

NORMAS MATRICIALES

ALAN REYES-FIGUEROA
MÉTODOS NUMÉRICOS II

(AULA 01) 03.JULIO.2025

Notaciones y algunas propiedades

Siempre vamos a escribir un vector $\mathbf{x} \in \mathbb{R}^n$ en formato vertical. Cuando necesitemos escribirlo en formato horizontal, escribiremos $\mathbf{x}^T$:

$$\mathbf{x} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}, \quad \mathbf{x}^T = (x_1, x_2, \dots, x_n).$$

Usaremos $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ para representar los vectores de la base canónica de $\mathbb{R}^n$.

Para dos vectores $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$, escribimos su producto punto en la forma vectorial

$$\mathbf{x}^T \mathbf{y} = (x_1, x_2, \dots, x_n) \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix} = \sum_{j=1}^n x_j y_j.$$

Notaciones y algunas propiedades

Sea $A \in \mathbb{R}^{m \times n}$ una matriz. Podemos escribir A a partir de sus columnas

$$A = \begin{pmatrix} \mathbf{a}_1 & \dots & \mathbf{a}_n \end{pmatrix},$$

En este caso, si $\mathbf{x} \in \mathbb{R}^n$ es el vector con componentes

$$\mathbf{x}^T = (x_1, x_2, \dots, x_n),$$

entonces $\mathbf{x} = \sum_{j=1}^n x_j \mathbf{e}_j$ (en la base canónica).

Luego,

$$A\mathbf{x} = A\left(\sum_{j=1}^n x_j \mathbf{e}_j\right) = \sum_{j=1}^n x_j (A\mathbf{e}_j) = \sum_{j=1}^n x_j \mathbf{a}_j,$$

de modo que la imagen $A\mathbf{x}$ siempre es una combinación lineal de las columnas de A .

Conceptos de Álgebra Lineal

Normas de Vectores:

Definición

Una **norma** (de vectores) en $\mathbb{R}^n$ es una función $\|\cdot\| : \mathbb{R}^n \rightarrow \mathbb{R}$ que satisface las siguientes propiedades

1. $\|\mathbf{x}\| \geq 0$, $\forall \mathbf{x} \in \mathbb{R}^n$, y $\|\mathbf{x}\| = 0$ si y sólo si $\mathbf{x} = \mathbf{0}$.
2. $\|\mathbf{x} + \mathbf{y}\| \leq \|\mathbf{x}\| + \|\mathbf{y}\|$, $\forall \mathbf{x}, \mathbf{y} \in \mathbb{R}^n$.
3. $\|\alpha \mathbf{x}\| = |\alpha| \|\mathbf{x}\|$, $\forall \alpha \in \mathbb{R}$, $\mathbf{x} \in \mathbb{R}^n$.

Las normas más populares en $\mathbb{R}^n$ corresponden a la familia de normas- p ó p -**normas**. A continuación se muestran algunos ejemplos de p -normas, así como del disco unitario $\mathbb{D} = \{\mathbf{x} \in \mathbb{R}^2 : \|\mathbf{x}\| \leq 1\}$:

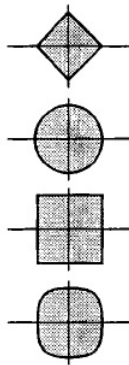
Conceptos de Álgebra Lineal

$$\|x\|_1 = \sum_{i=1}^m |x_i|,$$

$$\|x\|_2 = \left(\sum_{i=1}^m |x_i|^2 \right)^{1/2} = \sqrt{x^* x},$$

$$\|x\|_\infty = \max_{1 \leq i \leq m} |x_i|,$$

$$\|x\|_p = \left(\sum_{i=1}^m |x_i|^p \right)^{1/p} \quad (1 \leq p < \infty).$$



Aparte de las p -normas, otra de las más comunes es la familia de las **p -normas pesadas**, en donde en cada una de las coordenadas, las entradas se combinan mediante pesos.

