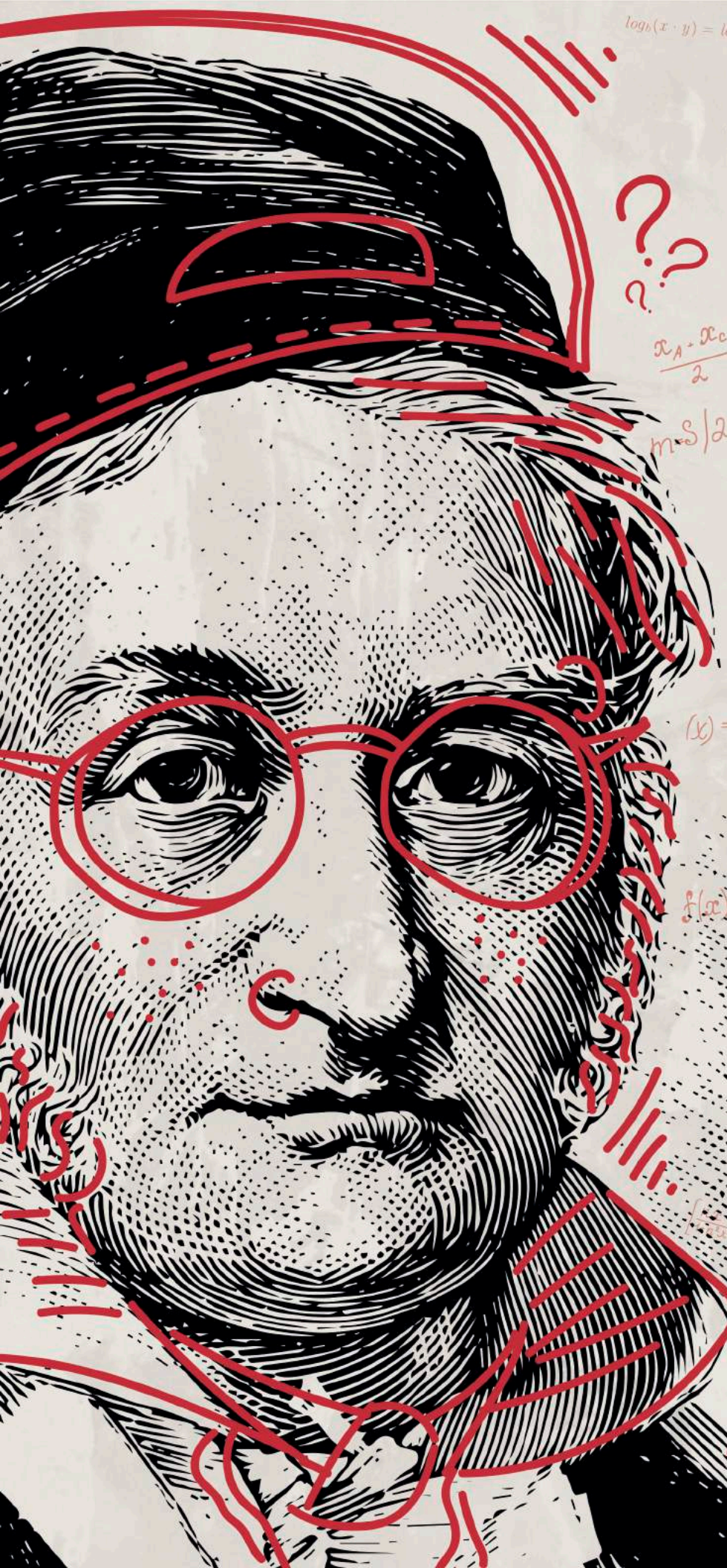




$$\log_b(x \cdot y) = \log_b(x) + \log_b y$$

$$y^3$$

$$\frac{df}{dt} = \lim_{h \rightarrow 0} \frac{f(t+h) - f(t)}{h}$$

$$\frac{\sin \alpha \cdot \cos \alpha}{\tan \alpha}$$

$$\sqrt{x^2 - x} - x^2$$

$$n! = \prod_{k=1}^n k$$

??
?..?

$$\frac{x_A - x_C}{2}$$

$$m \cdot S/2$$

$$(x) = \frac{d}{3}$$

$$f(x) \sim \frac{\alpha \pi}{2}$$

$$\frac{(n-1)x}{1-(n-1)2x^2}$$

$$\Delta x \rightarrow 0$$

$$(y) = d$$

$$\sum_{n=1}$$

$$\sqrt[3]{4}$$

$$m=S$$

$$V + E + F = 2$$

$$\sin \frac{3}{4}$$

KONRAD LORENZ
FUNDACIÓN UNIVERSITARIA

$$g(\xi) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} f(x) e^{-i\xi x} dx$$

$$f(x, \mu, \sigma) = \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$

e-ISSN
2665-2471

PAS XIN MATEMÁTICO

Vol. 5
No. 1

PASKÍN MATEMÁTICO

<https://editorial.konradlorenz.edu.co/paskin-matematico.html>

e-mail: paskin@konradlorenz.edu.co

Editor

John A. Arredondo

Fundación Universitaria Konrad Lorenz
alexander.arredondo@konradlorenz.edu.co

Comité Editorial

Ruth Alejandra Torres

Fundación Universitaria Konrad Lorenz
rutha.torresr@konradlorenz.edu.co

Julián Jiménez Cárdenas

Universidad de los Andes
jo.jimenezc1@uniandes.edu.co

Camilo Ramírez Maluendas

Universidad Nacional de Colombia
Sede Manizales
camramirezma@unal.edu.co

Andrés Mauricio Rivera

Pontificia Universidad Javeriana
Sede Cali
amrivera@javerianacali.edu.co

Leidy Catherine Sánchez

Fundación Cardio Infantil
lcascanio@lacardio.org

Jesús Muciño

Centro de Ciencias Matemáticas UNAM-Morelia
amrivera@javerianacali.edu.co

Esta publicación puede ser difundida y reproducida con fines académicos y científicos por todos aquellos que tengan a bien hacer un correcto uso de su contenido.

ISSN 2665-2471

Fundación Universitaria Konrad Lorenz: Tel: (57 1) 347 23 11, Carrera 9 Bis No. 62- 43 Bogotá – Colombia, email: info@konradlorenz.edu.co. **Carácter académico:** Institución Universitaria. Personería Jurídica por Resolución 18537 del 4 de noviembre de 1981 del Ministerio de Educación Nacional. Institución de Educación Superior sujeta a inspección y vigilancia por el Ministerio de Educación Nacional (Art. 2.5.3.2.10.2, Decreto 1075 de 2015).

SECCIONES DE POINCARÉ: DELATANDO EL CAOS EN EL PÉNDULO DOBLE

Dairo Andrés Sierra Zamora^{*}
dairoa.sierraz@konradlorenz.edu.co

Resumen.

En el presente artículo se presenta un estudio de la dinámica del péndulo doble y a partir de este sistema se introduce el concepto de caos en sistemas dinámicos deterministas. La estrategia empleada inicia con la deducción de las ecuaciones de movimiento, para lo cual se emplea el enfoque usado en la mecánica analítica mediante el planteamiento del lagrangiano del sistema, y a partir de este se deducen sus ecuaciones de movimiento usando las ecuaciones de *Euler-Lagrange*. Para comprender cómo abordar el problema mediante esta metodología se usa como ejemplo el péndulo simple y posteriormente se aplica la estrategia al péndulo doble. Finalmente, se introduce la noción de caos mediante el estudio numérico de la dinámica de las ecuaciones de movimiento para el péndulo doble y sus mapas de Poincaré.

1. Introducción.

El estudio del caos en los sistemas dinámicos se origina en el siglo XIX con Henri Poincaré, mientras estudiaba el problema de los tres cuerpos. Un sistema dinámico es un sistema que evoluciona en el tiempo. Un ejemplo de esto es el movimiento de la Tierra alrededor del Sol, o cómo cambia cierta población de conejos en el tiempo. Éstos son dos de los ejemplos más conocidos en la literatura, de hecho, también son ejemplos de sistema dinámico continuo y sistema dinámico discreto, respectivamente. La primera situación corresponde a un sistema dinámico continuo en donde la posición de la Tierra se puede modelar mediante una función que depende del tiempo. La segunda situación corresponde a un sistema dinámico discreto en la cual podemos pensar que la población de conejos este año es x_t y la del siguiente año será x_{t+1} ; la diferencia entre éstos es el dominio permitido para la variable independiente t . En el primero es claro que $t \in \mathbb{R}$ y en el segundo se considera que $t \in \mathbb{N}$. De acuerdo a [SSaRLD13] podemos definir formalmente un sistema dinámico de la siguiente forma.

Definición 1.1. *Un sistema dinámico continuo en $\mathbb{R}^n$ es una función continuamente diferenciable $\phi : \mathbb{R} \times \mathbb{R}^n \rightarrow \mathbb{R}^n$ llamada flujo, donde $\phi(t, X) = \phi_t(X)$ satisface las siguientes propiedades:*

1. $\phi_0 : \mathbb{R}^n \rightarrow \mathbb{R}^n$ es la función identidad: $\phi_0(X_0) = X_0$.

^{*}Estudiante de matemáticas Fundación Universitaria Konrad Lorenz.

2. $\phi_t \circ \phi_s = \phi_{t+s}$ para cada $t, s \in \mathbb{R}$.

El objetivo central en los sistemas dinámicos es encontrar cuando es posible la función $\phi(t, X)$, sin embargo, esto no es posible en todos los casos, para el caso del péndulo simple es sencillo determinar esta función pero para el caso del péndulo doble no es posible obtener esta función. Por esta razón en el estudio de los sistemas dinámicos existen estrategias cualitativas que permiten obtener información sobre el comportamiento del sistema dinámico; tales como los retratos de fase, el estudio de la estabilidad en los puntos de equilibrio o los diagramas de Poincaré.

A continuación se presentará el estudio del péndulo doble por medio de los mapas de Poincaré. El mapeo de Poincaré es una función definida en un subespacio de menor dimensión llamado sección de Poincaré [Nol15]. A continuación se representa la sección de Poincaré para un flujo en tres dimensiones en el cual notamos que es de dos dimensiones.

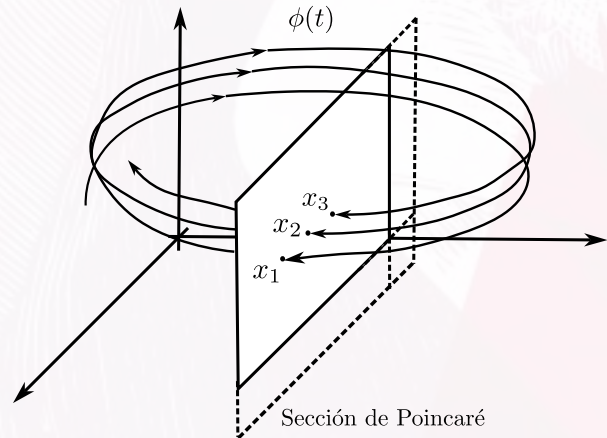


Figura 1: Sección de Poincaré para un flujo 3D, las trayectorias pueden pasar por encima o por debajo de cada uno de los puntos mostrados anteriormente.

Mediante el mapeo de Poincaré es posible estudiar de forma cualitativa la evolución de un sistema dinámico no lineal como lo es el péndulo doble. Al finalizar el estudio del péndulo doble se presentarán algunas secciones de Poincaré para ciertos valores de energía, y así poder estudiar la dinámica de este sistema.

2. Péndulo Simple.

El péndulo simple o péndulo matemático es uno de los sistemas dinámicos más sencillos de estudiar. Tal es su simplicidad que en cursos elementales de física se logra obtener información importante del sistema para oscilaciones pequeñas, es decir, donde $\sin \theta \approx \theta$; como que el período de oscilación (*tiempo que tarda en completar un ciclo completo de ida y vuelta*) es directamente proporcional a la longitud del péndulo y está dado por $T = 2\pi\sqrt{\ell/g}$. Este hecho es muy importante porque al leer con detalle esta expresión notamos que el periodo de oscilación T es independiente de la amplitud θ y de la masa m , afirmaciones que chocan con el sentido común, pues parece sensato pensar que si el péndulo se deja caer desde una amplitud $\theta_1 > \theta_2$ el periodo de oscilación $T_1 > T_2$ ver figura (2). Sin embargo, esto no es cierto como se puede comprobar experimentalmente. Además, el periodo de oscilación es independiente de la masa, es decir, no importa si la masa suspendida de éste es grande o pequeña su periodo de oscilación es el mismo; esto es realmente espectacular de verificar. Para ello sugiero al lector ver la siguiente clase magistral del profesor Walter Lewin en donde estos hechos se verifican de una manera sorprendente [LW11].

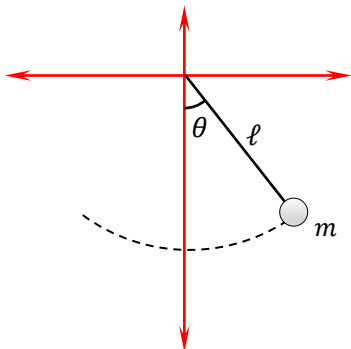


Figura 2: Péndulo simple de longitud ℓ y masa m .

Para el estudio de la dinámica de este sistema físico emplearemos el enfoque dado por la mecánica analítica o mecánica de Lagrange, es decir donde la evolución temporal del sistema se puede obtener a partir de cómo se cambia la energía cinética y la energía potencial del sistema. Posterior a ello es posible obtener las ecuaciones de movimiento mediante las ecuaciones de Euler-Lagrange y así tener la evolución temporal del sistema, el otro enfoque que se puede utilizar en este caso es el dado tradicionalmente por la segunda ley de Newton, estudiando las fuerzas que actúan sobre el sistema, para así determinar su comportamiento en función del tiempo, sin embargo, esta última estrategia no es recomendable cuando el sistema es complejo como lo es el péndulo doble.

Lagrangiano de un sistema:

El lagrangiano $\mathcal{L}$ para un sistema de partículas es una función escalar que depende de las coordenadas generalizadas

q_α que determinan la posición de la partícula de masa m_α y las velocidades asociadas $\dot{q}_\alpha$. El lagrangiano de un sistema de partículas carece de significado físico, pero es de gran utilidad matemática porque nos permite relacionar la energía cinética K_α y la energía potencial U_α del sistema en cada instante. El lagrangiano para un sistema de n partículas está definido por

$$\mathcal{L}(q_1, \dots, q_n, \dot{q}_1, \dots, \dot{q}_n) = \sum_{\alpha=1}^n [T_\alpha(\dot{q}_\alpha) - U_\alpha(q_\alpha)]. \quad (1)$$

En este sentido el lagrangiano es una función tal que

$$(q_1, \dots, q_n, \dot{q}_1, \dots, \dot{q}_n) \mapsto \mathcal{L}(q_\alpha, \dot{q}_\alpha, t) \in \mathbb{R}.$$

Por tanto, para el planteamiento del lagrangiano del péndulo simple es importante encontrar la energía cinética y la energía potencial para un sistema discreto de partículas.

Energía cinética:

La energía cinética es una cantidad escalar que para una masa puntual está dada por $K_\alpha = \frac{1}{2}m_\alpha \dot{r}_\alpha^2$. m_α representa la masa de la partícula α y $\dot{r}_\alpha$ la velocidad asociada a esta. De acuerdo al sistema de coordenadas planteado en la figura (2) el vector posición está dado por

$$\mathbf{r} = \ell(\sin \theta \hat{i} - \cos \theta \hat{j}),$$

y su derivada con respecto al tiempo es

$$\dot{\mathbf{r}} = \ell \dot{\theta}(\cos \theta \hat{i} + \sin \theta \hat{j}),$$

con lo cual podemos mostrar que la energía cinética para la partícula m está dada por

$$K = \frac{1}{2}m\ell^2\dot{\theta}^2. \quad (2)$$

Energía potencial:

Se denomina energía de configuración o simplemente energía potencial, a aquella que depende de la posición del objeto respecto a un punto de referencia, y está definida por la expresión $U(\mathbf{r}) = -\int \vec{F} \cdot d\vec{r}$. Antes de determinar los valores de energía potencial, es necesario establecer algunos puntos de referencia, y así determinar como está dada la energía potencial para este sistema (ver figura 2).

- $U = 0$, en la recta $y = 0$ de nuestro sistema de referencia.
- U es mínima cuando $\theta = 0$, es decir cuando el péndulo está en la posición más baja de la configuración mostrada en la figura (2).
- U es máxima cuando $\theta = \pi$, es decir cuando el sistema se encuentra en la posición más alta en relación al punto de referencia.

Observaremos entonces que el valor de la energía potencial dependerá de la variable θ como variable de movimiento, pero además también de la longitud ℓ , constante durante todo el movimiento. Luego, dado que el elemento de línea en coordenadas rectangulares es $d\vec{l} = dx\hat{i} + dy\hat{j}$ la energía potencial de cada partícula está dada por

$$U = -mg\ell \cos \theta. \quad (3)$$

2.1. Ecuaciones de movimiento.

Previamente determinamos la energía cinética y la energía potencial de la partícula, por lo que de acuerdo a la ecuación (1) tenemos que el lagrangiano para este sistema es

$$\mathcal{L}(\theta, \dot{\theta}) = \frac{1}{2}m\ell^2\dot{\theta}^2 + mg\ell \cos \theta. \quad (4)$$

Conociendo el lagrangiano de un sistema de partículas podemos deducir las ecuaciones de movimiento para poder determinar la evolución temporal del sistema. Las ecuaciones de movimiento las podemos obtener a partir de las ecuaciones de Euler-Lagrange

$$\frac{d}{dt} \left(\frac{\partial \mathcal{L}}{\partial \dot{q}_\alpha} \right) = \frac{\partial \mathcal{L}}{\partial q_\alpha}. \quad (5)$$

Este conjunto de ecuaciones diferenciales juegan un papel análogo al de la segunda ley de Newton. Con ellas podemos deducir cómo se comporta el sistema en el tiempo, es decir, cuál será su posición y su velocidad en cada instante de tiempo después de un valor $t = t_0$. Al aplicar estas ecuaciones obtenemos

$$\begin{aligned} \frac{d}{dt} \left(\frac{\partial \mathcal{L}}{\partial \dot{\theta}} \right) - \frac{\partial \mathcal{L}}{\partial \theta} &= 0, \\ m\ell^2\ddot{\theta} + mg\ell \sin \theta &= 0, \\ \ddot{\theta} + \frac{g}{\ell} \sin \theta &= 0. \end{aligned} \quad (6)$$

La última de estas ecuaciones es una ecuación diferencial no lineal de segundo orden muy conocida en la literatura de la física-matemática, y se suele denotar con $\omega^2 = \frac{g}{\ell}$ por lo cual toma la forma

$$\ddot{\theta} + \omega^2 \sin \theta = 0. \quad (7)$$

Además, es muy usual que dentro de la literatura se realice una linealización de esta ecuación tomando la forma de una ecuación diferencial ordinaria lineal de segundo orden; donde se suelen hacer aproximaciones de primer orden mediante la serie de Taylor de la función seno $\sin \theta = \theta - \frac{\theta^3}{3!} + \frac{\theta^5}{5!} - \dots$ y con lo cual se puede escribir la ecuación (7) de la forma $\ddot{\theta} + \omega^2 \theta = 0$.

2.2. Análisis de la dinámica del péndulo simple.

Sobre el estudio de un sistema dinámico usualmente nos debemos plantear un conjunto de preguntas que nos permitirán hacer una descripción del sistema a estudiar, como por ejemplo ¿Cuáles son los puntos de equilibrio de este sistema? Es decir, aquellos puntos en los cuales el sistema permanecerá completamente invariante en el tiempo. ¿Qué tipo de comportamiento tiene el sistema dinámico en el tiempo? Esta puede ser una de las preguntas más interesantes, ya que en el estudio de los sistemas dinámicos a partir del trabajo hecho por Lorenz en 1881 para el problema de los tres cuerpos [Lyn18], sabemos que algunos sistemas dinámicos se vuelven caóticos en el tiempo; sin embargo este no es el caso del péndulo simple, como se puede observar en la gráficas de la figura (2) estos sistemas bajo ausencia de fricción son muy estables y su comportamiento es periódico en el tiempo.

La solución de las ecuaciones diferenciales del péndulo simple son muy conocidas dentro de la literatura de la física matemática. Sin embargo, la intención de esta sección es mostrar el enfoque empleado en la mecánica de Lagrange para la deducción de las ecuaciones de movimiento del péndulo simple. Notemos que las ecuaciones diferenciales obtenidas en las ecuaciones (7) y (6) son exactamente las mismas que las que se pueden encontrar aplicando la segunda ley de Newton. A continuación se presenta una simulación hecha en Python mediante integración numérica para un péndulo simple con parámetros $\ell = 1.0 \text{ m}$ y $g = 9.8 \text{ m/s}^2$, luego se presenta la solución numérica de la posición angular θ y la velocidad angular ω para el sistema mostrado en la figura (3).

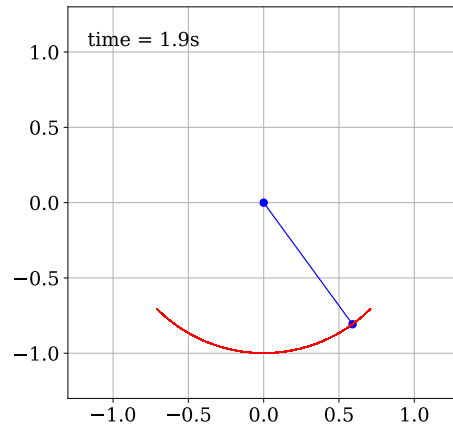


Figura 3: Simulación para el péndulo simple mediante la ecuación (7). Con condiciones iniciales $\theta_0 = \pi/4$ y $\omega_0 = 0.2 \text{ rad/s}$.

El objetivo al presentar esta sección es mostrar la metodología que se empleará para la deducción de las ecuaciones de movimiento para el péndulo doble en la siguiente sección. Notemos que en este caso el enfoque de la mecánica de Lagrange resulta mucho más sencillo que tratar de abordar el

problema mediante la aplicación de la segunda ley de Newton. Finalmente, una de las estrategias cualitativas para el estudio de los sistemas dinámicos es el diagrama de fase que, según [Nol15] se define como un caso especial de espacios de estados para sistemas Hamiltonianos o Lagrangianos compuestos de $2N$ dimensiones para N coordenadas generalizadas independientes.

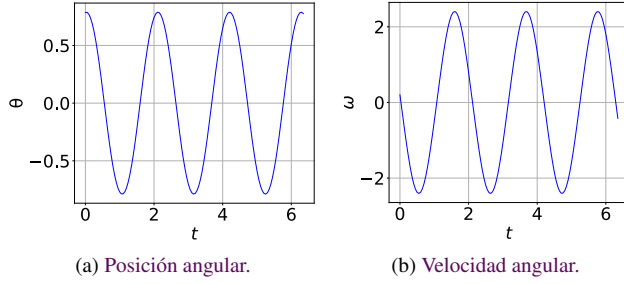


Figura 4: Solución numérica para el péndulo simple mediante la ecuación (7) y las condiciones dadas en la figura (3).

En la figura (5) podemos notar que tenemos puntos de equilibrio en $(-\pi, 0)$, $(0, 0)$ y $(\pi, 0)$ en este caso para el marco de referencia planteado, el primer punto de equilibrio resulta ser exactamente el tercer punto; por la trayectorias se puede notar que son puntos repulsivos o inestables, lo que quiere decir que cerca de estos puntos el sistema es altamente susceptible a cambios, como es de esperar desde el punto de vista físico. Por otra parte, tenemos que el punto de $(0, 0)$ es un atractor estable y vemos que las trayectorias alrededor de este punto son cíclicas, y se pueden interpretar como las oscilaciones periódicas para valores de θ pequeños. El hecho que sea un atractor nos dice que cerca a este punto el sistema tenderá a estar en equilibrio.

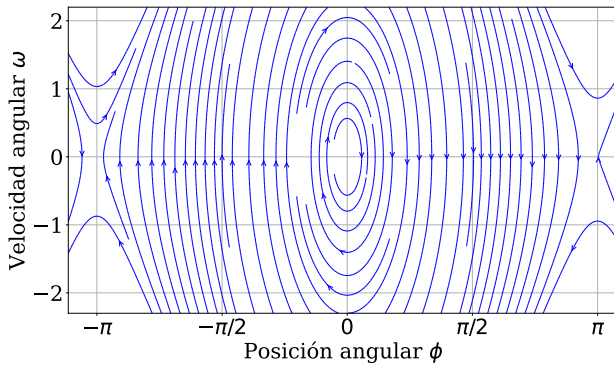


Figura 5: Diagrama de fase para el péndulo simple, con puntos de equilibrio $(-\pi, 0)$, $(0, 0)$ y $(\pi, 0)$.

3. Péndulo Doble.

A continuación aplicaremos un enfoque similar al usado en la sección anterior para el estudio de la dinámica del péndulo doble. Estudiar y comprender la dinámica del péndulo doble

consistirá en encontrar un conjunto de ecuaciones diferenciales que permitan describir su evolución temporal en el tiempo. En otras palabras, encontrar y “solucionar” un conjunto de ecuaciones diferenciales que nos permitan conocer su velocidad y su posición en cada instante de tiempo. Sin embargo, notaremos que en este caso, hallar las ecuaciones diferenciales que describen el movimiento es un poco más laborioso que lo visto en la sección anterior.

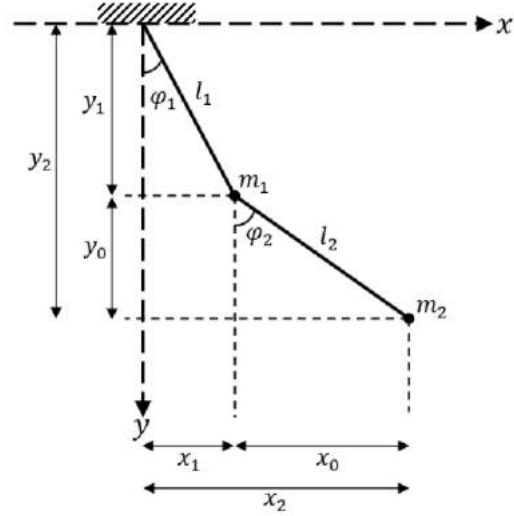


Figura 6: Péndulo doble de longitudes l_1, l_2 y masas m_1, m_2 .

El péndulo doble es un sistema compuesto de dos péndulos acoplados de longitudes l_1, l_2 y de masas m_1, m_2 como se ilustra en la figura (6). Cada una de las masas forma un ángulo respecto a la vertical de ϕ_1 y ϕ_2 , respectivamente. Aplicaremos la estrategia usada en la sección anterior para el caso del péndulo simple para hallar sus ecuaciones de movimiento, es decir, hallaremos la energía cinética y la energía potencial de cada masa y con ello determinaremos el *lagrangiano* del sistema para finalmente aplicar la ecuaciones de Euler-Lagrange.

Energía Cinética:

En este caso la posición de cada partícula de acuerdo a la figura (6) está dada por

$$r_1 = l_1 \sin \phi_1 \hat{i} - l_1 \cos \phi_1 \hat{j},$$

$$r_2 = (l_1 \sin \phi_1 + l_2 \sin \phi_2) \hat{i} - (l_1 \cos \phi_1 + l_2 \cos \phi_2) \hat{j}.$$

Donde r_1 y r_2 representan las posiciones de las masas m_1 y m_2 respectivamente. Teniendo en cuenta que $v_1 = \dot{r}_1 = \frac{dr_1}{dt}$ y $v_2 = \dot{r}_2 = \frac{dr_2}{dt}$ y empleando la definición de producto punto para calcular el cuadrado de la velocidad tenemos

$$\dot{r}_1 \cdot \dot{r}_1 = (\dot{r}_1)^2 = l_1^2 (\dot{\phi}_1)^2,$$

$$\dot{r}_2 \cdot \dot{r}_2 = (\dot{r}_2)^2 = l_1^2 (\dot{\phi}_1)^2 + l_2^2 (\dot{\phi}_2)^2 + 2l_1 l_2 \dot{\phi}_1 \dot{\phi}_2 \cos(\phi_1 - \phi_2).$$

De tal forma que la energía cinética para m_1 y m_2 , respectivamente es K_1 y K_2 , con lo cual tenemos que

$$K_1 = \frac{1}{2} m_1 l_1^2 (\dot{\varphi}_1)^2,$$

$$K_2 = \frac{1}{2} m_2 [l_1^2 \dot{\varphi}_1^2 + l_2^2 \dot{\varphi}_2^2 + 2l_1 l_2 \dot{\varphi}_1 \dot{\varphi}_2 \cos(\varphi_1 - \varphi_2)]. \quad (8)$$

Energía Potencial:

Observaremos que los valores de la energía potencial dependerán de las variables φ_1 y φ_2 como variables de movimiento, pero además también de las longitudes l_1 y l_2 constantes durante todo el movimiento. Dado que el elemento de línea en coordenadas rectangulares es $d\vec{l} = dx\hat{i} + dy\hat{j}$ la energía potencial de cada partícula está dada por

$$U_1 = -mgl_1 \cos \varphi_1,$$

$$U_2 = -mg(l_1 \cos \varphi_1 + l_2 \cos \varphi_2). \quad (9)$$

3.1. Ecuaciones de movimiento.

Previamente hallamos la energía cinética y la energía potencial para cada partícula, ahora podemos determinar el lagrangiano mediante la ecuación (1); teniendo en cuenta que en este caso debemos contemplar los valores de la energía cinética y la potencial obtenidos en las ecuaciones (8) y (9) respectivamente. De modo que el lagrangiano $\mathcal{L}(\varphi_1, \varphi_2, \dot{\varphi}_1, \dot{\varphi}_2)$ para este sistema es

$$\mathcal{L} = \frac{1}{2} (m_1 + m_2) l_1^2 \dot{\varphi}_1^2 + \frac{1}{2} l_2^2 \dot{\varphi}_2^2 + m_2 l_1 l_2 \dot{\varphi}_1 \dot{\varphi}_2 \cos(\varphi_1 - \varphi_2) + (m_1 + m_2) g l_1 \cos \varphi_1 + m_2 g l_2 \cos \varphi_2. \quad (10)$$

Como notamos en la anterior expresión, el lagrangiano de este sistema es mucho más complejo que el presentado en la ecuación (4), lo que permite intuir que la dinámica del péndulo doble es mucho más elaborada que la del péndulo simple. Ahora conociendo el lagrangiano de un sistema de partículas podemos deducir las ecuaciones de movimiento para así poder determinar la evolución temporal del sistema. En este caso, dadas las ecuaciones de Euler-Lagrange (5) podemos encontrar la ecuación de movimiento para cada partícula α .

$$\frac{d}{dt} \left(\frac{\partial \mathcal{L}}{\partial \dot{q}_1} \right) = \frac{\partial \mathcal{L}}{\partial q_1},$$

de modo que la ecuación de movimiento para m_1 quedaría

$$(m_1 + m_2) l_1^2 \ddot{\varphi}_1 + m_2 l_1 l_2 \ddot{\varphi}_2 \cos(\varphi_1 - \varphi_2) + m_2 l_1 l_2 \dot{\varphi}_2^2 \sin(\varphi_1 - \varphi_2) + (m_1 + m_2) g l_1 \sin \varphi_1 = 0. \quad (11)$$

Ahora, para m_2 tenemos que

$$\frac{d}{dt} \left(\frac{\partial \mathcal{L}}{\partial \dot{q}_2} \right) = \frac{\partial \mathcal{L}}{\partial q_2},$$

de modo que la ecuación de movimiento para m_2 es

$$m_2 l_2^2 \ddot{\varphi}_2 + m_2 l_1 l_2 \ddot{\varphi}_1 \cos(\varphi_1 - \varphi_2) - m_2 l_1 l_2 \dot{\varphi}_1^2 \sin(\varphi_1 - \varphi_2) + m_2 g l_2 \sin \varphi_2 = 0. \quad (12)$$

Con las ecuaciones (11) y (12) podemos describir la dinámica del péndulo doble. Notemos primero algo interesante: si en la ecuación (11) hacemos $m_2 = 0$ y $l_2 = 0$ tenemos la ecuación del péndulo simple (7), por lo cual podemos considerar que este modelo es una ampliación del mostrado en la figura (2). Por otro lado, hemos obtenido un sistema de ecuaciones diferenciales no lineales de segundo grado, para el cual no tenemos herramientas matemáticas para encontrar su solución de forma analítica. Sin embargo, disponemos de herramientas computacionales para comprender un poco la dinámica de este sistema. A continuación se presentan algunas simulaciones hechas en Python mediante la integración numérica de estas expresiones.

