

Capítulo 1

El plano euclidiano

1.1 La geometría griega

En el principio, la geometría era una colección de reglas de uso común para medir y construir casas y ciudades. No fue hasta el año 300 a.C. que Euclides de Alejandría, en sus *Elementos*, ordenó y escribió todo ese saber, imprimiéndole el sello de rigor lógico que caracteriza y distingue a las matemáticas. Se dio cuenta de que todo razonamiento riguroso (o *demostración*) debe basarse en ciertos principios previamente establecidos ya sea, a su vez, por otra demostración o bien por convención. Pero a final de cuentas, este método conduce a la necesidad ineludible de convenir en que ciertos principios básicos (*postulados* o *axiomas*) son válidos sin necesidad de demostrarlos, que están dados y son incontrovertibles para poder construir sobre ellos el resto de la teoría. Lo que hoy se conoce como Geometría Euclidiana, y hasta hace dos siglos simplemente como Geometría, está basada en los cinco postulados de Euclides:

- I Por cualesquiera dos puntos, se puede trazar el segmento de recta que los une.
- II Dados un punto y una distancia, se puede trazar el círculo con centro en el punto y cuyo radio es la distancia.
- III Un segmento de recta se puede extender en ambas direcciones indefinidamente.
- IV Todos los ángulos rectos son iguales.
- V Dadas dos rectas y una tercera que las corta, si los ángulos internos de un lado suman menos de dos ángulos rectos, entonces las dos rectas se cortan y lo hacen de ese lado.

Obsérvese que en estos postulados se describe el comportamiento y la relación entre ciertos elementos básicos de la geometría, como son “punto”, “trazar”, “segmento”, “distancia”, etc. De alguna manera, se le da la vuelta a su definición haciendo uso de

la noción intuitiva que se tiene de ellos y haciendo explícitas ciertas relaciones básicas que deben cumplir.

De estos postulados o axiomas, el quinto es el más famoso pues se creía que podría deducirse de los otros. Es más sofisticado, y entonces se pensó que debía ser un *teorema*, es decir, ser demostrable. De hecho, un problema muy famoso fue ése: demostrar el quinto postulado usando únicamente los otros cuatro. Y muchísimas generaciones de matemáticos lo atacaron sin éxito. No es sino hasta el siglo XIX, y para sorpresa de todos, que se demuestra que no se puede demostrar; que efectivamente hay que suponerlo como axioma, pues negaciones de él dan origen a “nuevas geometrías”, tan lógicas y bien fundamentadas como la euclidiana. Pero esta historia se verá más adelante (en los capítulos 6 y 8) pues por el momento nos interesa introducir la geometría analítica.

La publicación de la “*Géométrie*” de Descartes marca una revolución en las matemáticas. La introducción del álgebra a la solución de problemas de índole geométrico desarrolló una nueva intuición y, con ésta, un nuevo entender de la naturaleza de las “*lignes courbes*”.

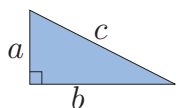
Para comprender lo que hizo René Descartes (1596–1650), se debe tener más familiaridad con el quinto postulado. Además del enunciado original, hay otras dos versiones que son equivalentes:

V.a (El Quinto) Dada una línea recta y un punto fuera de ella, existe una única recta que pasa por el punto y que es paralela a la línea.

V.b Los ángulos interiores de un triángulo suman dos ángulos rectos.

De las tres versiones que hemos dado del quinto postulado de Euclides, usaremos a lo largo de este libro a la **V.a**, a la cual nos referiremos simplemente como “*El Quinto*”.

Antes de entrar de lleno a la Geometría Analítica demostremos, a manera de homenaje a los griegos y con sus métodos, uno de sus teoremas más famosos e importantes.

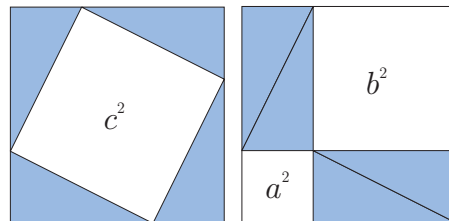


Teorema 1.1 (de Pitágoras) Dado un triángulo rectángulo, si los lados que se encuentran en un ángulo recto (llamados catetos) miden a y b , y el tercero (la hipotenusa) mide c , entonces

$$a^2 + b^2 = c^2.$$

Demostración. Considérese un cuadrado de lado $a + b$, y colóquense en él cuatro copias del triángulo de dos maneras diferentes como en la figura. Las áreas que quedan fuera son iguales. $\square$

Aunque la demostración anterior (hay otras) parece no usar nada, queda implícito el uso del quinto postulado; pues, por ejemplo, en el primer cuadrado, que el cuadradito interno de lado c tenga ángulo recto usa su versión **V.b**. Además, hace uso de nuevos conceptos como son el área y cómo calcularla; en fin, hace uso de nuestra intuición geométrica, que debemos creer y fomentar. Pero vayamos a lo nuestro.



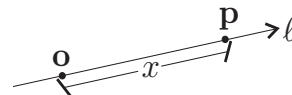
***EJERCICIO 1.1** Demuestra la equivalencia de **V**, **V.a** y **V.b**. Aunque no muy formalmente (como en nuestra demostración del Teorema de Pitágoras), convencerse con dibujos de que tienen todo que ver entre ellos.

EJERCICIO 1.2 Demuestra el Teorema de Pitágoras ajustando (sin que se traslapen) cuatro copias del triángulo en un cuadrado de lado c y usando que $(b-a)^2 = b^2 - 2ab + a^2$ (estamos suponiendo que $a < b$, como en la figura anterior).

1.2 Puntos y parejas de números

Para reinterpretar el razonamiento que hizo Descartes, supongamos por un momento que conocemos bien el *plano euclidiano* definido por los cinco axiomas de los *Elementos*; es el conjunto de puntos que se extiende indefinidamente a semejanza de un pizarrón, un papel, un piso o una pared y viene acompañado de nociones como rectas y distancia, entre otras; y en el cual se pueden demostrar teoremas como el de Pitágoras. Denotaremos por $\mathbb{E}^2$ a este plano; en este caso el exponente 2 se refiere a la dimensión y no es exponenciación (no es “ $\mathbb{E}$ al cuadrado”, sino que debe leerse “ $\mathbb{E}$ dos”). Da la idea de que puede haber espacios euclidianos de cualquier dimensión $\mathbb{E}^n$ (léase “ $\mathbb{E}$ ene”), aunque sería complicado definirlos axiomáticamente. De hecho los definiremos, pero de una manera más simple, que es con el uso de coordenadas: la idea genial de Descartes. Podríamos ahora, apoyados en el lenguaje moderno de los conjuntos,¹ recrear su razonamiento como sigue.

El primer paso es notar que los puntos de una recta $\ell \subset \mathbb{E}^2$ corresponden biunívocamente a los *números reales*, que se denotan por $\mathbb{R}$. Escogemos un punto $\mathbf{o} \in \ell$ que se llamará el *origen* y que corresponderá al número cero. El origen parte a la recta en dos mitades; a una de ellas se le asocian los números positivos de acuerdo con su distancia al origen, y a la otra mitad, los números negativos. Así, a cada número real x se le asocia un punto $\mathbf{p} \in \ell$, y a cada punto en ℓ le corresponde un número real.



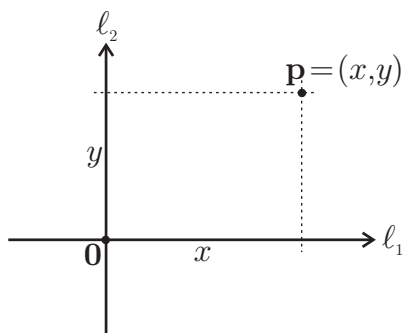
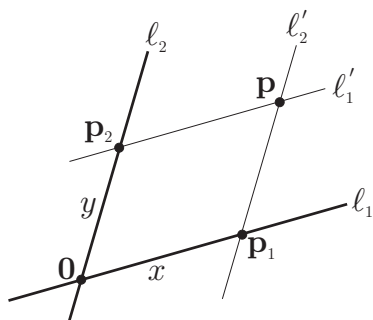
¹La notación básica de conjuntos se establece e introduce brevemente en el Apéndice.

De esta identificación natural surge el nombre de “recta de los números reales” y todo el desarrollo ulterior de la Geometría Analítica, e inclusive del Cálculo. Los números tienen un significado geométrico (los griegos lo sabían bien al entenderlos como distancias), y entonces problemas geométricos (y físicos) pueden atacarse y entenderse manipulando números.

El segundo paso para identificar los puntos del plano euclidiano con parejas de números hace uso esencial del quinto postulado. Tómense dos rectas ℓ_1 y ℓ_2 en el plano $\mathbb{E}^2$ que se intersecten en un punto $\mathbf{o}$ que será el *origen*. Orientamos las rectas para que sus puntos correspondan a números reales, como arriba. Entonces a cualquier punto $\mathbf{p} \in \mathbb{E}^2$ se le puede hacer corresponder una pareja de números de la siguiente

forma. Por el Quinto, existe una única recta ℓ'_1 que pasa por $\mathbf{p}$ y es paralela a ℓ_1 ; y análogamente, existe una única recta ℓ'_2 que pasa por $\mathbf{p}$ y es paralela a ℓ_2 . Las intersecciones $\ell_1 \cap \ell'_2$ y $\ell_2 \cap \ell'_1$ determinan los puntos $\mathbf{p}_1 \in \ell_1$ y $\mathbf{p}_2 \in \ell_2$, respectivamente, y por tanto, determinan dos números reales x y y ; es decir, una pareja ordenada (x, y) . Y al revés, dada la pareja de números (x, y) , tómense $\mathbf{p}_1$ como el punto que está sobre ℓ_1 a distancia x de $\mathbf{o}$, y $\mathbf{p}_2$ como el punto que está sobre ℓ_2 a distancia y de $\mathbf{o}$ (tomando en cuenta signos, por supuesto). Sea ℓ'_1 la recta que pasa por $\mathbf{p}_2$ paralela a ℓ_1 y sea

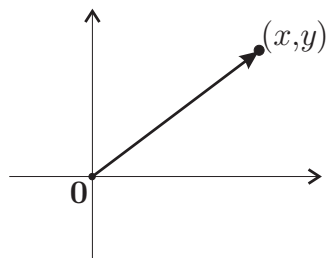
ℓ'_2 la recta que pasa por $\mathbf{p}_1$ paralela a ℓ_2 (¡acabamos de usar de nuevo al Quinto!). La intersección $\ell'_1 \cap \ell'_2$ es el punto $\mathbf{p}$ que corresponde a la pareja (x, y) .



La generalidad con que hicimos el razonamiento anterior fue para hacer explícito el uso del quinto postulado en el razonamiento clásico de Descartes, pero no conviene dejarlo tan ambiguo. Es costumbre tomar a las rectas ℓ_1 y ℓ_2 ortogonales: a ℓ_1 horizontal con su dirección positiva a la derecha (como vamos leyendo), conocida tradicionalmente como el *eje de las x*; y a ℓ_2 vertical con dirección positiva hacia arriba, el famoso *eje de las y*. No sólo es por costumbre que se utiliza esta convención, sino también porque simplifica el álgebra asociada a la geometría clásica.

Por ejemplo, permitirá usar el teorema de Pitágoras para calcular fácilmente las distancias a partir de las coordenadas. Pero esto se verá más adelante.

En resumen, al fijar los ejes coordenados, a cada pareja de números (x, y) le corresponde un punto $\mathbf{p} \in \mathbb{E}^2$, pero iremos más allá y los identificaremos diciendo $\mathbf{p} = (x, y)$. Además, se le puede hacer corresponder la flecha (segmento de recta dirigido llamado *vector*) que “nace” en el origen y termina en el punto. Así, las siguientes interpretaciones son equivalentes y se usarán de manera indistinta a lo largo de todo el libro.



- a. Un punto en el plano.
- b. Una pareja ordenada de números reales.
- c. Un vector que va del origen al punto.

Nótese que si bien en este texto las palabras *punto* y *vector* son equivalentes, tienen diferentes connotaciones. Mientras que un punto es pensado como un lugar (una posición) en el espacio, un vector es pensado como un segmento de línea dirigido de un punto a otro.

El conjunto de todas las parejas ordenadas de números se denota por $\mathbb{R}^2 = \mathbb{R} \times \mathbb{R}$. En este caso el exponente 2 sí se refiere a exponenciación, pero de conjuntos, lo que se conoce ahora como el *producto cartesiano* en honor de Descartes (véase el Apéndice A); aunque la usanza común es leerlo “ $\mathbb{R}$ dos”.

1.2.1 Geometría analítica

La idea genialmente simple $\mathbb{R}^2 = \mathbb{E}^2$ (es decir, las parejas ordenadas de números reales se identifican naturalmente con los puntos del plano euclidiano) logra que converjan las aguas del álgebra y la geometría. A partir de Descartes se abre la posibilidad de reconstruir la geometría de los griegos sobre la base de nuestra intuición numérica (básicamente la de sumar y multiplicar). Y a su vez, la luz de la geometría baña con significados y problemas al álgebra. Se hermanan, se apoyan, se entrelazan: nace la geometría analítica.

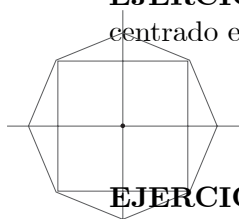
De aquí en adelante, salvo en contadas ocasiones donde sea inevitable y como comentarios, abandonaremos la línea del desarrollo histórico. De hecho, nuestro tratamiento algebraico está muy lejos del original de Descartes, pues usa ideas desarrolladas en el siglo XIX. Pero antes de entrar de lleno a él, vale la pena enfatizar que no estamos abandonando el método axiomático iniciado por Euclides, simplemente cambiaremos de conjunto de axiomas: ahora nos basaremos en los axiomas de los números reales (ver la sección 1.3.1). De tal manera que los objetos primarios de Euclides (línea, segmento, distancia, ángulo, etc.) serán definiciones (por ejemplo, *punto* ya es “pareja ordenada de números” y *plano* es el conjunto de todos los puntos), y los cinco postulados serán teoremas que hay que demostrar. Nuestros objetos primarios básicos serán ahora los números reales y nuestra intuición geométrica primordial es que forman una recta: La Recta Real, $\mathbb{R}$. Sin embargo, la motivación básica para las definiciones sigue siendo la intuición geométrica desarrollada por los griegos. No queremos hacernos los occisos como si no conociéramos la geometría

griega; simplemente la llevaremos a nuevos horizontes con el apoyo del álgebra —el enfoque analítico— y para ello tendremos que reconstruirla.

EJERCICIO 1.3 Encuentra diferentes puntos en el plano dados por sus coordenadas (por ejemplo, el $(3, 2)$, el $(-1, 5)$, el $(0, -4)$, el $(-1, -2)$, el $(1/2, -3/2)$).

EJERCICIO 1.4 Identifica los *cuadrantes* donde las parejas tienen signos determinados.

EJERCICIO 1.5 ¿Cuáles son las coordenadas de los vértices de un cuadrado de lado 2, centrado en el origen y con lados paralelos a los ejes?



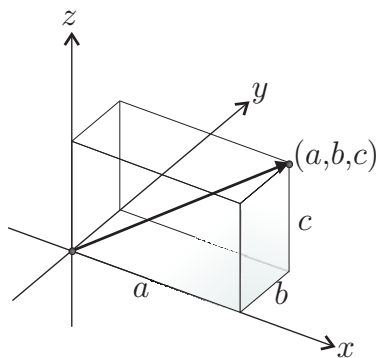
EJERCICIO 1.6 ¿Cuáles son las coordenadas de los vértices del octágono regular que incluye como vértices a los del cuadrado anterior? (*Tienes que usar el Teorema de Pitágoras.*)

EJERCICIO 1.7 ¿Puedes dar las coordenadas de los vértices de un triángulo equilátero centrado en el origen? Dibújalo. (*Usa de nuevo el Teorema de Pitágoras, en un triángulo rectángulo con hipotenusa 2 y un cateto 1.*)

1.2.2 El espacio de dimensión n

La posibilidad de extender con paso firme la geometría euclidiana a más de dos dimensiones es una de las aportaciones de mayor profundidad del método cartesiano.

Aunque nos hayamos puesto formales en la sección anterior para demostrar una correspondencia natural entre $\mathbb{E}^2$ y $\mathbb{R}^2$, intuitivamente es muy simple. Al fijar los



dos ejes coordenados —las dos direcciones preferidas— se llega de manera inequívoca a cualquier punto en el plano dando las distancias que hay que recorrer en esas direcciones para llegar a él. En el espacio, habrá que fijar tres ejes coordenados y dar tres números. Los dos primeros dan un punto en el plano (que podemos pensar horizontal, como el piso), y el tercero nos da la altura (que puede ser positiva o negativa). Si denotamos por $\mathbb{R}^3$ al conjunto de todas las ternas (x, y, z) de números reales, como éstas corresponden a puntos en el espacio una vez que se fijan los tres ejes (el eje vertical se conoce como *eje z*), podemos definir el espacio

euclidiano de tres dimensiones como $\mathbb{R}^3$, sin preocuparnos de axiomas. ¿Y por qué pararse en 3? ¿o en 4?

Dado un número natural $n \in \{1, 2, 3, \dots\}$, denotamos por $\mathbb{R}^n$ al conjunto de todas las n -adas (léase “eneadas”) ordenadas de números reales $(x_1, x_2, \dots, x_n)$. Formalmente,

$$\mathbb{R}^n := \{(x_1, x_2, \dots, x_n) \mid x_i \in \mathbb{R}, i = 1, 2, \dots, n\}.$$

Para valores pequeños de n , se pueden usar letras x, y, z e inclusive w para denotar las coordenadas; pero para n general esto es imposible y no nos queda más que

usar los subíndices. Así, tenemos un conjunto bien definido, $\mathbb{R}^n$, al que podemos llamar “*espacio euclidiano de dimensión n* ” y hacer geometría en él. A una n -ada $\mathbf{x} = (x_1, x_2, \dots, x_n) \in \mathbb{R}^n$ se le llama *vector* o *punto*.

El estudiante que se sienta incómodo con esto de “muchas dimensiones”, puede sustituir (pensar en) un 2 o un 3 siempre que se use n y referirse al plano o al espacio tridimensional que habitamos para su intuición. No pretendemos en este libro estudiar la geometría de estos espacios de muchas dimensiones. Nos concentraremos en dimensiones 2 y 3, así que puede pensarse la n como algo que puede tomar valores 2 o 3 y que resulta muy útil para hablar de los dos al mismo tiempo pero en singular. Sin embargo, es importante que el estudiante tenga una visión amplia y pueda preguntarse en cada momento ¿qué pasaría en 4 o más dimensiones? Como verá, y a veces señalaremos explícitamente, algunas cosas son muy fáciles de generalizar, y otras no tanto. Por el momento, basta con volver a enfatizar el amplísimo mundo que abrió el método de las coordenadas cartesianas para las matemáticas y la ciencia.

Vale la pena en este momento puntualizar las convenciones de notación que utilizaremos en este libro. A los números reales los denotamos por letras simples, por ejemplo x , y , z , o bien a , b y c o bien t , r y s ; e inclusive conviene a veces utilizar letras griegas, α (alfa), β (beta) y γ (gamma), o λ (lambda) y μ (mu) (por alguna razón hemos hecho costumbre de utilizarlas siempre en pequeñas familias). Por su parte, a los puntos o vectores los denotamos por letras en negritas, por ejemplo, $\mathbf{p}$ y $\mathbf{q}$, para señalar que estamos pensando en su carácter de puntos, o bien, usamos $\mathbf{u}$, $\mathbf{v}$, $\mathbf{w}$ o $\mathbf{d}$ (acrónimo elegante que usaremos para “dirección”) y $\mathbf{n}$ (acrónimo para “normal”) para subrayar su papel vectorial. Por supuesto, $\mathbf{x}$, $\mathbf{y}$ y $\mathbf{z}$ siempre son buenos comodines para vectores (puntos o eneadas) variables o incógnitos, que no deben confundirse con x , y , z , que son variables reales o numéricas.

1.3 El espacio vectorial $\mathbb{R}^2$

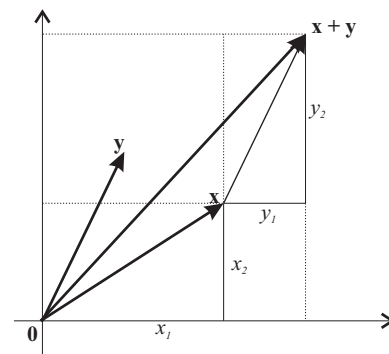
En esta sección se introduce la herramienta algebraica básica para hacer geometría con parejas, ternas o n -adas de números. Y de nuevo la idea central es muy simple: así como los números se pueden sumar y multiplicar, también los vectores tienen sus operacioncitas. Lo único que suponemos de los números reales es que sabemos o, mejor aún, que “saben ellos” sumarse y multiplicarse; a partir de ello extenderemos estas nociones a los vectores.

Definición 1.3.1 Dados dos vectores $\mathbf{u} = (x, y)$ y $\mathbf{v} = (x', y')$ en $\mathbb{R}^2$, definimos su *suma vectorial*, o simplemente su *suma*, como el vector $\mathbf{u} + \mathbf{v}$ que resulta de sumar coordenada a coordenada:

$$\mathbf{u} + \mathbf{v} := (x + x', y + y'),$$

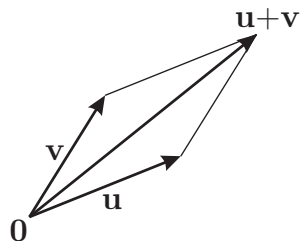
es decir,

$$(x, y) + (x', y') := (x + x', y + y').$$



Nótese que en cada coordenada, la suma que se usa es la suma usual de números reales. Así que al signo “+” de la suma le hemos ampliado el espectro de “saber sumar números” a “saber sumar vectores”; pero con una receta muy simple: “coordenada a coordenada”. Por ejemplo,

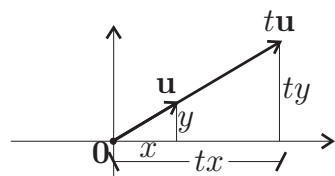
$$(3, 2) + (-1, 1) = (2, 3).$$



La suma vectorial corresponde geoméricamente a la regla del paralelogramo usada para encontrar la resultante de dos vectores. Esto es, se consideran los vectores como segmentos dirigidos que salen del origen, generan entonces un paralelogramo y el vector que va del origen a la otra esquina es la suma. También se puede pensar como la acción de dibujar un vector tras otro, pensando que los vectores son segmentos dirigidos que pueden moverse paralelos a sí mismos.

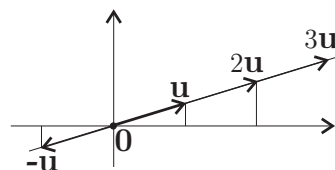
Definición 1.3.2 Dados un vector $\mathbf{u} = (x, y) \in \mathbb{R}^2$ y un número $t \in \mathbb{R}$ se define la *multiplicación escalar* $t\mathbf{u}$ como el vector que resulta de multiplicar cada coordenada del vector por el número:

$$t\mathbf{u} := (tx, ty).$$



Nótese que en cada coordenada, la multiplicación que se usa es la de los números reales.

La multiplicación escalar corresponde a la dilatación o contracción, y posiblemente al cambio de dirección, de un vector. Veamos. Está claro que



$$2\mathbf{u} = (2x, 2y) = (x + x, y + y) = \mathbf{u} + \mathbf{u},$$

así que $2\mathbf{u}$ es el vector “ $\mathbf{u}$ seguido de $\mathbf{u}$ ” o bien, “ $\mathbf{u}$ dilatado a su doble”. De la misma manera que $3\mathbf{u}$ es un vector que apunta en la misma dirección pero de tres veces su tamaño. O bien, es fácil deducir que $(\frac{1}{2})\mathbf{u}$, como punto, está justo a la mitad del camino del origen $\mathbf{0} = (0, 0)$ a $\mathbf{u}$. Así que $t\mathbf{u}$ para $t > 1$ es, estrictamente hablando, una *dilatación* de $\mathbf{u}$, y para $0 < t < 1$, una *contracción* del mismo. Por último, para $t < 0$, $t\mathbf{u}$ apunta en la dirección contraria ya que, en particular, $(-1)\mathbf{u} =: -\mathbf{u}$ es el vector que, como segmento dirigido, va del punto $\mathbf{u}$ al $\mathbf{0}$ (puesto que $\mathbf{u} + (-\mathbf{u}) = \mathbf{0}$) y el resto se obtiene como dilataciones o contracciones de $-\mathbf{u}$.

No está de más insistir en una diferencia esencial entre las dos operaciones que hemos definido. Si bien la suma vectorial es una operación que de dos ejemplares de la misma especie (vectores) nos da otro de ellos, la multiplicación escalar involucra a dos objetos de diferente índole, por un lado el “escalar”, un número real, y por el otro un vector, lo que da como resultado un nuevo vector. Los vectores no se multiplican (por el momento), sólo los escalares (los números reales) saben multiplicarlos (pegarles, podría decirse) y les cambian con ello su “escala”.

EJERCICIO 1.8 Sean $\mathbf{v}_1 = (2, 3)$, $\mathbf{v}_2 = (-1, 2)$, $\mathbf{v}_3 = (3, -1)$ y $\mathbf{v}_4 = (1, -4)$.

i) Calcula y dibuja: $2\mathbf{v}_1 - 3\mathbf{v}_2$; $2(\mathbf{v}_3 - \mathbf{v}_4) - \mathbf{v}_3 + 2\mathbf{v}_4$; $2\mathbf{v}_1 - 3\mathbf{v}_3 + 2\mathbf{v}_4$.

ii) ¿Qué vector $\mathbf{x} \in \mathbb{R}^2$ cumple que $2\mathbf{v}_1 + \mathbf{x} = 3\mathbf{v}_2$; $3\mathbf{v}_3 - 2\mathbf{x} = \mathbf{v}_4 + \mathbf{x}$?

iii) ¿Puedes encontrar $r, s \in \mathbb{R}$ tales que $r\mathbf{v}_2 + s\mathbf{v}_3 = \mathbf{v}_4$?

EJERCICIO 1.9 Dibuja el origen y tres vectores cualesquiera $\mathbf{u}$, $\mathbf{v}$ y $\mathbf{w}$ en un papel. Con un par de escuadras encuentra los vectores $\mathbf{u} + \mathbf{v}$, $\mathbf{v} + \mathbf{w}$ y $\mathbf{w} + \mathbf{u}$.

EJERCICIO 1.10 Dibuja el origen y dos vectores cualesquiera $\mathbf{u}$, $\mathbf{v}$ en un papel. Con regla (escuadras) y compás, encuentra los puntos $(1/2)\mathbf{u} + (1/2)\mathbf{v}$, $2\mathbf{u} - \mathbf{v}$, $3\mathbf{u} - 2\mathbf{v}$, $2\mathbf{v} - \mathbf{u}$. ¿Resultan colineales?

EJERCICIO 1.11 Describe el subconjunto de $\mathbb{R}^2$, o *lugar geométrico*, definido por $\mathcal{L}_{\mathbf{v}} := \{t\mathbf{v} \mid t \in \mathbb{R}\}$.

Estas definiciones se extienden a $\mathbb{R}^n$ de manera natural. Dados dos vectores $\mathbf{x} = (x_1, \dots, x_n)$ y $\mathbf{y} = (y_1, \dots, y_n)$ en $\mathbb{R}^n$ y un número real $t \in \mathbb{R}$, la *suma vectorial* (o simplemente la *suma*) y el *producto* (o *multiplicación*) *escalar* se definen como sigue:

$$\begin{aligned}\mathbf{x} + \mathbf{y} &:= (x_1 + y_1, \dots, x_n + y_n), \\ t\mathbf{x} &:= (tx_1, \dots, tx_n).\end{aligned}$$

Es decir, la suma de dos vectores (con el mismo número de coordenadas) se obtiene sumando coordenada a coordenada y el producto por un escalar (un número) se obtiene multiplicando cada coordenada por ese número.

Las propiedades básicas de la suma vectorial y la multiplicación escalar se reúnen en el siguiente teorema, donde el vector $\mathbf{0} = (0, \dots, 0)$ es llamado *vector cero* que corresponde al *origen*, y, para cada $\mathbf{x} \in \mathbb{R}^n$, el vector $-\mathbf{x} := (-1)\mathbf{x}$ se llama *inverso aditivo* de $\mathbf{x}$.

Teorema 1.2 Para todos los vectores $\mathbf{x}, \mathbf{y}, \mathbf{z} \in \mathbb{R}^n$ y para todos los números $s, t \in \mathbb{R}$ se cumple que:

- i) $(\mathbf{x} + \mathbf{y}) + \mathbf{z} = \mathbf{x} + (\mathbf{y} + \mathbf{z})$
- ii) $\mathbf{x} + \mathbf{y} = \mathbf{y} + \mathbf{x}$
- iii) $\mathbf{x} + \mathbf{0} = \mathbf{x}$
- iv) $\mathbf{x} + (-\mathbf{x}) = \mathbf{0}$
- v) $s(t\mathbf{x}) = (st)\mathbf{x}$
- vi) $1\mathbf{x} = \mathbf{x}$
- vii) $t(\mathbf{x} + \mathbf{y}) = t\mathbf{x} + t\mathbf{y}$
- viii) $(s + t)\mathbf{x} = s\mathbf{x} + t\mathbf{x}$

1.3.1 ¿Teorema o axiomas?

Antes de pensar en demostrar el teorema anterior, vale la pena reflexionar un poco sobre su carácter pues está muy cerca de ser un conjunto de axiomas y es sutil qué

quiere decir eso de demostrarlo.

Primero, debemos observar que $\mathbb{R}^n$ tiene sentido para $n = 1$, y es simplemente una manera rimbombante de referirse a los números reales. Así que del teorema, siendo $n = 1$ (es decir, si $\mathbf{x}$, $\mathbf{y}$, $\mathbf{z}$ están en $\mathbb{R}$), se obtiene parte de los *axiomas de los números reales*. Esto es, cada enunciado es una de las reglas elementales de la suma y la multiplicación que conocemos desde pequeños. Las propiedades (i) a (iv) son los axiomas que hacen a $\mathbb{R}$, con la operación suma, lo que se conoce como un “grupo conmutativo”: (i) dice que la suma es “asociativa”, (ii) que es “conmutativa”, (iii) que el $\mathbf{0}$ —o bien, el número 0— es su “neutro” (aditivo) y (iv) que todo elemento tiene “inverso” (aditivo). Por su parte, (v) y (vi) nos dicen que la multiplicación es asociativa y que tiene un neutro (multiplicativo), el 1; pero nos faltaría que es conmutativa y que tiene inversos (multiplicativos). Entonces tendríamos que añadir:

$$\begin{array}{ll} \text{ix)} & \mathbf{t} \mathbf{s} = \mathbf{s} \mathbf{t} \\ \text{x)} & \text{Si } \mathbf{t} \neq \mathbf{0}, \text{ existe } \mathbf{t}^{-1} \text{ tal que } \mathbf{t} \mathbf{t}^{-1} = \mathbf{1} \end{array}$$

para obtener que $\mathbb{R} - \{0\}$ (los reales sin el 0) son un grupo conmutativo con la multiplicación. Finalmente, (vii) y (viii) ya dicen lo mismo (en virtud de (ix)), que las dos operaciones se *distribuyen*. Estos axiomas definen lo que se llama un *campo*.

