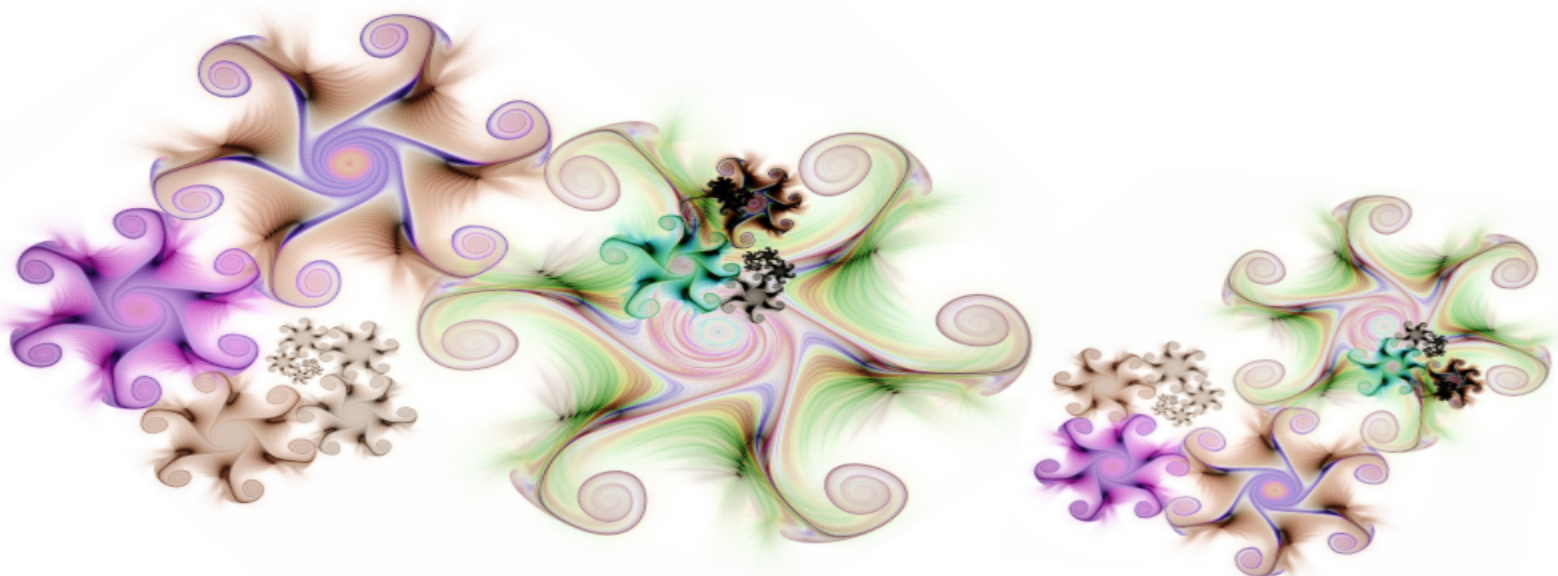




Resultados relevantes del álgebra lineal en modelos y aplicaciones

Lorenzo Héctor Juárez Valencia

**7° COLOQUIO DEL DEPARTAMENTO
DE MATEMÁTICAS UAM-I**

Del 27 al 31 de enero de 2025

Universidad Autónoma Metropolitana



UNIVERSIDAD AUTÓNOMA METROPOLITANA

Directorio

José Antonio de los Reyes Heredia
Rector General.

Verónica Medina Bañuelos
Rectora Unidad Iztapalapa.

Román Linares Romero
Director de CBI, UAM-Iztapalapa.

Raúl Montes de Oca Machorro
Jefe del Departamento de Matemáticas,
UAM-Iztapalapa.

Coordinador Editorial
Mario Pineda Ruelas
mpr@xanum.uam.mx

Comité Editorial

Elsa Baez Juárez
ebaez@cua.uam.mx

Shirley Bromberg Silverstein
stbsster@gmail.com

Judith Campos Cordero
judith@ciencias.unam.mx

Martín Celli Siboni,
celli@xanum.uam.mx

Pedro L. del Ángel Rodríguez
luis@cimat.mx

Begoña Fernández
bff@ciencias.unam.mx

Silvia Gavito Ticozzi
sgt@correo.azc.uam.mx

Jorge Ricardo Bolaños-Servín
jrbs@xanum.uam.mx

L. Héctor Juárez Valencia
hect@xanum.uam.mx

Jorge A. León Vázquez
jleon@ctrl.cinvestav.mx

Roberto Quezada Batalla
roqb@xanum.uam.mx

Edith Corina Sáenz Valadez
ecsv@ciencias.unam.mx

Martha L. Shaid Sandoval Miranda
marlisha@gmail.com

Ekaterina Todorova
todorova@cimat.mx

Luis Miguel Villegas Silva
villegas63@gmail.com

Editor web Pedro Iván Blanco Boa
ivanblc@gmail.com

Diseño logo Michael Rivera Arce
Portada archivo histórico de la UAM

MIXBA'AL. Vol. 16, No. 1, enero-diciembre de 2025, es una publicación anual de la Universidad Autónoma Metropolitana a través de la Unidad Iztapalapa, División de Ciencias Básicas e Ingeniería, Departamento de Matemáticas. Prolongación Canal de Miramontes 3855, Col. Ex Hacienda San Juan de Dios, Alcaldía Tlalpan, C.P. 14387, CDMX, México y Av. Ferrocarril San Rafael Atlixco, No. 186, Col. Leyes de Reforma 1a Sección, Alcaldía Iztapalapa, C.P. 09340, CDMX, México. Tel. 5804 4658. Página electrónica de la revista: <http://mat.izt.uam.mx/mat/index.php/revistamixba-al>. Correos electrónicos: mixbaal2009@gmail.com, mixb@xanum.uam.mx. Coordinador Editorial Mario Pineda Ruelas. Certificado de Reserva de Derechos al Uso Exclusivo de Título No. 04-2023-07031 1572300-102, ISSN: 2007-7874, ambos otorgados por el Instituto Nacional del Derecho de Autor. Responsable de la última actualización de este número Mario Pineda Ruelas, Departamento de Matemáticas, edificio AT, oficina 318. División de Ciencias Básicas e Ingeniería, Universidad Autónoma Metropolitana-Iztapalapa. Av. Ferrocarril San Rafael Atlixco No. 186, Colonia Leyes de Reforma 1a Sección, Alcaldía Iztapalapa, C.P. 09340, CDMX, México. Fecha de última modificación 30 de agosto de 2025. Tamaño del archivo 21.3 MB.

Las opiniones expresadas por los autores no necesariamente reflejan la postura del editor responsable de la publicación.

Queda estrictamente prohibida la reproducción total o parcial de los contenidos e imágenes de la publicación sin previa autorización de la Universidad Autónoma Metropolitana.

TALLER 1: ANÁLISIS DE DATOS CON UN ENFOQUE BAYESIANO

AUTOR: DR. ASael FABIAN MARTÍNEZ MARTÍNEZ

TALLER 2: UN BREVE RECORRIDO POR LA TEORÍA DE PRERRADICALES Y SUS RETÍCULAS

AUTORES: DR. ROGELIO FERNÁNDEZ-ALONSO, DRA. SILVIA GAVITO TICOZZI,
DRA. MARTHA LIZBETH SHAID SANDOVAL MIRANDA

TALLER 3: RESULTADOS DEL CÁLCULO Y ÁLGEBRA LINEAL RELEVANTES EN MODELOS Y APLICACIONES

AUTOR: DR. LORENZO HÉCTOR JUÁREZ

TALLER 4: ANÁLISIS GEOMÉTRICO DE SUPERFICIES: UNA INTRODUCCIÓN ELEMENTAL

AUTORES: DR. JOSUÉ MELENDEZ Y M. EN C. EDUARDO RODRÍGUEZ ROMERO

TALLER 5: UNA INTRODUCCIÓN A LOS TORNEOS Y SUS GENERALIZACIONES

AUTORES: DR. ILÁN GOLDFEDER ORTIZ Y DRA. NAHID YELENE JAVIER NOL

TALLER 6: DEL CERO AL QUANTUM: UN VIAJE POR EL MUNDO DE LOS CÓDIGOS

AUTORES: DR. JORGE BOLAÑOS SERVÍN, DRA. YURIKO PITONES AMARO Y DR. JOSUÉ RIOS CANGAS

TALLER 7: INTRODUCCIÓN A LA TEORÍA DE JUEGOS EPISTÉMICA

AUTOR: DR. RUBÉN BECERRIL BORJA

TALLER 8: CÓMO CONTAR MÁS ALLÁ DEL INFINITO Y PARA QUÉ SIRVE. INDUCCIÓN TRANSFINITA Y ALGUNAS APLICACIONES

AUTOR: DR. RODRIGO HERNÁNDEZ GUTIÉRREZ



7° Coloquio del Departamento de Matemáticas UAM-I 2025
Universidad Autónoma Metropolitana, Iztapalapa CDMX.

Introducción

En este curso se introducen tres temas del álgebra lineal que son relevantes en diversas aplicaciones. Los temas incluyen el concepto de ortogonalidad, proyecciones y factorizaciones matriciales, tales como la llamada factorización QR y la descomposición en valores singulares o SVD (singular value decomposition, por sus siglas en inglés). El curso es autocontenido y sólo se requieren conocimientos de un curso elemental de álgebra lineal que abarque los conceptos de independencia lineal, espacios vectoriales y subespacios, operaciones matriciales, operaciones elementales de fila para resolver sistemas de ecuaciones por medio de eliminación de Gauss y el cálculo de valores y vectores propios. Se mostrarán algunas aplicaciones elementales, principalmente la solución de problemas de mínimos cuadrados y manejo de compresión de imágenes. En la bibliografía de cada tema se dan referencias para estudiar los temas en forma exhaustiva y para tener idea del panorama de diversas aplicaciones, muchas de las cuales tienen un impacto importante en nuestra vida actual.

Lorenzo Héctor Juárez Valencia
hect@xanum.uam.mx
Departamento de Matemáticas
UAM-Iztapalapa

Índice general

Introducción	iii
1. Ortogonalidad y proyecciones	1
1.1. Producto interno y ortogonalidad	1
1.2. Proyecciones ortogonales	3
1.3. Aproximación de mínimos cuadrados	11
1.4. Bases ortogonales. Ortogonalización de Gram-Schmidt	14
1.5. Referencias	20
2. Factorización QR	21
2.1. Factorización reducida y completa	21
2.2. Aplicación a problemas de mínimos cuadrados	23
2.3. Reflexiones de Householder	24
2.4. Factorización QR mediante reflexiones matriciales	27
2.5. Referencias	32
3. Descomposición en valores singulares (SVD)	35
3.1. Simetría y ortogonalidad	35
3.2. Asimetría y simetrización. Descomposición SVD	37
3.3. Resumen de propiedades	43
3.4. Aplicaciones	46
3.5. Referencias	51

Capítulo 1

Ortogonalidad y proyecciones

El concepto de ortogonalidad es central en muchos campos del conocimiento científico y sus aplicaciones. Como sabemos, una base de un espacio vectorial es un conjunto de vectores linealmente independientes que generan dicho espacio. En el espacio euclídeo tridimensional, denotado por $\mathbb{R}^3$, la base más común es el conjunto de vectores canónicos, $\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3$, mutuamente perpendiculares. En geometría y mecánica, los ejes de coordenadas que mejor se ajustan a nuestra intuición son los que son perpendiculares. El concepto de ortogonalidad está fuertemente ligado a nuestra intuición geométrica en este caso y se sustenta en resultados básicos, como el concepto de producto interno, el teorema de Pitágoras y relaciones trigonométricas. Una primera generalización del concepto de ortogonalidad se introduce en el espacio n -dimensional $\mathbb{R}^n$. Esta generalización es muy útil para interpretar geométricamente los cálculos algebraicos, y permite simplificar los cálculos, como se mostrará en este capítulo.

1.1

Producto interno y ortogonalidad

Longitud de vectores

En $\mathbb{R}^2$ la longitud de un vector es consecuencia del teorema de Pitágoras: $\|\mathbf{x}\| = \sqrt{x_1^2 + x_2^2}$. En $\mathbb{R}^3$ se obtiene la longitud, aplicando el teorema de Pitágoras dos veces consecutivas: $\|\mathbf{x}\| = \sqrt{y^2 + x^2} = \sqrt{x_1^2 + x_2^2 + x_3^2}$, con $y^2 = x_1^2 + x_2^2$. La longitud de un vector $\mathbf{x} = (x_1, x_2, x_3)$ se puede pensar geométricamente como la longitud de la diagonal de una caja rectangular con lados de longitud igual al valor absoluto de x_1, x_2, x_3 como se muestra en la figura 1.1.

La generalización de la longitud de un vector en n dimensiones, $\mathbf{x} = (x_1, \dots, x_n)$ es inmediata:

$$\|\mathbf{x}\| = (x_1^2 + \dots + x_n^2)^{1/2} = \sqrt{\mathbf{x}^T \mathbf{x}}.$$

Esta expresión se puede interpretar ‘geométricamente’ como la aplicación del teorema de Pitágoras $n - 1$ veces consecutivas, agregando una dimensión en cada paso. La suma de cuadrados coincide

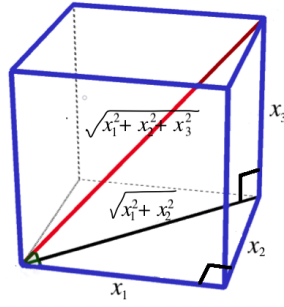


Figura 1.1: La longitud de un vector como la longitud de la diagonal de una caja.

con la multiplicación

$$\mathbf{x}^T \mathbf{x} = \begin{bmatrix} x_1 & x_2 & \cdots & x_n \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix} = x_1^2 + x_2^2 + \cdots + x_n^2 = \sum_{i=1}^n x_i^2.$$

Producto interno y ortogonalidad

La pregunta clave es ¿dados dos vectores $\mathbf{x}, \mathbf{y}$, cómo decidimos si son ortogonales? En el plano $\mathbb{R}^2$ y en el espacio tridimensional $\mathbb{R}^3$ puede responderse utilizando trigonometría y el teorema de Pitágoras. Por ejemplo, en $\mathbb{R}^2$ los dos vectores, $\mathbf{x}, \mathbf{y}$, son ortogonales si ellos forman un triángulo rectángulo con catetos representado por estos vectores, mientras que la hipotenusa se puede representar por $\mathbf{x} - \mathbf{y}$. Por lo tanto, de acuerdo al teorema de Pitágoras, $\mathbf{x}, \mathbf{y}$ son perpendiculares si

$$\|\mathbf{x}\|^2 + \|\mathbf{y}\|^2 = \|\mathbf{x} - \mathbf{y}\|^2 \implies (x_1^2 + x_2^2) + (y_1^2 + y_2^2) = (x_1 - y_1)^2 + (x_2 - y_2)^2.$$

Simplificando la última igualdad, se obtiene que $\mathbf{x}$ y $\mathbf{y}$ son ortogonales si

$$x_1 y_1 + x_2 y_2 = 0 \iff \mathbf{x}^T \mathbf{y} = 0.$$

Procediendo en forma análoga, los vectores $\mathbf{x}, \mathbf{y}$ en $\mathbb{R}^n$ son ortogonales si $\|\mathbf{x}\|^2 + \|\mathbf{y}\|^2 = \|\mathbf{x} - \mathbf{y}\|^2$. Calculando las normas y simplificando de nuevo obtenemos que los vectores son ortogonales si

$$x_1 y_1 + \cdots + x_n y_n = 0 \iff \mathbf{x}^T \mathbf{y} = 0.$$

Hacemos énfasis en que el producto $\mathbf{x}^T \mathbf{y}$ es el producto de una matriz de $1 \times n$ (el vector renglón $\mathbf{x}^T$) con otra matriz de $n \times 1$ (el vector columna $\mathbf{y}$):

$$\mathbf{x}^T \mathbf{y} = \begin{bmatrix} x_1 & x_2 & \cdots & x_n \end{bmatrix} \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{bmatrix} = x_1 y_1 + x_2 y_2 + \cdots + x_n y_n = \sum_{i=1}^n x_i y_i.$$

Este producto aparece en muchos contextos y es esencial para estudiar geometría en el espacio n -dimensional. Se le llama el **producto escalar** o **producto interno**. Nosotros utilizaremos el

término **producto interno** porque en el álgebra lineal existe otro producto con dos vectores, denominado el **producto externo** que da origen a una matriz de rango uno, dada por

$$\mathbf{xy}^T = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix} \begin{bmatrix} y_1 & y_2 & \cdots & y_n \end{bmatrix} = \begin{bmatrix} x_1 y_1 & x_1 y_2 & \cdots & x_1 y_n \\ x_2 y_1 & x_2 y_2 & \cdots & x_2 y_n \\ \vdots & \vdots & \ddots & \vdots \\ x_n y_1 & x_n y_2 & \cdots & x_n y_n \end{bmatrix}.$$

DEFINICIÓN 1.1. El **producto interno**. Dados dos vectores (columna) $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$, su producto interno se define como

$$\mathbf{x} \cdot \mathbf{y} = \mathbf{x}^T \mathbf{y} = x_1 y_1 + \cdots x_n y_n.$$

Este producto es cero si y solo si los vectores son ortogonales.

1.2

Proyecciones ortogonales

Ángulo entre vectores y la desigualdad de Schwarz

En los espacios vectoriales $\mathbb{R}^2$ o $\mathbb{R}^3$ sabemos que el producto interno de vectores está relacionado con el coseno del ángulo de dichos vectores. Utilizando trigonometría podemos encontrar fácilmente dicha relación. Supongamos que el vector $\mathbf{a} = (a_1, a_2)$ hace un ángulo α con el eje x y que $\mathbf{b} = (b_1, b_2)$ hace un ángulo β . El ángulo entre los dos vectores es $\theta = \beta - \alpha$, como se muestra en la figura 1.2.

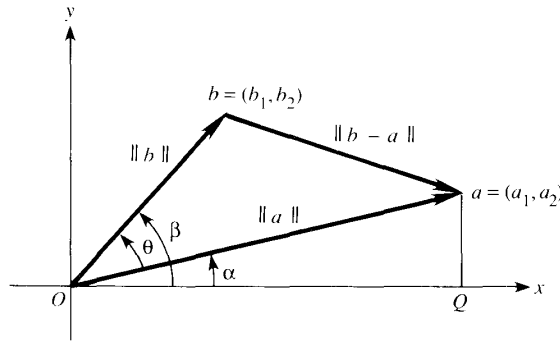


Figura 1.2: Relación entre $\cos \theta$ y el producto interno.

$$\cos \theta = \cos(\beta - \alpha) = \cos \beta \cos \alpha + \sin \beta \sin \alpha = \frac{a_1 b_1 + a_2 b_2}{\|\mathbf{a}\| \|\mathbf{b}\|} = \frac{\mathbf{a}^T \mathbf{b}}{\|\mathbf{a}\| \|\mathbf{b}\|},$$

debido a que

$$\cos(\alpha) = \frac{a_1}{\|\mathbf{a}\|}, \quad \sin(\alpha) = \frac{a_2}{\|\mathbf{a}\|}, \quad \cos(\beta) = \frac{b_1}{\|\mathbf{b}\|}, \quad \sin(\beta) = \frac{b_2}{\|\mathbf{b}\|}.$$

Por lo tanto, podemos introducir la siguiente definición.

DEFINICIÓN 1.2. El coseno del ángulo entre dos vectores cualquiera $\mathbf{a}$ y $\mathbf{b}$ es

$$\cos \theta = \frac{\mathbf{a}^T \mathbf{b}}{\|\mathbf{a}\| \|\mathbf{b}\|}.$$

Seguiremos considerando esta definición del ángulo entre vectores en $\mathbb{R}^n$. Así, dos vectores en $\mathbb{R}^n$ son ortogonales si y solo si hacen un ángulo $\theta = \pi/2$. Es decir, dos vectores son ortogonales sólo cuando su producto interno es cero.

Desigualdad de Schwarz. Debido a que $|\cos \theta| \leq 1$ para cualquier θ , utilizando la definición del ángulo entre vectores se obtiene la siguiente desigualdad.

$$|\mathbf{a}^T \mathbf{b}| \leq \|\mathbf{a}\| \|\mathbf{b}\| \quad \text{para cualquier par de vectores } \mathbf{a}, \mathbf{b} \in \mathbb{R}^n.$$

Proyección de un vector $\mathbf{b}$ sobre la línea definida por el vector $\mathbf{a}$

En la figura 1.3, el punto $\mathbf{p}$ define el vector que es la proyección ortogonal del vector $\mathbf{b} \in \mathbb{R}^n$ sobre la línea generada por el vector $\mathbf{a}$. Este vector es colineal con el vector $\mathbf{a}$, es decir $\mathbf{p} = \hat{x}\mathbf{a}$ con $\hat{x} \in \mathbb{R}$. El escalar $\hat{x}$ se puede encontrar utilizando la ortogonalidad del vector $\mathbf{p} - \mathbf{b}$ con el vector $\mathbf{a}$:

$$0 = \mathbf{a}^T (\mathbf{p} - \mathbf{b}) = \mathbf{a}^T (\hat{x}\mathbf{a}) - \mathbf{a}^T \mathbf{b} \implies \hat{x} = \frac{\mathbf{a}^T \mathbf{b}}{\mathbf{a}^T \mathbf{a}} \implies \mathbf{p} = \frac{\mathbf{a}^T \mathbf{b}}{\mathbf{a}^T \mathbf{a}} \mathbf{a}.$$

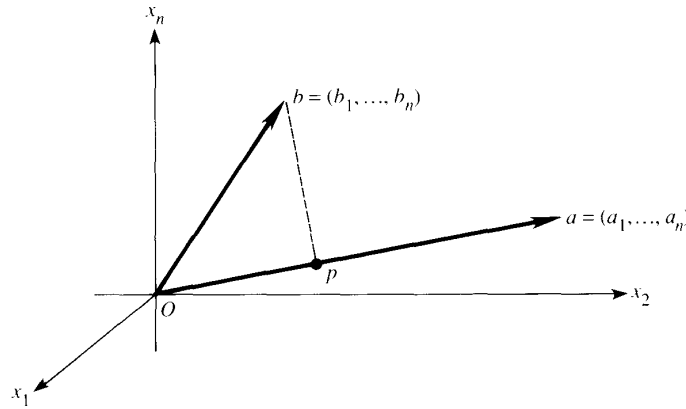


Figura 1.3: Proyección de $\mathbf{b}$ sobre $\mathbf{a}$.

DEFINICIÓN 1.3. La proyección de un vector $\mathbf{b}$ sobre la línea que pasa por el origen en la dirección del vector $\mathbf{a}$ es

$$\mathbf{p} = \hat{x}\mathbf{a} = \frac{\mathbf{a}^T \mathbf{b}}{\mathbf{a}^T \mathbf{a}} \mathbf{a}. \quad (1.1)$$

EJEMPLO 1.4. Proyectar $\mathbf{b} = (4, 1, 3)$ sobre la línea que pasa por el origen en la dirección de $\mathbf{a} = (1, 0, 2)$.

Solución. La proyección es

$$\hat{x} = \frac{\mathbf{a}^T \mathbf{b}}{\mathbf{a}^T \mathbf{a}} = 2 \implies \mathbf{p} = 2\mathbf{a} = (2, 0, 4).$$

Además, se satisface la desigualdad de Schwarz

$$10 = |\mathbf{a}^T \mathbf{b}| \leq \|\mathbf{a}\| \|\mathbf{b}\| = \sqrt{5} \sqrt{26} = \sqrt{130}.$$

Nota. Es posible también obtener la desigualdad de Schwarz partiendo de la fórmula de proyección $\mathbf{b}$ sobre la línea generada por $\mathbf{a}$. Una forma sencilla es observando que $\|\mathbf{b} - \mathbf{p}\|^2 \geq 0$. Es decir

$$0 \leq \left\| \mathbf{b} - \frac{\mathbf{a}^T \mathbf{b}}{\mathbf{a}^T \mathbf{a}} \mathbf{a} \right\|^2 = \mathbf{b}^T \mathbf{b} - 2 \frac{(\mathbf{a}^T \mathbf{b})^2}{\mathbf{a}^T \mathbf{a}} + \left(\frac{\mathbf{a}^T \mathbf{b}}{\mathbf{a}^T \mathbf{a}} \right)^2 \mathbf{a}^T \mathbf{a} = \frac{(\mathbf{b}^T \mathbf{b})(\mathbf{a}^T \mathbf{a}) - (\mathbf{a}^T \mathbf{b})^2}{\mathbf{a}^T \mathbf{a}}.$$

El denominador es positivo por lo que el numerador debe ser no negativo. Por lo tanto, del numerador se obtiene

$$(\mathbf{a}^T \mathbf{b})^2 \leq (\mathbf{b}^T \mathbf{b})(\mathbf{a}^T \mathbf{a}) \implies |\mathbf{a}^T \mathbf{b}| \leq \|\mathbf{a}\| \|\mathbf{b}\|.$$

Matriz de proyección de rango uno

La proyección de $\mathbf{b}$ sobre el espacio vectorial generado por $\mathbf{a}$ se puede expresar por medio de una matriz cuadrada llamada la **matriz de proyección**:

$$\mathbf{p} = P_a \mathbf{b}, \quad \text{con} \quad P_a = \frac{\mathbf{a} \mathbf{a}^T}{\mathbf{a}^T \mathbf{a}}.$$

Esta expresión se obtiene de la fórmula (1.1), asociando en forma diferente los productos en el numerador:

$$\mathbf{p} = \hat{x} \mathbf{a} = \frac{\mathbf{a}^T \mathbf{b}}{\mathbf{a}^T \mathbf{a}} \mathbf{a} = \mathbf{a} \frac{\mathbf{a}^T \mathbf{b}}{\mathbf{a}^T \mathbf{a}} = \frac{\mathbf{a} \mathbf{a}^T}{\mathbf{a}^T \mathbf{a}} \mathbf{b} = P_a \mathbf{b}.$$

La matriz de proyección es la matriz cuadrada P_a , la cual se forma del producto externo $\mathbf{a} \mathbf{a}^T$ dividida por el producto interno $\mathbf{a}^T \mathbf{a}$.

