 se anula exactamente una vez entre dos ceros sucesivos de y_2 , y recíprocamente.*

Además, si $p \in \mathcal{C}^1(\mathbb{R})$, un cambio de variable adecuado (véase [1]) permite reescribir la ecuación (1) en la forma

$$u'' + Q(t)u = 0. \quad (2)$$

PALABRAS CLAVE: Ecuación de Hill; teoría de oscilación; principios de comparación.

Dicho cambio de variable no afecta a los ceros de las soluciones y, por tanto, tampoco a la oscilación de las mismas. Estudiaremos pues las propiedades de (2):

Teorema 2. *Si $Q(t) < 0$, toda solución no trivial de (2) tiene a lo sumo un cero.*

No obstante, $Q(t) > 0$ no garantiza que las soluciones oscilen.

Una condición suficiente para la oscilación de las soluciones es la siguiente:

Teorema 3. *Sea u solución no trivial de (2), con $Q(t) > 0$ para todo $t > 0$. Si*

$$\int_1^\infty Q(t)dt = \infty,$$

entonces u tiene infinitos ceros en el semieje positivo.

Se tiene además el siguiente resultado:

Teorema 4. *Sea u una solución no trivial de (2) sobre un intervalo $[a, b]$. Entonces u tiene a lo sumo un número finito de ceros en dicho intervalo.*

Se deduce pues que si las soluciones oscilan, entonces sus ceros se extienden a lo largo de todo el semieje positivo. Por otra parte, se puede probar que la velocidad de oscilación de las soluciones aumenta al aumentar Q :

Teorema 5 (de comparación de Sturm). *Sean y y z soluciones no triviales de*

$$y'' + q(t)y = 0 \quad y \quad z'' + r(t)z = 0$$

respectivamente, con $q(t) > r(t)$. Entonces y se anula al menos una vez entre cada dos ceros consecutivos de z .

Problema de Dirichlet

Vamos a considerar ahora otra formulación del problema.

Fijemos un intervalo $[0, T]$ y una función $Q \in \mathcal{L}^\infty([0, T])$ y consideremos la ecuación

$$y''(t) + [Q(t) + \lambda]y(t) = 0, \quad t \in [0, T], \quad (3)$$

dependiente del parámetro $\lambda \in \mathbb{R}$.

Extendiendo con periodicidad Q a todo $\mathbb{R}$, se tiene que, a partir de cierto λ ,

$$\int_1^\infty [Q(t) + \lambda]dt = \infty,$$

por lo que, como consecuencia del Teorema 3, sabemos que las soluciones de (3) tienen infinitos ceros a lo largo del semieje positivo.

Sea y_λ la solución de (3) que verifica que $y_\lambda(0) = 0$ y denotemos por $t_0(\lambda)$ al primer cero positivo de y_λ . Como consecuencia del Teorema 5, sabemos que al aumentar λ , $t_0(\lambda)$ toma cada vez valores más pequeños.

La ecuación (3) junto con las condiciones de contorno

$$y(0) = y(T) = 0 \quad (4)$$

constituyen el llamado problema de Dirichlet.

El problema (3)-(4) tiene solución no trivial si y solo si existe algún λ para el cual $y_\lambda(T) = 0$, $y_\lambda \not\equiv 0$. Denotemos por S al conjunto de valores de λ que verifican dicha condición.

El Teorema 2 nos garantiza que si $Q(t) + \lambda < 0$, el problema (3)-(4) no tiene solución no trivial. Como consecuencia, S está acotado inferiormente.

Además, puesto que al aumentar λ , $t_0(\lambda)$ decrece y teniendo en cuenta la continuidad del problema respecto a los datos, tenemos garantizada la existencia de λ_0^D para el cual $t_0(\lambda_0^D) = T$. Se verificará pues que $\lambda_0^D = \min S$.

Podríamos repetir este proceso indefinidamente, por lo que llegamos a la conclusión de que existe una sucesión de infinitos valores

$$\lambda_0^D < \lambda_1^D < \lambda_2^D < \dots < \lambda_k^D < \dots$$

para los cuales el problema (3)-(4) con $\lambda = \lambda_k^D$ tiene una solución no trivial $y_{\lambda_k^D}$ que se anula exactamente k veces en el interior del intervalo $[0, T]$.

Problema periódico

Consideremos la ecuación de Hill en su forma estándar

$$y'' + [Q(t) + \lambda]y = 0, \quad t \in \mathbb{R}, \quad (5)$$

donde λ es un parámetro y $Q \in \mathcal{L}_{loc}^\infty(\mathbb{R})$ una función de periodo T (también podríamos considerar una función definida en el intervalo $[0, T]$ y tomar Q como una extensión periódica de la misma). Se introducen las siguientes definiciones:

Definición 6. *En las condiciones anteriores la ecuación (5) tiene dos soluciones y_1 e y_2 que quedan determinadas de forma única por las condiciones iniciales:*

$$y_1(0) = 1, \quad y_1'(0) = 0,$$

$$y_2(0) = 0, \quad y_2'(0) = 1.$$

Definición 7. *Se denomina ecuación característica asociada a (5) a*

$$\rho^2 - [y_1(T) + y_2'(T)]\rho + 1 = 0$$

y se denotan por ρ_1 y ρ_2 sus raíces, donde $\rho_1, \rho_2 \in \mathbb{C}$.

Definición 8. *Se denomina exponente característico al número $\alpha \in \mathbb{C}$ tal que*

$$e^{i\alpha T} = \rho_1, \quad e^{-i\alpha T} = \rho_2.$$

Definición 9. Si todas las soluciones de (5) están acotadas en $\mathcal{C}^1$ se dice que la solución trivial $y = 0$ es estable; en caso contrario, se denomina inestable.

El Teorema de Floquet aporta un primer resultado sobre la existencia de soluciones periódicas para la ecuación de Hill:

Teorema 10 (de Floquet). Si $\rho_1 \neq \rho_2$, entonces la ecuación de Hill tiene dos soluciones linealmente independientes

$$f_1(t) = e^{i\alpha t} p_1(t), \quad f_2(t) = e^{-i\alpha t} p_2(t),$$

con p_1 y p_2 funciones periódicas de periodo T y $\alpha \in \mathbb{C}$ dado en la Definición 8.

Si $\rho_1 = \rho_2$, entonces la ecuación de Hill tiene una solución periódica de periodo T (si $\rho_1 = \rho_2 = 1$) o $2T$ (si $\rho_1 = \rho_2 = -1$). Si denotamos por p dicha solución y si y es otra solución linealmente independiente de la anterior, entonces se verifica que

$$y(t+T) = \rho_1 y(t) + \theta p(t)$$

para alguna constante θ . Además, la condición $\theta = 0$ equivale a

$$y_1(T) + y_2'(T) = \pm 2, \quad y_2(T) = 0, \quad y_1'(T) = 0.$$

Observación 11. Si $\rho_1 \neq \rho_2$, toda solución no trivial de (5) puede expresarse como combinación lineal de f_1 y f_2 . Por tanto, si $\alpha \in \mathbb{R}$, cualquier solución de (5) tiene una cota superior en valor absoluto que no depende de t . En caso contrario, si α no es real, existe una solución no trivial de (5) que no está acotada.

Por otra parte, si $\rho_1 = \rho_2$, todas las soluciones no triviales de (5) son combinación lineal de una función p periódica y de otra función y que verifica que $y(t+T) = \rho_1 y(t) + \theta p(t)$. Luego toda solución está acotada si y solo si $\theta = 0$.

De acuerdo con la observación anterior, tenemos el siguiente test de estabilidad.

Corolario 12. La solución trivial de (5) es estable si y solo si se da uno de los siguientes casos:

- $y_1(T) + y_2'(T)$ toma un valor real e $|y_1(T) + y_2'(T)| < 2$.
- $y_1(T) + y_2'(T) = \pm 2$, $y_2(T) = y_1'(T) = 0$.

Por último, el teorema de oscilación garantiza la existencia de dos sucesiones de autovalores para los cuales los problemas T y $2T$ -periódicos tienen solución:

Teorema 13 (de Oscilación). Existen dos sucesiones de números reales

$$\lambda_0, \lambda_1, \lambda_2, \dots \quad \text{y} \quad \lambda'_1, \lambda'_2, \lambda'_3, \dots$$

tales que (5) tiene una solución de periodo T si y solo si $\lambda = \lambda_n$, $n = 0, 1, 2, \dots$, y una solución de periodo $2T$ si y solo si $\lambda = \lambda'_n$, $n = 1, 2, 3, \dots$. Los λ_n son las raíces de la ecuación $\Delta(\lambda) = 2$ y los λ'_n las de $\Delta(\lambda) = -2$, donde

$$\Delta(\lambda) = y_1(T, \lambda) + y_2'(T, \lambda).$$

Los λ_n y los λ'_n verifican las desigualdades

$$\lambda_0 < \lambda'_1 \leq \lambda'_2 < \lambda_1 \leq \lambda_2 < \lambda'_3 \leq \lambda'_4 < \lambda_3 \leq \lambda_4 \dots$$

y además se cumple que

$$\lim_{n \rightarrow \infty} \lambda_n^{-1} = 0, \quad \lim_{n \rightarrow \infty} (\lambda'_n)^{-1} = 0.$$

La solución trivial de (5) es estable si y solo si λ pertenece a los intervalos

$$(\lambda_0, \lambda'_1), (\lambda'_2, \lambda_1), (\lambda_2, \lambda'_3), (\lambda'_4, \lambda_3), \dots$$

Caso no homogéneo

Consideremos ahora el operador de Hill

$$L_Q u(t) \equiv u''(t) + Q(t)u(t), \quad t \in [0, T] \equiv I,$$

con

$$Q \in \mathcal{L}^\alpha(I), \quad \alpha \geq 1 \quad \text{y} \quad Q(t+T) = Q(t) \text{ c.t.p. } t \in \mathbb{R},$$

donde c.t.p. significa en todo punto salvo en un conjunto de medida nula. Definamos el espacio

$$X = \{u \in \mathcal{AC}^1(I), \quad u(0) = u(T), \quad u'(0) = u'(T)\}.$$

Se tienen las siguientes definiciones:

Definición 14. El operador L_Q verifica el principio del máximo (MP) si y solo si

$$u \in X, \quad L_Q u \geq 0 \quad \text{c.t.p. } t \in I \Rightarrow u \equiv 0 \text{ o } u < 0 \text{ en } I.$$

Definición 15. El operador L_Q verifica el principio del antimáximo (AMP) si y solo si

$$u \in X, \quad L_Q u \geq 0 \quad \text{c.t.p. } t \in I \Rightarrow u \equiv 0 \text{ o } u > 0 \text{ en } I.$$

Definición 16. El operador L_Q es no resonante en X si y solo si el problema

$$L_Q u = 0 \quad \text{c.t.p. } t \in I, \quad u(0) = u(T), \quad u'(0) = u'(T)$$

tiene únicamente la solución trivial.

Si L_Q verifica el MP o el AMP, entonces es no resonante. Además, si L_Q es no resonante, el problema periódico no homogéneo

$$L_Q u(t) = f(t) \quad \text{c.t.p. } t \in I, \quad u(0) = u(T), \quad u'(0) = u'(T),$$

con $f \in \mathcal{L}^\infty(I)$, tiene solución única dada por

$$u(t) = \int_0^T G_Q(t, s) f(s) ds \in \mathcal{AC}^1(I),$$

siendo G_Q la función de Green relativa al operador L_Q .

Teorema 17. *Se tienen las siguientes equivalencias entre los principios de comparación y el signo constante de la función de Green:*

- $G_Q \geq 0$ en $I \times I$ si y solo si L_Q verifica el AMP.
- $G_Q \leq 0$ en $I \times I$ si y solo si L_Q verifica el MP.

Además,

Lema 18. *Si $G_Q \leq 0$ en $I \times I$, entonces $G_Q < 0$ en $I \times I$.*

Cabe comentar que existe una fuerte relación entre los principios de comparación y los autovalores de la ecuación de Hill definidos en el Teorema de Oscilación. En efecto, se tiene lo siguiente:

- $L_{Q+\lambda}$ verifica el MP si y solo si $\lambda < \lambda_0$.
- $L_{Q+\lambda}$ verifica el AMP si y solo si $\lambda_0 < \lambda \leq \lambda'_1$.
- $L_{Q+\lambda}$ es no resonante en X si y solo si $\Delta(\lambda) \neq 2$.

El Teorema 17 y el Lema 18 permiten reescribir las equivalencias anteriores en términos de la función de Green:

- $G_{Q+\lambda} < 0$ para $\lambda < \lambda_0$.
- $G_{Q+\lambda} \geq 0$ para $\lambda_0 < \lambda \leq \lambda'_1$.
- $G_{Q+\lambda}$ cambia de signo para $\lambda > \lambda'_1$.
- $G_{Q+\lambda}$ no está definida para $\lambda = \lambda_n$, $n = 0, 1, \dots$

Además, se tienen diversas cotas para los primeros autovalores, entre las que se encuentran las siguientes:

Teorema 19. *Sea $Q \in \mathcal{L}^1(I)$. Entonces:*

- (i) $\lambda_0 \leq -\frac{1}{T} \int_0^T Q(s)ds$ y se tiene la igualdad si y solo si Q es constante.
- (ii) $\lambda'_1 = \min\{\lambda_1^D(Q_s), s \in \mathbb{R}\}$, con $Q_s(t) \equiv Q(t+s)$.

Bibliografía

- [1] Simmons, G. F. y Robertson, J. S. (1993). *Ecuaciones diferenciales con aplicaciones y notas históricas*, McGraw-Hill.

El camino para ganar la NBA: técnicas de predicción matemática

Marcos Matabuena Rodríguez

Departamento de Estadística e Investigación Operativa

11 de marzo de 2015

Resumen

Este texto pretende ser una breve introducción a la cuantificación en deportes de equipos, aplicado particularmente al baloncesto; y a su vez mostrar que el futuro en este campo pasa por la inclusión del estado físico del deportista en el modelo matemático. Además mostraremos como modelar la variación de la forma física de un deportista mediante un sistema de ecuaciones diferenciales.

Introducción

Las matemáticas pueden ser un elemento diferencial en la práctica deportiva, pese a todo estamos todavía muy lejos de alcanzar una transferencia real entre lo que quieren entrenador, director de equipo y deportista, y lo que en realidad se puede obtener con las técnicas matemáticas existentes. Algunas causas son las siguientes:

- Falta de desarrollo matemático a nivel teórico. Por ejemplo, en la teoría de juegos cooperativa con un número elevado de jugadores (con aplicación en ciclismo) o en las técnicas de análisis de cluster con series de tiempo para la clasificación del estado físico del deportista.
- Falta de conocimiento real del problema por parte del experto en matemáticas.
- Mentalidad cerrada por parte de los miembros de la comunidad deportiva a la utilización de las matemáticas como herramienta de ayuda.
- Exceso de simplificación en la modelización del problema y miedo a la innovación. Esto se hace latente en la aplicación de la lógica difusa para ciertos problemas, pese a ser necesaria su utilización tal y como vemos en la referencia por excelencia de la ciencia del entrenamiento deportivo [4].

Las métricas en deportes de equipo

Para cuantificar el rendimiento de un jugador o de un equipo, se han planteado los siguientes términos:

1. **Bottom up Metric.** Se centran en las estadísticas individuales de campo. Tienen el inconveniente de que hay jugadores que no registran grandes estadísticas individuales, pero que son imprescindibles en la victoria de su equipo tal y como se puede verificar al analizar los puntos que deja de conseguir un equipo cuando dicho jugador no está en pista. Un estudio riguroso de este hecho, puede encontrarse en [3].
2. **Top Down Measures.** Tienen en cuenta el rendimiento global del equipo y no se fijan en las estadísticas individuales de campo. La métrica mas conocida, es la llamada **Plus Minus Statistics**¹, que resulta de resolver un problema de optimización con norma cuadrática en donde se asigna a cada jugador el número de puntos por una determinada unidad de tiempo que aportará a su equipo estando presente en el terreno de juego. Existen diversas variaciones de este método, por ejemplo, ver [1] y que incluso han sido premiadas en el concurso del MIT de Boston (Mit Sloan Sport Conference). En cualquier libro clásico de análisis de datos se trata esta técnica ver [2], [3], [5].
3. **Módelos Mixtos.** Son modelos que combinan las dos metodologías anteriores.