Obsérvese que en el caso general ($n > 1$) (ix) y (x) ni siquiera tienen sentido; pues sólo cuando $n = 1$ la multiplicación escalar involucra a seres de la misma especie. En resumen, para el caso $n = 1$, el teorema es un subconjunto de los axiomas que definen las operaciones de los números reales. No hay nada que demostrar, pues son parte de los axiomas básicos que vamos a usar. El resto de los axiomas que definen los números reales se refieren a su orden y, para que el libro quede autocontenido, de una vez los enunciamos. Los números reales tienen una *relación de orden*, denotada por $\leq$ y que se lee “menor o igual que” que cumple, para $\mathbf{a}, \mathbf{b}, \mathbf{c}, \mathbf{d} \in \mathbb{R}$, con:

$$\begin{array}{ll} \text{Oi)} & \mathbf{a} \leq \mathbf{b} \quad \text{o} \quad \mathbf{b} \leq \mathbf{a} \quad (\text{es un orden total}) \\ \text{Oii)} & \mathbf{a} \leq \mathbf{b} \quad \text{y} \quad \mathbf{b} \leq \mathbf{a} \quad \Leftrightarrow \quad \mathbf{a} = \mathbf{b} \\ \text{Oiii)} & \mathbf{a} \leq \mathbf{b} \quad \text{y} \quad \mathbf{b} \leq \mathbf{c} \quad \Rightarrow \quad \mathbf{a} \leq \mathbf{c} \quad (\text{es transitivo}) \\ \text{Oiv)} & \mathbf{a} \leq \mathbf{b} \quad \text{y} \quad \mathbf{c} \leq \mathbf{d} \quad \Rightarrow \quad \mathbf{a} + \mathbf{c} \leq \mathbf{b} + \mathbf{d} \\ \text{Ov)} & \mathbf{a} \leq \mathbf{b} \quad \text{y} \quad \mathbf{0} \leq \mathbf{c} \quad \Rightarrow \quad \mathbf{ac} \leq \mathbf{bc} \end{array}$$

Además, los números reales cumplen con el *axioma del supremo* que, intuitivamente, dice que la recta numérica no tiene “hoyos”, que los números reales forman un continuo. Pero este axioma, fundamental para el cálculo pues hace posible formalizar lo “infinitesimal”, no se usa en este libro, así que lo dejamos de lado.

Regresando al Teorema 1.2, para $n > 1$ la cosa es sutilmente diferente. Nosotros definimos la suma vectorial, y al ser algo nuevo sí tenemos que verificar que cumple las propiedades requeridas. Vale la pena introducir a Dios, el de las matemáticas no el de las religiones, para que quede claro. Dios nos da los números reales con la suma y la multiplicación, de alguna manera nos “ilumina” y de golpe: ahí están, con todo y sus axiomas. Ahora, nosotros —simples mortales— osamos definir la suma vectorial

y la multiplicación de escalares a vectores, pero Dios ya no nos asegura nada, él ya hizo lo suyo: nos toca demostrarlo a nosotros.

Demostración. (del Teorema 1.2) Formalmente sólo nos interesa demostrar el teorema para $n = 2, 3$, pero esencialmente es lo mismo para cualquier n . La demostración de cada inciso es muy simple, tanto así que hasta confunde, y consiste en aplicar el axioma correspondiente que cumplen los números reales coordenada a coordenada y la definición de las operaciones involucradas. Sería muy tedioso, y aportaría muy poco al lector, demostrar los ocho incisos, así que sólo demostraremos con detalle el primero para $\mathbb{R}^2$, dejando los demás y el caso general como ejercicios.

i) Sean $\mathbf{x}, \mathbf{y}, \mathbf{z} \in \mathbb{R}^2$, entonces $\mathbf{x} = (x_1, x_2)$, $\mathbf{y} = (y_1, y_2)$ y $\mathbf{z} = (z_1, z_2)$, donde cada x_i, y_i y z_i ($i = 1, 2$) son números reales (nótese que conviene usar subíndices con la misma letra que en negritas denota al vector). Tenemos entonces, a partir de la definición de suma vectorial, que

$$(\mathbf{x} + \mathbf{y}) + \mathbf{z} = (x_1 + y_1, x_2 + y_2) + (z_1, z_2)$$

y usando nuevamente la definición se obtiene que esta última expresión es

$$= ((x_1 + y_1) + z_1, (x_2 + y_2) + z_2).$$

Luego, como *la suma de números reales es asociativa* (el axioma correspondiente de los números reales usado coordenada a coordenada) se sigue

$$= (x_1 + (y_1 + z_1), x_2 + (y_2 + z_2)).$$

Y finalmente, usando dos veces la definición de suma vectorial, se obtiene

$$\begin{aligned} &= (x_1, x_2) + (y_1 + z_1, y_2 + z_2) \\ &= (x_1, x_2) + ((y_1, y_2) + (z_1, z_2)) \\ &= \mathbf{x} + (\mathbf{y} + \mathbf{z}), \end{aligned}$$

lo que demuestra que $(\mathbf{x} + \mathbf{y}) + \mathbf{z} = \mathbf{x} + (\mathbf{y} + \mathbf{z})$; y entonces tiene sentido escribir $\mathbf{x} + \mathbf{y} + \mathbf{z}$. $\square$

Se conoce como *espacio vectorial* a un conjunto en el que están definidas dos operaciones (suma vectorial y multiplicación escalar) que cumplen con las ocho propiedades del Teorema 1.2. De tal manera que este teorema puede parafrasearse “ $\mathbb{R}^n$ es un espacio vectorial”. ¿Podría el lector mencionar un espacio vectorial distinto de $\mathbb{R}^n$?

EJERCICIO 1.12 Usando los axiomas de los números reales, así como las definiciones de suma vectorial y producto escalar, demuestra algunos incisos del Teorema 1.2 para el caso $n = 2$ y $n = 3$. ¿Verdad que lo único que cambia es la longitud de los vectores, pero los argumentos son exactamente los mismos? Demuestra (aunque sea mentalmente) el caso general.

EJERCICIO 1.13 Demuestra que si $a \leq 0$ entonces $0 \leq -a$. (Usa el axioma **(Oiv)**.)

EJERCICIO 1.14 Observa que los axiomas de orden no dicen nada sobre cómo están relacionados el 0 y el 1, pero se pueden deducir de ellos. Supón que $1 \leq 0$, usando el ejercicio anterior y los axiomas **(Ov)** y **(Oii)**, demuestra que entonces $-1 = 0$; como esto no es cierto se debe cumplir la otra posibilidad por **(Oi)**, es decir, que $0 \leq 1$.

EJERCICIO 1.15 Demuestra que si $a \leq b$ y $c \leq 0$ entonces $ac \geq bc$ (donde $\geq$ significa “es mayor o igual”).

Ejemplo. Como corolario de este teorema se tiene que el “álgebra” simple a la que estamos acostumbrados con números también vale con los vectores. Por ejemplo, en el ejercicio (1.8.ii.b) se pedía encontrar el vector $\mathbf{x}$ tal que

$$3\mathbf{v}_3 - 2\mathbf{x} = \mathbf{v}_4 + \mathbf{x}.$$

Hagámoslo. Se vale pasar con signo menos del otro lado de la ecuación (sumando el inverso aditivo a ambos lados), y por tanto esta ecuación equivale a

$$-3\mathbf{x} = \mathbf{v}_4 - 3\mathbf{v}_3$$

y multiplicando por $-1/3$ se obtiene que

$$\mathbf{x} = \mathbf{v}_3 - (1/3)\mathbf{v}_4.$$

Ya nada más falta sustituir por los valores dados. Es probable que el estudiante haya hecho algo similar y qué bueno: tenía la intuición correcta. Pero hay que tener cuidado, así como con los números no se puede dividir entre cero, no se le vaya a ocurrir tratar de dividir entre un vector! Aunque a veces se pueden “cancelar” (ver Ej. 1.17). Para tal efecto, el siguiente lema será muy usado y vale la pena verlo en detalle.

Lema 1.3 Si $\mathbf{x} \in \mathbb{R}^2$ y $t \in \mathbb{R}$ son tales que $t\mathbf{x} = \mathbf{0}$ entonces $t = 0$ o $\mathbf{x} = \mathbf{0}$.

Demostración. Suponiendo que $t \neq 0$, hay que demostrar que $\mathbf{x} = \mathbf{0}$ para completar el lema. Pero entonces t tiene inverso multiplicativo y podemos multiplicar por t^{-1} ambos lados de la ecuación $t\mathbf{x} = \mathbf{0}$ para obtener (usando **(v)** y **(vi)**) que $\mathbf{x} = (t^{-1})\mathbf{0} = \mathbf{0}$ por la definición del vector nulo $\mathbf{0} = (0, 0)$. $\square$

EJERCICIO 1.16 Demuestra que si $\mathbf{x} \in \mathbb{R}^3$ y $t \in \mathbb{R}$ son tales que $t\mathbf{x} = \mathbf{0}$ entonces $t = 0$ o $\mathbf{x} = \mathbf{0}$. ¿Y para $\mathbb{R}^n$?

EJERCICIO 1.17 Demuestra que si $\mathbf{x} \in \mathbb{R}^n$ es distinto de $\mathbf{0}$, y $t, s \in \mathbb{R}$ son tales que $t\mathbf{x} = s\mathbf{x}$, entonces $t = s$. (Es decir, si $\mathbf{x} \neq \mathbf{0}$ está permitido “cancelarlo” aunque sea vector.)

1.4 Líneas

En el estudio de la geometría clásica, las líneas o rectas son subconjuntos básicos descritos por los axiomas; se les reconoce intuitivamente y se parte de ellas para construir lo demás. Con el método cartesiano, esto no es necesario. Las líneas se pueden definir o construir, correspondiendo a nuestra noción intuitiva de ellas —que no varía en nada de la de los griegos—, para después ver que efectivamente cumplen con los axiomas que antes se les asociaban. En esta sección, definiremos las rectas y veremos que cumplen los axiomas primero y tercero de Euclides.

Nuestra definición de recta estará basada en la intuición física de una partícula que viaja en movimiento rectilíneo uniforme; es decir, con velocidad fija y en la misma dirección. Estas dos nociones (velocidad como magnitud y dirección) se han amalgamado en el concepto de vector.

Recordemos que en el Ejercicio 1.11, se pide describir el conjunto $\mathcal{L}_v := \{t\mathbf{v} \in \mathbb{R}^2 \mid t \in \mathbb{R}\}$ donde $\mathbf{v} \in \mathbb{R}^2$. Debía ser claro que si $\mathbf{v} \neq \mathbf{0}$, entonces $\mathcal{L}_v$, al constar de todos los múltiplos escalares del vector $\mathbf{v}$, se dibuja como una recta que pasa por el origen (pues $\mathbf{0} = 0\mathbf{v}$) con la dirección de $\mathbf{v}$. Si consideramos la variable t como el tiempo, la función $\varphi(t) = t\mathbf{v}$ describe el movimiento rectilíneo uniforme de una partícula que viene desde tiempo infinito negativo, pasa por el origen $\mathbf{0}$ en el tiempo $t = 0$ (llamada la posición inicial), y seguirá por siempre con velocidad constante $\mathbf{v}$.

Pero por el momento queremos pensar en las rectas como conjuntos (en el capítulo siguiente estudiaremos más a profundidad las funciones). Los conjuntos $\mathcal{L}_v$, al rotar $\mathbf{v}$, nos dan las rectas por el origen, y para obtener otras, bastará “empujarlas” fuera del origen (o bien, arrancar el movimiento con otra posición inicial).

Definición 1.4.1 Dados un punto $\mathbf{p}$ y un vector $\mathbf{v} \neq \mathbf{0}$, la recta que pasa por $\mathbf{p}$ con dirección $\mathbf{v}$ es el conjunto:

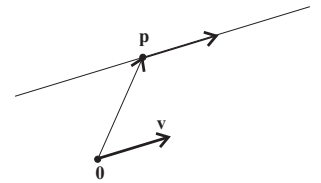
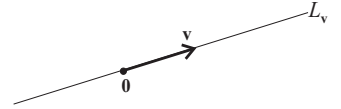
$$\ell := \{\mathbf{p} + t\mathbf{v} \mid t \in \mathbb{R}\}. \quad (1.1)$$

Una recta o línea en $\mathbb{R}^2$ es un subconjunto que tiene, para algún $\mathbf{p}$ y $\mathbf{v} \neq \mathbf{0}$, la descripción anterior.

A esta forma de definir una recta se le conoce como *representación o expresión paramétrica*. Esta representación trae consigo una función entre los números reales y la recta:

$$\begin{aligned} \varphi : \mathbb{R} &\rightarrow \mathbb{R}^2 \\ \varphi(t) &= \mathbf{p} + t\mathbf{v}. \end{aligned} \quad (1.2)$$

De hecho, esta función define una biyección entre $\mathbb{R}$ y ℓ . Como ℓ se definió por medio del parámetro $t \in \mathbb{R}$, es claro que *es* la imagen de la función φ ; es decir, φ es sobre. Demostremos (aunque de la intuición física parezca obvio) que es uno-a-uno. Supongamos para esto que $\varphi(t) = \varphi(s)$ para algunos $t, s \in \mathbb{R}$ y habrá que concluir



que $\mathbf{t} = \mathbf{s}$. Por la definición de φ se tiene

$$\mathbf{p} + \mathbf{t}\mathbf{v} = \mathbf{p} + \mathbf{s}\mathbf{v}.$$

De aquí, al sumar el inverso del lado izquierdo a ambos (que equivale a pasarlo de un lado al otro con signo contrario) se tiene

$$\begin{aligned} \mathbf{0} &= \mathbf{p} + \mathbf{s}\mathbf{v} - (\mathbf{p} + \mathbf{t}\mathbf{v}) \\ &= \mathbf{p} + \mathbf{s}\mathbf{v} - \mathbf{p} - \mathbf{t}\mathbf{v} \\ &= (\mathbf{p} - \mathbf{p}) + (\mathbf{s}\mathbf{v} - \mathbf{t}\mathbf{v}) \\ &= \mathbf{0} + (\mathbf{s} - \mathbf{t})\mathbf{v} \\ &= (\mathbf{s} - \mathbf{t})\mathbf{v}, \end{aligned}$$

donde hemos usado las propiedades (i), (ii), (iii), (iv) y (viii) del Teorema 1.2.² Del Lema 1.3, se concluye que $\mathbf{s} - \mathbf{t} = \mathbf{0}$ o bien que $\mathbf{v} = \mathbf{0}$. Pero por hipótesis $\mathbf{v} \neq \mathbf{0}$ (esto es la esencia), así que no queda otra que $\mathbf{s} = \mathbf{t}$. Esto demuestra que φ es inyectiva, y nuestra intuición se corrobora.

Hemos demostrado que cualquier recta está en biyección natural con $\mathbb{R}$, la recta modelo, así que todas las líneas son una “copia” de la recta real abstracta $\mathbb{R}$.

Observación 1.4.1 En la definición 1.4, así como en la demostración, no se usa de manera esencial que $\mathbf{p}$ y $\mathbf{v}$ estén en $\mathbb{R}^2$. Sólo se usa que hay una suma vectorial y un producto escalar bien definidos. Así que el punto y el vector podrían estar en $\mathbb{R}^3$ o en $\mathbb{R}^n$. Podemos entonces dar por establecida la noción de recta en cualquier dimensión al cambiar 2 por n en la Definición 1.4.

Observación 1.4.2 A la función 1.2 se le conoce como “movimiento rectilíneo uniforme” o “movimiento inercial”, pues, como decíamos al principio de la sección, describe la manera en que se mueve una partícula, o un cuerpo, que no está afectada por ninguna fuerza y tiene posición $\mathbf{p}$ y vector velocidad $\mathbf{v}$ en el tiempo $\mathbf{t} = 0$.

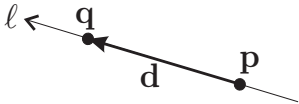
Ya podemos demostrar un resultado clásico, e intuitivamente claro.

Lema 1.4 *Dados dos puntos $\mathbf{p}$ y $\mathbf{q}$ en $\mathbb{R}^n$, existe una recta que pasa por ellos.*

Demostración. Si tuviéramos que $\mathbf{p} = \mathbf{q}$, como hay muchas rectas que pasan por $\mathbf{p}$, ya acabamos. Supongamos entonces que $\mathbf{p} \neq \mathbf{q}$, que es el caso interesante.

Si tomamos a $\mathbf{p}$ como punto base para la recta que buscamos, bastará encontrar un vector que nos lleve de $\mathbf{p}$ a $\mathbf{q}$, para tomarlo como dirección. Éste es la diferencia $\mathbf{q} - \mathbf{p}$, pues claramente

$$\mathbf{p} + (\mathbf{q} - \mathbf{p}) = \mathbf{q},$$



²Ésta es la última vez que haremos mención explícita del uso del Teorema 1.2; de aquí en adelante se aplicarán sus propiedades como si fueran lo más natural del mundo. Pero vale la pena que el estudiante se cuestione por algún tiempo qué es lo que se está usando en las manipulaciones algebraicas.

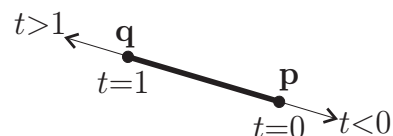
de tal manera que si definimos $\mathbf{d} := \mathbf{q} - \mathbf{p}$, como dirección, la recta

$$\ell = \{\mathbf{p} + t\mathbf{d} \mid t \in \mathbb{R}\},$$

(que sí es una recta pues $\mathbf{d} = \mathbf{q} - \mathbf{p} \neq \mathbf{0}$), es la que funciona. Con $t = 0$ obtenemos que $\mathbf{p} \in \ell$, y con $t = 1$ que $\mathbf{q} \in \ell$. $\square$

De la demostración del lema, se siguen los postulados I y III de Euclides. Obsérvese que cuando t toma valores entre 0 y 1, se obtienen puntos entre $\mathbf{p}$ y $\mathbf{q}$ (pues el vector direccional $(\mathbf{q} - \mathbf{p})$ se encoge al multiplicarlo por t), así que el *segmento* de $\mathbf{p}$ a $\mathbf{q}$, que denotaremos por $\overline{\mathbf{p}\mathbf{q}}$, se debe definir como

$$\overline{\mathbf{p}\mathbf{q}} := \{\mathbf{p} + t(\mathbf{q} - \mathbf{p}) \mid 0 \leq t \leq 1\}.$$



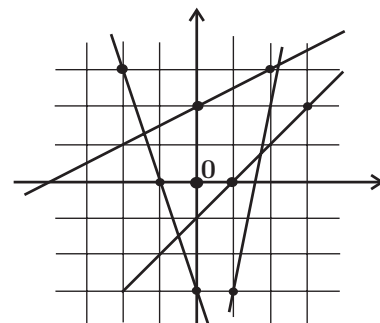
Y la recta ℓ que pasa por $\mathbf{p}$ y $\mathbf{q}$ se extiende “indefinidamente” a ambos lados del segmento $\overline{\mathbf{p}\mathbf{q}}$; para $t > 1$ del lado de $\mathbf{q}$ y para $t < 0$ del lado de $\mathbf{p}$.

***EJERCICIO 1.18** Aunque lo demostraremos más adelante, es un buen ejercicio demostrar formalmente en este momento que la recta ℓ , cuando $\mathbf{p} \neq \mathbf{q}$, es única.

EJERCICIO 1.19 Encuentra representaciones paramétricas de las rectas en la figura. Observa que su representación paramétrica no es única.

EJERCICIO 1.20 Dibuja las rectas

$$\begin{aligned} \{(2, 3) + t(1, 1) \mid t \in \mathbb{R}\}; \\ \{(-1, 0) + s(2, 1) \mid s \in \mathbb{R}\}; \\ \{(0, -2) + (-r, 2r) \mid r \in \mathbb{R}\}; \\ \{(t - 1, -2t) \mid t \in \mathbb{R}\}. \end{aligned}$$



EJERCICIO 1.21 Exhibe representaciones paramétricas para las seis rectas que pasan por dos de los siguientes puntos en $\mathbb{R}^2$: $\mathbf{p}_1 = (2, 4)$, $\mathbf{p}_2 = (-1, 2)$, $\mathbf{p}_3 = (-1, -1)$ y $\mathbf{p}_4 = (3, -1)$.

EJERCICIO 1.22 Muestra representaciones paramétricas para las tres rectas que genera el triángulo en $\mathbb{R}^3$: $\mathbf{q}_1 = (2, 1, 2)$, $\mathbf{q}_2 = (-1, 1, -1)$, $\mathbf{q}_3 = (-1, -2, -1)$.

EJERCICIO 1.23 Si una partícula viaja de $2\mathbf{q}_1$ a $3\mathbf{q}_3$ en movimiento inercial y tarda cuatro unidades de tiempo en llegar. ¿Cuál es su vector velocidad?

EJERCICIO 1.24 Dados dos vectores $\mathbf{u}$ y $\mathbf{v}$ en $\mathbb{R}^n$, el *paralelogramo* que definen tiene como vértices los puntos $\mathbf{0}$, $\mathbf{u}$, $\mathbf{v}$ y $\mathbf{u} + \mathbf{v}$ (como en la figura). Demuestra que sus *diagonales*, es decir, los segmentos de $\mathbf{0}$ a $\mathbf{u} + \mathbf{v}$ y de $\mathbf{u}$ a $\mathbf{v}$ se intersectan en su punto medio.

1.4.1 Coordenadas baricéntricas

Todavía podemos exprimirle más información a la demostración del Lema 1.4. Veremos cómo se escriben, en términos de $\mathbf{p}$ y $\mathbf{q}$, los puntos de su recta, y que esto tiene que ver con la clásica ley de las palancas. Además, de aquí surgirá una demostración muy simple del teorema clásico de concurrencia de medianas en un triángulo.

Supongamos que $\mathbf{p}$ y $\mathbf{q}$ son dos puntos distintos del plano. Para cualquier $t \in \mathbb{R}$, se tiene que

$$\begin{aligned} \mathbf{p} + t(\mathbf{q} - \mathbf{p}) &= \mathbf{p} + t\mathbf{q} - t\mathbf{p} \\ &= (1 - t)\mathbf{p} + t\mathbf{q} \end{aligned}$$

al reagrupar los términos. Y esta última expresión, a su vez, se puede reescribir como

$$s\mathbf{p} + t\mathbf{q} \quad \text{con} \quad s + t = 1,$$

donde hemos introducido la nueva variable $s = 1 - t$. De lo anterior se deduce que la recta ℓ que pasa por $\mathbf{p}$ y $\mathbf{q}$ puede también describirse como

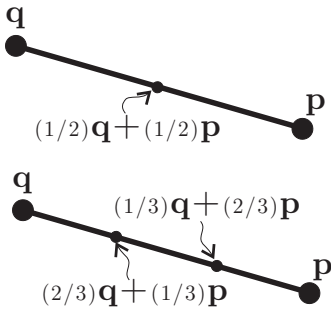
$$\ell = \{s\mathbf{p} + t\mathbf{q} \mid s + t = 1\},$$

que es una *expresión baricéntrica* de ℓ . A los números s, t (con $s + t = 1$, insistimos) se les conoce como *coordenadas baricéntricas* del punto $\mathbf{x} = s\mathbf{p} + t\mathbf{q}$ con respecto a $\mathbf{p}$ y $\mathbf{q}$.

Las coordenadas baricéntricas tienen la ventaja de que ya no distinguen entre los dos puntos. Como lo escribíamos antes —sólo con t —, $\mathbf{p}$ era el punto base (“el que la hace de 0”) y luego $\mathbf{q}$ era hacia donde íbamos (“el que la hace de 1”), y tenían un papel distinto. Ahora no hay preferencia por ninguno de los dos; las coordenadas baricéntricas son “democráticas”. Más aún, al expresar una recta con coordenadas baricéntricas no le damos una dirección preferida. Se usan simultáneamente los parámetros naturales para las dos direcciones (de $\mathbf{p}$ a $\mathbf{q}$ y de $\mathbf{q}$ a $\mathbf{p}$); pues si $s + t = 1$ entonces

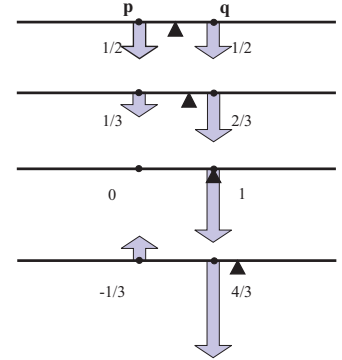
$$\begin{aligned} s\mathbf{p} + t\mathbf{q} &= \mathbf{p} + t(\mathbf{q} - \mathbf{p}) \\ &= \mathbf{q} + s(\mathbf{p} - \mathbf{q}). \end{aligned}$$

Nótese que $\mathbf{x} = s\mathbf{p} + t\mathbf{q}$ está en el segmento entre $\mathbf{p}$ y $\mathbf{q}$ si y sólo si sus dos coordenadas baricéntricas son no negativas, es decir, si y sólo si $t \geq 0$ y $s \geq 0$. Por ejemplo, el punto medio tiene coordenadas baricéntricas $1/2, 1/2$, es $(1/2)\mathbf{p} + (1/2)\mathbf{q}$; y el punto que divide al segmento de $\mathbf{p}$ a $\mathbf{q}$ en la proporción de $2/3$ a $1/3$ se escribe $(1/3)\mathbf{p} + (2/3)\mathbf{q}$, pues se acerca más a $\mathbf{q}$ que a $\mathbf{p}$. La extensión de la recta más allá de $\mathbf{q}$ tiene coordenadas baricéntricas s, t con $s < 0$ (y por lo tanto $t > 1$); así que los puntos de ℓ fuera del segmento de $\mathbf{p}$ a $\mathbf{q}$ tienen alguna coordenada baricéntrica negativa (la correspondiente al punto más lejano).



En términos físicos, podemos pensar la recta por $\mathbf{p}$ y $\mathbf{q}$ como una barra rígida. Si distribuimos una masa (que podemos pensar unitaria, es decir que vale 1) entre estos dos puntos, el punto de equilibrio tiene las coordenadas baricéntricas correspondientes a las masas. Si, por ejemplo, tienen el mismo peso, su punto de equilibrio está en el punto medio. Las masas negativas pueden pensarse como una fuerza hacia arriba y las correspondientes coordenadas baricéntricas nos dan entonces el punto de equilibrio para las palancas.

Veamos ahora un teorema clásico cuya demostración se simplifica enormemente usando coordenadas baricéntricas. Dado un triángulo con vértices $\mathbf{a}$, $\mathbf{b}$, $\mathbf{c}$, se definen las *medianas* como los segmentos que van de un vértice al punto medio del lado opuesto a éste.



Teorema 1.5 Dado un triángulo $\mathbf{a}$, $\mathbf{b}$, $\mathbf{c}$, sus tres medianas “concurren” en un punto que las parte en la proporción de $2/3$ (del vértice) a $1/3$ (del lado opuesto).

Demostración. Como el punto medio del segmento $\mathbf{b}$, $\mathbf{c}$ es $(\frac{1}{2}\mathbf{b} + \frac{1}{2}\mathbf{c})$, entonces la ...Dibujo mediana por $\mathbf{a}$ es el segmento

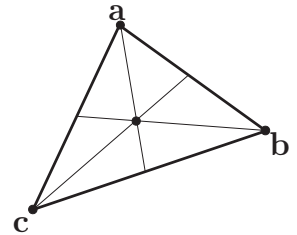
$$\{s\mathbf{a} + t\left(\frac{1}{2}\mathbf{b} + \frac{1}{2}\mathbf{c}\right) \mid s + t = 1, s \geq 0, t \geq 0\},$$

y análogamente se describen las otras dos medianas. Por suerte, el enunciado del teorema nos dice dónde buscar la intersección: el punto que describe en la mediana de $\mathbf{a}$ es precisamente

$$\frac{1}{3}\mathbf{a} + \frac{2}{3}\left(\frac{1}{2}\mathbf{b} + \frac{1}{2}\mathbf{c}\right) = \frac{1}{3}\mathbf{a} + \frac{1}{3}\mathbf{b} + \frac{1}{3}\mathbf{c}.$$

Entonces, de las igualdades

$$\begin{aligned} \frac{1}{3}\mathbf{a} + \frac{1}{3}\mathbf{b} + \frac{1}{3}\mathbf{c} &= \frac{1}{3}\mathbf{a} + \frac{2}{3}\left(\frac{1}{2}\mathbf{b} + \frac{1}{2}\mathbf{c}\right) \\ &= \frac{1}{3}\mathbf{b} + \frac{2}{3}\left(\frac{1}{2}\mathbf{c} + \frac{1}{2}\mathbf{a}\right) \\ &= \frac{1}{3}\mathbf{c} + \frac{2}{3}\left(\frac{1}{2}\mathbf{a} + \frac{1}{2}\mathbf{b}\right) \end{aligned}$$



se deduce que las tres medianas *concurren* en el punto $\frac{1}{3}(\mathbf{a} + \mathbf{b} + \mathbf{c})$, es decir, pasan por él. Este punto se llama el *baricentro* del triángulo, y claramente parte a las medianas en la proporción deseada. De nuevo, el baricentro corresponde al “centro de masa” o “punto de equilibrio” del triángulo.

□

EJERCICIO 1.25 Si tienes una barra rígida de un metro y con una fuerza de 10 kg quieres levantar una masa de 40 kg, ¿de dónde debes colgar la masa, estando el punto de apoyo al extremo de tu barra? Haz un dibujo.

EJERCICIO 1.26 Sean $\mathbf{a}$, $\mathbf{b}$ y $\mathbf{c}$ tres puntos distintos entre sí. Demuestra que si $\mathbf{a}$ está en la recta que pasa por $\mathbf{b}$ y $\mathbf{c}$, entonces $\mathbf{b}$ está en la recta por $\mathbf{a}$ y $\mathbf{c}$.