EJEMPLO 1.5. En el ejemplo anterior $\mathbf{a} = (1, 0, 2)$ y $\mathbf{b} = (4, 1, 3)$, la matriz de proyección es

$$P_a = \frac{\mathbf{a} \mathbf{a}^T}{\mathbf{a}^T \mathbf{a}} = \frac{1}{5} \begin{bmatrix} 1 & 0 & 2 \\ 0 & 0 & 0 \\ 2 & 0 & 4 \end{bmatrix}.$$

La proyección de $\mathbf{b}$ sobre el espacio vectorial generado por $\mathbf{a}$ es

$$\mathbf{p} = P_a \mathbf{b} = \frac{1}{5} \begin{bmatrix} 1 & 0 & 2 \\ 0 & 0 & 0 \\ 2 & 0 & 4 \end{bmatrix} \begin{bmatrix} 4 \\ 1 \\ 3 \end{bmatrix} = \begin{bmatrix} 2 \\ 0 \\ 4 \end{bmatrix},$$

obteniendo el mismo resultado que en el ejemplo 1.4.

Una característica de las proyecciones es

$$P_a^2 \mathbf{b} = P_a \mathbf{b} \text{ para todo } \mathbf{b} \in \mathbb{R}^3 \implies P_a^2 = P_a.$$

En el ejemplo anterior

$$P_a^2 = \frac{1}{25} \begin{bmatrix} 1 & 0 & 2 \\ 0 & 0 & 0 \\ 2 & 0 & 4 \end{bmatrix} \begin{bmatrix} 1 & 0 & 2 \\ 0 & 0 & 0 \\ 2 & 0 & 4 \end{bmatrix} = \frac{1}{5} \begin{bmatrix} 1 & 0 & 2 \\ 0 & 0 & 0 \\ 2 & 0 & 4 \end{bmatrix} = P_a.$$

Esta propiedad es consistente con el concepto de proyección: deja fijos los vectores en la imagen de la transformación matricial $T : \mathbb{R}^n \mapsto \mathbb{R}^n$ definida por

$$T(\mathbf{x}) = P_a \mathbf{x}.$$

Esta transformación es lineal por ser matricial. Su espacio imagen $Im(T)$ y su núcleo $Ker(T)$ coinciden con el espacio columna $\mathcal{R}(P_a)$ y con el espacio nulo $\mathcal{N}(P_a)$ de la matriz de proyección P_a , respectivamente:

$$Im(T) = \mathcal{R}(P_a) = Gen\{\mathbf{a}\}, \quad Ker(T) = \mathcal{N}(P_a) = \{\mathbf{x} \in \mathbb{R}^2 \mid \mathbf{a}^T \mathbf{x} = 0\},$$

Geométricamente el núcleo es el plano que pasa por el origen con vector normal $\mathbf{a} = [1 \ 0 \ 2]^T$, y el espacio imagen es la línea recta en la dirección de $\mathbf{a}$ como se ilustra en la figura 1.4. En dicha figura solo se muestra la perspectiva sobre los ejes x_1 - x_3 .

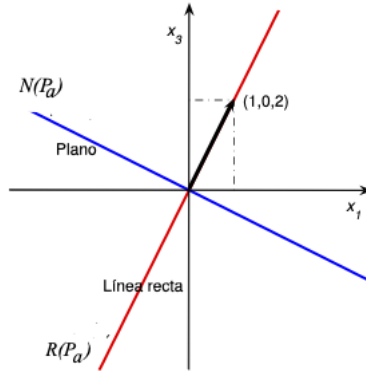


Figura 1.4: Ilustración del espacio imagen y el espacio nulo

La proyección se denomina **de rango uno** porque P_a es una matriz de rango uno ($r = 1$) ya que sus columnas son generadas por el vector $\mathbf{a}$. Por lo tanto, su espacio nulo es de dimensión $n - 1 = 3 - 1 = 2$, el plano cuya normal es $\mathbf{a}$. Las proyecciones de rango uno P_a en $\mathbb{R}^n$ se definen de manera análoga, pero ahora el vector $\mathbf{a}$ pertenece a $\mathbb{R}^n$:

TEOREMA 1.6. Para cualquier vector $\mathbf{a} \in \mathbb{R}^n$, $\mathbf{a} \neq \mathbf{0}$, existe una proyección ortogonal de rango uno, dada por

$$P_a = \frac{\mathbf{a}\mathbf{a}^T}{\mathbf{a}^T \mathbf{a}} = \hat{\mathbf{a}}\hat{\mathbf{a}}^T, \quad \text{con } \hat{\mathbf{a}} = \frac{\mathbf{a}}{\|\mathbf{a}\|}.$$

Como ya hemos visto anteriormente, esta matriz de proyección tiene las siguientes propiedades

- $P_a^2 = P_a$, $P_a = P_a^T$ (simétrica), P_a es de rango 1.
- $\mathcal{R}(P_a) = \{\mathbf{x} \in \mathbb{R}^n \mid \mathbf{x} = \alpha \mathbf{a}, \alpha \in \mathbb{R}\}$, la línea generada por el vector $\mathbf{a}$.
- $\mathcal{N}(P_a) = \{\mathbf{x} \in \mathbb{R}^n \mid \mathbf{a}^T \mathbf{x} = a_1 x_1 + a_2 x_2 + \dots + a_n x_n = 0\}$, hiperplano con vector normal $\mathbf{a} = (a_1, a_2, \dots, a_n)$.
- $\mathcal{R}(P_a)$ y $\mathcal{N}(P_a)$ son subespacios vectoriales de $\mathbb{R}^n$, mutuamente ortogonales y de dimensiones 1 y $n - 1$, respectivamente.

Debido a la propiedad de ortogonalidad de los subespacios anteriores, y a la suma de sus dimensiones, encontramos que $\mathbb{R}^n$ es la suma directa de los espacios columna y nulo de la matriz P_a , y escribimos

- $\mathcal{R}(P_a) \oplus \mathcal{N}(P_a) = \mathbb{R}^n$. Es decir
 - (a) $\mathcal{R}(P_a) \cup \mathcal{N}(P_a) = \mathbb{R}^n$ y $\mathcal{R}(P_a) \cap \mathcal{N}(P_a) = \{\mathbf{0}\}$.
 - (b) Todo $\mathbf{x} \in \mathbb{R}^n$ se puede descomponer en $\mathbf{x} = \mathbf{y} + \mathbf{z}$ con $\mathbf{y} \in \mathcal{R}(P_a)$ y $\mathbf{z} \in \mathcal{N}(P_a)$, donde $\mathbf{y} = P_a \mathbf{x}$ y $\mathbf{z} = (I - P_a) \mathbf{x}$.

Las proyecciones ortogonales de rango uno son casos especiales de las proyecciones en $\mathbb{R}^n$. Es posible construir proyecciones de cualquier rango entre 1 y $n - 1$. A continuación enlistamos algunos conceptos generales sobre las proyecciones en el espacio euclídeo n -dimensional y posteriormente estudiaremos proyecciones de rango mayor.

- Cualquier matriz cuadrada P que satisface $P^2 = P$ define una proyección, pues la acción de P es invariante sobre $\mathcal{R}(P)$. Es decir, $\mathbf{y} \in \mathcal{R}(P)$ implica $\mathbf{y} = P \mathbf{x}$ para algún $\mathbf{x} \in \mathbb{R}^n$ y $P \mathbf{y} = \mathbf{y}$.
- Si P es una matriz de proyección, entonces $I - P$, con I la matriz identidad, también es matriz de proyección, pues $(I - P)^2 = I - 2P + P^2 = I - 2P + P = I - P$. La matriz $I - P$ define la **proyección complementaria** de P .
- P proyecta sobre el espacio nulo de $I - P$ y viceversa. Es decir, $\mathcal{N}(P) = \mathcal{R}(I - P)$ y $\mathcal{N}(I - P) = \mathcal{R}(P)$.
- Una proyección P en $\mathbb{R}^n$, separa el espacio en dos subespacios $S_1 = \mathcal{R}(P)$ y $S_2 = \mathcal{N}(P)$ con $S_1 \cap S_2 = \{\mathbf{0}\}$ y $S_1 \cup S_2 = \mathbb{R}^n$. Dado $\mathbf{x} \in \mathbb{R}^n$, $\mathbf{x} = \mathbf{y} + \mathbf{z}$ con $\mathbf{y} = P \mathbf{x}$ y $\mathbf{z} = (I - P) \mathbf{x}$.
- Hay proyecciones oblicuas y proyecciones ortogonales. Las proyecciones más importantes en el álgebra lineal son las ortogonales. Una proyección P es ortogonal si además de satisfacer $P = P^2$ también satisface $P^T = P$. En este caso, para $\mathbf{x} = \mathbf{y} + \mathbf{z}$ con $\mathbf{y} = P \mathbf{x}$ y $\mathbf{z} = (I - P) \mathbf{x}$, se tiene

$$\mathbf{y}^T \mathbf{z} = (P \mathbf{x})^T (I - P) \mathbf{x} = \mathbf{x}^T P^T (I - P) \mathbf{x} = \mathbf{x}^T P (I - P) \mathbf{x} = \mathbf{0}.$$

Proyecciones de rango mayor

Es posible construir proyecciones de rango mayor a uno. Por ejemplo, dado el vector $\mathbf{a}$, P_a proyecta sobre la línea generada por $\mathbf{a}$, y el complemento ortogonal de esta línea es el hiperplano con vector normal $\mathbf{a}$ de dimensión $n - 1$. Para proyectar vectores sobre este hiperplano podemos utilizar la llamada **proyección complementaria**, dada por la matriz.

$$I - P_a = I - \frac{\mathbf{a} \mathbf{a}^T}{\mathbf{a}^T \mathbf{a}}, \quad \text{que es de rango } n - 1.$$

De modo que si $\mathbf{b} \in \mathbb{R}^n$, entonces $\mathbf{p} = (I - P_a)\mathbf{b}$ es su proyección sobre el hiperplano con vector normal $\mathbf{a}$.

EJEMPLO 1.7. En el ejemplo anterior encontramos que la matriz de proyección sobre la línea definida por $\mathbf{a} = (1, 0, 2)$ es

$$P_a = \frac{\mathbf{a}\mathbf{a}^T}{\mathbf{a}^T\mathbf{a}} = \frac{1}{5} \begin{bmatrix} 1 & 0 & 2 \\ 0 & 0 & 0 \\ 2 & 0 & 4 \end{bmatrix}.$$

La proyección complementaria $P_a^\perp = I - P_a$, proyecta sobre el plano ortogonal al vector $\mathbf{a} = (1, 0, 2)$ (ver figura 1.4). Su matriz es

$$P_a^\perp = I - P_a = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} - \frac{1}{5} \begin{bmatrix} 1 & 0 & 2 \\ 0 & 0 & 0 \\ 2 & 0 & 4 \end{bmatrix} = \frac{1}{5} \begin{bmatrix} 4 & 0 & -2 \\ 0 & 1 & 0 \\ -2 & 0 & 1 \end{bmatrix},$$

y es de rango dos, ya que tiene dos columnas linealmente independientes, la primera y la segunda (o la segunda y tercera). Esos dos vectores columna linealmente independientes generan el plano $\mathbf{a} \cdot \mathbf{x} = x_1 + 0x_2 + 2x_3 = 0$ (de hecho, cada columna satisface esta ecuación). Es fácil verificar que $P_a^\perp$ es una proyección ya que $(P_a^\perp)^2 = P_a^\perp$. Además como la matriz $P_a^\perp$ es simétrica la proyección es ortogonal.

Dejando de lado las proyecciones complementarias de la forma $I - P_a$, nuestro interés es construir proyecciones sobre planos o proyecciones sobre subespacios de mayor dimensión en $\mathbb{R}^n$. Por ejemplo, para construir una proyección sobre un plano que pasa por el origen (subespacio de dimensión dos) podemos tomar dos vectores diferentes en el plano, digamos $\mathbf{a}_1$ y $\mathbf{a}_2$. La proyección de cualquier vector $\mathbf{b} \in \mathbb{R}^n$ sobre el plano es una combinación lineal de los dos vectores generadores:

$$\mathbf{p} = \alpha \mathbf{a}_1 + \beta \mathbf{a}_2 \quad \text{con } \alpha, \beta \in \mathbb{R} \text{ constantes a determinar.}$$

Es decir

$$\mathbf{p} = A \hat{\mathbf{x}} \quad \text{con } A = \begin{bmatrix} | & | \\ \mathbf{a}_1 & \mathbf{a}_2 \\ | & | \end{bmatrix} \quad \text{y } \hat{\mathbf{x}} = \begin{bmatrix} \alpha \\ \beta \end{bmatrix},$$

en donde los vectores $\mathbf{a}_1, \mathbf{a}_2$, de la matriz A , son vectores columna. Para calcular α y β podemos auxiliarnos de la geometría, como se indica en la figura 1.5. En particular se debe satisfacer

$$\mathbf{b} - A \hat{\mathbf{x}} \perp \text{Gen}\{\mathbf{a}_1, \mathbf{a}_2\} = \mathcal{R}(A).$$

Lo cual se cumple

$$\begin{aligned} &\text{si y sólo si } \mathbf{a}_1^T(\mathbf{b} - A \hat{\mathbf{x}}) = 0, & \mathbf{a}_2^T(\mathbf{b} - A \hat{\mathbf{x}}) = 0, \\ &\text{si y sólo si } A^T(\mathbf{b} - A \hat{\mathbf{x}}) = \mathbf{0}, \\ &\text{si y sólo si } A^T A \hat{\mathbf{x}} = A^T \mathbf{b}. & \text{(ecuaciones normales)} \end{aligned}$$

Resolviendo el llamado **sistema de ecuaciones normales**, se obtiene el vector con los coeficientes buscados, $\hat{\mathbf{x}} = [\alpha \ \beta]^T$:

$$\hat{\mathbf{x}} = (A^T A)^{-1} A^T \mathbf{b},$$

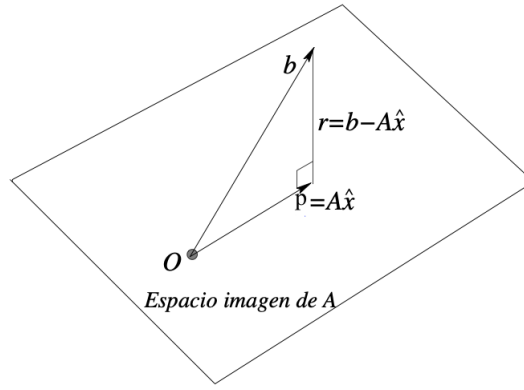


Figura 1.5: El espacio imagen o columna de A es ortogonal al residual.

en donde la matriz $A^T A$, de 2×2 , es invertible porque los vectores columna de A se suponen linealmente independientes. Por lo tanto, la proyección de $\mathbf{b} \in \mathbb{R}^n$ sobre el plano generado por $\mathbf{a}_1, \mathbf{a}_2$ es

$$\mathbf{p} = A \hat{\mathbf{x}} = A (A^T A)^{-1} A^T \mathbf{b}, \quad \text{con } A = \begin{bmatrix} | & | \\ \mathbf{a}_1 & \mathbf{a}_2 \\ | & | \end{bmatrix},$$

y la **matriz de proyección** es la matriz de $n \times n$

$$P_A = A (A^T A)^{-1} A^T.$$

EJEMPLO 1.8. Proyectar el vector $\mathbf{b} = [2 \ 3 \ 4]^T$ sobre el espacio generado por los dos vectores $\mathbf{a}_1 = [1 \ 0 \ 0]^T$ y $\mathbf{a}_2 = [1 \ 1 \ 0]^T$.

Solución. Podemos resolver el problema siguiendo dos procedimientos (ambos relacionados):

- Resolver el sistema de ecuaciones normales $A^T A \hat{\mathbf{x}} = A^T \mathbf{b}$ y luego proyectar $\mathbf{p} = A \hat{\mathbf{x}}$.
- Construir la matriz de proyección $P_A = A (A^T A)^{-1} A^T$ y luego proyectar $\mathbf{p} = P_A \mathbf{b}$.

Ambos procedimientos son equivalentes y producen el mismo resultado.

Procedimiento 1.

$$A = \begin{bmatrix} 1 & 1 \\ 0 & 1 \\ 0 & 0 \end{bmatrix}, \quad \mathbf{b} = \begin{bmatrix} 2 \\ 3 \\ 4 \end{bmatrix}.$$

Calculamos $A^T A$ y $A^T \mathbf{b}$

$$A^T A = \begin{bmatrix} 1 & 0 & 0 \\ 1 & 1 & 0 \end{bmatrix} \begin{bmatrix} 1 & 1 \\ 0 & 1 \\ 0 & 0 \end{bmatrix} = \begin{bmatrix} 1 & 1 \\ 1 & 2 \end{bmatrix}, \quad A^T \mathbf{b} = \begin{bmatrix} 1 & 0 & 0 \\ 1 & 1 & 0 \end{bmatrix} \begin{bmatrix} 2 \\ 3 \\ 4 \end{bmatrix} = \begin{bmatrix} 2 \\ 5 \end{bmatrix}.$$

El sistema de ecuaciones normales $A^T A \hat{\mathbf{x}} = A^T \mathbf{b}$ es entonces

$$\begin{bmatrix} 1 & 1 \\ 1 & 2 \end{bmatrix} \begin{bmatrix} \alpha \\ \beta \end{bmatrix} = \begin{bmatrix} 2 \\ 5 \end{bmatrix},$$

y su solución es $\hat{\mathbf{x}} = [-1 \ 3]^T$. Por lo tanto, el vector proyección sobre el plano generado por $\mathbf{a}_1$, $\mathbf{a}_2$ es

$$\mathbf{p} = A\hat{\mathbf{x}} = \begin{bmatrix} 1 & 1 \\ 0 & 1 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} -1 \\ 3 \end{bmatrix} = \begin{bmatrix} 2 \\ 3 \\ 0 \end{bmatrix}.$$

Procedimiento 2.

La matriz de proyección es

$$P_A = A(A^T A)^{-1}A^T = \begin{bmatrix} 1 & 1 \\ 0 & 1 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} 2 & -1 \\ -1 & 1 \end{bmatrix} \begin{bmatrix} 1 & 0 & 0 \\ 1 & 1 & 0 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0 \end{bmatrix}.$$

Calculamos la proyección $\mathbf{p} = P_A \mathbf{b}$:

$$\mathbf{p} = P_A \mathbf{b} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0 \end{bmatrix} \begin{bmatrix} 2 \\ 3 \\ 4 \end{bmatrix} = \begin{bmatrix} 2 \\ 3 \\ 0 \end{bmatrix}.$$

Nota. La matriz de proyección $P_A \mathbf{b} = A(A^T A)^{-1}A^T$ puede ser difícil de asimilar y/o memorizar, pero basta con recordar que es la generalización de una proyección de rango uno. Cuando A contiene un solo vector columna $\mathbf{a} \neq \mathbf{0}$ ($A = [\mathbf{a}]$), sabemos que $P_a = \frac{\mathbf{a}\mathbf{a}^T}{\mathbf{a}^T \mathbf{a}} = \mathbf{a}(\mathbf{a}^T \mathbf{a})^{-1}\mathbf{a}^T$, así que ahora es claro que cuando A tiene dos columnas linealmente independientes, en la fórmula $P_A = A(A^T A)^{-1}A^T$ se incluye A en lugar de $\mathbf{a}$.

Podemos repetir el procedimiento anterior para una colección de vectores linealmente independientes $\mathbf{a}_1, \dots, \mathbf{a}_n$ en $\mathbb{R}^m$ con $n < m$. Formamos la matriz A de $m \times n$ cuyos vectores columna son n los vectores dados. La proyección es ahora sobre el subespacio $\mathcal{R}(A) = \text{Gen}\{\mathbf{a}_1, \dots, \mathbf{a}_n\} = \mathcal{R}(A) \subset \mathbb{R}^m$.

TEOREMA 1.9. Dada la matriz A de $m \times n$ con $n < m$ de rango máximo $r = n$, la matriz cuadrada $A^T A$ de $n \times n$ es simétrica e invertible y la matriz $P_A = A(A^T A)^{-1}A^T$ de $m \times m$ (también simétrica) es una proyección ortogonal de $\mathbb{R}^n$ sobre $\mathcal{R}(A)$ (subespacio de $\mathbb{R}^m$).

DEMOSTRACIÓN 1.10. Es claro que $A^T A$ es simétrica de tamaño $n \times n$. Para verificar que $A^T A$ es invertible basta con verificar que su espacio nulo solo contiene el vector cero. Sea $\mathbf{x} \in \mathcal{N}(A^T A)$, entonces

$$A^T A \mathbf{x} = \mathbf{0} \Rightarrow 0 = \mathbf{x}^T A^T A \mathbf{x} = \|A \mathbf{x}\|^2 \Rightarrow A \mathbf{x} = \mathbf{0} \Rightarrow \mathbf{x} = \mathbf{0}.$$

La última implicación es válida porque las columnas de A son linealmente independientes. Por lo tanto, $\mathcal{N}(A^T A) = \{\mathbf{0}\}$ y $A^T A$ es invertible, y la matriz $P_A = A(A^T A)^{-1}A^T$ está bien definida. Para verificar que P_A es una proyección basta con verificar que $P_A^2 = P_A$

$$P_A^2 = A(A^T A)^{-1}A^T A(A^T A)^{-1}A^T = A(A^T A)^{-1}A^T = P_A,$$

y para verificar que es una proyección ortogonal basta verificar que $P_A^T = P_A$:

$$P_A^T = [A(A^T A)^{-1}A^T]^T = A[(A^T A)^{-1}]^T A^T = A(A^T A)^{-1}A^T = P_A.$$

Nota. La suposición que A sea de rango completo es crucial. De otra manera, la matriz $A^T A$ no será invertible y, por lo tanto, la expresión $P_A = A(A^T A)^{-1}A^T$ no tendrá sentido. La tabla 1.1 muestra dos matrices, una con columnas linealmente independientes y otra con columnas linealmente dependientes.

Columnas	A	$A^T A$	$\det(A^T A)$
L.I.	$\begin{bmatrix} 1 & 1 \\ 0 & 1 \\ 0 & 0 \end{bmatrix}$	$\begin{bmatrix} 1 & 1 \\ 1 & 2 \end{bmatrix}$	1 (invertible)
L.D.	$\begin{bmatrix} 1 & 2 \\ 1 & 2 \\ 0 & 0 \end{bmatrix}$	$\begin{bmatrix} 2 & 4 \\ 4 & 8 \end{bmatrix}$	0 (singular)

Tabla 1.1: $A^T A$ cuando A tiene columnas linealmente independientes y cuando no es así.

Para lograr una proyección en el segundo caso, se puede suprimir uno de los vectores columna de A , ya que $\mathcal{R}(A) = \text{Gen}\{\mathbf{a}_1, \mathbf{a}_2\} = \text{Gen}\{\mathbf{a}_1\} = \text{Gen}\{\mathbf{a}_2\}$, y entonces se trata de proyectar sobre la línea generada por cualquiera de las dos columnas. Se deja como ejercicio verificar que en ese caso:

$$P_{a_1} = \frac{1}{2} \begin{bmatrix} 1 & 1 & 0 \\ 1 & 1 & 0 \\ 0 & 0 & 0 \end{bmatrix} = P_{a_2}.$$

1.3

Aproximación de mínimos cuadrados

Consideremos el problema de ajustar una línea dada por $y = c_1 + c_2 t$ a un conjunto de m datos (puntos) $(t_1, y_1), (t_2, y_2), \dots, (t_m, y_m)$. Sabemos que para determinar una línea bastan dos puntos en el plano. En experimentos reales generalmente el número de puntos o datos excede al número $n = 2$ de parámetros a determinar: c_1, c_2 . Por lo tanto, si los m puntos (x_i, y_i) no son colineales, no existe línea que los contenga a todos. En estos casos, solo podemos aspirar a construir una línea que es cercana a los puntos dados. Para ilustrar consideremos el siguiente ejemplo.

EJEMPLO 1.11. Sean los puntos $(0, -1), (1, 0), (2, 4)$ en el plano. Encontrar la línea que mejor ajusta (más cercana) a los puntos dados.

