



**UNIVERSIDAD NACIONAL AUTÓNOMA
DE MÉXICO**

FACULTAD DE CIENCIAS

**SUPERFICIES EN EL ESPACIO DE MINKOWSKI
CON VECTOR DE CURVATURA MEDIA AFÍN**

T E S I S

QUE PARA OBTENER EL TÍTULO DE:

MATEMÁTICO

P R E S E N T A:

JOSÉ EDUARDO NÚÑEZ ORTIZ



**DIRECTOR DE TESIS:
DOCTOR GABRIEL RUIZ HERNÁNDEZ**

2015

1. Datos del alumno

Núñez

Ortiz

José Eduardo

56 02 90 32

Universidad Nacional Autónoma de

México

Facultad de Ciencias

Matemáticas

304158580

2. Datos del tutor

Dr.

Gabriel

Ruiz

Hernández

3. Datos del sinodal 1

Dr.

Oscar Alfredo

Palmas

Velasco

4. Datos del sinodal 2

Dra.

Laura

Ortiz

Bobadilla

5. Datos del sinodal 3

Dr.

Luis

Hernández

Lamonedá

6. Datos del sinodal 4

Dr.

Pablo

Suárez

Serrato

7. Datos del trabajo escrito.

Superficies en el espacio de Minkowski con vector de curvatura media afín

120 p

2015

Superficies en el espacio de Minkowski con vector de curvatura media afín

José Eduardo Núñez Ortiz

II

Quiero agradecer principalmente a mi familia por el apoyo recibido durante el tiempo que duro este proyecto. También le quiero gradecer al jurado que leyó ésta tesis por las correcciones y comentarios que por supuesto enriquecieron el resultado final aquí mostrado.

Investigación realizada gracias al Programa de Apoyo a Proyectos de Investigación e Innovación Tecnológica (PAPIIT) de la UNAM « PAPIIT IA100412 » « Geometría Diferencial de Subvariedades »

Índice general

Introducción	VII
1. Espacio de Minkowski	1
1.1. Métrica de Lorentz	2
1.2. Ortogonalidad	9
1.3. Vectores tipo tiempo	16
1.4. Isometrías	18
1.5. Operadores adjuntos y auto-adjuntos	24
2. Geometría Lorentziana	35
2.1. Variedades de Lorentz	35
2.2. La conexión de Levi-Civita	40
2.3. El tensor de curvatura	48
2.4. Operadores diferenciales	53
2.5. Ecuaciones de estructura	56
2.6. ¿Toda variedad puede ser de Lorentz?	59
3. Geometría Extrínseca	63
3.1. La conexión inducida	65
3.2. Ecuaciones fundamentales	70
3.3. Hipersuperficies	74
3.4. El espacio de De Sitter	78
4. Superficies tales que $\Delta\phi = A\phi + B$	83
4.1. El Laplaciano del vector H	84
4.2. Propiedades del operador A	93
4.3. Ejemplos de superficies en $\mathbb{R}_1^3$	97
4.4. Clasificación de superficies	106
4.5. Resultados afines	116
Bibliografía	119

Índice de figuras

1.1. Cono de luz	7
3.1. Espacio de De Sitter	76
3.2. Espacio Hiperbólico	77
4.1. Cilindro Circular	100
4.2. $H^1(r) \times \mathbb{R}$	101
4.3. $S_1^1(r) \times \mathbb{R}$	102

Introducción

La Física y las Matemáticas han caminado juntas a lo largo de siglos y es precisamente la Geometría una de las disciplinas en las que esta unión ha sido más notoria; esta tesis es una muestra de ello.

Aunque los temas aquí abordados son de naturaleza matemática, el lenguaje usado fue del interés de los matemáticos después que la Física mostrara su utilidad en la descripción de ciertos fenómenos físicos tales como la gravitación en el marco de la Teoría de la *Relatividad*. Así, si bien en este trabajo no se persigue ningún resultado de la Física se encontrarán nombres en las definiciones a hacer que nos recordarán la motivación de comprender el universo. Mas allá de lo anterior, el presente trabajo se basa en una serie de artículos de investigación.

Uno de los temas de mayor interés en la geometría de superficies es el de *superficies mínimas*, aquellas que cumplen con $H = 0$. Es un resultado conocido que las coordenadas de una superficie mínima encajada en $\mathbb{R}^3$ son armónicas, es decir, si ϕ es una inmersión isométrica de la superficie en el espacio Euclidiano, se tiene que $\Delta\phi = 0$. A lo largo de la historia los geómetras, curiosos como son, quisieron entender la geometría detrás de una condición aún más general y encontrar, de existir, las superficies tales que $\Delta\phi = \lambda\phi$ para alguna $\lambda \in \mathbb{R}$. Así es como se han generalizado las pesquisas de los geómetras hasta llegar al tema que nos ocupará aquí.

Esta tesis expone con cierto detalle los resultados obtenidos en [2] por Luis J. Alías, Ángel Ferrández y Pascual Lucas, usando además los resultados en [6] por los últimos dos y en [12] por Martin A. Magid. Por lo tanto, el objetivo de esta tesis es clasificar de manera local las superficies del espacio de Minkowski de dimensión tres, $\mathbb{R}_1^3$, que cumplan con la condición $\Delta\phi = A\phi + B$, donde ϕ es una inmersión isométrica de la superficie en $\mathbb{R}_1^3$, A es un operador lineal de $\mathbb{R}_1^3$ y B un vector en este espacio.

Para lograr esto me armaré con la herramienta necesaria a lo largo de tres capítulos para lograr el objetivo, el cual se expone en el cuarto y último.

En el primer capítulo se presenta el álgebra lineal básica de $\mathbb{R}_1^n$, pasando por la definición de *causalidad* en términos de la métrica, una nueva noción

de ortogonalidad y la versión *Lorentziana* del método de *Gram-Schmidt*, algunas propiedades de las isometrías en $\mathbb{R}_1^n$ y el *Grupo de Lorentz*. Al final de este capítulo sobresale una sección abocada a las propiedades de los operadores auto-adjuntos en el espacio de Minkowski, y se demuestra parcialmente un resultado sobre las representaciones canónicas que tienen este tipo de operadores.

El segundo capítulo es una introducción a las *variedades de Lorentz*, cuyo material reconocerá cualquier lector con un conocimiento básico del lenguaje de la Geometría Diferencial, rescatando la última sección en la que se dan condiciones suficientes para la existencia de *métricas de Lorentz* en una variedad C^∞ . Además, en este capítulo se introducen las *1-formas de la conexión* y la *ecuaciones de estructura*, que mostrarán su utilidad en el último capítulo adecuándolas a superficies en $\mathbb{R}_1^3$.

En el tercero se presenta la geometría de subvariedades y se muestran unas ecuaciones que son conocidas como fundamentales, las cuales se usarán mucho en el último capítulo.

El cuarto capítulo, donde se encuentra el resultado principal de la tesis, es el único en el cual se fija la dimensión de las variedades a dos, y se toma por espacio ambiente $\mathbb{R}_1^3$. Las primeras secciones de este capítulo preparan el material que se usará específicamente en la demostración de los teoremas principales, material que está basado principalmente en la herramienta desarrollada en los capítulos primero y tercero.

Con base en lo expuesto en los artículos [6] y [2], la exposición del resultado deseado empieza por encontrar una expresión explícita del Laplaciano del vector de curvatura media de una superficie inmersa en $\mathbb{R}_1^3$, lo cual se logra expresando en términos de la *función curvatura media*, h , la traza del operador $(\nabla_X S_\xi)(Y)$. Después de esto se exponen algunas propiedades de las superficies que cumplen con la condición deseada y así mostrar una serie de ejemplos de tales superficies. Es así como se prepara el camino para el Teorema de mayor importancia en este trabajo, el cual asegura que toda superficie que cumpla con la condición $\Delta\phi = A\phi + B$ es una superficie de curvatura media constante, y así, pasar a mostrar que tales superficies deben ser las *isoparamétricas* de $\mathbb{R}_1^3$, que junto con los resultados obtenidos en [12], logran la clasificación deseada.

La clasificación de superficies aquí mostrada es susceptible de una posterior generalización, donde los objetos a clasificar sean hipersuperficies en una variedad *semi-Riemanniana* conexa, completa y de curvatura constante. Este resultado se expone en [1] por los mismos tres autores del artículo [2]. Con esto, el presente trabajo puede ser tomado como una introducción a una clasificación geométrica mucho más rica y por supuesto, general.

Antes de comenzar con la exposición de esta tesis es necesario hacer una

observación. A lo largo de este trabajo se usa la convención de suma de índices de Einstein y se intenta seguir en todo caso la misma; en algunas partes no será del todo conveniente usarla y se regresará a la notación usual o se hará una modificación de la misma. En caso que el acomodo de los índices no implique una suma implícita se hará la observación conveniente. Esta convención determina que índices arriba y abajo se *contraen*, es decir, se suman; esto se logra numerando las *coordenadas* de un vector con índices superiores, y los vectores con índices inferiores. Por ejemplo, sea $v \in \mathbb{R}^n$ y $\beta = \{e_1, \dots, e_n\}$ la base usual; entonces v se puede expresar como combinación lineal de la base como:

$$v = x^i e_i = \sum_{i=1}^n x^i e_i.$$

Por último, en el primer capítulo también se conviene que los índices en letras griegas corren a partir del cero, mientras que los índices en letras latinas empiezan en el uno; esta última convención no se puede seguir del todo en los demás capítulos ya que la superficie puede ser tipo *tiempo* o tipo *espacio*.

Capítulo 1

Espacio de Minkowski

En el año de 1905 Einstein publicó una serie de artículos entre los cuales destacó uno que resolvía, entre otras cosas, las inconsistencias entre las *ecuaciones de Maxwell* y el *principio de relatividad de Galileo*; las *ecuaciones de Maxwell* parecían referirse a un sistema de referencia absoluto, es decir, uno respecto al cual los demás *sistemas inerciales* se encuentran en movimiento y en principio se podría calcular la velocidad absoluta de estos; *el principio de Galileo*, asegura que ningún experimento puede medir la velocidad absoluta de uno de éstos *sistemas inerciales*.

En la teoría propuesta por Einstein se preserva el *principio de relatividad de Galileo* para conservar la noción de que los fenómenos físicos son independientes del observador que los mida, es decir, no dependen de las coordenadas usadas para describirlos. Esta teoría, llamada Relatividad Especial, depende de otro principio:

- *Universalidad de la velocidad de la luz*: La velocidad de la luz para cualquier observador no *acelerado* (*observador* , *sistema inercial*) es $c = 3 \times 10^8 \text{ms}^{-1}$ sin importar el estado de movimiento de la fuente de luz. Esto quiere decir que dados dos *observadores inerciales* estos miden, para el mismo rayo de luz, que este se mueve con una velocidad c respecto a ellos independientemente del movimiento relativo entre los mismos.

En un principio, la exposición dada por Einstein del tema fue meramente algebraica, introduciendo primero las *transformaciones de Lorentz* para después entender fenómenos como la dilatación del tiempo, la contracción del espacio y sus consecuencias en la dinámica, particularmente, de objetos cargados en movimiento. A los pocos años de la publicación de este artículo por Einstein, Minkowski notó que era de mayor utilidad usar (t, x, y, z) como coordenadas en un espacio vectorial de cuatro dimensiones reales (en la base

usual de $\mathbb{R}^4$), al cual se llamó *espacio-tiempo* o espacio de Minkowski, en el cual uno puede desarrollar el lenguaje de la geometría diferencial.

Un *sistema inercial* es un *espacio-tiempo* tal que para cada t constante, la geometría es euclidiana en las restantes coordenadas. Para facilitar la notación, los físicos hacen un cambio de coordenadas al que llaman natural; con esto se refieren a expresar la velocidad de la luz como una constante adimensional e igual a uno, es decir, $c = 1$. Esto con el objetivo de medir el tiempo en metros, lo cual es útil en lo siguiente.

Usando los principios en los que se basa esta teoría supongamos dos *sistemas inerciales* (O y $\bar{O}$) en movimiento relativo el uno respecto del otro, y un rayo de luz tal que se medirá el tiempo que transcurre en llegar de un lugar A a otro B en las coordenadas O . Dado que la velocidad de la luz es igual a uno y se mide en metros, las diferencias $(\Delta t, \Delta x, \Delta y, \Delta z)$ entre la coordenadas de A a un tiempo t^A y de B a un tiempo t^B cumplen con la relación $(\Delta x)^2 + (\Delta y)^2 + (\Delta z)^2 - (\Delta t)^2 = 0$. Por lo *universalidad de la velocidad de la luz* se tiene que en el sistema de coordenadas $\bar{O}$ se cumple la relación $(\Delta \bar{x})^2 + (\Delta \bar{y})^2 + (\Delta \bar{z})^2 - (\Delta \bar{t})^2 = 0$. Por lo tanto el cambio de coordenadas $O \rightarrow \bar{O}$ debe ser tal que deje invariante la forma cuadrática:

$$-(\Delta t)^2 + (\Delta x)^2 + (\Delta y)^2 + (\Delta z)^2. \quad (1.1)$$

Con esto se puede ver que la forma cuadrática (1.1) va a funcionar en este caso como la *norma* del vector y, que los cambios de coordenadas de un *sistema inercial* a otro deberán ser *isometrías* del *espacio-tiempo*. Recordando que toda forma cuadrática viene de una forma *bilineal simétrica*, se puede ver la necesidad de construir un nuevo producto interno que admita vectores distintos del cero (como aquellos que unen distintos puntos sobre un mismo rayo de luz) tal que al hacer producto con ellos mismos este se haga cero. Esto nos lleva a definir el espacio de Minkowski como se hace a continuación, y desarrollar la herramienta matemática necesaria para la teoría de la Relatividad Especial, la cual no se expondrá aquí, pero que ha motivado la investigación de la geometría diferencial.

1.1. Métrica de Lorentz

En el caso clásico o euclidiano se tiene al espacio vectorial $\mathbb{R}^n$ acompañado de un *producto interno*, el cual es una forma *bilineal simétrica* tal que su forma cuadrática asociada es *positiva definida*; esto nos lleva a las nociones clásicas de medida, dadas por la trigonometría y el teorema de Pitágoras.

Siendo así, el único vector con *producto interno* consigo mismo igual a cero es el vector cero, y como se vio con anterioridad en la introducción, se

necesita de una forma cuadrática (la cual está asociada a una forma *bilineal simétrica*) que tome vectores distintos del vector cero y los lleve al cero.

Tomando la forma cuadrática (1.1) en las coordenadas usuales de $\mathbb{R}^4$ se obtiene la forma *bilineal*

$$\langle u, v \rangle := -u^0v^0 + u^1v^1 + u^2v^2 + u^3v^3 \quad (1.2)$$

cuya matriz asociada es

$$\eta = \begin{pmatrix} -1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

Es claro que esta función no define un *producto interno* ya que existe un subconjunto de $\mathbb{R}^4$ con vectores distintos del vector cero con la propiedad de que $\langle u, v \rangle = v^T \eta v = 0$, y no solo eso, sino que algunos vectores darán como resultado números negativos. Por lo tanto se necesita debilitar la condición de positivo definido, para que la forma *bilineal* definida en la ecuación (1.2) tome el papel del *producto interno*, alejándonos por supuesto, del caso clásico.

Definición 1.1. Una forma *bilineal simétrica* b , sobre un espacio vectorial V es:

1. *Positiva (negativa) definida* si para toda $u \in V$ distinta de cero se tiene que $b(u, u) > 0$ ($b(u, u) < 0$).
2. *Positiva (negativa) semidefinida* si para toda $u \in V$ se cumple que $b(u, u) \geq 0$ ($b(u, u) \leq 0$).
3. *No degenerada* si $b(u, v) = 0$ para toda $v \in V$, implica que $u = 0$

Para lograr que η tome el lugar del *producto interno* es necesario revisar, aunque sea someramente, la teoría de *formas bilineales*; dado que el tema principal de esta tesis no es el álgebra, no se demostrarán todos los detalles salvo los que se consideren necesarios. Cabe mencionar que a lo largo de toda la tesis todos los espacios vectoriales estarán definidos sobre el campo de los números reales.

Definición 1.2. Sea $\beta = \{v_0, \dots, v_n\}$ una base ordenada de un espacio vectorial V de dimensión finita y b una *forma bilineal* sobre V . Se define la *matriz asociada* o *representación matricial* de b en la base β como $[A_{\alpha\gamma}] =: A$ donde $A_{\alpha\gamma} := b(v_\alpha, v_\gamma)$.

Si definimos ahora el conjunto de todas las *formás bilineales* sobre un espacio vectorial V , el cual se denota $\mathcal{B}(V)$, con las siguientes operaciones:

- $(b_1 + b_2)(u, v) := b_1(u, v) + b_2(u, v)$
- $(cb_1)(u, v) := cb_1(u, v)$

Uno puede demostrar el siguiente resultado, donde $b_1, b_2 \in \mathcal{B}(V)$ y $c \in \mathbb{R}$.

Teorema 1.3. El conjunto $\mathcal{B}(V)$ tiene estructura de espacio vectorial.

El siguiente resultado es de mucha importancia, ya que nos permitirá usar la matriz asociada a una *forma bilineal* sin ser ambiguos en su uso.

Teorema 1.4. Sea V un espacio vectorial de dimensión finita y β una base ordenada de V . Entonces, la función $\psi_\beta : \mathcal{B}(V) \rightarrow M_{n \times n}(\mathbb{R})$ que asocia a cada *forma bilineal* una matriz, es un isomorfismo lineal.

De ahora en adelante se podrá usar de forma indistinta un solo símbolo para referirse tanto a la *forma bilineal* como a su representación matricial; lo mismo para todas las *formás bilineales* que se introduzcan a lo largo de la tesis.

Siendo un poco más formales, dada una base ordenada β en un espacio vectorial V de dimensión finita, existe una función $\phi_\beta : V \rightarrow \mathbb{R}^n$, la cual nos permite describir cada vector $u \in V$ en términos de sus coordenadas respecto a la base como una matriz de $1 \times n$. Siendo así, se obtiene el siguiente corolario al Teorema 1.4.

Corolario 1.5. Sea V un espacio vectorial de dimensión finita y sea β una base ordenada de V . Sean $b \in \mathcal{B}(V)$ y $A \in M_{n \times n}(\mathbb{R})$, entonces, $\psi_\beta(b) = A$ si y sólo si $b(u, v) = [\phi_\beta(u)]^T A [\phi_\beta(v)]$ para toda $u, v \in V$.

El punto 3 en la Definición 1.1 será de gran utilidad para lograr el objetivo deseado; vale la pena demostrar algo al respecto de este tipo de funciones.

Lema 1.6. Sea b una *forma bilineal simétrica* sobre un espacio vectorial V de dimensión finita; b es no degenerada si y sólo si su matriz asociada es invertible.

Demostración. Supongamos primero que b es no degenerada y $u \in V$ tal que $b(u, w) = 0$ para toda $w \in V$. Si $\beta = \{v_0, \dots, v_n\}$ es una base ordenada de V se obtiene que $b(u, v_\alpha) = 0$ para toda $\alpha = 0, \dots, n$.

En esta misma base $u = u^\gamma v_\gamma$. Entonces

$$b(u, v_\alpha) = b(u^\gamma v_\gamma, v_\alpha) = b_{\gamma\alpha} u^\gamma = 0 \text{ para toda } \alpha = 0, \dots, n$$

Como b es no degenerada, $u = 0$. Por lo tanto $u^\gamma = 0$ para toda $\gamma = 0, \dots, n$.

Dado que cada u^γ multiplica a toda una columna de la matriz asociada a b , se obtiene de lo anterior que las columnas de b son linealmente independientes; como b es *simétrica* también lo son los renglones de b . Por lo tanto, la matriz asociada a b es invertible.

Por un argumento similar se obtiene el regreso.

□

Teorema 1.7. Sea V un espacio vectorial de dimensión finita sobre $\mathbb{R}$. Entonces, toda *forma bilineal simétrica* en V es diagonalizable.

Una demostración del Teorema 1.7 se puede encontrar en [7] para campos que no sean de característica dos.

Lema 1.8. Sea b una *forma bilineal simétrica* sobre un espacio vectorial V de dimensión finita; b es positiva definida si y sólo si b es positiva semidefinida y no degenerada.

Demostración. Si b es positiva definida, claramente es positiva semidefinida y no degenerada. Supongamos ahora que b es positiva semidefinida y no degenerada. Haciendo uso del Teorema 1.7, existe una base de V tal que b es diagonal en esa base.

Como b es positiva semidefinida, se tiene que los elementos de la diagonal son mayores o iguales a cero; de no ser así $b_{\beta\beta} < 0$ para alguna β , lo cual implicaría que existe un subespacio $W \subset V$ tal que $b|_W$ es negativa definida, lo cual sería una contradicción.

Además, b es no degenerada y por el Lema 1.6 su matriz asociada es invertible; entonces, en su forma diagonal, los elementos de la diagonal son todos distintos de cero. Por lo tanto los elementos de la diagonal son mayores que cero, lo cual nos lleva a que b es positiva definida. □

A continuación se definirá el *índice* de una *forma bilineal simétrica*, el cual será una de las definiciones más importantes de esta sección.

Definición 1.9. El *índice* ν de una *forma bilineal* sobre V es la dimensión del mayor subespacio $W \subset V$ tal que $b|_W$ es negativa definida.

Entonces, $0 \leq \nu \leq \dim V$. Además, claramente $\nu = 0$ si y sólo si b es positiva semidefinida.

Con todo esto, ahora se puede dar la definición buscada para que se pueda usar a η como si fuera el *producto interno*.

Definición 1.10. Un *producto escalar* en un espacio vectorial V , es una *forma bilineal simétrica* no degenerada.

Sea g un producto escalar, si g tiene índice $\nu = 0$ éste es un *producto interno*. Por lo tanto, la Definición 1.10 generaliza la noción de *producto interno*, permitiendo distintas formas de *medir* en espacios vectoriales; si el índice de un producto escalar es distinto de cero, éste no induce una métrica en el sentido analítico; aunque η no induzca una *métrica* en $\mathbb{R}^n$, es común referirse a η como la métrica del *espacio de Minkowski*.

Definición 1.11. Sea $\langle \cdot, \cdot \rangle : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$ una *forma bilineal simétrica* sobre $\mathbb{R}^n$ que en la base canónica de este espacio se define como sigue:

$$\langle u, v \rangle := -u^0 v^0 + \cdots + u^{n-1} v^{n-1}$$

Así definida, su matriz asociada en dicha base es de la forma:

$$\eta = \begin{pmatrix} -1 & 0 & 0 & \cdots & 0 \\ 0 & 1 & 0 & \cdots & 0 \\ 0 & 0 & \ddots & & \vdots \\ \vdots & \vdots & & \ddots & \vdots \\ 0 & 0 & \cdots & \cdots & 1 \end{pmatrix}$$

Entonces, se define el *espacio de Minkowski* como la pareja $(\mathbb{R}^n, \langle \cdot, \cdot \rangle)$ para $n \geq 2$, donde $\langle \cdot, \cdot \rangle$ es un producto escalar de índice $\nu = 1$.

Para no cargar la notación, es común designar al espacio de Minkowski como $\mathbb{R}_1^n$; en el caso de un espacio vectorial V con un producto escalar de índice ν se suele denotar a este como V_ν ; si la dimensión de V es finita, $V_\nu \simeq \mathbb{R}_\nu^n$. Claramente la ecuación (1.2) es el producto escalar del espacio de Minkowski cuando $n = 4$; en este caso al espacio de Minkowski se le conoce mejor como *espacio-tiempo*.

Causalidad

Como se ha visto hasta ahora, con este nuevo producto escalar η tendremos vectores tales que su producto consigo mismos puede incluso ser negativo. Por ejemplo, tómese la base canónica $\{e_0, e_1, e_2, e_3\}$ de $\mathbb{R}^4$; entonces, para e_0 :

$$\langle e_0, e_0 \rangle = -1$$

para $e_0 + e_1$:

$$\langle e_0 + e_1, e_0 + e_1 \rangle = 0$$

y para e_i :

$$\langle e_i, e_i \rangle = 1 \text{ para toda } i = 1, 2, 3$$

Esto nos lleva a clasificar a los vectores en $\mathbb{R}_1^n$ por el signo de su producto consigo mismo; es a esta clasificación a la que llamamos causalidad.

Definición 1.12. Un vector $v \in \mathbb{R}_1^n$ es:

$$\begin{array}{lll} \text{tipo espacio} & \text{si} & \langle v, v \rangle > 0 \quad \text{o} \quad v = 0 \\ \text{tipo luz o nulo} & \text{si} & \langle v, v \rangle = 0 \quad \text{y} \quad v \neq 0 \\ \text{tipo tiempo} & \text{si} & \langle v, v \rangle < 0 \end{array}$$

Aquí se puede apreciar, por los nombres en la clasificación, que la *causalidad* que se acaba de definir refleja ciertas cualidades físicas del *espacio-tiempo*. Por ejemplo, que v sea *tipo espacio* quiere decir que éste une dos *eventos* (puntos en el *espacio-tiempo*) tales que uno no depende del otro causalmente, es decir, para que suceda uno no se necesita que suceda el otro; esto implica que si un *observador inercial* ve suceder primero el evento A que el evento B , existe otro que observará sucede B antes que A . Que un vector sea *tipo luz* significa que une dos *eventos* por los cuales pasa un mismo rayo de luz. Que sea *tipo tiempo* significa que no existen dos *sistemas inerciales* distintos tales que no se respete el orden temporal de los eventos que une el vector, es decir, estos eventos están relacionados causalmente.

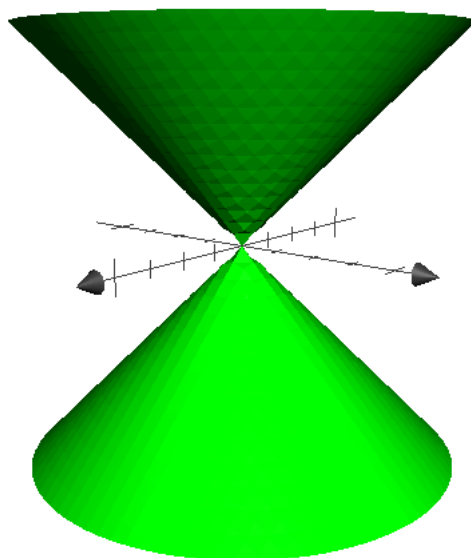


Figura 1.1: Cono de luz

Al conjunto de todos los vectores *nulos* (ver figura 1.1) se le conoce como el *cono de luz* $\mathcal{C}$. Los vectores *tipo tiempo* forman otro conjunto al cual se le da el nombre de *cono tipo tiempo* $\mathcal{T}$.

Como se vio en el ejemplo al inicio de esta sección, el vector e_0 es el único en la base $\{e_0, \dots, e_n\}$ que es tipo tiempo, esto quiere decir que en esta base la primera coordenada de un vector *tipo tiempo* siempre es distinta de cero. Al ser el vector cero *tipo espacio* se tiene que $\mathcal{T}$ es una unión disjunta de los conjuntos $\mathcal{T}^+$ y $\mathcal{T}^-$ definidos como sigue:

$$\mathcal{T}^+ := \{u \in \mathcal{T} \mid u^0 > 0\} ; \mathcal{T}^- := \{u \in \mathcal{T} \mid u^0 < 0\}$$

Se pueden definir los conjuntos $\mathcal{C}^+$ y $\mathcal{C}^-$ para $\mathcal{C}$ de la misma forma; siendo así, $\mathcal{C}$ es también una unión disjunta.

Proposición 1.13. Sean u y v dos vectores en $\mathcal{T}$. Si u, v se encuentran en la misma componente de $\mathcal{T}$ entonces $\langle u, v \rangle < 0$, de lo contrario $\langle u, v \rangle > 0$.

Demostración. Nótese que $\mathbb{R}_1^n$ puede ser visto como el producto $\mathbb{R}_1 \times \mathbb{R}^{n-1}$ tal que $\mathbb{R}^{n-1}$ es el factor que al restringir el producto escalar a este subespacio, éste se vuelve positivo definido. Por lo tanto, η restringida a $\mathbb{R}^{n-1}$ es el producto interno usual que coincide con la *norma* euclidiana. Sea $\tilde{u} := \pi_e(u)$ donde:

$$\pi_e: \mathbb{R} \times \mathbb{R}^{n-1} \rightarrow \mathbb{R}^{n-1} \text{ tal que } \pi_e(u^0, u^1, \dots, u^{n-1}) := (0, u^1, \dots, u^{n-1})$$

y defino $\|\tilde{u}\| := \sqrt{\langle \tilde{u}, \tilde{u} \rangle}$. Cuando u y v se encuentran en la misma componente de $\mathcal{T}$ se tiene que $u^0 v^0 > 0$. Por hipótesis:

$$\begin{aligned} \langle u, u \rangle &= -(u^0)^2 + \|\tilde{u}\|^2 < 0 \text{ y } \langle v, v \rangle = -(v^0)^2 + \|\tilde{v}\|^2 < 0 \\ (u^0)^2 &> \|\tilde{u}\|^2 ; (v^0)^2 > \|\tilde{v}\|^2 \end{aligned} \quad (1.3)$$

Por otro lado:

$$\begin{aligned} \langle u, v \rangle &= -u^0 v^0 + \eta_{ij} u^i v^j \\ &= -u^0 v^0 + \|\tilde{u}\| \|\tilde{v}\| \cos \theta \\ &\leq -u^0 v^0 + \|\tilde{u}\| \|\tilde{v}\| \end{aligned}$$

donde θ es el ángulo entre las proyecciones de los vectores. De las ecuaciones (1.3) se obtiene la siguiente desigualdad:

$$(u^0 v^0)^2 > (\|\tilde{u}\| \|\tilde{v}\|)^2$$

Como $u^0 v^0 > 0$ se tiene que $u^0 v^0 > \|\tilde{u}\| \|\tilde{v}\|$. Por lo tanto:

$$\langle u, v \rangle \leq -u^0 v^0 + \|\tilde{u}\| \|\tilde{v}\| < 0$$

En caso de que u y v estén en distintas componentes se tiene que $u^0 v^0 < 0$; en este caso, las ecuaciones (1.3) implican que $|u^0 v^0| > \|\tilde{u}\| \|\tilde{v}\|$. Por lo tanto:

$$\langle u, v \rangle \geq -u^0 v^0 - \|\tilde{u}\| \|\tilde{v}\| > 0$$

□

Así como se definió la causalidad para los vectores de $\mathbb{R}_1^n$, también se puede definir para los subespacios vectoriales de $\mathbb{R}_1^n$ en términos de la métrica restringida.

Definición 1.14. Sea η la métrica del espacio de Minkowski, y sea $U \subset \mathbb{R}_1^n$ un subespacio vectorial; decimos que:

$$\begin{aligned} U \text{ es tipo espacio} & \text{ si } \eta|_U \text{ es positiva definida} \\ U \text{ es tipo luz o nulo} & \text{ si } \eta|_U \text{ es degenerada} \\ U \text{ es tipo tiempo} & \text{ si } \eta|_U \text{ tiene índice uno} \end{aligned}$$

1.2. Ortogonalidad

En este caso al igual que en el euclidiano diremos que dos vectores u, v en $\mathbb{R}_1^n$ son *ortogonales* si $\langle u, v \rangle = 0$, y de la misma forma, diremos que un subconjunto de $\mathbb{R}_1^n$ es *ortogonal* si todos los vectores del subconjunto son ortogonales entre sí. La idea geométrica de ortogonalidad no coincidirá con la que se tenía antes, es decir, dos vectores que eran ortogonales en $\mathbb{R}^n$ no lo serán, en general, con el producto de $\mathbb{R}_1^n$.

El caso de los vectores nulos en $\mathbb{R}_1^n$ es el que más se aleja de la intuición euclidiana que poseemos puesto que un vector nulo siempre es ortogonal consigo mismo, y por lo tanto a todos los vectores que se encuentran en el subespacio generado por el mismo.

Proposición 1.15. Dos vectores nulos son ortogonales, si y sólo si, son linealmente dependientes

Demostración. Tomemos dos vectores nulos l y k , entonces:

$$(l^0)^2 = \|\tilde{l}\|^2 \text{ y } (k^0)^2 = \|\tilde{k}\|^2 \quad (1.4)$$

Por lo hecho en la demostración del Teorema 1.13 se tiene que $\langle l, k \rangle = -l^0 k^0 + \|\tilde{l}\| \|\tilde{k}\| \cos \theta$. Entonces, por las ecuaciones (1.4) se tiene:

$$(l^0 k^0)^2 = \left(\|\tilde{l}\| \|\tilde{k}\| \right)^2$$

Supongamos que l y k son ortogonales; para el caso que $l^0 k^0 > 0$ se tiene que $-l^0 k^0 + \|\tilde{l}\| \|\tilde{k}\| = 0$. Por lo tanto $\theta = 0$ y $\tilde{k} = c\tilde{l}$ con $c > 0$. Como l^0 y k^0 son números reales puedo afirmar que existe $r \in \mathbb{R}$ tal que $k^0 = r l^0$. Por lo tanto queda la siguiente ecuación:

$$-r (l^0)^2 + c \|\tilde{l}\|^2 = 0,$$

y finalmente que $r = c$. Con esto se tiene que l y k son linealmente dependientes; para el caso que $l^0 k^0 < 0$ se tiene que $\theta = \pm\pi$ y siguiendo de manera análoga lo anterior, pero con $c < 0$, se obtiene el mismo resultado.

El regreso es trivial; basta con recordar que si dos vectores son linealmente dependientes, uno es múltiplo escalar del otro. $\square$

Con todo esto ¿existen vectores no nulos en $\mathbb{R}_1^n$ que sean ortogonales a un vector nulo dado? La respuesta es sí, y no sólo eso, también se puede determinar la causalidad de estos.

Proposición 1.16. Dado un vector nulo l , no existe $u \in \mathcal{T}$ tal que $\langle l, u \rangle = 0$

Demostración. Hasta aquí debe ser claro que para un vector nulo como l , $l^0 \neq 0$, de no ser así $\|\tilde{l}\|$ sería igual a cero y el vector l sería el vector cero, el cual por definición es tipo espacio. Entonces, si u es cualquier vector tipo tiempo se tiene:

$$(u^0)^2 > \|\tilde{u}\|^2 ; (l^0)^2 = \|\tilde{l}\|^2$$

$$(u^0 l^0)^2 > \left(\|\tilde{u}\| \|\tilde{l}\| \right)^2$$

Como se ha visto en el desarrollo de la tesis, existen dos casos $u^0 l^0 > 0$ y $u^0 l^0 < 0$; tomemos el segundo caso. Entonces:

$$|u^0 l^0| = -u^0 l^0 > \|\tilde{u}\| \|\tilde{l}\|$$

y finalmente se obtiene:

$$\begin{aligned} \langle u, l \rangle &= -u^0 l^0 + \eta_{ij} u^i l^j \\ &= -u^0 l^0 + \|\tilde{u}\| \|\tilde{l}\| \cos \theta \\ &\geq -u^0 l^0 - \|\tilde{u}\| \|\tilde{l}\| > 0 \end{aligned}$$

En el caso restante se prosigue de manera análoga, como se ha visto en previas demostraciones a lo largo de este capítulo. $\square$

Se puede notar que gran parte de estos resultados dependen de que la primera coordenada de los vectores no se anule; lo mismo pasará para los vectores tipo espacio cuya primera coordenada no sea cero, es decir, si l es nulo y, v algún vector tipo espacio tal que $v^0 \neq 0$, entonces l y v no serán ortogonales.

Pensemos en un vector v tipo espacio tal que $v^0 = 0$, entonces $\langle v, l \rangle = \|\tilde{v}\| \|\tilde{l}\| \cos \theta$. Por lo tanto, los únicos vectores en $\mathbb{R}_1^n$ que son ortogonales a un vector nulo dado, son tipo espacio tales que su primera coordenada es cero y que sus imágenes bajo la proyección π_e son ortogonales en el sentido usual o euclidiano, es decir, el ángulo entre sus proyecciones es de $\pm\frac{\pi}{2}$.

Hasta aquí se han expuesto algunos de los resultados que se tienen con esta nueva concepción de ortogonalidad, pero algunos casos se mantienen igual, por ejemplo, la base $\{e_\alpha\}_{\alpha=0}^n$ conserva sus propiedades de ser un subconjunto ortogonal de $\mathbb{R}_1^n$. ¿Conserva su *ortonormalidad*? Para esto sería necesario volver a definir lo que significa que un conjunto sea ortonormal.

Definición 1.17. Decimos que una colección de vectores $\{u_\alpha\}_{\alpha=0}^k$ es *ortonormal* en $\mathbb{R}_1^n$ si:

$$|\langle u_\alpha, u_\beta \rangle| = \delta_{\alpha\beta}$$

Como se puede ver, esta definición excluye la posibilidad de que existan vectores nulos en un conjunto ortonormal. Aún así, en el caso de que un conjunto sea ortogonal no se excluye la posibilidad de que alguno de estos sea nulo, es más, no basta con la hipótesis de ortogonalidad para que el conjunto sea linealmente independiente.

Teorema 1.18. Sea $S = \{u_\alpha\}_{\alpha=0}^k$ un subconjunto ortogonal de $\mathbb{R}_1^n$ formado con vectores distintos de cero tal que $\text{span}(S) = U$ es un subespacio no nulo. Entonces u_α no es un vector nulo para ninguna α .

Demostración. Si U es tipo espacio no hay nada que demostrar ya que $\eta|_U$ es positiva definida y $\langle u_\alpha, u_\alpha \rangle \neq 0$ para toda α .

Si U es tipo tiempo existe $w = w^\alpha u_\alpha$ tipo tiempo en U tal que, tomando $u_\beta \in S$:

$$\langle w, u_\beta \rangle = \langle w^\alpha u_\alpha, u_\beta \rangle = w^\beta \langle u_\beta, u_\beta \rangle.$$

Si supongo que u_β es un vector nulo, por la Proposición 1.16 se tiene que $\langle w, u_\beta \rangle \neq 0$ y por lo anterior que $\langle u_\beta, u_\beta \rangle \neq 0$, lo cual es una contradicción. Por lo tanto u_β no puede ser un vector nulo, y como este vector se tomó de manera arbitraria, ningún elemento de S puede ser nulo. □

Corolario 1.19. Sea $S = \{u_\alpha\}_{\alpha=0}^k$ un subconjunto ortogonal de $\mathbb{R}_1^n$ formado con vectores distintos de cero tal que $\text{span}(S) = U$ es un subespacio no nulo. Entonces S es linealmente independiente.

Demostración. Por el Teorema 1.18 se tiene que ningún vector en S es nulo; así, si $a^\alpha u_\alpha = 0$ es una combinación lineal de los elementos de S se tiene que:

$$0 = \langle a^\alpha u_\alpha, u_\beta \rangle = a^\beta \langle u_\beta, u_\beta \rangle.$$

Como ningún vector en S es nulo, necesariamente $a^\beta = 0$ para toda β . Por lo tanto en una combinación lineal de elementos de S igualada a cero, los coeficientes de la combinación lineal son cero. □

Hasta aquí, los resultados expuestos parecen depender de que exista una base ortonormal para $\mathbb{R}_1^n$, la cual por construcción existe, pero no es la única base que se pudiera tener para este espacio. Para mostrar que estos resultados no dependen de la base canónica, necesito una forma de ortogonalizar toda base dada y así, al poder siempre ortogonalizar (y después normalizar) cualquiera, darle sentido a los resultados expuestos y seguir usando la base canónica. El siguiente teorema es el conocido *método de ortogonalización de Gram-Schmidt*, aplicado ahora a subconjuntos de $\mathbb{R}_1^n$.

Teorema 1.20. Sea $S = \{u_\alpha\}_{\alpha=0}^k$ un subconjunto linealmente independiente de $\mathbb{R}_1^n$ que consiste de vectores no nulos tal que $\text{span}(S) = U$ es un subespacio no nulo. Si $\bar{S} = \{v_\alpha\}_{\alpha=0}^k$ es otro subconjunto de $\mathbb{R}_1^n$ tal que $v_0 = u_0$ y:

$$v_i = u_i - \sum_{j=1}^{i-1} \frac{\langle u_i, v_j \rangle}{\langle v_j, v_j \rangle} v_j \quad ; \text{ para } 1 \leq i \leq k. \quad (1.5)$$

Entonces, $\bar{S}$ es un conjunto ortogonal de vectores no nulos tal que $\text{span}(\bar{S}) = U$.