3.2. Análisis de la dinámica del péndulo doble.

En la simulación presentada en la figura (7) se puede observar que bajo ciertas condiciones iniciales (*valores pequeños de $\varphi_1, \varphi_2, \omega_1$ y ω_2*) la dinámica del péndulo doble es *cuasiperiódica*, es decir, éste presenta un comportamiento repetitivo en el tiempo bajo la ausencia de fricción, en relación a las posiciones angulares y velocidades angulares las cuales son similares a un instante previo. Esta dinámica es mucho más evidente en las gráficas (8) y (9), en las cuales se presenta su evolución en un intervalo de tiempo de 12 s, donde observamos que tanto la posición angular y la velocidad angular de ambas masas es cuasiperiódica en el tiempo.

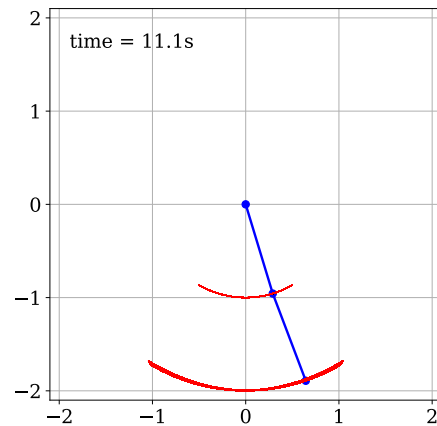


Figura 7: Simulación para el péndulo doble por medio de la ecuación (12). Con condiciones iniciales $\varphi_{01} = \pi/6, \omega_{01} = 0$ y $\varphi_{02} = \pi/6, \omega_{02} = 0$.

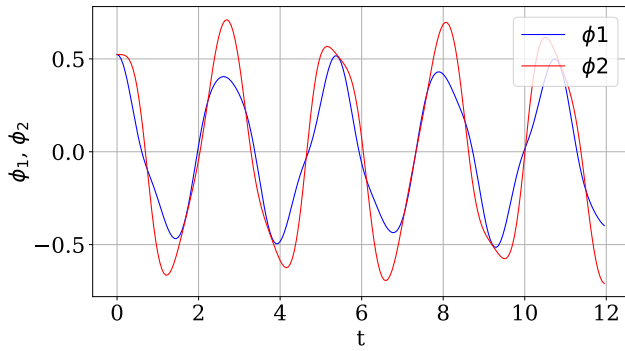


Figura 8: Comportamiento cuasiperiódico para la posición angular vs. tiempo de ambas masas. Con condiciones las iniciales dadas en la figura (7).

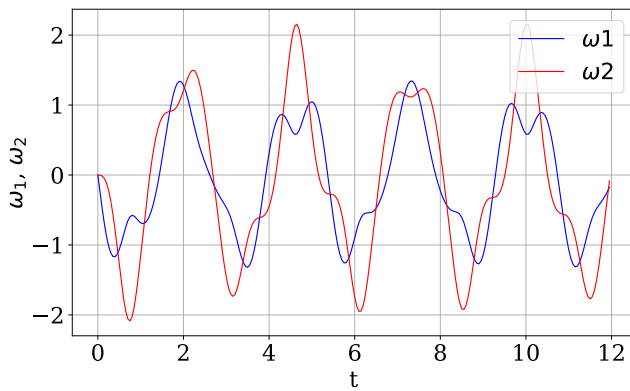


Figura 9: Comportamiento cuasiperiódico para la velocidad angular vs. tiempo de ambas masas. Con condiciones las iniciales dadas en la figura (7).

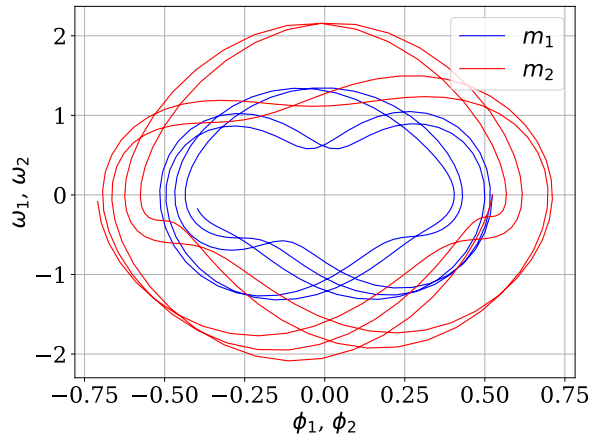


Figura 10: Diagrama de fase del péndulo doble para valores iniciales mostrados en la figura (7).

En la figura (10) se representa el diagrama de fase para el sistema, que resulta ser muy importante para el estudio cualitativo de la evolución del sistema, debido a que en este caso no es posible conocer la solución analítica de las ecuaciones de

movimiento obtenidas. A continuación se seleccionarán ciertos valores para las condiciones iniciales, donde la dinámica del péndulo doble es completamente caótica en el sentido que la posición angular de m_2 no es periódica y predecirla es una tarea compleja; es usual en la literatura llamar a este tipo de sistemas que presentan un *caos determinista*.

3.3. Caracterización del caos en el péndulo doble.

Para el péndulo doble tenemos que los puntos de equilibrio del sistema están dados de la forma $(\phi_1, \phi_2, \omega_1, \omega_2)$ son $(0, 0, 0, 0)$, $(\pi, 0, 0, 0)$, $(0, \pi, 0, 0)$ y $(\pi, \pi, 0, 0)$, conviene en este caso clasificar qué tipo de puntos críticos son de forma similar a lo hecho en el caso del péndulo simple. Desde el punto de vista físico es claro que el punto $(0, 0, 0, 0)$ es un punto estable como es de esperar bajo el sentido físico del problema y que los puntos $(\pi, 0, 0, 0)$, $(0, \pi, 0, 0)$ y $(\pi, \pi, 0, 0)$ son puntos críticos inestables, los cuales son muy susceptibles por la acción de la gravedad.

Según [Nol15] los sistemas caóticos poseen sensibilidad exponencial a las variaciones de las condiciones iniciales es decir, que un mismo sistema puede tener una evolución en el tiempo muy diferente bajo pequeños cambios en las condiciones. En las figuras (11), (12) y (13) se presenta este hecho mediante una serie de simulaciones hechas en Python a través de la integración numérica de sus ecuaciones de movimiento. Se observará que efectivamente es un sistema caótico.

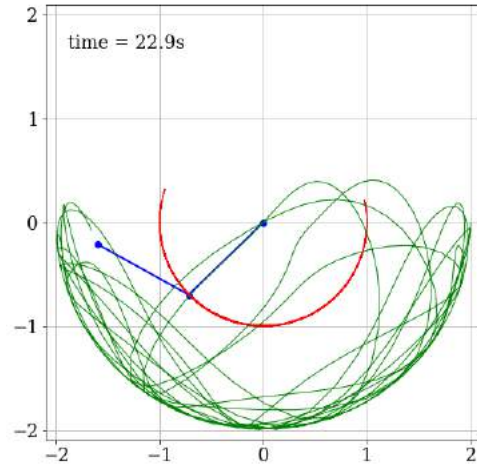


Figura 11: Simulación en Python para el péndulo doble mediante las ecuaciones (12) y (11), con parámetros $g = 9.8 \text{ m/s}^2$ y $l_1 = l_2 = 1 \text{ m}$.

En la figura (12) se representa la simulación de dos péndulos de forma simultánea con condiciones iniciales muy cercanas, con el fin de observar la sensibilidad que tiene la dinámica del sistema frente a pequeños cambios de las condiciones iniciales. Se puede observar que la evolución del sistema es similar los primeros 10 segundos, pero a partir de este momento el comportamiento del péndulo azul es muy diferente al péndulo cian. Este hecho se hace mucho más evidente en la

figura (13), en la cual se muestra cómo han evolucionado ambos péndulos pasados 22 segundos. La trayectoria de la masa m_2 para el péndulo de color azul se muestra de color verde y la trayectoria de la masa m_2 para el péndulo de color cyan se muestra de color negro, en este punto se observa que aunque las condiciones iniciales de ambos péndulos son muy cercanas entre sí, éstos tienen una evolución temporal muy diferente con lo cual podemos verificar el comportamiento caótico de este sistema dinámico.

Péndulo	Ángulo ($^{\circ}$)	Velocidad ($^{\circ}/seg$)
Azul	$\varphi_{01} = 90.0$ $\varphi_{02} = 90.0$	$\omega_{01} = 0.00$ $\omega_{02} = 1.00$
Cyan	$\varphi_{01} = 90.1$ $\varphi_{02} = 90.1$	$\omega_{01} = 0.00$ $\omega_{02} = 1.01$

Cuadro 1: Condiciones iniciales de las simulaciones mostradas en las figuras (12) y (13).

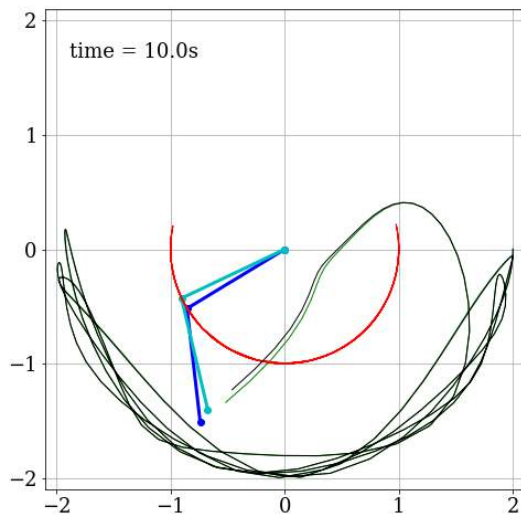


Figura 12: Representación de dos péndulos de forma simultánea pasados 10 s. Parámetros $g = 9.8 \text{ m/s}^2$ y $l_1 = l_2 = 1 \text{ m}$ para ambos péndulos.

Mediante integración numérica de las ecuaciones (11) y (12) podemos estudiar con detalle el comportamiento de este sistema, en la figura (14) se observa la evolución temporal de la posición angular para el péndulo azul con condiciones iniciales dadas en el cuadro (1). En esta gráfica se observa que la posición angular de m_2 es completamente errática y no corresponde a un movimiento periódico. También podemos evidenciar la sensibilidad del sistema a las condiciones iniciales dadas en el cuadro (1) en la figura (15) en la cual se puede ver que el comportamiento de la posición angular para m_2 es muy similar los diez primeros segundos, pero que después de este instante la evolución del sistema es muy diferente. La cual es una característica esencial del caos determinista mencionado

previamente.

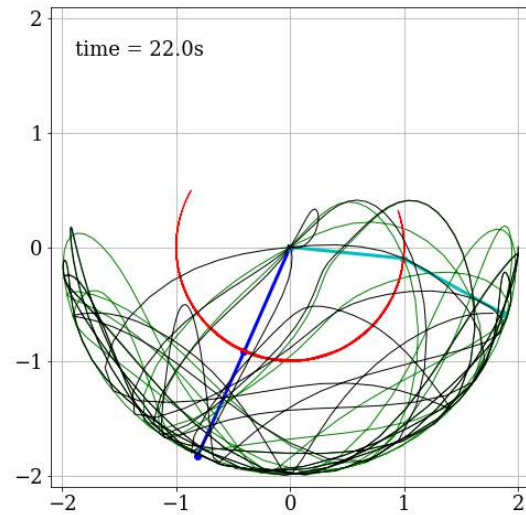


Figura 13: Simulación en Python para el péndulo doble mediante la ecuación (12).

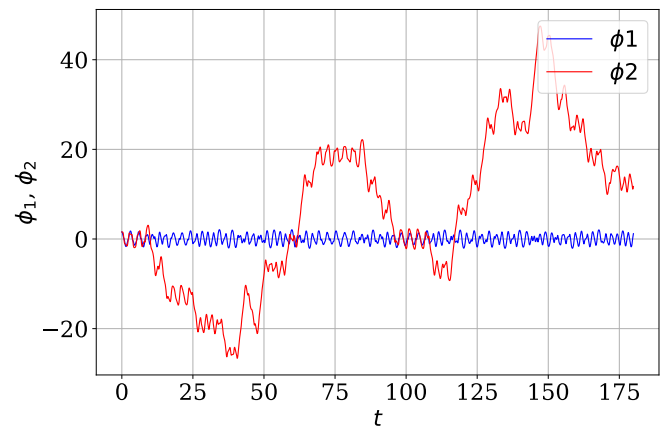


Figura 14: Posición angular vs. tiempo para el péndulo doble. Se puede evidenciar el comportamiento errático para φ_2 .

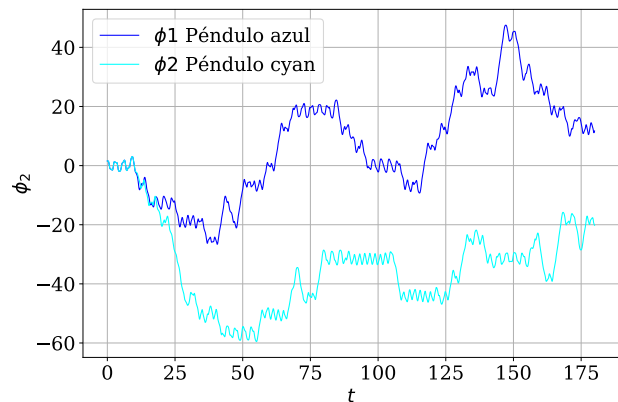


Figura 15: comparación de la evolución temporal de la posición angular para m_2 . Se puede evidenciar el comportamiento caótico para φ_2 y la sensibilidad a las condiciones iniciales.

3.4. Secciones de Poincaré para el péndulo doble.

La sección de Poincaré es una herramienta muy útil para el estudio de los sistemas dinámicos. Estas secciones nos permiten estudiar de forma cualitativa la evolución de un sistema dinámico complejo usualmente no lineal, y determinar si éste posee un comportamiento cuasiperiódico o caótico. Poincaré usó la idea de una sección (*Sección de Poincaré*) de menor dimensión al flujo $\phi(t)$ y transversal a este, de modo que, el flujo original del sistema corta esta sección en algunos puntos $x_1, x_2, \dots$ como se puede observar en la figura (1). Por ejemplo si el sistema dinámico modelado es el péndulo mostrado en la primera sección, debido a que su comportamiento es periódico en el tiempo la representación de la sección de Poincaré sería un punto ya que el $\phi(t)$ cortaría la sección siempre en el mismo punto. Sin embargo en sistemas no lineales como el caso del péndulo doble la sección de Poincaré es más interesante. A continuación se presentan algunos mapas de Poincaré para el péndulo doble para distintos valores de energía.

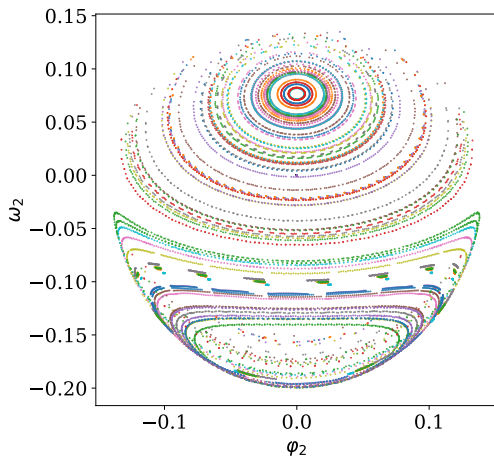


Figura 16: Sección de Poincaré para el péndulo doble para $E = 0.1$ se observa unas trayectorias cuasiperiódicas.

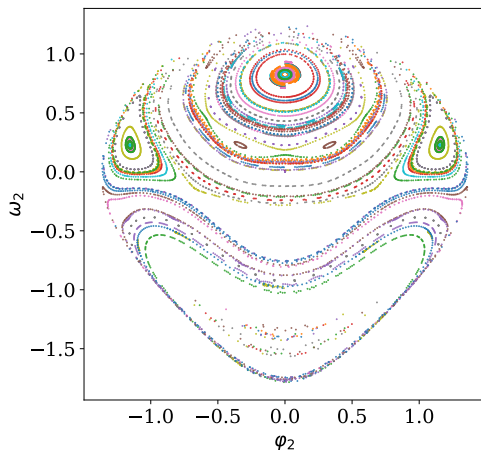


Figura 17: $E = 0.8$ las órbitas cerradas se pierden y observamos regiones más dispersas.

En las anteriores figuras se observa que para valores bajos de la energía del sistema el péndulo doble tiene un comportamiento cuasiperiódico, lo cual se puede evidenciar en la figura (16) en las regiones ovaladas. Cuando la energía del sistema aumenta se observa que los puntos en la sección de Poincaré la interceptan de forma dispersa lo cual caracteriza el comportamiento caótico lo cual se observa en la figura (18).

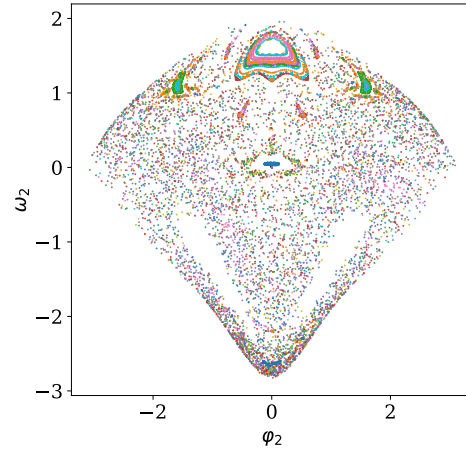


Figura 18: $E = 2.0$ para valores altos de energía se observa completamente el caos del sistema.

Como vemos el estudio de la dinámica del péndulo por medio de las secciones de Poincaré permite el análisis de cualitativo de un sistema dinámico, y determinar si éste posee un comportamiento caótico. Una ventaja de estudiar este sistema mediante esta estrategia es que un punto del flujo $\phi(t)$ está en cuatro dimensiones, con ayuda de la sección de Poincaré podemos hacer esta representación en una dimensión menor y proyectarla en el plano como se hizo en las últimas figuras.

Referencias.

- [Lyn18] S. Lynch, *Dynamical Systems with Applications using Python*, Springer International Publishing, 2018.
- [Nol15] D.D. Nolte, *Introduction to Modern Dynamics: Chaos, Networks, Space and Time*, Oxford University Press, 2015.
- [LLaEML94] L.D. Landau and E. M. Lifshitz, *Mecánica: Curso de física teórica: Curso de física teórica*, 1994.
- [SML70] S. M. Lee, *The double-simple pendulum problem*, American Journal of Physics **38** (1970), 536-537.
- [SSaRLD13] Stephen Smale and Robert L. Devaney, *Differential Equations, Dynamical Systems, and an Introduction to Chaos*, Third Edition, Academic Press, 2013.
- [LW11] Lewin W., *For the love of physics*, 2011. <https://www.youtube.com/watch?v=sJG-rXBmCc>.

Acerca del autor: Dairo actualmente es estudiante de octavo semestre de matemáticas en la Fundación Universitaria Konrad Lorenz, tiene una formación previa en licenciatura en física de la Universidad Pedagógica Nacional. Entre sus intereses se encuentran las ecuaciones diferenciales y los sistemas dinámicos. En su tiempo libre le gusta disfrutar de buena música y espera algún día aprender a tocar batería.

SOBRE EL PROBLEMA DE LOS CUATRO SUBESPACIOS

Gonzalo Medina*

gmedinaar@unal.edu.co

Resumen

Este artículo está dedicado al clásico Problema de los Cuatro Subespacios, que consiste en obtener la clasificación completa, salvo isomorfismo, de cuatro subespacios vectoriales de un espacio vectorial de dimensión finita sobre un campo. Nuestro objetivo es describir el problema en un contexto más general que permite exhibir y examinar algunas de las características que lo hacen interesante.

1. Introducción.

El célebre Problema de los Cuatro Subespacios, como ya dijimos, consiste en clasificar, salvo isomorfismo, los 4-sistemas de subespacios vectoriales de un espacio vectorial de dimensión finita sobre un campo k . El problema fue planteado e inicialmente resuelto por Gelfand y Ponomarev, en [GP70], para un campo algebraicamente cerrado. De manera independiente y utilizando métodos diferentes, Nazarova, en [Naz67, Naz75], resolvió también el problema, casi de manera simultánea a Gelfand y Ponomarev, pero ahora para un campo arbitrario. De hecho, tal como se demostró en [Naz75, págs 765-770], Nazarova resolvió en [Naz67] un problema matricial que es equivalente al problema de los cuatro subespacios.

Posteriormente, Brenner dio otra solución, primero parcial en [Bre67] en donde consideró el caso ahora denominado no regular, y posteriormente en [Bre74a] modificó y extendió los resultados de Gelfand y Ponomarev al caso de un campo arbitrario o incluso al de un anillo de división no conmutativo. Brenner fue quien primero describió completamente los anillos de endomorfismos asociados en los dos artículos mencionados.

En [MZ04], Medina y Zavadskij presentaron una solución mucho más corta y de carácter elemental para el problema sobre un campo arbitrario, que de hecho es válida sobre anillos de división arbitrarios y en la que también se describen los anillos de endomorfismos asociados.

Forbregd, en [For08] (tesis de maestría dirigida por Smalø), presenta una solución para campos algebraicamente cerrados haciendo uso de la teoría de Auslander-Reiten.

Las demostraciones originales de Gelfand y Ponomarev y de Nazarova, así como las correspondientes variantes de Brenner y de Forbregd, son largas y más bien complejas; la

solución dada por Medina y Zavadskij es corta y requiere tan solo herramientas de álgebra lineal, formas canónicas y de teoría de anillos.

Nuestro interés en este artículo es presentar algunas de las características que hacen interesante el problema de los cuatro subespacios. La estructura del artículo es la siguiente: comenzamos presentando, en la sección 2, las definiciones y notaciones iniciales necesarias; en la sección 3 mostramos cómo se resuelve el problema principal para los casos iniciales (en los que la solución es relativamente sencilla) y examinamos los casos que actualmente son intratables; finalmente, en la sección 4 nos ocupamos del caso especial que corresponde al clásico problema de los cuatro subespacios.

El autor agradece al revisor por sus valiosos y útiles comentarios y sugerencias.

2. Preliminares.

Vamos a considerar el problema más general de obtener la clasificación completa, salvo isomorfismo, de un número finito cualquiera de subespacios vectoriales de un espacio vectorial de dimensión finita sobre un campo. De manera precisa, dado un campo k , para cada entero positivo fijo r , un sistema de r -subespacios sobre k es una colección

$$U = (U_0, U_1, \dots, U_r),$$

en donde U_0 es un espacio vectorial de dimensión finita sobre k y cada U_i , para $i \in \{1, \dots, r\}$, es un subespacio vectorial de U_0 . La *dimensión* del r -sistema $U = (U_0, U_1, \dots, U_r)$ es el vector $\dim U = (d_0, d_1, \dots, d_r) \in \mathbb{Z}^{r+1}$ tal que, para todo $i \in \{0, \dots, r\}$, se tiene que $d_i = \dim U_i$.

Dados dos sistemas de r -subespacios $U = (U_0, U_1, \dots, U_r)$ y $V = (V_0, V_1, \dots, V_r)$, su *suma directa* es el sistema $U \oplus V$ dado, de manera natural, por

$$U \oplus V = (U_0 \oplus V_0, U_1 \oplus V_1, \dots, U_r \oplus V_r).$$

Para un entero positivo l y un sistema U de r -subespacios, definimos

$$U^l = \begin{cases} U, & \text{para } l = 1. \\ U^{l-1} \oplus U, & \text{para } l \geq 2. \end{cases}$$

Un *morfismo* $\mathcal{A}$ entre dos sistemas de r -subespacios $U = (U_0, U_1, \dots, U_r)$ y $V = (V_0, V_1, \dots, V_r)$ es una transformación k -lineal $\mathcal{A}: U_0 \rightarrow V_0$ tal que $\mathcal{A}(U_i) \subseteq V_i$, para todo $i \in \{1, \dots, r\}$. Dos sistemas de r -subespacios $U =$

*Departamento de Matemática y Estadística, Universidad Nacional de Colombia, Sede Manizales

$(U_0, U_1, \dots, U_r)$ y $V = (V_0, V_1, \dots, V_r)$ son *isomorfos*, lo cual denotamos $U \simeq V$, si existe un isomorfismo k -lineal $\mathcal{A}: U_0 \rightarrow V_0$ tal que $\mathcal{A}(U_i) = V_i$, para todo $i \in \{1, \dots, r\}$.

Un sistema de r -subespacios U es *indescomponible* si cada vez que se tenga $U = V \oplus W$, para sistemas de r -subespacios V y W , necesariamente $V = 0$ o $W = 0$. Aquí, 0 es el sistema nulo de r -subespacios: $0 = (0, 0, \dots, 0)$. Un sistema se denomina *descomponible* si no es indescomponible.

Para un sistema de r -subespacios $U = (U_0, U_1, \dots, U_r)$ sobre un campo k , su *sistema dual* es el sistema de r -subespacios $U^* = (U_0^*, U_1^*, \dots, U_r^*)$ sobre k dado por

$$\begin{aligned} U_0^* &= \text{Hom}_k(U_0, k) \\ &= \{\mathcal{A}: U_0 \rightarrow k \mid \mathcal{A} \text{ es funcional } k\text{-lineal}\}, \\ U_i^* &= U_i^\perp \\ &= \{\mathcal{B} \in U_0^* \mid \mathcal{B}(U_i) = 0\}, \text{ para todo } i \in \{1, \dots, r\}. \end{aligned}$$

Dado un morfismo $\mathcal{C}: U \rightarrow V$ entre dos sistemas de r -subespacios, su *morfismo dual* $\mathcal{C}^*: V^* \rightarrow U^*$ está definido por

$$\mathcal{C}^*(\mathcal{A}) = \mathcal{A}\mathcal{C}, \text{ para todo } \mathcal{A} \in V^*.$$

Nota. Ya que U_0 es de dimensión finita sobre k , tenemos un k -isomorfismo lineal $\text{Hom}_k(U_0, k) \simeq U_0$; en particular, $\dim U_0^* = \dim U_0$. Además, se cumple, para todo $i \in \{1, \dots, r\}$, que $\dim U_i^* = \dim U_0 - \dim U_i = \text{codim } U_i$.

Tomando como objetos los sistemas de r -subespacios y como morfismos los morfismos entre sistemas de r -subespacios, tenemos, para cada natural positivo fijo r , la categoría $(\mathcal{S}_r, k)$ de sistemas de r -subespacios sobre el campo k . Esta categoría es de Krull-Schmidt: es aditiva, todo sistema de r -subespacios o es indescomponible o se descompone en una suma directa de sistemas de r -subespacios indescomponibles y esta descomposición es única salvo isomorfismo y los anillos de endomorfismos asociados a los sistemas indescomponibles son anillos locales (una demostración puede consultarse en [Med20, págs. 25-29]). Para un entero positivo fijo r y un campo dado k , denotamos $\text{Ind}(\mathcal{S}_r, k)$ a un conjunto completo de representantes de las clases de isomorfismo de sistemas de r -subespacios indescomponibles sobre el campo k .

A cada sistema de r -subespacios $U = (U_0, U_1, \dots, U_r)$ le podemos asociar su *anillo de endomorfismos* $\text{End } U$ dado por el conjunto

$$\text{End } U = \{\mathcal{A} \mid \mathcal{A} \text{ es morfismo de } U \text{ en } U\},$$

con la estructura de anillo obtenida al considerar la suma y composición usuales de transformaciones k -lineales. Claramente, $\text{End } U$ es un subanillo del anillo $\text{End}_k(U_0)$ de los endomorfismos k -lineales de U_0 .

El problema central es entonces el obtener, para cada entero positivo r fijo y un campo k , la descripción de $\text{Ind}(\mathcal{S}_r, k)$, así como de los correspondientes anillos de endomorfismos.

Ya que estamos considerando espacios vectoriales de dimensión finita, tenemos la posibilidad de asignarles bases y de trabajar con las matrices asociadas, tal como veremos a continuación.

2.1. Presentaciones matriciales.

Una *presentación matricial* de un sistema de r -subespacios $U = (U_0, U_1, \dots, U_r)$ es una matriz con r franjas verticales

$$M_U = \begin{bmatrix} M_1 & \cdots & M_r \end{bmatrix}$$

en donde, para cada $i \in \{1, \dots, r\}$, las columnas de la franja M_i corresponden a las coordenadas (con respecto a una base fija para U_0) de un sistema minimal de generadores del espacio U_i . Dos presentaciones matriciales M_U, M_V son *equivalentes* si una de ellas puede transformarse en la otra mediante la aplicación de un número finito de las siguientes *transformaciones admitidas*:

- Transformaciones elementales con filas de toda la matriz.
- Transformaciones elementales con columnas dentro de cada franja vertical M_i .

Nótese que las transformaciones admitidas de tipo (a) corresponden a cambios de base para U_0 y las transformaciones admitidas de tipo (b) corresponden a cambios de base de los subespacios U_i lo cual nos lleva al siguiente resultado:

Proposición 2.1. *Dos sistemas de r -subespacios U y V son isomorfos si y solo si dos cualesquiera de sus presentaciones matriciales M_U y M_V son equivalentes.* $\square$

Si consideramos dos presentaciones matriciales

$$A = \begin{bmatrix} A_1 & \cdots & A_r \end{bmatrix} \quad \text{y} \quad B = \begin{bmatrix} B_1 & \cdots & B_r \end{bmatrix}$$

correspondientes a dos sistemas de r -subespacios U y V , respectivamente, a partir de las definiciones tenemos que una presentación matricial para $U \oplus V$ está dada por la matriz

$$A \oplus B = \begin{bmatrix} A_1 & 0 & \cdots & A_r & 0 \\ 0 & B_1 & \cdots & 0 & B_r \end{bmatrix}.$$

Teniendo en cuenta la proposición 2.1, resolver el problema central, es decir, clasificar los sistemas de r -subespacios indescomponibles equivale a resolver un *problema matricial*; es decir, resolver el problema fundamental equivale a hallar formas canónicas para las presentaciones matriciales utilizando las transformaciones admitidas. En la proposición 3.4 daremos un ejemplo de estos problemas matriciales.

Nota. Por supuesto, el problema central puede verse como un subproblema de la teoría de representaciones de conjuntos parcialmente ordenados ordinarios (sin estructura adicional) finitos ya que corresponde precisamente a obtener la clasificación de las representaciones indescomponibles sobre un campo dado k para una anticadena con r elementos. En este artículo, sin embargo, no asumimos familiaridad con esta teoría y utilizamos otro enfoque. Los lectores interesados en el enfoque desde la teoría de representaciones de conjuntos parcialmente ordenados, pueden consultar, por ejemplo [Arn00] o [Med20].

Hacemos ahora una observación simple de álgebra lineal que nos será útil:

Lema 2.2. *Cualquier matriz rectangular N puede llevarse, utilizando las transformaciones admitidas, a la forma siguiente (que denominaremos forma estándar):*

$$\begin{array}{c|c} 0 & I_n \\ \hline 0 & 0 \end{array}$$

en donde I_n es la matriz identidad de orden n (este n es precisamente el rango de N) y los bloques 0 representan matrices nulas; algunos bloques de la forma estándar pueden ser vacíos.