Solución. Los coeficientes c_1, c_2 en $y = c_1 + c_2 t$ son constantes que deben determinarse a partir de los tres puntos dados. Una primera aproximación al problema sería forzar la igualdad en cada punto:

$$\begin{array}{rcl} c_1 + c_2 t_1 & = & y_1 \\ c_1 + c_2 t_2 & = & y_2 \\ c_1 + c_2 t_3 & = & y_3 \end{array} \Rightarrow \begin{array}{rcl} c_1 + c_2(0) & = & -1 \\ c_1 + c_2(1) & = & 0 \\ c_1 + c_2(2) & = & 4 \end{array} \Rightarrow \begin{bmatrix} 1 & 0 \\ 1 & 1 \\ 1 & 2 \end{bmatrix} \begin{bmatrix} c_1 \\ c_2 \end{bmatrix} = \begin{bmatrix} -1 \\ 0 \\ 4 \end{bmatrix},$$

y debemos resolver el sistema lineal de tres ecuaciones con dos incógnitas $A\mathbf{x} = \mathbf{b}$, donde

$$A = \begin{bmatrix} 1 & 0 \\ 1 & 1 \\ 1 & 2 \end{bmatrix}, \quad \mathbf{x} = \begin{bmatrix} c_1 \\ c_2 \end{bmatrix}, \quad \mathbf{b} = \begin{bmatrix} -1 \\ 0 \\ 4 \end{bmatrix}.$$

Sin embargo, el sistema no tiene solución porque $\mathbf{b}$ no está en el espacio columna de A (el plano generado por las columnas de A). Para verificarlo, realizamos eliminación de Gauss:

$$[A \mid \mathbf{b}] = \left[\begin{array}{cc|c} 1 & 0 & -1 \\ 1 & 1 & 0 \\ 1 & 2 & 4 \end{array} \right] \sim \left[\begin{array}{cc|c} 1 & 0 & -1 \\ 1 & 1 & 0 \\ 0 & 2 & 5 \end{array} \right] \sim \left[\begin{array}{cc|c} 1 & 0 & -1 \\ 0 & 1 & 1 \\ 0 & 0 & 3 \end{array} \right].$$

El sistema escalonado reducido resultante es inconsistente (ver coeficientes en rojo). La inconsistencia proviene de que los tres puntos dados en el problema no son colineales. Por lo tanto, la estrategia planteada no es factible para resolver el problema. En su lugar se busca la recta más cercana a los puntos dados, aunque la recta no contenga los puntos. Para ello se busca minimizar la diferencia $A\mathbf{x} - \mathbf{b}$, que equivale a minimizar la función cuadrática $f: \mathbb{R}^2 \mapsto \mathbb{R}$, definida por

$$f(\mathbf{x}) = \|\mathbf{b} - A\mathbf{x}\|^2 = \mathbf{x}^T A^T A \mathbf{x} - 2\mathbf{x}^T A^T \mathbf{b} + \mathbf{b}^T \mathbf{b}$$

El gradiente y la Hessiana de esta función son:

$$\begin{aligned} \nabla f(\mathbf{x}) &= 2A^T A \mathbf{x} - 2A^T \mathbf{b}, \\ H_f(\mathbf{x}) &= 2A^T A. \end{aligned}$$

Como la Hessiana $H_f(\mathbf{x})$ es una matriz definida positiva, pues $\mathbf{x}^T A^T A \mathbf{x} = \|A\mathbf{x}\|^2 > 0$ si $\mathbf{x} \neq \mathbf{0}$, entonces el mínimo $\hat{\mathbf{x}}$ satisface $\nabla f(\mathbf{x}) = \mathbf{0}$, es decir satisface las **ecuaciones normales**:

$$A^T A \mathbf{x} = A^T \mathbf{b} \quad \text{con} \quad A^T A = \begin{bmatrix} 3 & 3 \\ 3 & 5 \end{bmatrix}, \quad A^T \mathbf{b} = \begin{bmatrix} 3 \\ 8 \end{bmatrix}.$$

Despejando $\hat{\mathbf{x}}$ se obtiene el mínimo:

$$\hat{\mathbf{x}} = (A^T A)^{-1} A^T \mathbf{b} \Rightarrow \hat{\mathbf{x}} = \begin{bmatrix} c_1 \\ c_2 \end{bmatrix} = \begin{bmatrix} -3/2 \\ 5/2 \end{bmatrix}.$$

Por lo tanto, la recta buscada está dada por la siguiente expresión y se muestra en la figura 1.6 junto con los tres puntos dados (no colineales).

$$y = -\frac{3}{2} + \frac{5}{2}t.$$

Nota. Para verificar la solución $\hat{\mathbf{x}} = (A^T A)^{-1} A^T \mathbf{b}$, multiplicamos por A ambos lados, obteniendo

$$A\hat{\mathbf{x}} = A(A^T A)^{-1} A^T \mathbf{b} \Rightarrow A\hat{\mathbf{x}} = P_A \mathbf{b} \quad \text{con} \quad P_A = A(A^T A)^{-1} A^T = \frac{1}{6} \begin{bmatrix} 5 & 2 & -1 \\ 2 & 2 & 2 \\ -1 & 2 & 5 \end{bmatrix}.$$

Ahora verificamos que ambos lados de $A\hat{\mathbf{x}} = P_A \mathbf{b}$ deben ser iguales:

$$A\hat{\mathbf{x}} = \begin{bmatrix} 1 & 0 \\ 1 & 1 \\ 1 & 2 \end{bmatrix} \begin{bmatrix} -3/2 \\ 5/2 \end{bmatrix} = \begin{bmatrix} -3/2 \\ 1 \\ 7/2 \end{bmatrix} \quad \text{y} \quad P_A \mathbf{b} = \frac{1}{6} \begin{bmatrix} 5 & 2 & -1 \\ 2 & 2 & 2 \\ -1 & 2 & 5 \end{bmatrix} \begin{bmatrix} -1 \\ 0 \\ 4 \end{bmatrix} = \begin{bmatrix} -3/2 \\ 1 \\ 7/2 \end{bmatrix},$$

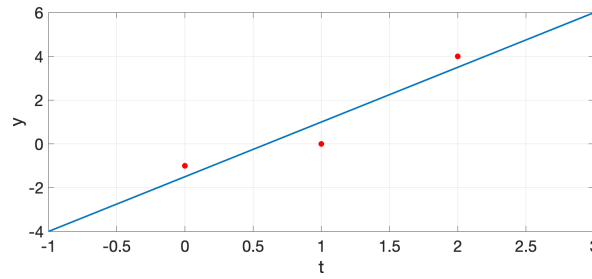


Figura 1.6: Ajuste de una recta a tres datos (puntos) dados.

los cuales coinciden.

Nota. Anteriormente verificamos que $\mathbf{b} = [-1 \ 0 \ 4]^T$ no se encuentra en el espacio columna de A y por eso el sistema $A\mathbf{x} = \mathbf{b}$ no tiene solución. Por otro lado, la expresión $A\hat{\mathbf{x}} = P_A\mathbf{b}$ indica que la solución encontrada por mínimos cuadrados es equivalente a proyectar primero $\mathbf{b}$ sobre el espacio columna de A , obteniendo $P_A\mathbf{b}$, y luego resolver la ecuación consistente $A\hat{\mathbf{x}} = P_A\mathbf{b}$.

El procedimiento anterior se puede aplicar para ajustar polinomios de grado mayor. Por ejemplo, un polinomio de grado dos $y = c_1 + c_2 t + c_3 t^2$ se puede ajustar a un conjunto de $m \geq 3$ puntos $(t_1, y_1), \dots, (t_m, y_m)$, la **matriz de diseño** A y el **vector de observación** $\mathbf{b}$ son

$$A = \begin{bmatrix} 1 & t_1 & t_1^2 \\ 1 & t_2 & t_2^2 \\ \vdots & \vdots & \vdots \\ 1 & t_m & t_m^2 \end{bmatrix}, \quad \mathbf{b} = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_m \end{bmatrix}.$$

Al igual que en caso anterior, **los coeficientes desconocidos forman el vector columna** $\mathbf{x} = [c_1 \ c_2 \ c_3]^T$. Para calcular estos coeficientes desconocidos debemos resolver el sistema de ecuaciones normales:

$$A^T A \mathbf{x} = A^T \mathbf{b}.$$

EJEMPLO 1.12. Ajustar el polinomio de grado dos que mejor aproxima a los cuatro puntos $(0, -1), (1, 0), (2, 4), (3, 6)$.

Solución. La matriz de diseño A y el vector de observación $\mathbf{b}$ son

$$A = \begin{bmatrix} 1 & 0 & 0 \\ 1 & 1 & 1 \\ 1 & 2 & 4 \\ 1 & 3 & 9 \end{bmatrix}, \quad \mathbf{b} = \begin{bmatrix} -1 \\ 0 \\ 4 \\ 6 \end{bmatrix}.$$

El sistema de ecuaciones normales es

$$A^T A \mathbf{x} = A^T \mathbf{b} \Rightarrow \begin{bmatrix} 4 & 6 & 14 \\ 6 & 14 & 36 \\ 14 & 36 & 98 \end{bmatrix} \begin{bmatrix} c_1 \\ c_2 \\ c_3 \end{bmatrix} = \begin{bmatrix} 9 \\ 26 \\ 70 \end{bmatrix}.$$

Hacemos eliminación para resolver el sistema (se muestran solo algunos pasos intermedios)

$$\begin{bmatrix} 4 & 6 & 14 & | & 9 \\ 6 & 14 & 36 & | & 26 \\ 14 & 36 & 98 & | & 70 \end{bmatrix} \sim \begin{bmatrix} 1 & 3/2 & 7/2 & | & 9/4 \\ 1 & 7/3 & 6 & | & 13/3 \\ 1 & 18/7 & 7 & | & 5 \end{bmatrix} \sim \begin{bmatrix} 1 & 3/2 & 7/2 & | & 9/4 \\ 0 & 1 & 3 & | & 5/2 \\ 0 & 1 & 49/15 & | & 77/30 \end{bmatrix}$$

$$\sim \begin{bmatrix} 1 & 3/2 & 7/2 & | & 9/4 \\ 0 & 1 & 3 & | & 5/2 \\ 0 & 0 & 1 & | & 1/4 \end{bmatrix} \sim \begin{bmatrix} 1 & 3/2 & 0 & | & 11/8 \\ 0 & 1 & 0 & | & 7/4 \\ 0 & 0 & 1 & | & 1/4 \end{bmatrix} \sim \begin{bmatrix} 1 & 0 & 0 & | & -5/4 \\ 0 & 1 & 0 & | & 7/4 \\ 0 & 0 & 1 & | & 1/4 \end{bmatrix}.$$

Por lo tanto, la solución es $c_1 = -5/4$, $c_2 = 7/4$, $c_3 = 1/4$, y el polinomio de segundo grado es

$$y = -\frac{5}{4} + \frac{7}{4}t + \frac{1}{4}t^2.$$

La figura 1.7 muestra la gráfica de la parábola junto a los puntos dados. Esta parábola es la mas cercana a dichos puntos en el plano.

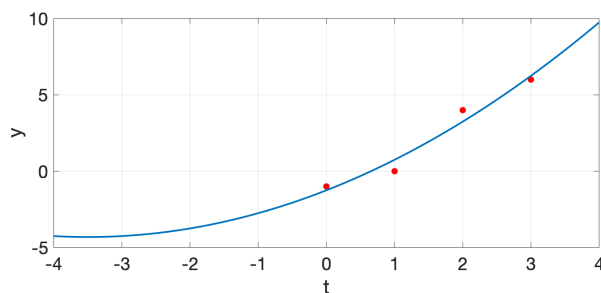


Figura 1.7: Ajuste de una parábola a cuatro puntos dados.

1.4

Bases ortogonales. Ortogonalización de Gram-Schmidt

En este apartado mostraremos ciertas ventajas de trabajar con conjuntos y bases de vectores ortogonales. Al final mostraremos el proceso de Gram-Schmidt para generar de manera sistemática esos conjuntos de vectores ortogonales a partir de conjuntos no ortogonales de vectores linealmente independientes.

Bases ortogonales y ortonormales

DEFINICIÓN 1.13. Conjuntos de vectores ortogonales y ortonormales. Un conjunto de vectores $\{\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_p\}$ se denomina **ortogonal** si cada vector del conjunto es ortogonal a los restantes. Entonces, para cada par de vectores distintos se satisface $\mathbf{a}_i \cdot \mathbf{a}_j = \mathbf{a}_i^T \mathbf{a}_j = 0$ si $i \neq j$. Si además los vectores son unitarios, entonces el conjunto se denomina **ortonormal**.

El ejemplo mas simple y familiar, es el conjunto de vectores canónicos $\{\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n\}$ del espacio $\mathbb{R}^n$. Sabemos que este conjunto de vectores ortonormales forman una base para $\mathbb{R}^n$.

Asimismo, cualquier otro conjunto ortonormal de n vectores $\{\mathbf{q}_1, \mathbf{q}_2, \dots, \mathbf{q}_n\}$ formará una base para $\mathbb{R}^n$, pues los vectores son linealmente independientes por ser mutuamente ortogonales. En ese caso, cualquier vector $\mathbf{b} \in \mathbb{R}^n$ se puede expresar como la combinación lineal

$$\mathbf{b} = c_1 \mathbf{q}_1 + c_2 \mathbf{q}_2 + \dots + c_n \mathbf{q}_n,$$

en donde cada coeficiente $c_j = \mathbf{q}_j^T \mathbf{b}$ se encuentra multiplicando escalarmente ambos lados de la igualdad por $\mathbf{q}_j$. Por lo tanto, debido a que $\mathbf{q}_i^T \mathbf{q}_j = \delta_{ij}$, obtenemos

$$\mathbf{b} = (\mathbf{q}_1^T \mathbf{b}) \mathbf{q}_1 + (\mathbf{q}_2^T \mathbf{b}) \mathbf{q}_2 + \dots + (\mathbf{q}_n^T \mathbf{b}) \mathbf{q}_n.$$

Obsérvese que cada sumando $(\mathbf{q}_j^T \mathbf{b}) \mathbf{q}_j = (\mathbf{q}_j \mathbf{q}_j^T) \mathbf{b} = P_j \mathbf{b}$ representa la proyección ortogonal de $\mathbf{b}$ sobre la ‘coordenada’ definida por el vector $\mathbf{q}_j$.

EJEMPLO 1.14. El conjunto de vectores $\mathbf{a}_1 = (-1, 2, 2)$, $\mathbf{a}_2 = (2, 2, -1)$, $\mathbf{a}_3 = (2, -1, 2)$, es ortogonal, dado que el producto interno de cualquier par de ellos es cero. Normalizando, encontramos el siguiente conjunto ortonormal:

$$\{\mathbf{q}_1, \mathbf{q}_2, \mathbf{q}_3\} = \left\{ \begin{bmatrix} -1/3 \\ 2/3 \\ 2/3 \end{bmatrix}, \begin{bmatrix} 2/3 \\ 2/3 \\ -1/3 \end{bmatrix}, \begin{bmatrix} 2/3 \\ -1/3 \\ 2/3 \end{bmatrix} \right\}.$$

Cualquier vector $\mathbf{b} \in \mathbb{R}^3$ es combinación lineal de ellos. Por ejemplo, $\mathbf{b} = (1, 2, -1)$ es la combinación lineal

$$\mathbf{b} = \frac{1}{3} \mathbf{q}_1 + \frac{7}{3} \mathbf{q}_2 - \frac{2}{3} \mathbf{q}_3, \quad \text{pues} \quad \mathbf{q}_1^T \mathbf{b} = \frac{1}{3}, \quad \mathbf{q}_2^T \mathbf{b} = \frac{7}{3}, \quad \mathbf{q}_3^T \mathbf{b} = -\frac{2}{3}.$$

La ventaja de las bases ortonormales sobre las bases que no lo son, es que para el cálculo de las coordenadas de cualquier vector $\mathbf{b} \in \mathbb{R}^n$ no es necesario resolver un sistema de ecuaciones, basta con utilizar las propiedades de ortonormalidad de la base, como se muestra en el ejemplo anterior ¹.

Proyecciones sobre coordenadas ortogonales

Otro ámbito en donde se obtiene simplificaciones con conjuntos ortogonales y ortonormales son las proyecciones. Por ejemplo, si A es una matriz con vectores columna mutuamente ortogonales, el cálculo de la matriz $A^T A$ en las ecuaciones normales, o bien el cálculo la matriz de proyección $P_A = A(A^T A)^{-1} A^T$, se simplificarán. Específicamente, si A de $m \times n$ tiene n columnas $\{\mathbf{a}_1, \dots, \mathbf{a}_n\}$ ortogonales en $\mathbb{R}^m$, entonces $A^T A$ será una matriz diagonal de $n \times n$, pues

$$A^T A = \begin{bmatrix} - & \mathbf{a}_1^T & - \\ - & \mathbf{a}_2^T & - \\ & \vdots & \\ - & \mathbf{a}_n^T & - \end{bmatrix} \begin{bmatrix} | & | & & | \\ \mathbf{a}_1 & \mathbf{a}_2 & \dots & \mathbf{a}_n \\ | & | & & | \end{bmatrix} = \begin{bmatrix} \mathbf{a}_1^T \mathbf{a}_1 & 0 & \dots & 0 \\ 0 & \mathbf{a}_2^T \mathbf{a}_2 & \dots & 0 \\ 0 & 0 & \ddots & 0 \\ 0 & 0 & \dots & \mathbf{a}_n^T \mathbf{a}_n \end{bmatrix},$$

¹Para una base arbitraria no ortogonal $\{\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n\}$ el cálculo de los coeficientes $c_1, c_2, \dots, c_n$ en la combinación

$$\mathbf{b} = c_1 \mathbf{a}_1 + c_2 \mathbf{a}_2 + \dots + c_n \mathbf{a}_n$$

requiere resolver el sistema de ecuaciones $A \mathbf{x} = \mathbf{b}$, donde $A = [\mathbf{a}_1 | \mathbf{a}_2 | \dots | \mathbf{a}_n]$ contiene los vectores $\mathbf{a}_j$ como columnas y $\mathbf{x} = (c_1, \dots, c_n)^T$ es el vector con las incógnitas.

ya que cada producto fuera de la diagonal es de la forma $\mathbf{a}_i^T \mathbf{a}_j = 0$ con $i \neq j$.

EJEMPLO 1.15. Consideremos la matriz cuyas columnas son los vectores de coordenadas $\mathbf{a}_1 = (-1, 2, 2)$, $\mathbf{a}_3 = (4, -2, 4)$. Estos vectores son ortogonales. Las matrices A y $A^T A$ son

$$A = \begin{bmatrix} -1 & 4 \\ 2 & -2 \\ 2 & 4 \end{bmatrix}, \quad A^T A = \begin{bmatrix} -1 & 2 & 2 \\ 4 & -2 & 4 \end{bmatrix} \begin{bmatrix} -1 & 4 \\ 2 & -2 \\ 2 & 4 \end{bmatrix} = \begin{bmatrix} 9 & 0 \\ 0 & 36 \end{bmatrix}.$$

La segunda matriz es diagonal y el cálculo de la matriz de proyección $P_A = A(A^T A)^{-1}A^T$ se simplifica porque

$$(A^T A)^{-1} = \begin{bmatrix} 1/9 & 0 \\ 0 & 1/36 \end{bmatrix}$$

y

$$P_A = A(A^T A)^{-1}A^T = \begin{bmatrix} -1 & 4 \\ 2 & -2 \\ 2 & 4 \end{bmatrix} \begin{bmatrix} 1/9 & 0 \\ 0 & 1/36 \end{bmatrix} \begin{bmatrix} -1 & 2 & 2 \\ 4 & -2 & 4 \end{bmatrix} = \frac{1}{9} \begin{bmatrix} 5 & -4 & 2 \\ -4 & 5 & 2 \\ 2 & 2 & 8 \end{bmatrix}.$$

Si los vectores columna son ortonormales, los cálculos se simplifican aún más.

EJEMPLO 1.16. En el ejemplo anterior si normalizamos los vectores $\mathbf{a}_1, \mathbf{a}_2$, obtenemos los vectores ortonormales $\mathbf{q}_1 = (-1/3, 2/3, 2/3)$, $\mathbf{q}_2 = (2/3, -1/3, 2/3)$. Por lo tanto, las matrices Q y $Q^T Q$ son

$$Q = \frac{1}{3} \begin{bmatrix} -1 & 2 \\ 2 & -1 \\ 2 & 2 \end{bmatrix}, \quad Q^T Q = \frac{1}{9} \begin{bmatrix} -1 & 2 & 2 \\ 2 & -1 & 2 \end{bmatrix} \begin{bmatrix} -1 & 2 \\ 2 & -1 \\ 2 & 2 \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}.$$

Por lo que la matriz de proyección sobre el subespacio $\text{Gen}\{\mathbf{a}_1, \mathbf{a}_2\} = \text{Gen}\{\mathbf{q}_1, \mathbf{q}_2\}$ en este caso será

$$P_Q = Q(Q^T Q)^{-1}Q^T = QQ^T = \frac{1}{9} \begin{bmatrix} 5 & -4 & 2 \\ -4 & 5 & 2 \\ 2 & 2 & 8 \end{bmatrix}.$$

La matriz de proyección coincide con la encontrada previamente, en donde los vectores ortogonales $\mathbf{a}_1, \mathbf{a}_2$ no están normalizados.

Matrices ortogonales y sus propiedades

La matriz Q del ejemplo anterior tiene vectores columna ortonormales y es muy fácil de manipular pues $Q^T Q = I$. Esta propiedad es general para cualquier matriz cuyas columnas sean ortonormales. Esquemáticamente, se tiene

$$Q^T Q = \begin{bmatrix} - & \mathbf{q}_1^T & - \\ - & \mathbf{q}_2^T & - \\ & \vdots & \\ - & \mathbf{q}_n^T & - \end{bmatrix} \begin{bmatrix} | & | & & | \\ \mathbf{q}_1 & \mathbf{q}_2 & \cdots & \mathbf{q}_n \\ | & | & & | \end{bmatrix} = \begin{bmatrix} 1 & 0 & \cdots & 0 \\ 0 & 1 & \cdots & 0 \\ 0 & 0 & \ddots & 0 \\ 0 & 0 & \cdots & 1 \end{bmatrix} = I,$$

en donde el producto del renglón i de la matriz Q^T con la columna j de la matriz Q es $\mathbf{q}_i^T \mathbf{q}_j = \delta_{ij}$. Por lo tanto, fuera de la diagonal esos productos son cero y sobre la diagonal $i = j$ y $\mathbf{q}_i^T \mathbf{q}_i = \|\mathbf{q}_i\|^2 = 1$.

En algunas aplicaciones la matriz Q es rectangular, de $m \times n$, con $m > n$, por ejemplo en problemas de mínimos cuadrados. En esos casos $\mathbf{q}_j \in \mathbb{R}^m$ y $Q^T Q = I$, con I la matriz identidad de tamaño n (ver ejemplo anterior). Cuando la matriz Q es cuadrada, $m = n$, y el producto $Q^T Q = I$ significa que $Q^{-1} = Q^T$.

DEFINICIÓN 1.17. Una **matriz cuadrada** Q con columnas ortonormales se denomina **matriz ortogonal**. En este caso $Q^T Q = I$, es decir $Q^{-1} = Q^T$.

Estas matrices tienen propiedades geométricas importantes debido a que preservan el producto interior, ya que

$$(Q\mathbf{x})^T (Q\mathbf{y}) = \mathbf{x}^T Q^T Q \mathbf{y} = \mathbf{x}^T \mathbf{y}, \quad \text{para todo } \mathbf{x}, \mathbf{y} \in \mathbb{R}^n.$$

La preservación del producto interno implica que los ángulos entre vectores se preservan bajo la acción de Q . De la misma manera es fácil verificar que se preservan las longitudes

$$\|Q\mathbf{x}\| = \|\mathbf{x}\|.$$

Por las razones anteriores las matrices ortogonales son muy importantes en geometría y aplicaciones. Ejemplos típicos de matrices ortogonales incluyen las matrices de rotación, las matrices de permutación y las de reflexión.

EJEMPLO 1.18. Matriz de rotación. La matriz

$$Q = \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix} = [\mathbf{q}_1 \mid \mathbf{q}_2] \quad \text{con} \quad Q^{-1} = \begin{bmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{bmatrix} = Q^T$$

gira cada punto en $\mathbb{R}^2$ un ángulo θ . Las columnas $\mathbf{q}_1, \mathbf{q}_2$, de Q son ortonormales, basta con calcular el producto interno $\mathbf{q}_1^T \mathbf{q}_2 = (\cos \theta)(-\sin \theta) + (\sin \theta)(\cos \theta) = 0$ y $\|\mathbf{q}_1\| = \|\mathbf{q}_2\| = (\cos^2 \theta + \sin^2 \theta)^{1/2} = 1$. La matriz inversa Q^{-1} rota vectores en sentido contrario un ángulo $-\theta$ y coincide con Q^T , pues $\cos(-\theta) = \cos \theta$ y $\sin(-\theta) = -\sin \theta$. Por supuesto, $Q^T Q = Q Q^T = I$.

EJEMPLO 1.19. Matriz de permutación. Una matriz de permutación se obtiene permutando las columnas o filas de la matriz identidad. Por ejemplo, la matriz

$$Q = \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 1 & 0 & 0 \end{bmatrix} \quad \text{con} \quad Q^{-1} = \begin{bmatrix} 0 & 0 & 1 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \end{bmatrix} = Q^T$$

permuta cada vector $\mathbf{x} = (x_1, x_2, x_3)^T \in \mathbb{R}^3$ en $Q\mathbf{x} = (x_2, x_3, x_1)^T$. La matriz inversa reordena el vector permutado a su orden original, es decir $Q^{-1}(x_2, x_3, x_1)^T = (x_1, x_2, x_3)^T$.

