

13.
2e)



UNIVERSIDAD NACIONAL AUTONOMA DE MEXICO

FACULTAD DE CIENCIAS

"INTRODUCCION A LA TEORIA DE MODELOS SIGMA NO LINEALES"

T E S I S
QUE PARA OBTENER EL TITULO DE:
F I S I C O
P R E S E N T A:

JERONIMO ALONSO CORTEZ QUEZADA



TESIS CON
FALLA DE ORIGEN

DIRECTOR DE TESIS:
DR. HERNANDO ORTIZ BILLOS



1227
FACULTAD DE CIENCIAS
SECCION ESCOLAR



Universidad Nacional
Autónoma de México



UNAM – Dirección General de Bibliotecas
Tesis Digitales
Restricciones de uso

DERECHOS RESERVADOS ©
PROHIBIDA SU REPRODUCCIÓN TOTAL O PARCIAL

Todo el material contenido en esta tesis esta protegido por la Ley Federal del Derecho de Autor (LFDA) de los Estados Unidos Mexicanos (México).

El uso de imágenes, fragmentos de videos, y demás material que sea objeto de protección de los derechos de autor, será exclusivamente para fines educativos e informativos y deberá citar la fuente donde la obtuvo mencionando el autor o autores. Cualquier uso distinto como el lucro, reproducción, edición o modificación, será perseguido y sancionado por el respectivo titular de los Derechos de Autor.

PAGINACION VARIA

COMPLETA LA INFORMACION



UNIVERSIDAD NACIONAL
AUTÓNOMA DE
MÉXICO

M. en C. Virginia Abrín Batule
Jefe de la División de Estudios Profesionales de la
Facultad de Ciencias
P r e s e n t e

Comunicamos a usted que hemos revisado el trabajo de Tesis: "Introducción a la Teoría de Modelos σ no lineales".

realizado por Jerónimo Alonso Cortez Quezada
con número de cuenta 9251787-9 , pasante de la carrera de Física
Dicho trabajo cuenta con nuestro voto aprobatorio.

Atentamente

Director de Tesis
Propietario

[Firma]
DR. HERNÁNDO QUEVEDO CUBILLOS

Propietario

[Firma]
DR. ALFREDO RAUL LUIS MACÍAS ALVAREZ

Propietario

[Firma]
DR. TONATIUH MATOS CHASSIN

Suplente

[Firma]
DR. GIL ALEJANDRO CORICHI RODRIGUEZ

Suplente

[Firma]
DR. HUGO AURELIO MORALES TECOTL

[Firma]
Coordinador Departamental de Física

P. 4 DR. ROBERTO ALEJANDRO RUELAS MAYORGA

A DIEGO, NOEMI y CLAUDE.

Agradecimientos

Quiero agradecer de manera muy especial al Dr. Hernando Quevedo Cubillos y al Dr. Darío Núñez Zúñiga el haberme brindado, además de la posibilidad de trabajar a su lado, su apoyo incondicional, confianza y amistad. Su sentido del humor y su calidez como personas son ejemplo de que investigadores de talla internacional también pueden ser estupendas personas, gracias por haberme enseñado que no hay como hacer las cosas con gusto y divertimento.

Agradezco a mis sinodales Dr. Tonatiuh Matos, Dr. Alfredo Macías, Dr. Hugo Morales y, de manera muy particular, al Dr. Alejandro Corichi sus valiosas sugerencias y comentarios.

A quien ha dado a mi vida un maravilloso vuelco con su amor. A quien me da alegría, aliento y ternura. Con quien comparto sueños, esperanzas y una vida plena e intensa. A tí, querida **Claudia**, dedico éste trabajo.

A la pandilla, a ese estupendo grupo de amigos con quienes he pasado momentos de gran felicidad y de quienes he aprendido mucho: A la Claus, responsable indirecta del ya famoso concepto del *pre-reven*, por su deliciosa sonrisa; al Mulato, entrañable amigo, cómplice de aventuras y con quien la alegría por la vida toma tintes desenfrenados; al Chalex, gran cuate, extraordinario reetiquetador de índices mudos y bailarín sobre mesas; a la Sub, maravillosa persona que descubrí tarde en la carrera, pero afortunadamente muy a tiempo en la vida; al Negro, con quien tengo la dicha de compartir una añeja y sólida amistad; a la familia Campuzano, quienes con su presencia siempre dieron un toque distinguido a las pachangas (Benja, gracias por los cursos de LaTeX); al Panek, el único que sabe bailar, por no sólo oír, sino también escuchar; a Felipe, extraordinario amigo, por su formidable manera de superar adversidades; a Galán, por su gran amistad y el recuerdo de ese viaje en tierras Oaxaqueñas; al Barberis, al Katz, al Marco y al Galáctico, menos aficionados al tendido, pero siempre con gran alborozo. Vale pues, esta tesis también está dedicada, con gran cariño, a vosotros.

Índice

Introducción	1
1 Topología	5
§1.1 Topologías y Vecindades	5
§1.2 Conjuntos cerrados, puntos de acumulación y cerradura	7
§1.3 Interior y frontera	8
§1.4 Espacio topológico de Hausdorff	9
§1.5 Bases y sub-bases	11
§1.6 Cubiertas y conjunto compacto	12
§1.7 Equivalentes e invariantes topológicos	13
§1.8 Conexidad	14
§1.9 Espacios métricos y pseudométricos	15
2 Geometría diferencial.	17
§2.1 Variedades diferenciables.	17
§2.2 Funciones y curvas.	18
§2.3 Vectores y campos vectoriales.	19
§2.3.1 Derivada direccional y vectores.	19
§2.3.2 Campos vectoriales.	20
§2.3.3 Curvas integrales.	21
§2.3.4 Paréntesis de Lie.	21
§2.4 Haces fibrados.	22

§2.5	Uno-formas.	26
§2.6	Tensores y campos tensoriales.	28
§2.7	Producto y componentes tensoriales.	28
§2.8	Transformaciones de base.	29
§2.9	El tensor métrico sobre un espacio vectorial.	31
§2.10	El campo tensorial métrico sobre una variedad	34
§2.11	Tensores antisimétricos.	35
§2.12	Formas diferenciales.	35
§2.13	Integración sobre variedades.	38
§2.14	n-vectores y duales.	39
§2.15	Derivada exterior.	41
§2.16	Formas cerradas y exactas.	42
§2.17	Teoría de cohomologías.	44
3	Derivada, grupos y álgebras de Lie.	47
§3.1	Arrastre de Lie.	47
§3.2	Derivada de Lie.	48
§3.3	Invariancia.	51
§3.4	Campos vectoriales de Killing.	52
§3.5	Definición de grupo de Lie.	52
§3.6	Subgrupos uniparamétricos.	54
§3.7	Ejemplos de grupos de Lie.	54
§3.8	Álgebras de Lie y sus grupos.	64
§3.9	Representaciones de un grupo de Lie.	67
4	Conexiones.	71
§4.1	Sobre la conexión y la curvatura.	71
§4.2	La forma de conexión y el potencial de norma.	75
§4.3	Transporte paralelo, derivada covariante y curvatura.	76
§4.4	Conexión en el haz tangente.	80

5 Modelo σ no lineal.	83
§5.1 Métricas axisimétricas estacionarias.	83
§5.2 Densidad lagrangiana axisimétrica estacionaria.	87
§5.3 Definición de un modelo σ no lineal.	90
§5.4 Construcción de la densidad lagrangiana para un modelo σ no lineal.	91
§5.5 El modelo σ no lineal $G = \text{Sl}(2, \mathbb{R})$, $H = \text{SO}(2)$	96
Conclusiones	103
A Variedades orientables y partición de unidad.	107
§A.1 Variedades orientables.	107
§A.2 Partición de unidad.	109
B Densidad de Routh y ecuaciones "principales" de Einstein axisimétricas estacionarias.	113
§B.1 Densidad de Routh.	113
§B.2 Ecuaciones "principales" de Einstein axisimétricas estacionarias.	115

x

Introducción

El dominio que cubre la física actual es sorprendentemente amplio, abarca desde escalas cosmológicas hasta escalas a nivel subatómico; pero todavía más sorprendente es el hecho de que todos los fenómenos que puedan producirse, en este vasto dominio, sólo tengan como origen a cuatro fuerzas físicas fundamentales: *la fuerza gravitacional, la electromagnética, la débil y la fuerte*. La primera de estas fuerzas, la gravitacional, gobierna el comportamiento de objetos astrofísicos, es de gran alcance y atractiva, es la fuerza dominante, la que "manda" a escalas astronómicas y cosmológicas. La fuerza electromagnética es la predominante en el mundo cotidiano y, también, la que rige todas las reacciones químicas, su acción es entonces preponderante tanto en fenómenos que se llevan a cabo a nuestra escala como en aquellos que ocurren a distancias moleculares o atómicas. Finalmente las otras dos fuerzas, llamadas simplemente interacción fuerte e interacción débil, son las que gobiernan el mundo a nivel de núcleos atómicos y de partículas elementales; su corto alcance no les permite actuar más que a distancias infinitesimales. La interacción débil es responsable de la desintegración de ciertas partículas como los neutrones y los muones, por otra parte, la interacción fuerte se "encarga" de enlazar a los protones y neutrones en el núcleo atómico y a los quarks en el interior de los mismos.

Uno de los problemas esenciales que se presenta en la física teórica actual, es el de la unificación de estas cuatro fuerzas fundamentales. A la fecha, los esfuerzos por conseguirlo están enmarcados en lo que se conoce como *Teorías de Norma*; la estructura geométrica de estas teorías describen de manera exacta las teorías de interacción entre las fuerzas electromagnética, débil y fuerte. En el contexto de estas teorías, el "acople" entre electricidad y magnetismo que logró Maxwell está geoméricamente descrito por el grupo $U(1)$, la descripción de la unión entre electromagnetismo y las interacciones débiles está dada por el producto de los grupos $SU(2)$ y $U(1)$ (Electro-Débil) y, finalmente, la conexión geométrica con la interacción fuerte es el producto $SU(3) \times SU(2) \times U(1)$ (Teorías de Gran Unificación). Sin embargo, la gravitación no es, en el sentido estricto, una Teoría de Norma,

consecuentemente no se ha logrado todavía cerrar el "círculo" que unifique a las cuatro fuerzas fundamentales. Es en este ámbito en el que surgen los *modelos σ no lineales*, que en cierto experimentos, donde intervienen la fuerza electromagnética y las interacciones fuertes y débiles, han resultado de gran utilidad, pues su estructura geométrica es muy similar a la estructura geométrica de las Teorías de Norma; adicionalmente, la importancia de estos modelos radica también en el hecho de que hay casos donde ciertos campos gravitacionales se pueden modelar como un *modelo σ no lineal* [1].

La similitud en estructura geométrica entre los *modelos σ no lineales* y la propia de las Teorías de Norma, además de su aplicación a ciertos campos gravitacionales, dan a estos modelos una particular relevancia en los intentos de unificación y motivan el estudio de los mismos. El presente trabajo de tesis se concretará *exclusivamente* a presentar de manera autocontenida e introductoria a dichos *modelos σ no lineales* y a exhibir el caso axisimétrico estacionario en vacío, como caso particular en Teoría General de la Relatividad. Con ello se pretende dar un primer paso que a la postre permita trazar objetivos mucho más ambiciosos.

La idea de este trabajo es que para cualquier lector, con un mínimo de conocimientos en matemáticas superiores, quede claro qué es un *modelo σ no lineal*. En consecuencia, para comprender la base geométrica de los *modelos σ no lineales*, que son los haces fibrados principales, haremos de estudiar con cierto detenimiento conceptos básicos de topología, geometría diferencial, conexiones, grupos y álgebras de Lie para, finalmente, definir y aplicar el modelo.

En el primer capítulo expondremos elementos básicos de topología, como lo son el de un *espacio topológico*, una *topología*, el de *equivalentes topológicos*, *conexidad* y *espacios métricos y pseudométricos* entre otros. Ello nos permitirá definir el concepto de *variedad diferenciable* (capítulo 2) y desarrollar sobre ella un cálculo que comprende temas tan importantes como las *funciones*, las *curvas*, los *tensores* y las *formas diferenciales* pero, sobre todo, a este nivel ya será posible introducir a los *haces fibrados*. El tercer capítulo consiste fundamentalmente en la definición de un *grupo de Lie* de dimensión finita; éste, al ser una variedad, cuenta con la estructura geométrica sobre la cual pueden definirse los diversos conceptos presentados en el segundo capítulo, con lo cual vamos ganando en estructura, pues contaremos con un grupo y una vasta estructura geométrica. Este capítulo es de gran importancia puesto que el grupo de simetría global en un *modelo σ no lineal* es, de facto, un grupo de Lie. Para notar la relevancia que los grupos de Lie tienen en física se introducen diversos ejemplos, tales como el grupo $SU(2)$ y el grupo de Poincaré. En el cuarto y penúltimo capítulo se expone, de manera somera, el tema de las *conexiones*; se manejan con familiaridad a los grupos de Lie y a los haces fibrados para derivar conceptos, como el *potencial de norma* y la *2-forma*

de curvatura, que en la construcción de la densidad lagrangiana para el modelo σ no lineal aparecen. En el quinto capítulo cinco se deduce con argumentos de plausibilidad el elemento de línea axisimétrico estacionario y se presenta a la densidad lagrangiana de Einstein-Hilbert correspondiente, a partir de la cual se obtienen las ecuaciones de campo para este caso (con lo cual se prepara el terreno para la aplicación). En este mismo capítulo, contando ya con todos los conocimientos matemáticos necesarios, se define al modelo σ no lineal, se construye la densidad lagrangiana para el mismo y se exhibe una aplicación.

Finalmente, en el apartado de conclusiones se proponen futuros temas de investigación que, entre otros puntos de partida, requieren del material expuesto en este trabajo.

Capítulo 1

Topología

La descripción de un evento, en física, es inherente a un espacio donde se lleva a cabo el mismo. El espacio de descripción es un *espacio topológico* (§1.1); es decir, un conjunto dotado con una estructura que, en el ámbito de la física, no es arbitraria puesto que esta íntimamente ligada a cantidades con contenido físico; la *topología*, a grandes rasgos, consiste en el análisis formal de la estructura del espacio. La topología nos dice, por ejemplo, que propiedades del espacio no cambian al variar de manera *continua* la masa de un objeto en reposo, respondiendo de tal forma a preguntas como si el espacio era conexo (§1.8) y metrizable (§1.9), ¿seguirá siéndolo durante el proceso de variación de la masa del objeto?

Lo anterior ilustra vagamente la importancia de la topología en el contexto de la física; sin embargo, deja entrever que el estudio de la topología, si nos interesa en un momento dado la estructura del espacio donde se llevan a cabo eventos, es inevitable. De hecho, como hemos mencionado, el espacio donde describimos a un determinado sistema físico es un espacio topológico, de tal suerte que resulta lógico pensar que la física siempre va acompañada de la topología. En el presente capítulo nos avocaremos a exponer, de manera relativamente formal, algunos de los conceptos más importantes de la topología.

§1.1 Topologías y Vecindades

Comencemos por definir qué es una topología [2]:

Una *topología* es una familia τ de conjuntos que satisfacen las siguientes condiciones:

- (i).- La intersección de cualesquiera dos miembros de τ es un miembro de τ .

(ii).- La unión de miembros de cada subfamilia de τ es un miembro de τ .

Sea X el conjunto dado por $X \equiv \cup\{U : U \in \tau\}$, entonces cada miembro de τ es un subconjunto de X y X mismo es un miembro de τ , pues τ es una subfamilia de sí mismo. El conjunto X es llamado el *espacio* de la topología τ , en tanto que τ es una *topología para X*. A la pareja (X, τ) se le conoce como *espacio topológico* (en casos en que no se preste a confusión denotaremos al espacio topológico simplemente por X).

Los miembros de la topología τ son llamados *abiertos* relativos a τ ó τ -abiertos (en caso de que solamente esté involucrada una topología en determinado análisis, entonces podemos eliminar la referencia a la topología y, sencillamente, llamar *conjuntos abiertos* a los miembros de la topología). En consecuencia, el espacio X y el conjunto vacío son abiertos (como el complemento de abiertos son cerrados (§1.2) entonces X y $\emptyset$ son tanto abiertos como cerrados). Si el espacio X y el conjunto vacío son los únicos conjuntos abiertos ($\tau = \{X, \emptyset\}$), entonces diremos que la topología es *trivial* (*indiscreta*) para X ; el espacio topológico (X, τ) es, por lo tanto, un espacio topológico *trivial* (*indiscreto*). La contraparte de ésta topología es aquella en la cual τ es la familia de todos los subconjuntos de X (cada subconjunto del espacio es abierto), tal topología es denominada topología *discreta* para X .

Las topologías discreta e indiscreta para un conjunto X son, respectivamente, la topología más grande y más pequeña para X . Una topología v es más pequeña que una u si cada v -abierto es u -abierto, en tal caso es usual decir que v es una topología más *gruesa* que u y que u es una topología más *fina* que v . Decimos que dos topologías arbitrarias para X , v y u , *no son comparables* si ocurre que v no es ni mayor ni menor que u .

Un conjunto U , en un espacio topológico (X, τ) , es una *vecindad* (τ -vecindad) de un punto x si U contiene un conjunto abierto que contenga a x . La vecindad de un punto no necesariamente es un abierto; sin embargo, cada conjunto abierto es una vecindad para cada uno de sus puntos. Por lo tanto, toda vecindad de un punto contiene una vecindad abierta del punto.

Teorema 1.1 *Sea A un conjunto, entonces A es abierto $\Leftrightarrow A$ contiene una vecindad (al menos) para cada punto en A .*

Demostración: $\Rightarrow$] Puesto que A es abierto, A es una vecindad para cada punto en A . Dado que $A \subset A$, entonces A contiene una vecindad (A) para cada punto en A .

$\Leftarrow$] Sea U el abierto que consiste en la unión de todos los subconjuntos abiertos de A . Si A contiene una vecindad para cada uno de sus puntos, entonces cada x en A pertenece a

algún subconjunto abierto de A y por lo tanto $x \in U$. En consecuencia $A = U$ y, entonces, A es abierto. $\square$

§1.2 Conjuntos cerrados, puntos de acumulación y cerradura

Definimos al complemento de $A \subset X$ relativo a X como el siguiente conjunto: $A \sim X \equiv \{x \mid x \in X \text{ y } x \notin A\}$. Un subconjunto A del espacio topológico (X, τ) es *cerrado* sii su complemento relativo a X ($X \sim A$) es abierto. Por lo tanto, un conjunto es abierto sii su complemento es cerrado. La unión de un número finito de conjuntos cerrados es un conjunto cerrado y la intersección de miembros de una familia arbitraria de conjuntos cerrados es, también, un conjunto cerrado.

Puesto que A es cerrado sii $X \sim A$ es abierto, entonces A es cerrado sii $X \sim A$ contiene una vecindad (al menos) para cada punto en $X \sim A$ (A y $X \sim A$ son disjuntos). Tenemos que, entonces, A es cerrado sii para cada x , tal que toda vecindad de x interseca a A , ocurre que $x \in A$. La última afirmación es relativamente sutil: es suficiente la existencia de un punto ajeno al conjunto A , tal que todas sus vecindades intersecten a A , para que A no sea cerrado. Precisemos esta idea con la definición de un punto de acumulación: Un punto x es un *punto de acumulación* de un subconjunto A de un espacio topológico (X, τ) sii cada vecindad de x contiene puntos de A diferentes a x . De tal manera, resulta evidente que cada vecindad de x interseca a A sii $x \in A$ ó x es un punto de acumulación de A . El siguiente teorema resume nuestro resultado:

Teorema 1.2 *Sea A un subconjunto de un espacio topológico, entonces A es cerrado $\Leftrightarrow A$ contiene a el conjunto de todos sus puntos de acumulación.*

Observemos que si x no es un punto de acumulación de A , entonces existe una vecindad abierta U de x la cual cumple que $A \cap U = \emptyset$. Puesto que U es una vecindad para cada uno de esos puntos, es claro que no existe punto en U que sea un punto de acumulación de A y, por ende, la unión de A y el conjunto de sus puntos de acumulación es el complemento de un conjunto abierto. Este razonamiento es, esencialmente, la demostración al siguiente teorema:

Teorema 1.3 *La unión de un conjunto y el conjunto de sus puntos de acumulación es un conjunto cerrado.*

Definamos a $\bar{A}$ (subconjunto del espacio topológico (X, τ)) como la intersección de los miembros de la familia de todos los conjuntos cerrados que contienen a un subconjunto A del espacio topológico. El conjunto $\bar{A}$ es conocido como *cerradura* (τ -cerradura) de A y es, en efecto, un conjunto cerrado puesto que es producto de la intersección de conjuntos cerrados que siempre es un cerrado. Evidentemente $\bar{A}$ es el cerrado más pequeño que contiene a A , por lo tanto A es cerrado si $A = \bar{A}$. El teorema que a continuación enunciaremos, describe a la cerradura de un conjunto en términos de los puntos de acumulación del mismo:

Teorema 1.4 *La cerradura de cualquier conjunto es la unión del conjunto y el conjunto de sus puntos de acumulación.*

Demostración: Sea x un punto de acumulación de A . Como $A \subset \bar{A}$, entonces x es un punto de acumulación de $\bar{A}$; puesto que $\bar{A}$ es un conjunto cerrado, entonces $x \in \bar{A}$. Por lo tanto, $\bar{A}$ contiene a A y a todos los puntos de acumulación de A .

Por otro lado, la unión de A y sus puntos de acumulación es un conjunto cerrado; como $\bar{A}$ es el más pequeño de los cerrados que contienen a A , entonces $\bar{A}$ está contenido en $A \cup \{\text{puntos de acumulación de } A\}$.

Por lo tanto: $\bar{A} = A \cup \{\text{puntos de acumulación de } A\}$ $\square$

§1.3 Interior y frontera

Un punto x , contenido en un subconjunto A de un espacio topológico, es un *punto interior* de A si A es una vecindad de x . Al conjunto de todos los puntos interiores de A se le llamará *interior* de A y se le denotará por $\overset{\circ}{A}$. Las propiedades más importantes de $\overset{\circ}{A}$ se pueden resumir en un teorema:

Teorema 1.5 *Sea A un subconjunto de un espacio topológico (X, τ) . Entonces:*

- (i).- $\overset{\circ}{A}$ es abierto y es el subconjunto abierto más grande de A .
- (ii).- A es abierto si $A = \overset{\circ}{A}$.
- (iii).- El conjunto de todos los puntos de A que no son puntos de acumulación de $X \sim A$ es, precisamente, $\overset{\circ}{A}$.
- (iv).- La cerradura de $X \sim A$ es $X \sim \overset{\circ}{A}$.

Demostración: (i): Sea $x \in \tilde{A}$, entonces existe un U abierto contenido en A tal que $x \in U$. $\forall x \in U \Rightarrow x \in \tilde{A}$ y entonces $\tilde{A}$ contiene una vecindad para cada uno de sus puntos. Por lo tanto, $\tilde{A}$ es abierto. Sea V un subconjunto abierto de A y tomemos z en V . De tal manera A es una vecindad de z y en consecuencia $z \in \tilde{A}$. Entonces $\forall z \in V \Rightarrow z \in \tilde{A}$. Es decir, $\tilde{A}$ contiene a todos los subconjuntos abiertos de A y es, por lo tanto, el subconjunto abierto más grande de A .

(ii): Si A es abierto $\Rightarrow A$ es el subconjunto abierto más grande de $A \Rightarrow A = \tilde{A}$. Por otro lado, si $A = \tilde{A}$ entonces como $\tilde{A}$ es abierto $\Rightarrow A$ es abierto.

(iii): Sea $B = \{x \mid x \in A \text{ y } x \text{ no es punto de acumulación de } X \sim A\}$. Sea $x \in B \Rightarrow$ existe una vecindad U de x tal que $(X \sim A) \cap U = \emptyset$. Entonces $U \subset A \Rightarrow A$ es vecindad de $x \Rightarrow x \in \tilde{A}$. Por lo tanto $B \subset \tilde{A}$. Por otro lado, tenemos que si $x \in \tilde{A} \Rightarrow \tilde{A}$ es una vecindad de x ; puesto que $\tilde{A} \cap (X \sim A) = \emptyset \Rightarrow x \in B$. Por lo tanto $\tilde{A} \subset B$. Dado que $B \subset \tilde{A}$ y $\tilde{A} \subset B \Rightarrow B = \tilde{A}$. Es decir, $\tilde{A}$ es el conjunto de todos los puntos de A que no son puntos de acumulación de $X \sim A$.

(iv): Por (iii) tenemos que $X \sim \tilde{A} = \{x \mid x \in X \sim A \text{ ó } x \text{ es punto de acumulación de } X \sim A\}$. Es decir, $X \sim \tilde{A} = X \sim A$. $\square$

Es importante notar que la topología determina por completo las características de cada subconjunto del espacio topológico; por ejemplo: Si (X, τ) es un espacio indiscreto, el interior de cada conjunto (excepto X mismo) es vacío. Si (X, τ) es un espacio discreto, entonces cada conjunto es abierto y cerrado, por lo tanto cada conjunto es idéntico con su interior y con su cerradura.

Sea $b(A)$ el subconjunto del espacio topológico (X, τ) que consiste en todos los puntos tales que no son puntos interiores ni de A ni de $X \sim A$. El conjunto $b(A)$ es conocido como **frontera** del subconjunto A del espacio topológico (X, τ) . De la definición se sigue que, entonces, un punto $x \in b(A)$ si y cada vecindad de x intersecta tanto a A como a $X \sim A$; es decir, $b(A) = (X \sim \tilde{A}) \cap \tilde{A} = \tilde{A} \sim \tilde{A}$. Nótese que $A \cup b(A) = \{x \mid x \in A \text{ ó } x \in (\tilde{A} \sim \tilde{A})\} = \tilde{A}$ y que $A \sim b(A) = \{x \mid x \in A \text{ y } x \notin (\tilde{A} \sim \tilde{A})\} = \tilde{A}$, en consecuencia, un conjunto es cerrado si contiene a su frontera y es abierto si es disjunto a su frontera.

§1.4 Espacio topológico de Hausdorff

Para poder establecer con exactitud cuando un espacio topológico es un espacio de Hausdorff, es necesario introducir primero el concepto de conjunto dirigido y, posteriormente, el de red.

Un conjunto dirigido es un par $(D, \geq)$ tal que $\geq$ (una relación binaria) dirige a D :

Una relación binaria $\geq$ dirige a un conjunto D si D es no vacío y:

- (a).- si m, n y p son miembros de D tales que $m \geq n$ y $n \geq p$, entonces $m \geq p$.
- (b).- si $m \in D$, entonces $m \geq m$.
- (c).- si m y n son miembros de D, entonces existe $p \in D$ tal que $p \geq m$ y $p \geq n$.

Entonces diremos que n precede a m y m sigue en el orden $\geq$ a n sii $m \geq n$.

Una red es una pareja $(S, \geq)$, donde S es una función y $\geq$ dirige al dominio de S . Si S es una función cuyo dominio contiene a D y D es dirigido por $\geq$, entonces $\{S_n, n \in D, \geq\}$ es la red $(S|D, \geq)$, donde $S|D$ denota S restringida a D. Una red $(S|D, \geq)$ está dentro de un conjunto A sii $S_n \in A$ para toda n ; la red estará eventualmente en A sii existe $m \in D$ tal que, si $n \in D$ y $n \geq m$, entonces $S_n \in A$. La red está frecuentemente en A sii para cada $m \in D$ existe $n \in D$ tal que $n \geq m$ y $S_n \in A$.

En un espacio topológico (X, τ) , la red $(S, \geq)$ converge a s relativo a τ sii $(S, \geq)$ está eventualmente en cada τ -vecindad de s . Observemos que una red no necesariamente converge a un solo punto: si X es el espacio indiscreto, entonces cada red en X converge a cada uno de los puntos en X.

Algunos resultados importantes, relativos a la convergencia, son los siguientes [2]:

- (a).- Un punto s es un punto de acumulación de un subconjunto A de X sii existe una red en $A \sim \{s\}$ que converge a s .
- (b).- Un punto $s \in \bar{A}$ sii existe una red en A ($A \subset X$) que converge a s .
- (c).- Un subconjunto A de X es cerrado sii no existe red en A que converja a un punto de $X \sim A$.

Anteriormente notamos que una red en un espacio topológico puede converger a puntos distintos. Es natural pensar, entonces, en la existencia de una topología para la cual, en el espacio topológico, la convergencia sea única: si una red S converge a un punto s y también a uno t , entonces $s = t$. La idea de convergencia única esta estrechamente ligada, como se muestra en el siguiente teorema, con el espacio topológico de Hausdorff: Un espacio topológico es un espacio de Hausdorff sii para cualesquiera puntos distintos x y z del espacio existen vecindades disjuntas de x y z .

Teorema 1.6 Un espacio topológico es un espacio topológico de Hausdorff sii cada red en el espacio converge a lo más a un punto.

Demostración: $\Rightarrow$ Si X es un espacio topológico de Hausdorff y s y t son puntos diferentes de X, entonces existen vecindades disjuntas de s y t . Dado que una red no puede estar

eventualmente en cada uno de los conjuntos disjuntos, es claro que entonces no existe una red en X que converja a s y t . La red converge a s ó a t (sólo a un punto), nunca a ambos.

⇐] Supongamos que cada red en el espacio topológico X converge a lo más a un punto y que ello implica que X no es un espacio topológico de Hausdorff.

Como X no es un espacio de Hausdorff, podemos decir que entonces existen puntos s y t distintos tales que cada vecindad de s intersecta a cada una de las vecindades de t . Sean U_s y U_t las familias de vecindades de s y t , respectivamente; ambas familias son dirigidas por la relación binaria $\subset$. Ordenemos el producto cartesiano $U_s \times U_t$ acordando que $(T, U) \geq (V, W)$ sii $T \subset V$ y $U \subset W$ (el producto cartesiano es dirigido por $\geq$). Para cada $(T, U) \in U_s \times U_t$ la intersección $T \cap U$ es no vacía, entonces podemos seleccionar un punto $S_{(T,U)}$ de $T \cap U$. Si $(V', W') \geq (T, U)$, entonces $S_{(V', W')} \in V' \cap W' \subset T \cap U$ y consecuentemente la red $\{S_{(T,U)}, (T, U) \in U_s \times U_t, \geq\}$ converge a s y a t . Esto último contradice la hipótesis de que cada red converge a lo más a un punto, la contradicción proviene de suponer que X no es Hausdorff. Por lo tanto, si cada red en el espacio converge a lo más a un punto, entonces el espacio es un espacio de Hausdorff. □

§1.5 Bases y sub-bases

Diremos que una familia B de conjuntos es una *base para la topología* τ sii B es una subfamilia de τ y para cada punto x del espacio, y cada vecindad U de x , existe un miembro V de B tal que $x \in V \subset U$. De manera equivalente, la subfamilia B de una topología τ es una base para τ sii cada miembro de τ es la unión de miembros de B .

Teorema 1.7 Una familia B de conjuntos es una base para alguna topología del conjunto $X = \{x : b \in B\} \Leftrightarrow \forall U, V \in B$ y $\forall x \in U \cap V \exists W \in B$ tal que $x \in W$ y $W \subset U \cap V$.

Demostración: ⇒] Si B es una base para alguna topología, U y V son miembros de B y $x \in U \cap V \Rightarrow$ dado que $U \cap V$ es abierto $\exists W \in B$ tal que $x \in W$ y $W \subset U \cap V$.

⇐] Sea τ la familia de todas las uniones de miembros de B . La unión de miembros de τ es, entonces, la unión de miembros de B y, en consecuencia, es un miembro de τ . Sean U y V en τ , si $x \in U \cap V$ podemos elegir A y C en B tales que $x \in A \subset U$ y $x \in C \subset V$ y, entonces, $W \in B$ es tal que $x \in W \subset A \cap C \subset U \cap V$. Entonces $U \cap V$ es la unión de miembros de $B \Rightarrow U \cap V \in \tau$. Por lo tanto τ es una topología. □

Una familia S de conjuntos es una *sub-base* para una topología τ sii la familia de intersecciones finitas de miembros de S es una base para τ (i.e: sii cada miembro de τ es la unión de intersecciones finitas de miembros de S).

§1.6 Cubiertas y conjunto compacto

Sea C una familia, entonces C es una *cubierta* de un conjunto B sii B es un subconjunto de la unión $\bigcup \{A \mid A \in C\}$; si cada miembro de C es un conjunto abierto, entonces C es una *cubierta abierta* de B . Una *subcubierta* de C es una subfamilia que también es cubierta.

Teorema 1.8 (Lindelöf) *Para cada cubierta abierta de un subconjunto de un espacio topológico, cuya topología tiene una base numerable, hay una subcubierta numerable.*

Demostración: Sea C una cubierta abierta de un conjunto A . Sea B una base numerable para la topología. Puesto que cada miembro de C es la unión de miembros de B , entonces hay una subfamilia K de B (que también cubre a A) tal que cada miembro de K es un subconjunto de algún miembro de C . Para cada miembro de K podemos seleccionar un miembro de C que lo contenga y así obtener una subfamilia numerable D de C . Entonces D es también una cubierta de A puesto que K cubre a A . Por lo tanto C tiene una subcubierta numerable. $\square$

Un espacio topológico es un espacio Lindelöf sii cada cubierta abierta del espacio tiene una subcubierta numerable.

Una cubierta C es *localmente finita* si para cada punto x existe una vecindad U tal que la intersección de U con miembros de C es no vacía para un número finito de miembros de C . Una cubierta es *finita* si es cubierta y esta constituida por un número finito de elementos.

Diremos que un AX es *compacto* si es Hausdorff y si cada cubierta de A tiene una subcubierta finita [3].

Teorema 1.9 *Un subespacio¹ compacto de un espacio topológico de Hausdorff es necesariamente cerrado. Cualquier subespacio cerrado de un espacio compacto es compacto.*

¹Un subespacio topológico de (X, τ_X) es el par (A, τ_A) , donde A es un subconjunto arbitrario de X y $\tau_A = \{A \cap V, V \in \tau_X\}$.

§1.7 Equivalentes e invariantes topológicos

La equivalencia e invariancia topológica son conceptos derivados a partir de las funciones, una *función* f de X a Y ($f: X \rightarrow Y$) es un operador que asocia a cada $x \in X$ un único elemento $y = f(x) \in Y$ ($f(x)$ es la imagen de x bajo f). El dominio de f , denotado con frecuencia por $D(f)$, es X . El rango de la función ($R(f)$) es tal que $R(f) \subseteq Y$: si $R(f) \subset Y$ entonces se dice que f "va dentro" de Y , si $R(f) = Y$ entonces f "va sobre" Y (f es suryectiva o sobreyectiva). Una función es un caso particular de lo que se conoce como *mapeos*: un mapeo es simplemente una relación entre espacios topológicos, $m: X \rightarrow Y$, y que se distingue de una función por el hecho de que a un elemento de X puede asociarle varios elementos de Y ².

Si B es un subconjunto de Y , entonces la inversa de B bajo f ($f^{-1}(B)$) es $\{x: f(x) \in B\}$. La inversa de la unión (intersección) de miembros de una familia de subconjuntos de Y bajo f es la unión (intersección) de las inversas de los miembros bajo f . La imagen de la unión de una familia de subconjuntos de X es la unión de las imágenes; sin embargo, en general, la imagen de la intersección no es la intersección de las imágenes.

La función f es una *función inyectiva* o *uno a uno* (1-1) sii para cualesquiera dos puntos distintos sus imágenes correspondientes son distintas (i.e. $x \neq y \Leftrightarrow f(x) \neq f(y)$). En este caso f^{-1} es la *función inversa* para f . La función f es *biyectiva* si es inyectiva y sobreyectiva.

El mapeo f que va del espacio topológico (X, τ) al espacio topológico (Y, u) es *continuo* sii la inversa de cada conjunto abierto es abierto. De manera más formal: f es continua respecto a τ y u (ó τ - u continua) sii $f^{-1}(V) \in \tau$ para cada $V \in u$. Una función $f: X \rightarrow Y$ es continua en un punto x sii la inversa bajo f de cada vecindad de $f(x)$ es una vecindad de x . Evidentemente f es continua sii f es continua en cada punto de su dominio.