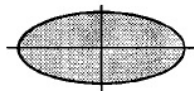
Conceptos de Álgebra Lineal

Estos pesos vienen dados por una matriz diagonal $W \in \mathbb{R}^{n \times n}$:

$$\|\mathbf{x}\|_W = \|W\mathbf{x}\|.$$

Por ejemplo, una 2-norma pesada en $\mathbb{R}^2$ se ve de la siguiente forma

$$\|x\|_W = \left(\sum_{i=1}^m |w_i x_i|^2 \right)^{1/2}.$$



Normas de matrices inducidas por normas vectoriales:

Una matriz $A \in \mathbb{R}^{m \times n}$ puede verse como un vector en $\mathbb{R}^{mn}$ mediante el mapa

$$\begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \dots & a_{mn} \end{pmatrix} \longrightarrow (a_{11}, \dots, a_{1n}, \dots, a_{m1}, \dots, a_{mn}) \in \mathbb{R}^{mn}.$$

Normas Matriciales

Tiene sentido entonces definir normas matriciales, para medir el “tamaño” de estas matrices.

Sin embargo, en los espacios de matrices, ciertas normas especiales son más útiles que las p -normas discutidas anteriormente. Estas son las normas inducidas, y se definen en términos del comportamiento de la matriz como operador $A : \mathbb{R}^n \rightarrow \mathbb{R}^m$.

Definición

Dadas normas $\|\cdot\|_{(n)}$ y $\|\cdot\|_{(m)}$ en el dominio y el rango de una matriz $A \in \mathbb{R}^{m \times n}$, decimos que $\|\cdot\|_{(m,n)}$ es la **norma inducida** por $\|\cdot\|_{(n)}$ y $\|\cdot\|_{(m)}$ si $\|\cdot\|_{(m,n)}$ es el menor número C que satisface

$$\|A\mathbf{x}\|_{(m)} \leq C\|\mathbf{x}\|_{(n)}, \quad \forall \mathbf{x} \in \mathbb{R}^n. \quad (1)$$

Normas Matriciales

Si $\mathbf{x} \neq \mathbf{0}$, podemos pasar el término $\|\mathbf{x}\|_{(n)}$ dividiendo al lado izquierdo de (1), de modo que

$$\frac{\|A\mathbf{x}\|_{(m)}}{\|\mathbf{x}\|_{(n)}} \leq C, \forall \mathbf{x} \in \mathbb{R}^n, \mathbf{x} \neq \mathbf{0}.$$

Así,

$$\|A\|_{(m,n)} = \sup_{\mathbf{x} \neq \mathbf{0}} \frac{\|A\mathbf{x}\|_{(m)}}{\|\mathbf{x}\|_{(n)}}, \quad (2)$$

de modo que $\|A\|_{(m,n)}$ es el mayor factor por el cual el operador A ampliado un vector. De la ecuación (3) en la definición de norma, la acción de A se determina por su acción sobre vectores unitarios. Luego (2) puede reducirse a

$$\|A\|_{(m,n)} = \sup_{\mathbf{x} \neq \mathbf{0}} \frac{\|A\mathbf{x}\|_{(m)}}{\|\mathbf{x}\|_{(n)}} = \sup_{\|\mathbf{x}\|_{(n)}=1} \|A\mathbf{x}\|_{(m)}. \quad (3)$$

Normas Matriciales

En el caso particular que $A \in \mathbb{R}^{m \times n}$, entonces la **p -norma matricial, inducida por la p -norma** en $\mathbb{R}^n$ es

$$\|A\|_p = \sup_{\mathbf{x} \neq \mathbf{0}} \frac{\|A\mathbf{x}\|_p}{\|\mathbf{x}\|_p} = \sup_{\|\mathbf{x}\|_p=1} \|A\mathbf{x}\|_p. \quad (4)$$