EJERCICIO 1.27 Encuentra el baricentro del triángulo $\mathbf{a} = (2, 4)$, $\mathbf{b} = (-1, 2)$, $\mathbf{c} = (-1, -1)$. Haz un dibujo.

EJERCICIO 1.28 Observa que en el Teorema 1.5, así como en la discusión que le precede, nunca usamos que estuviéramos en el plano. ¿Podría el lector enunciar y demostrar el teorema análogo para tetraedros en el espacio? (Un tetraedro está determinado por cuatro puntos en el espacio tridimensional. ¿Dónde estará su centro de masa?) ¿Puede generalizarlo a más dimensiones?

1.4.2 Planos en el espacio I

En la demostración del Teorema de las Medianas (1.5) se usó una idea que podemos utilizar para definir planos en el espacio. Dados tres puntos $\mathbf{a}$, $\mathbf{b}$ y $\mathbf{c}$ en $\mathbb{R}^3$ (aunque debe observar el lector que nuestra discusión se generaliza a $\mathbb{R}^n$), ya sabemos describir las tres líneas entre ellos; supongamos que son distintas. Entonces podemos tomar nuevos puntos en alguna de ellas y de estos puntos las nuevas líneas que los unen al vértice restante (como lo hicimos del punto medio de un segmento al vértice opuesto en el teorema). Está claro que la unión de todas estas líneas debe ser el plano que pasa por $\mathbf{a}$, $\mathbf{b}$ y $\mathbf{c}$; ésta es la idea que vamos a desarrollar.

Un punto $\mathbf{y}$ en la línea por $\mathbf{a}$ y $\mathbf{b}$ se escribe como

$$\mathbf{y} = s \mathbf{a} + t \mathbf{b} \quad \text{con} \quad s + t = 1.$$

Y un punto $\mathbf{x}$ en la línea que pasa por $\mathbf{y}$ y $\mathbf{c}$ se escribe entonces como

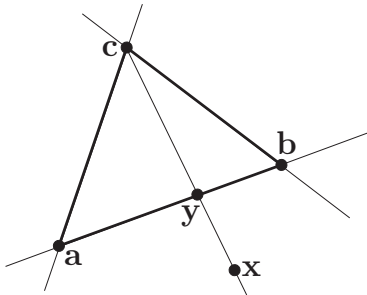
$$\mathbf{x} = r (s \mathbf{a} + t \mathbf{b}) + (1 - r) \mathbf{c},$$

para algún $r \in \mathbb{R}$; que es lo mismo que

$$\mathbf{x} = (rs) \mathbf{a} + (rt) \mathbf{b} + (1 - r) \mathbf{c}.$$

Observemos que los tres coeficientes suman uno:

$$(rs) + (rt) + (1 - r) = r(s + t) + 1 - r = r(1) + 1 - r = 1.$$



Y lo mismo hubiera pasado si en lugar de haber empezado con $\mathbf{a}$ y $\mathbf{b}$, hubiéramos comenzado con otra de las parejas. De lo que se trata entonces es que cualquier combinación de $\mathbf{a}$, $\mathbf{b}$ y $\mathbf{c}$ con coeficientes que sumen uno debe estar en su plano. Demostraremos que está en una línea por uno de los vértices y un punto en la línea que pasa por los otros dos.

Sean $\alpha, \beta, \gamma \in \mathbb{R}$ tales que $\alpha + \beta + \gamma = 1$. Consideremos al punto

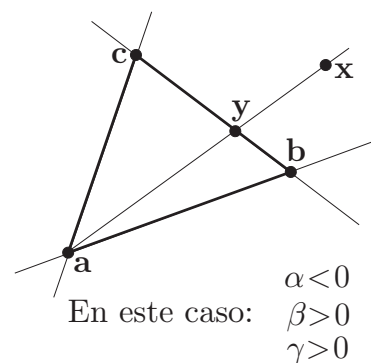
$$\mathbf{x} = \alpha \mathbf{a} + \beta \mathbf{b} + \gamma \mathbf{c}.$$

Como alguno de los coeficientes es distinto de 1 (si no, sumarían 3), podemos suponer sin pérdida de generalidad (esto quiere decir que los otros dos casos son análogos) que $\alpha \neq 1$. Entonces podemos dividir entre $1 - \alpha$ y se tiene

$$\mathbf{x} = \alpha \mathbf{a} + (1 - \alpha) \left(\frac{\beta}{1 - \alpha} \mathbf{b} + \frac{\gamma}{1 - \alpha} \mathbf{c} \right),$$

así que $\mathbf{x}$ está en la recta que pasa por $\mathbf{a}$ y el punto

$$\mathbf{y} = \frac{\beta}{1 - \alpha} \mathbf{b} + \frac{\gamma}{1 - \alpha} \mathbf{c}.$$



Como $\alpha + \beta + \gamma = 1$, entonces $\beta + \gamma = 1 - \alpha$, y los coeficientes de esta última expresión suman uno. Por lo tanto $\mathbf{y}$ está en la recta que pasa por $\mathbf{b}$ y $\mathbf{c}$. Hemos argumentado la siguiente definición:

Definición 1.4.2 Dados tres puntos $\mathbf{a}$, $\mathbf{b}$ y $\mathbf{c}$ en $\mathbb{R}^3$ *no colineales* (es decir, que no estén en una misma línea), el *plano* que pasa por ellos es el conjunto

$$\Pi = \{\alpha \mathbf{a} + \beta \mathbf{b} + \gamma \mathbf{c} \mid \alpha, \beta, \gamma \in \mathbb{R}; \alpha + \beta + \gamma = 1\}.$$

A una expresión de la forma $\alpha \mathbf{a} + \beta \mathbf{b} + \gamma \mathbf{c}$ con $\alpha + \beta + \gamma = 1$ la llamaremos *combinación afín* (o *baricéntrica*) de los puntos $\mathbf{a}$, $\mathbf{b}$, $\mathbf{c}$ y a los coeficientes α , β , γ las *coordenadas baricéntricas* del punto $\alpha \mathbf{a} + \beta \mathbf{b} + \gamma \mathbf{c}$.

EJERCICIO 1.29 Considera tres puntos $\mathbf{a}$, $\mathbf{b}$ y $\mathbf{c}$ no colineales, y sean α , β y γ las correspondientes coordenadas baricéntricas del plano que generan. Observa que cuando una de las coordenadas baricéntricas es cero, entonces el punto correspondiente está en una de las tres rectas por $\mathbf{a}$, $\mathbf{b}$ y $\mathbf{c}$. ¿En cuál?, ¿puedes demostrarlo? Haz un dibujo de los tres puntos y sus tres rectas. Éstas parten el plano en pedazos (¿cuántos?); en cada uno de ellos escribe los signos que toman las coordenadas baricéntricas (por ejemplo, en el interior del triángulo se tiene $(+, +, +)$ que corresponden respectivamente a (α, β, γ)).

EJERCICIO 1.30 Dibuja tres puntos $\mathbf{a}$, $\mathbf{b}$ y $\mathbf{c}$ no colineales. Con regla y compás encuentra los puntos $\mathbf{x} = (1/2)\mathbf{a} + (1/4)\mathbf{b} + (1/4)\mathbf{c}$, $\mathbf{y} = (1/4)\mathbf{a} + (1/2)\mathbf{b} + (1/4)\mathbf{c}$ y $\mathbf{z} = (1/4)\mathbf{a} + (1/4)\mathbf{b} + (1/2)\mathbf{c}$. Describe y argumenta tu construcción. Supón que puedes partir en tres el segmento $\overline{\mathbf{bc}}$ (hazlo midiendo con una regla o a ojo). ¿Puedes construir los puntos $\mathbf{u} = (1/2)\mathbf{a} + (1/3)\mathbf{b} + (1/6)\mathbf{c}$ y $\mathbf{v} = (1/2)\mathbf{a} + (1/6)\mathbf{b} + (1/3)\mathbf{c}$?

EJERCICIO 1.31 Demuestra que si $\mathbf{u}$ y $\mathbf{v}$ no son colineales con el origen, entonces el conjunto $\{s\mathbf{u} + t\mathbf{v} \mid s, t \in \mathbb{R}\}$ es un plano por el origen. (Usa que $s\mathbf{u} + t\mathbf{v} = r\mathbf{0} + s\mathbf{u} + t\mathbf{v}$ para cualquier $r, s, t \in \mathbb{R}$).

EJERCICIO 1.32 Demuestra que si $\mathbf{p}$ y $\mathbf{q}$ están en el plano Π de la Definición 1.4.2, entonces la recta que pasa por ellos está contenida en Π .

Definición lineal

Emocionado el autor por su demostración del teorema de medianas (1.5) usando coordenadas baricéntricas, se siguió de largo y definió planos en el espacio de una manera quizá no muy intuitiva. Que quede entonces la sección anterior como ejercicio en la manipulación de combinaciones de vectores; si no quedó muy clara puede releerse después de ésta. Pero conviene recapitular para definir los planos de otra manera, para aclarar la definición y así dejarnos la tarea de demostrar una equivalencia de dos definiciones.

La recta que genera un vector $\mathbf{u} \neq \mathbf{0}$ es el conjunto de todos sus “alargamientos” $\mathcal{L}_{\mathbf{u}} = \{s\mathbf{u} \mid s \in \mathbb{R}\}$. Si ahora tomamos un nuevo vector $\mathbf{v}$ que no esté en $\mathcal{L}_{\mathbf{u}}$,

tenemos una nueva recta $\mathcal{L}_{\mathbf{v}} = \{t\mathbf{v} \mid t \in \mathbb{R}\}$ que intersecta a la anterior sólo en el origen. Estas dos rectas generan un plano que consiste en todos los puntos (en $\mathbb{R}^3$, digamos, para fijar ideas) a los cuales se puede llegar desde el origen moviéndose únicamente en las direcciones $\mathbf{u}$ y $\mathbf{v}$. Este plano que pasa por el origen claramente se describe con dos parámetros independientes:

$$\Pi_0 = \{s\mathbf{u} + t\mathbf{v} \mid s, t \in \mathbb{R}\},$$

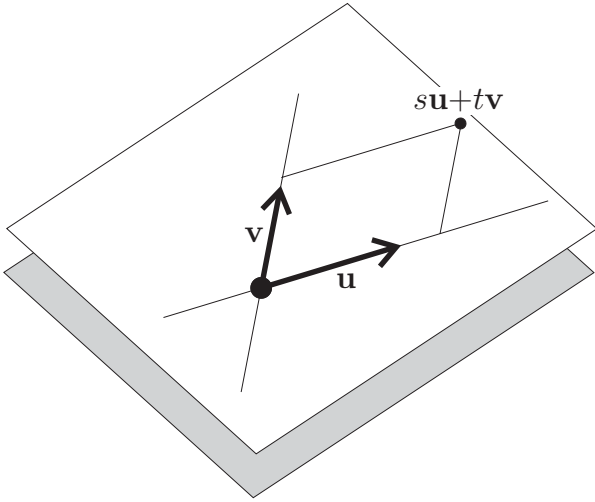
que por su propia definición está hecho a imagen y semejanza del plano $\mathbb{R}^2$ (ver Ejercicio 1.33, adelante). Como lo hicimos con

las rectas, podemos definir ahora un plano cualquiera como un plano por el origen (el que acabamos de describir) trasladado por cualquier otro vector constante.

Definición 1.4.3 Un *plano* en $\mathbb{R}^3$ es un conjunto de la forma

$$\Pi = \{\mathbf{p} + s\mathbf{u} + t\mathbf{v} \mid s, t \in \mathbb{R}\},$$

donde $\mathbf{u}$ y $\mathbf{v}$ son vectores no nulos tales que $\mathcal{L}_{\mathbf{u}} \cap \mathcal{L}_{\mathbf{v}} = \{\mathbf{0}\}$ y $\mathbf{p}$ es cualquier punto. A esta manera de describir un plano la llamaremos *expresión paramétrica*; a $\mathbf{u}$ y $\mathbf{v}$ se les llama *vectores direccionales* del plano Π y a $\mathbf{p}$ el *punto base* de la expresión paramétrica.



Para no confundirnos entre las dos definiciones de plano que hemos dado, llamemos *plano afín* a los que definimos en la sección anterior. Y ahora demostremos que las dos definiciones coinciden.

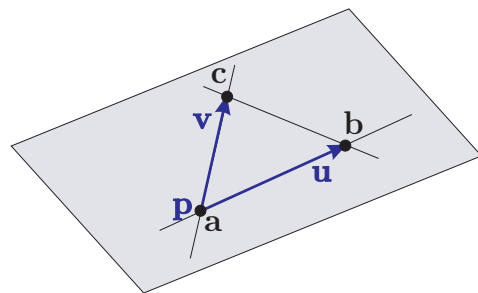
Lema 1.6 *En $\mathbb{R}^3$, todo plano es un plano afín y viceversa.*

Demostración. Sea Π como en la definición precedente. Debemos encontrar tres puntos en él y ver que todos los elementos de Π se expresan como combinación baricéntrica de ellos. Los puntos más obvios son

$$\mathbf{a} := \mathbf{p}; \mathbf{b} := \mathbf{p} + \mathbf{u}; \mathbf{c} := \mathbf{p} + \mathbf{v},$$

que claramente están en Π . Obsérvese que entonces se tiene que $\mathbf{u} = \mathbf{b} - \mathbf{a}$ y $\mathbf{v} = \mathbf{c} - \mathbf{a}$, de tal manera que para cualquier $s, t \in \mathbb{R}$ tenemos que

$$\begin{aligned} \mathbf{p} + s\mathbf{u} + t\mathbf{v} &= \mathbf{a} + s(\mathbf{b} - \mathbf{a}) + t(\mathbf{c} - \mathbf{a}) \\ &= \mathbf{a} + s\mathbf{b} - s\mathbf{a} + t\mathbf{c} - t\mathbf{a} \\ &= (1 - s - t)\mathbf{a} + s\mathbf{b} + t\mathbf{c}. \end{aligned}$$



Puesto que los coeficientes de esta última expresión suman 1, esto demuestra que Π está contenido en el plano afín generado por $\mathbf{a}$, $\mathbf{b}$, $\mathbf{c}$. E inversamente, cualquier combinación afín de $\mathbf{a}$, $\mathbf{b}$, $\mathbf{c}$ tiene esta última expresión, y por la misma igualdad se ve que está en Π . Obsérvese finalmente que si nos dan tres puntos $\mathbf{a}$, $\mathbf{b}$, $\mathbf{c}$, se pueden tomar como punto base a $\mathbf{a}$ y como vectores direccionales a $(\mathbf{b} - \mathbf{a})$ y $(\mathbf{c} - \mathbf{a})$ para expresar paramétricamente el plano afín que generan, pues las igualdades anteriores se siguen cumpliendo.

Ver que las condiciones que pusimos en ambas definiciones coinciden se deja como ejercicio. $\square$

Antes de seguir adelante, conviene establecer cierta **terminología** y **notación** para cosas, nociones y expresiones que estamos usando mucho:

- Dados los vectores $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_k$ en $\mathbb{R}^n$ (para incluir nuestros casos de interés $n = 2, 3$ de una buena vez), a una expresión de la forma

$$s_1\mathbf{u}_1 + s_2\mathbf{u}_2 + \dots + s_k\mathbf{u}_k,$$

donde $s_1, s_2, \dots, s_k$ son números reales (escalares), se le llama una *combinación lineal* de los vectores $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_k$ con *coeficientes* $s_1, s_2, \dots, s_k$. Obsérvese que toda combinación lineal da como resultado un vector, pero que un mismo vector tiene muchas expresiones tales.

- Como ya vimos, a una combinación lineal cuyos coeficientes suman 1 se le llama *combinación afín*. Y a una combinación afín de dos vectores distintos o de tres no colineales se le llama, además, *baricéntrica*.
- Al conjunto de *todas* las combinaciones lineales de los vectores $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_k$ se le llama el *subespacio generado* por ellos y se le denotará $\langle \mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_k \rangle$. Es decir,

$$\langle \mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_k \rangle := \{s_1 \mathbf{u}_1 + s_2 \mathbf{u}_2 + \dots + s_k \mathbf{u}_k \mid s_1, s_2, \dots, s_k \in \mathbb{R}\}.$$

Nótese que entonces, si $\mathbf{u} \neq \mathbf{0}$ se tiene que $\mathcal{L}_{\mathbf{u}} = \langle \mathbf{u} \rangle$ es la recta generada por $\mathbf{u}$; y ambas notaciones se seguirán usando indistintamente. Aunque ahora $\langle \mathbf{u} \rangle$ tiene sentido para $\mathbf{u} = \mathbf{0}$, en cuyo caso $\langle \mathbf{u} \rangle = \{\mathbf{0}\}$, y $\mathcal{L}_{\mathbf{u}}$ se usará para destacar que *es* una recta.

- Se dice que dos vectores $\mathbf{u}$ y $\mathbf{v}$ son *linealmente independientes* si son no nulos y tales que $\mathcal{L}_{\mathbf{u}} \cap \mathcal{L}_{\mathbf{v}} = \{\mathbf{0}\}$.
- Dado cualquier subconjunto $A \subset \mathbb{R}^n$, su *trasladado* por el vector (o al punto) $\mathbf{p}$ es el conjunto

$$A + \mathbf{p} = \mathbf{p} + A := \{\mathbf{x} + \mathbf{p} \mid \mathbf{x} \in A\}.$$

Podemos resumir entonces nuestras dos definiciones básicas como: una *recta* es un conjunto de la forma $\ell = \mathbf{p} + \langle \mathbf{u} \rangle$ con $\mathbf{u} \neq \mathbf{0}$; y un *plano* es un conjunto de la forma $\Pi = \mathbf{p} + \langle \mathbf{u}, \mathbf{v} \rangle$ con $\mathbf{u}$ y $\mathbf{v}$ linealmente independientes.

EJERCICIO 1.33 Demuestra que si $\mathbf{u}$ y $\mathbf{v}$ son linealmente independientes, entonces la función $f: \mathbb{R}^2 \rightarrow \mathbb{R}^3$ definida por $f(s, t) = s\mathbf{u} + t\mathbf{v}$ es inyectiva.

EJERCICIO 1.34 Demuestra que cualquier plano está en biyección con $\mathbb{R}^2$.

EJERCICIO 1.35 Demuestra que tres puntos $\mathbf{a}$, $\mathbf{b}$, $\mathbf{c}$ son no colineales si y sólo si los vectores $\mathbf{u} := (\mathbf{b} - \mathbf{a})$ y $\mathbf{v} := (\mathbf{c} - \mathbf{a})$ son linealmente independientes.

EJERCICIO 1.36 Da una expresión paramétrica para el plano que pasa por los puntos $\mathbf{a} = (0, 1, 2)$, $\mathbf{b} = (1, 1, 0)$ y $\mathbf{c} = (-1, 0, 2)$.

EJERCICIO 1.37 Demuestra que $\mathbf{0} \in \langle \mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_k \rangle$ para cualquier $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_k$ en $\mathbb{R}^n$.

EJERCICIO 1.38 Demuestra que dos vectores $\mathbf{u}$ y $\mathbf{v}$ son linealmente independientes si y sólo si la única combinación lineal de ellos que da $\mathbf{0}$ es la *trivial* (i.e., con ambos coeficientes cero).

EJERCICIO 1.39 Demuestra que

$$\mathbf{w} \in \langle \mathbf{u}, \mathbf{v} \rangle \iff \mathbf{w} + \langle \mathbf{u}, \mathbf{v} \rangle = \langle \mathbf{u}, \mathbf{v} \rangle \iff \langle \mathbf{u}, \mathbf{v}, \mathbf{w} \rangle = \langle \mathbf{u}, \mathbf{v} \rangle.$$

1.5 Medio Quinto

Regresemos al plano. Ya tenemos la noción de recta, y en esta sección veremos que nuestras rectas cumplen con la parte de existencia del quinto postulado de Euclides, que nuestra intuición va correspondiendo a nuestra formalización analítica de la geometría y que al cambiar los axiomas de Euclides por unos más básicos (los de los números reales) obtenemos los anteriores, pero ahora como teoremas demostrables. Ya vimos que por cualquier par de puntos se puede trazar un segmento que se extiende indefinidamente a ambos lados (axiomas I y III), es decir, que por ellos pasa una recta. El otro axioma que involucra rectas es el Quinto e incluye la noción de paralelismo, así que tendremos que empezar por ella.

Hay una definición conjuntista de rectas paralelas, así que formalicémosla. Como las rectas son, por definición, ciertos subconjuntos distinguidos del plano, tiene sentido la siguiente:

Definición 1.5.1 Dos rectas ℓ_1 y ℓ_2 en $\mathbb{R}^2$ son *paralelas*, que escribiremos $\ell_1 \parallel \ell_2$, si no se intersectan, es decir si

$$\ell_1 \cap \ell_2 = \emptyset,$$

donde $\emptyset$ denota al conjunto vacío (ver Apéndice A).

Pero además de rectas tenemos algo más elemental que son los vectores (segmentos dirigidos) y entre ellos también hay una noción intuitiva de paralelismo que corresponde al “alargamiento” o multiplicación por escalares.

Definición 1.5.2 Dados dos vectores $\mathbf{u}$, $\mathbf{v} \in \mathbb{R}^2$ distintos de $\mathbf{0}$, diremos que $\mathbf{u}$ es *paralelo a* $\mathbf{v}$, lo que escribiremos $\mathbf{u} \parallel \mathbf{v}$, si existe un número $t \in \mathbb{R}$ tal que $\mathbf{u} = t \mathbf{v}$.

Hemos eliminado el vector $\mathbf{0}$, el origen, de la definición para no complicar la situación. Si lo hubiéramos incluido no es cierto el siguiente ejercicio.

EJERCICIO 1.40 Demuestra que la relación “ser paralelo a” es una relación de equivalencia en $\mathbb{R}^2 - \{\mathbf{0}\}$ (en el plano menos el origen llamado el “plano agujerado”). Describe las clases de equivalencia. (En el Apéndice A se define relación de equivalencia.)

EJERCICIO 1.41 Sean $\mathbf{u}, \mathbf{v}$ dos vectores distintos de $\mathbf{0}$. Demuestra que:

$$\mathbf{u} \parallel \mathbf{v} \Leftrightarrow \mathcal{L}_{\mathbf{u}} \cap \mathcal{L}_{\mathbf{v}} \neq \{\mathbf{0}\} \Leftrightarrow \mathcal{L}_{\mathbf{u}} = \mathcal{L}_{\mathbf{v}}.$$

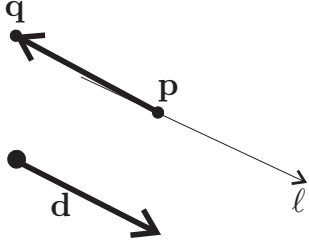
EJERCICIO 1.42 Demuestra que la relación entre rectas “ser paralelo a” es simétrica pero no reflexiva.

EJERCICIO 1.43 Demuestra que dos rectas *horizontales*, es decir, con vector direccional $\mathbf{d} = (1, 0)$, o son paralelas o son iguales.

Con la noción de paralelismo de vectores, podemos determinar cuándo un punto está en una recta.

Lema 1.7 Sea ℓ la recta que pasa por $\mathbf{p}$ con dirección $\mathbf{d}$ ($\mathbf{d} \neq \mathbf{0}$), y sea $\mathbf{q} \neq \mathbf{p}$, entonces

$$\mathbf{q} \in \ell \Leftrightarrow (\mathbf{q} - \mathbf{p}) \parallel \mathbf{d}.$$



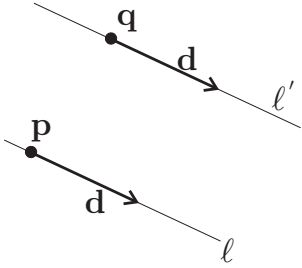
Demostración. Tenemos que $\mathbf{q} \in \ell$ si y sólo si existe una $t \in \mathbb{R}$ tal que

$$\mathbf{q} = \mathbf{p} + t \mathbf{d}.$$

Pero esto es equivalente a que $\mathbf{q} - \mathbf{p} = t \mathbf{d}$, que por definición es que $\mathbf{q} - \mathbf{p}$ es paralelo a $\mathbf{d}$. $\square$

Teorema 1.8 (1/2 Quinto) Sean ℓ una recta en $\mathbb{R}^2$ y $\mathbf{q}$ un punto fuera de ella, entonces existe una recta ℓ' que pasa por $\mathbf{q}$ y es paralela a ℓ .

Demostración. Nuestra definición de recta nos da un punto $\mathbf{p}$ y un vector dirección $\mathbf{d} \neq \mathbf{0}$ para ℓ , de tal manera que



$$\ell = \{\mathbf{p} + t \mathbf{d} \mid t \in \mathbb{R}\}.$$

Es intuitivamente claro —y a la intuición hay que seguirla pues es, de cierta manera, lo que ya sabíamos— que la recta paralela deseada debe tener la misma dirección, así que definamos

$$\ell' = \{\mathbf{q} + s \mathbf{d} \mid s \in \mathbb{R}\}.$$

Como $\mathbf{q} \in \ell'$, nos falta demostrar que $\ell \parallel \ell'$, es decir, que $\ell \cap \ell' = \emptyset$. Para lograrlo, supongamos que no es cierto, es decir, que existe $\mathbf{x} \in \ell \cap \ell'$. Por las expresiones paramétricas de ℓ y ℓ' , se tiene entonces que existen $t \in \mathbb{R}$ y $s \in \mathbb{R}$ para las cuales $\mathbf{x} = \mathbf{p} + t \mathbf{d}$ y $\mathbf{x} = \mathbf{q} + s \mathbf{d}$. De aquí se sigue que

$$\begin{aligned} \mathbf{q} + s \mathbf{d} &= \mathbf{p} + t \mathbf{d} \\ \mathbf{q} - \mathbf{p} &= t \mathbf{d} - s \mathbf{d} \\ \mathbf{q} - \mathbf{p} &= (t - s) \mathbf{d}, \end{aligned}$$

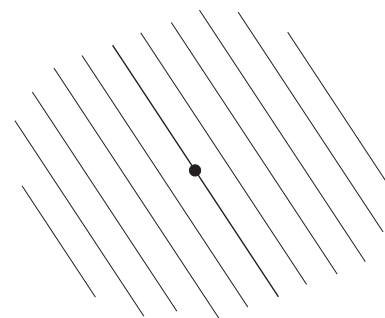
y entonces $\mathbf{q} \in \ell$ por el lema anterior, que es una contradicción a las hipótesis del teorema. Dicho de otra manera, como demostramos que $\ell \cap \ell' \neq \emptyset$ implica que $\mathbf{q} \in \ell$, podemos concluir que si $\mathbf{q} \notin \ell$, como en la hipótesis del teorema, no puede suceder que $\ell \cap \ell'$ sea no vacío, y por lo tanto $\ell \cap \ell' = \emptyset$ y $\ell \parallel \ell'$ por definición. Lo cual concluye con la parte de existencia del Quinto Postulado. $\square$

La parte que falta demostrar del Quinto es la unicidad, es decir, que cualquier otra recta que pase por $\mathbf{q}$ sí intersecta a ℓ ; que la única paralela es ℓ' . Esto se sigue de que rectas con vectores direccionales no paralelos siempre se intersectan, pero pospondremos la demostración formal de este hecho hasta tener más herramientas conceptuales. En particular, veremos en breve cómo encontrar la intersección de rectas para ejercitar la intuición, pero antes de entrarle, recapitulemos sobre la noción básica de esta sección, el paralelismo.

De la demostración del teorema se sigue que dos rectas con la misma dirección (o, lo que es lo mismo, con vectores direccionales paralelos) o son la misma o no se intersectan. Conviene entonces cambiar nuestra noción conjuntista de paralelismo a una vectorial.

Definición 1.5.3 Dos rectas ℓ_1 y ℓ_2 son *paralelas*, que escribiremos $\ell_1 \parallel \ell_2$, si tienen vectores direccionales paralelos.

Con esta nueva definición (que es la que se mantiene en lo sucesivo) una recta es paralela a sí misma, dos rectas paralelas distintas no se intersectan (son paralelas en el viejo sentido) y la relación es claramente transitiva. Así que es una relación de equivalencia y las clases de equivalencia corresponden a las clases de paralelismo de vectores (las rectas agujeradas por el origen). Llamaremos *haz de rectas paralelas* a una clase de paralelismo de rectas, es decir, al conjunto de todas las rectas paralelas a una dada (todas paralelas entre sí). Hay tantos haces de rectas paralelas como hay rectas por el origen, pues cada haz contiene exactamente una de estas rectas que, quitándole el origen, está formada por los posibles vectores direccionales para las rectas del haz.



Además, nuestra nueva noción de paralelismo tiene la gran ventaja de que se extiende a cualquier espacio vectorial. Nótese primero que la noción de paralelismo entre vectores no nulos se extiende sin problema a $\mathbb{R}^3$, pues está en términos del producto por escalares; y luego nuestra noción de paralelismo entre rectas se sigue de la de sus vectores direccionales. Por ejemplo, en el espacio $\mathbb{R}^3$ corresponde a nuestra noción intuitiva de paralelismo que no es conjuntista. Dos rectas pueden no intersectarse sin tener la misma dirección.

EJERCICIO 1.44 Da la descripción paramétrica de dos rectas en $\mathbb{R}^3$ que no se intersecten y que no sean paralelas.

EJERCICIO 1.45 (Quinto D3) Demuestra que dados un plano Π en $\mathbb{R}^3$ y un punto $\mathbf{q}$ fuera de él, existe un plano Π' que pasa por $\mathbf{q}$ y que no intersecta a Π .

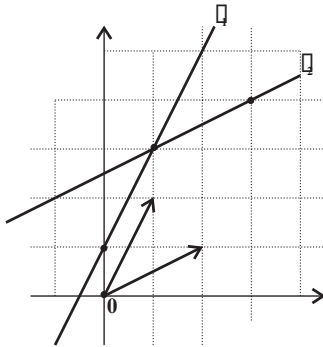
1.6 Intersección de rectas I

En esta sección se analiza el problema de encontrar la intersección de dos rectas. De nuevo, sabemos por experiencia e intuición que dos rectas no paralelas se deben intersectar en un punto. Cómo encontrar ese punto es la pregunta que responderemos en esta sección.