A una función $f: X \rightarrow Y$ continua, biyectiva y con inversa continua (bicontinua) la llamaremos *homeomorfismo* (transformación topológica). Si existe un homeomorfismo entre los espacios X y Y , entonces se dice que ambos espacios son *homeomórficos* y que uno es *homeomorfo* al otro. Es claro que la inversa de un homeomorfismo y la composición de homeomorfismos es un homeomorfismo. Por lo tanto, la colección de espacios topológicos puede dividirse dentro de clases de equivalencia tales que cada espacio topológico es homeomórfico a cada miembro de esta clase de equivalencia y sólo a cada uno de estos espacios. Es decir, la clase de equivalencia de un espacio topológico X es el conjunto de todos los espacios topológicos homeomórficos a X ; si Y es homeomórfico a X , entonces la clase de equivalencia

²Puesto que las funciones son casos particulares de mapeos, entonces la palabra función puede sustituirse por la de mapeo, pero no al revés.

para Y es la misma que para X . Dos espacios topológicos son *topológicamente equivalentes* si los espacios son homeomórficos; es por ello que, en términos topológicos, suele decirse que un triángulo, un círculo y un cuadrado, por ejemplo, son lo mismo.

Un *invariante topológico* es aquella propiedad que posee un espacio topológico y cada uno de los espacios homeomórficos a éste. En consecuencia, un invariante topológico es una propiedad común para los miembros que componen la clase de equivalencia de un determinado espacio topológico. Una propiedad definida en términos de los miembros del espacio y de la topología es, automáticamente, un invariante topológico.

§1.8 Conexidad

Si un espacio topológico X es la unión de dos subconjuntos ($A \cup B = X$) disjuntos, no vacíos y abiertos, entonces se dice que X debe ser *disconexo*. Puesto que B es el complemento de A , entonces B es abierto y cerrado (y viceversa). Lo anterior sugiere la validez del siguiente teorema:

Teorema 1.10 *Un espacio topológico X es conexo si y sólo si los únicos subconjuntos que son tanto abiertos como cerrados son el conjunto vacío y el espacio X .*

Un espacio topológico es *localmente conexo* si toda vecindad, para cualquier punto x , contiene una vecindad conexa. La conexidad de un espacio topológico implica la conexidad local del mismo; sin embargo, la conexidad local no implica la conexidad global.

El espacio topológico X es *arco-conexo* si dados cualesquiera dos puntos a y b en X , existe un camino continuo entre ellos: $C: [0,1] \rightarrow X$ tal que $C(0) = a$ y $C(1) = b$, con C continuo. El espacio topológico X es *localmente arco-conexo* si para cada $x \in X$, y cada vecindad V de x , existe una vecindad U de x , contenida en V , que es arco-conexa. Si un espacio topológico es arco-conexo entonces es conexo (el inverso no siempre es válido).

Se dice que dos caminos C_1 y C_2 son *homotópicos* si existe un mapeo continuo $F: [0,1] \times [0,1] \rightarrow X$ tal que $F(t,0) = C_1(t)$ y $F(t,1) = C_2(t)$. Si X es arco-conexo y localmente arco-conexo, entonces diremos que X es *simplemente conexo* si cada camino cerrado C en X es homotópico a un mapeo constante.

§1.9 Espacios métricos y pseudométricos

La *métrica* para un conjunto X es una función d , que actúa sobre el producto cartesiano $X \times X$ y tiene como imagen al conjunto (o a un subconjunto) de los números reales no negativos³, tal que para cualesquiera x, y y z de X cumple las siguientes condiciones:

- (i).- $d(x, y) = d(y, x)$ simetría
- (ii).- $d(x, y) + d(y, z) \geq d(x, z)$ desigualdad del triángulo
- (iii).- $d(x, y) = 0$ si $x = y$
- (iv).- si $d(x, y) = 0$ entonces $x = y$

Una función d que satisface sólo (i)-(iii) es denominada pseudométrica. En lo que sigue usaremos sólo el término métrica, pero hacemos notar que es igualmente válido al sustituir por pseudométrica:

Un *espacio métrico* es la pareja (X, d) , donde d es la métrica para X . Si x y z pertenecen a X , el número $d(x, z)$ es la *distancia* de x a z . Si r es un número real positivo, el conjunto $\{y | d(x, y) < r\}$ es la *esfera abierta* de d -radio r alrededor de x (o r -esfera abierta alrededor de x). La intersección de dos esferas abiertas no necesariamente es una esfera; no obstante, podemos construir una esfera (contenida en la intersección) para cada uno de los puntos de la intersección: sean $E_1 = \{x | d(y, x) < r\}$ y $E_2 = \{x | d(z, x) < s\}$, tomemos $u \in E_1 \cap E_2$; entonces $E_3 = \{x | d(u, x) < \min(r - d(y, u), s - d(z, u))\}$ es una esfera alrededor de u contenida en la intersección de E_1 y E_2 . Por lo tanto, la familia de todas las esferas abiertas es la base para una topología de X . A tal topología para X se le conoce como *topología métrica*.

Dada una topología τ para el espacio X , se dice que el espacio topológico formado por X y τ es *metrizable* si existe una métrica tal que la topología τ es la topología métrica. Esto implica, en particular, la existencia de espacios topológicos cuya topología puede derivarse a partir de la noción de distancia. Nótese que las propiedades topológicas no dependen de la forma explícita usada para definir la distancia.

³Esta restricción puede generalizarse a todos los números reales.

Capítulo 2

Geometría diferencial.

En este capítulo tan solo veremos algunos aspectos básicos que corresponden al amplio espectro de la geometría diferencial. El concepto fundamental del cual se parte es el de *variedad*: Como vimos en el capítulo anterior, la topología esencialmente estudia la "estructura" del espacio; sin embargo, en este espacio pueden tener lugar eventos dinámicos que, a la postre, requieren una descripción en términos del concepto de *diferenciabilidad*. Es entonces cuando surge la necesidad de definir a una variedad diferencial la cual, matemáticamente, es el "sustituto más preciso de la palabra espacio" [4]. Una variedad diferenciable, como la topología en el estudio de la continuidad, es la estructura matemática natural para estudiar la diferenciabilidad.

La geometría diferencial juega un papel relevante, por ejemplo, en la descripción de la Teoría General de la Relatividad, el hecho de que sobre una variedad que representa el mundo físico (espacio-tiempo) sea posible definir de manera directa múltiples conceptos matemáticos (llámense vectores, campos vectoriales, tensores y formas diferenciales por mencionar algunos) da pie a un detallado modelaje del evento, nos permite "visualizarlo" proporcionándonos una estrecha relación entre el rigor matemático, la idea geométrica y el concepto físico.

§2.1 Variedades diferenciables.

Denotaremos por M a una variedad y diremos que esta es una *variedad diferenciable* si [5]:

- (i).- M es un espacio topológico

(ii).- M cuenta con una familia de parejas $\{(M_\alpha, \Phi_\alpha)\}$. Donde $\{M_\alpha\}$ son una familia de conjuntos abiertos que cubren a M ($M = \cup_\alpha M_\alpha$) y $\{\Phi_\alpha\}$ son homeomorfismos que van de M_α a un subconjunto abierto O_α de $\mathbb{R}^n$ ($\Phi_\alpha: M_\alpha \rightarrow O_\alpha$).

(iii).- Dados M_α y M_β tales que $M_\alpha \cap M_\beta \neq \emptyset$, el mapeo $\Phi_\beta \circ \Phi_\alpha^{-1}$ que va del abierto $\Phi_\alpha(M_\alpha \cap M_\beta) \subset \mathbb{R}^n$ al abierto $\Phi_\beta(M_\alpha \cap M_\beta) \subset \mathbb{R}^n$, $\Phi_\beta \circ \Phi_\alpha^{-1}: \mathbb{R}^n \rightarrow \mathbb{R}^n$, es infinitamente diferenciable (C^∞) en el sentido usual de cálculo vectorial elemental. Es decir, $\Phi_\beta \circ \Phi_\alpha^{-1}(x) = (f_1(x), \dots, f_n(x))$ es C^∞ si cada f_i es C^∞ para toda $i = 1, \dots, n$; entonces, si el determinante de la matriz jacobiana es no nulo en $\bar{x}_0$, existe una vecindad V de $\bar{x}_0$ y una V' de $(f_1(\bar{x}_0), \dots, f_n(\bar{x}_0))$ donde el mapeo $\Phi_\beta \circ \Phi_\alpha^{-1}$ es invertible de manera única.

La familia $\{(M_\alpha, \Phi_\alpha)\}$ que satisface (ii) y (iii) es denominada un *atlas* para M , mientras que los miembros individuales de la familia son llamados *cartas*.

De (i) y (ii) se sigue que M es un espacio localmente euclidiano, para cada $M_\alpha \subset M$ existe un homeomorfismo Φ_α que le asigna coordenadas en $\mathbb{R}^n$. Es evidente que, globalmente, M no tiene porqué ser como $\mathbb{R}^n$. La condición (iii) afirma que si dos abiertos en M tienen intersección no nula, entonces existen dos posibles conjuntos de coordenadas en $\mathbb{R}^n$: $\Phi_\alpha(M_\alpha \cap M_\beta)$ y $\Phi_\beta(M_\alpha \cap M_\beta)$; además, si deseamos cambiar de un conjunto a otro, ello debe realizarse de una manera *muy suave* (C^∞).

§2.2 Funciones y curvas.

Una función f en M es un mapeo que le asigna un número real a cada punto p en M : $f: M \rightarrow \mathbb{R}$. Es claro que cuando una región de M es mapeada diferenciablemente sobre una región de $\mathbb{R}^n$ ($g: U \subseteq M \rightarrow \mathbb{R}^n$), la composición $(f \circ g^{-1})$ es una función de $\mathbb{R}^n$ a $\mathbb{R}$; si esta función $f \circ g^{-1}$ es diferenciable en $\mathbb{R}^n$, entonces se dice que f es una función diferenciable sobre M . Es decir, como el mapeo compuesto $f \circ g^{-1}$ es, justamente, la expresión de $f(p)$ en términos de las coordenadas en p , $f = f(x^1, x^2, \dots, x^n) = f(x^i)$, entonces si esta última expresión es diferenciable en sus argumentos, f es diferenciable.

Definimos a una *curva* en una variedad M como el mapeo de la forma $\gamma: A \subset \mathbb{R} \rightarrow M$, donde A es un intervalo cerrado en $\mathbb{R}$. Es importante señalar que dos curvas son iguales si $\gamma_1(\alpha \in \mathbb{R}) = \gamma_2(\beta \in \mathbb{R})$ y $\alpha = \beta$ (i.e: no basta que la imagen de las curvas sea la misma para que las curvas sean iguales, es necesaria también la igualdad entre parámetros).

Sean $g: U \subset M \rightarrow \mathbb{R}^n$ un difeomorfismo y $\gamma: A \subset \mathbb{R} \rightarrow U \subset M$, entonces diremos que γ es una *curva suave* si la función $g \circ \gamma: A \subset \mathbb{R} \rightarrow \mathbb{R}^n$ es C^∞ . Definimos a una *curva*

simple como una curva suave a trozos y que es 1-1 en el intervalo A (i.e.: una curva simple es aquella que no se intersecta a sí misma). Por una *curva cerrada y simple* entenderemos una curva suave por pedazos $\gamma: [a, b] \subset \mathbb{R} \rightarrow M$, 1-1 en $[a, b]$ y que satisface $\gamma(a) = \gamma(b)$. Si γ satisface que $\gamma(a) = \gamma(b)$ pero no necesariamente es 1-1 en $[a, b]$, entonces llamaremos a su imagen *curva cerrada*.

§2.3 Vectores y campos vectoriales.

§2.3.1 Derivada direccional y vectores.

Sea $U \subset M$ un abierto que es mapeado diferenciablemente por g en una región de $\mathbb{R}^n$. Consideremos una función f diferenciable en U y a una curva γ contenida en U ; la variación de f a lo largo de γ , en un sistema coordenado $\{x^i\}$, es la siguiente:

Dado que la composición de g con γ es la imagen de la curva en $\mathbb{R}^n$, entonces $f \circ g^{-1}|_{g(\gamma(\lambda))} = f(x^i(\lambda))$; por lo tanto, el cambio de f sobre γ en $\{x^i\}$ es, simplemente, la derivada direccional de f a lo largo de los vectores tangentes a la imagen de la curva en $\mathbb{R}^n$: $\frac{df}{d\lambda} = \frac{dx^i}{d\lambda} \frac{\partial f}{\partial x^i}$ para toda f diferenciable (se elimina el símbolo de suma y se sobreentiende la suma sobre índices repetidos). La última expresión es válida para toda f diferenciable y, por ende, se obtiene la siguiente igualdad entre operadores:

$$\frac{d}{d\lambda} = \frac{dx^i}{d\lambda} \frac{\partial}{\partial x^i} \quad (2.1)$$

Los coeficientes $\frac{dx^i}{d\lambda}$ de (2.1) son las componentes del vector tangente en el sistema coordenado $\{x^i\}$ de la curva $\tilde{X}(\lambda) = g(\gamma(\lambda))$ en $\mathbb{R}^n$.

Si consideramos ahora una segunda curva $\beta(\mu)$, contenida en $U \subset M$, que se intersecta con $\gamma(\lambda)$ en $p \in U$, la derivada direccional de f es $\frac{df}{d\mu} = \frac{dx^i}{d\mu} \frac{\partial f}{\partial x^i}$ y en p tenemos que:

$$\left\{ a \frac{d}{d\mu} + b \frac{d}{d\lambda} \right\} f|_{p(p)} = \left\{ a \frac{dx^i}{d\mu} + b \frac{dx^i}{d\lambda} \right\} \frac{\partial f}{\partial x^i} |_{p(p)} \text{ para toda } f \text{ diferenciable.} \quad (2.2)$$

Por lo tanto, obtenemos la derivada de f en la dirección del vector con componentes $a \frac{dx^i}{d\mu} + b \frac{dx^i}{d\lambda}$ en $\mathbb{R}^n$ el cual, a su vez, es el vector tangente a una cierta curva $\tilde{X}(\tau) \subset g(U) \subset \mathbb{R}^n$ que es la imagen de $\alpha(\tau) \subset U$ bajo g . Entonces, de (2.2) se sigue que: $\frac{d}{d\tau} f|_{p(p)} =$

$\{a \frac{\partial f}{\partial x^i} + b \frac{\partial f}{\partial x^j}\}_{\sigma(p)}$ para toda f diferenciable; en otras palabras, en $g(p)$ la combinación lineal de derivadas direccionales es, de nueva cuenta, una derivada direccional.

Observemos entonces que las derivadas direccionales a lo largo de curvas, de una función diferenciable, forman un campo vectorial cuya base está dada por el conjunto $\{\frac{\partial}{\partial x^i}; i = 1, \dots, n\}$. Puesto que f es cualquier función diferenciable, entonces los operadores $\frac{\partial}{\partial x^i}$ heredan la propiedad de campo vectorial y tienen como base del mismo al conjunto $\{\frac{\partial}{\partial x^i}; i = 1, \dots, n\}$. El "espacio de los operadores" y de las derivadas direccionales está en correspondencia 1-1, como además esta misma correspondencia ocurre entre el espacio de los vectores tangentes en $\mathbb{R}^n$ y las derivadas direccionales, entonces el "espacio de los operadores" y de los vectores tangentes están relacionados 1-1; es por esta razón que al operador $\frac{\partial}{\partial x^i}$ se le denominará, en adelante, *vector tangente* [4].

El vector $\frac{\partial}{\partial x^i}$ en el sistema coordenado $\{x^i\}$ tiene, en la base $\{\frac{\partial}{\partial x^i}\}$, componentes $\frac{\partial x^i}{\partial x^i}$. En una base arbitraria $\{\tilde{e}_i\}$ el vector es $\frac{\partial}{\partial x^i} = v^i \tilde{e}_i$, con v^i la i -ésima componente del vector. Estos vectores pertenecen al espacio tangente a $p \in M$, denotado por $T_p M$, de tal manera que el espacio vectorial formado por los vectores $\frac{\partial}{\partial x^i}$ en el punto p es, efectivamente, $T_p M$. Para todo punto p en M , $T_p M$ es un espacio vectorial con la misma dimensión (n) que la variedad y, por ende, toda colección de n vectores linealmente independientes en $T_p M$ será una base para $T_p M$.

Dado cualquier vector $\tilde{v}$ en $p \in M$, existe una única curva para la cual $\tilde{v}$ es un vector tangente, en al menos una vecindad de p (§2.3.3); en otras palabras, hablar de vectores es equivalente a hablar de vectores tangentes.

§2.3.2 Campos vectoriales.

En general, una transformación V es un *campo vectorial tangente* en un conjunto abierto $U \subset S$ de una superficie regular S si V asigna a cada p en U un vector $v_p \in T_p S$.

Si para todo p en M contamos con una base para $T_p M$, digamos $\{\tilde{e}_i\}$, entonces tendremos una base para campos vectoriales sobre M . Dado que los campos vectoriales sobre una variedad son, por definición, transformaciones del tipo $V : M \rightarrow T_p M$, entonces la imagen del campo está dada por $\tilde{v} = v^i \tilde{e}_i$, donde v^i es una función sobre la variedad. Diremos que un *campo vectorial es diferenciable* si cada una de las funciones del conjunto $\{v^i\}$ es diferenciable.

§2.3.3 Curvas integrales.

Supongamos que todo punto p , de un subconjunto abierto U de M , pertenece a una y sólo una curva suave contenida en U . Es evidente, en este caso, que el conjunto de los vectores tangentes a las curvas forman un campo vectorial. Resulta natural preguntarnos entonces si será cierto que, dado un campo vectorial, es posible partir de un $p \in M$ y hallar una curva cuyo vector tangente es siempre el campo vectorial en cualquier punto a través del cual pasa la curva. La respuesta a este requerimiento es afirmativa, siempre y cuando el campo tenga al menos primera derivada continua ($V \in C^k; k \geq 1$). A las curvas que tengan por vector tangente a elementos del campo se les llamará *curvas integrales*.

Notemos que en un sistema coordenado $\{x^i\}$ las componentes v_p^i del campo son $v^i(x^i)$, la condición para que el campo sea un vector tangente a una curva con parámetro λ se expresa, por consiguiente, como sigue: $\frac{dx^i(\lambda)}{d\lambda} = v^i(x^i(\lambda))$; si el campo es al menos C^1 entonces, por el teorema de existencia y unicidad para ecuaciones diferenciales ordinarias de primer orden, cada ecuación del sistema para $x^i(\lambda)$, siempre tiene solución y esta es única en alguna vecindad del punto inicial p . El conjunto de todas las curvas integrales que cubren una región de M es conocido como una congruencia (i.e. las curvas integrales que provienen de campos $C^{k \geq 1}$ forman una congruencia).

Es posible encontrar una relación entre dos puntos que pertenecen a una misma curva integral, para ello supongamos que M es analítica (C^∞) y que los valores coordenados $x^i(\lambda)$, de los puntos que pertenecen a la curva integral del campo $\tilde{Y} = \frac{d}{d\lambda}$, son funciones analíticas de λ . Entonces, los puntos con parámetro λ_0 y $\lambda_0 + \epsilon$ tienen la siguiente relación:

$$x^i(\lambda_0 + \epsilon) = x^i(\lambda_0) + \epsilon \frac{dx^i}{d\lambda} \Big|_{\lambda_0} + \frac{1}{2!} \epsilon^2 \frac{d^2 x^i}{d\lambda^2} \Big|_{\lambda_0} + \dots = \{1 + \epsilon \frac{d}{d\lambda} + \frac{1}{2!} \epsilon^2 \frac{d^2}{d\lambda^2} + \dots\} x^i \Big|_{\lambda_0} = \{e^{\epsilon \frac{d}{d\lambda}}\} x^i \Big|_{\lambda_0}$$

entonces:

$$x^i(\lambda_0 + \epsilon) = \{e^{\epsilon \tilde{Y}}\} x^i \Big|_{\lambda_0} \quad (2.3)$$

Al operador $e^{\epsilon \tilde{Y}}$, que da un "movimiento" finito, se le conoce como exponenciación del operador $\epsilon \tilde{Y}$, el cual es un "movimiento" infinitesimal a lo largo de la curva integral.

§2.3.4 Paréntesis de Lie.

Definimos al *paréntesis de Lie* de dos campos vectoriales $\frac{d}{dx^a}$ y $\frac{d}{dx^b}$ como el conmutador $[\frac{d}{dx^a}, \frac{d}{dx^b}] \equiv \frac{d}{dx^a} \frac{d}{dx^b} - \frac{d}{dx^b} \frac{d}{dx^a}$. La expresión para estos campos en el sistema de coordenadas $\{x^i\}$

es: $\frac{d}{dx} = v^i \frac{\partial}{\partial x^i}$ y $\frac{d}{d\mu} = w^j \frac{\partial}{\partial x^j}$, el paréntesis de Lie en este sistema coordenado es, entonces, el siguiente:

$$\left[\frac{d}{d\lambda}, \frac{d}{d\mu} \right] = \{v^i \frac{\partial w^j}{\partial x^i} - w^i \frac{\partial v^j}{\partial x^i}\} \frac{\partial}{\partial x^j} \quad (2.4)$$

Nótese que el paréntesis de Lie en cada punto es un vector que pertenece a $T_p M$, en consecuencia $\{\frac{d}{d\lambda}, \frac{d}{d\mu}\}$ es un campo vectorial con componentes $v^i \frac{\partial w^j}{\partial x^i} - w^i \frac{\partial v^j}{\partial x^i}$ en el sistema coordenado $\{x^i\}$.

Diremos que una base es coordenada si al desplazarnos sucesivamente un parámetro ϵ a lo largo de cada una de las curvas integrales de sus elementos (a partir de un punto p), sin importar el orden, llegamos al mismo punto. El paréntesis de Lie nos ayuda a discernir, justamente, cuándo una base es una base coordenada:

Si n campos $B = \{\tilde{A}_i = \frac{d}{dx^i}; i = 1, \dots, n\}$ son linealmente independientes en todo punto de un abierto U n -dimensional de M , entonces cualquier campo sobre U puede escribirse como $\tilde{v} = f^i \tilde{A}_i$, donde $f^i : U \subset M \rightarrow \mathbb{R}$ (i.e: B es una base para campos vectoriales sobre U). La condición que garantiza que B sea una base coordenada en U (i.e: $\{\lambda_i; i = 1, \dots, n\}$ son coordenadas para U) es la nulidad del paréntesis de Lie; es decir: $[\tilde{A}_i, \tilde{A}_j] = 0$ en $U \forall i, j = 1, \dots, n$.

§2.4 Haces fibrados.

Los conocimientos que hemos adquirido hasta el momento nos permiten ya definir a una de las estructuras matemáticas más importantes en geometría diferencial: los haces fibrados. Antes de ello, cabe advertir la necesidad de procurar un poco de mayor atención a este tema, pues, como constatará más adelante el lector (§5.4), estas estructuras son, en toda la extensión de la palabra, fundamentales en el estudio de los modelos σ no lineales.

Definimos a un *haz fibrado* (E, B, Π, F, G) como la colección de los siguientes requerimientos [5]:

- (i).- Un espacio topológico E llamado *espacio total*.
- (ii).- Un espacio topológico B denominado *espacio base*, y una *proyección* Π de E sobre B ($\Pi : E \rightarrow B$).
- (iii).- Un espacio topológico F llamado *fibra típica* o simplemente *fibra*.

(iv).- Un grupo G, denominado *grupo de estructura* del haz fibrado, de homeomorfismos de la fibra F.

(v).- Trivialidad local del haz: Existe un conjunto de vecindades coordenadas abiertas u_α que cubren a B, tales que para cada u_α esta dado el homeomorfismo Φ_α :

$$\Phi_\alpha : \Pi^{-1}(u_\alpha) \rightarrow u_\alpha \times F$$

donde Φ_α^{-1} satisface $\Pi\Phi_\alpha^{-1}(x, f) = x$, con $x \in u_\alpha$ y $f \in F$.

El grupo G surge de considerar la transición de un conjunto de coordenadas locales, determinadas por la pareja (u_α, Φ_α) , hasta otro conjunto de coordenadas, determinadas por (u_β, Φ_β) . Es decir, supongamos que $u_\alpha \cap u_\beta \neq \emptyset$, entonces $\Phi_\alpha \circ \Phi_\beta^{-1}$ es un homeomorfismo de $(u_\alpha \cap u_\beta) \times F$ a $(u_\alpha \cap u_\beta) \times F$. Sea $x \in u_\alpha \cap u_\beta$ fija y $f \in F$ variable, entonces $\Phi_\alpha \circ \Phi_\beta^{-1}|_x$ es un mapeo de F en F:

$$\Phi_\alpha \circ \Phi_\beta^{-1}|_x = g_{\alpha\beta}(x) : F \rightarrow F$$

$g_{\alpha\beta}(x)$ es conocida como *función de transición* y es un homeomorfismo de la fibra F, es una transición porque su dominio está definido por la pareja (u_β, Φ_β) y su imagen por (u_α, Φ_α) . Las funciones de transición tienen las siguientes propiedades (llamadas *relaciones de compatibilidad*):

(a).- $g_{\alpha\alpha}(x) = \text{identidad}$, $\forall x \in u_\alpha$.

(b).- $g_{\alpha\beta}(x) = \{g_{\beta\alpha}(x)\}^{-1}$, $\forall x \in u_\alpha \cap u_\beta$.

(c).- $g_{\alpha\beta}(x)g_{\beta\gamma}(x) = g_{\alpha\gamma}(x)$, $\forall x \in u_\alpha \cap u_\beta \cap u_\gamma$.

El conjunto de todos estos homeomorfismos $g_{\alpha\beta}$ para todas las elecciones de (u_α, Φ_α) forman un grupo y este es, justamente, el grupo de estructura G del haz fibrado E.

Es importante señalar que un haz fibrado puede reproducirse si son dados B, las funciones de transición $g_{\alpha\beta}(x)$, la fibra F y el grupo G. Todo E puede construirse por la aplicación de una relación de equivalencia para el conjunto $\tilde{E}$, donde $\tilde{E}$ consiste en todos los productos de la forma $u_\alpha \times F$: $\tilde{E} = \bigcup_\alpha u_\alpha \times F$.

Un elemento de $\tilde{E}$ es (x, f) . La relación de equivalencia ($\sim$) la definimos como sigue: sean $(x, f) \in u_\alpha \times F$ y $(y, h) \in u_\beta \times F \Rightarrow (x, f) \sim (y, h)$ si $x = y$ y $g_{\alpha\beta}(x)h = f$. El haz E es el conjunto de todas las clases de equivalencia bajo esta relación de equivalencia: $E = \tilde{E}/\sim$.

Si denotamos por $[(x, f)]$ a todas las clases de equivalencia que están o son de la forma $(x, f) \in u_\alpha \times F$ podemos entonces definir la proyección Π por $[(x, f)] \xrightarrow{\Pi} x$ y hallar Φ_α^{-1} como sigue:

$$\Phi_\alpha^{-1} : u_\alpha \times F \longrightarrow \Pi^{-1}(u_\alpha)$$

$$(x, f) \mapsto [(x, f)]$$

en consecuencia:

$$\Pi\Phi^{-1}(x, f) = \Pi[(x, f)] = x$$

Con lo cual concluimos la construcción del haz E , cuyas funciones de transición son $g_{\alpha\beta}(x)$.

Los haces fibrados pueden dividirse en dos categorías, a saber: haces fibrados *triviales*, donde $E = B \times F$, y *no triviales*, donde $E \neq B \times F$. La trivialidad de un haz (E, B, Π, F, G) puede derivarse al examinar el haz principal asociado a E : Siempre es posible construir a partir de E un *haz principal*, se mantiene el mismo espacio base B y las mismas funciones de transición $g_{\alpha\beta}(x)$, cambiando únicamente F por el grupo G . A un haz principal usualmente se le denota por $P(E)$.

Antes de enunciar el teorema que habrá de determinar cuándo un haz es trivial, definiremos el concepto de sección: Una *sección* de un haz E es un mapeo continuo $s : B \rightarrow E$ tal que $\Pi s(x) = x, \forall x \in B$.

Teorema 2.1 $P(E)$ y E son triviales $\iff P(E)$ tiene una sección.

Demostración: $\Rightarrow$.- Que $P(E)$ sea trivial $\Rightarrow \exists$ un homeomorfismo $\Phi : P(E) \rightarrow B \times G$, donde Φ^{-1} satisface $\Pi\Phi^{-1}(x, g) = x, \forall x \in B$ y $g \in G$.

Sea h el mapeo continuo $h(x) = (x, g)$ para alguna $g \in G$ y $\forall x \in B$ ($h : B \rightarrow B \times G$). Entonces $(\Phi^{-1} \circ h)(x) = \Phi^{-1}(x, g)$ para alguna $g \in G$, $\Phi^{-1}(x, g) \in P(E)$ y además $\Phi^{-1} \circ h$ es continua.

Sea $s(x) = (\Phi^{-1} \circ h)(x) \Rightarrow s$ es continua y $s : B \rightarrow P(E)$.

$\Pi s(x) = \Pi(\Phi^{-1} \circ h)(x) = \Pi(\Phi^{-1}(x, g)) = x, \forall x \in B$ y para alguna $g \in G \Rightarrow \Pi s(x) = x \forall x \in B$ por lo tanto $s : B \rightarrow P(E)$ es un mapeo continuo que satisface que $\Pi s(x) = x \forall x \in B \Rightarrow s$ es una sección del haz $P(E)$. Por lo tanto, $P(E)$ tiene una sección.

$\Leftarrow$.- Comencemos por mostrar que $P(E)$ es trivial. Puesto que $P(E)$ es un haz principal $\Rightarrow s(x) \in G$, también si g es cualquier elemento de $G \Rightarrow gs(x)$ pertenece a la fibra sobre el punto $x \in B$. De hecho, nótese que todos los elementos de la fibra sobre x son de la forma $gs(x)$ para alguna $g \in G$. Si x varía, dado que el haz es la unión de todas sus fibras, entonces todos los elementos $k \in P(E)$ pueden expresarse como $gs(x)$, donde $g \in G$ y $x \in B$. En consecuencia, hallamos la siguiente relación 1-1, continua y con inversa continua:

$$\Phi(k = gs(x)) = (x, g)$$

Φ es un homeomorfismo de $P(E)$ a $B \times G$, por lo tanto $P(E)$ es trivial.

Para mostrar la trivialidad de E , que se sigue de la de $P(E)$, observemos lo siguiente: Sean E y E' dos haces con el mismo espacio base, fibra y grupo, pero con coordenadas y cubiertas dadas por $\{\Phi_\alpha, u_\alpha\}$ y $\{\Psi_\alpha, u_\alpha\}$ respectivamente. El mapeo $\Psi_\alpha \circ \Phi_\alpha^{-1} : u_\alpha \times F \rightarrow u_\alpha \times F$ restringido a $x \in u_\alpha$ es un homeomorfismo $\lambda_\alpha(x)$ sobre la fibra F . Dado que $\lambda_\alpha(x) = \Psi_\alpha \circ \Phi_\alpha^{-1}(x)$ y $g'_{\alpha\beta}(x) = \Psi_\alpha \circ \Psi_\beta^{-1}(x) \forall x \in u_\alpha \cap u_\beta$, entonces:

$$\begin{aligned}\lambda_\alpha^{-1}(x)g'_{\alpha\beta}(x)\lambda_\beta(x) &= (\Psi_\alpha \circ \Phi_\alpha^{-1}(x))^{-1} \circ (\Psi_\alpha \circ \Psi_\beta^{-1}(x)) \circ (\Psi_\beta \circ \Phi_\beta^{-1}(x)) \\ &\Rightarrow \lambda_\alpha^{-1}(x)g'_{\alpha\beta}(x)\lambda_\beta(x) = \Phi_\alpha \circ \Phi_\beta^{-1}(x) = g_{\alpha\beta}(x) \forall x \in u_\alpha \cap u_\beta\end{aligned}$$

Entonces, la relación entre las correspondientes funciones de transición es:

$$g_{\alpha\beta}(x) = \lambda_\alpha^{-1}(x)g'_{\alpha\beta}(x)\lambda_\beta(x) \forall x \in u_\alpha \cap u_\beta$$

Diremos que E y E' son equivalentes si sus funciones de transición están relacionadas de acuerdo a la última expresión.

Puesto que $P(E)$ es trivial tenemos el homeomorfismo global $\Phi : P(E) \rightarrow B \times F$. Si las coordenadas y cubierta de $P(E)$ están dadas por $\{\sigma_\alpha, u_\alpha\}$, podemos construir, para x fija, el homeomorfismo $\lambda_\alpha(x)$ sobre F que pertenece al grupo de estructura G : $\lambda_\alpha(x) = \Phi \circ \sigma_\alpha^{-1}(x)$. Entonces:

$$\begin{aligned}\lambda_\alpha(x)g_{\alpha\beta}(x)\lambda_\beta^{-1}(x) &= (\Phi \circ \sigma_\alpha^{-1}(x)) \circ (\sigma_\alpha \circ \sigma_\beta^{-1}(x)) \circ (\Phi \circ \sigma_\beta^{-1}(x))^{-1} = \text{identidad} = e \\ &\Rightarrow g_{\alpha\beta}(x) = \lambda_\alpha^{-1}(x)\lambda_\beta(x)\end{aligned}$$

Como las funciones de transición de E y $P(E)$ son las mismas, entonces para E se satisface la relación anterior, puesto que $g_{\alpha\beta}(x) = \lambda_\alpha^{-1}(x)g'_{\alpha\beta}(x)\lambda_\beta(x)$ obtenemos que $g'_{\alpha\beta}(x) = e$ para un haz equivalente a E . Pero si las funciones de transición de un haz son la identidad sobre F , entonces es inmediato que el haz es un producto y por lo tanto es trivial: *las fibras son "pegadas" sin torsión*. Como el haz trivial es equivalente a E , entonces E es trivial. $\square$

Un ejemplo sencillo de un haz trivial es el cilindro C el cual es, por supuesto, un producto global. Globalmente $C = S^1 \times F$, con F (fibra típica) un segmento de línea y base $B = S^1$ un círculo. La banda de Möbius es, por otro lado, un ejemplo de haz fibrado no trivial; $B = S^1$ y F es también un segmento de línea, sin embargo no es globalmente un producto, a diferencia del cilindro para el cual G puede reducirse siempre a la identidad puesto que C es trivial, el grupo de estructura para la banda de Möbius es $G = \{e = \text{identidad}, g = \text{reflexión de}$

F alrededor del punto medio) con dos abiertos como cubierta de S^1 (la identidad actúa en una intersección y la reflexión en la otra). Otro ejemplo de haz fibrado es el espacio-tiempo Newtoniano: puesto que la liga entre observadores radica en que los eventos son simultáneos, la base habrá de ser entonces el tiempo t y las fibras el espacio $\mathbb{R}^3$.

Existen múltiples ejemplos de haces fibrados, por señalar uno más citemos el que consta de una variedad diferenciable M como base y sus espacios tangentes como fibras, este haz fibrado es conocido como *haz tangente* y es denotado por TM (el haz fibrado TM es una de las variedades abstractas más importantes en física); si la variedad base es n -dimensional, entonces TM es $2n$ -dimensional, dado que el espacio tangente tiene la misma dimensión que $M=B$.

§2.5 Uno-formas.

Una uno-forma es definida como una función lineal que actúa sobre vectores y cuya imagen es un real: Si v y w son vectores, entonces la uno-forma $\tilde{w}: V \rightarrow \mathbb{R}$ es tal que (para a y b en $\mathbb{R}$) $\tilde{w}(av + bw) = a\tilde{w}(v) + b\tilde{w}(w) \in \mathbb{R}$.

Podemos definir tanto la adición de uno-formas como su multiplicación por números reales:

$$(i).- (a\tilde{w})(v) = a(\tilde{w}(v)) \quad \forall v \in V.$$

$$(ii).- (\tilde{w} + \tilde{q})(v) = \tilde{w}(v) + \tilde{q}(v) \quad \forall v \in V.$$

Adicionalmente, cabe señalar que las uno-formas en $p \in M$ satisfacen los axiomas de un espacio vectorial. El espacio vectorial que conforman es llamado el *espacio dual* del espacio vectorial $T_p M$ y es denotado por $T_p^* M$. A la aplicación de una uno-forma $\tilde{w}$ sobre un vector v ($\tilde{w}(v)$) con frecuencia se le llama *contracción* de $\tilde{w}$ con v ; otras notaciones para la citada contracción son las siguientes: $\tilde{w}(v) \equiv v(\tilde{w}) \equiv \langle \tilde{w}, v \rangle \equiv \tilde{w} \lrcorner v$.

