

Una introducción a la geometría lorentziana¹

Oscar Palmas
Universidad Nacional Autónoma de México

Versión preliminar del 23 de julio de 2019

¹Notas para el curso del mismo nombre, impartido en la Escuela Internacional de Geometría Diferencial, julio 2019, Universidad de El Salvador.

Estas notas son apenas un esbozo del material necesario para un curso completo. En las referencias al final de estas notas hemos incluido una lista de libros donde el lector interesado puede consultar mayores detalles. Por supuesto, cualquier sugerencia es bienvenida, por ejemplo, escribiendo un correo a `oscar.palmas@ciencias.unam.mx`.

Agradezco a Daniel Brito por una primera lectura crítica de estas notas, verificando y proponiendo ejercicios, así como por la figura 2.1.

Índice general

1. Álgebra lineal	1
1.1. Hechos básicos	1
1.2. Subespacios y bases	5
1.3. Isometrías	8
1.4. Espacios vectoriales lorentzianos	10
1.5. Isometrías en espacios lorentzianos	13
1.6. Transformaciones lineales autoadjuntas	16
1.7. Ejercicios	18
2. Curvas y superficies en $\mathbb{R}_1^3$	21
2.1. Curvas	21
2.2. Superficies en $\mathbb{R}_1^3$	23
2.3. Superficies tipo espacio	26
2.4. Superficies tipo tiempo	28
2.5. La segunda forma fundamental	32
2.6. Ejercicios	35
3. Variedades lorentzianas	37
3.1. Introducción	37
3.2. Hipersuperficies umbílicas en formas espaciales	41
3.3. Hipersuperficies con curvatura media constante	43
3.4. Epílogo: Hipersuperficies no umbílicas	45
Índice alfabético	47
Bibliografía	49

Capítulo 1

Álgebra lineal

El estudio de la geometría lorentziana requiere algunas herramientas del álgebra lineal, algunas de las cuales no se abordan en un curso ni un libro de consulta usual sobre la materia, como el excelente [3]. Es por ello que en este capítulo daremos un panorama de tales herramientas. Por supuesto, supondremos que el lector tiene un conocimiento sólido de álgebra, como el desarrollado en la obra ya citada.

1.1. Hechos básicos

En todo este capítulo, V denotará un espacio vectorial de dimensión finita sobre los números reales. El lector conoce la definición usual del producto escalar (interior o punto), es decir, una forma bilineal simétrica positiva definida en V . Pero para nuestros fines, modificaremos la última condición de la manera siguiente.

Definición 1.1 (Producto escalar). Un *producto escalar en V* es una forma bilineal $\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{R}$ simétrica que es *no degenerada*, en el sentido de que si existe $u \in V$ tal que $\langle u, v \rangle = 0$ para todo $v \in V$, entonces $u = 0$.

Como de costumbre, un espacio vectorial dotado de un producto escalar se denotará por $(V, \langle \cdot, \cdot \rangle)$ cuando sea necesario hacer referencia a dicho producto, y en caso contrario se denotará simplemente por V .

Ejemplo 1.2. Consideremos dos enteros $m \geq \nu \geq 0$. Entonces $\mathbb{R}_\nu^m$ denotará al espacio vectorial $\mathbb{R}^m$ dotado del siguiente producto escalar:

$$\langle u, v \rangle = - \sum_{i=1}^{\nu} u_i v_i + \sum_{i=\nu+1}^m u_i v_i, \quad (1.1)$$

donde $u = (u_1, \dots, u_\nu, u_{\nu+1}, \dots, u_m)$ y $v = (v_1, \dots, v_\nu, v_{\nu+1}, \dots, v_m)$ con las coordenadas cartesianas usuales en $\mathbb{R}^m$. Si $\nu = 0$, la primera suma no aparece y recuperamos el *espacio euclidiano* usual con producto escalar positivo definido;

en otras palabras, $\mathbb{R}_0^m = \mathbb{R}^m$. En el otro extremo, si $\nu = m$, la segunda suma desaparece y tenemos un espacio con producto escalar negativo definido, que a veces se denota así: $\mathbb{R}_m^m = -\mathbb{R}^m$. En general, diremos que $\mathbb{R}_\nu^m$ es un *espacio semi-euclidiano*. Si $\nu = 1$ y $m \geq 2$, decimos que $\mathbb{R}_1^m$ es el *espacio de Lorentz-Minkowski de dimensión m* ; escribiremos $m = n + 1$ y las coordenadas cartesianas como $u = (u_0, u_1, \dots, u_n)$, de modo que el producto escalar (1.1) queda así:

$$\langle u, v \rangle = -u_0v_0 + \sum_{i=1}^n u_i v_i. \quad (1.2)$$

Observación 1.3. En la definición anterior hemos optado por colocar todos los términos negativos *al principio* de nuestra expresión. En algunos contextos, se acostumbra dejar estos términos *al final*. El lector deberá tener cuidado de verificar esta diferencia no esencial en los diversos contextos donde se apliquen.

A manera de publicidad (o *spoiler*, como se quiera ver), cualquier espacio vectorial de dimensión finita con producto escalar $(V, \langle \cdot, \cdot \rangle)$ será isométrico a uno de estos espacios $\mathbb{R}_\nu^m$ (véase el teorema 1.29). De hecho, a partir del siguiente capítulo trabajaremos fundamentalmente con los casos $\mathbb{R}_1^2$ y $\mathbb{R}_1^3$, donde es mucho más fácil visualizar las cosas. También debemos mencionar que en casos de dimensión baja, a veces se suele usar de manera alternativa las notaciones (t, x) para las coordenadas en $\mathbb{R}_1^2$ y (t, x, y) para las coordenadas en $\mathbb{R}_1^3$.

Puesto que hemos modificado nuestra definición de producto escalar, es de esperar que algunas de las propiedades de la antigua definición se cumplan y otras dejen de hacerlo. Tal vez las propiedades más fundamentales de este producto escalar sean las siguientes.

Proposición 1.4. Sean $u_1, u_2 \in V$ tales que

$$\langle u_1, v \rangle = \langle u_2, v \rangle$$

para todo $v \in V$. Entonces $u_1 = u_2$.

Proposición 1.5. Sean $m = \dim V$, $\{e_1, \dots, e_m\}$ una base de V y $g_{ij} = \langle e_i, e_j \rangle$. El producto escalar $\langle \cdot, \cdot \rangle$ es no degenerado si y sólo si la matriz (g_{ij}) es invertible.

Demostración. Sea $u \in V$ un vector arbitrario, que podemos escribir como

$$u = \sum_{i=1}^m u_i e_i.$$

El producto $\langle \cdot, \cdot \rangle$ es degenerado si y sólo si existe $u \neq 0$ tal que $\langle u, v \rangle = 0$ para todo $v \in V$; aplicando esto a los elementos de la base, tenemos que el producto escalar es degenerado si y sólo si existe $u \neq 0$ tal que

$$0 = \langle u, e_j \rangle = \sum_{i=1}^m u_i \langle e_i, e_j \rangle = \sum_{i=1}^m u_i g_{ij}$$

para todo $j = 1, \dots, m$. En otras palabras, el producto es degenerado si y sólo si el sistema de ecuaciones definido por la matriz (g_{ij}) admite una solución no trivial, lo que ocurre si y sólo si la matriz en cuestión no es invertible. $\square$

Dado que en nuestra definición el producto escalar puede dejar de ser positivo definido, habría vectores $v \in V$ tales que $\langle v, v \rangle$ sea negativo. Por ejemplo, esto ocurre con el vector $(1, 0) \in \mathbb{R}_1^2$. Esto sugiere una clasificación de los vectores, llamada su carácter causal.

Definición 1.6 (Carácter causal de un vector). Un vector $v \in V$ es

- *tipo espacio* o *espacial* si $v = 0$ o $\langle v, v \rangle > 0$;
- *tipo tiempo* o *temporal* si $\langle v, v \rangle < 0$;
- *tipo luz* o *nulo* si $v \neq 0$ y $\langle v, v \rangle = 0$.
- *causal* si $v \neq 0$ y $\langle v, v \rangle \leq 0$.

Aprovechando lo anterior, nos fijaremos en algunos conjuntos que usaremos con cierta frecuencia.

Definición 1.7 (Conos). Sea $(V, \langle \cdot, \cdot \rangle)$ un espacio vectorial real con producto escalar.

- El *cono de luz* o *cono nulo* de V es el conjunto $CN(V)$ de vectores nulos de V .
- El *cono de tiempo* de V es el conjunto $CT(V)$ de vectores tipo tiempo de V .
- Sea $u \in V$ un vector nulo. El *cono de luz* o *cono nulo* de u en V , denotado $CN(u)$, es el componente conexo de $CN(V)$ que contiene a u .
- Sea $u \in V$ un vector tipo tiempo. El *cono de tiempo* o *cono temporal* de u en V , denotado $CT(u)$, es el conjunto de vectores tipo tiempo $v \in V$ tales que $\langle u, v \rangle < 0$.
- Sea $u \in V$ un vector causal. El *cono causal* de u en V , denotado $CC(u)$, es la unión $CN(u) \cup CT(u)$.

El lector puede hacer el dibujo de $\mathbb{R}_1^2$ y de $\mathbb{R}_1^3$ y ubicar con facilidad los conjuntos de vectores de tipo espacio, de tipo tiempo, nulos y causales en cada caso.

Recordemos que un conjunto $A \subset V$ es *convexo* si para cualesquiera dos vectores $u, v \in A$, ocurre que el segmento determinado por ellos, es decir, el conjunto $\{tu + (1 - t)v \mid 0 \leq t \leq 1\}$, está contenido en A .

Proposición 1.8. *Los conos $CN(u)$, $CT(u)$ y $CC(u)$ son conjuntos convexos.*

Sabemos que cualquier producto escalar tiene asociada una norma. En nuestro caso, la definición de este concepto debe modificarse como sigue.

Definición 1.9. La *norma* de un vector $v \in V$ se define como

$$\|v\| = \sqrt{|\langle v, v \rangle|}. \quad (1.3)$$

Es fácil convencerse que algunas propiedades usuales de los espacios con producto escalar positivo definido, como la desigualdad del triángulo o la desigualdad de Cauchy-Schwarz, no se cumplen en general para cualquier pareja de vectores en V . Pero podemos decir más en algunos casos importantes. Si consideramos dos vectores $u, v \in V$ tipo espacio, entonces las pruebas usuales de estas desigualdades siguen valiendo. Pero es de sospechar que haya instancias en que valgan las desigualdades inversas. Éste es el caso del siguiente resultado.

Proposición 1.10 (Desigualdad inversa de Cauchy-Schwarz). *Sean $u, v \in V$ vectores tipo tiempo. Entonces*

$$|\langle u, v \rangle| \geq \|u\| \|v\|, \quad (1.4)$$

y la igualdad vale si y sólo si u y v son linealmente dependientes. Si además u, v están en el mismo cono de tiempo, entonces existe un único número real $\phi \geq 0$ tal que

$$\langle u, v \rangle = -\|u\| \|v\| \cosh \phi. \quad (1.5)$$

Demostración. Podemos adaptar la demostración usual de la desigualdad a este caso. Si u, v son linealmente dependientes, la desigualdad se cumple. Ahora, si u, v son linealmente independientes, consideramos la recta $u + \lambda v$, $\lambda \in \mathbb{R}$. La gráfica del polinomio de segundo grado en λ

$$P(\lambda) = \langle u + \lambda v, u + \lambda v \rangle = \langle u, u \rangle + 2\langle u, v \rangle \lambda + \langle v, v \rangle \lambda^2$$

es una parábola que abre hacia abajo, y cuyo vértice, correspondiente al valor $\lambda = -\langle u, u \rangle / \langle v, v \rangle$, está por arriba del eje λ ; por tanto, P tiene dos raíces. Esto dice que el discriminante asociado a la ecuación $P(\lambda) = 0$ es positivo; es decir,

$$4\langle u, v \rangle^2 - 4\langle u, u \rangle \langle v, v \rangle > 0$$

de donde se obtiene la desigualdad (1.4).

Para la segunda parte, si u, v están en el mismo cono de tiempo, la proposición 1.35 nos dice que $\langle u, v \rangle < 0$, de modo que

$$-\langle u, v \rangle \geq \|u\| \|v\|, \quad \text{o} \quad -\frac{\langle u, v \rangle}{\|u\| \|v\|} \geq 1.$$

Como la restricción del coseno hiperbólico a $[0, \infty)$ es biyectiva sobre $[1, \infty)$, existe un único número $\phi \in [0, \infty)$ que cumple (1.5). $\square$

Definición 1.11. El número ϕ que aparece en (1.5) es el *ángulo hiperbólico* entre u y v .

Corolario 1.12 (Desigualdad inversa del triángulo). *Sean $u, v \in V$ vectores tipo tiempo en el mismo cono de tiempo. Entonces*

$$\|u\| + \|v\| \leq \|u + v\|, \quad (1.6)$$

y la igualdad vale si y sólo si u, v son linealmente dependientes.

Demostración. Observemos que $u + v$ es tipo tiempo y, de hecho, está en el mismo cono de tiempo de u, v . Por las hipótesis, tenemos que $-\langle u, v \rangle \geq \|u\|\|v\|$. Por tanto,

$$\begin{aligned} (\|u\| + \|v\|)^2 &= \|u\|^2 + 2\|u\|\|v\| + \|v\|^2 \\ &\leq \|u\|^2 - 2\langle u, v \rangle + \|v\|^2 \\ &= -\langle u + v, u + v \rangle = \|u + v\|^2. \end{aligned}$$

El caso de la igualdad se sigue directamente de la proposición 1.10. $\square$

Hablemos un poco sobre conceptos conocidos en el caso euclidiano. Decimos que un vector $u \in V$ es *unitario* si $\|u\| = 1$. Por otro lado, dos vectores $u, v \in V$ son *ortogonales* si $\langle u, v \rangle = 0$.

Definición 1.13. Sea $A \subset V$. Denotamos por $A^\perp$ al conjunto de vectores ortogonales a A :

$$A^\perp = \{v \in V \mid \langle u, v \rangle = 0 \text{ para todo } u \in A\}.$$

Finalmente, en el caso particular de $\mathbb{R}_1^3$, podemos definir el producto cruz de dos vectores.

Definición 1.14 (Producto cruz lorentziano). Sean $u, v \in \mathbb{R}_1^3$, con coordenadas $u = (u_0, u_1, u_2)$ y $v = (v_0, v_1, v_2)$. Además, sean $i = (1, 0, 0)$, $j = (0, 1, 0)$, $k = (0, 0, 1)$. Entonces el *producto cruz lorentziano* de u y v en $\mathbb{R}_1^3$, denotado $u \times v$, se define como

$$u \times v = \begin{vmatrix} -i & j & k \\ u_0 & u_1 & u_2 \\ v_0 & v_1 & v_2 \end{vmatrix}.$$

Observación 1.15. De acuerdo con la observación 1.3, será importante notar dónde se coloca el signo negativo en la definición del producto escalar, ya que eso modificará también la definición del producto cruz.

Observación 1.16. El lector notará que hemos utilizado la notación $u \times v$, idéntica a la notación para el producto cruz en el caso euclidiano. Esto no debe llevar a confusión, pues sólo usaremos el caso lorentziano en estas notas. Por la misma razón, utilizaremos la expresión más corta *producto cruz*.

1.2. Subespacios y bases

Consideremos ahora un subespacio vectorial W de $(V, \langle \cdot, \cdot \rangle)$. En esta situación, podemos restringir el producto escalar de V a W . Esto no siempre induce un producto escalar en W ; es decir, no siempre es una forma no degenerada: Si $W = \text{span}\{(1, 1)\}$ en $\mathbb{R}_1^2$, la forma es degenerada.

Definición 1.17 (Subespacio no degenerado). Diremos que un subespacio W de $(V, \langle \cdot, \cdot \rangle)$ es *no degenerado* si la restricción del producto escalar de V a W es de nuevo un producto escalar.

En el caso euclidiano, para cualquier subespacio $W \subset V$ se tiene que V es la suma directa de W y $W^\perp$ (definición 1.13). El lector puede ver que esto no necesariamente ocurre en nuestro contexto; por ejemplo, si $W = \text{span}\{(1, 1)\} \subset \mathbb{R}_1^2$. Sin embargo, $W^\perp$ tiene otras propiedades que recuerdan al caso euclidiano.

Proposición 1.18. *Sea W un subespacio vectorial de $(V, \langle, \rangle)$.*

1. $\dim W + \dim W^\perp = \dim V$.
2. $(W^\perp)^\perp = W$.
3. W es no degenerado si y sólo si $V = W \oplus W^\perp$.
4. W es no degenerado si y sólo si $W^\perp$ es no degenerado.