Hagamos un ejemplo. Sean

$$\ell_1 = \{(0, 1) + t(1, 2) \mid t \in \mathbb{R}\}$$

$$\ell_2 = \{(3, 4) + s(2, 1) \mid s \in \mathbb{R}\}.$$



La intersección de estas dos rectas $\ell_1 \cap \ell_2$ es el punto que cumple con las dos descripciones. Es decir, buscamos una $t \in \mathbb{R}$ y una $s \in \mathbb{R}$ que satisfagan

$$(0, 1) + t(1, 2) = (3, 4) + s(2, 1),$$

donde es importante haber dado a los dos parámetros diferente nombre. Esta ecuación vectorial puede reescribirse como

$$t(1, 2) - s(2, 1) = (3, 4) - (0, 1)$$

$$(t - 2s, 2t - s) = (3, 3),$$

que nos da una ecuación lineal con dos incógnitas en cada coordenada. Es decir, debemos resolver el sistema:

$$t - 2s = 3$$

$$2t - s = 3.$$

Si despejamos t en la primera ecuación, se obtiene $t = 3 + 2s$. Y al sustituir en la segunda obtenemos la ecuación lineal en la variable s :

$$2(3 + 2s) - s = 3,$$

que se puede resolver directamente:

$$6 + 4s - s = 3$$

$$3s = -3$$

$$s = -1.$$

Y por lo tanto, al sustituir este valor de s en la descripción de ℓ_2 , el punto

$$(3, 4) - (2, 1) = (1, 3)$$

está en las dos rectas. Nótese que bastó con encontrar el valor de un solo parámetro, pues de la correspondiente descripción paramétrica se obtiene el punto deseado. Pero debemos encontrar el mismo punto si resolvemos primero el otro parámetro. Veámoslo. Otra manera de resolver el sistema es eliminar la variable s . Para esto nos conviene multiplicar por -2 la segunda ecuación y sumarla a la primera, para obtener

$$-3t = -3,$$

de donde $t = 1$. Al sustituir en la parametrización de ℓ_1 , nos da:

$$(0, 1) + (1, 2) = (1, 3),$$

como esperábamos.

Observemos primero que hay muchas maneras de resolver un sistema de dos ecuaciones lineales (se llaman así porque los exponentes de las incógnitas son uno —parece que no hay exponentes—) con dos incógnitas. Cualquiera de ellos es bueno y debe llevar a la misma solución.

Que en este problema particular aparezca un sistema de ecuaciones no es casualidad, pues en general encontrar la intersección de rectas se tiene que resolver así. Ya que si tenemos dos rectas cualesquiera en $\mathbb{R}^2$:

$$\begin{aligned}\ell_1 &= \{\mathbf{p} + t\mathbf{v} \mid t \in \mathbb{R}\} \\ \ell_2 &= \{\mathbf{q} + s\mathbf{u} \mid s \in \mathbb{R}\}.\end{aligned}$$

Su intersección está dada por la t y la s que cumplen

$$\mathbf{p} + t\mathbf{v} = \mathbf{q} + s\mathbf{u},$$

que equivale a

$$t\mathbf{v} - s\mathbf{u} = \mathbf{q} - \mathbf{p}. \quad (1.3)$$

Esta ecuación vectorial, a su vez, equivale a una ecuación lineal en cada coordenada. Como estamos en $\mathbb{R}^2$, y $\mathbf{p}$, $\mathbf{q}$, $\mathbf{v}$ y $\mathbf{u}$ son constantes, nos da entonces un sistema de dos ecuaciones lineales con las dos incógnitas t y s . En cada caso particular el sistema se puede resolver de diferentes maneras; por ejemplo, despejando una incógnita y sustituyendo, o bien tomando múltiplos de las ecuaciones y sumándolas para eliminar una variable y obtener una ecuación lineal con una sola variable (la otra) que se resuelve directamente.

Para hacer la teoría y entender qué pasa con estos sistemas en general, será bueno estudiar las ecuaciones lineales con dos incógnitas de manera geométrica; veremos que también están asociadas a las líneas rectas. Por el momento, de manera “algebraica”, o mejor dicho, mecánica, resolvamos algunos ejemplos.

EJERCICIO 1.46 Encuentra los seis puntos de intersección de las cuatro rectas del Ejercicio 1.20.

EJERCICIO 1.47 Encuentra las intersecciones de las rectas $\ell_1 = \{(2, 0) + t(1, -2) \mid t \in \mathbb{R}\}$, $\ell_2 = \{(2, 1) + s(-2, 4) \mid s \in \mathbb{R}\}$ y $\ell_3 = \{(1, 2) + r(3, -6) \mid r \in \mathbb{R}\}$. Dibújalas para entender qué está pasando.

1.6.1 Sistemas de ecuaciones lineales I

Veamos ahora un método general para resolver sistemas.

Lema 1.9 *El sistema de ecuaciones*

$$\begin{aligned}as + bt &= e \\cs + dt &= f\end{aligned}\tag{1.4}$$

en las dos incógnitas s, t tiene solución única si su determinante $ad - bc$ es distinto de cero.

Demostración. Hay que intentar resolver el sistema general, seguir un método que funcione en todos los casos independientemente de los valores concretos que puedan tener las constantes a, b, c, d, e y f . Y el método más general es el de “eliminar” una variable para despejar la otra.

Para eliminar la t del sistema (1.4), multiplicamos por d la primera ecuación, y por $-b$ la segunda, para obtener

$$\begin{aligned}ads + bdt &= ed \\-bcs - bdt &= -bf,\end{aligned}\tag{1.5}$$

de tal manera que al sumar estas dos ecuaciones se tiene

$$(ad - bc)s = ed - bf.\tag{1.6}$$

Si $ad - bc \neq 0$,³ entonces se puede despejar s :

$$s = \frac{ed - bf}{ad - bc}.$$

Y análogamente (¡hágalo como ejercicio!), o bien, sustituyendo el valor de s en cualquiera de las ecuaciones originales (¡convénzase!), se obtiene t :

$$t = \frac{af - ec}{ad - bc}.$$

Lo cual demuestra que si el determinante es distinto de cero, la solución es única; es precisamente la de las dos fórmulas anteriores. $\square$

³Al número $ad - bc$ se le llama el *determinante* del sistema, porque determina que en este punto podemos proseguir.

Veamos ahora qué pasa con el sistema (1.4) cuando su determinante es cero. Si $ad - bc = 0$, obsérvese que aunque la ecuación (1.6) parezca muy sofisticada porque tiene muchas letras, en realidad es o una contradicción (algo falso) o una trivialidad. El coeficiente de s es 0, así que el lado izquierdo es 0. Entonces, si el lado derecho no es 0 es una contradicción; y si sí es 0, es cierto para cualquier s y hay una infinidad (precisamente un $\mathbb{R}$) de soluciones.

Lo que sucedió en este caso ($ad - bc = 0$) es que al intentar eliminar una de las variables también se eliminó la otra (como debió haberle sucedido al estudiante que hizo el Ejercicio 1.47). Entonces, o las dos ecuaciones son múltiplos y comparten soluciones (cuando $ed = bf$ las ecuaciones (1.5) ya son una la negativa de la otra); o bien, sólo los lados izquierdos son múltiplos, pero los derechos no lo son por el mismo factor ($ed \neq bf$) y no hay soluciones comunes.

Podemos resumir esto con el siguiente teorema que ya hemos demostrado.

Teorema 1.10 *Un sistema de dos ecuaciones lineales*

$$\begin{aligned}as + bt &= e \\cs + dt &= f\end{aligned}\tag{1.7}$$

en las dos incógnitas s, t tiene solución única si y sólo si su determinante $ad - bc$ es distinto de cero. Además, si su determinante es cero (si $ad - bc = 0$) entonces no tiene solución o tiene una infinidad de ellas (tantas como $\mathbb{R}$).

Si recordamos que los sistemas de dos ecuaciones lineales con dos incógnitas aparecieron al buscar la intersección de dos rectas, este teorema corresponde a nuestra intuición de las rectas: dos rectas se intersectan en un solo punto o no se intersectan o son la misma. De él dependerá la demostración de la mitad del Quinto que tenemos pendiente, pero conviene entrarle al toro por los cuernos: estudiar y entender primero el significado geométrico de las ecuaciones lineales con dos incógnitas.

EJERCICIO 1.48 Para las rectas de los dos ejercicios anteriores (1.46, 1.47), determina cómo se intersectan las rectas, usando únicamente el determinante.

1.7 Producto interior

En esta sección se introduce un ingrediente central para el estudio moderno de la geometría euclidiana: el “*producto interior*”. Además de darnos la herramienta algebraica para el estudio geométrico de las ecuaciones lineales, en la sección siguiente nos dará mucho más. De él derivaremos después las nociones básicas de distancia y ángulo (de rigidez). Puede decirse que el producto interior es, aunque menos intuitivo, más elemental que las nociones de distancia y ángulo pues éstas se definirán en base a

aquél (aunque se verá también que estas dos juntas lo definen). El producto interior depende tan íntimamente de la idea cartesiana de coordenadas, que no tiene ningún análogo en la geometría griega. Pero tampoco viene de los primeros pininos que hizo la geometría analítica, pues no es sino hasta el siglo XIX cuando se le empezó a dar la importancia debida al desarrollarse las ideas involucradas en la noción general de espacio vectorial. Podría entonces decirse que el uso del producto interior (junto con el lenguaje de espacio vectorial) marca dos épocas en la geometría analítica. Pero es quizá la simpleza de su definición su mejor tarjeta de presentación.

Definición 1.7.1 Dados dos vectores $\mathbf{u} = (u_1, u_2)$ y $\mathbf{v} = (v_1, v_2)$, su *producto interno* (también conocido como producto *escalar* —que preferimos no usar para no confundir con la multiplicación por escalares— o bien como producto *punto*) es el número real:

$$\mathbf{u} \cdot \mathbf{v} = (u_1, u_2) \cdot (v_1, v_2) := u_1 v_1 + u_2 v_2.$$

En general, en $\mathbb{R}^n$ se define el *producto interior* (o el *producto punto*) de dos vectores $\mathbf{u} = (u_1, \dots, u_n)$ y $\mathbf{v} = (v_1, \dots, v_n)$ como la suma de los productos de sus coordenadas correspondientes, es decir,

$$\mathbf{u} \cdot \mathbf{v} := u_1 v_1 + \dots + u_n v_n = \sum_{i=1}^n u_i v_i.$$

Así, por ejemplo:

$$\begin{aligned}(4, 3) \cdot (2, -1) &= 4 \times 2 + 3 \times (-1) = 8 - 3 = 5, \\ (1, 2, 3) \cdot (4, 5, -6) &= 4 + 10 - 18 = -4.\end{aligned}$$

Obsérvese que el producto interior tiene como ingredientes dos vectores ($\mathbf{u}, \mathbf{v} \in \mathbb{R}^n$) y nos da un escalar ($\mathbf{u} \cdot \mathbf{v} \in \mathbb{R}$); no debe confundirse con la multiplicación escalar, que de un escalar y un vector nos da un vector, excepto en el caso $n = 1$, en el que ambos coinciden con la multiplicación de los reales.

Antes de demostrar las propiedades básicas del producto interior, observemos que nos será muy útil para el problema que dejamos pendiente sobre sistemas de ecuaciones. En efecto en una ecuación lineal con dos incógnitas, x y y digamos, aparece una expresión de la forma $ax + by$ con a y b constantes. Si tomamos un vector constante $\mathbf{u} = (a, b)$ y un vector variable $\mathbf{x} = (x, y)$, ahora podemos escribir

$$\mathbf{u} \cdot \mathbf{x} = ax + by$$

y el “paquete básico” de información está en $\mathbf{u}$. En particular, como la mayoría de los lectores ya deben saber, la ecuación lineal

$$ax + by = c$$

donde c es una nueva constante, define una recta. Con el producto interior esta ecuación se reescribe como

$$\mathbf{u} \cdot \mathbf{x} = c.$$

Veremos cómo en el vector $\mathbf{u}$ y en la constante c se almacena la información geométrica de esa recta. Pero no nos apresuremos.

Como en el caso de la suma vectorial y la multiplicación por escalares, conviene demostrar primero las propiedades elementales del producto interior para manejarlo después con más soltura.

Teorema 1.11 *Para todos los vectores $\mathbf{u}, \mathbf{v}, \mathbf{w} \in \mathbb{R}^n$, y para todo número $t \in \mathbb{R}$ se cumple que*

- i) $\mathbf{u} \cdot \mathbf{v} = \mathbf{v} \cdot \mathbf{u}$*
- ii) $\mathbf{u} \cdot (t \mathbf{v}) = t (\mathbf{u} \cdot \mathbf{v})$*
- iii) $\mathbf{u} \cdot (\mathbf{v} + \mathbf{w}) = \mathbf{u} \cdot \mathbf{v} + \mathbf{u} \cdot \mathbf{w}$*
- iv) $\mathbf{u} \cdot \mathbf{u} \geq 0$*
- v) $\mathbf{u} \cdot \mathbf{u} = 0 \Leftrightarrow \mathbf{u} = \mathbf{0}$*

Demostración. Nos interesa demostrarlo en $\mathbb{R}^2$ (y como ejercicio en $\mathbb{R}^3$), aunque el caso general es esencialmente lo mismo. Convendrá usar la notación de la misma letra con subíndices para las coordenadas, es decir, tomar $\mathbf{u} = (u_1, u_2)$, $\mathbf{v} = (v_1, v_2)$ y $\mathbf{w} = (w_1, w_2)$. El enunciado (i) se sigue inmediatamente de las definiciones y la conmutatividad de los números reales; (ii) se obtiene al factorizar:

$$\mathbf{u} \cdot (t \mathbf{v}) = u_1(tv_1) + u_2(tv_2) = t(u_1v_1 + u_2v_2) = t(\mathbf{u} \cdot \mathbf{v}).$$

De la distributividad y conmutatividad en los reales se obtiene (iii):

$$\begin{aligned} \mathbf{u} \cdot (\mathbf{v} + \mathbf{w}) &= u_1(v_1 + w_1) + u_2(v_2 + w_2) \\ &= u_1v_1 + u_1w_1 + u_2v_2 + u_2w_2 \\ &= (u_1v_1 + u_2v_2) + (u_1w_1 + u_2w_2) = \mathbf{u} \cdot \mathbf{v} + \mathbf{u} \cdot \mathbf{w}. \end{aligned}$$

Al tomar el producto interior de un vector consigo mismo se obtiene

$$\mathbf{u} \cdot \mathbf{u} = u_1^2 + u_2^2,$$

que es una suma de cuadrados. Como cada cuadrado es positivo la suma también lo es (iv); y si fuera 0 significa que cada sumando es 0, es decir, que cada coordenada es 0 (v). $\square$

EJERCICIO 1.49 Demuestra el Teorema 1.11 para $n = 3$.

EJERCICIO 1.50 Calcula el producto interior de algunos vectores del Ejercicio 1.3.

EJERCICIO 1.51 Demuestra sin usar (v), sólo la definición de producto interior, que dado $\mathbf{u} \in \mathbb{R}^2$ se tiene

$$\mathbf{u} \cdot \mathbf{x} = 0, \quad \forall \mathbf{x} \in \mathbb{R}^2 \Leftrightarrow \mathbf{u} = \mathbf{0}.$$

(Observa que sólo hay que usar el lado izquierdo para dos vectores muy sencillos.)

EJERCICIO 1.52 Demuestra el análogo del ejercicio anterior en $\mathbb{R}^3$. ¿Y en $\mathbb{R}^n$?

1.7.1 El compadre ortogonal

El primer uso geométrico que daremos al producto interior será para detectar la perpendicularidad. Es otra de las nociones básicas en los axiomas de Euclides que incluyen el concepto de ángulo recto.

Fijemos el vector $\mathbf{u} \in \mathbb{R}^2$, y para simplificar la notación digamos que $\mathbf{u} = (\mathbf{a}, \mathbf{b}) \neq \mathbf{0}$. Si tomamos $\mathbf{x} = (x, y)$ como un vector variable, vamos a ver que las soluciones de la ecuación

$$\mathbf{u} \cdot \mathbf{x} = 0, \quad (1.8)$$

es decir, los puntos en $\mathbb{R}^2$ cuyas coordenadas (x, y) satisfacen la ecuación

$$ax + by = 0,$$

forman una recta que es perpendicular a $\mathbf{u}$. Para esto consideraremos una solución particular (una de las más sencillas), $x = -b$, $y = a$; y será tan importante esta solución que le daremos nombre:

Definición 1.7.2 El *compadre ortogonal* del vector $\mathbf{u} = (\mathbf{a}, \mathbf{b})$, denotado $\mathbf{u}^\perp$ y que se lee “ $\mathbf{u}$ -perpendicular” o “ $\mathbf{u}$ -ortogonal”, es

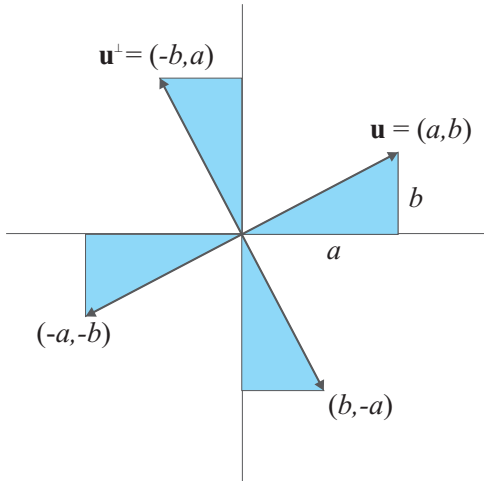
$$\mathbf{u}^\perp = (-b, a).$$

(Se intercambian coordenadas y a la primera se le cambia el signo.)

Se cumple que

$$\mathbf{u} \cdot \mathbf{u}^\perp = a(-b) + ba = 0;$$

y para justificar el nombre hay que mostrar que $\mathbf{u}^\perp$ se obtiene de $\mathbf{u}$ al girarlo 90° en sentido contrario a las manecillas del reloj (alrededor del origen). Para ver esto considérese la figura al margen, donde suponemos que el vector $\mathbf{u} = (\mathbf{a}, \mathbf{b})$ está en el primer cuadrante, es decir, que $\mathbf{a} > 0$ y $\mathbf{b} > 0$. Al rotar el triángulo rectángulo de catetos $\mathbf{a}$ y $\mathbf{b}$ (e hipotenusa $\mathbf{u}$) un ángulo de 90° en el origen se obtiene el correspondiente a $\mathbf{u}^\perp$. Se incluyen además las siguientes dos rotaciones para que quede claro que esto no depende de los signos de $\mathbf{a}$ y $\mathbf{b}$.



Podemos concluir entonces que el “compadre ortogonal”, pensado como la función de $\mathbb{R}^2$ en $\mathbb{R}^2$ que manda al vector $\mathbf{u}$ en su perpendicular $\mathbf{u}^\perp$ ($\mathbf{u} \mapsto \mathbf{u}^\perp$), es la rotación de 90° en sentido contrario a las manecillas del reloj alrededor del origen. Es fácil comprobar que $(\mathbf{u}^\perp)^\perp = -\mathbf{u}$, ya sea con la fórmula o porque la rotación de 180° (rotar y volver a rotar 90°) corresponde justo a tomar el inverso aditivo.

Como ya es costumbre, veamos algunas propiedades bonitas de la función “compadre ortogonal” respecto de las otras operaciones que hemos definido. La demostración, que consiste en dar coordenadas y aplicar definiciones, se deja al lector.

Lema 1.12 *Para todos los vectores $\mathbf{u}, \mathbf{v} \in \mathbb{R}^2$, y para todo número $t \in \mathbb{R}$ se cumple que*

$$\begin{aligned} i) \quad & (\mathbf{u} + \mathbf{v})^\perp = \mathbf{u}^\perp + \mathbf{v}^\perp \\ ii) \quad & (t\mathbf{u})^\perp = t(\mathbf{u}^\perp) \\ iii) \quad & \mathbf{u}^\perp \cdot \mathbf{v}^\perp = \mathbf{u} \cdot \mathbf{v} \\ iv) \quad & \mathbf{u}^\perp \cdot \mathbf{v} = -(\mathbf{u} \cdot \mathbf{v}^\perp). \end{aligned}$$

EJERCICIO 1.53 Demuestra el lema anterior.

Sistemas de ecuaciones lineales II

Usemos la herramienta del producto interior y el compadre ortogonal para revisar nuestro trabajo previo (sección 1.6.1) sobre sistemas de dos ecuaciones lineales con dos incógnitas. Si pensamos cada ecuación como la coordenada de un vector (y así fue como surgieron), un sistema tal se escribe

$$s\mathbf{u} + t\mathbf{v} = \mathbf{c}, \tag{1.9}$$

donde s y t son las incógnitas y $\mathbf{u}, \mathbf{v}, \mathbf{c} \in \mathbb{R}^2$ están dados. Para que este sistema sea justo el que ya estudiamos, (1.4), denotemos las coordenadas $\mathbf{u} = (\mathbf{a}, \mathbf{c})$, $\mathbf{v} = (\mathbf{b}, \mathbf{d})$; y para las constantes (con acrónimo $\mathbf{c}$), digamos que $\mathbf{c} = (\mathbf{e}, \mathbf{f})$. Como ésta es una ecuación vectorial, si la multiplicamos (con el producto interior) por un vector, obtendremos una ecuación lineal real. Y para eliminar una variable, digamos s , podemos multiplicar por el compadre ortogonal de $\mathbf{u}$, para obtener

$$\begin{aligned} \mathbf{u}^\perp \cdot (s\mathbf{u}) + \mathbf{u}^\perp \cdot (t\mathbf{v}) &= \mathbf{u}^\perp \cdot \mathbf{c} \\ s(\mathbf{u}^\perp \cdot \mathbf{u}) + t(\mathbf{u}^\perp \cdot \mathbf{v}) &= \mathbf{u}^\perp \cdot \mathbf{c} \\ s(0) + t(\mathbf{u}^\perp \cdot \mathbf{v}) &= \mathbf{u}^\perp \cdot \mathbf{c} \\ t(\mathbf{u}^\perp \cdot \mathbf{v}) &= \mathbf{u}^\perp \cdot \mathbf{c}. \end{aligned}$$

Definición 1.7.3 El *determinante* de los vectores $\mathbf{u}, \mathbf{v} \in \mathbb{R}^2$ es el número real

$$\det(\mathbf{u}, \mathbf{v}) := \mathbf{u}^\perp \cdot \mathbf{v}.$$

Así que si suponemos que $\det(\mathbf{u}, \mathbf{v}) \neq 0$, podemos despejar t . De manera análoga podemos despejar s , y entonces concluir que el sistema de ecuaciones tiene solución única:

$$t = \frac{\mathbf{u}^\perp \cdot \mathbf{c}}{\mathbf{u}^\perp \cdot \mathbf{v}}, \quad s = \frac{\mathbf{v}^\perp \cdot \mathbf{c}}{\mathbf{v}^\perp \cdot \mathbf{u}};$$

que es justo el Lema 1.9.

EJERCICIO 1.54 Escribe en coordenadas ($\mathbf{u} = (a, c)$, $\mathbf{v} = (b, d)$) cada uno de los pasos anteriores para ver que corresponden al método clásico de eliminación. (Quizá te convenga escribir las parejas como columnas, en vez de renglones, para que el sistema de ecuaciones sea claro.)

La ecuación lineal homogénea

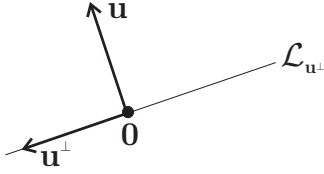
Ahora sí, describamos el conjunto de soluciones de la ecuación lineal homogénea:

$$\mathbf{u} \cdot \mathbf{x} = 0.$$

Proposición 1.13 Sea $\mathbf{u} \in \mathbb{R}^2 - \{\mathbf{0}\}$, entonces

$$\{\mathbf{x} \in \mathbb{R}^2 \mid \mathbf{u} \cdot \mathbf{x} = 0\} = \mathcal{L}_{\mathbf{u}^\perp},$$

donde, recuérdese, $\mathcal{L}_{\mathbf{u}^\perp} = \{t\mathbf{u}^\perp \mid t \in \mathbb{R}\}$.



Demostración. Tenemos que demostrar la igualdad de dos conjuntos. Esto se hace en dos pasos que corresponden a ver que cada elemento de un conjunto está también en el otro.

La contención más fácil es “ $\supseteq$ ”. Un elemento de $\mathcal{L}_{\mathbf{u}^\perp}$ es de la forma $t\mathbf{u}^\perp$ para algún $t \in \mathbb{R}$. Tenemos que demostrar que $t\mathbf{u}^\perp$ está en el conjunto de la izquierda. Para esto hay que ver que satisface la condición de pertenencia, que se sigue de (ii) en el Lema 1.12, y (ii) en el Teorema 1.11:

$$\mathbf{u} \cdot (t\mathbf{u}^\perp) = \mathbf{u} \cdot t(\mathbf{u}^\perp) = t(\mathbf{u} \cdot \mathbf{u}^\perp) = 0.$$

Para la otra contención, tomamos $\mathbf{x} \in \mathbb{R}^2$ tal que $\mathbf{u} \cdot \mathbf{x} = 0$ y debemos encontrar $t \in \mathbb{R}$ tal que $\mathbf{x} = t\mathbf{u}^\perp$. Sean $a, b \in \mathbb{R}$ las coordenadas de $\mathbf{u}$, es decir $\mathbf{u} = (a, b)$ y análogamente, sea $\mathbf{x} = (x, y)$. Entonces estamos suponiendo que

$$ax + by = 0.$$

Puesto que $\mathbf{u} \neq \mathbf{0}$ por hipótesis, tenemos que $a \neq 0$ o $b \neq 0$; y en cada caso podemos despejar una variable para encontrar el factor deseado:

$$a \neq 0 \Rightarrow x = -\frac{by}{a},$$

y por lo tanto

$$\mathbf{x} = (x, y) = \frac{y}{a}(-b, a) = \left(\frac{y}{a}\right)\mathbf{u}^\perp;$$

o bien,

$$b \neq 0 \Rightarrow y = -\frac{ax}{b} \Rightarrow \mathbf{x} = (x, y) = -\frac{x}{b}(-b, a) = \left(-\frac{x}{b}\right)\mathbf{u}^\perp.$$

□

Hemos demostrado que la ecuación lineal *homogénea* (se llama así porque la constante es 0) $\mathbf{a}\mathbf{x} + \mathbf{b}\mathbf{y} = 0$, con $\mathbf{a}$ y $\mathbf{b}$ constantes reales, tiene como soluciones las parejas $(\mathbf{x}, \mathbf{y})$ que forman una recta por el origen de $\mathbb{R}^2$. Si empaquetamos la información de la ecuación en el vector $\mathbf{u} = (\mathbf{a}, \mathbf{b})$, tenemos que $\mathbf{u} \neq \mathbf{0}$, pues de lo contrario no habría tal ecuación lineal (sería la tautología $0 = 0$ de lo que estamos hablando, y no es así). Pero además hemos visto que el vector $\mathbf{u}$ es ortogonal (también llamado *normal*) a la recta en cuestión, pues ésta está generada por el compadre ortogonal $\mathbf{u}^\perp$.

Se justifica entonces la siguiente definición, que era nuestro primer objetivo con el producto interior (donde, de una vez, estamos extrapolando a todas las dimensiones).

Definición 1.7.4 Se dice que dos vectores $\mathbf{u}$ y $\mathbf{v}$ en $\mathbb{R}^n$ son *perpendiculares* u *ortogonales*, y se escribe $\mathbf{u} \perp \mathbf{v}$, si $\mathbf{u} \cdot \mathbf{v} = 0$.

Podemos reformular entonces la proposición anterior como “dada $\mathbf{u} \in \mathbb{R}^2$, $\mathbf{u} \neq \mathbf{0}$, el conjunto de $\mathbf{x} \in \mathbb{R}^2$ que satisfacen la ecuación $\mathbf{u} \cdot \mathbf{x} = 0$ son la recta *perpendicular* a $\mathbf{u}$ ”.

Antes de estudiar la ecuación general (no necesariamente homogénea), notemos que si el producto interior detecta perpendicularidad, también se puede usar, junto con el compadre ortogonal, para detectar paralelismo en el plano.

Corolario 1.14 Sean $\mathbf{u}$ y $\mathbf{v}$ dos vectores no nulos en $\mathbb{R}^2$, entonces

$$\mathbf{u} \parallel \mathbf{v} \iff \det(\mathbf{u}, \mathbf{v}) = \mathbf{u}^\perp \cdot \mathbf{v} = 0.$$

Demostración. Por la Proposición 1.13, $\mathbf{u}^\perp \cdot \mathbf{v} = 0$ si y sólo si $\mathbf{v}$ pertenece a la recta generada por $(\mathbf{u}^\perp)^\perp$, que es la recta generada por $\mathbf{u}$ pues $(\mathbf{u}^\perp)^\perp = -\mathbf{u}$. $\square$

Hay que hacer notar que si les damos coordenadas a $\mathbf{u}$ y a $\mathbf{v}$, digamos $\mathbf{u} = (\mathbf{a}, \mathbf{b})$ y $\mathbf{v} = (\mathbf{c}, \mathbf{d})$, entonces el determinante, que detecta paralelismo, es

$$\mathbf{u}^\perp \cdot \mathbf{v} = \mathbf{ad} - \mathbf{bc},$$

que ya habíamos encontrado como determinante del sistema de ecuaciones (1.9), pero considerando $\mathbf{u} = (\mathbf{a}, \mathbf{c})$ y $\mathbf{v} = (\mathbf{b}, \mathbf{d})$.

EJERCICIO 1.55 Dibuja la recta definida como $\{\mathbf{x} \in \mathbb{R}^2 \mid \mathbf{u} \cdot \mathbf{x} = 0\}$ para $\mathbf{u}$ con cada uno de los siguientes vectores: a) $(1, 0)$, b) $(0, 1)$, c) $(2, 1)$, d) $(1, 2)$, e) $(-1, 1)$.

EJERCICIO 1.56 Describe el lugar geométrico definido como $\{\mathbf{x} \in \mathbb{R}^3 \mid \mathbf{u} \cdot \mathbf{x} = 0\}$ para $\mathbf{u}$ con cada uno de los siguientes vectores: a) $(1, 0, 0)$, b) $(0, 1, 0)$, c) $(0, 0, 1)$, d) $(1, 1, 0)$.

EJERCICIO 1.57 Sean $\mathbf{u} = (\mathbf{a}, \mathbf{b})$ y $\mathbf{v} = (\mathbf{c}, \mathbf{d})$ dos vectores no nulos. Sin usar el producto interior ni el compadre ortogonal —es decir, de la pura definición y con “álgebra elemental”— demuestra que $\mathbf{u}$ es paralelo a $\mathbf{v}$ si y sólo si $\mathbf{ad} - \mathbf{bc} = 0$. Compara con la última parte de la demostración de la Proposición 1.13.