Cualesquiera n uno-formas linealmente independientes constituyen una base para el espacio vectorial n dimensional de uno-formas ($T_p^* M$). Sin embargo, dada una base $\{\tilde{e}_i, i = 1, \dots, n\}$ para los vectores en $T_p M$, se induce una base preferente para $T_p^* M$ llamada *base dual* $\{\tilde{w}^i, i = 1, \dots, n\}$. La base dual se define como sigue:

Si $\tilde{v}$ es un vector cualquiera en $T_p M$ entonces $\tilde{w}^i$ produce la i -ésima componente de $\tilde{v}$: $\tilde{w}^i(\tilde{v}) = v^i$. Entonces, por la linealidad tenemos que $\tilde{w}^i(\tilde{v} = v^j \tilde{e}_j) = v^j \tilde{w}^i(\tilde{e}_j) = v^i$ (con la convención de suma sobre índices repetidos), en consecuencia $\tilde{w}^i(\tilde{e}_j) = \delta_j^i$. Sea $\tilde{q}$ cualquier uno-forma (1-forma), la contracción de ésta con un vector arbitrario $\tilde{v}$ es entonces

la siguiente: $\tilde{q}(\tilde{v}) = \tilde{q}(v^j \tilde{e}_j) = v^j \tilde{q}(\tilde{e}_j) = \tilde{w}^j(\tilde{v}) \tilde{q}(\tilde{e}_j)$. Puesto que $\tilde{v}$ es arbitrario obtenemos que: $\tilde{q} = \tilde{q}(\tilde{e}_j) \tilde{w}^j$; $\tilde{q}(\tilde{e}_j) = q_j$ es la j -ésima componente de la 1-forma $\tilde{q}$ en la base dual a $\{\tilde{e}_i\}$, por lo tanto la 1-forma $\tilde{q}$ queda escrita de manera natural en la base inducida: $\tilde{q} = q_j \tilde{w}^j$.

Por analogía con un campo vectorial, un *campo de 1-formas* es una regla que da una 1-forma para cada $p \in U \subseteq M$. Las reglas (i) y (ii), mencionadas con anterioridad, se extienden a campos con la única salvedad de que en este caso a es una función sobre $U \subseteq M$, no necesariamente constante. En cuanto a la diferenciabilidad de un campo de 1-formas diremos que el campo $\tilde{\Omega}$ es C^∞ si la función $\tilde{\Omega}(\tilde{V})$, construida con un campo vectorial $\tilde{V}$, es C^∞ para cualquier $\tilde{V} \in C^\infty$.

Una 1-forma que con frecuencia es utilizada es el gradiente de una función f , éste es denotado por $\tilde{d}f$ y actúa sobre un vector tangente arbitrario $\frac{d}{dx}$ como sigue: $\tilde{d}f(\frac{d}{dx}) = \frac{df}{dx}$ (i.e: el gradiente de f en cualquier punto p es el elemento de T_p^*M cuyo valor sobre un elemento $\tilde{v} \in T_pM$ es la derivada direccional de f a lo largo de una curva cuya tangente en p es $\tilde{v}$).

Nótese que si el conjunto de campos vectoriales $\{\tilde{e}_i\}$ es una base en cada $p \in U \subseteq M$, entonces los campos $\{\tilde{w}^i\}$ son también una base $\forall p \in U$. Sea $\{x^i\}$ un sistema coordenado sobre U , $\{x^i\}$ define una base natural $\{\frac{\partial}{\partial x^i}\}$ para campos vectoriales; en consecuencia, se induce una base coordenada natural para campos de 1-formas:

$$\tilde{w}^i(\frac{\partial}{\partial x^j}) = \delta_j^i = \frac{\partial x^i}{\partial x^j} = \tilde{d}x^i(\frac{\partial}{\partial x^j})$$

$$\Rightarrow \{\tilde{w}^i\} = \{\tilde{d}x^i\}$$

Por lo tanto, en el sistema coordenado $\{x^i\}$ un campo de 1-formas puede ser expresado de manera natural como $\tilde{\Omega} = \Omega_i \tilde{d}x^i$. Análogamente a un campo vectorial, el campo $\tilde{\Omega}$ es C^∞ sii las componentes $\{\Omega_i\}$ del campo asociadas con una base C^∞ para campos vectoriales son funciones C^∞ .

La introducción de 1-formas nos permite construir un haz fibrado, conocido como *haz cotangente*, que es el equivalente a TM para vectores: el haz cotangente T^*M tiene como base una variedad M y como fibra sobre $p \in M$ al espacio vectorial T_p^*M . Una sección en T^*M , evidentemente, es un campo de 1-formas.

§2.6 Tensores y campos tensoriales.

Comencemos por definir a un *tensor*, en un punto p en M , diciendo que un tensor del tipo $\binom{n}{m}$ es una función lineal con imagen en $\mathbb{R}$ que toma como argumentos n 1-formas y m vectores; es decir, un tensor es una aplicación multilineal sobre 1-formas y vectores cuya imagen es un real.

Como caso particular, tenemos que los vectores son tensores $\binom{1}{0}$, en tanto que las 1-formas son tensores $\binom{0}{1}$. Por convención, una función escalar sobre una variedad es el tensor $\binom{0}{0}$.

Obsérvese que un tensor $T \binom{1}{1}$ puede pensarse como una función lineal $T(\cdot; \vec{v})$ de vectores $\vec{v}$, también, como una función lineal $T(\vec{w}; \cdot)$ de 1-formas.

Un *campo tensorial* $\binom{n}{m}$ en $U \subseteq M$ es una regla que asigna un tensor $\binom{n}{m}$ a cada $p \in U$. La linealidad puntual se extiende al campo tensorial; en el campo las constantes multiplicativas en los argumentos del tensor son ahora funciones sobre M .

Se dice que un campo tensorial $F \binom{n}{m}$ es C^∞ si la función $F(\vec{w}_1, \vec{w}_2, \dots, \vec{w}_n; \vec{v}_1, \vec{v}_2, \dots, \vec{v}_m)$ es C^∞ para cualesquiera campos vectoriales $\{\vec{v}_1, \dots, \vec{v}_m\}$ y de 1-formas $\{\vec{w}_1, \dots, \vec{w}_n\}$ que sean C^∞ .

§2.7 Producto y componentes tensoriales.

Consideremos los tensores $A \binom{k}{m}$ y $B \binom{n}{r}$. Definimos al *producto tensorial* (también conocido como directo) de A y B como la operación que devuelve un tensor del tipo $\binom{k+n}{m+r}$ mediante la siguiente regla:

$$A \otimes B(\vec{a}_1, \dots, \vec{a}_k, \vec{b}_1, \dots, \vec{b}_n; \vec{v}_1, \dots, \vec{v}_m, \vec{w}_1, \dots, \vec{w}_r) \equiv A(\vec{a}_1, \dots, \vec{a}_k; \vec{v}_1, \dots, \vec{v}_m) B(\vec{b}_1, \dots, \vec{b}_n; \vec{w}_1, \dots, \vec{w}_r)$$

Las componentes de un tensor son los valores del mismo cuando se toman como argumentos a los elementos de las bases de 1-formas y de vectores. Si T es un tensor $\binom{n}{m}$ y los conjuntos $\{\vec{w}^i\}$ y $\{\vec{e}_i\}$ son la base para 1-formas y vectores respectivamente, entonces las componentes de T están dadas por la siguiente expresión:

$$T(\vec{w}^i, \vec{w}^j, \dots, \vec{w}^k; \vec{e}_l, \vec{e}_p, \dots, \vec{e}_q) \equiv T^{ij \dots k}_{l p \dots q}$$

donde $\{i, j, \dots, k\}$ son n índices y $\{l, p, \dots, q\}$ son m índices.

Dadas las definiciones de producto y componente tensorial es posible dar una expresión explícita de un tensor en términos de estos conceptos. Sea T un tensor $\binom{n}{m}$, entonces la

multilinealidad implica que:

$$T(w_{(1)i}\tilde{w}^i, \dots, w_{(n)k}\tilde{w}^k; v_{(1)}^i\tilde{e}_i, \dots, v_{(m)}^q\tilde{e}_q) = T^{i\dots k}_{i\dots q}w_{(1)i}\dots w_{(n)k}v_{(1)}^i\dots v_{(m)}^q$$

dado que $v_{(N)}^i = \tilde{w}^i(\tilde{v}_{(N)})$ y $w_{(N)i} = \tilde{w}_{(N)}(\tilde{e}_i) = \tilde{e}_i(\tilde{w}_{(N)})$ obtenemos lo siguiente:

$$T(\tilde{w}_{(1)}, \dots, \tilde{w}_{(n)}; \tilde{v}_{(1)}, \dots, \tilde{v}_{(m)}) = T^{i\dots k}_{i\dots q}\tilde{e}_i(\tilde{w}_{(1)})\dots\tilde{e}_k(\tilde{w}_{(n)})\tilde{w}^i(\tilde{v}_{(1)})\dots\tilde{w}^q(\tilde{v}_{(m)})$$

de la definición de producto tensorial se obtiene que:

$$T(\tilde{w}_{(1)}, \dots, \tilde{w}_{(n)}; \tilde{v}_{(1)}, \dots, \tilde{v}_{(m)}) = T^{i\dots k}_{i\dots q}\tilde{e}_i \otimes \dots \otimes \tilde{e}_k \otimes \tilde{w}^i \otimes \dots \otimes \tilde{w}^q(\tilde{w}_{(1)}, \dots, \tilde{w}_{(n)}; \tilde{v}_{(1)}, \dots, \tilde{v}_{(m)})$$

por lo tanto:

$$T = T^{i\dots k}_{i\dots q}\tilde{e}_i \otimes \dots \otimes \tilde{e}_k \otimes \tilde{w}^i \otimes \dots \otimes \tilde{w}^q$$

§2.8 Transformaciones de base.

Supongamos que tenemos en $T_p M$ una base vectorial $\{\tilde{e}_i, i = 1, \dots, n\}$ y que deseamos sustituirla por otra base $\{\tilde{e}'_k, k = 1, \dots, n\}$. En $T_p M$ existe una transformación lineal Λ que nos permite realizar el cambio de base: $\tilde{e}'_k = \Lambda^i_k \tilde{e}_i$.

Para hallar la expresión explícita de Λ^i_k notemos que $\Lambda^j_k \tilde{w}^i(\tilde{e}'_j) = \Lambda^j_k \delta^i_j = \Lambda^i_k$. Debido a la linealidad de las 1-formas $\Lambda^j_k \tilde{w}^i(\tilde{e}'_j) = \tilde{w}^i(\tilde{e}'_k)$ y, por lo tanto, $\Lambda^i_k = \tilde{w}^i(\tilde{e}'_k)$.

La matriz de transformación mapea vectores base en vectores base de manera unívoca, en consecuencia es no singular y, de acuerdo a la regla de transformación, es posible deducir que la inversa de Λ^i_k es Λ^k_i . En efecto:

$$\begin{aligned}\tilde{e}_i &= \Lambda^k_i \tilde{e}'_k = \Lambda^k_i \Lambda^l_k \tilde{e}_l \\ \Rightarrow \tilde{e}_i &= \Lambda^k_i \Lambda^l_k \tilde{e}_l \Rightarrow \Lambda^k_i \Lambda^l_k = \delta^l_i\end{aligned}$$

La transformación en la base de vectores induce una transformación en la base de 1-formas, hallemos tal transformación:

$$\begin{aligned}\Lambda^p_i \tilde{w}^i(\tilde{e}'_k) &= \Lambda^p_i \Lambda^l_k \tilde{w}^i(\tilde{e}_l) = \delta^p_k = \tilde{w}^p(\tilde{e}'_k) \\ \Rightarrow \tilde{w}^p &= \Lambda^p_i \tilde{w}^i\end{aligned}$$

Ahora es sencillo encontrar la transformación de las componentes:

$$v^i = \bar{w}^i(\bar{v}) = \Lambda^i_j \bar{w}^j(\bar{v}) = \Lambda^i_j v^j$$

$$q'_k = \bar{q}(\bar{e}'_k) = \bar{q}(\Lambda^j_k \bar{e}_j) = \Lambda^j_k \bar{q}(\bar{e}_j) = \Lambda^j_k q_j$$

Como se puede observar, las componentes de 1-formas se transforman como la base vectorial, es por esta razón que a las 1-formas se les suele llamar *vectores covariantes*; por otro lado, la transformación de las componentes vectoriales obedece la ley opuesta a la que gobierna en la base vectorial y, por ende, a los vectores se les conoce también como *vectores contravariantes*.

La generalización a un tensor T (n) es directa:

$$T^{i_1 \dots i_n}_{j_1 \dots j_m} = \Lambda^{i_1}_{r_1} \dots \Lambda^{i_n}_{r_n} \Lambda^{j_1}_{s_1} \dots \Lambda^{j_m}_{s_m} T^{r_1 \dots r_n}_{s_1 \dots s_m}$$

donde son n índices superiores y m índices inferiores.

Las componentes de un tensor evidentemente dependen de la elección de la base, sin embargo nótese que la cantidad $A^i B_i$ es independiente de la elección de la base. En efecto:

$$A^i B_i = A(\bar{w}^i) B(\bar{e}_i) = A(\Lambda^i_k \bar{w}^k) B(\Lambda^j_i \bar{e}_j) = \Lambda^i_k \Lambda^j_i A^k B_j = \delta^j_k A^k B_j = A^j B_j$$

A la cantidad $A^i B_i$ se le conoce como *contracción* de un tensor ($\binom{n}{0}$) con uno ($\binom{0}{1}$), el resultado es independiente de la base y el tensor que se obtiene es del tipo ($\binom{n}{0}$). Si ahora A es un tensor ($\binom{n}{m}$) y B uno ($\binom{k}{l}$), la contracción de q índices repetidos (donde $q < \{n, m, k, l\}$) es un tensor ($\binom{n+k-q}{m+l-q}$).

Definimos a un *escalar* como un tensor ($\binom{0}{0}$); es decir, un escalar puede pensarse como una función sobre M cuya definición no depende de la elección de la base. Por ejemplo, si A y B son campos definidos en M , la contracción $A^i B_i$ es un escalar; sin embargo, A^i y B_i , que son funciones sobre M , no son escalares puesto que su valor no es independiente de la elección de la base.

La contracción sobre pares de índices, la multiplicación de componentes de dos tensores, la multiplicación de todas las componentes por un número y la adición (resta) de componentes de tensores del mismo tipo, forman parte de las llamadas operaciones tensoriales: una *operación tensorial* es una operación sobre componentes que produce componentes del mismo tensor independientemente de la base.

Una ecuación que contenga componentes combinadas que usen sólo estas operaciones es llamada una *ecuación tensorial*. La ecuación tensorial es la misma en todas las bases.

§2.9 El tensor métrico sobre un espacio vectorial.

El tensor métrico g es un tensor ($\frac{9}{2}$) simétrico y está definido como sigue: $g(\vec{v}, \vec{u}) = g(\vec{u}, \vec{v}) \equiv \vec{u} \cdot \vec{v}$.

Por definición tenemos entonces que, en la base $\{\vec{e}_i\}$, las componentes del tensor métrico son: $g_{ij} = g(\vec{e}_i, \vec{e}_j) = \vec{e}_i \cdot \vec{e}_j$. Si $g_{ij} = \delta_{ij}$, entonces al tensor métrico se le denomina *Métrica Euclídiana* y al espacio vectorial se le llama *Espacio Euclídeo*.

La forma sencilla en las componentes métricas no siempre se presenta de manera natural, sin embargo el siguiente teorema garantiza que siempre tenemos la libertad de elegir una nueva base en la cual las componentes métricas son, de nueva cuenta, simples:

Teorema 2.2 (Teorema de diagonalización). *Sea V un espacio vectorial de dimensión finita sobre el cuerpo de los números reales, y sea f una forma bilineal simétrica (la métrica es una de ellas) sobre V . Entonces existe una base de V en la que f está representada por una matriz diagonal.*

Demostración: La demostración consiste en exhibir la existencia de una base $B = \{\alpha_1, \dots, \alpha_n\}$ tal que $f(\alpha_i, \alpha_j) = 0$ para $i \neq j$.

Evidentemente el teorema se cumple si $f = 0$ o $n = 1$; en consecuencia, habremos de suponer $f \neq 0$ y $n > 1$. Notemos que si $f(\alpha, \alpha) = 0$ para todo α en V , la forma cuadrática asociada q ($q(\alpha) = f(\alpha, \alpha)$) es idénticamente nula y la identidad de polarización ($f(\alpha, \beta) = \frac{1}{4}q(\alpha + \beta) - \frac{1}{4}q(\alpha - \beta)$) muestra que $f = 0$. Puesto que $f \neq 0$, entonces existe al menos un vector $\alpha \in V$ tal que $f(\alpha, \alpha) = q(\alpha) \neq 0$. Sea W el subespacio unidimensional de V que es generado por α y sea $W^\perp$ el conjunto de todos los vectores $\beta \in V$ tales que $f(\alpha, \beta) = 0$. Afirmamos que $V = W \oplus W^\perp$. En efecto: los subespacios W y $W^\perp$ son independientes. Un vector típico de W es $c\alpha$, donde c es un escalar. Si $c\alpha \in W^\perp$ también, entonces $f(c\alpha, c\alpha) = c^2 f(\alpha, \alpha) = 0$. Pero $f(\alpha, \alpha) \neq 0$, luego $c = 0$. De tal manera, cada vector en V es suma de un vector en W y un vector en $W^\perp$. En efecto, sea γ cualquier vector en V , consideremos la siguiente cantidad:

$$\beta = \gamma - \frac{f(\gamma, \alpha)}{f(\alpha, \alpha)} \alpha$$

Entonces

$$f(\alpha, \beta) = f(\alpha, \gamma) - \frac{f(\gamma, \alpha)}{f(\alpha, \alpha)} f(\alpha, \alpha)$$

y como f es simétrica, $f(\alpha, \beta) = 0$. Por consiguiente, β está en el subespacio $W^\perp$. La expresión $\gamma = \beta + \frac{f(\gamma, \alpha)}{f(\alpha, \alpha)}\alpha$ muestra que $V = W \oplus W^\perp$.

La restricción de f a $W^\perp$ es, obviamente, una forma bilineal simétrica. Como $W^\perp$ tiene dimensión $(n-1)$, por inducción suponemos que $W^\perp$ tiene una base $\{\alpha_2, \dots, \alpha_n\}$ tal que

$$f(\alpha_i, \alpha_j) = 0, \quad i \neq j \quad (i, j \geq 2)$$

Haciendo $\alpha = \alpha_1$, se obtiene una base $B = \{\alpha_1, \dots, \alpha_n\}$ de V tal que $f(\alpha_i, \alpha_j) = 0$ para $i \neq j$. $\square$

Denotemos por g_d a la matriz diagonal que representa al tensor métrico. A esta matriz, cuya existencia hemos probado, la podemos escribir explícitamente como $g_d = \text{diag}(g_1, \dots, g_n)$; si a esta última la multiplicamos por la izquierda y por la derecha por una matriz diagonal $D = \text{diag}(d_1, d_2, \dots, d_n)$, el resultado es una matriz diagonal $g' = \text{diag}(g_1 d_1^2, \dots, g_n d_n^2)$. Sea $d_k = |g_k|^{-1/2}$, entonces cada elemento en la diagonal de g' es $+1$ o -1 . Es decir, siempre existe una base en la cual el tensor métrico está representado por una matriz diagonal cuyos elementos son $+1$ o -1 , a esta representación se le denomina *canónica* y a la base para lo cual ello ocurre se le llama *base canónica*. La traza de la forma canónica es la *signatura* de la métrica: $\text{Tr}(g') = \text{signatura de la métrica}$.

Si la métrica es positiva (negativa) definida entonces todos los elementos de la forma canónica son $+1$ (-1) y el espacio es Euclídeo. Si la métrica no es de signo definido, entonces se le suele llamar *indefinida*; un caso importante es la forma canónica $(1, -1, -1, \dots, -1)$, que corresponde a una métrica usualmente conocida como métrica de Minkowski (la relatividad especial tiene tal métrica para $n=4$).

Antes de proseguir es importante observar cual es la ecuación matricial correspondiente a un cambio de base: Nosotros sabemos que al llevar a cabo el cambio de base $\{\vec{e}_i\} \rightarrow \{\vec{e}'_i\}$ las componentes de la métrica se transforman de acuerdo a la regla $g'_{ij} = \Lambda_i^k \Lambda_j^l g_{kl}$; la última expresión puede reescribirse como $g'_{ij} = \Lambda_i^k g_{kl} \Lambda_j^l$. En términos de matrices, el primer índice en g denota filas y el segundo índice columnas, entonces en Λ , l denota filas y k columnas. De lo anterior se sigue entonces que la matriz con componentes Λ_i^k es la transpuesta de la matriz con componentes Λ_j^l . Por lo tanto, la ecuación matricial es simplemente $g' = \Lambda^T g \Lambda$.

En el espacio Euclídeo las bases canónicas son las cartesianas. En estas bases el tensor métrico tiene componentes $g_{ij} = \delta_{ij}$ o bien $g = I = \text{matriz identidad}$. Sea Λ_c la matriz de transformación de una base canónica a otra, entonces:

$$I = \Lambda_c^T I \Lambda_c \Rightarrow \Lambda_c^T = \Lambda_c^{-1}$$

Lo anterior implica que las matrices ortogonales son las matrices de transformación entre bases canónicas en el espacio euclídeo. El conjunto de las matrices ortogonales conforman un grupo, éste es conocido como el grupo de simetría euclidiano $O(n)$.

En relatividad especial el tensor métrico en una base canónica es el de Minkowski ($\eta = \text{diag}(1, -1, -1, \dots, -1)$), si Λ_L es la matriz de transformación de una base canónica a otra entonces tenemos que:

$$\eta = \Lambda_L^T \eta \Lambda_L$$

La matriz de transformación Λ_L es, justamente, la transformación de Lorentz. El conjunto de las transformaciones de Lorentz constituyen un grupo llamado grupo de Lorentz $L(n) \equiv SO(3, 1)$.

El tensor métrico nos permite mapear vectores en 1-formas, es decir $g : T_p M \rightarrow T_p^* M$. En efecto:

$$v_i = \tilde{v}(\vec{e}_i) = g(\vec{v}, \vec{e}_i) = g(v^j \vec{e}_j, \vec{e}_i) = v^j g_{ji} = g_{ij} v^j$$

Si existe el inverso de g , entonces podemos decir que en términos de componentes el inverso de g_{ij} es g^{ij} ; la relación que habrá de cumplirse es, evidentemente, la siguiente: $g^{ij} g_{jk} = \delta_k^i$. De lo anterior se deduce entonces que $g^{ki} v_i = g^{ki} g_{ij} v^j = \delta_j^k v^j = v^k$.

La existencia de el inverso garantiza que la relación entre vectores y 1-formas es 1-1. Puesto que g es un tensor $\binom{2}{2}$, entonces $g(\vec{v}, \cdot)$ es una 1-forma $\tilde{v}$, si existe el inverso entonces la 1-forma $g(\vec{v}, \cdot)$ es única y, por lo tanto, el mapeo entre vectores y 1-formas es 1-1. Una vez establecida la transformación de vectores a 1-formas (y viceversa), es posible operar sobre diferentes tensores; por ejemplo, llevemos a cabo el mapeo que va de $\binom{1}{1} \rightarrow \binom{2}{2}$: Sea A el tensor $\binom{2}{2}$ que es el resultado del producto directo entre $\tilde{v}$ y $\tilde{w}$, y sea B el tensor $\binom{1}{1}$ dado por el producto directo de $\tilde{v}$ y $\tilde{w}$, entonces tenemos que

$$A = \tilde{v} \otimes \tilde{w} = C_{ij} \tilde{v}^i \otimes \tilde{w}^j = v_i w_j \tilde{v}^i \otimes \tilde{w}^j \Rightarrow C_{ij} = v_i w_j$$

y

$$B = \tilde{v} \otimes \tilde{w} = C_i^k \tilde{v}^i \otimes \tilde{w}_k = v_i w^k \tilde{v}^i \otimes \tilde{w}_k \Rightarrow C_i^k = v_i w^k$$

Entonces:

$$C_{ij} = v_i w_j = v_i g_{jk} w^k = g_{jk} v_i w^k \Rightarrow C_{ij} = g_{jk} C_i^k$$

§2.10 El campo tensorial métrico sobre una variedad

Un campo tensorial métrico sobre una variedad, es un campo de tensores métricos cuya inversa existe $\forall p \in M$ (con lo cual garantizamos un mapeo biyectivo entre vectores y 1-formas). Observemos que si el campo es al menos C^0 en $U \subseteq M$, entonces la forma canónica es constante en $U \subseteq M$ puesto que tal forma está compuesta sólo por enteros y, evidentemente, estos no cambian de manera continua; lo anterior implica que si, por ejemplo, alrededor (localmente) de $p \in M$ la variedad es Minkowskiana, entonces también será Minkowskiana alrededor (localmente) de $p' \in M$, donde $p' \neq p$. Para el caso de un campo tensorial métrico continuo, tenemos entonces que es posible hablar de la signatura del mismo.

Una propiedad importante de un campo tensorial métrico radica en el hecho de que nos permite definir una longitud sobre una variedad. Sea $\{x^i\}$ un sistema coordenado para $U \subset M$ y consideremos una curva $\gamma(\lambda)$ en U ; la tangente a $\gamma(\lambda)$ es entonces $\vec{v} = \frac{dx}{d\lambda}$. Un desplazamiento $d\lambda$ tiene longitud cuadrada $ds^2 = d\vec{x} \cdot d\vec{x}$, puesto que $d\vec{x} = \vec{v}d\lambda$ entonces $ds^2 = \vec{v} \cdot \vec{v}d\lambda^2 = g(\vec{v}, \vec{v})d\lambda^2$.

Si la métrica es positiva definida, entonces $g(\vec{v}, \vec{v}) > 0 \forall \vec{v} \neq 0$, lo cual implica que $ds^2 > 0$ y por lo tanto $ds = g(\vec{v}, \vec{v})^{1/2}d\lambda$, que no es otra cosa que la longitud de un elemento de la curva. Sin embargo, si la métrica es indefinida, ds^2 no tiene necesariamente definido el signo. Las curvas se distinguen entonces por tener ds^2 positivo (tipo tiempo), negativo (tipo espacio) o igual a cero (curvas nulas).

Podemos definir la cantidad $ds = |g(\vec{v}, \vec{v})|^{1/2}d\lambda \in \mathbb{R}$ para ser la *distancia propia* o el *tiempo propio* en caso de que la curva sea tipo espacio o tipo tiempo.

Es importante hacer notar que cuando la métrica es indefinida, un vector puede tener norma cero y ser distinto al vector cero. Un ejemplo ilustrativo podría ser la métrica de Minkowski (η):

$$\vec{v} \cdot \vec{v} = \eta_{\alpha\beta} v^\alpha v^\beta = (v^0)^2 - \sum_{i=1}^n (v^i)^2$$

Donde puede ocurrir que $\vec{v} \cdot \vec{v} = 0$ sin que necesariamente $v^\alpha = 0, \forall \alpha = 1, 2, \dots, n$.

§2.11 Tensores antisimétricos.

Consideremos un tensor $T \binom{p}{p}$, diremos que T es antisimétrico si $T(\vec{u}, \vec{v}) = -T(\vec{v}, \vec{u})$ para cualesquiera $\vec{u}$ y $\vec{v}$. Un tensor $\binom{p}{p}$, con $p \geq 3$, es *completamente antisimétrico* si cambia de signo bajo el intercambio de cualesquiera dos de sus argumentos.

Para hallar la componente totalmente antisimétrica de un tensor $\binom{p}{p}$ se llevan a cabo todas las permutaciones en los índices, de tales permutaciones las pares son acompañadas por un signo positivo, en tanto que las impares por uno negativo:

$$(T_A)_{ij\dots k} = \sum_{\sigma=0}^{p!-1} \frac{(-1)^\sigma}{p!} P_\sigma(T_{ij\dots k}) \equiv T_{[ij\dots k]}$$

donde P_σ denota permutación par si sigma es un número par o impar si sigma es impar (P_0 indica no permutar); evidentemente, cada P_σ es una permutación distinta en los índices.

De tal manera tenemos que si T es un tensor $\binom{p}{p}$, entonces:

$$T_{[ijk]} = \sum_{\sigma=0}^{3!-1} \frac{(-1)^\sigma}{3!} P_\sigma(T_{ijk}) = \frac{1}{p!} (P_0(T_{ijk}) - P_1(T_{ijk}) + P_2(T_{ijk}) - P_3(T_{ijk}) + P_4(T_{ijk}) - P_5(T_{ijk}))$$

$$T_{[ijk]} = \frac{1}{p!} (T_{ijk} - T_{jik} + T_{jki} - T_{kji} + T_{kij} - T_{ikj})$$

Notemos que sobre un espacio vectorial n -dimensional, un tensor completamente antisimétrico $\binom{p}{p}$, con $n \geq p$, tiene a lo más $C_n^p = \frac{n!}{p!(n-p)!}$ componentes independientes. Efectivamente: Cualquier componente queda definida al elegir p *diferentes* números del conjunto $\{1, \dots, n\}$; los p números deben ser distintos dado que las componentes de un tensor completamente antisimétrico son nulas si al menos dos de sus índices son iguales. El orden de los p números no es relevante ya que cambian exclusivamente el signo de la componente, en consecuencia el número de *componentes independientes* es igual a el número de conjuntos diferentes de p números (sin importar el orden) elegidos de n posibles, es decir, el coeficiente binomial.

§2.12 Formas diferenciales.

Definimos a una p -forma ($p \geq 2$) como un tensor completamente antisimétrico del tipo $\binom{p}{p}$. Los casos degenerados de tensores completamente antisimétricos son las 1-formas y las funciones escalares (0-formas). Diremos que el número p es el grado de la forma.

Dados dos tensores $(\overset{1}{T})$, mediante la operación $\otimes$ construimos un tensor $(\overset{2}{T})$ que, en general, no tiene porque ser antisimétrico. Definiremos entonces la operación $\wedge$ (*producto cuña*) entre los tensores $(\overset{1}{T})$ de manera tal que de por resultado un tensor $(\overset{2}{T})$ completamente antisimétrico:

$$\bar{p} \wedge \bar{q} \equiv \bar{p} \otimes \bar{q} - \bar{q} \otimes \bar{p}$$

donde $\bar{p}$ y $\bar{q}$ son 1-formas.

Nótese que efectivamente es un tensor $(\overset{2}{T})$ completamente antisimétrico: $\bar{p} \wedge \bar{q}(\bar{u}, \bar{v}) = \bar{p} \otimes \bar{q}(\bar{u}, \bar{v}) - \bar{q} \otimes \bar{p}(\bar{u}, \bar{v}) \Rightarrow \bar{p} \wedge \bar{q}(\bar{u}, \bar{v}) = \bar{p}(\bar{u})\bar{q}(\bar{v}) - \bar{q}(\bar{u})\bar{p}(\bar{v}) \Rightarrow \bar{p} \wedge \bar{q}(\bar{u}, \bar{v}) = -\{\bar{p}(\bar{v})\bar{q}(\bar{u}) - \bar{q}(\bar{v})\bar{p}(\bar{u})\} \Rightarrow \bar{p} \wedge \bar{q}(\bar{u}, \bar{v}) = -\bar{p} \wedge \bar{q}(\bar{v}, \bar{u})$.

El producto cuña es asociativo y anticonmutativo. Es posible extenderlo de manera natural para obtener una n-forma a partir de n 1-formas:

$$\bar{p}_1 \wedge \bar{p}_2 \wedge \dots \wedge \bar{p}_n = \sum_{\sigma=0}^{n!-1} (-1)^\sigma P_\sigma(\bar{p}_1 \otimes \bar{p}_2 \otimes \dots \otimes \bar{p}_n)$$

Por ejemplo, el producto cuña de tres 1-formas es:

$$\bar{p} \wedge \bar{q} \wedge \bar{r} \equiv \bar{p} \otimes \bar{q} \otimes \bar{r} + \bar{q} \otimes \bar{r} \otimes \bar{p} + \bar{r} \otimes \bar{p} \otimes \bar{q} - \bar{p} \otimes \bar{r} \otimes \bar{q} - \bar{q} \otimes \bar{p} \otimes \bar{r} - \bar{r} \otimes \bar{q} \otimes \bar{p}$$

Observemos que el conjunto de tensores completamente antisimétricos $\{\bar{w}^j \wedge \bar{w}^i; j, k = 1, \dots, n = \dim T_p M\}$ es una base para el espacio vectorial de todas las 2-formas. En efecto: sea $\bar{\alpha}$ una 2-forma cualquiera, entonces

$$\bar{\alpha}(\bar{v}, \bar{u}) = v^i u^j \bar{\alpha}(\bar{e}_i, \bar{e}_j) = \alpha_{ij} v^i u^j$$

puesto que $v^i u^j = \bar{w}^i(\bar{v})\bar{w}^j(\bar{u})$ y $\bar{\alpha}$ es por definición totalmente antisimétrica, obtenemos lo siguiente:

$$\bar{\alpha}(\bar{v}, \bar{u}) = \frac{1}{2!}(\alpha_{ij} - \alpha_{ji})\bar{w}^i(\bar{v})\bar{w}^j(\bar{u})$$

$$\bar{\alpha}(\bar{v}, \bar{u}) = \frac{1}{2!}(\alpha_{ij}\bar{w}^i(\bar{v})\bar{w}^j(\bar{u}) - \alpha_{ji}\bar{w}^j(\bar{v})\bar{w}^i(\bar{u}))$$

dado que los índices i y j son mudos, podemos intercambiarlos en el segundo sumando ($j \rightarrow i$ y $i \rightarrow j$):

$$\bar{\alpha}(\bar{v}, \bar{u}) = \frac{1}{2!}\alpha_{ij}(\bar{w}^i(\bar{v})\bar{w}^j(\bar{u}) - \bar{w}^j(\bar{v})\bar{w}^i(\bar{u}))$$

$$\bar{\alpha}(\bar{v}, \bar{u}) = \frac{1}{2!}\alpha_{ij}\bar{w}^i \wedge \bar{w}^j(\bar{v}, \bar{u})$$

lo cual es válido para cualesquiera $\vec{v}$ y $\vec{w}$; en consecuencia:

$$\tilde{\alpha} = \frac{1}{2!} \alpha_{ij} \tilde{w}^i \wedge \tilde{w}^j$$

En general, tendremos que la expresión para una h-forma es:

$$\tilde{h} = \frac{1}{h!} h_{ij\dots r} \tilde{w}^i \wedge \tilde{w}^j \wedge \dots \wedge \tilde{w}^r$$

donde el conjunto $\{i, j, \dots, r\}$ es un conjunto de h índices.

Sean $\tilde{q}$ una q-forma y $\tilde{p}$ una p-forma ($q + p \leq \dim(M)$), entonces el producto cuña de ambas es una $(p+q)$ -forma dada por:

$$(\tilde{p} \wedge \tilde{q}) = \frac{1}{(p+q)!} (\tilde{p} \wedge \tilde{q})_{is\dots rjk\dots m} \tilde{w}^i \wedge \dots \wedge \tilde{w}^r \wedge \tilde{w}^j \wedge \dots \wedge \tilde{w}^m$$

por otro lado, tenemos que:

$$(\tilde{p} \wedge \tilde{q}) = \left(\frac{1}{p!} p_{is\dots r} \tilde{w}^i \wedge \dots \wedge \tilde{w}^r \right) \left(\frac{1}{q!} q_{jk\dots m} \tilde{w}^j \wedge \dots \wedge \tilde{w}^m \right)$$

puesto que el producto cuña es asociativo tenemos entonces que:

$$(\tilde{p} \wedge \tilde{q}) = \frac{1}{p!q!} p_{is\dots r} q_{jk\dots m} \tilde{w}^i \wedge \dots \wedge \tilde{w}^r \wedge \tilde{w}^j \wedge \dots \wedge \tilde{w}^m$$

del hecho de que $(\tilde{p} \wedge \tilde{q})$ sea un tensor completamente antisimétrico se sigue que $p_{is\dots r} q_{jk\dots m} = p_{[is\dots r} q_{jk\dots m]}$; como ambas expresiones de la $(p+q)$ -forma tienen la misma base, entonces sus componentes habrán de ser idénticas:

$$\begin{aligned} \frac{1}{p!q!} p_{[is\dots r} q_{jk\dots m]} &= \frac{1}{(p+q)!} (\tilde{p} \wedge \tilde{q})_{is\dots rjk\dots m} \\ \Rightarrow (\tilde{p} \wedge \tilde{q})_{is\dots rjk\dots m} &= \frac{(p+q)!}{p!q!} p_{[is\dots r} q_{jk\dots m]} \\ \Rightarrow (\tilde{p} \wedge \tilde{q})_{is\dots rjk\dots m} &= C_p^{p+q} p_{[is\dots r} q_{jk\dots m]} \end{aligned}$$

Un álgebra sencilla nos proporciona la regla de conmutación entre una p-forma y una q-forma: $\tilde{p} \wedge \tilde{q} = (-1)^{pq} \tilde{q} \wedge \tilde{p}$.