Demostración. 1. Sean $m = \dim V$, $k = \dim W$ y $\{e_1, \dots, e_m\}$ una base de V adaptada a W ; es decir, de modo que $\{e_1, \dots, e_k\}$ sea una base de W . Entonces, un vector $v \in V$ se puede escribir como

$$v = \sum_{i=1}^m v_i e_i.$$

Observemos que $v \in W^\perp$ si y sólo si $\langle v, e_j \rangle = 0$ para todo $j = 1, \dots, k$; es decir,

$$0 = \langle v, e_j \rangle = \sum_{i=1}^m v_i \langle e_i, e_j \rangle = \sum_{i=1}^m v_i g_{ij};$$

para todo $j = 1, \dots, k$. Esta expresión representa un sistema de k ecuaciones con m incógnitas; pero recordando que la matriz (g_{ij}) es invertible (proposición 1.5), sabemos que el conjunto de soluciones de este sistema tiene dimensión $m - k$. Es decir, $\dim W^\perp = m - k$.

2. Por el primer inciso, tenemos que

$$\dim W + \dim W^\perp = \dim V = \dim W^\perp + \dim (W^\perp)^\perp,$$

de modo que $\dim W = \dim (W^\perp)^\perp$. Pero es claro que $W \subset (W^\perp)^\perp$, de modo que por un argumento sobre la dimensión, ambos conjuntos son iguales entre sí.

3. En general, si W, W' son subespacios de V , sabemos que

$$\dim(W + W') = \dim W + \dim W' - \dim(W \cap W'),$$

de modo que si $W' = W^\perp$, usando esto y el primer inciso, vemos que

$$\dim(W + W^\perp) = \dim V - \dim(W \cap W^\perp).$$

Ahora bien, si W es no degenerado, entonces $W \cap W^\perp = \{0\}$, lo que implica además que $W + W^\perp = V$, de modo que la suma es directa. Recíprocamente, si la suma es directa, entonces $W \cap W^\perp = \{0\}$ y W es no degenerado.

4. Consecuencia de los dos incisos anteriores. $\square$

Ahora estamos preparados para mostrar la existencia de cierto tipo importante de bases.

Teorema 1.19 (Existencia de una base ortonormal). *Dado $(V, \langle \cdot, \cdot \rangle)$, con $m = \dim V$, existe una base $\{e_1, \dots, e_m\}$ para V y un entero $k \in \{0, \dots, m\}$ tales que*

$$\langle e_i, e_j \rangle = \epsilon_i \delta_{ij},$$

para todo $i, j = 1, \dots, m$; aquí, $\epsilon_i = -1$ para $i = 1, \dots, k$, mientras que $\epsilon_i = +1$ para $i = k+1, \dots, m$. Una base con esta propiedad se llama una base ortonormal de V .

Demostración. Procedemos por inducción sobre m . Si $m = 1$, como el producto es no degenerado, existe $u \in V$ tal que $\langle u, u \rangle \neq 0$; entonces $\{u/\|u\|\}$ es una base ortonormal para V . Supongamos ahora que el resultado es cierto para todo espacio vectorial con dimensión menor que $m > 1$ y sea V un espacio de dimensión m . De nuevo, existe $u \in V$ tal que $\langle u, u \rangle \neq 0$. Por la proposición anterior, el subespacio $\{u\}^\perp$ es no degenerado, de modo que admite una base ortonormal $\{e_1, \dots, e_{m-1}\}$, y por tanto $\{e_1, \dots, e_{m-1}, e_m = u/\|u\|\}$ es una base ortonormal para V (salvo posiblemente el orden). $\square$

El siguiente resultado es muy similar a su correspondiente en los espacios con productos escalares positivo definidos.

Proposición 1.20 (Expresión de un vector con respecto de una base ortonormal). *Sea $\{e_1, \dots, e_m\}$ una base ortonormal de V . Si $v \in V$, entonces*

$$v = \sum_{i=1}^m \epsilon_i \langle v, e_i \rangle e_i,$$

donde los signos ϵ_i se ordenan como en el teorema 1.19.

Observación 1.21. De acuerdo con la notación establecida para $\mathbb{R}_1^{n+1}$, denotaremos por $\{e_0, e_1, \dots, e_n\}$ a su base canónica. Así, también los signos seguirán dicho orden: $\epsilon_0 = -1$, mientras que $\epsilon_1 = \dots = \epsilon_n = +1$.

Definición 1.22 (Definición del índice del producto escalar). *Dada una base ortonormal $\{e_1, \dots, e_n\}$ de V , el número de signos negativos en $\{\epsilon_1, \dots, \epsilon_n\}$ es el índice del producto escalar $\langle \cdot, \cdot \rangle$ de V . Cuando no haya posibilidad de confusión del producto escalar en cuestión, simplemente diremos que este número es el índice de V y lo denotaremos por $\text{ind } V$.*

Por supuesto, debemos mostrar que el índice no depende de la base ortonormal considerada. Para esto, daremos una caracterización alternativa de este número.

Proposición 1.23. *Sea W un subespacio maximal de V (con respecto de la contención) donde el producto escalar $\langle \cdot, \cdot \rangle$ es negativo definido. Entonces, para cualquier base ortonormal, el índice es igual a la dimensión de W .*

Demostración. Sean k el índice respecto de una base ortonormal. Si $k = 0$ o $k = \dim V$, el producto escalar es definido y tenemos que $W = \{0\}$ o $W = V$ respectivamente, de donde se sigue la afirmación. Así, supondremos que $0 < k < \dim V$. Denotemos $\nu = \dim W$.

Es claro que el producto escalar es negativo definido en $\text{span}\{e_1, \dots, e_k\}$ y por tanto $k \leq \nu$.

Para probar que $\nu \leq k$, consideremos la *proyección ortogonal* $\pi: W \rightarrow \text{span}\{e_1, \dots, e_k\}$ definida por

$$\pi(w) = - \sum_{i=1}^k \langle w, e_i \rangle e_i.$$

Queremos ver que π es inyectiva: Si $\pi(w) = 0$, entonces $\langle w, e_i \rangle = 0$ para todo $i = 1, \dots, k$ y

$$w = \sum_{i=k+1}^{\dim V} \langle w, e_i \rangle e_i,$$

lo que implica que

$$\langle w, w \rangle = \sum_{i=k+1}^{\dim V} \langle w, e_i \rangle^2;$$

como $w \in W$, el lado izquierdo de esta igualdad es menor o igual a cero, lo que implica que $\langle w, e_i \rangle = 0$ para todo $i = k+1, \dots, \dim V$. En resumen, $w = 0$ y π es inyectiva, de modo que $\nu \leq k$. Esto concluye la demostración. $\square$

Corolario 1.24. *Si W es un subespacio no degenerado de V , entonces*

$$\text{ind } V = \text{ind } W + \text{ind } W^\perp.$$

Observación 1.25. Es claro que el espacio $\mathbb{R}_\nu^m$ tiene índice ν . En concordancia con esta notación, y con el hecho (por demostrar) que todo espacio $(V, \langle, \rangle)$ de dimensión finita será isométrico a algún espacio semi-euclidiano, denotaremos el índice de V por la misma letra ν . En realidad, en estas notas nos interesará sobre todo el caso $\nu = 1$, pero antes de estudiar con detalle este importante caso, en la siguiente sección daremos la demostración del hecho antes mencionado.

1.3. Isometrías

Consideremos dos espacios con producto escalar $(V, \langle, \rangle)$ y $(W, \langle, \rangle)$. Observe que no utilizamos una notación especial para el producto de cada uno de estos espacios; de nuevo, para mayor sencillez de la notación.

Definición 1.26. Una transformación $T: V \rightarrow W$ *preserva el producto escalar* si para cualesquiera $v_1, v_2 \in V$,

$$\langle Tv_1, Tv_2 \rangle = \langle v_1, v_2 \rangle.$$

Entre todas estas transformaciones, nos fijaremos sobre todo en las que mandan el origen de V en el origen de W .

Proposición 1.27. *Sea $T: V \rightarrow W$ una transformación biyectiva que preserva el producto escalar. $T(0) = 0$ si y sólo si T es una transformación lineal.*

Demostración. La implicación “ $\Leftarrow$ ” es clara, de modo que supondremos que $T(0) = 0$ y mostraremos que T es lineal. Para ver que $T(v_1 + v_2) = Tv_1 + Tv_2$ para cualesquiera $v_1, v_2 \in V$, mostraremos que

$$\langle T(v_1 + v_2) - Tv_1 - Tv_2, w \rangle = 0$$

para cualesquiera $w \in W$. Como T es biyectiva, existe $u \in V$ tal que $T(u) = w$ y tenemos que

$$\begin{aligned} \langle T(v_1 + v_2) - Tv_1 - Tv_2, Tu \rangle &= \langle T(v_1 + v_2), Tu \rangle - \langle Tv_1, Tu \rangle - \langle Tv_2, Tu \rangle \\ &= \langle v_1 + v_2, u \rangle - \langle v_1, u \rangle - \langle v_2, u \rangle = 0, \end{aligned}$$

y usando que el producto escalar es no degenerado, obtenemos lo afirmado. La demostración de que $T(\lambda v) = \lambda T(v)$ es análoga. De nuevo, si $w \in W$ es arbitrario y $T(u) = w$, entonces

$$\begin{aligned} \langle T(\lambda v) - \lambda Tv, w \rangle &= \langle T(\lambda v) - \lambda Tv, Tu \rangle \\ &= \langle T(\lambda v), Tu \rangle - \lambda \langle Tv, Tu \rangle \\ &= \langle \lambda v, u \rangle - \lambda \langle v, u \rangle = 0, \end{aligned}$$

y de nuevo usamos que el producto escalar es no degenerado. $\square$

Definición 1.28. Una *isometría lineal* de V a W es una transformación biyectiva $T: V \rightarrow W$ que preserva el producto escalar.

Teorema 1.29. *Dos espacios con producto escalar $(V, \langle \cdot, \cdot \rangle)$ y $(W, \langle \cdot, \cdot \rangle)$ admiten una isometría lineal $T: V \rightarrow W$ entre ellos si y sólo si tienen la misma dimensión e índice.*

Demostración. Si existe una isometría lineal $T: V \rightarrow W$, entonces $\dim V = \dim W$. Además, si $\{e_1, \dots, e_m\}$ es una base ortonormal para V , es claro que $\{Te_1, \dots, Te_m\}$ es una base ortonormal para W y los índices son iguales.

Ahora, si $\{e_1, \dots, e_m\}$ es una base ortonormal para V y $\{e'_1, \dots, e'_m\}$ es una base ortonormal para W , podemos definir una transformación lineal $T: V \rightarrow W$ exigiendo que $Te_i = e'_i$ para todo $i = 1, \dots, m$ y extendiendo por linealidad. Entonces es fácil ver que T es una isometría lineal. $\square$

Corolario 1.30. *Sea $(V, \langle \cdot, \cdot \rangle)$ un espacio vectorial con producto escalar con índice ν y $\dim V = m$. Entonces existe una isometría lineal $T: V \rightarrow \mathbb{R}_\nu^m$.*

1.4. Espacios vectoriales lorentzianos

A partir de este punto, V será un espacio vectorial con producto escalar con índice 1 y dimensión $n + 1 \geq 2$. En este caso, diremos que el producto escalar es *lorentziano* y al espacio V lo llamaremos un *espacio vectorial lorentziano*. Como ya hemos mostrado, V es isométrico a $\mathbb{R}_1^{n+1}$, pero trataremos de dar las definiciones y proposiciones de manera general.

Observemos que si W es un subespacio vectorial de V , entonces puede ocurrir que W sea degenerado o no; además, en este segundo caso, el índice de W puede ser 0 o 1. Esto nos da una clasificación de los subespacios, como sigue.

Definición 1.31 (Carácter causal de un subespacio). Sea W un subespacio vectorial de V . Entonces

- W es *tipo espacio* o *espacial* si W es no degenerado y su índice es 0; es decir, el producto escalar restringido a W es positivo definido.
- W es *tipo tiempo* o *temporal* si W es no degenerado y su índice es 1; es decir, W es un espacio vectorial lorentziano.
- W es *tipo luz* o *nulo* si W es degenerado.

En los siguientes resultados reunimos varias caracterizaciones de estos tipos de subespacios.

Proposición 1.32. Sean $(V, \langle \cdot, \cdot \rangle)$ un espacio vectorial lorentziano y $W \subset V$ un subespacio.

1. W es tipo espacio si y sólo si $W^\perp$ es tipo tiempo.
2. Si W es tipo espacio, entonces W no contiene vectores nulos.
3. Si $\dim W \geq 2$ y W no contiene vectores nulos, entonces W es tipo espacio.

Demostración. El primer inciso es consecuencia inmediata del corolario 1.24, pues en este caso $\text{ind } W + \text{ind } W^\perp = 1$.

La afirmación del segundo inciso es clara. Por otro lado, si W no es tipo espacio, entonces existe $u \in W \setminus \{0\}$ tal que $\langle u, u \rangle \leq 0$. Si $\langle u, u \rangle = 0$, W contiene un vector nulo. Si $\langle u, u \rangle < 0$, como $\dim W \geq 2$, existe $v \in W$ linealmente independiente de u . Como V es lorentziano, $\langle v, v \rangle \geq 0$. De nuevo, si $\langle v, v \rangle = 0$, hemos terminado. Si $\langle u, u \rangle > 0$, consideramos los vectores en el segmento entre u y v , de la forma $tu + (1 - t)v$, $t \in [0, 1]$, y definimos la función

$$P(t) = \langle tu + (1 - t)v, tu + (1 - t)v \rangle;$$

que en realidad es un polinomio de segundo grado en t . Como $P(0) < 0$ y $P(1) > 0$, existe $t_0 \in (0, 1)$ tal que $P(t_0) = 0$, lo que implica la existencia del vector nulo buscado. $\square$

Proposición 1.33. Sean $(V, \langle \cdot, \cdot \rangle)$ un espacio vectorial lorentziano y $W \subset V$ un subespacio con $\dim W \geq 2$. Las siguientes afirmaciones son equivalentes:

1. W es tipo tiempo.
2. W contiene dos vectores nulos linealmente independientes.
3. W contiene un vector tipo tiempo.

Demostración. (1) $\Rightarrow$ (2): Como $\dim W \geq 2$, W admite una base ortonormal $\{e_0, e_1, \dots, e_k\}$ con $k \geq 1$, donde $\langle e_0, e_0 \rangle = -1$ y, digamos, $\langle e_1, e_1 \rangle = 1$. Entonces $e_0 \pm e_1$ son vectores nulos linealmente independientes.

(2) $\Rightarrow$ (3): Sean u, v vectores nulos linealmente independientes. Entonces $\langle u, v \rangle \neq 0$ (ejercicio 11). Si $\langle u, v \rangle > 0$, entonces $u - v$ es un vector tipo tiempo en W ; en caso contrario, elegimos el vector $u + v$.

(3) $\Rightarrow$ (1): Si $u \in W$ es tipo tiempo, por la proposición 1.33 tenemos que $\{u\}^\perp$ es tipo espacio. Observemos que $W^\perp \subset \{u\}^\perp$, de modo que $W^\perp$ es tipo espacio. Usamos de nuevo la citada proposición para obtener que $(W^\perp)^\perp = W$ es tipo tiempo. $\square$

Las afirmaciones para el caso degenerado son consecuencia de las proposiciones anteriores.

Proposición 1.34. Sean $(V, \langle \cdot, \cdot \rangle)$ un espacio vectorial lorentziano y $W \subset V$ un subespacio. Las siguientes afirmaciones son equivalentes:

1. W es nulo.
2. W contiene un vector nulo pero no un vector tipo tiempo.
3. $\dim(W \cap W^\perp) = 1$.

Para cerrar esta sección, veamos una caracterización del cono de tiempo (definición 1.7) en el contexto lorentziano.

Proposición 1.35. Sea $(V, \langle \cdot, \cdot \rangle)$ un espacio vectorial lorentziano. Sean $u, v \in V$ vectores tipo tiempo. Entonces $v \in CT(u)$ si y sólo si $CT(v) = CT(u)$.

Demostración. Observemos que cada vector v pertenece a su propio cono de modo que la implicación. " $\Leftarrow$ " es directa.

Sin pérdida de generalidad podemos suponer que $\|u\| = 1$. Como u es tipo tiempo, $V = \text{span}\{u\} \oplus \{u\}^\perp$ y podemos escribir

$$v = \lambda u + v', \quad w = \mu u + w', \quad \lambda, \mu \in \mathbb{R}, \quad v', w' \in \{u\}^\perp.$$

Como v es tipo tiempo, $\langle v, v \rangle < 0$, lo que implica $\|v'\|^2 < \lambda^2$ y $\|v'\| < \lambda$. Análogamente, w tipo tiempo implica que $\|w'\| < \mu$, mientras que $v \in CT(u)$ implica $\lambda > 0$.