La entropía de Shannon

Además de disponer de una medida que nos permita decidir si un equipo o jugador es bueno o malo en un determinado aspecto, debemos tener una herramienta para medir la distorsión de dicha medida o dicho en otros términos, determinar si el resultado obtenido es estable a lo largo del tiempo. La opción clásica en el mundo del deporte es utilizar la entropía de Shannon:

Definición 1. Sea X una variable aleatoria discreta con un espacio muestral Ω compuesto de n elementos. Definimos la entropía de Shannon asociado a la variable aleatoria X como:

$$-\sum_{i=1}^n p_i \log p_i,$$

donde p_i representa la probabilidad del evento i -ésimo de la variable X .

Ésta representa el grado de desorden de X , por ejemplo si $n=1$ la entropía vale 0, es decir no hay desorden al estar toda la informacion concentrada en un solo punto.

¹Tuvo sus inicios en el la NHL (National Hockey League), y actualmente se aplica en multitud de deportes de equipo. Por ejemplo, es muy popular en deportes como el baloncesto donde webs como <http://www.82games.com>, ofrecen actualizaciones en cada jornada sobre esta métrica.

Como aplicaciones, podemos citar las siguientes:

- Estudiar la variación de las anotaciones de un jugador respecto al número de compañeros que juegan con él.
- Estudiar la variación de las alineaciones de un equipo para por ejemplo determinar si un equipo tiene dependencia a grandes estrellas.
- Estudiar la variabilidad de las zonas de tiro de un jugador.

Una métrica espacial

La distribución de goles, canastas no es uniforme en todos los deportes de equipo, tanto desde el punto de vista espacial como temporal. Por ejemplo, en fútbol, se puede modelizar el número de goles entre dos periodos de tiempo bajo una distribución de Poisson. Veamos dos métricas de carácter espacial para analizar la habilidad de tiro de un jugador.

Sea R una región regular de $\mathbb{R}^2$ y dividamos R en 1284 rectángulos iguales. Denotemos mediante ij a cada subregión y definamos las siguientes métricas.

Definimos la métrica *Spread* mediante:

$$\sum_{ij \in R} FGA_{ij},$$

donde $FGA_{ij} = 1$ si lanza alguna canasta desde la zona ij y cero en caso contrario.

Por otro lado, definimos la métrica *Range* mediante:

$$\sum_{ij \in R} PPA_{ij},$$

donde $PPA_{ij} = 1$ si encesta alguna canasta desde la zona ij y cero en caso contrario.

El modelo de Banister

En las secciones anteriores, hemos visto unas pinceladas de un gran abánico de técnicas que permiten sacar conclusiones en un entorno no dinámico. Aunque puedan proporcionar una valiosa información, están fuertemente influenciadas por el estado físico del deportista, siendo precisamente este, el elemento diferencial para obtener la mayor ventaja posible frente al equipo rival, por ejemplo, el error continuado de un lanzador podría ser debido a un estado de fatiga.

Vamos a tratar brevemente la cuestión de modelizar la forma física de un deportista.

Sea $w(t)$ la carga física² ejecutada por el atleta en el instante de tiempo t y $g(t)$ y $h(t)$ los efectos positivos y negativos obtenidos por el atleta con el entrenamiento. Consideremos el siguiente sistema de ecuaciones diferenciales:

$$\frac{\partial g(t)}{\partial t} + \frac{1}{\tau_1} g(t) = w(t),$$

$$\frac{\partial h(t)}{\partial t} + \frac{1}{\tau_2} h(t) = w(t).$$

Utilizando el operador convolución y aplicando previamente la transformada de Laplace, obtenemos:

$$g(t) = w(t) * e^{\frac{-t}{\tau_1}} = \int_0^t w(s) e^{-\frac{t-s}{\tau_1}} ds,$$

$$h(t) = w(t) * e^{\frac{-t}{\tau_2}} = \int_0^t w(s) e^{-\frac{t-s}{\tau_2}} ds.$$

Discretizando:

$$g(n) = \sum_{i=1}^{n-1} w(i) e^{-\frac{n-i}{\tau_1}},$$

$$h(n) = \sum_{i=1}^{n-1} w(i) e^{-\frac{n-i}{\tau_2}}.$$

La forma física actual podemos modelizarla mediante:

$$p(n) = p(0) + k_1 \sum_{i=1}^{n-1} w(i) e^{-\frac{n-i}{\tau_1}} + k_2 \sum_{i=1}^{n-1} w(i) e^{-\frac{n-i}{\tau_2}}, \quad \forall n \in \mathbb{N}.$$

- Estimamos $p(0)$ con un test, además disponemos de muestra $p(1), \dots, p(s)$ con $s \in \mathbb{N}$, de la forma física del deportista en diferentes instantes de tiempo.
- El problema es estimar las constantes k_1, k_2, τ_1, τ_2 resolviendo un problema por mínimos cuadrados a través de la muestra de la forma física $p(0), \dots, p(s)$ y usando la ecuación de la forma física. Se necesitan por lo menos 5 muestras para obtener buenos resultados, aunque lo más importante es elegir una métrica fideligna.
- Las unidades dependen de la métrica elegida.

²Existen diversas métricas para medir la carga física realizada por el deportista, algunas se inspiran en medir la variación de la frecuencia cardiaca y otras a partir de la potencia desarrollada por el deportista, como la que se muestra en el gráfico inferior de la Figura 1. Además hay diversos test para evaluar el estado físico de un deportista, en el primer gráfico de la Figura 1 se muestra un test de potencia de 5 min.

En la Figura 1 vemos un ejemplo, donde el primer gráfico representa la variación de la forma física y el diagrama de barras representa la carga ejecutada en cada sesión por el deportista a lo largo del tiempo.

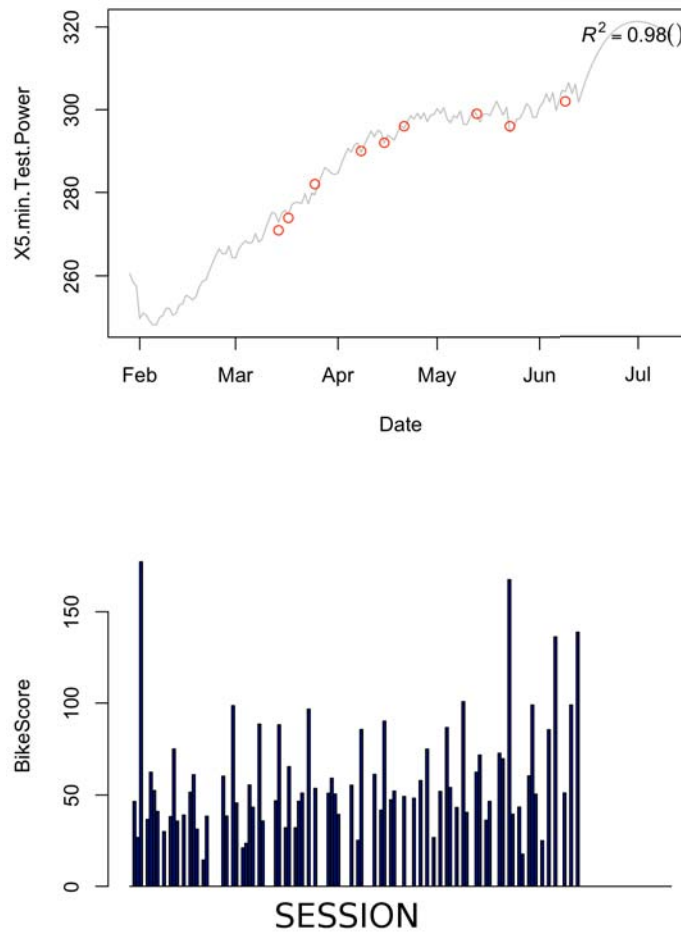


Figura 1: Aplicación del modelo banister

Conclusion

Como hemos visto en el texto, el futuro pasa por idear nuevas formulas matematicas que aumenten la precision de los modelos existentes y para ello es probable que muchas veces tengamos que recurrir a las matematicas mas vanguardistas. Tam-

bién quería destacar que el deporte, al igual que la física, puede ser una importante fuente de inspiración para la construcción de nuevas técnicas matemáticas. Muchos temas aquí citados, no han sido tratados con mucho detalle, una explicación visual más completa se puede encontrar en el siguiente enlace:

<https://www.youtube.com/watch?v=k5cc7irlA4k>.

Bibliografía

- [1] Brian Macdonald (2011). *An Improved Adjusted Plus-Minus Statistic for NHL Players*, in Proceedings of the MIT Sloan Sports Analytics Conference.
- [2] Schumaker, Robert P and Solieman, Osama K and Chen, Hsinchun (2010). *Predictive Modeling for Sports and Gaming*, Springer.
- [3] Shea, SM and Baker, CE (2013). *Basketball Analytics: Objective and Efficient Strategies for Understanding How Teams Win*, CreateSpace Independent Publishing Platform.
- [4] Siff, Mel C and Verkhoshansky, Yury (2000). *Superentrenamiento, volume 24*, Editorial Paidotribo.
- [5] Winston, Wayne L (2012). *Mathletics: How gamblers, managers, and sports enthusiasts use mathematics in baseball, basketball, and football*, Princeton University Press.

Geometría simpléctica superior motivada desde la teoría de campos lagrangiana

Néstor León Delgado

Instituto Max Planck de Matemáticas, Bonn

25 de marzo de 2015

La geometría simpléctica aparece de un modo natural en el estudio de la mecánica clásica. Gracias a esa estructura simpléctica podemos hablar del álgebra de Lie de los observables del sistema (variedad de Poisson). La teoría de campos lagrangiana es un formalismo matemático que describe una gran variedad de sistemas físicos. Cuando se intenta encontrar un análogo a la estructura simpléctica en este caso, es necesario hacer diversas elecciones, que pueden ser pensadas como condiciones iniciales. De forma universal (sin realizar ninguna elección) obtenemos una variedad simpléctica superior. Sucede entonces que el análogo a la variedad de Poisson no es un álgebra de Lie, sino un objeto que vive en geometría derivada: un álgebra L_∞ .

Geometría diferencial

Cálculo diferencial e integral más allá de $\mathbb{R}^n$

Las variedades diferenciales (de clase $\mathcal{C}^\infty$) son espacios topológicos en los cuales es posible realizar cálculo diferencial e integral. Los espacios euclídeos $\mathbb{R}^m$ son ejemplos de variedades diferenciales, pero también se incluyen otros espacios como las esferas o los toros. Básicamente una variedad diferencial M es un espacio topológico que localmente es como $\mathbb{R}^m$; aprovechamos esta estructura para definir sin ambigüedad los conceptos de aplicación diferenciable de clase $\mathcal{C}^\infty$ entre variedades: $E \rightarrow M$. El espacio de funciones (aplicaciones diferenciables de M a $\mathbb{R}$) se denota por $\mathcal{C}^\infty(M)$. Las aplicaciones diferenciables de $\mathbb{R}$ a M se denominan caminos.

La integración es un proceso global por lo que necesitamos algo más aparte de la descripción local de M . Necesitamos una noción de elemento infinitesimal de volumen n -dimensional en cada punto, estos medirán el volumen encerrado por n -tuplas de vectores. Para ello necesitamos el concepto de vector tangente a un punto de la variedad. Un elemento del espacio tangente $T_m M$, llamado vector tangente, puede ser pensado como una aproximación lineal a un camino que pasa por m .

PALABRAS CLAVE: Variedad simpléctica; variedad de Poisson; teoría clásica de campos; teoría de campos lagrangiana; variedad multisimpléctica; álgebras L_∞ .

Por tanto, la integración se ha de realizar a lo largo de M sobre aplicaciones que miden volúmenes. Por ejemplo, una aplicación lineal $T_m M \rightarrow \mathbb{R}$ mide longitudes; una aplicación multilinear totalmente antisimétrica $T_m M \times \cdots \times T_m M \rightarrow \mathbb{R}$ mide volúmenes n -dimensionales.

Queremos integrar a lo largo de M , para ello necesitamos unir todos los espacios tangentes a cada punto de m : $TM = \coprod_{m \in M} T_m M$. Esta variedad es un fibrado vectorial $TM \xrightarrow{\pi_{TM}} M$ donde $\pi_{TM}^{-1}(m) = T_m M$.

Un fibrado vectorial, en pocas palabras, es una aplicación diferenciable sobreyectiva $E \xrightarrow{\pi} M$ que localmente es una proyección $\mathbb{R}^k \times \mathbb{R}^m \rightarrow \mathbb{R}^m$ mediante la cual se asocia a cada $m \in M$ un espacio vectorial $\pi^{-1}(m) \cong \mathbb{R}^k$ de forma continuamente diferenciable.

Como ejemplo, tenemos el fibrado $\wedge^n T^* M \xrightarrow{\pi_n} M$ ($\wedge^1 T^* M$ se denota $T^* M$) tal que $\pi_n^{-1}(m) = \{T_m M \times \cdots \times T_m M \rightarrow \mathbb{R} \text{ multilinear y totalmente antisimétrica}\}$. Este fibrado asocia a cada punto m de M una aplicación que mide volúmenes en $T_m M$. Pero estamos interesados en un objeto que realice esta medición a lo largo de todo M : esto es una sección de nuestro fibrado vectorial. Dado un fibrado vectorial, $E \xrightarrow{\pi} M$, el espacio vectorial de secciones globales (o campos) es $\Gamma(E) = \Gamma^\infty(M, E) := \{M \xrightarrow{c} E: c(m) \in \pi^{-1}(m) \forall m \in M\}$.

$\Gamma(\wedge^n T^* M) =: \Omega^n(M)$ es el espacio de formas diferenciales de orden n . Estos son los objetos que podemos integrar a lo largo de una variedad $N \subset M$ de dimensión n . Obsérvese que $\pi_0^{-1}(m) = \text{Hom}(*, \mathbb{R}) \cong \mathbb{R}$, por tanto $\Omega^0(M) = \Gamma(\wedge^0 T^* M) = \Gamma(\mathbb{R} \times M) = \{M \xrightarrow{c} \mathbb{R} \times M: c(m) \in \mathbb{R} \times \pi^{-1}(m) \forall m \in M\} = \mathcal{C}^\infty(M)$.

La diferencial de una aplicación $M \rightarrow \mathbb{R}$ es un elemento de $\Omega^1 M$, veámoslo:

$$\begin{aligned} d: \mathcal{C}^\infty(M) &\rightarrow \Omega^1(M) \\ f &\mapsto \left(v = \sum v_i \frac{\partial}{\partial x_i} \mapsto \sum v_i \frac{\partial f}{\partial x_i} \right). \end{aligned}$$

Esta aplicación se extiende (exigiendo ciertas propiedades) a formas de grado superior produciendo la diferencial de De Rham ($d \circ d = 0$)

$$\Omega^m(M) \xleftarrow{d} \Omega^{m-1}(M) \xleftarrow{d} \cdots \xleftarrow{d} \Omega^1(M) \xleftarrow{d} \mathcal{C}^\infty(M).$$

Combinando la diferencial de De Rham con la integración sobre variedades con borde (pensemos en intervalos cerrados, discos, etcétera) se obtiene un análogo al Teorema fundamental del cálculo llamado Teorema de Stokes:

$$\int_{\partial N} \alpha = \int_N d\alpha,$$

donde ∂N denota el borde de N (y en este caso α es una forma de grado $(n-1)$). En particular si $\partial N = \emptyset$ o si α se anula a lo largo de ∂N tenemos que $\int_N d\alpha = 0$.

Por último, dadas dos formas $\alpha \in \Omega^n(M)$ y $\beta \in \Omega^n(M)$ queremos construir una forma de grado $n+k$ que mida volúmenes producto: esta nueva forma tiene que ser antisimétrica nuevamente, la denotaremos por $\alpha \wedge \beta$.