Sea $\tilde{\varphi}$ una p -forma, entonces $\tilde{\varphi}(\tilde{\xi}) \equiv \tilde{\varphi}(\tilde{\xi}^1, \dots, \tilde{\xi}^p)$ es una $(p-1)$ -forma que se obtiene al contraer $\tilde{\varphi}$ con el vector $\tilde{\xi}$. Puesto que $\tilde{\varphi}(\tilde{\xi}) = \{\frac{1}{p!} \varphi_{i_1 \dots i_p} \tilde{w}^{i_1} \wedge \dots \wedge \tilde{w}^{i_p}\}(\tilde{\xi})$ y la contracción de las p 1-formas $\tilde{w}^i$ con el vector $\tilde{\xi}$ es el tensor completamente antisimétrico $p \xi^{i_1} \tilde{w}^{j_1} \wedge \dots \wedge \tilde{w}^{j_p}$, se sigue entonces que:

$$\tilde{\varphi}(\tilde{\xi}) = \frac{1}{p!} \varphi_{i_1 \dots i_p} p \xi^{i_1} \tilde{w}^{j_1} \wedge \dots \wedge \tilde{w}^{j_p}$$

de lo cual se sigue que:

$$\tilde{\varphi}(\tilde{\xi}) = \frac{1}{(p-1)!} \xi^{i_1} \varphi_{i_1 \dots i_{p-1}} \tilde{w}^{j_1} \wedge \dots \wedge \tilde{w}^{j_{p-1}}$$

Si $\tilde{\alpha}$ es cualquier forma y $\tilde{\beta}$ es una p -forma, es sencillo mostrar (mediante un álgebra elemental pero tediosa) que la contracción del producto cuña de las formas con el vector $\tilde{\xi}$ está dada por:

$$(\tilde{\beta} \wedge \tilde{\alpha})(\tilde{\xi}) = \tilde{\beta}(\tilde{\xi}) \wedge \tilde{\alpha} + (-1)^p \tilde{\beta} \wedge \tilde{\alpha}(\tilde{\xi})$$

§2.13 Integración sobre variedades.

Definiremos el concepto de integral exclusivamente sobre variedades orientadas (§A.1).

Sea M una variedad n -dimensional orientada, $\{(U_r, \alpha_r)\}$ un atlas orientado tal que $\{U_r\}$ es una cubierta abierta localmente finita de M , y $\{h_r\}$ una partición de unidad (§A.2) subordinada a esta cubierta.

En primera instancia consideremos al dominio de integración de manera tal que $K \subset U_r$; además, pediremos que K sea suficientemente regular (por suficientemente regular habremos de entender que las integrales ordinarias están bien definidas). En términos de coordenadas orientadas $\{x^i\}$ sobre U_r un campo de n -formas puede escribirse como: $\tilde{\phi} = f dx^1 \wedge \dots \wedge dx^n$, donde f es una función suave sobre U_r .

Notemos que $\tilde{\phi}(dx^1 \frac{\partial}{\partial x^1}, \dots, dx^n \frac{\partial}{\partial x^n}) = f(\tilde{x}) dx^1 \dots dx^n$, entonces la integral de la forma sobre una región $Q \subset M^n$ no es mas que la integral usual de cálculo vectorial:

$$\int_Q \tilde{\phi} \equiv \int_Q \tilde{\phi}(dx^1 \frac{\partial}{\partial x^1}, \dots, dx^n \frac{\partial}{\partial x^n}) = \int \dots \int_Q f(\tilde{x}) dx^1 \dots dx^n$$

Puesto que $\{x^i\}$ puede ser interpretado como un sistema coordenado cartesiano, entonces consideraremos a éste como el sistema orientado de coordenadas sobre el conjunto $\alpha_r(K)$.

Dado que $\alpha_r(K)$ es orientado, entonces es posible llevar a cabo la integración sobre tal conjunto.

Consideremos el mapeo lineal $(\alpha_r^{-1})^*$ que va del espacio vectorial de las n -formas en $K \subset M$ al espacio vectorial de las n -formas en $\alpha_r(K) \subset \mathbb{R}^n$; entonces la forma $(\alpha_r^{-1})^*\tilde{\phi}$ puede describirse en un sistema coordenado orientado y en consecuencia $\int_{\alpha_r(K)} (\alpha_r^{-1})^*\tilde{\phi} = \int \dots \int_{\alpha_r(K)} f(\vec{x}) dx^1 \dots dx^n$. Por otro lado, la transformación $\alpha_r(K)$ nos proporciona la siguiente igualdad: $\int_K \tilde{\phi} = \int_{\alpha_r(K)} (\alpha_r^{-1})^*\tilde{\phi}$. Por lo tanto:

$$\int_K \tilde{\phi} = \int \dots \int_{\alpha_r(K)} f(\vec{x}) dx^1 \dots dx^n$$

Si la región de integración K no está contenida en un solo U_r , entonces habrá que recurrir a la partición de unidad para expresar a $\int_K \tilde{\phi}$ en términos de integrales ya definidas:

$$\int_K \tilde{\phi} = \sum_r \int_{K \cap U_r} h_r \tilde{\phi}$$

La partición de unidad se debe introducir puesto que en general $(K \cap U_r) \cap (K \cap U_s) \neq \emptyset$ (con $r \neq s$), si no empleamos la partición de unidad la integración da como resultado un "exceso" en las intersecciones.

§2.14 n-vectores y duales.

Por analogía con las n -formas, diremos que un tensor $\binom{0}{n}$ completamente antisimétrico será un n -vector. Obviamente, la dimensión del espacio vectorial de los p -vectores en cualquier punto de una variedad n -dimensional es C_n^p .

Notemos entonces que en un $p \in M$ hay cuatro espacios cuya dimensión es la misma, a saber: el espacio vectorial de las p -formas, el de las $(n-p)$ -formas, el de p -vectores y el constituido por los $(n-p)$ -vectores.

Sea $\tilde{w}$ una n -forma nunca nula definida en una región orientada de una variedad. La n -forma $\tilde{w}$ proporciona un mapeo (*mapeo dual*) entre p -formas y $(n-p)$ -vectores como sigue:

Dado un q -vector T , con componentes $T^{i_1 \dots i_q} = T^{[i_1 \dots i_q]}$, definimos un tensor $\tilde{A}$ mediante la expresión:

$$A_{j_1 \dots j_q} = \frac{1}{q!} w_{i_1 \dots i_q j_1 \dots j_q} T^{i_1 \dots i_q}$$

¹Para un análisis más detallado puede consultarse [6] por ejemplo.

simbólicamente escribimos $\tilde{A} = \tilde{\omega}(T)$ o simplemente $\tilde{A} = {}^*T$.

Diremos que $\tilde{A}$ es el dual de T respecto a $\tilde{\omega}$. Puesto que $A_{j...l}$ son las componentes de un tensor completamente antisimétrico de grado $(n-q)$, entonces $\tilde{A}$ es una $(n-q)$ -forma. El mapeo dual define una única $(n-q)$ -forma para cada q -vector.

El mapeo $T \rightarrow {}^*T$ es biyectivo dado que es 1-1 y sobre, en consecuencia es invertible (i.e: $\exists {}^*T \rightarrow T$).

Sean $w^{i...k}$ las componentes (de un n -vector) inversas a las de la n -forma $\tilde{\omega}$: $w^{i...k}\omega_{i...k} = n!$. Diremos que H es el dual de $\tilde{G}$ con respecto a $\tilde{\omega}$ ($H = {}^*\tilde{G}$) si (con $\tilde{G}$ una p -forma):

$$H^{i...k} = \frac{1}{p!} w^{i...mi...k} G_{i...m}$$

Definamos al $(n-p)$ -vector H como antes. Hallemos el dual de H :

$$\begin{aligned} ({}^*H)_{j...l} &= \frac{1}{(n-p)!} w_{i...kj...l} H^{i...k} \\ &= \frac{1}{(n-p)!p!} w_{i...kj...l} w^{r...si...k} G_{r...s} \\ &= \frac{(-1)^{p(n-p)}}{(n-p)!p!} w_{i...kj...l} w^{i...kr...s} G_{r...s} \end{aligned}$$

Fijemos los índices $(j...l)$ en, digamos, $(1...p)$. La suma para $(r...s)$ fijos es entonces: $w_{i...k1...p} w^{i...kr...s}$, lo cual implica que los índices $(i...k)$ deben elegirse del conjunto $\{p+1, \dots, n\}$. De tal manera, a lo mas, hay $(n-p)!$ términos no nulos en la suma. Entonces:

$$w_{i...k1...p} w^{i...kr...s} = (n-p)! w_{p+1...n1...p} w^{p+1...nr...s}$$

Observemos que la última expresión es cero a menos que $(r...s)$ sean una permutación de $(1...p)$. En la suma sobre $(r...s)$, $w^{p+1...nr...s} G_{r...s}$, hay a lo mas $p!$ términos no nulos. Por lo tanto:

$$w^{p+1...nr...s} G_{r...s} = p! w^{p+1...n1...p} G_{1...p}$$

en consecuencia:

$$({}^*H)_{1...p} = (-1)^{p(n-p)} G_{1...p}$$

dado que $w_{p+1...n1...p} w^{p+1...n1...p} = 1$.

Como $H = {}^*\tilde{G} \Rightarrow ({}^{**}\tilde{G})_{1\dots p} = (-1)^{p(n-p)}G_{1\dots p}$. Puesto que las etiquetas $(1\dots p)$ pueden ser para cualquier índice, entonces:

$$({}^{**}\tilde{G}) = (-1)^{p(n-p)}\tilde{G}$$

De manera similar se prueba para un q -vector T que:

$${}^{**}T = (-1)^{q(n-q)}T$$

En particular, tenemos que para una función f (0-forma) habrá de cumplirse que: ${}^{**}f = f$.

Puesto que el tensor métrico nos proporciona un mapeo 1-1 entre tensores $(\tilde{G})$ y tensores (G) , al componer este mapeo con el dual obtendremos un mapeo que va del espacio vectorial (EV) de las p -formas (p -vectores) al EV de las $(n-p)$ -formas ($(n-p)$ -vectores) [4]:

$$\text{EV } p\text{-formas} \xrightarrow{\text{dual}} \text{EV } (n-p)\text{-vectores} \xrightarrow{g} \text{EV } (n-p)\text{-formas}$$

$$* \equiv g \circ \text{dual} : \text{EV } p\text{-formas} \rightarrow \text{EV } (n-p)\text{-formas}$$

§2.15 Derivada exterior.

Con anterioridad mencionamos que la derivada direccional de una función f a lo largo de una curva con vector tangente $\frac{d}{dx}$ es: $\tilde{d}f(\frac{d}{dx})$. Del hecho de que la derivada direccional es un elemento en R , se sigue que $\tilde{d}f$ es una 1-forma.

Extendiendo el resultado, tendríamos que si $\tilde{\alpha}$ es una p -forma, entonces $\tilde{d}\tilde{\alpha}$ habrá de ser una $(p+1)$ -forma. La extensión citada se lleva a cabo de manera adecuada en tanto el operador $\tilde{d}$ cumpla con los siguientes requerimientos para las formas arbitrarias $\tilde{\alpha}$ (p -forma), $\tilde{\beta}$ y $\tilde{\gamma}$ (q -formas):

- (i).- $\tilde{d}(\tilde{\beta} + \tilde{\gamma}) = \tilde{d}\tilde{\beta} + \tilde{d}\tilde{\gamma}$.
- (ii).- $\tilde{d}(\tilde{\alpha} \wedge \tilde{\beta}) = \tilde{d}(\tilde{\alpha}) \wedge \tilde{\beta} + (-1)^p \tilde{\alpha} \wedge (\tilde{d}\tilde{\beta})$.
- (iii).- $\tilde{d}(\tilde{d}\tilde{\alpha}) = 0$.

A la aplicación $\tilde{d}\tilde{\alpha}$ se le conoce como *derivada exterior* de la p -forma $\tilde{\alpha}$.

En un sistema coordenado $\{x^i\}$ tendríamos que:

$$\tilde{d}\tilde{\alpha} = \frac{1}{p!} d\alpha_{i_1\dots i_p} \wedge \tilde{d}x^{i_1} \wedge \dots \wedge \tilde{d}x^{i_p}$$

$$\tilde{d}\alpha_{i\dots j} = \frac{\partial \alpha_{i\dots j}}{\partial x^k} \tilde{d}x^k \Rightarrow \tilde{d}\tilde{\alpha} = \frac{1}{p!} \frac{\partial \alpha_{i\dots j}}{\partial x^k} \tilde{d}x^k \wedge \tilde{d}x^i \wedge \dots \wedge \tilde{d}x^j$$

de lo cual se sigue que:

$$(\tilde{d}\tilde{\alpha})_{k i \dots j} = (p+1) \frac{\partial}{\partial x^k} \alpha_{i \dots j}$$

Con objeto de ilustrar el uso que tienen las formas diferenciales, en el ámbito de las ecuaciones diferenciales, consideremos el siguiente sistema de ecuaciones diferenciales parciales:

$$f_{,11} = g(x^1, x^2) \quad f_{,12} = h(x^1, x^2)$$

en donde se introdujo la notación $\frac{\partial f}{\partial x^i} \equiv f_{,i}$.

El sistema anterior puede escribirse de manera más compacta:

$$f_{,i} = a_i \quad \text{donde} \quad a_1 = g \quad \text{y} \quad a_2 = h$$

Puesto que en el sistema coordenado $\{x^i\}$ $\tilde{d}f = f_{,i} \tilde{d}x^i$ y las funciones g y h bien pueden ser las componentes de una 1-forma $\tilde{a}$ en dicho sistema, dado que $f_{,i} = a_i$ entonces habrá de cumplirse que $\tilde{d}f = \tilde{a}$ (notese que esta última expresión es independiente de las coordenadas).

Como $\tilde{d}f = \tilde{a} \Rightarrow \tilde{d}\tilde{d}f = \tilde{d}\tilde{a}$, pero $\tilde{d}\tilde{d}f = 0 \Rightarrow \tilde{d}\tilde{a} = 0$. De la igualdad $(\tilde{d}\tilde{a})_{ij} = 2a_{[i,j]}$ se sigue que $a_{[i,j]} = 0$ o bien que $a_{[i,j]} = 0$. Esta es una sola ecuación puesto, que $\tilde{d}\tilde{a}$ es una 2-forma en una variedad 2-dimensional y en consecuencia la componente $a_{[i,j]}$ es la única:

$$a_{[i,j]} = a_{1,2} - a_{2,1} = 0 \Rightarrow f_{,12} - f_{,21} = 0$$

es decir:

$$\frac{\partial f}{\partial y \partial x} = \frac{\partial f}{\partial x \partial y}$$

que es justamente la condición de integrabilidad del sistema inicial.

§2.16 Formas cerradas y exactas.

Diremos que una forma $\tilde{\alpha}$ es *cerrada* si satisface que $\tilde{d}\tilde{\alpha} = 0$; si además $\tilde{\alpha}$ es tal que $\tilde{\alpha} = \tilde{d}\tilde{\beta}$ entonces diremos que la forma es también *exacta*. Obsérvese que toda forma exacta es

²Esta notación se extiende a derivadas y tensores de mayor orden, por ejemplo: $\frac{\partial^2}{\partial x^i \partial x^j} T^k = T^k_{,ij}$.

cerrada; sin embargo, el recíproco no es cierto en general, se requieren mayores "restricciones" para su validez:

Consideremos una vecindad D de un punto p en la cual $\tilde{\alpha}$ está definida y es cerrada. Entonces existe una vecindad "suficientemente pequeña" alrededor de p en la cual una forma $\tilde{\beta}$ está siempre definida y cumple con la relación $\tilde{\alpha} = d\tilde{\beta}$. Es claro que $\tilde{\beta}$ no es única, también podríamos tener $\tilde{\beta} + d\tilde{\gamma}$ para cualquier $\tilde{\gamma}$ de grado correcto. De tal manera, la afirmación consiste en decir que al menos localmente una forma cerrada es exacta: dada una $D \subset M$ arbitraria en la cual $\tilde{\alpha}$ está definida y es cerrada, puede no existir una sola $\tilde{\beta}$ definida en D tal que $\tilde{\alpha} = d\tilde{\beta}$. Precisemos la afirmación en cuanto a la exactitud local de las formas cerradas se refiere:

Lema 2.1 (Poincaré) Sea $\tilde{\alpha}$ una p -forma cerrada definida en $U \subset M$ difeomorfo a la bola unitaria abierta en $\mathbb{R}^n$. Entonces existe al menos una $(p-1)$ -forma $\tilde{\beta}$ definida en U tal que $\tilde{\alpha} = d\tilde{\beta}$.

Demostración: La demostración consiste en exhibir a la forma $\tilde{\beta}$.

Puesto que existe un difeomorfismo entre U y la bola abierta unitaria en $\mathbb{R}^n$, entonces U está cubierto por un solo sistema coordenado cartesiano. Sean $\{x^i\}$ las coordenadas para U . En tal sistema, la p -forma $\tilde{\alpha}$ puede escribirse como sigue:

$$\tilde{\alpha} = \alpha_{i_1 \dots i_p}(x^1, \dots, x^n) dx^{i_1} \wedge \dots \wedge dx^{i_p}$$

Sea $\tilde{\mu} = \tilde{\alpha}(\vec{r})$, donde $\vec{r} = (x^1, \dots, x^n)$. La forma $\tilde{\mu}$ es una $(p-1)$ -forma cuya expresión en el sistema coordenado es simplemente:

$$\tilde{\mu} = \alpha_{ij \dots k}(x^1, \dots, x^n) x^i dx^j \wedge \dots \wedge dx^k$$

Definamos las siguientes funciones:

$$\beta_{j \dots k}(x^1, \dots, x^n) = \int_0^1 t^{p-1} \alpha_{ij \dots k}(tx^1, \dots, tx^n) x^i dt$$

Entonces:

$$\beta_{j \dots k, i} = \int_0^1 t^{p-1} \alpha_{ij \dots k}(tx^1, \dots, tx^n) dt + \int_0^1 t^p x^i \alpha_{ij \dots k, i}(tx^1, \dots, tx^n) dt$$

²La existencia de tal difeomorfismo no es otra cosa que imponer ciertas condiciones topológicas sobre U .

Puesto que $p\alpha_{[ij\dots k]i} - \alpha_{[ij\dots k]i} = 0$ dado que $\tilde{\alpha}$ es una p -forma cerrada, entonces:

$$\begin{aligned}(\tilde{d}\tilde{\beta})_{ij\dots k} &= p\beta_{[ij\dots k]i} = p \int_0^1 (t^{p-1}\alpha_{[ij\dots k]i} + t^p x^i \alpha_{i[j\dots k]i}) dt \\(\tilde{d}\tilde{\beta})_{ij\dots k} &= \int_0^1 (pt^{p-1}\alpha_{ij\dots k} + x^i t^p \alpha_{ij\dots k,i}) dt = \int_0^1 \frac{d}{dt} \{t^p \alpha_{ij\dots k}(tx^1, \dots, tx^n)\} dt \\&\Rightarrow (\tilde{d}\tilde{\beta})_{ij\dots k} = t^p \alpha_{ij\dots k}(tx^1, \dots, tx^n)|_0^1 = \alpha_{ij\dots k}(x^1, \dots, x^n) \quad \square\end{aligned}$$

§2.17 Teoría de cohomologías.

Consideremos ahora el conjunto de todas las p -formas cerradas y exactas sobre una variedad M : $C^p(M) \equiv \{\tilde{\alpha} \mid d\tilde{\alpha} = 0\}$ y $E^p(M) \equiv \{\tilde{\alpha} \mid \tilde{\alpha} = d\tilde{\beta}\}$. Evidentemente tanto $C^p(M)$ como $E^p(M)$ son espacios vectoriales sobre el campo de los números reales, adicionalmente notemos que si $\tilde{\beta} \in E^p(M)$ entonces $\tilde{\beta} \in C^p(M)$ (i.e. toda forma exacta es cerrada) y, en consecuencia, $E^p(M)$ es un subespacio de $C^p(M)$. Definamos al *Espacio Vectorial Cohomológico p -ésimo de de Rham* de M como el espacio cociente $H^p(M) \equiv C^p(M)/E^p(M)$, con la relación de equivalencia siguiente:

$$\tilde{\alpha}, \tilde{\beta} \in C^p(M) \text{ son equivalentes si } \tilde{\alpha} - \tilde{\beta} = \tilde{\gamma}, \text{ donde } \tilde{\gamma} \in E^p(M)$$

Sea $A \subset M$ una región difeomorfa a la bola unitaria abierta en $\mathbb{R}^n$ entonces, por el Lema de Poincaré, toda p -forma $\tilde{\alpha}$ ($p \geq 1$) cerrada definida en A es también exacta. Por lo tanto, todas las p -formas cerradas en A son equivalentes y, en consecuencia, obtenemos que: $H^p(A) = 0$ ($p \geq 1$), pues todas las p -formas cerradas son equivalentes a la p -forma nula. Para el caso $p = 0$ nótese que $H^0(A) = \mathbb{R}$; en efecto:

Una 0-forma es una función, entonces $C^0(A) = \{f \mid df = 0\} = \{f \mid f_i = 0, \forall i = 1, \dots, n\}$; puesto que A es conexo tenemos que C^0 es simplemente el conjunto de las funciones constantes: $C^0 = \{f : M \rightarrow \mathbb{R} \mid f \text{ es constante}\} = \{K \mid K \in \mathbb{R}\} = \mathbb{R}$. Ahora bien, $f \sim g$ si $f - g \in E^0(A)$, donde $E^0(A)$ es la cero-función; por consiguiente $f \sim g$ si $f = g$. Es decir, $H^0(A) = C^0(A) = \mathbb{R}$. Claramente este resultado se extiende a cualquier variedad M que sea conexa: $H^0(M) = \mathbb{R}$ si M es conexa. Por otra parte, si M no es conexa pero consta de k componentes conexas entonces $H^0(M) = C^0(M) = \mathbb{R}^k$.

Obsérvese como caso particular que $H^0(S^n) = \mathbb{R}$, pues S^n es una variedad conexa. Notemos que $H^1(S^1) = \mathbb{R}$ y que $H^1(S^2) = 0$; en efecto:

Sea la 1-forma $\tilde{w} = x\tilde{d}y - y\tilde{d}x = \tilde{d}\theta$ sobre S^1 , donde θ es el ángulo polar. Cualquier 1-forma $\tilde{\alpha}$ sobre S^1 puede escribirse entonces como $\tilde{\alpha} = h(\theta)\tilde{d}\theta$. Si $\lambda\tilde{w} \in [\tilde{\alpha}] \Rightarrow \tilde{\alpha} - \lambda\tilde{w} = \tilde{d}f$, donde f es continua, lo cual implica que:

$$\int_{S^1} \tilde{\alpha} = \lambda \int_{S^1} \tilde{w} = 2\pi\lambda \text{ p.a. } \lambda$$

lo cual siempre ocurre pues $h(\theta)$ es continua; en consecuencia:

$$\lambda\tilde{w} \sim \tilde{\alpha}, \text{ con } \lambda = \int_{S^1} \tilde{\alpha} / \int_{S^1} \tilde{w}$$

Entonces $[\tilde{\alpha}] = \lambda$, para alguna $\lambda \in \mathbb{R}$. Por lo tanto $H^1(S^1) = \{[\tilde{\alpha}]\} = \{\lambda \mid \lambda \in \mathbb{R}\} = \mathbb{R}$. La generalización al caso n -dimensional puede llevarse a cabo notando que:

$$\int_{S^n} \tilde{\alpha} = \lambda \int_{S^n} \tilde{w}, \text{ p.a. } \lambda \in \mathbb{R}$$

donde $\tilde{w} = \epsilon_{ij\dots k} x^i \tilde{d}x^j \wedge \dots \wedge \tilde{d}x^k$ es una n -forma definida sobre $\mathbb{R}^{n+1}$. Por lo tanto, $[\tilde{\alpha}] = \lambda \Rightarrow H^n(S^n) = \mathbb{R}$.

Por otra parte, si $\tilde{\alpha}$ es una 1-forma cerrada definida sobre S^2 , entonces integrando $\tilde{\alpha}\tilde{\alpha}$ en cualquier región de S^2 acotada por una curva cerrada y simple C obtenemos, según la versión para formas diferenciales del Teorema de Stokes⁴, que:

$$\oint_C \tilde{\alpha} = 0 \text{ para cualquier } C$$

Pero esto sólo ocurre si $\tilde{\alpha} = \tilde{d}g$ para alguna g , pues de no ser así podría hallarse una curva C para la cual la última integral es no nula. De tal manera, tenemos que toda 1-forma cerrada sobre S^2 es exacta y, por lo tanto, $H^1(S^2) = 0$. La generalización al caso n -dimensional es directa y tiene como consecuencia que $H^{n-1}(S^n) = 0$. De facto, esta generalización es todavía un caso especial del siguiente resultado: $H^p(S^n) = 0, 0 < p < n$. Este resultado puede probarse notando que toda p -forma cerrada sobre S^n es cerrada y exacta sobre S^{p+1} , que, a su vez, es una subvariedad de S^{p+2} ; entonces intuitivamente es claro que las p -formas también serán exactas en S^{p+2} y así sucesivamente hasta S^n .

⁴El teorema de Stokes para formas diferenciales estipula, a grandes rasgos, que la integral de la derivada exterior de cualquier $(n-1)$ -forma sobre un conjunto U n -dimensional con frontera ∂U es igual a la integral de la $(n-1)$ -forma sobre dicha frontera. En breve: $\int_U \tilde{d}\tilde{F} = \int_{\partial U} \tilde{F}$ [4].

En resumen, los Espacios Cohomológicos de de Raham para S^n son: $H^n(S^n) = \mathbb{R}$, $H^p(S^n) = 0$ ($0 < p < n$) y $H^0(S^n) = \mathbb{R}$. Estos espacios nos dan la relación que existe entre las formas cerradas y exactas sobre S^n para distintos ordenes. Pero incluso nos dicen más, por ejemplo $H^0(S^n) = \mathbb{R}$ implica que *cualquier* función $f : S^n \rightarrow \mathbb{R}$ tal que $df = 0$ es constante en todo S^n , pero ello implica que entonces S^n habrá de ser conexo; es decir, obtenemos información topológica de la variedad. Para ilustrar la importancia topológica que pueden tener los espacios cohomológicos, baste con mencionar el hecho de que una variedad M es simplemente conexa si $H^1(M) = 0$ [4].

El ejemplo anterior exhibe una dependencia del espacio cohomológico $H^1(M)$ de la estructura topológica de la variedad M , este hecho es de facto consecuencia del siguiente resultado general en teoría de cohomologías (Teorema de de Raham): Los espacios cohomológicos dependen solamente (aunque nuestra definición de $H^p(M)$ está basada en la estructura diferencial de la variedad) de la estructura topológica de M y no de su diferenciabilidad [4].

Capítulo 3

Derivada, grupos y álgebras de Lie.

Iniciaremos el presente capítulo definiendo con precisión el concepto de *derivada de Lie* para campos escalares y vectoriales sobre una variedad sin métrica y con curvatura no nula; para tal efecto, aprovecharemos algunos conocimientos adquiridos previamente. Adicionalmente, notaremos que la derivada de Lie nos ayuda a definir conceptos tales como la invariancia de un campo tensorial respecto a uno vectorial y, en caso de que la variedad este dotada con una métrica, los vectores de Killing.

Una vez expuesto lo anterior la discusión se centrará en torno a los grupos de Lie. Es importante que el lector no familiarizado con estos grupos procure especial atención a la definición de los mismos, a sus álgebras y a sus representaciones ya que los grupos de Lie, como veremos más adelante, tienen un importante papel en la construcción de un modelo σ -no lineal.

La presentación no soslaya lo fundamental y contempla lo necesario para nuestros propósitos ulteriores aunque, cabe advertir, los grupos de Lie conforman una extensa materia de estudio y el trabajo, consecuentemente, no esta exento de ciertas omisiones.

§3.1 Arrastre de Lie.

Sea $V \in C^\infty$ un campo vectorial definido en $U \subseteq M$ y λ el parámetro de la congruencia. Consideremos la transformación que consiste en mapear cada punto de una curva integral en otro una distancia paramétrica $\Delta\lambda$ sobre la misma curva de la congruencia. Si el mapeo existe para toda $\Delta\lambda$, entonces existe una familia unidimensional diferenciable de tales mapeos, a

saber: un grupo de Lie uniparamétrico con ley de composición $\Delta\lambda_1 + \Delta\lambda_2$. Un mapeo con estas características es denominado *arrastré a lo largo de la congruencia* o, simplemente, *arrastré de Lie*.

Consideremos una función f sobre una variedad M y α un arrastre. El arrastre de f a lo largo de $\frac{d}{ds}$ define una nueva función $f_{\Delta\lambda}$ tal que $f(q \in \text{curva}) = f_{\Delta\lambda}(\alpha(q) \in \text{curva})$. Si $f(s) = f_{\Delta\lambda}(s) \forall s$ en la región adecuada, entonces dado que $s = \alpha(q)$ tenemos que: $f(\alpha(q)) = f_{\Delta\lambda}(\alpha(q)) = f(q) \Rightarrow f(q) = f(\alpha(q))$. Es decir $f = f \circ \alpha$ y, por tanto, f es invariante bajo el arrastre α (f es una función periódica sobre cada curva de la congruencia). Si además f es invariante para toda $\Delta\lambda$, entonces f es invariante bajo todos los arrastres α y, en consecuencia, f es invariante bajo el grupo uniparamétrico de Lie: f es *Lie arrastrada*.

Del hecho de que f sea invariante bajo el grupo 1-paramétrico de Lie, se sigue que f es constante sobre cada una de las curvas de la congruencia: $\frac{d}{ds} f = 0$.

De tal manera, arrastrar una función consiste en construir una nueva $f_{\Delta\lambda}$ tal que $f(p) = f_{\Delta\lambda}(\alpha(p))$. El *arrastré de Lie de una función* ocurre cuando esta es invariante bajo el grupo 1-paramétrico de Lie: $f(\alpha(p)) = f_{\Delta\lambda}(\alpha(p)) \forall \Delta\lambda, \alpha(p)$. Por lo tanto, si f es Lie arrastrada por el campo $\frac{d}{ds}$, f es constante a lo largo de cada una de las curvas integrales de la congruencia.

El *arrastré de Lie de un campo vectorial* $\frac{d}{ds}$ por el campo $\frac{d}{ds}$ es el análogo vectorial del arrastre de Lie de una función:

Sea $(\frac{d}{ds})_0$ una curva integral del campo $\frac{d}{ds}$ elegida arbitrariamente. Arrastremos a tal curva una distancia paramétrica $\Delta\lambda'$ con α y en cada punto compáremosla con la curva integral del campo. Si sucede que para $\Delta\lambda'$ el arrastre de la curva coincide con la curva integral del campo, entonces el arrastre de la curva y las curvas integrales son iguales en "saltos" de α ; sin embargo, ello no implica aún que el campo sea Lie arrastrado. El campo $\frac{d}{ds}$ es Lie arrastrado si para toda α ($\forall \Delta\lambda$) coinciden el campo y el arrastrado; es decir, si el arrastre de $(\frac{d}{ds})_0$ coincide con la curva integral del campo $\forall \alpha$.¹

§3.2 Derivada de Lie.

Supongamos que sobre una variedad M hay un campo vectorial suave. Notemos que si la variedad no está dotada con una métrica y además cuenta con curvatura no nula, entonces no es posible recurrir a la definición de derivada de un campo que se da en cálculo vectorial elemental, ello se debe a la pérdida de la unicidad en los conceptos de distancia y paralelismo.

¹ Nótese entonces que $\mu = cte$ sobre $\frac{d}{ds}$.

En consecuencia, habremos de buscar sustitutos a tales conceptos que nos permitan definir a la derivada; dichos sustitutos serán, a grandes rasgos, los siguientes: la distancia entre dos puntos será simplemente la distancia paramétrica que existe entre los mismos dada una curva γ de una congruencia que los conecte y, por otro lado, el arrastre de Lie de un vector definido sobre $p \in \gamma$ a lo largo de la curva γ hasta el punto $q \in \gamma$ sustituirá al paralelismo.

Derivemos entonces la expresión analítica:

Comencemos considerando una función escalar f y una congruencia definida en alguna región de M dada por el campo $V = \frac{d}{d\lambda} \in C^\infty$. El arrastre de Lie de $f(\lambda_0 + \Delta\lambda)$ define un nuevo campo escalar f^* tal que $\frac{df^*}{d\lambda} = 0$. Definimos entonces a la derivada de f (derivada de Lie de f a lo largo del campo vectorial de arrastre $V = \frac{d}{d\lambda}$) como la siguiente operación:

$$\mathcal{L}_V f \equiv \lim_{\Delta\lambda \rightarrow 0} \frac{f^*(\lambda_0) - f(\lambda_0)}{\Delta\lambda} = \lim_{\Delta\lambda \rightarrow 0} \frac{f(\lambda_0 + \Delta\lambda) - f(\lambda_0)}{\Delta\lambda} = \frac{df}{d\lambda} \Big|_{\lambda_0}$$

Sea el campo vectorial $U = \frac{d}{d\mu}$ y f una función arbitraria. El arrastre de $U(\lambda_0 + \Delta\lambda)$ define un nuevo campo $U^* = \frac{d}{d\mu^*}$ tal que $[U^*, V] = 0$: $\frac{d}{d\lambda} \frac{d}{d\mu^*} f = \frac{d}{d\mu^*} \frac{d}{d\lambda} f$ donde sea. Por lo tanto tenemos, para campos vectoriales analíticos, lo siguiente:

$$\begin{aligned} \frac{d}{d\mu^*} f \Big|_{\lambda_0} &= \frac{d}{d\mu} f \Big|_{\lambda_0 + \Delta\lambda} - \frac{d}{d\mu^*} \frac{d}{d\lambda} f \Big|_{\lambda_0} \Delta\lambda + O(\Delta\lambda^2) \\ &= \frac{d}{d\mu} f \Big|_{\lambda_0} + \frac{d}{d\lambda} \frac{d}{d\mu} f \Big|_{\lambda_0} \Delta\lambda - \frac{d}{d\mu^*} \frac{d}{d\lambda} f \Big|_{\lambda_0} \Delta\lambda + O(\Delta\lambda^2) \\ \Rightarrow \frac{d}{d\mu^*} f \Big|_{\lambda_0} - \frac{d}{d\mu} f \Big|_{\lambda_0} &= \left\{ \frac{d}{d\lambda} \frac{d}{d\mu} f \Big|_{\lambda_0} - \frac{d}{d\mu^*} \frac{d}{d\lambda} f \Big|_{\lambda_0} \right\} \Delta\lambda + O(\Delta\lambda^2) \end{aligned}$$

Definimos la derivada de Lie $\mathcal{L}_V U$ como el campo vectorial que opera sobre f para dar:

$$\begin{aligned} [\mathcal{L}_V U](f) &= \lim_{\Delta\lambda \rightarrow 0} \left[\frac{U^*(\lambda_0) - U(\lambda_0)}{\Delta\lambda} \right](f) \\ &= \lim_{\Delta\lambda \rightarrow 0} \left\{ \frac{d}{d\lambda} \frac{d}{d\mu} f \Big|_{\lambda_0} - \frac{d}{d\mu^*} \frac{d}{d\lambda} f \Big|_{\lambda_0} \right\} \\ \Rightarrow \mathcal{L}_V U &= \lim_{\Delta\lambda \rightarrow 0} \left\{ \frac{d}{d\lambda} \frac{d}{d\mu} - \frac{d}{d\mu^*} \frac{d}{d\lambda} \right\} \end{aligned}$$

sustituyendo la identidad $\frac{d}{d\mu^*} = \frac{d}{d\mu} + \Delta\lambda \left(\frac{d}{d\lambda} \frac{d}{d\mu} - \frac{d}{d\mu^*} \frac{d}{d\lambda} \right) + O(\Delta\lambda^2)$ y llevando a cabo la operación de límite se obtiene, finalmente, que:

$$\mathcal{L}_V U = [V, U]$$

Este resultado no es sorprendente, pues de acuerdo a la definición de la derivada de Lie, un campo vectorial tiene derivada de Lie igual a cero si este es Lie arrastrado (i.e: si conmuta con el campo de arrastre).