Demostración. Inicialmente utilizamos transformaciones admitidas de tipo (a), tal como se hace en el proceso de reducción de Gauss-Jordan, para obtener la matriz R en forma escalonada reducida asociada a N y finalmente, aplicando transformaciones admitidas de tipo (b), anulamos cualquier posible elemento no nulo a la derecha de los unos principales de R y luego permutamos apropiadamente las columnas hasta obtener una matriz en la forma estándar. $\square$

Decimos que un sistema de r -subespacios $U = (U_0, U_1, \dots, U_r)$ es *trivial* si $\dim U_0 = 1$ o, equivalentemente, si U_0 es (isomorfo a) k .

Lema 2.3. *Sea U un sistema trivial de r -subespacios. Entonces, U es indescomponible y $\text{End } U \simeq k$.*

Demostración. Es claro que si $U_0 \simeq k$, entonces necesariamente U es indescomponible, pues $U = V \oplus W$ con $\dim U_0 = 1$ implica que o bien $\dim V_0 = 1$ y $\dim W_0 = 0$ o bien $\dim V_0 = 0$ y $\dim W_0 = 1$. Por otro lado, la multiplicación por escalar induce una inyección $k \rightarrow \text{End } U$. Además, $\text{End } U \subseteq \text{End } U_0 = \text{End } k \simeq k$; con esto, $\text{End } U \simeq k$. $\square$

El siguiente resultado muestra que los anillos de endomorfismos permiten obtener información sobre los sistemas a los que están asociados:

Lema 2.4. *Sea U un sistema de r -subespacios sobre un campo k . Entonces, U es indescomponible si y solo si 0 y 1 son los únicos idempotentes del anillo $\text{End } U$.*

Demostración. Si $U = (U_0, U_1, \dots, U_r)$ es un sistema de r -subespacios sobre un campo k y $\mathcal{A} \in \text{End } U$ es un idempotente, veamos que $U = \mathcal{A}(U) \oplus (1 - \mathcal{A})(U)$; a tal efecto, vamos a probar que, para todo $i \in \{0, 1, \dots, r\}$, se tiene que $U_i = \mathcal{A}(U_i) \oplus (1 - \mathcal{A})(U_i)$; si $\alpha \in U_i$, entonces

$$\begin{aligned} \alpha &= \mathcal{A}(\alpha) + (\alpha - \mathcal{A}(\alpha)) \\ &= \mathcal{A}(\alpha) + (1 - \mathcal{A})(\alpha) \end{aligned}$$

que muestra que $U_i = \mathcal{A}(U_i) + (1 - \mathcal{A})(U_i)$. Para ver que la suma es directa, nos falta ver que la intersección de los sumandos es trivial: supongamos entonces que $\beta \in \mathcal{A}(U_i) \cap (1 - \mathcal{A})(U_i)$. Tenemos entonces que $\beta = \mathcal{A}(\gamma)$, para algún $\gamma \in U_i$

y también $\beta = (1 - \mathcal{A})(\delta) = \delta - \mathcal{A}(\delta)$, para algún $\delta \in U_i$. De aquí, usando el que $\mathcal{A}^2 = \mathcal{A}$, obtenemos

$$\begin{aligned} \beta &= \mathcal{A}(\gamma) = \mathcal{A}^2(\gamma) = \mathcal{A}(\mathcal{A}(\gamma)) \\ &= \mathcal{A}(\delta - \mathcal{A}(\delta)) = \mathcal{A}(\delta) - \mathcal{A}^2(\delta) \\ &= \mathcal{A}(\delta) - \mathcal{A}(\delta) = 0. \end{aligned}$$

Por lo tanto, si $\mathcal{A}$ es un idempotente en $\text{End } U$ y U es indescomponible, necesariamente $\mathcal{A}(U) = 0$ o $(1 - \mathcal{A})(U) = 0$; es decir, $\mathcal{A} = 0$ o $\mathcal{A} = 1$ son los únicos idempotentes de $\text{End } U$. Recíprocamente, supongamos que $U = V \oplus W$, en donde U, V son sistemas no nulos de r -subespacios; el morfismo $\mathcal{B}: U \rightarrow U$ definido por

$$\mathcal{B}(v, w) = (v, 0), \text{ para todo } (v, w) \in V \oplus W,$$

es entonces un idempotente de $\text{End } U$ distinto de 0 y de 1 pues el núcleo de $\mathcal{B}$ es un subespacio propio no trivial de U ($\ker \mathcal{B} \simeq W$). $\square$

Un pequeño ejemplo para ilustrar algunas de las ideas expuestas hasta el momento:

Ejemplo 2.5. *Dado un campo k , consideremos los sistemas de 3-subespacios siguientes:*

$$U = (k, k, 0, k) \quad \text{y} \quad V = (k, 0, 0, k).$$

Por el lema 2.3, sabemos que U y V son indescomponibles y que $\text{End } U \simeq k$ y también $\text{End } V \simeq k$. Escogiendo bases apropiadas, tenemos las presentaciones matriciales M_U y M_V :

$$M_U = \begin{bmatrix} 1 & 0 & 1 \end{bmatrix} \quad \text{y} \quad M_V = \begin{bmatrix} 0 & 0 & 1 \end{bmatrix}$$

Además, $M_U \oplus M_U \oplus M_V$ corresponde a

$$\begin{array}{ccc|ccc|ccc} 1 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{array}.$$

Si $U = (k, U_1, \dots, U_r)$ es un sistema trivial de subespacios sobre el campo k , entonces su sistema dual $U^* = (U_0^*, U_1^*, \dots, U_r^*)$ es (isomorfo a) $(k, W_1, \dots, W_r)$, en donde, para todo $i \in \{1, \dots, r\}$, se tiene que

$$W_i = \begin{cases} k, & \text{si } U_i = 0. \\ 0, & \text{si } U_i = k. \end{cases}$$

En particular, para los dos sistemas de 3-subespacios dados al comienzo de este ejemplo

$$U = (k, k, 0, k) \quad \text{y} \quad V = (k, 0, 0, k),$$

se tiene que sus duales son (isomorfos a)

$$U^* \simeq (k, 0, k, 0) \quad \text{y} \quad V^* \simeq (k, k, k, 0)$$

y las correspondientes dimensiones están dadas por

$$\begin{aligned} \underline{\dim} U &= (1, 1, 0, 1), & \underline{\dim} V &= (1, 0, 0, 1), \\ \underline{\dim} U^* &= (1, 0, 1, 0), & \underline{\dim} V^* &= (1, 1, 1, 0). \end{aligned}$$

3. Los sistemas de r -subespacios.

Después de los preliminares anteriores, vamos inicialmente a examinar el problema central para los casos $r \in \{1, 2, 3\}$ en las tres proposiciones siguientes; como veremos, en estos casos el problema de clasificación de los sistemas de r -subespacios indescomponibles es relativamente sencillo:

Proposición 3.1. *Dado un campo arbitrario k , tenemos que $\text{Ind}(\mathcal{S}_1, k) = \{(k, 0), (k, k)\}$ y, en ambos casos, los anillos de endomorfismos asociados son triviales.*

Demostración. Si $U = (U_0, U_1)$ es un sistema arbitrario de 1-subespacio, entonces, para algún complemento X de U_1 en U_0 se tiene que

$$\begin{aligned} U &= (U_0, U_1) \\ &= (U_1 \oplus X, U_1) \\ &= (X, 0) \oplus (U_1, U_1) \\ &\simeq (k^m, 0) \oplus (k^n, k^n) \\ &= (k, 0)^m \oplus (k, k)^n, \end{aligned}$$

en donde $X \simeq k^m$ y $U_1 \simeq k^n$, para algunos enteros no negativos m y n . De aquí se obtienen los sumandos indescomponibles del resultado y el lema 2.3 garantiza que los anillos de endomorfismos son triviales. $\square$

Nota. El procedimiento para llevar cualquier matriz a su forma estándar, descrito en el lema 2.2, corresponde a la solución del problema matricial asociado a los sistemas de 1-subespacio sobre un campo k ; de hecho, la forma estándar no es otra cosa que una suma directa de representaciones matriciales de 1-sistemas indescomponibles de la forma (k, k) y $(k, 0)$.

Proposición 3.2. *Dado un campo arbitrario k ,*

$$\text{Ind}(\mathcal{S}_2, k) = \{(k, 0, 0), (k, k, 0), (k, 0, k), (k, k, k)\}$$

y los cuatro anillos de endomorfismos asociados son triviales.

Demostración. Consideremos $U = (U_0, U_1, U_2)$ un sistema arbitrario de 2-subespacios sobre k y sea $W = U_1 \cap U_2$; entonces $U_1 = W \oplus X$ y $U_2 = W \oplus Y$, para algunos subespacios complementarios X, Y y además $U_1 + U_2 = W \oplus X \oplus Y$. Ya que $U_0 = (U_1 + U_2) \oplus Z$ para algún complemento Z , tenemos que

$$\begin{aligned} U &= (W, W, W) \oplus (X, X, 0) \oplus (Y, 0, Y) \oplus (Z, 0, 0) \\ &\simeq (k^m, k^m, k^m) \oplus (k^n, k^n, 0) \oplus (k^p, 0, k^p) \oplus (k^q, 0, 0) \\ &= (k, k, k)^m \oplus (k, k, 0)^n \oplus (k, 0, k)^p \oplus (k, 0, 0)^q, \end{aligned}$$

de donde obtenemos los sumandos indescomponibles anunciados. Nuevamente el lema 2.3 garantiza que los anillos de endomorfismos son triviales. $\square$

Con la ayuda del siguiente resultado auxiliar, en la proposición 3.4 obtendremos la clasificación completa de los sistemas de 3-subespacios indescomponibles, y de sus anillos de

endomorfismos asociados. En la prueba de la proposición haremos uso de las presentaciones matriciales; una solución en usando el lenguaje de los espacios vectoriales puede consultarse en [Arn00, Ejemplo 1.1.5, págs. 4–7]

Lema 3.3. *Para los sistemas de 3-subespacios (U_0, U_1, U_2, U_3) con las relaciones adicionales $U_1 \subseteq U_3$ y $U_2 \subseteq U_3$ se tiene que todos los objetos indescomponibles vienen dados, salvo isomorfismo, por los elementos del conjunto*

$$\{(k, 0, 0, 0), (k, k, 0, k), (k, 0, k, k), (k, 0, 0, k), (k, k, k, k)\}$$

y todos los anillos de endomorfismos asociados son triviales.

Demostración. Consideremos $U = (U_0, U_1, U_2, U_3)$ un sistema de 3-subespacios sobre k con las relaciones adicionales $U_1 \subseteq U_3$ y $U_2 \subseteq U_3$ y sea $V = U_1 \cap U_2$; entonces $U_1 = V \oplus W$ y $U_2 = V \oplus X$, para algunos subespacios complementarios W, X y además $U_1 + U_2 = V \oplus W \oplus X$. Ya que $U_3 = (U_1 + U_2) \oplus Y$ y $U_0 = U_3 \oplus Z$ para algunos complementos Y y Z , tenemos que

$$\begin{aligned} U &= (V, V, V, V) \oplus (W, W, 0, W) \oplus (X, 0, X, X) \\ &\quad \oplus (Y, 0, 0, Y) \oplus (Z, 0, 0, 0), \end{aligned}$$

de donde obtenemos los sumandos indescomponibles anunciados. Ya que todos los objetos indescomponibles son triviales, el lema 2.3 garantiza que sus anillos de endomorfismos también lo son. $\square$

Proposición 3.4. *Dado un campo arbitrario k ,*

$$\begin{aligned} \text{Ind}(\mathcal{S}_3, k) &= \{(k, 0, 0, 0), (k, k, 0, 0), (k, 0, k, 0), (k, 0, 0, k), \\ &\quad (k, k, k, 0), (k, k, 0, k), (k, 0, k, k), (k, k, k, k), U\}, \end{aligned}$$

en donde U es (isomorfo a) el sistema de 3-subespacios dado de la manera siguiente:

$$\begin{aligned} U_0 &= k \oplus k, & U_2 &= 0 \oplus k, \\ U_1 &= k \oplus 0, & U_3 &= \{(x, x) \mid x \in k\}. \end{aligned} \quad (1)$$

y los nueve anillos de endomorfismos asociados son triviales.

Demostración. Consideremos una presentación matricial

$$M = \begin{bmatrix} M_1 & M_2 & M_3 \end{bmatrix}$$

de un sistema de 3-subespacios $U = (U_0, U_1, U_2, U_3)$. Inicialmente, aplicando transformaciones admitidas, llevamos la franja M_1 a la forma estándar (ver el lema 2.2), obteniendo así una matriz de la forma siguiente (los colores son utilizados únicamente como guía durante el proceso de reducción):

$$\begin{array}{c} \overbrace{M_1} \quad M_2 \quad M_3 \\ \begin{bmatrix} 0 & I & A & B \\ 0 & 0 & C & D \end{bmatrix} \end{array}$$

La franja horizontal inferior de esta matriz corresponde a cierta presentación matricial de un sistema de 2-subespacios ya

que aplicar en esta franja transformaciones admitidas de tipo (a) no altera los bloques nulos correspondientes a M_1 y, por otro lado, se tienen transformaciones admitidas independientes de tipo (b) para los bloques C y D que dejan invariante el bloque M_1 ; utilizando el resultado de la proposición 3.2, en los bloques $\begin{bmatrix} C & D \end{bmatrix}$ colocamos bloques correspondientes a sumandos de los sistemas indescomponibles de 2-subespacios $(k, k, 0)$, $(k, 0, k)$, (k, k, k) y $(k, 0, 0)$ y utilizamos adiciones de filas para anular en A y en B los elementos por encima de los bloques I de $(k, k, 0)$, $(k, 0, k)$ y de uno de los bloques I de (k, k, k) , para obtener la siguiente matriz (de ahora en adelante, los bloques que no aparezcan marcados en las matrices representan bloques nulos):

M_1			M_2			M_3		
	I	P				Q	R	
				I				
								I
				I			I	

Analizando ahora cuáles transformaciones admitidas no cambian los bloques I ni los bloques nulos de la matriz anterior, resulta que, aparte de las adiciones obvias de columnas de Q a columnas de R , tenemos ahora también adiciones de columnas de P a columnas de R (estas adiciones están indicadas por las flechas curvas en la matriz anterior): en efecto, cualquier columna de P puede adicionarse al bloque nulo justo a la derecha de P ; al restaurar los ceros de este bloque adicionamos filas del bloque I asociado a (k, k, k) , pero entonces también se deben realizar las mismas adiciones del otro bloque I de (k, k, k) al bloque R . Por lo tanto, los bloques $\begin{bmatrix} P & Q & R \end{bmatrix}$ corresponden a una presentación matricial de cierto sistema de 3-subespacios $U = (U_0, U_1, U_2, U_3)$ con relaciones adicionales $U_1 \subseteq U_3$ y $U_2 \subseteq U_3$. Gracias al lema 3.3, sabemos que los indescomponibles de este sistema con relaciones son $(k, k, 0, k)$, $(k, 0, k, k)$, $(k, 0, 0, k)$, (k, k, k, k) y $(k, 0, 0, 0)$; colocamos entonces, en lugar de los bloques $\begin{bmatrix} P & Q & R \end{bmatrix}$ bloques I correspondientes a sumandos de los anteriores sistemas y

obtenemos, finalmente, la siguiente matriz:

	M_1				M_2				M_3			
1		I				I				I		
2			I								I	
3				I								I
4					I					I		
5						I					I	
6							I					
7								I				
8									I			
9										I		
10											I	
11												I
12												

Con esto, tenemos que los sistemas de 3-subespacios indescomponibles son $(k, k, k, 0)$ (obtenido de la primera, de abajo hacia arriba, franja horizontal de la matriz), $(k, k, 0, k)$ (obtenido de la segunda franja horizontal de la matriz), (k, k, k, k) (obtenido de la cuarta franja horizontal de la matriz), $(k, k, 0, 0)$ (obtenido de la quinta franja horizontal de la matriz), $(k, 0, k, 0)$ (obtenido de la sexta y séptima franjas horizontales de la matriz), $(k, 0, 0, k)$ (obtenido de la octava y novena franjas horizontales de la matriz), $(k, 0, k, k)$ (obtenido de la décima franja horizontal de la matriz), $(k, 0, 0, 0)$ (obtenido de la última franja horizontal de la matriz) y el sistema no trivial con presentación matricial

$$\begin{bmatrix} 1 & 0 & 1 \\ 0 & 1 & 1 \end{bmatrix} \quad (2)$$

obtenido de la tercera y undécima franjas horizontales de la matriz.

Gracias al lema 2.3, tenemos que los anillos de endomorfismos asociados a los ocho 3-sistemas triviales de subespacios son también triviales; para terminar la demostración, examinemos el anillo de endomorfismos del sistema indescomponible no trivial: la presentación matricial dada en (2) corresponde al sistema de 3-subespacios $U = (U_0, U_1, U_2, U_3)$, en donde

$$\begin{aligned} U_0 &= k \oplus k, & U_2 &= 0 \oplus k, \\ U_1 &= k \oplus 0, & U_3 &= \{(x, x) \mid x \in k\}. \end{aligned}$$

Si hacemos $\varepsilon_1 = (1, 0)$ y $\varepsilon_2 = (0, 1)$, tenemos que los conjuntos $\{\varepsilon_1, \varepsilon_2\}$, $\{\varepsilon_1\}$, $\{\varepsilon_2\}$ y $\{\varepsilon_1 + \varepsilon_2\}$ son bases para los subespacios U_0 , U_1 , U_2 y U_3 , respectivamente. Ahora bien, si $\mathcal{A}: U \rightarrow U$ es un endomorfismo de U , la condición $\mathcal{A}(U_1) \subseteq U_1$ implica que $\mathcal{A}(\varepsilon_1) = a\varepsilon_1$, para algún $a \in k$; la condición $\mathcal{A}(U_2) \subseteq U_2$ implica que $\mathcal{A}(\varepsilon_2) = b\varepsilon_2$, para algún $b \in k$ y la condición $\mathcal{A}(U_3) \subseteq U_3$ implica que $\mathcal{A}(\varepsilon_1 + \varepsilon_2) = c(\varepsilon_1 + \varepsilon_2)$, para algún $c \in k$, pero entonces

$$\begin{aligned} c\varepsilon_1 + c\varepsilon_2 &= c(\varepsilon_1 + \varepsilon_2) \\ &= \mathcal{A}(\varepsilon_1 + \varepsilon_2) \\ &= \mathcal{A}(\varepsilon_1) + \mathcal{A}(\varepsilon_2) \\ &= a\varepsilon_1 + b\varepsilon_2, \end{aligned}$$

con lo que $a = b = c$; es decir, $\mathcal{A}$ es multiplicación por un escalar y esto nos lleva a que $\text{End} U \simeq k$ (teniendo en cuenta el lema 2.4, esto también reafirma el que este sistema de 3-subespacios es indescomponible). $\square$

Vale la pena notar que todos los indescomponibles para los sistemas de 1 y 2 subespacios son triviales, pero cuando pasamos a los sistemas de 3-subespacios ya aparece un indescomponible no trivial. En todos los casos, sin embargo, se tiene que $\dim U_0 \leq 2$ y todos los anillos de endomorfismos asociados son triviales. Como vamos a ver, cuando consideramos los sistemas de r -subespacios para $r \geq 4$ la situación cambia dramáticamente: el conjunto $\text{Ind}(\mathcal{S}_r, k)$ resulta infinito, para todo $r \geq 4$; además, aparecen sistemas indescomponibles en los que $\dim U_0$ puede ser tan grande como se quiera y también aparecen anillos de endomorfismos no triviales. Primero, un lema auxiliar:

Lema 3.5. *Para enteros positivos r y s , con $r \leq s$, se tiene una inyección $\iota: \text{Ind}(\mathcal{S}_r, k) \rightarrow \text{Ind}(\mathcal{S}_s, k)$.*

Demostración. Para $U = (U_0, U_1, \dots, U_r) \in \text{Ind}(\mathcal{S}_r, k)$, definamos $\iota(U) = \bar{U} = (\bar{U}_0, \bar{U}_1, \dots, \bar{U}_s)$ haciendo

$$\begin{aligned} \bar{U}_0 &= U_0 \quad \text{y, para } r \in \{1, \dots, s\}, \\ \bar{U}_i &= \begin{cases} U_i, & \text{si } i \in \{1, \dots, r\}. \\ 0, & \text{si } i \in \{r+1, \dots, s\}. \end{cases} \end{aligned}$$

Es inmediato, a partir de la definición dada, que ι es una inyección y que $\iota(U)$ es indescomponible si U lo es. $\square$

Proposición 3.6. *Para $r \geq 4$, el conjunto $\text{Ind}(\mathcal{S}_r, k)$ es infinito.*

Demostración. Dado un campo k , consideremos el sistema de 4-subespacios $U = (U_0, U_1, \dots, U_4)$ tal que, para cada entero $n \geq 0$, el k -espacio vectorial U_0 es $k^{n+1} \oplus k^{n+1}$ de dimensión $2n+2$ con base $\{\epsilon_1, \dots, \epsilon_{2n+2}\}$ y, para $i \in \{1, \dots, 4\}$, los subespacios U_i están dados de la siguiente manera:

$$\begin{aligned} U_1 &= k^{n+1} \oplus 0 = \langle \{\epsilon_1, \dots, \epsilon_{n+1}\} \rangle, \\ U_2 &= 0 \oplus k^{n+1} = \langle \{\epsilon_{n+2}, \dots, \epsilon_{2n+2}\} \rangle, \\ U_3 &= \langle \{\epsilon_1 + \epsilon_{n+2}, \dots, \epsilon_{n+1} + \epsilon_{2n+2}\} \rangle, \\ U_4 &= \langle \{\epsilon_2 + \epsilon_{n+2}, \dots, \epsilon_{n+1} + \epsilon_{2n+1}\} \rangle. \end{aligned}$$

Si $\mathcal{A} \in \text{End} U$, debe tenerse que $\mathcal{A}(U_i) \subseteq U_i$, para todo $i \in \{1, \dots, 4\}$, y esto nos lleva, después de un análisis sobre el efecto de $\mathcal{A}$ sobre los generadores de los espacios U_i , a que $\mathcal{A}$ es simplemente multiplicación por un elemento fijo de k . Tenemos así que $\text{End} U \simeq k$ y, en particular, los únicos idempotentes de $\text{End} U$ son 0 y 1; usando el lema 2.4, concluimos que U es indescomponible, para todo $n \geq 0$. Tenemos así que $\text{Ind}(\mathcal{S}_4, k)$ es un conjunto infinito y el lema 3.5 nos permite entonces concluir que $\text{Ind}(\mathcal{S}_r, k)$ es también un conjunto infinito, para todo $r \geq 5$. $\square$

Vamos ahora a examinar la situación para los sistemas de r -subespacios para $r \geq 5$; como veremos, en estos casos el problema es muy interesante pero surge una dificultad insalvable:

Proposición 3.7. *Para $r \geq 5$, el problema de clasificación de los sistemas indescomponibles de r -subespacios es un problema de tipo salvaje.*

Demostración. Teniendo en cuenta el lema 3.5, basta probar el resultado para $r = 5$: dado un campo k , consideremos el sistema de 5-subespacios $U = (U_0, U_1, \dots, U_5)$ dado por:

$$\begin{aligned} U_0 &= E \oplus E, & U_3 &= \{(\alpha, \alpha) \mid \alpha \in E\}, \\ U_1 &= E \oplus 0, & U_4 &= \{(\alpha, \mathcal{A}_1(\alpha)) \mid \alpha \in E\}, \\ U_2 &= 0 \oplus E, & U_5 &= \{(\alpha, \mathcal{A}_2(\alpha)) \mid \alpha \in E\}, \end{aligned}$$

en donde E es un espacio vectorial de dimensión finita n sobre k y $\mathcal{A}_1, \mathcal{A}_2$ son transformaciones k -lineales de E en E . Seleccionando bases apropiadas para U_0 y para cada uno de los subespacios U_i , obtenemos para U una presentación matricial M_U que tiene el aspecto siguiente:

$$M_U = \begin{array}{|c|c|c|c|c|} \hline I_n & 0 & I_n & I_n & I_n \\ \hline 0 & I_n & I_n & A_1 & A_2 \\ \hline \end{array},$$

en donde A_1 y A_2 son las matrices de los operadores $\mathcal{A}_1$ y $\mathcal{A}_2$, respectivamente.

Supongamos ahora que aplicamos una transformación admitida T sobre columnas de la franja vertical que contiene a A_1 ; esta operación cambia A_1 por cierta matriz B_1 y modifica el bloque I_n que está por encima de A_1 ; para restaurar el bloque identidad debemos aplicar ahora la operación inversa T^{-1} sobre filas de la primera franja horizontal; esto corrige el bloque identidad anterior pero modifica otros tres bloques identidad. En el diagrama de la figura 1 se ilustra la secuencia de transformaciones admitidas que deben aplicarse para restaurar los bloques identidad originales: cada transformación admitida se representa mediante una flecha horizontal, si es sobre filas, o vertical, si es sobre columnas; un bloque identidad en color rojo indica, que al aplicar la operación correspondiente, ese bloque sufre modificaciones.

Así pues, después de recomponer los bloques identidad, tenemos que aplicar la transformación T sobre las columnas de A_1 implica aplicar esa misma operación sobre las mismas columnas de A_2 y también aplicar la transformación inversa T^{-1} sobre filas tanto de A_1 como de A_2 ; de manera completamente análoga puede verse que algo similar sucede si se aplica inicialmente una transformación admitida sobre las columnas de A_2 . Teniendo en cuenta que aplicar transformaciones admitidas sobre columnas (sobre filas, respectivamente) equivale a multiplicación a derecha (a izquierda, respectivamente) por una matriz invertible, el análisis anterior nos muestra que las matrices A_1 y A_2 se están transformando, de manera simultánea, por multiplicación a derecha por cierta matriz invertible P y por multiplicación a izquierda precisamente por la matriz inversa P^{-1} .

Así pues, el sistema de 5-subespacios considerado corresponde al problema de la reducción simultánea, por semejanza, de un par de matrices cuadradas del mismo orden:

$$(A_1, A_2) \mapsto (P^{-1}A_1P, P^{-1}A_2P),$$

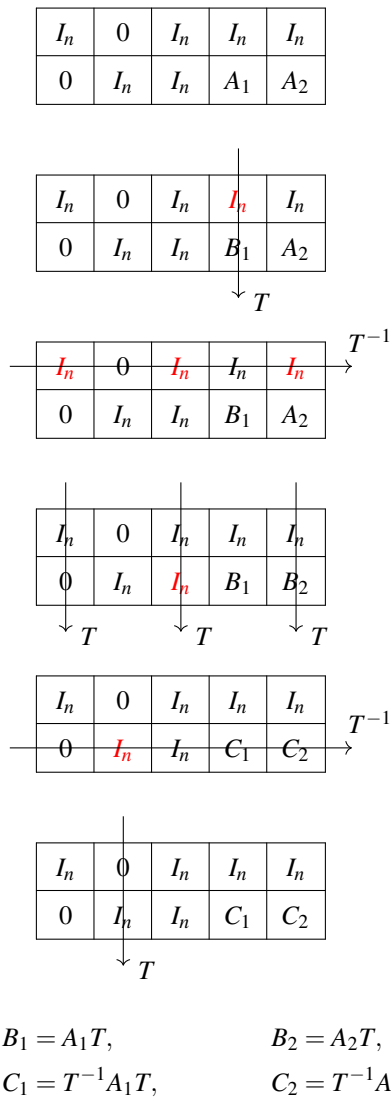


Figura 1: sucesión de transformaciones que ilustran el problema de la reducción simultánea de un par de matrices por semejanza como caso especial de un sistema de 5-subespacios.

en donde P es una matriz invertible del mismo orden que A_1 y A_2 . Este problema, que está asociado al carcaj



es de tipo salvaje. $\square$

El que el problema de la clasificación de los sistemas de r -subespacios para $r \geq 5$ sea un problema de tipo salvaje significa que, para $r \geq 5$, el conjunto $\text{Ind}(\mathcal{S}_r, k)$ contiene una copia de todos los R -módulos indecomponibles finito dimensionales para cada k -álgebra R de dimensión finita [Arn00, Bre74b],

así que es prácticamente imposible pensar en obtener una clasificación completa de los sistemas indecomponibles para $r \geq 5$.

Nota. Aunque hemos pedido que el entero r sea positivo, de hecho, podemos dar también la clasificación de los sistemas de 0-subespacios: en este caso, los sistemas tienen la forma (U_0) , en donde U_0 un espacio vectorial de dimensión finita sobre un campo k y es inmediato que, salvo isomorfismo, (k) es el único objeto indecomponible y que su anillo de endomorfismos es trivial.

4. Sistemas de 4-subespacios.

En esta sección vamos a centrarnos exclusivamente en el caso de los sistemas de 4-subespacios. Este caso es especialmente interesante por varias razones: por un lado, no es para nada sencillo (a diferencia de los casos para $r \in \{1, 2, 3\}$) y, por otro lado, aunque sabemos ya que en $\text{Ind}(\mathcal{S}_4, k)$ hay infinitos sistemas indecomponibles (ver la proposición 3.6), a diferencia de lo que sucede para $r \geq 5$, los infinitos sistemas en $\text{Ind}(\mathcal{S}_4, k)$ pueden ser descritos completamente así como sus anillos de endomorfismos. Adicionalmente, el problema de los cuatro subespacios contiene, como casos particulares, varios problemas clásicos e importantes, tal como ejemplificamos en la subsección siguiente.

Nota. En términos precisos, mientras que el problema de la clasificación de $\text{Ind}(\mathcal{S}_r, k)$, para $r \geq 5$, es un problema de tipo salvaje, el problema para $\text{Ind}(\mathcal{S}_4, k)$ es de tipo manso. Intuitivamente, el que un problema de clasificación de objetos sea salvaje significa que no admite solución completa y que el problema sea manso significa que, aunque involucra infinitos objetos, estos pueden ser clasificados completamente. Las definiciones precisas de mansedumbre y salvajismo, así como algunos de los resultados clásicos al respecto pueden consultarse en [Dro80].

4.1. El problema de los cuatro subespacios.

Antes de examinar el problema y su solución, mencionemos un par de cuestiones clásicas de álgebra lineal que aparecen como casos especiales del problema de la clasificación de los sistemas de 4-subespacios:

1. Sea $\mathcal{A}: E \rightarrow E$ una transformación lineal sobre un k -espacio vectorial E de dimensión finita sobre k . Consideremos la cuádrupla $U = (U_0, U_1, \dots, U_4)$, en donde

$$\begin{aligned} U_0 &= E \oplus E & U_3 &= \{(\alpha, \alpha) \mid \alpha \in E\}, \\ U_1 &= E \oplus 0, & U_4 &= \{(\alpha, \mathcal{A}(\alpha)) \mid \alpha \in E\}. \\ U_2 &= 0 \oplus E, \end{aligned}$$

Escogiendo bases apropiadas, tenemos para U la presentación matricial siguiente:

$$\begin{pmatrix} I_n & 0 & I_n & I_n \\ 0 & I_n & I_n & A \end{pmatrix},$$

en donde $n = \dim E$ y A es la matriz cuadrada de orden n asociada al operador lineal $\mathcal{A}$. Al hacer el análisis de las transformaciones admitidas que se pueden aplicar sin alterar los bloques identidad y nulos en la presentación matricial anterior, resulta que el bloque A se reduce por semejanza:

$$A \mapsto P^{-1}AP,$$

en donde P es una matriz invertible de orden n . Este problema es precisamente el problema del bucle



cuyas soluciones son las formas canónicas: por ejemplo, la de Jordan, introducida en 1870 en [Jor70], y la racional o de Frobenius.