Demostración. La prueba será por inducción sobre el número de elementos k de los conjuntos S_k y $\bar{S}_k$. El caso $k = 0$ es trivial ya que al existir un solo elemento en S_0 , $v_0 = u_0$ y $S_0 = \bar{S}_0$. Por las hipótesis se tiene que U_0 es no nulo.

En el caso del espacio de Minkowski se vuelve un poco necesario desarrollar el caso $k = 1$, ya que no es del todo evidente que los elementos del conjunto $\bar{S}_1$ sean vectores no nulos. Entonces, sea S_1 un subconjunto de $\mathbb{R}_1^n$ que cumpla con las hipótesis del teorema y sea $v_0 := u_0$; lo que se busca es construir un vector v_1 que sea ortogonal a v_0 , y que $\text{span}(\bar{S}_1) = U_1$. Siendo así, conviene expresar a v_1 como una combinación lineal de los elementos de S_1 .

Sea μ un escalar tal que:

$$v_1 := u_1 - \mu u_0.$$

Entonces:

$$0 = \langle v_0, v_1 \rangle = \langle v_0, u_1 \rangle - \mu \langle v_0, u_0 \rangle,$$

y tomando en cuenta que $v_0 = u_0$

$$\mu = \frac{\langle u_1, v_0 \rangle}{\langle v_0, v_0 \rangle}.$$

Por lo tanto:

$$v_1 = u_1 - \frac{\langle u_1, v_0 \rangle}{\langle v_0, v_0 \rangle} v_0.$$

Lo cual está bien definido ya que u_0 es no nulo; además, siempre y cuando S no sea un subconjunto ortogonal de $\mathbb{R}_1^n$ la constante real μ es distinta de cero.

Con lo anterior es fácil ver que la independencia lineal de S implica la independencia lineal del conjunto $\bar{S}_1$, el cual genera un subespacio de dimensión dos de $\mathbb{R}_1^n$, lo que implica que $\text{span}(\bar{S}_1) = U_1$.

Por último, para mostrar que $\bar{S}_1$ no contiene vectores nulos basta con recordar que U_1 es un subespacio no nulo y por lo tanto el determinante de $\eta|_U$ debe ser distinto de cero. La representación de $\eta|_U$ en la base $\bar{S}_1$ debe ser diagonal, lo cual restringe los valores de la diagonal a ser distintos de cero, es decir, $\langle v_0, v_0 \rangle$ y $\langle v_1, v_1 \rangle$ deben ser distintos de cero.

Ahora sí, supongamos el enunciado de este teorema válido para algún $k-1 > 2$. Entonces podemos asumir que existe el conjunto ortogonal $\bar{S}_{k-1} = \{v_\alpha\}_{\alpha=0}^{k-1}$ donde cada vector en este conjunto está dado por la ecuación (1.5), tal que $\text{span}(\bar{S}_{k-1}) = U_{k-1}$.

Para el conjunto $\bar{S}_k$, ya que por hipótesis de inducción se cumple el teorema para $k-1$ vectores en este conjunto, basta con ver que el vector v_k cumple con las propiedades deseadas así como para $\bar{S}_k$. Por la ecuación (1.5) se obtiene que $v_k \neq 0$ ya que el vector $u_k \notin U_{k-1}$ y por lo tanto la resta en (1.5) nunca se anula, así, el conjunto $\bar{S}_k$ consta de vectores distintos del vector cero.

Ahora bien, para $0 \leq l \leq k-1$ se tiene:

$$\begin{aligned} \langle v_k, v_l \rangle &= \langle u_k, v_l \rangle - \sum_{j=1}^{k-1} \frac{\langle u_k, v_j \rangle}{\langle v_j, v_j \rangle} \langle v_j, v_l \rangle, \\ &= \langle u_k, v_l \rangle - \frac{\langle u_k, v_l \rangle}{\langle v_l, v_l \rangle} \langle v_l, v_l \rangle = 0, \end{aligned}$$

asumiendo que $\bar{S}_{k-1}$ no contiene vectores no nulos, por la hipótesis de inducción. Entonces $\bar{S}_k$ es un conjunto ortogonal, y por (1.5) se tiene que $\text{span}(\bar{S}_k) \subseteq U_k$, faltando ver que son iguales.

Si al agregar al vector v_k , el conjunto $\bar{S}_k$ perdiera independencia lineal, se obtendría que $u_k \in \text{span}(\bar{S}_{k-1}) = U_{k-1}$ contradiciendo las hipótesis del teorema; siendo así, $\bar{S}_k$ es un conjunto linealmente independiente y tiene dimensión k . Por lo tanto $\text{span}(\bar{S}) = U_k$.

Como U_k es un subespacio no nulo de $\mathbb{R}_1^n$ el Teorema 1.18 asegura que no existen vectores nulos en $\bar{S}_k$.

□

Aún con esto, no queda del todo claro qué es lo que pasa cuando se tiene un espacio vectorial V de dimensión finita con un producto escalar arbitrario

de índice uno. Llamemos a tal espacio vectorial V_1 y ω a su producto escalar. Entonces, en V_1 existe una base tal que ω es diagonal. El siguiente teorema se conoce por el nombre de *ley de inercia de Sylvester*, [7].

Teorema 1.21. Sea H una forma bilineal simétrica en un espacio vectorial de dimensión finita W . Entonces el número de entradas positivas y negativas en la diagonal, en cualquier representación diagonal de H , es independiente de la representación diagonal

Por lo tanto, con este teorema en mano y el Teorema 1.7, se tiene que la forma en la que se han expuesto los resultados hasta ahora no dependen de la base usada, ya que son ciertos en cualquier base, es decir, sin importar cual sea el producto ω que se tenga en V_1 , al ser el índice un invariante se puede concluir que siempre existe un cambio de base Q tal que $\eta = Q^T \omega Q$, y basta con hacer el cambio de base para que las técnicas usadas hasta ahora cobren sentido.

Una consecuencia del Teorema 1.21 es que en toda base ortonormal de $\mathbb{R}_1^n$ existe solo un vector tipo tiempo, siendo los demás tipo espacio.

Definición 1.22. Sea U un subespacio de $\mathbb{R}_1^n$. Se define su *complemento ortogonal* $U^\perp$ como sigue

$$U^\perp := \{v \in \mathbb{R}_1^n \mid \langle u, v \rangle = 0 \ \forall \ u \in U\}$$

Es claro que $U^\perp$ es a su vez un subespacio de $\mathbb{R}_1^n$. ¿Pero es verdad que la suma de $U^\perp$ con U es directa e igual a $\mathbb{R}_1^n$? A diferencia del caso clásico, esto no sucederá siempre.

Lema 1.23. Sea U un subespacio de $\mathbb{R}_1^n$, entonces:

1. $\dim(U^\perp) = \dim(\mathbb{R}_1^n) - \dim(U)$.
2. $(U^\perp)^\perp = U$.
3. U es no degenerado si y sólo si $U^\perp$ es no degenerado.

Demostración. 1. Sea $\beta = \{v_\alpha\}_{\alpha=0}^k$ una base de U tal que $\bar{\beta} = \{v_\alpha\}_{\alpha=0}^{n-1}$ es la extensión de la base β a una base de $\mathbb{R}_1^n$. Si $w = w^\mu v_\mu$ es un elemento de $U^\perp$ en la base $\bar{\beta}$, se obtiene lo siguiente:

$$\langle w^\mu v_\mu, v_\nu \rangle = w^\mu \langle v_\mu, v_\nu \rangle = g_{\mu\nu} w^\mu = 0 \quad \text{donde } 1 \leq \nu \leq k.$$

Esto se puede expresar de forma matricial como sigue:

$$\begin{pmatrix} g_{00} & \cdots & g_{0n-1} \\ \vdots & \ddots & \vdots \\ g_{k0} & \cdots & g_{kn-1} \end{pmatrix} \begin{pmatrix} w^0 \\ \vdots \\ w^{n-1} \end{pmatrix} = \begin{pmatrix} 0 \\ \vdots \\ 0 \end{pmatrix}$$

Entonces, si $A = [g_{\mu\nu}]_{m \times n-1}$, A es un automorfismo de rango máximo ya que la métrica es no degenerada. Como $\ker(A) = U^\perp$, por el teorema de la dimensión se tiene que $\dim(U^\perp) = n - (k - 1) = \dim(\mathbb{R}_1^n) - \dim(U)$.

2. Dado que para todo $u \in (U^\perp)^\perp$ se cumple que $\langle u, v \rangle = 0$ si $v \in U^\perp$, se obtiene la contención $(U^\perp)^\perp \subseteq U$. Por 1 de este lema se tiene que estos conjuntos son iguales ya que tienen la misma dimensión.
3. Como U no es tipo luz existe una base ortonormal del mismo, además, esta se puede extender a una base ortonormal de $\mathbb{R}_1^n$. Siendo así, para todo $u \in U$ se tiene que $u^j = 0$ si $k + 1 \leq j \leq n - 1$ asumiendo que $\hat{\beta} = \{e_\alpha\}_{\alpha=0}^k$ es una base ortonormal de U y $\beta = \{e_\beta\}_{\beta=0}^{n-1}$ su extensión a una base ortonormal de $\mathbb{R}_1^n$. Entonces, para todo $v \in \mathbb{R}_1^n$ y para todo $u \in U$ se tiene que:

$$\begin{aligned}\langle u, v \rangle &= \pm u^0 v^0 \pm \dots \pm u^{n-1} v^{n-1} \\ &= \pm u^0 v^0 \pm \dots \pm u^k v^k\end{aligned}$$

En la base β , sea $W = \{w \in \mathbb{R}_1^n | w^i = 0 \text{ si } 0 \leq i \leq k\}$. Claramente W es un subespacio de $U^\perp$ y además, $\{e_\beta\}_{\beta=k+1}^{n-1}$ es una base para W . Así, por el número 1 de esta demostración se tiene que W y $U^\perp$ tienen la misma dimensión y por lo tanto son iguales, siendo $\{e_\beta\}_{\beta=k+1}^{n-1}$ una base ortonormal de $U^\perp$; entonces, la métrica en esta base tiene una representación matricial diagonal con ± 1 en la diagonal. Con esto, $U^\perp$ no es tipo luz. El regreso se obtiene usando el número 2 de este teorema y el mismo argumento que se hizo sobre U ahora para $U^\perp$.

□

El enunciado 3 del Lema 1.23 implica que si U es degenerado $U^\perp$ también lo será, lo cual nos lleva a preguntar en que caso la suma de U con su complemento ortogonal es directa. El siguiente teorema resuelve esta cuestión.

Teorema 1.24. Sea U un subespacio de $\mathbb{R}_1^n$, entonces, U es no nulo si y sólo si $U \cap U^\perp = 0$.

Demostración. Si U es no nulo existe una base ortonormal β para este espacio. Supongamos que existe un vector u distinto de cero en $U \cap U^\perp$. Como u es elemento del complemento ortogonal de U se cumple que:

$$\langle u, \beta_i \rangle = 0 ; \text{ para todo } \beta_i \in \beta.$$

Pero u también es elemento de U y se puede expresar en la base β como $u = u^j \beta_j$. Entonces:

$$\langle u^j \beta_j, \beta_i \rangle = u^i \langle \beta_i, \beta_i \rangle = 0. ; \text{ para todo } \beta_i \in \beta$$

Esto por ser β una base ortonormal. Pero ya que u es distinto de cero, para algún β_i el producto consigo mismo es cero, lo cual es una contradicción y por tanto u debe ser el vector cero.

Para demostrar la otra implicación supongamos que U es nulo. Por el Teorema 1.7 sabemos que existe una base de U tal que contiene solo un vector nulo v y que este es ortogonal a los demás elementos de dicha base, los cuales no pueden ser tipo tiempo por la Proposición 1.16; por el Lema 1.23 $U^\perp$ es nulo también. Así, si w es el vector nulo en la base de $U^\perp$ que tiene las mismas propiedades que la base en U , por la Proposición 1.15, v y w son linealmente dependientes y con esto $U \cap U^\perp \neq 0$. □

Con este resultado es claro que la suma de U con su complemento ortogonal es directa e igual a $\mathbb{R}_1^n$ si y sólo si U es no degenerado. Con esto se puede enunciar el siguiente resultado sin necesidad de demostración.

Teorema 1.25. U es tipo tiempo si y sólo si $U^\perp$ es tipo espacio.

1.3. Vectores tipo tiempo

Para los vectores tipo tiempo se cumple una desigualdad tipo *Cauchy-Schwarz*, la cual nos permitirá deducir una nueva forma de calcular el producto de dos vectores tipo tiempo y una desigualdad del triángulo invertida. Para esto serán necesarios los siguientes resultados.

Proposición 1.26. Sea U un subespacio tipo tiempo. Entonces existe una base de vectores nulos para U .

Demostración. Ya que η restringida a U es no degenerada existe una base ortonormal $\{e_\alpha\}_{\alpha=0}^k$ de U . Los vectores de la forma $e_0 + e_i$ son nulos y linealmente independientes para $1 \leq i \leq k$; entonces existen k vectores nulos linealmente independientes en U .

Para obtener una base de vectores nulos para U el vector nulo $e_1 - e_0$ puede servir; sean a y b números reales, entonces si:

$$a(e_0 + e_1) + b(e_1 - e_0) = 0$$

$$(a - b)e_0 + (a + b)e_1 = 0$$

Por lo tanto $(a - b) = (a + b) = 0$ de donde se sigue que $a = b = 0$. Es claro que los vectores $e_0 + e_i$ no generan al vector $e_1 - e_0$, construyendo así una base de vectores nulos para U . □

Definición 1.27. Dado cualquier vector u en $\mathbb{R}_1^n$, se define la *norma* de este como $\|u\| := \sqrt{|\langle u, u \rangle|}$.

Teorema 1.28. Sean u y v vectores tipo tiempo en $\mathbb{R}_1^n$. Entonces:

$$|\langle u, v \rangle| \geq \|u\| \|v\|. \quad (1.6)$$

Donde la igualdad se logra cuando uno de los vectores es múltiplo escalar del otro. En caso de que los vectores se encuentren en la misma componente de $\mathcal{T}$ existe un único número $\varphi \geq 0$ tal que

$$\langle u, v \rangle = -\|u\| \|v\| \cosh \varphi.$$

Al número φ se le conoce como el ángulo hiperbólico entre los vectores u y v .

Demostración. Considérense dos vectores tipo tiempo u y v linealmente independientes y el plano $U = \text{span}(u, v)$ que generan. Por la Proposición 1.26 se tiene que la siguiente ecuación tiene dos soluciones distintas con a y b distintos de cero.

$$\langle au + bv, au + bv \rangle = a^2 \langle u, u \rangle + 2ab \langle u, v \rangle + b^2 \langle v, v \rangle = 0$$

Dividiendo la ecuación entre a^2 se obtiene la siguiente ecuación de segundo grado con $\lambda = \frac{b}{a}$.

$$\lambda^2 \langle v, v \rangle + 2\lambda \langle u, v \rangle + \langle u, u \rangle = 0$$

Como a y b son reales el discriminante de la forma general que soluciona la ecuación de segundo grado debe ser mayor que cero. Entonces:

$$\langle u, v \rangle^2 > \langle u, u \rangle \langle v, v \rangle > (-\langle u, u \rangle) (-\langle v, v \rangle)$$

Lo cual demuestra la desigualdad deseada; el caso de la igualdad es sencillo al proponer que uno de los vectores, digamos u , es múltiplo escalar de v .

Para la segunda parte de la demostración, usando la ecuación (1.6) y que u y v se encuentran en la misma componente de $\mathcal{T}$, se tiene lo siguiente.

$$\frac{-\langle u, v \rangle}{\|u\| \|v\|} \geq 1$$

Ya que la función $\cosh \varphi : [0, \infty) \rightarrow (1, \infty)$ es inyectiva, existe un único número $\varphi \in [0, \infty)$ tal que:

$$\cosh \varphi = \frac{-\langle u, v \rangle}{\|u\| \|v\|}.$$

Demostrando el teorema. □

Corolario 1.29. Si u y v son dos vectores tipo tiempo que se encuentran en la misma componente de $\mathcal{T}$, entonces

$$\|u + v\| \geq \|u\| + \|v\|$$

y la igualdad se da si y sólo si u y v son proporcionales.

Demostración. Usando la definición de norma y que ambos vectores se encuentran en la misma componente de $\mathcal{T}$ se tiene lo siguiente.

$$\begin{aligned} \|u + v\|^2 &= |\langle u + v, u + v \rangle| \\ &= \|u\|^2 + \|v\|^2 + 2|\langle u, v \rangle| \\ &= \|u\|^2 + \|v\|^2 + 2\|u\|\|v\|\cosh \varphi \end{aligned}$$

De ser u y v linealmente dependientes, por lo visto en el Teorema 1.28 $\varphi = 0$. Por lo tanto $\|u + v\|^2 = \|u\|^2 + \|v\|^2 + 2\|u\|\|v\|$ lo cual demuestra la parte de la igualdad. Si los vectores son linealmente independientes, $\cosh \varphi > 1$ y se tiene que $\|u + v\|^2 > \|u\|^2 + \|v\|^2 + 2\|u\|\|v\|$, lo cual demuestra el corolario. $\square$

En $\mathbb{R}_1^n$ se puede seguir hablando de la orientación de una base ya que es una cuestión que no depende de ningún producto interno, sin embargo, ahora se puede hablar de una segunda orientación que dependerá del producto interno. Por ejemplo, en la base canónica de $\mathbb{R}_1^n$ el vector e_0 *apunta* en la dirección en la que *avanza* el tiempo; este vector será siempre tomado como referencia, lo cual se ve en la siguiente definición.

Definición 1.30. Dado un vector v tipo tiempo, se dice que este está orientado al *futuro* si $\langle v, e_0 \rangle < 0$; se dirá que está orientado al *pasado* si $\langle v, e_0 \rangle > 0$.

Por lo tanto, dada una base ortonormal de $\mathbb{R}_1^n$ diremos que esta se encuentra orientada al futuro si su vector tipo tiempo se encuentra en la misma componente de $\mathcal{T}$ que e_0 , de lo contrario diremos que tal base está orientada al pasado.

1.4. Isometrías

Una *isometría* es una función que preserva la noción de *distancia* de un espacio a otro, es decir, si (X, d) y $(Y, \tilde{d})$ son *espacios métricos*, se dice que una biyección $\phi : X \rightarrow Y$ es una isometría si

$$d(x, y) = \tilde{d}(\phi(x), \phi(y)) \quad \text{para toda } x, y \in X.$$

Si $X = Y$ y $d = \tilde{d}$, entonces se dice que ϕ es una isometría de X . Dado que toda *norma* induce una *métrica*, si X es ahora un *espacio normado*, una

isometría es una biyección tal que $\|x - y\|_X = \|\phi(x) - \phi(y)\|_Y$ para toda $x, y \in X$.

Ahora, si V es un espacio vectorial con producto interno, este induce una norma de la siguiente forma

$$\|v\|_V := \langle v, v \rangle_e^{\frac{1}{2}}.$$

Por lo tanto, una isometría en V es una biyección tal que

$$\langle u - v, u - v \rangle_e = \langle \phi(u) - \phi(v), \phi(u) - \phi(v) \rangle_e.$$

En el caso del producto escalar con el cual se ha trabajado a lo largo de esta tesis es claro que este no induce alguna *norma* y, mucho menos una métrica; es por eso que existe la necesidad de rastrear la métrica hasta el producto interno.

Definición 1.31. Sea $\phi : \mathbb{R}_1^n \rightarrow \mathbb{R}_1^n$ una función biyectiva. Decimos que esta es una *isometría* de $\mathbb{R}_1^n$ si para toda $u, v \in \mathbb{R}_1^n$ se cumple que

$$\langle u - v, u - v \rangle = \langle \phi(u) - \phi(v), \phi(u) - \phi(v) \rangle.$$

Antes de seguir adelante, demostraré un lema que será de importancia en lo siguiente.

Lema 1.32. Sean u, v cualesquiera dos vectores en $\mathbb{R}_1^n$ tales que para toda $w \in \mathbb{R}_1^n$ se cumple que $\langle u, w \rangle = \langle v, w \rangle$. Entonces $u = v$.

Demostración. Restando el término en el lado derecho de la ecuación a toda la ecuación se obtiene

$$\langle u - v, w \rangle = 0 \quad \text{para toda } w \in \mathbb{R}_1^n.$$

Por lo tanto $u - v = 0$ ya que el producto interno es no degenerado, lo cual demuestra el lema. $\square$

A continuación se mostrará que toda isometría de $\mathbb{R}_1^n$ se puede descomponer en la suma de una transformación lineal con una traslación al igual que en el caso clásico. Para lograr esto definiré la función $\Lambda : \mathbb{R}_1^n \rightarrow \mathbb{R}_1^n$ como $\Lambda(u) := \phi(u) - \phi_0$, donde $\phi_0 := \phi(0)$ tal que ϕ es una isometría.

Para esta función se cumplen tres propiedades muy importantes respecto del producto escalar.

- $\langle u, u \rangle = \langle \Lambda(u), \Lambda(u) \rangle$ para toda $u \in \mathbb{R}_1^n$.

Para toda $u \in \mathbb{R}_1^n$ se tiene

$$\begin{aligned}\langle \Lambda(u), \Lambda(u) \rangle &= \langle \phi(u) - \phi_0, \phi(u) - \phi_0 \rangle \\ &= \langle u - 0, u - 0 \rangle \\ &= \langle u, u \rangle\end{aligned}$$

Esto quiere decir que la función Λ respeta la causalidad de los vectores, así como preserva la *norma* de estos.

- $\langle u - v, u - v \rangle = \langle \Lambda(u) - \Lambda(v), \Lambda(u) - \Lambda(v) \rangle$ para cualesquiera $u, v \in \mathbb{R}_1^n$.

Siendo u, v vectores cualesquiera en $\mathbb{R}_1^n$ se cumple que:

$$\begin{aligned}\langle \Lambda(u) - \Lambda(v), \Lambda(u) - \Lambda(v) \rangle &= \\ \langle \phi(u) - \phi_0 - (\phi(v) - \phi_0), \phi(u) - \phi_0 - (\phi(v) - \phi_0) \rangle &= \\ \langle \phi(u) - \phi(v), \phi(u) - \phi(v) \rangle &= \\ \langle u - v, u - v \rangle.\end{aligned}$$

Es claro que Λ es una función biyectiva ya que ϕ lo es. Por lo tanto Λ es una isometría también.

- $\langle u, v \rangle = \langle \Lambda(u), \Lambda(v) \rangle$ para todo $u, v \in \mathbb{R}_1^n$

Para toda $u, v \in \mathbb{R}_1^n$ se tiene

$$\begin{aligned}\langle \Lambda(u) - \Lambda(v), \Lambda(u) - \Lambda(v) \rangle &= \langle \Lambda(u), \Lambda(u) \rangle - 2\langle \Lambda(u), \Lambda(v) \rangle + \langle \Lambda(v), \Lambda(v) \rangle \\ &= \langle u, u \rangle - 2\langle \Lambda(u), \Lambda(v) \rangle + \langle v, v \rangle\end{aligned}$$

Por otro lado

$$\langle u - v, u - v \rangle = \langle u, u \rangle - 2\langle u, v \rangle + \langle v, v \rangle$$

Lo cual demuestra la igualdad deseada. Este último punto junto con el teorema siguiente nos llevará a un resultado muy familiar.

Teorema 1.33. La función $\Lambda : \mathbb{R}_1^n \rightarrow \mathbb{R}_1^n$ que se define como $\Lambda(u) := \phi(u) - \phi_0$, es una isometría lineal.

Demostración. Como se ha visto hasta ahora la función Λ es una biyección; entonces, para cualquier $w \in \mathbb{R}_1^n$ se tiene que existe un único $\tilde{w}$ tal que $\Lambda(\tilde{w}) = w$. Entonces, siendo c cualquier número real se tiene que

$$\langle c\Lambda(u), w \rangle = c\langle \Lambda(u), \Lambda(\tilde{w}) \rangle = c\langle u, \tilde{w} \rangle$$

Y por otro lado se obtiene que

$$\langle \Lambda(cu), w \rangle = \langle \Lambda(cu), \Lambda(\tilde{w}) \rangle = c\langle u, \tilde{w} \rangle$$

Resultando que $\Lambda(cu) = c\Lambda(u)$ por el Lema 1.32. Para ver que Λ abre las sumas se procede de la misma forma

$$\langle \Lambda(u + v), w \rangle = \langle \Lambda(u + v), \Lambda(\tilde{w}) \rangle = \langle u + v, \tilde{w} \rangle$$

$$\begin{aligned} \langle \Lambda(u) + \Lambda(v), w \rangle &= \langle \Lambda(u) + \Lambda(v), \Lambda(\tilde{w}) \rangle \\ &= \langle \Lambda(u), \Lambda(\tilde{w}) \rangle + \langle \Lambda(v), \Lambda(\tilde{w}) \rangle \\ &= \langle u, \tilde{w} \rangle + \langle v, \tilde{w} \rangle \\ &= \langle u + v, \tilde{w} \rangle \end{aligned}$$

Una vez más por el Lema 1.32 se obtiene que $\Lambda(u + v) = \Lambda(u) + \Lambda(v)$. Por lo tanto Λ es una isometría lineal. $\square$

Ya que Λ respeta el producto escalar y es lineal, Λ toma una base ortonormal y la manda en otra base ortonormal, justo como los operadores ortogonales en el caso clásico. Esto quiere decir que las columnas de la matriz asociada a Λ , como vectores de $\mathbb{R}_1^n$, son vectores unitarios y ortogonales; en este caso la inversa de Λ no será Λ^T como se verá más adelante.

El siguiente corolario se sigue inmediatamente del teorema anterior.

Corolario 1.34. Dada una isometría ϕ de $\mathbb{R}_1^n$, existen una transformación lineal Λ y un vector v_0 tales que

$$\phi(u) = \Lambda(u) + v_0 \quad \text{para toda } u \in \mathbb{R}_1^n.$$

Al conjunto de todas las isometrías del espacio de Minkowski se le conoce como el *grupo de Poincaré*. A lo largo de la tesis sólo se estudiará el conjunto de las isometrías lineales al cual se nombra *grupo de Lorentz*. Es claro que el conjunto de las isometrías del espacio de Minkowski tiene estructura de grupo con la composición de funciones; ya que la composición de funciones lineales es también una función lineal, el conjunto de las isometrías lineales es a su vez un grupo al cual se denotará como $O_1(n)$, donde 1 es el índice del producto escalar y n la dimensión del espacio.

Expresando el producto escalar como producto de matrices, si utilizamos que $\langle u, v \rangle = \langle \Lambda(u), \Lambda(v) \rangle$ para toda $\Lambda \in O_1(n)$ se obtiene lo siguiente.

$$u^T \eta v = (\Lambda u)^T \eta (\Lambda v) = u^T (\Lambda^T \eta \Lambda) v \quad \text{para todo } u, v \in \mathbb{R}_1^n$$

Por lo tanto $\eta = \Lambda^T \eta \Lambda$; usando las propiedades del determinante se obtiene finalmente que $\det \Lambda = \pm 1$. Entonces se deduce que $O_1(n)$ tiene al menos dos componentes disjuntas.

Sea $\Lambda \in O_1(n)$ tal que su determinante es positivo; usando la notación de índices se tiene lo siguiente.

$$\Lambda_\lambda^\mu \eta_{\mu\nu} \Lambda_\sigma^\nu = \eta_{\lambda\sigma} \quad (1.7)$$

Si $\lambda = \sigma = 0$

$$\begin{aligned} \Lambda_0^\mu \eta_{\mu\nu} \Lambda_0^\nu &= \eta_{00} \\ -(\Lambda_0^0)^2 + (\Lambda_0^1)^2 + \cdots + (\Lambda_0^n)^2 &= -1 \end{aligned}$$

Entonces, $(\Lambda_0^0)^2$ debe ser mayor o igual a -1 . Por lo tanto

$$\Lambda_0^0 \geq 1 \quad \text{ó} \quad \Lambda_0^0 \leq -1$$

Lo cual muestra que la componente de isometrías lineales con determinante positivo tiene dos componentes disjuntas. Lo mismo pasará para aquellas que tienen determinante negativo. Por lo tanto $O_1(n)$ tiene cuatro componentes disjuntas las cuales se denotan como sigue.

$$\begin{aligned} \Lambda \in O_1(n)^{++} & \quad \text{si} \quad \det \Lambda = 1 \quad \text{y} \quad \Lambda_0^0 \geq 1 \\ \Lambda \in O_1(n)^{+-} & \quad \text{si} \quad \det \Lambda = 1 \quad \text{y} \quad \Lambda_0^0 \leq -1 \\ \Lambda \in O_1(n)^{-+} & \quad \text{si} \quad \det \Lambda = -1 \quad \text{y} \quad \Lambda_0^0 \geq 1 \\ \Lambda \in O_1(n)^{--} & \quad \text{si} \quad \det \Lambda = -1 \quad \text{y} \quad \Lambda_0^0 \leq -1 \end{aligned}$$

Teorema 1.35. El grupo de Lorentz $O_1(n)$ con $n \geq 2$, tiene cuatro componentes conexas.

Este teorema se demostrará más adelante en la tesis ya que se necesitará de mayor herramienta a la expuesta hasta ahora; sin embargo, podemos examinar lo que pasa cuando la dimensión del espacio de Minkowski es igual a dos.

Ejemplo 1.36. El grupo $O_1(2)$ tiene cuatro componentes conexas.

Fijémonos en la componente $O_1(2)^{++}$ del grupo de Lorentz. Todos los elementos de esta componente cumplen con lo siguiente; sean Λ_β^α las componentes de una representación matricial de Λ .

$$\Lambda_0^0 \Lambda_1^1 - \Lambda_1^0 \Lambda_0^1 = \det \Lambda = 1$$

Y por la ecuación (1.7)

$$-(\Lambda_0^0)^2 + (\Lambda_0^1)^2 = -1 \quad (1.8)$$

$$\Lambda_1^0 \Lambda_0^0 = \Lambda_1^1 \Lambda_0^1 \quad (1.9)$$

Si uno despeja $(\Lambda_0^1)^2$ en la ecuación (1.8), al sustituir este valor en la ecuación (1.9) multiplicada por Λ_0^1 se obtiene

$$\Lambda_0^1 \Lambda_1^0 \Lambda_0^0 = -\Lambda_1^1 + \Lambda_1^1 (\Lambda_0^0)^2$$

Como $\Lambda_0^0 \geq 1$ se puede dividir la ecuación anterior por este número resultando

$$\Lambda_1^0 \Lambda_0^1 = \frac{-\Lambda_1^1}{\Lambda_0^0} + \Lambda_1^1 \Lambda_0^0$$

Usando la hipótesis de que $\det \Lambda = 1$ se obtiene finalmente que $\Lambda_0^0 = \Lambda_1^1$. Análogamente se puede obtener que $\Lambda_1^0 = \Lambda_0^1$ suponiendo que Λ_1^0 ó Λ_0^1 son distintos de cero; digamos que $\Lambda_0^1 = 0$, entonces, por la ecuación (1.9) se obtiene que $\Lambda_1^0 = 0$ también.

Por lo tanto, existe $\varphi \in \mathbb{R}$ tal que

$$\Lambda_\varphi = \begin{pmatrix} \cosh \varphi & \sinh \varphi \\ \sinh \varphi & \cosh \varphi \end{pmatrix}$$

Para toda $\Lambda \in O_1(2)^{++}$. Para cada φ , a este tipo de matrices se les da el nombre de *boost* de $\mathbb{R}_1^2$ por un ángulo de *Lorentz* φ . En general se le llama *boost* a cualquier elemento de $O_1(n)^{++}$.

Ahora bien, definiremos la función $h(\varphi) : \mathbb{R} \rightarrow O_1(2)^{++}$ como $h(\varphi) = \Lambda_\varphi$. Haciendo uso de las identidades

$$\cosh(\varphi + \theta) = \cosh \varphi \cosh \theta + \sinh \varphi \sinh \theta$$

$$\sinh(\varphi + \theta) = \sinh \varphi \cosh \theta + \cosh \varphi \sinh \theta$$

se puede notar que h es un homeomorfismo de grupos. Además, dado que las funciones $\cosh$ y $\sinh$ son par e impar respectivamente, las identidades pasadas aplicadas a $\varphi - \theta$ nos permiten demostrar la inyectividad de h ; por lo dicho con anterioridad es obvio que esta función es sobre, y así, h es un isomorfismo.

Por otro lado, la función h determina una curva en $O_1(2)^{++} \subset M_{2 \times 2}(\mathbb{R}) \simeq \mathbb{R}^4$; de momento asumiré la existencia de una estructura diferenciable sobre el grupo de Lorentz, al final del tercer capítulo regresaré sobre esto.

Siendo así, h tiene una diferencial definida en todo su dominio

$$d_{\theta_0} h(c) = c \begin{pmatrix} \sinh \theta_0 & \cosh \theta_0 \\ \cosh \theta_0 & \sinh \theta_0 \end{pmatrix}$$

La diferencial $d_{\theta_0} h$ no es la función cero para toda $\theta_0 \in \mathbb{R}$. Por lo tanto h es un difeomorfismo, lo cual demuestra que $O_1(2)^{++}$ es conexo.

Sean $P := \text{diag}\{1, -1\}$ y $T := \text{diag}\{-1, 1\}$ los operadores de *inversión espacial* e *inversión temporal* respectivamente; se puede notar que $PT \in O_1(2)^{+-}$, $P \in O_1(2)^{-+}$, $T \in O_1(2)^{--}$. Con esto último, las restantes tres componentes de $O_1(2)$ se pueden expresar como cosets de $O_1(2)^{++}$ con el uso de los operadores PT, P, T .

Entonces se puede definir para el caso del conjunto $O_1(2)^{-+}$ la función $h_P : \mathbb{R} \rightarrow O_1(2)^{-+}$ como $h_P(\varphi) = (P)(h(\varphi)) = P\Lambda_\varphi$, la cual de manera análoga resultará ser un difeomorfismo. Prosiguiendo con funciones h_{PT} y h_T , definidas análogamente a h_P , se llega al mismo resultado para las restantes componentes de $O_1(2)$ demostrando así que el *grupo de Lorentz* de $\mathbb{R}_1^2$ tiene cuatro componentes conexas.

En general, el *grupo de Lorentz* $O_1(n)$ con la topología inducida de $\mathbb{R}^{n^2}$ no es compacto ya que la familia de elementos de la forma

$$\left(\begin{array}{c|c} \Lambda_\varphi & 0 \\ \hline 0 & I_{n-2} \end{array} \right)$$

Donde Λ_φ es la matriz de 2×2 descrita en la demostración anterior y la matriz I_{n-2} es la identidad de $(n-2) \times (n-2)$, no es acotada.

1.5. Operadores adjuntos y auto-adjuntos

Como se ha visto en el desarrollo de este capítulo, muchas de las propiedades que se tenían en un espacio vectorial con *producto interno* sobreviven en el caso de un producto escalar de índice uno, salvo ciertas modificaciones. El caso de los operadores *adjuntos* y *auto-adjuntos* en $\mathbb{R}_1^n$ sigue una línea similar, en general se seguirán valiendo gran parte de los teoremas del caso clásico salvo ciertas particularidades que se presentarán a lo largo de esta sección.

Teorema 1.37. Sea $g : \mathbb{R}_1^n \rightarrow \mathbb{R}$ una función lineal. Entonces, existe un único vector $v \in \mathbb{R}_1^n$ tal que $g(u) = \langle u, v \rangle$ para toda $u \in \mathbb{R}_1^n$.

Demostración. Sea $\{e_0, \dots, e_{n-1}\}$ una base ortonormal para $\mathbb{R}_1^n$ y definimos $v \in \mathbb{R}_1^n$ como sigue

$$v = -g^0 e_0 + g^i e_i \quad ; \text{ donde } g^\alpha := g(e_\alpha)$$

Definamos ahora la función $f : \mathbb{R}_1^n \rightarrow \mathbb{R}$ como $f(u) := \langle u, v \rangle$. Claramente esta función es lineal; además, si se valúa f en la base se obtiene que

$$\begin{aligned} f(e_\alpha) &= \langle e_\alpha, v \rangle \\ &= \langle e_\alpha, -g^0 e_0 + g^i e_i \rangle \\ &= -\langle e_\alpha, g^0 e_0 \rangle + \langle e_\alpha, g^i e_i \rangle \end{aligned}$$

Entonces, si $\alpha = 0$ se tiene que $f(e_0) = g(e_0)$, si α es mayor o igual que 1 se tiene que $f(e_i) = g(e_i)$; así, f y g coinciden en una base. Por lo tanto $f = g$.

Para ver que v es único supongamos que existe otro vector $\tilde{v} \in \mathbb{R}_1^n$ tal que $g(u) = \langle u, \tilde{v} \rangle$; entonces $\langle u, v \rangle = \langle u, \tilde{v} \rangle$ para toda $u \in \mathbb{R}_1^n$, y por el Lema 1.32, $v = \tilde{v}$. $\square$

Definición 1.38. Sea $T : \mathbb{R}_1^n \rightarrow \mathbb{R}_1^n$ un operador lineal. Se define el operador adjunto $T^* : \mathbb{R}_1^n \rightarrow \mathbb{R}_1^n$, de T como aquel que cumple

$$\langle T(u), v \rangle = \langle u, T^*(v) \rangle$$

Para todo $u, v \in \mathbb{R}_1^n$.

Como se puede observar la definición es la misma que en el caso de un *producto interno*, aunque falta ver que T^* es en realidad lineal y único.

Teorema 1.39. Sea T un operador lineal sobre $\mathbb{R}_1^n$. Entonces, el operador adjunto de T , T^* , es lineal y el único que cumple con la relación

$$\langle T(u), v \rangle = \langle u, T^*(v) \rangle ; \text{ para todo } u, v \in \mathbb{R}_1^n.$$

Demostración. Sea v un vector en $\mathbb{R}_1^n$ y $g : \mathbb{R}_1^n \rightarrow \mathbb{R}$ la función definida por $g(u) := \langle T(u), v \rangle$; dado que T es lineal y el producto escalar en una *forma bilineal*, la función g así definida es lineal.

Ahora, por el Teorema 1.37 existe un único vector de $\mathbb{R}_1^n$, w , tal que $\langle T(u), v \rangle = \langle u, w \rangle$ para todo $u \in \mathbb{R}_1^n$. Definamos ahora al operador T^* como aquel que toma v y lo lleva a w para todo $v \in \mathbb{R}_1^n$, entonces $\langle T(u), v \rangle = \langle u, T^*(v) \rangle$ para toda $u, v \in \mathbb{R}_1^n$.

Sean $v_1, v_2 \in \mathbb{R}_1^n$ y $c \in \mathbb{R}$, entonces

$$\begin{aligned} \langle u, T^*(cv_1 + v_2) \rangle &= \langle T(u), cv_1 + v_2 \rangle \\ &= c\langle T(u), v_1 \rangle + \langle T(u), v_2 \rangle \\ &= \langle u, cT^*(v_1) \rangle + \langle u, T^*(v_2) \rangle \\ &= \langle u, cT^*(v_1) + T^*(v_2) \rangle \end{aligned}$$

para toda $u \in \mathbb{R}_1^n$ y, por el Lema 1.32, se deduce que T^* es un operador lineal.

Para ver que T^* es único, supongamos que existe otro operador $U : \mathbb{R}_1^n \rightarrow \mathbb{R}_1^n$ tal que cumple con la propiedad de ser el adjunto de T , entonces

$$\langle u, T^*(v) \rangle = \langle u, U(v) \rangle$$

para toda $u, v \in \mathbb{R}_1^n$, y una vez más por el Lema 1.32 se obtiene que $T = U$. Por lo tanto, el operador T^* es lineal y es el único que cumple con ser el adjunto de T . $\square$

Además, usando la simetría del producto escalar se tiene lo siguiente:

$$\langle u, T(v) \rangle = \langle T(v), u \rangle = \langle v, T^*(u) \rangle = \langle T^*(u), v \rangle.$$

Con esto, la operación de *adjuntar* no depende del orden y solo depende del producto escalar.

Como se vio en la demostración del teorema anterior, la existencia del operador adjunto dependió por completo de la finitud de la dimensión de $\mathbb{R}_1^n$; en el caso que la dimensión no sea finita no se puede asegurar siempre la existencia de tal operador, pero eso es un detalle algebraico que no investigaré en esta tesis.