1.8 La ecuación normal de la recta

Demostraremos ahora que todas las rectas de $\mathbb{R}^2$ se pueden describir con una ecuación

$$\mathbf{u} \cdot \mathbf{x} = c,$$

donde $\mathbf{u} \in \mathbb{R}^2$ es un vector normal a la recta y la constante $c \in \mathbb{R}$ es escogida apropiadamente. Esta ecuación, vista en coordenadas, equivale a una ecuación lineal en dos variables ($ax + by = c$, al tomar $\mathbf{u} = (a, b)$ constante y $\mathbf{x} = (x, y)$ es el vector variable). Históricamente, el hecho de que las rectas tuvieran tal descripción fue una gran motivación en el inicio de la geometría analítica.

Tomemos una recta con dirección $\mathbf{d} \neq \mathbf{0}$

$$\ell = \{\mathbf{p} + t\mathbf{d} \mid t \in \mathbb{R}\}.$$

Si definimos $\mathbf{u} := \mathbf{d}^\perp$ y $c := \mathbf{u} \cdot \mathbf{p}$, afirmamos, es decir, demostraremos, que

$$\ell = \{\mathbf{x} \in \mathbb{R}^2 \mid \mathbf{u} \cdot \mathbf{x} = c\}.$$

Llamemos ℓ' a este último conjunto; demostrar la igualdad $\ell = \ell'$, equivale a demostrar las dos contenciones $\ell \subseteq \ell'$ y $\ell' \subseteq \ell$.

Si $\mathbf{x} = \mathbf{p} + t\mathbf{d} \in \ell$, entonces (usando las propiedades del producto interior y del compadre ortogonal) tenemos

$$\begin{aligned} \mathbf{u} \cdot \mathbf{x} &= \mathbf{u} \cdot (\mathbf{p} + t\mathbf{d}) = \mathbf{u} \cdot \mathbf{p} + t(\mathbf{u} \cdot \mathbf{d}) \\ &= \mathbf{u} \cdot \mathbf{p} + t(\mathbf{d}^\perp \cdot \mathbf{d}) = \mathbf{u} \cdot \mathbf{p} + 0 = c, \end{aligned}$$

lo cual demuestra que $\ell \subseteq \ell'$. Por el otro lado, dado $\mathbf{x} \in \ell'$ (i.e., tal que $\mathbf{u} \cdot \mathbf{x} = c$), sea $\mathbf{y} := \mathbf{x} - \mathbf{p}$ (obsérvese que entonces $\mathbf{x} = \mathbf{p} + \mathbf{y}$). Se tiene que

$$\mathbf{u} \cdot \mathbf{y} = \mathbf{u} \cdot (\mathbf{x} - \mathbf{p}) = \mathbf{u} \cdot \mathbf{x} - \mathbf{u} \cdot \mathbf{p} = c - c = 0.$$

Por la Proposición 1.13, esto implica que $\mathbf{y}$ es paralelo a $\mathbf{u}^\perp = -\mathbf{d}$; y por tanto que $\mathbf{x} = \mathbf{p} + \mathbf{y} \in \ell$. Esto demuestra que $\ell = \ell'$.

En resumen, hemos demostrado que toda recta puede ser descrita por una ecuación normal:

Teorema 1.15 *Sea $\ell = \{\mathbf{p} + t\mathbf{d} \mid t \in \mathbb{R}\}$ una recta en $\mathbb{R}^2$. Entonces ℓ consta de los vectores $\mathbf{x} \in \mathbb{R}^2$ que cumplen la ecuación $\mathbf{u} \cdot \mathbf{x} = c$, que escribimos*

$$\ell : \quad \mathbf{u} \cdot \mathbf{x} = c,$$

donde $\mathbf{u} = \mathbf{d}^\perp$ y $c = \mathbf{u} \cdot \mathbf{p}$.

□

Este teorema también podría escribirse:

Teorema 1.16 Sea $\mathbf{d} \in \mathbb{R}^2 - \{\mathbf{0}\}$, entonces

$$\{\mathbf{p} + t\mathbf{d} \mid t \in \mathbb{R}\} = \{\mathbf{x} \in \mathbb{R}^2 \mid \mathbf{d}^\perp \cdot \mathbf{x} = \mathbf{d}^\perp \cdot \mathbf{p}\}$$

Estos enunciados nos dicen cómo encontrar una ecuación normal para una recta descrita paramétricamente. Por **ejemplo**, la recta

$$\{(s-2, 3-2s) \mid s \in \mathbb{R}\}$$

tiene como vector direccional al $(1, -2)$. Por tanto tiene vector normal a su compadre ortogonal $(2, 1)$, y como pasa por el punto $(-2, 3)$ entonces está determinada por la ecuación

$$\begin{aligned}(2, 1) \cdot (x, y) &= (2, 1) \cdot (-2, 3) \\ 2x + y &= -1\end{aligned}$$

(sustituya el lector las coordenadas de la descripción paramétrica en esta última ecuación).

...Dibujo

Para encontrar una representación paramétrica de la recta ℓ dada por una *ecuación normal*

$$\ell: \mathbf{u} \cdot \mathbf{x} = c$$

(léase “ ℓ dada por la ecuación $\mathbf{u} \cdot \mathbf{x} = c$ ”), bastará con encontrar una solución particular $\mathbf{p}$ (que es muy fácil porque se puede dar un valor arbitrario a una variable, 0 es la más conveniente, y despejar la otra), pues sabemos que la dirección es $\mathbf{u}^\perp$ (o cualquier paralelo). Por **ejemplo**, la recta dada por la ecuación normal

$$2x - 3y = 2$$

tiene vector normal $(2, -3)$. Por tanto tiene vector direccional $(3, 2)$; y como pasa por el punto $(1, 0)$, es el conjunto

$$\{(1 + 3t, 2t) \mid t \in \mathbb{R}\}.$$

Esto también se puede obtener directamente de las coordenadas, pero hay que partirlo en casos. Si $\mathbf{u} = (a, b)$, tenemos

$$\ell: ax + by = c.$$

Supongamos que $a \neq 0$. Entonces se puede despejar x :

$$x = \frac{c}{a} - \frac{b}{a}y$$

y todas las soluciones de la ecuación se obtienen dando diferentes valores a y ; es decir, son

$$\left\{ \left(\frac{c}{a} - \frac{b}{a}y, y \right) \mid y \in \mathbb{R} \right\} = \left\{ \left(\frac{c}{a}, 0 \right) + y \left(-\frac{b}{a}, 1 \right) \mid y \in \mathbb{R} \right\},$$

que es una recta con dirección $(-b, a)$. Nos falta ver qué pasa cuando $a = 0$, pero entonces la recta es la horizontal $y = c/b$. Si hubiéramos hecho el análisis cuando $b \neq 0$, se nos hubiera confundido la notación del caso clásico que vemos en el siguiente párrafo.

La ecuación “funcional” de la recta

Es bien conocida la ecuación *funcional de la recta*

$$y = mx + b,$$

que se usa para describir rectas como gráficas de una función; aquí m es la *pendiente* y b es la llamada “ordenada al origen” o el “valor inicial”. En nuestros términos, esta ecuación se puede reescribir como

$$-mx + y = b$$

o bien, como

$$(-m, 1) \cdot (x, y) = b.$$

Por lo tanto, tiene vector normal $\mathbf{n} = (-m, 1)$. Podemos escoger a $\mathbf{d} = (1, m) = -\mathbf{n}^\perp$ como su vector direccional y a $\mathbf{p} = (0, b)$ como una solución particular. Así que su parametrización natural es (usando a x como parámetro):

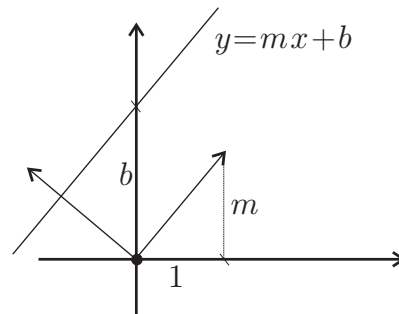
$$\{(0, b) + x(1, m) \mid x \in \mathbb{R}\} = \{(x, mx + b) \mid x \in \mathbb{R}\},$$

que es la gráfica de una función. Obsérvese que todas las rectas excepto las verticales (con dirección $(0, 1)$) se pueden expresar así.

No está de más observar que para cualquier función $f : \mathbb{R} \rightarrow \mathbb{R}$ se obtiene un subconjunto de $\mathbb{R}^2$, llamado la *gráfica de la función* f definida paramétricamente como

$$\{(x, f(x)) \mid x \in \mathbb{R}\};$$

y entonces hemos visto que todas las rectas, excepto las verticales, son la gráfica de funciones de la forma $f(x) = mx + b$, que en el capítulo 3 llamaremos funciones afines.



EJERCICIO 1.58 Para las siguientes rectas, encuentra una ecuación normal y, en su caso, su ecuación funcional:

- a) $\{(2, 3) + t(1, 1) \mid t \in \mathbb{R}\}$
- b) $\{(-1, 0) + s(2, 1) \mid s \in \mathbb{R}\}$
- c) $\{(0, -2) + (-r, 2r) \mid r \in \mathbb{R}\}$
- d) $\{(1, 3) + s(2, 0) \mid s \in \mathbb{R}\}$
- e) $\{(t - 1, -2t) \mid t \in \mathbb{R}\}$

EJERCICIO 1.59 Da una descripción paramétrica de las rectas dadas por las ecuaciones:

- a) $2x - 3y = 1$
- b) $2x - y = 2$
- c) $2y - 4x = 2$
- d) $x + 5y = -1$
- e) $3 - 4y = 2x + 2$

EJERCICIO 1.60 Encuentra una ecuación normal para la recta que pasa por los puntos:

- a) $(2, 1)$ y $(3, 4)$
- b) $(-1, 1)$ y $(2, 2)$
- c) $(1, -3)$ y $(3, 1)$
- d) $(2, 0)$ y $(1, 1)$

EJERCICIO 1.61 Sean $\mathbf{p}$ y $\mathbf{q}$ dos puntos distintos en $\mathbb{R}^2$. Demuestra que la recta ℓ que pasa por ellos tiene ecuación normal:

$$\ell: (\mathbf{q} - \mathbf{p})^\perp \cdot \mathbf{x} = \mathbf{q}^\perp \cdot \mathbf{p}.$$

EJERCICIO 1.62 Encuentra la intersección de las rectas (a) y (b) del Ejercicio 1.59.

EJERCICIO 1.63 Encuentra la intersección de la recta (α) del Ejercicio 1.58 con la de la recta (α) del Ejercicio 1.59, donde $\alpha \in \{a, b, c, d\}$.

EJERCICIO 1.64 Dibuja las gráficas de las funciones $f(x) = x^2$ y $g(x) = x^3$.

1.8.1 Intersección de rectas II

Regresemos ahora al problema teórico de encontrar la intersección de rectas, pero con la nueva herramienta de la ecuación normal. Consideremos dos rectas ℓ_1 y ℓ_2 en $\mathbb{R}^2$. Sean $\mathbf{u} = (a, b)$ y $\mathbf{v} = (c, d)$ vectores normales a ellas respectivamente, de tal manera que ambos son no nulos y existen constantes $e, f \in \mathbb{R}$ para las cuales

$$\begin{aligned}\ell_1 &: \mathbf{u} \cdot \mathbf{x} = e \\ \ell_2 &: \mathbf{v} \cdot \mathbf{x} = f.\end{aligned}$$

Es claro entonces que $\ell_1 \cap \ell_2$ consta de los puntos $\mathbf{x} \in \mathbb{R}^2$ que satisfacen ambas ecuaciones. Si llamamos, como de costumbre, x y y a las coordenadas de $\mathbf{x}$ (es decir, si hacemos $\mathbf{x} = (x, y)$), las dos ecuaciones anteriores son el sistema

$$\begin{aligned}ax + by &= e \\ cx + dy &= f,\end{aligned}$$

con a, b, c, d, e, f constantes y x, y variables o incógnitas. Éste es precisamente el sistema que estudiamos en la sección 1.6.1; aunque allá hablabamos de los “parámetros s y t ”, son esencialmente lo mismo. Pero ahora tiene el significado geométrico que buscábamos: resolverlo es encontrar el punto de intersección de dos rectas (sus coordenadas tal cual, y no los parámetros para encontrarlo, como antes). Como ya hicimos el trabajo abstracto de resolver el sistema, sólo nos queda por hacer la traducción al lenguaje geométrico. Como ya vimos, el tipo de solución (una única, ninguna o tantas como reales) depende del determinante del sistema $ad - bc$, y éste, a su vez, ya se nos apareció (nótese que $\det(\mathbf{u}, \mathbf{v}) = \mathbf{u}^\perp \cdot \mathbf{v} = ad - bc$) como el “numero mágico” que detecta paralelismo (Corolario 1.14). De tal manera que del Corolario 1.14 y del Teorema 1.10 podemos concluir:

Teorema 1.17 *Dadas las rectas*

$$\begin{aligned}\ell_1 &: \mathbf{u} \cdot \mathbf{x} = e \\ \ell_2 &: \mathbf{v} \cdot \mathbf{x} = f,\end{aligned}$$

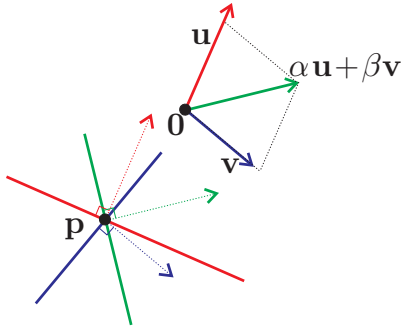
entonces:

- i) $\det(\mathbf{u}, \mathbf{v}) \neq 0 \Leftrightarrow \ell_1 \cap \ell_2$ es un único punto;
- ii) $\det(\mathbf{u}, \mathbf{v}) = 0 \Leftrightarrow \ell_1 \parallel \ell_2 \Leftrightarrow \ell_1 \cap \ell_2 = \emptyset$ o $\ell_1 = \ell_2$.

Consideremos con más detenimiento el segundo caso, cuando el determinante $\mathbf{u}^\perp \cdot \mathbf{v}$ es cero y por tanto los vectores normales (y las rectas) son paralelos. Tenemos entonces que para alguna $t \in \mathbb{R}$, se cumple que $\mathbf{v} = t\mathbf{u}$ (nótese que $t \neq 0$, pues ambos vectores son no nulos). La disyuntiva de si las rectas no se intersectan o son iguales corresponde a que las constantes e y f difieran por el mismo factor, es decir si $f \neq te$ o $f = te$ respectivamente. Podemos sacar dos conclusiones interesantes. Primero, dos ecuaciones normales definen la misma recta si y sólo si las tres constantes que las determinan (en nuestro caso $(\mathbf{a}, \mathbf{b}, e)$ y $(\mathbf{c}, \mathbf{d}, f)$) son vectores paralelos en $\mathbb{R}^3$. Y segundo, que al fijar un vector $\mathbf{u} = (\mathbf{a}, \mathbf{b})$ y variar un parámetro $c \in \mathbb{R}$, las ecuaciones $\mathbf{u} \cdot \mathbf{x} = c$ determinan el haz de rectas paralelas que es ortogonal a $\mathbf{u}$.

Otro punto interesante que vale la pena discutir es el método que usamos para resolver sistemas de ecuaciones. Se consideraron ciertos múltiplos de ellas y luego se sumaron. Es decir, para ciertas α y β en $\mathbb{R}$ se obtiene, de las dos ecuaciones dadas, una nueva de la forma

$$(\alpha \mathbf{u} + \beta \mathbf{v}) \cdot \mathbf{x} = \alpha e + \beta f.$$



Ésta es la ecuación de otra recta. La propiedad importante que cumple es que si $\mathbf{p}$ satisface las dos ecuaciones dadas, entonces también satisface esta última por las propiedades de nuestras operaciones. De tal manera que esta nueva ecuación define una recta que pasa por el punto de intersección $\mathbf{p}$. El método de eliminar una variable consiste en encontrar la horizontal o la vertical que pasa por el punto

(la mitad del problema). Pero en general, todas las ecuaciones posibles de la forma anterior definen el *haz de rectas concurrentes* por el punto de intersección $\mathbf{p}$ cuando $\mathbf{u}$ y $\mathbf{v}$ no son paralelos, pues los vectores $\alpha \mathbf{u} + \beta \mathbf{v}$ (al variar α y β) son más que suficientes para definir todas las posibles direcciones.

Por último, y aunque ya hayamos demostrado la parte de existencia, concluyamos con el Quinto usando al Teorema 1.17.

Teorema 1.18 *Dada una recta ℓ y un punto $\mathbf{p}$ fuera de ella en el plano, existe una única recta ℓ' que pasa por $\mathbf{p}$ y no intersecta a ℓ .*

Demostración. Existe un vector $\mathbf{u} \neq \mathbf{0}$ y una constante $c \in \mathbb{R}$, tales que ℓ está dada por la ecuación $\mathbf{u} \cdot \mathbf{x} = c$. Sea

$$\ell' : \quad \mathbf{u} \cdot \mathbf{x} = \mathbf{u} \cdot \mathbf{p}.$$

Entonces $\mathbf{p} \in \ell'$ porque satisface la ecuación, $\ell \cap \ell' = \emptyset$ porque $\mathbf{u} \cdot \mathbf{p} \neq c$ (pues $\mathbf{p} \notin \ell$), y cualquier otra recta que pasa por $\mathbf{p}$ intersecta a ℓ pues su vector normal no es paralelo a $\mathbf{u}$. $\square$

EJERCICIO 1.65 Encuentra la intersección de las rectas (a), (b), (c) y (d) del Ejercicio 1.59.

EJERCICIO 1.66 Demuestra que dos vectores $\mathbf{u}$ y $\mathbf{v}$ en $\mathbb{R}^2$ son linealmente independientes si y solo si $\mathbf{u}^\perp \cdot \mathbf{v} \neq 0$.

EJERCICIO 1.67 Demuestra que dos vectores $\mathbf{u}$ y $\mathbf{v}$ en $\mathbb{R}^2$ son linealmente independientes si y sólo si $\mathbb{R}^2 = \langle \mathbf{u}, \mathbf{v} \rangle$.

EJERCICIO 1.68 A cada terna de números reales (a, b, c) , asóciate la ecuación $ax + by = c$.

- i) ¿Cuáles son las ternas (como puntos en $\mathbb{R}^3$) a las que se le asocian rectas en $\mathbb{R}^2$ por su ecuación normal?
- ii) Dada una recta en $\mathbb{R}^2$, describe las ternas (como subconjunto de $\mathbb{R}^3$) que se asocian a ella.
- iii) Dado un haz de rectas paralelas en $\mathbb{R}^2$, describe las ternas (como subconjunto de $\mathbb{R}^3$) que se asocian a rectas de ese haz.
- iv) Dado un haz de rectas concurrentes en $\mathbb{R}^2$, describe las ternas (como subconjunto de $\mathbb{R}^3$) que se asocian a las rectas de ese haz.

1.8.2 Teoremas de concurrencia

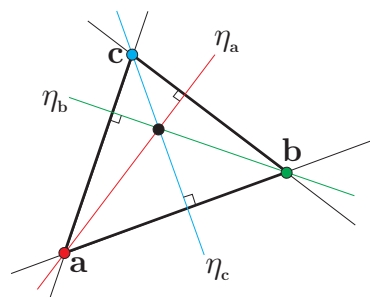
En esta sección, aplicamos la ecuación normal de las rectas para demostrar uno de los teoremas clásicos de concurrencia de líneas y dejamos el otro como ejercicio.

La *altura* de un triángulo es la recta que pasa por uno de sus vértices y es ortogonal al lado opuesto.

Teorema 1.19 *Las alturas de un triángulo son concurrentes.*

Demostración. Dado un triángulo con vértices $\mathbf{a}$, $\mathbf{b}$, $\mathbf{c}$, tenemos que la altura por el vértice $\mathbf{a}$, llamémosla η_a (“eta-sub-a”), está definida por

$$\eta_a : \quad (\mathbf{c} - \mathbf{b}) \cdot \mathbf{x} = (\mathbf{c} - \mathbf{b}) \cdot \mathbf{a},$$



pues es ortogonal al lado que pasa por $\mathbf{b}$ y $\mathbf{c}$, que tiene dirección $(\mathbf{c} - \mathbf{b})$, y pasa por el punto $\mathbf{a}$. Análogamente:

$$\begin{aligned}\eta_{\mathbf{b}} &: (\mathbf{a} - \mathbf{c}) \cdot \mathbf{x} = (\mathbf{a} - \mathbf{c}) \cdot \mathbf{b} \\ \eta_{\mathbf{c}} &: (\mathbf{b} - \mathbf{a}) \cdot \mathbf{x} = (\mathbf{b} - \mathbf{a}) \cdot \mathbf{c}.\end{aligned}$$

Obsérvese ahora que la suma (lado a lado) de dos de estas ecuaciones da precisamente el negativo de la tercera. Por ejemplo, sumando las dos primeras obtenemos

$$\begin{aligned}(\mathbf{c} - \mathbf{b}) \cdot \mathbf{x} + (\mathbf{a} - \mathbf{c}) \cdot \mathbf{x} &= (\mathbf{c} - \mathbf{b}) \cdot \mathbf{a} + (\mathbf{a} - \mathbf{c}) \cdot \mathbf{b} \\ (\mathbf{c} - \mathbf{b} + \mathbf{a} - \mathbf{c}) \cdot \mathbf{x} &= \mathbf{c} \cdot \mathbf{a} - \mathbf{b} \cdot \mathbf{a} + \mathbf{a} \cdot \mathbf{b} - \mathbf{c} \cdot \mathbf{b} \\ (\mathbf{a} - \mathbf{b}) \cdot \mathbf{x} &= (\mathbf{a} - \mathbf{b}) \cdot \mathbf{c}.\end{aligned}$$

Así que si $\mathbf{x} \in \eta_{\mathbf{a}} \cap \eta_{\mathbf{b}}$, entonces cumple las dos primeras ecuaciones y por tanto cumple su suma que es “menos” la ecuación de $\eta_{\mathbf{c}}$, y entonces $\mathbf{x} \in \eta_{\mathbf{c}}$. Así que las tres rectas pasan por el mismo punto. $\square$

EJERCICIO 1.69 Demuestra que las tres mediatrices de un triángulo son concurrentes (el punto en el que concurren se llama *circuncentro*). Donde la mediatriz de un segmento es su ortogonal que pasa por el punto medio.

EJERCICIO 1.70 Encuentra el circuncentro del triángulo con vértices $(1, 1)$, $(1, -1)$ y $(-2, -2)$. Haz el dibujo del triángulo y sus mediatrices.

1.8.3 Planos en el espacio II

Hemos definido las rectas en $\mathbb{R}^n$ por su representación paramétrica y, en $\mathbb{R}^2$, acabamos de ver que también tienen una representación normal (con base en una ecuación lineal). Es entonces importante hacer notar que este fenómeno sólo se da en dimensión 2. En $\mathbb{R}^3$, que es el otro espacio que nos interesa, los planos (¡no las rectas!) son las que se definen por la ecuación normal. Veamos esto con cuidado.

Dado un vector $\mathbf{n} \in \mathbb{R}^3$ distinto de $\mathbf{0}$; sea, para cualquier $d \in \mathbb{R}$,

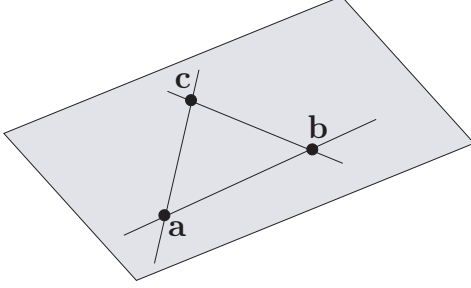
$$\Pi_d: \mathbf{n} \cdot \mathbf{x} = d$$

es decir, $\Pi_d := \{\mathbf{x} \in \mathbb{R}^3 \mid \mathbf{n} \cdot \mathbf{x} = d\}$. Llamamos a la constante d pues ahora el vector normal es de la forma $\mathbf{n} = (a, b, c)$ y entonces la ecuación anterior se escribe en coordenadas como

$$ax + by + cz = d$$

con el vector variable $\mathbf{x} = (x, y, z)$. Afirmamos que Π_d es un plano. Recuerdese que definimos plano como el conjunto de combinaciones afines (o baricéntricas) de tres puntos no colineales.

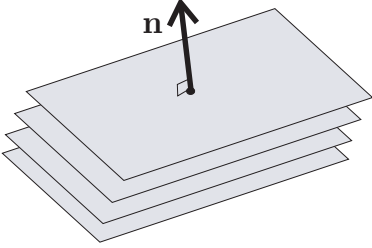
Veamos primero que si tres puntos $\mathbf{a}$, $\mathbf{b}$, $\mathbf{c}$ satisfacen la ecuación $\mathbf{n} \cdot \mathbf{x} = d$, entonces los puntos del plano que generan también la satisfacen. Para esto, tomamos $\alpha, \beta, \gamma \in \mathbb{R}$ tales que $\alpha + \beta + \gamma = 1$, y entonces se tiene



$$\begin{aligned} \mathbf{n} \cdot (\alpha \mathbf{a} + \beta \mathbf{b} + \gamma \mathbf{c}) &= \alpha (\mathbf{n} \cdot \mathbf{a}) + \beta (\mathbf{n} \cdot \mathbf{b}) + \gamma (\mathbf{n} \cdot \mathbf{c}) \\ &= \alpha d + \beta d + \gamma d \\ &= (\alpha + \beta + \gamma) d = d. \end{aligned}$$

Esto demuestra que si $\mathbf{a}, \mathbf{b}, \mathbf{c} \in \Pi_d$ entonces el plano (o si son colineales, la recta) que generan está contenido en Π_d . Faltaría ver que en Π_d hay tres puntos no colineales y luego demostrar que no hay más soluciones que las del plano que generan.

Pero mejor veámoslo desde otro punto de vista: el lineal en vez del baricéntrico. Afirmamos que los conjuntos Π_d , al variar d , son la familia de planos normales a $\mathbf{n}$. Para demostrar esto, hay que ver que Π_0 es un plano por el origen y que Π_d es un trasladado de Π_0 , es decir, Π_0 empujado por algún vector constante.



Esto último ya se hizo en esencia cuando vimos el caso de rectas en el plano, así que lo veremos primero para remarcar lo general de aquella demostración. Supongamos que $\mathbf{p} \in \mathbb{R}^3$ es tal que

$$\mathbf{n} \cdot \mathbf{p} = d,$$

es decir, que es una solución particular. Vamos a demostrar que

$$\Pi_d = \Pi_0 + \mathbf{p} := \{\mathbf{y} + \mathbf{p} \mid \mathbf{y} \in \Pi_0\}; \quad (1.10)$$

es decir, que si $\mathbf{x} \in \Pi_d$ entonces $\mathbf{x} = \mathbf{y} + \mathbf{p}$ para algún $\mathbf{y} \in \Pi_0$; y al revés, que si $\mathbf{y} \in \Pi_0$ entonces $\mathbf{x} = \mathbf{y} + \mathbf{p} \in \Pi_d$.

Dado $\mathbf{x} \in \Pi_d$ (i.e., tal que $\mathbf{n} \cdot \mathbf{x} = d$), sea $\mathbf{y} := \mathbf{x} - \mathbf{p}$. Como

$$\mathbf{n} \cdot \mathbf{y} = \mathbf{n} \cdot (\mathbf{x} - \mathbf{p}) = \mathbf{n} \cdot \mathbf{x} - \mathbf{n} \cdot \mathbf{p} = d - d = 0,$$

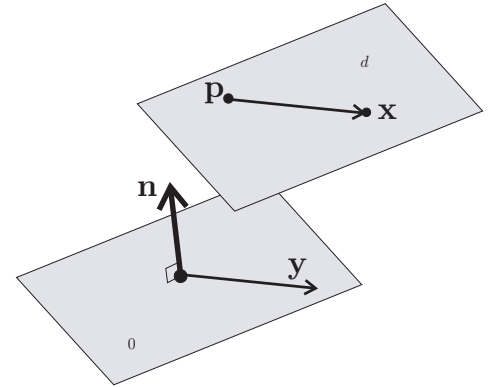
entonces $\mathbf{y} \in \Pi_0$ y es, por definición, tal que $\mathbf{x} = \mathbf{y} + \mathbf{p}$; por lo tanto $\Pi_d \subseteq \Pi_0 + \mathbf{p}$. Y al revés, si $\mathbf{y} \in \Pi_0$ y $\mathbf{x} = \mathbf{y} + \mathbf{p}$, entonces

$$\mathbf{n} \cdot \mathbf{x} = \mathbf{n} \cdot (\mathbf{y} + \mathbf{p}) = \mathbf{n} \cdot \mathbf{y} + \mathbf{n} \cdot \mathbf{p} = 0 + d = d$$

y por lo tanto $\mathbf{x} \in \Pi_d$. Hemos demostrado (1.10).

Nos falta ver que Π_0 es efectivamente un plano por el origen. Puesto que estamos suponiendo que $\mathbf{n} \neq \mathbf{0}$ entonces de la ecuación

$$ax + by + cz = 0$$



puede despejarse alguna variable en términos de las otras **dos**; sin pérdida de generalidad, digamos que es z , es decir, que $c \neq 0$. De tal forma que al dar valores arbitrarios a x y y , la fórmula nos da un valor de z , y por tanto un punto $\mathbf{x} \in \mathbb{R}^3$ que satisface la condición original $\mathbf{n} \cdot \mathbf{x} = 0$ —esto equivale a parametrizar las soluciones con $\mathbb{R}^2$, con dos grados de libertad—. En particular hay una solución $\mathbf{u}$ con $x = 1$ y $y = 0$ (a saber, $\mathbf{u} = (1, 0, -a/c)$) y una solución $\mathbf{v}$ con $x = 0$ y $y = 1$ (¿cuál es?). Como claramente los vectores $\mathbf{u}$ y $\mathbf{v}$ no son paralelos (tienen ceros en coordenadas distintas), generan linealmente todo un plano de soluciones (pues para cualquier $s, t \in \mathbb{R}$, se tiene que $\mathbf{n} \cdot (s\mathbf{u} + t\mathbf{v}) = s(\mathbf{n} \cdot \mathbf{u}) + t(\mathbf{n} \cdot \mathbf{v}) = 0$). Falta ver que no hay más soluciones que las que hemos descrito, pero esto se lo dejamos al estudiante para el siguiente ejercicio, que hay que comparar con este párrafo para entender cómo se llegó a su elegante planteamiento.