2. Sean E y F un par de k -espacios vectoriales de dimensión finita sobre k y consideremos dos transformaciones k -lineales $\mathcal{A}_1: E \rightarrow F$ y $\mathcal{A}_2: E \rightarrow F$. Tenemos la cuádrupla asociada $U = (U_0, U_1, \dots, U_4)$, en donde

$$\begin{aligned} U_0 &= E \oplus F & U_3 &= \{(\alpha, \mathcal{A}_1(\alpha)) \mid \alpha \in E\}, \\ U_1 &= E \oplus 0, & U_4 &= \{(\alpha, \mathcal{A}_2(\alpha)) \mid \alpha \in E\}, \\ U_2 &= 0 \oplus F, \end{aligned}$$

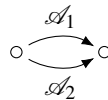
Después de escoger bases apropiadas para los espacios, resulta que a U le corresponde una presentación matricial que tiene la forma siguiente:

$$\begin{array}{|c|c|c|c|} \hline I_n & 0 & I_n & I_n \\ \hline 0 & I_m & A_1 & A_2 \\ \hline \end{array},$$

en donde $n = \dim E$, $m = \dim F$ y A_1, A_2 son las matrices de tamaño $m \times n$ asociadas a las transformaciones $\mathcal{A}_1, \mathcal{A}_2$, respectivamente. Analizando ahora cuáles de las transformaciones admitidas se pueden aplicar dejando invariantes los bloques identidad y nulos en la presentación matricial anterior, resulta que el par de matrices cuadradas A_1 y A_2 se reduce por equivalencia:

$$(A_1, A_2) \mapsto (PA_1Q, PA_2Q),$$

en donde P y Q son matrices invertibles de orden m y n , respectivamente. Este problema es el problema del haz de Kronecker:



que Weierstrass resolvió inicialmente para el caso regular en [Wei68] y, posteriormente, Kronecker extendió al caso singular (no regular) en [Kro90]. Una solución más reciente, con una extensión a los casos semilineal y pseudolineal puede consultarse en [Zav07, Zav08]

Nota. De hecho, el problema de los cuatro subespacios contiene, como casos especiales, muchos otros problemas clásicos e importantes de álgebra lineal; el autor está preparando un artículo en el que se incluye un estudio detallado de esta situación.

Dado un campo k , los sistemas de 4-subespacios sobre k se denominan también cuádruplas (de subespacios) sobre k . El problema de los cuatro subespacios no es otra cosa que el problema de obtener la clasificación completa, salvo isomorfismo, de todas las cuádruplas indescomponibles sobre un campo k .

4.2. Clasificación de las cuádruplas indescomponibles y de sus anillos de endomorfismos.

En la presentación de la clasificación completa de las cuádruplas indescomponibles y de los anillos de endomorfismos correspondientes seguimos las ideas de [MZ04]; allí puede consultarse una demostración de los dos teoremas siguientes:

Teorema 4.1. *Todas las cuádruplas indescomponibles $U = (U_0, U_1, \dots, U_4)$ sobre un campo arbitrario k están dadas, salvo isomorfismo, salvo dualidad y salvo permutaciones de los subespacios $U_1, \dots, U_4$, por las cuádruplas indescomponibles de los seis tipos $0, I, \dots, V$ presentados en forma matricial en los cuadros de la figura 2.*

Teorema 4.2. *Sea U una cuádrupla indescomponible sobre un campo arbitrario k , sean $d_0 = \dim U_0$ y $\mathfrak{E} = \text{End } U$ el anillo de endomorfismos de U . Entonces, se tiene lo siguiente:*

- (a) *Si U es regular de tipo 0 con $d_0 = 2n$ ($n \geq 1$), entonces $\mathfrak{E} \simeq k[t]/(p^s(t))$, en donde $p^s(t)$ es el polinomio minimal de la célula de Frobenius indescomponible X de orden n (mostrada en la figura 2) con $p(t)$ mónico irreducible y diferente de t y de $t-1$ (claramente $n = s \cdot \deg p(t)$).*
- (b) *Si U es regular de tipo I con $d_0 = 2n$ ($n \geq 1$), entonces $\mathfrak{E} \simeq k[t]/(t^n)$.*
- (c) *Si U es regular de tipo II con $d_0 = 2n+1$ ($n \geq 0$), entonces $\mathfrak{E} \simeq k[t]/(t^{n+1})$.*
- (d) *Si U es no regular, entonces $\mathfrak{E} \simeq k$.*

4.3. Cuadro de las cuádruplas indescomponibles.

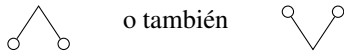
En esta sección explicamos las notaciones, las convenciones y toda la información contenida en los cuadros de la figura 2 en que se muestran todas las cuádruplas indescomponibles, salvo isomorfismo y permutaciones de los subespacios.

Si una matriz identidad I_k está equipada con una flecha de la forma $\leftarrow, \rightarrow, \uparrow$ o $\downarrow$, esto indica que debemos agregar a I_k una columna o fila de ceros a la izquierda, a la derecha, encima o abajo, respectivamente.

$F_n(p^s(t))$ denota la célula de Frobenius (también conocida como la célula de la forma canónica racional) de orden n que tiene como polinomio minimal $p^s(t)$, en donde $p(t)$ es mónico e irreducible. En otras palabras, $F_n(p^s(t))$ es la matriz acompañante de $p^s(t)$ (en particular, $n = s \cdot \deg p(t)$).

Denotamos $J_n^+(0)$ ($J_n^-(0)$) el bloque de Jordan de orden n con valor propio 0 y entradas 1 en la diagonal que está por encima (por debajo) de la diagonal principal.

Cada tipo de cuádrupla indescomponible viene acompañado con un útil diagrama característico formado por cuatro vértices (que representan los subespacios $U_1, \dots, U_4$) junto con símbolos adicionales de la forma



Estos símbolos significan que los subespacios correspondientes U_i y U_j satisfacen la relación $\dim(U_i \cap U_j) = 1$ o $\text{codim}(U_i + U_j) = 1$, respectivamente. En caso de ausencia de este símbolo, $U_i \cap U_j = 0$ o $U_i + U_j = U_0$, respectivamente. Los diagramas también incluyen las coordenadas del vector dimensión $d = \dim U$.

En el cuadro incluimos también los valores

$$\partial = \partial(d) = \sum_{i=1}^4 d_i - 2d_0$$

del defecto ∂ y

$$f = f(d) = d_0^2 + \sum_{i=1}^4 d_i^2 - d_0 \sum_{i=1}^4 d_i$$

de la forma cuadrática de Tits f en el vector de dimensión $d = \dim U$. Asimismo, damos para cada tipo su correspondiente anillo de endomorfismos $\mathfrak{E} = \text{End } U$.

Para un tipo T denotamos T^* el tipo dual (formado por las cuádruplas duales U^* en donde U es de tipo T).

En las matrices de tipos V y V^* la última franja horizontal es únicamente de una fila.

Nota. (a) De la construcción usada en [MZ04], se sigue que todas las cuádruplas listadas en el cuadro son, de hecho, indescomponibles y mutuamente no isomorfas. Además, cada cuádrupla U de uno de los tipos II, ..., V (o de los tipos duales) está determinada de manera unívoca, salvo isomorfismo, por su vector de dimensión $d = \dim U$ (que es una raíz de la forma f en el sentido que $f(d) = 1$). Cada cuádrupla de tipo 0 o I está determinada por el par $(d, p(t))$ o (d, t) respectivamente (en donde $p(t)$ y t son los polinomios mencionados antes y el vector de dimensión d es una raíz imaginaria de f en el sentido que $f(d) = 0$).

- (b) Si tomamos $p(t) = t$ o $p(t) = t - 1$ en el tipo 0, obtenemos el tipo I. Las cuádruplas de tipo I dadas por matrices con $J_n^+(0)$ y $J_n^-(0)$, son isomorfas.
- (c) Si sustituimos $J_n^+(0)$ por $J_n^-(0)$ en la forma de tipo V, debemos simultáneamente sustituir $(0 \dots 01)$ por $(10 \dots 0)$ en las últimas filas de las franjas M_3, M_4 (las últimas filas

en M_1, M_2 son filas nulas). Correspondientemente, en la forma V^* podemos invertir la dirección de todas las flechas, y simultáneamente sustituir $(0 \dots 01)$ por $(10 \dots 0)$ en la última fila de cada franja M_i , $i \in \{1, \dots, 4\}$.

Referencias.

- [Arn00] D.M. Arnold, *Abelian groups and representations of finite partially ordered sets*, Vol. 2, CMS Books in Mathematics, Canada, 2000.
- [Bre67] S. Brenner, *Endomorphism algebras of vector spaces with distinguished sets of subspaces*, J. Algebra **6** (1967), 100–114.
- [Bre74a] ———, *On four subspaces of a vector space*, J. Algebra **29** (1974), 587–599.
- [Bre74b] ———, *Decomposition properties of some small diagrams of modules*, 1974.
- [Dro80] Ju. A. Drozd, *Tame and wild matrix problems*, in Representation Theory II (Vlastimil Dlab and Gabriel, ed.), Springer, Berlin, Heidelberg, 1980.
- [For08] T.A. Forbregd, *The 4 Subspace Problem: Advisor: Sverre Smalø*, Master's Thesis, Norwegian University of Science and Technology, 2008., t. 1y/28hz.
- [Jor70] C. Jordan, *Traité des substitutions et des équations algébriques*, Paris, 1870.
- [GP70] I.M. Gelfand and V.A. Ponomarev, *Problems of linear algebra and classification of quadruples of subspaces in a finite dimensional vector space*, Colloq. Math. Soc. Janos Bolyai, Hilbert space operators, 1970, pp. 163–237.
- [Kro90] L. Kronecker, *Algebraische Reduktion der Scharen bilinearer Formen*, Sitzungsber. Akad. Berlin (1890), 763–776.
- [Med20] G. Medina, *Introducción a la teoría de representaciones de conjuntos parcialmente ordinarios*, Notas de clase, Departamento de Matemáticas, 2020.
- [MZ04] G. Medina and A.G. Zavatskij, *The four subspace problem: An elementary solution*, Linear Algebra and its Applications **392** (2004), 11–23.
- [Naz67] L.A. Nazarova, *Representations of a tetrad*, Izv. Akad. Nauk. SSSR **31** (1967), 1361–1377. Traducción al inglés en Math. USSR Izv. 1, 1967.
- [Naz75] ———, *Partially ordered sets of infinite type*, Izv. Akad. Nauk. SSSR **39** (1975), 963–991. Traducción al inglés en Math. USSR Izv. 9, 1975.
- [Wei68] K. Weierstrass, *Zur Theorie der quadratischen und bilinearen Formen*, Monatsber. Akad. Wiss. (1868), 311–338.
- [Zav07] A.G. Zavatskij, *On the Kronecker Problem and related problems of Linear Algebra*, Linear Algebra and its Applications **425** (2007), no. 1, 26–62, DOI <https://doi.org/10.1016/j.laa.2007.03.011>.
- [Zav08] ———, *A matrix problem over a central quadratic skew field extension*, Linear Algebra and its Applications **428** (2008), 393–399.

Acerca del autor: desde 1999 Gonzalo es profesor del Departamento de Matemáticas y Estadística de la Universidad Nacional de Colombia; su área de trabajo es la teoría de representaciones de estructuras algebraicas. Amante del álgebra, de la tipografía y de la historia y epistemología de las matemáticas así como apasionado por la música de Bach, las novelas de Heinrich Böll y los cubos de Rubik.

Regulares:

$0 = 0^*$	$\partial = 0$	$f = 0$	$\mathfrak{E} \simeq k[t]/(p^s(t))$									
$\circ$ n $\circ$ n $\circ$ n $\circ$ n $d_0 = 2n$	<table><tr><td>I_n</td><td>0</td><td>I_n</td><td>X</td></tr><tr><td>0</td><td>I_n</td><td>I_n</td><td>I_n</td></tr></table> $n \geq 1$	I_n	0	I_n	X	0	I_n	I_n	I_n	$X = F_n(p^s(t))$ $p(t) \neq t, t-1$		
I_n	0	I_n	X									
0	I_n	I_n	I_n									
$I = I^*$	$\partial = 0$	$f = 0$	$\mathfrak{E} \simeq k[t]/(t^n)$									
$\circ$ n $\circ$ n $\circ$ n $\circ$ n $d_0 = 2n$	<table><tr><td>I_n</td><td>I_n</td><td>0</td><td>$J_n^+(0)$</td></tr><tr><td>0</td><td>I_n</td><td>I_n</td><td>I_n</td></tr></table> $n \geq 1$	I_n	I_n	0	$J_n^+(0)$	0	I_n	I_n	I_n	$\Pi = \Pi^*$ $\partial = 0$ $f = 0$ $\mathfrak{E} \simeq k[t]/(t^{n+1})$		
I_n	I_n	0	$J_n^+(0)$									
0	I_n	I_n	I_n									
			$n+1$ $\circ$ $n+1$ $\circ$ $\circ$ n $\circ$ n $d_0 = 2n+1$	<table><tr><td>I_{n+1}</td><td>I_{n+1}</td><td>$I_n^\downarrow$</td><td>0</td></tr><tr><td>0</td><td>$I_n^{\leftarrow}$</td><td>I_n</td><td>I_n</td></tr></table> $n \geq 0$	I_{n+1}	I_{n+1}	$I_n^\downarrow$	0	0	$I_n^{\leftarrow}$	I_n	I_n
I_{n+1}	I_{n+1}	$I_n^\downarrow$	0									
0	$I_n^{\leftarrow}$	I_n	I_n									

No regulares:

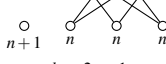
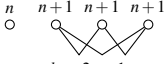
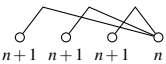
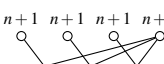
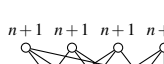
III	$\partial = -1$	$f = 1$	$\mathfrak{E} \simeq k$	III*	$\partial = 1$	$f = 1$	$\mathfrak{E} \simeq k$																				
 $d_0 = 2n + 1$	<table border="1"> <tr> <td>I_{n+1}</td><td>0</td><td>$I_n^\uparrow$</td><td>$I_n^\downarrow$</td></tr> <tr> <td>0</td><td>I_n</td><td>I_n</td><td>I_n</td></tr> </table> $n \geq 0$	I_{n+1}	0	$I_n^\uparrow$	$I_n^\downarrow$	0	I_n	I_n	I_n	 $d_0 = 2n + 1$	<table border="1"> <tr> <td>I_n</td><td>0</td><td>$I_{n+1}^\leftarrow$</td><td>$I_{n+1}^\rightarrow$</td></tr> <tr> <td>0</td><td>I_{n+1}</td><td>I_{n+1}</td><td>I_{n+1}</td></tr> </table> $n \geq 0$	I_n	0	$I_{n+1}^\leftarrow$	$I_{n+1}^\rightarrow$	0	I_{n+1}	I_{n+1}	I_{n+1}								
I_{n+1}	0	$I_n^\uparrow$	$I_n^\downarrow$																								
0	I_n	I_n	I_n																								
I_n	0	$I_{n+1}^\leftarrow$	$I_{n+1}^\rightarrow$																								
0	I_{n+1}	I_{n+1}	I_{n+1}																								
IV	$\partial = -1$	$f = 1$	$\mathfrak{E} \simeq k$	IV*	$\partial = 1$	$f = 1$	$\mathfrak{E} \simeq k$																				
 $d_0 = 2n + 2$	<table border="1"> <tr> <td>I_{n+1}</td><td>0</td><td>I_{n+1}</td><td>$I_n^\uparrow$</td></tr> <tr> <td>0</td><td>I_{n+1}</td><td>I_{n+1}</td><td>$I_n^\downarrow$</td></tr> </table> $n \geq 0$	I_{n+1}	0	I_{n+1}	$I_n^\uparrow$	0	I_{n+1}	I_{n+1}	$I_n^\downarrow$	 $d_0 = 2n + 2$	<table border="1"> <tr> <td>I_{n+1}</td><td>0</td><td>I_{n+1}</td><td>$I_{n+1}^\leftarrow$</td></tr> <tr> <td>0</td><td>I_{n+1}</td><td>I_{n+1}</td><td>$I_{n+1}^\rightarrow$</td></tr> </table> $n \geq 0$	I_{n+1}	0	I_{n+1}	$I_{n+1}^\leftarrow$	0	I_{n+1}	I_{n+1}	$I_{n+1}^\rightarrow$								
I_{n+1}	0	I_{n+1}	$I_n^\uparrow$																								
0	I_{n+1}	I_{n+1}	$I_n^\downarrow$																								
I_{n+1}	0	I_{n+1}	$I_{n+1}^\leftarrow$																								
0	I_{n+1}	I_{n+1}	$I_{n+1}^\rightarrow$																								
V	$\partial = -2$	$f = 1$	$\mathfrak{E} \simeq k$	V*	$\partial = 2$	$f = 1$	$\mathfrak{E} \simeq k$																				
 $d_0 = 2n + 1$	<table border="1"> <tr> <td>I_n</td><td>0</td><td>$J_n^+(0)$</td><td>I_n</td></tr> <tr> <td>0</td><td>I_n</td><td>I_n</td><td>$J_n^+(0)$</td></tr> <tr> <td>00...0</td><td>00...0</td><td>10...0</td><td>10...0</td></tr> </table> $n \geq 0$	I_n	0	$J_n^+(0)$	I_n	0	I_n	I_n	$J_n^+(0)$	00...0	00...0	10...0	10...0	 $d_0 = 2n + 1$	<table border="1"> <tr> <td>$I_n^{\leftarrow}$</td><td>$I_n^{\leftarrow}$</td><td>$I_n^{\rightarrow}$</td><td>0</td></tr> <tr> <td>0</td><td>$I_n^{\rightarrow}$</td><td>$I_n^{\leftarrow}$</td><td>$I_n^{\leftarrow}$</td></tr> <tr> <td>10...0</td><td>10...0</td><td>10...0</td><td>10...0</td></tr> </table> $n \geq 0$	$I_n^{\leftarrow}$	$I_n^{\leftarrow}$	$I_n^{\rightarrow}$	0	0	$I_n^{\rightarrow}$	$I_n^{\leftarrow}$	$I_n^{\leftarrow}$	10...0	10...0	10...0	10...0
I_n	0	$J_n^+(0)$	I_n																								
0	I_n	I_n	$J_n^+(0)$																								
00...0	00...0	10...0	10...0																								
$I_n^{\leftarrow}$	$I_n^{\leftarrow}$	$I_n^{\rightarrow}$	0																								
0	$I_n^{\rightarrow}$	$I_n^{\leftarrow}$	$I_n^{\leftarrow}$																								
10...0	10...0	10...0	10...0																								

Figura 2: Clasificación de todas las cuádruplas indecomponibles, salvo isomorfismo y permutaciones de los subespacios.

DE LA CUARTA DIMENSIÓN Y ESPACIOS PERFECTOS, A ROTAR UN ESPACIO TRIDIMENSIONAL

Daniel Esteban Galvis Sandoval *
daniele.galviss@konradlorenz.edu.co

1. Resumen.

El álgebra lineal, aunque se estudia a profundidad en matemáticas o física, a menudo parece tan abstracta que no se conoce claramente su grado de aplicación. En este artículo, se verá que los conceptos clave como transformaciones lineales y productos internos, llevarán a entender la lógica de la computación cuántica, un paradigma con una compleja estructura informática que pretende evolucionar la tecnología, y el futuro. Por eso, se introduce el espacio vectorial de los cuaterniones, y se interpreta como rotación de vectores en la esfera de radio 1, que se asocia en cierta forma con la esfera de Bloch. Esta última, entendida como la representación de estados de un sistema físico y en particular, de los qubits, que son los vectores del espacio de Hilbert cuyos estados vienen dados por un conjunto de transformaciones lineales o compuertas cuánticas. Adicionalmente, se presentará un pequeño código en Python con la idea de representar la rotación de vectores en la esfera de radio 1 del espacio de los cuaterniones. Por último, y con el fin de aplicar estas rotaciones, se explica brevemente el algoritmo de Shor para factorización de enteros positivos y su uso en la criptografía cuántica.

2. Introducción.

Desde que el matemático inglés Alan Turing diseñó el primer prototipo de ordenador clásico para descifrar el código secreto de la armada Nazi durante la segunda Guerra Mundial, el sistema computacional ha ido modernizándose, donde cada vez más, los ordenadores tienen mejores características de diseño, memoria, velocidad, entre otras, todas bajo el mismo sistema discreto cuya mínima unidad de información es el bit ¹ de unos y ceros. Estos dígitos representan dos estados: Encendido o apagado, cierto o falso, si o no [Lop16a]. Dentro de la circuitería electrónica de un sistema de computadora, estos valores significan presencia o ausencia de voltaje. Un bit es la mínima unidad de información, pues a partir de ellos se construye cantidades más grandes de información. Así, ocho bits conforman un octeto, también llamado byte [Lop16b], de modo que, como las computadoras están diseñadas para trabajar con bytes, y cada espacio del octeto tiene dos posibles

estados 1 o 0, entonces existen 2^8 combinaciones posibles, o sea 256 valores diferentes y se usan para medir el almacenamiento de memoria de un dato. Si digitamos un párrafo completo tendremos una serie de bytes que pueden tener un tamaño de un kilobyte (10^3 bytes) o un megabyte (10^6 bytes) etc., y en fracciones de segundo, el ordenador ya sabe lo que se le está escribiendo. Pero cuando el tamaño de las instrucciones al ordenador excede su capacidad, puede llegar a hacer que este se demore. Esto sucede mucho en matemáticas cuando se modela por ordenador problemas que requieren una enorme cantidad de cálculos. Pues bien, como la computadora clásica elabora estas instrucciones de forma serial, es decir de forma secuencial, paso a paso, una tras otra de forma independiente, entonces se abruma la memoria por el consumo de recursos.

A partir de esta dificultad, ha salido a la luz, la idea de construir un ordenador que realice operaciones de manera simultánea. En el año 2000, Enmanuel Knill, Raymond Laflamme, Rudy Martinez y Ching-Hua Tseng del MIT lograron desarrollar un computador cuántico de 7 Qubits [CC20]. Un Qubit, es la abreviación de Quantum bit, y representa la unidad mínima y básica de procesamiento de información de un computador cuántico y a comparación del bit, este puede estar en estado 0, 1 o en una combinación lineal de ambos. Más adelante, definiremos formalmente el qubit y su transformación en el espacio de Hilbert. El fin de esto, es reflejar la importancia de los estados que puede tomar un qubit en una operación dada. Con los bits convencionales, si teníamos un registro de tres bits, había ocho valores posibles, y el registro sólo podía tomar uno de esos valores. En cambio, si tenemos un vector de tres qubits, la partícula puede tomar ocho valores distintos a la vez gracias a lo que se denomina la superposición cuántica. De este modo, un vector de tres qubits computaría un total de ocho operaciones en paralelo. Para hacerse una idea del gran avance que esto supone, un computador cuántico de 30 qubits equivaldría a un procesador convencional de 10 teraflops (millones de millones de operaciones en coma flotante por segundo), es decir, una cantidad abrumadora de operaciones, que no es comparable con la capacidad de un procesador convencional de hoy día, ya que actualmente las computadoras trabajan en el orden de gigaflops (miles de millones de operaciones en coma flotante por segundo) [CC20]. En este artículo, se verá la forma de operar de los qubits mediante transformaciones lineales en un espacio de Hilbert, y cómo desde la perspectiva de las rotaciones

*Estudiante de matemáticas, Fundación Universitaria Konrad Lorenz, Bogotá-Colombia.

¹La palabra bit representa una abreviación de binary digit (dígito binario) [Lop16a]

en el espacio de cuaterniones, se puede interpretar las mismas que realiza en un espacio tridimensional de Hilbert. En la siguiente sección se generaliza una interpretación grosso modo de la física cuántica, y cómo esto se relaciona con el objeto de estudio de este escrito que son las rotaciones de un vector en un espacio complejo, posteriormente se introducen los cuaterniones, y luego los espacios de Hilbert. Finalmente, se presenta una aplicación en criptografía mediante el caso de factorización con el algoritmo de Shor.

3. Una nueva física.

De acuerdo con [Haw10] en su libro **El Gran Diseño**, el universo no tiene una existencia única sino que cada posible versión del universo existe simultáneamente en lo que se denomina la superposición cuántica; y con esta extraña introducción, nos preguntamos, ¿a quién se lo ocurrió hablar de algo como esto?, ¿por qué se habla de una nueva física del mundo subatómico?, ¿qué tiene que ver esto con la computación? y ¿por qué nos interesa para el estudio de rotaciones?

Bueno, hay que remontarse a comienzos del siglo XX, cuando el problema de saber en qué dirección o qué patrón seguía la emisión de luz de un cuerpo caliente, era irresoluble. Max Planck, descubrió la ley de la radiación electromagnética emitida por un cuerpo a cierta temperatura y dio solución a este problema pasando a ser considerado como el padre de la física cuántica [Anó19]. En este trabajo, intentó explicar que la luz se comportaba no solo como onda sino como corriente de partículas, que los electrones son simultáneamente partículas y ondas, y además, que los electrones ocupaban diferentes puntos de su órbita de manera simultánea saltando de una a otra, siendo imposible de conocer con precisión su trayectoria y llevando a pensar que era una cuestión del azar de la naturaleza.

Según [Anó20], el trabajo de Planck condujo a que Albert Einstein realizara su trabajo del efecto fotoeléctrico, basado en la idea del comportamiento de la luz no como onda continua sino como chorro de partículas denominadas cuantos de luz o fotones. En 1925, Werner Heisenberg desarrolló la matemática de la mecánica cuántica usando álgebra de matrices, le siguió el famoso físico austriaco Erwin Schrödinger que formuló en 1927 la función de onda y que condujo al famoso experimento del gato que se encuentra vivo y muerto a la vez debido a las propiedades cuánticas de la materia y propiamente de los electrones. Este último experimento, junto con el experimento de la doble rendija (que trataba de explicar el comportamiento en forma de onda de los electrones cuando interactuaba con un cristal de níquel) llevado a cabo en 1927 por Clinton Davisson y Lester Germer [Haw10] llevaron a la ciencia a otra escala y a pensar que ese comportamiento de onda de los fotones y electrones podían explicar el entrelazamiento cuántico. Esto motivó al físico Richard Feynman a estudiar el experimento de la doble rendija y a concluir que no se trataba de una posición indeterminada de la partícula sino que ella tomaba todos los caminos posibles a la vez, y ello depende de cuántas rendijas o compuertas se encuentren abier-

tas en dicho instante, es decir, que si una rendija se encuentra abierta, entonces todas las partículas pasaran por allí y harán una única historia, mientras que si hay más rendijas abiertas, entonces ellas pasaran por todas al mismo tiempo y harán historias alternativas. Se añade que el matemático y físico Paul Dirac, introdujo el proceso de unificación de la teoría de la mecánica cuántica con la teoría de la relatividad especial y en 1932 el matemático Húngaro John Von Neumann formuló la matemática para la mecánica cuántica mediante los espacios de Hilbert, y por esta razón, usamos notación de Dirac en las siguientes secciones.

Una vez sabemos de qué trata esta física, es importante cuestionarse por qué esto tuvo gran trascendencia en la computación del siglo XXI, pues bien, las ideas de entrelazamiento cuántico y número de posibles posiciones de una partícula se relacionan con el tamaño de las operaciones que realiza un ordenador y los recursos con los que dispone. Por lo tanto, si en la física cuántica una partícula que va del punto A a un punto B pasando por diferentes rendijas tomando todos los caminos posibles simultáneamente, en la computación cuántica, las operaciones de un algoritmo se realizan de forma simultánea y no se restringe a dos únicos estados de operación $(0,1)$ sino que pueden tener multitud de estados intermedios como resultado de la superposición de estos dos estados [Her10], es decir, se permite la evaluación de todos los posibles estados al mismo tiempo (algo que se conoce como paralelismo cuántico) y se traduce en una disminución de tiempo y recursos para el procesamiento de una operación [Cai10]. Por esta razón, consideramos que mediante el estudio de cuaterniones y el álgebra lineal, que es el caballito de batalla para entender esta computación, podríamos intentar explicar su funcionamiento a partir de simples rotaciones en un espacio tridimensional complejo.

4. Cuaterniones.

William Rowan Hamilton, un matemático nacido en Dublín en 1805, se preguntó por la forma de trasladar la noción geométrica de los números complejos del plano a dimensiones superiores, en particular, a la tercera dimensión. Lo cierto es que no pudo, pero cuando lo intentó con cuatro dimensiones le fue posible. De acuerdo con [Rod12] Hamilton se encontraba paseando con su esposa un 16 de Octubre de 1843 por el canal Real de Dublín, cuando de repente le llegó una idea a su cabeza, como si fuese un tipo de alucinación y para que no se le olvidase, escribió la identidad $i^2 = j^2 = k^2 = ijk = -1$ en una piedra del puente de Broom, donde $i, j, k \in \mathbb{C}$, es decir, forman un espacio tridimensional complejo, y la cuarta dimensión la denotó con una componente real a , por lo tanto, denominó a q como cuaternión si $q = a + bi + cj + dk$ donde $a, b, c, d \in \mathbb{R}$.

Tras este descubrimiento, el estudio de las ciencias evolucionó, en particular, en aplicaciones álgebras de Lie o matrices de Pauli, o incluso en informática, con las rotaciones en el espacio tridimensional usada en la robótica e inteligencia artificial.

4.1. Teoría.

Algebraicamente el conjunto de los complejos forma una estructura de cuerpo con ciertas propiedades y que serán de gran utilidad en el conjunto de los cuaterniones. Ahora bien, un cuaternión, como bien ya se dijo, es un número o un objeto de cuatro dimensiones de la forma:

$$q = a + bi + cj + dk, \quad (1)$$

donde $i, j, k \in \mathbb{C}$ y $a, b, c, d \in \mathbb{R}$, o sea que si $a = b = c = d = 0$ se tiene el cuaternión nulo, o si $a = 0$ y $b = c = d \neq 0$ entonces se tiene un objeto en el espacio tridimensional complejo. Así, podríamos seguir dándole distintas formas al cuaternión y usarlo como bien convenga. A continuación, se dan unas definiciones indispensables para abordar propiedades de los cuaterniones, estas definiciones se basan en las propiedades de los números complejos.

Definición 4.1. Dado el cuaternión $q = a + bi + cj + dk$, su cuaternión conjugado es $\bar{q} = a - bi - cj - dk$ [Fav08].

Definición 4.2. Dado el cuaternión $q = a + bi + cj + dk$, su cuaternión opuesto es $-q = -a - bi - cj - dk$ [Fav08].

Definición 4.3. Dado el cuaternión $q = a + bi + cj + dk$, se define su norma como $|q| = \sqrt{a^2 + b^2 + c^2 + d^2}$, además, q se puede normalizar mediante la siguiente operación: $q' = \frac{q}{|q|}$ donde q' es q normalizado [Fav08].

Definición 4.4. Dado el cuaternión $q = a + bi + cj + dk$, se dice que q es unitario si su norma es igual a uno [Fav08].

Definición 4.5. Dado el cuaternión $q = a + bi + cj + dk$, su inverso es $q^{-1} = \frac{\bar{q}}{|q|^2}$ [Fav08].

El conjunto que contiene a los cuaterniones se denota $\mathbb{H}$, con $\mathbb{H} = \{q : q = a + bi + cj + dk, i, j, k \in \mathbb{C} \wedge a, b, c, d \in \mathbb{R}\}$.

De forma análoga a los números complejos, se define la suma de los números cuaterniones:

Dados $q = a_1 + b_1i + c_1j + d_1k$, $p = a_2 + b_2i + c_2j + d_2k$, se define la suma de cuaterniones como:

$$q + p = (a_1 + a_2) + (b_1 + b_2)i + (c_1 + c_2)j + (d_1 + d_2)k,$$

donde $(a_1 + a_2), (b_1 + b_2), (c_1 + c_2), (d_1 + d_2) \in \mathbb{R}$

Debido a las propiedades de los números complejos y los números reales, se puede ver fácilmente que los cuaterniones cumplen la asociatividad y conmutatividad, además, por la definición 4.2, existe el neutro aditivo y cuando todas las componentes reales son cero, se asegura la existencia del neutro aditivo, o sea que $(\mathbb{H}, +)$ es un grupo abeliano.