Para probar que $CT(v) = CT(u)$, consideremos primero $w \in CT(u)$; entonces $\langle u, w \rangle < 0$ y por tanto, $\mu > 0$, de modo que

$$\langle v, w \rangle = -\lambda\mu + \langle v', w' \rangle \leq -\lambda\mu + \|v'\|\|w'\| < 0,$$

donde hemos usado la desigualdad de Cauchy-Schwarz para el espacio vectorial de tipo espacio $\{u\}^\perp$. Así, $w \in CT(v)$ y $CT(u) \subset CT(v)$. Podemos parafrasear lo que demostramos, diciendo que si dos vectores v, w están en un mismo cono $CT(u)$, entonces uno está en el cono del otro.

Para demostrar que $CT(v) \subset CT(u)$, sea $w \in CT(v)$. Como $v \in CT(u)$ si y sólo si $u \in CT(v)$, tenemos de nuevo dos vectores u, w que están en un mismo cono $CT(v)$, de modo que por la observación anterior obtenemos que $w \in CT(u)$. $\square$

Con la mira puesta en una aplicación posterior, recordemos que aunque en muchos casos lo usual es utilizar bases ortonormales, también es posible trabajar con otros tipos de bases. A continuación definimos otro tipo de base que son útiles en el contexto lorentziano.

Definición 1.36 (Base pseudo-ortonormal). Sea $(V, \langle \cdot, \cdot \rangle)$ un espacio vectorial lorentziano, con $\dim V = n + 1 \geq 2$. Una base $\{v_1, \dots, v_{n+1}\}$ de V es *pseudo-ortonormal* si se cumplen las siguientes condiciones:

$$\langle v_1, v_1 \rangle = \langle v_2, v_2 \rangle = 0, \quad \langle v_1, v_2 \rangle = 1; \quad \langle v_1, v_i \rangle = \langle v_2, v_i \rangle = 0,$$

y $\langle v_j, v_i \rangle = \delta_{ij}$ para $i, j = 3, \dots, n + 1$.

En otras palabras, una base pseudo-ortonormal es casi ortonormal, excepto por los dos primeros elementos de la base, que son vectores nulos. La condición $\langle v_1, v_2 \rangle = 1$ se puede considerar como una condición de normalización.

Gracias al teorema 1.19, es fácil mostrar la existencia de una base pseudo-ortonormal:

Proposición 1.37 (Existencia de una base pseudo-ortonormal). *Todo espacio vectorial lorentziano $(V, \langle \cdot, \cdot \rangle)$ admite una base pseudo-ortonormal.*

Demostración. El teorema 1.19 implica que V admite una base ortonormal $\{e_1, \dots, e_{n+1}\}$; en particular, $\langle e_1, e_1 \rangle = -1$, $\langle e_2, e_2 \rangle = 1$ y $\langle e_1, e_2 \rangle = 0$. Definimos $v_i = e_i$ para $i = 3, \dots, n + 1$ y

$$v_1 = \frac{1}{\sqrt{2}}(e_1 + e_2) \quad \text{y} \quad v_2 = -\frac{1}{\sqrt{2}}(e_1 - e_2). \quad (1.7)$$

Es fácil ver que $\{v_1, \dots, v_{n+1}\}$ es una base pseudo-ortonormal para V . $\square$

Por supuesto, también se puede obtener una base ortonormal a partir de una pseudo-ortonormal; simplemente resolvemos el sistema de ecuaciones (1.7) para obtener e_1, e_2 en términos de v_1, v_2 . Como es el caso de todas las situaciones que implican el uso de bases, podemos escoger una base según nuestra conveniencia.

Proposición 1.38 (Expresión de un vector con respecto de una base pseudo-ortonormal). *Sea $\{v_1, \dots, v_{n+1}\}$ una base pseudo-ortonormal de V . Si $v \in V$, entonces*

$$v = \langle v, v_2 \rangle v_1 + \langle v, v_1 \rangle v_2 + \sum_{i=3}^{n+1} \langle v, e_i \rangle e_i.$$

1.5. Isometrías en espacios lorentzianos

Anteriormente vimos que las isometrías biyectivas que dejan fijo al origen de $(V, \langle, \rangle)$ son en realidad transformaciones lineales. Es fácil ver que este conjunto de transformaciones es un grupo bajo la operación de composición. Pero antes de entrar en más detalles, veamos dos ejemplos sencillos de estas transformaciones.

Ejemplo 1.39 (Rotaciones). Sean $V = \mathbb{R}^2$ y $\theta \in \mathbb{R}$. Recordemos que la rotación R_θ en torno del eje t con ángulo θ tiene asociada la siguiente matriz con respecto de la base canónica:

$$\begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}.$$

Podemos usar esta matriz para construir una isometría lineal en el espacio $\mathbb{R}_1^3$, que también llamamos la rotación R_θ con ángulo θ :

$$\begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & -\sin \theta \\ 0 & \sin \theta & \cos \theta \end{pmatrix}.$$

Es fácil convencerse de que R_θ es una isometría lineal. En el caso general de $\mathbb{R}_1^{n+1}$, si pensamos a $SO(n)$ como el conjunto de matrices ortogonales $n \times n$ con determinante 1 y $R \in SO(n)$, entonces la matriz

$$\bar{R} = \begin{pmatrix} 1 & 0 \\ 0 & R \end{pmatrix}.$$

representa una isometría lineal llamada *rotación propia* en $\mathbb{R}_1^{n+1}$.

Ejemplo 1.40 (Empujones¹). Sean $V = \mathbb{R}_1^2$ y $\theta \in \mathbb{R}$. El empujón E_θ con ángulo θ es el operador en $\mathbb{R}_1^2$ con la siguiente matriz asociada con respecto de la base canónica:

$$\begin{pmatrix} \cosh \theta & \sinh \theta \\ \sinh \theta & \cosh \theta \end{pmatrix}.$$

Como en el caso de las rotaciones, podemos usar esta matriz para definir una isometría lineal en $\mathbb{R}_1^{n+1}$, llamada también empujón, usando la matriz

$$\begin{pmatrix} \cosh \theta & \sinh \theta & 0 \\ \sinh \theta & \cosh \theta & 0 \\ 0 & 0 & I_{n-1} \end{pmatrix},$$

donde I_{n-1} denota la matriz identidad de $(n-1) \times (n-1)$.

Consideremos el estudio de las isometrías lineales desde el punto de vista matricial. Si V es un espacio vectorial lorentziano, fijemos una base ortonormal $\{e_0, e_1, \dots, e_n\}$ de V e $I_{1,n} = \text{diag}(-1, 1, \dots, 1)$ la matriz asociada al producto

¹La palabra en inglés es *boost*, que también se puede traducir como estímulo, incremento, incentivo, impulso.

escalar; es decir, si $u, v \in V$ y $[u], [v]$ son sus expresiones en coordenadas con respecto de la base dada (como vectores columna), entonces

$$\langle u, v \rangle = [u]^t I_{1,n} [v],$$

donde t denota la transpuesta de la matriz correspondiente.

Sea $A = (a_{ij})$ la matriz $(n+1) \times (n+1)$ asociada a la isometría lineal $T: V \rightarrow V$ con respecto de esta base, que se calcula de la manera usual:

$$Te_j = \sum_{i=0}^n a_{ij} e_i.$$

T es una isometría lineal si y sólo si $\langle Tu, Tv \rangle = \langle u, v \rangle$, o en términos matriciales, $(A[u])^t I_{1,n} A[v] = [u]^t I_{1,n} [v]$; como esto debe valer para todo u, v , T es una isometría lineal si y sólo si $A^t I_{1,n} A = I_{1,n}$. Calculando el determinante a ambos lados de esta ecuación, obtenemos que $\det A = \pm 1$.

Definición 1.41. Sea $O(V)$ el grupo de isometrías lineales de V ; matricialmente, $O(V)$ es el conjunto de matrices invertibles A de $(n+1) \times (n+1)$ tales que $A^t I_{1,n} A = I_{1,n}$. Llamamos a este grupo el *grupo ortogonal* de V , que por la observación anterior a esta definición es la unión de dos conjuntos:

$$SO(V) = \{A \in O(V) : \det A = 1\} \quad \text{y} \quad \{A \in O(V) : \det A = -1\};$$

el primero de estos conjuntos es un subgrupo de $O(V)$, llamado el *grupo ortogonal especial* de V .

La diferencia entre los casos $+1$ y -1 tiene que ver con el concepto de orientación, que recordamos a continuación.

Definición 1.42 (Orientación). Sea V un espacio vectorial de dimensión finita. Definimos una relación entre las bases ordenadas de V de la siguiente manera: Dos bases están relacionadas si y sólo si la matriz de cambio de base tiene determinante positivo. Ésta es una relación de equivalencia con dos clases de equivalencia; cada clase es una *orientación* para V . La elección de una de estas clases determina si una base es *positiva*, cuando pertenece a la clase elegida; o bien si la base es *negativa*, cuando ocurre lo contrario.

Así, en el caso de un espacio vectorial general con producto escalar V , el grupo $O(V)$ tiene al menos dos componentes, pues podemos separar a este conjunto en dos subconjuntos, dependiendo si las transformaciones preservan o invierten la orientación.

Para el caso en que V es un espacio lorentziano, podemos estudiar con más precisión la estructura de $O(V)$, pues en este caso tenemos una dirección distinguida, la llamada dirección hacia el futuro (o hacia el pasado). Como en el caso del concepto de orientación, podemos distinguir de entre las bases aquellas cuyo vector tipo tiempo *apunta hacia el futuro* y aquellas tales que dicho vector *apunta hacia el pasado*.

Definición 1.43 (Orientación del tiempo o temporal). Sea V un espacio vectorial de dimensión finita. Definimos una relación entre las bases de V de la siguiente manera: Sean $\{u_0, u_1, \dots, u_n\}$ y $\{v_0, v_1, \dots, v_n\}$ dos bases de modo que u_0, v_0 sean tipo tiempo. Entonces las bases están relacionadas si y sólo si $v_0 \in CT(u_0)$; es decir, si y sólo si $\langle u_0, v_0 \rangle < 0$. Como en el caso de la orientación, ésta es una relación de equivalencia con dos clases de equivalencia; cada clase es una *orientación del tiempo o temporal* para V . La elección de una de estas clases determina si una base *apunta hacia el futuro*, cuando pertenece a la clase elegida; o bien si *apunta hacia el pasado*, cuando ocurre lo contrario.

Ejemplo 1.44 (Orientaciones canónicas de $\mathbb{R}_1^{n+1}$). Como de costumbre, la orientación canónica de $\mathbb{R}_1^{n+1}$ es aquella donde la base canónica $\{e_0, e_1, \dots, e_n\}$ es positiva. Por otro lado, la orientación canónica del tiempo en $\mathbb{R}_1^{n+1}$ es aquella donde e_0 apunta hacia el futuro.

Lo anterior sugiere que $O(V)$ podría tener cuatro componentes, dependiendo de cómo actúe una transformación T con respecto de estos dos tipos de orientación. Antes de mostrar que éste es el caso, daremos dos definiciones.

Definición 1.45 (Grupos ortogonal ortócrono y ortogonal especial ortócrono). Una isometría lineal $T: V \rightarrow V$ es *ortócrona* si T preserva la orientación del tiempo; es decir, manda una base que apunta hacia el futuro en una base que apunta hacia el futuro.

El *grupo ortogonal ortócrono* de V es el subgrupo $O^+(V)$ de $O(V)$ formado por todas las isometrías lineales ortogonales ortócronas.

El *grupo ortogonal especial ortócrono* de V es el subgrupo $SO^+(V)$ de $O^+(V)$ formado por las isometrías lineales ortogonales que preservan tanto la orientación como la orientación del tiempo.

Observación 1.46. En la literatura hay varias notaciones para estos conjuntos de transformaciones. Para el caso de espacios vectoriales lorentzianos a menudo se usa la notación $O(1, n)$ para $O(V)$, así como $SO^+(1, n)$ para $SO^+(V)$. Puesto que en estas notas estamos distinguiendo este caso lorentziano, de ahora en adelante usaremos estas notaciones.

En breve mostraremos que $O(1, n)$ tiene cuatro componentes conexos y que $SO^+(1, n)$ se puede caracterizar como el componente conexo de $O(1, n)$ que contiene a la identidad. Pero primero nos detendremos un poco para analizar el caso de $SO^+(1, n)$.

Teorema 1.47 (Descomposición en valores singulares de $SO^+(1, n)$). Sea $A \in SO^+(1, n)$. Entonces existen $P, Q \in SO(n)$, $\alpha \in \mathbb{R}$ tales que

$$A = \begin{pmatrix} 1 & 0 \\ 0 & P \end{pmatrix} \begin{pmatrix} \cosh \alpha & \sinh \alpha & 0 \\ \sinh \alpha & \cosh \alpha & 0 \\ 0 & 0 & I_{n-1} \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & Q^t \end{pmatrix}, \quad (1.8)$$

donde Q^t denota la transpuesta de Q , como es usual.

Para la demostración, véase, por ejemplo, [4, p. 169]. Con esto a la mano, podemos ya dar una idea de la demostración del teorema que cierra esta sección.

Teorema 1.48. *El grupo ortogonal $O(1, n)$ tiene cuatro componentes:*

- *El grupo ortogonal especial ortócrono $SO^+(1, n)$.*
- *El conjunto de isometrías lineales que preservan la orientación pero invierten la orientación del tiempo.*
- *El conjunto de isometrías lineales que invierten la orientación pero preservan la orientación del tiempo.*
- *El conjunto de isometrías lineales que invierten tanto la orientación como la orientación del tiempo.*

Demostración. Primero veamos que los cuatro conjuntos definidos en el enunciado del teorema son homeomorfos entre sí. Por ejemplo, si definimos una transformación de $O(1, n)$ en sí mismo mandando una matriz $A \in O(1, n)$ en la matriz $I_{1,n}A$, donde $I_{1,n} = \text{diag}(-1, 1, \dots, 1)$, entonces la transformación $A \mapsto I_{1,n}A$ restringida a $SO^+(1, n)$ es un homeomorfismo de este conjunto con el último conjunto de la lista. Los demás casos son análogos.

Basta entonces mostrar que $SO^+(1, n)$ es conexo. La idea es mostrar que es conexo por trayectorias. Más precisamente, veamos que para cualquiera $A \in SO^+(1, n)$ existe una curva $c: [a, b] \rightarrow SO^+(1, n)$ que le une con la matriz identidad I_{n+1} . Por el teorema de descomposición 1.47, A se puede escribir en la forma (1.8) para ciertas $P, Q \in SO(n)$ y $\alpha \in \mathbb{R}$. Ahora, usamos el hecho de que la exponencial $\exp: \mathfrak{so}(1, n) \rightarrow SO(1, n)$ es suprayectiva² para escribir

$$A = \exp(P') \begin{pmatrix} \cosh \alpha & \sinh \alpha & 0 \\ \sinh \alpha & \cosh \alpha & 0 \\ 0 & 0 & I_{n-1} \end{pmatrix} \exp(Q').$$

Definimos entonces la curva $c: [0, 1] \rightarrow SO(1, n)$ como

$$c(t) = \exp(tP') \begin{pmatrix} \cosh t\alpha & \sinh t\alpha & 0 \\ \sinh t\alpha & \cosh t\alpha & 0 \\ 0 & 0 & I_{n-1} \end{pmatrix} \exp(tQ');$$

claramente la curva c es la trayectoria buscada. □

1.6. Transformaciones lineales autoadjuntas

Puesto que más adelante tendremos que considerar este tipo de transformaciones, es conveniente reunir aquí algunos hechos básicos.

²Ver el Corolario 11.10 en [5].

Como de costumbre, $(V, \langle, \rangle)$ denota un espacio vectorial lorentziano. También sabemos que una transformación lineal $T: V \rightarrow V$ es *autoadjunta* con respecto del producto escalar dado en V si y sólo si para cada $u, v \in V$,

$$\langle Tu, v \rangle = \langle u, Tv \rangle.$$

Cuando es claro a cuál producto escalar nos referimos, basta decir que T es autoadjunta. Recordemos también que si T es autoadjunta con respecto de un producto escalar positivo definido, la transformación en realidad es diagonalizable, es decir, existe una base ortonormal de vectores propios de T . Podemos ver que esto es diferente en presencia de un producto escalar diferente.