Sea A la representación matricial de de una transformación lineal $T : \mathbb{R}_1^n \rightarrow \mathbb{R}_1^n$ y B la representación matricial de T^* en una base ortogonal; entonces $\langle Au, v \rangle = \langle u, Bv \rangle$ para toda $u, v \in \mathbb{R}_1^n$. Por lo tanto $A^T \eta = \eta B$. Con esto uno tiene una forma clara de saber quién es el operador adjunto T^* , ya que su matriz asociada es $B = \eta A^T \eta$ en una base ortonormal.

En el caso de tener un *producto interno*, ya que la matriz asociada al producto en una base ortonormal resulta ser la matriz identidad, hubiera bastado con transponer A para obtener B . Si defino $A^* := \eta A^T \eta$ el siguiente teorema queda demostrado en una notación más simple.

Teorema 1.40. Sea $T : \mathbb{R}_1^n \rightarrow \mathbb{R}_1^n$ una transformación lineal y, $[T]_\beta$ su representación matricial en una base ortonormal β . Entonces:

$$[T^*]_\beta = [T]_\beta^*.$$

Si A es una matriz de $n \times n$, la transformación lineal asociada a A en una base ortonormal se denota como L_A . Entonces, el siguiente corolario se demuestra trivialmente ya que $[L_A]_\beta = A$.

Corolario 1.41. Sea A una matriz de $n \times n$. Entonces $L_{A^*} = (L_A)^*$.

A continuación demostraré una serie de resultados básicos sobre operadores auto-adjuntos que son de mucha utilidad.

Teorema 1.42. Sean T y U operadores lineales sobre $\mathbb{R}_1^n$. Entonces se cumplen las siguientes propiedades:

1. $(T + U)^* = T^* + U^*$
2. $(cT)^* = cT^*$ para cualquier $c \in \mathbb{R}$
3. $(TU)^* = U^*T^*$
4. $T^{**} = T$

5. $I^* = I$ donde I es el operador identidad

Demostración. Los cinco puntos se demuestran de forma similar por lo que solo se demostrará con detalle el primero, mientras que en los puntos restantes solo se mostrará parte de la demostración. Sean $u, v \in \mathbb{R}_1^n$.

1.

$$\begin{aligned}
 \langle u, (T + U)^*(v) \rangle &= \langle (T + U)(u), v \rangle \\
 &= \langle T(u) + U(u), v \rangle \\
 &= \langle T(u), v \rangle + \langle U(u), v \rangle \\
 &= \langle u, T^*(v) \rangle + \langle u, U^*(v) \rangle \\
 &= \langle u, T^*(v) + U^*(v) \rangle
 \end{aligned}$$

Para todo $u, v \in \mathbb{R}_1^n$; entonces, haciendo uso del Lema 1.32 se obtiene que $(T + U)^* = T^* + U^*$.

2.

$$\langle u, (cT)^*(v) \rangle = c\langle T(u), v \rangle = \langle u, cT^*(v) \rangle$$

3.

$$\langle u, (TU)^*(v) \rangle = \langle T(U(u)), v \rangle = \langle U(u), T^*(v) \rangle$$

4.

$$\langle u, (T^*)^*(v) \rangle = \langle T^*(u), v \rangle = \langle u, T(v) \rangle$$

5.

$$\langle u, I^*(v) \rangle = \langle u, v \rangle = \langle u, I(v) \rangle$$

□

Corolario 1.43. Sean A y B matrices de $n \times n$. Entonces:

1. $(A + B)^* = A^* + B^*$
2. $(cA)^* = cA^*$ para cualquier $c \in \mathbf{R}$
3. $(AB)^* = B^*A^*$
4. $A^{**} = A$
5. $A^* = A$ donde A es la matriz identidad

Este corolario es una consecuencia directa del teorema anterior y del Corolario 1.41, por lo cual se omitirá la demostración.

Antes de dar la definición de lo que es un operador *auto-adjunto*, se necesitarán de algunos resultados previos referentes a los *vectores propios* de un operador y su adjunto.

Lema 1.44. Sea T un operador lineal sobre $\mathbb{R}_1^n$. Si T tiene un vector propio no nulo, entonces T^* también tiene un vector propio no nulo.

Demostración. Sea v un vector propio de T con *valor propio* λ . Entonces, para todo $w \in \mathbb{R}_1^n$ se cumple lo siguiente

$$0 = \langle 0, w \rangle = \langle (T - \lambda I)v, w \rangle = \langle v, (T - \lambda I)^* w \rangle = \langle v, (T^* - \lambda I)w \rangle$$

Por lo tanto w es ortogonal a la imagen de $(T^* - \lambda I)$; con esto, la imagen de $(T^* - \lambda I)$ es un subconjunto de $\langle v \rangle^\perp$ el cual tiene dimensión $n - 1$, y por el teorema de la dimensión se tiene que el núcleo de $(T^* - \lambda I)$ no es trivial. Cualquier elemento del núcleo es un vector propio de $(T^* - \lambda I)$ con valor propio λ ; dado que v no es nulo este no se encuentra en la imagen de $(T^* - \lambda I)$ y es un elemento del núcleo, demostrando el lema. □

A continuación demostraré una versión del *teorema de Schur* para nuestro caso, pero antes necesitare de otro teorema el cual solo enunciaré.

Teorema 1.45. Sea T un operador lineal sobre un espacio vectorial de dimensión finita V y, sea W un subespacio T -invariante de V . Entonces, el *polinomio característico* de $T|_W$ divide al *polinomio característico* de T .

Recordando, se dice que un espacio es T -invariante si la imagen de $T|_W$ está contenida en W . Por otro lado, para el teorema que sigue es necesario recordar que un polinomio $P(t)$ se *factoriza* si este se escribe en términos de la forma $(\lambda - t)$; ahora se tiene lo necesario para llevar a cabo una demostración del *teorema de Schur* en $\mathbb{R}_1^n$.

Teorema 1.46. Sea T un operador lineal sobre $\mathbb{R}_1^n$ tal que no tiene vectores propios nulos. Supongamos que el polinomio característico de T se *factoriza*, entonces existe una base ortonormal β tal que en ella $[T]_\beta$ es triangular superior.

Demostración. La demostración se elaborará sobre la dimensión de $\mathbb{R}_1^n$. Empezaré con el caso $n = 2$; ya que el polinomio característico de T se factoriza, por el Lema 1.44 puedo suponer que T^* tiene un vector propio unitario v ; ya que este no es nulo, su complemento ortogonal tampoco lo será, el cual

tiene dimensión uno y $v \notin \langle v \rangle^\perp = W$. Si tomo $v^\perp \in W$ unitario, el conjunto $\beta = \{v, v^\perp\}$ es una base ortonormal para $\mathbb{R}_1^2$. Además, para todo $u \in W$ y $c \in \mathbb{R}$ se tiene lo siguiente

$$\langle T(u), cv \rangle = \langle u, cT^*(v) \rangle = \langle u, c\lambda v \rangle = c\lambda \langle u, v \rangle = 0 \quad (1.10)$$

Por lo tanto, W es un subespacio T -invariante lo cual demuestra que $[T]_\beta$ es triangular superior.

Ahora puedo suponer cierto el resultado para operadores lineales que cumplan con las hipótesis en un espacio de dimensión $n - 1$. Por el Lema 1.44 puedo tomar una vez más un vector propio $v \in \mathbb{R}_1^n$ de T^* unitario; al igual que en la base de inducción, llamaré W al complemento ortogonal de v y por el mismo argumento expresado en la ecuación (1.10) se tiene que W es T -invariante. Ahora, por el Teorema 1.45 el polinomio característico de $T|_W$ divide al polinomio característico de T , y como este último se factoriza, el primero se factoriza también; ya que $\dim W = n - 1$ por el Lema 1.23 y $T|_W$ no tiene vectores propios nulos, existe una base ortonormal γ en la cual $[T|_W]_\gamma$ es triangular superior. Esto se sigue de la hipótesis de inducción cuando v es tipo espacio y $W \simeq \mathbb{R}_1^{n-1}$; si v es tipo tiempo, $W \simeq \mathbb{R}^{n-1}$ y la afirmación anterior se sigue de utilizar el *teorema de Schur* clásico.

Por lo tanto, el conjunto $\beta = \gamma \cup v$ es una base ortonormal de $\mathbb{R}_1^n$ en la cual, al ser W un espacio T -invariante, $[T]_\beta$ es triangular superior. $\square$

Definición 1.47. sea T un operador lineal sobre $\mathbb{R}_1^n$. Decimos que T es *Auto-adjunto* si $T = T^*$; a su vez, se dice que una matriz A de $n \times n$ es *Auto-adjunta* si $A = A^*$.

La mayoría de las propiedades de los operadores auto-adjuntos son claras cuando el campo es el de los números complejos, así que por un momento *complejizaremos* los operadores lineales a los que nos refiramos.

Pero el espacio en el que hemos estado trabajando no es el clásico $\mathbb{R}^n$ si no el espacio de Minkowski y, por eso debemos llevar al caso complejo el producto escalar. La forma de llevar un *producto interno* al espacio $\mathbb{C}^n$ es a través de un *producto Hermitiano*, entonces, el equivalente al producto Lorentziano se define de la siguiente manera

$$\langle u, v \rangle_H := -u^0 \bar{v}^0 + \dots + u^{n-1} \bar{v}^{n-1} \text{ para todo } u, v \in \mathbb{C}^n$$

Donde la barra significa como siempre conjugación compleja; entonces, la pareja $(\mathbb{C}, \langle \cdot, \cdot \rangle_H) \approx \mathbb{C}_1^n$ es la versión compleja del espacio de Minkowski. A la matriz asociada en este caso le llamaremos η_H . Además, viendo el producto

en forma matricial este se expresaría como $u^T \eta \bar{v}$, entonces, si A es una matriz con coeficientes en los complejos se tiene lo siguiente:

$$(Au)^T \eta \bar{v} = u^T A^T \eta \bar{v} = \langle Au, v \rangle = \langle u, A^* v \rangle = u^T \eta \overline{A^* v} = u^T \eta \bar{A}^* \bar{v}$$

Para toda $u, v \in \mathbb{C}_1^n$. Entonces $A^T \eta = \eta \bar{A}^*$ y $A^* = \eta \bar{A}^T \eta$. Por lo tanto, en este contexto adjuntar un operador lineal se asocia no solo a la operación de transponer, sino también a la de conjugar.

Ya con esta adecuación, es necesario volver a demostrar el punto número dos del Teorema 1.42, cabe mencionar que la definición de operador adjunto no cambia para este caso. Entonces, sea $c \in \mathbb{C}$ y T un operador sobre $\mathbb{C}_1^n$

$$\langle u, (cT)^*(v) \rangle_H = \langle cT(u), v \rangle_H = c \langle u, T^*(v) \rangle_H = \langle u, \bar{c} T^*(v) \rangle_H$$

para todo $u, v \in \mathbb{C}_1^n$. Por lo tanto, $(cT)^* = \bar{c}T$ para todo $c \in \mathbb{C}$.

Teorema 1.48. Sea T un operador lineal sobre $\mathbb{C}_1^n$ tal que $TT^* = T^*T$, entonces las siguientes afirmaciones son verdaderas:

1. $\langle T(u), T(u) \rangle = \langle T^*(u), T^*(u) \rangle$ para todo $u \in \mathbb{C}_1^n$.
2. Si $U = T - cI$, entonces $UU^* = U^*U$ para toda $c \in \mathbb{C}$.

Demostración. 1. Para todo $u \in \mathbb{C}_1^n$ se cumple lo siguiente

$$\langle T(u), T(u) \rangle = \langle T^*T(u), u \rangle = \langle TT^*(u), u \rangle = \langle T^*(u), T^*(u) \rangle$$

Lo cual demuestra el enunciado.

2. Sea $v \in \mathbb{C}_1^n$ un vector arbitrario, entonces

$$\begin{aligned} UU^*(v) &= (T - cI)(T^*(v) - \bar{c}I(v)) \\ &= (T - cI)(T^*(v)) - (T - cI)(\bar{c}v) \\ &= TT^*(v) - cT^*(v) - \bar{c}T(v) + c\bar{c}v \end{aligned}$$

$$\begin{aligned} U^*U(v) &= (T^* - \bar{c}I)(T(v) - cI(v)) \\ &= (T^* - \bar{c}I)(T(v)) - (T^* - \bar{c}I)(cv) \\ &= T^*T(v) - \bar{c}T(v) - cT^*(v) + c\bar{c}v \end{aligned}$$

Como T y T^* conmutan se tiene finalmente que $UU^* = U^*U$. □

Teorema 1.49. Sea T un operador lineal sobre $\mathbb{C}_1^n$ tal que conmuta con T^* y no tiene vectores propios nulos. Entonces se cumplen los siguientes enunciados:

1. Si v es vector propio de T , entonces también es vector propio de T^* . Además, si $T(v) = \lambda v$, entonces $T^*(v) = \bar{\lambda}v$.
2. Sean λ_1 y λ_2 valores propios distintos con vectores propios v_1 y v_2 respectivamente; entonces v_1 y v_2 son ortogonales entre si.

Demostración. 1. Sea v un vector propio de T con valor propio λ y U el operador lineal definido como $U := (T - \lambda I)$; ya que v es vector propio de T , se tiene que $U(v) = 0$ y $U^*U(v) = 0$ y por el Teorema 1.48 $UU^*(v) = 0$, lo cual quiere decir que $U^*(v)$ es vector propio de T ó $U^*(v) = 0$.

Entonces, por el número 1 del Teorema 1.48 se tiene lo siguiente

$$\langle U(v), U(v) \rangle = \langle U^*(v), U^*(v) \rangle = 0$$

Por lo tanto, como T no tiene vectores propios nulos, si $U^*(v)$ es vector propio de T por la ecuación anterior se tiene que $U^*(v) = 0$, lo cual implica que v es vector propio de T^* con valor propio $\bar{\lambda}$.

2. Sean λ_1 y λ_2 dos valores propios distintos de T con sus correspondientes vectores propios v_1 y v_2 . Entonces, usando 1 se tiene lo siguiente

$$\begin{aligned} \lambda_1 \langle v_1, v_2 \rangle &= \langle \lambda_1 v_1, v_2 \rangle = \langle T(v_1), v_2 \rangle \\ &= \langle v_1, T^*(v_2) \rangle = \langle v_1, \bar{\lambda}_2 v_2 \rangle \\ &= \lambda_2 \langle v_1, v_2 \rangle \end{aligned}$$

Ya que los valores propios son distintos por hipótesis, se tiene que v_1 y v_2 son ortogonales necesariamente. □

En el siguiente lema se hace evidente por qué la necesidad de llevar nuestro problema al espacio $\mathbb{C}_1^n$; en principio, un operador lineal en un espacio vectorial real solo tiene valores propios reales y el polinomio característico no se *abre* necesariamente, lo cual como, se ha visto, no se espera que suceda.

Lema 1.50. Sea T un operador auto-adjunto sobre un espacio vectorial de dimensión finita V_1 tal que no tiene vectores propios nulos. Entonces:

1. Todos los valores propios de T son reales.
2. Si $V_1 \simeq \mathbb{R}_1^n$, es decir, es un espacio vectorial real de índice uno, el polinomio característico de T se factoriza.

Demostración. 1. Sea $v \in V_1$ un vector propio de T distinto de cero. Ya que T es auto-adjunto, claramente conmuta con T^* ; por el número 1 del Teorema 1.49 se tiene lo siguiente

$$\lambda v = T(v) = T^*(v) = \bar{\lambda}v$$

Por lo tanto $\lambda = \bar{\lambda}$, lo cual implica que $\lambda \in \mathbb{R}$.

2. Sea $n = \dim V_1$, β una base ortonormal de V_1 y T un operador auto-adjunto sobre V_1 . Entonces, la matriz $A = [T]_\beta$ es auto-adjunta.

Sea T_A el operador lineal sobre $\mathbb{C}_1^n$ definido como $T_A(z) := Az$ para todo $z \in \mathbb{C}_1^n$. Siendo γ la base ortonormal canónica de $\mathbb{C}_1^n$ (que es la misma que β), se tiene claramente que T_A es auto-adjunto ya que $[T_A]_\gamma = A$; entonces por 1 se tiene que los valores propios de T_A son todos reales. Por lo tanto, haciendo uso del *teorema fundamental del álgebra* se tiene que el polinomio característico de T_A se factoriza y, ya que sus valores propios son reales, el polinomio característico de T_A se factoriza sobre los reales.

Como el polinomio característico de T_A es el mismo que el de A , y a su vez el mismo que el de T , podemos afirmar que el polinomio característico de T se factoriza sobre los reales.

□

Teorema 1.51. Sea T un operador lineal sobre $\mathbb{R}_1^n$ tal que no tiene vectores propios nulos. Entonces, T es auto-adjunto si y sólo si existe una base ortonormal β de $\mathbb{R}_1^n$ tal que los elementos de la base son vectores propios de T .

Demostración. Primero supongamos que T es auto-adjunto, como este no tiene valores propios nulos por el Lema 1.50 puedo aplicar el Teorema 1.46 para obtener una base ortonormal β de $\mathbb{R}_1^n$ tal que $A = [T]_\beta$ es triangular superior. Entonces

$$A = [T]_\beta = [T^*]_\beta = [T]_\beta^* = A^*$$

es decir, al ser A auto-adjunta A^* será triangular superior también lo cual hace que A sea diagonal. Siendo T diagonal en esta base se tiene que $T(\beta_i) = a_i\beta_i$ para todo $\beta_i \in \beta$ ya que de no ser así T no sería diagonal. Por lo tanto la base β está formada por vectores propios de T .

Ahora voy a suponer que existe una base ortonormal β que consta de vectores propios de T , claramente se está suponiendo también que este operador tiene vectores propios no nulos. Como los vectores propios de T forman una base de $\mathbb{R}_1^n$ el polinomio característico de T se factoriza ya que de no

ser así, al complejificar el problema podrían aparecer más vectores propios de T tal que se tendría un subconjunto linealmente independiente de más de n vectores en $\mathbb{C}_1^n$, lo cual sería una clara contradicción. Además, al ser β base y el polinomio característico factorizarse, se tiene que la multiplicidad de cada valor propio iguala la dimensión de su respectivo subespacio propio. Con esto se puede afirmar que T es diagonalizable y por tanto auto-adjunto.

□

Pero ¿qué sucede cuando un operador lineal es auto-adjunto y posee vectores propios nulos?

Como se puede ver este caso resulta diferente del caso clásico donde todo operador auto-adjunto es diagonalizable. En el caso del espacio de Minkowski un operador auto-adjunto tiene cuatro representaciones canónicas; dos de ellas en una base ortonormal, las dos restantes en lo que a continuación definiré como una base *pseudo-ortonormal*.

Definición 1.52. Sea β una base de $\mathbb{R}_1^n$. Se dice que β es una base *pseudo-ortonormal* si esta posee dos vectores nulos u_1 y u_2 tales que $\langle u_1, u_2 \rangle = -1$,

$$\langle u_1, u_j \rangle = \langle u_2, u_j \rangle = 0 \text{ ; para } 3 \leq j \leq n,$$

y $\langle u_i, u_j \rangle = \delta_{ij}$ para los demás casos.

Con esto se puede enunciar el siguiente lema.

Lema 1.53. Un operador auto-adjunto T en un espacio vectorial de dimensión finita con métrica de Lorentz se puede representar de una de las

siguientes formas:

$$\text{I.}T \sim \begin{pmatrix} a_1 & & & 0 \\ & a_2 & & \\ & & \ddots & \\ 0 & & & a_n \end{pmatrix},$$

$$\text{II.}T \sim \begin{pmatrix} a_0 & 0 & & 0 \\ 1 & a_0 & & \\ & & a_3 & \\ 0 & & & \ddots & \\ & & & & a_n \end{pmatrix},$$

$$\text{III.}T \sim \begin{pmatrix} a_0 & 0 & 0 & & 0 \\ 0 & a_0 & 1 & & \\ -1 & 0 & a_0 & & \\ & & & a_4 & \\ 0 & & & & \ddots & \\ & & & & & a_n \end{pmatrix},$$

$$\text{IV.}T \sim \begin{pmatrix} a_0 & b_0 & & 0 \\ -b_0 & a_0 & & \\ & & a_3 & \\ & & & \ddots & \\ 0 & & & & a_n \end{pmatrix},$$

donde $b_0 \neq 0$. Las representaciones I y IV se encuentran en una base ortonormal mientras que las representaciones II y III se encuentran en una base pseudo-ortonormal.

Es claro que la representación I del Lema 1.53 ha quedado demostrada con el Teorema 1.51. El resto de las representaciones quedan sin demostración en esta tesis, y se puede consultar sobre este tema en [4] y [12] de la bibliografía.

Capítulo 2

Geometría Lorentziana

En la *geometría Riemanniana* se tiene una noción de *medir* en variedades C^∞ , a través de definir en cada espacio tangente un *producto interno*, el cual varía suavemente sobre la variedad. Como se mostró en el capítulo anterior, esta noción de medir (que nada tiene que ver con una medida en un espacio métrico) puede ser generalizada a través de definir un producto escalar con índice arbitrario, y al igual que en el caso Riemanniano esto se puede generalizar para variedades C^∞ .

El objetivo de este capítulo es mostrar la herramienta que se usará posteriormente en esta tesis. Se asumirán todos los resultados básicos respecto de las variedades C^∞ , las cuales se asumirán a todo momento Hausdorff y paracompactas. Cabe mencionar que los vectores tangentes a una variedad se usarán como derivaciones puntuales y los campos vectoriales como derivaciones del espacio de funciones C^∞ en un abierto de la variedad.

2.1. Variedades de Lorentz

Como se usarán variedades C^∞ a lo largo de toda la tesis, omitiré de ahora en adelante la clasificación C^∞ para referirme a éstas sólo por el nombre de variedades.

Definición 2.1. Una *métrica de Lorentz* $g_p = \langle \cdot, \cdot \rangle_p$ en una variedad M es un campo tensorial simétrico no degenerado de rango $(0, 2)$, con índice constante $\nu = 1$.

Por lo general se suprimirá el índice p para designar a la métrica; en este caso lo que se quiere es hacer énfasis en que, para cada $p \in M$, $\langle \cdot, \cdot \rangle_p$ es una forma bilineal simétrica de índice $\nu = 1$ en $T_p M$, es decir, un producto escalar en el espacio tangente a M en p el cual hace de $T_p M$ un espacio de

Minkowski. Por campo tensorial en este contexto nos referimos a uno que sea C^∞ , es decir, $\langle \cdot, \cdot \rangle_p$ varía suavemente respecto de p .

Esta definición se podría formular para un índice arbitrario ν , con esto, g_p tendría índice ν para cada espacio tangente.

No toda variedad admite una métrica de Lorentz, a diferencia del caso Riemanniano donde toda variedad admite una métrica Riemanniana, es decir, una métrica con índice constante cero.

Definición 2.2. Una *variedad de Lorentz* es una variedad M que admite una métrica de lorentz.

En el caso de tener una métrica de índice arbitrario, al par (M, g_p) se le conoce como una *variedad semi-Riemanniana*; entonces, de tener g_p índice cero M sería una *variedad Riemanniana* en el sentido de la teoría clásica de geometría diferencial.

Sean X y Y campos vectoriales sobre una variedad de Lorentz M , entonces, dado $p \in M$ el campo asocia un vector a este punto en T_p y $\langle X, Y \rangle_p \in \mathbb{R}$; ya que los campos están definidos sobre todo M al igual que la métrica se tiene que $\langle X, Y \rangle$ es una función C^∞ de M en los reales.

Si $\{x^\alpha\}$ es un sistema de coordenadas locales en un abierto $U \subset M$ los coeficientes de la métrica se definen de la siguiente forma

$$g_{\alpha\beta} := \langle \partial_\alpha, \partial_\beta \rangle$$

Entonces, en U los campos se escriben como $X = X^\alpha \partial_\alpha$ $Y = Y^\beta \partial_\beta$ y

$$\langle X, Y \rangle = g_{\alpha\beta} X^\alpha Y^\beta$$

Dado que la métrica es no degenerada la matriz $[g_{\alpha\beta}(p)]$ es invertible cuya matriz inversa se denota como $[g^{\alpha\beta}(p)]$, donde los coeficientes $g_{\alpha\beta}(p)$ son funciones C^∞ . A través de la formula de la matriz inversa se obtiene que los coeficientes $g^{\alpha\beta}$ son también funciones C^∞ sobre la variedad.

Como en el caso clásico, una forma de construir nuevas variedades de Lorentz es a través del producto de variedades.

Lema 2.3. Sea M una variedad de Lorentz y N una variedad Riemanniana con métricas g_M y g_N respectivamente. Si π y σ son las proyecciones de $M \times N$ en M y N respectivamente, sea

$$g := \pi^*(g_M) + \sigma^*(g_N)$$

donde $\pi^*(g_M)$ y $\sigma^*(g_N)$ denotan el *pullback* de g_M y g_N respectivamente. Entonces g es una métrica de Lorentz en $M \times N$ haciendo de esta una *variedad de Lorentz*.

Demostración. Para todo $v, w \in T_{(p,q)}(M \times N)$ se tiene lo siguiente

$$g(v, w) = g_M(d\pi(v), d\pi(w)) + g_N(d\sigma(v), d\sigma(w)).$$

Ya que el *pullback* está bien definido se tiene de entrada que g es un campo tensorial de rango $(0, 2)$ sobre $M \times N$, además, claramente es simétrico. Ahora, supongamos que para todo $w \in T_{(p,q)}(M \times N)$ $g(v, w) = 0$, entonces para todo $w \in T_{(p,q)}M$ se cumple $g_M(d\pi(v), d\pi(w)) = 0$ ya que $d\sigma(w) = 0$; como $d\pi$ es sobre se puede asegurar que $d\pi(v) = 0$ ya que g_M es no degenerada. Siguiendo el mismo argumento se tiene que $d\sigma(v) = 0$, lo cual implica que $v = 0$.

En ambos espacios T_pM y T_qN existen bases ortonormales, las cuales al juntarse formarán una base ortonormal para $T_{(p,q)}(M \times N)$ en la cual sólo existirá un vector tipo tiempo. Por lo tanto el índice de g es uno y constante, lo cual termina la demostración. $\square$

Por lo tanto, siendo $\mathbb{R}_1^1$ el conjunto de los números reales con su métrica negativa, basta hacer producto de este espacio con cualquier variedad Riemanniana M para obtener una variedad de Lorentz $\mathbb{R}_1^1 \times M$.

Isometrías

Al igual que en el caso Riemanniano, en este caso las *isometrías* juegan un papel importante ya que si dos variedades de Lorentz son *isométricas*, es decir, existe una *isometría* entre estas, se les puede tomar por la misma variedad en términos geométricos. A continuación definiré lo que es una *isometría* ya que es claro que la definición usada en el capítulo anterior no tiene sentido en este contexto a no ser que la variedad de la que se hable sea un espacio vectorial.

Definición 2.4. Sean M y N variedades de Lorentz con métricas g_M y g_N respectivamente. Un difeomorfismo $\phi : M \rightarrow N$ es una *isometría* si $\phi^*(g_N) = g_M$, es decir, siendo $\phi(p) = q$:

$$\langle d\phi(u), d\phi(v) \rangle_q = \langle u, v \rangle_p \quad \text{para todo } u, v \in T_pM \text{ y } p \in M.$$

Dado que ϕ es un difeomorfismo, la diferencial $d\phi_p$ es una función lineal de rango máximo, por lo cual la definición anterior implica que $d\phi_p$ es una isometría lineal para toda p . Con esta nueva definición de isometría se puede ver que las isometrías de las que se habló en el capítulo anterior son también isometrías en el contexto de variedades de Lorentz.

Existen ciertas propiedades básicas de las isometrías las cuales mencionaremos a continuación:

- La función identidad es una isometría.

Esta propiedad es clara ya que la diferencial de la función identidad en una variedad M es siempre la función identidad en $T_p M$ para todo $p \in M$.

- composición de isometrías es una isometría.

Sean $\phi : M \rightarrow N$ y $\psi : N \rightarrow P$ isometrías, entonces

$$\langle d\phi_p(u), d\phi_p(v) \rangle_q = \langle u, v \rangle_p ; \quad \langle d\psi_q(\tilde{u}), d\psi_q(\tilde{v}) \rangle_r = \langle \tilde{u}, \tilde{v} \rangle_q$$

donde $q = \phi(p)$ y $r = \psi(q)$. Como ϕ y ψ son difeomorfismos sus diferenciales son sobre y podemos escribir siempre a $\tilde{u}$ y $\tilde{v}$ como $d\phi_p(u)$ y $d\phi_p(v)$ respectivamente, entonces:

$$\langle d\psi_q(d\phi_p(u)), d\psi_q(d\phi_p(v)) \rangle_r = \langle u, v \rangle_p$$

Por lo tanto, haciendo uso de la regla de la cadena se obtiene que composición de isometrías, es isometría.

- La inversa de una isometría es también una isometría.

Ya que toda isometría ϕ es también un difeomorfismo, su diferencial es invertible y $d\phi_q^{-1} = (d\phi_p)^{-1}$ para todo $p \in M$. Entonces:

$$\langle d\phi_p(u), d\phi_p(v) \rangle_q = \langle u, v \rangle_p ; \quad \langle d\phi_q^{-1}d\phi_p(u), d\phi_q^{-1}d\phi_p(v) \rangle_p = \langle u, v \rangle_p$$

Por lo tanto, haciendo $\tilde{u} = d\phi(u)$ y $\tilde{v} = d\phi(v)$ queda demostrado que ϕ^{-1} es una isometría.

1-formas

Antes de pasar a la siguiente sección es prudente hacer ciertas definiciones que serán de gran utilidad. Primero cabe recordar lo que es el *espacio dual* en el contexto del álgebra lineal; siendo V un espacio vectorial real el *espacio dual* a este, el cual se denota V^* , es el espacio de todas las funciones lineales $g : V \rightarrow \mathbb{R}$. Este claramente es un espacio vectorial real a su vez.

Usando el Teorema 1.37 del capítulo anterior, se tiene que la función $V \rightarrow V^*$ que asocia a cada vector un único elemento del espacio dual, es un isomorfismo de espacios vectoriales (cuando la dimensión de V es finita). Por lo tanto V y V^* tienen la misma dimensión.

En el contexto de lo diferenciable, a los elementos del espacio dual se les conoce como *1-formas*, aunque en un sentido un tanto más amplio.

Definición 2.5. Una 1-forma ω sobre una variedad M es una función que asigna a $p \in M$ un elemento del espacio dual T_p^*M .

A lo largo de esta tesis se supondrá que toda 1-forma es C^∞ , esto quiere decir que, expresado esta en términos de la *base dual*, sus coeficientes son funciones C^∞ sobre M . Al espacio T_p^*M se le conoce como el espacio *cotangente* y a sus elementos se les conoce como *covectores*. Un ejemplo clásico de una 1-forma diferenciable es la diferencial df de una función $f : M \rightarrow \mathbb{R}$.

De ahora en adelante se denotará al conjunto de todas las funciones C^∞ sobre M como $C^\infty(M)$, al de los campos vectoriales C^∞ como $\mathfrak{X}(M)$, y al de las 1-formas C^∞ se le denotará $\mathfrak{X}^*(M)$.

A continuación demostraré un par de teoremas, los cuales son generalizaciones del Lema 1.32 y el Teorema 1.37 del capítulo anterior a este contexto.

Lema 2.6. Sean X y Y campos vectoriales sobre una variedad de Lorentz M tales que $\langle X, Z \rangle = \langle Y, Z \rangle$ para toda $Z \in \mathfrak{X}(M)$. Entonces $X = Y$.

Demostración. Si $\langle X, Z \rangle = \langle Y, Z \rangle$ para toda $Z \in \mathfrak{X}(M)$, entonces en cada espacio tangente se tiene que $\langle X_p, Z_p \rangle = \langle Y_p, Z_p \rangle$ para todo $Z_p \in T_pM$. Por el Lema 1.32 se tiene que $X_p = Y_p$ y, ya que esto es para toda $p \in M$, $X = Y$. $\square$

Teorema 2.7. Sea ω una 1-forma en $\mathfrak{X}^*(M)$ donde M es una variedad de Lorentz. Entonces, existe un único campo vectorial $Y \in \mathfrak{X}(M)$ tal que $\omega(X) = \langle X, Y \rangle$ para toda $X \in \mathfrak{X}(M)$.

Demostración. Sea $U \subset M$ una vecindad coordenada, entonces en U , tomando por base del tangente a $\{\partial_\alpha\}$ se obtiene la base dual $\{dx^\alpha\}$. En estas bases se obtiene la siguiente expresión local para $\omega \in \mathfrak{X}^*(M)$, esta se escribe como $\omega = \omega_\alpha dx^\alpha$; ahora, sea $Y \in \mathfrak{X}(M)$ tal que $Y = \omega^\beta \partial_\beta$ donde $\omega^\beta := g^{\alpha\beta} \omega_\alpha$. Entonces:

$$\langle Y, \partial_\sigma \rangle = \omega^\beta \langle \partial_\beta, \partial_\sigma \rangle = \omega_\alpha g^{\alpha\beta} g_{\beta\sigma}$$

Como $[g^{\alpha\beta}]$ es la inversa de $[g_{\beta\sigma}]$ finalmente se obtiene

$$\langle Y, \partial_\sigma \rangle = \omega_\alpha \delta_\sigma^\alpha = \omega_\sigma$$

Ya que $\omega(\partial_\sigma) = \omega_\sigma$ se tiene que estas dos funciones coinciden en una base en cada espacio tangente. Por lo tanto $\omega(X) = \langle X, Y \rangle$ para toda $X \in \mathfrak{X}(M)$. Para ver que Y es en realidad único, suponiendo que existe otro campo $\tilde{Y}$ que cumple con esta propiedad, usando el Lema 2.6, se prueba que que $\tilde{Y} = Y$. $\square$

En esta demostración se usó un método que es conocido en el lenguaje tensorial como *subir* índices; para esto se usaron las componentes del operador inverso de la métrica y se contrajo uno de sus índices con el índice de las componentes de la 1-forma para así obtener las componentes de un campo vectorial. El proceso inverso conocido como *bajar* índices, que convierte las componentes de un vector en los de una 1-forma, se lleva a cabo con las componentes de la métrica.

Para el siguiente teorema, que es una consecuencia del anterior, es necesario resaltar que los conjuntos $\mathfrak{X}(M)$ y $\mathfrak{X}^*(M)$ son módulos sobre el anillo $C^\infty(M)$. La demostración del siguiente teorema puede ser entendida también si uno toma en cuenta el proceso de subir y bajar índices.

Teorema 2.8. Sea M una variedad de Lorentz. Si $Z \in \mathfrak{X}(M)$, sea Z^* la 1-forma definida como

$$Z^*(X) := \langle X, Z \rangle \text{ para todo } X \in \mathfrak{X}(M)$$

Entonces, la función $Z \rightarrow Z^*$ es un isomorfismo $C^\infty(M)$ -lineal entre los espacios $\mathfrak{X}(M)$ y $\mathfrak{X}^*(M)$.

Demostración. Primero voy a probar que la función es $C^\infty(M)$ -lineal. Sean $Z, Y \in \mathfrak{X}(M)$ y $f \in C^\infty(M)$ entonces

$$\langle X, Z + fY \rangle = \langle X, Z \rangle + \langle X, fY \rangle = \langle X, Z \rangle + f\langle X, Y \rangle$$

Por lo tanto la función es $C^\infty(M)$ -lineal.

Por el Teorema 2.7 se tiene que esta función es sobre y uno a uno, y tiene una inversa. Para esta función $Z^* \rightarrow Z$, su $C^\infty(M)$ -linealidad queda demostrada también como una consecuencia del Teorema 2.7; basta ver el procedimiento en la demostración.

□

2.2. La conexión de Levi-Civita

En el caso de la geometría diferencial clásica donde se estudian superficies regulares como subconjuntos de $\mathbb{R}^3$ uno de los problemas principales es el siguiente. Sea $\alpha : I \rightarrow \mathbb{R}^3$ una curva suave en una superficie regular S y Y un campo vectorial sobre α , el vector $\frac{dY}{dt}(t)$ para todo $t \in I$ por lo regular no se encuentra en $T_{\alpha(t)}S$. En la geometría diferencial uno de los objetivos principales es hablar de objetos o nociones que sean intrínsecas al espacio en el cual se está trabajando, es decir en este caso, describir la geometría de S desde el punto de vista de S ; si la derivada del campo vectorial sobre α

no siempre es parte del espacio tangente a S esto quiere decir que no es un concepto intrínseco a S y por lo tanto no toma parte en la geometría de la superficie. En la teoría clásica para resolver este problema se recurre a definir la *derivada covariante* de Y , que es una proyección del vector $\frac{dY}{dt}(t)$ que se encuentra en el espacio ambiente, sobre el espacio tangente a S . Por lo tanto es necesario tener una noción intrínseca a toda variedad, que nos permita derivar campos vectoriales en la dirección de otro sin tener que recurrir a un espacio ambiente.

Existe una forma natural de hacer esto en $\mathbb{R}^n$. Sean $\{x^\alpha\}$ las coordenadas naturales de $\mathbb{R}^n$; entonces, sean X y $Y = Y^\beta \partial_\beta$ dos campos vectoriales sobre $\mathbb{R}^n$, se define el campo vectorial $D_X Y$ como

$$D_X Y := X(Y^\beta) \partial_\beta$$

Al cual se le llama la *derivada covariante natural* de Y respecto de X . El termino $X(Y^\beta)$ no es otra cosa que la derivada direccional de las funciones coordenadas del campo Y en la dirección de X , así que se está derivando el campo Y en la dirección de X . Aún así, no es claro como esto se puede llevar de una buena forma a una variedad arbitraria; es por eso que se necesita definir una *conexión* en la variedad.

Definición 2.9. Una *conexión* en una variedad M es una función C^∞ $D : \mathfrak{X}(M) \times \mathfrak{X}(M) \rightarrow \mathfrak{X}(M)$ tal que:

1. $D_{fX+gY}Z = fD_X Z + gD_Y Z$
2. $D_X(Y + Z) = D_X Y + D_X Z$
3. $D_X(fY) = fD_X Y + X(f)Y$

Para todo $f, g \in C^\infty(M)$ y $X, Y, Z \in \mathfrak{X}(M)$.

El punto número 1 quiere decir que $D_X Y$ es $C^\infty(M)$ -lineal en X , mientras que los puntos 2 y 3 que es $\mathbb{R}$ -lineal en Y . El número 3 en la definición anterior es una especie de regla de Leibniz, es decir que en Y una conexión actúa parecido a una derivación, lo que se busca para tener una forma de derivar campos vectoriales. De ser así esta noción debe ser local. Tomando a las *cartas* de la variedad como transformaciones $\phi : U \subset M \rightarrow \mathbb{R}_1^n$, un sistema coordenado en M es una colección de funciones $\{x^\alpha\}$ tales que $x^\alpha := r^\alpha \circ \phi$ donde $\{r^\alpha\}$ son las coordenadas usuales de $\mathbb{R}_1^n$. Entonces, escogiendo un sistema coordenado sobre M al rededor de un punto p , dos campos vectoriales X y Y se expresan localmente como sigue.

$$X = X^\alpha \partial_\alpha ; Y = Y^\beta \partial_\beta$$

Usando el número 3 de la Definición 2.9 se obtiene lo siguiente

$$D_X Y = X^\alpha D_{\partial_\alpha} (Y^\beta \partial_\beta) = X^\alpha (\partial_\alpha (Y^\beta) \partial_\beta + Y^\beta D_{\partial_\alpha} \partial_\beta)$$

Si defino $D_{\partial_\alpha} \partial_\beta := \Gamma_{\beta\alpha}^\mu \partial_\mu$, claramente las funciones $\Gamma_{\beta\alpha}^\mu$, que se conocen como los *símbolos de Christoffel*, tienen que ser C^∞ ; finalmente el campo $D_X Y$ se expresa como

$$D_X Y = (X^\alpha \partial_\alpha (Y^\mu) + X^\alpha Y^\beta \Gamma_{\beta\alpha}^\mu) \partial_\mu. \quad (2.1)$$

Por lo tanto, $D_X Y(p)$ sólo depende de los valores $X^\alpha(p)$, $Y^\beta(p)$ y de las derivaciones ∂_α sobre las funciones Y^μ en el punto p , lo cual deja ver que la noción de conexión es local.