EJEMPLO 1.20. Matriz de reflexión. Dado un vector unitario $\mathbf{u}$, la matriz $Q = I - 2\mathbf{u}\mathbf{u}^T$ es una matriz ortogonal, pues $Q^2 = Q$ y $Q^T = Q$. Específicamente

$$Q^T Q = Q Q^T = Q^2 = (I - 2\mathbf{u}\mathbf{u}^T)(I - 2\mathbf{u}\mathbf{u}^T) = I - 4\mathbf{u}\mathbf{u}^T + 4(\mathbf{u}\mathbf{u}^T)(\mathbf{u}\mathbf{u}^T) = I,$$

pues $(\mathbf{u}\mathbf{u}^T)(\mathbf{u}\mathbf{u}^T) = \mathbf{u}(\mathbf{u}^T\mathbf{u})\mathbf{u}^T = \mathbf{u}\|\mathbf{u}\|^2\mathbf{u}^T = \mathbf{u}\mathbf{u}^T$. Por lo tanto $Q^{-1} = Q^T = Q$. La matriz Q actúa sobre el vector unitario $\mathbf{u}$, produciendo su reflexión sobre el origen:

$$Q\mathbf{u} = (I - 2\mathbf{u}\mathbf{u}^T)\mathbf{u} = \mathbf{u} - 2\mathbf{u}(\mathbf{u}^T\mathbf{u}) = \mathbf{u} - 2\mathbf{u} = -\mathbf{u}.$$

La reflexión $Q\mathbf{x}$ de cualquier otro vector $\mathbf{x}$ se ilustra en la figura 1.8 para el caso bidimensional. En el caso n -dimensional el ‘espejo’ es un hiperplano con vector normal $\mathbf{u}$ y que pasa por el origen.

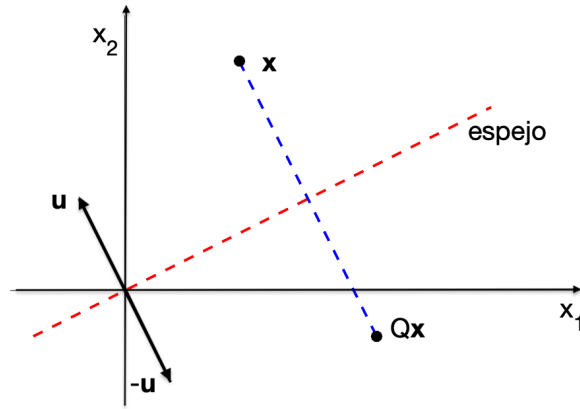


Figura 1.8: Acción de la matriz de reflexión Q definida por el vector unitario $\mathbf{u}$.

Por supuesto si el vector que define la reflexión es un vector $\mathbf{v}$ no unitario, entonces normalizando tenemos $\mathbf{u} = \mathbf{v}/\|\mathbf{v}\|$ y la matriz de proyección es

$$Q = I - 2\frac{\mathbf{v}\mathbf{v}^T}{\mathbf{v}^T\mathbf{v}}, \quad \text{pues } \mathbf{v}^T\mathbf{v} = \|\mathbf{v}\|^2.$$

Proceso de ortogonalización de Gram–Schmidt

Una vez convencidos de la conveniencia de los conjuntos ortogonales y ortonormales, ahora nos centraremos en la construcción de una base ortonormal partiendo de una base de vectores linealmente independientes.

EJEMPLO 1.21. Comencemos con un ejemplo sencillo en $\mathbb{R}^3$. Sean $\{\mathbf{a}_1, \mathbf{a}_2, \mathbf{a}_3\}$, tres vectores linealmente independientes, no necesariamente ortogonales. Queremos construir con estos vectores otra base $\{\mathbf{q}_1, \mathbf{q}_2, \mathbf{q}_3\}$ de $\mathbb{R}^3$ que sea ortonormal. Seguimos los siguientes pasos:

Paso 1. Escogemos un primer vector, digamos $\mathbf{v}_1 = \mathbf{a}_1$.

Paso 2. Ahora escogemos $\mathbf{a}_2$ y sustraemos su proyección a lo largo de $\mathbf{v}_1$, obteniendo

$$\mathbf{v}_2 = \mathbf{a}_2 - \frac{\mathbf{v}_1^T \mathbf{a}_2}{\mathbf{v}_1^T \mathbf{v}_1} \mathbf{v}_1, \quad \text{ortogonal a } \mathbf{v}_1.$$

Paso 3. Finalmente, tomamos el vector $\mathbf{a}_3$. Este vector no es combinación lineal de $\mathbf{v}_1$ y $\mathbf{v}_2$ ya que $\mathbf{a}_3$ no es combinación lineal de $\mathbf{a}_1$ y $\mathbf{a}_2$. Por lo tanto, sustraemos sus componentes en las dos direcciones $\mathbf{v}_1$ y $\mathbf{v}_2$ para obtener el vector

$$\mathbf{v}_3 = \mathbf{a}_3 - \frac{\mathbf{v}_1^T \mathbf{a}_3}{\mathbf{v}_1^T \mathbf{v}_1} \mathbf{v}_1 - \frac{\mathbf{v}_2^T \mathbf{a}_3}{\mathbf{v}_2^T \mathbf{v}_2} \mathbf{v}_2, \quad \text{ortogonal a ambos, } \mathbf{v}_1, \mathbf{v}_2.$$

A continuación verificamos que efectivamente los vectores $\mathbf{v}_1, \mathbf{v}_2, \mathbf{v}_3$ son ortogonales entre sí. Los vectores $\mathbf{v}_1$ y $\mathbf{v}_2$ son ortogonales, pues

$$\mathbf{v}_1^T \mathbf{v}_2 = \mathbf{v}_1^T \mathbf{a}_2 - \frac{\mathbf{v}_1^T \mathbf{a}_2}{\mathbf{v}_1^T \mathbf{v}_1} \mathbf{v}_1^T \mathbf{v}_1 = \mathbf{v}_1^T \mathbf{a}_2 - \mathbf{v}_1^T \mathbf{a}_2 = 0.$$

Asimismo, $\mathbf{v}_3$ es ortogonal a $\mathbf{v}_1$ y $\mathbf{v}_2$ ya que

$$\mathbf{v}_1^T \mathbf{v}_3 = \mathbf{v}_1^T \mathbf{a}_3 - \frac{\mathbf{v}_1^T \mathbf{a}_3}{\mathbf{v}_1^T \mathbf{v}_1} \mathbf{v}_1^T \mathbf{v}_1 - \frac{\mathbf{v}_2^T \mathbf{a}_3}{\mathbf{v}_2^T \mathbf{v}_2} \mathbf{v}_1^T \mathbf{v}_2 = \mathbf{v}_1^T \mathbf{a}_3 - \mathbf{v}_1^T \mathbf{a}_3 - 0 = 0,$$

$$\mathbf{v}_2^T \mathbf{v}_3 = \mathbf{v}_2^T \mathbf{a}_3 - \frac{\mathbf{v}_1^T \mathbf{a}_3}{\mathbf{v}_2^T \mathbf{v}_1} \mathbf{v}_2^T \mathbf{v}_1 - \frac{\mathbf{v}_2^T \mathbf{a}_3}{\mathbf{v}_2^T \mathbf{v}_2} \mathbf{v}_2^T \mathbf{v}_2 = \mathbf{v}_2^T \mathbf{a}_3 - 0 - \mathbf{v}_2^T \mathbf{a}_3 = 0.$$

Por lo tanto, los tres vectores $\mathbf{v}_1, \mathbf{v}_2, \mathbf{v}_3$ son ortogonales entre sí. Finalmente, normalizamos estos vectores, obteniendo el siguiente conjunto ortonormal:

$$\{\mathbf{q}_1, \mathbf{q}_2, \mathbf{q}_3\} \quad \text{con} \quad \mathbf{q}_1 = \frac{\mathbf{v}_1}{\|\mathbf{v}_1\|}, \quad \mathbf{q}_2 = \frac{\mathbf{v}_2}{\|\mathbf{v}_2\|}, \quad \mathbf{q}_3 = \frac{\mathbf{v}_3}{\|\mathbf{v}_3\|}.$$

EJEMPLO 1.22. Considerar los siguientes tres vectores linealmente independientes

$$\mathbf{a}_1 = \begin{bmatrix} 1 \\ -1 \\ 0 \end{bmatrix}, \quad \mathbf{a}_2 = \begin{bmatrix} 2 \\ 0 \\ -2 \end{bmatrix}, \quad \mathbf{a}_3 = \begin{bmatrix} 3 \\ -3 \\ 3 \end{bmatrix}.$$

Encontrar una base ortonormal de $\mathbb{R}^3$ a partir de estos vectores.

Solución. En efecto, es posible verificar que los tres vectores son linealmente independientes. Aplicamos el procedimiento:

Paso 1.

$$\mathbf{v}_1 = \mathbf{a}_1 = \begin{bmatrix} 1 \\ -1 \\ 0 \end{bmatrix}.$$

Paso 2.

$$\mathbf{v}_2 = \mathbf{a}_2 - \frac{\mathbf{v}_1^T \mathbf{a}_2}{\mathbf{v}_1^T \mathbf{v}_1} \mathbf{v}_1 = \begin{bmatrix} 2 \\ 0 \\ -2 \end{bmatrix} - \frac{2}{2} \begin{bmatrix} 1 \\ -1 \\ 0 \end{bmatrix} = \begin{bmatrix} 1 \\ 1 \\ -2 \end{bmatrix}.$$

Paso 3.

$$\mathbf{v}_3 = \mathbf{a}_3 - \frac{\mathbf{v}_1^T \mathbf{a}_3}{\mathbf{v}_1^T \mathbf{v}_1} \mathbf{v}_1 - \frac{\mathbf{v}_2^T \mathbf{a}_3}{\mathbf{v}_2^T \mathbf{v}_2} \mathbf{v}_2 = \begin{bmatrix} 3 \\ -3 \\ 3 \end{bmatrix} - \frac{6}{2} \begin{bmatrix} 1 \\ -1 \\ 0 \end{bmatrix} - \frac{-6}{6} \begin{bmatrix} 1 \\ 1 \\ -2 \end{bmatrix} = \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix}.$$

Normalizando los vectores $\mathbf{v}_1, \mathbf{v}_2, \mathbf{v}_3$, se obtiene

$$\mathbf{q}_1 = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ -1 \\ 0 \end{bmatrix}, \quad \mathbf{q}_2 = \frac{1}{\sqrt{6}} \begin{bmatrix} 1 \\ 1 \\ -2 \end{bmatrix}, \quad \mathbf{q}_3 = \frac{1}{\sqrt{3}} \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix}.$$

El procedimiento anterior se puede generalizar para conjuntos de vectores ortogonales en espacios generales $\mathbb{R}^n$.

TEOREMA 1.23. *El proceso Gram-Schmidt. Dada una base $\{\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_p\}$, $p \leq n$, para un subespacio $V \subset \mathbb{R}^n$ de dimensión p , se puede construir una base ortonormal $\{\mathbf{q}_1, \mathbf{q}_2, \dots, \mathbf{q}_p\}$ de V por medio de la siguiente secuencia de operaciones:*

$$\begin{aligned} \mathbf{q}_1 &= \frac{\mathbf{v}_1}{\|\mathbf{v}_1\|} & \text{con } \mathbf{v}_1 &= \mathbf{a}_1, \\ \mathbf{q}_2 &= \frac{\mathbf{v}_2}{\|\mathbf{v}_2\|} & \text{con } \mathbf{v}_2 &= \mathbf{a}_2 - (\mathbf{q}_1^T \mathbf{a}_2) \mathbf{q}_1, \\ \mathbf{q}_3 &= \frac{\mathbf{v}_3}{\|\mathbf{v}_3\|} & \text{con } \mathbf{v}_3 &= \mathbf{a}_3 - (\mathbf{q}_1^T \mathbf{a}_3) \mathbf{q}_1 - (\mathbf{q}_2^T \mathbf{a}_3) \mathbf{q}_2, \\ &\vdots & &\vdots \\ \mathbf{q}_p &= \frac{\mathbf{v}_p}{\|\mathbf{v}_p\|} & \text{con } \mathbf{v}_p &= \mathbf{a}_p - (\mathbf{q}_1^T \mathbf{a}_p) \mathbf{q}_1 - (\mathbf{q}_2^T \mathbf{a}_p) \mathbf{q}_2 - \dots - (\mathbf{q}_{p-1}^T \mathbf{a}_p) \mathbf{q}_{p-1}, \end{aligned}$$

y cualquier vector $\mathbf{x} \in V$ tiene las expansión

$$\mathbf{x} = (\mathbf{q}_1^T \mathbf{x}) \mathbf{q}_1 + (\mathbf{q}_2^T \mathbf{x}) \mathbf{q}_2 + \dots + (\mathbf{q}_p^T \mathbf{x}) \mathbf{q}_p.$$

Además

$$V = \text{Gen} \{\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_p\} = \text{Gen} \{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_p\} = \text{Gen} \{\mathbf{q}_1, \mathbf{q}_2, \dots, \mathbf{q}_p\}.$$

1.5

Referencias

1. Gilbert Strang, *Introduction to Linear Algebra*, 5th Edition, Wellesley - Cambridge Press, 2016. ISBN 978-0-9802327-7-6.

Los conceptos importantes introducidos en este capítulo (y los siguientes) se encontrarán sin excepción en este libro de texto. Es un libro moderno de álgebra lineal con orientación a las aplicaciones. Muy recomendable también como libro de texto para enseñar álgebra lineal, en mi opinión.

2. <http://ulaff.net/>

Sitio en internet contiene abundante material sobre álgebra lineal y aplicaciones, en forma de videos, notas, programación, entre otros materiales.

3. David C. Lay, *Algebra Lineal y sus Aplicaciones*, tercera edición, Pearson Addison - Wesley, 2006 (2007 en México). ISBN: 978-970-26-0906-3.

Libro de texto que contiene un enfoque básico pero en forma didáctica, clara y extensa. Los últimos dos capítulos del libro contienen la mayoría de los conceptos incluidos en estas notas. Se puede encontrar archivo en formato pdf en internet de las primeras ediciones. La quinta edición apareció en 2016.

Capítulo 2

Factorización QR

Uno de los métodos más exitosos para resolver problemas de mínimos cuadrados, en cómputo científico, es el método de factorización QR . Este es un método moderno, popularizado desde 1960 por Gene Golub (ver referencias). Actualmente se considera que este método representa una de las ideas algorítmicas mas importante en el álgebra lineal numérica. El método también es ampliamente utilizado para calcular computacionalmente los valores y vectores propios de una matriz. Asimismo, en el presente juega un papel importante en diversos campos de conocimiento como ciencia de datos, estadística y '*machine learning*' (aprendizaje de máquina, redes neuronales). En este capítulo presentaremos de manera detallada los aspectos matemáticos y algorítmicos que dan origen a esta factorización y mostraremos su aplicación a la solución de problemas de mínimos cuadrados.

2.1

Factorización reducida y completa

Factorización reducida

Considérese la matriz $A \in \mathbb{R}^{m \times n}$ con $m \geq n$ y sean $\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n \in \mathbb{R}^m$ sus vectores columna. Escribimos

$$A = \begin{bmatrix} | & | & \cdots & | \\ \mathbf{a}_1 & \mathbf{a}_2 & \cdots & \mathbf{a}_n \\ | & | & \cdots & | \end{bmatrix}.$$

Si la matriz A es de rango completo n , entonces los vectores columna de A son linealmente independientes. Por lo que podemos aplicar el proceso de ortogonalización de Gram–Schmidt para generar n vectores ortonormales $\mathbf{q}_1, \mathbf{q}_2, \dots, \mathbf{q}_n$ en $\mathbb{R}^m$. Por lo tanto, el espacio columna de A será igual al espacio columna de la siguiente matriz de $m \times n$

$$\hat{Q} = \begin{bmatrix} | & | & \cdots & | \\ \mathbf{q}_1 & \mathbf{q}_2 & \cdots & \mathbf{q}_n \\ | & | & \cdots & | \end{bmatrix},$$

con $\|\mathbf{q}_i\|_2 = 1$ y $\mathbf{q}_i^T \mathbf{q}_j = \delta_{ij}$. Como hemos visto en el anterior capítulo, estos vectores son

$$\begin{aligned}\mathbf{q}_1 &= \frac{\mathbf{v}_1}{\|\mathbf{v}_1\|_2} \quad \text{con} \quad \mathbf{v}_1 = \mathbf{a}_1, \\ \mathbf{q}_2 &= \frac{\mathbf{v}_2}{\|\mathbf{v}_2\|_2} \quad \text{con} \quad \mathbf{v}_2 = \mathbf{a}_2 - (\mathbf{q}_1^T \mathbf{a}_2) \mathbf{q}_1, \\ \mathbf{q}_3 &= \frac{\mathbf{v}_3}{\|\mathbf{v}_3\|_2} \quad \text{con} \quad \mathbf{v}_3 = \mathbf{a}_3 - (\mathbf{q}_1^T \mathbf{a}_3) \mathbf{q}_1 - (\mathbf{q}_2^T \mathbf{a}_3) \mathbf{q}_2.\end{aligned}$$

En general, suponiendo conocidos $\mathbf{q}_1, \mathbf{q}_2, \dots, \mathbf{q}_{j-1}$, el siguiente vector $\mathbf{q}_j$, ortonormal a los anteriores, está dado por

$$\mathbf{q}_j = \frac{\mathbf{v}_j}{\|\mathbf{v}_j\|_2} \quad \text{con} \quad \mathbf{v}_j = \mathbf{a}_j - (\mathbf{q}_1^T \mathbf{a}_j) \mathbf{q}_1 - \dots - (\mathbf{q}_{j-1}^T \mathbf{a}_j) \mathbf{q}_{j-1} = \mathbf{a}_j - \sum_{i=1}^{j-1} (\mathbf{q}_i^T \mathbf{a}_j) \mathbf{q}_i.$$

Si definimos los escalares $r_{ij} \equiv \mathbf{q}_i^T \mathbf{a}_j$ para $i > j$ y $r_{jj} = \|\mathbf{v}_j\|_2$, entonces

$$\begin{aligned}\mathbf{q}_1 &= \frac{\mathbf{a}_1}{r_{11}}, \\ \mathbf{q}_2 &= \frac{\mathbf{a}_2 - r_{12} \mathbf{q}_1}{r_{22}}, \\ &\vdots \\ \mathbf{q}_n &= \frac{\mathbf{a}_n - \sum_{i=1}^{n-1} r_{in} \mathbf{q}_i}{r_{nn}}.\end{aligned} \quad \Rightarrow \quad \begin{aligned}\mathbf{a}_1 &= r_{11} \mathbf{q}_1, \\ \mathbf{a}_2 &= r_{12} \mathbf{q}_1 + r_{22} \mathbf{q}_2, \\ &\vdots \\ \mathbf{a}_n &= r_{1n} \mathbf{q}_1 + r_{2n} \mathbf{q}_2 + \dots + r_{nn} \mathbf{q}_n.\end{aligned}$$

Este último conjunto de ecuaciones tienen la siguiente representación matricial

$$\begin{aligned}A &= \underbrace{\begin{bmatrix} | & | & & | \\ \mathbf{a}_1 & \mathbf{a}_2 & \dots & \mathbf{a}_n \\ | & | & & | \end{bmatrix}}_{\text{matriz } m \times n} \\ &= \underbrace{\begin{bmatrix} | & | & & | \\ \mathbf{q}_1 & \mathbf{q}_2 & \dots & \mathbf{q}_n \\ | & | & & | \end{bmatrix}}_{\text{matriz } m \times n} \underbrace{\begin{bmatrix} r_{11} & r_{12} & \dots & r_{1n} \\ 0 & r_{22} & \dots & r_{2n} \\ \vdots & & \ddots & \vdots \\ 0 & 0 & \dots & r_{nn} \end{bmatrix}}_{\text{matriz } n \times n} \\ &= \hat{Q} \hat{R}\end{aligned}$$

Por lo tanto, la factorización $\hat{Q}\hat{R}$ se puede construir con el siguiente algoritmo:

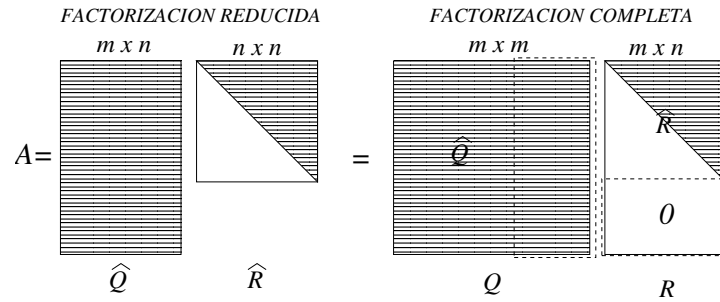
Algoritmo de Gram-Schmidt clásico

- Para $j = 1, \dots, n$
 - $\mathbf{v}_j = \mathbf{a}_j$
 - Para $i = 1, \dots, j-1$
 - $r_{ij} = \mathbf{q}_i^T \mathbf{a}_j$
 - $\mathbf{v}_j = \mathbf{v}_j - r_{ij} \mathbf{q}_i$
 - Fin
 - $r_{jj} = \|\mathbf{v}_j\|_2$
 - $\mathbf{q}_j = \mathbf{v}_j / r_{jj}$
- Fin

TEOREMA 2.1. Si $A \in \mathbb{R}^{m \times n}$ con $m \geq n$ es una matriz de rango completo, existe una única factorización reducida QR , dada por $A = \hat{Q}\hat{R}$ con $r_{ii} > 0$, $i = 1, \dots, n$ (Ver Trefethen–Bau).

Factorización completa

Una factorización completa QR de $A \in \mathbb{R}^{m \times n}$ ($m \geq n$) va más allá agregando $m - n$ columnas ortonormales a $\hat{Q}$, y agregando $m - n$ renglones de ceros a $\hat{R}$ de manera tal que obtenemos una matriz ortogonal $Q \in \mathbb{R}^{m \times m}$ y otra matriz $R \in \mathbb{R}^{m \times n}$ triangular superior. Esquemáticamente



En la factorización completa las columnas adicionales $\mathbf{q}_j$, $j = n + 1, \dots, m$, son ortogonales al espacio columna de A . Por supuesto, la matriz Q es una matriz ortogonal, pues $Q^T Q = I$.

TEOREMA 2.2. Cualquier matriz $A \in \mathbb{R}^{m \times n}$ ($m \geq n$) tiene una factorización completa QR , dada por $A = QR$ con $Q \in \mathbb{R}^{m \times m}$ matriz ortogonal y $R \in \mathbb{R}^{m \times n}$ matriz triangular superior (ver Trefethen–Bau).

2.2

Aplicación a problemas de mínimos cuadrados

Habiendo obtenido una factorización completa QR de $A \in \mathbb{R}^{m \times n}$, el problema sobredeterminado $A\mathbf{x} = \mathbf{b}$, que aparece en el problema de ajuste lineal por medio de mínimos cuadrados, con $\mathbf{b} \in \mathbb{R}^m$ se puede expresar en la forma $QR\mathbf{x} = \mathbf{b}$, que a su vez es equivalente al sistema triangular superior (después de multiplicar ambos lados por $Q^{-1} = Q^T$):

$$R\mathbf{x} = Q^T \mathbf{b}. \quad (2.1)$$

Este sistema se puede escribir en forma detallada como

$$\begin{bmatrix} r_{11} & r_{12} & \cdots & r_{1n} \\ 0 & r_{22} & \cdots & r_{2n} \\ \vdots & & \ddots & \vdots \\ 0 & 0 & \cdots & r_{nn} \\ 0 & 0 & \cdots & 0 \\ \vdots & & \ddots & \vdots \\ 0 & 0 & \cdots & 0 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix} = \begin{bmatrix} f_1 \\ f_2 \\ \vdots \\ f_n \\ f_{n+1} \\ \vdots \\ f_m \end{bmatrix}$$

con $f_i = (Q^T \mathbf{b})_i$, $i = 1, \dots, m$.