De la antisimetría del paréntesis de Lie se sigue que:

$$\mathcal{L}_V U = -\mathcal{L}_U V$$

La derivada de Lie de un campo de 1-formas se halla análogamente a la de un campo vectorial: se arrastra la 1-forma en $\lambda_0 + \Delta\lambda$ hasta λ_0 y se lleva a cabo la operación de límite correspondiente. El resultado es que la derivada de Lie de la 1-forma $\bar{k}$ a lo largo de V está dada por:

$$\mathcal{L}_V[\bar{k}(W)] = (\mathcal{L}_V \bar{k})(W) + \bar{k}(\mathcal{L}_V W)$$

donde W es cualquier campo vectorial ².

La extensión de la derivada de Lie para tensores de mayor orden se sigue de manera natural:

$$\mathcal{L}_V(T(\bar{w}, \dots; \bar{u}, \dots)) = (\mathcal{L}_V T)(\bar{w}, \dots; \bar{u}, \dots) + T(\mathcal{L}_V \bar{w}, \dots; \bar{u}, \dots) + \dots + T(\bar{w}, \dots; \mathcal{L}_V \bar{u}, \dots) + \dots$$

y

$$\mathcal{L}_V(A \otimes B) = (\mathcal{L}_V A) \otimes B + A \otimes (\mathcal{L}_V B)$$

donde los tensores, 1-formas y vectores son arbitrarios.

En un sistema coordenado $\{x^i\}$ tendríamos por ejemplo que:

$$\mathcal{L}_V(T(\bar{dx}^i; \frac{\partial}{\partial x^j})) = (\mathcal{L}_V T)(\bar{dx}^i; \frac{\partial}{\partial x^j}) + T(\mathcal{L}_V \bar{dx}^i; \frac{\partial}{\partial x^j}) + T(\bar{w}; \mathcal{L}_V \frac{\partial}{\partial x^j})$$

Es fácil verificar que:

$$\begin{aligned} \mathcal{L}_V \bar{dx}^i &= v^i_{,k} \bar{dx}^k \\ \mathcal{L}_V \frac{\partial}{\partial x^j} &= -v^k_{,j} \frac{\partial}{\partial x^k} \end{aligned}$$

donde v^j son las componentes del campo V ; por lo tanto:

² W juega el mismo papel que f en la deducción de la expresión para la derivada de Lie de un campo vectorial

$$\mathcal{L}_V T_j^i = (\mathcal{L}_V T)_j^i + T(v^i, {}_k \bar{d}x^k; \frac{\partial}{\partial x^j}) + T(\bar{d}x^i; -v^k, {}_j \frac{\partial}{\partial x^k})$$

es decir:

$$\mathcal{L}_V T_j^i = (\mathcal{L}_V T)_j^i + v^i, {}_k T_j^k - v^k, {}_j T_k^i$$

§3.3 Invariancia.

Diremos que un campo tensorial T es invariante bajo un campo vectorial V si $\mathcal{L}_V T = 0$.

Supongamos que contamos con el conjunto de campos tensoriales $F = \{T_1, T_2, \dots\}$ y deseamos estudiar sus propiedades de invariancia. Entonces, para tal efecto, es necesario recurrir al concepto de álgebra de Lie de campos vectoriales:

Definimos a un *álgebra de Lie de campos vectoriales* sobre una región U de M , como el conjunto $\mathcal{A}$ de campos vectoriales sobre U que conforman un espacio vectorial bajo la adición y que es cerrado bajo la operación del paréntesis de Lie.

Teorema 3.1 *El conjunto de todos los campos vectoriales V , bajo los cuales todos los campos en F son invariantes, forman un álgebra de Lie.*

Demostración: Sean V y W cualesquiera dos campos que satisfacen que $\mathcal{L}_V T_i = \mathcal{L}_W T_i = 0 \quad \forall T_i \in F$. Sea $X = aV + bW$ (con a y b constantes), entonces:

$$\mathcal{L}_X T_i = \mathcal{L}_{aV} T_i + \mathcal{L}_{bW} T_i = a\mathcal{L}_V T_i + b\mathcal{L}_W T_i = 0$$

$$\Rightarrow \mathcal{L}_{aV+bW} T_i = 0 \quad \forall T_i \in F$$

por lo tanto, los V tales que $\mathcal{L}_V T_i = 0$ para todo $T_i \in F$ forman un espacio vectorial bajo la adición.

Dado que $\mathcal{L}_V T_i = \mathcal{L}_W T_i = 0 \Rightarrow [\mathcal{L}_V, \mathcal{L}_W] T_i = 0$. Por otra parte, sobre tensores es válida la siguiente expresión: $[\mathcal{L}_V, \mathcal{L}_W] = \mathcal{L}_{[V, W]}$. En consecuencia:

$$\mathcal{L}_{[V, W]} T_i = 0 \quad \forall T_i \in F$$

entonces $[V, W]$ está en el conjunto y por lo tanto éste es cerrado bajo los paréntesis de Lie.

□

§3.4 Campos vectoriales de Killing.

Sea M una variedad riemanniana³ con métrica g .

Diremos que un campo vectorial V sobre M es un *campo vectorial de Killing* si V es tal que $\mathcal{L}_V g = 0$. De acuerdo a esta definición, es posible probar [7] que en una variedad n -dimensional hay, a lo mas, $\frac{n(n+1)}{2}$ campos vectoriales de Killing independientes.

Sabemos que en un sistema coordenado $\{x^i\}$ tenemos la siguiente igualdad:

$$v^k g_{ij,k} = (\mathcal{L}_V g)_{ij} - v^k_{,i} g_{kj} - v^k_{,j} g_{ik}$$

en consecuencia, si V es un campo vectorial de Killing, las componentes de la derivada de Lie del tensor métrico habrán de ser nulas:

$$(\mathcal{L}_V g)_{ij} = v^k g_{ij,k} + v^k_{,i} g_{kj} + v^k_{,j} g_{ik} = 0$$

Si elegimos el sistema coordenado de tal manera que las curvas integrales de V sean una familia de líneas coordenadas; digamos que $V = \frac{\partial}{\partial x^1}$, entonces $v^i = \delta^i_1$ y en consecuencia:

$$(\mathcal{L}_V g)_{ij} = g_{ij,1} = 0$$

Es decir, las componentes de la métrica son independientes de la coordenada x^1 .

El recíproco es también cierto: si existe un sistema coordenado en el cual las componentes de la métrica son independientes de cierta coordenada, entonces el vector base para tal coordenada es un vector de Killing.

§3.5 Definición de grupo de Lie.

Definimos a un *grupo de Lie de dimensión finita* como una variedad G , de dimensión n y C^∞ , que cuenta con los siguientes difeomorfismos:

Cualquier $g \in G$ mapea a $h \in G$ como sigue:

$$h \mapsto gh \quad \text{traslación izquierda por } g$$

6

³Una variedad dotada con una métrica es llamada una variedad riemanniana [6].

$h \mapsto hg$ traslación derecha por g

Cabe destacar que, en general, el grupo no es abeliano $hg \neq gh$.

Cualquier vecindad de la identidad (e) es entonces mapeada por una traslación izquierda, dada por una g particular, en una vecindad de g . Puesto que el mapeo lleva curvas en curvas, entonces éste induce un mapeo de vectores tangentes en e ($T_e G$) a vectores tangentes en g ($T_g G$). Este mapeo suele denotarse como $L_g : T_e G \rightarrow T_g G$.

Se dice que un campo vectorial V definido en G es *invariante-izquierdo* si el campo V en g coincide con el mapeo L_g de V en e para toda g (i.e: V es un campo invariante-izquierdo si $L_g V(e) = V(g) \forall g \in G$). Por la ley de composición en grupos se sigue que $L_g V(h) = V(gh)$ para cualquier $h \in G$.

Cada $W \in T_e G$ define un único campo invariante-izquierdo. En consecuencia, los campos vectoriales invariantes-izquierdos forman un espacio vectorial n -dimensional. Si V y W son cualesquiera dos campos vectoriales invariantes-izquierdos, entonces L_g mapea $[V, W]$ en e a $[V, W]$ en g (i.e: $[V, W]$ es también un invariante-izquierdo). Por consiguiente, tenemos que los campos vectoriales invariantes-izquierdos *forman un álgebra de Lie* que habremos simplemente de llamar *álgebra de Lie de G* y denotar por $\mathcal{G}$. Es importante señalar que el álgebra de Lie es completamente caracterizada por las llamadas *constantes de estructura* (c_{kl}^i) definidas como sigue:

Sea $\{V_i; i = 1, \dots, n\}$ una base para el álgebra de Lie; esto es, un conjunto linealmente independiente de campos vectoriales invariantes-izquierdos. Entonces, siempre es posible escribir $[V_k, V_l] = c_{kl}^i V_i$. Si se da el caso en que todas las constantes de estructura son nulas, entonces diremos que el álgebra de Lie es *abeliana*.

Es obvio que la base $\{V_i\}$ no es única, bajo un cambio de base las componentes c_{kl}^i se transforman como las componentes de un tensor ($\frac{1}{2}$). Sin embargo, cada grupo y álgebra de Lie tiene un único *tensor de estructura* C .

Notemos que cualquier base $\{V_i(e); i = 1, \dots, n\}$ para $T_e G$ define un conjunto linealmente independiente de campos vectoriales invariantes-izquierdos. En efecto:

Como g es un difeomorfismo $\Rightarrow L_g : T_e G \rightarrow T_g G$ es 1-1 e invertible $\Rightarrow L_g V_i(e)$ es linealmente independiente de $L_g V_j(e)$ para toda g e $i \neq j$. En consecuencia, el conjunto $\{V_i(g), \dots, V_n(g)\}$ es un conjunto de vectores linealmente independientes en $T_g G$ para toda g .

§3.6 Subgrupos uniparamétricos.

Inherente a la definición de campo vectorial invariante izquierdo, esta bien definido, en al menos una vecindad de la identidad e de G , el concepto de curvas integrales suaves y simples. Por ende, es posible considerar, en particular, a la curva integral de un campo V invariante-izquierdo que pasa a través de e . Puesto que hay un único vector tangente en e y un único parámetro t para el cual corresponde a $t = 0$, por analogía con el análisis llevado a cabo en la sección §2.3.3 podemos denotar entonces a los puntos de G sobre la citada curva como:

$$g_{V(e)}(t) = e^{tV}|_e$$

Puesto que por definición tenemos que:

$$e^{t_1}V e^{t_2}V|_e = e^{(t_1+t_2)V}|_e$$

entonces los puntos sobre la curva forman un grupo:

$$g_{V(e)}(t_1 + t_2) = e^{(t_1+t_2)V}|_e = e^{t_1}V e^{t_2}V|_e = g_{V(e)}(t_1)g_{V(e)}(t_2)$$

Este grupo es denominado *subgrupo uniparamétrico* de G y es, obviamente, abeliano puesto que $g_{V(e)}(t_1 + t_2) = g_{V(e)}(t_2 + t_1)$.

Para cada $V \in T_e G$ hay un único subgrupo. Inclusive, dado que cada subgrupo 1-paramétrico debe ser una curva C^∞ en G , que pasa a través de e , entonces hay una correspondencia 1-1 entre los subgrupos 1-paramétricos de G y los elementos del álgebra de Lie de G .

Antes de pasar a la siguiente sección, hagamos un breve paréntesis para definir a un campo vectorial invariante-derecho:

Por analogía con los campos invariantes-izquierdos, diremos que un campo V es invariante-derecho si el campo vectorial V coincide con el mapeo vectorial inducido en la acción por la derecha $\forall g \in G$. Los campos invariantes-derechos también forman un álgebra de Lie, las curvas integrales correspondientes a estos campos que pasan por e coinciden con las curvas integrales de los campos vectoriales invariantes-izquierdos; sin embargo, las curvas de la congruencia que no contienen a e no coinciden, en general, a menos que el grupo sea abeliano.

§3.7 Ejemplos de grupos de Lie.

i).- El primer ejemplo, y el mas sencillo de todos, es el espacio vectorial $\mathbb{R}^n$. Este espacio vectorial es una variedad y un grupo abeliano bajo la adición de vectores. Los subgrupos

1-paramétricos son las líneas rectas que pasan por el origen (rayos). Los campos vectoriales invariantes-izquierdos son paralelos a los rayos, de tal manera que conmutan entre sí. El álgebra de Lie es entonces el espacio vectorial $T_e \mathbb{R}^n$ con el paréntesis trivial abeliano $[V, W] = 0 \forall V, W \in T_e \mathbb{R}^n$.

ii).- El siguiente ejemplo consiste en presentar a uno de los grupos más importantes en física: el grupo de las matrices reales de $n \times n$, con determinante no nulo, llamado *grupo lineal general en n dimensiones reales* y denotado por $GL(n, \mathbb{R})$. Observemos que, efectivamente, es un grupo de Lie:

1).- Notemos que este es un grupo bajo la operación de multiplicación de matrices y que la matriz unitaria es el elemento identidad del grupo ⁴.

2).- El grupo es una variedad. Cualquier $A \in GL(n, \mathbb{R})$ con entradas $\{a'_{ij}; i, j = 1, \dots, n\}$ tiene una vecindad de radio ϵ definida por aquellas matrices B para las cuales sucede que $|b'_{ij} - a'_{ij}| < \epsilon \forall i, j$, donde ϵ puede ser tan pequeño como sea necesario para asegurar que cada B tenga determinante no nulo.

Los números $x^i_j = b^i_j - a^i_j$ son las coordenadas para esta vecindad; puesto que hay n^2 coordenadas, todas independientes, la dimensión de $GL(n, \mathbb{R})$ es entonces n^2 . De hecho, ésta es una subvariedad de $\mathbb{R}^{n^2}$. Dado que $\mathbb{R}^{n^2}$ es idéntico al espacio tangente de cualquiera de sus puntos, entonces el espacio tangente a la identidad e de $GL(n, \mathbb{R})$ es $\mathbb{R}^{n^2}$, y cualquier vector tangente puede ser representado como una matriz.

Por ejemplo, la curva en $GL(n, \mathbb{R})$ que corresponde a las matrices $\text{diag}(1 + e^\lambda, 1, 1, \dots, 1)$, cuyo parámetro es λ , tiene vector tangente $\text{diag}(1, 0, 0, \dots, 0)$ en $\lambda = 0$. Este último tiene determinante nulo, ilustrando el hecho de que cualquier matriz está en $T_e M$ y, por lo tanto, cualquier matriz genera un subgrupo 1-paramétrico y un campo vectorial invariante-izquierdo [4].

El subgrupo 1-paramétrico generado por cualquier matriz A , es la curva integral de un campo vectorial invariante-izquierdo que pasa por e y cuya tangente en e es A . Si denotamos a las componentes del subgrupo como las matrices $g_A(t)$, con $\frac{dg_A(t)}{dt}|_{t=0} = A$ (lo cual sólo significa que $\frac{dg_A(t)}{dt}|_{t=0} = a'_{ij} \forall i, j$) entonces, dado que los puntos de la curva integral forman un grupo, tenemos que:

$$g_A(t + \Delta t) = g_A(t)g_A(\Delta t)$$

⁴La restricción de determinante no nulo es necesaria para garantizar la existencia del elemento inverso para cualquier matriz

$$\begin{aligned}
\Rightarrow \lim_{\Delta t \rightarrow 0} \frac{g_A(t + \Delta t) - g_A(t)}{\Delta t} &= \lim_{\Delta t \rightarrow 0} g_A(t) \frac{g_A(\Delta t) - g_A(0)}{\Delta t} \\
\Rightarrow \frac{dg_A(t)}{dt} &= g_A(t) \left\{ \frac{dg_A(t)}{dt} \right\}_{t=0} = g_A(t) A \\
\Rightarrow g_A(t) &= e^{tA} = e + tA + \frac{1}{2!}(tA)^2 + \frac{1}{3!}(tA)^3 + \dots
\end{aligned}$$

De tal manera, los subgrupos 1-paramétricos de $GL(n, \mathbb{R})$ son la exponenciación de matrices arbitrarias de $n \times n$. La matriz A es, con frecuencia, llamada el *generador infinitesimal* del subgrupo $g_A(t)$.

Es importante notar que no todo elemento de $GL(n, \mathbb{R})$ es un miembro de un subgrupo 1-paramétrico. Un argumento de plausibilidad es el siguiente:

El subgrupo 1-paramétrico es una curva continua en $GL(n, \mathbb{R})$ y sobre tal curva el determinante debe cambiar continuamente. Dado que el determinante es 1 en e , y no puede ser nulo, entonces no hay una curva continua que enlace a e con una matriz de determinante negativo.

Los elementos que pueden ser unidos a e de manera continua (sin que necesariamente pertenezcan a un subgrupo uniparamétrico) son denominados componentes de la identidad del grupo. Por ejemplo, la siguiente matriz

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

es un elemento de la componente de la identidad de $GL(2, \mathbb{R})$ y, no obstante, no pertenece a ningún subgrupo 1-paramétrico [4].

Ahora bien, dado un vector tangente $A(e)$ en e y su subgrupo 1-paramétrico $g_{A(e)}(t)$, la acción por la izquierda $f g_{A(e)}(t)$ sobre esta curva por la matriz $f \in GL(n, \mathbb{R})$ produce una curva de la congruencia del campo vectorial invariante-izquierdo correspondiente a $A(e)$. Si f está en la curva $g_{B(e)}(t)$, que pasa por e , generada por $B(e) \in T_e GL(n, \mathbb{R})$, entonces el paréntesis de Lie de los dos campos vectoriales en e es:

$$\begin{aligned}
g_{A(e)}(t)g_{B(e)}(t) - g_{B(e)}(t)g_{A(e)}(t) &= t^2[A, B]|_e + O(\geq t^3) \\
\Rightarrow [A, B]|_e &= \lim_{t \rightarrow 0} \frac{1}{t^2} (g_{A(e)}(t)g_{B(e)}(t) - g_{B(e)}(t)g_{A(e)}(t))
\end{aligned}$$

por el desarrollo exponencial tenemos entonces que:

$$[A, B]|_e = \lim_{t \rightarrow 0} \frac{1}{t^2} \{t^2(A(e)B(e) - B(e)A(e)) + O(\geq t^3)\}$$

$$\Rightarrow [A, B]_* = A(e)B(e) - B(e)A(e)$$

Es decir, el paréntesis de Lie de cualesquiera dos campos vectoriales invariantes-izquierdos en e es simplemente el conmutador de las matrices que generan los campos. En consecuencia, el campo vectorial invariante-izquierdo generado por este conmutador es el elemento del álgebra de Lie de $GL(n, \mathbb{R})$, que no es mas que el paréntesis de los campos originales.

iii).- El siguiente ejemplo consiste en un grupo que surgió en una discusión previa (§ 2.9): el grupo de simetría Euclidiano $O(n)$ (donde $O(n)$ significa *grupo ortogonal en n dimensiones*).

Comencemos notando una característica representativa para este grupo:

Puesto que $\det(A) = (\det(A^{-1}))^{-1}$ y $\det(A) = \det(A^T)$, entonces si $A \in O(n)$ tenemos entonces que: $(\det(A))^2 = 1 \Rightarrow \det(A) = \pm 1$. En consecuencia, se sigue entonces que $O(n)$ es disconexo.

El conjunto de matrices $A \in O(n)$ tales que $\det(A) = 1$, forman un subgrupo (grupo de rotaciones) llamado *grupo especial ortogonal* y denotado por $SO(n)$. Obsérvese que el conjunto de las matrices $A \in O(n)$ tales que $\det(A) = -1$ no son un subgrupo dado que no contienen al elemento identidad e .

El álgebra de Lie de $O(n)$ consta de todas las matrices antisimétricas y la dimensión del grupo es $\frac{1}{2}n(n-1)$. En efecto:

Sea $A \in T_e O(n)$, el subgrupo 1-paramétrico correspondiente está dado entonces por:

$$g_A(t) = e^{tA}$$

Puesto que $g_A(t) \in O(n)$, entonces $(g_A(t))^{-1} = (g_A(t))^T \forall t$, en consecuencia:

$$(e + tA + \frac{1}{2!}(tA)^2 + \dots)^{-1} = (e + tA + \frac{1}{2!}(tA)^2 + \dots)^T$$

pero como sucede que:

$$(e + tA + \frac{1}{2!}(tA)^2 + \dots)^{-1} = (e - tA + \frac{1}{2!}(-tA)^2 + \dots)$$

y

$$(e + tA + \frac{1}{2!}(tA)^2 + \dots)^T = (e + tA^T + \frac{1}{2!}(tA^T)^2 + \dots)$$

entonces

$$t(A^T + A) + \frac{1}{2!}t^2((A^T)^2 - A^2) + \frac{1}{3!}t^3((A^T)^3 + A^3) + \dots = 0, \quad \forall t$$

$$\Rightarrow A^T + A = 0 \Rightarrow A = -A^T$$

por lo tanto, A es una matriz antisimétrica y se sigue entonces que el álgebra de $O(n)$ está compuesta por el conjunto de las matrices antisimétricas.

La dimensión de $O(n)$ es, en consecuencia, $\frac{1}{2}n(n-1)$ dado que la dimensión está dada por el número máximo de matrices antisimétricas linealmente independientes.

Notemos que una base para el álgebra de Lie en el caso particular $O(3)$ ($n=3$) podría ser el siguiente conjunto de matrices L_i antisimétricas y linealmente independientes:

$$L_1 = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & -1 \\ 0 & 1 & 0 \end{pmatrix}, \quad L_2 = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ -1 & 0 & 0 \end{pmatrix}, \quad L_3 = \begin{pmatrix} 0 & -1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}$$

que satisfacen las siguientes reglas de conmutación:

$$[L_1, L_2] = L_3, \quad [L_2, L_3] = L_1, \quad [L_3, L_1] = L_2$$

iv).- Ejemplos importantes de grupos de Lie son el grupo homogéneo de Lorentz y el grupo inhomogéneo de Lorentz o grupo de Poincaré.

Grupo homogéneo de Lorentz:

Sean S y S' dos sistemas inerciales con ejes espaciales paralelos y supongamos que S' se mueve respecto a S en la dirección positiva del eje x , con velocidad $v < c$. Adicionalmente pedimos que al tiempo $t = 0$ en S , los orígenes espaciales O' de S' y O de S coincidan (y viceversa desde S'). Si algún evento en el espacio tiempo tiene coordenadas (x, y, z, t) en S y (x', y', z', t') en S' , entonces los sistemas se conectan a través de las ecuaciones de transformación de Lorentz [8]:

$$x' = \gamma(x - vt), \quad y' = y, \quad z' = z, \quad t' = \gamma\left(t - \frac{vx}{c^2}\right) \quad (3.1)$$

donde $\gamma = (1 - v^2/c^2)^{-1/2}$.

Ahora bien, si un objeto se desplaza a velocidad constante u a lo largo del eje x en S , entonces su velocidad u' en S' a lo largo de x' es:

$$u' = \frac{u - v}{1 - uv/c^2}$$

Nótese que si el "objeto" fuese una onda de luz, entonces $u = c$ y $u' = c$.

Si S'' se mueve respecto a S' a velocidad v' , como S' lo hace respecto a S a velocidad v , entonces la velocidad de S'' relativa a S es v'' :

$$v'' = \frac{v + v'}{1 + vv'/c^2} \quad (3.2)$$

Las variables x', t' están relacionadas con x, t como en (3.1) y x'', t'' con x', t' de manera análoga pero con v' en lugar de v ; entonces, las variables x'', t'' están relacionadas con x, t también de la misma manera pero con v'' . Por lo tanto, la composición de transformaciones de Lorentz en una dimensión da de nueva cuenta una transformación de Lorentz del mismo tipo y, juntas, constituyen un grupo abeliano 1-paramétrico. La simetría del miembro derecho de (3.2) respecto a v y v' garantiza la naturaleza abeliana de éste grupo de transformaciones. Si la velocidad v es el parámetro que etiqueta a las transformaciones, nótese que entonces la composición de dos transformaciones no corresponde a la adición de parámetros; para conseguir que la adición de parámetros corresponda a la composición de transformaciones es necesario definir un nuevo parámetro:

Sea $v = c \tanh(\xi)$, $-\infty < \xi < \infty$, donde ξ es el parámetro de "rapidez". Evidentemente $|v| < |c|$, pues $-1 < \tanh(\xi) < 1$. Entonces, $\cosh(\xi) = \gamma$ y $\sinh(\xi) = \gamma v/c$. Si contamos con la cadena anterior, i.e:

$$S \xrightarrow{v} S' \xrightarrow{v'} S''$$

tenemos, por (3.2), en términos del nuevo parámetro que:

$$\xi'' = \xi + \xi'$$

La transformación (3.1) es ahora:

$$x' = x \cosh(\xi) - ct \sinh(\xi), \quad ct' = ct \cosh(\xi) - x \sinh(\xi), \quad y' = y, \quad z' = z$$

En general: si los sistemas inerciales S' y S tienen ejes espaciales paralelos, S' se mueve respecto a S uniformemente en la dirección $\hat{\xi}$ con rapidez $|\hat{\xi}| = \xi$ y, como antes, los orígenes O y O' coinciden al tiempo $t = 0$ en S así como al tiempo $t' = 0$ en S' , entonces la transformación de Lorentz entre los sistemas es [8]:

$$\begin{aligned} x'_j &= (\delta_{jk} + \xi_j \xi_k \frac{\cosh(\xi) - 1}{\xi^2}) x_k - ct \xi_j \frac{\sinh(\xi)}{\xi} \\ ct' &= ct \cosh(\xi) - x_j \xi_j \frac{\sinh(\xi)}{\xi} \end{aligned} \quad (3.3)$$

donde ξ_j son las coordenadas de $\vec{\xi}$ y se sobreentiende suma sobre índices repetidos.

La familia de transformaciones de Lorentz en la misma dirección $\vec{\xi}$, con parámetros de rapidez $\xi, \xi', \dots$ forman un grupo 1-paramétrico abeliano, donde la adición de los parámetros corresponde a la composición de las transformaciones.

Consideremos ahora no paralelos a los ejes de los sistemas inerciales S y S' , pero con la misma orientación y manteniendo las demás condiciones. Entonces la *transformación de Lorentz homogénea* consiste en aplicar una rotación en S para obtener un nuevo sistema S_1 con ejes paralelos a S' :

$$S_1 = e^{\alpha_i L^i} S, \quad e^{\alpha_i L^i} \in SO(3)$$

Aplicamos una transformación "pura" de Lorentz a S_1 , correspondiente al vector de rapidez $\vec{\xi}$, para obtener finalmente la relación entre S' y S . Denotemos a esta secuencia de operaciones como el par $(\vec{\xi}, \vec{\alpha})$, entonces S' es la aplicación del par sobre S : $S' = (\vec{\xi}, \vec{\alpha})S$. Es decir, en notación 4-vectorial, tendríamos que [8]:

$$S \rightarrow S' : x'^{\mu} = \Lambda^{\mu}_{\nu}(\vec{\xi}, \vec{\alpha}) x^{\nu}, \quad (\mu, \nu = 0, 1, 2, 3)$$

donde:

$$\Lambda^i_k(\vec{\xi}, \vec{\alpha}) = A_{jk}(\vec{\alpha}) + \xi_j \xi_m A_{mk}(\vec{\alpha}) \frac{\cosh(\xi) - 1}{\xi^2}$$

$$\Lambda^0_j(\vec{\xi}, \vec{\alpha}) = -\xi_j \frac{\sinh(\xi)}{\xi} \quad (3.4)$$

$$\Lambda^0_0(\vec{\xi}, \vec{\alpha}) = -\xi_k A_{kj}(\vec{\alpha}) \frac{\sinh(\xi)}{\xi}$$

$$\Lambda^0_0(\vec{\xi}, \vec{\alpha}) = \cosh(\xi)$$

donde el elemento de línea $ds^2 = \eta_{\mu\nu} dx^{\mu} dx^{\nu}$, con $\eta_{\mu\nu} = 0$ si $\mu \neq \nu$, $\eta_{00} = -1$, $\eta_{11} = \eta_{22} = \eta_{33} = 1$, es invariante bajo estas transformaciones.

Si S y S' son cualesquiera sistemas inerciales cuyos orígenes espacio-temporales coinciden y para los cuales el sentido de incremento temporal y orientación del sistema coordinado espacial coincide, entonces las variables espacio-temporales x^{μ} y x'^{μ} están relacionadas por (3.4). El conjunto de todas las transformaciones $S \rightarrow S' = (\vec{\xi}, \vec{\alpha})S \forall \vec{\xi}, \vec{\alpha}$, constituyen el *grupo homogéneo de Lorentz* (GHL), este grupo es 6-paramétrico (i.e. es de seis dimensiones) y puesto que la forma cuadrática invariante contiene tres términos del mismo signo y uno del signo opuesto, el grupo de transformaciones Λ puede ser pensado como un grupo de rotaciones "pseudo-ortogonal" en el espacio real 4-dimensional y denotado por $SO(3, 1)$.

Si tenemos tres sistemas inerciales S, S', S'' , entonces:

$$S \xrightarrow{\Delta} S' \xrightarrow{\Delta} S''$$

i.e:

$$S'' = \Lambda' \Lambda S : x''^\mu = \Lambda'^\mu_\nu x'^\nu = \Lambda'^\mu_\nu \Lambda^\nu_\sigma x^\sigma \quad (3.5)$$

Los vectores base del subgrupo $(0, \vec{\alpha})$ son los vectores $\{L_i\}$ (como en el anterior ejemplo), en tanto que $\{K_i; i = 1, 2, 3\}$ corresponden al subgrupo 1-paramétrico de transformaciones "puras" de Lorentz. Como antes,

$$[L_i, L_j] = \epsilon_{ijm} L^m$$

donde ϵ_{ijm} son las componentes del tensor completamente antisimétrico de Levi-Civita.

Usando (3.4) y (3.5) se obtiene que:

$$e^{\alpha_i L^i} e^{\xi_j K^j} e^{-\alpha_m L^m} = e^{\xi'_k K^k}, \quad \xi'_j = A_{jk}(\vec{\alpha}) \xi_k \quad (3.6)$$

Comparando el primer orden del desarrollo de las exponenciales para el vector $\vec{K}$ en ambos miembros de (3.6), y desarrollando hasta primer orden en las componentes de $\vec{\alpha}$, tomando en cuenta la igualdad [8]:

$$A_{jk}(\vec{\alpha}) = \delta_{jk} \cos(\alpha) + \alpha_j \alpha_k \frac{1 - \cos(\alpha)}{\alpha^2} + \epsilon_{jkl} \alpha_l \frac{\sin(\alpha)}{\alpha}$$

obtenemos que:

$$[L_j, K_m] = \epsilon_{jmn} K^n$$

Para hallar $[K_j, K_m]$ basta con considerar el producto $e^{aK_1} e^{bK_2}$ entre subgrupos 1-paramétricos de transformaciones "puras" de Lorentz y considerar sólo los términos lineales en a y b , y de segundo orden solamente con ab , en el desarrollo de la exponencial. Entonces, es fácil probar que:

$$e^{-aK_1 - bK_2} e^{aK_1} e^{bK_2} = e^{ab[K_1, K_2]/2 + \dots}$$

Introduciendo la siguiente "cadena" de sistemas inerciales:

$$S_1 = e^{bK_2} S, \quad S_2 = e^{aK_1} S_1, \quad S' = e^{-aK_1 - bK_2} S_2$$

usando (3.3) en cada caso y combinando los resultados para la relación entre S' y S :

$$S' = e^{ab[K_1, K_2]/2 + \dots} S$$

se verifica directamente que S' se obtiene de S sólo por una rotación espacial [8]. En consecuencia el conmutador $[K_1, K_2]$ debe ser una combinación de los vectores L_i ; de hecho, si escribimos $[K_1, K_2] = \beta_j L^j$ se prueba, desarrollando $A_{jk}(\beta ab/2)$ como antes, que $\beta_i = -\epsilon_{12i}$. Por consiguiente, la permutación cíclica de los índices nos proporciona la siguiente regla de conmutación:

$$[K_i, K_m] = -\epsilon_{imn} L^n$$

El grupo de Poincaré ($\mathcal{P}$) o grupo inhomogéneo de Lorentz es un caso más general, que surge al combinar al GHL con el grupo de traslaciones espacio-temporales. Es decir, los elementos de $\mathcal{P}$ corresponden a todas las posibles transformaciones que conectan dos sistemas de referencia inerciales S y S' en el espacio-tiempo cuyos orígenes espacio-temporales y ejes pueden diferir. Los elementos de $\mathcal{P}$ pueden construirse como sigue: Relacionamos a S con S_1 vía una transformación homogénea de Lorentz ($S_1 = \Lambda S$). Después vamos de S_1 a S' con la aplicación de una traslación espacial dada por una cantidad $\vec{a}$ y, por último, una traslación en el tiempo por una cantidad b :

$$S' = e^{b\partial} e^{a^i \partial_i} \Lambda S, \quad x'^j = \Lambda^j_\mu x^\mu + a_j, \quad x'^0 = \Lambda^0_\mu x^\mu - b$$

Introduciendo notación tetra vectorial: a^μ , con $a^j = a_j$ y $a^0 = -a_0 = -b$; d_μ , con $-d_0 = a^0 = h$. Entonces:

$$S' = e^{a^\mu d_\mu} \Lambda S : \quad x'^\mu = \Lambda^\mu_\nu x^\nu + a^\mu$$

En consecuencia un elemento de $\mathcal{P}$ puede ser denotado por la pareja (Λ, a^μ) ; $\Lambda \in SO(3, 1)$ y a^μ alguna traslación espacio-temporal. Entonces:

$$S'' = (\Lambda', a'^\mu) S' = (\Lambda', a'^\mu) (\Lambda, a^\mu) S :$$

$$x''^\mu = \Lambda'^\mu_\nu x'^\nu + a'^\mu = \Lambda'^\mu_\nu (\Lambda^\nu_\sigma x^\sigma + a^\nu) + a'^\mu$$

$$(\Lambda', a'^\mu) (\Lambda, a^\mu) = (\Lambda' \Lambda, a'^\mu + \Lambda'^\mu_\nu a^\nu)$$

Para especificar a un elemento de $\mathcal{P}$ son necesarios diez parámetros ($\mathcal{P}$ es un grupo 10-paramétrico): seis de $SO(3, 1)$ y cuatro de las traslaciones espacio-temporales. Por lo tanto, $\dim \mathcal{P} = 10$.

Las reglas de conmutación para el grupo de Poincaré son [8]:

$$[L_j, L_m] = -[K_j, K_m] = \epsilon_{jmn} L^n, \quad [L_j, K_m] = \epsilon_{jmn} K^n$$

$$[d_j, d_m] = [d_j, h] = 0, \quad [L_j, d_m] = \epsilon_{jmn} d^n, \quad [L_j, h] = 0$$

$$[K_j, d_m] = \delta_{jm} h, \quad [K_j, h] = d^j$$

v).- Por último, concluiremos los ejemplos presentando al grupo especial unitario en n dimensiones $SU(n)$.

El grupo $SU(n)$ es un subgrupo de $U(n)$ (grupo unitario) que, a su vez, es un subgrupo de $GL(n, C)$; es decir, el grupo compuesto por todas las matrices complejas de $n \times n$ con determinante no nulo (grupo lineal general en n dimensiones complejas).

Dado que cada entrada puede ser un número complejo, entonces $GL(n, C)$ tiene $2n^2$ dimensiones reales. Hallemos la dimensión del subgrupo $U(n)$:

El grupo unitario está compuesto por las matrices u tales que $u^{-1} = u^\dagger$ (donde $u^\dagger$ es la matriz transpuesta conjugada de u). Por analogía con $O(n)$, el álgebra de $U(n)$ está compuesta por todas las $n \times n$ matrices antihermitianas ⁵. La dimensión real del álgebra es n^2 puesto que hay $n(n-1)$ reales en las partes triangulares y n en la diagonal. Como la dimensión del álgebra es igual a la dimensión del grupo, entonces $\dim U(n) = n^2$.

El subgrupo $SU(n)$ de $U(n)$ está compuesto por matrices cuyo determinante es igual a 1, esta condición extra implica que $SU(n)$ habrá de tener dimensión $n^2 - 1$. El álgebra de Lie de $SU(n)$ es el conjunto de todas las $n \times n$ matrices antihermitianas con traza nula.

Para el caso específico dos-dimensional, una base para el álgebra de Lie de $SU(2)$ podría ser la siguiente:

$$J_1 = \frac{1}{2} \begin{pmatrix} 0 & i \\ i & 0 \end{pmatrix}, \quad J_2 = \frac{1}{2} \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}, \quad J_3 = \frac{1}{2} \begin{pmatrix} i & 0 \\ 0 & -i \end{pmatrix}$$

Nótese que las reglas de conmutación son idénticas a las de $\{L_1, L_2, L_3\}$:

$$[J_1, J_2] = J_3, \quad [J_2, J_3] = J_1, \quad [J_3, J_1] = J_2$$

⁵A es una matriz antihermitiana si $A^\dagger = -A$.