Lema 2.10. Sea M una variedad con una conexión D . Entonces, existe una única correspondencia que a cada campo vectorial Y sobre una curva suave $\alpha : I \rightarrow M$ le asocia otro campo vectorial $\frac{DY}{dt}$ sobre α , al cual se le nombra la *derivada covariante* de Y en la dirección de α , y que cumple con lo siguiente:

1. $\frac{D}{dt}(Y + Z) = \frac{DY}{dt} + \frac{DZ}{dt}$.
2. $\frac{D}{dt}(fY) = \frac{df}{dt}Y + f\frac{DY}{dt}$, donde $Y, Z \in \mathfrak{X}(\alpha)$ y $f \in C^\infty(I)$.
3. Si Y es la restricción de un campo $\tilde{Y} \in \mathfrak{X}(M)$ en α , entonces $\frac{DY}{dt} = D_{\frac{d\alpha}{dt}} \tilde{Y}$.

Con esto es obvio que la razón principal para definir una conexión en una variedad, es para poder derivar campos en la dirección de otro.

Definición 2.11. Sea M una variedad con una conexión D . Se dice que un campo vectorial Y sobre una curva $\alpha : I \rightarrow M$ es *paralelo* si $\frac{DY}{dt} = 0$, para todo $t \in I$.

Este concepto, junto con el de *transporte paralelo*, es otro que llevó a los matemáticos a desarrollar el concepto de conexión. Esto está ligado fuertemente al estudio de curvas *geodésicas* en una variedad.

Teorema 2.12. Sea M una variedad con una conexión D , $\alpha : I \rightarrow M$ una curva C^∞ y X_0 un vector en $T_{\alpha(t_0)}M$. Entonces existe un único campo vectorial paralelo $X \in \mathfrak{X}(\alpha)$ tal que $X(t_0) = X_0$. A X se le conoce como el *transporte paralelo* de X_0 a lo largo de α .

El pasado Lema 2.10 y Teorema 2.12 se dejarán sin demostración ya que de el objetivo de esta sección es el teorema de *Levi-Civita*, la demostración de estos se puede encontrar en [5]. Antes necesitaré una definición.

Definición 2.13. Sea M una variedad. El *corchete* es una función $[\cdot, \cdot] : \mathfrak{X}(M) \times \mathfrak{X}(M) \rightarrow \mathfrak{X}(M)$ definida como:

$$[X, Y]f := XY(f) - YX(f)$$

Para toda $f \in C^\infty(M)$.

El siguiente teorema es en realidad válido para toda variedad semi-Riemanniana así como su demostración; si en las hipótesis se pide que la variedad sea de Lorentz es tan sólo porque en toda la tesis sólo estudiaremos esta clase de variedades.

Teorema 2.14. Sea M una variedad de Lorentz. Existe una única conexión D en M que cumple con las siguientes propiedades:

1. $[X, Y] = D_X Y - D_Y X$; D es simétrica o libre de torsión.
2. $X\langle Y, Z \rangle = \langle D_X Y, Z \rangle + \langle Y, D_X Z \rangle$; D es compatible con la métrica.

Para todo $X, Y, Z \in \mathfrak{X}(M)$. A esta conexión se le conoce como la *conexión de Levi-Civita*.

Demostración. Primero supondré que tal conexión existe, entonces:

$$X\langle Y, Z \rangle = \langle D_X Y, Z \rangle + \langle Y, D_X Z \rangle \quad (2.2)$$

$$Y\langle Z, X \rangle = \langle D_Y Z, X \rangle + \langle Z, D_Y X \rangle \quad (2.3)$$

$$Z\langle X, Y \rangle = \langle D_Z X, Y \rangle + \langle X, D_Z Y \rangle \quad (2.4)$$

Sumando las ecuaciones (2.2) y (2.3), y restando la ecuación (2.4) se obtiene

$$\langle D_Y Z - D_Z Y, X \rangle + \langle D_X Z - D_Z X, Y \rangle + \langle D_X Y, Z \rangle + \langle D_Y X, Z \rangle$$

$$\langle [Y, Z], X \rangle + \langle [X, Z], Y \rangle + \langle D_X Y, Z \rangle + \langle D_Y X, Z \rangle$$

Entonces, sumando el término $\langle D_Y X, Z \rangle - \langle D_Y X, Z \rangle$ se obtiene finalmente

$$2\langle D_Y X, Z \rangle = X\langle Y, Z \rangle + Y\langle Z, X \rangle - Z\langle X, Y \rangle - \langle [Y, Z], X \rangle - \langle [X, Z], Y \rangle - \langle [X, Y], Z \rangle \quad (2.5)$$

A esta ecuación se le conoce como la *fórmula de Koszul*. Ahora voy a definir $F(X, Y, Z)$ como la parte derecha de la ecuación (2.5); si X y Y son campos fijos sucede que $F(X, Y, Z + W) = F(X, Y, Z) + F(X, Y, W)$ para todo campo Z y W en $\mathfrak{X}(M)$ por la linealidad de la métrica y la propiedad del corchete

$[X, Y + Z] = [X, Y] + [X, Z]$. Otra propiedad que usaré del corchete es, dada $f \in C^\infty(M)$ se tiene que $[fX, Y] = -YfX + f[X, Y]$. Entonces:

$$\begin{aligned} F(X, Y, fZ) = & X(f\langle Y, Z \rangle) + Y(f\langle Z, X \rangle) - fZ(\langle X, Y \rangle - \langle [X, Y], Z \rangle) \\ & - \langle YfZ + f[Y, Z], X \rangle - \langle XfZ + f[X, Z], Y \rangle \end{aligned}$$

Usando la regla de Leibniz en los primeros dos términos de la función F se obtiene

$$Xf\langle Y, Z \rangle + fX\langle Y, Z \rangle ; Yf\langle Z, X \rangle + fY\langle Z, X \rangle$$

Estas dos expresiones sumadas con los dos últimos términos de la función F demuestran que $F(X, Y, fZ) = fF(X, Y, Z)$ y con esto, que $X \mapsto F(X, Y, Z)$ es $C^\infty(M)$ -lineal y por lo tanto una 1-forma. Por el Teorema 2.8 se tiene que existe un único campo vectorial en $\mathfrak{X}(M)$, al cual ventajosamente nombraré $D_Y X$, tal que $F(X, Y, Z) = 2\langle D_Y X, Z \rangle$; ahora sólo falta ver que este campo representa en realidad una conexión con las cualidades deseadas.

Empezaré por demostrar que este campo es una conexión, cumpliendo con las tres propiedades enunciada en la Definición 2.9.

1. Sean $f, g \in C^\infty(M)$ y $W, X, Y, Z \in \mathfrak{X}(M)$; lo que se quiere demostrar es que $F(X, fY + gW, Z) = fF(X, Y, Z) + gF(X, Y, Z)$. Entonces:

$$\begin{aligned} F(X, fY + gW, Z) = & X(f\langle Y, Z \rangle + g\langle W, Z \rangle) + (fY + gW)\langle Z, X \rangle \\ & - Z(f\langle X, Y \rangle + g\langle X, W \rangle) \\ & - \langle [fY + gW, Z], X \rangle - \langle [X, Z], fY + gW \rangle \\ & - \langle [X, fY + gW], Z \rangle \end{aligned}$$

Como se puede ver, el segundo término así como el quinto ya son parte del resultado al que se quiere llegar; usando las técnicas que se han usado para reescribir los términos de la forma $[fY, Z]$ se demuestra que $F(X, fY + gW, Z) = fF(X, Y, Z) + gF(X, Y, Z)$.

2. Este punto sólo depende de que la métrica como los brackets abran las sumas de campos, lo cual sí hacen.
3. En este caso se tiene que demostrar que $F(fX, Y, Z) = 2\langle YfX, Z \rangle + fF(X, Y, Z)$; usando las mismas técnicas se obtiene lo siguiente:

$$\begin{aligned} F(fX, Y, Z) = & fX\langle Y, Z \rangle + Y(f\langle Z, X \rangle) - Z(f\langle X, Y \rangle) \\ & - f\langle [Y, Z], X \rangle - \langle -ZfX + f[X, Z], Y \rangle \\ & - \langle -YfX + f[X, Y], Z \rangle \end{aligned}$$

Aplicando la regla de Leibniz una vez más y sumando los términos restantes se obtiene el resultado esperado.

En todas estas demostraciones se ha usado de manera implícita el Lema 2.6. Con esto queda claro que el campo $D_Y X$ es en realidad una conexión, ahora sólo falta ver que es la deseada.

Para ver que la conexión es libre de torsión hay que empezar con lo siguiente.

$$2\langle D_X Y - D_Y X, Z \rangle = F(Y, X, Z) - F(X, Y, Z)$$

Haciendo la resta en el lado derecho de la ecuación queda:

$$\langle [X, Y], Z \rangle - \langle [Y, X], Z \rangle = 2\langle [X, Y], Z \rangle$$

Para ver que es compatible con la métrica es necesario hacer la suma

$$F(Y, X, Z) + F(Z, X, Y) = 2\langle D_X Y, Z \rangle + 2\langle D_X Z, Y \rangle$$

cuyo resultado es $2X\langle Y, Z \rangle$ como se esperaba.

□

Derivación de tensores

El Teorema 2.8 nos permite ver a los campos vectoriales como campos tensoriales de orden $(1, 0)$, es decir, un vector actúa sobre una 1-forma obteniendo una función en $C^\infty(M)$. Siendo así, para X fijo en $\mathfrak{X}(M)$ la función DX es un tensor de rango $(1, 1)$ sobre M , ya que $D_Y X$ es $C^\infty(M)$ -lineal en Y y $D_Y X$ es un campo vectorial sobre M .

Un tensor de rango $(0, 0)$ es por definición cualquier función C^∞ sobre M ; a diferencia de los vectores, estos no tienen una base y al *derivarlos* uno no tendría que derivar la base como se hace con vectores. Por lo tanto, del punto 3 de la Definición 2.9 podemos pensar que si $f \in C^\infty(M)$ entonces, $D_Y f := Yf = df(Y)$. Esto es tan sólo un método heurístico de justificar la siguiente definición.

Definición 2.15. Sea $Y \in \mathfrak{X}(M)$. Se define la *derivada covariante* (de Levi-Civita) D_Y como la única *derivación tensorial* sobre M tal que

$$D_Y f = Y(f) = df(Y) \text{ para todo } f \in C^\infty(M),$$

y $D_Y X$ es la conexión (derivada) de Levi-Civita para todo $X \in \mathfrak{X}(M)$.

Hasta aquí vale la pena hacer dos comentarios. El primero es una cosa trivial; de manera alternada se empezarán a usar los términos derivada covariante y conexión de Levi-Civita no para referirnos al mismo objeto u operación pero sí para la noción de derivar campos tensoriales arbitrarios

sobre M . El segundo es referente al término *derivación tensorial*; una forma de construir vectores, el espacio tangente, y campos vectoriales sobre una variedad es haciendo uso de derivaciones sobre el conjunto $C^\infty(M)$, o de haces tensoriales; esta teoría de derivaciones tensoriales es lo que al final justifica formalmente la pasada definición y algunas otras por venir. Ya que entrar en detalles de esta teoría sería un considerable desvío, además de ser parte de la teoría básica de variedades C^∞ que se ha obviado desde un principio, esta no se expondrá aquí.

Antes de dar la definición de derivada covariante sobre un tensor arbitrario me permitiré hacer lo siguiente. Sean $\omega \in \mathfrak{X}^*(M)$ y $Y \in \mathfrak{X}(M)$; si f es una función en $C^\infty(M)$ entonces su derivada covariante en la dirección del vector ∂_β no es más que su derivada parcial $\frac{\partial f}{\partial x^\beta}$, entonces, ya que $\omega(Y)$ es una función escalar sobre M , en términos de sus componentes existe una función $h \in C^\infty(M)$ tal que $h = \omega_\alpha Y^\alpha$. Usando la regla de Leibniz usual se tiene lo que sigue.

$$D_{\partial_\beta} h = \frac{\partial h}{\partial x^\beta} = \frac{\partial \omega_\alpha}{\partial x^\beta} Y^\alpha + \omega_\alpha \frac{\partial Y^\alpha}{\partial x^\beta}$$

Usando la ecuación (2.1) se puede sustituir el término $\frac{\partial Y^\alpha}{\partial x^\beta}$ en la ecuación anterior obteniendo

$$D_{\partial_\beta} h = \frac{\partial \omega_\alpha}{\partial x^\beta} Y^\alpha + \omega_\alpha D_{\partial_\beta} Y^\alpha - \omega_\alpha Y^\mu \Gamma_{\mu\beta}^\alpha$$

donde $D_{\partial_\beta} Y^\alpha$ es la componente α del vector $D_{\partial_\beta} Y$. Cambiando algunos índices, lo cual siempre se puede hacer cuando estos se encuentran repetidos, finalmente se obtiene la siguiente ecuación.

$$D_{\partial_\beta} h = \left(\frac{\partial \omega_\alpha}{\partial x^\beta} - \omega_\mu \Gamma_{\alpha\beta}^\mu \right) Y^\alpha + \omega_\alpha D_{\partial_\beta} Y^\alpha$$

Ya que $D_{\partial_\beta} Y$ es un campo vectorial y $D_{\partial_\beta} g$ es una función, la ecuación anterior sugiere que el término entre paréntesis sean las componentes de una 1-forma lo cual nos lleva a definir la derivada covariante de una 1-forma ω en la dirección de un vector X como

$$D_X \omega := (X^\beta \partial_\beta (\omega_\alpha) - X^\beta \omega_\mu \Gamma_{\alpha\beta}^\mu) dx^\alpha \quad (2.6)$$

la cual es una 1-forma a su vez, obteniendo así la relación

$$D_X \omega(Y) = D_X(\omega(Y)) - \omega(D_X Y)$$

donde $D_X(\omega(Y)) = X(\omega(Y))$.

Usando el mismo argumento que nos llevó a la ecuación (2.6) y tomando en cuenta que un campo vectorial $Y \in \mathfrak{X}(M)$ es un tensor actuando en $\mathfrak{X}^*(M)$ se obtiene la relación

$$D_X Y(\omega) = X(Y(\omega)) - Y(D_X \omega).$$

Se puede ver hasta ahora que la *derivada covariante* de un vector es un vector así como la derivada de una 1-forma es de nuevo una 1-forma. Como se dijo al inicio de esta subsección, D_X es $C^\infty(M)$ -lineal en X y por lo tanto se tiene que $D_X Y$ es un tensor de rango $(1, 1)$ para cualquier vector Y ; entonces $D_X \omega$ es un tensor de rango $(0, 2)$. Todo esto nos lleva a la siguiente definición.

Definición 2.16. Sea A un tensor de orden (r, s) sobre M . Definimos la *derivada covariante* de A como el tensor de rango (r, s)

$$\begin{aligned} (D_X A)(\omega^1, \dots, \omega^r, Y_1, \dots, Y_s) := & X(A(\omega^1, \dots, \omega^r, Y_1, \dots, Y_s)) \\ & A(D_X \omega^1, \dots, \omega^r, Y_1, \dots, Y_s) - \dots \\ & \dots - A(\omega^1, \dots, \omega^r, Y_1, \dots, D_X Y_s) \end{aligned}$$

para todo $\omega^\alpha \in \mathfrak{X}^*(M)$ y $Y_\beta, X \in \mathfrak{X}(M)$.

Ya que $D_X A$ es $C^\infty(M)$ -lineal en X , como había mencionado antes, este se puede ver siempre como un tensor de rango $(r, s + 1)$. Cabe que observar que en todo esto siempre se ha supuesto que D es la conexión de Levi-Civita y no cualquier conexión; esto nos lleva al siguiente lema.

Lema 2.17. Sea M una variedad de Lorentz, g su métrica y D su conexión de Levi-Civita. Entonces, para todo $X \in \mathfrak{X}(M)$ $D_X g = 0$.

Demostración. Como D es compatible con la métrica se tiene lo siguiente:

$$Xg(Y, Z) = g(D_X Y, Z) + g(Y, D_X Z) ; \text{ para todo } X, Y, Z \in \mathfrak{X}(M).$$

Por otro lado, si derivo la métrica g obtengo:

$$(D_X g)(Y, Z) = Xg(Y, Z) - g(D_X Y, Z) - g(Y, D_X Z),$$

para todo $X, Y, Z \in \mathfrak{X}(M)$. Usando la compatibilidad de la conexión con la métrica y substituyendo el término $Xg(Y, Z)$ en la última ecuación se tiene que $D_X g = 0$. □

De aquí en adelante me referiré a la conexión de Levi-Civita de una variedad, simplemente como la conexión. En caso de usar una conexión que no sea esta se especificará cuando sea necesario.

2.3. El tensor de curvatura

En la geometría clásica la *curvatura* de una superficie regular se define por medios que no son enteramente intrínsecos, de ahí la importancia del Teorema *egregio* de Gauss el cual hace ver que tal curvatura es en realidad una cantidad intrínseca para cada superficie regular. En la actualidad, para llegar a tal noción de *curvatura* es necesario desarrollar una maquinaria que a primera vista no proporciona la información geométrica del caso clásico, pero que la engloba con la ventaja de expresarse en un lenguaje enteramente intrínseco.

Lema 2.18. Sea M una variedad de Lorentz con conexión de Levi-Civita D . La función $R : \mathfrak{X}(M)^3 \rightarrow \mathfrak{X}(M)$ dada por la expresión

$$R(X, Y)Z := D_{[X, Y]}Z - D_X D_Y Z + D_Y D_X Z$$

define un tensor de orden $(1, 3)$ al cual se le conoce como el *tensor de curvatura* de M .

Demostración. Ver que R abre sumas de campos vectoriales en cada uno de sus argumentos es trivial; basta con demostrar que es $C^\infty(M)$ -lineal en cada uno de sus argumentos para ver que es en realidad un tensor.

$$\begin{aligned} R(fX, Y)Z &= D_{[fX, Y]}Z - D_{fX}D_Y Z + D_Y D_{fX}Z \\ &= D_{-YfX + f[X, Y]}Z - fD_X D_Y Z + D_Y(fD_X Z) \\ &= -YfD_X Z + fD_{[X, Y]}Z - fD_X D_Y Z \\ &\quad + YfD_X Z + fD_Y D_X Z \\ &= fR(X, Y)Z \end{aligned}$$

Para fY se procede de manera análoga; sólo falta demostrarlo para fZ .

$$\begin{aligned} R(X, Y)(fZ) &= D_{[X, Y]}(fZ) - D_X D_Y(fZ) + D_Y D_X(fZ) \\ &= [X, Y]fZ + fD_{[X, Y]}Z - D_X(YfZ) \\ &\quad - D_X(fD_Y Z) + D_Y(XfZ) + D_Y(fD_X Z) \end{aligned} \tag{2.7}$$

Sumando el tercer y quinto término de la ecuación (2.7) se obtiene lo siguiente

$$\begin{aligned} -D_X(YfZ) + D_Y(XfZ) &= -XYfZ - YfD_X Z + YXfZ + XfD_Y Z \\ &= -[X, Y]fZ - YfD_X Z + XfD_Y Z \end{aligned}$$

mientras que sumando el cuarto y el sexto

$$D_Y(fD_X Z) - D_X(fD_Y Z) = YfD_X Z + fD_Y D_X Z - XfD_Y Z - fD_X D_Y Z.$$

Sumando todo se obtiene que $R(X, Y)(fZ) = fR(X, Y)Z$. Por lo tanto R es $C^\infty(M)$ -lineal; ya que todo elemento de $\mathfrak{X}(M)$ es un tensor de orden $(1, 0)$ se tiene finalmente que R es un tensor de orden $(1, 3)$. $\square$

Ademas, este tensor sólo depende de los valores de X, Y, Z en $p \in M$ y de la conexión D , por lo cual el valor que tome R es una cuestión local.

El tensor de curvatura nos provee una medida de qué tanto los operadores D_X y D_Y no conmutan. Sea $M = \mathbb{R}^n$, entonces para un campo vectorial $Z = (z^0, \dots, z^{n-1})$ expresado en las coordenadas usuales de $\mathbb{R}^n$ se tiene que

$$D_X Z = (Xz^0, \dots, Xz^{n-1}) ; \text{ para todo } X, Z \in \mathfrak{X}(\mathbb{R}^n)$$

ya que $\Gamma_{\beta\alpha}^\mu = 0$ para todo α y β en la base usual de $\mathbb{R}^n$. Entonces:

$$D_Y D_X Z = (YXz^0, \dots, YXz^{n-1}) ; \text{ para todo } X, Y, Z \in \mathfrak{X}(\mathbb{R}^n)$$

Esto quiere decir que el tensor de curvatura se anula en todos los puntos de $\mathbb{R}^n$. En cierta forma, el tensor de curvatura también nos dice qué tanto nuestra variedad se aleja de ser un espacio euclidiano.

Sea $y \in T_p M$, es un resultado familiar que este se puede extender a un campo vectorial Y sobre M que sea C^∞ de muchas formas; expresado a Y como $y^\alpha \partial_{\alpha|q}$ en un sistema coordenado, donde $q \in U$ y U es un abierto de M , se obtiene un campo vectorial C^∞ sobre U , lo cual para cuestiones locales sigue teniendo importancia. Por ejemplo, el corchete $[X, Y]$ se anula en U para esta clase de extensiones sobre dos vectores $x, y \in T_p M$, simplificando la expresión del tensor de curvatura en U .

El tensor de curvatura posee distintas *simetrías* respecto a sus argumentos las cuales numeraré a continuación y dejaré sin demostrar, ya que son resultados bastante familiares en este contexto.

Sean $X, Y, Z, W \in \mathfrak{X}(M)$ campos vectoriales cualesquiera, entonces:

1. $R(X, Y) = -R(Y, X)$
2. $\langle R(X, Y)Z, W \rangle = -\langle R(X, Y)W, Z \rangle$
3. $R(X, Y)Z + R(Y, Z)X + R(Z, X)Y = 0$
4. $\langle R(X, Y)Z, W \rangle = \langle R(Z, W)X, Y \rangle$

A la tercera ecuación se le conoce como la *identidad de Bianchi*.

Curvatura seccional

Algunas veces el tensor de curvatura, a pesar de toda su información, se vuelve un aparato complicado de manejar. Es por eso que se define la *curvatura seccional* la cual, como demostraré más adelante, determina al tensor de curvatura.

Sea $T_p M$ el espacio tangente a M , se llama *plano tangente* Π de M a cualquier subespacio de $T_p M$ que sea de dimensión dos. La siguiente función a definir guarda cierto significado geométrico:

$$Q(v, w) := \langle v, v \rangle \langle w, w \rangle - \langle v, w \rangle^2 ; \text{ para todo } v, w \in T_p M.$$

Si $T_p M$ fuera un espacio con producto interno la pasada función Q no es más que el área del paralelogramo formado por los vectores v y w en el plano Π generado por estos dos, la cual siempre es distinto de cero si estos son linealmente independientes. En nuestro caso $T_p M$ es un espacio con producto escalar de índice uno; por el Lema 1.6 se tiene que $Q(\Pi) = 0$ si y sólo si Π es nulo. Así, si Π es tipo tiempo $Q(\Pi) < 0$ y, si Π es tipo espacio $Q(\Pi) > 0$. En estos dos últimos casos $Q(\Pi)$ puede seguirse viendo como el área del paralelogramo formado por los vectores usados, sólo que tendrá signo dependiendo de la causalidad de Π .

Lema 2.19. Sea Π un plano tangente no nulo a una variedad M en p y sean v, w una base para de este. El número

$$K(v, w) := \frac{\langle R(v, w)v, w \rangle}{Q(v, w)}$$

es independiente de la base. A este se le conoce como la *curvatura seccional* $K(\Pi)$ de Π .

Demostración. Cualesquiera dos bases en Π se encuentran relacionadas por la siguientes ecuaciones

$$v = ax + by ; w = cx + dy$$

tales que el determinante $ad - bc \neq 0$. Por la definición del tensor de curvatura se tiene $R(X, X)$ es la función cero; además, haciendo uso de las simetrías del tensor se tiene también que $\langle R(X, Y)Z, Z \rangle = 0$, entonces:

$$\langle R(v, w)v, w \rangle = ad\langle R(x, y)v, w \rangle - bc\langle R(x, y)v, w \rangle,$$

$$\langle R(x, y)v, w \rangle = ad\langle R(x, y)x, y \rangle - bc\langle R(x, y)x, y \rangle.$$

Por lo tanto:

$$\langle R(v, w)v, w \rangle = (ad - bc)^2 \langle R(x, y)x, y \rangle.$$

Como la función Q es el determinante del producto escalar de $T_p M$ restringido a Π , se tiene que

$$Q(v, w) = (ad - bc)^2 Q(x, y).$$

□

Entonces, la curvatura seccional está definida en todos los planos tangentes a M que sean no nulos. Además, queda claro por la definición que K está determinado por R , pero; ¿Está R determinado por K ?

Lema 2.20. Sean v y w dos vectores en un espacio con producto escalar de índice uno. Entonces existen vectores $\bar{v}$ y $\bar{w}$ tan cercanos a v y w respectivamente tales que generan un plano no nulo.

Este lema nos ayudará a resolver la pregunta planteada con anterioridad. Una demostración de este se puede ver en [13].

Teorema 2.21. Si $K(\Pi) = 0$ para todo plano no nulo en $T_p M$, entonces $R(v, w)z = 0$ para todo $v, w, z \in T_p M$.

Demostración. Supongamos primero que v y w en $T_p M$ generan un plano no nulo, por la definición de curvatura seccional se tiene que $\langle R(v, w)v, w \rangle = 0$. Usando el Lema 2.20 se tiene que existen siempre dos vectores que generan un plano no nulo tan cerca como se quiera de cualquier par arbitrario de vectores en $T_p M$; ya que R es una función multilineal en $T_p M^4$ es continua, y por tanto R también se anula en los vectores límite, es decir, $\langle R(v, w)v, w \rangle = 0$ para todo $v, w \in T_p M$.

Ahora, para x arbitrario en $T_p M$:

$$\begin{aligned} \langle R(v, w + x)v, w + x \rangle &= \langle R(v, w)v, w \rangle + \langle R(v, x)v, w \rangle \\ &\quad + \langle R(v, w)v, x \rangle + \langle R(v, x)v, x \rangle. \end{aligned}$$

Ya que $\langle R(v, w)v, w \rangle = 0$, usando las simetrías de R se obtiene que:

$$\langle R(v, w)v, x \rangle = 0 ; \text{ para toda } x \in T_p M$$

Por lo tanto $R(v, w)v = 0$ para todo $v, w \in T_p$. Sea ahora z arbitraria:

$$R(v + z, w)v + z = R(v, w)v + R(z, w)v + R(v, w)z + R(z, w)z.$$

Otra vez son tres términos de la ecuación los que se anulan, y por la anti-simetría de R en las primeras dos entradas se tiene finalmente que $R(v, w)z =$

$R(w, z)v$. Por lo tanto se tiene que R se mantiene invariante al permutar las primeras tres entradas de manera cíclica, entonces usando la identidad de Bianchi se obtiene que $R(v, w)z = 0$ para todo $v, w, z \in T_p M$ y con esto $R = 0$ donde $K = 0$. □

Sea $F : T_p M^4 \rightarrow \mathbb{R}$ una función multilinear tal que posee todas las simetrías del tensor de curvatura; en la demostración del teorema anterior jamás se tomó en cuenta la definición del tensor R , sólo sus simetrías, entonces, siendo F una función con las mismas simetrías que R siempre que $F(v, w, v, w) = 0$ para vectores $v, w \in T_p M$ tales que estos generen un plano no nulo, $F = 0$. Ya que la función $F - R$ claramente es multilinear y posee las simetrías del tensor de curvatura, el siguiente corolario queda automáticamente demostrado.

Corolario 2.22. Sea $F : T_p M^4 \rightarrow \mathbb{R}$ una función multilinear con las simetrías del tensor de curvatura tal que

$$K(v, w) = \frac{F(v, w, v, w)}{Q(v, w)}$$

para vectores $v, w \in T_p M$ que generen un plano no nulo. Entonces

$$\langle R(v, w)x, y \rangle = F(v, w, x, y)$$

para todo $v, w, x, y \in T_p M$.

Por lo tanto se tiene que la curvatura seccional tiene la misma información que el tensor de curvatura, es decir, K determina a R .

Se dice que una variedad M tiene *curvatura constante* si su curvatura seccional es constante, en este caso el tensor de curvatura tiene una expresión más sencilla.

Corolario 2.23. Sea M una variedad con curvatura constante C , entonces

$$R(x, y)z = C [\langle z, x \rangle y - \langle z, y \rangle x].$$

Demostración. Voy a definir la función F como sigue:

$$F(x, y, v, w) := C [\langle v, x \rangle \langle y, w \rangle - \langle v, y \rangle \langle x, w \rangle].$$

Esta función claramente es multilinear y también posee las simetrías del tensor de curvatura. Solo demostraré que cumple con la identidad de Bianchi.

$$F(x, y, v) = C [\langle v, x \rangle y - \langle v, y \rangle x]$$

$$F(y, v, x) = C [\langle x, y \rangle v - \langle x, v \rangle y]$$

$$F(v, x, y) = C [\langle y, v \rangle x - \langle y, x \rangle v]$$

Dada la simetría del producto escalar en $T_p M$ basta con sumar estas tres ecuaciones para verificar que F cumple con la identidad de Bianchi. Entonces, $F(x, y, x, y) = CQ(x, y)$ y siempre que x y y generen un plano no nulo se tiene que:

$$C = K(x, y) = \frac{F(x, y, x, y)}{Q(x, y)}.$$

Aplicando el Corolario 2.22 se obtiene el resultado deseado. $\square$

Antes de terminar con esta sección vale la pena mencionar cierta simplificación para la expresión de la curvatura seccional. Sea $\{u, v\}$ una base ortonormal para el plano tangente Π ; si Π es tipo tiempo

$$K(u, v) = -R(u, v, u, v);$$

si Π es tipo espacio

$$K(u, v) = R(u, v, u, v).$$

2.4. Operadores diferenciales

En esta sección definiré operadores diferenciales que son clásicos del cálculo, tales como el *gradiente* o el *laplaciano*.

En cálculo el gradiente de una función $f : \mathbb{R}^n \rightarrow \mathbb{R}$ es un vector cuyas componentes son las derivadas parciales de f en la base canónica y las entradas de este vector cambian si se cambia de coordenadas. En caso de estar en las coordenadas usuales de $\mathbb{R}^n$ a este vector también se le puede ver como la diferencial de f o, mejor dicho, confundir con la diferencial de f . En el contexto de las variedades, queda muy claro que df es en realidad una 1-forma, y por lo tanto esta no depende de las coordenadas, sin embargo es evidente que existe una fuerte relación entre la diferencial de f y su gradiente.

Definición 2.24. Sea $f \in C^\infty(M)$. Se define el *gradiente* de f , $\text{grad}f$ como el único campo vectorial tal que:

$$\langle \text{grad}f, X \rangle = df(X) = Xf ; \text{ para todo } X \in \mathfrak{X}(M).$$

Esto quiere decir que $\text{grad}f \in \mathfrak{X}(M)$ es un campo vectorial que es *métricamente equivalente* a la diferencial de f , es más, el único.

Para definir la *divergencia* de un campo vectorial necesitaré hablar de que es un *marco* sobre una variedad M . Un *marco* no es más que una colección de campos vectoriales $\{E_0, \dots, E_{n-1}\}$, si la dimensión de M es n , tales que para cada p en M estos forman un conjunto linealmente independiente y ortonormal en $T_p M$. Esta clase de marcos no existe siempre de manera global, pero sí localmente.

En cálculo la *divergencia* de un campo Y se define como la suma de las derivadas $\frac{\partial Y^\alpha}{\partial x^\alpha}$; como ya se ha mencionado antes, derivar sólo parcialmente en este contexto no nos da una noción completa de lo que es derivar un campo vectorial y, por eso, la *divergencia* debe incluir a la conexión.

Definición 2.25. Sea $Y \in \mathfrak{X}(M)$. Para un marco, se define la *divergencia* de Y , $\text{div}Y$ como:

$$\text{div}Y := \epsilon^\alpha \langle D_{E_\alpha} Y, E_\alpha \rangle.$$

Donde $\epsilon^\alpha := \langle E_\alpha, E_\alpha \rangle$.

Antes de seguir adelante me gustaría hacer una observación. El operador $F[Z] : \mathfrak{X}(M) \times \mathfrak{X}(M) \rightarrow C^\infty(M)$ definido como $F[Z](X, Y) := \langle D_X Z, Y \rangle$ claramente es $C^\infty(M)$ -lineal para cada $Z \in \mathfrak{X}(M)$. Por lo tanto un tensor de rango $(0, 2)$.

En general, para un tensor sobre M de rango (r, s) se dice que este tiene r entradas *contravariantes* y s entradas *covariantes*, donde las primeras r corresponden a elementos de $\mathfrak{X}^*(M)$ y las últimas s a elementos de $\mathfrak{X}(M)$.

Sea $\{X_\alpha\}$ una base (local) de $\mathfrak{X}(M)$ y $\{\theta^\beta\}$ su base dual. Se definen los *símbolos* de un tensor A de rango (r, s) como:

$$A_{\alpha_1 \dots \alpha_s}^{\beta_1 \dots \beta_r} := (\theta^{\beta_1}, \dots, \theta^{\beta_r}, X_{\alpha_1}, \dots, X_{\alpha_s})$$

En este lenguaje, usando el tensor métrico y su operador inverso, se puede modificar el rango de un tensor dado al subir y bajar índices, es decir, se puede convertir índices covariantes en índices contravariantes y viceversa. Los tensores obtenidos así no son iguales al original, pero se dice que son *métricamente equivalentes*.

Regresando a la Definición 2.25 y usando la función $F[Z]$ definida arriba, la divergencia de un campo Y se puede reescribir como $\epsilon^\alpha F[Y]_{\alpha\alpha}$, pero esto es en realidad una abreviación de un método mucho más natural de calcular la *traza* de un operador lineal. La *traza* de un operador lineal T en términos del álgebra lineal es la suma de los elementos de la diagonal (de su matriz asociada); en el lenguaje tensorial un operador lineal es un tensor de orden $(1, 1)$, por lo tanto en términos de sus símbolos este tendría un índice covariante y otro contravariante. Así, la *traza* es la *contracción* de sus índices repetidos, es decir, la suma T^α_α . En general los tensores tienen más de un índice covariante

y contravariante y se tiene que especificar que índice se contrae con algún otro, lo cual generaliza el concepto de traza de un operador lineal.

En la Definición 2.25 la función $F[Y]$ tiene dos índices covariantes, así que no podrían contraerse sus índices a no ser que uno de ellos se suba. Entonces se define el tensor cuyos símbolos en una base son $F[Y]^\alpha_\beta := g^{\alpha\gamma} F[Y]_{\gamma\beta}$. Por lo tanto, la divergencia de un campo vectorial Y se puede reescribir como:

$$\operatorname{div} Y = \operatorname{tr} F[Y] = F[Y]^\alpha_\alpha$$

Expresión que se encuentra en concordancia con la convención de suma de Einstein. Si se calcula la divergencia de un campo Y en términos de un marco local, al ser el marco ortogonal, el tensor métrico obtiene una representación diagonal con ± 1 , recuperándose así la expresión en la Definición 2.25.

Comúnmente se usará la notación en la Definición 2.25 para calcular la contracción o traza de un tensor, salvo que se necesite mayor claridad.

Definición 2.26. Se define el *Hessiano* de una función $f \in C^\infty(M)$ como la segunda derivada covariante de f . $H^f := D(Df)$.

Ya que Df deja de ser una simple función sobre M , debe intervenir la derivada covariante en la definición del Hessiano, más allá de sólo derivadas parciales de segundo orden.

Lema 2.27. El Hessiano de $f \in C^\infty(M)$ es un tensor simétrico sobre M de orden $(0, 2)$ tal que:

$$H^f(X, Y) = XYf - (D_X Y)f = \langle D_X(\operatorname{grad} f), Y \rangle.$$

Demostración. Por definición se tiene que

$$\begin{aligned} H^f(X, Y) &= D_X(D_Y f) = D_X(df(Y)) \\ &= X(df(Y)) - df(D_X Y) \\ &= XYf - (D_X Y)f \end{aligned}$$

Es claro que H^f es tensor en X ; para ver que es un tensor en Y basta con aplicar H^f a gY , donde $g \in \mathfrak{X}(M)$, y usar la regla de Leibniz para $X(gYf)$ en el primer término y la conexión a gY en el segundo.

Ya que D es libre de torsión $XY - YX = D_X Y - D_Y X$, entonces, sustituyendo $D_X Y$ en la ecuación anterior se tiene que H^f es simétrico. Finalmente, como D es compatible con la métrica:

$$\langle D_X(\operatorname{grad} f), Y \rangle = X\langle \operatorname{grad} f, Y \rangle - \langle \operatorname{grad} f, D_X Y \rangle = H^f(X, Y).$$

□

Definición 2.28. El *Laplaciano* Δf de una función $f \in C^\infty(M)$ se define como el negativo de la divergencia del gradiente de f , es decir, $\Delta f := -\text{div}(\text{grad} f)$.

De la definición de divergencia para un campo vectorial se puede ver que el Laplaciano de f es la contracción del Hessiano multiplicada por un menos uno.

2.5. Ecuaciones de estructura

En esta sección derivaré un par de ecuaciones conocidas como *ecuaciones de estructura*, las cuales se escriben en términos de unas 1-formas particulares que definiré a continuación.

Definición 2.29. Sea M una variedad de Lorentz con su conexión de Levi-Civita D ; además, sea $\{X_\alpha\}$ un base local para $\mathfrak{X}(M)$ y Z cualquier campo vectorial. Entonces:

$$D_Z X_\beta = \omega_\beta^\alpha(Z) X_\alpha.$$

Claramente las funciones ω_β^α son $C^\infty(M)$ -lineales. A estas se les conoce como las *1-formas de la conexión*.

Al estar definidas en términos de la conexión, estas 1-formas cargan consigo mismas una noción de la geometría intrínseca de la variedad, en términos de la base local usada para $\mathfrak{X}(M)$. Geométricamente las 1-formas ω_β^α representan que tanto rota X_β inicialmente hacia X_α , en la dirección de Z .

Hasta este punto no he mencionado el producto *exterior* para formas diferenciales ni la *derivada exterior* para estas, es más, sólo he definido 1-formas y no k -formas en general. Esto porque toda esa maquinaria no se usará exhaustivamente, sin embargo, es necesario que defina estas nociones para el caso de 1-formas.

Definición 2.30. Sean ω y η 1-formas sobre una variedad M . El producto *exterior* entre dos 1-formas se define como sigue.

$$(\omega \wedge \eta)(X, Y) := \omega(X)\eta(Y) - \omega(Y)\eta(X) ; \text{ para todo } X, Y \in \mathfrak{X}(M).$$

Donde $\omega \wedge \eta$ es un $(0, 2)$ tensor antisimétrico, es decir, una *2-forma*.

Definición 2.31. La *derivada exterior* de una 1-forma ω se define como:

$$d\omega(X, Y) := X\omega(Y) - Y\omega(X) - \omega([X, Y]) ; \text{ para todo } X, Y \in \mathfrak{X}(M).$$

Donde $d\omega$ es una *2-forma*.

Con esto es posible derivar las *ecuaciones de estructura*.

Teorema 2.32. Sea $\{X_\alpha\}$ una base local de $\mathfrak{X}(M)$, $\{\omega^\alpha\}$ la base dual y $\{\omega_\alpha^\beta\}$ las 1-formas de la conexión en esta base. Entonces se cumplen las siguientes relaciones:

$$d\omega^\alpha = \omega^\gamma \wedge \omega_\gamma^\alpha \quad (2.8)$$

$$d\omega_\beta^\alpha = \omega_\beta^\mu \wedge \omega_\mu^\alpha - \Omega_\beta^\alpha \quad (2.9)$$

Donde $\Omega_\beta^\alpha(X, Y) := \omega^\alpha(R(X, Y)X_\beta)$.