El sistema triangular superior (2.1) se puede resolver utilizando sustitución regresiva, sin perder de vista que las últimas $m - n$ componentes de $\mathbf{f} = Q^T \mathbf{b}$ no intervienen en la solución. Es decir, la solución del problema de mínimos cuadrados es solución del sistema reducido

$$\hat{R}\mathbf{x} = \hat{Q}^T \mathbf{b} \Rightarrow \begin{bmatrix} r_{11} & r_{12} & \dots & r_{1n} \\ 0 & r_{22} & \dots & r_{2n} \\ \vdots & & \ddots & \vdots \\ 0 & 0 & \dots & r_{nn} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix} = \begin{bmatrix} f_1 \\ f_2 \\ \vdots \\ f_n \end{bmatrix}. \quad (2.2)$$

Las últimas $m - n$ componentes de $\mathbf{f} = Q^T \mathbf{b}$ están relacionadas con el residual $\mathbf{r} = \mathbf{b} - A\mathbf{x}$, pues

$$Q^T \mathbf{r} = Q^T \mathbf{b} - R\mathbf{x} = \begin{bmatrix} 0 \\ \vdots \\ 0 \\ f_{n+1} \\ \vdots \\ f_m \end{bmatrix} \Rightarrow \|\mathbf{r}\|_2 = \|Q^T \mathbf{r}\|_2 = \left\| \begin{bmatrix} f_{n+1} \\ \vdots \\ f_m \end{bmatrix} \right\|_2.$$

Es fácil verificar que la solución $\hat{\mathbf{x}}$ del sistema (2.2) resuelve el problema de mínimos cuadrados, pues

$$\hat{R}\mathbf{x} = \hat{Q}^T \mathbf{b} \iff \hat{Q}\hat{R}\mathbf{x} = \hat{Q}\hat{Q}^T \mathbf{b} \iff A\mathbf{x} = P_{\hat{Q}} \mathbf{b},$$

donde $P_{\hat{Q}} = \hat{Q}\hat{Q}^T$ proyecta vectores de $\mathbb{R}^m$ en el espacio generado por los vectores columna de $\hat{Q}$, los cuales a su vez también generan el espacio columna de A . Es decir $P_{\hat{Q}} = P_A = A(A^T A)^{-1} A^T$ y

$$A\mathbf{x} = P_{\hat{Q}} \mathbf{b} \iff A\mathbf{x} = A(A^T A)^{-1} A^T \mathbf{b} \iff A^T A\mathbf{x} = A^T \mathbf{b}.$$

Concluimos que $\hat{\mathbf{x}}$ resuelve (2.2) si y solo si resuelve las ecuaciones normales. Por lo tanto, es la solución del problema de mínimos cuadrados.

Advertencia. Las fórmulas de Gram-Schmidt no se aplican en cómputo científico para encontrar la factorización QR , ya que la sucesión de operaciones resulta numéricamente inestable (sensible a errores de redondeo). Esta inestabilidad se produce debido a las sustracciones presentes en el algoritmo y a que los vectores $\mathbf{q}_j$ no son estrictamente ortogonales debido a errores de redondeo. Se pueden utilizar métodos de estabilización, cambiando el orden en que se realizan las operaciones. Afortunadamente, hay un método más efectivo, estable por supuesto, para encontrar la factorización QR . El nuevo método hace uso de las propiedades de las *proyecciones ortogonales* y, en particular de las matrices de *reflexión*, como se muestra en las siguientes secciones.

2.3

Reflexiones de Householder

Ya hemos mencionado que el proceso de ortogonalización de Gram-Schmidt no es un método estable cuando se utiliza aritmética de punto flotante en cómputo científico. Por tal motivo, no se

utiliza como método general para realizar la factorización QR . Afortunadamente es posible construir la factorización QR por medio de la aplicación sucesiva de transformaciones que involucran matrices ortogonales. A continuación explicamos la idea principal y posteriormente formalizaremos la construcción de la factorización mediante matrices de reflexión, ya introducidas previamente en la última sección del capítulo 1.

Idea principal

La **triangularización** es el proceso de construcción que permite transformar una matriz A en una matriz triangular superior R , mediante la aplicación sucesiva de transformaciones matriciales, semejante al proceso de eliminación de Gauss. La diferencia es que en este caso se multiplica por **matrices ortogonales** Q_k , de tal manera que al final del proceso

$$Q_n \cdots Q_2 Q_1 A = R,$$

resulta en una matriz triangular superior R . Cada matriz Q_k se escoge para introducir ceros debajo de la diagonal en la k -ésima columna. Por ejemplo, para una matriz A de $m \times n = 5 \times 3$, las operaciones Q_k se aplican como se muestra a continuación:

$$\underbrace{\begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \\ a_{41} & a_{42} & a_{43} \\ a_{51} & a_{52} & a_{53} \end{bmatrix}}_{A^{(0)} \quad A} \rightarrow \underbrace{\begin{bmatrix} a_{11}^{(1)} & a_{12}^{(1)} & a_{13}^{(1)} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & a_{32}^{(1)} & a_{33}^{(1)} \\ 0 & a_{42}^{(1)} & a_{43}^{(1)} \\ 0 & a_{52}^{(1)} & a_{53}^{(1)} \end{bmatrix}}_{Q_1 A^{(1)}} \rightarrow \underbrace{\begin{bmatrix} a_{11}^{(1)} & a_{12}^{(1)} & a_{13}^{(1)} \\ 0 & a_{22}^{(2)} & a_{23}^{(2)} \\ 0 & 0 & a_{33}^{(2)} \\ 0 & 0 & a_{43}^{(2)} \\ 0 & 0 & a_{53}^{(2)} \end{bmatrix}}_{Q_2 Q_1 A^{(2)}} \rightarrow \underbrace{\begin{bmatrix} a_{11}^{(1)} & a_{12}^{(1)} & a_{13}^{(1)} \\ 0 & a_{22}^{(2)} & a_{23}^{(2)} \\ 0 & 0 & a_{33}^{(3)} \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix}}_{Q_3 Q_2 Q_1 A^{(3)} \quad R},$$

en donde la última matriz es la matriz triangular superior R . A este proceso se le denomina *triangularización de Householder* (Alston Householder, 1958), y es el método más usado actualmente para encontrar la factorización QR debido a que es numéricamente más estable que el método de Gram-Schmidt. El paso clave es la construcción de las matrices Q_k . Hay dos métodos principales utilizados para construir estas matrices: el *método de rotaciones de Givens* y el *método de reflexiones de Householder*. En este curso solo consideraremos las reflexiones de Householder.

Reflexiones de Householder

Una **reflexión de Householder**, denotada por H , se construye a partir de un vector fijo dado $\mathbf{x}$ en un espacio euclídeo $\mathbb{R}^p$ con $p \geq 2$, buscando que su vector reflejado sea $H\mathbf{x} = \|\mathbf{x}\| \mathbf{e}_1 = (\|\mathbf{x}\|, 0, \dots, 0)^T$, como se muestra en la figura 2.1. Esta reflexión refleja sobre un hiperplano $H^\perp$ con vector normal $\mathbf{v}$ y está dada por

$$H = I - 2 \frac{\mathbf{v} \mathbf{v}^T}{\mathbf{v}^T \mathbf{v}} \quad \text{con} \quad \mathbf{v} = \|\mathbf{x}\| \mathbf{e}_1 - \mathbf{x}.$$

Hacemos énfasis que, dado el vector $\mathbf{x}$, la proyección H realiza la siguiente transformación

$$\mathbf{x} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_p \end{bmatrix} \mapsto H\mathbf{x} = \begin{bmatrix} \|\mathbf{x}\| \\ 0 \\ \vdots \\ 0 \end{bmatrix} = \|\mathbf{x}\| \mathbf{e}_1.$$

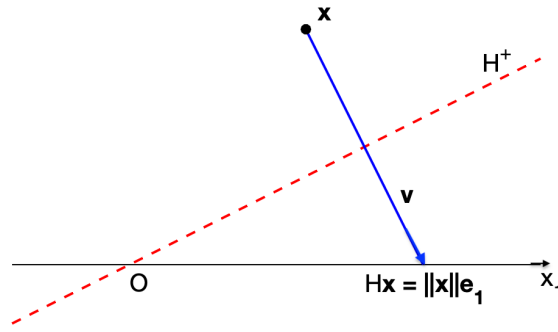


Figura 2.1: Reflexión de Householder con vector de reflexión $\mathbf{v} = \|\mathbf{x}\| \mathbf{e}_1 - \mathbf{x}$.

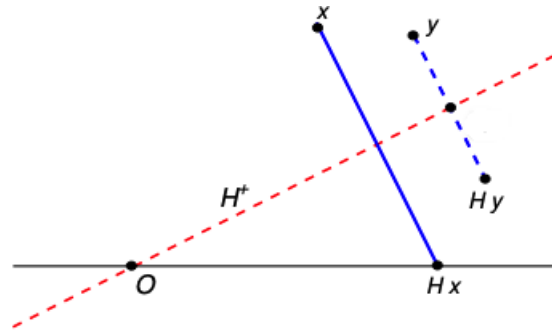


Figura 2.2: La reflexión de un vector $\mathbf{y}$ con la reflexión de Householder $H = I - 2\mathbf{v}\mathbf{v}^T/\mathbf{v}^T\mathbf{v}$.

Cuando la reflexión se aplica a cada punto $\mathbf{y}$ en algún lado del plano H^+ , el resultado es la imagen reflejada en el otro lado de H^+ , como se indica en la figura 2.2

EJEMPLO 2.3. Dado el vector $\mathbf{x} = (3, 4, 0)^T$, encontrar la reflexión de Householder H tal que $H\mathbf{x} = \|\mathbf{x}\| \mathbf{e}_1 = (5, 0, 0)^T$.

Solución.

$$\mathbf{v} = \|\mathbf{x}\| \mathbf{e}_1 - \mathbf{x} = \begin{bmatrix} 5 \\ 0 \\ 0 \end{bmatrix} - \begin{bmatrix} 3 \\ 4 \\ 0 \end{bmatrix} = \begin{bmatrix} 2 \\ -4 \\ 0 \end{bmatrix},$$

$$\mathbf{v}\mathbf{v}^T = \begin{bmatrix} 2 \\ -4 \\ 0 \end{bmatrix} \begin{bmatrix} 2 & -4 & 0 \end{bmatrix} = \begin{bmatrix} 4 & -8 & 0 \\ -8 & 16 & 0 \\ 0 & 0 & 0 \end{bmatrix} \quad \text{y} \quad \mathbf{v}^T\mathbf{v} = 20.$$

Luego

$$H = I - 2\frac{\mathbf{v}\mathbf{v}^T}{\mathbf{v}^T\mathbf{v}} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} - \frac{1}{10} \begin{bmatrix} 4 & -8 & 0 \\ -8 & 16 & 0 \\ 0 & 0 & 0 \end{bmatrix} = \begin{bmatrix} 3/5 & 4/5 & 0 \\ 4/5 & -3/5 & 0 \\ 0 & 0 & 1 \end{bmatrix}.$$

Se puede verificar que H efectivamente satisface $H\mathbf{x} = \|\mathbf{x}\| \mathbf{e}_1 = (5, 0, 0)^T$. En la figura 2.3 se ilustra el resultado de este ejemplo. El plano H^+ tiene ecuación $v_1x_1 + v_2x_2 + v_3x_3 = 0$, es decir, $2x_1 - 4x_2 = 0$ con x_3 arbitrario. Obsérvese además que

$$H^TH = HH^T = \begin{bmatrix} 3/5 & 4/5 & 0 \\ 4/5 & -3/5 & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 3/5 & 4/5 & 0 \\ 4/5 & -3/5 & 0 \\ 0 & 0 & 1 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} = I.$$

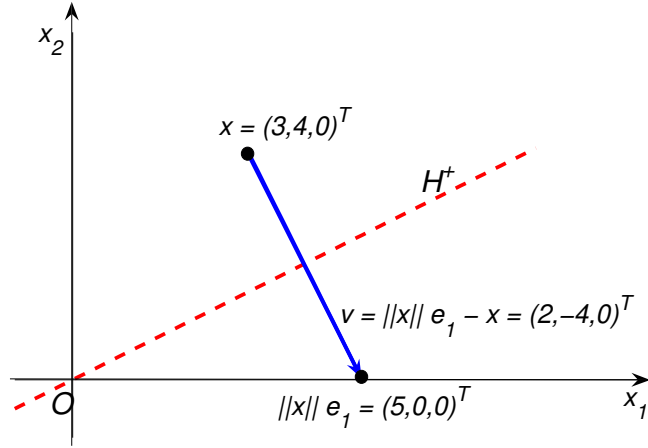


Figura 2.3: Reflexión de Householder del ejemplo 2.3.

En general, se puede verificar fácilmente que

$$H^T H = H H^T = \left[I - 2 \frac{\mathbf{v} \mathbf{v}^T}{\mathbf{v}^T \mathbf{v}} \right] \left[I - 2 \frac{\mathbf{v} \mathbf{v}^T}{\mathbf{v}^T \mathbf{v}} \right]^T = I.$$

2.4

Factorización QR mediante reflexiones matriciales

Volviendo a la triangularización de Householder al inicio de la sección, los pasos para triangularizar la matriz $A \in \mathbb{R}^{m \times n}$ son los siguientes:

- **Paso 1.** La primera matriz ortogonal es $Q_1 = H$, en donde la matriz de reflexión H de orden m se construye con la primera columna de $A^{(0)} = A$, es decir $\mathbf{x} = \mathbf{a}_1$. Entonces $A^{(1)} = Q_1 A^{(0)}$.
- **Paso 2.** La matriz ortogonal Q_2 se construye como

$$Q_2 = \begin{bmatrix} 1 & \mathbf{0}^T \\ \mathbf{0} & H \end{bmatrix},$$

con H la matriz de reflexión de orden $(m-1)$. En esta ocasión el vector $\mathbf{x} \in \mathbb{R}^{m-1}$ para construir la matriz H , se extrae de la primera columna de la submatriz $A^{(1)}(2:m, 2:n)$ de $A^{(1)}$. El vector $\mathbf{0}$ es el vector columna nulo de tamaño $m-1$. Entonces $A^{(2)} = Q_2 A^{(1)}$.

- **Paso k -ésimo.** En general, cada $Q_k \in \mathbb{R}^{m \times m}$ se escoge como una matriz ortogonal de la forma

$$Q_k = \begin{bmatrix} I_{k-1} & \mathbb{O}^T \\ \mathbb{O} & H \end{bmatrix},$$

en donde I_{k-1} es la matriz identidad de orden $(k-1)$, H es la matriz de reflexión de orden $(m-k+1)$ basada en el vector $\mathbf{x} \in \mathbb{R}^{m-k+1}$, que se toma de la primera columna de la submatriz $A^{(k-1)}(k:m, k:n)$. La submatriz rectangular nula $\mathbb{O}$ es de tamaño $(m-k+1) \times (k-1)$ y $\mathbb{O}^T$ su transpuesta.

El k -ésimo paso con detalle

La siguiente sucesión de matrices produce la matriz triangular superior R

$$A^{(0)} = A \rightarrow A^{(1)} = Q_1 A^{(0)} \quad \dots \quad A^{(k)} = Q_k A^{(k-1)} \quad \dots \quad A^{(n)} = Q_n A^{(n-1)} = R,$$

en donde el k -ésimo paso en forma detallada (quitando superíndices a los coeficientes a_{ij}), es

$$A^{(k)} = Q_k A^{(k-1)} = \begin{bmatrix} I_{k-1} & \mathbb{O}^T \\ \mathbb{O} & H \end{bmatrix}_{m \times m} \begin{bmatrix} a_{11} & \dots & a_{1k-1} & a_{1k} & \dots & a_{1n} \\ \vdots & \ddots & \vdots & \vdots & \ddots & \vdots \\ 0 & \dots & a_{k-1k-1} & a_{k-1k} & \dots & a_{k-1n} \\ 0 & \dots & 0 & a_{kk} & \dots & a_{kn} \\ \vdots & \ddots & \vdots & \vdots & \ddots & \vdots \\ 0 & \dots & 0 & a_{mk} & \dots & a_{mn} \end{bmatrix}_{m \times n}$$

La matriz $A^{(k)} \in \mathbb{R}^{m \times n}$ mantiene los mismos bloques que $A^{(k-1)}$, excepto el bloque inferior derecho que debe cambiar por

$$H \begin{bmatrix} a_{kk} & \dots & a_{kn} \\ \vdots & \ddots & \vdots \\ a_{mk} & \dots & a_{mn} \end{bmatrix} = \begin{bmatrix} a'_{kk} & a'_{kk+1} & \dots & a'_{kn} \\ 0 & a'_{k+1k+1} & \dots & a'_{k+1n} \\ \vdots & \vdots & \ddots & \vdots \\ 0 & a'_{mk+1} & \dots & a'_{mn} \end{bmatrix}.$$

Por supuesto, aquí el paso fundamental es la multiplicación de H por la primera columna de la submatriz de la derecha:

$$H \begin{bmatrix} a_{kk} \\ a_{k+1k} \\ \vdots \\ a_{mk} \end{bmatrix} = \begin{bmatrix} a'_{kk} \\ 0 \\ \vdots \\ 0 \end{bmatrix},$$

donde $a'_{kk} = \|\mathbf{x}_k\|$, con $\mathbf{x}_k = (a_{kk}, \dots, a_{mk})^T$.

Las matrices Q_k en el proceso de triangularización de Householder son ortogonales, pues

$$Q_k Q_k^T = \begin{bmatrix} I_{k-1} & \mathbb{O}^T \\ \mathbb{O} & H \end{bmatrix} \begin{bmatrix} I_{k-1} & \mathbb{O}^T \\ \mathbb{O} & H^T \end{bmatrix} = \begin{bmatrix} I_{k-1} & \mathbb{O}^T \\ \mathbb{O} & HH^T \end{bmatrix} = \begin{bmatrix} I_{k-1} & \mathbb{O}^T \\ \mathbb{O} & I_{m-k+1} \end{bmatrix} = I_m.$$

La mejor de las dos reflexiones

En la exposición anterior se ha simplificado el procedimiento. La reflexión del vector $\mathbf{x}$ sobre el eje x_1 puede dar como resultado $\|x\|e_1$ o bien $-\|x\|e_1$, dando lugar a dos reflexiones, una a través del hiperplano H^+ , y otra a través del hiperplano H^- , como se muestra en la Figura 2.4. Matemáticamente, cualquiera de las dos elecciones es satisfactoria. Sin embargo, para asegurar estabilidad en el algoritmo numérico se debe escoger la reflexión en la que el punto reflejado esté más alejado de $\mathbf{x}$. Para tal fin escogemos $\mathbf{v} = -\text{signo}(x_1)\|\mathbf{x}\|\mathbf{e}_1 - \mathbf{x}$, donde x_1 es la primera coordenada de $\mathbf{x}$ y

$$\text{signo}(x_1) = \begin{cases} 1 & \text{si } x_1 \geq 0, \\ -1 & \text{si } x_1 < 0. \end{cases}$$

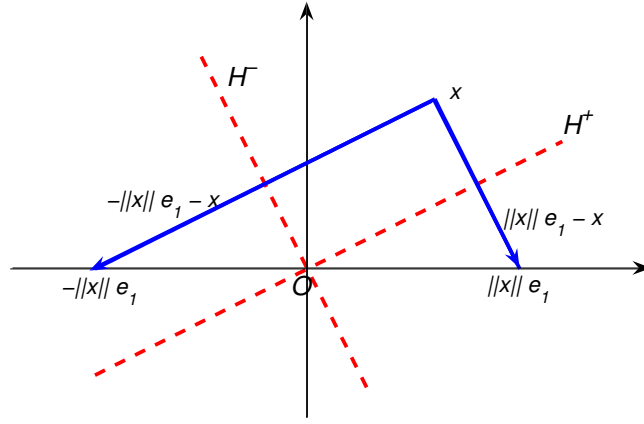
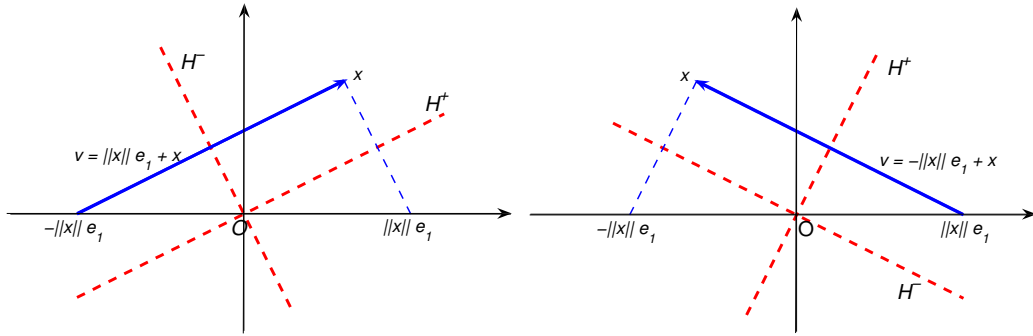


Figura 2.4: Las dos reflexiones de Householder.

La reflexión no depende del signo del vector $\mathbf{v}$ y también es posible escoger $\mathbf{v} = \text{signo}(x_1) \|\mathbf{x}\| \mathbf{e}_1 + \mathbf{x}$. La figura 2.5 muestra ambos casos. Es evidente que la elección $\mathbf{v} = \text{signo}(x_1) \|\mathbf{x}\| \mathbf{e}_1 + \mathbf{x}$ permite que $\|\mathbf{v}\|$ nunca sea más pequeño que $\|\mathbf{x}\|$, evitando cancelación por sustracción al dividir por $\mathbf{v}^T \mathbf{v}$ en la construcción de la matriz de proyección H , con lo cual aseguramos estabilidad en el método.

Figura 2.5: Izquierda: $x_1 > 0 \Rightarrow \text{signo}(x_1) = +1$. Derecha: $x_1 < 0 \Rightarrow \text{signo}(x_1) = -1$.

Algoritmo de factorización de Householder

Utilizando las ideas anteriores podemos construir el algoritmo que factoriza A en la forma $A = QR$, es decir $Q^T A = R$, con $Q = (Q_n \cdots Q_2 Q_1)^T$ matriz ortogonal y R matriz triangular superior. Sin embargo, en el algoritmo que escribimos a continuación solamente calculamos el factor R y lo almacenamos en el mismo espacio de memoria que ocupa A con el objeto de ahorrar memoria. También, si se desea se pueden almacenar los n vectores de reflexión $\mathbf{v}_1, \dots, \mathbf{v}_n$, para posibles usos posteriores.

Algoritmo de factorización

Para $k = 1, \dots, n$
 . $\mathbf{x} = A(k : m, k)$
 . $\mathbf{v}_k = \text{signo}(x_1) \|\mathbf{x}\| \mathbf{e}_1 + \mathbf{x}$
 . $\mathbf{v}_k = \mathbf{v}_k / \|\mathbf{v}_k\|_2$
 . $A(k : m, k : n) = A(k : m, k : n) - 2\mathbf{v}_k (\mathbf{v}_k^T A(k : m, k : n))$
 Fin

Obsérvese que en el algoritmo se han realizado dos simplificaciones:

1. Se ha denotado por $\mathbf{v}_k$ a $\mathbf{v}_k / \|\mathbf{v}_k\|_2$.
2. Al multiplicar o aplicar la reflexión por un vector $\mathbf{y}$ obtenemos

$$(I - 2\mathbf{v}_k\mathbf{v}_k^T)\mathbf{y} = \mathbf{y} - 2\mathbf{v}_k(\mathbf{v}_k^T\mathbf{y}).$$

de tal manera que cuando $\mathbf{y}$ es una de las columnas de la submatriz $A(k:m, k:n)$, entonces la forma económica de aplicar la reflexión a cada una de estas columnas se indica en el último renglón del algoritmo. La matriz ortogonal $Q^T = Q_n \cdots Q_2 Q_1$ no se construye explícitamente, puesto que tomaría trabajo adicional. Afortunadamente en muchas aplicaciones no es necesario obtener esta matriz. Por ejemplo, si se quiere resolver el sistema sobredeterminado $A\mathbf{x} = \mathbf{b}$ y sabemos que $Q^T A = R$, entonces $R\mathbf{x} = Q^T A\mathbf{x} = Q^T \mathbf{b}$ y basta con resolver el sistema $R\mathbf{x} = Q^T \mathbf{b}$ utilizando sustitución regresiva. El lado derecho $Q^T \mathbf{b}$ puede calcularse aplicando a $\mathbf{b}$ la misma sucesión de operaciones aplicadas a la matriz A :

Para $k = 1, \dots, n$
 $\vdots$ $\mathbf{b}(k:m) = \mathbf{b}(k:m) - 2\mathbf{v}_k(\mathbf{v}_k^T \mathbf{b}(k:m))$
 Fin

Hemos supuesto que los vectores de reflexión $\mathbf{v}_1, \dots, \mathbf{v}_n$ se almacenan al encontrar R en el algoritmo anterior. Sin embargo, en el problema de la solución del problema sobredeterminado $A\mathbf{x} = \mathbf{b}$ no es necesario almacenar los vectores de reflexión $\mathbf{v}_k$. Además, el cálculo de $Q^T \mathbf{b}$ puede realizarse simultáneamente junto con el cálculo de R . En este caso el algoritmo completo, incluyendo sustitución regresiva, es como se muestra a continuación:

Algoritmo de triangularización de Householder para resolver $A\mathbf{x} = \mathbf{b}$

Para $k = 1, \dots, n$
 $\cdot$ $\mathbf{x} = A(k:m, k)$
 $\cdot$ $\mathbf{v} = \text{signo}(x_1)\|\mathbf{x}\| \mathbf{e}_1 + \mathbf{x}$
 $\cdot$ $\mathbf{v} = \mathbf{v} / \|\mathbf{v}\|_2$
 $\cdot$ $A(k:m, k:n) = A(k:m, k:n) - 2\mathbf{v}(\mathbf{v}^T A(k:m, k:n))$
 $\cdot$ $\mathbf{b}(k:m) = \mathbf{b}(k:m) - 2\mathbf{v}(\mathbf{v}^T \mathbf{b}(k:m))$
 Fin
 $\mathbf{x}(n) = \mathbf{b}(n) / A(n, n)$
 Para $k = n-1 : -1 : 1$
 $\vdots$ $\mathbf{x}(k) = (\mathbf{b}(k) - A(k, k+1:n) \cdot \mathbf{b}(k+1:n)) / A(k, k)$
 Fin

EJEMPLO 2.4 El *NIST* (National Institute of Standards and Technology) es una rama del Departamento de Comercio de los EU, que se responsabiliza de establecer estándares nacionales e internacionales. El *NIST* mantiene conjuntos de datos de referencia, para su uso en la calibración y certificación de software en estadística. En su página Web

<http://www.itl.nist.gov/div898/strd>

en la liga *Dataset Archives*, y después bajo *Linear Regression*, se encuentra el conjunto de datos *Filip*, que consiste de 82 observaciones de una variable y para diferentes valores de x . El problema

es modelar este conjunto de datos por medio de un polinomio de grado 10. Este conjunto de datos es controversial: algunos paquetes pueden reproducir valores cercanos a los coeficientes que el *NIST* ha declarado como los **valores certificados**; otros paquetes dan advertencias o mensajes de error de que el problema es mal condicionado.