Variedades simplécticas y de Poisson: el espacio cotangente

Dada una variedad diferencial M podemos añadir diversas estructuras para obtener objetos interesantes: una métrica riemanniana, una estructura compleja, etcétera. En nuestro caso vamos a centrarnos en dos: variedades simplécticas y variedades de Poisson. Como ejemplo podemos pensar en el espacio cotangente a una variedad.

Una variedad simpléctica es un par (M, ω) donde $\omega \in \Omega^2(M)$ cumple que $d\omega = 0$ y que la siguiente aplicación es un isomorfismo:

$$\begin{aligned} \tilde{\omega}: TM &\rightarrow T^*M \\ u &\mapsto (v \mapsto \omega(u, v)). \end{aligned}$$

Una variedad de Poisson es un par $(M, \{-, -\})$ donde $\{-, -\}$, llamado corchete de Poisson, es una aplicación bilinear antisimétrica $\mathcal{C}^\infty(M) \times \mathcal{C}^\infty(M) \rightarrow \mathcal{C}^\infty(M)$ que satisface las identidades de Jacobi y de Leibniz (en otras palabras es un álgebra de Lie que actúa como derivación del producto de funciones).

$$\{f, \{g, h\}\} = \{\{f, g\}, h\} + \{\{h, f\}, g\}. \quad \text{Identidad de Jacobi.}$$

$$\{f \cdot g, h\} = f \cdot \{g, h\} + g \cdot \{f, h\}. \quad \text{Identidad de Leibniz.}$$

Dada una variedad simpléctica (M, ω) y una función cualquiera $f \in \mathcal{C}^\infty(M)$ obtenemos un campo vectorial asociado, llamado campo vectorial hamiltoniano:

$$df \in \Omega^1(M) = \Gamma(T^*M) \xleftarrow{\tilde{\omega}} v_f \in \mathfrak{X} = \Gamma(TM).$$

De esta forma podemos definir un corchete, que resulta ser de Poisson, en $\mathcal{C}^\infty(M)$ mediante la fórmula $\{f, g\} := \omega(v_g, v_f)$.

Como ejemplo de variedad simpléctica, y por tanto de variedad de Poisson, tenemos el espacio cotangente T^*X a una variedad X . Si $p_1, \dots, p_n$ son coordenadas entorno a $x \in X$, denotamos por $q_i \in T_x^*X$ la aplicación $q_i(\frac{\partial}{\partial p_j}) = \delta_{i,j}$. La forma $\omega := \sum dp_i \wedge dq_i \in \Omega^2(T^*X)$ hace del espacio cotangente una variedad simpléctica.

Mecánica clásica: formulación lagrangiana y variedad simpléctica asociada

Olvidamos por un momento las variedades diferenciales, para hacer una pequeña introducción a la teoría de campos lagrangiana. Comencemos con un ejemplo típico

en mecánica clásica: sea $Q = \mathbb{R}^3$ la variedad en la que modelamos el espacio (podría ser otra variedad de dimensión menor o incluso compacta dependiendo del fenómeno que queramos estudiar). Generalmente estamos interesados en la evolución de un punto en Q , esta evolución será descrita como un camino $c: \mathbb{R} \rightarrow Q$.

El principio de acción mínima formula que dicha evolución ha de ser tal que minimice una acción: $\mathcal{A}(c) = \int_{\mathbb{R}} L(c) dt$. Asumimos otro principio, el de localidad, que dice que $L(c(t))$ depende solo del valor de $c(t)$ y de sus derivadas parciales hasta cierto orden (finito). Como ejemplo, podemos pensar en minimizar longitudes (en realidad, energías) $L(c(t)) = |c'(t)|^2$ para así hallar las geodésicas.

En ocasiones, como en el ejemplo anterior, $L(c(t))$ solo depende de las derivadas parciales de $c(t)$ de primer orden. En estos casos L puede ser identificado como una aplicación $TQ \times \mathbb{R} \rightarrow \mathbb{R}$.

Una forma de interpretar el principio de acción mínima es pensando que si $\mathcal{A}(c)$ es mínimo, entonces $\partial(\mathcal{A}(c)) = 0$. Demos sentido a esta última expresión: sea $\epsilon: \mathbb{R} \rightarrow Q$ tal que $c + \epsilon$ es una variación de c . Ahora podemos definir $\partial(\mathcal{A}(c))$ y hacer algunos cálculos:

$$\begin{aligned} \partial(\mathcal{A}(c))(\epsilon) &:= \int_{\mathbb{R}} \partial L(c(t), \epsilon(t)) dt = \int_{\mathbb{R}} \sum \left(\frac{\partial L}{\partial q_i} \epsilon_i + \frac{\partial L}{\partial \dot{q}_i} \dot{\epsilon}_i \right) dt = \\ &= \int_{\mathbb{R}} \sum \left(\frac{\partial L}{\partial q_i} - \frac{d}{dt} \frac{\partial L}{\partial \dot{q}_i} \right) \epsilon_i dt + \int_{\mathbb{R}} \sum \frac{d}{dt} \left(\frac{\partial L}{\partial \dot{q}_i} \epsilon_i \right) dt \end{aligned}$$

Los cálculos realizados han utilizado, en primer lugar la regla de la cadena ($df(x, y) = \frac{df}{dx}(x, y)dx + \frac{df}{dy}(x, y)dy$) y para la última igualdad, integración por partes ($\int \frac{d}{dt}(fg)dt = \int \frac{d}{dt}(f)gdt + \int f \frac{d}{dt}(g)dt$) sobre $f = \frac{\partial L}{\partial \dot{q}_i}$ y $g = \epsilon_i$. El último término generalmente se descarta considerando intervalos en vez de $\mathbb{R}$ y variaciones que se anulen en los extremos del intervalo (véase la anotación al teorema de Stokes).

$\partial(\mathcal{A}(c)) = 0$ si para toda variación el primer término de la ecuación anterior se anula: es decir si $\frac{\partial L}{\partial q_i} - \frac{d}{dt} \frac{\partial L}{\partial \dot{q}_i} = 0$. Esta ecuación en derivadas parciales, cuyas soluciones son las evoluciones de nuestro sistema, se denomina ecuación de Euler-Lagrange.

El segundo término, si bien no es interesante en la búsqueda de soluciones del sistema, es importante para otro tipo de cuestiones que no serán explicadas aquí. Sin entrar en más detalles, podemos pensar que $\epsilon_i = dq_i$ dado que ϵ_i mide la variación en la dirección marcada por q_i . Igualmente, llamando p_i a $\frac{\partial L}{\partial \dot{q}_i}$ obtenemos que $\sum d(\frac{\partial L}{\partial \dot{q}_i} \epsilon_i) = \sum d(p_i \wedge dq_i) = \sum dp_i \wedge dq_i$: es decir, recuperamos la forma simpléctica estándar en T^*Q . (Observación: hemos asumido implícitamente una condición de no degeneración de la Hessiana de L , conocida como la condición de Legendre).

Teoría de campos lagrangiana: una invitación a la geometría simpléctica superior

Pasamos ahora a describir una teoría de campos lagrangiana en general. Los principios van a ser los mismos, habrá una acción local que actúa sobre un espacio de campos (en el caso anterior, caminos) que queremos minimizar. Advertimos en este momento que lo dicho a continuación es una introducción a la teoría que dista mucho de ser formalmente correcta.

Observemos en primer lugar que los caminos en Q son secciones del fibrado vectorial $Q \times \mathbb{R} \rightarrow \mathbb{R}$. Ahora tendremos un fibrado vectorial $E \rightarrow M$ sobre una variedad de dimensión m y el espacio de campos será $\mathcal{E} := \Gamma(E)$.

La acción en este caso es $\mathcal{A}(c) = \int_M \tilde{L}(c)$ donde $\tilde{L}(c) \in \Omega^m(M)$ (para poder integrar a lo largo de M). Además $\tilde{L}$ es local en el mismo sentido que antes, $\tilde{L}(c(m))$ depende solo de las derivadas parciales de $c(m)$ hasta cierto orden (finito).

$\partial\tilde{L}$ tiene ahora una interpretación sencilla. Recordemos que $\mathcal{E} = \Gamma(E)$ es un espacio vectorial, por ello podemos tomar formas diferenciales en $\mathcal{E}$. En realidad, solo estamos interesados en formas locales: $\Omega_{\text{loc}}^n(\mathcal{E})$. Se puede entender $\tilde{L}$ como un elemento de $\mathcal{C}_{\text{loc}}^\infty(\mathcal{E}) \otimes \Omega^m(M)$. Para distinguir la diferencial de De Rham en $\mathcal{E}$ y en M denotamos a la primera por ∂ y a la segunda por d . $\partial\tilde{L}$ es sencillamente el elemento de $\Omega_{\text{loc}}^1(\mathcal{E}) \otimes \Omega^m(M)$ que se obtiene al aplicar la diferencial en la dirección de $\mathcal{E}$.

Escribamos, como antes, $\partial(\mathcal{A}(c))(\epsilon) := \int_M \partial\tilde{L}(c, \epsilon)$. El cálculo realizado en la sección previa toma ahora forma de teorema:

Teorema 1 (Zuckermann [2]). $\partial\tilde{L} = EL + d\gamma$ donde $EL(-, \epsilon(m))$ solo depende del valor de $\epsilon(m)$ y no de sus derivadas parciales y donde $\gamma \in \Omega_{\text{loc}}^1(\mathcal{E}) \otimes \Omega^{m-1}(M)$.

Tomando como antes $N \subset M$ una variedad con borde y centrándonos en variaciones ϵ que se anulan en ∂N obtenemos que los campos que minimizan la acción son aquellos que satisfacen las ecuaciones de Euler-Lagrange $EL(c, -) = 0$.

Observemos qué sucede en el otro término: $\partial\gamma =: \omega \in \Omega_{\text{loc}}^2 \otimes \Omega^{m-1}(M)$. Cualquier forma α de grado $m-1$ en M tal que $d\alpha = 0$ se denomina corriente conservada. $\omega(*, -)$ es una corriente conservada, llamada universal ([2]), y por tanto $\int_N \omega$ es una forma cerrada (tal vez degenerada, en este caso diríamos que pre-simpléctica) que no es independiente de la elección de la hipersuperficie N (de hecho depende de la clase de homología de N , este problema no sucedía cuando $M = \mathbb{R}$).

Para eludir la pérdida de generalidad que conlleva la elección de la hipersuperficie N se puede estudiar ω en sí misma. $\omega \in \Omega_{\text{loc}}^2 \otimes \Omega^{m-1}(M) \subset \Omega_{\text{loc}}^{2+(m-1)=m+1}(\mathcal{E} \times M)$, este nuevo espacio tiene como diferencial $D = \partial + d$, $D\omega = 0$.

Geometría simpléctica superior y álgebras L_∞

Podemos ahora generalizar a formas de grado superior la definición de variedad simpléctica e intentar ver qué obtenemos como análogo al corchete de Poisson. Como ejemplo fundamental tenemos a $\mathcal{E} \times M$ y la corriente conservada universal ω .

Una variedad simpléctica de orden n (o variedad n -pléctica [1]) es un par (M, ω) donde $\omega \in \Omega^{n+1}(M)$ cumple que $d\omega = 0$ y que la siguiente aplicación es inyectiva:

$$\begin{aligned} \tilde{\omega}: \mathfrak{X}(M) &\rightarrow \Omega^n(M) \\ u &\mapsto (v_1, \dots, v_n \mapsto \omega(u, v_1, \dots, v_n)). \end{aligned}$$

Dado que $\tilde{\omega}$ no es sobreyectiva en general, no toda forma $\alpha \in \Omega^{n-1}(M)$ es tal que $d\alpha \in \text{im}(\tilde{\omega})$ (al contrario que en el caso simpléctico, $n = 1$). Denominamos forma hamiltoniana a las formas $\alpha \in \Omega^{n-1}(M)$ que sí estén en ese caso, es decir $d\alpha = \omega(v_\alpha, -)$ para cierto campo vectorial v_α , también llamado hamiltoniano.

Dadas dos formas hamiltonianas α y β , $\{\alpha, \beta\} := \omega(v_\beta, v_\alpha, -)$ es nuevamente una forma hamiltoniana. Este corchete es antisimétrico pero no satisface la identidad de Jacobi:

$$\{\alpha, \{\beta, \gamma\}\} - \{\{\alpha, \beta\}, \gamma\} - \{\{\gamma, \alpha\}, \beta\} = d\omega(v_\gamma, v_\beta, v_\alpha) \neq 0.$$

En cambio tenemos el siguiente resultado:

Teorema 2 (Rogers, [1]). *Dada una variedad simpléctica de orden n la familia de aplicaciones $\{l_k(\alpha_1, \dots, \alpha_k) := \omega(v_{\alpha_1}, \dots, v_{\alpha_k}, -)\}_{k=1}^{m+1}$ dota de una estructura de álgebra L_∞ a $\Omega_{\text{ham}}(M) = \left(\Omega_{\text{ham}}^{n-1}(M) \xleftarrow{d} \Omega^{n-2}(M) \xleftarrow{d} \dots \xleftarrow{d} \Omega^1(M) \xleftarrow{d} \mathcal{C}^\infty(M)\right)$.*

En geometría derivada, los espacios vectoriales son sustituidos por cadenas de complejos (como $\Omega_{\text{ham}}(M)$). Igualmente, las álgebras L_∞ son los análogos a las álgebras de Lie. De manera poco precisa, las aplicaciones l_k toman valores en $\Omega_{\text{ham}}(M)$, pero no de forma arbitraria, sino que respetando el orden de las formas involucradas en cierta manera. La identidad de Jacobi involucra ahora más sumandos:

$$\sum_{i+j=k+1} \sum_{\sigma \in S_i^{j-1}} \varepsilon(\sigma) l_j(l_i(v_{\sigma(1)}, \dots, v_{\sigma(i)}), v_{\sigma(i+1)}, \dots, v_{\sigma(k+1)}) = 0.$$

Bibliografía

- [1] Rogers, C. L. (2011). *Higher symplectic geometry*, Tesis doctoral, Universidad de California en Riverside.
- [2] Zuckerman, G. J. (1987). *Action principles and global geometry*, in Proceedings of the Conference on Mathematical Aspects of String Theory, California, USA, August 1986, pp. 259-284.

Incendios forestales y estadística

M^a Isabel Borrajo García

Departamento de Estadística e Investigación Operativa

15 de abril de 2015

Resumen

La modelización de incendios forestales puede abordarse utilizando la metodología de procesos puntuales, que combina los procesos estocásticos (en lo que a eventos se refiere) con la inferencia estadística asociada a diversas características (marcas) de esos eventos.

En la caracterización de este tipo de procesos es crucial el análisis de su estructura de 1er orden a través de la función de intensidad, que será uno de nuestros principales objetivos. Se hará una breve introducción a la teoría de procesos puntuales, centrándonos en la función de intensidad y en las posibilidades existentes para su estimación. Finalmente se aplica esta metodología a los datos de incendios forestales registrados en Galicia en el año 2006.

Introducción a los procesos puntuales

Los procesos puntuales pueden definirse intuitivamente como modelos matemáticos que describen la disposición de objetos distribuidos de manera irregular o aleatoria en el espacio. Esta metodología se ha aplicado a lo largo de los años en numerosos campos de investigación como ecología (distribución de especies vegetales), biología molecular (distribución de un determinado tipo de células), astronomía (distribución de estrellas o galaxias), epidemiología (distribución de casos de enfermedades)..., pueden verse estos y otros ejemplos en [3].

Formalmente, un proceso puntual, X , es un modelo estocástico que genera un número finito de eventos en un conjunto $W \subset \mathbb{R}^d$ (denominado **región de observación**). Lo habitual es disponer de una realización de ese proceso que denotamos por $\{X_1, \dots, X_N\}$ y a cada una de esas localizaciones se le denomina **evento**. Tengamos en cuenta que el número de eventos, N , varía con cada realización, por consiguiente, tanto dicho número como las localizaciones son aleatorios.