Algunos Grupos de Lie.		
Grupo.	Dimensión.	álgebra de Lie.
$\mathbb{R}^n$	n	$T_x \mathbb{R}^n$
$GL(n, \mathbb{R})$	n^2	$C_k A^k, A^k \in M_{n \times n}$ **
$SL(2, \mathbb{R})$ *	3	$aT + b_k S^k$
$O(n)$	$\frac{1}{2}n(n-1)$	$\{A \in M_{n \times n} A = -A^T\}$
$SU(2)$	3	$C_k J^k$
$SO(3)$	3	$C_k L^k$
$SO(3, 1)$	6	$C_i L^i + D_j K^j$
$\mathbb{P}$	10	$C_i L^i + D_j K^j + E_l d^l + Fh$

(*) ver §5.5. (**) $M_{n \times n}$ denota a el conjunto de las matrices de $n \times n$.

§3.8 Álgebras de Lie y sus grupos.

Empezaremos esta sección dando una definición, un tanto más general, para el álgebra de Lie:

Definiremos a un *álgebra de Lie* como un espacio vectorial real V sobre el cual esta definida una regla de multiplicación bilineal, denotada por $[,]$, que produce a partir de cualesquiera dos campos vectoriales A y B un nuevo campo vectorial $[A, B] \in V$ que satisface lo siguiente:

- 1).- $[A, B] = -[B, A]$
- 2).- $[A, [B, C]] + [B, [C, A]] + [C, [A, B]] = 0$

La diferencia entre esta definición y la previa consiste en el hecho de que el conmutador de campos vectoriales efectivamente satisface las citadas reglas, pero no es el único. Por ejemplo, podríamos considerar igualmente al espacio vectorial $\mathbb{R}^3$ con el producto cruz usual: $[a, b] \equiv a \times b$.

A continuación haremos un breve paréntesis sólo para definir dos conceptos y con ello preparar adecuadamente el terreno para presentar un importante teorema en cuanto a las álgebras de Lie se refiere:

En primer término, sobre las variedades, diremos que una variedad conexa M cubre a otra N si hay un mapeo Π de M sobre N (todo elemento de N proviene de al menos uno de M)

tal que la imagen inversa de alguna vecindad v de cualquier punto $p \in N$ es la unión disjunta de vecindades abiertas que contienen a $\Pi^{-1}(p) \in M$. Por ejemplo, S^1 (círculo unitario) es cubierto por la línea real $\mathbb{R}$ una infinidad de veces a través del mapeo $\Pi: \mathbb{R} \rightarrow S^1$, donde $\Pi = (\cos(x), \sin(x))$.

La segunda definición que nos atañe concierne al estudio de grupos: Dos grupos, G_1 y G_2 , con operaciones binarias $(\cdot)$ y $(*)$ respectivamente, son isomorfos (idénticos en sus propiedades de grupo) si existe un mapeo f , 1-1 de G_1 sobre G_2 , tal que: $f(x \cdot y) = f(x) * f(y)$. Un *homomorfismo* entre grupos es semejante a un isomorfismo con la salvedad de que el mapeo puede ser "muchos a uno" y puede solo ser *dentro* y no *sobre* (i.e.: el mapeo esta definido para todos los elementos del dominio).

Teorema 3.2 *Cada álgebra de Lie es el álgebra de Lie de uno y sólo un grupo de Lie simplemente conexo⁶. Mas aún, cualquier otro grupo de Lie con la misma álgebra, pero sin ser simplemente conexo, es cubierto por el simplemente conexo correspondiente. La cubierta debe ser un homomorfismo entre los dos grupos [4].*

Como un ejemplo ilustrativo de este teorema consideremos a los grupos $SO(3)$ y $SU(2)$:

Observemos que existe un difeomorfismo φ entre la esfera S^3 , que es una variedad simplemente conexa, y $SU(2)$: $\varphi(\alpha_1, \alpha_2, \alpha_3, \alpha_4) = 2(\alpha_1 J_1 + \alpha_2 J_2 + \alpha_3 J_3) + \alpha_4 e$. En consecuencia, $SU(2)$ es una variedad simplemente conexa.

Notemos ahora lo siguiente:

$$e^{yJ_1} = \begin{pmatrix} \cos(\frac{y}{2}) & i \sin(\frac{y}{2}) \\ i \sin(\frac{y}{2}) & \cos(\frac{y}{2}) \end{pmatrix}$$

y

$$e^{sL_1} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos(s) & -\sin(s) \\ 0 & \sin(s) & \cos(s) \end{pmatrix}$$

entonces podemos decir que $\Pi: SU(2) \rightarrow SO(3)$ es tal que:

$$\Pi: \begin{pmatrix} \cos(\frac{t}{2}) & i \sin(\frac{t}{2}) \\ i \sin(\frac{t}{2}) & \cos(\frac{t}{2}) \end{pmatrix} \mapsto \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos(t) & -\sin(t) \\ 0 & \sin(t) & \cos(t) \end{pmatrix}$$

⁶Recordemos que una variedad es simplemente conexa si cada curva cerrada contenida en ella puede contraerse suavemente hasta un punto.

donde claramente Π es un homomorfismo entre los dos subgrupos 1-paramétricos.

Notemos que los elementos de $SU(2)$ correspondientes a los valores paramétricos t y $t+2\pi$ tienen la misma imagen en $SO(3)$. Los valores paramétricos $\{t+4n\pi : n \in \mathbb{Z}\}$ corresponden al mismo elemento de $SU(2)$, por lo tanto podemos afirmar que e^{tL_1} es una cubierta doble de e^{tL_1} .

Lo anterior naturalmente se extiende a todo el grupo: El mapeo

$$\Upsilon : e^{(t_1 J_1 + t_2 J_2 + t_3 J_3)} \longmapsto e^{(t_1 L_1 + t_2 L_2 + t_3 L_3)}$$

es una doble cubierta de $SO(3)$ por $SU(2)$.

El subgrupo uniparamétrico e^{tL_1} de $SU(2)$ comienza en e con $t = 0$ y vuelve a e en $t = 4\pi$. Los puntos que corresponden a t y $t + 2\pi$ son diametralmente opuestos uno del otro⁷. Estos puntos, bajo Π , son el mismo punto en $SO(3)$, de tal manera que habremos de identificar a $SO(3)$ como una mitad de S^3 , notando que los puntos ubicados en los extremos del diámetro que atraviesa al ecuador son, en realidad, el mismo punto. Esta mitad de S^3 , con la identificación de puntos citada, no es simplemente conexa; por ejemplo, la curva e^{tL_1} no puede ser contraída a un punto.

En consecuencia, los grupos $SO(3)$ y $SU(2)$ son idénticos *solamente* en alguna vecindad de e y esta última es, justamente, la razón por la cual sus álgebras de Lie coinciden.

Una vez ilustrado el Teorema 3.2, notemos ahora que este nos permite asegurar que un álgebra de Lie abeliana es el álgebra de Lie de un grupo abeliano. Veamos por qué:

Recordemos que (§3.7) el álgebra de Lie de $\mathbb{R}^n$ es el espacio vectorial $T_e \mathbb{R}^n$ dotado con el paréntesis trivial abeliano $[V, W] = 0 \forall V, W \in T_e \mathbb{R}^n$. Es decir, un álgebra de Lie abeliana n -dimensional es el álgebra de Lie de $\mathbb{R}^n$.

Ahora bien, como $\mathbb{R}^n$ es simplemente conexo se sigue entonces por el Teorema 3.2 que cualquier otro grupo de Lie con álgebra abeliana, que no sea simplemente conexo, es cubierto por $\mathbb{R}^n$; adicionalmente, sabemos que este grupo de Lie debe ser idéntico a $\mathbb{R}^n$ en alguna vecindad de e . Puesto que $\mathbb{R}^n$ es abeliano, entonces lo es cualquier otro grupo de Lie con álgebra de Lie abeliana:

Sea G un grupo que tiene álgebra de Lie abeliana $\Rightarrow G$ es cubierto por $\mathbb{R}^n$ via un homomorfismo:

$$\Pi(x + y) = \Pi(x) * \Pi(y)$$

⁷Nunca hablamos de métrica sobre $SU(2)$, es por esta razón que podemos decir que los valores correspondientes a t y $t + 2\pi$ están diametralmente opuestos. Sólo nos interesa la topología global.

y como R^n es abeliano:

$$\Pi(x + y) = \Pi(y + x) \Rightarrow \Pi(x) * \Pi(y) = \Pi(y) * \Pi(x) \Rightarrow G \text{ es abeliano}$$

En consecuencia, el álgebra de Lie abeliana es el álgebra de Lie de un grupo abeliano.

§3.9 Representaciones de un grupo de Lie.

Considérese a un grupo de Lie G y a una variedad M . Una *representación de G sobre M* [6] es un mapeo diferenciable $G \times M \rightarrow M$, donde $(g, x) \mapsto \Gamma_g x$, tal que para cada elemento g del grupo $\Gamma_g : M \rightarrow M$ es un difeomorfismo de M que satisface lo siguiente: $\Gamma_e = id_M$ y $\Gamma_g \circ \Gamma_{g'} = \Gamma_{gg'}$. Es decir, una representación Γ es un mapeo de G a $Diff(M)$ ⁸.

Cabe señalar que a una representación suele llamársele también realización, acción o acción por la izquierda. Una acción por la derecha se construye de igual manera que la izquierda, salvo que $\Gamma_g \circ \Gamma_{g'} = \Gamma_{g'g}$; en consecuencia, una acción derecha Γ_g no es una representación en el contexto de nuestra definición inicial.

Para cualquier $x \in M$ definamos el mapeo $b_x : G \rightarrow M$, mediante la regla $g \mapsto b_x g = \Gamma_{g^{-1}} x$. La órbita O de $x \in M$ esta dada por el conjunto $O(x) \equiv \{\Gamma_g x = x' \in M \mid g \in G\}$. Notemos entonces que $b_x(G) = \{b_x g = \Gamma_{g^{-1}} x \mid g \in G\} = O(x)$.

Diremos que una representación es *transitiva* si M consiste de sólo una órbita (ie: para cualesquiera x y y en $M \exists g \in G$ tal que $y = \Gamma_g x$).

Una representación es *fiel* si el mapeo $g \mapsto \Gamma_g$ es inyectivo.

Afirmación 3.1 Si b_x es inyectivo para alguna (p.a.) $x \in M \Rightarrow$ la representación es fiel.

Demostración: Supongamos que existen h y g en G tales que $\Gamma_{h^{-1}} = \Gamma_{g^{-1}} \Rightarrow \Gamma_{h^{-1}x_0} = \Gamma_{g^{-1}x_0}$.

Sea x_0 donde el mapeo b_x es inyectivo. Entonces, que $\Gamma_{h^{-1}x_0} = \Gamma_{g^{-1}x_0} \Rightarrow h^{-1} = g^{-1}$. Por lo tanto, $\Gamma_{h^{-1}} = \Gamma_{g^{-1}} \Rightarrow h^{-1} = g^{-1} \Rightarrow g \mapsto \Gamma_g$ es inyectivo y, en consecuencia, la representación es fiel. $\square$

Se dice que una representación es *libre* si Γ_e es el único que fija a un punto: si para algún $x \in M$ resulta que $\Gamma_g x = x \Rightarrow g = e \in G$.

⁸Por $Diff(M)$ denotamos al grupo de difeomorfismos de M en si misma.

Afirmación 3.2 Una representación es libre sii b_x es inyectiva $\forall x \in M$.

Demostración: $\Rightarrow$] Supongamos que existe $x' \in M$ y $h, g \in G$ tales que $b_{x'}h = b_{x'}g$ y $h \neq g$ (ie: $b_{x'}$ no es inyectiva en x').

Como $b_{x'}h = b_{x'}g \Rightarrow \Gamma_{h^{-1}}x' = \Gamma_{g^{-1}}x' \Rightarrow (\Gamma_h \circ \Gamma_{h^{-1}})x' = (\Gamma_h \circ \Gamma_{g^{-1}})x' \Rightarrow x' = \Gamma_{hg^{-1}}x'$. Puesto que la representación es libre, tenemos que $hg^{-1} = e \Rightarrow h = g$!

La contradicción proviene de suponer que b_x no es inyectiva en x' . Por lo tanto, b_x es inyectiva $\forall x \in M$.

$\Leftarrow$] Que b_x sea inyectiva $\forall x \in M \Rightarrow$ si $\Gamma_{h^{-1}}x = \Gamma_{g^{-1}}x \Rightarrow h = g$. Si p.a. $x \in M$ ocurre que $x = \Gamma_g x$, tenemos entonces que $\Gamma_g x = \Gamma_e x \Rightarrow b_x g^{-1} = b_x e^{-1} \Rightarrow g^{-1} = e^{-1} \Rightarrow g = e \in G$. En consecuencia $x = \Gamma_g x \Rightarrow g = e \in G$, por lo tanto la representación es libre. $\square$

Es posible derivar un tercer resultado observando lo siguiente: Como la representación libre implica que el mapeo b_x es inyectivo $\forall x \in M$, en particular lo es entonces para $x^* \in M$ y, en consecuencia, la representación es fiel. Por lo tanto, si la representación es libre, entonces también es fiel.

Dado $x \in M$ definimos al grupo de isotropía como sigue:

$$I(x) \equiv \{g \in G \mid \Gamma_g x = x\}$$

Afirmación 3.3 Una representación es libre sii todos los grupos de isotropía son triviales ($I(x) = \{e\}$).

Demostración: $\Rightarrow$] Si la representación es libre entonces, dada cualquier x , tenemos que si $x = \Gamma_g x \Rightarrow g = e \Rightarrow \{g \in G \mid \Gamma_g x = x\} = \{e\} \Rightarrow I(x) = \{e\}$.

$\Leftarrow$] Si todos los grupos de isotropía son triviales $\Rightarrow \{g \mid \Gamma_g x = x\} = \{e\}, \forall x \in M$. Entonces $\Gamma_g x = x \Rightarrow g = e \in G$ y, por lo tanto, la representación es libre. $\square$

Consideremos dos representaciones de $SO(3)$ como ejemplos que ilustren lo expuesto hasta el momento:

- La representación de este grupo de Lie sobre $\mathbb{R}^3$ no es transitiva, sus órbitas son el origen y esferas de diferente radio; es fiel y no libre. Sus grupos de isotropía son $SO(2)$ $\forall x \neq 0$ y $SO(3)$ para $x = 0$ [6].

- La representación de $SO(3)$ sobre S^2 es, por otro lado, transitiva, fiel y no libre. Todos sus grupos de isotropía son $SO(2)$ [6].

Procedamos ahora a presentar algunas de las representaciones mas importantes de G en si mismo (ie: $G = M$):

Dentro de el contexto formal de las representaciones, una traslación izquierda es un mapeo $G \times G \rightarrow G$, donde $(g, h) \mapsto l_g h = gh$ (g actúa por la izquierda sobre h). Esta representación es transitiva y libre, en efecto:

Para probar la transitividad consideremos dos elementos de G arbitrarios, a saber g_1 y g_2 . Puesto que G es un grupo, entonces existe $g \in G$ tal que $g_1 = gg_2 \Rightarrow g_1 = l_g g_2$ y, por lo tanto, la representación es transitiva.

La representación es libre ya que si suponemos que $h = l_g h = gh \Rightarrow h = gh \Rightarrow g = e$.

Adicionalmente tenemos que por ser libre es, automáticamente, fiel.

Otra representación de G en sí mismo es la conocida bajo el nombre de *automorfismo interno*. Esta representación es un mapeo $G \times G \rightarrow G$, donde $(g, h) \mapsto \text{Aut}_g h = ghg^{-1}$. Notemos que en esta representación el grupo de isotropía de la identidad e es el mismo G y que la representación, por tanto, no es libre.

Diremos que una representación Γ es *lineal* si la variedad M sobre la cual actúa es un espacio vectorial V y si todos los difeomorfismos Γ_g son lineales. Observemos que ninguna representación lineal es transitiva:

En efecto. Si la representación es lineal entonces

$$\Gamma_g(x) = \Gamma_g(x + 0) = \Gamma_g(x) + \Gamma_g(0) \Rightarrow \Gamma_g(0) = 0$$

Por lo tanto, la órbita del 0 es el 0 y, entonces, no es transitiva la representación.

Se dice que una representación lineal es *irreducible* si V no contiene a algún subespacio invariante W , donde W es no trivial: $W \neq V$ y $W \neq \{0\}$. De manera tal que tendríamos, adicionalmente, que la representación de $SO(3)$ sobre $\mathbb{R}^3$ es también lineal e irreducible.

Una representación más surge a partir del automorfismo interno Aut_g . Dado que Aut_g mapea a en e , entonces la imagen bajo Aut_g de una curva que pasa a través de e es también una curva que pasa por la identidad e ; la imagen de la curva es idéntica a la misma si el grupo G es abeliano, pero como ello no tiene porque ser así, entonces las curvas, por lo general, son distintas. De tal manera, el mapeo Aut_g induce un mapeo de cualquier vector tangente en $T_e G$ en otro también en $T_e G$. Este mapeo inducido es llamado *mapeo adjunto* de $T_e G$ y es denotado como Ad_g .

Notemos que si la curva inicial es un subgrupo 1-paramétrico, digamos e^{tX} (con $X \in T_e G$), entonces $\text{Aut}_g e^{tX}$ es también un subgrupo 1-paramétrico:

Sea $h = e^{tsX}$ y $f = e^{tX}$. Como $g(fh)g^{-1} = (gfg^{-1})(ghg^{-1})$, entonces tenemos que:

$$\text{Aut}_g(fh) = \text{Aut}_g f \text{Aut}_g h$$

$$\Rightarrow \text{Aut}_g e^{(t_2+t_1)X} = \text{Aut}_g e^{t_2X} \text{Aut}_g e^{t_1X}$$

entonces:

$$\text{Aut}_g e^{tX} = e^{t \text{Ad}_g X}$$

$$\Rightarrow \frac{d}{dt} (\text{Aut}_g e^{tX})|_{t=0} = \frac{d}{dt} (e^{t \text{Ad}_g X})|_{t=0} \Rightarrow \frac{d}{dt} (g e^{tX} g^{-1})|_{t=0} = \text{Ad}_g X$$

por lo tanto:

$$\text{Ad}_g X = gXg^{-1}$$

Nótese que la representación Ad_g es una representación lineal en el álgebra de Lie de G .

Capítulo 4

Conexiones.

¿Por qué estudiar conexiones? Al menos como será abordada en este trabajo, existe una razón obvia: La conexión, como observaremos en este capítulo, está estrechamente vinculada al concepto de curvatura, éste, a su vez, es un "ingrediente" fundamental e insoslayable en la Teoría General de la Relatividad; en consecuencia, el estudio de estas estructuras matemáticas justifica y proporciona el desarrollo formal, desde una perspectiva meramente geométrica, de resultados ligados a la curvatura en el ámbito de esta teoría.

Ahora bien, ¿por qué estudiarlas en este trabajo? En este capítulo notaremos, a grandes rasgos, la íntima relación entre el *transporte paralelo* y la conexión. Veremos que esta relación conlleva a la asociación de la conexión con ciertas formas diferenciales, definidas en un haz principal, cuyas transformaciones de norma están caracterizadas como cambios de sección en el haz. Puesto que más adelante (siguiente capítulo) dichas formas tomarán un papel central en la construcción de densidades lagrangianas para modelos σ no lineales, como podrá percibirse el lector, es entonces clara la necesidad de una discusión previa.

§4.1 Sobre la conexión y la curvatura.

Consideremos una variedad M con curvatura no nula, intuitivamente queda claro que entonces los espacios tangentes $T_p M$ y $T_{p'} M$ sobre M cambian conforme nos movemos de p a p' . Una *conexión Riemanniana* es, esencialmente, aquella estructura que nos permite comparar los espacios tangentes en un par de puntos infinitesimalmente separados [5]; i.e: la estructura que cuantifica el *cambio*, debido a la curvatura, entre los espacios tangentes. La conexión,

en este contexto y de manera formal, es la regla que usualmente se conoce como *transporte paralelo o traslación en M* [5]; en consecuencia, para puntualizar matemáticamente la idea intuitiva de conexión es necesario definir primero al transporte paralelo:

Sea γ cualquier curva suave y simple que conecte a los puntos p y p' . Sea $\vec{x} \in T_p M$ tal que $\vec{x}$ es tangente a γ en p . Diremos que $\vec{x}$ es *transportado paralelamente a lo largo de γ* si $\vec{x}$ es llevado de p a p' de manera tal que siempre permanece paralelo a sí mismo [5]. Si λ es el parámetro de la curva γ , entonces la derivada covariante de $\vec{x}$ es la razón de cambio de $\vec{x}$ respecto a λ . La derivada covariante, en general, difiere de la derivada parcial ordinaria; la cantidad que mide tal discrepancia es, justamente y de nueva cuenta, la conexión. Cabe señalar que la citada discrepancia tiene como consecuencia, en general, la no conmutatividad de la diferenciación covariante [5]; puesto que la discrepancia es inducida por la curvatura, es coherente afirmar entonces que la curvatura es la medida de la no conmutatividad de la diferenciación covariante.

El transporte paralelo nos permite, además de comparar vectores en distintos puntos, definir el concepto de curvas geodésicas: Diremos que una curva $\gamma \in M$ es una geodésica si sus vectores tangentes son paralelos a lo largo de γ .

Formalicemos el escenario matemático de nuestra discusión. Sea el haz fibrado principal P con fibra G (donde G es un grupo de Lie) sobre la variedad M ; i.e. el haz $(E = P, F = G, B = M, \Pi)$. Del hecho de que G y M sean estructuras diferenciables se sigue que P es una variedad diferenciable y, como consecuencia de ello, es posible considerar al haz tangente TP y al haz cotangente T^*P de P .

Concentremos nuestra atención en TP y cualquier $q \in P$. Evidentemente en q hay un espacio tangente $T_q P$, los vectores tangentes a la fibra G que pasa a través de q constituyen un subespacio del espacio tangente $T_q P$ que habremos de llamar, en lo que sigue, *subespacio vertical* en q y denotar por $V_q P$. Definamos al *subespacio horizontal*¹ $H_q P$ de tal manera que:

$$T_q P = V_q P \oplus H_q P$$

y bajo la transformación $q \mapsto q' = qg$, con $g \in G$, suceda que $H_{q'} = H_{qg} = R_g H_q$, donde $R_g H_q$ denota la acción sobre H_q del mapeo lineal inducido sobre el espacio tangente $T_q P$ por el mapeo $q \mapsto qg$, traslación derecha, de los elementos de la fibra [5]. Es decir, podemos obtener un espacio horizontal a partir de otro a través de la acción por la derecha del grupo

¹ Hacemos notar desde este momento que para construir al subespacio horizontal es necesaria una conexión, como veremos más adelante

G.

Definamos ahora una noción general de transporte paralelo y notemos la existencia de una conexión asociada con éste:

Sea γ cualquier curva en M que pasa por p y p' . Con el transporte paralelo definimos transportes o mapeos de la fibra $\Pi^{-1}(p)$ sobre la fibra $\Pi^{-1}(p')$. La regla para el transporte paralelo de fibras a lo largo de la curva γ es la siguiente [5]:

Sea $\tilde{\gamma}$ una curva en P tal que $\Pi(\tilde{\gamma}) = \gamma$ y que sus vectores tangentes en cualquier punto pertenezcan al subespacio horizontal. Nótese que dada γ y un punto $q_0 \in P$ hay solamente una de tales $\tilde{\gamma}$ con punto inicial q_0 (§4.3); las curvas $\tilde{\gamma}$ son llamadas *levantamiento horizontal* de γ .

Para llevar a cabo el transporte paralelo de fibras a lo largo de γ de un punto $p \in M$ a un punto $p' \in M$ tomaremos cada $q \in \Pi^{-1}(p)$, construiremos el levantamiento único $\tilde{\gamma}$ de γ con punto inicial q y mapeamos q a $q' \in \Pi^{-1}(p')$, donde q' es el punto de $\tilde{\gamma}$ que "vive" sobre p' . Asociada a este transporte hay una conexión, esta surge a partir de la necesidad de especificar primero al subespacio horizontal en la construcción previamente descrita. Veamos pues, cuál es la conexión que nos permite llevar a cabo el transporte paralelo.

Sean (x, g) , $x \in M$ y $g \in G$, las coordenadas locales del haz P . Sea $w \in T^*P$ la siguiente 1-forma²:

$$w = g^{-1}dg + g^{-1}Ag$$

donde

$$A = A_a^\alpha(x) \left(\frac{\lambda_a}{2i} \right) dx^\alpha$$

y $\{\frac{\lambda_a}{2i}; a = 1, \dots, \dim G = [G]\}$ son la base del álgebra de Lie de G que satisface la siguiente regla de conmutación:

$$\left[\frac{\lambda_a}{2i}, \frac{\lambda_b}{2i} \right] = f_{ab}^c \frac{\lambda_c}{2i}$$

con f_{ab}^c las constantes de estructura del álgebra de Lie de G .

Puesto que $T_qP = V_qP \oplus H_qP$ y $\dim T_qP = [G] + n$, entonces $\dim V_qP = [G]$ y $\dim H_qP = n$. El vector base del subespacio horizontal no necesariamente tiene componentes nulas en la dirección vertical, por consiguiente podemos expresar al mismo como sigue:

²En nuestra notación las 1-formas se distinguen de los vectores por llevar consigo una tilde, por no prestarse a confusión en esta ocasión las suprimiremos para simplificar notación.

$$\frac{\partial}{\partial x^\mu} + C_{\mu ij} \frac{\partial}{\partial g^{ij}} ; \mu = 1, \dots, n \text{ y } i, j = 1, \dots, [G]$$

en consecuencia, el subespacio horizontal queda especificado una vez que sean determinadas las componentes $C_{\mu ij}$. Especifiquemos a las citadas componentes exigiendo que para cualquier $\vec{x} \in T_q P$ suceda que $\langle w, \vec{x} \rangle = 0$ si $\vec{x} \in H_q P$.

Sea $\vec{x} \in T_q P$, entonces en las coordenadas locales para P tenemos que:

$$\vec{x} = \alpha_{ij} \frac{\partial}{\partial g^{ij}} + b_\mu \left(\frac{\partial}{\partial x^\mu} + C_{\mu ij} \frac{\partial}{\partial g^{ij}} \right)$$

Como w en términos de este sistema coordenado para el haz principal P es simplemente³:

$$w = g_{ij}^{-1} dg_{ij} + g_{ij}^{-1} A_\mu^a(x) \frac{\lambda_{aik}}{2i} dx^\mu g_{kj}$$

entonces

$$\langle w, \vec{x} \rangle = g_{ij}^{-1} (\alpha_{ij} + b_\mu C_{\mu ij}) + g_{ij}^{-1} A_\mu^a(x) \frac{\lambda_{aik}}{2i} g_{kj} b_\mu$$

Si $\vec{x} \in H_q P \Rightarrow \alpha_{ij} = 0$ y $\langle w, \vec{x} \rangle = 0$; entonces:

$$g_{ij}^{-1} b_\mu C_{\mu ij} + g_{ij}^{-1} A_\mu^a(x) \frac{\lambda_{aik}}{2i} g_{kj} b_\mu = 0 \quad \forall b_\mu$$

$$\Rightarrow g_{ij}^{-1} C_{\mu ij} + g_{ij}^{-1} A_\mu^a(x) \frac{\lambda_{aik}}{2i} g_{kj} = 0$$

por lo tanto:

$$C_{\mu ij} = -A_\mu^a(x) \frac{\lambda_{aik}}{2i} g_{kj}$$

con lo cual quedan determinadas las componentes $C_{\mu ij}$ y, por ende, el subespacio horizontal $H_q P$. Puede verificarse que el requerimiento $H_{gg} = R_g H_g$ se satisface; por lo tanto, el subespacio horizontal está bien especificado. Es decir, el subespacio que anula a la 1-forma w cumple con los requerimientos de ser un subespacio horizontal.

³En este caso se abusa en la notación de índices, no obstante el lector debe tener claro que índices repetidos siguen significando suma sobre los mismos.

Usualmente la cantidad $A^a_\mu(x)$ se le denomina *conexión asociada con el subespacio horizontal*, en tanto que a la forma A se le llama con frecuencia *forma de conexión* [5]. Ahora podemos fijar la idea general de mejor manera: El transporte paralelo de fibras, que nos permite comparar fibras en diferentes puntos, está definido en términos del levantamiento horizontal que, a su vez, sólo puede llevarse a cabo especificando al subespacio horizontal. La introducción de la 1-forma w hace posible tal especificación a través de determinadas cantidades que habremos de llamar, de manera genérica, conexiones. De tal manera es claro que para llevar a cabo un transporte paralelo de fibras es necesario contar con una conexión que permita especificar al subespacio horizontal.

§4.2 La forma de conexión y el potencial de norma.

Supongamos que llevamos a cabo un cambio de coordenadas en el haz P , puesto que la forma A es una componente de la forma w su ley de transformación es inducida por esta última. Hagamos un cambio de coordenadas exclusivamente en la fibra (i.e.: $x = x'$): $g' = hg$, con $h \in G$.

La invariancia de w bajo tal transformación significa simplemente que:

$$g^{-1}dg + g^{-1}Ag = (g')^{-1}dg' + (g')^{-1}A'g'$$

como $dg' = (dh)g + h(dg)$, entonces:

$$g^{-1}dg + g^{-1}Ag = g^{-1}h^{-1}(dh)g + g^{-1}dg + g^{-1}h^{-1}A'hg$$

$$\Rightarrow g^{-1}Ag = g^{-1}h^{-1}(dh)g + g^{-1}h^{-1}A'hg$$

$$\Rightarrow A = h^{-1}(dh) + h^{-1}A'h$$

dado que $(dh)h^{-1} + h dh^{-1} = d(hh^{-1}) = 0$, tenemos que:

$$A' = h dh^{-1} + h A h^{-1} \quad (4.1)$$

La expresión (4.1) es la ley de transformación de la forma de conexión A . A esta ley de transformación se le denomina, en el ámbito de las teorías electromagnética y de Yang-Mills, ley de transformación de norma [5]: Una transformación de norma es un cambio en

las coordenadas de la fibra de un haz principal. Debido a este símil, cabe señalar que a la conexión $A_\mu^a(x)$ también se le conoce como *potencial de norma*.

§4.3 Transporte paralelo, derivada covariante y curvatura.

Como vimos con antelación nuestra definición de transporte paralelo de fibras esta formulada a través del levantamiento horizontal $\tilde{\gamma}(t) \in P$ que existe para cualquier curva $\gamma(t) \in M$. Recalquemos que para realizar dicho transporte es necesario especificar al subespacio horizontal. La especificación consiste en notar que el subespacio anulador de la 1-forma w cumple con las exigencias para ser horizontal. La especificación del citado subespacio horizontal nos permite definir también el transporte paralelo de un vector a lo largo de $\tilde{\gamma}(t)$: diremos que un vector $\tilde{x} \in T_{\tilde{\gamma}(t)}P$ es transportado paralelamente a lo largo de $\tilde{\gamma}(t)$ si este permanece siempre en el subespacio horizontal conforme es transportado a lo largo de $\tilde{\gamma}(t)$. Anteriormente mencionamos que el levantamiento horizontal es único dada una curva $\gamma(t) \in M$ y un punto en una fibra sobre tal curva. Observemos que esto es efectivamente así:

Sea (x^μ, g) un sistema local de coordenadas para el haz P . Entonces $\gamma(t) = x^\mu(t)$ y $\tilde{\gamma}(t) = (x^\mu(t), g(t))$, por lo tanto la tangente a $\tilde{\gamma}(t)$ es el vector $\frac{d}{dt} = \dot{x}^\mu \frac{\partial}{\partial x^\mu} + \dot{g} \frac{\partial}{\partial g}$.

Como $\frac{d}{dt} \in HP$, entonces:

$$\dot{x}^\mu(t) \frac{\partial}{\partial x^\mu} + \dot{g}(t) \frac{\partial}{\partial g} = b_\mu \left(\frac{\partial}{\partial x^\mu} - A_\mu^a \frac{\lambda_a}{2i} g \frac{\partial}{\partial g} \right) \quad \text{p.a. } b_\mu$$

En consecuencia:

$$\dot{x}^\mu(t) = b_\mu$$

$$\dot{g}(t) = -b_\mu A_\mu^a \frac{\lambda_a}{2i} g$$

entonces:

$$\dot{g}(t) + \dot{x}^\mu(t) A_\mu^a \frac{\lambda_a}{2i} g = 0 \quad (4.2)$$

A la ecuación (4.2) se le conoce como *ecuación de transporte paralelo*, nótese que esta es una ecuación diferencial de primer orden para $g(t)$ y, en al menos una vecindad del punto inicial, siempre tiene una solución y esta es única. La solución de (4.2) nos da el transporte paralelo para los vectores tangentes a $\tilde{\gamma}(t)$.

Denotemos al operador de derivada covariante por D_μ , recordando la definición dada para esta derivada en la sección §4.1 notamos entonces que $x^\mu D_\mu$ es la razón de cambio respecto a t a lo largo del levantamiento horizontal $\tilde{\gamma}(t)$. Es decir:

$$D_\mu = \frac{\partial}{\partial x^\mu} - A_\mu^a \frac{\lambda_a}{2i} g \frac{\partial}{\partial g}$$

Obsérvese que si la conexión $A_\mu^a(x)$ es nula, entonces la derivada covariante D_μ es simplemente la derivada parcial ordinaria $\frac{\partial}{\partial x^\mu} \equiv \partial_\mu$.

Una vez definido el operador D_μ es posible hallar la expresión para la operación $[D_\mu, D_\nu]$:

$$\begin{aligned} [D_\mu, D_\nu] &= [\partial_\mu - A_\mu^a \frac{\lambda_a}{2i} g \frac{\partial}{\partial g}, \partial_\nu - A_\nu^b \frac{\lambda_b}{2i} g \frac{\partial}{\partial g}] \\ &= -\partial_\mu A_\nu^b \frac{\lambda_b}{2i} g \frac{\partial}{\partial g} + \partial_\nu A_\mu^a \frac{\lambda_a}{2i} g \frac{\partial}{\partial g} + A_\mu^a A_\nu^b (\frac{\lambda_a}{2i} g \frac{\partial}{\partial g})(\frac{\lambda_b}{2i} g \frac{\partial}{\partial g}) - A_\nu^b A_\mu^a (\frac{\lambda_b}{2i} g \frac{\partial}{\partial g})(\frac{\lambda_a}{2i} g \frac{\partial}{\partial g}) \end{aligned}$$

Sea $R_k = \frac{\lambda_k}{2i} g \frac{\partial}{\partial g} = \frac{\lambda_k}{2i} g_{JI} \frac{\partial}{\partial g^I}$, puede entonces verificarse que $[R_a, R_b] = -f_{ab}^c R_c$. En consecuencia:

$$[D_\mu, D_\nu] = (-\partial_\mu A_\nu^c + \partial_\nu A_\mu^c) R_c + A_\mu^a A_\nu^b [R_a R_b, R_b R_a]$$

$$\Rightarrow [D_\mu, D_\nu] = (-\partial_\mu A_\nu^c + \partial_\nu A_\mu^c) R_c - A_\mu^a A_\nu^b f_{ab}^c R_c$$

sea $F_{\mu\nu}^c = \partial_\mu A_\nu^c - \partial_\nu A_\mu^c + f_{ab}^c A_\mu^a A_\nu^b$, entonces:

$$[D_\mu, D_\nu] = -F_{\mu\nu}^c R_c$$

Al tensor $F_{\mu\nu}^c$ se le conoce como *tensor de curvatura* [5], esto se debe a que la curvatura es la medida de la no conmutatividad de la diferenciación covariante (§4.1).

Sean $F_{\mu\nu}^c \frac{\lambda_c}{2i}$ las componentes de la 2-forma F de curvatura:

$$F = \frac{1}{2} F_{\mu\nu}^c \frac{\lambda_c}{2i} dx^\mu \wedge dx^\nu$$

$$\begin{aligned}\Rightarrow F &= \frac{1}{2}(\partial_\mu A_\nu^c - \partial_\nu A_\mu^c) \frac{\lambda_c}{2i} dx^\mu \wedge dx^\nu + \frac{1}{2}(f_{ab}^c A_\mu^a A_\nu^b \frac{\lambda_c}{2i}) dx^\mu \wedge dx^\nu \\ \Rightarrow F &= \partial_\mu A_\nu^c \frac{\lambda_c}{2i} dx^\mu \wedge dx^\nu + \frac{1}{2}[A_\mu^a \frac{\lambda_a}{2i}, A_\nu^b \frac{\lambda_b}{2i}] dx^\mu \wedge dx^\nu\end{aligned}$$

Por lo tanto:

$$F = dA + A \wedge A \equiv DA$$

donde DA denota a la *derivada exterior covariante* de A .