Demostración. Empezaré con la ecuación (2.8). Recordando la definición de derivada exterior se tiene que:

$$\begin{aligned} d\omega^\alpha(Z, Y) &= Z(\omega^\alpha(Y)) - Y(\omega^\alpha(Z)) - \omega^\alpha([Z, Y]) \\ &= ZY^\alpha - YZ^\alpha - \omega^\alpha(D_Z Y) + \omega^\alpha(D_Z X) \end{aligned}$$

Voy a calcular los dos últimos términos de la ecuación por separado.

$$\begin{aligned} \omega^\alpha(D_Z Y) &= \omega^\alpha(ZY^\gamma X_\gamma + Y^\gamma D_Z X_\gamma) \\ &= \omega^\alpha(ZY^\gamma X_\gamma) + \omega^\alpha(Y^\gamma(\omega_\gamma^\mu(Z)X_\mu)) \\ &= XY^\alpha + Y^\gamma \omega_\gamma^\alpha(Z) \\ \omega^\alpha(D_Y Z) &= YZ^\alpha + Z^\gamma \omega_\gamma^\alpha(Z) \end{aligned}$$

Entonces, sustituyendo esto en la expresión de la derivada exterior $d\omega^\alpha$ se obtiene lo siguiente.

$$\begin{aligned} d\omega^\alpha(Z, Y) &= ZY^\alpha - YZ^\alpha - (ZY^\alpha + Y^\gamma \omega_\gamma^\alpha(Z)) \\ &\quad + (YZ^\alpha + Z^\gamma \omega_\gamma^\alpha(Y)) \\ &= \omega^\gamma(Z)\omega_\gamma^\alpha(Y) - \omega^\gamma(Y)\omega_\gamma^\alpha(Z) \end{aligned}$$

Por definición de producto exterior la ecuación (2.8) queda demostrada. Para la demostración de la ecuación (2.9) es necesario que lleve a cabo los siguientes cálculos:

$$\begin{aligned} D_Z D_Y X_\beta &= D(\omega_\beta^\mu(Z)X_\mu) \\ &= Z(\omega_\beta^\mu(Y))X_\mu + \omega_\beta^\mu(Y)D_Z X_\mu \\ &= Z(\omega_\beta^\gamma(Y))X_\gamma + \omega_\beta^\mu(Y)\omega_\mu^\gamma(Z)X_\gamma \\ D_Y D_Z X_\beta &= Y(\omega_\beta^\gamma(Z))X_\gamma + \omega_\beta^\mu(Z)\omega_\mu^\gamma(Y)X_\gamma \\ D_{[Z, Y]} X_\beta &= \omega_\beta^\gamma([Z, Y])X_\gamma \end{aligned}$$

Entonces, usando el tensor de curvatura:

$$\begin{aligned}
 R(Z, Y)X_\beta &= D_{[Z, Y]}X_\beta - D_Z D_Y X_\beta + D_Y D_Z X_\beta \\
 &= [\omega_\beta^\gamma([Z, Y]) - Z(\omega_\beta^\gamma(Y)) - \omega_\beta^\mu(Y)\omega_\mu^\gamma(Z) \\
 &\quad + Y(\omega_\beta^\gamma(Z)) + \omega_\beta^\mu(Z)\omega_\mu^\gamma(Y)]X_\gamma \\
 \omega^\alpha(R(Z, Y)X_\beta) &= \omega_\beta^\alpha([Z, Y]) - Z(\omega_\beta^\alpha(Y)) + Y(\omega_\beta^\alpha(Z)) \\
 &\quad - \omega_\beta^\mu(Y)\omega_\mu^\alpha(Z) + \omega_\beta^\mu(Z)\omega_\mu^\alpha(Y)
 \end{aligned}$$

Se puede apreciar que los primeros tres términos de la ecuación anterior corresponden a la definición de derivada exterior, mientras que los últimos dos a la definición de producto exterior. Con esto la ecuación anterior se reescribe como:

$$\Omega_\beta^\alpha(Z, Y) = -d\omega_\beta^\alpha(Z, Y) + (\omega_\beta^\mu \wedge \omega_\mu^\alpha)(Z, Y)$$

Lo cual termina con la demostración. □

Dado un tensor métrico sobre la variedad M , las 1-formas de la conexión adquieren una propiedad respecto de sus índices. Sea $\{E_\alpha\}$ un marco local ortonormal en M y Y cualquier campo vectorial en $\mathfrak{X}(M)$; entonces:

$$\langle D_Y E_\alpha, E_\beta \rangle = \epsilon^\beta \omega_\alpha^\beta.$$

Usando ahora que D es compatible con la métrica, con $\alpha \neq \beta$ se obtiene:

$$\begin{aligned}
 Y\langle E_\alpha, E_\beta \rangle &= \langle D_Y E_\alpha, E_\beta \rangle + \langle E_\alpha, D_Y E_\beta \rangle \\
 \langle D_Y E_\alpha, E_\beta \rangle &= -\langle E_\alpha, D_Y E_\beta \rangle
 \end{aligned}$$

Por lo tanto se tiene que:

$$\epsilon^\beta \omega_\alpha^\beta = -\epsilon^\alpha \omega_\beta^\alpha. \quad (2.10)$$

Se puede ver que en el caso Riemanniano la ecuación (2.10) implica una antisimetría en las 1-formas de la conexión respecto de sus índices. En este caso dicha antisimetría está supeditada a la causalidad de los vectores en la base aunque en nuestro caso, siendo $\{E_\alpha\}$ un marco ortonormal, existe un sólo vector tipo tiempo en la base siendo los demás tipo espacio; así, siempre que no se tome en cuenta dicho vector la antisimetría se sigue cumpliendo, y en otro caso, se puede hablar de una simetría respecto de sus índices. Además, con esto se tiene que en términos de un marco ortonormal, las 1-formas ω_α^α son iguales a cero.

Cabe mencionar que esta simetría o antisimetría en los índices de las 1-formas de la conexión no es de carácter tensorial, y es más, depende de la base en la que se expresen dichas 1-formas y de la métrica. Por ejemplo, si se tuviera una base nula la ecuación (2.10) no se cumple.

2.6. ¿Toda variedad puede ser de Lorentz?

A lo largo de este capítulo se ha supuesto la existencia de una variedad de Lorentz pero en ningún momento se ha demostrado su existencia. A diferencia del caso Riemanniano, donde toda variedad admite una métrica Riemanniana, no toda variedad admite una métrica de Lorentz y su existencia depende de ciertas características que posea la variedad.

Muchas de las técnicas requeridas para demostrar la existencia de métricas de Lorentz en una variedad quedan fuera del interés de esta tesis y el enunciar este teorema aquí es por dar una exposición un tanto más completa del tema, así que en realidad esta sección está dedicada a hacer un bosquejo de la demostración del teorema de existencia de métricas de Lorentz, más que una demostración completa y formal.

Antes de enunciar el teorema, quisiera explicar que significa que una variedad de Lorentz tenga una *orientación temporal*. Una pregunta usual en el contexto de las variedades es saber si todos los espacios tangentes tienen la misma orientación; así, tratándose de variedades de Lorentz es usual preguntarse si todos los espacios tangentes comparten la misma orientación temporal.

Entonces, sea $\mathcal{T}$ una función sobre una variedad de Lorentz M tal que a cada $p \in M$ le asocia un cono tipo tiempo $\mathcal{T}_p$ en $T_p M$. Se dice que tal función $\mathcal{T}$ es C^∞ si para todo $p \in M$ existe un abierto U que lo contenga y un campo vectorial Y definido en U tal que Y es C^∞ y $Y(q) \in \mathcal{T}_q$ para todo $q \in U$; a esta función $\mathcal{T}$ se dice que es una *orientación temporal* sobre M . Por lo tanto, a toda variedad que admita una orientación de este tipo se dice que es *orientable temporalmente*.

Lema 2.33. Una variedad de Lorentz es orientable temporalmente si y sólo si existe un campo vectorial $Y \in \mathfrak{X}(M)$ que sea tipo tiempo.

Demostración. Supongamos primero que tal campo $Y \in \mathfrak{X}(M)$ existe, entonces, tomando en cada espacio tangente el cono tipo tiempo en el cual se encuentra $Y(p)$ uno obtiene una orientación temporal.

Supongamos ahora que M admite una orientación temporal $\mathcal{T}$; por definición, para cada $p \in M$ existe un abierto U tal que $p \in U$ y existe un campo vectorial C^∞ tal que $Y_U(q)$ se encuentra en $\mathcal{T}_q$ para todo $q \in U$. Con esto, los abiertos que salen de la definición constituyen una cubierta de M . Sea $\{\theta_\alpha | \alpha \in A\}$ una partición de la unidad sobre M subordinada a la cubierta antes dicha $\{U_\alpha\}$; entonces, como las funciones θ_α son positivas y los conos tipo tiempo son conexos el campo vectorial $Y := \sum \theta_\alpha Y_{U_\alpha}$ es elemento de $\mathfrak{X}(M)$.

□

Con esto es suficiente preámbulo para enunciar el teorema de existencia de métricas de Lorentz, el cual es el tema principal de esta sección.

Teorema 2.34. Sea M una variedad. Las siguientes afirmaciones son equivalentes:

1. Existe una métrica de Lorentz en M .
2. M es orientable temporalmente.
3. Existe un campo vectorial sobre M tal que este no se anula.
4. M es una variedad no compacta, o M es compacta y tiene característica de Euler $\chi(M) = 0$.

Para que una variedad sea orientable temporalmente, claramente se requiere de antemano que sea una variedad de Lorentz y por tanto se tiene que 2 implica 1, lo cual denotaré de aquí en adelante $2 \rightarrow 1$.

Por el Lema 2.33 se tiene ahora que $2 \rightarrow 3$ ya que un campo vectorial tipo tiempo nunca se hace cero (hay que recordar que el vector cero es tipo espacio), pero no se puede aplicar el si y sólo si de este lema ya que el campo al cual se hace referencia en el punto número tres puede tener cualquier causalidad. Para demostrar que $3 \rightarrow 2$ necesito del siguiente lema.

Lema 2.35. Sea Z un campo unitario sobre una variedad Riemanniana con métrica $\langle X, Y \rangle_R$ positiva definida para todo campo vectorial X y Y sobre M . Entonces, el tensor

$$\langle X, Y \rangle := \langle X, Y \rangle_R - 2\langle X, Z \rangle_R \langle Y, Z \rangle_R ; \text{ para todo } X, Y \in \mathfrak{X}(M)$$

es una métrica de Lorentz para M . Además, M es orientable temporalmente.

Demostración. Por un lado se tiene que Z es ya un campo unitario, entonces, localmente existen campos $\{E_j\}_{j=1}^n$ tal que $\{Z, E_j\}$ en un marco respecto de la métrica de Riemann. Ya que $\langle E_j, Z \rangle_R = 0$ para toda E_j , entonces

$$\langle E_i, E_j \rangle = \langle E_i, E_j \rangle_R = \delta_{ij},$$

$$\langle E_i, Z \rangle = \langle E_i, Z \rangle_R = 0.$$

Además $\langle Z, Z \rangle = -1$. Claramente esta función es un tensor debido a la $C^\infty(M)$ -linealidad de la métrica de Riemann. Por lo tanto este tensor es una métrica de Lorentz sobre M , siendo Z tipo tiempo. Haciendo uso del Lema 2.33 se tiene que M es orientable temporalmente.

□

Como toda variedad M admite una métrica Riemanniana, siendo $X \in \mathfrak{X}(M)$ cualquier campo se puede normalizar como $\tilde{X} := \frac{X}{\|X\|}$ y, aplicándole el Lema 2.35 al campo $\tilde{X}$ se tiene finalmente que $2 \leftrightarrow 3$.

La equivalencia $3 \leftrightarrow 4$ no la expondré aquí, este es un resultado de la topología algebraica cuya referencia se puede encontrar en [15, p. 188].

Por lo tanto sólo falta hacer la implicación $1 \rightarrow 4$ para cerrar el círculo de equivalencias. De esta implicación tan sólo haré un bosquejo, para los detalles se puede ver [13].

Proposición 2.36. Sea $\kappa : \Theta \rightarrow M$ una función que sea *dos a uno* del conjunto Θ sobre la variedad M . Además, sea $\mathcal{A}$ el conjunto de todas las funciones $\lambda : U \rightarrow \Theta$, donde U es un abierto de M , tales que:

- $\kappa \circ \lambda = id_\Theta$ para toda λ en $\mathcal{A}$.
- Si $\lambda(p) = \mu(p)$ con $\lambda, \mu \in \mathcal{A}$. Entonces $\lambda = \mu$ en un abierto de M alrededor de p .
- Todo elemento de Θ es la imagen de alguna función $\lambda \in \mathcal{A}$.

Entonces existe una única forma de hacer de Θ una variedad tal que κ sea un *doble cubriente* C^∞ de M y cada función en $\mathcal{A}$ una *sección local* C^∞ de κ .

Entonces, sea $\tilde{M}$ el conjunto de todos los conos tipo tiempo en los espacios tangentes de una variedad de Lorentz M , así, la función $\kappa : \tilde{M} \rightarrow M$ que asocia a cada punto $p \in M$ ambos conos tipo tiempo en $T_p M$ es sobre y dos a uno. Ahora, sean $\mathcal{T}_Y$ y $\mathcal{T}_Z$ dos orientaciones temporales sobre M y Y, Z campos C^∞ definidos en un abierto $U \subset M$ tales que se encuentran en dichas orientaciones respectivamente; si $\mathcal{T}_Y(p) = \mathcal{T}_Z(p)$ por la Proposición 1.13 se tiene que $\langle Y_p, Z_p \rangle < 0$ y, como esta expresión representa una función C^∞ en U , las orientaciones son iguales en este abierto. Por lo tanto, la función κ y las orientaciones de M cumplen con los requisitos de la Proposición 2.36 haciendo de $\tilde{M}$ una variedad y de κ un doble cubriente de M .

La función κ restringida a un abierto es un difeomorfismo, por tanto un difeomorfismo local. Esta propiedad hace que el pullback de la métrica en M sea una métrica de Lorentz a su vez en $\tilde{M}$ haciendo de esta una variedad de Lorentz.

Lema 2.37. Sea M una variedad de Lorentz. Entonces su espacio cubriente $\tilde{M}$, definido arriba, es orientable temporalmente.

Regresando a la implicación que se quiere mostrar, si M es orientable temporalmente se implica 4. Si M no es orientable temporalmente, el Lema 2.37 dice que la doble cubierta de M , $\tilde{M}$, es orientable temporalmente y por lo tanto se cumple 4 para esta variedad.

Como κ (la función cubriente) es C^∞ , si $\tilde{M}$ es compacta M lo será también. Si M es compacta toda cubierta abierta de M será también una cubierta abierta de M y lo mismo para sus subcubiertas; como M es compacto siempre existe la subcubierta finita de este y como la imagen inversa de cada uno de los elementos de la subcubierta tendrá sólo dos copias en $\tilde{M}$, la subcubierta finita también existirá para $\tilde{M}$ en todos los casos. Por lo tanto, M es compacta si y sólo si $\tilde{M}$ es compacta.

Entonces, si $\tilde{M}$ es no compacta no se tiene nada que demostrar. Si $\tilde{M}$ es compacta M también lo será y $\chi(\tilde{M}) = 0$; además $2\chi(M) = \chi(\tilde{M})$ (véase [10, p. 256]), lo cual termina por demostrar que $1 \rightarrow 4$.

Capítulo 3

Geometría Extrínseca

Hasta ahora sólo se ha expuesto la geometría intrínseca de una variedad, es decir, todo lo que se puede decir de la geometría de esta en términos de ella misma sin asumir que esta se pueda encontrar *inmersa* en otra variedad.

Que una variedad M sea *subvariedad* de una variedad N deja de ser una propiedad intrínseca de la variedad M , y su geometría depende de como esta se encuentra *inmersa* en la variedad N . Un ejemplo clásico es el de la *botella de Klein* que es una variedad de dimensión 2 la cual es imposible *encajar* en $\mathbb{R}^3$ a diferencia de la esfera S^2 , lo cual hace que su geometría desde la perspectiva del espacio ambiente sea distinta.

Entonces uno puede hablar de dos tipos de subvariedades; aquellas que se encuentran *inmersas*, como la *botella de Klein*, las cuales no se puede asegurar que no tengan intersecciones consigo mismas, y aquellas que están *encajadas* en su espacio ambiente y forman la noción más común de *subvariedad*.

Definición 3.1. Sea $\phi : M \rightarrow N$ una función entre variedades. Se dice que ϕ es una *inmersión* si la diferencial $d\phi_p$ es inyectiva para todo $p \in M$; además, si ϕ es un *homeomorfismo* con la imagen $\phi(M)$ se dice que esta función es un *encaje*.

Definición 3.2. Sea $\phi : M \rightarrow N$ una inmersión entre variedades de Lorentz. Se dice que ϕ es una inmersión *isométrica* si

$$\langle u, v \rangle_p = \langle d\phi_p(u), d\phi_p(v) \rangle_{\phi(p)}$$

para todo $u, v \in T_p M$ y $p \in M$.

De ahora en adelante siempre que se hable de una inmersión o un encaje, siempre debe entenderse que se habla de uno que sea isométrico.

Hasta ahora no se ha tomado en cuenta si M es un subconjunto de N , y no se ha especificado cuando M será una *subvariedad semi-Riemanniana*.

Si ϕ fuera un encaje entonces la imagen $\phi(M)$ es un objeto geométrico tal que este no tiene intersecciones consigo mismo, el cual vendría siendo nuestro ideal a definir, pero ya que nuestro estudio de la geometría siempre es local se puede ser un poco más laxo en esto recordando que toda inmersión es un encaje local.

Definición 3.3. Sea ϕ una inmersión isométrica de M en N y g la métrica de N . Si el *pullback* $\phi^*(g)$ es una métrica *semi-Riemanniana* sobre M se tienen los siguientes casos:

- M es *tipo tiempo* si $\phi^*(g)$ es no degenerado y tiene índice $\nu = 1$.
- M es *tipo espacio* si $\phi^*(g)$ es no degenerado y tiene índice $\nu = 0$.

Si $\phi^*(g)$ es una métrica degenerada para todo $p \in M$ se dice que M es *tipo luz* o *nulo*.

Tomemos de aquí en adelante el caso no degenerado a no ser que se especifique lo contrario, es decir, entiéndase de aquí en adelante por *sudvariedad* una que sea *semi-Riemanniana*.

Como toda inmersión ϕ es un encaje local, existe un abierto U alrededor de $p \in M$ tal que $\phi(U)$ es un objeto geométrico razonable. Entonces se puede asociar a cada vector $u \in T_p M$ un único vector $d\phi_p(u) \in T_p N$ para todo p ; así, $T_p M$ será siempre un subespacio no nulo de $T_p N$. Con esto, si M es tipo tiempo el espacio tangente $T_p M$ es un subespacio tipo tiempo de $T_p N$ para toda $p \in M$, lo mismo para los dos casos restantes; si M es tipo espacio $T_p M$ es tipo espacio y si M es nulo $T_p M$ es nulo para toda $p \in M$.

Entonces, ya que $T_p M$ será un subespacio no degenerado de $T_p N$ el complemento ortogonal tampoco será degenerado, obteniéndose que:

$$T_p N = T_p M \oplus T_p^\perp M.$$

Al complemento ortogonal $T_p^\perp M$ se le conoce como el *subespacio normal* a M en p y a los vectores en este conjunto se les conoce como vectores normales a M en p , claramente los vectores que se encuentren en $T_p M$ son los vectores tangentes a M .

Para todo vector $v \in T_p N$ se tiene la descomposición

$$v = \tan v + \text{nor} v.$$

Donde las funciones $\tan : T_p N \rightarrow T_p M$ y $\text{nor} : T_p N \rightarrow T_p^\perp M$ son las proyecciones usuales, las cuales son $\mathbb{R}$ -lineales. Estas funciones se extienden de manera natural sobre $\mathfrak{X}(N)|_M$ tomando el conjunto de los campos C^∞ tangentes a M en un caso, y el conjunto de los campos C^∞ normales a M en el

otro, $\mathfrak{X}(M)$ y $\mathfrak{X}^\perp(M)$ respectivamente. Estas funciones extendidas en los conjuntos mencionados son $C^\infty(M)$ -lineales, esto se logra extendiendo funciones C^∞ sobre M a funciones C^∞ en un abierto de N .

3.1. La conexión inducida

Sea M una subvariedad en una variedad de Lorentz N y sean X y Y dos campos vectoriales tangentes a M , es decir, elementos de $\mathfrak{X}(M)$, además, sea U un abierto en N y $V = M \cap U$ abierto en M ; se dice que un campo $\tilde{X} \in \mathfrak{X}(N)$ es una extensión de $X \in \mathfrak{X}(M)$ si la restricción $\tilde{X}|_M = X$.

Aunque la subvariedad M es una variedad en sí misma, esta podría tener su conexión de Levi-Civita y no guardar relación con la conexión de N . Ya que esta está jugando precisamente el papel de subvariedad de N se espera que su conexión, y así muchas entidades geométricas, se mantenga en relación con la conexión de N . Es por eso que se necesita hablar de las extensiones de campos para poder así aplicar primero las operaciones entre estos en M ; por ejemplo el corchete $[\tilde{X}, \tilde{Y}]|_M$, este sólo depende de los valores de $\tilde{X}$ y $\tilde{Y}$ en un punto $p \in M$. Por lo tanto el resultado no dependerá de la extensión usada, es decir.

$$[\tilde{X}, \tilde{Y}]|_M = [X, Y] \quad (3.1)$$

Además, ya que la construcción del corchete necesita solamente de las direcciones de los vectores X_p y Y_p el campo resultante será tangente a M , es decir, $[X, Y] \in \mathfrak{X}(M)$. Para ver que la conexión D de N , aplicada a las extensiones $\tilde{X}$ y $\tilde{Y}$ de los campos $X, Y \in \mathfrak{X}(M)$, no depende de la extensión usada observemos lo siguiente. Sean $\tilde{X}_1$ y $\tilde{X}_2$ extensiones distintas del campo X , ya que estas tienen el mismo valor restringidas a M se tiene que $\tilde{X}_1 - \tilde{X}_2 = 0$, con esto $D_{\tilde{X}_1 - \tilde{X}_2} \tilde{Y}|_M = 0$. Ahora sean $\tilde{Y}_1$ y $\tilde{Y}_2$ dos extensiones distintas del campo Y , lo que se quiere demostrar es que la diferencia $D_{\tilde{X}} \tilde{Y}_1 - D_{\tilde{X}} \tilde{Y}_2$ se anula al restringirla a M . Entonces, sumando y restando los factores $D_{\tilde{Y}_1} \tilde{X}$ y $D_{\tilde{Y}_2} \tilde{X}$ se tiene, usando que D es libre de torsión, que

$$D_{\tilde{X}} \tilde{Y}_1 - D_{\tilde{X}} \tilde{Y}_2 = [\tilde{X}, \tilde{Y}_1] - [\tilde{X}, \tilde{Y}_2] + D_{\tilde{Y}_1 - \tilde{Y}_2} \tilde{X}.$$

Por lo tanto, restringiendo esto a M , por la ecuación (3.1) y el argumento anterior esta expresión se anula, así la conexión no depende de las extensiones usadas y definir esta en $\mathfrak{X}(M)$ tiene sentido, aunque esto no asegura que $D_{\tilde{X}} \tilde{Y}$ sea un vector tangente a M , por eso es necesaria la siguiente definición.

Definición 3.4. Sean $\tilde{X}$ y $\tilde{Y}$ las extensiones sobre N de los campos $X, Y \in \mathfrak{X}(M)$ respectivamente, donde M es subvariedad de N y D es la conexión de

la última. Entonces, restringiendo D a M :

$$D_{\tilde{X}}\tilde{Y} = \nabla_X Y + II(X, Y).$$

Donde $\nabla_X Y := \tan D_{\tilde{X}}\tilde{Y}$ y $II(X, Y) := \text{nor} D_{\tilde{X}}\tilde{Y}$. A esta ecuación se le conoce como la *fórmula de Gauss*.

Lo importante de la fórmula de Gauss es que nos muestra las dos partes importantes sobre la geometría de la subvariedad M . La parte tangente estará relacionada con la geometría intrínseca de M , mientras que la parte normal con la geometría de esta desde el punto de vista de N .

Teorema 3.5. Sea M una subvariedad de N y D la conexión de esta última. Entonces:

1. La función $\nabla : \mathfrak{X}(M) \times \mathfrak{X}(M) \rightarrow \mathfrak{X}(M)$ definida en 3.4 es la conexión de Levi-Civita de M .
2. La función $II : \mathfrak{X}(M) \times \mathfrak{X}(M) \rightarrow \mathfrak{X}^\perp(M)$ definida en 3.4 es $C^\infty(M)$ -bilineal y simétrica. A esta se le conoce como la *segunda forma fundamental* de M .

Demostración. Para usar la fórmula de Gauss es necesario, como se ha hecho hasta ahora, extender los campos tangentes a M y las funciones definidas sobre esta; todo esto será válido en un abierto de N lo cual no invalida la demostración por hacer ya que tratamos con una cuestión local.

Sean $X, Y, Z \in \mathfrak{X}(M)$ y $f \in C^\infty(M)$. Extendiendo los campos y las funciones a abiertos de N , demostrar los dos primeros puntos de la Definición 2.9 se reduce primero a restringir la conexión a M para que las funciones se encuentren en $C^\infty(M)$, y después proyectar sobre la parte tangente a M . Para el tercer punto se tiene lo siguiente:

$$D_X(fY) = XfY + fD_XY.$$

Proyectando esta ecuación sobre la parte tangente a M se obtiene finalmente el resultado. Por lo tanto ∇ es en verdad una conexión sobre M , sólo falta ver que es la conexión de Levi-Civita.

Usando la fórmula de Gauss y el hecho de que D es la conexión de Levi-Civita en N se obtiene:

$$[\tilde{X}, \tilde{Y}] = \nabla_X Y - \nabla_Y X + II(X, Y) - II(Y, X).$$

Por la ecuación (3.1) se tiene que el corchete en la parte izquierda de la pasada ecuación es tangente a M , entonces el lado derecho de la pasada ecuación debe ser tangente a M también, obteniendo así

$$[X, Y] = \nabla_X Y - \nabla_Y X.$$

Esto quiere decir que ∇ es libre de torsión, además $II(X, Y) - II(Y, X) = 0$, lo cual implica que II es simétrica.

Como D es compatible con la métrica se tiene:

$$\tilde{X}\langle\tilde{Y}, \tilde{Z}\rangle = \langle D_{\tilde{X}}\tilde{Y}, \tilde{Z}\rangle + \langle\tilde{Y}, D_{\tilde{X}}\tilde{Z}\rangle.$$

Restringiendo esta ecuación a M , como el producto escalar sólo depende de $p \in M$, el resultado no depende de las extensiones usadas. Por otro lado:

$$\langle D_{\tilde{X}}\tilde{Y}, Z\rangle = \langle \nabla_X Y, Z\rangle + \langle II(X, Y), Z\rangle = \langle \nabla_X Y, Z\rangle.$$

Entonces, por un proceso similar para $\langle Y, D_{\tilde{X}}\tilde{Z}\rangle$ se obtiene lo siguiente.

$$X\langle Y, Z\rangle = \langle \nabla_X Y, Z\rangle + \langle Y, \nabla_X Z\rangle.$$

Y así, ∇ es compatible con la métrica y por el Teorema 2.14 se tiene que esta es la conexión de Levi-Civita de M . Solo falta mostrar que II es $C^\infty(M)$ -bilineal.

Para la conexión D se tiene que $D_{\tilde{f}\tilde{X}}\tilde{Y} = \tilde{f}D_{\tilde{X}}\tilde{Y}$; restringiendo a M y usando la fórmula de Gauss en ambos lados de la ecuación

$$\nabla_{fX}Y + II(fX, Y) = f(\nabla_X Y + II(X, Y)).$$

Por lo tanto, usando la simetría de II , esta función es $C^\infty(M)$ -bilineal. $\square$

El operador de forma

Sea ξ un campo vectorial normal a M , es decir, $\xi \in \mathfrak{X}^\perp(M)$; extendiendo ξ sobre un abierto de N se le puede aplicar el operador D_X , donde $X \in \mathfrak{X}(M)$ y D es la conexión de N , como se ha hecho hasta ahora. De la misma forma el resultado no dependerá de las extensiones usadas y el resultado al restringir $D_X\xi$ a M sólo dependerá de estos campos. Al igual que el caso de la fórmula de Gauss, esta operación se puede descomponer en una parte tangente y una normal de la siguiente forma.

$$D_X\xi = -S_\xi(X) + \nabla_X^\perp\xi \quad (3.2)$$

Donde $-S_\xi(X) := \tan(D_X\xi)$ y $\nabla_X^\perp\xi := \text{nor}(D_X\xi)$. A la ecuación (3.2) se le conoce como la *fórmula de Weingarten*.

La fórmula de Weingarten tiene información acerca de la variación del espacio tangente respecto de una dirección normal, es decir, como se curva la subvariedad desde el punto de vista no del espacio ambiente en general, sino

en una dirección de éste que además es normal a la subvariedad. En el caso que el espacio ambiente sea $\mathbb{R}^3$ y la subvariedad una superficie, la fórmula de Weingarten tiene la misma información que la diferencial de la función esférica de Gauss.

Teorema 3.6. Sea M una subvariedad de N ; de la ecuación (3.2) se derivan los siguientes resultados:

1. $S_\xi(X)$ es $C^\infty(M)$ -bilineal en ambas entradas. A este operador se le conoce como el *operador de forma*.
2. Sean $\xi \in \mathfrak{X}^\perp(M)$ y $X, Y \in \mathfrak{X}(M)$. Entonces:

$$\langle II(X, Y), \xi \rangle = \langle S_\xi(X), Y \rangle. \quad (3.3)$$

3. La función $\nabla^\perp$ es una conexión sobre el haz normal $T^\perp M$, y es compatible con la métrica inducida de N sobre $T^\perp M$.

Demostración. 1. Claramente el operador de forma cumplirá con abrir sumas en sus dos argumentos ya que la conexión D de N lo hace. Sean f y g funciones en $C^\infty(M)$ y $X \in \mathfrak{X}(M)$, $\xi \in \mathfrak{X}^\perp(M)$; usando la fórmula de Weingarten se tiene lo siguiente.

$$\begin{aligned} D_{fX}(g\xi) &= fD_X(g\xi) \\ &= f(Xg\xi + gD_X\xi) \\ &= fXg\xi - fgS_\xi(X) + fg\nabla_X^\perp\xi \end{aligned}$$

Ya que $D_{fX}(g\xi) = -S_{g\xi}(fX) + \nabla_{fX}^\perp(g\xi)$, sumando las partes tangente y normal de la ecuación se obtiene el siguiente par de ecuaciones.

$$\begin{aligned} S_{g\xi}(fX) &= fgS_\xi(X) \\ \nabla_{fX}^\perp(g\xi) &= f(Xg\xi + g\nabla_X^\perp\xi) \end{aligned}$$

Con esto queda demostrado no sólo que el operador de forma es $C^\infty(M)$ -bilineal, también que el operador $\nabla^\perp$ es una conexión.

2. Sean $\xi \in \mathfrak{X}^\perp(M)$ y $X, Y \in \mathfrak{X}(M)$ campos arbitrarios; como $\langle Y, \xi \rangle = 0$ se tiene lo siguiente.

$$\begin{aligned} 0 = X\langle Y, \xi \rangle &= \langle D_X Y, \xi \rangle + \langle Y, D_X \xi \rangle \\ &= \langle \nabla_X Y, \xi \rangle + \langle II(X, Y), \xi \rangle - \langle Y, S_\xi(X) \rangle + \langle Y, \nabla_X^\perp \xi \rangle \\ &= \langle II(X, Y), \xi \rangle - \langle Y, S_\xi(X) \rangle \end{aligned}$$

Usando la simetría de la métrica se obtiene finalmente la ecuación (3.3).

3. Ya se demostró con anterioridad que el operador $\nabla^\perp$ tiene las propiedades de una conexión, pero ¿dónde se encuentra definida esta? De acuerdo con la Definición 2.9 $\nabla^\perp$ no podría ser una conexión, sin embargo, la Definición 2.9 es un caso particular de una definición más general de conexión la cual se puede ver en [11].

Entonces, sean $\xi, \zeta \in \mathfrak{X}^\perp(M)$ y $X \in \mathfrak{X}(M)$, por la fórmula de Weingarten:

$$D_X \xi = -S_\xi(X) + \nabla_X^\perp \xi ; \quad D_X \zeta = -S_\zeta(X) + \nabla_X^\perp \zeta.$$

Para cada una de estas formulas se tiene lo siguiente.

$$\begin{aligned} \langle D_X \xi, \zeta \rangle &= -\langle S_\xi(X), \zeta \rangle + \langle \nabla_X^\perp \xi, \zeta \rangle = \langle \nabla_X^\perp \xi, \zeta \rangle \\ \langle \xi, D_X \zeta \rangle &= -\langle \xi, S_\zeta(X) \rangle + \langle \xi, \nabla_X^\perp \zeta \rangle = \langle \xi, \nabla_X^\perp \zeta \rangle \end{aligned}$$

Sumando ambas ecuaciones y usando que D es compatible con la métrica:

$$\langle D_X \xi, \zeta \rangle + \langle \xi, D_X \zeta \rangle = X \langle \xi, \zeta \rangle = \langle \nabla_X^\perp \xi, \zeta \rangle + \langle \xi, \nabla_X^\perp \zeta \rangle.$$

Ya que la métrica de N coincide claramente con la métrica inducida sobre $T^\perp M$, se puede decir ahora que $\nabla^\perp$ es compatible con la métrica. $\square$

De la ecuación (3.3) se desprende el siguiente resultado.

Corolario 3.7. Para cada $p \in M$, el operador de forma es auto-adjunto.

Demostración. El teorema anterior asegura que $\langle S_\xi(X), Y \rangle = \langle II(X, Y), \xi \rangle$ donde $\xi \in \mathfrak{X}^\perp(M)$ y $X, Y \in \mathfrak{X}(M)$. Entonces:

$$\langle X, S_\xi(Y) \rangle = \langle S_\xi(Y), X \rangle = \langle II(Y, X), \xi \rangle = \langle II(X, Y), \xi \rangle.$$

Por lo tanto $\langle S_\xi(X), Y \rangle = \langle X, S_\xi(Y) \rangle$, es decir, S_ξ es un operador auto-adjunto. $\square$

El operador de forma es un tensor de rango $(1, 2)$ en un sentido más general al que se ha manejado hasta ahora, ya que su entrada contravariante no actúa sobre el espacio dual a $\mathfrak{X}(M)$, este actúa sobre el espacio dual a $\mathfrak{X}^\perp(M)$. Por lo tanto, si se contraen los dos índices covariantes de la segunda forma fundamental nos quedamos con las componentes de un vector, el cual por supuesto es normal a M .

Definición 3.8. Sea M una subvariedad y sea $\text{tr}II$ la contracción de ambos índices covariantes de II . Al campo vectorial $H \in \mathfrak{X}(M)$ definido como:

$$H := \left(\frac{1}{n} \right) \text{tr}II$$

se le conoce como el *vector de curvatura media* de M .

Definición 3.9. Sea M una subvariedad. Se dice que M es mínima si su vector de curvatura media se anula, es decir, $H \equiv 0$.

3.2. Ecuaciones fundamentales

A lo largo de esta tesis serán de utilidad un par de ecuaciones, las cuales se conocen como *ecuaciones fundamentales*.

Antes de empezar es prudente hacer un pequeño cambio en la notación. Al tensor de curvatura descrito en términos de la conexión de la variedad ambiente N se le denotará R^D , mientras que al tensor de curvatura intrínseco a la subvariedad M , el cual está determinado por la conexión ∇ , se le denotará R^∇ .

Proposición 3.10. El tensor de curvatura cumple con la siguiente ecuación:

$$\begin{aligned} \langle R^D(X, Y)Z, W \rangle &= \langle R^\nabla(X, Y)Z, W \rangle - \langle II(Y, W), II(X, Z) \rangle \\ &\quad + \langle II(X, W), II(Y, Z) \rangle \end{aligned} \quad (3.4)$$

La cual es conocida como la *ecuación de Gauss*.

Demostración. Sean $X, Y, Z \in \mathfrak{X}(M)$; usando la fórmula de Gauss se puede hacer el siguiente desarrollo.

$$\begin{aligned} R^D(X, Y)Z &= \nabla_{[X, Y]}Z + II([X, Y], Z) - D_X(\nabla_Y Z + II(Y, Z)) \\ &\quad + D_Y(\nabla_X Z + II(X, Z)) \end{aligned}$$

Utilizando ahora la fórmula de Weingarten, la fórmula de Gauss y agrupando los términos con el operador ∇ se obtiene la siguiente relación.

$$\begin{aligned} R^D(X, Y)Z &= R^\nabla(X, Y)Z + S_{II(Y, Z)}(X) - S_{II(X, Z)}(Y) \\ &\quad + II([X, Y], Z) - II(X, \nabla_Y Z) + II(Y, \nabla_X Z) \\ &\quad + \nabla_Y^\perp II(X, Z) - \nabla_X^\perp II(Y, Z) \end{aligned} \quad (3.5)$$

Sea W otro campo en $\mathfrak{X}(M)$; entonces:

$$\begin{aligned} \langle R^D(X, Y)Z, W \rangle &= \langle R^\nabla(X, Y)Z, W \rangle - \langle S_{II(X, Z)}(Y), W \rangle \\ &\quad + \langle S_{II(Y, Z)}(X), W \rangle \end{aligned}$$

Usando la ecuación (3.3) se obtiene finalmente la ecuación deseada.

$$\begin{aligned} \langle R^D(X, Y)Z, W \rangle &= \langle R^\nabla(X, Y)Z, W \rangle - \langle II(Y, W), II(X, Z) \rangle \\ &\quad + \langle II(X, W), II(Y, Z) \rangle \end{aligned} \quad (3.6)$$

□

Antes de seguir adelante, es necesario que defina una nueva operación sobre la segunda forma fundamental; esta nace de la necesidad de derivar tensores usando la conexión normal y la conexión de Levi-Civita de la sub-variedad M .

$$(\bar{\nabla}_Y II)(X, Z) := \nabla_Y^\perp II(X, Z) - II(\nabla_Y X, Z) - II(X, \nabla_Y Z)$$

Proposición 3.11. El tensor de curvatura cumple con la siguiente ecuación:

$$(R^D(X, Y)Z)^\perp = (\bar{\nabla}_Y II)(X, Z) - (\bar{\nabla}_X II)(Y, Z) \quad (3.7)$$

La cual se conoce como la *ecuación de Codazzi*.

Demostración. La ecuación (3.5) se puede reescribir como sigue.

$$\begin{aligned} R^D(X, Y)Z &= R^\nabla(X, Y)Z + S_{II(Y, Z)}(X) - S_{II(X, Z)}(Y) \\ &\quad + (\bar{\nabla}_Y II)(X, Z) - (\bar{\nabla}_X II)(Y, Z) \end{aligned}$$

Si nos quedamos sólo con la componente normal de la pasada ecuación, se obtiene la ecuación de Codazzi.

$$(R^D(X, Y)Z)^\perp = (\bar{\nabla}_Y II)(X, Z) - (\bar{\nabla}_X II)(Y, Z) \quad (3.8)$$

□

A las ecuaciones (3.4) y (3.7) se les conoce como las ecuaciones fundamentales.

Del Corolario 2.23, si la variedad N tiene curvatura constante C su tensor de curvatura obtiene la expresión:

$$R^D(X, Y)Z = C [\langle Z, X \rangle Y - \langle Z, Y \rangle X]$$

Claramente $R^D(X, Y)Z$ es un campo vectorial tangente a M ya que este es combinación lineal de elementos en $\mathfrak{X}(M)$. Por lo tanto $(R^D(X, Y)Z)^\perp = 0$ si la variedad ambiente tiene curvatura constante, y la ecuación de Codazzi se reescribe como sigue.

$$(\bar{\nabla}_Y II)(X, Z) = (\bar{\nabla}_X II)(Y, Z)$$

Tomando en cuenta las técnicas usadas para derivar tensores, se puede definir la derivada del operador de forma como:

$$(\nabla_Y S_\xi)(X) := \nabla_Y (S_\xi(X)) - S_\xi(\nabla_Y X) \quad (3.9)$$

Queda claro que este nuevo operador es un tensor en Y , pero ¿será también un operador en X ? Sea $f \in C^\infty(M)$.

$$\begin{aligned} (\nabla_Y S_\xi)(fX) &= \nabla_Y (S_\xi(fX)) - S_\xi(\nabla_Y fX) \\ &= (Yf)S_\xi(X) + f\nabla_Y (S_\xi(X)) - S_\xi(YfX + f\nabla_Y X) \end{aligned}$$

El primer término se cancelará con uno de los términos que surgen del tercero usando la linealidad de S_ξ en la última ecuación. Con esto queda claro que el operador definido por la ecuación 3.9 es un tensor.