La resta de cuaterniones se hace de forma análoga y se puede probar fácilmente. Así también, podemos pensar en la suma de un cuaternión con su conjugado:

$$q + \bar{q} = (a + a) + (b - b)i + (c - c)j + (d - d)k = 2a.$$

Por su parte, el producto de cuaterniones es un poco diferente a la multiplicación de complejos puesto que se debe tener en cuenta la identidad arriba descrita $i^2 = j^2 = k^2 = ijk = -1$ de donde se extrae que:

$$ij = -ij = k, jk = -kj = i, ki = -ik = j. \quad (2)$$

Con esto en mente, se define el producto de cuaterniones así:

$q \cdot p = (a_1a_2 - b_1b_2 - c_1c_2 - d_1d_2) + (a_1b_2 + a_2b_1 + c_1d_2 - d_1c_2)i + (a_2c_1 + a_1c_2 + d_1b_2 - b_1d_2)j + (a_2d_1 + a_1d_2 - c_1b_2 + b_1c_2)k$ [Fav08]. De hecho, este resultado es muy complejo, ya que usa las propiedades de 2, por lo que tiene una mejor visualización de la siguiente manera:

q/p	a_2	b_2i	c_2j	d_2k
a_1	a_1a_2	a_1b_2i	a_1c_2j	a_1d_2k
b_1i	a_2b_1i	$b_1b_2i^2$	b_1c_2ij	b_1d_2ik
c_1j	a_2c_1j	b_2c_1ji	$c_1c_2j^2$	c_1d_2jk
d_1k	a_2d_1k	b_2d_1ki	c_2d_1kj	$d_1d_2k^2$

Cuadro 1: Multiplicación de cuaterniones.

realizando los productos definidos:

q/p	a_2	b_2i	c_2j	d_2k
a_1	a_1a_2	a_1b_2i	a_1c_2j	a_1d_2k
b_1i	a_2b_1i	$-b_1b_2$	b_1c_2j	b_1d_2k
c_1j	a_2c_1j	b_2c_1i	$-c_1c_2$	c_1d_2k
d_1k	a_2d_1k	b_2d_1j	c_2d_1i	$-d_1d_2$

Cuadro 2: Multiplicación de cuaterniones.

y organizando por columnas y por términos semejantes se tiene:

q/p	a_2	b_2i	c_2j	d_2k
a_1	a_1a_2	a_1b_2i	a_1c_2j	a_1d_2k
b_1i	$-b_1b_2$	a_2b_1i	b_1c_2j	b_1d_2k
c_1j	$-c_1c_2$	c_1d_2j	a_2c_1j	b_2c_1i
d_1k	$-d_1d_2$	c_2d_1i	b_2d_1j	a_2d_1k

Cuadro 3: Multiplicación de cuaterniones.

Es decir, tomando por columnas y factorizando tenemos que $q \cdot p = (a_1a_2 - b_1b_2 - c_1c_2 - d_1d_2) + (a_1b_2 + a_2b_1 + c_1d_2 - d_1c_2)i + (a_2c_1 + a_1c_2 + d_1b_2 - b_1d_2)j + (a_2d_1 + a_1d_2 - c_1b_2 + b_1c_2)k$, como ya nos había dado antes. Veamos el siguiente teorema:

Teorema 4.1. Sea $q = t + xi + yj + zk \in \mathbb{H}$, entonces la matriz asociada $L(q)$ es:

$$L(q) = \begin{pmatrix} t & -x & -y & -z \\ x & t & -z & y \\ y & z & t & -x \\ z & -y & x & t \end{pmatrix}. \quad (3)$$

El teorema anterior nos dice que existe una matriz asociada para cualquier cuaternión y eso no es más que lo que vimos en las tres tablas anteriores, solo que allí organizamos por columnas y sumamos términos semejantes. De hecho, esta forma es útil cuando los productos son mucho más complejos. Además,

de esta forma también puede probarse que el producto de cuaterniones es asociativo.

Por otra parte, nos podemos dar cuenta fácilmente, que dada la identidad 1 y los productos 2, la conmutatividad de cuaterniones no se cumple y se puede probar de la misma forma. Ahora bien, igual que con la suma, podemos mirar qué pasa con el producto de un cuaternión con su conjugado, esto es:

$$q \cdot \bar{q} = a^2 + b^2 + c^2 + d^2. \quad (4)$$

Este resultado será muy importante puesto que de aquí surge la idea de la rotación tridimensional de puntos respecto a un vector.

Así también, otra propiedad que vale resaltar es el conjugado del producto de dos cuaterniones:

$$\overline{q \cdot p} = \bar{p} \cdot \bar{q}.$$

Sintetizando un poco, los cuaterniones son un conjunto con ciertas propiedades para la suma y para la multiplicación escalar. De hecho, teniendo en cuenta que satisface todas las propiedades para la suma y casi todas para la multiplicación escalar a excepción de la conmutatividad, podríamos decir que los cuaterniones $(\mathbb{H}, +, \cdot)$ forman una estructura algebraica de anillo no abeliano y con elemento neutro. Ahora bien, observemos que si generalizamos este anillo, nos damos cuenta que los cuaterniones son un espacio vectorial de $\mathbb{R}^4$ que no representa el mismo $\mathbb{R}^n$ sino que lo denotamos $\mathbb{R}'^4$ considerando que la primera componente es real y las otras tres son complejas. Se puede probar que este conjunto es un espacio vectorial peculiar probando cada uno de los 10 axiomas de espacio vectorial; sin embargo no lo hacemos porque no es el objetivo de este escrito. De todas formas, este resultado es demasiado importante porque como veremos más adelante, nos servirá para interpretar un vector bidimensional de Hilbert en este espacio, para realizar operaciones sobre este vector en una esfera de radio uno. Por ahora, nos ayuda bastante tener la idea clara de la forma de un vector en $\mathbb{H}$.

Hasta aquí hemos visto tan solo las propiedades de estos conjuntos; sin embargo son suficientes, así que, sin más que decir, vamos a comenzar con uno de los resultados clave de este artículo:

Consideremos por ejemplo que cualquier punto $P(x, y, z)$ del espacio tridimensional conocido puede ser representado como un cuaternión cuya parte real es cero: $P(x, y, z) \rightarrow P = 0 + xi + yj + zk$. Así, si $P = (1, 2, 3) \rightarrow P = (i + 2j + 3k)$. Ahora, consideremos otro punto $S = (x_1, y_1, z_1)$, representémoslo como cuaternión y luego normalicémoslo. De esta forma, podemos imaginarnos al cuaternión S como la esfera de radio 1 en el espacio tridimensional, donde (x, y, z) son las componentes de los ejes. Por lo tanto, esta esfera la podemos rotar y todos los puntos del espacio con ella, pero ¿cómo? Consideremos ahora un ángulo θ y dejemos que el vector unitario S sea el eje de rotación, es decir, el vector referencia respecto de quien se realiza la rotación del espacio en un ángulo θ dado. Entonces, esta operación de rotación se puede representar como un cuaternión q :

$$q = (s, v), \quad (5)$$

donde $s = \cos \frac{\theta}{2}$ y $v = S \cdot \sin \frac{\theta}{2}$ [Mar16].

El punto P' final después de la rotación se obtiene de realizar la siguiente operación:

$$P' = q \cdot P \cdot \bar{q}. \quad (6)$$

En esencia, necesitamos un punto, un vector unitario que actúa como el eje de referencia de la rotación y un ángulo de rotación. Veamos la siguiente ilustración:

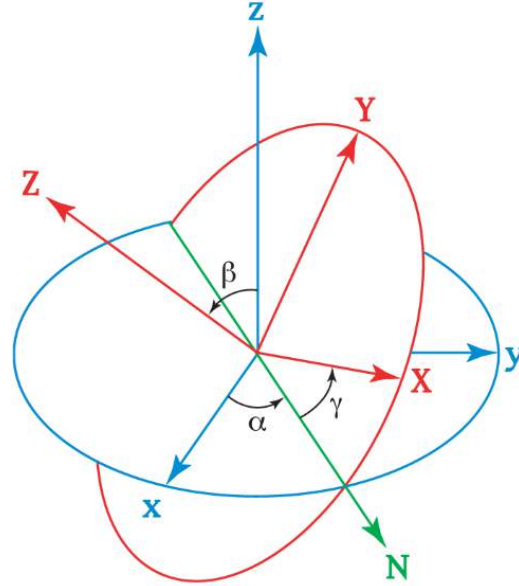


Figura 1: Rotación en el espacio

Como se puede observar, el vector referente o eje de rotación es n y todos los puntos del plano original xy (azul) rota un ángulo β hacia el plano rojo. Ahora bien, recordando que estamos trabajando con el punto P y el vector normalizado q y su conjugado, notemos que si el vector unitario q lo tomamos respecto de algún eje, es decir, respecto a x , y o z , el cálculo se hace demasiado sencillo. Por ejemplo, si queremos rotar el espacio respecto al eje z , entonces tenemos que calcular $P' = q \cdot P \cdot \bar{q}$, que equivale a:

$$\left(\cos \frac{\theta}{2} + 1 \cdot \sin \frac{\theta}{2}\right) \cdot (0 + x_1 i + y_1 j + z_1 k) \cdot \left(\cos \frac{\theta}{2} - 1 \cdot \sin \frac{\theta}{2}\right).$$

Con todo, no siempre se quiere una rotación respecto a uno de los ejes sino que en la mayoría de los casos, las rotaciones son sobre distintos vectores referencia del espacio, y en este, hay infinitos. Además, uno de los mayores problemas sería el hecho de rotaciones sucesivas sobre distintos ejes. Es claro que para ello haríamos una matriz con cuantos cuaterniones consideremos necesarios para hacer las rotaciones del espacio necesarias, pero haciendo esta tarea con distintos ejes del espacio, el resultado final o posición final de cada punto no podría ser resultado de un número mayor a uno de vectores referencia sobre el que se rotó el espacio. Esto da cabida al siguiente teorema propuesto por Leonhard Euler antes de que Hamilton descubriera los cuaterniones:

Teorema 4.2. Si R es una matriz que representa una rotación en $\mathbb{R}^3$, entonces R tiene un vector propio $n \in \mathbb{R}^3$ tal que $Rn = n$, esto es, n es un vector propio con valor propio 1 [Anó10].

Este teorema lo que nos dice es que, con toda seguridad, cualquier rotación por un cierto ángulo en el espacio tridimensional deja fija una línea recta, que es el eje de rotación, o sea que dada una secuencia de rotaciones sobre distintos ejes, existe una rotación que generaliza toda la secuencia de rotaciones, y esta, a su vez tiene asociada una sola matriz de rotación sobre un eje de referencia o eje de rotación.

4.2. Rotaciones.

Ya vimos que un vector cualquiera en el espacio se puede expresar como cuaternión, y este, normalizado, es una esfera de radio 1 que se puede rotar. Así pues, podemos usar la geometría analítica para deducir una expresión sumamente importante y es la siguiente:

Consideremos el plano $\mathbb{R}^2$ con vector v de coordenadas (v_x, v_y) que se quiere rotar θ grados en sentido antihorario hasta el vector w de coordenadas w_x, w_y .

Digamos que la norma del vector $v = r$, así, $|v| = |w| = r$, la coordenada en x destino $w_x = r \cdot \cos \theta + \alpha = r \cdot [\cos \alpha \cos \theta - \sin \alpha \sin \theta] = r \cos \alpha \cos \theta - r \sin \alpha \sin \theta$, pero $v_x = r \cdot \cos \alpha$ y $v_y = r \cdot \sin \alpha$, o sea que reescribiendo se tiene: $w_x = v_x \cos \theta - v_y \sin \theta$. Análogamente, la coordenada en y destino $w_y = v_x \sin \theta + v_y \cos \theta$.

Si expresamos esto en forma matricial:

$$\begin{bmatrix} w_x \\ w_y \end{bmatrix} = \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix} \cdot \begin{bmatrix} v_x \\ v_y \end{bmatrix}$$

Extrapolar este resultado a $\mathbb{R}^3$ no es complicado, para empezar, podríamos tomar la siguiente matriz cuyo eje de rotación es el eje z , entonces se tendría:

$$\begin{bmatrix} w_x \\ w_y \\ w_z \end{bmatrix} = \begin{bmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} v_x \\ v_y \\ v_z \end{bmatrix}$$

Para entender la relación entre el resultado anterior y la rotación usando cuaterniones, es necesario asociar las rotaciones tridimensionales a la fórmula de Rodrigues [Anó10] en honor al matemático Olinde Rodrigues,² que consiste en rotar un vector v alrededor de otro unitario n y cuyas coordenadas se obtienen de $v \cos \alpha + (nxv) \sin \alpha + n(n \cdot v)(1 - \cos \alpha)$, donde (nxv) es el producto cruz de vectores y $(n \cdot v)$ el producto punto de ambos vectores. Así, con cuaterniones, se definen dos cuaterniones entre las partes imaginarias de dos cuaterniones $q = (s_1, v_1), p = (s_2, v_2)$: El producto punto y el producto cruz $(v_1 xv_2)$ y $(v_1 \cdot v_2)$. Luego, se tiene que el producto de cuaterniones es:

$$q \cdot p = (s_1, v_1) \cdot (s_2, v_2) = s_1 s_2 - (v_1 \cdot v_2), s_1 v_2 + s_2 v_1 + (v_1 xv_2).$$

²A quien se le atribuyen trabajos memorables previos al descubrimiento de Hamilton de los cuaterniones.

Ahora, si $q = (s_1, \gamma v_2)$ con $s_1 = \cos \frac{\theta}{2}$ y $\gamma = \sin \frac{\theta}{2}$ se puede llegar a demostrar que dado el punto P que se quiere rotar, 6 es equivalente.

El resultado de Rodrigues nos da paso a buscar un vector característico que podamos rotar y nos arroje algún resultado interesante. En la siguiente sección, introducimos los espacios de Hilbert, ya que las rotaciones de vectores que nos interesa, se realizan en una esfera de radio 1 cuyos valores posibles representan los estados de un qubit.

5. Espacios de Hilbert y transformaciones lineales.

Definición 5.1. Sean V, W dos espacios vectoriales de dimensión finita, F un campo y $f: V \rightarrow W$ una función, diremos que f es una transformación lineal de V en W si se cumple:

- $f(x + y) = f(x) + f(y), \forall x, y \in V$.
- $f(\alpha x) = \alpha f(x) \forall x \in V, \alpha \in F$.

Lo anterior es lo mismo que $f(\alpha x + \alpha y) = \alpha f(x) + \alpha f(y)$. Ahora, sabemos que el producto interno es una función lineal que va del espacio vectorial al campo sobre el que se define dicho espacio vectorial.

Definición 5.2. Sea V un espacio vectorial de dimensión finita y $\mathbb{R}$ el campo sobre el que se define V . El producto interno o producto escalar definido sobre V es una aplicación o función $f: V \times V \rightarrow \mathbb{R}$, denotado $\langle u, v \rangle$ que satisface las siguientes propiedades para todo $u, v, w \in V$ y todo escalar $\alpha \in \mathbb{R}$:

- $\langle u, v \rangle = \langle v, u \rangle$.
- $\alpha \langle u, v \rangle = \langle \alpha u, v \rangle = \langle u, \alpha v \rangle$.
- $\langle u + v, w \rangle = \langle u, w \rangle + \langle v, w \rangle$.
- $\langle u, u \rangle \geq 0 \wedge \langle u, u \rangle = 0 \Leftrightarrow u = 0$.

Es decir, el producto interno entre dos vectores del espacio devuelve un número del campo de los reales. Además, se puede demostrar fácilmente que la propiedad 3 de 5.2 es lineal en la segunda componente como en la primera componente. Es decir, aplicando 5.1 $\langle u, \alpha v + \alpha w \rangle = \alpha \langle u, v \rangle + \alpha \langle u, w \rangle$. Lo mismo con la primera componente. Así, se dice que el producto interno es una forma bilineal simétrica y definida positiva por las anteriores propiedades; sin embargo, ¿qué sucede con esta forma cuando se tienen vectores de un cuerpo complejo?

Supongamos un vector $u = (1, 0, i) \in \mathbb{C}^3$ y probemos la cuarta propiedad de 5.2. Así, $(1, 0, i) \cdot (1, 0, i) = 0$, o sea que no se cumple la propiedad de que $u = 0$. Por lo tanto, se debe hacer una pequeña modificación. Si hacemos que la propiedad sea $\langle u^*, u \rangle$, es decir, el complejo conjugado por el vector, entonces podemos observar lo siguiente: $(1, 0, -i) \cdot (1, 0, i) =$

$2 \neq 0$. Ahora, la propiedad simétrica no se cumple y necesitamos que $\langle u, v \rangle = \langle v, u \rangle^*$, o sea al complejo conjugado, y ahora se dice que cumple la propiedad hermítica. Por último, y no menos importante, se puede probar que bajo las condiciones impuestas en este párrafo, ahora, la condición de bilinealidad no se cumpliría sino solo para la segunda componente, mientras que para la primera, los escalares saldrían complejos conjugados de la siguiente forma: $\langle \alpha_1 u + \alpha_2 w, v \rangle = \alpha_1^* \langle u, v \rangle + \alpha_2^* \langle w, v \rangle$. Por consiguiente, tenemos la siguiente definición:

Definición 5.3. Sea V un espacio vectorial de dimensión finita sobre el cuerpo $\mathbb{K}$, el producto interno o producto escalar definido sobre V es una aplicación o función $f : V \times V \rightarrow \mathbb{K}$, denotado $\langle u, v \rangle$ que satisface las siguientes propiedades para todo $u, v, w \in V$ y todo escalar $\alpha \in \mathbb{K}$:

- $\langle u, v \rangle = \langle v, u \rangle^*$.
- $\langle \alpha_1 u + \alpha_2 w, v \rangle = \alpha_1^* \langle u, v \rangle + \alpha_2^* \langle w, v \rangle$, y $\langle u, \alpha_1 v + \alpha_2 w \rangle = \alpha_1 \langle u, v \rangle + \alpha_2 \langle u, w \rangle$.
- $\langle u, u \rangle \geq 0 \wedge \langle u, u \rangle = 0 \Leftrightarrow u = 0$.

De la anterior definición se sigue que:

Definición 5.4. Un espacio pre-Hilbertiano es un espacio vectorial sobre $\mathbb{K}$ con producto interno y normado con $\|v\| = \sqrt{\langle v, v \rangle}$ [AD15].

Definición 5.5. Un espacio vectorial V es completo para la norma, sí y sólo si toda sucesión de cauchy³ converge con esa norma [AD15].

De lo anterior, tenemos la siguiente definición:

Definición 5.6. Un espacio de Hilbert $\mathcal{H}$ es un espacio pre-Hilbertiano cuya norma es completa [AD15].

Una vez hemos introducido el espacio de Hilbert, definiremos formalmente el qubit:

Definición 5.7. Un qubit o bit cuántico es un vector normalizado del espacio de Hilbert $\mathbb{C}^2$ [AD15].

Así, vemos como nos ha venido de bien introducir previamente los números complejos y el producto interno sobre ellos. Ahora bien, y a manera de sintaxis únicamente, definiremos la notación de Dirac que se emplea en la mecánica cuántica para vectores:

Definición 5.8. Llamamos ket a un vector de la forma:

$$|\psi\rangle = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}. \quad (7)$$

Definición 5.9. Llamamos bra a un vector de la forma:

$$\langle\psi| = (x_1^* \ x_2^* \ \dots \ x_n^*). \quad (8)$$

Es decir, los kets son vectores columna y los bras son vectores fila donde cada una de sus entradas es el complejo conjugado de un número cualquiera del campo de los complejos. Ahora, dijimos anteriormente que un qubit podría ser expresado como una combinación lineal de los vectores de la base de $\mathcal{H}$. En este caso, definimos como vectores de la base fundamental de la computación cuántica a los vectores $|0\rangle$ y $|1\rangle$, que en notación convencional:

$$|0\rangle = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, |1\rangle = \begin{pmatrix} 0 \\ 1 \end{pmatrix} \quad (9)$$

Así, cualquier vector del espacio dos dimensional de Hilbert puede ser generado por el ket cero y el ket uno, en particular, el qubit $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ con $\alpha, \beta \in \mathbb{C}$.

Con la base fundamental, podemos ver qué sucede con el producto interno de las mismas y que se nota de la siguiente manera: Se llama bracket a $\langle\psi|\psi\rangle$. En particular, teniendo en cuenta la multiplicación de matrices:

$$\langle 0|0\rangle = \begin{pmatrix} 1 & 0 \end{pmatrix} \cdot \begin{pmatrix} 1 \\ 0 \end{pmatrix} = 1, \quad (10)$$

$$\langle 1|1\rangle = \begin{pmatrix} 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 0 \\ 1 \end{pmatrix} = 1, \quad (11)$$

$$\langle 0|1\rangle = \begin{pmatrix} 1 & 0 \end{pmatrix} \cdot \begin{pmatrix} 0 \\ 1 \end{pmatrix} = 0, \quad (12)$$

$$\langle 1|0\rangle = \begin{pmatrix} 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 \\ 0 \end{pmatrix} = 0. \quad (13)$$

Hasta aquí, definir un qubit necesitó de la introducción de un espacio vectorial perfecto, pero ahora, queremos operar con él, y claro, para que una computadora interprete un qubit, debe existir una operación denominada medición, y a su vez, un operador que actúe sobre el qubit para colapsarlo, es decir, para que lo lleve a otra base del mismo espacio o cambie de estado. Esta medición tiene un comportamiento peculiar siguiendo las leyes de la física cuántica. Es por eso que a continuación presentamos tres postulados de la mecánica cuántica que servirán para definir el comportamiento de un qubit.

1. Estado del espacio: Nos dice que a cada sistema físico aislado, se asocia un espacio de Hilbert que es el espacio de estados del sistema. Hay que tener en cuenta, que en el plano complejo, los vectores no representan posiciones sino estados, o sea que cada eje del plano o espacio complejo explican un estado de la cosa o vector. El sistema físico está completamente descrito por el estado de sus vectores, y la dimensión del espacio depende de los grados de libertad de la propiedad física, es decir, el número de posibles valores que puede tomar el objeto dentro del sistema.

³ $\langle u, v \rangle = \langle v, u \rangle^*$ para $u, v \in V$ un espacio vectorial cualquiera.

Este postulado implica que una combinación lineal de un vector genera otro vector del mismo espacio, y esto se conoce como el principio de superposición.

- 2. Evolución cuántica:** Nos dice que el estado de un sistema en el tiempo t_2 de acuerdo al estado en el tiempo t_1 está dado por $|\psi_{t_2}\rangle = U|\psi_{t_1}\rangle$.

O sea que el estado de un sistema se altera con la aplicación de un operador U que se definirá como una transformación lineal sobre el vector en el espacio $\mathcal{H}$. Así, aplicando un operador de estos a un qubit, encontraremos el estado de este y será equivalente a correr un programa en una computadora.

- 3. Medida cuántica** Nos dice que todos los resultados de un experimento son generados de acuerdo a cierta distribución probabilística. Además, antes de la medición el estado cuántico de un sistema es diferente al estado cuántico del objeto después de la medición.

Del tercer postulado se extrae que la medición es un proceso probabilístico, los estados pre y post medición no son iguales, y medir un sistema equivale a proyectar el vector sobre las bases del espacio de Hilbert sobre el que emana.

Así pues, al medir un qubit, nos referimos a la evaluación probabilística del estado pre y post. Es por eso que para la preparación de un qubit, se deben elegir adecuadamente los escalares α, β de tal manera que $|\alpha|^2 + |\beta|^2 = 1$, y estas mediciones las podemos conseguir mediante operadores o básicamente, transformaciones lineales en el mismo espacio de Hilbert. A continuación, veremos que estas transformaciones están definidas por ciertas compuertas denominadas operadores.

5.1. Operadores y compuertas cuánticas.

A las transformaciones lineales se les conoce como operadores o como puertas cuánticas, que significan una multiplicación de una matriz por los vectores de la base del qubit y poder encontrar los diferentes estados. Esta matriz, generalmente es unitaria, es decir, una matriz cuya inversa es igual a su conjugada transpuesta. Si hacemos una analogía de una compuerta cuántica con una compuerta lógica de la computación clásica, las puertas lógicas AND, OR y NOT se aplican a un bit de tal forma que si el estado de este es 0 y se aplica NOT entonces se obtiene 1; sin embargo, al aplicar una compuerta cuántica sobre un qubit no se obtiene un resultado discreto y entero sino probabilidades de obtener uno de los dos estados y por ende multitud de estados intermedios como resultado de la superposición de estos. Así, podemos empezar a pensar en un conjunto de matrices que nos van a permitir obtener más resultados de un evento físico, en particular, nos interesamos en buscar compuertas cuánticas definidas como operadores sobre qubits. Pero antes de ello, es importante enterarnos de las siguientes definiciones.

Definición 5.10. Un Operador $\hat{A}$ de $\mathbb{C}^n$ es una matriz cuadrada de dimensión n con entradas complejas [AD15].

Definición 5.11. El adjunto de un operador $\hat{A}$ se nota $\hat{A}^\dagger$ conocido como A dagger y significa el operador transpuesto conjugado de $\hat{A}$ [1007AD15].

El operador $\hat{A}^\dagger$ lo que hace es transformar un bra, mientras que el operador $\hat{A}$ induce sobre un ket así:

$$\hat{A} \cdot |\psi\rangle = |\psi'\rangle, \quad (14)$$

$$\hat{A}^\dagger \cdot \langle\psi| = \langle\psi'|. \quad (15)$$

Definición 5.12. Sea $\hat{A}$ un operador lineal tal que $\hat{A} : \mathcal{H} \rightarrow \mathcal{H}$, $\hat{A}$, entonces $|\psi\rangle \rightarrow |\psi'\rangle$.

Esta definición es muy importante porque estamos diciendo que existe una transformación lineal en el espacio n – dimensional de Hilbert que lleva un ket a otro ket del mismo espacio mediante un operador lineal. Así, un qubit se transforma y como resultado de esta transformación lineal, se obtiene otro qubit de diferente estado.

En teoría de la computación cuántica, se han definido unas compuertas muy especiales como el operador de Hadamard, que es una compuerta lógica sobre la que se cimientan muchos circuitos cuánticos. La siguiente definición, ayuda a entender este operador.

Definición 5.13. A los operadores de la forma $P = |\phi\rangle\langle\phi|$ se les llama proyectores ⁴ [AD15].

De acuerdo con 7, 8, es claro que un proyector devuelve una matriz a diferencia del producto interno de ambos que devolvía un número del campo complejo. Por ello, una representación matricial

$$|0\rangle\langle 0| = \begin{pmatrix} 1 \\ 0 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}, \quad (16)$$

$$|1\rangle\langle 1| = \begin{pmatrix} 0 \\ 1 \end{pmatrix} \cdot \begin{pmatrix} 0 & 1 \end{pmatrix} = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}, \quad (17)$$

$$|0\rangle\langle 1| = \begin{pmatrix} 1 \\ 0 \end{pmatrix} \cdot \begin{pmatrix} 0 & 1 \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}, \quad (18)$$

$$|1\rangle\langle 0| = \begin{pmatrix} 0 \\ 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 \end{pmatrix} = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}. \quad (19)$$

Ahora bien, con base en lo anterior podemos obtener que $(|0\rangle\langle 0|)(|0\rangle\langle 0|) = |0\rangle\langle 0|$ y $(|1\rangle\langle 1|)(|1\rangle\langle 1|) = |1\rangle\langle 1|$ como evidentemente se podría probar haciendo el producto de la misma matriz. Lo mismo sucederá con el operador dagger si $(|0\rangle\langle 0|)^\dagger(|0\rangle\langle 0|)^\dagger = |0\rangle\langle 0|$ y $(|1\rangle\langle 1|)^\dagger(|1\rangle\langle 1|)^\dagger = |1\rangle\langle 1|$ Esto será importante para minimizar el número de operaciones en una compuerta cuántica y más aún, para el cálculo de probabilidades.

⁴Se llaman proyectores porque proyectan ortogonalmente un ket $|\psi\rangle$ cualquiera sobre el ket $|\phi\rangle$: $P|\phi\rangle = |\phi\rangle\langle\phi|\psi\rangle = c|\phi\rangle$.

Definición 5.14. Se define como el operador de Hadamard a [VM13]:

$$\hat{H} = \frac{1}{\sqrt{2}} (|0\rangle\langle 0| + |0\rangle\langle 1| + |1\rangle\langle 0| - |1\rangle\langle 1|). \quad (20)$$

Por ejemplo, la acción que toma el operador 5.14 sobre el ket cero es la siguiente:

$$\begin{aligned} \hat{H}|0\rangle &= \frac{1}{\sqrt{2}} (|0\rangle\langle 0| + |0\rangle\langle 1| + |1\rangle\langle 0| - |1\rangle\langle 1|)|0\rangle \\ &= \left(\frac{1}{\sqrt{2}}|0\rangle\langle 0| + \frac{1}{\sqrt{2}}|0\rangle\langle 1| + \frac{1}{\sqrt{2}}|1\rangle\langle 0| - \frac{1}{\sqrt{2}}|1\rangle\langle 1| \right)|0\rangle \\ &= \frac{\langle 0|0\rangle}{\sqrt{2}}|0\rangle + \frac{\langle 1|0\rangle}{\sqrt{2}}|0\rangle + \frac{\langle 0|0\rangle}{\sqrt{2}}|1\rangle - \frac{\langle 1|0\rangle}{\sqrt{2}}|1\rangle, \end{aligned}$$

pero por 10, 11, 12, 13 se obtiene que $\hat{H}|0\rangle = \frac{1}{\sqrt{2}}|0\rangle + \frac{1}{\sqrt{2}}|1\rangle$.

Esto viene siendo lo mismo que

$$\hat{H}|0\rangle = \begin{pmatrix} \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} & -\frac{1}{\sqrt{2}} \end{pmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix}. \quad (21)$$

Teniendo en cuenta el resultado obtenido, nos damos cuenta que un qubit tendría un estado intermedio dado por la compuerta cuántica de Hadamard, puesto que el resultado involucra tanto al ket cero como al ket uno, y bastaría no más con multiplicar la matriz obtenida con cualquier otro ket o qubit.

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

Recordemos que hemos encontrado cómo es la forma de operar un qubit. Así pues, para describir o medir cualquier sistema cuántico se requiere definir el sistema que deseamos analizar (en este caso un qubit en un espacio dos dimensional de Hilbert), tener en cuenta los posibles resultados que puede tomar el sistema (estos también vienen dados por la dimensión del espacio), tener una base ortonormal $B = \{|i\rangle : |i\rangle \in \mathcal{H}^n\}$, un conjunto de operadores de medición $\hat{A}$ dados por la proyección de cada par de vectores de la base, y una distribución probabilística en función de los experimentos sobre el sistema. Para este último requisito, sabemos que el estado pre y post medición son distintos. Por lo tanto, se tiene la siguiente definición:

Definición 5.15. De acuerdo con [VM13], se dice que un sistema representado por un ket $|\psi\rangle$ evoluciona al sistema $|\phi\rangle$, cuando se realiza una de las siguientes operaciones:

- 1 Se premultiplica por una compuerta cuántica $\hat{A}$: $|\phi\rangle = \hat{A}|\psi\rangle$.
- 2 Se aplica un operador de medición $\hat{M}$: $|\phi\rangle = \frac{\hat{M}|\psi\rangle}{\sqrt{\langle\psi|\hat{M}^\dagger\hat{M}|\psi\rangle}}$.

Donde $\sqrt{\langle\psi|\hat{M}^\dagger\hat{M}|\psi\rangle}$ es un estado que no se conoce, mas sí su probabilidad. Así pues, buscamos la probabilidad de obtener un estado del qubit usando la compuerta cuántica de Hadamard.