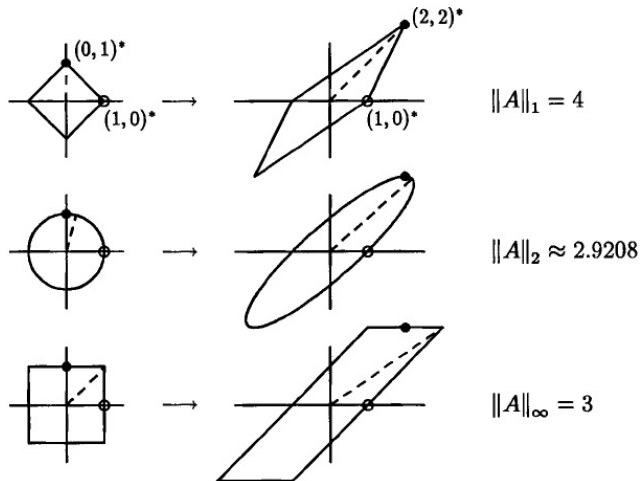
Ejemplo: Considere la matriz

$$A = \begin{pmatrix} 1 & 2 \\ 0 & 2 \end{pmatrix},$$

la cual corresponde a una transformación lineal $A : \mathbb{R}^2 \rightarrow \mathbb{R}^2$. (También la podemos ver como un mapa $A : \mathbb{C}^2 \rightarrow \mathbb{C}^2$, pero $\mathbb{R}$ es más conveniente para hacer ilustraciones).

La siguiente imagen muestra el resultado de algunas p -normas para A :

Normas Matriciales



1-norma de una matriz

Teorema

Sea $A \in \mathbb{R}^{m \times n}$, y sean $\mathbf{a}_1, \dots, \mathbf{a}_n \in \mathbb{R}^m$ sus columnas. Entonces la 1-norma inducida de A es

$$\|A\|_1 = \max_{1 \leq j \leq n} \|\mathbf{a}_j\|_1.$$

Prueba: Sea $A \in \mathbb{R}^{m \times n}$. Escribimos A en sus columnas

$$A = \begin{pmatrix} \mathbf{a}_1 & \dots & \mathbf{a}_n \end{pmatrix},$$

donde cada $\mathbf{a}_j \in \mathbb{R}^m$ es un vector. Considere la bola unitaria en la 1-norma, dada por $B_1(\mathbf{0}) = \{\mathbf{x} \in \mathbb{R}^n : \sum_{i=1}^n |x_i| \leq 1\}$. Entonces, de la ecuación (4), para todo

$\mathbf{x} = (x_1, \dots, x_n)^T \in \mathbb{R}^n$, se tiene

1-norma de una matriz

$$\|A\mathbf{x}\|_1 = \left\| \sum_{j=1}^n x_j \mathbf{a}_j \right\|_1 \leq \sum_{j=1}^n |x_j| \|\mathbf{a}_j\|_1 \leq \sum_{j=1}^n |x_j| \max_{1 \leq j \leq n} \|\mathbf{a}_j\|_1 \leq \max_{1 \leq j \leq n} \|\mathbf{a}_j\|_1, \quad \forall \|\mathbf{x}\|_1 \leq 1.$$

Esto muestra que $\sup_{\|\mathbf{x}\|_1=1} \|A\mathbf{x}\|_1 \leq \max_{1 \leq i \leq n} \|\mathbf{a}_i\|_1$.

Sea $i = \arg \max_{1 \leq j \leq n} \|\mathbf{a}_j\|_1$, esto es $\|\mathbf{a}_i\|_1 = \max_{1 \leq j \leq n} \|\mathbf{a}_j\|_1$.