Con esto se obtiene el siguiente teorema.

Teorema 3.12. Sea M una subvariedad de N , donde esta última tiene curvatura constante. Entonces:

$$(\nabla_X S_\xi)(Y) = (\nabla_Y S_\xi)(X) \quad (3.10)$$

Demostración. Por definición se tiene que:

$$(\bar{\nabla}_Y II)(X, Z) = \nabla_Y^\perp II(X, Z) - II(\nabla_Y X, Z) - II(X, \nabla_Y Z)$$

Entonces, sea ξ un campo normal unitario a M .

$$\begin{aligned} \langle \nabla_Y^\perp II(X, Z), \xi \rangle &= \langle II(\nabla_Y X, Z), \xi \rangle - \langle II(X, \nabla_Y Z), \xi \rangle \\ \langle \nabla_Y^\perp II(X, Z), \xi \rangle &= \langle S_\xi(\nabla_Y X), Z \rangle - \langle S_\xi(X), \nabla_Y Z \rangle \end{aligned} \quad (3.11)$$

Como ∇ es compatible con la métrica se obtiene la siguiente relación.

$$\begin{aligned} -\langle S_\xi(X), \nabla_Y Z \rangle &= \langle \nabla_Y S_\xi(X), Z \rangle - Y \langle S_\xi(X), Z \rangle \\ &= \langle \nabla_Y S_\xi(X), Z \rangle - Y \langle II(X, Z), \xi \rangle \end{aligned}$$

Sustituyendo esto en la ecuación (3.11) y usando la ecuación (3.9), la primera se reescribe como:

$$\langle \nabla_Y^\perp II(X, Z), \xi \rangle - Y \langle II(X, Z), \xi \rangle + \langle (\nabla_Y S_\xi)(X), Z \rangle$$

Los primeros dos términos de la ecuación anterior, como $\nabla^\perp$ es compatible con la métrica, se pueden cambiar por $\langle II(X, Z), \nabla_Y^\perp \xi \rangle$. Como ξ es un campo unitario la derivada $\nabla_Y^\perp \xi$ se hace cero, anulando el producto anterior. Haciendo el mismo procedimiento para $(\bar{\nabla}_X II)(Y, Z)$ se obtiene:

$$\langle (\nabla_Y S_\xi)(X), Z \rangle = \langle (\nabla_X S_\xi)(Y), Z \rangle$$

□

Esta ecuación se derivó de la ecuación de Codazzi, lo cual la hace equivalente en el caso que la variedad ambiente tenga curvatura constante.

La fórmula de Beltrami

La *fórmula de Beltrami* es una ecuación que cumplen todas las subvariedades del espacio de Minkowski, así como en el caso clásico las subvariedades de $\mathbb{R}^n$.

Antes es necesario definir el Laplaciano para una función $F : M \rightarrow \mathbb{R}_1^n$. Aprovechando que uno tiene la noción de coordenadas en un espacio vectorial, definiré el Laplaciano de F en términos de éstas. Así, sea $\{v_\beta\}$ una base de $\mathbb{R}_1^n$ tal que $F(p) = (F^1(p), \dots, F^n(p))$ en esta base, el Laplaciano de F se define como:

$$\Delta F := (\Delta F^1, \dots, \Delta F^n)$$

Antes de demostrar la *fórmula de Beltrami* es conveniente demostrar el siguiente lema.

Lema 3.13. Sea $\phi : M \rightarrow \mathbb{R}_1^n$ una inmersión isométrica, y sea $\{E_\alpha\}$ un marco local de M . Entonces:

$$D_{E_\alpha} \phi = E_\alpha ; \text{ donde } D \text{ es la conexión usual de } \mathbb{R}_1^n.$$

Demostración. Tomando a la subvariedad M como un subconjunto de N la inmersión es una función inclusión, es decir, ϕ es la restricción de la función identidad a M . Siendo así, $\phi = \phi^\beta \partial_\beta$ en la base usual de $\mathbb{R}_1^n$ y, $E_\alpha = E_\alpha^\gamma \partial_\gamma$. Entonces:

$$\begin{aligned} D_{E_\alpha} \phi &= E_\alpha \phi^\beta \partial_\beta = E_\alpha^\gamma (\partial_\gamma \phi^\beta) \partial_\beta \\ &= E_\alpha^\beta \partial_\beta = E_\alpha \end{aligned}$$

□

Teorema 3.14. Sea $\phi : M \rightarrow \mathbb{R}_1^n$ una inmersión isométrica. Entonces:

$$\Delta_M \phi = -mH \quad (3.12)$$

Donde H es el vector de curvatura media de M . A esta ecuación se le conoce como la *fórmula de Beltrami*.

Demostración. Sea $\{e_\alpha\}$ la base usual de $\mathbb{R}_1^n$ y $\{E_\beta\}$ un marco local sobre M en una vecindad de $p \in M$ tal que $\nabla_{E_\gamma} E_\beta(p) = 0$ para todos los elementos del marco.

En términos de la base $\{e_\alpha\}$ las funciones coordenadas de ϕ se pueden obtener como $\phi^\alpha := \langle \phi, e_\alpha \rangle$. Entonces:

$$\begin{aligned} \Delta_M \phi^\alpha(p) &= -\epsilon^\beta E_\beta E_\beta \phi^\alpha(p) + (\nabla_{E_\beta} E_\beta(p)) \phi^\alpha(p) = -\epsilon^\beta E_\beta E_\beta \phi^\alpha(p) \\ &= -\epsilon^\beta E_\beta [E_\beta \langle \phi, e_\alpha \rangle] (p) = -\epsilon^\beta E_\beta [\langle D_{E_\beta} \phi, e_\alpha \rangle + \langle \phi, D_{E_\beta} e_\alpha \rangle] (p) \end{aligned}$$

Por el Lema 3.13, el hecho de que e_α es un campo vectorial constante sobre $\mathbb{R}_1^n$ y la fórmula de Gauss, la expresión del Laplaciano de ϕ^α se simplifica a lo siguiente.

$$\begin{aligned}\Delta_M \phi^\alpha(p) &= -\epsilon^\beta E_\beta \langle E_\beta, e_\alpha \rangle(p) = -\epsilon^\beta \langle D_{E_\beta} E_\beta, e_\alpha \rangle(p) \\ &= -\langle \epsilon^\beta II(E_\beta, E_\beta), e_\alpha \rangle(p) = \langle -mH, e_\alpha \rangle(p)\end{aligned}$$

Donde m es la dimensión de M .

Como $\Delta_M \phi^\alpha(p) = \langle \Delta_M \phi, e_\alpha \rangle(p)$ se tiene finalmente que $\langle \Delta_M \phi, e_\alpha \rangle = \langle -mH, e_\alpha \rangle$ en p , demostrando la fórmula de Beltrami en una vecindad de p , y por tanto en M . □

3.3. Hipersuperficies

Sea M una subvariedad de N ; muchas veces es más cómodo hablar de la *codimensión* de la variedad M en N , la cual se define como:

$$\text{codim}_N M := \dim N - \dim M$$

Con esto puedo hacer la siguiente definición.

Definición 3.15. Se dice que una subvariedad *semi-Riemanniana* M de una variedad de Lorentz N es una *hipersuperficie* si ésta tiene codimensión uno en N .

En el caso de las hipersuperficies, al tener estas codimensión 1 el espacio $\mathfrak{X}^\perp(M)$ queda determinado localmente por un sólo campo normal unitario de tal forma que, si N es un campo normal unitario sobre M , cualquier campo $X \in \mathfrak{X}^\perp(M)$ se puede escribir como $X = fN$ donde $f \in C^\infty(M)$.

En este caso se puede hacer la siguiente definición.

Definición 3.16. El *signo* ε de una hipersuperficie M se define como:

- $\varepsilon = 1$ si $\langle u, u \rangle > 0$ para todo $u \in T_p^\perp M$ y para todo $p \in M$.
- $\varepsilon = -1$ si $\langle u, u \rangle < 0$ para todo $u \in T_p^\perp M$ y para todo $p \in M$.

De la definición se puede apreciar que si $\varepsilon = 1$, M tendrá el mismo índice que N y por lo tanto será tipo tiempo; si el signo de M es $\varepsilon = -1$ se tiene entonces que esta sera tipo espacio.

El siguiente teorema es una versión, en este contexto, de un resultado muy conocido de topología diferencial el cual dice que la imagen inversa de un valor regular para una función entre variedades es una subvariedad en el dominio.

Teorema 3.17. Sea f una función en $C^\infty(N)$. Entonces, $M = f^{-1}(c)$ es una hipersuperficie si $\langle \text{grad} f, \text{grad} f \rangle_p \neq 0$ para todo $p \in M$. Además, $N = \frac{\text{grad} f}{\|\text{grad} f\|}$ es un campo unitario normal sobre M .

Demostración. Como $\text{grad} f$ es métricamente equivalente a $d_p f$, la condición $\langle \text{grad} f, \text{grad} f \rangle \neq 0$ quiere decir, entre otras cosas, que $\text{grad} f$ no es el vector cero y por lo tanto que la diferencial $d_p f$ es sobre para todo $p \in M$, cumpliéndose así las hipótesis del teorema de la imagen inversa de un valor regular. Por lo tanto M es una subvariedad de $\text{codim}_N M = 1$.

Antes de demostrar que M es semi-Riemanniana necesito ver que $\text{grad} f$ es en realidad un vector normal a M . Sea $v \in T_p M$, entonces:

$$\langle \text{grad} f, v \rangle = v(f) = v(f|_M) = 0$$

Esto último porque f es constante por definición en M .

Entonces $\text{grad} f$ se encuentra en el complemento ortogonal de $T_p M$ para cada p , lo cual implica que $T_p M$ es no degenerado para todo p en M . $\square$

Este teorema nos permite construir muchas hipersuperficies de una manera sencilla, aunque no todas pueden ser construidas a partir de este teorema, como la banda de Möbius.

Hipercuádricas

En el caso Riemanniano el Teorema 3.17 sirve muy bien para determinar hipersuperficies de $\mathbb{R}^n$ tales como la esfera, o distintos hiperboloides a través de *formas cuadráticas*. En este caso también se pueden construir hipersuperficies de $\mathbb{R}_1^n$ usando la misma maquinaria.

Claramente toda forma cuadrática sobre un espacio vectorial es C^∞ . Sea $q \in C^\infty(\mathbb{R}_1^{n+1})$ una forma cuadrática definida por $q(x) := \langle x, x \rangle$, que en la base usual de $\mathbb{R}_1^{n+1}$ se escribe como:

$$q(x) = \eta_{\alpha\beta} x^\alpha x^\beta = -(x^0)^2 + (x^1)^2 + \dots + (x^{n+1})^2$$

Sea P el vector de posición en $\mathbb{R}_1^{n+1}$; entonces, la función $q(P)$ cumple con lo siguiente:

$$\langle \text{grad} q, v \rangle = v(q) = v\langle P, P \rangle = 2\langle D_V P, P \rangle = 2\langle V, P \rangle ; \text{ para todo } V \in \mathbb{R}_1^{n+1}$$

Por lo tanto $\text{grad} q = 2P$ y $\langle \text{grad} q, \text{grad} q \rangle = 4q$. Por el Teorema 3.17 se tiene que siempre y cuando $\langle P, P \rangle \neq 0$, la imagen inversa de P a través de q será una hipersuperficie de $\mathbb{R}_1^{n+1}$.

En la imagen inversa del cero a través de q se encuentran el cono de luz $\mathcal{C}$ y el origen. Sin embargo $\mathcal{C}$ si es una subvariedad de $\mathbb{R}^{n+1}$ con codimensión igual a uno, ya que P no es cero en el cono de luz y por lo tanto su diferencial no se anula, aunque esta no será semi-Riemanniana.

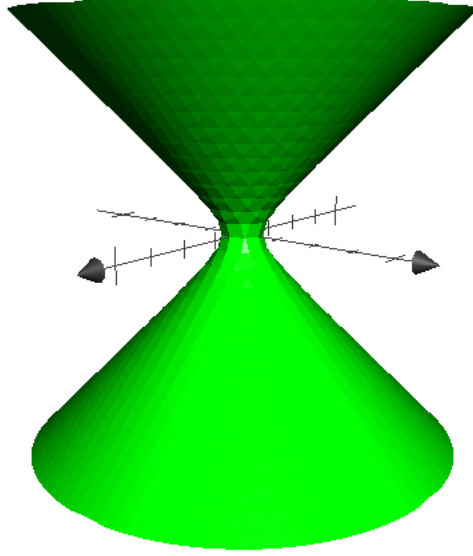


Figura 3.1: Espacio de De Sitter

Definición 3.18. Sea $n \geq 1$. Entonces se define lo siguiente:

- La *pseudo-esfera* de dimensión n e índice uno ,con radio $r > 0$ en $\mathbb{R}_1^{n+1}$ como (figura 3.1):

$$S_1^n(r) := q^{-1}(r^2) = \{p \in \mathbb{R}_1^{n+1} | \langle p, p \rangle = r^2\}$$

- El *espacio hiperbólico* de dimensión n e índice cero, con radio $r > 0$ en $\mathbb{R}_1^{n+1}$ como (figura 3.2):

$$H^n(r) := q^{-1}(-r^2) = \{p \in \mathbb{R}_1^{n+1} | \langle p, p \rangle = -r^2\}$$

Estos espacios, junto con el cono de luz y el origen, llenan todo $\mathbb{R}_1^{n+1}$. El espacio de estos se puede simplificar al estudio de aquellos que tengan radio uno, ya que sólo se diferencian por un factor escalar. Por ejemplo, todas las pseudo-esferas son homotéticas al espacio $S_1^n = S_1^n(1)$, el cual es conocido como el espacio de *De Sitter*.

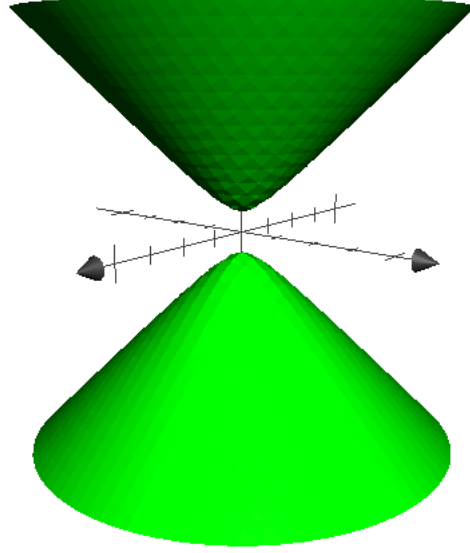


Figura 3.2: Espacio Hiperbólico

Teorema 3.19. El espacio S_1^n es difeomorfo a $\mathbb{R}_1^1 \times S^{n-1}$.

Demostración. Sean $x \in \mathbb{R}_1^1$ y $p \in S^{n-1} \subset \mathbb{R}^n$, y sea ϕ la siguiente función:

$$\phi(x, p) := (x, (1 + |x|^2)^{1/2} p) \in \mathbb{R}_1^1 \times \mathbb{R}^n \approx \mathbb{R}_1^{n+1}.$$

Donde $|x|$ es el valor absoluto de x . La razón de definir esta función como se acaba de hacer es que:

$$\langle \phi(x, p), \phi(x, p) \rangle = -|x|^2 + (1 + |x|^2) \langle p, p \rangle = 1.$$

Entonces la función ϕ lleva $\mathbb{R}_1^1 \times S^{n-1}$ en S_1^n . Esta función claramente es uno a uno; para que esta sea una biyección tengo que asegurar que todos los elementos de S_1^n se pueden escribir como imagen de ϕ . Voy a llamar p a todos los elementos de S_1^n tales que su primera coordenada sea cero, es decir, $p \in S^{n-1}$; ahora, sea q un elemento de S_1^n en la dirección de p y sea x la primera coordenada de q (todo esto en la base usual), entonces $q = (x, (1 + |x|^2)^{1/2} p)$. Por lo tanto ϕ es una biyección.

Claramente ϕ es C^∞ ya que el factor $1 + |x|^2$ no se anula para ninguna $x \in \mathbb{R}_1^1$ y la función en su segunda entrada es básicamente la inclusión;

además ϕ tiene una inversa dada por:

$$\phi^{-1}(x, q) = (x, (1 + |x|^2)^{-1/2} q).$$

Por lo tanto ϕ es un difeomorfismo y $\mathbb{R}_1^1 \times S^{n-1} \approx S_1^n$.

□

Este resultado se puede generalizar fácilmente a una pseudo-esfera de índice arbitrario tomando $|x|$ por la norma euclidiana de $x \in \mathbb{R}^\nu$, obteniendo así que $\mathbb{R}_\nu^\nu \times S^{n-\nu} \approx S_\nu^n$.

Es una convención que $\mathbb{R}^0$ represente un solo punto y S^0 dos puntos; entonces $S_1^1 \subset \mathbb{R}_1^2$ no es conexo y es el único caso en el que S_1^n no lo será.

Una demostración similar nos lleva a que $H^n \approx S^0 \times \mathbb{R}^n$. Por lo tanto el espacio hiperbólico tiene dos componentes conexas difeomorfas a $\mathbb{R}^n$.

3.4. El espacio de De Sitter

En esta sección aplicaré la herramienta expuesta al principio del capítulo a la pseudo-esfera de radio r ; claramente esta servirá para el espacio de De Sitter también.

El operador de forma en $S_1^n(r)$

Sea $\alpha(t)$ una curva en $S_1^n(r) \subset \mathbb{R}_1^{n+1}$. Entonces $\langle \alpha(t), \alpha(t) \rangle = r^2$ y usando la compatibilidad de la métrica con la conexión, derivando en la dirección de $\dot{\alpha}$, se obtiene que:

$$\langle \dot{\alpha}(t), \alpha(t) \rangle = 0 ; \text{ para todo } t \in \mathbb{R}.$$

Como α es cualquier curva en $S_1^n(r)$, se deduce de esto que el vector de posición de un punto $p \in S_1^n(r)$ es ortogonal a todo el espacio tangente $T_p S_1^n(r)$. Por lo tanto, el campo vectorial P dado por todos los vectores de posición para la pseudo-esfera es un campo normal para esta y, ya que la codimensión es uno, todo campo normal será un múltiplo de éste.

Entonces, puedo definir un campo normal unitario ξ a $S_1^n(r)$ como sigue:

$$\xi(p) := \frac{P}{r} ; \text{ para todo } p \in S_1^n.$$

Antes de continuar me gustaría hacer la siguiente aclaración. Sea $\hat{\xi}$ cualquier campo normal unitario sobre una subvariedad M de N , y sea $\nabla^\perp$ la conexión normal; como esta es compatible con la métrica se tiene que para todo campo

vectorial X tangente a M , $\langle \nabla_X^\perp \xi, \xi \rangle = 0$, lo cual implica que $\nabla_X^\perp \xi = 0$ para todo $X \in \mathfrak{X}(M)$.

Por lo dicho con anterioridad y usando la fórmula de Weingarten, se tiene que:

$$S_\xi(X) = -D_X \xi ; \text{ para todo } X \in \mathfrak{X}(S_1^n(r)).$$

Por lo tanto

$$S_\xi(X) = -\frac{X}{r},$$

por el Lema 3.13.

La segunda forma fundamental de $S_1^n(r)$

Como la codimensión de $S_1^n(r)$ es uno, se tiene que para todo $X, Y \in \mathfrak{X}(S_1^n(r))$ la segunda forma fundamental tiene la forma $II(X, Y) = \beta \xi$, donde β es una función en $C^\infty(S_1^n(r))$ que depende de los vectores X_p y Y_p . Entonces:

$$\beta = \langle II(X, Y), \xi \rangle = \langle S_\xi(X), Y \rangle = -\frac{1}{r} \langle X, Y \rangle.$$

Por lo tanto:

$$II(X, Y) = -\frac{1}{r} \langle X, Y \rangle \xi.$$

Entonces, para S_1^n estos operadores se escriben de la siguiente forma:

$$S_\xi(X) = -X ; II(X, Y) = -\langle X, Y \rangle \xi ; \text{ para todo } X, Y \in \mathfrak{X}(S_1^n).$$

Las cuatro componentes conexas de $O_1(n)$

En el primer capítulo de esta tesis quedó pendiente la demostración del Teorema 1.35, donde se afirma que el conjunto $O_1(n)$ tiene cuatro componentes conexas. Empezaré por demostrar que el grupo de Lorentz es en verdad una variedad.

Sean $M_{n \times n}(\mathbb{R})$ y $S_{n \times n}^L(\mathbb{R})$ los conjuntos de todas las matrices de $n \times n$ con entradas en los reales y el de las matrices L -simétricas, respectivamente. Una matriz $A \in M_{n \times n}(\mathbb{R})$ es L -simétrica si $A = A^L := \eta A^T \eta$.

Claramente el conjunto $M_{n \times n}(\mathbb{R})$ tiene una estructura de variedad heredada de $\mathbb{R}^{n^2}$, ya que estos dos espacios son isomorfos como espacios vectoriales. Dicho isomorfismo φ puede ser tomado por atlas junto con $M_{n \times n}(\mathbb{R})$, y así obtener una estructura C^∞ para este conjunto; restringiendo φ al conjunto de matrices L -simétricas, su imagen será un subespacio vectorial de $\mathbb{R}^{n^2}$ el cual a su vez será isomorfo a $\mathbb{R}^k$, donde $k = \frac{n(n+1)}{2}$, haciendo de $S_{n \times n}^L(\mathbb{R})$ una variedad.

Entonces, sea $F : M_{n \times n}(\mathbb{R}) \rightarrow S_{n \times n}^L(\mathbb{R})$ una función tal que $F(A) : AA^L$. La imagen inversa de la matriz identidad bajo esta función coincide con el conjunto $O_1(n)$, de tal forma que si la diferencial de F es sobre en $O_1(n)$, por el teorema de la imagen inversa de un valor regular se tendrá que el grupo de Lorentz es una variedad. Usando la estructura de $M_{n \times n}(\mathbb{R})$, la diferencial está dada por:

$$\begin{aligned} d_A F(B) &= \lim_{s \rightarrow 0} \frac{F(A + sB) - F(A)}{s} \\ &= \lim_{s \rightarrow 0} \frac{(A + sB)(A + sB)^L - AA^L}{s} \\ &= \lim_{s \rightarrow 0} \frac{sAB^L + sBA^L + s^2BB^L}{s} \\ &= AB^L + BA^L \end{aligned}$$

Ya que a estas dos variedades se les está pensando como espacios euclidianos, sus espacios tangentes en cada punto coinciden con ellos mismos; entonces, que la diferencial de F sea sobre en $O_1(n)$ quiere decir que para toda $C \in S_{n \times n}^L(\mathbb{R})$ existe $B \in M_{n \times n}(\mathbb{R})$ tal que $C = AB^L + BA^L$. Sea A un elemento de $O_1(n)$ y C una matriz L -simétrica, y defino $B = \frac{1}{2}CA$. Entonces:

$$\begin{aligned} d_A F(B) &= A\left(\frac{1}{2}CA\right)^L + \left(\frac{1}{2}CA\right)A^L \\ &= \frac{1}{2}(AA^L C^L + CAA^L) \\ &= \frac{1}{2}(C^L + C) = C \end{aligned}$$

Por lo tanto $d_A F$ es sobre para toda $A \in O_1(n)$ probando así que el grupo de Lorentz es una variedad de dimensión $\frac{n(n-1)}{2}$.

La demostración de las cuatro componentes conexas del grupo de Lorentz será por inducción sobre n y la componente $O_1^{++}(n+1)$ del mismo. La base de inducción se hace para $n = 2$ ya que el espacio de Minkowski no se encuentra bien definido para $n = 1$, este caso se encuentra demostrado en el ejercicio 1.36. Entonces, suponiéndolo para n queda demostrarlo para $n + 1$.

Antes de seguir es necesario mencionar, para continuar con esta demostración, que existe una correspondencia uno a uno entre el conjunto de todos los marcos ortonormales en un punto de S_1^n y $O_1(n+1)$. Sea $\{e_\gamma\}_{\gamma=0}^n$ la base usual de $\mathbb{R}_1^{n+1}$, usando la ecuación (1.7) se puede ver que las columnas de las matrices $\Lambda \in O_1(n+1)$ son vectores unitarios en $\mathbb{R}_1^n$. Si tomo Λe_n como el vector posición de un punto en S_1^n , las restantes columnas Λe_γ forman un marco ortonormal en el espacio tangente a ese punto.

Ahora, sea β una curva C^∞ en S_1^n que una a los puntos Λe_n y e_n . Usando el Teorema 1.35 se puede mover un marco arbitrario en Λe_n , por transporte paralelo, a un marco en e_n . Ya que S_1^n es conexa para $n > 1$, el argumento anterior define una curva α en $O_1(n+1)$; esta será C^∞ y, ya que el determinante es una función continua, ésta se encontrará contenida en $O_1^{++}(n+1)$ si Λ se encuentra en el mismo. Por lo tanto se tiene que α une a Λ con la matriz:

$$B = \begin{pmatrix} b & 0 \\ 0 & 1 \end{pmatrix}$$

Por lo tanto $b \in O_1^{++}(n)$, y por hipótesis de inducción existe una curva C^∞ en $O_1^{++}(n)$ que une esta matriz con la identidad. Con esto queda demostrado que todo par de elementos de $O_1^{++}(n+1)$ se puede unir por trayectorias y por tanto este conjunto como variedad, es conexo. Por el mismo argumento llevado a cabo en el ejercicio 1.36 se tiene finalmente que el grupo de Lorentz tiene cuatro componentes conexas.

Capítulo 4

Superficies en $\mathbb{R}_1^3$ tales que $\Delta\phi = A\phi + B$

De ahora en adelante trataré el caso de superficies en $\mathbb{R}_1^3$, es decir, subvariedades de dimensión dos en el espacio de Minkowski de dimensión tres, ya que en este capítulo demostraré un teorema de clasificación local de superficies. Sea ϕ una inmersión isométrica de la superficie M en $\mathbb{R}_1^3$; el teorema central de este capítulo clasifica localmente las superficies tales que su inmersión isométrica cumple con la ecuación:

$$\Delta_M\phi = A\phi + B \quad (4.1)$$

Donde A es un operador lineal sobre $\mathbb{R}_1^3$ y B es un vector en $\mathbb{R}_1^3$.

Esta clasificación se logrará primero por demostrar que una superficie que cumpla con la ecuación anterior tenga curvatura media constante; después se deducirá que estas superficies deberán ser *isoparamétricas* y con esto usar los resultados de [12] para completar la clasificación. Antes de empezar a demostrar los resultados que nos llevarán lentamente hasta el objetivo buscado definiré lo que es una superficie *isoparamétrica*.

Definición 4.1. Se dice que una superficie M es *isoparamétrica* si, en caso de ser su operador de forma diagonalizable sus valores propios son constantes. De no ser así, se pide que los coeficientes de su polinomio mínimo sean constantes.

Como se puede ver, el caso diagonalizable está pidiendo que las curvaturas principales de la superficie sean constantes; de no ser así, el equivalente es pedir que los coeficientes del polinomio mínimo sean constantes ya que de existir las curvaturas principales (los valores propios de S_ξ pueden ser complejos) estas estarían determinadas por tales coeficientes.

Como se vio en el capítulo anterior, una de las ventajas de tener codimensión uno en una superficie M es que todo campo vectorial normal a ésta

se puede escribir localmente como un múltiplo de un campo normal unitario ξ . Entonces puedo expresar el vector de curvatura media de toda superficie como $H = h\xi$, donde $h \in C^\infty(M)$ es una función a la cual me referiré como curvatura media, a secas.

4.1. El Laplaciano del vector H

Para abrir camino hacia el tema central de este capítulo es necesario obtener una forma explícita del Laplaciano del vector de curvatura media H , lo cual haré en esta sección. El siguiente teorema será nuestro primer paso.

Lema 4.2. El Laplaciano del vector de curvatura media H se expresa como:

$$\Delta_M H = 2S_\xi(\text{grad}h) + h\text{tr}(\nabla S_\xi) + [\Delta_M h + \varepsilon h|S_\xi|^2]\xi \quad (4.2)$$

Demostración. Sea $\{E_\alpha\}$ un marco local ortonormal sobre un abierto U de una superficie $M \subset \mathbb{R}_1^3$ tal que en un punto $p \in U$ se tiene que $\nabla_{E_\beta} E_\alpha(p) = 0$. Entonces, siendo $\{e_\beta\}$ la base usual de $\mathbb{R}_1^3$ se tiene lo siguiente:

$$\begin{aligned} \Delta_M H^\beta(p) &= \tilde{\epsilon}^\beta \langle \Delta_M H, e_\beta \rangle_p = -\epsilon^\gamma E_\gamma E_\gamma H^\beta(p) + \nabla_{E_\gamma} E_\gamma H^\beta(p) \\ &= -\tilde{\epsilon}^\beta \epsilon^\gamma E_\gamma (E_\gamma \langle H, e_\beta \rangle_p) \\ &= -\tilde{\epsilon}^\beta \epsilon^\gamma E_\gamma (\langle D_{E_\gamma} H, e_\beta \rangle_p - \langle H, D_{E_\gamma} e_\beta \rangle_p) \\ &= -\tilde{\epsilon}^\beta \epsilon^\gamma E_\gamma \langle D_{E_\gamma} H, e_\beta \rangle_p \\ &= -\tilde{\epsilon}^\beta \epsilon^\gamma (\langle D_{E_\gamma} D_{E_\gamma} H, e_\beta \rangle_p - \langle D_{E_\gamma} H, D_{E_\gamma} e_\beta \rangle_p) \\ &= \tilde{\epsilon}^\beta \langle -\epsilon^\gamma D_{E_\gamma} D_{E_\gamma} H, e_\beta \rangle_p, \end{aligned}$$

donde no se hace contracción en el índice β , pero sí en γ . Como el elemento de la base es arbitrario y la métrica es no degenerada se tiene finalmente que

$$\Delta_M H(p) = -\epsilon^\gamma D_{E_\gamma} D_{E_\gamma} H(p).$$

Pero la expresión de esta operación puede ser aún más explícita si se toma en cuenta la fórmula de Weingarten y que $H = h\xi$.

$$\begin{aligned} D_{E_\gamma} H &= E_\gamma h\xi + hD_{E_\gamma} \xi \\ &= E_\gamma \xi - hS_\xi(E_\gamma) \end{aligned}$$

Volviendo a derivar en la dirección de E_γ y usando además la fórmula de Gauss se obtiene lo siguiente.

$$\begin{aligned} D_{E_\gamma} D_{E_\gamma} H &= D_{E_\gamma} (E_\gamma h\xi - hS_\xi(E_\gamma)) \\ &= D_{E_\gamma} (E_\gamma h\xi) - D_{E_\gamma} (hS_\xi(E_\gamma)) \\ &= E_\gamma E_\gamma h\xi - E_\gamma hS_\xi(E_\gamma) - \\ &\quad \nabla_{E_\gamma} (hS_\xi(E_\gamma)) - II(E_\gamma, hS_\xi(E_\gamma)) \end{aligned}$$

y con esto finalmente se obtiene la expresión.

$$D_{E_\gamma} D_{E_\gamma} H = E_\gamma E_\gamma h \xi - 2E_\gamma h S_\xi(E_\gamma) - h \nabla_{E_\gamma}(S_\xi(E_\gamma)) - h II(E_\gamma, S_\xi(E_\gamma)) \quad (4.3)$$

Si la pasada ecuación (4.3) la evalúo en el punto $p \in U$ tal que $\nabla_{E_\alpha} E_\beta(p) = 0$, puedo cambiar el término $\nabla_{E_\gamma}(S_\xi(E_\gamma))$ por $(\nabla_{E_\gamma} S_\xi)(E_\gamma)$. Entonces, contrayendo la ecuación (4.3) en sus dos entradas se obtiene, evaluando ésta en p , lo siguiente.

$$-\epsilon^\gamma D_{E_\gamma} D_{E_\gamma} H = \Delta_M h \xi + 2S_\xi((\epsilon^\gamma E_\gamma h) E_\gamma) + h \text{tr}(\nabla S_\xi) + h \epsilon^\gamma II(E_\gamma, S_\xi(E_\gamma))$$

El segundo término así como el cuarto de la última ecuación se pueden simplificar aún más.

Para simplificar el segundo término basta con recordar que en la base local $\{E_\alpha\}$ la diferencial de h se escribe como $dh = (E_\gamma h) \omega^\alpha$, donde $\{\omega^\alpha\}$ es la base dual; con esto, el término en el argumento del operador de forma es el gradiente en la base local $\{E_\alpha\}$.

Para el cuarto término se tiene que $\langle II(E_\gamma, S_\xi(E_\gamma)), \xi \rangle = \varepsilon g$, donde ε es el signo de la superficie y g es una función en $C^\infty(M)$. Entonces:

$$\langle II(E_\gamma, S_\xi(E_\gamma)), \xi \rangle = \langle S_\xi(E_\gamma), S_\xi(E_\gamma) \rangle.$$

Lo cual hace que $g = \langle S_\xi(E_\gamma), S_\xi(E_\gamma) \rangle$. Si defino:

$$|S_\xi|^2 := \epsilon^\gamma \langle S_\xi(E_\gamma), S_\xi(E_\gamma) \rangle, \quad (4.4)$$

en la base $\{E_\alpha\}$, la ecuación (4.3) se reescribe como:

$$\Delta_M H = 2S_\xi(\text{grad} h) + h \text{tr}(\nabla S_\xi) + [\Delta_M h + \varepsilon h |S_\xi|^2] \xi$$

Ya que la traza de un tensor es un invariante del mismo ésta no depende de la base usada para calcularla, particularmente hablando de la derivada del operador de forma el cual es un tensor; esto hace que la ecuación (4.2) se cumpla en todo abierto de M .

□

La traza de $(\nabla_X S_\xi)(Y)$

Como se expresa en el Lema 1.53, existen cuatro formas distintas de representar a un operador auto-adjunto en el espacio de Minkowski. En este caso, ya que el espacio tangente a toda superficie (tipo tiempo) es isomorfo

al espacio de Minkowski, se tienen cuatro formas distintas de representar al operador de forma. Aún así, ya que en este caso la dimensión de la superficie es dos, los cuatro casos se pueden reducir a tres ya que existen dos casos de representación diagonal, uno en una base ortonormal y otro en una base pseudo-ortonormal.

Lema 4.3. Un operador auto-adjunto T en un espacio vectorial de dimensión dos con métrica de Lorentz se puede representar de una de las siguientes formas:

$$\begin{aligned} \text{I. } T &\sim \begin{pmatrix} a_1 & 0 \\ 0 & a_2 \end{pmatrix}, \quad \text{II. } T \sim \begin{pmatrix} a_0 & 0 \\ 1 & a_0 \end{pmatrix}, \\ \text{III. } T &\sim \begin{pmatrix} a_0 & b_0 \\ b_0 & a_0 \end{pmatrix} \end{aligned}$$

donde $b_0 \neq 0$. Las representaciones I y III se encuentran en una base ortonormal, mientras que la representación II se encuentra en una base pseudo-ortonormal o nula.

Demostración. Del Lema 1.53 las representaciones II y III quedan demostradas; sin embargo, aún queda por demostrar que las dos formas diagonales que se obtienen del Lema 1.53 en el caso de dimensión dos son equivalentes.

Si existe un cambio de base tal que la representación diagonal en la base pseudo-ortonormal $\{X_\alpha\}$ se mantiene diagonal en la nueva base ortonormal $\{Y_\beta\}$, se tendrá en realidad una reducción a tres casos.

El cambio de coordenadas dado por:

$$X_1 = -\frac{1}{\sqrt{2}}Y_1 - \frac{1}{\sqrt{2}}Y_2; \quad X_2 = -\frac{1}{\sqrt{2}}Y_1 + \frac{1}{\sqrt{2}}Y_2.$$

Donde Y_1 es tipo tiempo, es justo el cambio de coordenadas tal que en la base $\{Y_\beta\}$ la representación del operador de forma también es diagonal. Por lo tanto se puede estudiar el caso diagonal como un sólo caso.

□

Cada una de estas representaciones tiene un polinomio mínimo asociado, y así, la clasificación se puede hacer en términos del polinomio mínimo para todo operador auto-adjunto sobre la superficie, como se puede ver en [12].

El Lema 1.53 así como el pasado hablan de representaciones de un operador auto-adjunto en un espacio vectorial, pero lo que realmente se necesita en este caso es que ambos Lemas sean válidos en un abierto de la superficie (una variedad en general). Los coeficientes del polinomio mínimo dependen de los coeficientes de dicho operador auto-adjunto T , los cuales son funciones

reales continuas sobre la superficie; justo la continuidad de estas funciones hace que cada una de estas representaciones sea válida en un abierto de la superficie.

Teorema 4.4. La traza del operador $(\nabla_X S_\xi)(Y)$ es igual a $2\varepsilon \text{grad} h$, donde $\varepsilon = \langle \xi, \xi \rangle$.