Como vimos previamente, $\frac{1}{\sqrt{2}}|0\rangle + \frac{1}{\sqrt{2}}|1\rangle$ es el estado de un qubit luego de multiplicarlo por la matriz u operador de

Hadamard. Sea pues $|\psi\rangle = \hat{H}|0\rangle$, vemos que este qubit está definido sobre la base computacional $B = \{|0\rangle, |1\rangle\}$. Además, como vive en el espacio dos dimensional de Hilbert, tiene dos eventos posibles, a saber $\{0, 1\}$. Para calcular $\sqrt{\langle\psi|\hat{M}^\dagger\hat{M}|\psi\rangle}$, necesitamos saber cómo se comporta el operador para cada estado posible. Así, nos interesa $\hat{M}_0^\dagger, \hat{M}_0, \hat{M}_1^\dagger$ y $\hat{M}_1$.

$$\hat{M}_0 = |0\rangle\langle 0|$$

$$\hat{M}_1 = |1\rangle\langle 1|$$

$$\hat{M}_0^\dagger = (|0\rangle\langle 0|)^\dagger = |0\rangle\langle 0|$$

$$\hat{M}_1^\dagger = (|1\rangle\langle 1|)^\dagger = |1\rangle\langle 1|$$

pero ya vimos que $\hat{M}_0^\dagger\hat{M}_0 = (|0\rangle\langle 0|)(|0\rangle\langle 0|) = |0\rangle\langle 0|$ y $\hat{M}_1^\dagger\hat{M}_1 = (|1\rangle\langle 1|)(|1\rangle\langle 1|) = |1\rangle\langle 1|$. Así, tenemos el siguiente resultado:

$$\begin{aligned} \langle\psi|\hat{M}_0^\dagger\hat{M}_0|\psi\rangle &= \left(\frac{1}{\sqrt{2}}\langle 0| + \frac{1}{\sqrt{2}}\langle 1|\right)(|0\rangle\langle 0|)\left(\frac{1}{\sqrt{2}}|0\rangle + \frac{1}{\sqrt{2}}|1\rangle\right) \\ &= \left(\frac{1}{\sqrt{2}}\langle 0| + \frac{1}{\sqrt{2}}\langle 1|\right)\left(\frac{1}{\sqrt{2}}\langle 0|0\rangle|0\rangle + \frac{1}{\sqrt{2}}\langle 0|1\rangle|0\rangle\right) \\ &= \left(\frac{1}{\sqrt{2}}\langle 0| + \frac{1}{\sqrt{2}}\langle 1|\right)\left(\frac{1}{\sqrt{2}}|0\rangle\right) \\ &= \left(\frac{1}{\sqrt{2}}\right)^2\langle 0|0\rangle + \left(\frac{1}{\sqrt{2}}\right)^2\langle 1|0\rangle \\ &= \frac{1}{2}. \end{aligned}$$

Es decir, la probabilidad de obtener el ket cero es del 50%.

Lo mismo sucederá si repetimos el procedimiento para el estado uno. Esto lo podemos corroborar si usamos una simulación de la computadora cuántica de IBM. A continuación, veremos una aplicación simulada sobre un qubit usando la compuerta de Hadamard:

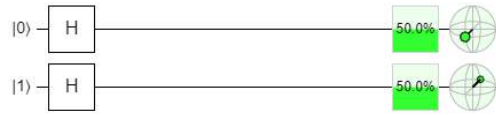


Figura 2: Aplicación de la compuerta de Hadamard en el simulador cuántico de IBM.

Vemos que tanto para el ket 0 como el ket 1, la probabilidad de estar en superposición en una operación es del 50%. Veamos ahora qué sucede si cambiamos de base y volvemos a aplicar la compuerta de Hadamard sobre $|\psi\rangle$:

Sea $B_2 = \{|+\rangle, |-\rangle\}$ la base del espacio $\mathcal{H}^2$, conocida como la base diagonal donde:

$|+\rangle = \frac{1}{\sqrt{2}}|0\rangle + \frac{1}{\sqrt{2}}|1\rangle$ y $|-\rangle = \frac{1}{\sqrt{2}}|0\rangle - \frac{1}{\sqrt{2}}|1\rangle$. Así, tendríamos la siguiente operación:

$$\begin{aligned} \langle\psi|\hat{M}_-^\dagger\hat{M}_+|\psi\rangle &= \left(\frac{1}{\sqrt{2}}\langle 0| + \frac{1}{\sqrt{2}}\langle 1|\right)(|+\rangle\langle +|)\left(\frac{1}{\sqrt{2}}|0\rangle + \frac{1}{\sqrt{2}}|1\rangle\right) \\ &= (\langle +|)(|+\rangle\langle +|)(|+\rangle) \\ &= (\langle +|)(\langle +|+\rangle)(|+\rangle) \\ &= \langle +|\left(\frac{1}{2} - \frac{1}{2}\right)|+\rangle \\ &= \langle +|0|+\rangle \\ &= 0. \end{aligned}$$

Es decir, el qubit no refleja posibilidad alguna de tomar el estado + cuando se le aplica la compuerta de Hadamard a $|\psi\rangle$. Si lo replicamos en el simulador, vemos que confirma

una probabilidad del 100% del qubit en función de la base escogida.

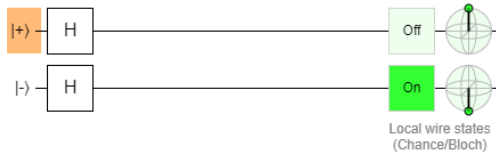


Figura 3: Aplicación de la compuerta de Hadamard en el simulador cuántico de IBM.

A continuación, vemos algunas otras compuertas que se pueden emplear para la construcción de circuitos o algoritmos cuánticos:

Compuerta	Símbolo	Matriz
X/NOT		$\begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$
H		$\frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}$
Phase shift (ϕ)		$\begin{bmatrix} 1 & 0 \\ 0 & e^{i\phi} \end{bmatrix}$
$\sqrt{X}/\sqrt{NOT}$		$\frac{1}{2} \begin{bmatrix} 1+i & 1-i \\ 1-i & 1+i \end{bmatrix}$
$CX/CNOT$		$\begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{bmatrix}$

Figura 4: Compuertas cuánticas.

Así, hemos observado la esencia de las compuertas cuánticas sobre un qubit. Pero, ¿cómo se están viendo estas probabilidades gráficamente? Bueno, el turno es de la esfera de radio 1, esa misma que rota un cuaternión en el espacio tridimensional. Resulta que si nos fijamos bien, el simulador ejecuta una compuerta sobre un vector que rota en una esfera de radio uno. Este vector no es más que un qubit, cuyas rotaciones representan un estado mas no una posición, en mecánica cuántica, se usa la esfera de Bloch, como una representación geométrica de los estados del sistema. La interpretación que se sigue es que cuando el vector o qubit se encuentra en el norte de la esfera, el qubit tiene valor $|0\rangle$, y cuando se encuentra en el polo sur de la esfera, su valor es $|1\rangle$ [MP19].

Los extremos del eje Z representan los estados de la base (en este caso la base computacional). De esta forma, se trabaja rotando el vector a través del espacio tridimensional. Antes de proponer, resaltamos que se puede representar cualquier estado o punto en la esfera de Bloch como una combinación lineal de $|0\rangle, |1\rangle$ así: $|\psi\rangle = \cos \frac{\theta}{2} |0\rangle + \sin \frac{\theta}{2} e^{i\phi} |1\rangle$.

Ahora bien, de acuerdo con [MP19] las transformaciones en la esfera de Bloch se pueden dividir en tres grupos: Puertas de Pauli, Puerta de Hadamard y Puertas de fase.

Definición 5.16. Las puertas de Pauli son aquellas que aplican a un qubit un giro o ángulo de π con respecto a los ejes X, Y, Z . Dependiendo del eje reciben el nombre de puertas

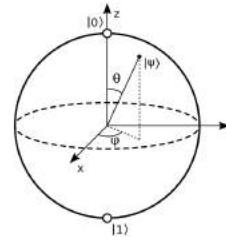


Figura 5: Esfera de Bloch.

X, Y, Z . El efecto que causan sobre los estados básicos del qubit descritos anteriormente se pueden observar a continuación:

	X	Y	Z
$ 0\rangle$	$ 1\rangle$	$ 1\rangle$	$ 0\rangle$
$ 1\rangle$	$ 0\rangle$	$ 0\rangle$	$ 1\rangle$
$ +\rangle$	$ +\rangle$	$ -\rangle$	$ -\rangle$
$ -\rangle$	$ -\rangle$	$ +\rangle$	$ +\rangle$
MATRIZ	$\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$	$\begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}$	$\begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$

Figura 6: Efecto de la transformación con Puertas de Pauli.

Definición 5.17. Las puertas de fase son giros sobre el eje Z. Estas puertas no modifican el valor final del qubit, pero sí afectan a futuras operaciones sobre el mismo. Dentro de este grupo se encuentran las puertas S y St , que corresponden a un giro de ángulo $\frac{\pi}{2}$ y de $-\frac{\pi}{2}$ respectivamente. También se encuentran las puertas T y Tt , que realizan un giro de ángulo de $\frac{\pi}{4}$ y $-\frac{\pi}{4}$ respectivamente. a continuación se observa su efecto sobre el qubit:

	S	S'	T	T'
MATRIZ	$\begin{pmatrix} 1 & 0 \\ 0 & i \end{pmatrix}$	$\begin{pmatrix} 1 & 0 \\ 0 & -i \end{pmatrix}$	$\begin{pmatrix} 1 & 0 \\ 0 & e^{i\frac{\pi}{4}} \end{pmatrix}$	$\begin{pmatrix} 1 & 0 \\ 0 & e^{-i\frac{\pi}{4}} \end{pmatrix}$

Figura 7: Efecto de la transformación con Puertas de fase.

Hemos llegado a ver que el estado de un qubit se comporta de acuerdo con la transformación o compuerta que se utilice sobre él. Nuestra intuición nos lleva a pensar que podemos llevar este vector al espacio de los cuaterniones y mirar qué pasa allí si se rota respecto al eje Z e interpretar cada punto como un estado del sistema.

6. Del espacio perfecto al espacio de los cuaterniones.

6.1. Intuición sobre la transformación.

Para migrar el qubit, debemos llevar el vector dos dimensional a la base cuaterniónica, por eso introducimos la siguiente definición:

Definición 6.1. El producto tensorial de dos matrices P y Q se define como la matriz [AD15]:

$$P \otimes Q = \begin{pmatrix} p_{11}Q & \cdot & \cdot & \cdot & p_{1m}Q \\ \cdot & & & & \cdot \\ \cdot & & & & \cdot \\ \cdot & & & & \cdot \\ p_{n1}Q & \cdot & \cdot & \cdot & p_{nm}Q \end{pmatrix}. \quad (22)$$

Adicionalmente tengamos en cuenta las siguientes definiciones:

Definición 6.2. [AD15] Sean $B_1 = \{v_1, v_2, \dots, v_n\}$ una base del espacio vectorial V y $B_2 = \{w_1, w_2, \dots, w_m\}$ una base del espacio vectorial W . El producto tensorial entre dichas bases se define como:

$$V \otimes W = \{v_1 \otimes w_1, \dots, v_n \otimes w_1, \dots, v_1 \otimes w_m, \dots, v_n \otimes w_m\}.$$

Es decir, el producto tensorial entre espacios vectoriales se define como el espacio generado por el producto tensorial de todos los vectores de una base del primero con los vectores de una base del segundo.

Definición 6.3. Sean B_1 una base del espacio vectorial V , y B_2 una base del espacio vectorial W , entonces: $V \otimes W = \text{span}(B_1, B_2)$ [AD15].

Conforme lo anterior, podemos realizar el siguiente ejercicio y probar qué sucedería si los vectores de la base computacional que generan un qubit en $\mathbb{C}^2$ les aplicamos el producto tensorial con $\mathbb{C}^2$. Es decir, consideremos el espacio generado $\mathbb{C}^2 \otimes \mathbb{C}^2$ ya que:

$$\mathbb{C}^2 \otimes \mathbb{C}^2 = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}. \quad (23)$$

Claramente el espacio generado dado por el producto tensorial de $\mathbb{C}^2$ es $\mathbb{C}^4$. Dejemos que $\mathbb{C}^4 = \{l, i, j, k\}$ sea la base, y que $l = 0$. Por lo tanto, el qubit del espacio de Hilbert, de alguna manera puede ser interpretado en el espacio de los cuaterniones cuando $l = 0$.

6.2. Una rotación del vector en un plano de la esfera de radio uno.

De acuerdo con lo visto previamente cuando presentamos los cuaterniones, vimos que era posible rotar vectores en este espacio sobre una esfera de radio uno, y conforme vimos en 4.1 supimos cómo hacerlo sobre el eje z . Así pues, podríamos simular una de las compuertas cuánticas en este espacio. Por ejemplo, vamos a simular una compuerta de Pauli. De acuerdo con 5.16, el qubit rota ángulos de π en el eje X , Y o Z . Veamos el siguiente código en python que replica una rotación respecto al eje Z de un vector, es decir, una rotación del mismo en el eje Y . Como vamos a ver diferentes rotaciones, dejaremos a continuación explicado el código que usamos para ello:

Importamos librerías:

```
1 import matplotlib.pyplot as plt
2 import pandas as pd
3 from mpl_toolkits.mplot3d import axes3d
4 import matplotlib.cm as cm
5 import numpy as np
```

A continuación definimos el vector V que queremos rotar en un array que contendrá sus coordenadas.

```
1 # Coordenadas del vector unitario
2 V = np.array([(0,0,0), (1/sqrt(2), 1/sqrt(2), 0)])
3
4 Ax = [] #Lista de primeras componentes
5 By = [] #Lista de segundas componentes
6 Cz = [] #Lista de terceras componentes
```

las listas vacías A_x, B_y, C_z se usarán a continuación para almacenar las nuevas coordenadas rotadas, es decir, las coordenadas que resultan al multiplicar el vector por la matriz de rotación. A continuación, definimos la función que llamamos **imp** que contiene la matriz de rotación de los vectores iniciales:

```
1 # Funcion que rota las coordenadas del vector V
  en
2 # funcion del ngulo theta
3 def imp(theta):
4 # Matriz de Rotacion
5 R=np.array([[np.cos(theta), np.sin(theta),
6             0.],
7            [-np.sin(theta), np.cos(theta), 0.],
8            [0., 0., 1.]])
9 for row in V:
10     output = (R.dot(row)) #Multiplicacion de la
11     # matriz por el punto respectivo
12     x, y, z= output[0],output[1], output[2]
13     #Coordenadas transformadas
14     Ax.append(x) #Se guardan las coordenadas
15     By.append(y)
16     Cz.append(z)
17 return Ax,By,Cz
```

La matriz R trabaja en función del ángulo que se reciba como parámetro y el bucle tiene la finalidad de multiplicar cada vector. En el siguiente ciclo le indicamos a la función **imp** que queremos que la transformación se haga de polo a polo sobre el eje $x-y$.

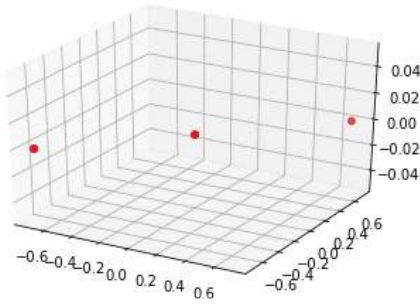
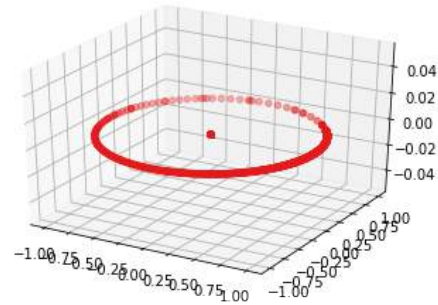
```
1 fig=plt.figure()
2 for i in range(1,10):
3     ax = plt.gca() #gca=obtiene los ejes actuales
4     A,B,C=imp((np.pi)*i) #rotamos un ngulo Pi
5
6 ax=fig.add_subplot(111,projection='3d')
7 ax.scatter3D(Ax,By,Cz,c=Cz,cmap='Set1')
8 plt.show()
```

Tener en cuenta que así se reiteren n rotaciones, siempre veremos como se muestra a continuación:

Cabe resaltar que los vectores están normalizados. En el caso particular, tenemos un vector con coordenadas $(\frac{1}{\sqrt{2}}, \frac{1}{\sqrt{2}}, 0)$.

Mirando el simulador:

El mismo vector usando puertas diferentes se puede rotar sucesivamente:

Figura 8: 1 Rotación de V .Figura 10: Rotaciones de V .Figura 9: Compuerta de Pauli Y . Rotación de V .

La razón de estas rotaciones se justifica en la búsqueda de la probabilidad de obtener uno de los estados del qubit. Es claro que así de simple no se consigue esto, por el contrario, surge como un problema; sin embargo, es la forma como se interpretaría un qubit y claro, ahora nos interesa saber por qué las probabilidades son tan importantes para resolver problemas y por qué esto supone que una máquina haría millones de operaciones más que una máquina de turing convencional. Hasta aquí, nos emociona hablar de cuaterniones y qubits, pero veamos que una de las más grandes aplicaciones de las rotaciones de un vector en esta esfera radica en la criptografía.

7. Aplicaciones.

De antemano, debemos empezar a pensar que un bit es un estado que lo define el paso o no de electricidad por un cable de un transistor de un computador. Cuando se ejecuta un algoritmo, el paso de electricidad actúa mediante un circuito electrónico o lógico bajo ciertas compuertas y que permiten la interacción con el ordenador. O sea que si se quiere evaluar una función con 10 valores, se tendría que hacer un circuito donde pase la función y genere un resultado para cada valor de forma independiente y serial. Ahora, si lo quisiéramos hacer en paralelo, necesitaríamos por cada valor una función, un circuito, o diferentes ordenadores y eso significaría ineficiencia en tiempo y recursos. Con esto en mente, la importancia de la superposición de estados de la computación cuántica es que se contempla la totalidad de estados posibles entre cero y uno, y en esta misma función se evaluarían los 10 valores de forma simultánea. Lo interesante es que cuando se evalúa una función, esta es sometida a un sistema de probabilidades que depende del colapso de la función de onda y por lo tanto hay unos estados que se vuelven más probables que otros.

A nivel de algoritmos, hay problemas que cuestan demasiado resolverlos como el hecho de factorizar un número. En respuesta a esto, y al desarrollo de teoría cuántica

computacional de las últimas dos décadas, Peter Shor, un matemático estadounidense y actualmente profesor del MIT,⁵ diseñó un algoritmo pensando en la eficiencia computacional de una máquina de Turing comparada con la de una máquina cuántica. Este algoritmo cimienta su potencia en determinar el período de una función mediante el uso de teoría de números. Aunque su estudio presenta un grado de complejidad relativamente alto, es muy interesante analizar el nuevo enfoque que la mecánica cuántica ofrece para solucionar el problema de factorización. [Cai10] Antes de explicar brevemente este algoritmo, debemos tener en cuenta que en la teoría cuántica computacional, se tienen por ahora tres clases de algoritmos:

- Algoritmos basados en la transformada cuántica de Fourier.
- Algoritmos de búsqueda. (Ejemplo: Algoritmo de Grover)
- Algoritmos de simulación de experimentos físicos

De la clasificación anterior, el algoritmo de Shor se basa en la transformada cuántica de Fourier, y esto es básicamente porque al tenerse tantas operaciones, aplicar la transformada ayuda a eliminar ruido, es decir, reprocesos. En este ejemplo simple de aplicación, no nos centraremos en indagar el funcionamiento exacto de la transformada en un circuito cuántico sino simplemente explicar por qué las rotaciones y las probabilidades son una ventaja en la computación cuántica respecto a la clásica.

En términos simples, el algoritmo de Shor busca factorizar un número natural. Es decir, dado $n \in \mathbb{N}$, queremos encontrar p, q primos $\in \mathbb{N}$ tal que $p \cdot q = n$. Este interés surge porque muchos sistemas de encriptación trabajan con números muy grandes producto de dos o más primos que se convierten en las claves de encriptación y desencriptación. Así, si logramos factorizar estos enteros, tendríamos la posibilidad de romper estos sistemas y tendríamos una gran oportunidad de hacernos con el control del sistema económico mundial, los bancos, sistemas de seguridad o hasta del mercado de las criptomonedas; sin embargo, esto no es tan fácil puesto que ya hay teoría cuántica desarrollada para evitar estas situaciones.

⁵Massachusetts Institute of Technology.

A continuación, introducimos algunos teoremas necesarios para la comprensión del algoritmo:

Teorema 7.1. Sean $n, n+1$ dos enteros positivos, entonces $\text{mcd}(n, n+1) = 1$.

Esto es gracias a las propiedades de la divisibilidad, ya que si $a, b, c \in \mathbb{Z}$ y $c|a \wedge c|b \rightarrow c|ax + by, x, y \in \mathbb{Z}$. En particular, si $x = -1, y = 1$ entonces $c|(n+1) - n$ y esto es posible si $c = 1$.

Teorema 7.2. Sea $n > 1 \in \mathbb{N}$, entonces $\forall n, n$ posee una única factorización (salvo el orden) en números primos, es decir $\exists m \in \mathbb{N}$ tal que $n = p_1 p_2 \dots p_m = \prod_{i=1}^m p_i$ con p_i primo y $1 \leq i \leq m$ y $p_1 \leq p_2 \leq \dots \leq p_m$.

Teorema 7.3. Dado $n \in \mathbb{N}$, $x \equiv_n y$ si y sólo si $n|x - y$.

Teorema 7.4. Dados $n, m, r \in \mathbb{Z}$, $m > 0$ entonces $m \equiv_n r_i$. Más aún, $m^{i+1} \equiv_n m \cdot r_i$.

Es decir, que todo número es congruente con su residuo módulo n . De aquí se deduce que si existe $r_j = 0$, entonces $r_i = 0 \forall i > j$. O sea que cuando un resto sea cero, a partir de allí, todos los restos potenciales en adelante serán cero. De igual manera, a partir del primer resto que se repita, toda la secuencia de residuos previos se reproduce en igual orden infinitamente.

Con lo anterior es suficiente para entender el algoritmo. Ahora bien, ¿qué pretende Shor?, bueno, en primer lugar, introduce el concepto de la raíz cuadrada no trivial módulo n . Es decir, si $x \in \mathbb{N}$, entonces x es una raíz cuadrada no trivial módulo n si $x^2 \equiv_n 1$ y $x \in [2, n-2]$. Para entender el por qué de este intervalo, podemos hacer uso de 7.3 de la siguiente forma:

Si $x^2 \equiv_n 1$, entonces $x^2 - 1 \equiv_n 0$, de aquí se tiene que $(x+1)(x-1) = kn$ ya que $n|x^2 - 1$. Así, por 7.1 ni $(x+1)$ ni $(x-1)$ pueden ser $n-1$ ni n . Por lo tanto, $2 \leq x < n-2$. El algoritmo supone que encontrar la raíz cuadrada no trivial módulo n equivale a encontrar los números primos p, q tales que factorizan n . Supongamos ahora que los factores $(n+1), (n-1)$ poseen los siguientes factores:

$$P = \{p_i \in \mathbb{N} : p_i|(x+1)\}$$

$$Q = \{q_i \in \mathbb{N} : q_i|(x-1)\}$$

Se deduce que p, q pertenecen a P, Q respectivamente y no existe la posibilidad de que p, q estén en el mismo factor ya sea $(x+1)$ o $(x-1)$. De esta forma, $\text{mcd}((x+1), n) = p, \text{mcd}((x-1), n) = q$. Por lo tanto, si encontramos x y calculando el máximo común divisor como acabamos de señalar, habremos encontrado p, q y el número quedaría factorizado. Pero bien, ¿cómo encontrar x ? Bueno, para ello recordamos que introducimos el teorema 7.4. Usémoslo en el siguiente ejemplo para entenderlo:

Supongamos que queremos encontrar el orden r módulo n . Siendo $m = 11, n = 7$. Buscamos que la función $f(i) = 11 \bmod(7)$ sea periódica.

$$i = 0 \rightarrow 11^0 \equiv_7 1.$$

$$i = 1 \rightarrow 11^1 \equiv_7 4.$$

$$i = 2 \rightarrow 11^2 \equiv_7 11 \cdot 4 \equiv_7 2.$$

$$i = 3 \rightarrow 11^3 \equiv_7 11 \cdot 2 \equiv_7 1.$$

Como el uno vuelve y se repite, se arma un bucle infinito y hemos encontrado el período u orden de la función f , en este caso es 3. Muy bien, ahora podríamos pensar en coger un número al azar entre 2 y $n-2$, construimos la función y calculamos el período esperando que el período sea en algún momento 2 tal y como lo define Shor para el cálculo de la raíz cuadrada no trivial módulo n $x^2 \equiv_n 1$. Para este caso, hemos visto que nuestro período u orden fue 3, por lo tanto no tuvimos suerte de encontrar 2 y este caso no nos serviría para descomponer un número. Pero ¿qué pasa si buscamos el período de $f(i) = 3^i \bmod(7)$?

Bueno, replicando el anterior ejercicio nos daríamos cuenta que el período es 6 y no 2; sin embargo, si esto lo expresamos como $(3^4)^2 \equiv_7 1$ claramente habremos encontrado la raíz buscada que es 3^4 y el problema se reduciría a simplemente encontrar x tales que su período sea un número par. Por lo anterior, Shor define el siguiente teorema:

Teorema 7.5. Dado un número aleatorio $x \in [2, n-2]$, la probabilidad de que x tenga período par es mayor a 0,5.

Veamos ahora sí, cómo descomponer 15 en factores de números primos. Puesto que $n = 15$, entonces por el teorema de Shor, escogemos un número al azar $x \in (2, 13)$. Escogemos 4 y buscamos el período:

$$i = 0, 4^0 \equiv_1 51.$$

$$i = 1, 4^1 \equiv_1 54.$$

$$i = 2, 4^2 \equiv_1 51.$$

Así, la raíz cuadrada no trivial módulo n sería 4. Ahora, calculamos $\text{mcd}((4+1), 15)$ y $\text{mcd}((4-1), 15)$, dando como resultados 5 y 3 respectivamente.

En este punto, claramente todo este algoritmo se puede hacer en un ordenador clásico; sin embargo, cuando el número es muy grande, probar cada x al azar entre 2 y $n-2$ es una tarea titánica y muy difícil de resolver en tiempo polinomial, o sea, debería tomar cada elemento del intervalo, calcular sus potencias, verificar sus residuos, y eso se haría en forma serial para finalmente compararlos. El algoritmo se diseñó para computarlo en un ordenador cuántico, porque esto significa que si ponemos un valor x que tenga período par y otro que sea impar, entonces el sistema va a colapsar hacia el x que tenga probabilidad más alta de ser período par. Además, este proceso de tomar una x , evaluarla en una cantidad finita de números, evaluar sus potencias y residuos, se realiza mediante la transformada cuántica de Fourier usando los principios de superposición y paralelismo. Finalmente, en cuanto a tiempo polinomial se refiere, en computación clásica, solo hacer el paso de la transformada de Fourier toma $n2^n$ pasos para 2^n datos, mientras que en computación cuántica tan solo toma n^2 pasos. O sea que pasamos de un costo exponencial a polinomial, haciendo que problemas que antes se solucionaban en siglos, en un futuro se puedan resolver en minutos.

8. Conclusiones.

El conjunto de los cuaterniones es el conjunto de los denominados hipercomplejos que representa un espacio vectorial.

rial abstracto bajo la operación suma usual de complejos y multiplicación escalar definidas en la identidad descubierta por Hamilton. Abstraer el movimiento de vectores unitarios alrededor de un eje en el espacio tridimensional nos llevó a presentar los espacios pre Hilbertianos y posteriormente los espacios de Hilbert, que por sus características matemáticas, son muy ricos sobre todo para trabajar con vectores que representan estados de un sistema físico. Encontramos que las operaciones realizadas sobre un qubit que es la mínima unidad de información de un computador cuántico, representa toda la lógica que hoy podemos asociar con la lógica de las computadoras digitales de la computación tradicional. Estas operaciones a su vez se representan gráficamente mediante una esfera de radio 1 cuyo vector de dimensión dos rota sobre un plano asociándole a cada posible punto del mismo una probabilidad de ocurrencia. Finalmente, vimos que rotar este vector tiene aplicaciones significativas en la criptografía, porque es así como se determina el período de la función o la raíz cuadrada no trivial módulo n mediante el uso del algoritmo de Shor.

Agradecimientos.

Al profesor John Alexander Arredondo, por la dinámica implementada en la clase de álgebra lineal incentivando la investigación. Además, por su disposición y sus valiosos comentarios en cada revisión para hacer posible esta publicación.

Referencias

- [Rod12] V Rodríguez, *Sobre los cuaterniones, álgebras de Lie, y matrices de Pauli. Teoría básica y aplicaciones físicas.*, Universidad de Oviedo, 2012.
- [Bae11] J Baez, *Octoniones y teoría de cuerdas*, Investigación y Ciencia, 2011.
- [Mar16] A Markelov, *Uso de Cuaterniones para Representar Rotaciones*, Anónimo, 2016.
- [Anó19] Anónimo, *Max Planck, el padre de la teoría cuántica que intentó convencer a Hitler de que permitiera trabajar a los científicos judíos*, BBC news, 2019.
- [Anó20] ———, *El largo camino para entender la física cuántica*, OpenMindBBVA, 2020.
- [Pas19] J Pastor, *El cuaternión está más vivo que nunca: así es como la NASA y los videojuegos lo usan más de un siglo después de su descubrimiento*, Xataca, 2019.
- [Her10] C Hernando, *Algoritmo de factorización para un computador cuántico*, Edvcatio physicorvm qvo non ascendam, 2010.
- [Fav08] A Favieri, *Introducción a los Cuaterniones*, Editorial de la Universidad Tecnológica Nacional – U.T.N. - Argentina edUTecNe, Haedo, Argentina, 2008.
- [Haw10] S and Mlodinow Hawking L, *El gran Diseño*, Bantam Books, Nueva York, EEUU, 2010.
- [Cai10] H Caicedo, *Algoritmo de factorización para un computador cuántico*, Unidad Profesional “Adolfo López Mateos”, Zacatenco, Delegación Gustavo A. Madero, 2010.
- [Anó10] Anónimo, *Rotaciones y Cuaternios*, Unidad Profesional “Adolfo López Mateos”, Zacatenco, Delegación Gustavo A. Madero, 2010.
- [Lop16a] E Lopategui, *Historia de las computadoras*, 2016.
- [Lop16b] ———, *Manejo de la información y uso de la computadora.*, 2016.
- [CC20] Carlos Coy, *Implementación de Circuitos Cuánticos y Propuesta Experimental Para Computas Cuánticas*, Universidad del Valle, 2020.
- [AD15] Alejandro Díaz, *Introducción a la computación cuántica*, Universidad de Rosario, 2015.
- [VM13] Vicente Moret, *Principios Fundamentales de la Computación Cuántica.*, Universidad de la Coruña, 2013.
- [MP19] Maria Prieto, *Algoritmos Cuánticos para la Resolución del Problema de Satisfacibilidad Booleana.*, Universidad Autónoma de Madrid, 2019.

Acerca del autor: Daniel Esteban Galvis Sandoval es estudiante de maestría en inteligencia artificial y big data de la universidad de Valencia. Le gusta la musica ska, blues, reggae. Es un apasionado por la programación, las chaquetas de cuero, las teorías de vida extraterrestre, quiere aprender a tocar el piano y a preparar cocteles.

SI QUIERO INVERTIR ¿CÓMO ME DEBO CUBRIR?: UNA APROXIMACIÓN A LAS MEDIDAS COHERENTES DE RIESGO EN R Y PYTHON

Marco Blanco Ariza *

marco.blanco01@correo.usa.edu.co

1. Introducción.

La evolución de los mercados financieros, generada entre otras cosas por la aparición de instrumentos *derivados*¹, evidenció la necesidad de formalizar las *finanzas* empleando conceptos y teorías matemáticas importantes, como el cálculo y la probabilidad. Dicha sinergia tuvo dos consecuencias directas: la primera, facilitó encontrar expresiones que permitieran *valorar* dichos instrumentos y *estimar* el riesgo que supone mantener inversiones en ellos; la segunda, dotó a las instituciones reguladoras de herramientas para establecer mínimos de capital, con el propósito de que las empresas emisoras de dichos instrumentos pudieran responder ante eventuales pérdidas.