Existen diversas condiciones de regularidad que se pueden exigir a los procesos puntuales y que, en caso de cumplirse, facilitan enormemente los desarrollos teóricos. Las principales son **estacionariedad** e **isotropía**; la primera de ellas se

traduce en que las propiedades del proceso son invariantes por traslaciones; mientras que la isotropía consiste en que dichas propiedades sean invariantes por rotaciones. Por tanto, las propiedades de un proceso que sea simultáneamente estacionario e isotrópico, dependerán únicamente de la longitud del vector diferencia de posiciones (un valor escalar), en lugar de las localizaciones.

Otro concepto importante es el de CSR (*Complete Spatial Randomness*), que es como el “ruido blanco” de un proceso puntual, ya que se caracteriza por la ausencia de estructura en los datos. Habitualmente suele emplearse como hipótesis nula de test estadísticos para determinar si hay o no estructura en un determinado patrón.

Definición 1. *Un proceso puntual cumple la condición de **CSR** si verifica los siguientes postulados:*

- *El número de eventos, N , en una región W , sigue una distribución $Pois(\lambda|W|)$ donde λ es la intensidad del proceso, es decir, el número de eventos por unidad de área y $|W|$ denota la medida de la región W .*
- *Condicionado a una realización con n eventos, $\{X_1, \dots, X_n\}$, los $\{X_i\}_{i=1}^N$ son una muestra aleatoria simple de una distribución $Unif(W)$.*

Las funciones que caracterizan a los procesos puntuales se agrupan en dos clases, las de primer orden, entre las que destaca fundamentalmente la función de intensidad, y las de segundo orden.

Definición 2. *Dado un proceso puntual $X \subset W$ y siendo N la medida de contar, i.e $N(A)$, con $A \subset W$ indica el número de eventos que hay en A ; se define la función de **intensidad** o intensidad de primer orden como*

$$\lambda(x) = \lim_{|dx| \rightarrow 0} \frac{\mathbb{E}[N(dx)]}{|dx|}.$$

Definición 3. *Se define la **intensidad de segundo orden** como*

$$\lambda_2(x, y) = \lim_{|dx|, |dy| \rightarrow 0} \frac{\mathbb{E}[N(dx)N(dy)]}{|dx||dy|}. \quad (1)$$

La función (1) es una medida de la estructura de dependencia de los eventos en W , y está directamente relacionada en el caso de procesos estacionarios e isotrópicos con la K función

$$K(t) = \lambda^{-1} \mathbb{E}[N_0(t)],$$

siendo $N_0(t)$ el número de eventos a una distancia a lo sumo t de un evento escogido arbitrariamente. Esta función ha sido especialmente útil en el contraste de modelos paramétricos ya que presenta formas muy características y es sencilla de estimar.

Haciendo un paralelismo con la estadística clásica, en la que la hipótesis de normalidad está muy extendida y permite simplificar notablemente los desarrollos metodológicos, en procesos puntuales se definen los procesos de Poisson.

Definición 4. Un proceso puntual X se dice que es Poisson homogéneo si verifica las mismas condiciones que CSR (ver Definición 1).

Definición 5. Un proceso puntual se dice que es Poisson no homogéneo si cumple:

- El n^o de eventos en W sigue $\text{Pois}(\int_W \lambda(x)dx)$.
- Dados n eventos, estos forman una muestra independiente de una distribución en W con densidad proporcional a $\lambda(x)$.

Bajo la condición de proceso de Poisson no homogéneo se ha desarrollado la mayor parte de la metodología estadística en este campo, aunque debido a su simplicidad, algunos autores emplean modelos paramétricos un poco más complejos:

Definición 6. Un proceso de Cox se define a través de los siguientes postulados:

- $\{\Lambda(x)/x \in \mathbb{R}^2\}$ es un proceso estocástico no negativo.
- Condicionalmente a una realización de $\{\Lambda(x)/x \in \mathbb{R}^2\}$, $\lambda(x)$, los eventos forman un proceso de Poisson no homogéneo con función de intensidad $\lambda(x)$.

Esta última clase de modelos se denomina doblemente estocásticos, debido a la existencia del campo Λ que añade un nuevo nivel de aleatoridad al proceso. En general dan lugar a patrones agregados que suelen estar presentes en la modelización de distribución de individuos en competencia, como por ejemplo algunas especies vegetales que “compiten” por nutrientes o luz solar.

La función de intensidad de primer orden

En *Definición 2* hemos definido formalmente el concepto de intensidad de primer orden; en el caso más sencillo en el que dicha función sea constante sobre W indica el número medio de eventos por unidad de área y su estimación sería trivial, $\hat{\lambda} = n/|W|$.

Esta idea intuitiva no es válida en el caso más general en el que realmente λ sea función de las localizaciones. Inicialmente se usaron hipótesis paramétricas en su estimación, sin embargo es conocido que estas técnicas podrían dar lugar a estimadores poco precisos en el caso en que la intensidad teórica se desviase del modelo asumido, es por ello que se emplean técnicas no paramétricas.

La primera aportación a este respecto aparece en [2], que propuso un estimador tipo núcleo para la función de intensidad, (DIE):

$$\hat{\lambda}_h(x) = \frac{1}{p_h(x)} \sum_{i=1}^N k_h(x - X_i) = \frac{1}{p_h(x)h^2} \sum_{i=1}^N k((x - X_i)/h), \quad (2)$$

donde k es un núcleo bivalente radialmente simétrico, $h > 0$ es la ventana o parámetro de suavizado, k_h es el núcleo reescalado y $p_h(x) = \int_W h^{-2}k((x - y)/h)dy$ es la corrección de efecto frontera.

El problema es que este estimador es inconsistente, esto es, al llevar al límite el tamaño muestral no conseguimos que el error de estimación converja a cero. Debido a este hecho, el uso de este estimador a lo largo de los años se ha reducido casi exclusivamente a nivel exploratorio.

La cuestión de la inconsistencia perduró varias décadas, no fue hasta 2006 en [1] que se propuso una solución. La idea es sencilla y consiste en aplicar los conceptos de estimación no paramétrica de densidad (que sí proporcionan estimadores consistentes), al caso de la función de intensidad. Para ello se define una densidad artificial $\lambda_0(x) = \lambda(x) / \int_W \lambda(u) du$, se estima mediante el estimador tipo núcleo y se recupera finalmente una estimación para la intensidad, (KDE):

$$\hat{\lambda}_0(x) = \frac{1}{N} \sum_{i=1}^N \frac{1}{h} k\left(\frac{x - x_i}{h}\right) \mathbb{I}_{\{N \neq 0\}}.$$

Otra solución al problema de la inconsistencia de (2) es la propuesta de [4], que se sustenta en asumir $\lambda(\cdot) = \rho[f(\mathbf{Z}(\cdot))]$, donde $\mathbf{Z}$ representa un campo aleatorio (covariables) y ρ es una función continua no negativa. El estimador propuesto, (CIE), es

$$\hat{\lambda}_{Gh}(x) = \frac{\sum_{i=1}^N k\left(\frac{\|\mathbf{Z}(x) - \mathbf{Z}(x_i)\|}{h}\right)}{\int_W k\left(\frac{\|\mathbf{Z}(u) - \mathbf{Z}(x)\|}{h}\right) du}, \quad (3)$$

siendo k un núcleo univariante y $\|\cdot\|$ la norma euclídea usual.

Además del hecho de ser consistente, (3) presenta la ventaja de trabajar en dimensión uno aun cuando el proceso esté definido en dimensión superior.

Estudio de simulación

Presentados los principales estimadores de la función de intensidad es interesante comparar su comportamiento; para ello se ha realizado un completo estudio de simulación del que hemos seleccionado una parte representativa para incluir en este resumen.

Se han generado $B = 100$ realizaciones de cuatro procesos puntuales con funciones de intensidad $\lambda_i(x) = \exp(\beta_0 + \beta_1 Y_i(x))$, donde $Y_1 = Z$, $Y_2 = Z + e$, $Y_3 = d_R(x)$ y $Y_4 = d_R(x)^2$ y e el término de error. Para cada una de las realizaciones de los procesos aplicamos DIE, KDE y CIE (en el cálculo de CIE usamos la covariable Z para los modelos 1 y 2, y d_R para los modelos 3 y 4).

El criterio de error que empleamos es una versión relativa del error cuadrático integrado (ISE; *Integrated Squared Error*):

$$ISE = \int_W \left(\frac{\hat{\lambda}(x) - \lambda(x)}{\lambda(x)} \right)^2 dx.$$

Los resultados obtenidos pueden verse en la Tabla 1 en la que se presentan la media de los ISE's para cada modelo y estimador, así como la correspondiente desviación típica muestral de dichos errores.

Se observa que KDE mejora los resultados de DIE tal y como era esperable. Por otra parte, CIE obtiene mejores resultados que los otros dos estimadores gracias a la información adicional que le proporciona la covariable, aunque si bien es cierto que en el modelo 2, en donde la información de la covariable está perturbada por la presencia del término de error, esta mejoría no es tan notable. En los modelos 3 y 4 la ganancia del CIE no es significativa respecto al KDE, posiblemente ligado al carácter determinístico de la covariable.

	Model 1	Model 2	Model 3	Model 4
DIE	0.1640 0.0374	0.3312 0.0563	0.0691 0.0175	0.0612 0.0206
KDE	0.1459 0.0221	0.3282 0.0539	0.0411 0.0112	0.0297 0.0097
CIE	0.0233 0.0194	0.2218 0.0505	0.0324 0.0131	0.0296 0.0116

Tabla 1: Media y desviación típica de los ISEs de los cuatro modelos resultados de aplicar DIE, KDE y CIE.

Análisis de los incendios forestales en Galicia

Los incendios forestales son un gran problema socio-económico y medioambiental en la actualidad; particularmente en Galicia los incendios intencionados son la principal causa de destrucción de superficie forestal, de hecho, más de la mitad de los incendios ocurridos anualmente en España, están localizados en Galicia.

Conocer la distribución espacial de los incendios forestales es un factor clave en la prevención y desarrollo de planes de extinción. Si asociamos los incendios registrados en un determinado período de tiempo con las coordenadas espaciales de su punto de ignición, éstos pueden ser considerados como un patrón espacial generado por un proceso estocástico, y la distribución espacial de los puntos de ignición se identifica con la intensidad.

Hemos considerado los datos registrados el día 09/08/2006 en el que hubo un total de 190 incendios; y para aplicar el CIE hemos considerado la covariable temperatura media, obtenida mediante un procedimiento de suavizado con núcleo gaussiano a partir de los datos de 44 estaciones meteorológicas repartidas por toda la comunidad.

En la Figura 1 podemos ver los puntos de ignición para ese día (izq.), junto con la temperatura media (dcha.) que hemos empleado como covariable, en donde los puntos representan las 44 estaciones de medición.

Observando la Figura 2 podemos concluir que DIE y KDE aportan estimaciones muy similares y razonables con los datos proporcionados aunque si bien es cierto que el primero de ellos presenta cierta tendencia a la sobresuavización.



Figura 1: Puntos de ignición (izq.) y T^a media con las estaciones de medición (dcha.).

El resultado de CIE difiere de los anteriores y no reproduce adecuadamente el patrón puntual observado, de hecho notamos que este estimador se deja llevar en exceso por la información de la covariable, en el sentido de que asigna intensidades más altas en regiones con temperaturas elevadas. El hecho de que CIE no replique el patrón, se debe a que el proceso estocástico que genera los incendios en Galicia no depende de la temperatura, es decir, ésta no es el principal factor de riesgo. En efecto, se sabe que la mayoría de dichos incendios están fuertemente vinculados con el factor humano (incendios provocados, negligencias...) y estas estimaciones lo corroboran.

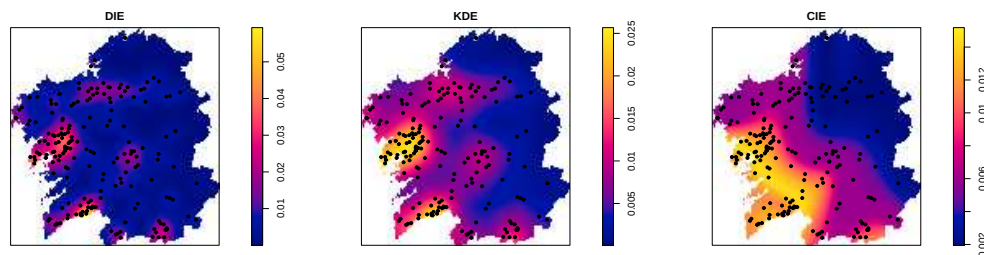


Figura 2: DIE, KDE y CIE (de izq. a dcha.) aplicados a los datos del 09/08/2006.

Bibliografía

- [1] Cucala, L. (2006) *Espacements bidimensionnels et données entachées d'erreurs dans l'analyse des processus ponctuels spatiaux*, PhD Thesis, Université des Sciences de Toulouse I.
- [2] Diggle, P.J. (1985) *A Kernel Method for Smoothing Point Process Data*, Journal of the Royal Statistical Society (Series C), **34**, pp. 138–147.
- [3] Diggle, P. (2013) *Statistical analysis of spatial and spatio-temporal point patterns*, CRC Press.
- [4] Guan, Y. (2008) *On consistent nonparametric intensity estimation for inhomogeneous spatial point processes*, Journal of the American Statistical Association, **103** (483), 1238–1247.

Dinámica de exoplanetas y exosatélites

Pedro Pablo Campo Díaz

Observatorio Astronómico Ramón María Aller

29 de abril de 2015

Resumen

El estudio de los planetas extrasolares es uno de los campos que mayor crecimiento ha experimentado en los últimos años en el ámbito de la Astronomía. El desarrollo de las técnicas de observación ha permitido la detección de objetos de masa subestelar que anteriormente quedaban fuera de los límites de la instrumentación disponible.

Desde el punto de vista de la Mecánica Celeste, estos sistemas planetarios plantean problemas muy interesantes, sobre todo en sistemas binarios o múltiples, dado que a las consideraciones puramente dinámicas hay que añadir otras de índole física como la habitabilidad de los planetas (o satélites) del sistema.

Introducción

En Mecánica Celeste se denomina problema de n cuerpos al estudio del movimiento de n objetos sujetos a sus mutuas atracciones gravitatorias. Solamente en el caso de $n = 2$ existe solución analítica general del problema, ya conocida desde el siglo XVIII. Para $n \geq 3$ cada problema debe ser tratado de diferente forma.

Por ejemplo, los primeros trabajos que estudiaron la dinámica de exoplanetas utilizaron el problema restringido de $n + \nu$ cuerpos, en el que las masas de los ν cuerpos se consideran despreciables con respecto a las de los n cuerpos y por lo tanto su presencia no afecta al movimiento de las otras [9]. Sin embargo este modelo no es el adecuado, pues varios métodos de detección se basan precisamente en la influencia que ejercen los planetas sobre la estrella o sobre otros planetas.

Formulación hamiltoniana

Dado un sistema de coordenadas canónicas $(\mathbf{p}, \mathbf{q})$, se llama función hamiltoniana o simplemente hamiltoniano $\mathcal{H} = \mathcal{H}(\mathbf{p}, \mathbf{q}, t)$ a una función $\mathcal{H} : \mathbb{R}^{2n} \rightarrow \mathbb{R}$ tal que cumple las ecuaciones de Hamilton:

PALABRAS CLAVE: Exoplanetas; exosatélites; Mecánica Celeste.

$$\begin{aligned}\frac{d\mathbf{p}}{dt} &= -\frac{\partial \mathcal{H}}{\partial \mathbf{q}}, \\ \frac{d\mathbf{q}}{dt} &= \frac{\partial \mathcal{H}}{\partial \mathbf{p}}.\end{aligned}$$

En Mecánica Celeste $\mathbf{p}$ son las posiciones y $\mathbf{q}$ los momentos. En general se utiliza esta formulación para representar las ecuaciones del movimiento, por lo que cuando se trata de resolver un problema de n cuerpos, lo habitual es hallar en primer lugar el hamiltoniano asociado.