Demostración. Usando el Lema 4.3 se puede calcular explícitamente la traza de $(\nabla_X S_\xi)(Y)$ en tres casos.

$$\blacksquare \quad S_\xi = \begin{pmatrix} \mu_1 & 0 \\ 0 & \mu_2 \end{pmatrix}$$

En este caso tenemos los operadores con polinomio mínimo $(\lambda - \mu_1)(\lambda - \mu_2)$, y $(\lambda - \mu_0)$. Usando las 1-formas de la conexión, en un marco local $\{E_\beta\}$, se obtiene lo siguiente.

$$\begin{aligned} (\nabla_{E_1} S_\xi)(E_1) &= \nabla_{E_1} S_\xi(E_1) - S_\xi(\nabla_{E_1} E_1) \\ &= \nabla_{E_1}(\mu_1 E_1) - S_\xi(\omega_1^\gamma(E_1)E_\gamma) \\ &= E_1 \mu_1 E_1 + \mu_1 \omega_1^2(E_1)E_2 - \mu_2 \omega_1^2(E_1)E_2 \end{aligned}$$

En una base como ésta las 1-formas $\omega_\beta^\beta = 0$. Entonces, la traza del operador $(\nabla_X S_\xi)(Y)$ se puede expresar de la siguiente manera.

$$\begin{aligned} \text{tr} \nabla S_\xi &= \epsilon^1 (\nabla_{E_1} S_\xi)(E_1) + \epsilon^2 (\nabla_{E_2} S_\xi)(E_2) \\ &= \epsilon^1 [E_1 \mu_1 E_1 + \mu_1 \omega_1^2(E_1)E_2 - \mu_2 \omega_1^2(E_1)E_2] \\ &\quad + \epsilon^2 [E_2 \mu_2 E_2 + \mu_2 \omega_2^1(E_2)E_1 - \mu_1 \omega_2^1(E_2)E_1] \\ &= [\epsilon^1 E_1 \mu_1 + \epsilon^2 (\mu_2 - \mu_1) \omega_2^1(E_2)] E_1 \\ &\quad + [\epsilon^2 E_2 \mu_2 + \epsilon^1 (\mu_1 - \mu_2) \omega_1^2(E_1)] E_2 \end{aligned}$$

Además.

$$\begin{aligned} (\nabla_{E_1} S_\xi)(E_2) &= \nabla_{E_1} S_\xi(E_2) - S_\xi(\nabla_{E_1} E_2) \\ &= \nabla_{E_1}(\mu_2 E_2) - S_\xi(\omega_2^1(E_1)E_1) \\ &= E_1 \mu_2 E_2 + \mu_2 \omega_2^1(E_1)E_1 - \mu_1 \omega_2^1(E_1)E_1 \\ (\nabla_{E_2} S_\xi)(E_1) &= E_2 \mu_1 E_1 + \mu_1 \omega_1^2(E_2)E_2 - \mu_2 \omega_1^2(E_2)E_2 \end{aligned}$$

Entonces, haciendo uso de la ecuación de Codazzi (3.10) se obtiene que:

$$\begin{aligned} E_1 \mu_2 &= (\mu_1 - \mu_2) \omega_1^2(E_2) ; \text{ en } E_2 \\ E_2 \mu_1 &= (\mu_2 - \mu_1) \omega_2^1(E_1) ; \text{ en } E_1. \end{aligned}$$

Usando las simetrías que obtienen las 1-formas de la conexión en una base ortonormal, expresadas en la ecuación 2.10, las ecuaciones anteriores se reescriben como sigue.

$$\begin{aligned}\epsilon^1 E_1 \mu_2 &= \epsilon^2 (\mu_2 - \mu_1) \omega_2^1(E_2) ; \text{ en } E_2 \\ \epsilon^2 E_2 \mu_1 &= \epsilon^1 (\mu_1 - \mu_2) \omega_1^2(E_1) ; \text{ en } E_1\end{aligned}$$

Por lo tanto:

$$\text{tr} \nabla S_\xi = \epsilon^1 E_1 (\mu_1 + \mu_2) E_1 + \epsilon^2 E_2 (\mu_1 + \mu_2) E_2$$

Sean g_h^β los coeficientes de $\text{grad}h$ estando este representado en términos del marco $\{E_\beta\}$, y sea X cualquier campo vectorial tangente a M con coeficientes X^β en términos del mismo marco local. Entonces:

$$\langle \text{grad}h, X \rangle = \epsilon^1 g_h^1 X^1 + \epsilon^2 g_h^2 X^2,$$

y por definición de gradiente se tiene que:

$$\epsilon^1 g_h^1 X^1 + \epsilon^2 g_h^2 X^2 = Xh = X^1 E_1 h + X^2 E_2 h.$$

Haciendo que:

$$X^1 (\epsilon^1 g_h^1 - E_1 h) + X^2 (\epsilon^2 g_h^2 - E_2 h) = 0 ; \text{ para todo } X \in \mathfrak{X}(M).$$

Así, en términos de un marco ortonormal el gradiente de h se representa como:

$$\text{grad}h = (\epsilon^1 E_1 h) E_1 + (\epsilon^2 E_2 h) E_2.$$

Lo que quiero demostrar es que la traza del operador $(\nabla_X S_\xi)(Y)$ se puede expresar como el gradiente de la curvatura media; para esto es necesario encontrar una forma de expresar la función h en términos de las funciones μ_α . Recordando que $H = h\xi$ se tiene lo siguiente.

$$\begin{aligned}2\varepsilon h &= \langle 2H, \xi \rangle = \langle \epsilon^\beta II(E_\beta, E_\beta), \xi \rangle \\ &= \epsilon^\beta \langle S_\xi(E_\beta), E_\beta \rangle = \epsilon^\beta \mu_\beta \langle E_\beta, E_\beta \rangle \\ &= \mu_1 + \mu_2\end{aligned}$$

Por lo tanto:

$$\begin{aligned}\text{tr} \nabla S_\xi &= \text{grad}(2\varepsilon h), \\ \text{tr} \nabla S_\xi &= 2\varepsilon \text{grad}(h).\end{aligned}$$

$$\blacksquare S_\xi = \begin{pmatrix} \mu & 0 \\ 1 & \mu \end{pmatrix}$$

Este es el caso con polinomio mínimo $(\lambda - \mu)^2$. De acuerdo al Lema 4.3, esta representación se encuentra expresada en términos de una base pseudo-ortonormal $\{X_\alpha\}$. Siendo así, la expresión para la traza del operador de forma debe ser distinta de la usada hasta ahora.

Sean $\nabla S_{\alpha\beta}^\gamma$ los símbolos del tensor $(\nabla_X S_\xi)(Y)$ en la base $\{X_\alpha\}$ y su dual; entonces, la traza que se quiere conocer del tensor es la respectiva a los símbolos covariantes, es decir, los símbolos α y β . En esta base la métrica queda representada como:

$$[g_{\alpha\beta}] = [g^{\alpha\beta}] = \begin{pmatrix} 0 & -1 \\ -1 & 0 \end{pmatrix}.$$

Entonces, en términos de los símbolos queda lo siguiente.

$$\text{tr} \nabla S_\xi = g^{\alpha\nu} \nabla S_{\nu\alpha}^\gamma$$

Por lo tanto, la expresión de la traza de la derivada del operador de forma en una base pseudo-ortonormal queda como sigue.

$$\text{tr} \nabla S_\xi = -(\nabla_{X_1} S_\xi)(X_2) - (\nabla_{X_2} S_\xi)(X_1)$$

Es claro que las simetrías de las 1-formas de la conexión usadas hasta ahora dependen por completo de la base en la que se represente el tensor métrico y, por tanto estas deben ser distintas en una base pseudo-ortonormal.

$$\begin{aligned} \langle \nabla_Z X_1, X_1 \rangle &= \langle \omega_1^\gamma(Z) X_\gamma, X_1 \rangle = -\omega_1^2(Z) \\ \langle \nabla_Z X_1, X_2 \rangle &= \langle \omega_1^\gamma(Z) X_\gamma, X_2 \rangle = -\omega_1^1(Z) \\ \langle \nabla_Z X_2, X_1 \rangle &= \langle \omega_2^\gamma(Z) X_\gamma, X_1 \rangle = -\omega_2^2(Z) \\ \langle \nabla_Z X_2, X_2 \rangle &= \langle \omega_2^\gamma(Z) X_\gamma, X_2 \rangle = -\omega_2^1(Z) \end{aligned}$$

Ya que $Z\langle X_\alpha, X_\alpha \rangle = 0$, se tiene que $\langle \nabla_Z X_\alpha, X_\alpha \rangle = 0$. Por lo tanto:

$$\omega_1^2 = 0 \text{ y } \omega_2^1 = 0.$$

Aplicando el mismo procedimiento a $\langle X_1, X_2 \rangle$ se obtiene que $\langle \nabla_Z X_1, X_2 \rangle = -\langle \nabla_Z X_2, X_1 \rangle$. Por lo tanto:

$$-\omega_1^1 = \omega_2^2$$

Con esto ya se puede empezar con el cálculo que se requiere.

$$\begin{aligned}
(\nabla_{X_1} S_\xi)(X_2) &= \nabla_{X_1} S_\xi(X_2) - S_\xi(\nabla_{X_1} X_2) \\
&= \nabla_{X_1} \mu X_2 - S_\xi(\omega_2^2(X_1)X_2) \\
&= X_1 \mu X_2 + \mu \omega_2^2(X_1)X_2 - \omega_2^2(x_1)S_\xi(X_2) \\
&= X_1 \mu X_2 \\
(\nabla_{X_2} S_\xi)(X_1) &= \nabla_{X_2} S_\xi(X_1) - S_\xi(\nabla_{X_2} X_1) \\
&= \nabla_{X_2} (\mu X_1 + X_2) - S_\xi(\omega_1^1(X_2)X_1) \\
&= X_2 \mu X_1 + \mu \omega_1^1(X_2)X_1 \\
&\quad + \omega_2^2(X_2)X_2 - \omega_1^1(X_2)[\mu X_1 + X_2] \\
&= X_2 \mu X_1 + [\omega_2^2(X_2) - \omega_1^1(X_2)]X_2 \\
&= X_2 \mu X_1 - 2\omega_1^1(X_2)X_2
\end{aligned}$$

Usando la ecuación de Codazzi se obtiene lo siguiente.

$$X_2 \mu = 0 ; \quad -X_1 \mu = 2\omega_1^1(X_2)X_2$$

Entonces:

$$\begin{aligned}
\text{tr} \nabla S_\xi &= -X_1 \mu X_2 + 2\omega_1^1(X_2)X_2, \\
\text{tr} \nabla S_\xi &= -2X_1 \mu X_2.
\end{aligned}$$

Usando el mismo argumento que se usó para calcular la traza del operador $(\nabla_X S_\xi)(Y)$ en una base pseudo-ortonormal, se tiene para la segunda forma fundamental lo siguiente:

$$\text{tr} II = -II(X_1, X_2) - II(X_2, X_1) = 2II(X_1, X_2).$$

En términos de una base pseudo-ortonormal $\{X_\alpha\}$. Para que exista una base de este tipo en un abierto de la superficie, ésta debe ser tipo tiempo, por lo que $\varepsilon = \langle \xi, \xi \rangle = 1$. Entonces.

$$\begin{aligned}
h &= \langle H, \xi \rangle = \left\langle \frac{1}{2} \text{tr} II, \xi \right\rangle = -\langle II(X_1, X_2), \xi \rangle \\
&= -\langle S_\xi(X_1), X_2 \rangle = -\langle \mu X_1 + X_2, X_2 \rangle \\
&= -\langle \mu X_1, X_2 \rangle = \mu
\end{aligned}$$

Expresando en esta base a cualquier campo tangente a M se tiene que:

$$\langle \text{grad} h, Z \rangle = -Z^1 g_h^2 - Z^2 g_h^1$$

y por definición de gradiente:

$$-Z^1 g_h^2 - Z^2 g_h^1 = Z^1(X_1 h) + Z^2(X_2 h).$$

Entonces en la base $\{X_\alpha\}$ el gradiente de la curvatura media se expresa como:

$$\text{grad}h = (-X_2 h)X_1 + (-X_1 h)X_2.$$

Por lo tanto en este caso, ya que $h = \mu$, se tiene finalmente que:

$$\text{tr} \nabla S_\xi = 2 \text{grad}h.$$

$$\blacksquare \quad S_\xi = \begin{pmatrix} \mu & \nu \\ -\nu & \mu \end{pmatrix}$$

Este es el caso con polinomio mínimo $(\lambda - \mu)^2 + \nu^2$. En este caso el operador de forma se encuentra representado en un marco local $\{E_\beta\}$, así que la expresión de la traza es ya conocida. Además, la superficie debe ser tipo tiempo ya que de no ser así el operador de forma siempre es diagonalizable; con esto, las simetrías de las 1-formas de la conexión son $\omega_1^2 = \omega_2^1$.

$$\begin{aligned} (\nabla_{E_1} S_\xi)(E_1) &= \nabla_{E_1} S_\xi(E_1) - S_\xi(\nabla_{E_1} E_1) \\ &= \nabla_{E_1}(\mu E_1 - \nu E_2) - S(\omega_1^2(E_1)E_2) \\ &= E_1 \mu E_1 - E_1 \nu E_2 + \mu \nabla_{E_1} E_1 \\ &\quad - \nu \nabla_{E_1} E_2 - \omega_1^2(E_1)[\nu E_1 + \mu E_2] \\ &= E_1 \mu E_1 - E_1 \nu E_2 - 2\nu \omega_2^1(E_1)E_1 \\ (\nabla_{E_1} S_\xi)(E_1) &= \nabla_{E_2}(\nu E_1 + \mu E_2) - S_\xi(\omega_2^1(E_2)E_1) \\ &= E_2 \nu E_1 + E_2 \mu E_2 + \nu \nabla_{E_2} E_1 \\ &\quad + \mu \nabla_{E_2} E_2 - \omega_2^1(E_2)[\mu E_1 - \nu E_2] \\ &= E_2 \nu E_1 + E_2 \mu E_2 + 2\nu \omega_2^1(E_2)E_2 \end{aligned}$$

Entonces, la traza del operador $(\nabla_X S_\xi)(Y)$ se expresa como sigue.

$$\begin{aligned} \text{tr} \nabla S_\xi &= [\epsilon^1 E_1 \mu + \epsilon^2 E_2 \nu - \epsilon^1 2\nu \omega_2^1(E_1)]E_1 \\ &\quad + [\epsilon^2 E_2 \mu - \epsilon^1 E_1 \nu + \epsilon^2 2\nu \omega_2^1(E_2)]E_2 \end{aligned}$$

Como se ha hecho en los anteriores casos, se usará la ecuación de Codazzi.

Antes calcularé ambos términos de esta ecuación.

$$\begin{aligned}
(\nabla_{E_1} S_\xi)(E_2) &= \nabla_{E_1}(\nu E_1 + \mu E_2) - S_\xi(\omega_2^1(E_1)E_1) \\
&= E_1\nu E_1 + E_1\mu E_2 + \nu\nabla_{E_1}E_1 \\
&\quad + \beta\nabla_{E_1}E_2 - \omega_2^1(E_1)[\mu E_1 - \nu E_2] \\
&= E_1\nu E_1 + [E_1\mu + 2\nu\omega_2^1(E_1)]E_2 \\
(\nabla_{E_2} S_\xi)(E_1) &= \nabla_{E_2}(\mu E_1 - \nu E_2) - S_\xi(\omega_1^2(E_2)E_2) \\
&= E_2\mu E_1 - E_2\nu E_2 + \mu\nabla_{E_2}E_1 \\
&\quad - \nu\nabla_{E_2}E_2 - \omega_1^2(E_2)[\nu E_1 + \mu E_2] \\
&= [E_2\mu - 2\nu\omega_1^2(E_2)]E_1 - E_2\nu E_2
\end{aligned}$$

Entonces:

$$\begin{aligned}
E_1\nu &= E_2\mu - 2\nu\omega_2^1(E_2), \\
E_2\nu &= -E_1\mu - 2\nu\omega_2^1(E_1).
\end{aligned}$$

Ya que la superficie M es tipo tiempo se puede suponer que $\epsilon^1 = -\epsilon^2$. Así al sustituir las ecuaciones anteriores, resultantes de la ecuación de Codazzi, en la expresión de la traza del operador $(\nabla_X S_\xi)(Y)$ se obtiene:

$$\text{tr}\nabla S_\xi = 2(\epsilon^1 E_1\mu E_1 + \epsilon^2 E_2\mu E_2).$$

Solo falta justificar que $h = \langle H, \xi \rangle = \mu$. Esto se hace de forma similar a como se hizo en el caso de la representación diagonal del operador de forma. Por lo tanto:

$$\text{tr}\nabla S_\xi = 2\text{grad}h.$$

□

Concluidos todos los posibles casos nos damos cuenta que la traza del operador $(\nabla_X S_\xi)(Y)$ siempre se puede escribir de la misma forma, salvo un signo, el cual dependerá de la causalidad de la superficie. Por lo tanto, se tiene una expresión explícita del Laplaciano del vector de curvatura media.

$$\Delta_M H = 2S_\xi(\text{grad}h) + \varepsilon 2h\text{grad}h + [\Delta_M h + \varepsilon h|S_\xi|^2]\xi \quad (4.5)$$

Ya que el gradiente de una función cumple con una regla de Leibniz, el término $2h\text{grad}h$ se puede reescribir como $\text{grad}h^2$.

4.2. Propiedades del operador A

Para las superficies que cumplen con la ecuación (4.1), el operador lineal A de $\mathbb{R}_1^3$ que se muestra en la ecuación cumple con diversas propiedades las cuales demostraré en esta sección.

Para empezar, usando la fórmula de Beltrami (3.12) se puede reescribir la ecuación (4.1) como:

$$-2H = A\phi + B. \quad (4.6)$$

Derivando esta ecuación en la dirección de un vector X tangente a la superficie se obtiene, por un lado.

$$\begin{aligned} -2D_X H &= -2D_X h\xi = -2Xh\xi - 2hD_X \xi \\ &= -2Xh\xi + 2hS_\xi(X) \end{aligned}$$

Por otro lado, usando el Lema 3.13.

$$D_X A\phi + D_X B = D_X A\phi = A(X)$$

Por lo tanto:

$$A(X) = 2hS_\xi(X) - 2Xh\xi. \quad (4.7)$$

Con este resultado se puede plantear el siguiente lema.

Lema 4.5. El operador A restringido al espacio tangente de la superficie M es auto-adjunto.

Demostración. La demostración de este lema es sencilla con el uso de la ecuación (4.7). Sean $X, Y \in \mathfrak{X}(M)$ campos cualesquiera, entonces:

$$\begin{aligned} \langle A(X), Y \rangle &= \langle 2hS_\xi(X), Y \rangle = 2h\langle S_\xi(X), Y \rangle \\ &= 2h\langle X, S_\xi(Y) \rangle = \langle X, 2hS_\xi(Y) \rangle \end{aligned}$$

En las pasadas ecuaciones usé de manera implícita que ξ es un campo normal a la superficie; siguiendo con ese razonamiento se tiene finalmente.

$$\langle A(X), Y \rangle = \langle X, A(Y) \rangle \quad (4.8)$$

□

Derivando la ecuación (4.8) en la dirección de $Z \in \mathfrak{X}(M)$, y utilizando que el operador A es auto-adjunto en el espacio tangente a la superficie se obtiene lo siguiente.

$$\begin{aligned} Z\langle A(X), Y \rangle &= \langle AD_Z X, Y \rangle + \langle A(X), D_Z Y \rangle \\ &= \langle \nabla_Z X, A(Y) \rangle + \langle A(II(Z, X)), Y \rangle \\ &\quad + \langle A(X), \nabla_Z Y \rangle + \langle A(X), II(Z, Y) \rangle \end{aligned}$$

Haciendo lo mismo para $\langle X, A(Y) \rangle$ e igualando se obtiene la siguiente relación.

$$\begin{aligned} \langle A(II(Z, X)), Y \rangle - \langle A(II(Z, Y)), X \rangle = \\ \langle II(Z, X), A(Y) \rangle - \langle II(Z, Y), A(X) \rangle \end{aligned} \quad (4.9)$$

Regresando a la ecuación (4.6), aplicándole el Laplaciano a esta ecuación y usando una vez más la fórmula de Beltrami se obtiene:

$$\Delta_M H = A(H).$$

Por lo tanto, de la ecuación (4.5) se obtiene la siguiente expresión.

$$A(H) = 2S_\xi(\text{grad}h) + \varepsilon 2h\text{grad}h + [\Delta_M h + \varepsilon h|S_\xi|^2]\xi \quad (4.10)$$

Hasta aquí con los resultados previos a las demostraciones principales de esta tesis referentes al operador A , ya que aún existen un par de resultados que serán de gran utilidad en lo que sigue y que están relacionados con otros aspectos de dichas demostraciones.

Lema 4.6. Sea M una superficie en $\mathbb{R}_1^3$ con curvatura media constante tal que $|S_\xi|^2$ es constante también. Entonces M es isoparamétrica.

Demostración. Como ya se vio con anterioridad, el operador de forma en este caso tiene tres representaciones canónicas, por lo cual la demostración de este lema se realizará en tres casos.

$$\blacksquare \quad S_\xi = \begin{pmatrix} \mu_1 & 0 \\ 0 & \mu_2 \end{pmatrix}$$

En este caso el polinomio mínimo del operador de forma es:

$$(\lambda - \mu_1)(\lambda - \mu_2) = \lambda^2 - \lambda(\mu_1 + \mu_2) + \mu_1\mu_2.$$

De la expresión (4.4) se obtiene que $\mu_1^2 + \mu_2^2$ es constante; además, como la curvatura media de la superficie es constante se tiene que $\mu_1 + \mu_2$ es constante también. Con todo esto, la siguiente expresión es constante.

$$(\mu_1 + \mu_2)^2 = \mu_1^2 + \mu_2^2 + 2\mu_1\mu_2$$

Con lo cual se obtiene que $\mu_1\mu_2$ es constante y por tanto los coeficientes del polinomio mínimo también lo serán. Ya que μ_1 y μ_2 son raíces del polinomio mínimo del operador de forma, estos se encuentran determinados por los coeficientes del mismo, lo cual hace que las curvaturas principales de la superficie sean constantes.

$$\blacksquare S_\xi = \begin{pmatrix} \mu & \nu \\ -\nu & \mu \end{pmatrix}$$

En este caso el polinomio mínimo es:

$$(\lambda - \mu)^2 + \nu^2.$$

Como esta representación se encuentra en una base ortonormal, calcular $|S_\xi|^2$ es relativamente sencillo. Sea $\{E_1, E_2\}$ un marco local, entonces:

$$S_\xi(E_1) = \mu E_1 - \nu E_2 ; S_\xi(E_2) = \nu E_1 + \mu E_2.$$

Así, que $|S_\xi|^2$ sea constante implica que $\mu^2 - \nu^2$ es constante y, que la curvatura media sea constante implica que μ es constante. Estas dos condiciones implican por su parte que ν^2 es constante, lo cual demuestra que el polinomio mínimo de S_ξ tiene coeficientes constantes.

$$\blacksquare S_\xi = \begin{pmatrix} \mu & 0 \\ 1 & \mu \end{pmatrix}$$

El polinomio mínimo en este caso es:

$$(\lambda - \mu)^2.$$

Entonces, basta con que μ sea constante para que la superficie sea isoparamétrica. Si la curvatura media de la superficie es constante entonces la traza de esta lo será y, como se puede ver de la representación que se tiene en este caso, esto implica que μ es constante

□

Los siguientes resultados son referentes a las 1-formas de la conexión definidas en una superficie, y sus respectivas ecuaciones de estructura. Para esto es necesaria la siguiente definición.

Definición 4.7. Sea M una superficie de $\mathbb{R}_1^3$. Se dice que $\{E_1, E_2, E_3\}$ es un *marco local adaptado* a la superficie M si $\{E_1, E_2\}$ es un marco local ortonormal de M , y E_3 es un campo normal a la misma.

Esta noción de marco adaptado a una superficie es muy útil para encontrar la forma que tendrán las ecuaciones de estructura, dadas por el Teorema 2.32, en este caso. Tomando la Definición 2.29 y restringiendo esta a vectores tangentes a una superficie de $\mathbb{R}_1^3$ se obtienen 1-formas sobre dicha superficie. Siendo así, la 1-forma dual $\omega^3 = 0$ ya que ningún campo normal se proyecta (con la métrica dada) sobre el tangente de la superficie.

Es claro que dado un marco adaptado a una superficie, los campos que forman dicho marco se extienden a campos C^∞ en un abierto de $\mathbb{R}_1^3$, lo cual le da sentido a la siguiente definición.

Definición 4.8. Sea M una superficie en $\mathbb{R}_1^3$, $\{E_1, E_2, E_3\}$ un marco adaptado a M y X un campo tangente a dicha superficie. Entonces, las 1-formas de la conexión de M quedan definidas como:

$$D_X E_\beta = \omega_\beta^\alpha(X) E_\alpha.$$

Donde D es la conexión de $\mathbb{R}_1^3$.

Definir así las 1-formas de la conexión de la superficie M cobra mucho sentido ya que se está ganando información sobre su geometría extrínseca, es decir, la forma en que M se ve desde la geometría de $\mathbb{R}_1^3$.

¿Pero, qué sucede con las 1-formas de la conexión inducida? El siguiente resultado muestra que las 1-formas $\{\omega_\beta^\alpha\}$ conservan la información de la geometría intrínseca a M .

Teorema 4.9. Sean $\{\tilde{\omega}_\beta^\alpha\}$ las 1-formas de la conexión inducida ∇ en M . Entonces $\tilde{\omega}_\beta^\alpha(X) = \omega_\beta^\alpha(X)$ si $\beta \neq 3$, donde $X \in \mathfrak{X}(M)$.

Demostración. Si $\beta \neq 3$ se cumple la fórmula de Gauss.

$$D_X E_\beta = \nabla_X E_\beta + II(X, E_\beta)$$

Como la segunda forma fundamental se encuentra en el haz normal existe $g_\beta^3(X) \in C^\infty(M)$ tal que $II(X, E_\beta) = g_\beta^3(X) E_3$. Entonces:

$$(\omega_\beta^1(X) - \tilde{\omega}_\beta^1(X))E_1 + (\omega_\beta^2(X) - \tilde{\omega}_\beta^2(X))E_2 + (\omega_\beta^3(X) - g_\beta^3(X))E_3 = 0.$$

Por lo tanto, como el marco local es linealmente independiente, se cumple el enunciado. □

¿Qué pasa ahora si $\beta = 3$? Justo se obtiene la información extrínseca de la superficie como se puede ver en el siguiente teorema.

Teorema 4.10. Sea M una superficie de $\mathbb{R}_1^3$ y $\{E_1, E_2, E_3\}$ un marco adaptado a ésta. Entonces el operador de forma de la superficie se puede escribir como sigue.

$$S_{E_3}(X) = -\omega_3^1(X)E_1 - \omega_3^2(X)E_2$$

Demostración. Como E_3 es un campo normal unitario a M se tiene que $S_{E_3}(X) = -D_X E_3$ para todo $X \in \mathfrak{X}(M)$. Además como me encuentro trabajando en un marco adaptado, el cual es ortonormal, se tiene que $\omega_3^3 = 0$. Estas dos observaciones juntas terminan la demostración. □

Tomando en cuenta que restringiendo las 1-formas duales $\{\omega^1, \omega^2, \omega^3\}$ al marco adaptado $\{E_1, E_2, E_3\}$ a campos tangentes hace que ω^3 se anule, es obvio que las ecuaciones de estructura deberán cambiar a una forma más simple. El siguiente teorema se encarga de mostrar esto.

Teorema 4.11. Sea $\{E_1, E_2, E_3\}$ un marco local ortonormal adaptado a una superficie $M \subset \mathbb{R}_1^3$ con su respectivo marco dual. Entonces, las siguientes ecuaciones se derivan de la primera ecuación de estructura (2.8):

$$\begin{aligned} d\omega^1 &= \omega^2 \wedge \omega_2^1, \\ d\omega^2 &= \omega^1 \wedge \omega_1^2, \\ \omega^1 \wedge \omega_1^3 + \omega^2 \wedge \omega_2^3 &= 0, \end{aligned}$$

mientras que de la segunda ecuación de estructura (2.9), se derivan las siguientes:

$$\begin{aligned} d\omega_2^1 &= \omega_2^3 \wedge \omega_3^1 \\ d\omega_3^1 &= \omega_3^2 \wedge \omega_2^1 \\ d\omega_3^2 &= \omega_3^1 \wedge \omega_1^2 \end{aligned}$$

Demostración. Ya que estas ecuaciones se están derivando en un marco ortonormal se tiene que $\omega_\alpha^\alpha = 0$ para todo $\alpha = 1, 2, 3$.

Para las primeras ecuaciones se toma la primera ecuación de estructura (2.8), y que la 1-forma $\omega^3 = 0$ para todo campo en $\mathfrak{X}(M)$ ya que estos no tienen componente en E_3 .

Para las segundas ecuaciones se toma la segunda ecuación de estructura (2.9) la cual, al considerar $\mathbb{R}_1^3$ como una variedad plana se reescribe como:

$$d\omega_\beta^\alpha = \omega_\beta^\mu \wedge \omega_\mu^\alpha.$$

Además, ya que el marco local sobre el cual se trabaja es ortonormal, las 1-formas de la conexión cumplen con la ecuación (2.10). Esto hace que las últimas tres ecuaciones que se muestran en este teorema sean todos los casos posibles que se derivan de la segunda ecuación de estructura. □

4.3. Ejemplos de superficies en $\mathbb{R}_1^3$

A continuación mostraré algunas superficies cuya inmersión isométrica en $\mathbb{R}_1^3$ cumple con la ecuación 4.1, así como una que no cumple con ésta.

Cuadráticas

Sea $f : \mathbb{R}_1^3 \rightarrow \mathbb{R}$ una función definida como

$$f(t, x, y) := -\delta_1 t^2 + x^2 + \delta_2 y^2,$$

tal que $\delta_1, \delta_2 \in \{0, 1\}$ y no se anulan al mismo tiempo. Entonces:

$$df = -2\delta_1 t dt + 2x dx + 2\delta_2 y dy.$$

Por lo tanto:

$$\text{grad} f = (2\delta_1 t, 2x, 2\delta_2 y).$$

Ya que $\delta_i^2 = \delta_i$ se tiene finalmente que

$$\langle \text{grad} f, \text{grad} f \rangle = 4f(t, x, y).$$

Sea $r > 0$ y $\varepsilon \in \{-1, 1\}$. Entonces, si $f(t, x, y) = \varepsilon r$, por el Teorema 3.17 se tiene que $f^{-1}(\varepsilon r^2)$ es una superficie de $\mathbb{R}_1^3$ siempre que $(\delta_1, \delta_2, \varepsilon) \neq (0, 1, -1)$, donde ε es el signo de la superficie. Además, el mismo Teorema 3.17 nos enseña como calcular el vector normal a la superficie el cual queda expresado como:

$$\xi = \frac{1}{r}(\delta_1 t, x, \delta_2 y).$$

Para calcular las curvaturas principales μ_j de la superficie se deben calcular los valores propios del operador de forma. Si tomo el operador de forma respecto del campo ξ se tiene que:

$$S_\xi(V) = -D_V \xi ; \text{ para todo } V \in \mathfrak{X}(M).$$

Ya que D es la conexión usual de $\mathbb{R}_1^3$, en las coordenadas usuales se tiene lo siguiente.

$$D_V \xi = (V\xi^\alpha)\partial_\alpha = (V^\beta\partial_\beta\xi^\alpha)\partial_\alpha$$

Por lo tanto:

$$S_\xi(V) = \frac{1}{r}(V^t\delta_1, V^x, V^y\delta_2) \quad (4.11)$$

De la función f , cuando esta es igual a εr^2 para alguna r se puede despejar a y como función de (t, x) ; así, se puede definir una parametrización de la superficie como:

$$\varphi(t, x) := (t, x(t, x), y),$$

tal que:

$$\varphi_t = (1, x_t, 0) ; \varphi_x = (0, x_x, 1),$$

forman una base local del espacio tangente donde:

$$x_t = \frac{\delta_1 t}{\sqrt{\varepsilon r^2 + \delta_1 t^2 - \delta_2 y^2}} ; \quad x_y = \frac{-\delta_1 y}{\sqrt{\varepsilon r^2 + \delta_1 t^2 - \delta_2 y^2}}.$$

Con esto, la ecuación (4.11), y usando que $\delta_i^2 = \delta_i$ se obtiene finalmente que:

$$S_\xi(\varphi_t) = \frac{-\delta_1}{r} \varphi_t ; \quad S_\xi(\varphi_y) = \frac{-\delta_2}{r} \varphi_y.$$

Conociendo el valor de las curvaturas principales μ_j de la superficie se puede calcular ahora la curvatura media de la misma.

$$h = \frac{\varepsilon}{2}(\mu_1 + \mu_2) = \frac{-1}{2r}(\delta_1 + \delta_2)$$

Usando la fórmula de Beltrami (3.12) se tiene lo siguiente.

$$\begin{aligned} \Delta_M \phi &= -2H = -2h\xi = \frac{-2h}{r}(\delta_1 t, x, \delta_2 y) \\ &= \frac{\varepsilon}{r^2}(\delta_1 + \delta_2)(\delta_1 t, x, \delta_2 y) \end{aligned}$$

Ya que la función inmersión ϕ se ve como la identidad en $\mathbb{R}_1^3$ se tiene que:

$$A = \frac{\varepsilon}{r^2}(\delta_1 + \delta_2) \begin{pmatrix} \delta_1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & \delta_2 \end{pmatrix}$$

Por lo tanto, la familia de superficies descrita por la función f cumple con una ecuación de la forma $\Delta_M \phi = A\phi + B$, donde $B = 0$.

Existen cuatro combinaciones posibles para $\{\delta_1, \delta_2\}$ de las cuales una está descartada ya que se pidió de inicio que estas constantes no se anularan al mismo tiempo. En este argumento no he tomado en cuenta el signo ε de la superficie, entonces, tomando en cuenta las tres combinaciones posibles de $\{\delta_1, \delta_2\}$ se tienen seis casos, pero como se vio anteriormente, se descarta el caso $(\delta_1 = 0, \delta_2 = 1, \varepsilon = -1)$. Por lo tanto, la función describe los siguientes cinco casos.

$$\blacksquare \quad \delta_1 = 0, \delta_2 = 1, \varepsilon = 1$$

En este caso la ecuación de la superficie es:

$$x^2 + y^2 = r^2.$$

Esta es la ecuación de una circunferencia en un plano donde la coordenada x es constante. Ya que todos estos planos son tipo espacio los vectores tangentes

a estas circunferencias también lo serán. Pero esta ecuación es válida para toda $t \in \mathbb{R}$. Entonces, a cada punto en $S^1(r)$ le tendría que pegar una copia de $\mathbb{R}$ que se encuentra en el complemento ortogonal al plano que la contiene, el cual es tipo tiempo. Por lo tanto la superficie es (figura 4.1):

$$\mathbb{R}_1 \times S^1(r) ; \text{ con } A = \begin{pmatrix} 0 & 0 & 0 \\ 0 & \frac{1}{r^2} & 0 \\ 0 & 0 & \frac{1}{r^2} \end{pmatrix}$$

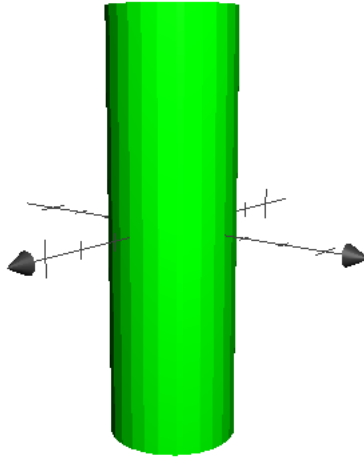


Figura 4.1: Cilindro Circular

$$\blacksquare \delta_1 = 1, \delta_2 = 0, \varepsilon = -1$$

En este caso la ecuación de la superficie es:

$$-t^2 + x^2 = -r^2.$$

Esta es la ecuación de una hipérbola en el plano. Ya que la suma es negativa, en el plano $y = 0$ esta ecuación representa el conjunto de todos los vectores tipo tiempo con producto escalar consigo mismo igual a $-r^2$ al cual se denota como $H^1(r)$; esta ecuación se cumple para todo $y \in \mathbb{R}$, entonces, a cada punto

de $H^1(r)$ se le pega una copia de $\mathbb{R}$ la cual se encuentra en el complemento ortogonal al plano que contiene a esta curva. Por lo tanto la superficie es (figura 4.2):

$$H^1(r) \times \mathbb{R} ; \text{ con } A = \begin{pmatrix} \frac{-1}{r^2} & 0 & 0 \\ 0 & \frac{-1}{r^2} & 0 \\ 0 & 0 & 0 \end{pmatrix}$$

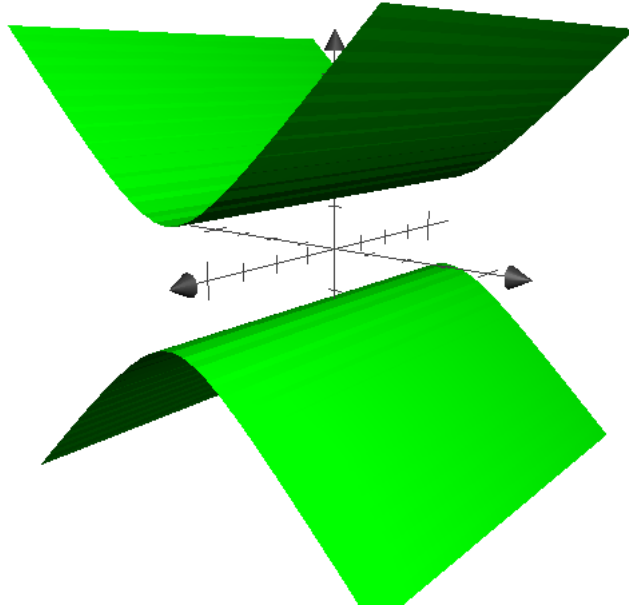


Figura 4.2: $H^1(r) \times \mathbb{R}$

■ $\delta_1 = 1, \delta_2 = 0, \varepsilon = 1$

En este caso la ecuación de la superficie es:

$$-t^2 + x^2 = r^2.$$

Una vez más se tiene la ecuación de una hipérbola en el plano. Aquí la suma resulta en un número positivo lo cual implica que esta ecuación describe en el plano $y = 0$ el conjunto de los vectores tipo espacio tales que el producto escalar consigo mismo es igual a r^2 , a este conjunto se le conoce como el espacio de De Sitter $S_1^1(r)$. Ya que la ecuación se cumple para todo $y \in \mathbb{R}$ se

tiene finalmente que la superficie es (figura 4.3):

$$S_1^1(r) \times \mathbb{R} ; \text{ con } A = \begin{pmatrix} \frac{1}{r^2} & 0 & 0 \\ 0 & \frac{1}{r^2} & 0 \\ 0 & 0 & 0 \end{pmatrix}$$

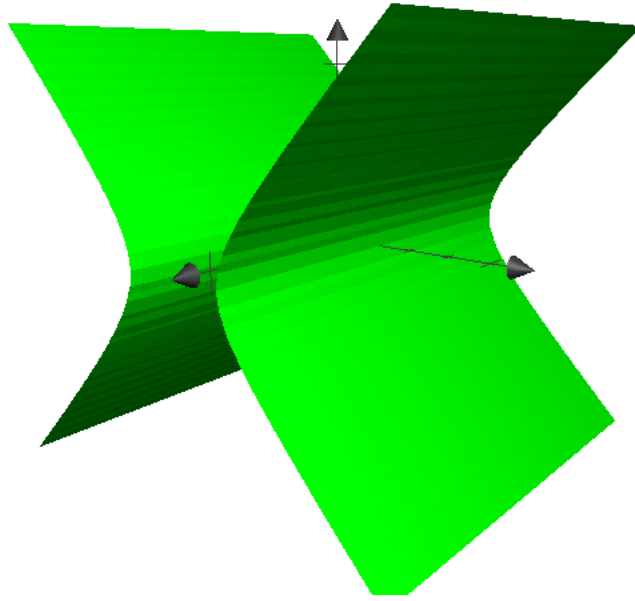


Figura 4.3: $S_1^1(r) \times \mathbb{R}$

$$\blacksquare \delta_1 = 1, \delta_2 = 1, \varepsilon = -1$$

La ecuación de la superficie es:

$$-t^2 + x^2 + y^2 = -r^2.$$

Esta ecuación representa el conjunto de vectores tipo tiempo en $\mathbb{R}_1^3$ tales que el producto escalar consigo mismos es igual a $-r^2$. Por definición esta superficie es el espacio hiperbólico de radio r :

$$H^2(r) ; \text{ con } A = \begin{pmatrix} \frac{-2}{r^2} & 0 & 0 \\ 0 & \frac{-2}{r^2} & 0 \\ 0 & 0 & \frac{-2}{r^2} \end{pmatrix}$$

$$\blacksquare \delta_1 = 1, \delta_2 = 1, \varepsilon = 1$$

La ecuación de la superficie es:

$$-t^2 + x^2 + y^2 = r^2.$$

Esta ecuación representa el conjunto de vectores tipo espacio en $\mathbb{R}_1^3$ tales que el producto escalar consigo mismos es igual a r^2 . Por definición esta superficie es el espacio de De Sitter de radio r :

$$S_1^2(r) ; \text{ con } A = \begin{pmatrix} \frac{2}{r^2} & 0 & 0 \\ 0 & \frac{2}{r^2} & 0 \\ 0 & 0 & \frac{2}{r^2} \end{pmatrix}$$

B-scroll

Para definir esta superficie necesitaré de una curva nula γ en $\mathbb{R}_1^3$ y un marco nulo $\{X, Y, Z\}$ sobre esta tal que X es el vector velocidad de la curva γ .

$$\begin{aligned} \langle X, X \rangle = \langle Y, Y \rangle = 0 & \quad ; \langle X, Y \rangle = -1, \\ \langle X, Z \rangle = \langle Y, Z \rangle = 0 & \quad ; \langle Z, Z \rangle = 1, \end{aligned}$$

y con ecuaciones de Frenet:

$$\begin{aligned} \frac{dX}{ds} &= k(s)Z \\ \frac{dY}{ds} &= \omega_0 Z \\ \frac{dZ}{ds} &= \omega_0 X + k(s)Y \end{aligned}$$

donde ω_0 es una constante distinta de cero. Entonces, la superficie que llamaremos B-scroll ([6],[8]) queda definida por la parametrización:

$$\varphi(s, u) = \gamma(s) + uY(s).$$

Con esta parametrización se obtiene la base

$$\varphi_s = X + u\omega_0 Z ; \varphi_u = Y,$$

para el espacio tangente del B-scroll. Restringiendo el producto escalar de $\mathbb{R}_1^3$ al espacio tangente de la superficie se tiene que esta debe ser tipo tiempo y, por lo tanto se puede definir un campo normal ξ sobre la superficie. Localmente se tiene que:

$$\langle \varphi_u, \xi \rangle = 0 ; \langle \xi, \xi \rangle = 1.$$

Entonces para $\lambda_1, \lambda_2 \in \mathbb{R}$, en términos de la base definida por el marco de Frenet, el vector normal se puede expresar como:

$$\xi = \lambda_1 Y + \lambda_2 Z.$$

Además.

$$\begin{aligned} \langle \xi, \xi \rangle &= \lambda_1^2 \langle Y, Y \rangle + 2\lambda_1 \lambda_2 \langle Y, Z \rangle + \lambda_2^2 \langle Z, Z \rangle \\ &= \lambda_2^2 = 1 \end{aligned}$$

Tomando $\lambda_2 = -1$ y usando ahora que φ_s y ξ son ortogonales se obtiene lo siguiente.

$$\begin{aligned} \langle \varphi_s, \xi \rangle &= \lambda_1 \langle X, Y \rangle - \langle X, Z \rangle \\ &\quad + u\omega_0 \lambda_1 \langle Z, Y \rangle - u\omega_0 \langle Z, Z \rangle \\ &= -\lambda_1 - u\omega_0 = 0 \end{aligned}$$

Por lo tanto, el campo normal a la superficie se puede representar como:

$$\xi(s, u) = -u\omega_0 Y(s) - Z(s).$$

Para conocer la curvatura media del B-scroll calcularé el operador de forma en la base $\{\varphi_s, \varphi_u\}$. Respecto del campo normal ξ se tiene que $S_\xi(X) = -D_X \xi$, entonces:

$$\begin{aligned} D_{\varphi_u} \xi &= D_Y \xi \\ &= -D_Y u\omega_0 Y - D_Y Z \\ &= -(Y u\omega_0) Y - u\omega_0 D_Y Y - D_Y Z \\ D_{\varphi_s} \xi &= D_X \xi + u\omega_0 D_Z \xi \\ &= -D_X u\omega_0 Y - D_X Z + u\omega_0 [-D_Z u\omega_0 Y - D_Z Z] \\ &= -(X u\omega_0) Y - u\omega_0 D_X Y - D_X Z \\ &\quad + u\omega_0 [-(Z u\omega_0) Y - u\omega_0 D_Z Y - D_Z Z] \end{aligned}$$

Usando que $X = \frac{d}{ds}$ y que D es la conexión de $\mathbb{R}_1^3$, $D_{\varphi_s} \xi$ se puede reescribir como sigue.

$$\begin{aligned} D_{\varphi_s} \xi &= -u\omega_0^2 Z - \omega_0 X - k(s) Y \\ &\quad + u\omega_0 [-(Z u\omega_0) Y - u\omega_0 D_Z Y - D_Z Z] \end{aligned}$$

Ya que el marco $\{X, Y, Z\}$ se encuentra definido sobre la curva γ , la variación de estos vectores sobre la curva se encuentra determinada por las ecuaciones

de Frenet, pero cómo varían estos respecto de otro vector en el marco no se especifica. Para esto habría que extenderlos a un abierto de $\mathbb{R}_1^3$.