Solución. Se resuelve el problema utilizando el método de ecuaciones normales y el método de factorización QR , con el objeto de comparar el desempeño de cada uno de ellos. Tenemos $m = 82$ datos (t_i, y_i) y debemos encontrar $n = 11$ coeficientes c_j para el polinomio de ajuste de grado 10. La matriz de diseño A tiene coeficientes $a_{ij} = t_i^{j-1}$, $i = 1, \dots, 82$, $j = 1, \dots, 11$. Para dar una idea de la complejidad de esta matriz obsérvese que el coeficiente con mínimo valor absoluto es 1 y el coeficiente con mayor valor absoluto es un poco mayor a 2.7×10^9 . Esta matriz es *mal condicionada*, siendo su número de condición $\kappa(A) \approx \mathcal{O}(10^{15})$. La matriz de ecuaciones normales $A^T A$ es todavía más mal comportada, desde el punto de vista numérico, pues tiene un coeficiente mínimo en valor absoluto de 82 y uno máximo de 5.1×10^{19} , su número de condición es proporcional a 10^{30} . En los cálculos se utilizó aritmética de punto flotante de doble precisión.

La Figura 2.6 muestra la gráfica de los datos junto con la gráfica de los polinomios de ajuste obtenidos por medio de los dos métodos. En esta figura no se aprecia la diferencia entre la gráfica del polinomio obtenido por medio del método QR y el certificado por el *NIST*. La gráfica del polinomio de ajuste obtenido por medio de ecuaciones normales difiere de los dos anteriores, sobre todo en la mitad derecha, intervalo $(-6, -3)$, donde presenta pequeñas oscilaciones alrededor de los datos.

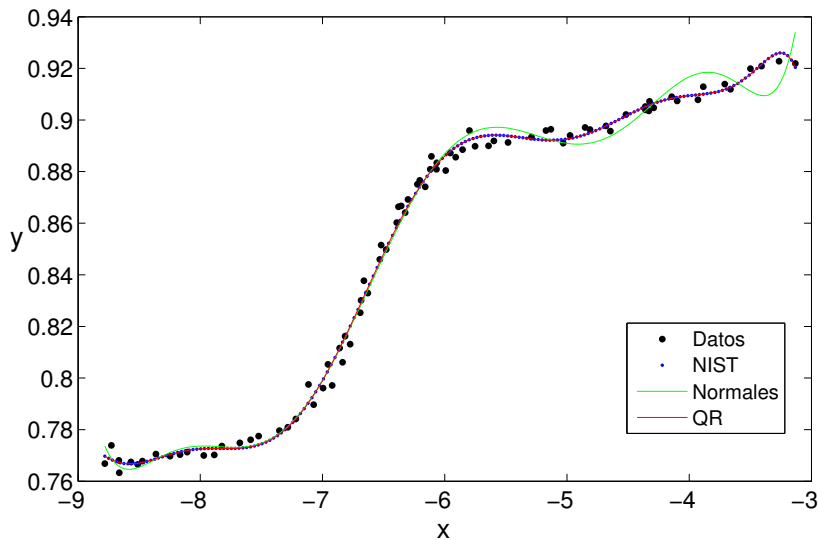


Figura 2.6: Gráfica de los datos *FILIP* y los polinomios de ajuste.

La tabla 2.1 muestra los coeficientes obtenidos con los dos métodos junto con los certificados por el *NIST*. Los resultados obtenidos con el método QR son muy buenos, tomando en cuenta que los resultados certificados por el *NIST* fueron obtenidos con cálculos de muy alta precisión, a saber precisión múltiple de 500 dígitos (lo cual representa una idealización de lo que se alcanzaría si los cálculos se hicieran sin error de redondeo). Por el contrario, los coeficientes obtenidos por medio del método de ecuaciones normales son totalmente diferentes a los certificados, incluso en el signo.

Finalmente, en la tabla 2.2 se muestra la diferencia relativa del vector de coeficientes obtenido

Coef.	NIST ($\times 10^3$)	QR ($\times 10^3$)	Ec. Normales ($\times 10^2$)
c_1	-1.467489614229800	-1.467489582388863	3.508390286748969
c_2	-2.772179591933420	-2.772179531420028	5.458039227995218
c_3	-2.316371081608930	-2.316371030602281	3.674467707503117
c_4	-1.127973940983720	-1.127973915874232	1.398316266894402
c_5	-0.354478233703349	-0.354478225708109	0.330592755393584
c_6	-0.075124201739376	-0.075124200018508	0.050170726293308
c_7	-0.010875318035534	-0.010875317781920	0.004860547121394
c_8	-0.001062214985889	-0.001062214960612	0.000287297102007
c_9	-0.000067019115459	-0.000067019113828	0.000009255550238
c_{10}	-0.000002467810783	-0.000002467810721	0.000000120446348
c_{11}	-0.000000040296253	-0.000000040296251	0.000000000044876

Tabla 2.1: Comparación de los coeficientes del polinomio de ajuste obtenidos con los certificados por el *NIST*.

con los dos métodos con respecto al valor certificado del *NIST*. Esta cantidad se define como:

$$\frac{\|\mathbf{c}_{nist} - \mathbf{c}\|_2}{\|\mathbf{c}_{nist}\|_2} \times 100 \quad (\mathbf{c} = \mathbf{x} \text{ denota el vector de coeficientes de la tabla}).$$

Se observa que la diferencia relativa de la solución obtenida por medio del método de factorización *QR* es muy pequeña, a diferencia de la obtenida por medio del método de ecuaciones normales. En esta última tabla también se incluye la magnitud del residual $\|\mathbf{b} - \mathbf{Ax}\|_2$ en cada caso, en donde $\mathbf{b} = \{y_i\}_{i=1}^m$ es el vector de observaciones y $\mathbf{x} = \{c_j\}_{j=1}^n$ el vector de coeficientes encontrados. Obsérvese que el residual obtenido por medio de *QR* coincide con el del *NIST* hasta la onceava cifra, mientras que el obtenido por medio del método de ecuaciones normales es casi el doble que el certificado. En conclusión, los resultados obtenidos en este problema muestran la superioridad del método de factorización *QR* sobre el método de ecuaciones normales para la solución de problemas de mínimos cuadrados cuando el problema es mal condicionado.

Cantidad	NIST	QR	Ec. Normales
Diferencia relativa	0	$2.2 \times 10^{-6} \%$	118 %
Residual	0.028210838148692	0.028210838142565	0.041705784558465

Tabla 2.2: Diferencias relativas y los residuales obtenidos por cada método.

2.5

Referencias

1. Gene H. Golub, Charles F. Van Loan, *Matrix Computations*, The Johns Hopkins University Press, ediciones en 1983, 1989, 1996, 2013.

La primera referencia sobre el tema es sin duda este libro, el cual es considerado por muchos investigadores como 'la biblia' en álgebra lineal numérica, convertido actualmente en un libro clásico.

2. Lloyd N. Trefethen, David Bau III, *Numerical Linear Algebra*, SIAM, 1997.

Un libro de texto que presenta las ideas fundamentales del álgebra lineal numérica en forma elegante pero detallada. Contiene un capítulo dedicado exclusivamente a la factorización QR y mínimos cuadrados, además de otros temas muy utilizados en las aplicaciones. Incluye el análisis de los métodos y algoritmos cuando se utiliza aritmética de punto flotante (aritmética finita en dispositivos de cómputo).

3. Householder, A. S. (1958). *A Class of Methods for Inverting Matrices*. Journal of the Society for Industrial and Applied Mathematics, 6(2), 189-195. doi:10.1137/0106012.
4. Gene H. Golub, Frank Uhlig, *The QR algorithm: 50 years later its genesis by John Francis and Vera Kublanovskaya and subsequent developments*, IMA Journal of Numerical Analysis (2009) 29, 467-485 doi:10.1093/imanum/drp012.

Una reconstrucción de las ideas e influencias que llevaron a la génesis del algoritmo QR .

5. Strang, Gilbert (2019), *Linear Algebra and Learning from Data* (1st edition). Wellesley Cambridge Press. ISBN 978-0-692-19638-0.

Un libro básico para profundizar en el álgebra lineal y sus aplicaciones a ciencia de datos y aprendizaje automático.

Capítulo 3

Descomposición en valores singulares (SVD)

En álgebra lineal la descomposición en valores singulares de una matriz (real o compleja) es una factorización que revela sus componentes intrínsecas. Cualquier matriz real A de $m \times n$ puede factorizarse como

$$A = U \Sigma V^T$$

donde U es una matriz ortogonal de $m \times m$ cuyas columnas son los vectores propios de AA^T , V es una matriz ortogonal de $n \times n$ cuyas columnas son los vectores propios de $A^T A$ y Σ es una matriz diagonal de $m \times n$ con los valores singulares sobre la diagonal, $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_n \geq 0$.

En este capítulo explicaremos dicha descomposición y su utilidad en aplicaciones modernas, incluida la solución de sistemas de ecuaciones lineales y mínimos cuadrados. Por supuesto que sus aplicaciones son bastas, principalmente en optimización, estadística, ciencia de datos e inteligencia artificial, procesamiento de señales e imágenes, manejo de información en internet, entre muchas otras.

3.1

Simetría y ortogonalidad

La simetría es una característica muy común en la naturaleza. De hecho, hay evidencia de que las estructuras simétricas necesitan menos información para codificarse, y por lo tanto es mucho más probable que aparezcan como variación potencial en el proceso de evolución biológica, de acuerdo a la investigación publicada en <https://doi.org/10.1073/pnas.2113883119>. La simetría es una característica igualmente deseable, y muy útil, en las matemáticas y otras disciplinas básicas o aplicadas. En álgebra lineal esta característica viene acompañada frecuentemente con propiedades de ortogonalidad. Ya hemos visto, al final del capítulo 1, que una reflexión sobre un plano (que da origen a simetría axial) se define por una matriz ortogonal $H = I - 2\mathbf{u}\mathbf{u}^T$ ($\mathbf{u}$ vector unitario) que además es simétrica $H^T = H$. Asimismo, las rotaciones generan simetrías y ellas pueden describirse por medio de matrices ortogonales Q . Por lo tanto, las matrices ortogonales están fuertemente vinculadas con alguna propiedad de simetría.

Por otro lado, y en sentido opuesto, las matrices simétricas están relacionadas con las matrices ortogonales por medio de transformaciones de similaridad. El teorema espectral establece con claridad dicha relación.

TEOREMA 3.1. *Cualquier matriz simétrica real A de $n \times n$ tiene n valores propios reales $\lambda_1, \lambda_2, \dots, \lambda_n$*

(no necesariamente distintos) y un conjunto correspondiente de vectores propios unitarios $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_n$ mutuamente ortogonales.

Las relaciones entre valores y vectores propios en el teorema anterior

$$A\mathbf{u}_1 = \lambda_1 \mathbf{u}_1, \quad A\mathbf{u}_2 = \lambda_2 \mathbf{u}_2, \quad \dots, \quad A\mathbf{u}_n = \lambda_n \mathbf{u}_n,$$

se pueden compactar en el producto

$$A \begin{bmatrix} | & | & \cdots & | \\ \mathbf{u}_1 & \mathbf{u}_2 & \cdots & \mathbf{u}_n \\ | & | & \cdots & | \end{bmatrix} = \begin{bmatrix} | & | & \cdots & | \\ \lambda_1 \mathbf{u}_1 & \lambda_2 \mathbf{u}_2 & \cdots & \lambda_n \mathbf{u}_n \\ | & | & \cdots & | \end{bmatrix}.$$

Si denotamos como U a la matriz que tiene como columnas los vectores propios $\mathbf{u}_i$ y D es la matriz diagonal con los valores propios correspondientes λ_i , la anterior igualdad se puede expresar como

$$AU = UD.$$

Por último, como U es una matriz ortogonal se satisface $U^{-1} = U^T$, y entonces

$$A = \underbrace{\begin{bmatrix} | & | & \cdots & | \\ \mathbf{u}_1 & \mathbf{u}_2 & \cdots & \mathbf{u}_n \\ | & | & \cdots & | \end{bmatrix}}_U \underbrace{\begin{bmatrix} \lambda_1 & 0 & \cdots & 0 \\ 0 & \lambda_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \lambda_n \end{bmatrix}}_D \underbrace{\begin{bmatrix} - & \mathbf{u}_1^T & - \\ - & \mathbf{u}_2^T & - \\ \vdots & \vdots & \vdots \\ - & \mathbf{u}_n^T & - \end{bmatrix}}_{U^T} = A^T \quad (3.1)$$

Los vectores propios $\{\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_n\}$ forman una base ortonormal de $\mathbb{R}^n$, entonces cualquier vector $\mathbf{x} \in \mathbb{R}^n$ se puede expresar como

$$\mathbf{x} = c_1 \mathbf{u}_1 + c_2 \mathbf{u}_2 + \cdots + c_n \mathbf{u}_n = (\mathbf{u}_1^T \mathbf{x}) \mathbf{u}_1 + (\mathbf{u}_2^T \mathbf{x}) \mathbf{u}_2 + \cdots + (\mathbf{u}_n^T \mathbf{x}) \mathbf{u}_n,$$

y, en consecuencia

$$\begin{aligned} A\mathbf{x} &= (\mathbf{u}_1^T \mathbf{x}) \lambda_1 \mathbf{u}_1 + (\mathbf{u}_2^T \mathbf{x}) \lambda_2 \mathbf{u}_2 + \cdots + (\mathbf{u}_n^T \mathbf{x}) \lambda_n \mathbf{u}_n \\ &= \lambda_1 \mathbf{u}_1 \mathbf{u}_1^T \mathbf{x} + \lambda_2 \mathbf{u}_2 \mathbf{u}_2^T \mathbf{x} + \cdots + \lambda_n \mathbf{u}_n \mathbf{u}_n^T \mathbf{x} \\ &= (\lambda_1 \mathbf{u}_1 \mathbf{u}_1^T + \lambda_2 \mathbf{u}_2 \mathbf{u}_2^T + \cdots + \lambda_n \mathbf{u}_n \mathbf{u}_n^T) \mathbf{x}, \quad \forall \mathbf{x} \in \mathbb{R}^n. \end{aligned}$$

Por lo tanto, cualquier matriz simétrica A se puede expresar como combinación de proyecciones ortogonales $P_i = \mathbf{u}_i \mathbf{u}_i^T$ de rango uno:

$$A = \lambda_1 \mathbf{u}_1 \mathbf{u}_1^T + \lambda_2 \mathbf{u}_2 \mathbf{u}_2^T + \cdots + \lambda_n \mathbf{u}_n \mathbf{u}_n^T = \lambda_1 P_1 + \lambda_2 P_2 + \cdots + \lambda_n P_n, \quad (3.2)$$

y cada proyección P_i proyecta sobre la línea generada por el vector propio correspondiente $\mathbf{u}_i$.

Conclusión. Toda matriz simétrica tiene un conjunto completo de vectores propios ortonormales y por lo tanto puede diagonalizarse. Mas aún, esta propiedad permite expresar la matriz como combinación de proyecciones ortogonales de rango uno. Así pues, la simetría y ortogonalidad aparecen ligadas en el álgebra lineal.

EJEMPLO 3.2. La matriz $A = \begin{bmatrix} 3/2 & -1/2 \\ -1/2 & 3/2 \end{bmatrix}$ es simétrica tiene los valores propios $\lambda_1 = 1$, $\lambda_2 = 2$ con vectores propios $\mathbf{u}_1 = (1/\sqrt{2}, 1/\sqrt{2})^T$, $\mathbf{u}_2 = (-1/\sqrt{2}, 1/\sqrt{2})^T$. Su descomposición espectral (diagonalización) es

$$A = \begin{bmatrix} \frac{3}{2} & -\frac{1}{2} \\ -\frac{1}{2} & \frac{3}{2} \end{bmatrix} = \begin{bmatrix} \frac{1}{\sqrt{2}} & -\frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix} \begin{bmatrix} \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \\ -\frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \end{bmatrix} = UDU^T = A^T.$$

La factorización permite realizar la transformación $\mathbf{x} \mapsto A\mathbf{x}$ en tres pasos:

$$\mathbf{x} \mapsto U^T \mathbf{x} \mapsto DU^T \mathbf{x} \mapsto UDU^T \mathbf{x} = A\mathbf{x}.$$

Geoméricamente las matrices ortogonales U y U^T producen rotaciones o reflexiones (dependiendo de si $\det(U) = 1$ ó -1), y D es una matriz de deformación o estiramiento a lo largo de las coordenadas cartesianas $\mathbf{e}_1, \mathbf{e}_2$, como se muestra en la figura 3.1.

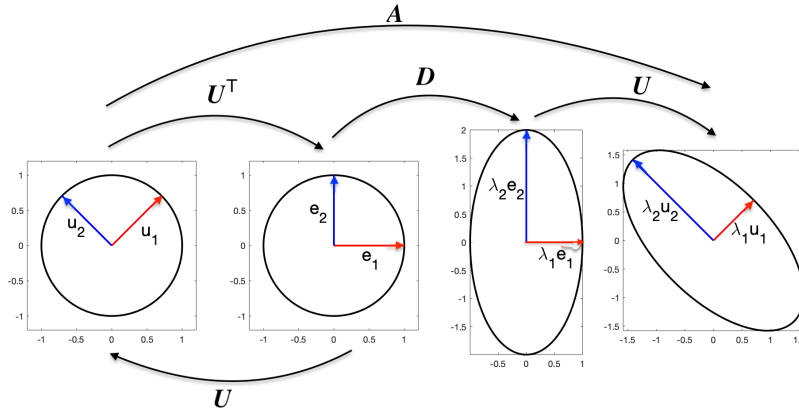


Figura 3.1: La aplicación $\mathbf{x} \mapsto A\mathbf{x}$ utilizando la descomposición $A = UDU^T$ en tres pasos $\mathbf{x} \mapsto U^T \mathbf{x} \mapsto DU^T \mathbf{x} \mapsto UDU^T \mathbf{x}$.

3.2

Asimetría y simetrización. Descomposición SVD

Desafortunadamente la diagonalización $A = UDU^T$ no es siempre posible. La condición para su existencia es que la matriz A de $n \times n$ sea simétrica. Entonces la pregunta es si es posible lograr alguna descomposición de la matriz A cuando no es simétrica. Más aún, nos preguntamos si es posible lograr una factorización análoga cuando la matriz A no es cuadrada. La respuesta ya la hemos anticipado al principio de este capítulo, la descomposición es posible y se denomina **descomposición en valores singulares** (SVD, ‘singular value decomposition’).

Simetrizando desde la asimetría

La clave para lograr la descomposición de valores singulares es ‘simetrizar partiendo de la asimetría’. Es decir, si A es una matriz de tamaño $m \times n$, podemos considerar los productos $A^T A$ y AA^T , los cuales proporcionan matrices simétricas cuadradas de tamaños $n \times n$ y $m \times m$,

respectivamente. Por el teorema espectral 3.1 para matrices simétricas, estas dos matrices se pueden diagonalizar. Por ejemplo, suponiendo que los valores propios de $A^T A$ son $\lambda_1, \dots, \lambda_n$ y que los vectores propios correspondientes son los vectores ortonormales $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n$, entonces

$$A^T A = V D V^T, \quad \text{con } V = \begin{bmatrix} | & | & & | \\ \mathbf{v}_1 & \mathbf{v}_2 & \cdots & \mathbf{v}_n \\ | & | & & | \end{bmatrix}, \quad D = \begin{bmatrix} \lambda_1 & 0 & \cdots & 0 \\ 0 & \lambda_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \lambda_n \end{bmatrix}.$$

Debido a la estructura de la matriz simétrica $A^T A$ se tienen propiedades adicionales. Por ejemplo, podemos verificar que los valores propios son no negativos, debido a que la matriz $A^T A$ es positiva semidefinida, pues

$$\mathbf{x} A A^T \mathbf{x} = \|A \mathbf{x}\|^2 \geq 0, \quad \text{para todo } \mathbf{x} \in \mathbb{R}^n,$$

y en consecuencia, para cualquier valor propio λ con vector propio $\mathbf{v}$, se tiene

$$A^T A \mathbf{v} = \lambda \mathbf{v} \Rightarrow \mathbf{v}^T A^T A \mathbf{v} = \lambda \mathbf{v}^T \mathbf{v} \Rightarrow \lambda = \frac{\|A \mathbf{v}\|^2}{\|\mathbf{v}\|^2} \geq 0. \quad (3.3)$$

Por lo tanto, podemos ordenar los valores propios. Sin pérdida de generalidad, suponemos que

$$\lambda_1 \geq \lambda_2 \geq \cdots \geq \lambda_n \geq 0. \quad (\text{algunos pueden repetirse})$$

Valores propios nulos. Si $\lambda = 0$ es un valor propio de $A^T A$, entonces su espacio propio asociado es el espacio nulo $\mathcal{N}(A)$ de A , pues debido a (3.3) los vectores propios asociados satisfacen $A \mathbf{v} = 0$.

Matrices con rango deficiente. El número de columnas linealmente independientes de cualquier matriz es igual al número de renglones linealmente independientes. Así que el número de columnas (o renglones) linealmente independientes de una matriz A es una característica muy importante que se denomina **rango de la matriz**, y usualmente se denota por $r = \text{rango}(A)$. Si $r < \min\{m, n\}$, entonces $A^T A$ es de rango deficiente y los primeros r valores propios son positivos y el resto son nulos.

Igualdad de valores propios de $A^T A$ y $A A^T$. Los valores propios de ambas matrices son los mismos pues si λ es un valor propio de $A^T A$ con vector propio $\mathbf{v}$, entonces

$$A^T A \mathbf{v} = \lambda \mathbf{v} \Rightarrow (A A^T) A \mathbf{v} = \lambda A \mathbf{v},$$

es decir λ también es valor propio de $A A^T$ con vector propio $\mathbf{u} = A \mathbf{v}$. En forma análoga, se puede verificar que si λ es valor propio de $A A^T$ con vector propio $\mathbf{u}$, entonces λ también es valor propio de $A^T A$ con vector propio $\mathbf{v} = A^T \mathbf{u}$.

Las anteriores propiedades se pueden resumir en el siguiente resultado:

TEOREMA 3.3. *Dada A matriz de $m \times n$, se satisfacen las siguientes aseveraciones:*

- Las matrices $A^T A$ de $n \times n$ y $A A^T$ de $m \times m$ son positivas semidefinidas y tienen los mismos valores propios no-negativos.
- Si $r = \text{rango}(A)$, entonces los primeros r valores propios son positivos y los restantes son cero.

- Si $\lambda = 0$ es valor propio de $A^T A$, entonces el espacio propio asociado es $\mathcal{N}(A)$ (el espacio nulo de A).
- Si $\mathbf{v}$ es valor propio de $A^T A$ asociado al valor propio λ , entonces $\mathbf{u} = A\mathbf{v}$ es vector propio de AA^T con el mismo valor propio.
- Si $\mathbf{u}$ es valor propio de AA^T asociado al valor propio λ , entonces $\mathbf{v} = A^T \mathbf{u}$ es vector propio de $A^T A$ con el mismo valor propio.

Descomposición espectral de matrices rectangulares

De acuerdo al teorema previo, las matrices $A^T A$ y AA^T tienen los mismos valores propios

$$\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n \geq 0,$$

con vectores propios correspondientes

$$\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n \in \mathbb{R}^n \quad \text{y} \quad A\mathbf{v}_1, A\mathbf{v}_2, \dots, A\mathbf{v}_n \in \mathbb{R}^m,$$

respectivamente. Los vectores propios $\mathbf{v}_j$ de $A^T A$ son ortonormales. Los vectores propios $\mathbf{u}_j$ de AA^T son ortogonales, pues

$$(A\mathbf{v}_i)^T A\mathbf{v}_j = \mathbf{v}_i^T A^T A\mathbf{v}_j = \mathbf{v}_i^T \lambda_j \mathbf{v}_j = \lambda_j \delta_{ij}.$$

Por lo tanto, normalizando estos vectores, obtenemos el siguiente conjunto ortonormal de n vectores propios en $\mathbb{R}^m$ para AA^T :

$$\mathbf{u}_j = \frac{A\mathbf{v}_j}{\|A\mathbf{v}_j\|} = \frac{A\mathbf{v}_j}{(\lambda_j)^{1/2}}, \quad j = 1, 2, \dots, n. \quad (3.4)$$

DEFINICIÓN 3.4. Los valores no negativos

$$\sigma_1 = \sqrt{\lambda_1} \geq \sigma_2 = \sqrt{\lambda_2} \geq \dots \geq \sigma_n = \sqrt{\lambda_n} \geq 0,$$

se denominan los *valores singulares* de la matriz A .