Ejemplo 1.49. Sea $V = \mathbb{R}_1^2$ y $T: V \rightarrow V$ dada en forma matricial por

$$\begin{pmatrix} a & b \\ -b & a \end{pmatrix}$$

con respecto de la base canónica. Es un sencillo ejercicio mostrar que T es autoadjunta con respecto del producto escalar en $\mathbb{R}_1^2$. Supongamos ahora que $b \neq 0$ y veamos que T no es diagonalizable. De hecho, T no tiene vectores propios con valores propios reales: Si $Tv = \lambda v$ y $v = (v_1, v_2)$ con respecto de la base canónica, tenemos el sistema

$$\begin{aligned} av_1 + bv_2 &= \lambda v_1, \\ -bv_1 + av_2 &= \lambda v_2; \end{aligned}$$

y queremos que su determinante sea igual a cero: $(a - \lambda)^2 + b^2 = 0$, pero esto es imposible debido a nuestra hipótesis sobre b .

Ejemplo 1.50. Recordemos que el bloque de Jordan más sencillo tiene la forma

$$\begin{pmatrix} a & 0 \\ 1 & a \end{pmatrix}.$$

Si consideramos a esta matriz como la asociada a una transformación lineal en $\mathbb{R}_1^2$ por medio de la base canónica, el lector se puede convencer que ni siquiera es autoadjunta. Sin embargo, recordemos que hay otras bases importantes en los espacios lorentzianos, como por ejemplo, la base formada por los vectores nulos $\{(1, 1), (-1, 1)\}$. Veamos que la matriz anterior sí representa a una transformación autoadjunta T , cuando nos referimos a esta base.

Como $T(1, 1) = a(1, 1) + 1(-1, 1)$, tenemos que

$$\langle T(1, 1), (-1, 1) \rangle = a\langle (1, 1), (-1, 1) \rangle = 2a.$$

Por otro lado, como $T(-1, 1) = 0(1, 1) + a(-1, 1)$, tenemos que

$$\langle T(-1, 1), (1, 1) \rangle = a\langle (-1, 1), (1, 1) \rangle = 2a.$$

Puesto que bastaba verificar este caso, tenemos que T es autoadjunta. Por otro lado, si buscamos los vectores propios de T , sabemos que satisfacen $Tv = \lambda v$, lo que se puede escribir como el sistema

$$\begin{aligned} av_1 &= \lambda v_1, \\ v_1 + av_2 &= \lambda v_2, \end{aligned}$$

cuyo polinomio característico es $(a - \lambda)^2$, el cual tiene a $\lambda = a$ como una raíz doble. Pero al regresar al sistema, esto implica que $v_1 = 0$, lo que nos dice que la transformación sólo tiene un subespacio propio de dimensión 1 y por tanto no es diagonalizable.

Estos ejemplos nos sugieren preguntarnos cuáles serán las formas canónicas adecuadas para las transformaciones autoadjuntas en espacios vectoriales lorentzianos.

Teorema 1.51 (Clasificación de transformaciones lineales autoadjuntas en espacios lorentzianos; ver [11] y [7]). *Sea A autoadjunta. Entonces A se puede escribir de alguna de las siguientes cuatro formas. Las dos primeras formas se calculan con respecto de una base ortonormal (la numeración sigue el orden de [7]):*

$$I: \begin{pmatrix} a_0 & 0 & \cdots & 0 \\ 0 & a_1 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & a_n \end{pmatrix}, \quad IV: \begin{pmatrix} a_0 & b_0 & 0 & \cdots & 0 \\ -b_0 & a_0 & 0 & \cdots & 0 \\ 0 & 0 & a_1 & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \cdots & a_{n-1} \end{pmatrix},$$

y las otras dos formas se calculan con respecto de una base pseudo-ortonormal:

$$II: \begin{pmatrix} a_0 & 0 & 0 & \cdots & 0 \\ 1 & a_0 & 0 & \cdots & 0 \\ 0 & 0 & a_1 & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \cdots & a_{n-1} \end{pmatrix}, \quad III: \begin{pmatrix} a_0 & 0 & 0 & 0 & 0 & \cdots & 0 \\ 0 & a_0 & 1 & 0 & 0 & \cdots & 0 \\ -1 & 0 & a_0 & 0 & 0 & \cdots & 0 \\ 0 & 0 & 0 & a_1 & 0 & \cdots & 0 \\ 0 & 0 & 0 & 0 & a_2 & \cdots & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & 0 & 0 & \cdots & a_{n-2} \end{pmatrix},$$

1.7. Ejercicios

1. Muestre que la ecuación (1.1) realmente define un producto escalar en $\mathbb{R}^n$.
2. Demuestre la proposición 1.4.
3. Demuestre la proposición 1.8.
4. Sea $(V, \langle \cdot, \cdot \rangle)$ un espacio vectorial con producto escalar. Muestre que para todo $A \subset V$, el conjunto $A^\perp$ (definición 1.13) es un subespacio de V .

5. Verifique que el producto cruz de $\mathbb{R}_1^3$ (definición 1.14) satisface

$$\langle u, v \times w \rangle = \det(u, v, w).$$

para todo $u, v, w \in \mathbb{R}_1^3$. Concluya que $v \times w$ es ortogonal a v, w .

6. Demuestre la proposición 1.20.
7. Demuestre la proposición 1.38.
8. Pruebe el corolario 1.24.
9. Concluya la demostración del teorema 1.29, probando que T es una isometría lineal.
10. Sea $(V, \langle \cdot, \cdot \rangle)$ un espacio vectorial lorentziano.
- Sean $u, v \in V$ vectores tipo tiempo. Muestre que $\langle u, v \rangle \neq 0$. Sugerencia: Escriba $v = \lambda u + v'$ con $v' \in \{u\}^\perp$ y suponga que $\langle u, v \rangle = 0$.
 - Sean $u, v \in V$ tales que u es tipo tiempo y v es nulo. Muestre que $\langle u, v \rangle \neq 0$.
 - Concluya que si $u, v \in V$ son vectores causales con $\langle u, v \rangle = 0$, entonces u, v son vectores nulos.
11. Sea $(V, \langle \cdot, \cdot \rangle)$ un espacio vectorial lorentziano. Demuestre que dos vectores nulos u, v son linealmente independientes si y sólo si $\langle u, v \rangle \neq 0$. Esto se usó en la demostración de la proposición 1.33.
12. Demuestre la proposición 1.34.
13. Pruebe que si u es un vector tipo nulo entonces existen vectores nulos v, w tal que $\langle u, v \rangle = -1$ y $\langle u, w \rangle = 1$.
14. Pruebe que la transformación dada en el Ejemplo 1.40 es en realidad una isometría lineal en $\mathbb{R}_1^3$.

Capítulo 2

Curvas y superficies en $\mathbb{R}_1^3$

Aquí estudiaremos los objetos geométricos más fáciles de visualizar para los espacios ambiente más sencillos, $\mathbb{R}_1^2$ y $\mathbb{R}_1^3$. En realidad, aunque la teoría de curvas es amplia (ver, por ejemplo, [6]), nos centraremos más en el caso de las superficies.¹

2.1. Curvas

Como de costumbre, una *curva parametrizada diferenciable* en $\mathbb{R}_1^3$ es la imagen de una transformación diferenciable $\alpha: I \rightarrow \mathbb{R}_1^3$, donde I es un intervalo en $\mathbb{R}$. Aquí usamos la misma estructura diferenciable de $\mathbb{R}^3$ para definir la diferenciable de la curva. También decimos que α es una *curva parametrizada regular* si $\alpha'(s) \neq 0$ para todo $s \in I$. Usualmente, trabajaremos con curvas regulares, a menos que se indique lo contrario.

Ejemplo 2.1. Consideremos la circunferencia $\alpha: \mathbb{R} \rightarrow \mathbb{R}_1^2$ parametrizada por

$$\alpha(s) = (\cos s, \sin s), \quad s \in \mathbb{R}.$$

Es claro que α es una curva regular, pues $\alpha'(s) = (-\sin s, \cos s) \neq 0$ para todo s . También es fácil ver que el carácter causal de α' va cambiando, pues por ejemplo, $\alpha'(0) = (0, 1)$ es un vector de tipo espacio, $\alpha'(\pi/4) = (1/\sqrt{2}, 1/\sqrt{2})$, que es nulo, y $\alpha'(\pi/2) = (1, 0)$, que es tipo tiempo.

Para nuestros fines, será mejor estudiar curvas tales que el carácter causal de α' no varíe. Así, nos fijaremos en los siguientes tipos de curvas:

Definición 2.2 (Carácter causal de una curva). Sea $\alpha: I \rightarrow \mathbb{R}_1^3$ una curva parametrizada regular. Entonces

- α es *tipo tiempo* si $\alpha'(s)$ es tipo tiempo para todo $s \in I$;

¹Recordemos que en casos de dimensión baja denotaremos a las coordenadas cartesianas como (t, x) o (t, x, y) según sea el caso.

- α es *nula* si $\alpha'(s)$ es un vector nulo para todo $s \in I$;
- α es *tipo espacio* si $\alpha'(s)$ es tipo espacio para todo $s \in I$;

La circunferencia arriba mencionada queda entonces eliminada de las curvas interesantes. Sin embargo, no sólo estamos quitando de nuestra lista a esta curva, sino que también quitaremos a muchas curvas cerradas, como veremos a continuación. Por comodidad, diremos que una curva parametrizada $\alpha: \mathbb{R} \rightarrow \mathbb{R}_1^3$ es *cerrada* si es periódica; es decir, existe $T > 0$ tal que $\alpha(s+T) = \alpha(s)$ para todo $s \in I$. Por lo general, consideraremos que T denota el menor número positivo con esta propiedad.

Proposición 2.3. *No existen curvas parametrizadas regulares cerradas tipo tiempo o nulas en $\mathbb{R}_1^3$.*

Demostración. Sea $\alpha: \mathbb{R} \rightarrow \mathbb{R}_1^3$ una curva regular cerrada, con coordenadas $(t(s), x(s), y(s))$. Como α es periódica, su imagen es compacta, de modo que existe un punto $s_0 \in \mathbb{R}$ donde $t(s)$ alcanza su máximo. En este caso, $\alpha'(s_0)$ tiene la forma $(0, x'(s_0), y'(s_0))$, que es un vector diferente de cero y por tanto $\langle \alpha'(s_0), \alpha'(s_0) \rangle > 0$, lo cual no puede ocurrir si α es tipo tiempo o nula. $\square$

En el caso de las curvas tipo espacio o tipo tiempo, podemos definir la *longitud de arco* de una curva $\alpha: I \rightarrow \mathbb{R}_1^3$ entre dos puntos $a, b \in I$ como

$$\text{long}(\alpha)|_a^b = \int_a^b \|\alpha'(s)\| ds,$$

y usar esta función para *parametrizar por longitud de arco*, obteniendo una nueva curva β con la misma imagen que α , pero tal que $\langle \beta', \beta' \rangle = \epsilon$, donde $\epsilon = 1$ si α era de tipo espacio y $\epsilon = -1$ si α era de tipo tiempo.

Sin embargo, el caso de las curvas nulas es muy diferente. Como $\alpha'(s)$ es un vector nulo, $\langle \alpha'(s), \alpha'(s) \rangle = 0$, de modo que al derivar esta expresión obtenemos $\langle \alpha'(s), \alpha''(s) \rangle = 0$, lo que nos dice que $\alpha''(s)$ vive en el plano nulo $(\alpha'(s))^\perp$. Esto separa el análisis en dos casos, cuando el vector $\alpha''(s)$ es nulo o cuando es tipo espacio; ver la proposición 1.34.

Si $\alpha''(s)$ es un vector nulo, entonces $\alpha'(s)$ y $\alpha''(s)$ son linealmente dependientes (ejercicio 11 del capítulo anterior) y entonces α es una *pregeodésica*, término que no definiremos por el momento; ver el ejercicio 19, capítulo 3 de [10].

Si $\alpha''(s)$ es tipo espacio, entonces es posible reparametrizar la curva y obtener una nueva curva β de modo que $\|\beta''\| = 1$, que es llamada una *parametrización por pseudo-longitud de arco*; ver por ejemplo [6].

No nos extenderemos mucho en el estudio de las curvas $\alpha: I \rightarrow \mathbb{R}_1^3$. Sólo diremos que en analogía con el caso euclidiano, es posible construir, para cada $s \in I$, una base $\{T(s), N(s), B(s)\}$ de $\mathbb{R}_1^3$ llamada *triedro de Frenet*, que se puede usar para definir la curvatura y la torsión de la curva. Un buen recuento de los resultados en esta dirección aparece en [6].

2.2. Superficies en $\mathbb{R}_1^3$

La definición de superficie en $\mathbb{R}_1^3$ es idéntica a la definición para $\mathbb{R}^3$, puesto que la definición sólo depende de la estructura diferenciable del espacio ambiente y no de su métrica (producto escalar). Por completez, la incluimos.

Definición 2.4. Una *superficie diferenciable* en $\mathbb{R}_1^3$ es un conjunto $S \subset \mathbb{R}_1^3$ tal que para cada punto $p \in S$ existen una vecindad U de p en S (con la topología inducida por $\mathbb{R}_1^3$) y una transformación diferenciable $\varphi: V \rightarrow U$, donde V es un conjunto abierto en $\mathbb{R}^2$, cuya diferencial $d\varphi_q$ es inyectiva para todo $q \in V$. φ es una *parametrización* de la vecindad U .

En otras palabras, una superficie diferenciable es localmente la imagen de una inmersión de un abierto de $\mathbb{R}^2$ en $\mathbb{R}_1^3$. Para construir algunos ejemplos de superficies diferenciables, los siguientes dos métodos son los más usuales.

Ejemplo 2.5 (Gráfica de una función). Sea $V \subset \mathbb{R}^2$ un conjunto abierto, y $f: V \rightarrow \mathbb{R}$ una función diferenciable. Entonces definimos la *gráfica de f* como el conjunto

$$\text{Gráf}(f) = \{ (t, x, y) \in \mathbb{R}_1^3 \mid t = f(x, y), (x, y) \in V \};$$

nótese que el orden de las coordenadas es el inverso al orden usual; es decir, los puntos de la gráfica tienen las coordenadas $(f(x, y), x, y)$. Por supuesto, también es posible considerar funciones de (t, x) o (t, y) , lo que haremos más adelante.

Para la siguiente manera de construir ejemplos de superficies, hay que introducir cierta terminología.

Definición 2.6 (Puntos/valores críticos/regulares). Sea W un conjunto abierto en $\mathbb{R}_1^3$, $F: W \rightarrow \mathbb{R}$ una función diferenciable, $p \in W$ y $s \in \mathbb{R}$.

- p es un *punto crítico* de F si $dF_p = 0$.
- p es un *punto regular* de F si no es un punto crítico.
- s es un *valor crítico* de F si existe un punto crítico $p \in W$ tal que $f(p) = s$.
- s es un *valor regular* de F si no es un valor crítico; en particular, si $s \notin \text{Im}(F)$, entonces s es un valor regular.

El siguiente resultado es consecuencia del teorema de la función inversa.

Proposición 2.7 (Imagen inversa de un valor regular). Sean W un conjunto abierto en $\mathbb{R}_1^3$, $F: W \rightarrow \mathbb{R}$ una función diferenciable y $s \in \mathbb{R}$ un valor regular de F . Entonces $F^{-1}(s)$ es una superficie diferenciable.

Observemos que hasta el momento no hemos utilizado el hecho de la existencia de una métrica lorentziana. Al igual que en el caso de las curvas, algunas superficies se pueden clasificar en términos del carácter causal de sus planos tangentes.

Definición 2.8 (Clasificación causal de las superficies). Sea S una superficie diferenciable en $\mathbb{R}_1^3$. Decimos que

- S es *tipo espacio* si para todo $p \in S$, el espacio tangente $T_p S$ es tipo espacio;
- S es *tipo tiempo* si para todo $p \in S$, el espacio tangente $T_p S$ es tipo tiempo;
- S es *no degenerada* si es tipo espacio o tipo tiempo;
- S es *degenerada o nula* si para todo $p \in S$, el espacio tangente $T_p S$ es nulo.

Ejemplo 2.9. Sea $v \in \mathbb{R}_1^3 \setminus \{0\}$ un vector fijo y S el plano $\text{span}\{v\}^\perp$. Para todo $p \in S$, $T_p S = S$ y S es tipo espacio (tipo tiempo, nulo, respectivamente) si v es tipo tiempo (tipo espacio, nulo, respectivamente).

Debido a la correspondencia entre el carácter causal de un plano con el de su conjunto ortogonal indicada en el ejemplo anterior, si una superficie S admite un campo vectorial ξ que nunca se anula y genera a $T_p S^\perp$ para cada $p \in S$, podemos analizar el carácter causal de S por medio del carácter causal de ξ . Veamos algunos ejemplos.

Ejemplo 2.10 (La esfera). Consideremos $\mathbb{S}^2$, el conjunto de puntos $(t, x, y) \in \mathbb{R}_1^3$ tales que $t^2 + x^2 + y^2 = 1$; es decir, la esfera unitaria con respecto de la métrica euclidiana. Podemos usar la proposición 2.7, puesto que $\mathbb{S}^2 = F^{-1}(1)$, donde $F(t, x, y) = t^2 + x^2 + y^2$. El número 1 es un valor regular de F y por tanto $F^{-1}(1)$ es una superficie diferenciable.