EJERCICIO 1.71 Sea $\mathbf{n} = (a, b, c)$ tal que $a \neq 0$. Demuestra que

$$\{\mathbf{x} \in \mathbb{R}^3 \mid \mathbf{n} \cdot \mathbf{x} = 0\} = \{s(-b, a, 0) + t(-c, 0, a) \mid s, t \in \mathbb{R}\}.$$

EJERCICIO 1.72 Describe el plano en $\mathbb{R}^3$ dado por la ecuación $x + y + z = 1$.

***EJERCICIO 1.73** Sea $\mathbf{n}$ un vector no nulo en $\mathbb{R}^n$.

- Demuestra que $V_0 := \{\mathbf{x} \in \mathbb{R}^n \mid \mathbf{n} \cdot \mathbf{x} = 0\}$ es un *subespacio vectorial* de $\mathbb{R}^n$; es decir, que la suma de vectores en V_0 se queda en V_0 y que el “alargamiento” de vectores en V_0 también está en V_0 .
 - Demuestra que para cualquier $d \in \mathbb{R}$, el conjunto $V_d := \{\mathbf{x} \in \mathbb{R}^n \mid \mathbf{n} \cdot \mathbf{x} = d\}$ es un trasladado de V_0 ; es decir, que existe un $\mathbf{p} \in \mathbb{R}^n$ tal que $V_d = V_0 + \mathbf{p} = \{\mathbf{y} + \mathbf{p} \mid \mathbf{y} \in V_0\}$.
 - ¿Qué dimensión dirías que tiene V_d ? Argumenta un poco tu respuesta.
-

1.9 Norma y ángulos

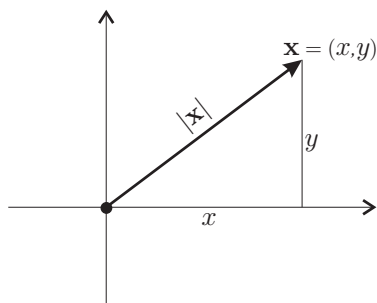
Del producto interior, obtendremos la noción de *norma* o *magnitud* de los vectores, que corresponde a la *distancia* del punto al origen.

Definición 1.9.1 Dado un vector $\mathbf{v} \in \mathbb{R}^n$, su *norma* (o *magnitud*) es el número real:

$$|\mathbf{v}| := \sqrt{\mathbf{v} \cdot \mathbf{v}},$$

de tal manera que la norma es una función de $\mathbb{R}^n$ en $\mathbb{R}$.

Como $\mathbf{v} \cdot \mathbf{v} \geq 0$ (Teorema 1.11), tiene sentido tomar su raíz cuadrada (que a su vez se define como el número positivo tal que al elevarlo al cuadrado nos da el dado). Se tiene entonces la siguiente fórmula que es una definición equivalente de la norma y que será usada con mucho más frecuencia

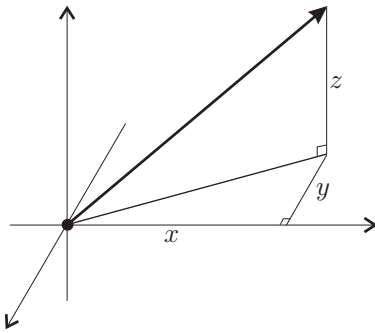


$$|\mathbf{v}|^2 = \mathbf{v} \cdot \mathbf{v}.$$

En $\mathbb{R}^2$ la norma se escribe, con coordenadas $\mathbf{v} = (x, y)$, como

$$|\mathbf{v}|^2 = x^2 + y^2.$$

Entonces $|\mathbf{v}|$ corresponde a la distancia euclidiana del origen al punto $\mathbf{v} = (x, y)$, pues de acuerdo con el Teorema de Pitágoras, x y y son lo que miden los catetos del triángulo rectángulo con hipotenusa $\mathbf{v}$. Aquí es donde resulta importante (por primera vez) que los ejes coordenados se tomen ortogonales, pues entonces la fórmula para calcular la distancia euclidiana al origen a partir de las coordenadas se hace sencilla.



En $\mathbb{R}^3$, con coordenadas $\mathbf{v} = (x, y, z)$, la norma se escribe

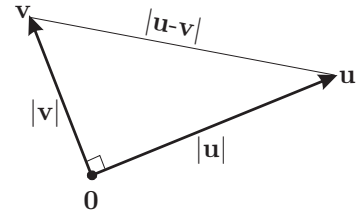
$$|\mathbf{v}|^2 = x^2 + y^2 + z^2,$$

que de nuevo, usando dos veces el Teorema de Pitágoras (en los triángulos de la figura) y que la nueva dirección z es ortogonal al plano x, y , corresponde a la *magnitud* del vector $\mathbf{v}$.

Demostremos primero el Teorema de Pitágoras para vectores en general. Este teorema fue la motivación básica para la definición de norma; sin embargo su uso ha sido sólo ése: como motivación, pues no lo hemos usado formalmente sino para ver que la norma corresponde a la noción euclidiana de distancia al origen (magnitud de vectores) cuando los ejes se toman ortogonales entre sí. Ahora veremos que al estar en el trasfondo de nuestras definiciones, éstas le hacen honor al hacerlo verdadero.

Teorema 1.20 (Pitágoras vectorial) Sean $\mathbf{u}$ y $\mathbf{v}$ dos vectores en $\mathbb{R}^n$. Entonces son perpendiculares si y sólo si

$$|\mathbf{u}|^2 + |\mathbf{v}|^2 = |\mathbf{u} - \mathbf{v}|^2.$$



Demostración. De las propiedades del producto interior (y la definición de norma) se obtiene

$$\begin{aligned} |\mathbf{u} - \mathbf{v}|^2 &= (\mathbf{u} - \mathbf{v}) \cdot (\mathbf{u} - \mathbf{v}) \\ &= \mathbf{u} \cdot \mathbf{u} + \mathbf{v} \cdot \mathbf{v} - \mathbf{u} \cdot \mathbf{v} - \mathbf{v} \cdot \mathbf{u} \\ &= \mathbf{u} \cdot \mathbf{u} + \mathbf{v} \cdot \mathbf{v} - 2(\mathbf{u} \cdot \mathbf{v}) \\ &= |\mathbf{u}|^2 + |\mathbf{v}|^2 - 2(\mathbf{u} \cdot \mathbf{v}). \end{aligned}$$

Por definición, $\mathbf{u}$ y $\mathbf{v}$ son perpendiculares si y sólo si $\mathbf{u} \cdot \mathbf{v} = 0$ y entonces el teorema se sigue de la igualdad anterior. $\square$

Así que nuestra definición de perpendicularidad (que el producto punto se anule) corresponde a que el Teorema de Pitágoras se vuelva cierto. Y entonces ahora tiene sentido nuestra manera informal de llamar, en $\mathbb{R}^3$, a las soluciones de la ecuación $\mathbf{u} \cdot \mathbf{x} = 0$ “el plano normal a $\mathbf{u}$ ”. Pues $\mathbf{u} \cdot \mathbf{x} = 0$ si y solo si los vectores $\mathbf{u}$ y $\mathbf{x}$ son los catetos de un triángulo que cumple el Teorema de Pitágoras y por tanto es rectángulo.

Las propiedades básicas de la norma se resumen en el siguiente teorema.

Teorema 1.21 *Para todos los vectores $\mathbf{v}, \mathbf{u} \in \mathbb{R}^n$ y para todo número real $t \in \mathbb{R}$ se tiene que:*

- i) $|\mathbf{v}| \geq 0$
- ii) $|\mathbf{v}| = 0 \Leftrightarrow \mathbf{v} = \mathbf{0}$
- iii) $|t\mathbf{v}| = |t| |\mathbf{v}|$
- iv) $|\mathbf{u}| + |\mathbf{v}| \geq |\mathbf{u} + \mathbf{v}|$
- v) $|\mathbf{u}| |\mathbf{v}| \geq |\mathbf{u} \cdot \mathbf{v}|$

Demostración. La primera afirmación es consecuencia inmediata de la definición de raíz cuadrada, y la segunda del Teorema 1.11. La tercera, donde también se usa ese teorema, es muy simple pero con una sutileza:

$$|t\mathbf{v}|^2 = (t\mathbf{v}) \cdot (t\mathbf{v}) = t^2(\mathbf{v} \cdot \mathbf{v}) = t^2 |\mathbf{v}|^2$$

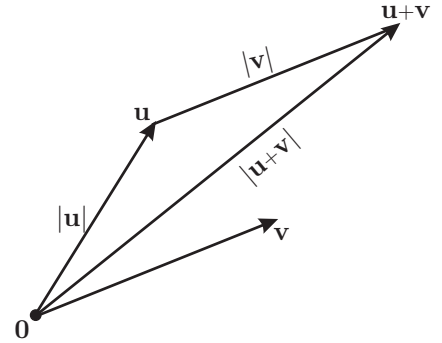
de donde, tomando raíz cuadrada, se deduce $|t\mathbf{v}| = |t| |\mathbf{v}|$. Puesto que $\sqrt{t^2}$ no siempre es t sino su *valor absoluto* $|t|$, que es positivo siempre (y que podría definirse $|t| := \sqrt{t^2}$); nótese además que como $|\mathbf{v}| \geq 0$ entonces $\sqrt{|\mathbf{v}|^2} = |\mathbf{v}|$. Obsérvese que para $n = 1$ (es decir, en $\mathbb{R}$) la norma y el valor absoluto coinciden, así que usar la misma notación para ambos no causa ningún conflicto.

El inciso (iv) se conoce como “**la desigualdad del triángulo**”, pues dice que en el triángulo con vértices $\mathbf{0}$, $\mathbf{u}$ y $\mathbf{u} + \mathbf{v}$, es más corto ir directamente de $\mathbf{0}$ a $\mathbf{u} + \mathbf{v}$ ($|\mathbf{u} + \mathbf{v}|$) que pasar primero por $\mathbf{u}$ ($|\mathbf{u}| + |\mathbf{v}|$). Para demostrarla, nótese primero que como ambos lados de la desigualdad son no negativos, entonces ésta es equivalente a que la misma desigualdad se cumpla para los cuadrados, *i.e.*,

$$|\mathbf{u}| + |\mathbf{v}| \geq |\mathbf{u} + \mathbf{v}| \quad \Leftrightarrow \quad (|\mathbf{u}| + |\mathbf{v}|)^2 \geq |\mathbf{u} + \mathbf{v}|^2.$$

Y esta última desigualdad es equivalente a que $(|\mathbf{u}| + |\mathbf{v}|)^2 - |\mathbf{u} + \mathbf{v}|^2 \geq 0$. Así que debemos desarrollar el lado izquierdo:

$$\begin{aligned} (|\mathbf{u}| + |\mathbf{v}|)^2 - |\mathbf{u} + \mathbf{v}|^2 &= (|\mathbf{u}| + |\mathbf{v}|)^2 - (\mathbf{u} + \mathbf{v}) \cdot (\mathbf{u} + \mathbf{v}) \\ &= |\mathbf{u}|^2 + |\mathbf{v}|^2 + 2|\mathbf{u}||\mathbf{v}| - (\mathbf{u} \cdot \mathbf{u} + \mathbf{v} \cdot \mathbf{v} + 2(\mathbf{u} \cdot \mathbf{v})) \\ &= 2(|\mathbf{u}||\mathbf{v}| - (\mathbf{u} \cdot \mathbf{v})). \end{aligned}$$



Queda entonces por demostrar que $|\mathbf{u}||\mathbf{v}| \geq (\mathbf{u} \cdot \mathbf{v})$, pero esto se sigue del inciso (v), que es una afirmación más fuerte que demostraremos a continuación independientemente.

El inciso (v) se conoce como “**la desigualdad de Schwartz**”. Por un razonamiento análogo al del inciso anterior, bastará demostrar que $(|\mathbf{u}||\mathbf{v}|)^2 - |\mathbf{u} \cdot \mathbf{v}|^2$ es positivo. Lo haremos para $\mathbb{R}^2$ con coordenadas, dejando el caso de $\mathbb{R}^3$, y de $\mathbb{R}^n$, como ejercicios. Supongamos entonces que $\mathbf{u} = (a, b)$ y $\mathbf{v} = (\alpha, \beta)$ para obtener

$$\begin{aligned} |\mathbf{u}|^2 |\mathbf{v}|^2 - |\mathbf{u} \cdot \mathbf{v}|^2 &= (a^2 + b^2)(\alpha^2 + \beta^2) - (a\alpha + b\beta)^2 \\ &= a^2\alpha^2 + b^2\beta^2 + a^2\beta^2 + b^2\alpha^2 \\ &\quad - (a^2\alpha^2 + b^2\beta^2 + 2a\alpha b\beta) \\ &= a^2\beta^2 - 2(a\beta)(b\alpha) + b^2\alpha^2 \\ &= (a\beta - b\alpha)^2 \geq 0; \end{aligned}$$

donde la última desigualdad se debe a que el cuadrado de cualquier número es no negativo. Lo cual demuestra la desigualdad de Schwartz, la del Triángulo y completa el Teorema. $\square$

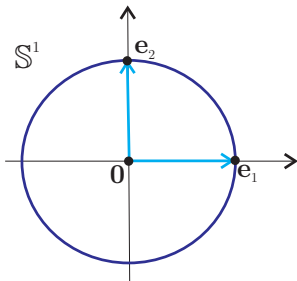
EJERCICIO 1.74 Demuestra que dos vectores $\mathbf{u}$ y $\mathbf{v}$ son perpendiculares si y sólo si $|\mathbf{u} + \mathbf{v}|^2 = |\mathbf{u}|^2 + |\mathbf{v}|^2$. Haz el dibujo.

EJERCICIO 1.75 Demuestra que dos vectores $\mathbf{u}$ y $\mathbf{v}$ son perpendiculares si y sólo si $|\mathbf{u} + \mathbf{v}| = |\mathbf{u} - \mathbf{v}|$. Haz el dibujo. Observa que este enunciado corresponde a que un paralelogramo es un rectángulo si y sólo si sus diagonales miden lo mismo.

EJERCICIO 1.76 Sean $\mathbf{u}$ y $\mathbf{v}$ dos vectores no nulos. Demuestra que tienen la misma norma si y sólo si $\mathbf{u} + \mathbf{v}$ y $\mathbf{u} - \mathbf{v}$ son ortogonales.

EJERCICIO 1.77 Demuestra la desigualdad de Schwartz en $\mathbb{R}^3$. ¿Puedes dar, o describir, la demostración en $\mathbb{R}^n$?

1.9.1 El círculo unitario



A los vectores que tienen norma igual a uno, se les llama *vectores unitarios*. Y al conjunto de todos los vectores unitarios en $\mathbb{R}^2$ se le llama el *círculo unitario* y se denota con $\mathbb{S}^1$. La notación viene de la palabra *sphere* en inglés, que significa esfera; pues en general se puede definir a la *esfera de dimensión n* como

$$\mathbb{S}^n = \{ \mathbf{x} \in \mathbb{R}^{n+1} \mid |\mathbf{x}| = 1 \}.$$

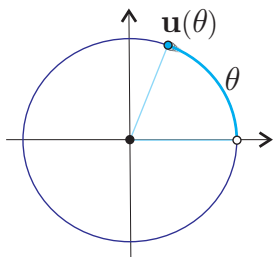
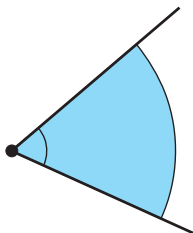
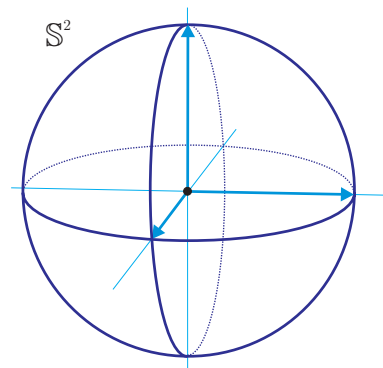
Nótese que el exponente se refiere a la dimensión de la esfera en sí, y que ésta necesita una dimensión más para “vivir”. Así, la esfera de dimensión 2 vive en $\mathbb{R}^3$, y es la representación abstracta de las pompas de jabón. Más adelante la estudiaremos con cuidado. Por lo pronto nos interesa el círculo unitario.

Los puntos de $\mathbb{S}^1$ corresponden a los *ángulos*, y a estos los mediremos con radianes. La definición formal o muy precisa de ángulo involucra necesariamente nociones de cálculo que no entran en este libro. Sin embargo, es un concepto muy intuitivo y de esa intuición nos valdremos. Un *ángulo* es un sector radial del círculo unitario, aunque se denota gráficamente con un arco pequeño cerca del centro para indicar que no depende realmente del círculo de referencia, sino que es más bien un sector de cualquier círculo concéntrico. Es costumbre partir el círculo completo en 360 sectores iguales llamados *grados*. Resulta conveniente usar el número 360, pues $360 = 2^3 3^2 5$ tiene muchos divisores, así que el círculo se puede partir en cuatro sectores iguales de 90 grados (denotado 90° y llamado *ángulo recto*) o en 12 de 30° que corresponden a las horas del reloj, etc. Pero la otra manera de medirlos, que no involucra la convención de escoger 360 para la vuelta entera, es por la longitud del arco de círculo que abarcan; a esta medida se le conoce como

en radianes. Un problema clásico que resolvieron los griegos con admirable precisión fue medir la *circunferencia* del círculo unitario (de radio uno), es decir, cuánto mide un hilo “untado” en el círculo. El resultado, como todos sabemos, es el doble del famoso número $\pi = 3.14159\dots$, donde los puntos suspensivos indican que su expresión decimal sigue infinitamente pues no es racional. Otra manera de entender los radianes es cinemática y es la que usaremos a continuación para establecer notación; es la versión teórica del movimiento de la piedra en una honda justo antes de lanzarla.

Si una partícula viaja dentro de $\mathbb{S}^1$ a velocidad constante 1, partiendo del punto $\mathbf{e}_1 = (1, 0)$ y en dirección contraria a las manecillas del reloj (es decir, saliendo hacia arriba con vector velocidad $\mathbf{e}_2 = (0, 1)$), entonces en un tiempo θ estará en un punto que llamaremos $\mathbf{u}(\theta)$; en el tiempo $\pi/2$ estará en el $\mathbf{e}_2 = (0, 1)$ (es decir, $\mathbf{u}(\pi/2) = \mathbf{e}_2$), en el tiempo π en el $(-1, 0)$ y en el 2π estará de regreso para empezar de nuevo. El tiempo aquí, que estamos midiendo con el parametro θ , corresponde precisamente a los *radianes*, pues al viajar a velocidad constante 1 el tiempo es igual a la distancia recorrida. Podemos también dar sentido a $\mathbf{u}(\theta)$ para θ negativa pensando que la partícula viene dando vueltas desde siempre (tiempo infinito negativo) de tal manera que en el tiempo 0 pasa justo por $\mathbf{e}_1$, es decir, que

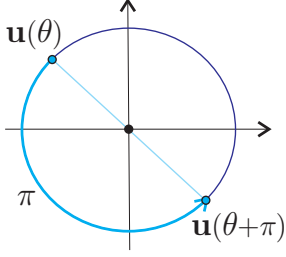
$$\mathbf{u}(0) = (1, 0).$$



Nuestra suposición básica es que la función $\mathbf{u} : \mathbb{R} \rightarrow \mathbb{S}^1$, que hemos descrito, está bien definida. Se cumple entonces que

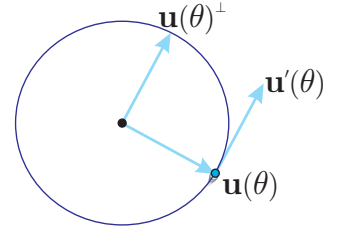
$$\mathbf{u}(\theta) = \mathbf{u}(\theta + 2\pi),$$

pues en el tiempo 2π la partícula da justo una vuelta (la longitud del círculo mide lo mismo independientemente de dónde empecemos a medir). Y también que $\mathbf{u}(\theta + \pi) = -\mathbf{u}(\theta)$, pues π es justo el ángulo que da media vuelta.



Observemos ahora que al pedir que la partícula viaje en el círculo unitario con velocidad constante 1 entonces su *vector velocidad*, que incluye ahora dirección además de magnitud, es tangente a $\mathbb{S}^1$ (la piedra de la honda sale por la tangente), y es fácil ver que entonces es perpendicular al vector de posición $\mathbf{u}(\theta)$, pero más precisamente es justo su compadre ortogonal porque va girando en “su” dirección, “hacia él”. Usando la notación de derivadas, podríamos escribir

$$\mathbf{u}'(\theta) = \mathbf{u}(\theta)^\perp.$$



Hay que remarcar que al medir ángulos con radianes, si bien ganamos en naturalidad matemática, es inevitable la ambigüedad de que a los ángulos no corresponda un número único. Pues θ y $\theta + 2n\pi$, para cualquier $n \in \mathbb{Z}$, determinan al mismo ángulo, es decir $\mathbf{u}(\theta) = \mathbf{u}(\theta + 2n\pi)$. Podemos exigir que θ esté en el intervalo entre 0 y 2π , o bien, que resulta más agradable geométricamente, en el intervalo de $-\pi$ a π , para reducir la ambigüedad sólo a los extremos. Pero al sumar o restar ángulos, inevitablemente nos saldremos de este intervalo y habrá que volver a ajustar.

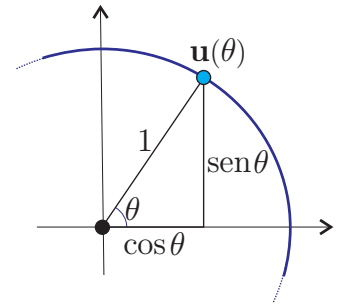
Funciones trigonométricas

Recordemos ahora que $\mathbb{S}^1$ es un subconjunto de $\mathbb{R}^2$, así que el punto $\mathbf{u}(\theta)$ tiene dos coordenadas precisas. Éstas están dadas por las importantísimas *funciones trigonométricas* coseno y seno. Más precisamente, podemos definir

$$(\cos \theta, \sin \theta) := \mathbf{u}(\theta),$$

es decir, $\cos \theta$ y $\sin \theta$ son las coordenadas cartesianas del punto $\mathbf{u}(\theta) \in \mathbb{S}^1$. Entonces el coseno y el seno corresponden respectivamente al cateto adyacente y al cateto opuesto de un triángulo rectángulo con hipotenusa 1 y ángulo θ , como se ve en trigonometría de secundaria. También pensarse que si damos por establecidas las funciones seno y coseno, entonces la función $\mathbf{u}(\theta)$ está dada por la ecuación anterior. Hay que remarcar que la pertenencia $\mathbf{u}(\theta) \in \mathbb{S}^1$ equivale entonces a la ecuación

$$\sin^2 \theta + \cos^2 \theta = 1.$$

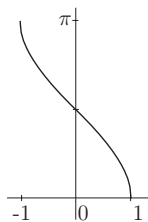
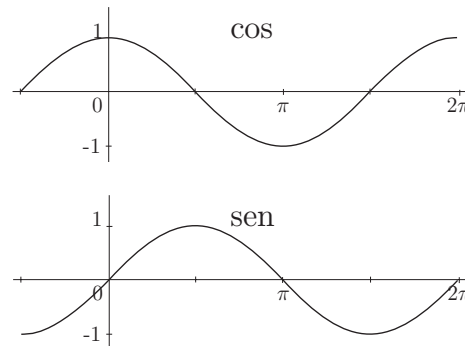


Las propiedades de la función $\mathbf{u}$ que hemos enumerado, traducidas a sus coordenadas (por ejemplo, la del vector velocidad se traduce a $\cos' \theta = -\sin \theta$ y $\sin' \theta = \cos \theta$) son suficientes para definir a las funciones $\cos$ y $\sin$, y en adelante las daremos por un hecho. Sus bien conocidas gráficas aparecen en la figura.

Observemos, por último, que cualquiera de las funciones coseno o seno casi determinan al ángulo, pues de la ecuación anterior se pueden despejar con la ambigüedad de un signo al tomar raíz cuadrada, es decir,

$$\sin \theta = \pm \sqrt{1 - \cos^2 \theta}.$$

Dicho de otra manera, dado x en el intervalo de -1 a 1 (escrito $x \in [-1, 1]$,⁴ o bien $-1 \leq x \leq 1$) tenemos dos posibles puntos en $\mathbb{S}^1$ con esa ordenada, a saber $(x, \sqrt{1-x^2})$ y $(x, -\sqrt{1-x^2})$. Entonces con la ambigüedad de un signo podemos determinar el ángulo del cuál x es el coseno. Si escogemos la parte superior del círculo unitario obtenemos una función, llamada *arcocoseno* y leída “el arco cuyo coseno es ...”



$$\arccos : [-1, 1] \rightarrow [0, \pi],$$

y definida por

$$\cos(\arccos(x)) = x.$$

La ambigüedad surge al componer en el otro sentido, pues para $\theta \in [-\pi, \pi]$ se tiene $\arccos(\cos(\theta)) = \pm\theta$. Obsérvese que entonces la mitad superior del círculo unitario se describe como la gráfica de la función $x \mapsto \sin(\arccos(x))$.

EJERCICIO 1.78 Demuestra (usando que definimos $\mathbf{u}(\theta) \in \mathbb{S}^1$) que para cualquier $\theta \in \mathbb{R}$ se cumplen

$$\begin{aligned} -1 &\leq \cos \theta \leq 1 \\ -1 &\leq \sin \theta \leq 1 \end{aligned}$$

EJERCICIO 1.79 Demuestra que para cualquier $\theta \in \mathbb{R}$ se cumplen

$$\begin{aligned} \cos(\theta + \pi/2) &= -\sin \theta & \cos(\theta + \pi) &= -\cos \theta & \cos(-\theta) &= \cos \theta \\ \sin(\theta + \pi/2) &= \cos \theta & \sin(\theta + \pi) &= -\sin \theta & \sin(-\theta) &= -\sin \theta \end{aligned}$$

EJERCICIO 1.80 Construye una tabla con los valores de las funciones $\cos$ y $\sin$ para los ángulos $0, \pi/6, \pi/4, \pi/3, \pi/2, 2\pi/3, 3\pi/4, 5\pi/6, \pi, 5\pi/4, 3\pi/2, -\pi/3, -\pi/4$, y dibuja los correspondientes vectores unitarios.

⁴Dados $a, b \in \mathbb{R}$, con $a \leq b$, denotamos por $[a, b] := \{x \in \mathbb{R} \mid a \leq x \leq b\}$ al *intervalo cerrado* entre a y b .

EJERCICIO 1.81 ¿Cuál es el conjunto de números que usan las coordenadas de las horas del reloj?

1.9.2 Coordenadas polares

El círculo unitario $\mathbb{S}^1$ (cuyos puntos hemos identificado con los ángulos tomando su ángulo con el vector base $\mathbf{e}_1 = (1, 0)$) tiene justo un representante de cada posible dirección en $\mathbb{R}^2$, donde ahora una dirección y su opuesta son diferentes (aunque sean, según nuestra definición, paralelas). De tal manera que a cualquier punto de $\mathbb{R}^2$ que no sea el origen se llega viajando en una (y exactamente una) de estas direcciones. Ésta es la idea central de las coordenadas polares, que en muchas situaciones son más naturales, o útiles, para identificar los puntos del plano. Por ejemplo, son las que intuitivamente usa un cazador: apuntar —decidir una dirección— es la primera coordenada y luego, dependiendo de la distancia a la que vuela el pato —la segunda coordenada— tiene que ajustar el tiro para que las trayectorias de pato y perdigones se intersecten.

Sea $\mathbf{x}$ cualquier vector no nulo en $\mathbb{R}^2$. Entonces $|\mathbf{x}| \neq 0$ y tiene sentido tomar el vector $(|\mathbf{x}|^{-1})\mathbf{x}$, que es unitario pues como $|\mathbf{x}| > 0$ implica que $|\mathbf{x}|^{-1} > 0$, entonces tenemos

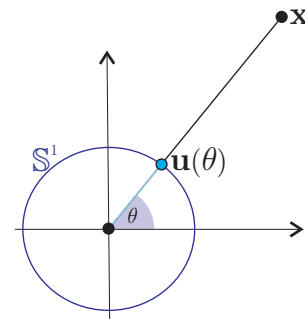
$$| |\mathbf{x}|^{-1} \mathbf{x} | = | |\mathbf{x}|^{-1} | |\mathbf{x}| = |\mathbf{x}|^{-1} |\mathbf{x}| = 1.$$

De aquí que exista un único $\mathbf{u}(\theta) \in \mathbb{S}^1$ (aunque θ no es único como real, sí lo es como ángulo), tal que

$$\mathbf{u}(\theta) = |\mathbf{x}|^{-1} \mathbf{x}$$

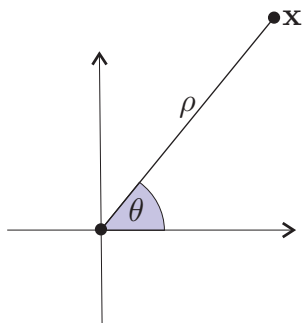
y claramente se cumple que

$$\mathbf{x} = |\mathbf{x}| \mathbf{u}(\theta).$$



A $\mathbf{u}(\theta)$ se le referirá como la *dirección* de $\mathbf{x}$. A la pareja $(\theta, |\mathbf{x}|)$ se le llama las *coordenadas polares* del vector $\mathbf{x}$. Se usa el término “coordenadas” pues determinan al vector, es decir, si nos dan una pareja (θ, ρ) con $\rho \geq 0$, obtenemos un vector

$$\mathbf{x} = \rho \mathbf{u}(\theta)$$



con magnitud ρ y dirección $\mathbf{u}(\theta)$, o bien, *ángulo* θ . Obsérvese que sólo con el origen hay ambigüedad: si $\rho = 0$ en la fórmula anterior, para cualquier valor de θ obtenemos el origen; éste no tiene un ángulo definido. Así que al hablar de coordenadas polares supondremos que el vector en cuestión es no nulo. Por nuestra convención de $\mathbf{u}(\theta)$, lo que representa θ es el ángulo respecto al eje x en su dirección positiva.

EJERCICIO 1.82 Da las coordenadas (cartesianas) de los vectores cuyas coordenadas polares son $(\pi/2, 3)$, $(\pi/4, \sqrt{2})$, $(\pi/6, 2)$, $(\pi/6, 1)$.

EJERCICIO 1.83 Si tomamos (θ, ρ) como coordenadas polares, describe los subconjuntos de $\mathbb{R}^2$ definidos por las ecuaciones $\theta = \text{cte}$ (“ θ igual a una constante”) y $\rho = \text{cte}$.

1.9.3 Ángulo entre vectores

Podemos usar coordenadas polares para definir el ángulo entre vectores en general.