Por lo tanto, de acuerdo a (3.4), se obtiene la siguiente relación entre los dos conjuntos de vectores ortonormales $\{\mathbf{u}_1, \dots, \mathbf{u}_n\} \subset \mathbb{R}^m$ y $\{\mathbf{v}_1, \dots, \mathbf{v}_n\} \subset \mathbb{R}^n$:

$$A\mathbf{v}_j = \sigma_j \mathbf{u}_j, \quad j = 1, 2, \dots, n. \quad (3.5)$$

Estas relaciones pueden expresarse como el producto matricial:

$$A \begin{bmatrix} | & | & \cdots & | \\ \mathbf{v}_1 & \mathbf{v}_2 & \cdots & \mathbf{v}_n \\ | & | & & | \end{bmatrix} = \begin{bmatrix} | & | & \cdots & | \\ \sigma_1 \mathbf{u}_1 & \sigma_2 \mathbf{u}_2 & \cdots & \sigma_n \mathbf{u}_n \\ | & | & & | \end{bmatrix}$$

y que finalmente puede expresarse como se indica a continuación.

Factorización reducida SVD

$$A = \underbrace{\begin{bmatrix} | & | & & | \\ \mathbf{u}_1 & \mathbf{u}_2 & \cdots & \mathbf{u}_n \\ | & | & & | \end{bmatrix}}_{\widehat{U} \ (m \times n)} \underbrace{\begin{bmatrix} \sigma_1 & 0 & \cdots & 0 \\ 0 & \sigma_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \sigma_n \end{bmatrix}}_{\widehat{\Sigma} \ (n \times n)} \underbrace{\begin{bmatrix} \text{---} & \mathbf{v}_1^T & \text{---} \\ \text{---} & \mathbf{v}_2^T & \text{---} \\ & \vdots & \\ \text{---} & \mathbf{v}_n^T & \text{---} \end{bmatrix}}_{V^T \ (n \times n)} \quad (3.6)$$

Esta factorización también puede expresarse como **suma de matrices de rango uno**:

$$A = \sigma_1 \mathbf{u}_1 \mathbf{v}_1^T + \sigma_2 \mathbf{u}_2 \mathbf{v}_2^T + \cdots + \sigma_n \mathbf{u}_n \mathbf{v}_n^T, \quad \text{con } \sigma_1 \geq \sigma_2 \geq \cdots \geq \sigma_n \geq 0. \quad (3.7)$$

Nota. En la ecuación matricial $AV = \widehat{U}\widehat{\Sigma}$ o $A = \widehat{U}\widehat{\Sigma}V^T$, la matriz $\widehat{U}$ es una matriz rectangular de tamaño $m \times n$ con columnas ortonormales en $\mathbb{R}^m$, $\widehat{\Sigma}$ es una matriz diagonal de $n \times n$ con los valores singulares y V es una matriz ortogonal (i.e. $V^{-1} = V^T$) de $n \times n$. La descomposición es también válida para matrices con entradas o coeficientes complejos, en cuyo caso la matriz cuadrada V es hermitiana y en la factorización se utiliza V^* , la matriz compleja conjugada, en lugar de V^T .

La descomposición SVD completa

En la mayoría de las aplicaciones la descomposición SVD se utiliza en la forma reducida. Sin embargo, en los textos y muchas publicaciones se utiliza la llamada descomposición SVD ‘**completa**’. Consideramos dos casos: $m > n$ y $m < n$. El caso $m = n$ se obtendrá fácilmente de cualquiera de los dos casos anteriores.

- **Caso $m > n$:** las columnas de la matriz $\widehat{U}$ no forman una base de $\mathbb{R}^m$. En este caso la idea es aumentar la matriz, agregando $m - n$ columnas ortonormales para extenderla a una matriz ortogonal de $m \times m$ (unitaria, en el caso complejo), que llamamos U . Al reemplazar $\widehat{U}$ por U también debemos reemplazar la matriz diagonal $\widehat{\Sigma}$ por la matriz rectangular Σ de $m \times n$, agregando $m - n$ renglones con ceros debajo de $\widehat{\Sigma}$. La siguiente ilustración muestra en color azul las columnas agregadas a $\widehat{U}$ y los renglones agregados a $\widehat{\Sigma}$:

$$A = \underbrace{\begin{bmatrix} | & | & & | & | & | \\ \mathbf{u}_1 & \mathbf{u}_2 & \cdots & \mathbf{u}_n & \mathbf{u}_{n+1} & \cdots & \mathbf{u}_m \\ | & | & & | & | & | \end{bmatrix}}_{U \ (m \times m)} \underbrace{\begin{bmatrix} \sigma_1 & 0 & \cdots & 0 \\ 0 & \sigma_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \sigma_n \\ 0 & 0 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & 0 \end{bmatrix}}_{\Sigma \ (m \times n)} \underbrace{\begin{bmatrix} \text{---} & \mathbf{v}_1^T & \text{---} \\ \text{---} & \mathbf{v}_2^T & \text{---} \\ & \vdots & \\ \text{---} & \mathbf{v}_n^T & \text{---} \end{bmatrix}}_{V^T \ (n \times n)} \quad (3.8)$$

- **Caso $m < n$:** la descomposición reducida es de la forma $A = \widehat{U}\widehat{\Sigma}\widehat{V}^T$. En este caso U es matriz ortogonal de $m \times m$, $\widehat{\Sigma}$ matriz diagonal de $m \times m$ con los valores singulares y $\widehat{V}^T$ matriz de $m \times n$ con filas ortonormales. Los valores singulares $\sigma_i = \sqrt{\lambda_i}$, se obtienen de los valores propios $\{\lambda_i\}_{i=1}^m$ de la matriz simétrica AA^T . Además, las m filas de $\widehat{V}^T$ no forman

una base de $\mathbb{R}^n$. Por lo tanto, la descomposición SVD completa, se obtiene agregando $n - m$ columnas cero a $\hat{\Sigma}$ y $n - m$ filas ortogonales a las de $\hat{V}^T$. La siguiente ilustración muestra la descomposición SVD completa en este caso.

$$A = \underbrace{\begin{bmatrix} | & | & & | \\ \mathbf{u}_1 & \mathbf{u}_2 & \cdots & \mathbf{u}_m \\ | & | & & | \end{bmatrix}}_{U \ (m \times m)} \underbrace{\begin{bmatrix} \sigma_1 & 0 & \cdots & 0 & 0 & \cdots & 0 \\ 0 & \sigma_2 & \cdots & 0 & 0 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & \cdots & \sigma_m & 0 & \cdots & 0 \end{bmatrix}}_{\Sigma \ (m \times n)} \underbrace{\begin{bmatrix} \text{---} & \mathbf{v}_1^T & \text{---} \\ \text{---} & \mathbf{v}_2^T & \text{---} \\ & \vdots & \\ \text{---} & \mathbf{v}_m^T & \text{---} \\ \text{---} & \mathbf{v}_{m+1}^T & \text{---} \\ & \vdots & \\ \text{---} & \mathbf{v}_n^T & \text{---} \end{bmatrix}}_{V^T \ (n \times n)} \quad (3.9)$$

Los anteriores resultados se resumen en el siguiente teorema.

TEOREMA 3.5. Toda matriz $A \in \mathbb{R}^{m \times n}$ ($\mathbb{C}^{m \times n}$, en el caso complejo) tiene una descomposición en valores singulares de la forma

$$A = U \Sigma V^T, \quad (A = U \Sigma V^*, \text{ en el caso complejo}) \quad (3.10)$$

con las matrices ortogonales U, V (o unitarias, en el caso complejo), y la matriz Σ como se indica en el desarrollo anterior.

EJEMPLO 3.6. Encontrar la descomposición SVD de la matriz $A = \begin{bmatrix} -3 & 1 \\ 6 & -2 \\ 6 & -2 \end{bmatrix}$ (caso $m > n$).

Solución. Las matrices simétricas son

$$A^T A = \begin{bmatrix} 81 & -27 \\ -27 & 81 \end{bmatrix}, \quad A A^T = \begin{bmatrix} 10 & -20 & -20 \\ -20 & 40 & 40 \\ -40 & 40 & 40 \end{bmatrix}.$$

Los valores propios de $A^T A$ son $\lambda_1 = 90$, $\lambda_2 = 0$, por lo que $\sigma_1 = 3\sqrt{10}$, $\sigma_2 = 0$.
Vector propio unitario para $\lambda_1 = 90$:

$$A^T A - \lambda_1 I = \begin{bmatrix} -9 & -27 \\ -27 & -81 \end{bmatrix} \sim \begin{bmatrix} 1 & 3 \\ 0 & 0 \end{bmatrix} \Rightarrow \mathbf{v}_1 = \frac{1}{\sqrt{10}} \begin{bmatrix} -3 \\ 1 \end{bmatrix}.$$

Vector propio unitario para $\lambda_2 = 0$:

$$A^T A - \lambda_2 I = \begin{bmatrix} 81 & -27 \\ -27 & 9 \end{bmatrix} \sim \begin{bmatrix} 1 & -1/3 \\ 0 & 0 \end{bmatrix} \Rightarrow \mathbf{v}_2 = \frac{1}{\sqrt{10}} \begin{bmatrix} 1 \\ 3 \end{bmatrix}.$$

Ahora consideremos los vectores propios de $A A^T$:

$$\mathbf{u}_1 = \frac{1}{\sigma_1} A \mathbf{v}_1 = \frac{1}{3\sqrt{10}} \begin{bmatrix} -3 & 1 \\ 6 & -2 \\ 6 & -2 \end{bmatrix} \frac{1}{\sqrt{10}} \begin{bmatrix} -3 \\ 1 \end{bmatrix} = \frac{1}{3} \begin{bmatrix} 1 \\ -2 \\ -2 \end{bmatrix}.$$

Los otros vectores propios no pueden calcularse de la misma manera ya que $\sigma_2 = 0$. Sin embargo, pertenecen a $\mathcal{N}(AA^T)$.

$$AA^T \sim \begin{bmatrix} 1 & -2 & -2 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix} \Rightarrow \mathbf{u}_2 = \frac{1}{\sqrt{5}} \begin{bmatrix} 2 \\ 1 \\ 0 \end{bmatrix}, \quad \mathbf{u}_3 = \frac{1}{\sqrt{5}} \begin{bmatrix} 2 \\ 0 \\ 1 \end{bmatrix}.$$

Por lo tanto, $A = U\Sigma V^T$ con

$$U = \begin{bmatrix} 1/3 & 2/\sqrt{5} & 2/\sqrt{5} \\ -2/3 & 1/\sqrt{5} & 0 \\ -2/3 & 0 & 1/\sqrt{5} \end{bmatrix}, \quad \Sigma = \begin{bmatrix} 3\sqrt{10} & 0 \\ 0 & 0 \\ 0 & 0 \end{bmatrix}, \quad V = \begin{bmatrix} -3/\sqrt{10} & 1/\sqrt{10} \\ 1/\sqrt{10} & 3/\sqrt{10} \end{bmatrix}.$$

Por último, la descomposición SVD reducida es $A = \hat{U}\hat{\Sigma}\hat{V}^T$ con

$$\hat{U} = \begin{bmatrix} 1/3 \\ -2/3 \\ -2/3 \end{bmatrix}, \quad \hat{\Sigma} = [3\sqrt{10}], \quad \hat{V} = \begin{bmatrix} -3/\sqrt{10} \\ 1/\sqrt{10} \end{bmatrix}.$$

EJEMPLO 3.7. Encontrar la descomposición SVD de la matriz $A = \begin{bmatrix} 3 & 2 & 2 \\ 2 & 3 & -2 \end{bmatrix}$ (caso $m < n$).

Solución. En este ejemplo preferimos obtener los valores propios de AA^T , los cuales son $\lambda_1 = 25$, $\lambda_2 = 9$. Los valores singulares son $\sigma_1 = 5$, $\sigma_2 = 3$.

Vector propio unitario para $\lambda_1 = 25$,

$$A^T A - \lambda_1 I = \begin{bmatrix} -8 & 8 \\ 8 & -8 \end{bmatrix} \sim \begin{bmatrix} 1 & -1 \\ 0 & 0 \end{bmatrix} \Rightarrow \mathbf{u}_1 = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ 1 \end{bmatrix}.$$

Vector propio unitario para $\lambda_2 = 9$,

$$A^T A - \lambda_2 I = \begin{bmatrix} 8 & 8 \\ 8 & 8 \end{bmatrix} \sim \begin{bmatrix} 1 & 1 \\ 0 & 0 \end{bmatrix} \Rightarrow \mathbf{u}_2 = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ -1 \end{bmatrix}.$$

Los correspondientes dos vectores propios de $A^T A = \begin{bmatrix} 13 & 12 & 2 \\ 12 & 13 & -2 \\ 2 & -2 & 8 \end{bmatrix}$ son

$$\mathbf{v}_1 = \frac{1}{\sigma_1} A^T \mathbf{u}_1 = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ 1 \\ 0 \end{bmatrix}, \quad \mathbf{v}_2 = \frac{1}{\sigma_2} A^T \mathbf{u}_2 = \frac{1}{3\sqrt{2}} \begin{bmatrix} 1 \\ -1 \\ 4 \end{bmatrix}.$$

El tercer vector propio se obtiene del espacio nulo de $A^T A$:

$$A^T A = \begin{bmatrix} 13 & 12 & 2 \\ 12 & 13 & -2 \\ 2 & -2 & 8 \end{bmatrix} \sim \begin{bmatrix} 1 & 0 & 2 \\ 0 & 1 & -2 \\ 0 & 0 & 0 \end{bmatrix} \Rightarrow \mathbf{v}_3 = \frac{1}{3} \begin{bmatrix} -2 \\ 2 \\ 1 \end{bmatrix}.$$

Por lo tanto, la descomposición SVD completa es $A = U\Sigma V^T$ con

$$U = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}, \quad \Sigma = \begin{bmatrix} 5 & 0 \\ 0 & 3 \\ 0 & 0 \end{bmatrix}, \quad V = \frac{1}{3\sqrt{2}} \begin{bmatrix} 3 & 1 & -2\sqrt{2} \\ 3 & -1 & 2\sqrt{2} \\ 0 & 4 & \sqrt{2} \end{bmatrix}.$$

La descomposición reducida es $A = U\hat{\Sigma}\hat{V}^T$ con

$$U = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}, \quad \hat{\Sigma} = \begin{bmatrix} 5 & 0 \\ 0 & 3 \end{bmatrix}, \quad \hat{V} = \frac{1}{3\sqrt{2}} \begin{bmatrix} 3 & 1 \\ 3 & -1 \\ 0 & 4 \end{bmatrix}.$$

Caso cuando A tiene rango deficiente

En el caso en que la matriz $A \in \mathbb{R}^{m \times n}$ sea de rango deficiente $r = \text{rango}(A) < \min(m, n)$, los valores singulares mayores a cero son $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_r > 0$ y el resto son cero. Por lo tanto, también podemos expresar la descomposición SVD completa o reducida como la suma de matrices de rango uno:

$$A = \sigma_1 \mathbf{u}_1 \mathbf{v}_1^T + \sigma_2 \mathbf{u}_2 \mathbf{v}_2^T + \dots + \sigma_r \mathbf{u}_r \mathbf{v}_r^T. \quad (3.11)$$

En el ejemplo anterior la matriz A tiene rango $r = 1$, pues sus columnas son colineales y

$$A = \begin{bmatrix} -3 & 1 \\ 6 & -2 \\ 6 & -2 \end{bmatrix} = 3\sqrt{10} \begin{bmatrix} 1/3 \\ -2/3 \\ -2/3 \end{bmatrix} \begin{bmatrix} -3/\sqrt{10} & 1/\sqrt{10} \end{bmatrix} = \sigma_1 \mathbf{u}_1 \mathbf{v}_1^T.$$

3.3

Resumen de propiedades

Propiedades básicas

En la descomposición $A = U\Sigma V^T \in \mathbb{R}^{m \times n}$ se satisface:

- Las matrices $A^T A$ y AA^T tienen los mismo valores propios no negativos λ_i .
- Las matrices $U \in \mathbb{R}^{m \times m}$ y $V \in \mathbb{R}^{n \times n}$ son ortogonales.
- Las columnas de U son los vectores propios de AA^T .
- Las columnas de V son los vectores propios de $A^T A$.
- La matriz Σ de $m \times n$ contiene los valores singulares $\sigma_i = \sqrt{\lambda_i}$ en la diagonal.
- Los vectores propios normalizados $\mathbf{u}_i$ y $\mathbf{v}_i$ satisfacen

$$A\mathbf{v}_i = \sigma_i \mathbf{u}_i, \quad A^T \mathbf{u}_i = \sigma_i \mathbf{v}_i, \quad \|A\mathbf{v}_i\| = \sigma_i = \|A^T \mathbf{u}_i\|.$$

- Si $A \in \mathbb{R}^{n \times n}$, $|\det(A)| = |\det(U) \det(\Sigma) \det(V^T)| = \det(\Sigma) = \prod_{i=1}^n \sigma_i$.

- Si A es simétrica, entonces $U = V$ y sus valores singulares son los valores absolutos de los valores propios de A .
- $\text{rango}(A) = \text{rango}(\Sigma) = r$ es el número de valores singulares no cero.
- A es la suma de r matrices de rango uno:

$$A = \sigma_1 \mathbf{u}_1 \mathbf{v}_1^T + \sigma_2 \mathbf{u}_2 \mathbf{v}_2^T + \cdots + \sigma_r \mathbf{u}_r \mathbf{v}_r^T.$$

- $\|A\|_2 = \|\Sigma\|_2 = \sigma_1$ y $\|A\|_F = \sqrt{\sigma_1^2 + \sigma_2^2 + \cdots + \sigma_r^2}$, son dos expresiones para calcular la norma de una matriz A : la norma-2 y la norma de Frobenius.

Los cuatro subespacios fundamentales de una matriz

Si la matriz A es de rango r , y suponemos que el orden de los vectores propios en U y V corresponde al orden de los valores singulares de mayor a menor, entonces

- El espacio columna de A es $\mathcal{R}(A) = \text{Gen}\{\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_r\}$.
- El espacio nulo de A es $\mathcal{N}(A) = \text{Gen}\{\mathbf{v}_{r+1}, \mathbf{v}_{r+2}, \dots, \mathbf{v}_n\}$.
- El espacio fila de A es $\mathcal{R}(A^T) = \text{Gen}\{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r\}$.
- El espacio nulo de A^T es $\mathcal{N}(A^T) = \text{Gen}\{\mathbf{u}_{r+1}, \mathbf{u}_{r+2}, \dots, \mathbf{u}_m\}$.

Esquemáticamente:

$$A = \left[\begin{array}{ccc|ccc} | & & & | & & \\ | & & & | & & \\ | & & & | & & \\ \hline | & & & | & & \\ | & & & | & & \\ | & & & | & & \end{array} \right] \left[\begin{array}{cccccc} \sigma_1 & \cdots & 0 & 0 & \cdots & 0 \\ \vdots & \ddots & \vdots & \vdots & \ddots & \vdots \\ 0 & \cdots & \sigma_r & 0 & \cdots & 0 \\ 0 & \cdots & 0 & 0 & \cdots & 0 \\ \vdots & \ddots & \vdots & \vdots & \ddots & \vdots \\ 0 & \cdots & 0 & 0 & \cdots & 0 \end{array} \right] \left\{ \begin{array}{l} \left[\begin{array}{ccc} - & \mathbf{v}_1^T & - \\ & \vdots & \\ - & \mathbf{v}_r^T & - \\ - & \mathbf{v}_{r+1}^T & - \\ & \vdots & \\ - & \mathbf{v}_n^T & - \end{array} \right] \mathcal{R}(A^T) \\ \left[\begin{array}{ccc} - & \mathbf{v}_1^T & - \\ & \vdots & \\ - & \mathbf{v}_r^T & - \\ - & \mathbf{v}_{r+1}^T & - \\ & \vdots & \\ - & \mathbf{v}_n^T & - \end{array} \right] \mathcal{N}(A) \end{array} \right\}$$

$\underbrace{\hspace{10em}}_{\mathcal{R}(A)} \quad \underbrace{\hspace{10em}}_{\mathcal{N}(A^T)}$

Interpretación geométrica

Si consideramos la matriz $A \in \mathbb{R}^{m \times n}$ como la matriz de la transformación lineal $\mathbf{x} \mapsto A\mathbf{x}$ del espacio $\mathbb{R}^n$ al espacio $\mathbb{R}^m$, podemos pensar a las matrices ortogonales U y V como transformaciones que representan rotaciones o reflexiones, y a la matriz Σ como matriz de deformación o estiramiento.

EJEMPLO 3.8. Consideremos la matriz cuadrada $A = \begin{bmatrix} 1 & 1 \\ 2 & -2 \end{bmatrix}$. Esta matriz transforma una circunferencia unitaria en $\mathbb{R}^2$ sobre una elipse en $\mathbb{R}^2$. La transformación se puede realizar en tres pasos utilizando la descomposición SVD, como se muestra a continuación.

Los valores propios de

$$A^T A = \begin{bmatrix} 5 & -3 \\ -3 & 5 \end{bmatrix} \quad \text{y} \quad A A^T = \begin{bmatrix} 2 & 0 \\ 0 & 8 \end{bmatrix},$$

son $\lambda_1 = 8$, $\lambda_2 = 2$. Por lo tanto, los valores singulares en orden decreciente son $\sigma_1 = 2\sqrt{2}$, $\sigma_2 = \sqrt{2}$. Los vectores propios de $A^T A$ se obtienen de

$$\begin{aligned} A^T A - \lambda_1 I &= \begin{bmatrix} -3 & -3 \\ -3 & -3 \end{bmatrix} \sim \begin{bmatrix} 1 & 1 \\ 0 & 0 \end{bmatrix} \Rightarrow \mathbf{v}_1 = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ -1 \end{bmatrix}, \\ A^T A - \lambda_2 I &= \begin{bmatrix} 3 & -3 \\ -3 & 3 \end{bmatrix} \sim \begin{bmatrix} 1 & -1 \\ 0 & 0 \end{bmatrix} \Rightarrow \mathbf{v}_2 = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ 1 \end{bmatrix}. \end{aligned}$$

Los vectores propios de $A A^T$ son

$$\begin{aligned} \mathbf{u}_1 &= \frac{1}{\sigma_1} A \mathbf{v}_1 = \frac{1}{2\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 2 & -2 \end{bmatrix} \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ -1 \end{bmatrix} = \begin{bmatrix} 0 \\ 1 \end{bmatrix}, \\ \mathbf{u}_2 &= \frac{1}{\sigma_2} A \mathbf{v}_2 = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 2 & -2 \end{bmatrix} \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ 1 \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \end{bmatrix}. \end{aligned}$$

La descomposición SVD de A es

$$A = U \Sigma V^T = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} 2\sqrt{2} & 0 \\ 0 & \sqrt{2} \end{bmatrix} \begin{bmatrix} 1/\sqrt{2} & -1/\sqrt{2} \\ 1/\sqrt{2} & 1/\sqrt{2} \end{bmatrix}.$$

La transformación $\mathbf{x} \mapsto \mathbf{b} = A\mathbf{x}$, cuando se utiliza la descomposición SVD, involucra los tres pasos, como se muestra en la figura 3.2:

$$\mathbf{x}' = V^T \mathbf{x} \implies \mathbf{b}' = \Sigma \mathbf{x}' \implies \mathbf{b} = U \mathbf{b}',$$

en donde hay dos cambios de base (a bases ortogonales), uno en el primer paso y otro en el tercer paso. Es decir

$$A \mathbf{x} = \mathbf{b} \iff \Sigma \mathbf{x}' = \mathbf{b}' \quad \text{con} \quad \begin{cases} \mathbf{x}' = V^T \mathbf{x} \\ \mathbf{b}' = U^T \mathbf{b} \end{cases}$$

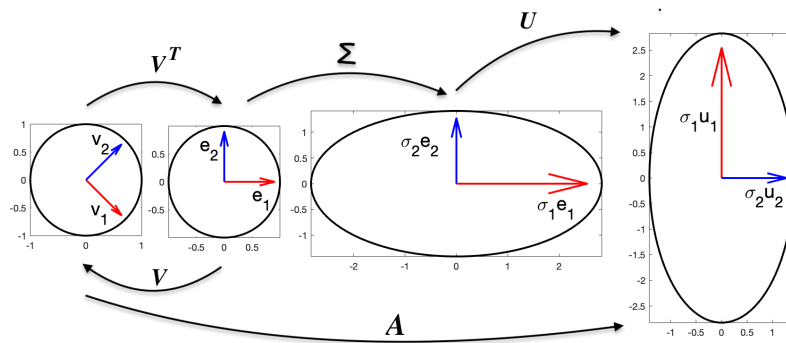


Figura 3.2: La aplicación $\mathbf{x} \mapsto A\mathbf{x}$ utilizando la descomposición $A = U\Sigma V^T$ en tres pasos $\mathbf{x} \mapsto V^T \mathbf{x} \mapsto \Sigma V^T \mathbf{x} \mapsto U \Sigma V^T \mathbf{x}$.