Afirmamos que para cada punto $(t_0, x_0, y_0) \in \mathbb{S}^2$, el vector $N = (-t_0, x_0, y_0)$ es un vector normal a $\mathbb{S}^2$ en (t_0, x_0, y_0) . Para ver esto, probaremos que este vector es ortogonal a cualquier vector tangente v a $\mathbb{S}^2$. Si $\alpha: (-\epsilon, \epsilon) \rightarrow \mathbb{S}^2$ es una curva diferenciable tal que $\alpha(0) = p$ y $\alpha'(0) = v$, con $\alpha(s) = (t(s), x(s), y(s))$, entonces $t(s)^2 + x(s)^2 + y(s)^2 = 1$, de modo que derivando obtenemos

$$2(t(s)t'(s) + x(s)x'(s) + y(s)y'(s)) = 0, \quad s \in (-\epsilon, \epsilon),$$

y al evaluar en 0, podemos deducir de lo anterior que

$$\langle (t'(s_0), x'(s_0), y'(s_0)), (-t_0, x_0, y_0) \rangle = \langle v, N \rangle = 0.$$

Como el carácter causal de N varía, $\mathbb{S}^2$ no es una superficie de alguno de los tres tipos definidos anteriormente.

Ejemplo 2.11 (El espacio de de Sitter). Un mejor análogo de la esfera unitaria, pero con respecto de la métrica en $\mathbb{R}_1^3$, es el conjunto

$$\mathbb{S}_1^2 = \{ (t, x, y) \in \mathbb{R}_1^3 \mid -t^2 + x^2 + y^2 = 1 \},$$

llamado el *espacio de de Sitter de dimensión 2*. Como en el caso anterior, este conjunto es la imagen inversa $G^{-1}(1)$, donde ahora $G(t, x, y) = -t^2 + x^2 + y^2$. De manera análoga al ejemplo anterior, observemos que un vector normal es $N = (t, x, y)$ y $\langle N, N \rangle = 1$, de modo que este vector normal es tipo espacio para todo punto de $\mathbb{S}_1^2$ y $\mathbb{S}_1^2$ es una superficie tipo tiempo.

Ejemplo 2.12 (El modelo del hiperboloide para el plano hiperbólico). Definimos

$$\mathbb{H}_0^2 = \{ (t, x, y) \in \mathbb{R}_1^3 \mid -t^2 + x^2 + y^2 = -1, t > 0 \};$$

que geoméricamente corresponde a la parte superior de un hiperboloide de dos hojas. En este caso, un vector normal es $N = (t, x, y)$, que es tipo tiempo. Así, $\mathbb{H}_0^2$ es una superficie tipo espacio. $\mathbb{H}_0^2$ tiene curvatura constante negativa, lo que lo hace un modelo para la geometría hiperbólica, llamado el modelo del hiperboloide.

En el ejemplo 2.5 consideramos una superficie S dada por la gráfica de una función $f: V \rightarrow \mathbb{R}$, como

$$\text{Gráf}(f) = \{ (t, x, y) \in \mathbb{R}_1^3 \mid t = f(x, y), (x, y) \in V \};$$

¿Cuándo la gráfica de la función es tipo espacio, tipo tiempo o nula? Fijemos un punto $(x_0, y_0) \in V$ y observemos que los vectores

$$\partial_x(x_0, y_0) = (f_x(x_0, y_0), 1, 0) \quad \text{y} \quad \partial_y(x_0, y_0) = (f_y(x_0, y_0), 0, 1)$$

forman una base del plano tangente $T_p S$, donde $p = (f(x_0, y_0), (x_0, y_0))$ y f_x, f_y denotan derivadas parciales en el punto correspondiente. El producto cruz de ellos es normal a $T_p S$, por lo que el campo vectorial dado por

$$\xi_{(x_0, y_0)} = -\partial_x(x_0, y_0) \times \partial_y(x_0, y_0) = (1, f_y(x_0, y_0), f_x(x_0, y_0))$$

es un campo vectorial normal a p para cada $p \in S$. Observando que

$$\langle \xi, \xi \rangle = -1 + f_x^2 + f_y^2 = -1 + \|\text{grad } f\|^2,$$

tenemos que

- S es tipo espacio si y sólo si $\|\text{grad } f\| < 1$ para todo $p \in S$;
- S es tipo tiempo si y sólo si $\|\text{grad } f\| > 1$ para todo $p \in S$;
- S es nula si y sólo si $\|\text{grad } f\| = 1$ para todo $p \in S$;

Un ejemplo sencillo de esta última condición consiste en considerar la función $f: \mathbb{R}^2 \setminus \{0\} \rightarrow \mathbb{R}$ dada por $f(x, y) = \sqrt{x^2 + y^2}$. Como sabemos, la gráfica de f es un cono (o la mitad de un cono, según se considere) sin el vértice.

Al igual que cuando analizamos a las curvas, no cualquier superficie cae en la clasificación de la definición 2.8:

Proposición 2.13. *Una superficie diferenciable compacta $S \subset \mathbb{R}_1^3$ no es tipo espacio, ni tipo tiempo, ni nula.*

Este hecho genera algunos problemas para el estudio de las superficies en $\mathbb{R}_1^3$; sin duda, el lector recordará muchas instancias en las que la compacidad de un objeto juega un papel primordial en algunas cuestiones globales. Pero eso no debe espantarnos... al menos, no por ahora, pues nos dedicaremos a algunos aspectos locales de las superficies.

En las siguientes secciones estudiaremos los aspectos básicos de la geometría diferencial de superficies en $\mathbb{R}_1^3$, empezando por aquellas que no sean nulas (o degeneradas). Por conveniencia, a partir de este momento supondremos que todas las superficies consideradas son conexas.

2.3. Superficies tipo espacio

Una superficie S de este tipo es, digamos, la más parecida a las superficies en $\mathbb{R}^3$ ya conocidas: Por supuesto, S es una superficie desde el punto de vista diferenciable, pero además y por definición, la restricción de la métrica del espacio ambiente a cada plano tangente $T_p S$ es un producto escalar positivo definido. Esto hace que la geometría diferencial de S se pueda estudiar con los mismos métodos conocidos del caso de $\mathbb{R}^3$. Por otro lado, esto también impone condiciones sobre S . Veamos unos ejemplos de estas restricciones, no sin recordar nuestra última observación de la sección anterior: Supondremos que todas las superficies son diferenciables y conexas.

Proposición 2.14. *Sea S una superficie tipo espacio en $\mathbb{R}_1^3$. Entonces S se puede ver localmente como la gráfica de una función diferenciable $f: V \rightarrow \mathbb{R}$ definida en un conjunto abierto $V \subset \mathbb{R}^2$, como en el ejemplo 2.5.*

Demostración. Sea $p \in S$. Por definición, existe una vecindad U de p en S y una parametrización $\varphi: V \rightarrow U$ dada en coordenadas por

$$(u, v) \mapsto (t(u, v), x(u, v), y(u, v)).$$

Como S es tipo espacio, los vectores tangentes (los subíndices indican derivadas parciales)

$$\varphi_u = (t_u, x_u, y_u) \quad \text{y} \quad \varphi_v = (t_v, x_v, y_v)$$

son tipo espacio y por tanto su producto cruz

$$\varphi_u \times \varphi_v = \begin{vmatrix} -i & j & k \\ t_u & x_u & y_u \\ t_v & x_v & y_v \end{vmatrix}$$

es tipo tiempo. Esto implica que la primera coordenada de $\varphi_u \times \varphi_v$ es diferente de cero.

Para mostrar que S se ve localmente como una gráfica, consideramos la proyección $\pi: \mathbb{R}_1^3 \rightarrow \mathbb{R}^2$ dada por $\pi(t, x, y) = (x, y)$ y la composición $\pi \circ \varphi: V \rightarrow \mathbb{R}^2$; es decir, $\pi \circ \varphi(u, v) = (x(u, v), y(u, v))$. Observe que por la observación sobre la primera coordenada de $\varphi_u \times \varphi_v$, podemos aplicar el teorema de la función

inversa a esta transformación, de modo que existe una vecindad (posiblemente más pequeña que V) y una función invertible ψ tal que

$$\psi(x, y) = \psi \circ \pi \circ \varphi(u, v) = (u, v), \quad \text{o} \quad (u, v) = \psi^{-1}(x, y),$$

lo que nos dice que en una vecindad de p , los puntos de S tienen la forma $(t(\psi^{-1}(x, y)), x, y)$; o en otras palabras, localmente la superficie S es la gráfica de la función $t \circ \psi^{-1}$. $\square$

Proposición 2.15. *Sea S una superficie tipo espacio en $\mathbb{R}_1^3$. Entonces S es orientable.*

Demostración. Sea $\{U_\alpha\}$ una cubierta de S , donde para cada α existe una parametrización $\varphi_\alpha: V_\alpha \rightarrow U_\alpha$ y cada U_α es un abierto conexo de S . En U_α podemos elegir como vector normal (unitario) a

$$N_\alpha = \frac{(\varphi_\alpha)_u \times (\varphi_\alpha)_v}{\|(\varphi_\alpha)_u \times (\varphi_\alpha)_v\|} \quad (2.1)$$

o bien a $-N_\alpha$. Elegimos entre ellos a aquel que está en el mismo cono de tiempo de $e_0 = (1, 0, 0)$. Es claro que esto define un campo vectorial normal unitario que varía de manera diferenciable en toda la superficie S , y por tanto S es orientable. $\square$

Definición 2.16. El campo vectorial normal unitario N dado en la proposición anterior define una transformación diferenciable $N: S \rightarrow \mathbb{H}_0^2$ (ver el ejemplo 2.12) llamada la *transformación de Gauss*.

La expresión local de la transformación de Gauss, dada por (2.1) (o por su negativo), muestra que N es una transformación diferenciable. La diferencial de $N: S \rightarrow \mathbb{H}_0^2$ está definida como una transformación lineal de $T_p S$ en $T_{N(p)} \mathbb{H}_0^2$, pero, al igual que en el caso de las superficies en $\mathbb{R}^3$, los planos anteriores son en realidad el mismo, pues ambos tienen el mismo vector normal $N(p)$ y pasan por el origen de $\mathbb{R}_1^3$. Podemos pensar entonces que dN_p va de $T_p S$ en sí mismo; esta transformación nos dice esencialmente cómo varía el vector normal y por tanto cómo se curva la superficie cerca de un punto.

Lema 2.17. *Para cada $p \in S$, $dN_p: T_p S \rightarrow T_p S$ es autoadjunta con respecto del producto escalar en $T_p S$.*

Demostración. Sea $\varphi = \varphi(u, v)$ una parametrización de una vecindad de p . Entonces tenemos para cada (u, v) ,

$$\langle N \circ \varphi, \varphi_u \rangle = 0 \quad \text{y} \quad \langle N \circ \varphi, \varphi_v \rangle = 0;$$

derivando la primera igualdad con respecto de v y la segunda con respecto de u , obtenemos que

$$\langle (N \circ \varphi)_v, \varphi_u \rangle = -\langle N \circ \varphi, \varphi_{vu} \rangle = -\langle N \circ \varphi, \varphi_{uv} \rangle = \langle (N \circ \varphi)_u, \varphi_v \rangle;$$

pero por la regla de la cadena, $(N \circ \varphi)_u(p) = dN_p(\varphi_u)$ y $(N \circ \varphi)_v(p) = dN_p(\varphi_v)$, de modo que lo anterior se escribe como

$$\langle dN_p(\varphi_v), \varphi_u \rangle = \langle dN_p(\varphi_u), \varphi_v \rangle;$$

como $\{\varphi_u, \varphi_v\}$ es base de $T_p S$, lo anterior basta para afirmar que dN_p es autoadjunta. $\square$

Como S es tipo espacio, el producto $\langle \cdot, \cdot \rangle$ de $\mathbb{R}_1^3$ restringido a cada $T_p S$ es positivo definido y tenemos el siguiente resultado.

Corolario 2.18. *Sea S una superficie tipo espacio en $\mathbb{R}_1^3$ con transformación de Gauss $N: S \rightarrow \mathbb{H}_0^2$. Entonces, para cada $p \in S$, dN_p es diagonalizable.*

Por comodidad, estudiaremos a la transformación $-dN_p$ en vez de dN_p , que sigue siendo autoadjunta. Como $T_p S$ es un espacio vectorial con producto escalar positivo definido, $-dN_p$ tiene valores propios reales.

Definición 2.19. Sea S una superficie tipo espacio en $\mathbb{R}_1^3$ y $N: S \rightarrow \mathbb{H}_0^2$ un campo vectorial unitario normal a S en cada punto. Dado $p \in S$, el *operador de forma* asociado a N es la transformación $A_p: T_p S \rightarrow T_p S$ dado por

$$A_p(v) = -dN_p(v).$$

Los valores propios κ_1, κ_2 de A_p se llaman las *curvaturas principales* de S en p .

Con base en lo anterior podemos definir conceptos completamente análogos al caso de superficies en $\mathbb{R}^3$.

Definición 2.20. Sea S una superficie tipo espacio en $\mathbb{R}_1^3$, $N: S \rightarrow \mathbb{H}_0^2$ un campo vectorial unitario normal a S en cada punto, A_p el operador de forma asociado a N en el punto $p \in S$ y κ_1, κ_2 las curvaturas principales de S en p . Entonces decimos que S es una *superficie isoparamétrica* si las funciones κ_1, κ_2 son constantes en S .

Por el momento no analizaremos con detalle este concepto, sino sólo hasta después de introducir la otra clase de superficies en la que estamos interesados.

2.4. Superficies tipo tiempo

En una superficie S de tipo tiempo, el producto escalar inducido por el producto de $\mathbb{R}_1^3$ en cada $T_p S$ ahora tiene índice uno. Ya hemos visto un ejemplo de esto, el espacio de de Sitter $\mathbb{S}_1^2$ del ejemplo 2.11. Ahora estudiaremos las propiedades básicas de estas superficies.

En la proposición 2.14 vimos que una superficie tipo espacio se puede ver localmente como la gráfica de una función de las coordenadas espaciales (x, y) . La siguiente proposición y su demostración son sus análogos para el caso que estamos analizando.

Proposición 2.21. Sea S una superficie tipo tiempo en $\mathbb{R}_1^3$. Entonces S se puede ver localmente como la gráfica de una función diferenciable $x = f(t, y)$ o $y = f(t, x)$, donde $f: V \rightarrow \mathbb{R}$ está definida en un conjunto abierto $V \subset \mathbb{R}^2$.

Localmente una superficie es orientable, pues para cada $p \in S$ podemos elegir una parametrización $\varphi: V \rightarrow U$, con U abierto y conexo en S y definir el campo vectorial normal definido como en (2.1), a saber,

$$N = \frac{\varphi_u \times \varphi_v}{\|\varphi_u \times \varphi_v\|}; \quad (2.2)$$

pero ahora el campo normal N es de tipo espacio, de modo que toma sus valores justamente en el espacio de de Sitter, $N: U \rightarrow \mathbb{S}_1^2$. De nuevo, podemos definir el operador de forma $A_p = -dN_p$, y un vistazo a la demostración del lema 2.17 convencerá al lector que, de nuevo, A_p es una transformación lineal autoadjunta.

Por el teorema 1.51 (p. 18) y dado que $\dim S = 2$, la forma canónica de la matriz de A_p puede ser de alguna de las formas I, II o IV, a saber,

$$\text{I: } \begin{pmatrix} a_0 & 0 \\ 0 & a_1 \end{pmatrix}, \quad \text{IV: } \begin{pmatrix} a_0 & b_0 \\ -b_0 & a_0 \end{pmatrix}, \quad \text{II: } \begin{pmatrix} a_0 & 0 \\ 1 & a_0 \end{pmatrix},$$

donde, recordemos, las primeras dos matrices se calculan con respecto de una base ortonormal y la última con respecto de una base pseudo-ortonormal. Veamos algunos ejemplos de superficies en las que el operador realmente asume una de estas formas.

Ejemplo 2.22. Consideremos de nuevo $\mathbb{S}_1^2$. Recordemos que el llamado *vector de posición* $N(p) = p$ es normal a $\mathbb{S}_1^2$ en cada punto p ; además, es claro que $\langle N, N \rangle = 1$. Calculemos entonces el operador de forma A_p . Con la misma notación anterior,

$$A_p(v) = -dN_p(v) = -dN_p(\alpha'(0)) = -(N \circ \alpha)(0) = -\alpha'(0) = -v,$$

de modo que $A_p = -I$, que por supuesto es diagonalizable. Un razonamiento análogo con el *espacio de de Sitter de radio $r > 0$*

$$\mathbb{S}_1^2(r) = \{ (t, x, y) \in \mathbb{R}_1^3 \mid -t^2 + x^2 + y^2 = r^2 \}$$

muestra que para todo $p \in \mathbb{S}_1^2(r)$ se tiene $A_p = -I/r$.