Definición 1.9.2 Dados $\mathbf{x}$ y $\mathbf{y}$, dos vectores no nulos en $\mathbb{R}^2$, sean $(\alpha, |\mathbf{x}|)$ y $(\beta, |\mathbf{y}|)$ sus coordenadas polares respectivamente. El *ángulo de $\mathbf{x}$ a $\mathbf{y}$* es

$$\overrightarrow{\text{ang}}(\mathbf{x}, \mathbf{y}) = \beta - \alpha.$$

Obsérvese que el ángulo tiene dirección (de ahí que le hayamos puesto la flechita de sombrero), ya que claramente

$$\overrightarrow{\text{ang}}(\mathbf{x}, \mathbf{y}) = -\overrightarrow{\text{ang}}(\mathbf{y}, \mathbf{x}).$$

Pues $\overrightarrow{\text{ang}}(\mathbf{x}, \mathbf{y})$ corresponde al movimiento angular que nos lleva de la dirección de $\mathbf{x}$ a la de $\mathbf{y}$, y $\overrightarrow{\text{ang}}(\mathbf{y}, \mathbf{x})$ es justo el movimiento inverso. Así que podemos resumir diciendo que las coordenadas polares de un vector $\mathbf{x}$ (recuérdese que al aplicar el término ya suponemos que es no nulo) es la pareja

$$(\overrightarrow{\text{ang}}(\mathbf{e}_1, \mathbf{x}), |\mathbf{x}|).$$

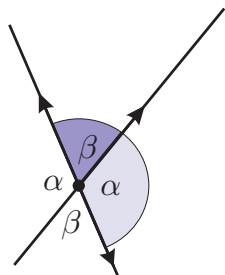
Con mucha frecuencia van a aparecer ángulos no dirigidos, pensados como un sector del círculo unitario sin principio ni fin determinados. Para ellos usaremos la notación $\text{ang}(\mathbf{x}, \mathbf{y})$ (sin la flecha), llamándolo el *ángulo entre $\mathbf{x}$ y $\mathbf{y}$* . Aunque intuitivamente es claro lo que queremos (piénsese en una porción de una pizza), su definición es un poco más laboriosa. Si tomamos dos vectores unitarios en $\mathbb{S}^1$, estos parten el círculo en dos sectores, uno de ellos es menor que medio círculo y ése es el ángulo entre ellos. Así que debemos pedir que

$$0 \leq \text{ang}(\mathbf{x}, \mathbf{y}) \leq \pi.$$

Y es claro que escogiendo a α y a β (los ángulos de $\mathbf{x}$ y $\mathbf{y}$, respectivamente) de manera adecuada se tiene que $\text{ang}(\mathbf{x}, \mathbf{y}) := |\beta - \alpha| \in [0, \pi]$. Así que, por ejemplo, $\text{ang}(\mathbf{x}, \mathbf{y}) = 0$ significa que $\mathbf{x}$ y $\mathbf{y}$ apuntan en la misma dirección, mientras que $\text{ang}(\mathbf{x}, \mathbf{y}) = \pi$ significa que $\mathbf{x}$ y $\mathbf{y}$ apuntan en direcciones contrarias, y se tiene que $\text{ang}(\mathbf{x}, \mathbf{y}) = \pi/2$ cuando son perpendiculares.

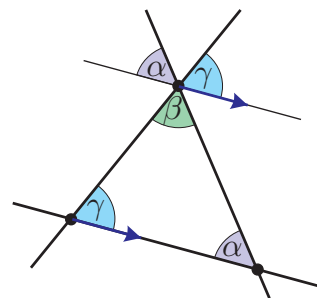
EJERCICIO 1.84 ¿Cuál es el ángulo de $\mathbf{x}$ a $\mathbf{y}$, cuando a) $\mathbf{x} = (1, 0)$, $\mathbf{y} = (0, -2)$; b) $\mathbf{x} = (1, 1)$, $\mathbf{y} = (0, 1)$; c) $\mathbf{x} = (\sqrt{3}, 2)$, $\mathbf{y} = (1, 0)$.

Suma de ángulos



Si ya definimos ángulo entre vectores, resulta natural definir *ángulo entre dos rectas* como el ángulo entre sus vectores direccionales. Pero hay una ambigüedad pues podemos escoger direcciones opuestas para una misma recta. Esta ambigüedad equivale a que dos rectas que se intersectan definen cuatro sectores o “ángulos” que se agrupan naturalmente en dos parejas opuestas por el vértice que miden lo mismo; y estos dos ángulos son *complementarios*, es decir, suman π . No vale la pena entrar en los detalles formales de esto usando vectores direccionales pues sólo se confunde lo obvio, que desde muy temprana edad sabemos. Pero sí podemos delinear rápidamente la demostración de que la suma de los ángulos de un triángulo siempre es π , que para Euclides tuvo que ser axioma (disfrazado o no, pues recuérdese que es equivalente al Quinto).

Un triángulo son tres rectas que se intersectan dos a dos (pero no las tres), y en cada vértice (la intersección de dos de las rectas) define un ángulo interno; llamémoslos α , β y γ . Al trazar la paralela a un lado del triángulo que pasa por el vértice opuesto (usando el Quinto), podemos medir ahí los tres ángulos, pues por nuestra demostración del Quinto esta paralela tiene el mismo vector direccional que el lado original, y nuestra definición de ángulo entre líneas usa la de ángulo entre vectores direccionales. Finalmente, hay que observar que si en ese vértice cambiamos uno de los ángulos por su opuesto, los tres ángulos se juntan para formar el ángulo de un vector a su opuesto, es decir, suman π . Hemos delineado la demostración del siguiente teorema clásico.

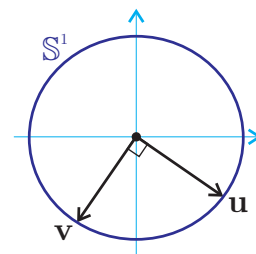


Teorema 1.22 *Los ángulos internos de un triángulo suman π .* □

1.10 Bases ortonormales

Definición 1.10.1 Una *base ortonormal* de $\mathbb{R}^2$ es una pareja de vectores unitarios perpendiculares.

Por ejemplo, la *base canónica* $\mathbf{e}_1 = (1, 0)$, $\mathbf{e}_2 = (0, 1)$ es una base ortonormal, y para cualquier $\mathbf{u} \in \mathbb{S}^1$ se tiene que $\mathbf{u}, \mathbf{u}^\perp$ es una base ortonormal, pues también $|\mathbf{u}^\perp| = 1$. Obsérvese además que si $\mathbf{u}$ y $\mathbf{v}$ son una base ortonormal entonces ambos están en $\mathbb{S}^1$ y además $\mathbf{v} = \mathbf{u}^\perp$ o $\mathbf{v} = -\mathbf{u}^\perp$. Por el momento, su importancia radica en que es muy fácil escribir cualquier vector como combinación lineal de una base ortonormal:



Teorema 1.23 Sean $\mathbf{u}$ y $\mathbf{v}$ una base ortonormal de $\mathbb{R}^2$, entonces para cualquier $\mathbf{x} \in \mathbb{R}^2$ se tiene que

$$\mathbf{x} = (\mathbf{x} \cdot \mathbf{u}) \mathbf{u} + (\mathbf{x} \cdot \mathbf{v}) \mathbf{v}.$$

Demostración. Suponemos que $\mathbf{x}$ está dado, y vamos a resolver el sistema de ecuaciones

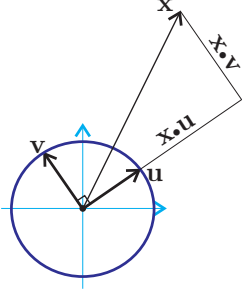
$$s \mathbf{u} + t \mathbf{v} = \mathbf{x} \quad (1.11)$$

con incógnitas t, s . Tomando el producto interior con $\mathbf{u}$ en la ecuación anterior obtenemos

$$s (\mathbf{u} \cdot \mathbf{u}) + t (\mathbf{v} \cdot \mathbf{u}) = \mathbf{x} \cdot \mathbf{u}.$$

Pero $\mathbf{u} \cdot \mathbf{u} = 1$ pues $\mathbf{u}$ es unitario y $\mathbf{v} \cdot \mathbf{u} = 0$ ya que $\mathbf{u}$ y $\mathbf{v}$ son perpendiculares, así que

$$s = \mathbf{x} \cdot \mathbf{u}.$$



Análogamente, tomando el producto interior por $\mathbf{v}$, se obtiene que $t = \mathbf{x} \cdot \mathbf{v}$. De tal manera que el sistema (1.11) sí tiene solución y es la que asegura el teorema. $\square$

Vale la pena observar que el truco que acabamos de usar para resolver un sistema de dos ecuaciones con dos incógnitas no es nuevo, lo usamos en la Sección 1.7.1 al tomar el producto interior con el compadre ortogonal; la diferencia es que aquí el otro vector de la base ortonormal está dado.

Como corolario de este teorema obtenemos la interpretación geométrica del producto interior de vectores unitarios.

Corolario 1.24 Si $\mathbf{u}, \mathbf{v} \in \mathbb{S}^1$ y α es el ángulo entre ellos, i.e., $\alpha = \text{ang}(\mathbf{u}, \mathbf{v})$, entonces

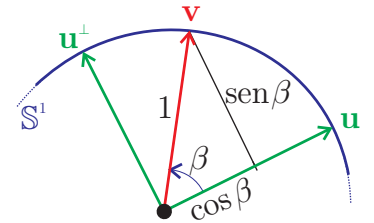
$$\mathbf{u} \cdot \mathbf{v} = \cos \alpha.$$

Demostración. Sea $\beta = \overrightarrow{\text{ang}}(\mathbf{u}, \mathbf{v})$ el ángulo de $\mathbf{u}$ a $\mathbf{v}$, de tal manera que, escogiéndolos adecuadamente, se tiene que $\alpha = \pm\beta$. Usando el teorema anterior con la base ortonormal $\mathbf{u}, \mathbf{u}^\perp$ y el vector $\mathbf{v}$ (en vez de $\mathbf{x}$) se obtiene

$$\mathbf{v} = (\mathbf{v} \cdot \mathbf{u}) \mathbf{u} + (\mathbf{v} \cdot \mathbf{u}^\perp) \mathbf{u}^\perp.$$

Pero también es claro geométricamente que

$$\mathbf{v} = \cos \beta \mathbf{u} + \sin \beta \mathbf{u}^\perp.$$



Puesto que $\cos \beta = \cos \alpha$, pues $\alpha = \pm\beta$, el corolario se sigue de que la solución de un sistema con determinante no cero es única (en nuestro caso el determinante del sistema $\mathbf{v} = s \mathbf{u} + t \mathbf{u}^\perp$ es $\det(\mathbf{u}, \mathbf{u}^\perp) = \mathbf{u}^\perp \cdot \mathbf{u} = 1$). $\square$

EJERCICIO 1.85 Demuestra que si $\mathbf{u}$, $\mathbf{v}$ y $\mathbf{w}$ son una *base ortonormal* de $\mathbb{R}^3$ (es decir, si los tres son unitarios y dos a dos son perpendiculares), entonces para cualquier $\mathbf{x} \in \mathbb{R}^3$ se tiene que

$$\mathbf{x} = (\mathbf{x} \cdot \mathbf{u}) \mathbf{u} + (\mathbf{x} \cdot \mathbf{v}) \mathbf{v} + (\mathbf{x} \cdot \mathbf{w}) \mathbf{w}.$$

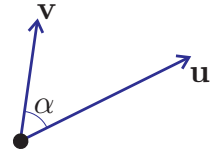
EJERCICIO 1.86 Escribe los vectores $(1, 0)$, $(0, 1)$, $(2, 1)$ y $(-1, 3)$ como combinación lineal de $\mathbf{u} = (3/5, 4/5)$ y $\mathbf{v} = (4/5, -3/5)$, es decir, como $t\mathbf{u} + s\mathbf{v}$ con t y s apropiadas.

1.10.1 Fórmula geométrica del producto interior

Como consecuencia del corolario anterior obtenemos un resultado importante que nos dice que el producto interior puede definirse en términos de la norma y el ángulo entre vectores.

Teorema 1.25 Sean $\mathbf{u}$ y $\mathbf{v}$ dos vectores en $\mathbb{R}^2$ y α el ángulo entre ellos, entonces

$$\mathbf{u} \cdot \mathbf{v} = |\mathbf{u}| |\mathbf{v}| \cos \alpha.$$



Demostración. Si $\mathbf{u} = \mathbf{0}$ (o $\mathbf{v} = \mathbf{0}$) la igualdad anterior se cumple (pues ambos lados son 0); podemos suponer entonces que ambos son distintos de cero. Ahora podemos reescalar a $\mathbf{u}$ (y a $\mathbf{v}$) para ser unitarios (pues tiene sentido $(|\mathbf{u}|^{-1})\mathbf{u}$ que es un vector unitario) manteniendo el ángulo entre ellos. Y por el Corolario 1.24,

$$\cos \alpha = \left(\frac{1}{|\mathbf{u}|}\right)\mathbf{u} \cdot \left(\frac{1}{|\mathbf{v}|}\right)\mathbf{v} = \frac{\mathbf{u} \cdot \mathbf{v}}{|\mathbf{u}| |\mathbf{v}|},$$

de donde se sigue inmediatamente el teorema. $\square$

De esta fórmula geométrica para el producto interior, se obtiene que $|\mathbf{u} \cdot \mathbf{v}| = |\cos \alpha| |\mathbf{u}| |\mathbf{v}|$, así que la desigualdad de Schwartz ($|\mathbf{u} \cdot \mathbf{v}| \leq |\mathbf{u}| |\mathbf{v}|$) es equivalente a $|\cos \alpha| \leq 1$.

EJERCICIO 1.87 Demuestra la *ley de los cosenos* para vectores, es decir, que dados dos vectores $\mathbf{u}, \mathbf{v} \in \mathbb{R}^2$ con ángulo α entre ellos, se cumple que

$$|\mathbf{u} - \mathbf{v}|^2 = |\mathbf{u}|^2 + |\mathbf{v}|^2 - 2 |\mathbf{u}| |\mathbf{v}| \cos \alpha.$$

Observa que el Teorema de Pitágoras es un caso especial.

EJERCICIO 1.88 Encuentra la distancia de los puntos $\mathbf{p}_1 = (0, 5)$, $\mathbf{p}_2 = (4, 0)$ y $\mathbf{p}_3 = (-1, 1)$ a la recta $\ell = \{t(4, 3) \mid t \in \mathbb{R}\}$.

1.10.2 El caso general

Puesto que el Teorema 1.23 y su Corolario 1.24 se demostraron para $\mathbb{R}^2$, puede dudarse de la validez de los resultados en el caso general de $\mathbb{R}^n$, pero nos interesa establecer el caso de $n = 3$. Vale la pena entonces hacer unas aclaraciones y ajustes a manera de repaso. Lo que realmente se usó del Teorema 1.23 fue que cuando $\mathbf{u}$ es un vector unitario (en cualquier dimensión ahora) entonces $\mathbf{x} \cdot \mathbf{u}$ tiene un significado geométrico muy preciso: es lo que hay que viajar en la dirección $\mathbf{u}$ para que de ahí, $\mathbf{x}$ quede en dirección ortogonal a $\mathbf{u}$. Podemos precisarlo de la siguiente manera.

Lema 1.26 Sea $\mathbf{u} \in \mathbb{R}^n$ un vector unitario (es decir, tal que $|\mathbf{u}| = 1$). Entonces para cualquier $\mathbf{x} \in \mathbb{R}^n$ existe un vector $\mathbf{y}$ perpendicular a $\mathbf{u}$ y tal que

$$\mathbf{x} = (\mathbf{x} \cdot \mathbf{u}) \mathbf{u} + \mathbf{y}.$$

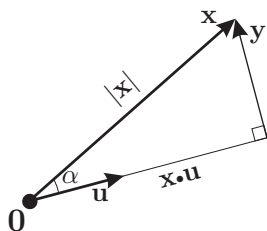
Demostración. Es muy fácil, declaremos $\mathbf{y} := \mathbf{x} - (\mathbf{x} \cdot \mathbf{u}) \mathbf{u}$ y basta ver que $\mathbf{y} \cdot \mathbf{u} = 0$:

$$\begin{aligned} \mathbf{y} \cdot \mathbf{u} &= (\mathbf{x} - (\mathbf{x} \cdot \mathbf{u}) \mathbf{u}) \cdot \mathbf{u} \\ &= \mathbf{x} \cdot \mathbf{u} - (\mathbf{x} \cdot \mathbf{u})(\mathbf{u} \cdot \mathbf{u}) = \mathbf{x} \cdot \mathbf{u} - \mathbf{x} \cdot \mathbf{u} = 0. \end{aligned}$$

□

Si pensamos ahora en el plano generado por $\mathbf{u}$ y $\mathbf{x}$ (en el que también se encuentra $\mathbf{y}$) y dentro de él en el triángulo rectángulo $\mathbf{0}, (\mathbf{x} \cdot \mathbf{u}) \mathbf{u}, \mathbf{x}$, se obtiene que

$$\cos \alpha = \frac{\mathbf{x} \cdot \mathbf{u}}{|\mathbf{x}|},$$



donde α es el ángulo entre los vectores $\mathbf{u}$ y $\mathbf{x}$. Y de aquí, el Teorema 1.25 para $\mathbb{R}^n$ se sigue inmediatamente al permitir que $\mathbf{u}$ sea no necesariamente unitario.

Si nos ponemos quisquillosos con la formalidad, se notará una falla en lo anterior (pasamos un *strike* al lector), pues no hemos definido “ángulo entre vectores” más allá de $\mathbb{R}^2$. Tómese entonces como la motivación intuitiva para la definición general que hará que todo cuadre bien.

Definición 1.10.2 Dados dos vectores no nulos $\mathbf{u}$ y $\mathbf{v}$ en $\mathbb{R}^n$ el *ángulo entre ellos* es

$$\text{ang}(\mathbf{u}, \mathbf{v}) = \arccos \left(\frac{\mathbf{u} \cdot \mathbf{v}}{|\mathbf{u}| |\mathbf{v}|} \right).$$

Nótese que por la desigualdad de Schwartz (dejada al lector para $\mathbb{R}^3$ en el Ejercicio 1.77) el argumento de la función arccos está en el intervalo $[-1, 1]$, así que está bien definido el ángulo, y queda en el intervalo $[0, \pi]$.

Hay que remarcar que a partir de $\mathbb{R}^3$ ya no tiene sentido hablar de ángulos dirigidos; esto sólo se puede hacer en el plano. Pues si tomamos dos vectores no nulos

$\mathbf{u}$ y $\mathbf{v}$ en $\mathbb{R}^3$, no hay manera de decir si el viaje angular de $\mathbf{u}$ hacia $\mathbf{v}$ es positivo o negativo. Esto dependería de escoger un lado del plano para “verlo” y ver si va contra o con el reloj, pero cualquiera de los dos lados son “iguales” y desde cada uno se “ve” lo contrario del otro; no hay una manera coherente de decidir. Lo que sí se puede decir es cuándo *tres* vectores están orientados positivamente, pero esto se verá mucho más adelante. De tal manera que a partir de $\mathbb{R}^3$, sólo el ángulo *entre* vectores está definido, y tiene valores entre 0 y π ; donde los extremos corresponden a paralelismo pero con la misma dirección (ángulo 0) o la contraria (ángulo π).

EJERCICIO 1.89 Sean $\mathbf{u}$ y $\mathbf{v}$ dos vectores unitarios y ortogonales en $\mathbb{R}^3$. Demuestra que para cualquier $\mathbf{x} \in \mathbb{R}^3$ existe un único $\mathbf{w} \in \langle \mathbf{u}, \mathbf{v} \rangle$ (el plano generado por $\mathbf{u}$ y $\mathbf{v}$) tal que $\mathbf{x} - \mathbf{w}$ es ortogonal a $\mathbf{u}$ y a $\mathbf{v}$; al punto $\mathbf{w}$ se le llama la *proyección ortogonal* de $\mathbf{x}$ al plano $\langle \mathbf{u}, \mathbf{v} \rangle$. ¿Puedes concluir que la distancia del punto $\mathbf{x}$ al plano $\langle \mathbf{u}, \mathbf{v} \rangle$ es

$$\sqrt{|\mathbf{x}|^2 - (\mathbf{x} \cdot \mathbf{u})^2 - (\mathbf{x} \cdot \mathbf{v})^2} ?$$

1.11 Distancia

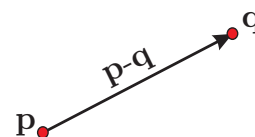
Esta sección cierra el capítulo con una breve discusión sobre el concepto euclidiano de distancia, que se deduce naturalmente de la norma

Dados dos puntos $\mathbf{p}, \mathbf{q} \in \mathbb{R}^n$ se puede definir su *distancia euclidiana*, o simplemente su *distancia*, $d(\mathbf{p}, \mathbf{q})$, como la norma de su diferencia, es decir, como la magnitud del vector que lleva a uno en otro, *i.e.*,

$$d(\mathbf{p}, \mathbf{q}) = |\mathbf{p} - \mathbf{q}|,$$

que explícitamente en coordenadas da la fórmula (en $\mathbb{R}^2$)

$$d((x_1, y_1), (x_2, y_2)) = \sqrt{(x_1 - x_2)^2 + (y_1 - y_2)^2}.$$



Las propiedades básicas de la distancia se reúnen en el siguiente:

Teorema 1.27 Para todos los vectores $\mathbf{x}, \mathbf{y}, \mathbf{z} \in \mathbb{R}^n$ se cumple que

- i) $d(\mathbf{x}, \mathbf{y}) \geq 0$
- ii) $d(\mathbf{x}, \mathbf{y}) = 0 \Leftrightarrow \mathbf{x} = \mathbf{y}$
- iii) $d(\mathbf{x}, \mathbf{y}) = d(\mathbf{y}, \mathbf{x})$
- iv) $d(\mathbf{x}, \mathbf{y}) \leq d(\mathbf{x}, \mathbf{z}) + d(\mathbf{z}, \mathbf{y})$.

El *espacio métrico* es un conjunto en el que está definida una distancia que cumple con las cuatro propiedades del Teorema 1.27.

EJERCICIO 1.90 Da explícitamente con coordenadas la fórmula de la distancia en $\mathbb{R}^3$ y en $\mathbb{R}^n$.

EJERCICIO 1.91 Demuestra el Teorema 1.27.

EJERCICIO 1.92 ¿Podrías mencionar algún espacio métrico distinto del que se definió arriba?

1.11.1 El espacio euclidiano (primera misión cumplida)

Cuando $\mathbb{R}^n$ se considera junto con la distancia antes definida (llamada la *distancia euclidiana*), se dice que es el espacio euclidiano de dimensión n —que a veces se denota $\mathbb{E}^n$, y que formalmente podríamos definir $\mathbb{E}^n := (\mathbb{R}^n, d)$ — ya que éste cumple con los postulados de Euclides para $n = 2$. Veamos.

Hemos demostrado que $\mathbb{R}^2$ cumple los tres postulados que se refieren a las rectas, y aún nos falta analizar dos. El II se refiere a “trazar” círculos. Ya hemos trabajado con el círculo unitario y claramente se generaliza. Dado un punto $\mathbf{p}$ y una distancia r (un número positivo) definimos el *círculo con centro $\mathbf{p}$ y radio r* como el conjunto de puntos a distancia r de $\mathbf{p}$ (que se estudiará en el siguiente capítulo). Se trata de cierto subconjunto de $\mathbb{R}^2$, que además es no vacío (es fácil dar explícitamente un punto en él usando coordenadas); y nuestra noción de “definir” un conjunto es la análoga de “trazar” para los griegos. Se cumple entonces el axioma II.

Nos falta únicamente discutir el IV, “todos los ángulos rectos son iguales”. Esta afirmación es sutil. Tenemos, desde hace un buen rato, la noción de ángulo recto, perpendicularidad u ortogonalidad, pero ¿a qué se refiere la palabra “igualdad” en el contexto griego? Si fuera a la igualdad del numerito que mide los ángulos, la afirmación es obvia, casi vacua —“ $\pi/2$ es igual a $\pi/2$ ”—, y no habría necesidad de enunciarla como axioma. Entonces se refiere a algo mucho más profundo. En este postulado está implícita la noción de *movimiento*; se entiende la “igualdad” como que podemos *mover* el plano para llevar un ángulo recto formado por dos rectas a cualquier otro. Esto es claro al ver otros teoremas clásicos como “dos triángulos son *iguales* si sus lados miden lo mismo”: no se puede referir a la igualdad estricta de conjuntos, quiere decir que se puede *llevar* a uno sobre el otro para que, entonces sí, coincidan como conjuntos. Y este *llevar* es el *movimiento* implícito en el uso de la palabra “igualdad”. En términos modernos, esta noción de *movimiento* del plano se formaliza con la noción de función. Nosotros lo haremos en el capítulo 3, y corresponderá formalmente a la noción de *isometría* (una función de $\mathbb{R}^2$ en $\mathbb{R}^2$ que preserva distancias). Pero ya no queremos distraernos con la discusión de la axiomática griega. Creamos de buena fe que al desarrollar formalmente los conceptos necesarios el Postulado IV será cierto, o tomémoslo en su acepción numérica trivial, para concluir esta discusión con la que arrancamos el libro.

A partir de los axiomas de los números reales construimos el Plano Euclidiano,

$\mathbb{E}^2$, con todo y sus nociones básicas que cumplen los postulados de Euclides. Se tiene entonces que para estudiarlo son igualmente válidos el método *sintético* (el que usaban los griegos pues sus axiomas ya son ciertos) o el *analítico* (que inauguró Descartes y estamos siguiendo). En este último método siempre hay elementos del primero, no siempre es más directo irse a las coordenadas; no hay un divorcio y en general es muy difícil trazar la línea que los separa; además no vale la pena. No hay que preocuparse de eso, en cada momento toma uno lo que le conviene. Pero sí hay que hacer énfasis en la enorme ventaja que dio el método cartesiano, al construir de golpe (y sólo con un poco de esfuerzo extra) espacios euclidianos para todas las dimensiones. Es un método que permitió generalizar y abrir, por tanto, nuevos horizontes. Y no sólo en cuestión de dimensiones.

Como se verá en capítulos posteriores, siguiendo el método analítico pueden construirse espacios de dimensión 2 con formas “raras” de medir ángulos y distancias que los hacen cumplir todos los axiomas menos el Quinto, y haciéndolos entonces espacios *no euclidianos* . Pero esto vendrá a su tiempo; por el momento y para cerrar el círculo de esta discusión, reescribiremos con la noción de distancia el Teorema 1.1, cuya demostración se deja como ejercicio, y donde, nótese, tenemos el extra del “si y sólo si”.

Se puede definir ángulo en tercia ordenadas de puntos. De nuevo, refiriéndonos al ángulo que ya definimos entre vectores (haciendo que el punto de enmedio sea como el origen):

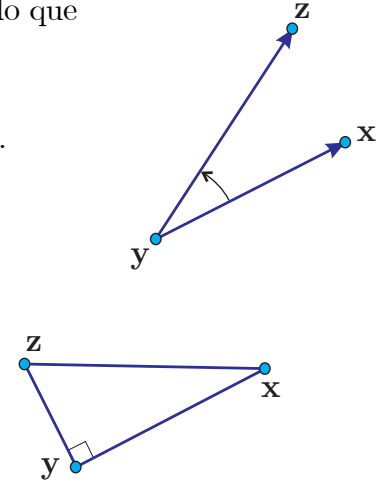
$$\angle xyz := \text{ang}((\mathbf{x} - \mathbf{y}), (\mathbf{z} - \mathbf{y})) = \arccos \frac{(\mathbf{x} - \mathbf{y}) \cdot (\mathbf{z} - \mathbf{y})}{|\mathbf{x} - \mathbf{y}| |\mathbf{z} - \mathbf{y}|}.$$

Y entonces tenemos:

Teorema 1.28 (Pitágoras) *Dados tres puntos $\mathbf{x}$, $\mathbf{y}$, $\mathbf{z}$ en $\mathbb{E}^n$, se tiene que*

$$d(\mathbf{x}, \mathbf{z})^2 = d(\mathbf{x}, \mathbf{y})^2 + d(\mathbf{y}, \mathbf{z})^2 \quad \Leftrightarrow \quad \angle xyz = \pi/2.$$

□



Para retomar el hilo de la corriente principal del texto, en los siguientes ejercicios vemos cómo con la noción de distancia se pueden reconstruir objetos básicos como segmentos y rayos, que juntos dan las líneas; y delineamos un ejemplo de un espacio métrico que da un plano con una geometría “rara”.

EJERCICIO 1.93 Demuestra esta versión del Teorema de Pitágoras.

EJERCICIO 1.94 Demuestra que el segmento de $\mathbf{x}$ a $\mathbf{y}$ es el conjunto

$$\overline{xy} = \{\mathbf{z} \mid d(\mathbf{x}, \mathbf{y}) = d(\mathbf{x}, \mathbf{z}) + d(\mathbf{z}, \mathbf{y})\}.$$

EJERCICIO 1.95 Demuestra que el rayo que parte de $\mathbf{y}$ desde $\mathbf{x}$ (es decir la continuación de la recta por $\mathbf{x}$ y $\mathbf{y}$ más allá de $\mathbf{y}$) es el conjunto $\{\mathbf{z} \mid d(\mathbf{x}, \mathbf{y}) + d(\mathbf{y}, \mathbf{z}) = d(\mathbf{x}, \mathbf{z})\}$.

***EJERCICIO 1.96** (Un ejemplo cotidiano de otro espacio métrico). El plano con la métrica de “Manhattan” (en honor de la famosísima y bien cuadrículada ciudad) se define como $\mathbb{R}^2$ con la función de distancia

$$d_m((x_1, x_2), (y_1, y_2)) := |y_1 - x_1| + |y_2 - x_2|;$$

pensando que para ir de un punto a otro sólo se puede viajar en dirección horizontal o vertical (como en las calles de una ciudad).

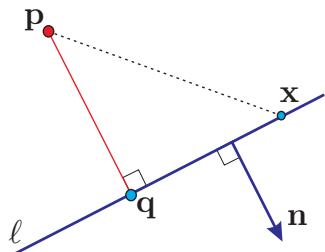
- Demuestra que la métrica de Manhattan cumple el Teorema 1.27.
- Dibuja el conjunto de puntos con distancia 1 al origen; ¿cómo son los círculos?
- ¿Cómo son los segmentos (definidos por la distancia como en el ejercicio anterior)?
- ¿Cómo son los rayos?