Una manera de hacer esto es transportando paralelamente los vectores en las direcciones de los demás (que no sea la dirección X), es decir, extenderlos como campos constantes en las direcciones restantes. Definiendo las siguientes curvas

$$\alpha_0(u) := \varphi(s_0, u) ; \beta_0(v) := \gamma(s_0) + vZ.$$

Se tiene que $Y = \frac{d}{du}$ y $Z = \frac{d}{dv}$. Como las componentes de los vectores Y y Z son funciones de s las ecuaciones anteriores quedan como:

$$\begin{aligned} D_{\varphi_u} \xi &= -\omega_0 Y = -\omega_0 \varphi_u \\ D_{\varphi_s} \xi &= -u\omega_0^2 Z - \omega_0 X - k(s)Y \\ &= -\omega_0 \varphi_s - k(s) \varphi_u \end{aligned}$$

Con esto el operador de forma de la superficie se representa en la base $\{\varphi_s, \varphi_u\}$ como sigue.

$$S_\xi = \begin{pmatrix} \omega_0 & 0 \\ k(s) & \omega_0 \end{pmatrix}$$

Por lo tanto la curvatura media del B-scroll es igual a ω_0 y es constante. Usando ahora la fórmula de Beltrami se tiene finalmente que:

$$\Delta_M \phi = 2u\omega_0^2 Y(s) + 2\omega_0 Z(s).$$

Ya que ϕ es una inmersión isométrica del B-scroll, esta se puede ver como la inclusión de la imagen de la parametrización del mismo. Por lo tanto, si el B-scroll cumple con la ecuación (4.1) se tiene que:

$$2u\omega_0^2 Y(s) + 2\omega_0 Z(s) = A(\gamma) + uA(Y(s)) + B. \quad (4.12)$$

Anteriormente se mostró la representación del operador de forma en la base dada por la parametrización de la superficie; es fácil ver que el polinomio mínimo del operador de forma es $(\lambda - \omega_0)^2$, con lo cual se conocería la representación canónica que le corresponde.

En $u = 0$ la parametrización induce la base $\{X, Y\}$ en el espacio tangente de la superficie, la cual es una base pseudo-ortonormal. Como el tangente es isomorfo a $\mathbb{R}_1^2$ existe una familia de bases con la misma propiedad; esta es $\{\lambda X, \frac{1}{\lambda} Y\}$ donde λ es una función que no se anula y por tanto es siempre positiva o negativa. Usando el Lema 4.3 se deduce que para alguna de estas bases se obtiene la representación canónica del operador de forma y, conjugando la representación conocida con la matriz cambio de base se obtiene finalmente que:

$$\begin{pmatrix} \omega_0 & 0 \\ \lambda^2 k(s) & \omega_0 \end{pmatrix} = \begin{pmatrix} \omega_0 & 0 \\ 1 & \omega_0 \end{pmatrix}$$

Con lo cual se concluye que la función k no se anula para toda s en un intervalo.

Por otro lado, en $u = 0$ el campo X es tangente a la superficie; usando la ecuación (4.7) y el hecho de que la curvatura media es constante se tiene la siguiente expresión.

$$A(X) = 2\omega_0^2 X + 2\omega_0 k Y \quad (4.13)$$

Para $u \neq 0$, usando la misma ecuación (4.7) ahora para φ_s se obtiene que

$$\begin{aligned} A(\varphi_s) &= 2\omega_0 S_\xi(\varphi_s) = 2\omega_0(\omega_0 \varphi_s + k \varphi_u) \\ A(X) + u\omega_0 A(Z) &= 2\omega_0^2(X + u\omega_0 Z) + 2\omega_0 k Y \end{aligned}$$

Los valores obtenidos para $A(X)$ o $A(Z)$ dependen de estos vectores y no de u . Por lo tanto, usando el valor ya conocido para $A(X)$ y que tanto u como ω_0 son distintos de cero.

$$A(Z) = 2\omega_0^2 Z \quad (4.14)$$

Derivando respecto de s en la ecuación (4.13) se obtiene, usando que $k \neq 0$ como se demostró con anterioridad, la siguiente expresión para $A(Z)$

$$A(Z) = 4\omega_0^2 Z + \frac{2\omega_0 \dot{k}}{k} Y \quad (4.15)$$

Las diferentes expresiones para $A(Z)$ deben ser iguales; de igualar las ecuaciones (4.14) y (4.15) se obtiene:

$$\frac{2\omega_0 \dot{k}}{k} Y + 2\omega_0^2 Z = 0.$$

Como Y y Z son linealmente independientes sus coeficientes deben ser cero, lo cual implica que $\omega_0 = 0$. Esto es una contradicción con las hipótesis en la construcción del B-scroll. Por lo tanto esta superficie no cumple con la ecuación (4.1).

4.4. Clasificación de superficies en $\mathbb{R}_1^3$ que satisfacen la ecuación $\Delta\phi = A\phi + B$

En la sección anterior se dieron ejemplos de superficies que cumplen con la ecuación (4.1) y todas ellas resultaron ser superficies de curvatura media constante. Entonces, cabe preguntarse si todas las superficies de $\mathbb{R}_1^3$ que cumplan con la ecuación (4.1) tienen curvatura media constante, dicho de otra forma, quiero saber si existe una superficie de $\mathbb{R}_1^3$ tal que su curvatura media no es constante.

A continuación enunciaré y demostraré los dos Teoremas principales de esta tesis, de los cuales el primero resuelve la pregunta hecha en el párrafo anterior y el segundo hace la clasificación de superficies que tanto se ansía.

Teorema 4.12. Sea $\phi : M \rightarrow \mathbb{R}_1^3$ una inmersión isométrica que cumple con la ecuación $\Delta_M \phi = A\phi + B$. Entonces M tiene curvatura media constante.

Demostración. La demostración de este teorema la realizaré por contradicción. Empezaré por definir el subconjunto de U de M como:

$$U := \{p \in M \mid \text{grad} h^2 \neq 0\}.$$

Donde h es la curvatura media de la superficie. Claramente el subconjunto U es un abierto de M ya que si $\text{grad} h^2 \neq 0$ en $p \in M$, será distinto de cero en un abierto alrededor de p . El objetivo de la demostración será mostrar que tal abierto U de M es vacío, así que estaré trabajando dentro de tal abierto.

Ya que $\text{grad} h^2 = 2h \text{grad} h$ se tiene que en U la curvatura media no se anula, además de no ser constante. Por lo tanto se puede escribir un campo normal a la superficie en U como $\xi = \frac{1}{h}H$ y:

$$\begin{aligned} II(X, Y) &= \varepsilon \langle II(X, Y), \xi \rangle \xi, \\ &= \frac{\varepsilon}{h} \langle S_\xi(X), Y \rangle H, \end{aligned} \quad (4.16)$$

para cualesquiera campos $X, Y \in \mathfrak{X}(M)$.

Aplicando el operador A a la ecuación (4.16) y usando (4.10) se obtiene la siguiente expresión.

$$\begin{aligned} A(II(X, Y)) &= \frac{\varepsilon}{h} \langle S_\xi(X), Y \rangle [2S_\xi(\text{grad} h) \\ &\quad + 2\varepsilon h \text{grad} h + \{\Delta_M h + \varepsilon h |S_\xi|^2\} \xi] \end{aligned}$$

y siendo $Z \in \mathfrak{X}(M)$.

$$\langle A(II(X, Y)), Z \rangle = \frac{2}{h} \langle S_\xi(X), Y \rangle \langle \varepsilon S_\xi(\text{grad} h) + h \text{grad} h, Z \rangle$$

Desglosando los términos de la ecuación anterior:

$$\begin{aligned} \langle \varepsilon S_\xi(\text{grad} h), Z \rangle &= \varepsilon \langle \text{grad} h, S_\xi(Z) \rangle = \varepsilon dh(S_\xi(Z)), \\ &= \varepsilon S_\xi(Z)h, \\ \langle h \text{grad} h, Z \rangle &= h \langle \text{grad} h, Z \rangle = hZh. \end{aligned}$$

Entonces:

$$\langle A(II(X, Y)), Z \rangle = \frac{2}{h} \langle S_\xi(X), Y \rangle (\varepsilon S_\xi(Z)h + hZh),$$

y

$$\begin{aligned} & \langle A(II(X, Y)), Z \rangle - \langle A(II(Z, Y)), X \rangle = \\ & \frac{2}{h} [\langle S_\xi(X), Y \rangle (\varepsilon S_\xi(Z)h + hZh) - \langle S_\xi(Z), Y \rangle (\varepsilon S_\xi(X)h + hXh)]. \end{aligned}$$

Además, usando las ecuaciones (4.7) y (4.16) se puede reescribir la siguiente expresión.

$$\begin{aligned} & \langle II(X, Y), A(Z) \rangle - \langle II(Z, Y), A(X) \rangle = \\ & \langle \frac{\varepsilon}{h} \langle S_\xi(X), Y \rangle H, -2Zh\xi \rangle - \langle \frac{\varepsilon}{h} \langle S_\xi(Z), Y \rangle H, -2Xh\xi \rangle = \\ & -2Zh \langle S_\xi(X), Y \rangle + 2Xh \langle S_\xi(Z), Y \rangle \end{aligned}$$

Por lo tanto, usando la ecuación (4.9) y eliminando términos se obtiene:

$$\langle (2hZh + \varepsilon S_\xi(Z)h)S_\xi(X), Y \rangle = \langle (2hXh + \varepsilon S_\xi(X)h)S_\xi(Z), Y \rangle.$$

Ya que esta última ecuación se derivó para todo $X, Y, Z \in \mathfrak{X}(M)$, si defino el operador $T(X)$ para todo $X \in \mathfrak{X}(M)$ como:

$$T(X) := 2hX + \varepsilon S_\xi(X), \quad (4.17)$$

se obtiene finalmente la ecuación

$$(T(Z)h)S_\xi(X) = (T(X)h)S_\xi(Z) \quad (4.18)$$

Con esto se tiene que si M es una superficie de $\mathbb{R}_1^3$ que cumple con la ecuación (4.1), cualesquiera dos campos tangentes a M deben cumplir con la ecuación (4.18). Así, la demostración de este teorema se divide en dos casos.

Caso 1 $T(\text{grad}h) \neq 0$ en U .

Ya que $T(\text{grad}h) \neq 0$, puedo afirmar que existe un campo en U en la dirección de $\text{grad}h$ tal que $T(X)$ no se anula. Además $\text{grad}h \neq 0$, entonces h es una función en U que no se anula y que se encuentra variando, es decir, no es constante lo cual me permite afirmar que $T(X)h \neq 0$. Siendo X tal campo en U y Y otro campo distinto a X , la ecuación (4.18) se puede expresar como sigue.

$$S_\xi(Y) = \frac{T(Y)h}{T(X)h} S_\xi(X)$$

Esta última ecuación hace ver que la imagen de Y es un múltiplo de la imagen de X , para toda Y . Esto quiere decir que el rango del operador de forma es cero o uno.

Si el rango de S_ξ fuera cero, S_ξ sería el operador cero y $h = 0$, lo cual implica una contradicción con el hecho de estar trabajando en el abierto U . Entonces el rango del operador de forma es uno.

Tomando en cuenta las posibles representaciones de S_ξ acorde al Lema 4.10 sólo la representación I es posible; ya que el rango de S_ξ es uno su determinante se anula, lo que hace que la representación II sea el caso de curvatura media nula, mientras que en la representación III S_ξ se convierte en el operador nulo.

Siendo así la situación del operador de forma en U puedo tomar un marco adaptado $\{E_\alpha\}$ tal que:

$$S_\xi(E_1) = \lambda E_1 ; S_\xi(E_2) = 0 ; E_3 = \xi. \quad (4.19)$$

Para determinar el valor de λ basta con recordar que $2h = \varepsilon \text{tr}(S_\xi)$. Así:

$$\lambda = 2\varepsilon h.$$

Además, usando la ecuación (4.18) en los campos E_1 y E_2 se obtiene que $T(E_2)h = 0$, implicando que:

$$E_2 h = 0.$$

Sea $\{\omega^\alpha\}$ la base dual al marco adaptado y $\{\omega_\beta^\alpha\}$ las 1-formas de la conexión en términos del mismo. Si se toman los elementos de este par de conjuntos como 1-formas sobre la superficie M , por el Lema 4.10 se obtienen las siguientes ecuaciones.

$$\omega_3^1 = -2\varepsilon h \omega^1 \quad (4.20)$$

$$\omega_3^2 = 0 \quad (4.21)$$

Ahora, sea $X \in \mathfrak{X}(M)$. Recordando que $dh(X) = Xh$, si expreso X en términos de $\{E_\alpha\}$ se obtiene que:

$$dh = (E_1 h) \omega^1. \quad (4.22)$$

A partir de estas tres ecuaciones empezaré a trazar el camino hacia el resultado deseado.

Empezaré por calcular la derivada exterior de la ecuación (4.20), en la cual usaré las ecuaciones de estructura dadas en el Teorema 4.11. Así, por un lado:

$$\begin{aligned} d\omega_3^1 &= -2\varepsilon dh \wedge \omega^1 - 2\varepsilon h d\omega^1, \\ &= -2\varepsilon h d\omega^1, \end{aligned}$$

y por el otro, usando la ecuación (4.21):

$$d\omega_3^1 = \omega_3^2 \wedge \omega_2^1 = 0.$$

Entonces $d\omega^1 = 0$. Esto último implica que localmente existe una función C^∞ tal que $\omega^1 = dx$, es decir, localmente existe una función coordenada sobre M tal que su diferencial es igual a ω^1 .

Sea $\{x, y\}$ un sistema coordenado local en M tal que $dx = \omega^1$. Usando la ecuación (4.22) y la expresión de dh en este sistema coordenado:

$$(E_1 h)dx = \frac{\partial h}{\partial x}dx + \frac{\partial h}{\partial y}dy.$$

Al asumir que $\{x, y\}$ es un sistema coordenado, dx y dy serán linealmente independientes. Esto implica que:

$$\frac{\partial h}{\partial x} = E_1 h ; \quad \frac{\partial h}{\partial y} = 0.$$

Por lo tanto la curvatura media de M en U puede ser vista como una función de una sola variable en las coordenadas dadas por el marco adaptado con el cual se trabaja. Con esto la curvatura media tiene una sola derivada, a saber:

$$h' := E_1 h = \frac{dh}{dx}.$$

Derivando la ecuación (4.21) se tiene que:

$$\omega_3^1 \wedge \omega_1^2 = -2\varepsilon h \omega^1 \wedge \omega_1^2 = 0.$$

Entonces existe $f \in C^\infty(M)$ tal que $\omega_1^2 = f\omega^1$; además, ya que $d\omega^1 = 0$, usando las ecuaciones de estructura y la ecuación (2.10):

$$0 = \omega^2 \wedge \omega_2^1 = -\epsilon^2 \epsilon^1 \omega^2 \wedge \omega_1^2 = -\epsilon^2 \epsilon^1 f \omega^2 \wedge \omega^1.$$

Por lo tanto $f = 0$ y $\omega_1^2 = 0$.

Recordando que $E_2 h = 0$ se puede calcular el Laplaciano de h en la superficie.

$$\Delta_M h = -\epsilon^\gamma (E_\gamma E_\gamma h - \nabla_{E_\gamma} E_\gamma h) = -\epsilon^1 h''$$

Y así reescribir la ecuación (4.10), tomando en cuenta que $\text{grad}h = \epsilon^1 h' E_1$ y las ecuaciones (4.19), como sigue.

$$A(\xi) = 6\varepsilon \epsilon^1 h' E_1 + \left(-\epsilon^1 \frac{h''}{h} + 4\varepsilon h^2 \right) \xi \quad (4.23)$$

Además, aplicando la ecuación (4.7) a los campos E_1 y E_2 :

$$A(E_1) = 4\epsilon h^2 E_1 - 2h'\xi ; \quad A(E_2) = 0. \quad (4.24)$$

Entonces, las ecuaciones (4.23) y (4.24) nos permiten expresar el operador A en términos del marco adaptado como:

$$A = \begin{pmatrix} 4\epsilon h^2 & 0 & 6\epsilon\epsilon^1 h' \\ 0 & 0 & 0 \\ -2h' & 0 & -\epsilon^1 \frac{h''}{h} + 4\epsilon h^2 \end{pmatrix}$$

Como A es un operador lineal de $\mathbb{R}_1^3$ su traza, así como sus menores, serán constantes. De la representación matricial de A en esta base sólo un menor no se anula; calculando este menor y la traza de A se deducen las siguientes ecuaciones.

$$\epsilon^1 h'' = 8\epsilon h^3 - \lambda_1 h \quad (4.25)$$

$$-4\epsilon\epsilon^1 h h'' + 16h^4 + 12\epsilon\epsilon^1 (h')^2 = \lambda_2 \quad (4.26)$$

Donde λ_1 es la traza de A y λ_2 es el único menor de A que no se anula.

Antes de continuar con la demostración haré el siguiente cambio de variable.

$$\beta := (h')^2 ; \quad \frac{d\beta}{dh} = 2h'' \quad (4.27)$$

Se mostró un poco antes que h es una función que depende de una sola variable u en U ; entonces usando que h depende de una sola variable y la regla de la cadena se tiene que:

$$\frac{d\beta}{dx} = \frac{d\beta}{dh} \frac{dh}{dx} = 2h' h''.$$

Como $\frac{dh}{dx} = h'$, las ecuaciones (4.27) quedan justificadas. Entonces la ecuación (4.25) se puede reescribir como:

$$\frac{d\beta}{dh} = 16\epsilon\epsilon^1 h^3 - 2\epsilon^1 \lambda h,$$

e integrando se obtiene:

$$\beta = 4\epsilon\epsilon^1 h^4 - \epsilon^1 \lambda h^2 + C,$$

donde C es cualquier constante real. De las ecuaciones (4.25) y (4.26) se obtiene:

$$12\beta = \epsilon\epsilon^1 \lambda_2 + 16\epsilon\epsilon^1 h^4 - 4\lambda_1 \epsilon^1 h^2,$$

y substituyendo el valor de β en esta última ecuación:

$$32\varepsilon h^4 - 16\lambda_1 h^2 = \lambda_2 \varepsilon - 12C\varepsilon^1.$$

El lado izquierdo de la ecuación es un polinomio en h con coeficientes constantes, y es igual a una constante para todo $p \in U$. Por lo tanto h debe ser una función constante, contradiciendo la hipótesis y mostrando que U es el conjunto vacío.

Caso 2 Existe $p \in U$ tal que $T(\text{grad}h) = 0$.

Como $T(\text{grad}h)_p = 0$ para algún $p \in U$, por (4.17) se tiene que:

$$S_\xi(\text{grad}h)_p = -2\varepsilon h \text{grad}h_p,$$

donde todas las funciones se encuentran evaluadas en p . Entonces se puede reescribir (4.10) como:

$$A(H)_p = -2\varepsilon h h_p + (\Delta_M h + \varepsilon h |S_\xi|^2) \xi_p.$$

Sea $X \in \mathfrak{X}(M)$, usando (4.7) y la ecuación anterior:

$$\begin{aligned} \langle H, A(X) \rangle &= -2Xh \langle H, \xi \rangle \\ &= -2\varepsilon h Xh, \\ \langle A(H), X \rangle_p &= -2\varepsilon h \langle \text{grad}h, X \rangle_p \\ &= -2\varepsilon h Xh_p. \end{aligned}$$

Por lo tanto:

$$\langle A(H), X \rangle_p = \langle H, A(X) \rangle_p.$$

El Lema 4.5 asegura que A es auto-adjunto actuando sobre campos tangentes a M , y la última ecuación que existe $p \in U$ tal que A es auto-adjunto en H_p y cualquier vector tangente a M en p . Como cualquier vector de $\mathbb{R}_1^3$ puede expresarse como una combinación lineal de H_p , X_p y Y_p , siendo X y Y elementos de $\mathfrak{X}(M)$, el operador A es auto-adjunto en $\mathbb{R}_1^3$ y la pasada ecuación resulta válida en todo U . Entonces:

$$\begin{aligned} \langle A(H), X \rangle &= \langle -2\varepsilon h \text{grad}h, X \rangle \text{ en } U \text{ y,} \\ \langle 2S_\xi(\text{grad}h) + 2\varepsilon h \text{grad}h, X \rangle &= \langle -2\varepsilon h \text{grad}h, X \rangle \end{aligned}$$

para todo $X \in \mathfrak{X}(M)$, con lo cual:

$$S_\xi(\text{grad}h) = -2\varepsilon h \text{grad}h$$

en todo U . Por definición de curvatura media se tiene que $\text{tr}S_\xi = 2\varepsilon h$ y como los valores propios del operador de forma al sumarse siempre coinciden con su traza, se puede determinar el otro valor propio, que adquiere el valor $4\varepsilon h$. Como los valores propios de S_ξ son distintos este es diagonalizable y existe un marco adaptado (Lema 4.3) tal que:

$$S_\xi(E_1) = -2\varepsilon h E_1 ; S_\xi(E_2) = 4\varepsilon h E_2 ; E_3 = \xi,$$

es decir, E_3 es un campo normal. Entonces, si $\{\omega^\alpha\}$ es su marco dual y $\{\omega_\beta^\alpha\}$ son las 1-formas de la conexión en esta base, del Lema 4.10 se deducen las siguientes ecuaciones.

$$\omega_3^1 = 2\varepsilon h \omega^1 \quad (4.28)$$

$$\omega_3^2 = -4\varepsilon h \omega^2 \quad (4.29)$$

Además, como $\text{grad}h$ es paralelo a E_1 :

$$dh = (E_1 h) \omega^1 \quad (4.30)$$

Derivando exteriormente (4.28) y usando las ecuaciones de estructura para $d\omega_3^1$:

$$d\omega^1 = -2\omega^2 \wedge \omega_2^1.$$

Usando ahora las ecuaciones de estructura para $d\omega^1$:

$$\omega^2 \wedge \omega_2^1 = 0,$$

y así, $d\omega^1 = 0$. Esto último junto con (4.30) implica la existencia de un sistema coordenado en U tal que $\omega^1 = dx$, y al igual que en el caso anterior, h depende sólo de una coordenada tal que $h' = E_1 h$.

Derivando ahora (4.29) y usando (4.28) y las ecuaciones de estructura:

$$\begin{aligned} \omega_3^1 \wedge \omega_1^2 &= -4\varepsilon(h'\omega^1 \wedge \omega^2 + h\omega^1 \wedge \omega_1^2) \\ h\omega^1 \wedge \omega_1^2 &= -2\omega^1 \wedge (h'\omega^2 + h\omega_1^2) \end{aligned}$$

y

$$\omega^1 \wedge (3h\omega_1^2 + 2h'\omega^2) = 0$$

Como $\omega_1^2 = -\varepsilon^2 \omega^1 \omega_2^1$ y $d\omega^1 = 0$, el término entre paréntesis en la ecuación anterior es un múltiplo de ω^2 el cual debe ser igual a cero ya que $\omega^1 \wedge \omega^2 \neq 0$ en M ; entonces:

$$3h\varepsilon^2 \omega_2^1 = 2h'\omega^2. \quad (4.31)$$

Antes de seguir es importante encontrar una expresión para la diferencial de h' . De (4.30), derivando exteriormente y tomando en cuenta que $d\omega^1 = 0$ se obtiene que $dh' \wedge \omega^1 = 0$, implicando que dh' es paralelo a ω^1 . Entonces:

$$dh' = (E_1 h')\omega^1 = h''\omega^1.$$

Siguiendo con la demostración seguiré empleando la derivada exterior como técnica usual y la usaré ahora en (4.31), de la cual resulta lo siguiente después de usar las ecuaciones de estructura y la expresión obtenida para dh' .

$$3\epsilon^2\epsilon^1 h'\omega^1 \wedge \omega_2^1 + 3\epsilon^2\epsilon^1 h\omega_2^3 \wedge \omega_3^1 = 2h''\omega^1 \wedge \omega^2 + 2h'\omega^1 \wedge \omega_1^2$$

Usando ahora las simetrías en los índices de las 1-formas de la conexión, (4.28) y (4.29), la pasada ecuación se reescribe:

$$3\epsilon^2\epsilon^1 h'\omega^1 \wedge \omega_2^1 + 24\epsilon\epsilon^1 h^3\omega^2 \wedge \omega^1 = 2h''\omega^1 \wedge \omega^2 - 2\epsilon^2\epsilon^1 h'\omega^1 \wedge \omega_2^1,$$

y agrupando se tiene finalmente:

$$\omega^1 [(5\epsilon^2\epsilon^1 h')\omega_2^1 - (2h'' + 24\epsilon\epsilon^1 h^3)\omega^2] = 0$$

Como ω_2^1 es paralelo a ω^2 el término entre paréntesis debe ser igual a cero; además, (4.31) me provee de una forma explícita para representar ω_2^1 como múltiplo de ω^1 . Así, el coeficiente de ω^1 debe ser cero, lo que resulta en la siguiente ecuación:

$$3hh'' = 5(h')^2 - 36\epsilon\epsilon^1 h^4.$$

Haciendo el cambio de variable $\beta = (h')^2$ al igual que en el caso anterior, la pasada ecuación se reescribe como:

$$\frac{3h}{2} \frac{d\beta}{dh} = 5\beta - 36\epsilon\epsilon^1 h^4. \quad (4.32)$$

La ecuación diferencial (4.32) es inhomogénea y su solución es la suma de la solución homogénea y una solución particular de la inhomogénea. Entonces:

$$\beta = Ch^{\frac{10}{3}} - 36\epsilon\epsilon^1 h^4, \quad (4.33)$$

donde C es una constante.

Al igual que en el caso anterior el objetivo será dar otra expresión para β y obtener una expresión que dependa sólo de h .

De (4.30) se tiene que $E_2 h = 0$. Entonces, usando (4.31):

$$\begin{aligned} \Delta_M h &= -\epsilon^\alpha (E_\alpha E_\alpha h - \nabla_{E_\alpha} E_\alpha h) \\ &= -\epsilon^1 h'' + \epsilon^2 \omega_2^1 (E_2) h' \\ &= -\epsilon^1 h'' + \epsilon^1 \frac{2(h')^2}{3h} \end{aligned}$$

y así:

$$h\Delta_M h = -\epsilon^1 h h'' + \epsilon^1 \frac{2(h')^2}{3} \quad (4.34)$$

En este caso, al conocer la representación explícita del operador de forma, se puede reescribir la ecuación (4.10) como:

$$A(\xi) = -2\epsilon \text{grad} h + \left[\frac{\Delta_M h}{h} + 20\epsilon h^2 \right] \xi,$$

y por (4.7):

$$A(E_1) = -4\epsilon h^2 E_1 - 2h'\xi ; \quad A(E_2) = 8\epsilon h^2 E_2.$$

Con esto se puede calcular la traza de A que, al ser un operador lineal de $\mathbb{R}_1^3$, es igual a una constante λ . Entonces:

$$\lambda = -4\epsilon h^2 + 8\epsilon h^2 + \frac{\Delta_M h}{h} + 20\epsilon h^2,$$

y:

$$h\Delta_M h = \lambda h^2 - 24\epsilon h^4. \quad (4.35)$$

Entonces, igualando (4.34) y (4.35) se obtiene, en términos de β y su derivada la siguiente ecuación.

$$\frac{3h}{2} \frac{d\beta}{dh} = 72\epsilon \epsilon^1 h^4 - 3\lambda \epsilon^1 h^2 + 2\beta \quad (4.36)$$

Ahora, igualando las ecuaciones (4.32) y (4.36) se obtiene finalmente una expresión para β en términos de la curvatura media.

$$\beta = 36\epsilon \epsilon^1 h^4 - \lambda \epsilon^1 h^2 \quad (4.37)$$

Igualando ahora las diferentes expresiones para β en (4.33) y (4.37):

$$Ch^{\frac{4}{3}} - 72\epsilon \epsilon^1 h^2 = -\lambda \epsilon^1. \quad (4.38)$$

La ecuación (4.38) presenta un polinomio en h tal que es igual a una constante, y la única forma que esta ecuación sea cierta en U es si h es constante. Esto es una contradicción con la construcción del abierto U , implicando que el U es el conjunto vacío y demostrando el teorema. □

A continuación enunciaré y demostraré el teorema principal de esta tesis. Cabe advertir que la demostración de este teorema se basa, en parte, en el artículo [12] el cual clasifica localmente todas las hipersuperficies isoparamétricas de $\mathbb{R}_1^n$; así que la clasificación que se busca en este trabajo no se hará de manera explícita en esta demostración.

Teorema 4.13. Sea $\phi : M \rightarrow \mathbb{R}_1^3$ una inmersión isométrica. Entonces $\Delta_M\phi = A\phi + B$ si y sólo si una de las siguientes afirmaciones son ciertas:

1. M tiene curvatura media nula globalmente.
2. M es localmente una de las siguientes superficies:

$$\mathbb{R}_1 \times S^1(r) ; H^1(r) \times \mathbb{R} ; S_1^1(r) \times \mathbb{R} ; H^2(r) ; S_1^2(r).$$

Demostración. Sea M una superficie de $\mathbb{R}_1^3$ que cumple con la ecuación (4.1), entonces, por el Teorema 4.12 esta tiene curvatura media constante h . Así, si $h = 0$ el teorema queda demostrado y queda el caso en que $h \neq 0$.

Entonces, siendo la curvatura media distinta de cero las ecuaciones (4.7) y (4.10) se reescriben como:

$$\begin{aligned} A(X) &= 2hS_\xi(X), \\ A(\xi) &= \varepsilon|S_\xi|^2\xi, \end{aligned}$$

donde X es cualquier campo tangente y ξ un campo normal, ambos sobre M .

Como A es un operador lineal de $\mathbb{R}_1^3$ su traza es constante; entonces usando las ecuaciones anteriores ésta se puede expresar como sigue.

$$\begin{aligned} \text{tr}A &= 2h\text{tr}S_\xi + \varepsilon|S_\xi|^2 \\ &= 4\varepsilon h^2 + \varepsilon|S_\xi|^2 \end{aligned}$$

La última igualdad implica que $|S_\xi|^2$ es constante y por el Lema 4.6 se obtiene que M es una superficie isoparamétrica de M .

Entonces, usando la clasificación hecha en [12] y aplicándola a nuestro caso se obtiene que, si M es tipo espacio, M es localmente $H^2(r)$ o $H^1(r) \times \mathbb{R}$; por otro lado si M es tipo tiempo, M es localmente $S_1^2(r)$, $S_1^1(r) \times \mathbb{R}$, $\mathbb{R}_1 \times S^1(r)$ o el B-scroll.

En la Sección 4.3 se mostró que el B-scroll no cumple con la ecuación (4.1), lo cual concluye con la demostración, ya que si M es localmente alguna de las cuadráticas mencionadas en la Sección 4.3 o una superficie con curvatura media nula, ésta cumplirá con la ecuación (4.1).

□

4.5. Resultados afines

En el Teorema 4.12 se utilizó el lenguaje de 1-formas y derivada exterior para probar que las superficies que cumplen con la condición (4.1) tienen

curvatura media constante. Este lenguaje es independiente de la métrica ya que se construye sin la necesidad de una, y hace pensar que la demostración de este teorema puede ser válida en $\mathbb{R}^3$. Para esto habría que ver en la demostración del Teorema 4.12 en dónde se usa la hipótesis de que el espacio ambiente sea $\mathbb{R}_1^3$ y ver si éste se puede cambiar por $\mathbb{R}^3$ sin cambiar el resultado.

La respuesta es que sí se puede cambiar y sólo habría que hacer un par de modificaciones en la demostración. La primera modificación es que las constantes ϵ^1 , ϵ^2 y ε serían todas iguales a uno; la segunda es que en el caso Riemanniano el operador de forma tiene una única representación y es diagonal. Esto hace que la demostración, cuando el espacio ambiente se cambia por $\mathbb{R}^3$, sea menos complicada y por supuesto, válida.

La ecuación (4.7) depende sólo de la conexión, siendo esta válida en $\mathbb{R}^3$; por otro lado, (4.10) también es válida tomando en cuenta que en $\mathbb{R}^3$ el signo ε de la superficie siempre es igual a uno, y que el operador de forma siempre tiene una representación diagonal. Con esto se obtiene que $|S_\xi|^2$ es constante, implicando que la superficie debe ser isoparamétrica. Conociendo las superficies isoparamétricas de $\mathbb{R}^3$ el siguiente resultado se puede ver como un corolario del Teorema 4.13.

Corolario 4.14. Sea $\phi : M \rightarrow \mathbb{R}^3$ una inmersión isométrica. Entonces $\Delta_M \phi = A\phi + B$ si y sólo si M es localmente una superficie mínima, $S^2(r)$ o $S^1(r) \times \mathbb{R}$.

Las superficies que cumplen con la condición (4.1) en $\mathbb{R}_1^3$ también cumplen con la condición $\Delta_M H = AH$ y a su vez con la condición $\Delta_M H = \lambda H$, donde λ es una constante real. Pero ¿son éstas todas las superficies de $\mathbb{R}^3$ tales que su vector de curvatura media cumple con la condición $\Delta_M H = \lambda H$? La respuesta es no.

Teorema 4.15. Sea M una superficie de $\mathbb{R}_1^3$. Entonces su vector de curvatura media cumple con la condición $\Delta_M H = \lambda H$ si y sólo si se cumple una de los siguientes enunciados:

1. Para todo $p \in M$ la curvatura media es igual a cero.
2. M es localmente un B-scroll.
3. M es localmente una de las siguientes superficies:

$$\mathbb{R}_1 \times S^1(r) ; H^1(r) \times \mathbb{R} ; S_1^1(r) \times \mathbb{R} ; H^2(r) ; S_1^2(r).$$

Como se puede ver en el artículo [6], la demostración de este teorema es similar a la de los Teoremas 4.12 y 4.13, usando también el lenguaje de

1-formas y su derivada exterior. A diferencia del Teorema 4.13, en esta lista se suma el B-scroll, el cual se puede ver en lo demostrado en la sección 4.3 que cumple con dicha condición.

Por último, el Teorema 4.13 se puede generalizar a hipersuperficies en una variedad semi-Riemanniana siempre y cuando la variedad ambiente sea *completa*, conexa, y tenga curvatura constante. Sea N_ν^{n+1} la variedad ambiente de dimensión $n+1$ con índice ν , y M_μ^n una hipersuperficie de N_ν^{n+1} con índice μ . Entonces el Teorema 4.13 es un caso particular del siguiente teorema.

Teorema 4.16. Sea $\phi : M_\mu^n \rightarrow N_\nu^{n+1}$ una inmersión isométrica. Entonces, $\Delta_M\phi = A\phi + B$ si y sólo si se cumple una de las siguientes afirmaciones:

1. M_μ^n es una hipersuperficie mínima.
2. M_μ^n es una hipersuperficie totalmente umbílica.
3. M_μ^n es un producto semi-Riemanniano, que no es mínimo.
4. M_μ^n es una hipersuperficie cuadrática de la forma $q(X) = \langle L(X), X \rangle$, donde L es un operador auto-adjunto de $\mathbb{R}_s^{n+2}$ con polinomio mínimo $p_L(x) = x^2 + at + b$ tal que $a, b \in \mathbb{R}$, con operador de forma no diagonalizable.

La demostración de este teorema se puede consultar en [1], así como la lista de hipersuperficies que cumplen con la clasificación.

Bibliografía

- [1] Luis J. Alías, Angel Ferrández, Pascual Lucas, *Hypersurfaces in space forms satisfying the condition $\Delta x = Ax + B$* . Trans. Amer. Math. Soc., 347-5, 1995, p. 1793-1801.
- [2] Luis J. Alías, A. Ferrández, P. Lucas, *Surfaces in the 3-dimensional Lorentz-Minkowski Space Satisfying $\Delta x = Ax + B$* . Pacific J. Math., 156, 1992, p. 201-208.
- [3] I. Chavel, *Riemannian Geometry: A Modern Introduction*. Cambridge University Press, New York, 2nd Edition, 2006.
- [4] Bang-Yen Chen, *Pseudo-Riemannian Geometry, δ -Invariants and Applications*. World Scientific, New Jersey, 2011.
- [5] M. P. do Carmo, *Riemannian Geometry*. Birkhäuser, Boston, 2nd Edition, 1993.
- [6] A. Ferrández, P. Lucas, *On Surfaces in the 3-dimensional Lorentz-Minkowski Space*. Pacific J. Math., 152, 1992, p. 93-100.
- [7] S. H. Friedberg, A. J. Insel, L. E. Spence, *Linear Algebra*. Prentice Hall, New Jersey, 4th Edition, 2003.
- [8] L. K. Graves, *Codimension One Isometric Immersions Between Lorentz Spaces*. Trans. Amer. Math. Soc., 252, 1979, p. 367-392.
- [9] V. Guillemin, A. Pollack, *Differential Topology*. Prentice-Hall, Inc. Englewood Cliffs, New Jersey, 1974.
- [10] John M. Lee, *Introduction to Topological Manifolds*. Springer-Verlag, New York, 2000.
- [11] John M. Lee, *Riemannian Manifolds: An Introduction to Curvature*. Springer-Verlag, New York, 1997.

- [12] Martin A. Magid, *Lorentzian Isoparametric Hypersurfaces*. Pacific J. Math., 118, 1985, p. 165-197.
- [13] Barrett O'Neill, *Semi-Riemannian Geometry: With Applications to Relativity*. Academic Press, New York, 1983.
- [14] B. F. Schutz, *A First Course in General Relativity*. Cambridge University Press, New York, 2nd Edition, 2009.
- [15] James W. Vick, *Homology Theory: An Introduction to Algebraic Topology*. Springer, New York, 2nd Edition, 1994.