Ejemplo 2.23 (Cilindro). Sea $r > 0$ y definamos el conjunto

$$C(r) = \{ (t, x, y) \in \mathbb{R}_1^3 \mid -t^2 + x^2 = r^2 \},$$

que es un cilindro sobre la hipérbola $-t^2 + x^2 = r^2$ del plano tx ; ver figura 2.1.

Como el cilindro $C(r)$ es la imagen inversa $F^{-1}(r^2)$ bajo la función $F(t, x, y) = -t^2 + x^2$ y r^2 es un valor regular de F , y el campo vectorial $(t, x, 0)$ es normal a $C(r)$ en cada punto, de tipo espacio. Así,

$$N = \frac{1}{r}(t, x, 0)$$

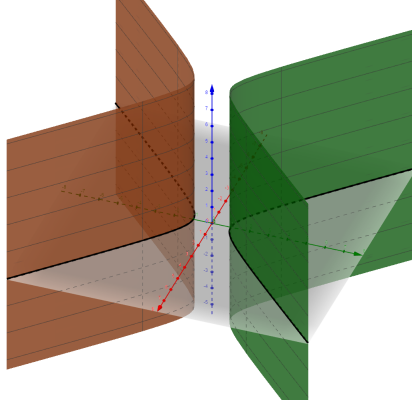


Figura 2.1: Cilindro sobre la hipérbola $-t^2 + x^2 = r^2$.

es un campo unitario normal a $C(r)$. Para calcular la matriz asociada a A_p , debemos elegir una base. Podemos parametrizar las hojas de $C(r)$ por

$$\varphi(u, v) = (r \sinh u, \pm r \cosh u, v), \quad (u, v) \in \mathbb{R}^2,$$

y tomar como base de cada plano tangente a las parciales φ_u, φ_v ; haremos los cálculos eligiendo el signo $+$ en la expresión anterior, pues el otro caso es completamente análogo:

$$\varphi_u = (r \cosh u, r \sinh u, 0), \quad \varphi_v = (0, 0, 1);$$

Entonces

$$A_p(\varphi_u) = -dN_p(\varphi_u) = -(N \circ \varphi)_u = -(\cosh u, \sinh u, 0) = -\frac{1}{r}\varphi_u,$$

y

$$A_p(\varphi_v) = -dN_p(\varphi_v) = -(N \circ \varphi)_v = -(0, 0, 0) = 0 \cdot \varphi_v,$$

lo que nos dice que φ_u, φ_v son vectores propios de A_p y por tanto su matriz con respecto de esta base es del tipo I:

$$\begin{pmatrix} -\frac{1}{r} & 0 \\ 0 & 0 \end{pmatrix}.$$

En los dos ejemplos anteriores el operador de forma fue diagonalizable; es decir, del tipo I. Ahora veremos un ejemplo del tipo II.

Ejemplo 2.24 (Cilindro generalizado). En [7], M. Magid desarrolla el caso de los llamados *cilindros generalizados de tipo I* (ver p. 175) en espacios vectoriales lorentzianos de dimensión mayor o igual a 3. Para el caso de dimensión 3, podemos simplificar su análisis como sigue: Se considera una curva nula $\alpha: I \rightarrow \mathbb{R}_1^3$ y

un marco pseudo-ortonormal $\{A(s), B(s), C(s)\}$, es decir, una terna de campos vectoriales definidos para $s \in I$, que cumplen las siguientes restricciones:

$$\begin{aligned} \langle A, A \rangle = \langle B, B \rangle = \langle A, C \rangle = \langle B, C \rangle = 0, \quad \langle A, B \rangle = \langle C, C \rangle = 1, \\ A(s) = \alpha'(s), \quad C'(s) = -g(s)B(s), \quad g \neq 0. \end{aligned} \quad (2.3)$$

Con estas condiciones, se define el *cilindro generalizado*² como la superficie parametrizada por

$$f(s, u) = \alpha(s) + uB(s), \quad s \in I, u \in \mathbb{R}.$$

Como ejemplo concreto, sean

$$\alpha(s) = \left(\frac{s^3}{6} + \frac{s}{2}, \frac{s^2}{2}, \frac{s^3}{6} - \frac{s}{2} \right), \quad B(s) = (-1, 0, -1), \quad C(s) = (s, 1, s). \quad (2.4)$$

Dejaremos como ejercicio verificar que se satisfacen las condiciones (2.3), con $g(s) \equiv -1$. Al considerar la superficie parametrizada $f(s, u)$ definida arriba, tenemos que $C = f_s \times f_u$ (donde como de costumbre, los subíndices indican derivadas parciales), de modo que un campo vectorial unitario normal a la superficie es precisamente $N = C$. Para obtener la expresión matricial de $A_p = -dN_p$ en términos de la base pseudo-ortonormal $\{f_s = \alpha' = A, f_u = B\}$, tenemos

$$dN_p(f_s) = C'(s) = B(s), \quad \text{y} \quad dN_p(f_u) = 0,$$

donde la última igualdad se debe a que el vector B es constante. Así, la matriz de A_p es

$$\begin{pmatrix} 0 & 0 \\ -1 & 0 \end{pmatrix},$$

que no es diagonalizable.

Los tres ejemplos anteriores de superficies tipo tiempo no son casuales. En los tres casos, las curvaturas principales (los valores propios de $A_p = -dN_p$) son reales y constantes: En el ejemplo 2.22, ambas son iguales a $-1/r$; en el ejemplo 2.23, una es $-1/r$ y la otra 0; mientras que en el ejemplo 2.24, ambas son iguales a cero. Recordemos que en el caso de superficies tipo espacio (definición 2.20), dijimos que una superficie es isoparamétrica si sus curvaturas principales son constantes. Por razones que no expondremos con detalle aquí, debemos modificar nuestra definición para el caso tipo tiempo.

Definición 2.25. Sea S una superficie tipo tiempo en $\mathbb{R}_1^3$, $N: S \rightarrow S_1^2$ un campo vectorial unitario normal a S en cada punto, A_p el operador de forma asociado a N en el punto $p \in S$.

²Este es el nombre utilizado por M. Magid en [7], pero actualmente se usa más el término *B-scroll* en inglés. Por desgracia, este vocablo alguna vez elegante que se refiere a un largo papel enrollado usado antiguamente para leyes, tratados y otras importantes cuestiones, se traduce actualmente en la sencilla palabra *rollo*.

- Si A_p es diagonalizable, decimos que S es una *superficie isoparamétrica* si los valores propios de A_p (sus *curvaturas principales*) son constantes en S .
- Si A_p no es diagonalizable, decimos que S es una *superficie isoparamétrica* si el polinomio mínimo de A_p es constante en S .

El lector puede verificar que en el ejemplo 2.24 el polinomio mínimo de A_p es igual a x^2 para todo $p \in S$, de modo que la superficie en dicho ejemplo es isoparamétrica. En cualquier superficie de este tipo, las curvaturas principales son constantes en S .

Antes de finalizar esta introducción a las superficies en $\mathbb{R}_1^3$, el lector podrá observar que no hemos dado ejemplos de superficies tipo tiempo cuyo operador de forma con respecto de una base ortonormal sea del tipo IV, es decir,

$$\begin{pmatrix} a_0 & b_0 \\ -b_0 & a_0 \end{pmatrix},$$

donde $b_0 \neq 0$. En este caso, los valores propios son complejos. Tenemos el siguiente resultado:

Teorema 2.26. *No existen superficies isoparamétricas tipo tiempo en $\mathbb{R}_1^3$ cuyos operadores de forma A_p sean del tipo IV.*

2.5. La segunda forma fundamental

Ahora estudiaremos algunos conceptos importantes en el caso de las superficies, reuniendo los casos de tipo espacio y tipo tiempo. Sea S una superficie no degenerada en $\mathbb{R}_1^3$, $p \in S$, N un campo vectorial unitario normal a S , definido al menos en una vecindad U de p , y ϵ el *signo* de N , es decir,

$$\epsilon = \langle N, N \rangle = \begin{cases} -1, & \text{si } S \text{ es tipo espacio,} \\ 1, & \text{si } S \text{ es tipo tiempo.} \end{cases}$$

Finalmente, sea A_p el operador de forma asociado a N en el punto $p \in S$.

Definición 2.27. Con las notaciones anteriores, definimos:

- La *segunda forma fundamental* de S en p , $h_p: T_p S \times T_p S \rightarrow \mathbb{R}$ como

$$h_p(u, v) = \langle A_p u, v \rangle.$$

- La *curvatura media* H de S en p como

$$H(p) = \frac{\epsilon}{2} \text{Traza}(A_p).$$

- La *curvatura gaussiana* K de S en p como

$$K(p) = \epsilon \det A_p.$$

Observación 2.28. De ahora en adelante y por brevedad, omitiremos la evaluación en el punto p .

Como A es un operador lineal autoadjunto, tenemos lo siguiente.

Proposición 2.29. *La segunda forma fundamental de S es una forma bilineal simétrica.*

Observación 2.30. El término *segunda forma fundamental* también se refiere a la siguiente transformación vectorial:

$$II(u, v) = \epsilon \langle Au, v \rangle N;$$

en otras palabras, $h(u, v)$ es la coordenada de $II(u, v)$ con respecto del vector normal N . En la práctica, es indistinto trabajar con una u otra definición, aunque en el contexto lorentziano hay que recordar la existencia del signo ϵ .

Definición 2.31. Con las notaciones anteriores, decimos que un punto $p \in S$ es *umbílico* si existe un número $\kappa = \kappa(p) \in \mathbb{R}$ tal que

$$\langle II(u, v), N(p) \rangle = \kappa(p) \langle u, v \rangle$$

para todo $u, v \in T_p S$. Una superficie S es *totalmente umbílica* si y sólo si todo punto de S es umbílico.

La siguiente proposición es inmediata.

Proposición 2.32. *Si A es diagonalizable y κ_1, κ_2 son las curvaturas principales de S en p , entonces*

- *La curvatura media de S en p es*

$$H = \frac{\epsilon}{2} (\kappa_1 + \kappa_2).$$

- *La curvatura gaussiana de S en p es*

$$K = \epsilon \kappa_1 \kappa_2.$$

- *Un punto $p \in S$ es umbílico si y sólo si $A = \kappa I$ en p .*

El siguiente resultado es otra caracterización de los puntos umbílicos.

Proposición 2.33. *Sea S una superficie no degenerada en $\mathbb{R}_1^3$ y $p \in S$ tal que su operador de forma A sea diagonalizable. Entonces*

$$H^2(p) - \epsilon K(p) \geq 0,$$

y la igualdad vale si y sólo si p es umbílico.

Ejemplo 2.34. Ya vimos que $\mathbb{S}_1^2$ es una superficie totalmente umbílica (ejemplo 2.22). De manera análoga, se puede ver que el espacio hiperbólico $\mathbb{H}_0^2$ definido en el ejemplo 2.12 también es totalmente umbílico. Los planos no degenerados son un ejemplo sencillo de este tipo de superficies.

La siguiente proposición muestra que, esencialmente, los anteriores son los únicos ejemplos posibles. Recordemos que nuestras superficies son siempre conexas.

Proposición 2.35. *Sea S una superficie totalmente umbílica en $\mathbb{R}_1^3$. Entonces, salvo una transformación rígida, S es un conjunto abierto en un plano, en $\mathbb{S}_1^2(r)$ o en el plano hiperbólico de radio r :*

$$\mathbb{H}_0^2(r) = \{ (t, x, y) \in \mathbb{R}_1^3 \mid -t^2 + x^2 + y^2 = -r^2, t > 0 \};$$

Demostración. Sabemos que para cada punto en S , $AX = \kappa X$. Primero mostraremos que la función κ es constante en S . Para cada p , sea $\varphi = \varphi(u, v)$ una parametrización de una vecindad U de este punto. Al evaluar A en los elementos de la base $\{\varphi_u, \varphi_v\}$, tenemos

$$\kappa\varphi_u = A\varphi_u = -dN(\varphi_u) = -(N \circ \varphi)_u; \quad \text{y, análogamente,} \quad \kappa\varphi_v = -(N \circ \varphi)_v.$$

Derivando la primera igualdad con respecto de v y la segunda con respecto de u , para obtener

$$\kappa_v\varphi_u + \kappa\varphi_{vu} = -(N \circ \varphi)_{vu} \quad \text{y} \quad \kappa_u\varphi_v + \kappa\varphi_{uv} = -(N \circ \varphi)_{uv};$$

Simplificando tenemos

$$\kappa_u\varphi_v - \kappa_v\varphi_u = 0;$$

como $\{\varphi_u, \varphi_v\}$ es una base, $\kappa_u = \kappa_v = 0$, de donde κ es constante en una vecindad de p . Como S es conexa, κ es constante en S .

Ahora tenemos dos casos: $\kappa = 0$ y $\kappa \neq 0$. En el primer caso, es fácil mostrar que entonces el vector normal N es constante y que S está contenida en un plano ortogonal a N .

Si $\kappa \neq 0$, entonces se tiene que el punto

$$P(u, v) = \varphi(u, v) + \frac{1}{\kappa}(N \circ \varphi)(u, v)$$

es un punto fijo P_0 , pues $P_u = \varphi_u - \frac{1}{\kappa}(N \circ \varphi)_u = 0$ y análogamente $P_v = 0$. Así, de la ecuación anterior tenemos

$$\varphi(u, v) - P_0 = -\frac{1}{\kappa}N;$$

Sin pérdida de generalidad, podemos hacer una transformación rígida de modo que $P_0 = 0$ y entonces

$$-t^2 + x^2 + y^2 = \langle \varphi(u, v), \varphi(u, v) \rangle = \frac{\epsilon}{\kappa^2},$$

donde como de costumbre $\epsilon = \pm 1 = \langle N, N \rangle$. Entonces, dependiendo del carácter causal de N , tenemos que los puntos cubiertos por la parametrización están en $\mathbb{S}_1^2(r)$, si $\epsilon = +1$; o en $\mathbb{H}_0^2(r)$, si $\epsilon = -1$; en ambos casos $r = 1/|\kappa|$. De nuevo, un argumento de conexidad muestra que S queda contenida en alguna de estas variedades. $\square$

Para cerrar este capítulo, mencionaremos algunos resultados básicos acerca de las superficies con curvatura media constante o con curvatura gaussiana constante. En realidad, estas superficies tienen una historia muy rica, pues han sido estudiadas desde hace siglos y aún ahora se les encuentran muchas propiedades y aplicaciones interesantes.

Teorema 2.36. *Sea S una superficie no degenerada en $\mathbb{R}_1^3$ con curvaturas media y gaussiana constantes. Entonces S es totalmente umbílica, o un cilindro circular recto, o una superficie reglada.*

Definición 2.37. Una superficie S es de *traslación* si es la gráfica de una función de la forma $\phi(u) + \psi(v)$.

Teorema 2.38. *Una superficie tipo espacio, de traslación y con curvatura media constante es un plano, una superficie de Scherk o un cilindro.*

2.6. Ejercicios

1. Sea $\alpha: \mathbb{R} \rightarrow \mathbb{R}_1^3$ una curva regular cerrada, contenida en un plano P . Si α es tipo espacio, muestre que P debe ser también de tipo espacio. [SUGERENCIA: Use un argumento análogo al de la demostración de la Proposición 2.3.]
2. Verifique que la gráfica de una función, definida en el ejemplo 2.5 es en realidad una superficie diferenciable.
3. Demuestre la proposición 2.7.
4. Pruebe la proposición 2.13.
5. Muestre la proposición 2.21.
6. Demuestre que una superficie S es tipo tiempo si y sólo si S admite un campo vectorial tangente tipo tiempo.
7. Complete los detalles del ejemplo 2.23.
8. ¿Puede haber una superficie tipo tiempo no orientable en $\mathbb{R}_1^3$? SUGERENCIA: Puesto que la banda de Möbius es no orientable, ¿es posible colocarla en $\mathbb{R}_1^3$ de una manera suficientemente *vertical*?
9. Compruebe que los campos $A(s) = \alpha'(s)$, $B(s)$ y $C(s)$ de la ecuación (2.24) satisfacen las condiciones (2.3).

10. Pruebe que el hiperboloide del ejemplo 2.12 tiene curvatura gaussiana negativa.
11. La *superficie de Scherk* es la gráfica de la función

$$t = f(x, y) = \log \frac{\cosh y}{\cosh x} = \log \cosh y - \log \cosh x.$$

Muestre que la curvatura media de esta superficie es idénticamente igual a cero.