1.11.2 Distancia de un punto a una recta

Consideremos el siguiente problema. Nos dan una recta

$$\ell: \mathbf{n} \cdot \mathbf{x} = c$$

y un punto $\mathbf{p}$ (que puede estar fuera o dentro de ella); y se nos pregunta cuál es la distancia de $\mathbf{p}$ a ℓ , que podemos denotar con $d(\mathbf{p}, \ell)$.



Es más fácil resolverlo si lo pensamos junto con un problema aparentemente más complicado: ¿cuál es el punto de ℓ más cercano a $\mathbf{p}$? Llamémoslo $\mathbf{q} \in \ell$, aunque sea incógnito. El meollo del asunto es que la recta que pasa por $\mathbf{p}$ y $\mathbf{q}$ debe ser ortogonal a ℓ . Es decir, si el segmento $\overline{\mathbf{p}\mathbf{q}}$ es ortogonal a ℓ entonces $\mathbf{q}$ es el punto en ℓ más cercano a $\mathbf{p}$. Para demostrarlo considérese cualquier otro punto $\mathbf{x} \in \ell$ y el triángulo rectángulo $\mathbf{x}, \mathbf{q}, \mathbf{p}$ y aplíquese el Teorema de Pitágoras para ver que $d(\mathbf{x}, \mathbf{p}) \geq d(\mathbf{q}, \mathbf{p})$. Como la dirección ortogonal a ℓ es precisamente $\mathbf{n}$, entonces debe existir $t \in \mathbb{R}$ tal que $\mathbf{q} - \mathbf{p} = t\mathbf{n}$. La nueva incógnita t nos servirá para medir la distancia de $\mathbf{p}$ a ℓ . Tenemos entonces que nuestras incógnitas $\mathbf{q}$ y t cumplen las siguientes ecuaciones:

$$\begin{aligned} \mathbf{n} \cdot \mathbf{q} &= c \\ \mathbf{q} - \mathbf{p} &= t\mathbf{n}, \end{aligned}$$

donde la primera dice que $\mathbf{q} \in \ell$ y la segunda que el segmento $\overline{\mathbf{p}\mathbf{q}}$ es ortogonal a ℓ .

Tomando el producto interior de $\mathbf{n}$ con la segunda ecuación, obtenemos

$$\mathbf{n} \cdot \mathbf{q} - \mathbf{n} \cdot \mathbf{p} = t(\mathbf{n} \cdot \mathbf{n}).$$

De donde, al sustituir la primera ecuación, se puede despejar $\mathbf{t}$ (pues $\mathbf{n} \neq \mathbf{0}$) en puros términos conocidos:

$$\begin{aligned} \mathbf{c} - (\mathbf{n} \cdot \mathbf{p}) &= \mathbf{t} (\mathbf{n} \cdot \mathbf{n}) \\ \mathbf{t} &= \frac{\mathbf{c} - (\mathbf{n} \cdot \mathbf{p})}{\mathbf{n} \cdot \mathbf{n}}. \end{aligned}$$

De aquí se deduce que

$$\begin{aligned} \mathbf{q} &= \mathbf{p} + \left(\frac{\mathbf{c} - (\mathbf{n} \cdot \mathbf{p})}{\mathbf{n} \cdot \mathbf{n}} \right) \mathbf{n} \\ d(\mathbf{p}, \ell) &= |\mathbf{t} \mathbf{n}| = |\mathbf{t}| |\mathbf{n}| \\ &= \frac{|\mathbf{c} - (\mathbf{n} \cdot \mathbf{p})|}{|\mathbf{n} \cdot \mathbf{n}|} |\mathbf{n}| = \frac{|\mathbf{c} - (\mathbf{n} \cdot \mathbf{p})|}{|\mathbf{n}|}. \end{aligned}$$

Puesto que sólo hemos usado las propiedades del producto interior y la norma, sin ninguna referencia a las coordenadas, podemos generalizar a $\mathbb{R}^3$; o bien a $\mathbb{R}^n$ pensando que un *hiperplano* (definido por una ecuación lineal $\mathbf{n} \cdot \mathbf{x} = \mathbf{c}$) es una recta cuando $n = 2$ y es un plano cuando $n = 3$.

Proposición 1.29 *Sea Π un hiperplano en $\mathbb{R}^n$ dado por la ecuación $\mathbf{n} \cdot \mathbf{x} = \mathbf{c}$, y sea $\mathbf{p} \in \mathbb{R}^n$ cualquier punto. Entonces*

$$d(\mathbf{p}, \Pi) = \frac{|\mathbf{c} - (\mathbf{n} \cdot \mathbf{p})|}{|\mathbf{n}|},$$

y la proyección ortogonal de $\mathbf{p}$ sobre Π (es decir, el punto más cercano a $\mathbf{p}$ en Π) es

$$\mathbf{q} = \mathbf{p} + \left(\frac{\mathbf{c} - \mathbf{n} \cdot \mathbf{p}}{\mathbf{n} \cdot \mathbf{n}} \right) \mathbf{n}.$$

EJERCICIO 1.97 Evalúa las dos fórmulas de la Proposición 1.29 cuando $\mathbf{p} \in \Pi$.

EJERCICIO 1.98 Encuentra una expresión para $d(\mathbf{p}, \ell)$ cuando $\ell = \{\mathbf{q} + \mathbf{t}\mathbf{v} \mid \mathbf{t} \in \mathbb{R}\}$.

EJERCICIO 1.99 Demuestra que la fórmula de la distancia de un punto a un plano obtenida en el Ejercicio 1.89 coincide con la de la Proposición 1.29.

EJERCICIO 1.100 Encuentra la distancia del punto ... a la recta ... y su proyección ortogonal.

1.11.3 El determinante como área dirigida

No estamos en posición de desarrollar una teoría general de la noción de área; para ello se requieren ideas de cálculo. Pero sí podemos trabajarla para figuras simples

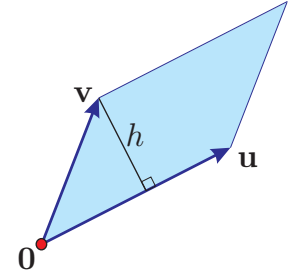
como triángulos y paralelogramos. Se define el *área* de un paralelogramo de manera euclidiana (como en secundaria), es decir, $(\text{base}) \times (\text{altura})$. Veremos que se calcula con un viejo conocido.

Nos encontramos primero el determinante asociado a sistemas de dos ecuaciones lineales con dos incógnitas. Después, definimos el determinante de dos vectores en $\mathbb{R}^2$, otra vez por sistemas de ecuaciones y vimos que detecta (cuando es cero) el paralelismo. Ahora podemos darle una interpretación geométrica en todos los casos.

Teorema 1.30 Sean $\mathbf{u}$ y $\mathbf{v}$ dos vectores cualesquiera en $\mathbb{R}^2$. El área del paralelogramo que definen $\mathbf{u}$ y $\mathbf{v}$ es $|\det(\mathbf{u}, \mathbf{v})|$. Además $\det(\mathbf{u}, \mathbf{v}) \geq 0$ cuando el movimiento angular más corto de $\mathbf{u}$ a $\mathbf{v}$ es positivo ($\overrightarrow{\text{ang}}(\mathbf{u}, \mathbf{v}) \in [0, \pi]$) y $\det(\mathbf{u}, \mathbf{v}) \leq 0$ cuando $\overrightarrow{\text{ang}}(\mathbf{u}, \mathbf{v}) \in [-\pi, 0]$.

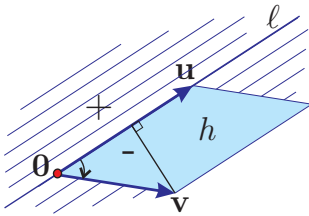
Demostración. Podemos tomar como base del paralelogramo en cuestión al vector $\mathbf{u}$, es decir, (base) en la fórmula $(\text{base}) \times (\text{altura})$ es $|\mathbf{u}|$, y suponemos que $\mathbf{u} \neq \mathbf{0}$. La altura, h digamos, es entonces la distancia de $\mathbf{v}$, pensado como punto, a la recta ℓ generada por $\mathbf{u}$. Esta recta tiene ecuación normal $\mathbf{u}^\perp \cdot \mathbf{x} = 0$. Entonces, por la Proposición 1.29 se tiene que

$$h = \frac{|0 - (\mathbf{u}^\perp \cdot \mathbf{v})|}{|\mathbf{u}^\perp|} = \frac{|\mathbf{u}^\perp \cdot \mathbf{v}|}{|\mathbf{u}|}.$$



Al multiplicar por la base se obtiene que el área del paralelogramo es el valor absoluto del determinante

$$\frac{|\mathbf{u}^\perp \cdot \mathbf{v}|}{|\mathbf{u}|} |\mathbf{u}| = |\mathbf{u}^\perp \cdot \mathbf{v}| = |\det(\mathbf{u}, \mathbf{v})|.$$

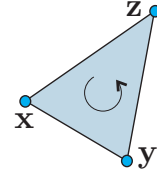


el lado *negativo*.

□

De esta manera, el determinante es un “área signada”: no sólo nos da el área sino que también nos dice si los vectores están orientados positiva o negativamente. Obsérvese que nuestra manera de detectar el paralelismo corresponde entonces a que el área es cero (que el paralelogramo en cuestión es unidimensional).

EJERCICIO 1.101 El área signada de un triángulo *orientado* $\mathbf{x}, \mathbf{y}, \mathbf{z}$ se debe definir como $\det((\mathbf{y} - \mathbf{x}), (\mathbf{z} - \mathbf{x}))$. Demuestra que sólo depende del orden cíclico en que aparecen los vértices, es decir, que da lo mismo si tomamos como ternas a $\mathbf{y}, \mathbf{z}, \mathbf{x}$ o a $\mathbf{z}, \mathbf{x}, \mathbf{y}$. Pero que es su inverso aditivo si tomamos la otra orientación $\mathbf{z}, \mathbf{y}, \mathbf{x}$ para el triángulo.

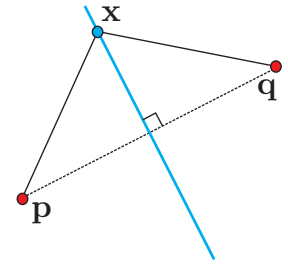


1.11.4 La mediatriz

Cuando empecemos a estudiar las curvas cónicas, las definiremos como *lugares geométricos* de puntos que cumplen cierta propiedad que involucra distancias. Así que para terminar este capítulo veremos los ejemplos más sencillos.

Dados dos puntos $\mathbf{p}$ y $\mathbf{q}$ ($\in \mathbb{R}^2$), ¿cuál es el lugar geométrico de los puntos que equidistan de ellos? Es decir, hay que describir al conjunto

$$\{\mathbf{x} \in \mathbb{R}^2 \mid d(\mathbf{p}, \mathbf{x}) = d(\mathbf{q}, \mathbf{x})\}.$$



Al desarrollar la ecuación $d(\mathbf{p}, \mathbf{x})^2 = d(\mathbf{q}, \mathbf{x})^2$, que equivale a la anterior pues las distancias son positivas, tenemos

$$\begin{aligned} (\mathbf{p} - \mathbf{x}) \cdot (\mathbf{p} - \mathbf{x}) &= (\mathbf{q} - \mathbf{x}) \cdot (\mathbf{q} - \mathbf{x}) \\ \mathbf{p} \cdot \mathbf{p} - 2(\mathbf{p} \cdot \mathbf{x}) + \mathbf{x} \cdot \mathbf{x} &= \mathbf{q} \cdot \mathbf{q} - 2(\mathbf{q} \cdot \mathbf{x}) + \mathbf{x} \cdot \mathbf{x} \\ \mathbf{p} \cdot \mathbf{p} - \mathbf{q} \cdot \mathbf{q} &= 2(\mathbf{p} - \mathbf{q}) \cdot \mathbf{x} \end{aligned}$$

que se puede reescribir como

$$(\mathbf{p} - \mathbf{q}) \cdot \mathbf{x} = (\mathbf{p} - \mathbf{q}) \cdot \left(\frac{1}{2}(\mathbf{p} + \mathbf{q})\right).$$

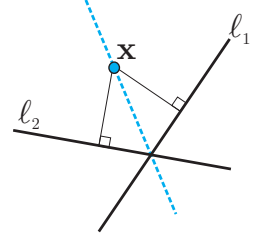
Y ésta es la ecuación normal de la recta ortogonal al segmento $\overline{\mathbf{p}\mathbf{q}}$ y que lo intersecta en su punto medio, llamada la *mediatriz* de $\mathbf{p}$ y $\mathbf{q}$.

EJERCICIO 1.102 ¿Cuál es el lugar geométrico de los puntos en $\mathbb{R}^3$ que equidistan de dos puntos? Describe la demostración.

1.11.5 Bisectrices y ecuaciones unitarias

Preguntamos ahora cuál es el lugar geométrico de los puntos que equidistan de dos rectas ℓ_1 y ℓ_2 . De todas las ecuaciones normales que definen a estas rectas, conviene escoger una donde el vector normal es unitario para que las distancias sean más fáciles de expresar. Tal ecuación se llama *unitaria* y se obtiene de cualquier ecuación normal dividiendo —ambos lados de la ecuación, por supuesto— entre la norma del vector normal. Sean entonces $\mathbf{u}_1$, $\mathbf{u}_2 \in \mathbb{S}^1$ y $c_1, c_2 \in \mathbb{R}$ tales que para $i = 1, 2$

$$\ell_i : \mathbf{u}_i \cdot \mathbf{x} = c_i.$$



Por la fórmula de la sección 1.11.2 se tiene entonces que los puntos $\mathbf{x}$ que equidistan de ℓ_1 y ℓ_2 son precisamente los que cumplen con la ecuación

$$|\mathbf{u}_1 \cdot \mathbf{x} - c_1| = |\mathbf{u}_2 \cdot \mathbf{x} - c_2|.$$

Ahora bien, si el valor absoluto de dos números reales coincide entonces o son iguales o son inversos aditivos, lo cual se expresa elegantemente de la siguiente manera:

$$|a| = |b| \Leftrightarrow (a + b)(a - b) = 0,$$

que, aplicado a nuestra ecuación anterior, nos dice que el lugar geométrico que estamos buscando está determinado por la ecuación

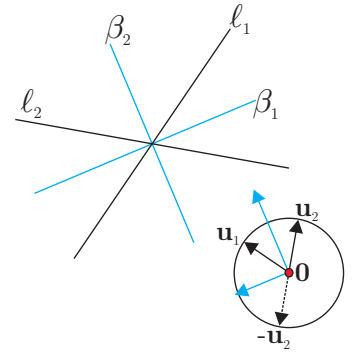
$$\begin{aligned} (\mathbf{u}_1 \cdot \mathbf{x} - c_1 + \mathbf{u}_2 \cdot \mathbf{x} - c_2)(\mathbf{u}_1 \cdot \mathbf{x} - c_1 - \mathbf{u}_2 \cdot \mathbf{x} + c_2) &= 0 \\ ((\mathbf{u}_1 + \mathbf{u}_2) \cdot \mathbf{x} - (c_1 + c_2))((\mathbf{u}_1 - \mathbf{u}_2) \cdot \mathbf{x} - (c_1 - c_2)) &= 0. \end{aligned}$$

Como cuando el producto de dos números es cero alguno de ellos lo es, entonces nuestro lugar geométrico es la unión de las dos rectas:

$$\begin{aligned} \beta_1 : (\mathbf{u}_1 + \mathbf{u}_2) \cdot \mathbf{x} &= (c_1 + c_2) \\ \beta_2 : (\mathbf{u}_1 - \mathbf{u}_2) \cdot \mathbf{x} &= (c_1 - c_2), \end{aligned} \quad (1.12)$$

que son las *bisectrices* de las líneas ℓ_1 y ℓ_2 . Observemos que son ortogonales pues

$$\begin{aligned} (\mathbf{u}_1 + \mathbf{u}_2) \cdot (\mathbf{u}_1 - \mathbf{u}_2) &= |\mathbf{u}_1|^2 - |\mathbf{u}_2|^2 \\ &= 1 - 1 \\ &= 0, \end{aligned}$$

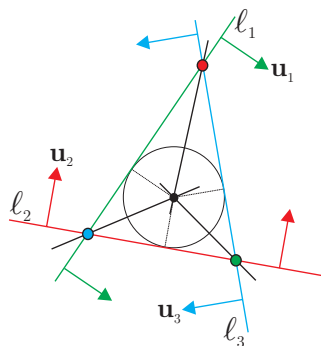


y, geoméricamente, bisectan a los dos sectores o ángulos que forman las rectas; también pasan por el punto de intersección pues sus ecuaciones se obtienen sumando y restando las ecuaciones originales (el caso en que sean paralelas se deja como ejercicio). En el círculo unitario también es claro que la suma y la diferencia de dos vectores (aunque ya no sean necesariamente unitarios) tienen los ángulos adecuados.

Como corolario, obtenemos otro de los teoremas clásicos de concurrencia.

Teorema 1.31 *Las bisectrices (internas) de un triángulo son concurrentes.*

Demostración. Sólo hay que tener cuidado con cuáles ecuaciones unitarias trabajamos, pues hay dos posibles para cada recta. Si escogemos los vectores unitarios de las tres rectas, ℓ_1 , ℓ_2 y ℓ_3 digamos, apuntando hacia adentro del triángulo, $\mathbf{u}_1$, $\mathbf{u}_2$ y $\mathbf{u}_3$ respectivamente, entonces está claro que sus sumas son normales a las bisectrices exteriores. El teorema habla entonces de las bisectrices que se obtienen como diferencias, que son, numerándolas por el índice de la recta opuesta:



$$\begin{aligned}\beta_1 &: (\mathbf{u}_2 - \mathbf{u}_3) \cdot \mathbf{x} = (c_2 - c_3) \\ \beta_2 &: (\mathbf{u}_3 - \mathbf{u}_1) \cdot \mathbf{x} = (c_3 - c_1) \\ \beta_3 &: (\mathbf{u}_1 - \mathbf{u}_2) \cdot \mathbf{x} = (c_1 - c_2).\end{aligned}$$

El negativo de cualquiera de ellas se obtiene sumando las otras dos.

□

EJERCICIO 1.103 Un argumento aparentemente más elemental para el teorema anterior es que si un punto está en la intersección de dos bisectrices, entonces ya equidista de las tres rectas y por tanto también está en la bisectriz del vértice restante. ¿Qué implica este argumento cuando tomamos bisectrices exteriores? ¿En dónde concurren dos bisectrices exteriores? ¿Una interior y una exterior? Da enunciados y demuéstralos con ecuaciones normales. Completa el dibujo.

EJERCICIO 1.104 Demuestra que el lugar geométrico de los puntos que equidistan a dos rectas paralelas es la paralela a ambas que pasa por el punto medio de sus intersecciones con una recta no paralela a ellas, a partir de las ecuaciones (1.12).

EJERCICIO 1.105 ¿Cuál es el lugar geométrico de los puntos que equidistan de dos planos en $\mathbb{R}^3$? Describe una demostración de tu aseveración.

1.12 *Los espacios de rectas en el plano

Esta sección “asterisco” o extra es, hasta cierto punto, independiente del texto principal; será relevante como motivación hasta el capítulo 6, donde se estudia el plano proyectivo y cristaliza en la llamada “dualidad”. Estudiaremos dos ejemplos que ya están a la mano de cómo la noción de *el Espacio*, que por milenios se consideró única y amorfa —“el espacio en que vivimos”—, se pluraliza *los espacios*. Ya vimos que el Espacio teórico que consideraron los griegos es sólo la instancia en dimensión 3

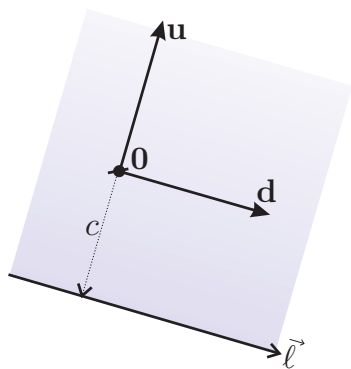
de la familia de espacios euclidianos $\mathbb{E}^n$; pero además la *pluralización* (pasar de uno a muchos) se da en otro sentido. En matemáticas, y en particular en la geometría, surgen naturalmente otros conjuntos, además de los $\mathbb{E}^n$, con plenos derechos de ser llamados “espacios”; en la física se les llama “espacios fase”. Los dos “espacios”⁵ que nos ocupan ahora son los de rectas en el plano. Consideraremos cada recta como un punto de un nuevo “espacio” y, puesto que intuitivamente podemos decir si dos rectas son cercanas o no, en este nuevo “espacio”, sabremos si dos puntos lo están. La idea clave está en que así como Descartes asoció parejas de números a los puntos del plano, a las rectas del plano también hemos asociado números (los involucrados en sus ecuaciones) de tal manera que si esos números varían un poco la recta asociada también se mueve poco; dan lugar a una noción de “continuidad” que es la noción básica de la topología.

En el Ejercicio 1.68 de la sección 1.8, se hizo notar que a cada terna de números $(a, b, c) \in \mathbb{R}^3$, con $(a, b) \neq (0, 0)$, se le puede asociar una recta en $\mathbb{R}^2$, a saber, la definida por la ecuación $ax + by = c$; el estudiante debió haber demostrado que las rectas por el origen en $\mathbb{R}^3$ dan la misma recta en $\mathbb{R}^2$ y que los planos por el origen en $\mathbb{R}^3$ corresponden a los haces de rectas en el plano. Dando esto por entendido, en esta última sección abundaremos en tales ideas. Se puede decir que la pregunta que nos guía es: ¿cuántas rectas hay en $\mathbb{R}^2$?, pero el término “cuántas” es demasiado impreciso. Veremos que las rectas de $\mathbb{R}^2$ forman algo que llamaremos el “espacio” de rectas (un caso especial de las llamadas “Grassmannianas” pues quien primero desarrolló estas ideas fue Grassmann), y para esto les daremos “coordenadas” explícitas, muy a semejanza de las coordenadas polares para los puntos.

1.12.1 Rectas orientadas

Consideremos primero las rectas orientadas en $\mathbb{R}^2$, pues a ellas se les puede asociar un vector normal (y por tanto una ecuación) de manera canónica. Una *recta orientada*

es una recta $\vec{\ell}$ junto con una orientación preferida, que nos dice la manera *positiva* de viajar en ella. Entonces, de todos sus vectores direccionales podemos escoger el unitario $\mathbf{d}$ con la orientación positiva y, como vector normal, escogemos el compadre ortogonal de éste; es decir, el vector unitario tal que después de la orientación de la recta nos da la orientación positiva del plano. Así, cada recta orientada está definida por una ecuación *unitaria* única

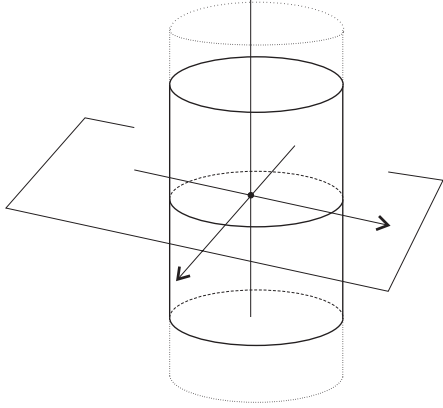


$$\mathbf{u} \cdot \mathbf{x} = c,$$

⁵Mantenemos las comillas pues la noción aún no está bien definida: va en la dirección de lo que ahora se llama topología.

donde $\mathbf{u} \in \mathbb{S}^1$, y $c \in \mathbb{R}$ es su distancia orientada al origen.

Nótese que a cada recta orientada se le puede asociar su semiplano positivo $\mathbf{u} \cdot \mathbf{x} \geq c$; así que también podemos pensar a las rectas orientadas como semiplanos.



Puesto que esta ecuación unitaria es única (por la orientación) entonces hay tantas rectas orientadas como parejas

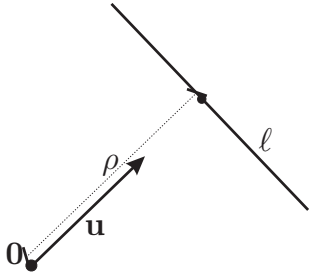
$$(\mathbf{u}, c) \in \mathbb{S}^1 \times \mathbb{R} \subset \mathbb{R}^2 \times \mathbb{R} = \mathbb{R}^3,$$

y este conjunto es claramente un cilindro vértical en $\mathbb{R}^3$. Los haces de líneas orientadas paralelas corresponden a líneas verticales en el cilindro, pues en estos haces el vector normal está fijo. Y los haces de líneas (orientadas) concurrentes corresponden a lo que corta un plano por el origen al cilindro $\mathbb{S}^1 \times \mathbb{R}$ (que, como veremos en el capítulo siguiente, se llama elipse).

EJERCICIO 1.106 Demuestra la última afirmación.

1.12.2 Rectas no orientadas

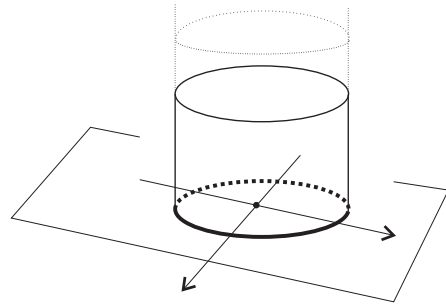
Las dos orientaciones de una recta ℓ tienen ecuaciones unitarias $\mathbf{u} \cdot \mathbf{x} = c$ y $(-\mathbf{u}) \cdot \mathbf{x} = -c$. Si la recta no pasa por el origen (si $\mathbf{0} \notin \ell$ y por tanto $c \neq 0$), de estas dos ecuaciones unitarias podemos tomar la ecuación con constante estrictamente positiva, $\mathbf{u} \cdot \mathbf{x} = \rho$, con $\rho > 0$, digamos. Así que las “coordenadas” $(\mathbf{u}, \rho)$ corresponden a las coordenadas polares del punto en la recta ℓ más cercano al origen; y viceversa, a cada punto $\mathbf{x} \in \mathbb{R}^2 - \{\mathbf{0}\}$ se le asocia la recta que pasa por el punto $\mathbf{x}$ y es normal al vector $\mathbf{x}$. O dicho de otra manera, si la recta ℓ no contiene al origen, podemos escoger el semiplano que define que no contiene a $\mathbf{0}$.



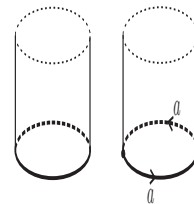
Nos queda una ambigüedad en el origen: para las rectas que pasan por él tenemos dos vectores unitarios normales que definen la misma recta. Podemos resumir que las rectas en $\mathbb{R}^2$ corresponden a los puntos de un cilindro “semicerrado”

$$\mathbb{S}^1 \times \mathbb{R}_+,$$

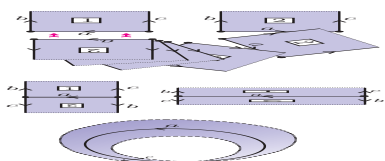
(donde $\mathbb{R}_+ = \{\rho \in \mathbb{R} \mid \rho \geq 0\}$), donde hay que identificar su frontera por antípodas, es decir, en este conjunto $(\mathbf{u}, 0)$ y $(-\mathbf{u}, 0)$ aún representan la misma recta. Si identificamos $(\mathbf{u}, 0)$ con $(-\mathbf{u}, 0)$ en el cilindro $\mathbb{S}^1 \times \mathbb{R}_+$, se obtiene una banda de Moebius abierta. Veámos:



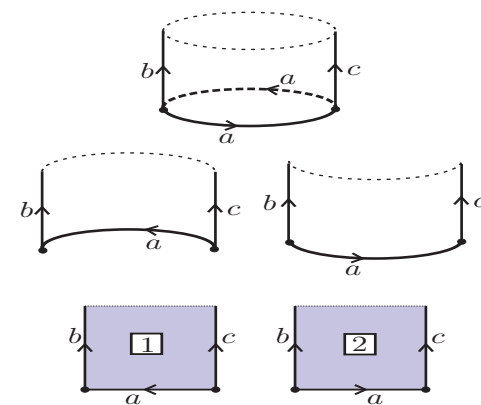
Hay que pensar que el cilindro $\mathbb{S}^1 \times \mathbb{R}_+$ está hecho de un material flexible (pensarlo como “espacio topológico”, se dice actualmente), y hacer muy pequeño el factor $\mathbb{R}_+$, para tener un cilindrito con una de sus fronteras inexistente (la que corresponde al infinito), indicada en las figuras con línea punteada. En la otra frontera (que corresponde al 0) debemos pegar puntos antípodos. Llamemos $\mathbf{a}$ al segmento dirigido que es la mitad de ese círculo y entonces la otra mitad, con la misma orientación del círculo, vuelve a ser $\mathbf{a}$.



Si tratamos de pegar directamente los segmentos $\mathbf{a}$, la superficie del cilindro nos estorba y no se ve qué pasa con claridad (de hecho no se puede lograr en $\mathbb{R}^3$ sin que la superficie se autointersecte). Conviene entonces cortar el cilindro por los segmentos verticales, denotados con $\mathbf{b}$ y $\mathbf{c}$ en la figura, sobre el principio y fin de $\mathbf{a}$, para obtener dos pequeños cuadra-



dos (que hemos bautizado 1 y 2). Los segmentos $\mathbf{b}$ y $\mathbf{c}$ están dirigidos y debemos recordar que hay que volverlos a pegar cuando



se pueda.

Ahora sí podemos pegar los segmentos $\mathbf{a}$ recordando su dirección. Basta voltear el cuadrado 2 de cabeza para lograrlo, y obtener un rectángulo. Si alargamos el rectángulo en su dirección horizontal, obtenemos una banda en la cual hay que identificar las fronteras verticales hechas, cada una, con los pares de segmentos $\mathbf{b}$ y $\mathbf{c}$ (la frontera horizontal sigue siendo punteada). Y esta identificación, según indi-

cación, según indican nuestros segmentos dirigidos, es invirtiendo la orientación; hay que dar media vuelta

para pegarlos, así que el resultado es una banda de Moebius a la que se le quita su frontera, y ya

no queda nada por identificar.

Hemos demostrado que las rectas del plano corresponden biyectivamente a los puntos de una banda de Moebius abierta, y que esta correspondencia preserva “cercanía” o “continuidad”.

Obsérvese además que las rectas que pasan por el origen (el segmento α con sus extremos identificados) quedan en el “corazón” (el círculo central) de la banda de Moebius; los haces paralelos corresponden a los intervalos abiertos transversales al corazón que van de la frontera a sí misma (el punto de intersección con el corazón es su representante en el origen). ¿A qué corresponden los haces concurrentes?

EJERCICIO 1.107 Demuestra que dos haces de rectas, no ambos haces paralelos, tienen una única recta en comun.