Tomando el vector $\mathbf{x} = \mathbf{e}_i$, obtenemos que

$$\|A\mathbf{x}\|_1 = \|A\mathbf{e}_i\|_1 = \|\mathbf{a}_i\|_1 = \max_{1 \leq j \leq n} \|\mathbf{a}_j\|_1,$$

y el supremo se alcanza en el vector unitario $\mathbf{e}_i$. Luego, $\|A\|_1 = \max_{1 \leq j \leq n} \|\mathbf{a}_j\|_1$. $\square$

∞ -norma de una matriz

Teorema

Sea $A \in \mathbb{R}^{m \times n}$, y sean $\mathbf{a}_1^T, \dots, \mathbf{a}_m^T \in \mathbb{R}^n$ sus filas. Entonces la ∞ -norma inducida de A es

$$\|A\|_{\infty} = \max_{1 \leq i \leq m} \|\mathbf{a}_i^T\|_1.$$

Prueba: Si ahora denotamos A por sus filas, $A = \begin{pmatrix} \mathbf{a}_1^T \\ \vdots \\ \mathbf{a}_m^T \end{pmatrix}$, $\mathbf{a}_i \in \mathbb{R}^n$. Un argumento muy similar al anterior se utiliza para mostrar que

$$\|A\|_{\infty} = \sup_{\|\mathbf{x}\|_{\infty}=1} \|A\mathbf{x}\|_{\infty} = \max_{1 \leq i \leq m} \|\mathbf{a}_i^T\|_1. \quad \square$$

Normas Matriciales

Calcular p -normas matriciales, con $p \neq 1, \infty$ es más difícil. Para ello, observemos que los productos internos acotan el producto de normas vectoriales.

Sean $1 \leq p, q \leq \infty$, tales que $\frac{1}{p} + \frac{1}{q} = 1$. Tenemos las siguientes desigualdades:

$$|\mathbf{x}^T \mathbf{y}| \leq \|\mathbf{x}\|_p \|\mathbf{y}\|_q \quad (\text{desigualdad de Hölder})$$

$$|\mathbf{x}^T \mathbf{y}| \leq \|\mathbf{x}\|_2 \|\mathbf{y}\|_2 \quad (\text{desigualdad de Cauchy-Schwarz})$$

Con esta información, vamos a calcular la

2-norma de una matriz fila:

Sea $A \in \mathbb{R}^{1 \times n}$ una matriz de una sola fila. Podemos escribir $A = \mathbf{a}^T$, con $\mathbf{a}^T \in \mathbb{R}^n$.

Entonces, de la desigualdad de Cauchy-Schwarz

$$\|A\mathbf{x}\|_2 = \|\mathbf{a}^T \mathbf{x}\|_2 \leq \|\mathbf{a}\|_2 \|\mathbf{x}\|_2,$$

$$\text{luego } \sup_{\|\mathbf{x}\|_2=1} \frac{\|A\mathbf{x}\|_2}{\|\mathbf{x}\|_2} \leq \|\mathbf{a}\|_2.$$

Esta cota se alcanza. Haciendo $\mathbf{x} = \mathbf{a}$, se tiene que $\|A\mathbf{a}\|_2 = \|\mathbf{a}\|_2^2$, y se tiene que $\|A\|_2 = \|\mathbf{a}\|_2$.

2-norma de un producto exterior:

Sea $A \in \mathbb{R}^{m \times n}$ una matriz de rango 1, digamos $A = \mathbf{u}\mathbf{v}^T$, con $\mathbf{u} \in \mathbb{R}^m$, $\mathbf{v} \in \mathbb{R}^n$. Para todo $\mathbf{x} \in \mathbb{R}^n$,

$$\|A\mathbf{x}\|_2 = \|\mathbf{u}\mathbf{v}^T \mathbf{x}\|_2 = \|\mathbf{u}\|_2 \|\mathbf{v}^T \mathbf{x}\|_2 \leq \|\mathbf{u}\|_2 \|\mathbf{v}\|_2 \|\mathbf{x}\|_2.$$

De ahí que $\|A\|_2 \leq \|\mathbf{u}\|_2 \|\mathbf{v}\|_2$. De nuevo, esta desigualdad se alcanza haciendo $\mathbf{x} = \mathbf{v}$, de modo que $\|A\|_2 = \|\mathbf{u}\|_2 \|\mathbf{v}\|_2$.