12. Sea $\alpha: I \rightarrow \mathbb{R}_1^3$ una curva nula y v_0 un vector constante tipo tiempo, con $\langle \alpha'(s), v_0 \rangle \neq 0$ para todo $s \in I$. Defina la superficie reglada

$$\varphi(s, t) = \alpha(s) + tv_0.$$

Pruebe que las curvaturas media y gaussiana de esta superficie son idénticamente nulas.

Capítulo 3

Variedades lorentzianas

3.1. Introducción

Ahora generalizaremos lo dicho en el capítulo anterior, al caso de las variedades abstractas, de cualquier dimensión. Usaremos la notación y algunos resultados de la referencia [10]. Para comenzar, recordemos algunas definiciones.

Definición 3.1. Una *variedad lorentziana* es una variedad diferenciable $\bar{M}^{n+1}$, de dimensión $n+1 \geq 2$, dotada de una métrica semi-riemanniana $\langle \cdot, \cdot \rangle$ de índice 1.

Ejemplo 3.2. El ejemplo más sencillo de variedad lorentziana es justamente el espacio de Lorentz-Minkowski $\mathbb{R}_1^{n+1}$. En analogía con las superficies definidas en el capítulo anterior, definimos el *espacio de de Sitter de dimensión $n+1$* como

$$\mathbb{S}_1^{n+1} = \{ p \in \mathbb{R}_1^{n+2} \mid \langle p, p \rangle = 1 \},$$

y el *espacio hiperbólico de dimensión $n+1$* es el conjunto

$$\mathbb{H}_1^{n+1} = \{ p \in \mathbb{R}_2^{n+2} \mid \langle p, p \rangle = -1, u_0 > 0 \};$$

observe el cambio de índice del espacio ambiente en esta última definición.

Sea M una variedad riemanniana con métrica g . Si damos a $\mathbb{R}$ la métrica negativo definida (lo que, recordemos, denotamos por $-\mathbb{R}$), entonces es natural darle al producto $\mathbb{R} \times M$ una métrica lorentziana producto. Éste es un ejemplo particular de la siguiente construcción.

Definición 3.3 (Espacios de Robertson-Walker). Sea F una variedad riemanniana con métrica g_F , $I \subset \mathbb{R}$ un intervalo y $\varrho: I \rightarrow \mathbb{R}^+$ una función diferenciable positiva. El *espacio generalizado de Robertson-Walker* $\bar{M} = -I \times_{\varrho} F$ es la variedad producto dotada de la métrica

$$g_{\bar{M}} = -\pi_1^* g_{\mathbb{R}} + (\varrho \circ \pi_1)^2 \pi_2^* g_F,$$

donde $\pi_1: I \times F \rightarrow I$, $\pi_2: I \times F \rightarrow M$ son las proyecciones naturales. En otras palabras, dado un punto $(t, p) \in \bar{M}$ y un vector tangente $(a, v) \in T_t\mathbb{R} \times T_pF = T_{(t,p)}(\bar{M})$, entonces

$$\langle (a, v), (a, v) \rangle = -a^2 + \varrho^2(t) \langle v, v \rangle$$

Se dice que $\bar{M} = -I \times_{\varrho} F$ es un *espacio clásico de Robertson-Walker* si la variedad F tiene curvatura seccional constante.

Como en el caso de las superficies, ahora estudiaremos un tipo particular de subvariedades.

Definición 3.4. Sea $\bar{M}$ una variedad lorentziana de dimensión $n + 1$. Una *hipersuperficie* es una subvariedad diferenciable $M \subset \bar{M}$ de dimensión n . Decimos que M es una *hipersuperficie no degenerada* si y sólo si para cada $p \in M$, la métrica g de $\bar{M}$ restringida a T_pM es no degenerada.

Recordemos que estamos suponiendo que todos los conjuntos que consideramos son conexos. Así, una hipersuperficie no degenerada es tipo espacio o tipo tiempo. En ambos casos, para cada $p \in M$ tenemos la siguiente suma directa:

$$T_p\bar{M} = T_pM + (T_pM)^{\perp} \quad (3.1)$$

El espacio $(T_pM)^{\perp}$ se llama el *espacio normal a M en p* y en ocasiones se denota como N_pM . La descomposición (3.1) será fundamental para estudiar la geometría de M . La idea general es que separaremos los conceptos geométricos en intrínsecos a M y los extrínsecos: Los primeros sólo dependerán de M y no de su colocación dentro de $\bar{M}$.

El teorema fundamental de la geometría diferencial nos dice que, dada una métrica (no degenerada) en una variedad, sólo existe una conexión afín, llamada la conexión de Levi-Civita de la variedad. En nuestro caso, tenemos la métrica del espacio ambiente $\bar{M}$ y su conexión de Levi-Civita, que denotamos por $\bar{D}$. Por otro lado, dada una hipersuperficie M no degenerada, tenemos la restricción de la métrica y su correspondiente conexión de Levi-Civita. Uno de los teoremas fundamentales de la geometría de subvariedades nos dice que

$$D_X Y = (\bar{D}_X Y)^T;$$

es decir, la conexión de Levi-Civita en M no es más que la proyección de la conexión de Levi-Civita de $\bar{M}$ a la subvariedad, usando la descomposición (3.1). Por otro lado, la parte normal $(\bar{D}_X Y)^{\perp}$ también es muy importante.

Definición 3.5. La *segunda forma fundamental* de M en $\bar{M}$ es la forma bilineal

$$II(X, Y) = (\bar{D}_X Y)^{\perp} = \bar{D}_X Y - D_X Y.$$

Como en el caso de las superficies, esta segunda forma fundamental toma valores vectoriales; más precisamente, para cada $p \in M$, $II_p(X, Y) \in N_pM$. Para definir la segunda forma fundamental con valores reales, hacemos el procedimiento siguiente.

Definición 3.6. Sea $\bar{M}$ una variedad lorentziana, M una hipersuperficie no degenerada y $p \in M$. En una vecindad U de p existe definido un campo vectorial normal unitario ξ , con $\langle \xi, \xi \rangle = \epsilon = \pm 1$. Entonces la *segunda forma fundamental* de M en $\bar{M}$ es la función real dada por

$$II(X, Y) = \epsilon \cdot h(X, Y)\xi. \quad (3.2)$$

Proposición 3.7. La segunda forma fundamental (en cualquiera de sus definiciones) es simétrica.

Demostración. Basta mostrar que II es simétrica:

$$\begin{aligned} II(X, Y) - II(Y, X) &= (\bar{D}_X Y - D_X Y) - (\bar{D}_Y X - D_Y X) \\ &= (\bar{D}_X Y - \bar{D}_Y X) - (D_X Y - D_Y X) \\ &= [X, Y] - [X, Y] = 0. \quad \square \end{aligned}$$

Por el álgebra lineal sabemos que existe una transformación lineal autoadjunta A tal que $h(X, Y) = \langle AX, Y \rangle$. Por conveniencia, daremos la forma explícita de este operador.

Proposición 3.8. Sea $\bar{M}$ una variedad lorentziana, M una hipersuperficie no degenerada, $p \in M$ y ξ un campo vectorial unitario normal a M . Definimos el operador de forma $A : T_p M \rightarrow T_p M$ como

$$AX = -\bar{D}_X \xi(p).$$

Entonces

$$h(X, Y) = \langle AX, Y \rangle.$$

Demostración. De la ecuación (3.2), tomando el producto escalar con ξ , tenemos

$$h(X, Y) = \langle II(X, Y), \xi \rangle = \langle \bar{D}_X Y, \xi \rangle = -\langle \bar{D}_X \xi, Y \rangle = \langle AX, Y \rangle. \quad \square$$

Ahora veremos la relación entre los tensores de curvatura de $\bar{M}$ y de M . Recordemos que dada una conexión $\bar{D}$, definimos el *tensor de curvatura* R como

$$\bar{R}_{XY}Z = -\bar{D}_X \bar{D}_Y Z + \bar{D}_Y \bar{D}_X Z + \bar{D}_{[X, Y]}Z. \quad (3.3)$$

Teorema 3.9 (Ecuación de Gauss). Sean $\bar{M}$ una variedad lorentziana y M una hipersuperficie no degenerada con $\bar{R}, R$ sus tensores de curvatura, respectivamente. Sea II la segunda forma fundamental de M en $\bar{M}$. Entonces

$$\langle \bar{R}_{VW}X, Y \rangle = \langle R_{VW}X, Y \rangle + \langle II(V, Y), II(W, X) \rangle - \langle II(V, X), II(W, Y) \rangle.$$

Corolario 3.10. Si X, Y son una base para un plano no degenerado, entonces las curvaturas seccionales $\bar{K}$ y K de $\bar{M}$ y M satisfacen la también llamada ecuación de Gauss:

$$\bar{K}(X, Y) = K(X, Y) - \frac{\langle II(X, X), II(Y, Y) \rangle - \langle II(X, Y), II(X, Y) \rangle}{\langle X, X \rangle \langle Y, Y \rangle - \langle X, Y \rangle^2} \quad (3.4)$$

Como en el caso de las superficies, tenemos las definiciones de la curvatura media, gaussiana, las hipersuperficies isoparamétricas, totalmente geodésicas, totalmente umbílicas. Por completez, las incluimos.

Definición 3.11. Sea $\bar{M}$ una variedad lorentziana, M una hipersuperficie no degenerada, $p \in M$, ξ un campo vectorial unitario normal a M , $\langle \xi, \xi \rangle = \epsilon$, h la segunda forma fundamental de M en $\bar{M}$ y A el operador de forma asociado a ξ .

- La *curvatura media* H de M en p como

$$H(p) = \frac{\epsilon}{2} \text{Traza}(A).$$

- La *curvatura de Gauss-Kronecker* K de S en p como

$$K(p) = \epsilon \det A.$$

- p es *umbílico* si existe un número $\kappa = \kappa(p) \in \mathbb{R}$ tal que

$$\langle II(X, Y), \xi \rangle = \kappa(p) \langle X, Y \rangle$$

en p . M es totalmente umbílica si y sólo si todo punto de M es umbílico. En particular, si $\kappa \equiv 0$, decimos que M es *totalmente geodésica*.

- Si A es diagonalizable para todo punto en M , decimos que M es una *hipersuperficie isoparamétrica* si los valores propios de A son constantes en M .
- Si A no es diagonalizable, decimos que M es una *hipersuperficie isoparamétrica* si el polinomio mínimo de A es constante en M .

Ejemplo 3.12. Ya hemos visto que $\mathbb{S}_1^2$ y $\mathbb{H}_0^2$ son superficies totalmente umbílicas de $\mathbb{R}_1^3$. Exactamente la misma técnica se utiliza para ver que $\mathbb{S}_1^{n+1}$ y $\mathbb{H}_1^{n+1}$ son totalmente umbílicas en sus espacios ambiente correspondientes: En cada caso elegimos el campo normal ξ como el vector de posición, $\xi(p) = p$, entonces

$$AX = -\bar{D}_X \xi = -X,$$

de modo que $h(X, Y) = -\langle X, Y \rangle$ y $II(X, Y) = -\epsilon \langle X, Y \rangle \xi$. Esto muestra que las hipersuperficies ya mencionadas son totalmente umbílicas.

Al sustituir en la ecuación de Gauss (3.4),

$$\begin{aligned} \bar{K}(X, Y) &= K(X, Y) - \frac{\langle II(X, X), II(Y, Y) \rangle - \langle II(X, Y), II(X, Y) \rangle}{\langle X, X \rangle \langle Y, Y \rangle - \langle X, Y \rangle^2} \\ &= K(X, Y) - \epsilon. \end{aligned}$$

Como $\mathbb{S}_1^{n+1}$ es tipo tiempo, $\epsilon = 1$ y $K(X, Y) = 1$. Así, el espacio de de Sitter es un espacio con curvatura seccional constante positiva. Por otro lado, $\mathbb{H}_1^{n+1}$ es tipo espacio, de modo que $\epsilon = -1$ y $K(X, Y) = -1$. Así, el espacio hiperbólico tiene curvatura seccional constante negativa.

Definición 3.13 (Formas espaciales). Una *forma espacial lorentziana* o *espacio modelo lorentziano* es una variedad lorentziana con curvatura seccional constante.

Las formas espaciales lorentzianas *simplemente conexas* son precisamente $\mathbb{R}_1^{n+1}$, $\mathbb{H}_1^{n+1}$ y $\mathbb{S}_1^{n+1}$; en este último caso pedimos que $n \geq 2$.

En las siguientes secciones estudiaremos algunas hipersuperficies, mostrando a la par algunas técnicas que se utilizan con dicho fin.

3.2. Hipersuperficies umbílicas en formas espaciales

La ecuación de Gauss del teorema 3.9 analiza la relación entre el tensor de curvatura de $\bar{M}$ y M y nos dice que la parte tangente de $\bar{R}_{VW}X$ está directamente relacionada con $R_{VW}X$. La ecuación de Codazzi se encarga de la parte normal.

Definición 3.14. Sean $\bar{M}$ una variedad lorentziana con conexión de Levi-Civita $\bar{D}$, M una hipersuperficie no degenerada, X un campo tangente a M y Z un campo normal a M .

- Definimos la *conexión normal* de M en $\bar{M}$ como

$$D_X^\perp Z = (\bar{D}_X Z)^\perp.$$

- Sea II la segunda forma fundamental de M en $\bar{M}$. Si $X, Y, V \in \mathfrak{X}(M)$, entonces

$$(D_V II)(X, Y) = D_V^\perp(II(X, Y)) - II(D_V X, Y) - II(X, D_V Y).$$

Teorema 3.15 (Ecuación de Codazzi). Sean $\bar{M}$ una variedad lorentziana con conexión de Levi-Civita $\bar{D}$, M una hipersuperficie no degenerada, y $V, W, X \in \mathfrak{X}(M)$. Entonces

$$(\bar{R}_{VW}X)^\perp = -(D_V II)(W, X) + (D_W II)(V, X).$$

Corolario 3.16. Si $\bar{M}$ tiene curvatura seccional constante, entonces

$$(D_V II)(W, X) = (D_W II)(V, X).$$

Si A es el operador de forma de M en $\bar{M}$, entonces

$$(D_V A)(W) = (D_W A)(V), \tag{3.5}$$

donde, por definición,

$$(D_V A)(W) = D_V(AW) - A(D_V W).$$

Usaremos la ecuación de Codazzi para caracterizar a las hipersuperficies totalmente umbílicas de $\mathbb{R}_1^{n+1}$. Compare con la proposición 2.35.

Proposición 3.17. *Sean $n \geq 0$, M una hipersuperficie conexa, no degenerada, totalmente umbílica en $\mathbb{R}_1^{n+2}$, pero no totalmente geodésica. Entonces, salvo una transformación rígida, M es un conjunto abierto en un espacio de de Sitter $\mathbb{S}_1^{n+1}(r)$ o en un espacio hiperbólico $\mathbb{H}_0^{n+1}(r)$ de radio r .*

Demostración. La demostración es muy similar a la de la proposición 2.35. Como M es totalmente umbílica, $AX = \kappa X$. Para mostrar que κ es constante en M , usamos la ecuación de Codazzi en su forma (3.5):

$$(D_V A)(W) = (D_W A)(V),$$

de donde

$$D_V(AW) - A(D_V W) = D_W(AV) - A(D_W V),$$

y como $AW = \kappa W$, $AV = \kappa V$,

$$V(\kappa)W + \kappa D_V W - \kappa D_V W = W(\kappa)V + \kappa D_W V - \kappa D_W V.$$

o $V(\kappa)W - W(\kappa)V = 0$. Aquí es donde usamos la condición sobre la dimensión: Para cada V podemos elegir W linealmente independiente de V , de modo que $V(\kappa) = 0$ para todo V y κ es constante en M .

Como M no es totalmente geodésica, $\kappa \neq 0$. Entonces si P denota el vector de posición de cada punto en $\mathbb{R}_1^{n+2}$, tenemos que

$$P + \frac{1}{\kappa}\xi$$

es constante en M , pues

$$\bar{D}_X P + \frac{1}{\kappa}\bar{D}_X \xi = X - \frac{1}{\kappa}AX = 0.$$

Si llamamos P_0 a este punto, podemos suponer que $P_0 = 0$ y entonces

$$\langle P, P \rangle = \frac{1}{\kappa^2} \langle \xi, \xi \rangle = \frac{\epsilon}{\kappa^2},$$

lo que nos dice que P está en $\mathbb{S}_1^{n+1}(1/|\kappa|)$ o en $\mathbb{H}_0^{n+1}(1/|\kappa|)$, dependiendo del signo ϵ . $\square$

Ahora veremos cuáles son las hipersuperficies umbílicas de nuestros otros modelos de variedades lorentzianas, $\mathbb{S}_1^{n+1}$ y $\mathbb{H}_1^{n+1}$. Enunciaremos sólo el primer caso, pues la demostración para el otro caso es completamente análoga.