Una vez hecho esto, hay diferentes técnicas, tanto analíticas como numéricas, que se pueden utilizar para resolver el sistema.

Resolución analítica

Para la resolución analítica de problemas de n cuerpos se suele usar la teoría de perturbaciones. Ésta consiste en desarrollar en serie el hamiltoniano con respecto a un pequeño parámetro, de forma que:

$$\mathcal{H}(\mathbf{p}, \mathbf{q}, t) = \sum_{i=0}^{\infty} \epsilon^i \mathcal{H}_i(\mathbf{p}, \mathbf{q}, t).$$

Esto permite truncar el desarrollo a partir de un cierto valor n , de forma que se eliminan partes del hamiltoniano que tienen muy poca influencia en la resolución del mismo. También se utilizan transformaciones de Lie, basándose en el siguiente teorema:

Teorema 1. Sean ξ_j, η_j un conjunto de $2n$ coordenadas canónicas y $f(\boldsymbol{\xi}, \boldsymbol{\eta})$ y $S(\boldsymbol{\xi}, \boldsymbol{\eta})$ funciones arbitrarias de $\boldsymbol{\xi}$ y $\boldsymbol{\eta}$. Se definen los operadores D_S^n ($n = 0, 1, 2, \dots$) como:

$$D_S^0 f = f, \quad D_S^1 f = \{f, S\}, \quad D_S^n f = D_S^{n-1}(D_S^1 f), \quad n \geq 2,$$

siendo $\{ , \}$ los corchetes de Poisson. Si ϵ es un pequeño parámetro independiente de $\boldsymbol{\xi}$ y $\boldsymbol{\eta}$, el conjunto de $2n$ coordenadas x_j, y_j , definidas por:

$$f(\mathbf{x}, \mathbf{y}) = \sum_{n=0}^{\infty} \frac{\epsilon^n}{n!} D_S^n f(\boldsymbol{\xi}, \boldsymbol{\eta}), \quad (1)$$

es canónico si la serie en el lado derecho de (1) es convergente.

Hori [8] utiliza este resultado para formular su método paramétrico, que se basa en transformaciones de coordenadas canónicas para, a partir de un hamiltoniano $\mathcal{H} = \sum_{n=0}^{\infty} \mathcal{H}_n$, obtener sucesivamente nuevos hamiltonianos $\mathcal{H}^* = \sum_{n=0}^{\infty} \mathcal{H}_n^*$ que resulten más sencillos de integrar, generalmente eliminando variables.

Abad y Ribera [3] generalizaron este método para dos parámetros, y posteriormente Andrade [4] para un número arbitrario.

Un ejemplo de aplicación del método biparamétrico aplicado a exoplanetas puede verse en [6], donde se obtiene la integración analítica de un sistema formado por una binaria (P_3, P_4) con un planeta (P_2) girando en torno a una de las estrellas (P_3) y un satélite (P_1) en torno a este último. La masa de cada cuerpo se denotará por m_i .

Para ello, se formula el problema en un sistema de coordenadas escalonadas, dado por Abad [1] para tratar sistemas jerarquizados. En este sistema las coordenadas son:

- $\vec{r}_2$: vector de posición del satélite al planeta.
- $\vec{r}_1$: vector de posición de la estrella (P_3) al centro de masas del satélite y el planeta.
- $\vec{r}_0$: vector de posición de la estrella (P_4) al centro de masas de los otros tres.

Se obtiene el Hamiltoniano de la forma

$$\mathcal{F} = \mathcal{T} + \mathcal{V},$$

donde $\mathcal{T}$ es la energía cinética dada por

$$\mathcal{T} = \frac{1}{2}(\mathcal{M}_0 \dot{\vec{r}}_0 + \mathcal{M}_1 \dot{\vec{r}}_1 + \mathcal{M}_2 \dot{\vec{r}}_2),$$

con

$$\begin{aligned}\mathcal{M}_2 &= \frac{m_1 m_2}{m_1 + m_2}, \\ \mathcal{M}_1 &= \frac{(m_1 + m_2)m_3}{m_1 + m_2 + m_3}, \\ \mathcal{M}_0 &= \frac{(m_1 + m_2 + m_3)m_4}{m_1 + m_2 + m_3 + m_4}.\end{aligned}$$

El potencial $\mathcal{V}$ se formula originalmente en un sistema baricéntrico inercial, ξ_i :

$$\mathcal{V} = -\tilde{G} \sum_{1 \leq i < j \leq 4} \frac{m_i m_j}{|\vec{\xi}_j - \vec{\xi}_i|},$$

donde $\tilde{G}$ es la constante gravitatoria. Al transformarlo al sistema de coordenadas escalonadas se obtiene:

$$\begin{aligned}
\mathcal{V} = & -\tilde{G} \left(\frac{m_1 m_2}{r_2} + \frac{(m_1 + m_2) m_3}{r_1} + \frac{(m_1 + m_2 + m_3) m_4}{r_0} + \sum_{n=2}^{\infty} N_n \frac{r_2^n}{r_1^{n+1}} P_n(\sigma_0) \right. \\
& + \sum_{n=1}^{\infty} N_n^* \frac{r_1^n}{r_0^{n+1}} P_n(\sigma_2) + \frac{1}{r_0} \sum_{n=0}^{\infty} \binom{-1/2}{n} \\
& \left. + \sum_{j1, \dots, j5=1, \sum ji=n}^n \frac{n!}{j1! \dots j5!} \sigma_2^{j3} \sigma_1^{j4} \sigma_0^{j5} N_n^{**} \left(\frac{r_1}{r_0} \right)^{2j1+j3+j5} \left(\frac{r_2}{r_0} \right)^{2j2+j4+j5} \right)
\end{aligned}$$

$P_n(x)$ es el polinomio de Legendre de orden n y N_n , N_n^* y N_n^{**} son constantes que dependen de las masas. σ_1 , σ_2 y σ_0 son los cosenos de los ángulos entre $\vec{r}_2$ y $\vec{r}_0$, $\vec{r}_1$ y $\vec{r}_0$ y $\vec{r}_1$ y $\vec{r}_2$, respectivamente y $r_i = |\vec{r}_i|$, $i = 0, 1, 2$.

Se hace un cambio a coordenadas de Delaunay, que es una transformación canónica:

$$\begin{aligned}
\bar{L} &= \sqrt{\mu a}, \\
\bar{G} &= \bar{L} \sqrt{1 - e^2}, \\
\bar{H} &= \bar{G} \cos I, \\
l &= M, \\
g &= \omega, \\
h &= \Omega,
\end{aligned}$$

I es la inclinación de la órbita, ω el argumento del periastro, Ω el ángulo del nodo, a el semieje mayor de la órbita, e la excentricidad, y M la anomalía media. $\mu = \frac{\mathcal{M}_A}{\mathcal{M}_A + \mathcal{M}_B}$ (siendo $\mathcal{M}_A \geq \mathcal{M}_B$ las masas que intervienen en la órbita, aquí serían las $\mathcal{M}_i$). Se obtiene el hamiltoniano:

$$\begin{aligned}
\mathcal{F} = & \frac{A_0}{L_0^2} + \frac{A_1}{L_1^2} + \frac{A_2}{L_2^2} + \frac{\epsilon_0^2}{2!} B_1 \frac{L_1^2}{r_0^2} \frac{r_1}{a_1} P_1(\sigma_2) + \frac{\epsilon_0^3}{3!} B_2 \frac{L_1^4}{r_0^3} \frac{r_1^2}{a_1^2} P_2(\sigma_2) \\
& + \frac{\epsilon_1^3}{3!} C_2 \frac{L_2^4}{r_1^3} \left(\frac{r_2}{a_2} \right)^2 P_2(\sigma_0) + \frac{\epsilon_1^4}{4!} C_3 \frac{L_2^6}{r_1^4} \left(\frac{r_2}{a_2} \right)^3 P_3(\sigma_0) \\
& + \frac{\epsilon_2^3}{3!} D_2 \frac{L_2^4}{r_0^3} \left(\frac{r_2}{a_2} \right)^2 P_2(\sigma_1) + \dots
\end{aligned}$$

A_i , B_i , C_i y D_i son de nuevo constantes, y la dependencia en las coordenadas viene dada de forma implícita por los r_i y los ángulos σ_i . Por el momento se mantienen algunas coordenadas antiguas para que la formulación del hamiltoniano sea más sencilla. Aquí $\epsilon_0 = a_2^{(0)}/q_1^{(0)}$, $\epsilon_1 = a_1^{(0)}/q_0^{(0)}$ y $\epsilon_2 = a_2^{(0)}/q_0^{(0)}$, donde $a_i^{(0)}$, $i=0,1,2$, son los semiejes mayores y $q_i^{(0)}$, $i=0,1,2$, las distancias al periastro. Dado que es un sistema jerarquizado, se pueden utilizar como pequeños parámetros.

Utilizando unas fórmulas empíricas dadas por Holman y Wiegert (REF) para la estabilidad en este tipo de sistemas, se hace un truncamiento del hamiltoniano a orden 4 en ϵ_0 , 3 en ϵ_1 y 2 en ϵ_2 . Este último aparece por vez primera en orden 3, por lo que quedan dos parámetros. Físicamente, al eliminar los términos de ϵ_2 , estamos eliminando la influencia del satélite sobre la estrella más lejana.

Finalmente, se realizan transformaciones sucesivas utilizando el método biparamétrico para eliminar la dependencia en las variables angulares, con lo que se obtiene el hamiltoniano final:

$$\begin{aligned} \mathcal{F} = & \frac{A_0}{L_0^2} + \frac{A_1}{L_1^2} + \frac{A_2}{L_2^2} + \frac{\epsilon_0^3}{3!} \frac{B_2(5 + 3(\frac{G_1}{L_1})^2)(-1 + 3(\frac{H_0}{G_0})^2)L_1^4}{16 \frac{L_0^4 G_0^2}{M_0^6 \mu_0^3}} + \\ & + \frac{\epsilon_1^3}{3!} \frac{C_2(5 + 3(\frac{G_2}{L_2})^2)(-1 + 3(\frac{H_2}{G_2})^2)L_2^4}{16 \frac{L_1^4 G_1^2}{M_1^6 \mu_1^3}}. \end{aligned}$$

Y las ecuaciones del movimiento:

$$\begin{aligned} \frac{dL_i}{dt} &= -\frac{\partial \mathcal{F}}{\partial l_i} = 0; & \frac{dl_i}{dt} &= -\frac{\partial \mathcal{F}}{\partial L_i}, & i &= 0, 1, 2, \\ \frac{dG_i}{dt} &= -\frac{\partial \mathcal{F}}{\partial g_i} = 0; & \frac{dg_i}{dt} &= -\frac{\partial \mathcal{F}}{\partial G_i}, & i &= 0, 1, 2, \\ \frac{dH_i}{dt} &= -\frac{\partial \mathcal{F}}{\partial h_i} = 0; & \frac{dh_i}{dt} &= -\frac{\partial \mathcal{F}}{\partial H_i}, & i &= 0, 1, 2. \end{aligned}$$

Técnicas numéricas

A menudo la integración analítica de un problema es muy complicada o no resulta eficiente computacionalmente. Por ello se utilizan también métodos numéricos. Por una parte están los métodos genéricos para la resolución de sistemas de Ecuaciones Diferenciales Ordinarias, como pueden ser los desarrollos en series de Taylor [2].

Sin embargo, dado que la integración de los problemas dinámicos se realiza a lo largo de escalas de tiempo enormes (del orden de millones de años, en el mejor de los casos), métodos clásicos como los de Runge-Kutta sufren problemas en la conservación de la energía. De hecho se produce un incremento lineal de ésta.

Para solucionar esto, se han desarrollado métodos específicos para integrar este tipo de ecuaciones, los llamados métodos simplécticos. Se basan en dividir el hamiltoniano en dos partes $\mathcal{H} = \mathcal{H}_1 + \mathcal{H}_2$, de forma que cada uno de ellos sea integrable. Con un paso lo suficientemente pequeño (τ), se hacen avanzar consecutivamente cada uno de ellos en 2n-1 subetapas, con n el orden del algoritmo. Por ejemplo, para orden 1 se avanza en cada etapa $\mathcal{H}_1$ y después $\mathcal{H}_2$ el paso τ . Para orden 2, por ejemplo se avanzaría $\mathcal{H}_2$ en $\frac{\tau}{2}$, luego $\mathcal{H}_1$ en τ y de nuevo $\mathcal{H}_2$ en $\frac{\tau}{2}$.

Hay en la actualidad un buen número de algoritmos de este tipo, como el MVS [10], o el HJS [5].

Bibliografía

- [1] Abad, A. (1984). *Estudio de sistemas estelares múltiples*, Tesis doctoral, Publicaciones del Seminario Matemático García de Galdeano, Universidad de Zaragoza.
- [2] Abad, A., Barrio, R., Blesa, F. y Rodríguez, M. (2011). *TIDES tutorial: Integrating ODEs by using the Taylor Series Method.*, Monografías de la Real Academia de Ciencias Exactas, Físicas, Químicas y Naturales de Zaragoza.
- [3] Abad, A. y Ribera, J. (1984). *Método biparamétrico de perturbaciones del tipo Hori*, Publicaciones del Seminario Matemático García de Galdeano, Universidad de Zaragoza.
- [4] Andrade, M. (2008). *N-Parametric Canonical Perturbation Method Based on Lie Transforms*, The Astronomical Journal, **136** (3), pp. 1030–1038.
- [5] Beust, H. (2003). *Symplectic integration of hierarchical stellar systems*, Astronomy and Astrophysics, **400**, pp. 1129–1144.
- [6] Campo, P. P. y Docobo, J. A. (2014). *Analytical study of a four-body configuration in exoplanet scenarios*, Astronomy Letters, **14** (11), pp. 737–748.
- [7] Holman, M. J. y Wiegert, P. A. (1999). *Long-Term Stability of Planets in Binary Systems*, The Astronomical Journal, **117** (1), pp. 621–628.
- [8] Hori, G. (1966). *Theory of General Perturbation with Unspecified Canonical Variable*, Publications of the Astronomical Society of Japan, **18**, pp. 287–296.
- [9] Whipple, A. L. y Szebehely, V. (1984). *The restricted problem of $N + \nu$ bodies*, Celestial Mechanics, **32**, pp. 137–144.
- [10] Wisdom, J. y Holman, M. (1991). *Symplectic maps for the n -body problem*, Astronomical Journal, **102** (4), pp. 1528–1538.

Dando una vuelta a la estadística: los datos circulares

Jose Ameijeiras Alonso

Departamento de Estadística e Investigación Operativa

13 de mayo de 2015

En diversos campos científicos, las medidas son direcciones: la dirección de vuelo migratorio de una especie de pájaro, la dirección en la que sopla el viento o la dirección en la que se mueven las placas tectónicas. En otras ocasiones es de interés analizar “medidas de reloj”: cuándo llegan los pacientes a una unidad hospitalaria (diariamente), cuándo se escucha una canción (semanalmente) o cuándo se producen incendios en una región (anualmente). Este tipo de datos, que se pueden representar en un círculo, se conocen como datos circulares.

Esta estructura periódica hace que estos datos presenten características especiales, que diferenciará su tratamiento estadístico del realizado en datos lineales. Nuestro objetivo será el de analizar las limitaciones de los métodos lineales clásicos, así como presentar algunas ideas básicas de esta área de la estadística.

Introducción

Tanto en las medidas de “compás” como en las de “reloj” surgen naturalmente los datos circulares. Un ejemplo, donde aparecen este tipo de datos, surge cuando queremos estudiar el patrón de la gripe a través de las búsquedas que se hacen en Google de esta palabra en Galicia, se muestra dicho ejemplo en la Figura 1. En este ejemplo, se puede ver la necesidad de diferenciarlos de los datos lineales, ya que un dato como puede ser el día 30 de diciembre está “más cerca” de un dato del inicio del soporte, como puede ser el día 2 de enero, que de otro dato que también esté al final del soporte, como el 20 de diciembre.