Proposición

Sean $A \in \mathbb{R}^{\ell \times m}$, $B \in \mathbb{R}^{m \times n}$ matrices, y sean $\|\cdot\|_{(\ell)}$, $\|\cdot\|_{(m)}$, $\|\cdot\|_{(n)}$, normas matriciales en $\mathbb{R}^{\ell}$, $\mathbb{R}^m$, $\mathbb{R}^n$, respectivamente.

Entonces

$$\|AB\|_{(\ell,n)} \leq \|A\|_{(\ell,m)} \|B\|_{(m,n)}.$$

Prueba:

Por definición, para cualquier $\mathbf{x} \in \mathbb{R}^n$ vale

$$\|AB\mathbf{x}\|_{(\ell)} \leq \|A\|_{(\ell,m)} \|B\mathbf{x}\|_{(m)} \leq \|A\|_{(\ell,m)} \|B\|_{(m,n)} \|\mathbf{x}\|_{(n)}.$$

Luego,

$$\|AB\|_{(\ell,n)} = \sup_{\|\mathbf{x}\|_{(n)} \neq 0} \frac{\|AB\mathbf{x}\|_{(\ell)}}{\|\mathbf{x}\|_{(n)}} \leq \|A\|_{(\ell,m)} \|B\|_{(m,n)}. \quad \square$$

Normas Matriciales

Esta desigualdad no es una igualdad. Por ejemplo, para matrices cuadradas vale en general $\|A^n\|_p \leq \|A\|_p^n$, para todo $n \geq 1$.

Sin embargo, no se cumple en general que $\|A^n\|_p = \|A\|_p^n$, para $n \geq 2$.

Ejemplo:

Para $A = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$, $p = 1$, tenemos que $A^n = \begin{pmatrix} 1 & n \\ 0 & 1 \end{pmatrix}$, $\forall n \in \mathbb{N}$. Luego,

$$\|A^n\|_1 = n + 1, \quad \|A\|_1^n = 2^n.$$

Ejemplo

Ejemplo: Calcular la norma 1 inducida de $A = \begin{pmatrix} 1 & 2 & 3 \\ 0 & -3 & 1 \end{pmatrix}$.

En este caso, tenemos tres vectores columna

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

Las normas 1 de estos vectores columna son, respectivamente,

$$\|\mathbf{a}_1\|_1 = |1| + |0| = 1, \quad \|\mathbf{a}_2\|_1 = |2| + |-3| = 5, \quad \|\mathbf{a}_3\|_1 = |3| + |1| = 4.$$

Finalmente, el máximo de estas normas 1 vectoriales es 6, y tenemos

$$\|A\|_1 = \max_{1 \leq j \leq 3} \|\mathbf{a}_j\|_1 = \max\{1, 5, 4\} = 5.$$

Ejemplo

Ejemplo: Calcular la norma ∞ inducida de $A = \begin{pmatrix} 1 & 2 & 3 \\ 0 & -3 & 1 \end{pmatrix}$.

En este caso, tenemos dos vectores fila

$$\mathbf{a}_1^T = (1 \ 2 \ 3), \quad \mathbf{a}_2^T = (0 \ -3 \ 1).$$

Las normas 1 de estos vectores fila son, respectivamente,

$$\|\mathbf{a}_1^T\|_1 = |1| + |2| + |3| = 6, \quad \|\mathbf{a}_2^T\|_1 = |0| + |-3| + |1| = 4.$$

Finalmente, el máximo de estas normas 1 vectoriales es 6, y tenemos

$$\|A\|_\infty = \max_{1 \leq i \leq 2} \|\mathbf{a}_i^T\|_1 = \max\{6, 4\} = 6.$$

Teorema

Sea $A \in \mathbb{R}^{m \times n}$ una matriz, y sean $\lambda_1, \lambda_2, \dots, \lambda_m$ los autovalores de $A^T A \in \mathbb{R}^{m \times m}$. Entonces

$$\|A\|_2 = \left(\max_{1 \leq i \leq m} |\lambda_i| \right)^{1/2} = \sqrt{\max_{1 \leq i \leq m} \lambda_i}.$$

Prueba: La prueba se dará más adelante.