Teorema 3.18. *Sea M una hipersuperficie no degenerada, totalmente umbílica de $\mathbb{S}_1^{n+1}$. Entonces M es la intersección de $\mathbb{S}_1^{n+1}$ con un hiperplano de dimensión $(n+1)$ en $\mathbb{R}_1^{n+2}$.*

3.3. Hipersuperficies con curvatura media constante

La historia de estas hipersuperficies abarca ya varios siglos, pues en el siglo XVIII Lagrange planteó el siguiente problema: Dada una curva cerrada C en $\mathbb{R}^3$, ¿existe una superficie S que tenga área mínima entre todas las superficies cuya frontera sea C ?

Pronto se descubrió una relación de este problema con la curvatura: Meusnier mostró que el hecho de tener área mínima implica que la curvatura media de S debe ser igual a cero. Lagrange mismo llegó a plantear la ecuación diferencial que debe satisfacer S para tener curvatura media cero, si está dada como la gráfica de una función $z = u(x, y)$:

$$(1 + u_x^2)u_{yy} - 2u_x u_y u_{xy} + (1 + u_y^2)u_{xx} = 0,$$

la cual se puede escribir en forma de divergencia:

$$\operatorname{div} \left(\frac{\operatorname{grad} u}{\sqrt{1 + \|\operatorname{grad} u\|^2}} \right) = 0. \quad (3.6)$$

Así, comenzaron varios esfuerzos por encontrar soluciones a esta ecuación. Durante muchos años, fueron conocidas escasas soluciones: el plano, la catenoide, la helicoides y la superficie de Scherk (ver ejercicio 11 del capítulo 2).

Podemos mostrar lo que fue un gran avance en la teoría de las superficies mínimas en $\mathbb{R}^3$ de la manera siguiente: En la ecuación (3.6), supongamos por un momento que u es una función *eikonal*; es decir, que $\|\operatorname{grad} u\|^2$ es constante. Entonces la ecuación se reduce a $\Delta u = 0$, donde Δ es el operador laplaciano. Esto quiere decir que u es una función *armónica*. Recordando que en un dominio simplemente conexo, una función armónica siempre es la parte real de una función analítica (desde el punto de vista complejo), esto nos dice que:

- Podemos encontrar muchos ejemplos de superficies mínimas usando las funciones analíticas;
- Podemos usar la teoría de funciones analíticas para obtener resultados acerca de superficies mínimas.

Por ejemplo, tenemos el siguiente resultado clásico.

Teorema 3.19 (Bernstein). *Si S es una superficie mínima dada como la gráfica de una función u definida en todo el plano $\mathbb{R}^2$, entonces S es un plano.*

La idea de la demostración consiste en considerar a $\mathbb{R}^2$ como $\mathbb{C}$ y a la aplicación de Gauss, que asocia a cada punto de S su vector normal, como una aplicación de $\mathbb{C}$ en $\mathbb{C}$ via la proyección estereográfica. Se muestra que esta aplicación de Gauss es analítica, y que puede considerarse como una aplicación acotada, de modo que por el teorema de Liouville es constante. Así, el vector normal es constante y S es un plano.

Este teorema es un ejemplo del tipo de resultados que buscamos: Bajo una hipótesis geométrica, podemos caracterizar a algunas superficies distinguidas en nuestro espacio ambiente. Debido a limitaciones de tiempo, dedicaremos la última parte de esta sesión a enunciar algunos resultados de tipo Bernstein en espacios lorentzianos. En el espacio de Lorentz-Minkowski tenemos lo siguiente.

Teorema 3.20 (Calabi, Cheng-Yau). *Sea M una hipersuperficie tipo espacio en $\mathbb{R}_1^{n+1}$, con curvatura media cero, dada como la gráfica de una función $u: \mathbb{R}^n \rightarrow \mathbb{R}$ definida en todo $\mathbb{R}^n$, entonces M es un hiperplano.*

Para el caso del espacio de de Sitter $\mathbb{S}_1^{n+1}$, comenzaremos nuestra historia en 1977. En ese año, Goddard estudió las hipersuperficies completas, tipo espacio y con curvatura media constante en $\mathbb{S}_1^{n+1}$ y conjeturó que una especie de teorema de Bernstein debía ser cierto: Todas estas hipersuperficies deberían ser totalmente umbílicas.

Muy rápidamente, se vio que la conjetura era falsa: Dajczer y Nomizu [2] mostraron en 1981 un ejemplo de una superficie en $\mathbb{S}_1^3$ con curvatura media constante que no es totalmente umbílica. Entonces, la pregunta que surgió naturalmente es,

¿Bajo qué condiciones una hipersuperficie con curvatura media constante en $\mathbb{S}_1^{n+1}$ es totalmente umbílica?

Las primeras respuestas surgieron también rápidamente.

Teorema 3.21 (Akutagawa, Montiel). *Sea M una hipersuperficie compacta, tipo espacio, con curvatura media constante en $\mathbb{S}_1^{n+1}$. Entonces M es totalmente umbílica.*

En particular, el trabajo de Montiel (1988) mostró que este tipo de resultados se podría extender a otros espacios. En la definición 3.3 mencionamos a los espacios de Robertson-Walker, cuya definición recordamos rápidamente.

Definición 3.22 (Espacios de Robertson-Walker). *Sea F una variedad riemanniana con métrica g_F , $I \subset \mathbb{R}$ un intervalo y $\varrho: I \rightarrow \mathbb{R}^+$ una función diferenciable positiva. El espacio generalizado de Robertson-Walker $\bar{M} = -I \times_{\varrho} F$ es la variedad producto dotada de la métrica*

$$g_{\bar{M}} = -\pi_1^* g_{\mathbb{R}} + (\varrho \circ \pi_1)^2 \pi_2^* g_F,$$

donde $\pi_1: I \times F \rightarrow I$, $\pi_2: I \times F \rightarrow F$ son las proyecciones naturales.

Estos espacios han adquirido relevancia por ser modelos adecuados en relatividad general. Nosotros no entraremos en la discusión de estos espacios desde el punto de vista físico, sino que daremos un esbozo de su estudio desde el punto de vista geométrico. Empecemos con algunos ejemplos.

Ejemplo 3.23. Las formas espaciales lorentzianas se pueden ver como productos alabeados:

- El espacio de Lorentz-Minkowski $\mathbb{R}_1^{n+1}$ se puede escribir como:

$$\mathbb{R}_1^{n+1} = -\mathbb{R} \times_{\text{id}} \mathbb{R}_1^n,$$

donde $\text{id}: \mathbb{R} \rightarrow \mathbb{R}$ denota a la transformación identidad.

- El espacio de de Sitter $\mathbb{S}_1^{n+1}$ es isométrico al producto alabeado lorentziano

$$-\mathbb{R} \times_{\cosh} \mathbb{S}^n.$$

- Un subconjunto abierto del espacio hiperbólico $\mathbb{H}_1^{n+1}$ se puede ver como

$$-\left(-\frac{\pi}{2}, \frac{\pi}{2}\right) \times_{\cos} \mathbb{H}^n.$$

Para la demostración de estos dos últimos hechos, ver por ejemplo [8].

Veamos un ejemplo de resultado en los espacios RW.

Teorema 3.24 (Alías-Romero-Sánchez). *Sean F una variedad riemanniana compacta tal que $\bar{M} = -I \times_{\varrho} F$ sea una variedad de Einstein. Sea M una hipersuperficie tipo espacio en $\bar{M}$ con curvatura media constante, dada como la gráfica de una función $u: F \rightarrow I$. Entonces M es totalmente umbílica en $\bar{M}$.*

La teoría de las hipersuperficies con curvatura media constante es una teoría viva y en constante movimiento. En nuestra última sección hablaremos de algunos trabajos más recientes en esta dirección.

3.4. Epílogo: Hipersuperficies no umbílicas

En general, el estudio de las hipersuperficies de un determinado espacio ambiente comienza con las hipersuperficies más sencillas posibles: Las hipersuperficies totalmente geodésicas o las totalmente umbílicas. Pero es claro que si nuestro espacio ambiente es rico en isometrías, entonces puede haber otras hipersuperficies igual de interesantes y que merecen un estudio aparte.

Hace poco mencionamos un contraejemplo de Dajzcer y Nomizu a la conjetura de Goddard, es decir, una superficie no umbílica en $\mathbb{S}_1^3$. En este caso de baja dimensión, no ser umbílica implica tener dos curvaturas principales diferentes. Pero en dimensión mayor, ¿cuáles serían las hipersuperficies no umbílicas más sencillas?

Pensemos en el caso fácil, una hipersuperficie tipo espacio M , digamos, en $\mathbb{R}_1^{n+1}$. En este caso, sabemos que su operador de forma A es diagonalizable y por tanto M tiene n curvaturas principales $\kappa_1, \dots, \kappa_n$. En el caso umbílico, todas estas curvaturas son iguales. Pero el siguiente caso más sencillo sería, por ejemplo, que $(n-1)$ de estas curvaturas fuesen iguales. En este caso, ¿qué podríamos decir de M ?

Regresemos un poco a $\mathbb{R}^3$ y recordemos la definición de una superficie de rotación S . En este caso, tenemos una curva generatriz α , que podemos suponer

parametrizada por longitud de arco, y hacemos actuar sobre ella al grupo de $SO(2)$, pensado como un grupo de rotaciones.

Cuando tenemos más dimensiones, podemos seguir un procedimiento similar: Dada α una curva de $\mathbb{R}_1^{n+1}$, podemos hacer actuar sobre ella a un subgrupo de $SO(1, n)$; por ejemplo, a $SO(1, n-1)$. Al conjunto obtenido le llamaremos una *hipersuperficie de rotación*. Al realizar los cálculos, se puede ver que esta hipersuperficie cumple lo dicho en los párrafos anteriores; es decir, tiene $(n-1)$ curvaturas principales iguales.

En general, a las hipersuperficies de dimensión n que tienen $(n-1)$ curvaturas principales iguales se les llama $(n-1)$ -*umbílicas*. Es fácil convencerse que no todas las hipersuperficies $(n-1)$ -umbílicas son de rotación: De hecho, se puede pensar en una superficie en $\mathbb{R}^3$ que esté foliada por esferas, pero que no sea de rotación.

Así, un problema que puede considerarse es el estudio de este tipo de hipersuperficies y las condiciones bajo las cuales sean de rotación. Este problema ha sido abordado por el autor de estas notas, junto con varios colegas, en el contexto de las formas espaciales semi-riemannianas, obteniendo una clasificación de todas las hipersuperficies de rotación de las formas espaciales.

Sin embargo, el estudio de éste y otros tipos de hipersuperficies similares está lejos de ser completado en espacios ambiente que no sean formas espaciales. Por ejemplo, se pueden plantear preguntas como las siguientes:

- En un producto alabeado $\bar{M} = -I \times_{\varrho} F$, donde F admite una acción por $SO(n-1)$, ¿es posible caracterizar a todas las hipersuperficies de rotación, o a todas las hipersuperficies G -invariantes, donde G es un subgrupo de las isometrías de $\bar{M}$?
- Más en general, si una variedad $\bar{M}$ admite una acción de un grupo G por isometrías, ¿es posible caracterizar a todas las hipersuperficies G -invariantes?

Índice alfabético

- ángulo
 - hiperbólico, 4
- índice
 - de un espacio vectorial, 7
 - de un producto escalar, 7
- base
 - adaptada a un subespacio, 6
 - negativa, 14
 - ortonormal, 7
 - positiva, 14
 - pseudo-ortonormal, 12
 - que apunta hacia el futuro, 15
 - que apunta hacia el pasado, 15
- carácter causal
 - de un subespacio, 10
 - de un vector, 3
- conexión
 - normal, 41
- conjunto convexo, 3
- cono
 - causal de un vector, 3
 - de luz o nulo, 3
 - de luz o nulo de un vector, 3
 - de tiempo, 3
 - de tiempo o temporal de un vector, 3
- curva, 22
 - nula, 22
 - parametrizada diferenciable, 21
 - parametrizada regular, 21
 - tipo espacio, 22
 - tipo tiempo, 21
- curvatura
 - de Gauss-Kronecker, 40
 - gaussiana, 32
 - media, 32, 40
- curvaturas principales, 28
 - de una superficie tipo tiempo, 32
- descomposición en valores singulares, 15
- desigualdad
 - inversa de Cauchy-Schwarz, 4
 - inversa del triángulo, 4
- ecuación de Gauss, 39
- empujón (boost), 13
- espacio
 - de Lorentz-Minkowski, 2
 - euclidiano, 1
 - semi-euclidiano, 2
 - vectorial lorentziano, 10
- espacio de de Sitter
 - de dimensión 2, 24
 - de dimensión $n + 1$, 37
 - de radio r , 29
- espacio de Robertson-Walker
 - clásico, 38
 - generalizado, 37, 44
- espacio modelo lorentziano, 41
- forma espacial lorentziana, 41
- función
 - armónica, 43
 - eikonal, 43
- gráfica de una función como superficie, 23
- grupo
 - ortogonal, 14
 - ortogonal especial, 14
 - ortogonal especial ortócrono, 15

- ortogonal ortócrono, 15
- hipersuperficie, 38
 - $(n - 1)$ -umbílica, 46
 - de rotación, 46
 - isoparamétrica, 40
 - no degenerada, 38
 - totalmente umbílica, 40
- imagen inversa de un valor regular, 23
- isometría
 - lineal, 9
- laplaciano, 43
- longitud de arco, 22
- modelo del hiperboloide para el plano
 - hiperbólico, 25
- norma, 3
- operador de forma, 28
- orientación, 14
 - del tiempo o temporal, 15
- orientaciones
 - canónicas de $\mathbb{R}_1^{n+1}$, 15
- parametrización, 23
 - por longitud de arco, 22
 - por pseudo-longitud de arco, 22
- pregeodésica, 22
- producto
 - cruz lorentziano, 5
 - escalar, 1
 - escalar lorentziano, 10
- proyección
 - ortogonal, 8
- punto
 - crítico, 23
 - regular, 23
 - umbílico, 33, 40
- rotación, 13
 - propia, 13
- segunda forma fundamental, 32, 38
- subespacio
 - no degenerado, 5
 - tipo espacio o espacial, 10
 - tipo luz o nulo, 10
 - tipo tiempo o temporal, 10
- superficie
 - de Scherk, 36
 - de traslación, 35
 - degenerada o nula, 24
 - diferenciable, 23
 - isoparamétrica, 28, 32
 - no degenerada, 24
 - tipo espacio, 24
 - tipo tiempo, 24
 - totalmente umbílica, 33
- tensor de curvatura, 39
- transformación
 - de Gauss, 27
 - ortócrona, 15
 - que preserva el producto escalar, 8
- transformación lineal autoadjunta, 17
- valor
 - crítico, 23
 - regular, 23
- vector
 - causal, 3
 - tipo espacio o espacial, 3
 - tipo luz o nulo, 3
 - tipo tiempo o temporal, 3
 - unitario, 5
- vectores
 - ortogonales, 5

Bibliografía

- [1] L. J. Alías, A. Romero, M. Sánchez, *Spacelike hypersurfaces of constant mean curvature and Calabi-Bernstein type problems*, Tôhoku Math. Journal 49, pp. 337–345, 1997.
- [2] M. Dajczer, K. Nomizu, *On flat surfaces in $\mathbb{S}_1^3$ and $\mathbb{H}_1^3$* , in Manifolds and Lie groups, volumen 14, Progr. Math., pp. 71–108. Birkhäuser Boston, Mass., 1981.
- [3] S. H. Friedberg, A. J. Insel, L. E. Spence, *Linear Algebra*, Fourth Edition, Pearson Education, 2003.
- [4] J. Gallier, J. Quaintance, *Differential Geometry and Lie Groups for Computing*, libro en proceso.
- [5] B. Hall, *Lie groups, Lie algebras, and representations*, Springer, 2015.
- [6] R. López, *Differential Geometry of Curves and Surfaces in Lorentz-Minkowski Space*, International Electronic Journal of Geometry 7, número 1, pp. 44–107, 2014.
- [7] M. Magid, *Isoparametric Lorentzian hypersurfaces*, Pacific Journal of Mathematics 118, número 1, pp. 165–197, 1985.
- [8] S. Montiel, *Uniqueness of spacelike hypersurfaces of constant mean curvature in foliated spacetimes*, Mathematische Annalen 314, pp. 529–553, 1999.
- [9] S. Montiel, *An integral inequality for compact spacelike hypersurfaces in the de Sitter space and applications to the case of constant mean curvature*, Indiana Univ. Math. J. 37, pp. 909–917, 1988.
- [10] B. O’Neill, *Semi-Riemannian Geometry, with Applications to Relativity*, Academic Press, 1983.
- [11] Z. Petrov, *Einstein spaces*, Pergamon Press, 1969.