En general este tipo de datos surgen cada vez que queremos estudiar el patrón de una variable cíclica ya que, en este caso, emplear medias y varianzas muestrales como uno haría en el análisis de series temporales o en la regresión estándar podría dar lugar a conclusiones erróneas. Para elaborar la presente introducción a los datos circulares, se han seguido dos referencias clásicas de este tipo de datos como son [1] y [2].

Primeras nociones

Este tipo de datos se pueden representar como ángulos medidos con respecto a alguno convenientemente elegido, es decir, una vez elegido un punto de partida

PALABRAS CLAVE: Datos circulares.

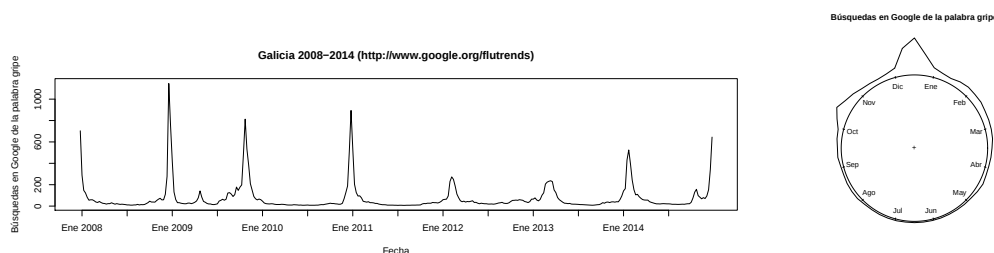


Figura 1: Búsquedas en Google de la palabra gripe en Galicia entre el 1 de enero de 2008 y el 31 de diciembre de 2014. Izquierda: representación lineal. Derecha: representación circular de las búsquedas acumuladas.

y un sentido de rotación (por ejemplo, tomando como dirección positiva el sentido de las agujas del reloj o el sentido antihorario). Esto hace que su representación numérica no sea necesariamente única. Por tanto, es importante asegurarse de que las conclusiones extraídas de los datos en términos descriptivos o inferenciales están en función de las observaciones dadas y no dependen de los valores arbitrarios en los que nos referimos a ellos.

Herramientas exploratorias

La forma más sencilla de representar los datos circulares es mostrar las observaciones como puntos en el círculo unidad. Una forma de agrupar los datos, para analizar de forma visual dónde aparecen con más frecuencia, es el diagrama de rosa. Esta representación se realiza de forma análoga a la del histograma, simplemente cambiando las barras por sectores circulares. El área de cada sector debería ser proporcional a la frecuencia de cada grupo. Una forma de realizar estos diagramas cuando los grupos son de la misma amplitud, es tomar el radio igual a la raíz cuadrada de la frecuencia relativa. Con el fin de visualizar esta herramienta, se muestra en la Figura 2 (izquierda) el diagrama de rosa para una muestra de observaciones de direcciones.

Medidas de localización y escala

Una vez que se ha visto cómo representar los datos, es útil poder resumirlos a través de un análisis descriptivo apropiado. En la medición de la localización aparece la primera muestra de la necesidad de un tratamiento estadístico especial para este tipo de datos. Con el fin de entender esta necesidad, se va a suponer que se tienen las direcciones de vuelo de dos pájaros, que se muestran en la Figura 2 (centro) y son $\pi/6$ y $(2\pi - \pi/6)$ con respecto al Norte. Si se tomase la media muestral, se llegaría a la conclusión de que la dirección media de vuelo es el Sur (π), lo cual contradice la idea intuitiva de dirección media de vuelo de estas dos aves, que debería ser el

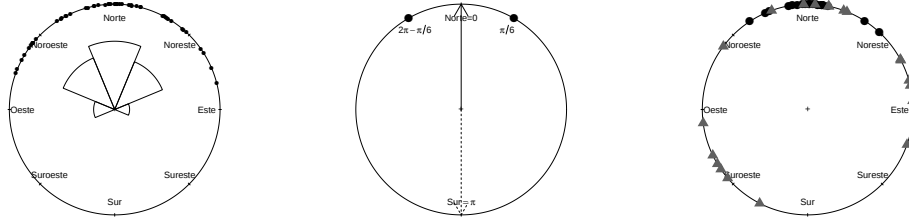


Figura 2: Izquierda: Diagrama de rosa para una muestra de observaciones de direcciones. Centro: Media muestral escalar (flecha continua) y circular (flecha discontinua) de dos direcciones. Derecha: Observaciones con una concentración baja (triángulos grises) y una alta (puntos negros)

Norte. En este ejemplo, se ve la necesidad de redefinir la media muestral para los datos circulares.

Para poder definir la media muestral direccional, supongamos que $X_i = (\cos \theta_i, \sin \theta_i)$, con $i = 1, \dots, n$ son las observaciones de la muestra. Entonces el centro de masas de la muestra será $(\bar{C}, \bar{S}) = \left(\frac{1}{n} \sum_{i=1}^n \cos \theta_i, \frac{1}{n} \sum_{i=1}^n \sin \theta_i \right)$ y por tanto la dirección donde se alcanza este centro vendrá dada por el punto $\bar{\theta} = \arg(\bar{C} + i\bar{S})$. Además, la dirección media así definida es equivariante por rotación, lo cual garantiza que esta media está bien definida ya que no depende del punto que se tome como 0.

Una forma de medir la escala de la muestra es a través de la longitud del centro de masas, que viene dada por $\bar{R} = (\bar{C}^2 + \bar{S}^2)^{1/2} \in [0, 1]$, ya que si los datos están muy concentrados alrededor de la media muestral circular, como ocurre en los puntos negros de la Figura 2 (derecha), entonces $\bar{R}$ estaría próxima a 1, mientras que si están muy dispersos, como ocurre en los triángulos grises de la Figura 2 (derecha), esta $\bar{R}$ sería casi 0. Así, se puede ver que, de forma natural, surge la medida de concentración como forma de medir la escala y juega el papel inverso al parámetro escala que aparece en el caso lineal, que es la varianza, la cual mide la dispersión de los datos.

Distribuciones circulares

Una distribución circular es una distribución de probabilidad cuya probabilidad total está concentrada en la circunferencia de un círculo unidad. Como cada punto en la circunferencia representa una dirección, tal distribución es una forma de asignar probabilidades a diferentes direcciones. El soporte habitual de una variable aleatoria circular Θ , medida en radianes, suele ser $[0, 2\pi)$. Las distribuciones circulares son esencialmente de dos tipos: pueden ser discretas, asignando masa de probabilidad

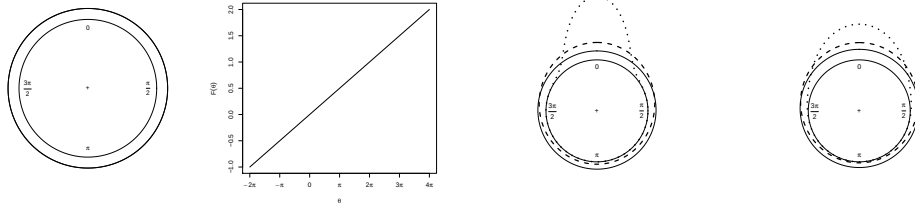


Figura 3: De izquierda a derecha: función de densidad y función de distribución de la distribución uniforme, función de densidad de la von Mises con $\mu = 0$ y $\kappa = 0,1$ (trazo continuo), $\kappa = 1$ (trazo de guiones) y $\kappa = 10$ (línea punteada) y función de densidad de la normal enrollada con $\mu = 0$ y $\kappa = 0,15$ (trazo continuo), $\kappa = 0,5$ (trazo de guiones) y $\kappa = 0,85$ (línea punteada).

solamente en un número de direcciones numerable, o pueden ser absolutamente continuas (con respecto a la medida de Lebesgue en el círculo). En este segundo caso, la función de densidad $f(\theta)$ existe y además de las propiedades que tiene esta función en el caso lineal: $f(\theta) \geq 0, \forall \theta \in [0, 2\pi)$ y $\int_0^{2\pi} f(\theta)d\theta = 1$, también se verifica que $f(\theta) = f(\theta + 2l\pi), \forall \theta \in [0, 2\pi)$ y $\forall l \in \mathbb{Z}$. Un primer ejemplo de una función de densidad circular sería la que se muestra en la Figura 3 (izquierda) que es la distribución uniforme, la cual asigna a todas las direcciones la misma probabilidad, esto es, su función de densidad es $f(\theta) = 1/(2\pi), \theta \in [0, 2\pi)$. A continuación, en este documento, se tratarán algunas de las distribuciones unimodales y simétricas más importantes. Para obtener otras distribuciones véase [1], por ejemplo.

El definir de esta forma la función de densidad hace que se tenga que redefinir la función de distribución F , donde además de verificarse que $F(\phi) = \mathbb{P}(0 \leq \theta \leq \phi) = \int_0^\phi f(\theta)d\theta$ si $\phi \in [0, 2\pi)$, también se cumple que $F(\phi + 2\pi) - F(\phi) = 1, \forall \phi \in \mathbb{R}$, haciendo que esta tenga soporte en toda la recta real, como se puede ver en la Figura 3 (centro-izquierda).

Distribución von Mises

Una de las distribuciones circulares más empleadas es la von Mises. La obtención de esta distribución se basa en la idea de obtener una distribución para la cual el estimador de máxima verosimilitud del parámetro de localización coincida con la media muestral circular, al igual que pasa en el caso lineal con la distribución normal y la media muestral. La función de densidad que cumple esa propiedad es

$$f(\theta; \mu, \kappa) = \frac{\exp(\kappa \cos(\theta - \mu))}{2\pi I_0(\kappa)}, \text{ con } \theta \in [0, 2\pi),$$

donde $\mu \in [0, 2\pi)$ es la dirección media y $\kappa \geq 0$ el parámetro de concentración. La función $I_p(\kappa) = \frac{1}{2\pi} \int_0^{2\pi} \cos(p\theta) \exp(\kappa \cos \theta) d\theta$ denota la función de Bessel modificada

de orden p . Se muestra en la Figura 3 (centro-derecha) la representación de esta densidad para $\mu = 0$ y distintos valores de κ (0,1; 1 y 10).

Como se comentó anteriormente el estimador de máxima verosimilitud de la dirección media es la media muestral circular, mientras que la estimación de κ por máxima verosimilitud se consigue despejando $\bar{\kappa}$ en la siguiente expresión $I_1(\bar{\kappa})/I_0(\bar{\kappa}) = \frac{1}{n} \sum_{i=1}^n \cos(\theta_i - \bar{\theta}) = \bar{R}$. Se obtiene el mismo estimador de ambos parámetros cuando se calculan por el método de los momentos.

Distribución normal enrollada

Una forma de generar distribuciones circulares a partir de distribuciones lineales, es “enrollando” una distribución lineal alrededor de un círculo de radio unidad. Si X es una variable aleatoria en la recta real con función de densidad g , esta puede transformarse en una variable aleatoria circular realizando $\Theta = X(\bmod 2\pi)$ y su función de densidad asociada sería $f(\theta) = \sum_{l=-\infty}^{\infty} g(\theta + 2l\pi)$, $\theta \in [0, 2\pi)$.

Con este método se obtiene otra de las distribuciones más empleadas en el contexto circular, que es la normal enrollada, la cual es generada a partir de la distribución $N(\mu, \sigma^2)$. Su función de densidad puede escribirse como

$$f(\theta; \mu, \kappa) = \frac{1}{2\pi} \left(1 + 2 \sum_{p=1}^{\infty} \kappa^{p^2} \cos p(\theta - \mu) \right),$$

donde $\mu \in [0, 2\pi)$ es la dirección media y $\kappa = \exp(-\sigma^2/2) \in [0, 1]$ es el parámetro de concentración. Se muestra en la Figura 3 (derecha) la representación de esta densidad para $\mu = 0$ y distintos valores de κ (0,15; 0,5 y 0,85).

Test

Una vez vistas las principales distribuciones, en diversas ocasiones es de interés contrastar si una muestra dada sigue una distribución como puede ser la uniforme o la von Mises o si dos muestras poseen la misma distribución. En este documento se pretende dar una pequeña pincelada de algunos de los test presentes en la literatura. Aunque aquí solo se presentan tres test, hay que tener en cuenta que existen diversas propuestas, tanto para realizar los contrastes mencionados anteriormente como para realizar otros contraste clásicos como pudieran ser, por ejemplo, los test para la dirección media (véase [1] o [2]).

Cuando el objetivo es contrastar la hipótesis nula $H_0 : F = F_0$, con F_0 la distribución circular uniforme, lo lógico es pensar en el test clásico de Kolmogorov–Smirnov, donde dada una muestra ordenada $\theta_{(1)} \leq \dots \leq \theta_{(n)}$, el estadístico que se emplea es el $\max(D_n^+, D_n^-)$, con $D_n^+ = \max_{1 \leq i \leq n} (i/n - U_i)$ y $D_n^- = \max_{1 \leq i \leq n} (U_i - (i-1)/n)$, siendo $U_i = F(\theta_i) = \theta_{(i)}/2\pi$. El problema de este estadístico, en el caso circular,

α	0.10	0.05	0.01	α	0.10	0.05	0.01	α	0.10	0.05	0.01
V_n^*	1.620	1.747	2.001	U_n^2	0.152	0.187	0.268	U_{n_1, n_2}^2	0.152	0.187	0.268

Tabla 2: De izquierda a derecha: Cuantiles del estadístico de Kuiper para contrastar uniformidad, del estadístico de Watson para contrastar si la muestra sigue una von Mises de parámetros conocidos y para contrastar si dos muestras siguen la misma distribución.

es que depende del punto que se escoja como 0. Para solucionar este inconveniente Kuiper define como estadístico para el contraste de uniformidad $V_n = D_n^+ + D_n^-$. A partir de este estadístico, tomando la corrección $V_n^* = \sqrt{n}V_n \left(1 + \frac{0.155}{\sqrt{n}} + \frac{0.24}{n}\right)$, para determinar cuando se debe rechazar H_0 , su distribución puede ser aproximada a través de los cuantiles que se dan en la Tabla 2 cuando $n \geq 8$.

En general, si se quiere realizar el contraste $H_0 : F = F_0$, con F_0 una distribución circular continua se puede emplear el test U^2 de Watson, cuyo estadístico asociado es

$$U_n^2 = n \int_0^{2\pi} \left[F_n(\theta) - F_0(\theta) - \int_0^{2\pi} (F_n(\vartheta) - F_0(\vartheta)) dF_0(\vartheta) \right]^2 dF_0(\theta),$$

donde $F_n(\theta) = (1/n) \sum_{i=1}^n \mathcal{I}(\theta_i \leq x)$ es la función de distribución empírica e $\mathcal{I}$ la función indicatriz.

En el caso de querer contrastar si la muestra sigue una distribución von Mises de parámetros (μ, κ) , los cuantiles de U_n^2 se pueden aproximar como se muestra en la Tabla 2 cuando $n \geq 20$.

Finalmente, el test de Watson también se puede emplear para realizar el contraste $H_0 : F_1 = F_2$, con F_1, F_2 distribuciones circulares y continuas. El estadístico empleado en este contraste es

$$U_{n_1, n_2}^2 = \frac{n_1 n_2}{n} \int_0^{2\pi} \left[F_{n_1}(\theta) - F_{n_2}(\theta) - \int_0^{2\pi} (F_{n_1}(\vartheta) - F_{n_2}(\vartheta)) dF_n(\vartheta) \right]^2 dF_n(\theta),$$

siendo $F_n(\theta) = (n_1 F_{n_1} + n_2 F_{n_2})/n$, n_1 el tamaño muestral de una muestra, n_2 el de la otra y $n = n_1 + n_2$. De nuevo, la distribución de este estadístico puede ser aproximada cuando $n \geq 17$ y sus cuantiles se muestran en las Tabla 2.

Bibliografía

- [1] Jammalamadaka, S. R. y SenGupta, A. (2001). *Topics in Circular Statistics*, World Scientific Publishing. Great Britain.
- [2] Mardia, K. V. y Jupp, P. E. (2000). *Directional Statistics*, John Wiley & Sons. Great Britain.