Observe que:

- $A^T A$ es una matriz cuadrada simétrica, portanto sus autovalores son siempre números reales.
- Casi siempre ocurre que $A^T A$ es semi-definida positiva, de modo que sus autovalores son ≥ 0 .
- $\max_{1 \leq i \leq m} \lambda_i$ es el mayor autovalor de $A^T A$. Este se llama el **radio espectral** de $A^T A$.

Ejemplo

Ejemplo: Calcular la norma 2 inducida para $A = \begin{pmatrix} 1 & 2 \\ 0 & 2 \end{pmatrix}$.

En este caso, comenzamos calculando la matriz $A^T A$

$$A^T A = \begin{pmatrix} 1 & 0 \\ 2 & 2 \end{pmatrix} \begin{pmatrix} 1 & 2 \\ 0 & 2 \end{pmatrix} = \begin{pmatrix} 1 & 2 \\ 2 & 8 \end{pmatrix}.$$

Luego, el polinomio característico de $A^T A$ es

$$\det(\lambda I - A^T A) = \lambda^2 - \operatorname{tr}(A^T A) \lambda + \det(A^T A) = \lambda^2 - 9\lambda + 4 = 0,$$

cuyas raíces son $\lambda_{\max} = \frac{9 + \sqrt{65}}{2}$ y $\lambda_{\min} = \frac{9 - \sqrt{65}}{2}$.

Luego, $\|A\|_2 = \sqrt{\lambda_{\max}} = \sqrt{\frac{9 + \sqrt{65}}{2}} \approx 2.92080962648$.

Ejemplo

Ejemplo: Calcular la norma 2 inducida para $A = \begin{pmatrix} 1 & 2 & 3 \\ 0 & -3 & 1 \end{pmatrix}$.

Respuesta: $\|A\|_2 \approx 3.95038622$ (Ejercicio!)

Normas Matriciales

En general, no todas las normas matriciales son inducidas por una norma vectorial. Damos la siguiente definición general.

Definición

Una **norma matricial** es una función $\|\cdot\| : \mathbb{R}^{m \times n} \rightarrow \mathbb{R}$ que satisface:

1. $\|A\| \geq 0, \forall A \in \mathbb{R}^{m \times n}$.
2. $\|A + B\| \leq \|A\| + \|B\|, \forall A, B \in \mathbb{R}^{m \times n}$.
3. $\|\alpha A\| = |\alpha| \|A\|, \forall \alpha \in \mathbb{R}, \forall A \in \mathbb{R}^{m \times n}$.

La norma matricial más importante que no proviene de una norma vectorial es la **norma de Frobenius** o **norma de Hilbert-Schmidt**

$$\|A\|_F = \left(\sum_{i=1}^m \sum_{j=1}^n |a_{ij}|^2 \right)^{1/2}. \quad (5)$$

Normas Matriciales

Si $\mathbf{a}_1, \dots, \mathbf{a}_n \in \mathbb{R}^m$ denotan las columnas de A , entonces

$$\|A\|_F = \left(\sum_{j=1}^n \|\mathbf{a}_j\|_2^2 \right)^{1/2}. \quad (6)$$

Proposición

Para $A \in \mathbb{R}^{m \times n}$, vale $\|A\|_F = \sqrt{\text{tr}(A^T A)} = \sqrt{\text{tr}(A A^T)}$.