3.4

Aplicaciones

Sistemas de ecuaciones

Cuando se tiene disponible la descomposición SVD de la matriz A de $m \times n$, es posible describir completamente el conjunto de soluciones del sistema

$$A\mathbf{x} = \mathbf{b}, \quad \text{con} \quad \begin{cases} \mathbf{x} \in \mathbb{R}^n, \\ \mathbf{b} \in \mathbb{R}^m. \end{cases} \quad (3.12)$$

De acuerdo al último párrafo de la sección anterior, el sistema se puede representar como un sistema diagonal

$$\Sigma \mathbf{x}' = \mathbf{b}', \quad \text{con los vectores 'rotados'} \quad \begin{cases} \mathbf{x}' = V^T \mathbf{x} & \text{en } \mathbb{R}^n, \\ \mathbf{b}' = U^T \mathbf{b} & \text{en } \mathbb{R}^m. \end{cases} \quad (3.13)$$

Si el rango de la matriz es r , el sistema diagonal obtenido nos permite escribir

$$\begin{aligned} \sigma_i x'_i &= b'_i, & i = 1, \dots, r, \\ 0 &= b'_i, & i = r+1, \dots, m. \end{aligned}$$

Podemos concluir entonces que existen dos posibilidades:

1. El sistema es inconsistente si $b_i \neq 0$ para alguna $i = r+1, \dots, m$. Esto ocurre cuando $\mathbf{b} \notin \mathcal{R}(A) = \text{Gen}\{\mathbf{u}_1, \dots, \mathbf{u}_r\}$. En este caso el conjunto solución es vacío.
2. Cuando $\mathbf{b} \in \mathcal{R}(A)$, entonces el sistema es consistente y se obtiene

$$\begin{aligned} x'_i &= \frac{b'_i}{\sigma_i}, & i = 1, \dots, r, \\ x'_i &\text{ son variables libres para } & i = r+1, \dots, m. \end{aligned}$$

Las variables libres proporcionan elementos de $\mathcal{N}(A)$. Por otro lado, si A tiene n columnas linealmente independientes (tiene rango completo por columnas), entonces $\mathcal{N}(A) = \{\mathbf{0}\}$, y la solución $\mathbf{x} = V\mathbf{x}'$ es única.

El conjunto solución del sistema lineal (3.12) se puede describir de manera conveniente mediante la llamada pseudo-inversa de A , la cual puede definirse mediante la descomposición SVD.

DEFINICIÓN 3.9. Sea $A = U\Sigma V^T$ de $m \times n$, su pseudo-inversa, denotada por $A^\dagger$, se define mediante

$$A^\dagger = V\Sigma^\dagger U^T, \quad \text{con} \quad \Sigma^\dagger = \text{diag}(1/\sigma_1, \dots, 1/\sigma_r, 0, \dots, 0) \in \mathbb{R}^{n \times m} \quad (3.14)$$

Se observan las siguientes propiedades

- $A^\dagger$ está bien definida para toda matriz A .

- $A^\dagger$ es del mismo tamaño que A^T .
- Si A tiene columnas linealmente independientes, entonces $A^\dagger = (A^T A)^{-1} A^T$, se denomina inversa izquierda de A , pues $A^\dagger A = I_n$.
- Si A tiene filas linealmente independientes, entonces $A^\dagger = A^T (A A^T)^{-1}$, se denomina inversa derecha de A , pues $A A^\dagger = I_m$.
- Cuando A es cuadrada e invertible $A^\dagger = A^{-1}$.

Volviendo al sistema de ecuaciones (3.12) y suponiendo $\mathbf{b} \in \mathcal{R}(A)$, debido a (3.13) y (3.14) se obtiene la solución particular

$$\mathbf{x}' = \Sigma^\dagger \mathbf{b}' \Rightarrow V \mathbf{x}' = V \Sigma^\dagger \mathbf{b}' \Rightarrow V \mathbf{x}' = V \Sigma^\dagger U^T \mathbf{b} \Rightarrow \mathbf{x} = A^\dagger \mathbf{b}.$$

Agregando las soluciones de la ecuación homogénea, $A\mathbf{x} = \mathbf{0}$, se obtiene el conjunto solución:

$$S = \left\{ A^\dagger \mathbf{b} + \mathbf{z} : \mathbf{z} \in \mathcal{N}(A) \right\},$$

Mínimos cuadrados lineales con SVD

En los anteriores capítulos se ha encontrado que los problemas de mínimos cuadrados lineales llevan a problemas lineales sobredeterminados de la forma $A\mathbf{x} = \mathbf{b}$ con A de $m \times n$ y $m \geq n$. Estos problemas usualmente son inconsistentes debido a que $\mathbf{b} \notin \mathcal{R}(A)$. Por este motivo se busca el mínimo $\hat{\mathbf{x}} \in \mathbb{R}^n$ del residual $\mathbf{r} = \mathbf{b} - A\mathbf{x}$. Sabemos que (ver capítulo 1)

$$\hat{\mathbf{x}} = \arg \min_{\mathbf{x} \in \mathbb{R}^n} \frac{1}{2} \|\mathbf{b} - A\mathbf{x}\|_2^2 \iff A^T A \hat{\mathbf{x}} = A^T \mathbf{b}.$$

Además si A es de rango completo $r = n$, entonces $A^T A$ es simétrica definida positiva y la solución de mínimos cuadrados es

$$\hat{\mathbf{x}} = (A^T A)^{-1} A^T \mathbf{b}$$

Via SVD, $A = U \Sigma V^T$, y utilizando que $V^{-1} = V^T$, $U^{-1} = U^T$, $\Sigma^T \Sigma$ invertible, se obtiene

$$(A^T A)^{-1} A^T = (V \Sigma^T U^T U \Sigma V^T)^{-1} V \Sigma^T U^T = V \Sigma^\dagger U^T = A^\dagger.$$

Por lo tanto, la solución de mínimos cuadrados del problema inconsistente $A\mathbf{x} = \mathbf{b}$ es

$$\hat{\mathbf{x}} = V \Sigma^\dagger U^T \mathbf{b} = A^\dagger \mathbf{b}. \quad (3.15)$$

EJEMPLO 3.10. Considerar el conjunto de puntos FILIP del ejemplo 2.4, en donde el ajuste se realizó utilizando el método de ecuaciones normales y el método de factorización *QR*. Realizar el ajuste utilizando la descomposición SVD.

Solución. Dado el conjunto de puntos $(t_1, y_1), \dots, (t_m, y_m)$. El algoritmo para calcular $\hat{\mathbf{x}} \in \mathbb{R}^{n+1}$, con los coeficientes del polinomio de grado n , consiste de los siguientes pasos:

1. Construir la matriz de diseño A de $m \times (n+1)$, con coeficientes $a_{ij} = t_i^{j-1}$.

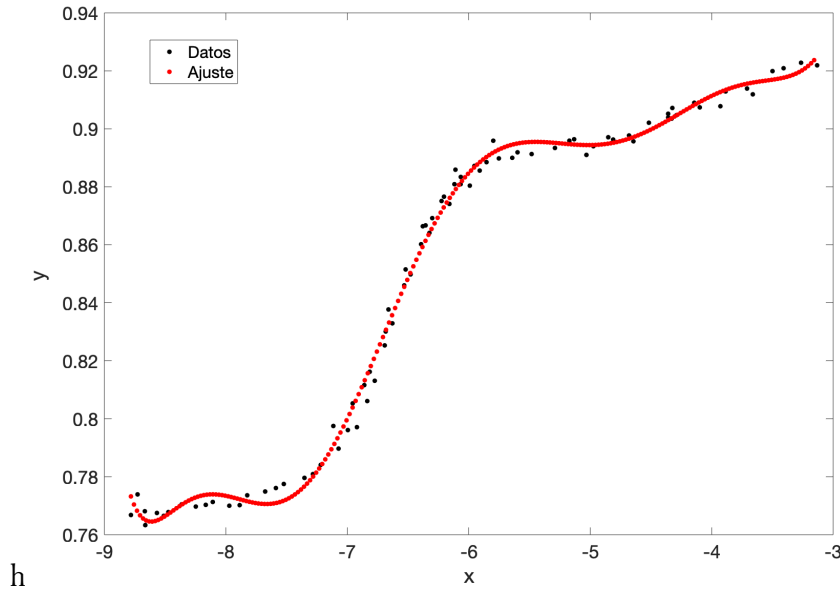


Figura 3.3: Ajuste de los datos FILIP a un polinomio de grado 10 utilizando SVD.

2. Utilizar SVD para obtener $A = U \Sigma V^T$. En realidad basta con la SVD reducida $A = \hat{U} \hat{\Sigma} V^T$.
3. Calcular la inversa generalizada $A^\dagger = V \Sigma^\dagger U^T$.
4. Calcular $\hat{\mathbf{x}} = A^\dagger \mathbf{b}$, con $\mathbf{b} = \{y_j\}_{j=1}^m$ el vector de observaciones.

Se muestra el resultado en la figura 3.3. Se ha utilizado el ambiente MATLAB para los cálculos, utilizando la rutina de descomposición SVD (disponible también en otros ambientes de programación como Phytion).

El resultado del ajuste es muy bueno pero no mejora el que se encontró previamente mediante la factorización QR en el capítulo anterior. Sin embargo, la descomposición SVD ofrece ventajas que permiten su aplicación en muchos campos y problemas. Por ejemplo, una de sus mayores ventajas es que SVD permite realizar aproximaciones de rango bajo.

Aproximaciones de rango bajo

Dada A , matriz de $m \times n$ con $m \geq n$, sabemos que SVD proporciona $A = \hat{U} \hat{\Sigma} V^T = U \Sigma V^T$, y si la matriz es de rango completo n , entonces la descomposición también se puede representar como la suma de n matrices de rango uno

$$A = \sigma_1 \mathbf{u}_1 \mathbf{v}_1^T + \sigma_2^2 \mathbf{u}_2 \mathbf{v}_2^T + \cdots + \sigma_n \mathbf{u}_n \mathbf{v}_n^T,$$

en donde $\sigma_1 \geq \sigma_2 \geq \cdots \geq \sigma_n$. En muchos problemas y aplicaciones la diferencia entre σ_1 y σ_n es muy grande. De hecho, el cociente de estos valores singulares, se denomina el número de condición de la matriz:

$$\text{Cond}(A) = \frac{\sigma_1}{\sigma_n}.$$

En álgebra lineal numérica se conoce que el número de condición de una matriz importa mucho cuando se realizan operaciones en computadora (es decir, en aritmética finita). Un valor enorme indica que la matriz es casi singular (o de rango deficiente). De hecho, algunos valores singulares pequeños pueden ser consecuencia de ruido en los datos para construir la matriz de diseño o ruido

numérico al realizar operaciones en computadora para construir A . Por lo que varios singulares muy pequeños deberían ser cero y/o en realidad no son tan relevantes. Por lo tanto, los primeros valores singulares (los mayores) son los que tienen mayor preponderancia y en muchas aplicaciones basta con aproximar A mediante una matriz de rango menor:

$$A_r = \sigma_1 \mathbf{u}_1 \mathbf{v}_1^T + \sigma_2 \mathbf{u}_2 \mathbf{v}_2^T + \cdots + \sigma_r \mathbf{u}_r \mathbf{v}_r^T, \quad r < n \Rightarrow \|A - A_r\|_2 = \|\Sigma - \Sigma_r\|_2 = \sigma_{r+1}.$$

EJEMPLO 3.11. En el ejemplo anterior, sobre el ajuste de un polinomio de grado 10 con 82 datos. La matriz de diseño A es de tamaño $m \times n = 82 \times 11$ y de rango completo $n = 11$. El máximo y mínimo valores singulares son aproximadamente

$$\sigma_1 = 7.1969e + 09 \quad \sigma_{11} = 4.0707e - 06 \Rightarrow \text{Cond}(A) = 1.76796e + 15,$$

y la matriz es cercana a una de rango deficiente. Es posible descartar algunos valores singulares muy pequeños.

La figura (3.4) muestra los valores singulares. Se utiliza una escala semi-logarítmica para una mejor visualización.

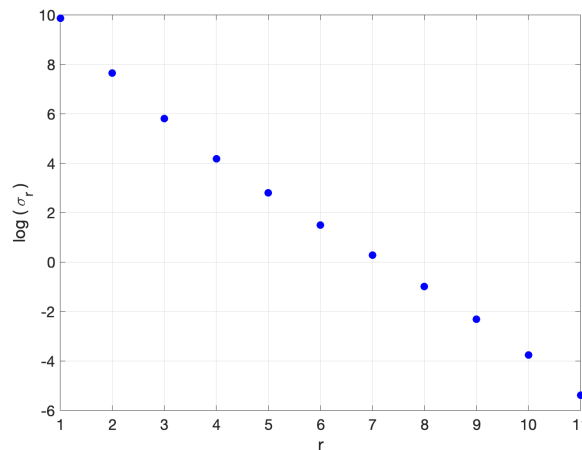


Figura 3.4: Los valores singulares de la matriz de diseño A .

Descartando los últimos valores singulares de A , se obtienen las matrices A_9 , A_7 , A_6 , de rango 9, 7, 6, respectivamente. Utilizando cada una de ellas, se muestran los resultados de ajuste respectivos en la figura 3.5. Se observa que la aproximación es mala para A_6 (rango 6), aceptable para A_7 y muy buena para A_9 (rango 9).

Compresión o reducción de datos

La descomposición SVD tiene múltiples aplicaciones. En particular en cómputo científico, matemáticas, inteligencia artificial, y otras disciplinas, es útil en la reducción de datos redundantes o no relevantes. La información contenida en una matriz A es relevante en la vida cotidiana (Netflix, Google, redes sociales, fotos digitales, etc...). En muchas aplicaciones estas matrices son densas de gran tamaño y ocupan una gran cantidad de memoria para su almacenamiento digital y es conveniente aproximarlas por una de rango bajo, descartando los valores singulares menos

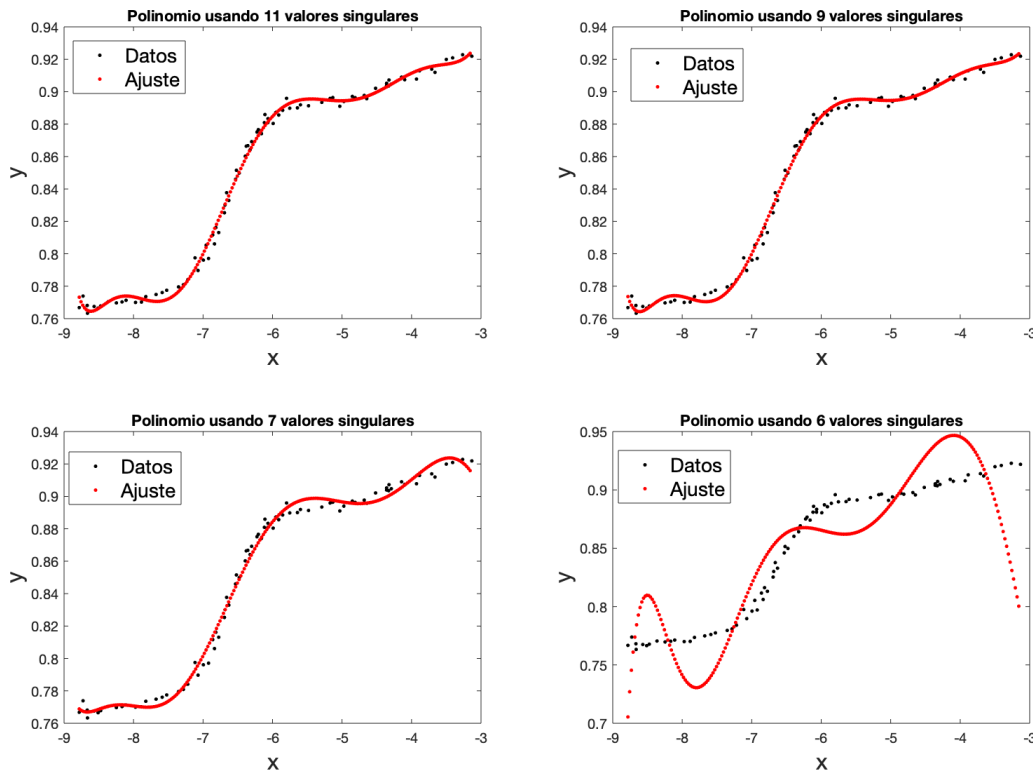


Figura 3.5: Ajuste con la matriz de rango completo A y sus aproximaciones A_r con $r = 9, 7, 6$.

significativos. Los primeros valores singulares, al ser los mayores, son los que generalmente contienen más información de la matriz. Para medir de manera precisa cuántos valores singulares σ_i son relevantes, una posible medida es el porcentaje acumulado por la suma parcial de los valores singulares respecto de la suma total. Es decir

$$A_r = \sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T \approx A = \sum_{i=1}^n \sigma_i \mathbf{u}_i \mathbf{v}_i^T \Rightarrow \% \text{ Aprox} = 100 \times \sum_{i=1}^r \sigma_i / \sum_{i=1}^n \sigma_i \quad (3.16)$$

Por ejemplo, si en la reducción de A queremos mantener el 80% de la información, entonces sumamos los valores singulares de 1 a r hasta alcanzar ese porcentaje, obteniendo la reducción de la matriz original a la aproximada A_r .

EJEMPLO 3.12. La figura 3.6 muestra la imagen de un gato a color, en formato jpg, que se representa por medio de tres matrices en el modelo RGB. En la misma figura se muestra también la imagen de alta resolución en escala de grises, la cual se representa por medio de una matriz A de $m \times n = 1600 \times 2000$ (píxeles). Cada píxel representa la intensidad de gris en la imagen. El objetivo es reducir o compactar la información de esa imagen digital.

Se calcula la descomposición SVD de A utilizando la rutina estándar en la mayoría de ambientes de programación (Matlab o Python, por ejemplo). El número de condición de la matriz es aproximadamente $\text{Cond}(A) = \sigma_1 / \sigma_{1600} = 2.35489e + 16$. El mayor valor singular es $\sigma_1 = 697.809$ y el menor es $\sigma_{1600} = 2.963236e - 14$, aproximadamente. La figura 3.7 muestra los valores singulares en escala logarítmica junto con la figura del porcentaje acumulado conforme se agregan

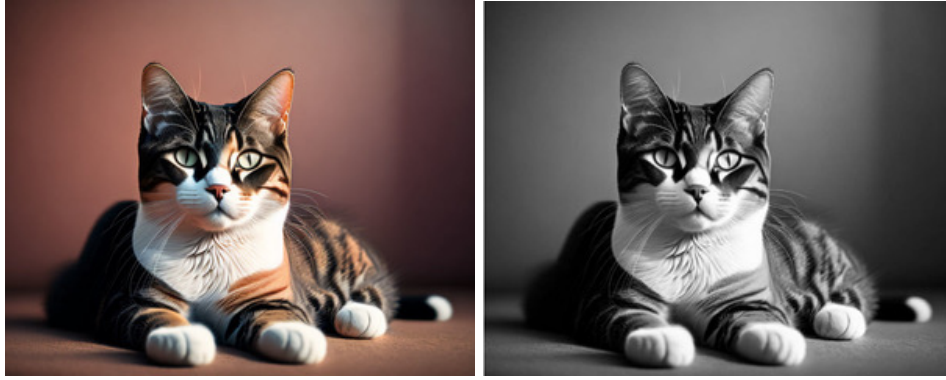


Figura 3.6: Imagen de alta resolución de un gato a color y en blanco y negro.

valores singulares a la aproximación de rango bajo A_r , $1 \leq r < 1600$. La figura 3.8 muestra las imágenes aproximadas para varios valores de r . Es sorprendente que con cuarenta valores singulares obtenemos una buena calidad de imagen, y que con cien valores singulares la imagen resultante es excelente. Aún con los primeros diez valores singulares podemos discernir que la imagen corresponde a un gato. Agregando más valores singulares es posible distinguir más detalles, como las franjas y manchas del gato, sus ojos, nariz y boca. Es posible observar que para $r = 10$ y $r = 20$ las manchas borrosas en el fondo detrás del gato, son más visibles y se van desvaneciendo conforme se incorporan más valores singulares.

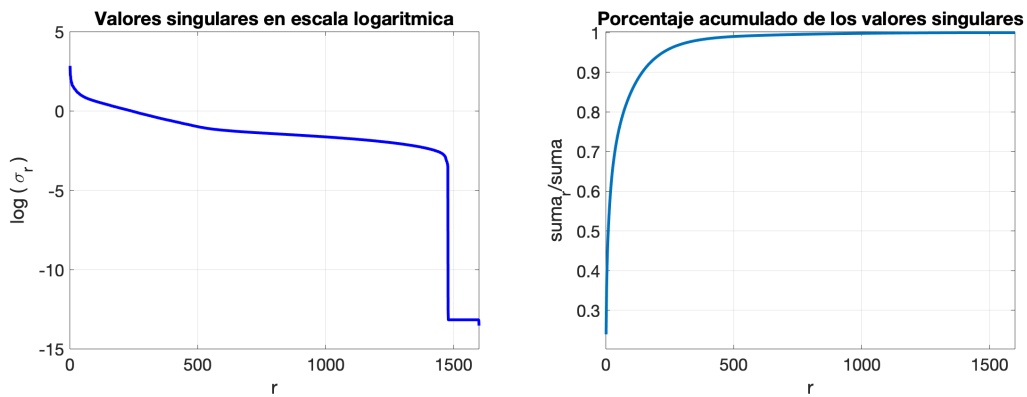


Figura 3.7: Valores singulares y su porcentaje acumulado.

3.5

Referencias

Existe una lista muy extensa de trabajos de investigación sobre la descomposición en valores singulares y aplicaciones interesantes. La mayoría de las referencias ya incluidas en los capítulos anteriores abordan en tema en sus diferentes aspectos y con mayor o menor profundidad. Principalmente, aunque no únicamente, los libros de Gene Golub, Lloyd Trefethen y Gilbert Strang. Las siguientes referencias agregan información que en mi opinión es pertinente incluir en estas notas.

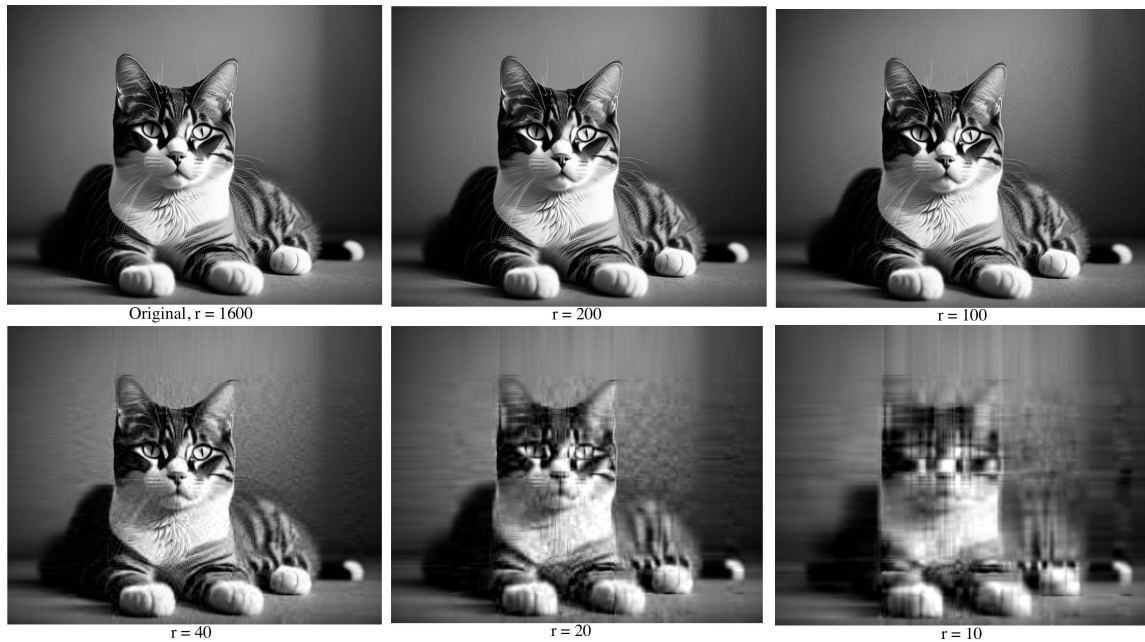


Figura 3.8: Imágenes de los gatos obtenidos con matrices de rango bajo .

1. G. W. Stewart, *On the Early History of the Singular Value Decomposition*, SIAM Review, Vol. 35, Issue 4, 1993.

Este artículo es un recuento de la contribución de cinco matemáticos por su responsabilidad en probar la existencia de la descomposición de valores singulares y desarrollar la teoría subyacente.

2. WH Press, SA Teukolsky, WT Vetterling, BP Flannery (2007), *Numerical Recipes: The Art of Scientific Computing* (3rd ed.), New York: Cambridge University Press, ISBN 978-0-521-88068-8.

Libro clásico, ampliamente reconocido como el libro de cómputo científico con la lista más completa, accesible y práctica de algoritmos más importantes. Contiene más de 400 rutinas, incluyendo la descomposición QR y SVD.

3. <https://github.com/mohammedAljadd/svd-image-compression>
Código para compresión de imágenes utilizando SVD con *Matlab*.
4. <https://dmicz.github.io/machine-learning/svd-image-compression>
Código para compresión de imágenes utilizando SVD con *Phyton*.

5. Michael W. Berry and Murray Browne, *Understanding Search Engines: Mathematical Modeling and Text Retrieval*, SIAM Book Series: Software, Environments, and Tools, Second Edition (May 2005), ISBN: 0-89871-581-4.

Un libro que incluye la modelación y algoritmos para obtener la recuperación de texto en las maquinas de búsqueda en internet, en particular *Google*. Muestra la aplicación de la factorización QR y la descomposición SVD, entre otras herramientas.