En ese sentido, se han construido medidas de riesgo financiero, entre las que se destacan el *Valor en Riesgo* o *VaR* y el *Valor en Riesgo Condicional* o *CVaR*, las cuales permiten cuantificar, estableciendo un nivel de confianza y un horizonte temporal, la exposición al riesgo de mercado. Sin embargo, como en cualquier teoría económica y financiera, es necesario establecer supuestos que faciliten representar una situación real en lenguaje matemático. En ese sentido, en el presente artículo se exponen las medidas antes mencionadas y las condiciones que se deben cumplir para que sean *coherentes*; asimismo, se presentan dos expresiones analíticas, una para el *VaR* y otra para el *CVaR*, que facilitan aterrizar los cálculos a los rendimientos diarios de la acción de Ecopetrol entre el 01 de enero de 2018 y el 18 de marzo de 2022.

2. Riesgo, tipos de riesgo y medidas coherentes de riesgo.

A continuación se presenta una definición de riesgo y los tipos de riesgo que componen el conjunto de *riesgos financieros*, con el fin de introducir al lector en los conceptos fundamentales. Asimismo, se expone formalmente el concepto de *medida de riesgo* y los atributos que debe tener para ser *coherente*.

*Estudiante de Matemáticas, Universidad Sergio Arboleda.

¹Son aquellos instrumentos financieros cuyo valor depende de otro activo que se denomina *subyacente*. Algunos ejemplos de subyacentes son: acciones, bonos, commodities, entre otros.

2.1. Riesgo y tipos de riesgo.

El riesgo, según la Real Academia Española, se define como la “contingencia o proximidad de un daño” [RAE19]. Por otra parte, en [Kni21] se diferencia el riesgo de la incertidumbre, debido a que al riesgo se le puede asociar una probabilidad pero a la incertidumbre no. Con base en lo anterior, el riesgo puede definirse como la probabilidad de que ocurra un evento dañino; en el ámbito financiero, dicho daño se entiende como la *pérdida* de valor de un activo².

En ese sentido, en [AB10] se expone que existen diferentes fuentes de riesgo, lo cual se deriva del hecho de que diferentes variables, tanto internas (por ejemplo: la estrategia de ventas) como externas (por ejemplo: la inflación), pueden afectar negativamente el valor de una entidad o de una inversión. Dentro del conjunto de riesgos, existen tres que componen el riesgo financiero³:

- **Riesgo de mercado:** hace referencia a la pérdida de valor de un activo como consecuencia de las fluctuaciones que pueden experimentar los precios de mercado.
- **Riesgo de liquidez:** hace referencia al riesgo que supone realizar transacciones en un mercado con liquidez baja (bajo volumen de transacciones).
- **Riesgo de crédito:** hace referencia al riesgo que se deriva del impago total o parcial por parte de un deudor.

Teniendo en cuenta lo anterior, el presente artículo se centra en *medir* la exposición al *riesgo de mercado*, es decir, la pérdida máxima de valor que puede experimentar un activo en un horizonte de tiempo, debido a fluctuaciones en sus precios de mercado.

2.2. Medidas coherentes de riesgo.

Los sistemas de administración de riesgo de mercado⁴, en general, se componen de cuatro etapas: identificación, medición, control y monitoreo. Dichas etapas siguen un proceso

²Algunos ejemplos de activos son: acciones, bonos, derivados, CDT, cuentas por cobrar, entre otros.

³Para mayor detalle consultar [AB10].

⁴Para el caso colombiano, consultar el capítulo XXI de la Circular Externa 100 de 1995 de la Superintendencia Financiera de Colombia.

secuencial, es decir, no se puede medir algo que previamente no ha sido identificado, así como tampoco se puede controlar algo que no ha sido medido, por lo cual cada etapa constituye un pilar fundamental dentro del proceso.

Para comprender la importancia de medir adecuadamente el riesgo, suponga que una compañía mide su exposición al riesgo de mercado en \$100 y que ocurre alguno de los siguientes escenarios: primero, el riesgo se materializa y la compañía pierde \$150; segundo, el riesgo se materializa y la compañía pierde \$35. En el primer escenario la compañía no contaba con los recursos para atender la pérdida, por lo cual afectó su rentabilidad más allá de lo pronosticado y comprometió su estabilidad financiera en el corto plazo; en el segundo escenario, la compañía inmovilizó una cantidad considerable de recursos que excedieron la pérdida real, lo cual generó para la compañía un costo de oportunidad, ya que los \$65 que inmovilizó de más los hubiese podido invertir para obtener una rentabilidad adicional.

Del ejemplo anterior se concluye que la medición del riesgo que hizo la compañía fue errónea, ya que en el primer escenario la subestimó y en el segundo escenario la sobreestimó considerablemente. En ese sentido, es importante medir adecuadamente el riesgo y para ello se deben emplear medidas *coherentes* de riesgo.

Formalmente, una medida de riesgo es *el resultado de un mapeo sobre una variable aleatoria X , definida sobre un espacio de probabilidad fijo $(\Omega, \mathcal{F}, \mathbb{P}_\theta)$, donde θ es un vector de parámetros asociados con la distribución de X , que tiene como resultado $\rho = \mathcal{A} \rightarrow \mathbb{R}$, donde $\mathcal{A} = \{X|X : \Omega \rightarrow \mathbb{R}\}$ representa el conjunto de retornos posibles de un portafolio o un activo [Gar15]. No obstante, dicha medida debe cumplir cuatro axiomas para ser *coherente* [ADEH99]:*

- **Homogeneidad positiva:** $\rho(\lambda u) = \lambda \rho(u)$
Este axioma garantiza que si se incrementa en λ la posición en un activo, entonces el riesgo aumenta proporcionalmente en λ .
- **Monotonicidad:** $u \leq v$ implica $\rho(u) \leq \rho(v)$
Este axioma garantiza que se cumpla el principio *riesgo/rentabilidad*, es decir, una mayor rentabilidad implica un mayor riesgo.
- **Invariante en traslación:** $\rho(u + a) = \rho(u) + a$
Este axioma garantiza que al adicionar a la posición una cantidad inicial segura a el riesgo se disminuye en dicho monto adicional.
- **Subaditividad:** $\rho(u + v) \leq \rho(u) + \rho(v)$
Este axioma establece que el riesgo del portafolio no aumenta por la composición del portafolio. En otras palabras, se espera que con la inclusión de activos en el portafolio el riesgo disminuya o, a lo sumo, permanezca constante.

Gracias a los cuatro axiomas enunciados, se garantiza que las mediciones del riesgo sean coherentes, lo cual permite controlarlos de una manera adecuada y responsable. Es importante

resaltar que el ser *coherente* no es sinónimo de *exacta*, por lo cual, en la mayoría de los casos, existirán desviaciones entre las pérdidas que se materializan y la medida de riesgo obtenida, es decir, la medida de riesgo es una *estimación*, y por lo tanto, el objetivo de que una medida de riesgo sea coherente es aproximar dicha estimación, en la medida de lo posible, a la realidad. Con base en lo anterior, a continuación se presentan dos medidas de riesgo de mercado: el *Value at Risk* - *VaR*, la cual sólo es *coherente* bajo ciertas distribuciones de probabilidad; y el *Conditional Value at Risk* - *CVaR*, la cual se calcula a partir del *VaR* y es una medida *coherente* de riesgo.

3. Valor en riesgo (*VaR*).

El valor en riesgo o *VaR* (por las iniciales del término *Value at Risk*) es una medida de riesgo ampliamente utilizada por los participantes del mercado financiero, bien sean inversionistas, emisores, entidades reguladoras, entre otros. La popularidad de esta medida se debe a que intenta resumir el riesgo total de un activo, o una cartera de activos, en una sola cifra, lo cual resulta útil para la toma de decisiones ya que, en esencia, da cuenta de la pérdida máxima que podría experimentar un activo, o una cartera de activos, en un horizonte de tiempo y un nivel de confianza establecidos.

Formalmente, *el valor en riesgo de X al nivel de confianza $1 - q$ denotado por $-VaR_{1-q}^X$, se define como el peor valor del activo (o portafolio), en un período de tiempo dado $[t, T]$, para un intervalo de confianza del $(1 - q)100\%$ (ver [Ven06]). Es decir:*

$$\mathbb{P}_\theta(-VaR_{1-q}^X \leq X) = 1 - q \quad \text{y} \quad \mathbb{P}_\theta(X \leq -VaR_{1-q}^X) = q.$$

Dado lo anterior, es posible definir el *VaR* mediante el *supremo*⁵ o el *ínfimo*⁶ de los siguientes conjuntos:

$$\begin{aligned} VaR_{1-q}^X &= -\inf\{x \in \mathbb{R} \mid \mathbb{P}_\theta(X \leq x) \geq q\} \\ &= -\sup\{x \in \mathbb{R} \mid \mathbb{P}_\theta(X \leq x) \leq q\}. \end{aligned} \quad (1)$$

De manera equivalente:

$$VaR_{1-q}^X = -\inf\{x \in \mathbb{R} \mid \mathbb{P}_\theta(X > x) \leq 1 - q\}. \quad (2)$$

Dado que el conjunto de referencia es $\mathbb{R}$, el *ínfimo* y el *supremo* existen debido a que los conjuntos en (1) y (2) no son vacíos y están acotados⁷.

⁵La menor de las cotas superiores.

⁶La mayor de las cotas inferiores.

⁷Al respecto consultar [Rud76] - axioma del supremo.

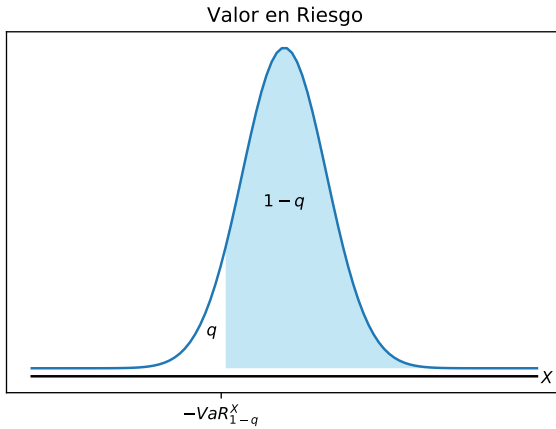


Figura 1: VaR al $(1 - q)$ % de confianza.

A manera de resumen, la Figura 3.1 permite ver gráficamente tanto la expresión (1) como la (2) así: la región no sombreada corresponde al conjunto de todos los x tal que la probabilidad de que la variable aleatoria X tome un valor menor o igual que x , es menor o igual que q , es decir, $\{x \in \mathbb{R} | \mathbb{P}_\theta(X \leq x) \leq q\}$. Ahora note que cualquier valor a la derecha del $-VaR_{1-q}^X$ (incluyéndolo), es una cota superior de dicho conjunto, por lo cual, al tomar la más pequeña de ellas obtenemos la expresión (1).

Por otra parte, la región sombreada corresponde al conjunto de todos los x tal que la probabilidad de que la variable aleatoria X tome un valor mayor o igual que x , es menor o igual que $1 - q$, es decir, $\{x \in \mathbb{R} | \mathbb{P}_\theta(X > x) \leq 1 - q\}$. Ahora note que cualquier valor a la izquierda del $-VaR_{1-q}^X$ (incluyéndolo), es una cota inferior de dicho conjunto, por lo cual, al tomar la más grande de ellas obtenemos la expresión (2).

Una vez definido el VaR , es importante tener *expresiones* que faciliten su cálculo. En ese sentido, existen tres enfoques para ello: paramétrico, histórico y Monte Carlo. Para el presente artículo se empleará un enfoque paramétrico y se supondrá que los rendimientos financieros siguen una distribución normal de probabilidad.

3.1. Método paramétrico y distribución Normal.

Como se infiere de lo expuesto anteriormente, el valor en riesgo es aquel valor que divide el conjunto de rendimientos de un activo en dos: $\{x \in \mathbb{R} | \mathbb{P}_\theta(X \leq x) \leq q\}$ y $\{x \in \mathbb{R} | \mathbb{P}_\theta(X > x) \leq 1 - q\}$. Por lo tanto, el VaR_{1-q}^X representa un *cuantil* de la distribución de rendimientos financieros. De esta manera, y partiendo de la función de distribución acumulada (*cdf*) de una distribución normal de probabilidad, se obtiene la siguiente expresión para el $VaR - normal$:

$$VaR_q(X) = \mu + \sigma z_q. \quad (3)$$

En efecto, si tomamos la *cdf* de una distribución normal:

$$F_X(x) = \int_{-\infty}^x f_X(u) du = \int_{-\infty}^x \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{u-\mu}{\sigma}\right)^2} du.$$

Sea:

$$z = \frac{u - \mu}{\sigma}, \quad (4)$$

luego:

$$F_X(x) = \int_{-\infty}^{\frac{x-\mu}{\sigma}} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}z^2} dz.$$

Si $F_X(x_q) = q$, entonces:

$$F_X(x_q) = \Phi\left(\frac{x_q - \mu}{\sigma}\right) = q,$$

donde $\Phi(\cdot)$ es la *cdf* de una distribución normal estándar. Despejando x_q , se llega a:

$$x_q = \mu + \sigma z_q. \quad (5)$$

La expresión (5) corresponde al $VaR_q(X)$. A manera de ejemplo, se calculará el VaR para los rendimientos diarios (r_t) de la acción de Ecopetrol entre el 1 de enero de 2018 y el 18 de marzo de 2022. Para el cálculo de dichos rendimientos se empleó la siguiente expresión:

$$r_t = \frac{p_t - p_{t-1}}{p_{t-1}}. \quad (6)$$

Donde p_t corresponde al precio de cierre de un activo financiero en el momento t y p_{t-1} al precio de cierre en el momento inmediatamente anterior $t - 1$. Para facilitar la obtención y el tratamiento de los datos, se emplearon R y Python, siguiendo los siguientes pasos:

Paso 1a: importación de los paquetes en **R**.

```
install.packages('quantmod')
library('quantmod')
```

Paso 1b: importación de los paquetes en **Python**.

```
import numpy as np
import pandas as pd
import pandas_datareader
from pandas_datareader import data as wb
from pylab import *
import scipy.stats as st
```

Paso 2a: obtención de los datos de *Yahoo finance* con **R**.

```
getSymbols('EC', from = '2018-01-1')
```

Paso 2b: obtención de los datos de *Yahoo finance* con **Python**.

```
precio_Ecopetrol =
wb.DataReader('EC',
data_source = 'yahoo',
start = '01-01-2018',
end = '18-03-2022')
```

Paso 3: se asignan los datos al objeto *Ecopetrol* y se eliminan datos faltantes en **R**.

```
Ecopetrol = Ad(EC)
Ecopetrol = na.omit(Ecopetrol)
```

Paso 4: se crea una función que calcula los retornos diarios en **R**.

```
retornos_aritmeticos = function(datos)\{
  return(datos/lag(datos, k = 1) - 1)\}
```

Paso 5a: se calculan los retornos diarios en **R**.

```
Rendimientos = na.omit(retornos_aritmeticos(Ecopetrol))
```

Paso 5b: se calculan los retornos diarios en **Python**.

```
retorno_Ecopetrol = ((precio_Ecopetrol['Adj
Close']/precio_Ecopetrol['Adj
Close'].shift(1))-1).dropna()
```

Paso 6a: se grafican los rendimientos financieros de *Ecopetrol* en **R**.

```
plot(Rendimientos, main = "Rendimientos
Ecopetrol", col = 'blue')
```

Paso 6b: se grafican los rendimientos financieros de *Ecopetrol* en **Python**.

```
precio_Ecopetrol['retorno_Ecopetrol'].plot(
  figsize=(9, 6.5), color = "deepskyblue",
  label = 'Retorno diario')
plt.legend(loc='lower right')
plt.grid(color='grey',
  linestyle='dotted',
  linewidth=0.6)
plt.title('Retornos diarios Ecopetrol')
```

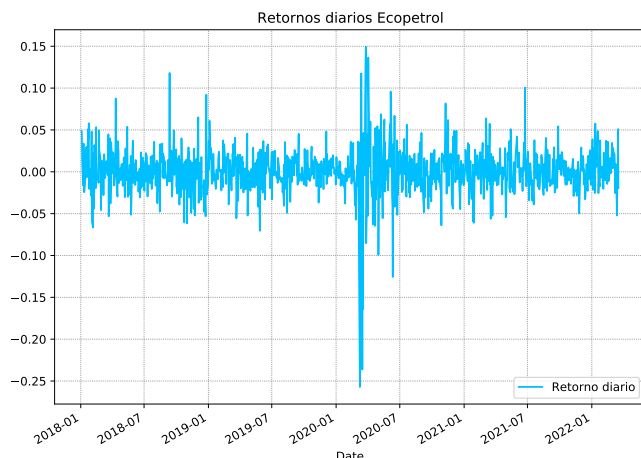


Figura 2: Rendimientos Financieros de *Ecopetrol*.

Paso 7: se crea una función que gráfica el histograma de los retornos, la función de densidad empírica (natural de los datos) y la función de densidad normal en **R**.

```
histograma <- function(Rendimientos, nbreks=30){
  hist(Rendimientos, breaks=nbreks, freq=FALSE,
  main="Ecopetrol")
  rug(jitter(Rendimientos), col="purple")
  curve(dnorm(x,mean=mean(Rendimientos),
  sd=sd(Rendimientos)), add=TRUE, col="blue", lwd=2)
  lines(density(Rendimientos)$x,
  density(Rendimientos)$y, col="green", lwd=2, lty=1)
  legend("topright",inset = c(-0.15, 0),
```

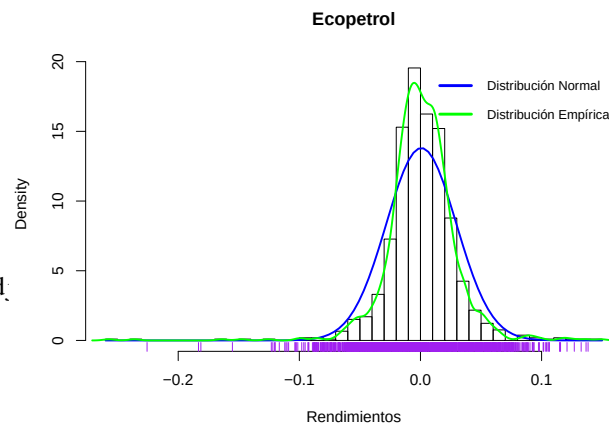


Figura 3: Densidad Empírica y Normal de los Rendimientos Financieros de *Ecopetrol*.

En este paso es importante notar que los rendimientos financieros presentan exceso de curtosis (colas pesadas). Al comparar la distribución empírica con la distribución normal, se evidencia que la gráfica verde no se ajusta a la gráfica azul, lo cual permite inferir que los rendimientos financieros no siguen una distribución normal de probabilidad. Por otra parte, es importante notar que los puntos de color morado en la Figura 3 reflejan la concentración de los datos (rendimientos), en ese sentido, los puntos aislados representan rendimientos extremos (ganancias o pérdidas) que por definición, el *VaR* no tiene en cuenta en su totalidad.

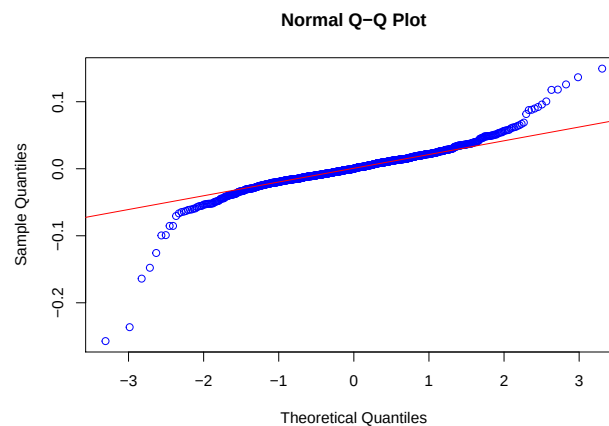


Figura 4: Prueba de Normalidad de los Rendimientos Financieros.

Paso 8: se realiza la gráfica Q-Q plot en R.

```
qqnorm(Rendimientos)
qqline(Rendimientos, col = "red")
```

A partir de la Figura 4, se demuestran dos hechos importantes: que los rendimientos financieros no siguen una distribución normal de probabilidad y que dichos datos tienen colas pesadas. Cabe resaltar que existen datos que no están sobre la línea roja, lo cual indica que son datos extremos que no siguen una distribución normal.

Cabe resaltar que una de las características de la distribución normal es que es simétrica, lo cual implica que la frecuencia de los eventos centrales es mayor a la de los eventos extremos. Al respecto, se ha demostrado que el comportamiento de los rendimientos de activos en mercados financieros no es normal, evidenciando la existencia de colas pesadas y asimetrías negativas [CB12]. La existencia de asimetrías negativas implica la existencia de altas probabilidades de que ocurran *pérdidas excesivas o extremas* [Her13]. Por lo tanto, se deben emplear distribuciones de colas pesadas, con el fin de conformar portafolios óptimos que tengan en cuenta los posibles eventos extremos [HdV07].

Paso 9a: cálculo del VaR para Ecopetrol en R.

```
VaR_Ecopetrol = mean(Rendimientos) +
sd(Rendimientos)*qnorm(0.95)
```

```
VaR_Ecopetrol*100
```

Paso 9b: cálculo del VaR para Ecopetrol en Python.

```
VaR_normal = (retorno_Ecopetrol.mean
+retorno_Ecopetrol.std()*
st.norm.ppf(0.95))
VaR_normal*100
```

Este cálculo permite concluir que el VaR para Ecopetrol es de 4.82 %, es decir, con el 95 % de confianza se concluye que la máxima pérdida que podría experimentar la acción de Ecopetrol es de 4.82 %. Sin embargo, con el 5 % de significancia se concluye que la mínima pérdida que podría experimentar dicha acción sería de 4.82 %. En este paso es natural que el lector se pregunte: ¿qué sucede con los valores que puede tomar la variable aleatoria (*rendimientos*) que exceden el VaR?.

La respuesta a la pregunta planteada es que el VaR no da cuenta del riesgo que lo excede, es decir, no le da importancia al comportamiento de la variable aleatoria en las colas de la distribución. Asimismo, cabe resaltar que el VaR es coherente en el caso en que los rendimientos financieros siguen una distribución elíptica de probabilidad (por ejemplo: la normal o la t-Student), en otro caso el VaR no es subaditivo y por lo tanto deja de ser coherente (ver [ADEH99]).

Debido a estas debilidades que presenta el VaR, en [AT02] se propone una medida de riesgo que las supera y, por lo tanto, es coherente, la cual denominan Expected Shortfall (ES) o Valor en Riesgo Condicional (CVaR).

4. Valor en riesgo condicional (CVaR).

El ES, también conocido como “Valor en Riesgo Condicional” (CVaR), “Pérdida Esperada en la Cola” (ETL, por sus siglas en inglés “Expected Tail Loss”) o “Valor Medio en Riesgo” (AVaR, por sus siglas en inglés “Average Value at Risk”) [Jim14]. Son medidas que buscan tener en cuenta las pérdidas que están más allá del VaR. Formalmente, en [Jim14] se define el ES como:

$$ES_q(X) = CVaR_q(X) = \mathbb{E}[X|X > VaR_q(X)]. \quad (7)$$

Es decir, el CVaR se calcula como el valor esperado de los rendimientos que son mayores que el VaR_{1-q}^X o, de manera equivalente, menores que el $-VaR_{1-q}^X$. Para hallar una expresión paramétrica del CVaR, se sigue el siguiente procedimiento.

Empleando la definición de CVaR, la expresión (5) y la función de densidad de probabilidad (*pdf*) de una distribución normal, se obtiene:

$$\begin{aligned} CVaR_q &= \frac{1}{1-q} \int_{x_q}^{\infty} u f_X(u) du \\ &= \frac{1}{1-q} \int_{\mu+\sigma z_q}^{\infty} \frac{u}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{u-\mu}{\sigma}\right)^2} du. \end{aligned} \quad (8)$$

Sustituyendo las expresiones (4) y (??) en (8),

$$\begin{aligned} CVaR_q &= \frac{1}{1-q} \int_{z_q}^{\infty} \frac{\mu + \sigma z}{\sqrt{2\pi}} e^{-\frac{1}{2}z^2} dz, \\ &= \frac{1}{1-q} \int_{z_q}^{\infty} \left(\frac{\mu}{\sqrt{2\pi}} + \frac{\sigma z}{\sqrt{2\pi}} \right) e^{-\frac{1}{2}z^2} dz, \\ &= \frac{1}{1-q} \left(\int_{z_q}^{\infty} \frac{\mu}{\sqrt{2\pi}} e^{-\frac{1}{2}z^2} dz + \int_{z_q}^{\infty} \frac{\sigma z}{\sqrt{2\pi}} e^{-\frac{1}{2}z^2} dz \right), \\ &= \frac{1}{1-q} \left(\mu \int_{z_q}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}z^2} dz + \sigma \int_{z_q}^{\infty} \frac{z}{\sqrt{2\pi}} e^{-\frac{1}{2}z^2} dz \right), \\ &= \frac{1}{1-q} \left(\mu(1 - \Phi(z_q)) + \sigma \int_{z_q}^{\infty} \frac{z}{\sqrt{2\pi}} e^{-\frac{1}{2}z^2} dz \right), \\ &= \frac{1}{1-q} \left(\mu(1 - q) - \frac{\sigma}{\sqrt{2\pi}} \int_{z_q}^{\infty} -ze^{-\frac{1}{2}z^2} dz \right), \\ &= \mu - \frac{\sigma}{1-q} \frac{1}{\sqrt{2\pi}} \left(e^{-\frac{1}{2}z^2} \right) \Big|_{z=z_q}^{z \rightarrow \infty}, \\ &= \mu + \frac{\sigma}{1-q} \frac{1}{\sqrt{2\pi}} \left(e^{-\frac{1}{2}z_q^2} \right), \\ &= \mu + \frac{\sigma}{1-q} \varphi(z_q), \end{aligned} \quad (9)$$

donde $\varphi(\cdot)$ denota la *pdf* de una normal estándar⁸. Continuando con el ejemplo de la acción de Ecopetrol, el CVaR al 95 % de confianza se calcula a partir de la expresión (9), así:

Paso 10a: cálculo del CVaR para Ecopetrol en R.

⁸Tener en cuenta que: $\int_{-\infty}^{x_q} f_X(u) du + \int_{x_q}^{\infty} f_X(u) du = 1$, $\int_{x_q}^{\infty} f_X(u) du = 1 - \int_{-\infty}^{x_q} f_X(u) du$ y $\int_{x_q}^{\infty} f_X(u) du = 1 - \Phi(z_q)$ donde $\Phi(\cdot)$ es la *cdf* de una distribución normal estándar.


```
CVaR_Ecopetrol = mean(Rendimientos) +
    sd(Rendimientos)
    *pnorm(0.95)*(1/(1-0.95)) CVaR_Ecopetrol*100
```

Paso 10b: cálculo del *CVaR* para Ecopetrol en **Python**.

```
CVaR_normal = (retorno_Ecopetrol.mean()+
    ((retorno_Ecopetrol.std()/(1-0.95))*
    st.norm.pdf(st.norm.ppf(0.95)))
CVaR_normal
```

Una vez computado el *CVaR* se concluye que, con 95 % de confianza, la pérdida esperada que podría experimentar la acción de Ecopetrol es de 6.04 %. Es importante notar que el $CVaR \geq VaR$, debido a que el *CVaR* tiene en cuenta pérdidas que exceden al *VaR*, por lo cual es de esperarse que el *promedio* aumente. Lo anterior tiene implicaciones importantes, tanto para los inversionistas como para los reguladores, debido a que permite establecer coberturas al riesgo de mercado adecuadas que generen confianza en la sociedad; establecer coberturas basadas en estimaciones incorrectas pone en riesgo a las organizaciones que las emplean, ya que pueden establecer coberturas que sean insuficientes a la hora de atender una pérdida excesiva.

5. VaR vs. CVaR: ¿cuál es mejor?

Tanto el *VaR* como el *CVaR* representan medidas de riesgo, aunque la primera, como se indicó anteriormente, no es coherente y la segunda sí. El efecto de la coherencia de las medidas de riesgo, a juicio del autor, se evidencia en el nivel de *ajuste* que tiene la información histórica con base en la cual se calcularon. En otras palabras, qué tanta información está teniendo en cuenta la medida de *VaR* o *CVaR* del total de rendimientos financieros históricos.

Para estimar dicho *ajuste* se propone una metodología *gráfica* y otra *cuantitativa*, las cuales se presentan a continuación:

5.1. Método gráfico.

Una vez realizado el cálculo del *VaR* y el *CVaR* para la acción de Ecopetrol se procede a graficarlos, junto con los rendimientos financieros de la acción, tal como se evidencia en la Figura 5:

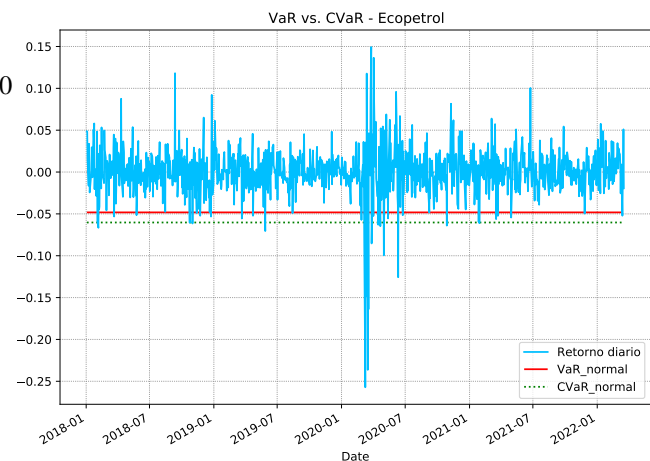


Figura 5: Rendimientos Financieros de Ecopetrol.

El *VaR* (color rojo) y el *CVaR* (color verde) se grafican como líneas rectas, con el fin de observar si los rendimientos financieros diarios (color azul) los exceden. En ese sentido, cada día en que la recta roja es intersecada por la gráfica azul se produjo una pérdida que excedió o fue igual al *VaR*. Del mismo modo, cada día en que la recta punteada verde es intersecada por la gráfica azul se produjo una pérdida que excedió o fue igual al *CVaR*.

Por lo tanto, gráficamente se aprecia que la medida *CVaR* presenta mejor ajuste que el *VaR* debido a que la gráfica azul interseca más a la gráfica roja que a la punteada verde, es decir, hubo menos pérdidas que excedieron el *CVaR* en comparación con el *VaR*, debido a que el *CVaR* captura en su cálculo el promedio de los datos que exceden al *VaR*.

De lo anterior se desprende inmediatamente que el *CVaR* permite realizar una cobertura adecuada del riesgo de mercado.

5.2. Método cuantitativo o del porcentaje de ajuste.

Este método⁹ intuitivo busca establecer un porcentaje de *ajuste* de la siguiente manera:

$$\% \text{ ajuste} = \left(1 - \frac{\text{Excepciones}}{\text{Total observaciones}} \right) * 100\%.$$

En ese sentido, se busca que el porcentaje de efectividad sea mayor o igual a 95 %, ya que este fue el nivel de confianza que se estableció. Para realizar dicho cálculo, se proponen dos algoritmos en Python, uno para el *VaR* y otro para el *CVaR*, respectivamente, los cuales se presentan a continuación:

Paso 11: cálculo del porcentaje de ajuste *VaR* para Ecopetrol en **Python**.