Prueba: Si $A = (\mathbf{a}_1 \quad \mathbf{a}_2 \quad \dots \quad \mathbf{a}_n)$ son las columnas de A , $\mathbf{a}_j \in \mathbb{R}^m$, entonces

$$A^T A = \begin{pmatrix} \mathbf{a}_1^T \\ \vdots \\ \mathbf{a}_n^T \end{pmatrix} \begin{pmatrix} \mathbf{a}_1 & \dots & \mathbf{a}_n \end{pmatrix} = (\mathbf{a}_i^T \mathbf{a}_j).$$

$$\text{Luego, } \text{tr}(A^T A) = \sum_{j=1}^n \mathbf{a}_j^T \mathbf{a}_j = \sum_{j=1}^n \|\mathbf{a}_j\|_2^2 = \sum_{j=1}^n \sum_{i=1}^m a_{ij}^2 = \sum_{i=1}^m \sum_{j=1}^n |a_{ij}|^2 = \|A\|_F^2. \quad \square$$

Ejemplo

Ejemplo: Calcular la norma de Frobenius para $A = \begin{pmatrix} 1 & 2 & 3 \\ 0 & -3 & 1 \end{pmatrix}$.

Tenemos que

$$\sum_{i=1}^2 \sum_{j=1}^3 a_{ij}^2 = 1^2 + 2^2 + 3^2 + 0^2 + (-3)^2 + 1^2 = 24.$$

Luego, $\|A\|_F = \sqrt{24} = 4.89897948$.

Otra manera de calcular la norma de Frobenius sería usando $\|A\|_F = \sqrt{\text{tr}(A^T A)}$. Observe que

$$A^T A = \begin{pmatrix} 1 & 2 & 3 \\ 2 & 13 & 3 \\ 3 & 3 & 10 \end{pmatrix} \implies \text{tr}(A^T A) = 1 + 13 + 10 = 24.$$

Portanto, $\|A\|_F = \sqrt{24} = 4.89897948$.

¿Para qué sirven las normas matriciales?

Tiene sentido entonces definir normas matriciales, para medir el “tamaño” de estas matrices.

Recordemos que en el contexto de métodos numéricos, las normas vectoriales (y matriciales nos sirven para calcular distancias entre vectores (o matrices).

$$\text{dist}(\mathbf{x}, \mathbf{y}) = \|\mathbf{y} - \mathbf{x}\|_p.$$

y las distancias nos dan un mecanismo para medir el error de aproximación. Por ejemplo, si $\mathbf{x} \in \mathbb{R}^n$ es la solución de un problema, y tenemos un algoritmo numérico que nos devuelve una solución aproximada $\hat{\mathbf{x}} \in \mathbb{R}^n$, entonces el error de aproximación es

$$\text{error de aproximación} = \|\mathbf{x} - \hat{\mathbf{x}}\|_p.$$

Lo mismo ocurre para un algoritmo que calcula matrices:

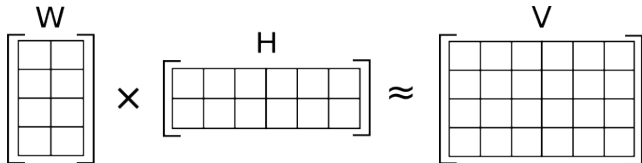
$$\text{error de aproximación} = \|A - \hat{A}\|_p.$$

Aplicación

La factoración de matrices es una estrategia muy usada para obtener variables latentes, útiles para revelar información no observable dentro de una matriz de datos $\mathbb{X}$ (e.g. un sistema de recomendación).

Sea $\mathbb{X} \in \mathbb{R}^{n \times d}$ una matriz de datos. El objetivo es descomponer $\mathbb{X}$ como el producto de dos matrices con entradas no-negativas $W \in \mathbb{R}^{n \times r}$ y $H \in \mathbb{R}^{r \times d}$

$$\mathbb{X} = WH,$$



La solución de este algoritmo consiste en encontrar la combinación de matrices $\hat{W}$ y $\hat{H}$ que minimiza el error de aproximación $\|\mathbb{X} - \hat{W}\hat{H}\|_F$.