Parte II

10º Aniversario

Simetría e forma

J. Carlos Díaz Ramos

Departamento de Xeometría e Topoloxía

17 de decembro de 2014

Intuitivamente, a simetría é a distribución regular das partes dunha cousa ou de obxectos semellantes a unha e á outra banda dun eixe, arredor dun centro, etc. Son exemplos disto a simetría bilateral dun humano, a simetría rotacional dunha peza de barro obtida mediante xiro nun torno, ou as simetrías dos poliedros regulares.

En xeral, queda claro que a anterior definición non é moi precisa, e que é necesario especificar que tipo de simetría ten un obxecto. Para aclarar mellor esta definición a teoría de grupos aparece de xeito natural, e dirase, de novo de xeito informal, que un obxecto ten unha G -simetría, onde G é un grupo, se as transformacións de G deixan o obxecto en cuestión invariante.

Por exemplo, un polígono regular de n lados dise que ten unha D_n -simetría, onde a D vén de diédrico, que é o nome co que se coñecen estes grupos. O grupo diédrico inclúe as rotacións de ángulos que sexan múltiplos de $2\pi/n$, e tamén outras como as reflexións con respecto ós eixos que dividen ó polígono en dúas partes iguais.

Para ir introducindo algunha terminoloxía, un obxecto que ten unha simetría rotacional dise que ten unha $O(2)$ -simetría. Sería unha $O(3)$ -simetría as que producen os movementos ríxidos do espazo que deixan fixo a orixe de coordenadas. En xeral, o grupo $O(n)$ (grupo ortogonal) é o grupo dos movementos ríxidos que deixan fixo un punto do espazo n -dimensional.

O emprego da teoría de grupos é ubicuo hoxe en día en Matemáticas, por exemplo porque é unha estrutura subxacente a moitos dos obxectos de estudo, tales como espazos vectoriais ou estruturas alxébricas ou xeométricas máis complexas. Felix Klein chegou incluso a dicir que a xeometría non é outra cousa máis có estudo das propiedades dun espazo que permanecen invariantes pola acción do grupo de transformacións dese espazo.

A simetría tamén é moi importante noutras ciencias como a Física. Emmy Noether probou un teorema fundamental a este respecto: as propiedades de simetría dun sistema físico correspóndense con leis de conservación. Este resultado é importante no só polo interese físico que ten, senón tamén por unha cuestión máis profunda: pódense atopar simetrías non soamente en figuras e obxectos xeométricos, senón tamén en elementos máis abstractos como as ecuacións diferenciais, que en definitiva son as que rixen as leis da mecánica newtoniana. De feito, o premio Nóbel P. W. Anderson expresou no seu día o feito de que só é un pouco esaxerado dicir que a Física é o estudo da simetría.

PALABRAS CLAVE: Simetría; superfice; curvaturas principais constantes.

O concepto de forma ten un problema similar ó da simetría. Todo o mundo ten unha idea intuitiva do que é pero na práctica é moi difícil de definir. Un diccionario di algo así como que a forma é a realización concreta con que se presenta un fenómeno, con características propias que os diferencian doutros. Esta definición poderémola concretar máis adiante no caso da xeometría diferencial de subvariedades. En todo caso, hai algo que é claro: a simetría influencia de xeito decisivo a forma, facendo que a última sexa moito máis regular e autosemellante, ou en sentido vago, “constante”. O problema que abordamos neste traballo, en sentido moi xeral, é o seguinte:

¿Ata que punto a forma constante dun obxecto implica que ese obxecto posúe unha simetría?

Trataremos de afondar nesta cuestión para un exemplo relativamente sinxelo que estará ó alcance de calquera graduado en Matemáticas, e que ten como concepto preliminar as superficies regulares de $\mathbb{R}^3$.

Superficies homoxéneas

Recordemos brevemente que unha superficie M de $\mathbb{R}^3$ é un subconxunto que localmente se pode describir mediante unha parametrización $\mathbf{x}: U \subset \mathbb{R}^2 \rightarrow M \subset \mathbb{R}^3$, onde U é un aberto, e tal que a diferencial $D\mathbf{x}$ ten rango 2. Esíxese ademais que a topoloxía inducida por estas parametrizacións coincida coa topoloxía de M como subespacio de $\mathbb{R}^3$.

As superficies herdán por tanto a topoloxía relativa do espacio euclidiano e por tanto o seu carácter métrico. Dise que unha superficie é (extrinsecamente) *homoxénea* se para calquera par de puntos da superficie existe unha isometría de $\mathbb{R}^3$ que deixa invariante a superficie e que leva un punto no outro. Este será o noso concepto de “superficie simétrica” que trataremos de caracterizar mediante un concepto de “forma”.

Para cada punto $p \in M$ podemos construír o espacio tanxente T_pM a M en p . Sexa ξ_p un vector normal unitario en p . Suporemos que ξ_p se pode estender a un campo unitario normal ξ polo menos nunha veciñanza de p . Como cada espacio tanxente é un subconxunto de $\mathbb{R}^3$ podemos dotar este do produto interior restrinxido a T_pM , que denotaremos por $\langle \cdot, \cdot \rangle$. Este produto interior tamén se coñece como *primeira forma fundamental*.

Se X é un campo de vectores tanxente a M defínese o *operador forma* (ou segunda forma fundamental) S de M como $SX = -D_X\xi$, onde D é a derivada usual de $\mathbb{R}^3$. Nótese que SX é tanxente pois ξ é unitario. Máis intuitivamente, nótese que este operador é o que precisamente dá a forma da superficie M xa que mide como varía o espacio tanxente (ou equivalentemente o normal) de punto a punto. Isto coincide coa idea intuitiva de “ve-la forma”, xa que o xeito en que o noso cerebro a interpreta é precisamente ó percibi-la diferente tonalidade de cor que depende do ángulo co que a luz se reflicte (no plano tanxente) en cada punto.

O operador S é autoadxunto con respecto á primeira forma fundamental, e por tanto, polo teorema espectral é diagonalizable con autovalores reais e autoespacios

ortogonais. Estes autovalores son os que se coñecen como *curvaturas principais*. A súa suma é a *curvatura media* $H = \text{tr } S$, e o seu produto a *curvatura de Gauss* $K = \det S$. Tamén denotaremos por ∇ a *derivada covariante* definida como $\nabla_X Y = (D_X Y)^\top$, a proxección tanxente á superficie da derivada usual de $\mathbb{R}^3$.

Trataremos de facer fincapé no que segue nun feito fundamental na xeometría: a capacidade de expresa-los diferentes conceptos dun xeito independente da parametrización. Por iso, as expresións que presentaremos a continuación estarán escritas sen facer referencia algunha a unha elección de coordenadas.

Nesta exposición recordaremos (aínda que non demostraremos) as ecuacións fundamentais dunha superficie. Empezamos pola *fórmula de Gauss*

$$D_X Y = \nabla_X Y + \langle SX, Y \rangle \xi,$$

e o *Teorema Egregium de Gauss*

$$K = \frac{R(X, Y, Y, X)}{\langle X, X \rangle \langle Y, Y \rangle - \langle X, Y \rangle^2},$$

que establece que a curvatura de Gauss se pode escribir en función do tensor de Riemann (ou símbolos de Riemann), que depende só da primeira forma fundamental. O tensor de Riemann é xustamente o lado esquerdo da fórmula que aparece a continuación. Este é un resultado instrumental da xeometría diferencial, xa que implica que o concepto de curvatura se pode definir independentemente do feito de que a superficie sexa un subconxunto do espacio euclidiano. Ademais deste importante resultado, o Teorema Egregium é equivalente á chamada *ecuación de Gauss* que se pode escribir como segue:

$$\langle \nabla_X \nabla_Y X - \nabla_Y \nabla_X X - \nabla_{\nabla_X Y - \nabla_Y X} X, Y \rangle = \langle SX, X \rangle \langle SY, Y \rangle - \langle SX, Y \rangle^2.$$

Tamén empregarémola *ecuación de Codazzi*, que expresa o feito de que a derivada de S é simétrica, e que tamén se pode escribir de xeito equivalente como

$$\langle \nabla_X (SY) - S(\nabla_X Y) - \nabla_Y (SX) + S(\nabla_Y X), Z \rangle = 0.$$

As ecuacións de Gauss e Codazzi son instrumentais en xeometría de superficies xa que expresan as condicións de compatibilidade do teorema fundamental das superficies, que basicamente di que, dadas unha primeira e segunda forma fundamental de xeito que se satisfagan as ecuacións de Gauss e Codazzi, entón existe unha superficie que ten como primeira e segunda forma fundamental precisamente as prescritas; ademais, tal superficie é única salvo movemento ríxido do espacio.

No resto do artigo probaremos o resultado fundamental deste traballo, que foi obxecto de numerosas xeneralizacións e liñas de investigación en xeometría diferencial. Antes de enunciálo fagamos notar que baixo a hipótese de que M é homoxénea, é bastante claro que a xeometría de M como subconxunto de $\mathbb{R}^3$ é a mesma en cada punto. En particular, séguese que os operadores forma en dous puntos distintos son conxugados e por tanto teñen os mesmos autovalores. A constancia dos autovalores vén a querer dicir que a forma desta superficie é “constante”. Vemos a continuación que o recíproco a este resultado tamén é certo:

Teorema 1. *Unha superficie conexas con curvaturas principais constantes en $\mathbb{R}^3$ é un aberto dunha superficie homoxénea. Ademais, ten que ser un aberto dun plano, dunha esfera ou dun cilindro.*

Dedúcese por tanto do enunciado deste teorema a clasificación das superficies homoxéneas de $\mathbb{R}^3$, xa que claramente os planos, esferas e cilindros son homoxéneos e o teorema di que non pode haber máis.

Demostración. Para fixa-la notación supoñamos que M ten dúas curvaturas principais constantes λ e μ e tomemos X e Y dous campos de vectores tales que $SX = \lambda X$, $SY = \mu Y$ e $\langle X, X \rangle = \langle Y, Y \rangle = 1$, $\langle X, Y \rangle = 0$.

A demostración baséase na análise de tres casos: $\lambda = \mu = 0$, $\lambda = \mu \neq 0$ e $\lambda \neq \mu$. Tratamos cada caso a continuación.

Supoñamos primeiro que $\lambda = \mu = 0$. Isto vén a dicir, xa que S é diagonalizable, que $S = 0$, ou equivalentemente, que ξ , o vector normal, é constante. Por tanto, M é un plano.

Supoñamos agora que $\lambda = \mu \neq 0$, e denotemos $r = 1/\lambda$. Así o operador forma é un múltiplo da identidade, $S = \frac{1}{r}I$. Definímo-la seguinte aplicación

$$\Phi_r: M \rightarrow \mathbb{R}^3, \quad p \mapsto p + r\xi_p.$$

O significado xeométrico desta aplicación é que Φ_r deslaza a superficie orixinal M unha distancia r na dirección do vector normal. Empregando a definición do operador forma resulta que $D\Phi_r = I - rS = 0$. En consecuencia Φ_r é constante, o cal significa que tódolos puntos de M , despois de ser desplazados unha distancia r acabaron no mesmo punto c . Por tanto, M debe ser un aberto dunha esfera centrada en c e de radio r .

Finalmente supoñamos $\lambda \neq \mu$. Escribindo a ecuación de Codazzi, poñendo $Z = X$ e $Z = Y$, resulta $(\mu - \lambda)\langle \nabla_X Y, X \rangle = 0$ e $(\mu - \lambda)\langle \nabla_Y X, Y \rangle = 0$, co cal $\langle \nabla_X Y, X \rangle = \langle \nabla_Y X, Y \rangle = 0$. Nótese que $\langle \nabla_X X, Y \rangle = -\langle \nabla_X Y, X \rangle$ pois $\langle X, Y \rangle = 0$ e simplemente aplicamos a regra de Leibnitz para a derivada do produto escalar. Análogamente, $\langle \nabla_Y Y, X \rangle = 0$ pois $\langle X, Y \rangle = 0$. Como $\nabla_X X$ é tanxente a M e $\{X, Y\}$ é unha base do tanxente, concluímos $\nabla_X X = 0$. Procedendo de xeito similar obtemos

$$\nabla_X X = \nabla_X Y = \nabla_Y X = \nabla_Y Y = 0.$$

Agora empregamos estas ecuacións para substituír na ecuación de Gauss e obtemos facilmente $\lambda\mu = 0$. Así que de aquí en diante podemos supoñer $\lambda \neq 0$ e $\mu = 0$.

A fórmula de Gauss di que $D_Y Y = \nabla_Y Y + \langle SY, Y \rangle \xi = 0$, co que as curvas integrais de Y son rectas de $\mathbb{R}^3$. (Recordemos que unha curva integral dun campo de vectores é unha curva que ten por vector tanxente o valor do campo do que é curva integral nese punto; tal curva existe pois é solución dunha ecuación diferencial ordinaria.)

Por outra banda, defimos coma antes $r = 1/\lambda$. Considerando de novo a aplicación Φ_r definida anteriormente, neste caso vemos que $D\Phi_r(X) = X - r\lambda X = 0$, o cal significa que Φ_r é constante ó longo das curvas integrais de X .

Combinando este resultado co anterior, chegamos a que para un punto dunha recta determinada por Y , o desprazamento na dirección normal a M dunha curva integral de X é un punto. En definitiva, o que probamos é que M ten que ser un aberto dun cilindro de radio r ó longo dunha recta paralela á que determina o campo de vectores Y . $\square$

Votando secuencialmente

Julio González Díaz

Departamento de Estadística e Investigación Operativa

17 de diciembre de 2014

Introducción

En esta charla presentamos una aplicación de la teoría de juegos al estudio de modelos de votación. En particular, queremos estudiar un modelo de votación en el cual los resultados intermedios de la elección se van haciendo públicos en tiempo real, a medida que la gente vota. Cada votante puede elegir estratégicamente si prefiere votar pronto o esperar a que la elección avance.

Un problema habitual de muchos sistemas electorales es que los votantes no pueden coordinarse de manera efectiva. Un caso típico es el fenómeno del “voto inútil”, en el que un candidato no llega al porcentaje mínimo de votos y todos los votos que ha recibido pasan a ser inútiles. De saber que esta candidata no llegaría al mínimo, probablemente muchos votantes habrían decidido votar a otro candidato. Además, cuando el ganador final de la elección se elige por mayoría simple, también pueden aparecer problemas de coordinación si hay, por ejemplo, un partido grande de derechas y varios partidos no tan grandes de izquierdas (es un fenómeno habitual que la izquierda esté más disgregada que la derecha).

En este tipo de situaciones puede ser que un *perdedor de Condorcet* acabe siendo elegido, donde un *perdedor de Condorcet* es un candidato que perdería en un mano a mano con cualquiera de los otros candidatos. Pensemos por ejemplo en una situación como la descrita en la Tabla 3. Con la distribución de votos ahí representada, es claro que la mayoría del electorado quiere un gobierno de izquierdas. Sin embargo, a menos que consigan coordinarse, será el candidato C el que resulte ganador. En caso de un mano a mano entre A y C , los votantes de B apoyarían a A , con lo que C perdería. De modo análogo, C también perdería en un mano a mano contra B .

Situaciones similares a la que acabamos de describir ya han tenido lugar en el pasado, con las elecciones generales de los Estados Unidos en el año 2000 como ejemplo más destacado. En ese año el resultado estuvo muy reñido entre Al Gore (demócrata) y George Bush (republicano). Además, había un tercer candidato, el independiente Ralph Nader, que obtuvo alrededor del 1% de los votos. Dado lo igualado del recuento final y que Al Gore terminó perdiendo, lo más probable es que una gran mayoría de los votantes de Ralph Nader hubieran preferido votar por

Candidatos	A	B	C
Ideología	izq.	izq.	der.
Votantes	30 %	30 %	40 %

Tabla 3: Una situación en la que un perdedor de Condorcet puede ganar.

Al Gore (cuya ideología era más cercana a la de Ralph Nader de lo que lo era la de George Bush).