⁹Cabe resaltar que el método que aquí se presenta no constituye una metodología de *backtesting*.

```

excepciones_VaR = 0
for i in precio_
    Ecopetrol['retorno_Ecopetrol']:
        if i < VaR_normal:
            excepciones_VaR += 1
ajuste_VaR = (1-(excepciones_VaR /
    len(precio_Ecopetrol['retorno_
    Ecopetrol']))) * 100
ajuste_VaR

```

Paso 12: cálculo del porcentaje de ajuste CVaR para Ecopetrol en Python.

```

excepciones_CVaR = 0
for i in precio_
    Ecopetrol['retorno_Ecopetrol']:
        if i < CVaR_normal:
            excepciones_CVaR += 1
ajuste_CVaR = (1-(excepciones_CVaR /
    len(precio_Ecopetrol['retorno_
    Ecopetrol']))) * 100
ajuste_CVaR

```

Una vez ejecutados los algoritmos de los pasos 11 y 12 se llega que el ajuste del VaR y CVaR es de 96.5 % y 98.4 %, respectivamente. Lo cual indica que el CVaR se ajusta mejor a la información histórica facilitando realizar una gestión adecuada del riesgo.

6. Conclusiones.

Del desarrollo anterior se destacan cuatro conclusiones:

1. Las series financieras, como la de los rendimientos diarios de Ecopetrol, presentan exceso de curtosis, por lo cual suponer que se distribuyen normalmente constituye un error que conlleva a la subestimación del riesgo de mercado.
2. El VaR, como medida de riesgo no coherente, subestima la exposición al riesgo de mercado por dos razones: en primer lugar no tiene en cuenta las pérdidas que lo exceden, es decir, no le da importancia a los datos que se ubican en la cola de la distribución; en segundo lugar, el VaR está construido para funcionar adecuadamente bajo distribuciones elípticas (entre ellas la normal), sin embargo, las series financieras, en general, no siguen dicha distribución.
3. El CVaR, como medida coherente de riesgo, permite cuantificar de manera adecuada la exposición al riesgo de mercado al promediar las pérdidas que exceden el VaR. En este sentido, logra capturar los datos que se alojan en la cola de la distribución, los cuales representan pérdidas extremas que podría sufrir el activo en cuestión.
4. El cálculo y la probabilidad cobran especial importancia en el mundo financiero, ya que permiten dotarlo de cierto grado de formalidad lo cual genera mayor confianza en

los inversionistas, reguladores y demás participantes de los mercados.

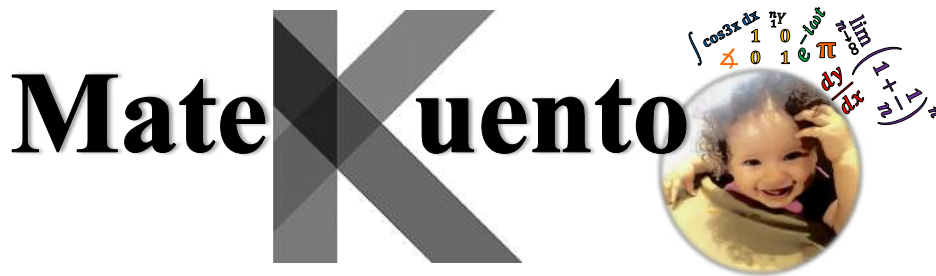
Agradecimientos.

A los profesores Oscar Martínez y Sergio Carrillo, quienes me enseñaron la potencia de la integración. Al profesor John Arredondo por su disposición y direccionamiento. A Mafe, por ser mi motivación.

Referencias.

- [RAE19] RAE, *Definición de riesgo*, 2019.
- [Kni21] F. Knight, *Risk, Uncertainty and Profit*, Houghton Mifflin Co, Boston, MA, 1921.
- [AB10] J. Alonso and L. Berggrun, *Introducción al análisis de riesgo financiero*, 2nd ed., Universidad Icesi, Bogotá, D.C., 2010.
- [ADEH99] P. Artzner, F. Delbaen, J. Eber, and D. Heath, *Coherent Measures of Risk*, *Mathematical Finance* **9** (1999), no. 3, 203-228.
- [Gar15] J. García, *Medidas de riesgo coherentes y su aplicación sobre la asignación estratégica de activos bajo el esquema Risk Parity: caso reservas internacionales de petróleo.*, Universidad EAFIT, 2015.
- [Ven06] F. Venegas, *Riesgos financieros y económicos: productos derivados y decisiones económicas bajo incertidumbre*, 2nd ed., Cengage Learning, 2006.
- [Rud76] W. Rudin, *Principles of Mathematical Analysis*, International series in pure and applied mathematics, McGraw-Hill, 1976.
- [AT02] C. Acerbi and D. Tasche, *On the coherence of expected shortfall*, *Journal of Banking & Finance* **26** (2002), no. 7, 1487-1503.
- [Jim14] J. Jiménez, *Distribuciones de probabilidad alternativas para la gestión de riesgo en mercados financieros*, Universidad de Valencia, Valencia, España, 2014.
- [CB12] M Constantino and C Brebbia, *Computational Finance and its Applications*, WIT Press, 2012.
- [Her13] L Herrera, *Acerca de la Teoría de los valores extremos y su aplicabilidad a la estimación del riesgo financiero*, *Anuario de la Facultad de Ciencias Económicas del Rosario* **9** (2013), 31-57.
- [HdV07] N Hyung and C de Vries, *Portfolio selection with heavy tails*, *Journal of Empirical Finance* **14** (2007), no. 3, 383-400.

Acerca del autor: Marco es un estudiante de sexto semestre de Matemáticas de la Universidad Sergio Arboleda. Entre los campos de la matemática, encuentra especial interés en la teoría de la medida, la probabilidad, el cálculo y su relación con el mundo financiero. Practica ciclismo con regularidad y espera algún día enseñar matemáticas y finanzas cuantitativas a nivel universitario.



Teoría de catástrofes y sistemas políticos

Las singularidades en la guerra y el amor

Su mirada, como las estepas de nuestra patria, en las que de niña yo jugaba justo antes del inicio del otoño: reverdecida, profunda, y serena, me tomó por sorpresa; más que el tibio clima de la ciudad a mitad del verano de 1967, la impresionante arquitectura medieval, la fastuosidad del complejo que sería mi hogar por los próximos años, o la estricta pesquisa que nos practicaban a quienes allí entrábamos, cada vez que lo hacíamos; fue su penetrante mirada lo que más desconcertó me produjo en mi primer día en Praga. Son hombres como Vuk Baranov los que, sin proponérselo, causan una intensa impresión en quienes con él se cruzan; no fue su garbo de diplomático, sus facciones escrupulosamente definidas, su lenguaje corporal masculino ni su estilo sutil y distante, no, no era eso, es que era en cada aspecto, en cada gesto, hasta en el más baladí: enteramente él, enteramente Vuk Baranov.

Tenía muchas expectativas por mi nombramiento como criptóloga adscrita al Ministerio de Relaciones Exteriores de la URSS, y asignada al Departamento de Comunicaciones de la Embajada Soviética en Checoslovaquia. Mi deber era garantizar una parte de la seguridad en la transmisión de los mensajes entre el Kremlin y el destacamento consular al que ahora pertenecía, el alcance específico de mi servicio: hacer inútil -mediante cifrado algorítmico- cualquier mensaje que cayera en poder de un tercero distinto al destinatario previsto, sobre todo al haber sido entrenada en la Unión Soviética, antes del inicio de mi viaje, para repeler el programa anglosajón de descifrado Venona.

Siempre quise sumar al propósito de alcanzar las formas de organización más avanzadas para la humanidad, como matemática pensé que entender la progresión de la sociedad a través de una ciencia objetiva sería un aporte valioso, por eso, el estudio de los Sistemas Dinámicos -y su eventual aplicación a los Sistemas Sociales- captó mi atención durante buena parte de mi vida académica.

Un sistema dinámico con determinada estabilidad estructural es propenso a contener regiones inestables que provocan cambios sustanciales, cualitativos: singularidades. En la teoría de catástrofes las singularidades poseen propiedades como la discontinuidad, la divergencia y la no reversibilidad, es decir que, aunque las transformaciones en el estado del sistema obedezcan únicamente a determinadas variaciones, al invertirse éstas, no necesariamente retorna el sistema a su estado inicial. Las catástrofes en la teoría son eventos en los que se da un cambio disruptivo que generalmente no se deshace de ma-

nera trivial, su estudio involucra los conceptos de estabilidad y bifurcaciones, y es útil -entre otras cosas- para identificar las condiciones en las que el sistema presenta diferentes estados, aún sin que las ecuaciones se encuentren completamente resueltas.

Me dediqué al estudio de la Teoría de Catástrofes al terminar la carrera en ciencias fisicomatemáticas en la Universidad Estatal Lomónosov de Moscú, tema que, luego de concluir la Cátedra de Investigación en Sistemas Dinámicos, sería medular en mi Candidatura en Ciencias.

Replanteé la materia de mi Candidatura al recibir el consejo de uno de mis profesores, quien notó en mí habilidades sobresalientes para la criptología, esto por mi destacada participación en un concurso que, realmente, era una prueba de aptitud que por encargo del Politburó, la Academia de Ciencias realizó para identificar funcionarios capaces para el servicio al pueblo socialista; la criptología era una materia útil en las condiciones políticas del momento y que me permitiría el honor de servir con eficiencia al socialismo, la instancia más evolucionada de la sociedad.

En mi primer desayuno en el restaurante de la embajada, el día siguiente a mi arribo, Vuk Baranov se dirigió a mi mesa, saludando a quienes nos encontrábamos en medio de una conversación marcada por los rituales propios de una bienvenida entre compatriotas, incursionó sutilmente en la plática, en el tercer comentario ya puso el ritmo de la conversación, captó la atención de los más activos partícipes de la charla, les devolvió a ellos el protagonismo en el discurso y, antes de darme cuenta, ya había dos conversaciones distintas, una en la que participaban los demás comensales, y otra en la que apenas estábamos él y yo; como una fiera en cacería pero con la elegancia propia de la esgrima, me apartó de la manada.

Supo de mí y supe de él, de su cargo como Agregado Cultural, su marcado sentido de compromiso político por el que se negaba a aprender un idioma distinto a los oficiales en los países del Pacto de Varsovia (a parte de su ruso natal), su mala condición física por sus problemas respiratorios congénitos, su incapacidad con los números, pero sobre todo pude notar su aguda sensibilidad artística, manifiesta en la forma en que cualquier comentario -por superficial que fuera- le servía para citar referencias en la literatura, la arquitectura, la pintura, la música y el cine soviéticos, a través de sus labios tuve una perspectiva más profunda de Dostoyevski, Trifonov, Rudnev, Vladimírski, Musorgsky y Eisenstein, entre muchos otros; lo

que fue una constante: conversar con él acerca de cualquier cosa, por banal que fuera, era un paseo por la arquitectura bizantina, el modernismo, el realismo socialista, el constructivismo y en general, por la cultura del y para el nuevo hombre.

Inevitablemente, antes de que finalizara el año 1967, ya estaba profundamente enamorada de Vuk, pasaba más tiempo en los aposentos que se le habían asignado a él al interior de la Embajada que en los que a mí me correspondían; el limitado alcance de las competencias de Vuk lo hacían un candidato poco interesante para las misiones en otros países diferentes a Checoslovaquia, por lo que pasamos ese primer semestre mucho tiempo juntos.

Modelar matemáticamente eventos sociales no es una tarea sencilla (y si se es riguroso, puede ser imposible), pero se puede partir de axiomas elementales y de fácil aceptación de la sociología y de la psicología de masas, así como de simplificaciones y generalizaciones que no eviten la caracterización de los matices relevantes del sistema.

A pesar de que Antonín Novotný como Primer Secretario del Partido Comunista de Checoslovaquia (en la práctica, su Presidente), inició el proceso de desestalinización encabezado por el Secretario General del Partido Comunista de la Unión Soviética: Nikita Jruschov desde 1956, había desde que llegué a Praga -en 1967- una fuerte oposición a tales reformas, principalmente por la lentitud de las mismas y por el impacto económico que habían generado; en los informes de inteligencia, la KGB afirmaba que tales efectos habían sido magnificados por elementos instigadores al interior de la sociedad Checoslovaca, y con mayor ahínco, en los colectivos de los estudiantes. El Secretario General del Partido Comunista de la Unión Soviética de la época, Leonid Brézhnev había sido informado de tal situación y promovió un cambio en el liderazgo del país y el desarrollo de reformas para que fueran más agudas y aceleradas.

Alexander Dubček fue nombrado el nuevo Secretario General del Partido Comunista de Checoslovaquia en enero de 1968 y desarrolló los principios de una apertura que podrían darle más estabilidad al sistema comunista, de alguna manera un modelo que formulé para el sistema socialista parecía llegar a esa conclusión al ubicar el estado actual del sistema más lejos del vecindario de la singularidad (la catástrofe) que bajo el mando de Novotný, seguramente empleando razonamientos muy distintos, Brézhnev llegaba a las mismas conclusiones y apoyaba los cambios.

El modelo que formulé mostraba una alta inestabilidad del sistema en general, de todo el bloque oriental en el caso en que se dieran dos eventos: que los movimientos por los derechos civiles en occidente perdieran fuerza de contrapeso en sus sistemas de gobierno capitalista, y a la vez, se desataran represiones violentas masivas en los países del Pacto de Varsovia, sobre todo en la periferia, como el caso de Checoslovaquia; estas conclusiones las guardé en mi habitación y con justificada paranoia, las encripté empleando elementos matemáticos y gramaticales anglo germánicos, hasta que días después sorprendí a Vuk hurgando en mi archivo, cuando le encontré le miré de forma atónita, me mostró de inmediato que intentaba,

de manera inocente, esconder entre mis papeles de consulta cotidiana (en los que me veía invertir horas casi a diario), un extraño cristal de centro oscuro, era una hermosa piedra que había conseguido Vuk y que me entregaba para mostrarme sus nuevos intereses; desconocía su afición por la mineralogía y su habilidad para ingresar ese elemento extraño a la Embajada (¿cómo lo hizo?), aunque nada era sorprendente en el genio de Vuk.

La total ignorancia de Vuk de las más elementales matemáticas, así como de las lenguas anglosajonas, me tranquilizaron con respecto a que hubiera visto algunos de los muy agudos resultados de mi modelización, de la propensión a singularidades del sistema socialista, sin embargo, evité volver a hacer registro de mis investigaciones y realicé un esfuerzo por guardar en mi memoria los nuevos resultados, al menos hasta que pudiera mostrar una teoría sólida a mis superiores.

El Primer Alto Directorio de la KGB (encargado de las operaciones en el extranjero y del contraespionaje) nos sorprendió al otro día, fuimos interrogados por aparte a causa del ingreso del extraño mineral a la Embajada por parte de Vuk, nuestras versiones fueron registradas, comparadas, el cristal fue analizado, se concluyó que había sido un obsequio sin trascendencia para la seguridad de la Embajada, y todo no pasó de ser un evento si se quiere, anecdótico, del que preferimos no hablar de nuevo; Vuk se ofreció para deshacerse de la roca y así olvidarnos eventualmente del tema.

En mis investigaciones amplié el alcance del sistema dinámico analizado en un segundo sistema: el comunista – capitalista, llegué a la conclusión de que la desestabilización de uno de los dos regímenes en el estado actual de las cosas, conllevaría a una singularidad, la catástrofe se daría por el uso desmedido de las fuerzas de una de las partes; una interpretación del modelo podría ser que la discontinuidad del sistema, la variación divergente del estado del mismo y la no reversibilidad podría darse por el uso del arsenal nuclear por la parte que se viera gravemente desestabilizada, y la respuesta de la otra parte.

Para mi terror, empezó a darse antes de terminado el primer semestre de 1968, una de las dos condiciones que, en el primer modelo, causarían la singularidad en el sistema comunista: la pérdida del contrapeso que al sistema capitalista daba el movimiento por los derechos civiles: los asesinatos de Martin Luther King Jr. y de Robert Kennedy en Estados Unidos, y los preparativos para las represiones de los movimientos del 68 en Estados Unidos, Francia y España.

La segunda condición estaba a punto de darse, y según el modelo, con ella la desestabilización del sistema comunista, lo que, en ese momento de la historia, podría causar el uso de fuerzas incontenibles en busca de mantener el estatus quo. Brézhnev y Dubček habían sido desinformados y se agudizaban las tensiones entre ellos, se había notificado a la Embajada de forma secreta que entre el 20 al 21 de agosto arribaría a Praga una enorme cantidad de fuerza militar del Pacto de Varsovia, para contener el descontento popular, ¿ocurriría una masacre de la población civil, provocada por fuerzas oscuras? ¿era acertado el modelo al advertir que esto causaría la deses-

tabilización del modelo comunista? ¿era acertado el segundo modelo al predecir que la desestabilización sería general y que ocurriría la catástrofe del uso de las fuerzas a disposición de las dos superpotencias?

Jamás había visto a Vuk respirar con dificultad como esa mañana del sábado 17 de agosto, a pesar de sus antecedentes de asma, estos no le impedían un desempeño vigoroso en las actividades físicas, pero esa mañana se desmayó justo al abrirme la puerta de su habitación, de inmediato pedí auxilio y le pusimos en manos de un médico de la Embajada, que en principio encontró sus vías respiratorias inflamadas, ¿algo le habría causado alergia?, no sabía qué, pero antes de devolverle a su habitación decidí abrir las ventanas y cortinas para que se ventilara.

Al hacerlo toqué sin querer su escritorio y un sistema mecánico cambió la disposición de las gavetas, cerrando violentamente una de ellas en el proceso, la cual contenía un dispositivo que pude ver unos segundos antes de que se retrajera al interior del mueble, miré hacia el piso por un momento, luego me recompuse y en ese instante llegó Vuk, esforzándose por respirar con tranquilidad, pero con la mitad de la camisa aún por fuera del pantalón, le pregunté como seguía y le ayudé a recostarse, le comenté que había escuchado a un mueble estremecerse y me explicó que era un dispositivo normal en los escritorios de funcionarios de su rango, para el resguardo de documentos clave, le noté verme a los ojos más que de costumbre, me tomó las muñecas para recostarse, buscaba tal vez una dilatación de mis pupilas o un incremento de mi ritmo cardíaco que revelara un comportamiento anormal de mi parte, jamás sabré si lo encontró.

Casi no quisiera haberlo visto, casi que quisiera haber perdido la capacidad de asociación de conceptos, pero un vistazo de apenas segundos bastó para formular una hipótesis terriblemente coherente con los hechos.

Por mi formación en el Ministerio de Comunicaciones pude reconocer un muy compacto transmisor de onda corta, fabricado artesanalmente, seguramente aquí mismo, para burlar los controles de acceso a la Embajada, empleaba casi en su totalidad materiales de uso cotidiano con excepción de una extraña piedra, del mismo tipo que el cristal oscuro que me obsequió, que estaba puesta en un tubo adyacente al comunicador, lo que debería ser el filtro de guía de onda, y el mineral: octaedrita, un tipo de óxido de titanio que, calibrado con total precisión, causa la resonancia de las ondas de interés y difracción de las que no lo son, el posicionamiento requiere de conocimientos muy sólidos de ingeniería y matemáticas y su inhalación puede causar la inflamación de las vías respiratorias. La octaedrita en forma de cristales oscuros y grandes corresponde a vetas muy específicas de Noruega, un país que no pertenece al Pacto de Varsovia, las marcas de posicionamiento se encontraban dispuestas en fracciones, no en décimas (como en el sistema internacional), era un aparato de comunicaciones construido con técnicas estadounidenses y un material clave de occidente, era un artefacto de contraespionaje.

Vuk claramente sabía más de lo que decía conocer, ¿qué

tanto conocía de mi hipótesis de la desestabilización del bloque oriental?, ¿participó en el desarrollo de los hechos motivándolos o informado al enemigo para que lo hiciera?, ¿faltaba el paso final de la masacre a la población civil en Praga?, ¿colaboraba con las facciones radicales que podrían motivar una confrontación?

Una catástrofe u otra, denunciar a Vuk Baranov, el traidor al que amaba, el desertor quien, sin duda alguna, también me había entregado su corazón; o arriesgar la ocurrencia -si los modelos eran acertados- de la singularidad en el estable sistema comunismo - capitalismo, pasando por la masacre en Praga que marcaría ese punto discontinuo, divergente e irreversible, que abriría la puerta a las dos superpotencias para que pudieran actuar sin autocensura, usando todas sus capacidades de ataque.

Sé que Vuk me amó tan intensamente como yo a él, pero en el estable sistema de nuestras vidas juntos, la catástrofe: discontinua, divergente e irreversible, se daba por estar en este momento histórico, por nuestras convicciones contrarias, por pertenecer a bandos distintos, esa sería la devastadora singularidad que marcaría nuestras vidas.

Salí de la habitación de Vuk y al terminar el pasillo y girar, de inmediato ingresé a un apartamento distinto al mío y frente a la mirada sorprendida del camarada responsable de las traducciones oficiales -quien era el ocupante de la habitación- llamé a la operadora de la Embajada, notifiqué mi hallazgo y en segundos un grupo de 10 oficiales del KGB ingresó a la habitación de Vuk derribando la puerta; lo encontramos sentado en su escritorio, a sus pies el dispositivo de comunicaciones hecho pedazos, en sus manos una copa de brandi y en sus labios un puro; el comandante del grupo le pidió a Vuk que se abstuviera de usar el diente, y que podrían conversar en favor de la patria. ¿Tenía Vuk una capsula de cianuro en su diente?, ¿era entonces un espía soviético?, ¿transmitía mensajes también a occidente y era entonces un agente doble?

Luego de mirarme a los ojos, nuevamente, de la misma forma que ese primer día, escuchamos un sutil sonido de algo quebrándose; los oficiales se abalanzaron hacia él y lo sacaron de la habitación intentando hacerle escupir; fui interrogada por varios días y luego de contar todo lo que sabía y todo acerca de mi modelo, la ocupación de la ciudad se realizó por parte de las fuerzas del Pacto de Varsovia, con un balance que, aunque trágico, no alcanzó a ser ni siquiera una fracción de lo que pudo haber sido, ni de lo que fueron las represiones que en ese mismo año se dieron en occidente, sobre todo la ocurrida más tarde en Tlatelolco, en las vecindades y bajo la influencia de la contraparte. Fui reasignada a Moscú y se me prohibió hablar del tema.

En 1975 encontré un obituario en el diario Pravda con la fotografía de Vuk Baranov en el que se anunciaba la muerte por causas naturales de un tal Vladímir Aleksándrovich Pozner, en Moscú, ¿habría sobrevivido Vuk? ¿realmente era Vuk Baranov?, ¿el obituario buscaba legalizar las acciones de inteligencia de la KGB de hacía siete años? Como iba yo a saberlo, si durante mi tiempo de servicio me presenté como Brana Sokolova, el que tampoco era mi nombre.

Dayron Fabián Achury-Calderón

dayronf.achuryc@konradlorenz.edu.co

Acerca del autor: Fabián Achury-Calderón es ingeniero de la Universidad Distrital Francisco José de Caldas y matemático de la Fundación Universitaria Konrad Lorenz. Se ha interesado por estudiar la interacción de las matemáticas con otras ciencias. Es Auxiliar de la Justicia y de manera indepen-

diente asesora compañías en asuntos financieros y de dirección estratégica. Le apasionan la ópera, el ballet y las ciencias sociales. Sus pasatiempos favoritos son la lectura y la colección de registros de radio internacional. Disfruta recorrer museos cuando viaja y admite que se conmueve con facilidad. Algún día espera aprender a tocar piano.

PASKÍN-CHALLENGE

CONCURSO PARA ESTUDIANTES DE COLEGIO

Un recorrido óptimo por el centro de Bogotá

Una turista visita Bogotá con el objetivo de conocer el centro de la ciudad. Después de leer en internet sobre la ciudad, creó la siguiente lista de lugares que quería conocer, e incluyó el tiempo que deseaba permanecer en cada uno:

- La Plaza de Bolívar (45 minutos).
- El Museo del Oro (6 horas).
- La terraza de la Torre Colpatria (30 minutos).
- La plaza de la Perseverancia (2 horas).
- La exposición del Museo Botero (6 horas).
- La Biblioteca Luis Ángel Arango (7 horas).
- El cerro de Monserrate (5 horas).
- El Museo Nacional (6 horas).
- El Museo de Arte Moderno (150 minutos).

La turista se hospedará en el Hotel Tequendama, ubicado en la Carrera 10 # 26-21. Estando ahí, planea iniciar su recorrido a los lugares de interés a las 8:00 a.m. Después, almorzará todos los días entre las 12:00 y la 1:00 p.m., continuará su recorrido y regresará al hotel máximo a las 5:00 p.m. Dado que la turista realizará todos los recorridos caminando, consideraremos que se desplazará a una velocidad constante de 4.5 km/hora.

Con el objetivo de conocer todos lugares en el menor tiempo posible, ¿qué lugares deberá conocer cada día y en qué orden?

Para resolver el problema tenga en cuenta la siguiente información:

- El objetivo consiste en diseñar un cronograma detallado que minimice lo más posible el tiempo que la turista está por fuera del hotel.

- Para los recorridos podrá utilizar cualquier camino entre la carrera décima (sobre la que está ubicado el hotel) y la carrera segunda este (sobre la que está ubicada la entrada al cerro de Monserrate), así como entre las calles 31 (sobre la que está ubicada la Plaza de la Perseverancia) y la calle 11 (sobre la que está ubicada el inicio de la Plaza de Bolívar).

- Utilice la aplicación Google Maps para diseñar el mapa que la turista utilizará.

- No se deben considerar tiempos de desplazamiento para almorzar ni buscar restaurantes, ya que la turista almorzará siempre en alguna de las atracciones turísticas. Almorzar se puede hacer antes, durante o después de una visita. En cualquier caso, esa hora de almuerzo no se debe incluir como el tiempo de visita a un sitio.

- La turista puede regresar al hotel en cualquier momento del día, pero nunca después de las 5 pm. Por eso, el cronograma detallado de su visita debe también considerar los tiempos que toman los desplazamientos.

- Los tiempos de cada desplazamiento se deben redondear hacia abajo. Por ejemplo, si ir del punto A al punto B toma 3 minutos y 45 segundos, entonces se considera que el tiempo del desplazamiento es de 3 minutos.

- Cada lugar debe visitarse en un solo día y de manera continua. Por ejemplo, si la visita al sitio A es de 5 horas, la turista no puede pasar 3 horas un día y luego 2 horas otro día. El almuerzo, sin embargo, puede interrumpir la visita y es el único caso en el cual la visita a un sitio no es continua.

Bases del Concurso:

Elegibilidad:

Estudiantes menores de 18 años, de cualquier nacionalidad y pertenecientes a un colegio colombiano.

Duración:

El concurso finaliza un año después de la publicación del Paskín que presenta el problema, o hasta que haya una persona ganadora.

Documento Solución:

La solución debe presentarse de manera individual en un documento formal, escrito en español, en el que se explique de manera detallada la solución al problema planteado.

La primera página del documento debe contener por lo menos los siguientes cuatro elementos:

- (1) El nombre completo de la persona participante.
- (2) El curso de la persona participante.
- (3) La edad de la persona participante.
- (4) El nombre completo del colegio de la persona.
- (5) Municipio o ciudad de residencia.
- (6) Telefono y correo electrónico de contacto.

El cuerpo del documento, que sigue a esta primera página, puede tener cualquier extensión para informar sobre la solución.

Los y las participantes pueden usar libros, blogs, computadores, internet, programas computacionales, pero no pueden consultar o interactuar con ninguna otra persona durante la solución. Ninguna contribución puede ser hecha por otra persona que no sea la persona que envía la solución.

Aunque los problemas estarán destinados a un análisis teórico, los participantes pueden realizar experimentos relevantes y presentar los datos resultantes en su trabajo si lo desean.

En caso de que se utilicen, cada solución debe incluir una lista de todas las referencias.

Las soluciones pueden usar algoritmos y herramientas computacionales existentes (que incluyen no solo herramientas disponibles gratuitamente, sino también herramientas dentro de sistemas como Matlab o Mathematica), siempre que se mencionen y citen correctamente y los métodos se expliquen claramente en la solución del problema. Cualquier programa de computadora escrito, si se escribió uno, debe incluirse como un apéndice de la solución. Sin embargo, todos los algoritmos, métodos y resultados deben explicarse en el texto principal del documento para recibir consideración durante la evaluación.

Las soluciones enviadas pueden incluir ecuaciones, gráficos, figuras y tablas. Cada solución debe incluir como mínimo:

- (1) Una corta introducción del problema tal como lo interpretó la persona participante.
- (2) Una explicación de todas las suposiciones y aproximaciones realizadas.
- (3) Justificación de todo el trabajo realizado.
- (4) Una breve discusión de las fortalezas y debilidades del enfoque adoptado.

Envío de la solución:

Las soluciones pueden enviarse como archivos .pdf, .odf, .doc o docx. Cada participante debe enviar su solución por correo electrónico a paskin@konradlorenz.edu.co y los documentos deben recibirse antes de que finalice el concurso. Si la solución no se recibe por correo electrónico durante la duración del concurso, la solución no será considerada. El documento completo debe incluirse como un archivo adjunto al correo electrónico. No envíe un enlace al documento en un servicio en la nube como Google Docs, Dropbox, etc. Una vez recibida la solución se le notificará la recepción en un plazo menor a cinco días hábiles. Si Ud. no recibe confirmación, comuníquese con el editor de la revista.

Resultados:

Toda solución recibida se irá evaluando en orden de llegada hasta la finalización del concurso. Cuando una solución sea correcta, la o el participante será invitada o invitado a exponer y justificar de manera verbal su solución de manera remota sincrónica utilizando alguna plataforma virtual como Teams, Zoom o Google Meets.

Premiación:

- La primera solución correcta será publicada en una siguiente edición del Paskín. Será una versión simplificada de la solución realizada en colaboración con el comité editorial, que incluirá una corta reseña de la persona ganadora.
- Certificado. Cada persona que participe y presente una solución correcta recibirá un certificado. El certificado se enviará por correo electrónico a la dirección utilizada para el envío de la solución. Espere varias semanas después de enviar su solución para recibir su certificado.
- La persona que envíe la primera solución correcta recibirá un bono virtual de regalo por un valor de **300.000 COP** (trescientos mil pesos colombianos). La persona ganadora puede escoger el establecimiento comercial que desee, siempre y cuando ofrezca bonos electrónicos y el comité editorial pueda adquirirlos.

Política Editorial

El Paskín Matemático es una producción del Programa de Matemáticas de la Fundación Universitaria Konrad Lorenz. Está abierto a todos y todas. Si tienes interés en escribir un artículo con cualquier tipo de contenido matemático, reseñar a un personaje, comentar un problema o tienes comentarios sobre nuestra publicación, envía un mensaje a paskin@konradlorenz.edu.co

$$\frac{dp}{dt} = r\left(1 - \frac{p}{k}\right)$$

El Programa de Matemáticas de la Fundación Universitaria Konrad Lorenz fue creado en 1998. Conoce más acerca de nuestro programa en nuestra página institucional

<http://www.konradlorenz.edu.co/es/aspirantes/carreras-universitarias/carrera-de-matematicas/presentacion.html>

Bogotá, Colombia. 2023

PASKÍN MATEMÁTICO

Volumen 5, No. 1, 2023

Contenido

- Dairo Andrés Sierra
Secciones de Poincare: delatando el caos en el péndulo doble1
- Gonzalo Medina
Sobre el problema de los cuatro subespacios9
- Daniel Esteban Galvis
De la cuarta dimensión y espacios perfectos, a rotar un espacio tridi-
mensional19
- Marco Blanco Ariza
Si quiero invertir ¿cómo me debo cubrir? Una aproximación a las medi-
das coherentes de riesgo en R y Python32
- Dayron Fabian Achurry
MateKuento.....39
- John A. Arredondo y Alejandro Cárdenas
Paskín Challenge43

La esencia de las matemáticas reside en su libertad.

George Cantor