Nótese que la 2-forma de curvatura puede escribirse también en términos de g y w :

$$F = g(dw + w \wedge w)g^{-1} = g\Omega g^{-1} \text{ donde } \Omega = dw + w \wedge w$$

Es claro que entonces, bajo una transformación de norma $g \mapsto g' = hg$, la 2-forma de curvatura cumple la siguiente ley de transformación:

$$F' = hFh^{-1}$$

Notemos que el par (Ω, w) se relaciona de igual manera que el par (F, A) y aunque F y A son formas en la variedad base M mientras que Ω y w son formas en el haz P , es posible pasar de unas a otras. En efecto:

Sea $K : U \rightarrow V$ suave, donde U y V son abiertos de $\mathbb{R}^n$ y $\mathbb{R}^m$ respectivamente. Sea $\beta : [a, b] \rightarrow U$ una curva en U , $K \circ \beta$ es entonces una curva en V , evidentemente:

$$\frac{d\beta^i}{dt} \frac{\partial}{\partial x^i} |_{\beta(t)} \in T_{\beta(t)}U ; i = 1, \dots, n$$

y

$$\frac{d}{dt}(K \circ \beta)^j \frac{\partial}{\partial y^j} |_{K(\beta(t))} \in T_{K(\beta(t))}V ; j = 1, \dots, m$$

como

$$\frac{d}{dt}(K \circ \beta)^j = K^j_{,i}(\beta(t)) \frac{d\beta^i}{dt}$$

ESTA TESIS NO DEBE SALIR DE LA BIBLIOTECA

79

si $T_{\beta(v)}K$ es tal que:

$$T_{\beta(v)}U \xrightarrow{T_{\beta(v)}K} T_{K(\beta(v))}V$$

entonces tenemos que:

$$(T_x K)(v) = v^i K^j_{,i} \frac{\partial}{\partial y^j}(K(x))$$

para $v = v^i \frac{\partial}{\partial x^i} \in T_x U$.

El mapeo $T_x K$ es llamado *mapeo tangente* y generalmente es denotado por K_* . [6]. Sea K^* el mapeo que va de las p -formas en V a las p -formas en U , $K^*: \Lambda^p V \rightarrow \Lambda^p U$, a esta operación se le conoce bajo el nombre de pullback de formas diferenciales [6].

Para $\vartheta \in \Lambda^p V$, $K^*\vartheta$ esta dada por

$$(K^*\vartheta)_x(v_1, \dots, v_p) = \vartheta_{K(x)}((K_*v_1), \dots, (K_*v_p))$$

con $v_1, \dots, v_p \in T_x U$.

El pullback de formas diferenciales tiene las siguientes propiedades:

- (i).- K^* es lineal.
- (ii).- $K^*(\alpha \wedge \beta) = (K^*\alpha) \wedge (K^*\beta)$.
- (iii).- $(S \circ K)^* = K^* \circ S^*$.

Ahora bien, debido a la trivialidad local de P podemos dar una sección local en esa región, sea $s(x)$ tal sección:

$$s(x) : U_\alpha \rightarrow P \text{ con } U_\alpha \in M$$

Dada $s(x)$ construimos $s^*(x)$ como el pullback de formas sobre P a formas sobre U_α . El mapeo $s^*\Omega = g\Omega g^{-1}$ cumple con las propiedades del pullback $\Lambda^2 P \rightarrow \Lambda^2 U_\alpha$; por lo tanto $s^*\Omega = F$ y $s^*w = A$. i.e: vía una sección es posible pasar de formas en P a formas en M .

Sea U_β tal que $U_\alpha \cap U_\beta \neq \emptyset$ y s' una sección local en este abierto. Entonces $s' = hs$ en la intersección y tanto las conexiones como las curvaturas estarán sometidas a una transformación de norma bajo la acción de h . Es por esta razón que se dice que las transformaciones de norma están también caracterizadas como cambios locales de secciones en un haz principal P ; notemos entonces que fijar una norma, en este contexto, es fijar una sección.

Para finalizar este apartado derivemos las llamadas *identidades de Bianchi*. Consideremos la siguiente identidad nula:

$$dA \wedge A - A \wedge dA - dA \wedge A + A \wedge dA + A \wedge A \wedge A - A \wedge A \wedge A = 0$$

$$\Rightarrow dA \wedge A - A \wedge dA + A \wedge (dA + A \wedge A) - (dA + A \wedge A) \wedge A = 0$$

Como $dF = d(A \wedge A) = dA \wedge A - A \wedge dA$, entonces:

$$DF \equiv dF + A \wedge F - F \wedge A = 0 \quad (4.3)$$

La identidad (4.3) es, justamente, la *identidad de Bianchi*, esta identidad es muy conocida en el contexto de la Teoría General de la Relatividad en donde la curvatura toma un rol central.

§4.4 Conexión en el haz tangente.

Para dar la conexión en el haz tangente es necesario definir primero al *haz de referencia*:

Sea FM el haz que tiene por base a una variedad diferenciable M y como fibra sobre cada $p \in M$ a la colección de todas las bases ordenadas del espacio tangente $T_p M$. A este haz se le denomina *haz de referencia* [5]. Puesto que la representación de $GL(n, \mathbb{R})$ es transitiva en las bases, entonces $F \simeq GL(n, \mathbb{R})$. Como el grupo de estructura es $GL(n, \mathbb{R})$, entonces el haz FM es un haz principal. Dado que el grupo de estructura del haz TM es también $GL(n, \mathbb{R})$, entonces FM es el haz principal asociado a TM .

Sea $\tilde{x} \in T_x M$ fijo y $R \in GL(n, \mathbb{R})$. Entonces $R\tilde{x}$ es, de nueva cuenta, un vector en $T_x M$.

Sea $\tilde{\gamma}(t)$ un levantamiento horizontal en el haz FM . Dicho levantamiento horizontal induce uno propio en el haz TM según la siguiente relación [5]:

$$\gamma_{TM}(t) = \tilde{\gamma}(t)\tilde{x}$$

De tal manera obtenemos que conforme uno circula en $\tilde{\gamma}(t) \in FM$ también lo hace en $\gamma_{TM}(t)$.

Halleemos el vector tangente a $\gamma_{TM}(t)$:

$$\frac{d}{dt}\gamma_{TM}(t) = \frac{d}{dt}(\tilde{\gamma}(t)\tilde{x})$$

$$\Rightarrow \frac{d}{dt} \gamma_{TM}(t) = \dot{x}^\mu(t) D_\mu(\dot{\gamma}(t) \bar{x})$$

dónde $\dot{x}^\mu(t)$ son las componentes del vector tangente a la curva $\gamma(t) \in M$. Sea $\bar{Y}$ dicho vector tangente: $\bar{Y} = \dot{x}^\mu \frac{\partial}{\partial x^\mu}$.

Evidentemente $\gamma_{TM}(t) = \dot{\gamma}(t) \bar{x}$ es para cada t un cierto vector tangente $\bar{Z} \in T_{\gamma(t)} M$. En el haz TM , la derivada covariante de $\bar{Z}$ en la dirección de $\bar{Y}$ (o respecto a $\bar{Y}$) es denotada por $\nabla_{\bar{Y}} \bar{Z}$ y definida como:

$$\nabla_{\bar{Y}} \bar{Z} \equiv \frac{d}{dt} \gamma_{TM}(t) = \dot{x}^\mu(t) D_\mu(\dot{\gamma}(t) \bar{x})$$

Entonces, de acuerdo con nuestra definición de transporte paralelo de fibras, la cantidad $\nabla_{\bar{Y}} \bar{Z}$ es la razón de cambio de $\bar{Z}$ en la dirección $\bar{Y}$ bajo un transporte paralelo.

De la definición se siguen las siguientes propiedades para la derivada covariante:

$$(i). - \nabla_{\bar{X} + \bar{Y}}(\bar{Z}) = \nabla_{\bar{X}}(\bar{Z}) + \nabla_{\bar{Y}}(\bar{Z}).$$

$$(ii). - \nabla_{\bar{X}}(\bar{Y} + \bar{Z}) = \nabla_{\bar{X}}(\bar{Y}) + \nabla_{\bar{X}}(\bar{Z}).$$

$$(iii). - \nabla_{\bar{X}}(f\bar{Y}) = f\nabla_{\bar{X}}(\bar{Y}) + (\bar{X}f)\bar{Y}.$$

$$(iv). - \nabla_{\bar{X}}(\bar{Y}) = f\nabla_{\bar{X}}(\bar{Y}). \text{ donde } f \in C^\infty(M) \text{ y } \bar{X}, \bar{Y}, \bar{Z} \text{ son campos vectoriales sobre } M.$$

Los puntos $g(t)$ sobre el levantamiento horizontal $\dot{\gamma}(t)$ pertenecen a $Gl(n, \mathbb{R})$, de modo que es posible representarlos por matrices de $n \times n$ cuyas coordenadas locales son X^μ_ν ; $\mu, \nu = 1, \dots, n$. Dado que la combinación $A^\mu_{\lambda\alpha} \frac{\partial}{\partial x^\mu} \in Gl(n, \mathbb{R})$, entonces el operador D_μ en términos de las coordenadas locales es simplemente el siguiente:

$$D_\mu = \frac{\partial}{\partial x^\mu} - A^\nu_{\mu\lambda} X^\lambda_\sigma \frac{\partial}{\partial X^\sigma_\nu}$$

donde $A^\nu_{\mu\lambda}$ son funciones que expresan la combinación $A^\alpha_{\lambda\alpha} \frac{\partial}{\partial x^\mu}$ en coordenadas locales.

En consecuencia:

$$\nabla_{\bar{Y}} \bar{Z} = \dot{x}^\lambda (\partial_\lambda - A^\rho_{\lambda\sigma} X^\sigma_\tau \frac{\partial}{\partial X^\tau_\rho}) X^\mu_\alpha \dot{\gamma}^\alpha \partial_\mu$$

$$\Rightarrow \nabla_{\bar{Y}} \bar{Z} = \dot{x}^\lambda b^\mu{}_{,\lambda} \partial_\mu - \dot{x}^\lambda A^\mu_{\lambda\alpha} b^\alpha \partial_\mu$$

donde $b^\lambda = X_\sigma^\lambda a^\sigma$ son las componentes de $\vec{Z}$ en la base coordenada $\{\frac{\partial}{\partial x^\sigma}\}$.

Podemos escribir a la derivada covariante de manera mas compacta:

$$\nabla_{\vec{\gamma}} \vec{Z} = \dot{x}^\lambda b^\mu{}_{;\lambda} \partial_\mu$$

con

$$b^\mu{}_{;\lambda} = b^\mu{}_{,\lambda} - A^\mu_{\lambda\sigma} b^\sigma$$

La cantidad $-A^\mu_{\lambda\sigma}$ es usualmente llamada *símbolo de Christoffel* y denotada por $\Gamma^\mu_{\lambda\sigma}$. El símbolo de Christoffel es entonces una conexión dada su íntima relación con la conexión asociada al subespacio horizontal $H(FM)$.

Dentro de este formalismo podemos ahora presentar con precisión el concepto de geodésicas. De acuerdo a la definición dada en la sección §4.1 una geodésica es una curva $\gamma \in M$ tal que sus vectores tangentes son paralelos a lo largo de γ , dado que $\nabla_{\vec{X}} \vec{X}$ es la razón de cambio de $\vec{X}$ en la dirección $\vec{X}$ bajo un transporte paralelo, entonces si $\vec{X}$ es el vector tangente a una geodésica habrá de cumplirse entonces la siguiente ecuación:

$$\nabla_{\vec{X}} \vec{X} = 0$$

es decir:

$$\dot{x}^\lambda \dot{x}^\mu{}_{,\lambda} + \Gamma^\mu_{\lambda\sigma} \dot{x}^\lambda \dot{x}^\sigma = 0$$

entonces:

$$\frac{d^2 x^\mu}{dt^2} + \Gamma^\mu_{\lambda\sigma} \frac{dx^\lambda}{dt} \frac{dx^\sigma}{dt} = 0 \quad (4.4)$$

De tal manera que las curvas $x^\mu(t)$ que satisfagan (4.4) son geodésicas.

Capítulo 5

Modelo σ no lineal.

En este capítulo mostraremos que las ecuaciones que del modelo σ no lineal $G = Sl(2, \mathbb{R})$, $H = SO(2)$ se derivan, son las ecuaciones "principales" de Einstein axisimétricas estacionarias en el vacío. Para tal efecto hemos dividido el presente capítulo en dos bloques, uno donde se expone someramente la axisimetría estacionaria (§5.1 y §5.2) y otro donde se define a un modelo σ no lineal (§5.3), se construye la densidad lagrangiana para el mismo (§5.4) y finalmente se muestra, como aplicación, que uno de tales modelos reproduce las ecuaciones "principales" de Einstein axisimétricas estacionarias en vacío (§5.5).

Es importante que el lector, antes de abordar éste capítulo, tenga una idea global de los anteriores, pues en las secciones §5.3, §5.4 y §5.5 se manejan sin mayor hincapié diversos conceptos expuestos en aquellos.

§5.1 Métricas axisimétricas estacionarias.

Como vimos en la sección §2.10, la expresión para el cuadrado de la distancia entre dos eventos infinitesimalmente próximos en el espacio-tiempo esta dada, en el sistema coordenado $\{x^i\}$, por $ds^2 = g_{ij}dx^i dx^j$, con $i, j = 0, 1, 2, 3$. Puesto que g es un tensor simétrico entonces hay diez componentes métricas independientes y por lo tanto $ds^2 = g_{ii}dx^i dx^i + 2g_{ij}dx^i dx^j$, donde en el segundo sumando $i < j$; sin embargo, generalmente las simetrías que caracterizan al espacio-tiempo reducen el número de coeficientes métricos independientes y determinan la dependencia de los mismos en las coordenadas.

En consecuencia, la forma de la métrica y la dependencia específica en las coordenadas

de los coeficientes métricos son un reflejo de las simetrías del espacio-tiempo; para hallar tal reflejo en el caso axisimétrico estacionario consideraremos el caso de una estrella con rotación fija (i.e: independiente del tiempo) [13]. Intuitivamente es claro que tanto la estrella como el campo en torno a ella poseen simetría axial respecto al eje de rotación, eje z , que pasa por el centro de la estrella (origen de nuestro sistema coordenado). La independencia temporal y la simetría axial del sistema en consideración sugieren la existencia de una coordenada angular $x^3 = \phi$ y una tipo temporal $x^0 = t$ para las cuales $\mathcal{L}_{\frac{\partial}{\partial t}} g_{ij} = 0$ y $\mathcal{L}_{\frac{\partial}{\partial \phi}} g_{ij} = 0$. Por lo tanto, $g_{ij} = g_{ij}(x^1, x^2)$. Nótese que $(t, x^1, x^2, \phi) = (t, x^1, x^2, \phi + 2\pi)$ puesto que ϕ es la coordenada angular alrededor del eje de rotación.

Las componentes de la 4-velocidad de la estrella rotando en la dirección del vector $\frac{\partial}{\partial \phi}$ son, en el sistema coordenado $\{(t, x^1, x^2, x^3)\}$, las siguientes:

$$U^0 = \frac{dt}{ds} = U^0(x^1, x^2), \quad U^1 = \frac{dx^1}{ds} = 0, \quad U^2 = \frac{dx^2}{ds} = 0 \quad y \quad (5.1)$$

$$U^3 = \frac{d\phi}{ds} = \frac{d\phi}{dt} \frac{dt}{ds} = \Omega(x^1, x^2) U^0(x^1, x^2)$$

donde $\Omega(x^1, x^2)$ es la velocidad angular medida en unidades de la coordenada temporal t de nuestro sistema coordenado. Evidentemente (5.1) muestra que una partícula material en la estrella corresponde a valores fijos de x^1 y x^2 (i.e: una partícula material solamente gira alrededor del eje de simetría). Para rotaciones rígidas Ω es una constante, en tanto que para una rotación diferencial Ω es en general una función de las coordenadas x^1 y x^2 .

Obsérvese que el campo *no es invariante* bajo la transformación $t \mapsto -t$ ó $\phi \mapsto -\phi$. En efecto, cada una de estas transformaciones invierte el sentido de rotación de la estrella y , por ende, se obtiene una distinta geometría espacio-temporal. No obstante, el campo de la estrella *es invariante* si las anteriores transformaciones son aplicadas simultáneamente. Con este resultado y dada la invariancia de ds^2 , se infiere que:

$$g_{01} = g_{02} = g_{13} = g_{23} = 0.$$

Entonces $ds^2 = g_{00}dt^2 + 2g_{03}dtd\phi + g_{33}d\phi^2 + g_{AB}dx^A dx^B$, donde $A, B = 1, 2$. Notemos que de facto el sistema coordenado introducido desde un principio es el que corresponde a las coordenadas cilíndricas, de tal manera que al elemento ds^2 lo podemos escribir como:

$$ds^2 = f dt^2 - 2k dt d\phi - l d\phi^2 - A d\rho^2 - 2B \rho dp dz - C dz^2 \quad (5.2)$$

donde f, k, l, A, B y C son funciones de ρ y z .

Sea la transformación $(\rho, z) \mapsto (\rho', z')$ dada por $F(\rho, z) = \rho'$ y $G(\rho, z) = z'$. Entonces:

$$d\rho = J^{-1}(G_z d\rho' - F_z dz') \quad \text{y} \quad dz = J^{-1}(F_\rho dz' - G_\rho d\rho')$$

donde $J = F_\rho G_z - F_z G_\rho$ y asumimos que es no nulo ¹.

Sustituyendo en (5.2) se obtiene que:

$$\begin{aligned} ds^2 = & f dt^2 - 2k dt d\phi - l d\phi^2 - J^{-2} \{ (AG_z^2 - 2BG_\rho G_z + CG_\rho^2) d\rho'^2 + \\ & 2(BG_z F_\rho - AG_z F_z + BF_z G_\rho - CF_\rho G_\rho) d\rho' dz' + (AF_z^2 - 2BF_z F_\rho + CF_\rho^2) dz'^2 \} \end{aligned}$$

Si exigimos ahora que F y G satisfagan el siguiente conjunto de ecuaciones diferenciales acopladas:

$$\begin{aligned} AG_z^2 - 2BG_\rho G_z + CG_\rho^2 &= AF_z^2 - 2BF_z F_\rho + CF_\rho^2 \\ BG_z F_\rho - AG_z F_z + BF_z G_\rho - CF_\rho G_\rho &= 0 \end{aligned} \quad (5.3)$$

donde asumimos que para A, B y C dadas el sistema (5.3) tiene solución no trivial con $J \neq 0$. Entonces:

$$ds^2 = f dt^2 - 2k dt d\phi - l d\phi^2 - J^{-2} (AG_z^2 - 2BG_\rho G_z + CG_\rho^2) (d\rho'^2 + dz'^2)$$

donde $J^{-2} (AG_z^2 - 2BG_\rho G_z + CG_\rho^2) > 0$ y, por lo tanto, podemos escribir (eliminando a las primas) lo siguiente:

$$ds^2 = f dt^2 - 2k dt d\phi - l d\phi^2 - e^\mu (d\rho^2 + dz^2)$$

donde f, k, l y μ son funciones de las nuevas variables ρ y z .

Las componentes contravariantes no nulas del tensor métrico son:

$$g^{00} = P^{-2} l, \quad g^{03} = -P^{-2} k, \quad g^{11} = g^{22} = -e^{-\mu}, \quad g^{33} = -P^{-2} f$$

con $P^2 = fl + k^2$.

¹En adelante el subíndice en funciones denotará derivada parcial respecto a la coordenada correspondiente, en tanto ello no se preste a confusiones.

Con objeto de simplificar aún más la expresión para ds^2 centraremos momentáneamente nuestra discusión en torno a tres de las ecuaciones de Einstein en el vacío dadas genéricamente por $R_{\mu\nu} = 0$ (ver §B.2):

$$2e^\mu P^{-1}R_{00} = (P^{-1}f_\rho)_\rho + (P^{-1}f_z)_z + P^{-3}f(f_\rho l_\rho + f_z l_z + k_\rho^2 + k_z^2) = 0 \quad (5.4)$$

$$-2e^\mu P^{-1}R_{03} = (P^{-1}k_\rho)_\rho + (P^{-1}k_z)_z + P^{-3}k(f_\rho l_\rho + f_z l_z + k_\rho^2 + k_z^2) = 0 \quad (5.5)$$

$$-2e^\mu P^{-1}R_{33} = (P^{-1}l_\rho)_\rho + (P^{-1}l_z)_z + P^{-3}l(f_\rho l_\rho + f_z l_z + k_\rho^2 + k_z^2) = 0 \quad (5.6)$$

De 5.4, 5.5 y 5.6 se sigue que:

$$e^\mu P^{-1}(lR_{00} - 2kR_{03} - fR_{33}) = P_{\rho\rho} + P_{zz} = 0 \quad (5.7)$$

De acuerdo a (5.7) P satisface la ecuación de Laplace en las variables ρ y z . Por consiguiente P es una función armónica y puede ser considerada como la parte real de una función analítica $F = F(\rho + iz)$ [16]. Sea I la función armónica conjugada de P , entonces $F(\rho + iz) = P(\rho, z) + iI(\rho, z)$ y debe satisfacerse que:

$$P_\rho = I_z \quad y \quad P_z = -I_\rho \quad (\text{Ecuaciones de Cauchy-Riemann}). \quad (5.8)$$

Sea la transformación $(\rho, z) \mapsto (\tilde{\rho}, \tilde{z})$ dada por P e I . Entonces:

$$d\tilde{\rho}^2 + d\tilde{z}^2 = P_\rho^2 d\rho^2 + P_z^2 dz^2 + 2P_\rho P_z d\rho dz + I_\rho^2 d\rho^2 + I_z^2 dz^2 + 2I_\rho I_z d\rho dz$$

por (5.8) obtenemos que:

$$d\tilde{\rho}^2 + d\tilde{z}^2 = (P_\rho^2 + I_\rho^2)(d\rho^2 + dz^2)$$

y por lo tanto:

$$ds^2 = f dt^2 - 2k dt d\phi - l d\phi^2 - e^\mu (P_\rho^2 + I_\rho^2)^{-1} (d\tilde{\rho}^2 + d\tilde{z}^2)$$

sea $e^\mu = e^\mu (P_\rho^2 + I_\rho^2)^{-1}$, entonces:

$$ds^2 = f dt^2 - 2k dt d\phi - l d\phi^2 - e^\mu (d\tilde{\rho}^2 + d\tilde{z}^2) \quad (5.9)$$

donde f, k, l y $\bar{\mu}$ son funciones de las nuevas variables $\bar{\rho}$ y $\bar{z}$. Sea $\omega = f^{-1}k$, puesto que $f l + k^2 = \bar{\rho}^2 \Rightarrow l = f^{-1}(\bar{\rho}^2 - f^2 \omega^2)$; sustituyendo en (5.9) tenemos que:

$$ds^2 = f(dt - \omega d\phi)^2 - f^{-1} \bar{\rho}^2 d\phi^2 - e^{\bar{\mu}} (d\bar{\rho}^2 + d\bar{z}^2) \quad (5.10)$$

A la forma (5.10) se le conoce como la forma de Weyl-Lewis-Papapetrou [13]. Finalmente hagamos $f = e^{2\psi}$ y $\bar{\mu} = 2(\gamma - \psi)$; entonces (eliminando las tildes) obtenemos que:

$$ds^2 = e^{2\psi} (dt - \omega d\phi)^2 - e^{-2\psi} \{ e^{2\gamma} (d\rho^2 + dz^2) + \rho^2 d\phi^2 \} \quad (5.11)$$

donde ψ, ω y γ son funciones solamente de ρ y z .

Las componentes covariantes y contravariantes no nulas del tensor métrico son entonces:

$$g_{tt} = e^{2\psi}, \quad g_{t\phi} = -e^{2\psi}\omega, \quad g_{\rho\rho} = g_{zz} = -e^{2(\gamma-\psi)}, \quad (5.12)$$

$$g_{\phi\phi} = (e^{2\psi}\omega^2 - e^{-2\psi}\rho^2)$$

$$g^{tt} = \rho^{-2}(e^{-2\psi}\rho^2 - e^{2\psi}\omega^2), \quad g^{t\phi} = -e^{2\psi}\omega\rho^{-2}, \quad (5.13)$$

$$g^{\rho\rho} = g^{zz} = -e^{2(\psi-\gamma)}, \quad g^{\phi\phi} = -e^{2\psi}\rho^{-2}$$

En tanto que la raíz cuadrada del valor absoluto del determinante de la métrica es:

$$\sqrt{-g} = \rho e^{2(\gamma-\psi)} \quad (5.14)$$

§5.2 Densidad lagrangiana axisimétrica estacionaria.

Las ecuaciones de Einstein en el vacío, cuya solución determina a los coeficientes métricos y en consecuencia describe a un determinado espacio-tiempo exterior, se obtienen a partir del principio de mínima acción para la acción [11]:

$$S = \int R \sqrt{-g} d\Omega \quad (5.15)$$

Observemos que la densidad lagrangiana² de (5.15) puede descomponerse en dos sumandos, uno de los cuales contiene solamente al tensor g_{ik} y sus primeras derivadas, en tanto que el otro es la divergencia de una cierta cantidad:

$$\mathcal{L}_{EH} = R\sqrt{-g} = \sqrt{-g}g^{ik}R_{ik} = \sqrt{-g}(g^{ik}\Gamma_{ik,i}^i - g^{ik}\Gamma_{ii,k}^i + g^{ik}\Gamma_{ik}^i\Gamma_{lm}^m - g^{ik}\Gamma_{ii}^m\Gamma_{km}^i) \quad (5.16)$$

Los dos primeros sumandos de (5.16) pueden escribirse como:

$$\begin{aligned}\sqrt{-g}g^{ik}\Gamma_{ik,i}^i &= (\sqrt{-g}g^{ik}\Gamma_{ik}^i)_{,i} - \Gamma_{ik}^i(\sqrt{-g}g^{ik})_{,i} \\ \sqrt{-g}g^{ik}\Gamma_{ii,k}^i &= (\sqrt{-g}g^{ik}\Gamma_{ii}^i)_{,k} - \Gamma_{ii}^i(\sqrt{-g}g^{ik})_{,k}\end{aligned} \quad (5.17)$$

y sustituyendo (5.17) en (5.16) obtenemos lo que anunciábamos en un inicio:

$$\sqrt{-g}R = \sqrt{-g}G + \{\sqrt{-g}(g^{ii}\Gamma_{ii}^k - g^{ik}\Gamma_{ii}^k)\}_{,k} \quad (5.18)$$

donde $\sqrt{-g}G = \Gamma_{ii}^i(\sqrt{-g}g^{ik})_{,k} - \Gamma_{ik}^i(\sqrt{-g}g^{ik})_{,i} + g^{ik}\sqrt{-g}(\Gamma_{ik}^i\Gamma_{lm}^m - \Gamma_{ii}^m\Gamma_{km}^i)$. Esta expresión para G puede llevarse a una forma más compacta [11]:

$$G = g^{ik}(\Gamma_{mk}^i\Gamma_{ii}^m - \Gamma_{ik}^i\Gamma_{lm}^m)$$

La ecuación (5.18) implica que la acción dada en (5.15) puede escribirse como:

$$S = \int G\sqrt{-g}d\Omega + \int \{\sqrt{-g}(g^{ii}\Gamma_{ii}^k - g^{ik}\Gamma_{ii}^k)\}_{,k}d\Omega \quad (5.19)$$

Aplicando el teorema de Gauss en la segunda integral de (5.19) y variando la acción obtenemos que:

$$\delta S = \int \delta(G\sqrt{-g})d\Omega + \int \delta(\sqrt{-g}(g^{ii}\Gamma_{ii}^k - g^{ik}\Gamma_{ii}^k))d\Sigma \quad (5.20)$$

donde la segunda integral es la integral de superficie de la hipersuperficie que rodea al cuadrivolumen en el cual se llevan a cabo las otras integrales.

Puesto que son nulas las variaciones del campo en los límites de la región de integración, de acuerdo con el principio de mínima acción [11], entonces la segunda integral en (5.20)

²A esta densidad se le conoce en la literatura como densidad lagrangiana de Einstein-Hilbert y la denotaremos por $\mathcal{L}_{EH}$.

es cero³. En consecuencia, $\int \delta(R\sqrt{-g})d\Omega = \int \delta(G\sqrt{-g})d\Omega$; es decir, las ecuaciones que se derivan por el principio de mínima acción para la densidad $\mathcal{L}_{EH}$ son las mismas que se derivan a partir de $G\sqrt{-g}$. Por esta razón, aunque G no sea un escalar, es posible dar a $G\sqrt{-g}$ el carácter de densidad lagrangiana y denotarla por $\mathcal{L}'$.

La densidad lagrangiana $\mathcal{L}_{EH}$ para la métrica (5.11), que reproduce las ecuaciones de Einstein para el caso axisimétrico estacionario⁴, es:

$$\mathcal{L}_{EH} = -\frac{\sqrt{-g}}{2\rho^2} \{4e^{2(\psi-\gamma)}\rho^2(\psi_\rho^2 + \psi_z^2) - e^{6\psi-2\gamma}(\omega_\rho^2 + \omega_z^2)\}$$

$$-2\sqrt{-g}e^{2(\psi-\gamma)}\{\gamma_{\rho\rho} + \gamma_{zz} - \psi_{\rho\rho} - \psi_{\rho\rho} - \psi_{\rho\rho}^{-1}\}$$

sumando el cero de la forma $2\sqrt{-g}\rho^{-1}e^{2(\psi-\gamma)}\gamma_\rho - 2\sqrt{-g}\rho^{-1}e^{2(\psi-\gamma)}\gamma_\rho$ obtenemos que:

$$\mathcal{L}_{EH} = -\frac{\sqrt{-g}}{2\rho^2} \{4e^{2(\psi-\gamma)}\rho^2(\psi_\rho^2 + \psi_z^2) - e^{6\psi-2\gamma}(\omega_\rho^2 + \omega_z^2) - 4\rho e^{2(\psi-\gamma)}\gamma_\rho\} \quad (5.21)$$

$$-\sqrt{-g}e^{2(\psi-\gamma)}\{\gamma_{\rho\rho} + \gamma_{zz} - \psi_{\rho\rho} - \psi_{\rho\rho} - \psi_{\rho\rho}^{-1} + \gamma_{\rho\rho}^{-1}\}$$

sea $\vec{\chi} = (\psi_\rho - \gamma_\rho, \psi_z - \gamma_z, 0)$, entonces el segundo término de (5.21) es igual a $2\sqrt{-g}e^{2(\psi-\gamma)}\nabla \cdot \vec{\chi}$ (donde $\nabla \cdot \vec{\chi}$ es la divergencia usual del campo χ) substituyendo $\sqrt{-g}$ por el valor dado en (5.14) tenemos que $2\sqrt{-g}e^{2(\psi-\gamma)}\nabla \cdot \vec{\chi} = 2\rho\nabla \cdot \vec{\chi}$ y por lo tanto:

$$\mathcal{L}_{EH} = -\frac{\sqrt{-g}}{2\rho^2} \{4e^{2(\psi-\gamma)}\rho^2(\psi_\rho^2 + \psi_z^2) - e^{6\psi-2\gamma}(\omega_\rho^2 + \omega_z^2) - 4\rho e^{2(\psi-\gamma)}\gamma_\rho\} + 2\rho\nabla \cdot \vec{\chi}$$

Entonces, la densidad lagrangiana $\mathcal{L}'$ está dada por:

$$\mathcal{L}' \equiv \sqrt{-g}G = \frac{\sqrt{-g}}{2\rho^2} \{e^{6\psi-2\gamma}(\omega_\rho^2 + \omega_z^2) + 4\rho e^{2(\psi-\gamma)}\gamma_\rho - 4e^{2(\psi-\gamma)}\rho^2(\psi_\rho^2 + \psi_z^2)\}$$

y substituyendo el valor de $\sqrt{-g}$:

$$\mathcal{L}' = \frac{1}{2\rho} \{e^{4\psi}(\omega_\rho^2 + \omega_z^2) + 4\rho\gamma_\rho - 4\rho^2(\psi_\rho^2 + \psi_z^2)\} \quad (5.22)$$

³En términos estrictos esto no es cierto, pues aunque las variaciones del campo en la frontera son nulas, ello no necesariamente debe cumplirse para las derivadas del campo. Es necesario agregar a la acción un término extra que cancele en la frontera las variaciones en la derivada del campo, el término es $2\int_{\partial U} K$, donde ∂U es la frontera de la región de integración y K es la traza de la curvatura extrínseca de dicha frontera [12].

⁴En este caso conmuta el imponer la simetría en la densidad lagrangiana y luego variar, que variar y luego imponer la simetría.

Introduciendo el operador $D \equiv (\partial_\rho, \partial_z)$, (5.22) queda como:

$$\mathcal{L}' = 2D\rho D\gamma + \frac{e^{4\psi}}{2\rho}(D\omega)^2 - 2\rho(D\psi)^2 \quad (5.23)$$

donde γ y ω son ahora coordenadas cíclicas [17] de la densidad lagrangiana (5.23). La densidad Routhiana (ver §B.1) es entonces la siguiente:

$$\mathcal{R}_{EH} = \frac{\partial \mathcal{L}'}{\partial D\gamma} D\gamma + \frac{\partial \mathcal{L}'}{\partial D\omega} D\omega - \mathcal{L}' = \frac{1}{2}\rho e^{-4\psi} \Pi_\omega^2 + 2\rho(D\psi)^2 \quad (5.24)$$

donde Π_ω es el momento canónico conjugado asociado a la "coordenada" generalizada ω [17].

Las ecuaciones de campo que de (5.24) se derivan (ver §B.1) son:

$$D\left(\frac{\partial \mathcal{R}_{EH}}{\partial D\psi}\right) - \frac{\partial \mathcal{R}_{EH}}{\partial \psi} = 0 \quad (5.25)$$

$$D\omega = \frac{\partial \mathcal{R}_{EH}}{\partial \Pi_\omega}, \quad D\Pi_\omega = -\frac{\partial \mathcal{R}_{EH}}{\partial \omega} = 0$$

Si ahora definimos a $\Omega = \Omega(\rho, z)$ como $\Pi_\omega = \tilde{D}\Omega$, donde $\tilde{D} = (-\partial_z, \partial_\rho)$ [17], entonces la densidad Routhiana queda como:

$$\mathcal{R}_{EH} = 2\rho(D\psi)^2 + \frac{\rho}{2}e^{-4\psi}(D\Omega)^2 \quad (5.26)$$

y las ecuaciones de campo (5.25) como:

$$D\left(\frac{\partial \mathcal{R}_{EH}}{\partial D\psi}\right) - \frac{\partial \mathcal{R}_{EH}}{\partial \psi} = 0 \quad (5.27)$$

$$D\omega = \frac{\partial \mathcal{R}_{EH}}{\partial \Pi_\omega}, \quad D(\tilde{D}\Omega) = 0$$

El sistema de ecuaciones diferenciales acopladas (5.27) es equivalente a las ecuaciones "principales" de Einstein axisimétricas estacionarias para el vacío (ver §B.2).

§5.3 Definición de un modelo σ no lineal.

Por un modelo σ no lineal (o simplemente modelo no lineal) habremos de entender, a grandes rasgos, una teoría de campo con las siguientes dos propiedades [14]:

- Los campos se encuentran sujetos a constricciones no lineales.
- El lagrangiano L (o equivalentemente la densidad lagrangiana $\mathcal{L}$) y las constricciones son invariantes bajo la acción de un grupo de simetría global G .

Puntualicemos el concepto de no linealidad citado en nuestra definición. La descripción no lineal se refiere, de manera más precisa, a aquellos modelos donde los campos físicos, para todos los puntos $\mathcal{X}$ del espacio-tiempo X , toman valores en una variedad M que no es un espacio vectorial.

Cabe señalar que en la gran mayoría de estos modelos el grupo G actúa transitivamente sobre la citada variedad M . Si $H \subset G$ es el grupo de isotropía (o grupo de estabilidad) de un $p \in M$: $H(p) = \{h \in G \mid hp = p\}$ ⁵, entonces la transitividad de G sobre M garantiza que el estabilizador es el mismo para todos los puntos, si además H es un divisor normal⁶ (o subgrupo invariante) de G entonces es posible identificar a M con el espacio G/H ; las clases de equivalencia en G/H tendrán entonces una correspondencia 1-1 con los puntos en M : $\Phi: G/H \rightarrow M, [g] \mapsto p$.