Teniendo en cuenta que la victoria de un perdedor de Condorcet puede verse como algo socialmente ineficiente e indeseable, ¿cómo podrían resolverse este tipo de problemas de coordinación? Esta pregunta ha sido ampliamente estudiada en el campo de la economía política, y en esta charla presentamos una nueva idea que puede servir para mejorar la coordinación de los votantes, poniéndole las cosas más difíciles a este tipo de alternativas socialmente no deseables.

Antes de presentar el modelo es interesante mencionar que una solución natural para la situación anterior hubiera sido llevar a cabo una segunda vuelta entre Al Gore y George Busch. Sin embargo, una segunda vuelta tampoco es siempre garantía de resultados socialmente eficientes. De hecho, el sistema de segunda vuelta estaba presente en Francia en 2002, cuando tuvieron unas elecciones generales muy controvertidas. En ellas se partía de una situación social con una gran notable mayoría de votantes que iban a votar a opciones de izquierdas, pero al mismo tiempo había una gran cantidad de dichas opciones. El resultado después de la primera vuelta fue el siguiente: Jean Marie le Pen (extrema derecha) 20 %, Jacques Chirac (derecha) 17 %, Lionel Jospin (izquierda) 16 %, y otros partidos de izquierdas sumaron más del 30 %. Como consecuencia, a pesar de que Lionel Jospin era el claro favorito en las encuestas, en la segunda vuelta los franceses tuvieron que elegir entre dos candidatos de derechas. En general, tener un sistema con $(n - 1)$ vueltas cuando hay n candidatos podría ser una buena forma de minimizar este tipo de resultados, pero uno no puede esperar que los votantes acudan a votar tantas veces. A medida que aumente el número de vueltas de la elección bajará la participación en la misma.

El modelo

Aunque idealmente nos gustaría estudiar modelos con una cantidad arbitraria de candidatos y una cantidad arbitraria de periodos, la complejidad de los modelos de votación secuencial obliga a asumir importantes simplificaciones.

Consideraremos una elección entre tres candidatos (o alternativas), A , B y C , y con $N \geq 4$ votantes. El sistema de votación elegirá exactamente a un candidato por mayoría simple: cada votante emite un voto y el candidato con más votos es elegido.

En caso de empate, el ganador es elegido mediante un sorteo en el que todos los candidatos empatados tienen igual probabilidad de salir elegidos.

La principal desviación con respecto a los modelos de votación habituales es que la elección es secuencial, en dos periodos. Cada votante puede elegir estratégicamente si votar en el periodo 1 o en el periodo 2. En este último caso, el votante conocería el *recuento* de votos recibidos por cada candidato en el periodo 1. Nótese que este sistema de votación respeta el principio fundamental de “una persona un voto”.

Decimos que cada votante es de un cierto *tipo*, que viene caracterizado por la utilidad que le reporta que cada uno de los candidatos salga elegido. Por comodidad asumimos que las utilidades han sido normalizadas de tal manera que el votante i asigna utilidad 1 a su candidato preferido, utilidad 0 al candidato que menos le gusta y utilidad $v_i \in [0, 1]$ al candidato intermedio. Por tanto, si denotamos por Π al conjunto de todas las ordenaciones posibles de $\{A, B, C\}$, el tipo de un votante i se compone de dos elementos: i) una ordenación $\pi \in \Pi$ de los candidatos y ii) la utilidad v_i asociada a su candidato intermedio. Normalmente nos referiremos a los candidatos con un valor bajo de v_i como *partidistas* y a los votantes con valores altos de v_i como *antagonistas*, ya que son totalmente contrarios a un candidato pero le gustan los otros dos. Durante la exposición diremos que un votante es un votante AB , para indicar que A es su candidato preferido y B su candidato intermedio.

Inicialmente podemos pensar que, antes de que la elección comience, los tipos de los votantes se generan de modo *i.i.d.* (independientes e idénticamente distribuidos) de una cierta distribución de probabilidad. Por tanto, saber el tipo de un grupo de votantes no da ninguna información acerca del tipo de los restantes.

Definición 1. *Dos candidatos son simétricos (ex ante) si la distribución de probabilidad de la que se obtienen los tipos de los votantes trata de modo idéntico a estos candidatos.*

Estrategias

Dado un votante i , una estrategia σ_i específica, para cada posible tipo, cuál sería su estrategia si finalmente ése es su tipo. Siendo más precisos, la estrategia ha de especificar la probabilidad de que i vote en el periodo 1 y en el periodo 2, además de la probabilidad con la que elegiría a cada candidato en el periodo 1 y también dichas probabilidades para cada posible resultado del recuento que se pueda encontrar tras el periodo 1. Un perfil de estrategias para todos los votantes se denotará por σ .

Definición 2. *Un perfil de estrategias es simétrico si todos los votantes del mismo tipo siguen la misma estrategia.*

La ventaja de trabajar con estrategias simétricas es que basta especificar el comportamiento de cada tipo de votante. Durante la mayor parte del análisis nos centraremos en el caso en el que las estrategias son simétricas. Eso sí, cuando se estudie si una estrategia es un equilibrio y se consideren desviaciones de jugadores

con respecto a la misma, estas desviaciones sí que podrán romper la simetría (esto implica que los equilibrios que estudiaremos no serán más débiles que los equilibrios no simétricos, ya que se permite el mismo tipo de desviaciones).

Ahora introduciremos una propiedad de anonimato, que requiere que candidatos simétricos sean tratados de modo idéntico. Aunque la idea es bastante natural, en este contexto es algo compleja de formalizar. Es importante recordar aquí que toda la información que un votante recibe durante la elección, además de su propio tipo, es el resultado del recuento después del primer periodo.

Definición 3. *Un perfil de estrategias σ es anónimo si, para cada votante i , cada par de candidatos simétricos, digamos D_1 y D_2 , y cada par de tipos θ y θ' que simplemente difieren en que los roles de D_1 y D_2 se han intercambiado, σ_i cumple lo siguiente:*

- *Si bajo el tipo θ en un determinado momento de la elección y dada una cierta información, el votante i votaría por el candidato D_1 con probabilidad p , entonces, bajo el tipo θ' , en una situación análoga en la que la información relativa a D_1 y D_2 se haya intercambiado, el votante i votaría por el candidato D_2 con probabilidad p .*

La siguiente propiedad captura la idea natural de que los votantes en el periodo 2 se pueden ver atraídos hacia los candidatos más fuertes (aumentando así la probabilidad de que su voto sea un voto útil).

Definición 4. *Un perfil de estrategias se dice que es débilmente monótono si, para cada candidato D , una vez fijado el número de votos que en el periodo 1 han llevado el resto de candidatos, el porcentaje de votos para D al final de la elección es débilmente creciente en el número de votos que D recibe en el periodo 1.*

Un objetivo importante de nuestro estudio es entender hasta qué punto los votantes del segundo periodo se ven influidos por el resultado que observan al final del periodo 1. Las siguientes definiciones capturan dos grados extremos de reacción ante el resultado del periodo 1.

Definición 5. *Un perfil de estrategias es insensible si, para cada candidato D y cada votante i , ninguna desviación de i cambia el número esperado de votos de D al final de la elección más allá del propio voto del votante i .*

Definición 6. *Un perfil de estrategias es totalmente sensible si es débilmente monótono y, además, en el periodo 2 un votante vota por el candidato que va de primero entre aquellos que le dan una utilidad positiva.*

El concepto de totalmente sensible es una forma muy fuerte de monotonía, en la que los votantes reaccionan yendo por el candidato que parece más fuerte después del periodo 1 (siempre y cuando su victoria les reporte una utilidad positiva). Aunque esta forma de monotonía puede no ser natural en general, en nuestro análisis se presentan dos casos particulares en los que sí que es relativamente natural (aunque

no los describiremos en este resumen). También es importante destacar que, en general, las estrategias totalmente sensibles no son compatibles con las condiciones de equilibrio. Esto lo ilustramos en el siguiente ejemplo.

Ejemplo 7. *Consideremos una situación en la que tenemos 100 votantes y, en el periodo 1, 50 han votado por B y 49 han votado por A. Asumamos, además, que el único votante restante, el votante i , tiene a A como su candidato favorito y a B como su candidato intermedio con $v_i \in (0, 1)$. Entonces, bajo una estrategia totalmente sensible el votante i debería votar por B, pero su utilidad esperada sería mayor si votase por A.*

Situaciones como la que acabamos de describir, donde tras el periodo 1 un votante sabe que es el único votante que queda por votar y que su voto decide la elección son muy improbables, pero pueden hacer el análisis de los equilibrios del juego subyacente muy complejos sin añadir comprensión del modelo.

Uno de los aspectos más delicados de hacer el análisis de equilibrios en modelos de votación es que la utilidad de los jugadores con cada estrategia depende crucialmente de las situaciones en las que su voto hace *pivotar* el resultado final hacia uno u otro candidato. Caracterizar analíticamente las situaciones en las que cada votante puede ser *pivote* es muy complejo. Debido a esto, los resultados se obtienen para casos especiales en los que puede haber estrategias totalmente sensibles en equilibrio. Esto simplifica sustancialmente el análisis, ya que el comportamiento de los votantes en el segundo periodo está prácticamente caracterizado por esta propiedad.

Potencial del modelo

A continuación discutimos informalmente las principales propiedades del modelo de votación secuencial que acabamos de describir, que lo hacen apropiado para estudiar los problemas de coordinación, que eran nuestra motivación inicial.

Primero, es importante destacar que, con bastante generalidad, habrá gente que votará haciendo uso de más información de la que tenía al principio de la elección. Siendo más precisos: habrá gente que votará en ambos periodos. La intuición es bastante sencilla. Por un lado, si yo sé que todo el mundo va a votar en el periodo 1, entonces prefiero esperar al periodo 2 y hacer una decisión más informada (al fin y al cabo, mi voto en el periodo 1 no iba a servir para cambiar el voto de nadie, pues nadie está esperando al periodo 2). Por otro lado, si sé que todo el mundo va a votar en el periodo 2, entonces preferiré votar en el periodo 1, ya que en ningún caso voy a tener información a la hora de votar y votando en el periodo 1 puedo influir en el comportamiento de algún votante en el periodo 2.

El argumento que acabamos de hacer captura uno de los principales incentivos que intentamos entender con nuestro modelo: la tensión entre i) votar en el periodo 1, para así hacer que tu candidato parezca más fuerte y conseguir coordinación a su favor y ii) votar en el periodo 2 para tomar una decisión más informada. Para ilustrar, pensemos en un votante AB bajo estrategias totalmente sensibles:

1. Votando por A en el periodo 1, un votante AB será pivote principalmente cuando rompa un empate entre A y otro candidato (aumentando la posibilidad de que la coordinación final sea en A) o induce un empate (aumentando la coordinación en A y reduciendo la coordinación en el candidato con el que A ha empatado).
2. Votando en el periodo 2, el votante AB puede ser pivote si B está por delante de A después del periodo 1 y B y C están muy igualados, pues un voto adicional por B puede decidir la elección en su favor.

El primer punto es el efecto “hacer más fuerte a tu candidato” y el segundo es el efecto “evitar el voto inútil”. Uno de nuestros objetivos es entender cómo estos dos efectos entran en juego. Esto sugiere ya una implicación natural en nuestro modelo: cuanto más partidista sea un votante, más importante será el primer efecto y más probable será que prefiera votar pronto.

De cara a nuestros objetivos es todavía más importante el hecho de que, una vez que hay “voto informado” en equilibrio, se abre la puerta para estudiar hasta qué punto esto puede dar lugar a suficiente coordinación como para que se reduzcan significativamente las opciones de que un perdedor de Condorcet gane unas elecciones.

Resumen de resultados

Para terminar este resumen, presentamos un breve esquema de los principales resultados matemáticos y numéricos obtenidos en el trabajo en el que se ha basado este resumen y la charla impartida:

- En el caso de que los distintos candidatos que participan en la elección sean bastante simétricos entre sí, el efecto “hacer más fuerte a tu candidato” tiende a ser el dominante, con lo que la mayor parte de los votantes votan en el primer periodo.
- Además, en el caso en el que los candidatos no sean simétricos y se sepa de antemano que hay un perdedor de Condorcet (situaciones como la descrita en la Tabla 3), se observa como en el modelo con dos etapas se consigue una gran coordinación entre los votantes. Cuanto más fuerte es el perdedor de Condorcet, más votantes de los otros dos candidatos esperan hasta el segundo periodo para votar. Como consecuencia, la probabilidad de que el perdedor de Condorcet sea elegido se reduce significativamente con respecto al caso en el que la votación se lleva a cabo de modo clásico (y sin posibilidad de coordinación entre los votantes que no quieren que salga elegido el perdedor de Condorcet).

Coloquio: Indicadores de producción científica

Ariadna Arias

Facultade de Periodismo

9 de febrero de 2015

Moderadora:

Ariadna Arias

Participantes:

Javier Tarrío-Saavedra (UdC)

Salvador Naya (UdC)

Jesús R. Aboal (USC)

Juan José Nieto (USC)

El lunes 9 de febrero a las cinco de la tarde se celebró en la Facultad de Matemáticas un coloquio sobre los índices de productividad científica, esto es, indicadores que miden diferentes aspectos sobre las publicaciones que realiza un investigador. Para desarrollar el tema se contó con la participación de Javier Tarrío-Saavedra y Salvador Naya, autores de “Estudo bibliométrico do SUG”, junto con Jesús R. Aboal Viñas, investigador de la USC y Juan José Nieto Roig, investigador de la USC y editor de la revista *Fixed Point Theory and Applications*.

Cierto es que los indicadores bibliométricos no son un tema muy conocido fuera de la comunidad científica, tal y como explicó Javier Tarrío: “ni siquiera mis estudiantes están familiarizados este término y es algo muy importante si quieres ser investigador”. Ante un público formado principalmente por alumnos de matemáticas, era preciso definir el asunto. Según Salvador Naya “los índices bibliométricos miden la productividad y el impacto que tiene un investigador”. Además, existen diferentes tipos, como el índice H o el índice G, que calculan las citas que recibe un autor en diversos trabajos. Sin embargo, esto no tiene en cuenta la calidad del artículo. Juan José Nieto ejemplificó este problema citando una anécdota de su revista, y es que cada día le llegan cientos de artículos cuyo contenido es, la mayoría de las veces, carente de valor. Nos encontramos ante la dicotomía de producir y producir con calidad. Durante la charla, las intervenciones de todos los participantes giraron en torno a ello. Según Naya, un índice de impacto es mayor cuanto más se publique,

PALABRAS CLAVE: Indicador de producción científica.

por lo que ya no es posible dedicar todo un año a un proyecto, ya que se exigen varios en un corto período de tiempo. Durante el turno de preguntas, un alumno sacó a relucir este tema ya que, según sus propias palabras, la revista *Hola* tiene un gran índice de impacto. Claro que esto no es comparable a investigaciones científicas.

En China existe una urgencia cada vez mayor a producir trabajos como si de una cadena de montaje se tratara. Ante la proliferación de artículos de este país, ¿cómo puede competir España? El miedo al contagio de este método empieza a acrecentarse porque incluso para solicitar una *Starting Grant* de la Unión Europea se exige la estipulación del índice H. Perdemos calidad, pero también perdemos prestigio entre nuestros investigadores. Haciendo referencia a un artículo del periódico *El Mundo*, Nieto explicó que, según el índice H, sólo cuatro rectores universitarios de toda España dan la nota para considerarse investigadores notables.

Durante el coloquio también se mencionó el estudio bibliométrico llevado a cabo por Naya y Tarrío-Saavedra, en el cual se recoge un ranking de las universidades gallegas con respecto a los índices de impacto. Según este estudio, Santiago está en la vanguardia de la investigación, con Vigo pisándole los talones y Coruña en un tercer puesto. Hay que destacar que Vigo, a pesar de ser más nueva que las otras, produce muchos más artículos que el resto.

Si no existieran los índices bibliométricos no tendríamos una forma objetiva de valorar el trabajo de los investigadores. No obstante, es precisamente esa objetividad, esa frialdad de cálculo, la que supone una traba al futuro de la investigación: los jóvenes que quieren empezar de cero.