Para esta teoría de campo particular existen métodos generales enfocados a la construcción de lagrangianos. A continuación presentaremos uno de estos métodos que, a la postre, nos permitirá mas adelante mostrar que las ecuaciones "principales" de Einstein axisimétricas estacionarias para el vacío están relacionadas con las correspondientes al modelo σ no lineal $G = SI(2, \mathbb{R})$, $H = SO(2)$ (§5.5).

§5.4 Construcción de la densidad lagrangiana para un modelo σ no lineal.

Comenzaremos la construcción pidiendo que los campos $\{g\}$ en nuestra densidad lagrangiana $\mathcal{L}$ tengan valores en el grupo de Lie G ; i.e.: $g(\mathcal{X}) \in G$ [14]. Adicionalmente, la densidad lagrangiana $\mathcal{L}$ habrá de ser invariante bajo las transformaciones de norma⁷:

⁵El lector debe tomar en cuenta que ahora G es una representación del grupo y, por ende, h corresponde a Γ_a .

⁶Para que un subgrupo H de un grupo G sea un divisor normal de éste es necesario y suficiente que para todo elemento $g \in G$ se cumpla la igualdad $g^{-1}Hg = H$.

⁷Es pertinente aclarar que la discusión previa sobre transformaciones de norma, donde la acción se tomaba por la izquierda, es igualmente válida si la acción es por la derecha.

$$g(\vec{x}) \mapsto g(\vec{x})h(\vec{x}), \quad h(\vec{x}) \in H \subset G$$

donde H es un divisor normal de G .

En consecuencia, los campos invariantes de norma ó campos físicos tienen valores en G/H .

Llevaremos a cabo la construcción de la densidad lagrangiana a partir de una forma diferencial simple, a saber la forma de Maurer-Cartan, definida sobre G [14]. Para tal efecto, identificamos al grupo de Lie G con cualquiera de sus representaciones fieles y de nueva cuenta etiquetamos a la representación por G :

Sea $\{L_\rho\}$ la base para el álgebra de Lie $\underline{G}$ de G en alguna representación fiel de G , con la siguiente propiedad:

$$[Tr\{L_\rho L_\sigma\}] = \delta_{\rho\sigma}; \quad \text{con } \rho, \sigma = \{1, 2, \dots, [G] = \dim G\} \quad (5.28)$$

Para $\alpha \leq [H]$ los generadores L_α dan como resultado el álgebra de Lie $\underline{H}$ de H y serán denotados por T_α :

$$L_\alpha = T_\alpha; \quad \alpha \leq [H].$$

A los generadores complementarios los llamaremos S_i :

$$L_i = S_i; \quad [H] + 1 \leq i \leq [G].$$

Halleemos las relaciones de conmutación:

a).- $[T_\alpha, T_\beta] = iC_{\alpha\beta}^\gamma T_\gamma$, dado que el conjunto $\{T_\alpha; \alpha \leq [H]\}$ es un generador del álgebra de Lie $\underline{H}$ de H .

b).- $[T_\alpha, S_i] = iA_{\alpha i}^\sigma T_\sigma + iC_{\alpha i}^j S_j$. Pero nótese que esta expresión en realidad es más compacta; en efecto:

$$\begin{aligned} \text{Como } Tr\{T_\gamma[T_\alpha, S_i]\} &= Tr\{T_\gamma T_\alpha S_i\} - Tr\{T_\gamma S_i T_\alpha\} \\ &= Tr\{S_i T_\gamma T_\alpha\} - Tr\{S_i T_\alpha T_\gamma\} \\ &= Tr\{S_i [T_\gamma, T_\alpha]\} = Tr\{S_i iC_{\gamma\alpha}^\beta T_\beta\} = 0 \quad \text{en virtud de (5.28).} \\ \Rightarrow [T_\alpha, S_i] &= iC_{\alpha i}^j S_j \end{aligned}$$

c).- $[S_i, S_j] = i\{D_{ij}^\alpha T_\alpha + D_{ij}^* S_k\}$, con $D_{ij}^\alpha = \bar{C}_{\alpha i}^j$. En efecto:

$$\begin{aligned} \text{Tr}\{T_\alpha[S_i, S_j]\} &= \text{Tr}\{S_i T_\alpha S_j\} - \text{Tr}\{T_\alpha S_i S_j\} \\ &= \text{Tr}\{[S_i, T_\alpha] S_j\} = -\text{Tr}\{[T_\alpha, S_i] S_j\} \\ &= -\text{Tr}\{i\bar{C}_{\alpha i}^k S_k S_j\} = -i\bar{C}_{\alpha i}^j \delta_{kj} = -i\bar{C}_{\alpha i}^j \\ \Rightarrow \text{Tr}\{T_\alpha[S_i, S_j]\} &= i\bar{C}_{\alpha i}^j \end{aligned}$$

Por otro lado:

$$\begin{aligned} \text{Tr}\{T_\alpha[S_i, S_j]\} &= i\text{Tr}\{D_{ij}^\alpha T_\beta T_\alpha\} + i\text{Tr}\{D_{ij}^* S_k T_\alpha\} \\ &= iD_{ij}^\alpha \delta\beta\alpha = iD_{ij}^\alpha \end{aligned}$$

En consecuencia tenemos que $D_{ij}^\alpha = \bar{C}_{\alpha i}^j$.

Sea $\tilde{w}$ una 1-forma definida en G con componentes $w_\mu(g) = g^{-1}\partial_\mu g$, $g = g(\vec{x})$. Bajo la transformación de norma $g \mapsto gh$ tenemos que:

$$\begin{aligned} w_\mu(gh) &= h^{-1}g^{-1}g\partial_\mu h + h^{-1}g^{-1}\partial_\mu gh \\ \Rightarrow w_\mu(gh) &= h^{-1}\partial_\mu h + h^{-1}w_\mu(g)h \end{aligned}$$

Sean las 1-formas $\tilde{A}$ y $\tilde{B}$ con componentes $A_\mu(g) = T_\alpha \text{Tr}\{T^\alpha w_\mu(g)\}$ y $B_\mu(g) = S_i \text{Tr}\{S^i w_\mu(g)\}$ respectivamente. Notemos entonces que:

$$\begin{aligned} \text{Tr}\{T_\beta A_\mu(g)\} &= \text{Tr}\{T_\beta w_\mu(g)\} \\ \Rightarrow w_\mu(g) &= A_\mu(g) + \xi_\mu(g) \text{ con } \xi_\mu(g) \text{ tal que } \text{Tr}\{T_\beta \xi_\mu(g)\} = 0 \end{aligned}$$

y

$$\begin{aligned} \text{Tr}\{S_j B_\mu(g)\} &= \text{Tr}\{S_j w_\mu(g)\} \\ \Rightarrow w_\mu(g) &= B_\mu(g) + \eta_\mu(g) \text{ con } \eta_\mu(g) \text{ tal que } \text{Tr}\{S_j \eta_\mu(g)\} = 0 \end{aligned}$$

En consecuencia $A_\mu(g) + \xi_\mu(g) = B_\mu(g) + \eta_\mu(g)$ y por lo tanto:

$$\text{Tr}\{T_\beta A_\mu(g)\} = \text{Tr}\{T_\beta \eta_\mu(g)\}, \quad \text{Tr}\{S_j \xi_\mu(g)\} = \text{Tr}\{S_j B_\mu(g)\}$$

De lo cual se deriva que una solución al sistema es $A_\mu(g) = \eta_\mu(g)$ y $B_\mu(g) = \xi_\mu(g)$. Por lo tanto obtenemos que $\tilde{w} = \tilde{A} + \tilde{B}$ con $w_\mu(g) = A_\mu(g) + B_\mu(g)$.

Bajo la transformación de norma $g \mapsto gh$ las componentes de $\tilde{A}$ y $\tilde{B}$ se transforman como sigue [14]:

$$A_\mu(gh) = h^{-1} \partial_\mu h + h^{-1} A_\mu(g) h$$

$$B_\mu(gh) = h^{-1} B_\mu(g) h$$

Debe notarse que la estructura con la cual contamos es la de un haz fibrado con $E = G$, $B = G/H$, $F \simeq H$, grupo de estructura H y proyector $\Pi: G \rightarrow G/H$, donde $g(x) \mapsto [g(x)]$ y $[g(x)]$ denota la clase de equivalencia del elemento $g(x) \in G$. Dentro de esta estructura tenemos que por construcción *todos* los campos están en G en tanto que *los campos físicos* se hallan en la base G/H del haz. Debido a esta estructura tenemos entonces que la 1-forma $\tilde{A}$ se transforma como un potencial de norma para el grupo de norma H y, por ende, actúa como una 1-forma de conexión (§4.2). Por lo tanto, a partir de $\tilde{A}$ podemos construir la 2-forma de curvatura $F_{\mu\nu} = \partial_\mu A_\nu - \partial_\nu A_\mu + [A_\mu, A_\nu]$ que, bajo una transformación de norma, tiene la siguiente ley (§4.3): $F' = h^{-1} F h$.

Observemos que entonces las cantidades $\text{Tr}\{B_\mu B_\nu\}$ y $\text{Tr}\{F_{\mu\nu} F_{\tau\rho}\}$ son invariantes bajo una transformación de norma. En virtud de que la densidad lagrangiana habrá de ser invariante bajo transformaciones de norma, resulta natural asociar entonces una densidad a cada una de las cantidades anteriormente citadas:

$$\mathcal{L}_\sigma = \sqrt{-g} g^{\mu\nu} \text{Tr}\{B_\mu B_\nu\} \quad (5.29)$$

y

$$\mathcal{L}_\sigma = \sqrt{-g} g^{\mu\tau} g^{\nu\rho} \text{Tr}\{F_{\mu\nu} F_{\tau\rho}\} \quad (5.30)$$

Todas las densidades lagrangianas para modelos no lineales son construidas a partir de A_μ y B_μ . En principio no hay ninguna prohibición para mezclarlos; por ejemplo, la densidad

$$\mathcal{L}_\sigma = \sqrt{-g} g^{\alpha\beta} g^{\mu\tau} g^{\nu\rho} \text{Tr}\{F_{\mu\nu} F_{\tau\rho}\} \text{Tr}\{B_\alpha B_\beta\}$$

es igual de válida que las anteriores.

Puesto que $\tilde{u}$ es una sección en el haz cotangente T^*G entonces es claro que la densidad $\mathcal{L}_\sigma$ es función de una sección en T^*G y, además, es invariante bajo aquellos cambios de sección en T^*G inducidos por el grupo de norma H . Estos cambios de sección en T^*G son la consecuencia de una transformación del tipo $g \mapsto gh$ en G , por consiguiente la densidad $\mathcal{L}_\sigma$ es invariante también bajo cambios de sección en G . En consecuencia basta fijar una sección en G , una norma (§4.3), para hallar a la densidad lagrangiana $\mathcal{L}_\sigma$.

Sea Υ un campo invariante de norma, entonces $\Upsilon \in G/H$ y puede expresarse como sigue:

$$\Upsilon(\vec{x}) = g(\vec{x}) \left(\sum_{\alpha} T_{\alpha} \right) g^{-1}(\vec{x}) \quad (5.31)$$

y dado que el lado derecho de (5.31) pertenece a $\underline{G}$, entonces:

$$g(\vec{x}) \left(\sum_{\alpha} T_{\alpha} \right) g^{-1}(\vec{x}) = \sum_{\rho} v_{\rho}(\vec{x}) L_{\rho}$$

donde los campos $\{v_{\rho}(\vec{x})\}$ evidentemente también son campos físicos. Notemos que la construcción para estos campos es, en efecto, no lineal:

$$Tr\{ (g \left(\sum_{\alpha} T_{\alpha} \right) g^{-1}) (g \left(\sum_{\beta} T_{\beta} \right) g^{-1}) \} = Tr\{ g \left(\sum_{\alpha, \beta} T_{\alpha} T_{\beta} \right) g^{-1} \}$$

$$\Rightarrow Tr\{ \sum_{\alpha, \beta} T_{\alpha} T_{\beta} \} = Tr\{ \sum_{\rho, \sigma} v_{\rho}(\vec{x}) L_{\rho} v_{\sigma}(\vec{x}) L_{\sigma} \}$$

por la propiedad (5.28) tenemos entonces que:

$$\sum_{\alpha} Tr\{T_{\alpha} T_{\alpha}\} = \sum_{\rho} v_{\rho}^2 Tr\{L_{\rho} L_{\rho}\} \quad (5.32)$$

que obviamente es una construcción no lineal para los campos v_{ρ} .

La construcción (5.32) es justamente la que define a la variedad M que, dado el carácter no lineal de la construcción, no es un espacio vectorial. Notemos también que de acuerdo a (5.32) hay solamente $[G] - 1$ campos físicos independientes.

Es importante notar que la construcción presentada en este apartado efectivamente corresponde a un modelo σ no lineal. Para tal efecto obsérvese, en primera instancia, un hecho obvio: el espacio G/H es invariante bajo G . Es decir, la aplicación de cualquier $g \in G$ sobre

todos los elementos de G/H es de nueva cuenta el espacio $G/H: G/H \stackrel{p \in G}{\rightarrow} G/H$. Como por construcción $G/H \simeq M$ y M es la variedad definida por la constricción (5.32), entonces la constricción para los campos físicos también es invariante bajo G .

Por otra parte, la densidad lagrangiana $\mathcal{L}_\sigma$ depende de $(v_1, \dots, v_i(q)) \equiv M$ dado que basta fijar una sección en G para hallar a dicha densidad $\mathcal{L}_\sigma$. De la invariancia de M bajo G se sigue que $\mathcal{L}_\sigma$ es invariante bajo G .

§5.5 El modelo σ no lineal $G = \text{Sl}(2, \mathbb{R})$, $H = \text{SO}(2)$.

En la presente sección nos abocaremos a construir la densidad correspondiente a (5.29) con $g^{\mu\nu}$ dada como en (5.12) para el grupo de simetría $G = \text{Sl}(2, \mathbb{R})$ con divisor $H = \text{SO}(2)$. Comencemos hallando los generadores de $G = \text{Sl}(2, \mathbb{R})$:

Consideremos la representación matricial de G sobre si mismo. Sea $M \in G \Rightarrow M$ es una matriz de 2×2 con $\det M = 1$. Entonces $\det M = \det(e^{tA})$ puesto que M pertenece a un subgrupo 1-paramétrico generado por alguna $A \in \text{Sl}(2, \mathbb{R})$. Como

$$\det(e^{tA}) = e^{t \text{Tr}(A)} = 1 \quad \forall t \Rightarrow \text{Tr}(A) = 0.$$

Sea $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$, como A tiene traza nula entonces $a = -d$. En consecuencia a la matriz se le puede expresar como sigue:

$$A = a \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} + \frac{b+c}{2} \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} + \frac{b-c}{2} \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}$$

Puesto que:

$$\alpha \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} + \beta \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} + \delta \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} = 0 \Rightarrow \alpha = \beta = \delta = 0$$

los generadores de $\text{Sl}(2, \mathbb{R})$ son:

$$L_1 = \frac{1}{\sqrt{2}} \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}; \quad L_2 = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}; \quad L_3 = \frac{1}{\sqrt{2}} \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$$

y satisfacen que $|\text{Tr}\{L_\rho L_\sigma\}| = \delta_{\rho\sigma}$ con $\rho, \sigma = \{1, 2, 3 = [G]\}$.

Nótese que L_1 es el generador de $\text{SO}(2)$. En efecto:

Sea $s = e^{tB} \in SO(2)$, como $s^T = s^{-1}$ entonces $B = -B^T$ y por lo tanto tenemos que:

$$B = \frac{a}{\sqrt{2}} \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} = aL_1$$

De acuerdo a la notación dada en §5.3 tenemos que: $L_1 = T_1$ y $L_{i+1} = S_i$ con $i = 1, 2$.

Observemos que solamente dos campos invariantes de norma son independientes. En efecto, según (5.32) tenemos que:

$$v_1^2 - (v_2^2 + v_3^2) = 1$$

En consecuencia es natural fijar una norma en G en términos de dos campos independientes. Con tal finalidad, consideremos a los campos η y ξ y fijemos la norma como sigue [15]:

$$g(\rho, z) = \begin{pmatrix} e^{\eta(\rho, z)} & \xi(\rho, z)e^{-\eta(\rho, z)} \\ 0 & e^{-\eta(\rho, z)} \end{pmatrix} \in SI(2, \mathbb{R}) \quad (5.33)$$

donde imponemos que la dependencia de los campos η y ξ sea exclusivamente en ρ y z para garantizar una condición indispensable de axisimetría estacionaria.

Por construcción $g(\rho, z)$ es una sección en $SI(2, \mathbb{R})$ a partir de la cual es posible especificar a la densidad lagrangiana $\mathcal{L}_\sigma$ del modelo σ no lineal. Hallemos entonces a las componentes de la 1-forma w :

$$\begin{aligned} w_\mu(g) &= g^{-1} \partial_\mu g = \begin{pmatrix} e^{-\eta} & -\xi e^{-\eta} \\ 0 & e^\eta \end{pmatrix} \begin{pmatrix} e^\eta \eta_\mu & e^{-\eta}(\xi_\mu - \xi \eta_\mu) \\ 0 & -e^{-\eta} \eta_\mu \end{pmatrix} \\ \Rightarrow w_\mu(g) &= \begin{pmatrix} \eta_\mu & e^{-2\eta} \xi_\mu \\ 0 & -\eta_\mu \end{pmatrix} \end{aligned}$$

Puesto que $B_\mu(g) = S_i Tr\{S^i w(g)\}$, un cálculo sencillo muestra que $B_\mu(g)$ esta dada por:

$$B_\mu(g) = \begin{pmatrix} \eta_\mu & \frac{1}{2} e^{-2\eta} \xi_\mu \\ \frac{1}{2} e^{-2\eta} \xi_\mu & -\eta_\mu \end{pmatrix}$$

Entonces:

$$B_\mu(g)B_\nu(g) = \begin{pmatrix} \eta_\mu\eta_\nu + \frac{1}{2}e^{-4\eta}\xi_\mu\xi_\nu & \frac{1}{2}e^{-2\eta}(\eta_\mu\xi_\nu - \eta_\nu\xi_\mu) \\ \frac{1}{2}e^{-2\eta}(\eta_\nu\xi_\mu - \eta_\mu\xi_\nu) & \eta_\mu\eta_\nu + \frac{1}{2}e^{-4\eta}\xi_\mu\xi_\nu \end{pmatrix}$$

$$\Rightarrow \text{Tr}\{B_\mu B_\nu\} = 2\eta_\mu\eta_\nu + \frac{1}{2}e^{-4\eta}\xi_\mu\xi_\nu$$

En consecuencia:

$$\mathcal{L}_\sigma = \sqrt{-g}g^{\mu\nu}(2\eta_\mu\eta_\nu + \frac{1}{2}e^{-4\eta}\xi_\mu\xi_\nu)$$

Puesto que $\eta = \eta(\rho, z)$ y $\xi = \xi(\rho, z)$ (i.e: $B_t = B_\phi = 0$), los únicos coeficientes métricos contravariantes cuyo factor multiplicativo es no nulo en la última expresión son $g^{\rho\rho}$, $g^{\phi\phi}$, g^{zz} y g^{zz} ; sustituyendo a estos coeficientes por su expresión dada en (5.13) y recurriendo al valor explícito para $\sqrt{-g}$ según (5.14) se obtiene:

$$\mathcal{L}_\sigma = 2\rho(\eta_\rho^2 + \eta_z^2) + \frac{\rho}{2}e^{-4\eta}(\xi_\rho^2 + \xi_z^2)$$

que en términos del operador D queda como:

$$\mathcal{L}_\sigma = 2\rho(D\eta)^2 + \frac{\rho}{2}e^{-4\eta}(D\xi)^2 \quad (5.34)$$

Cabe advertir que cuando se habla de una densidad lagrangiana, en principio, ello no excluye a priori la posibilidad de que alguna de sus variables sea hamiltoniana. Por consiguiente, habrá que contemplar esto para el caso de la densidad (5.34); comencemos entonces considerando a todas las variables como lagrangianas:

Si las variables son lagrangianas, entonces las ecuaciones de campo estarán dadas por la ecuación de Euler-Lagrange [9] $D((\mathcal{L}_\sigma)_{,DX^i}) - (\mathcal{L}_\sigma)_{,X^i} = 0$, donde $X^i = \{\eta, \xi\}$:

$$\eta_\rho + \rho(\eta_{\rho\rho} + \eta_{zz}) + \xi e^{-4\eta}(\xi_\rho^2 + \xi_z^2) = 0$$

$$\xi_{\rho\rho} + \xi_{zz} - 4(\eta_\rho\xi_\rho + \eta_z\xi_z) + \rho^{-1}\xi_\rho = 0 \quad (5.35)$$

En caso de que alguna de las variables sea hamiltoniana, obtendremos entonces uno de los siguientes sistemas:

Supongamos primero que $\mathcal{L}_\sigma = 2\rho(D\eta)^2 + \frac{g}{2}e^{-4\eta}(\Pi_\beta)^2$, donde Π_β es el momento asociado a la "coordenada" β (con β distinto a η y a ξ). Entonces obtenemos el siguiente sistema de ecuaciones:

$$D((\mathcal{L}_\sigma), D\eta) - (\mathcal{L}_\sigma),_\eta = 0$$

$$D\beta = (\mathcal{L}_\sigma),_{\Pi_\beta} , \quad D\Pi_\beta = -(\mathcal{L}_\sigma),_\beta = 0$$

es decir:

$$\eta_\rho + \rho(\eta_{\rho\rho} + \eta_{xx}) + \frac{g}{2}e^{-4\eta}(\xi_\rho^2 + \xi_x^2) \quad (5.36)$$

$$D\beta = \rho e^{-4\eta} \Pi_\beta , \quad D\Pi_\beta = 0 \Rightarrow \Pi_\beta = \tilde{D}\xi$$

Ahora supongamos que $\mathcal{L}_\sigma = 2\rho(\Pi_\alpha)^2 + \frac{g}{2}e^{-4\eta}(D\xi)^2$, donde Π_α es el momento asociado a la "coordenada" α (con α distinto a η y a ξ). Las ecuaciones para esta densidad son entonces las siguientes:

$$D((\mathcal{L}_\sigma), D\xi) - (\mathcal{L}_\sigma),_\xi = 0$$

$$D\alpha = (\mathcal{L}_\sigma),_{\Pi_\alpha} , \quad D\Pi_\alpha = -(\mathcal{L}_\sigma),_\alpha = 0$$

y por lo tanto:

$$\xi_{\rho\rho} + \xi_{xx} - 4(\eta_\rho \xi_\rho + \eta_x \xi_x) + \rho^{-1} \xi_\rho = 0 \quad (5.37)$$

$$D\alpha = 4\rho \Pi_\alpha , \quad D\Pi_\alpha = 0 \Rightarrow \Pi_\alpha = \tilde{D}\eta$$

Para llevar a cabo la comparación entre nuestro modelo y el caso axisimétrico estacionario, basta confrontar las ecuaciones de campo. Comencemos pues con el sistema (5.35):

Hagamos $\eta = \psi$ y $\xi = \Omega$, donde ψ y Ω están dadas como en (§5.2). Entonces (5.35) da:

$$\begin{aligned} \psi_\rho + \rho(\psi_{\rho\rho} + \psi_{xx}) + \frac{\rho}{2e^{-4\psi}}(\Omega_\rho^2 + \Omega_x^2) &= 0 \\ \Omega_{\rho\rho} + \Omega_{xx} - 4(\psi_\rho \Omega_\rho + \psi_x \Omega_x) + \rho^{-1} \Omega_\rho &= 0 \end{aligned}$$

Sustituyendo a las derivadas de Ω por las de ω obtenemos que la primera ecuación es una de las ecuaciones "principales" de Einstein (ver §B.2):

$$\psi_\rho + \rho(\psi_{\rho\rho} + \psi_{zz}) + \frac{e^{4\psi}}{2\rho}(\omega_\rho^2 + \omega_z^2) = 0$$

Por otro lado, con la citada sustitución, la segunda de las ecuaciones es simplemente una identidad ($0 = 0$). Como $\xi \in C^2$, entonces $\Omega_{\rho z} = \Omega_{z\rho}$ que es, justamente, la segunda de las ecuaciones "principales" de Einstein. De tal manera, el sistema (5.35), con $\eta = \psi$ y $\xi = \Omega$, tomando en cuenta que $\xi \in C^2$, es equivalente al sistema de ecuaciones "principales" axisimétricas estacionarias.

El sistema (5.36) también es equivalente al hacer $\eta = \psi$ y $\xi = \Omega$; de facto es el más evidente y, en un lenguaje coloquial, el más puro. Según (5.36), la densidad del modelo proviene de la transformación $(\beta, \eta, D\beta, D\eta) \mapsto (\beta, \eta, \Pi_\beta, D\eta)$; entonces la sustitución de los campos η y ξ implica que $\beta = \omega$. El sistema (5.36) queda entonces como:

$$\psi_\rho + \rho(\psi_{\rho\rho} + \psi_{zz}) + \frac{\rho}{2e^{-4\psi}}(\Omega_\rho^2 + \Omega_z^2) = 0$$

$$D\omega = \rho e^{-4\psi} \Pi_\omega, \quad \Pi_\omega = \tilde{D}\Omega$$

donde $\Pi_\omega = \tilde{D}\Omega$ se sigue del hecho de que $\Omega \in C^2$; entonces:

$$D\omega = \rho e^{-4\psi} \tilde{D}\Omega \Rightarrow \omega_\rho = -\rho e^{-4\psi} \Omega_z, \quad \omega_z = \rho e^{-4\psi} \Omega_\rho$$

Tenemos entonces que las ecuaciones "principales" de Einstein axisimétricas estacionarias son las del sistema (5.36).

Al sistema (5.37) lo descartamos por la siguiente razón. La densidad de la cual provienen las ecuaciones de campo (5.37), a grosso modo, está en términos de un impulso + campo + derivada de un campo, estructura que no se presenta en ninguna de las densidades asociadas al caso axisimétrico estacionario.

En resumen, los sistemas (5.35) y (5.36) son equivalentes a las ecuaciones "principales" axisimétricas estacionarias de Einstein. El sistema (5.36), que nos proporciona la relación entre Ω y ω (hay que pedir menos para comparar), es de facto el que proviene de la densidad Routhiana axisimétrica estacionaria; por lo tanto, elegimos a este último para modelar el caso axisimétrico estacionario. En otras palabras, *las ecuaciones de campo del modelo σ no lineal $G = SI(2, \mathbb{R})$, $H = SO(2)$ son las ecuaciones "principales" de Einstein axisimétricas estacionarias en el vacío, cuando $\eta = \psi$, $\xi = \Omega$ y $(D\Omega)^2$ es el cuadrado de un impulso.*

Nótese que una solución de las ecuaciones "principales" de Einstein en el vacío se puede entonces representar geoméricamente como una sección $g(\rho, z) \in SI(2, \mathbb{R})$. En consecuencia, un cambio de ésta sección bajo el grupo de norma $H = SO(2)$ (transformación de norma) nos proporciona una nueva sección s' que también representa una solución para las ecuaciones "principales" de Einstein axisimétricas estacionarias. Una familia de tales secciones es entonces una familia de soluciones a las ecuaciones "principales" axisimétricas estacionarias.

Conclusiones.

En el presente trabajo de tesis hemos estudiado con cierto detalle algunos de los aspectos fundamentales en topología, geometría diferencial, conexiones, grupos y álgebras de Lie, con el objetivo de adquirir los conocimientos necesarios para comprender la definición de un modelo σ no lineal, ser capaces de construir, con estas nociones básicas, la densidad lagrangiana para un modelo de este estilo y dar una aplicación enmarcada en el ámbito de la Teoría General de la Relatividad.

El punto de partida fue el estudio de la topología, pues dentro de ésta rama de las matemáticas es que se define, dotando a un conjunto con determinadas propiedades, a un espacio de la manera más primitiva (espacio topológico) y se investiga la estructura del mismo. La topología nos proporciona entonces el escenario básico de representación y el estudio de su estructura; es a este nivel que el concepto de continuidad adquiere sentido y queda determinado de manera general. Imponiendo todavía más estructura a este espacio primitivo es que surgen las variedades diferenciables, sobre éstas se extienden de manera generalizada los conceptos de cálculo vectorial y tensorial; y, en consecuencia, es en este punto donde radica su importancia en el contexto de la física, ya que es posible representar y analizar eventos dinámicos desde una perspectiva estrictamente geométrica, pues las fuerzas que los generan son "traducidas" a deformaciones en nuestra variedad. A partir de estas estructuras geométricas es que construimos a los haces fibrados, que son la base geométrica de un modelo σ no lineal.

Notamos que como casos particulares de variedades se encuentran los grupos de Lie. Al estudiarlos, observamos el importante papel que desempeñan sus álgebras en la construcción y propiedades de estos grupos; se expuso el concepto de una representación para grupos de Lie, definiéndose que significa que éstas sean libres, transitivas, fieles y/o lineales. Aprovechando la enorme herramienta matemática con la que hasta ese momento contábamos, desarrollamos el tema de conexiones. Éstas nos permitieron formalizar la idea intuitiva de curvatura y de

transporte paralelo, en este desarrollo surgió la forma de conexión (ingrediente fundamental en la construcción de la densidad lagrangiana para nuestro modelo σ no lineal) y dimos sentido geométrico a las transformaciones de norma. Finalmente, construimos el elemento de línea general para el caso axisimétrico estacionario, presentamos la densidad lagrangiana a partir de la cual se derivan las ecuaciones de campo para éste caso y expusimos a la densidad Routhiana que nos proporciona las ecuaciones "principales" de Einstein axisimétricas estacionarias en vacío; todo ello con el objetivo de preparar el terreno para nuestra aplicación. Enseguida definimos al modelo σ no lineal y construimos la densidad lagrangiana para el mismo; en la construcción se hace patente la necesidad del estudio matemático previo, pues la misma requiere de prácticamente todo el material expuesto, como podemos observar en síntesis:

La base geométrica es un haz fibrado principal, que tiene como espacio total un grupo de Lie G (cuya representación se dio en sí mismo para garantizar la transitividad) que es el grupo de simetría global de la densidad lagrangiana del modelo, y como base el espacio cociente G/H , con H un divisor normal de G , donde toman valores los campos invariantes de norma. Este espacio base es equivalente a una variedad que no es un espacio vectorial, y que queda definida por las constricciones de los campos físicos derivadas de la representación adjunta de G en el álgebra de Lie $\mathcal{H}$. La descomposición de la 1-forma de Maurer-Cartan en una 1-forma que toma valores en el álgebra de Lie del divisor normal, $\tilde{A}$, y en otra que toma valores en el complemento a esta álgebra, $\tilde{B}$, nos permitió construir cantidades invariantes de norma, simplemente tomando la traza del producto de las componentes para la 1-forma $\tilde{B}$ y la 2-forma de curvatura F , ésta última se construye a partir de la 1-forma $\tilde{A}$, que se transforma como el potencial de norma, para el grupo de norma H , de nuestro haz fibrado principal. La densidad lagrangiana del modelo σ no lineal debe ser un escalar, en consecuencia la definimos como la contracción de la citada traza con el tensor métrico del problema particular sujeto a investigación, módulo la raíz del determinante del tensor métrico. Se probó que esta construcción efectivamente corresponde a un modelo σ no lineal y que, por lo tanto, cualquier densidad derivada de esta manera queda enmarcada dentro de estos modelos; i.e. los campos que la componen están sujetos a constricciones no lineales y tanto las constricciones como la densidad lagrangiana son invariantes bajo un grupo de simetría global.

Posteriormente se aplicó el modelo en Teoría General de la Relatividad, para ello elegimos como grupo de simetría global al grupo de Lie $Sl(2, \mathbb{R})$, con divisor normal $SO(2)$, introduciendo campos con dependencia exclusivamente en ρ y z para garantizar una condición indispensable de axisimetría y se fijó una norma en el haz fibrado principal con espacio to-

tal $SI(2, \mathbb{R})$. Llevamos a cabo la construcción de la densidad a través de la 1-forma $\tilde{B}$ y como resultado de un estudio comparativo de las ecuaciones de campo obtuvimos que las ecuaciones del modelo σ no lineal son las ecuaciones "principales" de Einstein axisimétricas estacionarias en vacío.

Con esta breve recapitulación hemos resumido la estrategia general empleada para abordar y cumplir nuestro principal propósito: el estudio de la estructura geométrica de los modelos σ no lineales. Cabe señalar que este estudio es sólo un primer paso, pues las componentes del modelo, que en sí mismas representan una enorme materia de estudio, y el modelo mismo, pueden ser extendidos. Sin embargo, esta introducción a los modelos σ no lineales ya nos permite trazar futuras líneas de investigación; por ejemplo, podemos pensar en una cuantización basada en álgebras no conmutativas ó en ciertas deformaciones del modelo que podrían conducirnos a una posible "cuantización geométrica". Podemos estudiar también de manera independiente sistemas integrables y aplicarlos a los modelos σ no lineales, con la idea de conseguir una posible cuantización de los modelos σ no lineales como sistemas integrables [15]. En sí, prácticamente todas las futuras líneas de investigación, en cuanto a modelos σ no lineales se refiere, están enfocadas a la cuantización; con ello, queda claro que este es apenas el comienzo de un largo, y tal vez fructífero, camino por recorrer en la física teórica actual.

Apéndice A

Variedades orientables y partición de unidad.

§A.1 Variedades orientables.

Consideremos un espacio vectorial V de dimensión n . Las bases ordenadas $B_1 = \{v_1, \dots, v_n\}$ y $B_2 = \{v'_1, \dots, v'_n\}$ para V determinan un isomorfismo $f: V \rightarrow V$ con $f(v_i) = v'_i$; la matriz A de f con componentes a_{ij} nos proporciona la siguiente ecuación matricial: $v'^i = a_{ij}v^j$.

Diremos que las bases B_1 y B_2 están *igualmente orientadas* si $\det(A) > 0$ y *opuestamente orientadas* si $\det(A) < 0$. La relación de ser igualmente orientadas es una relación de equivalencia, dividiendo la colección de todas las bases orientadas en dos clases de equivalencia.

De lo anterior se infiere que una clase de equivalencia es entonces una orientación para V . La clase de equivalencia para la base $\{v_1, \dots, v_n\}$ la podemos denotar como $[v_1, \dots, v_n]$, de manera tal que si ξ es una orientación de V , entonces $\{v_1, \dots, v_n\} \in \xi$ sii $[v_1, \dots, v_n] = \xi$. La orientación opuesta a ξ la denotaremos simplemente como $-\xi$. La orientación $[e_1, \dots, e_n]$ para $\mathbb{R}^n$ será llamada orientación estándar [18].

Si (V, ξ) y (W, η) son dos espacios vectoriales n -dimensionales con orientaciones ξ y η respectivamente, un isomorfismo f entre estos dos espacios ($f: V \rightarrow W$) se dice que preserva orientación (respecto a ξ y η) si $[f(v_1), \dots, f(v_n)] = \eta$ siempre que $[v_1, \dots, v_n] = \xi$.

Sea el haz trivial $E = X \times \mathbb{R}^n$ ($p \in E \Rightarrow p = (x \in X, v \in \mathbb{R}^n)$). En cada fibra podemos poner la orientación estándar: $\{(x, e_1), \dots, (x, e_n)\}$. Si $f: E \rightarrow E$ es una equivalencia y X es conexo, entonces f preserva la orientación o la revierte sobre cada fibra al tomar funciones

$\alpha_{ij} : X \rightarrow \mathbb{R}$ como sigue:

$$f(x, e_i) = \sum_{j=1}^n \alpha_{ij}(x) (x, e_j)$$

lo cual implica que $\det(A) : X \rightarrow \mathbb{R}$ es continua y nunca nula.

Si el haz no es trivial, una orientación ξ de E es una colección de orientaciones ξ_p , para $\Pi^{-1}(p)$, que satisface la siguiente condición de compatibilidad para cualquier conjunto abierto y conexo $U \subset B$ [18]:

Si $t : \Pi^{-1}(U) \rightarrow U \times \mathbb{R}^n$ es una equivalencia ($\det(t) > 0$ o $\det(t) < 0$) y las fibras de $U \times \mathbb{R}^n$ tienen la orientación estándar, entonces t preserva o revierte la orientación sobre todas las fibras.

Notemos entonces que las orientaciones ξ_p definen una orientación de E si la condición de compatibilidad se cumple para una colección de conjuntos U que cubren a la base B :

Sea $C = \{U_i, i = 1, \dots, n\}$ una cubierta de B , donde para cada elemento de C se cumple la condición de compatibilidad. Sobre U_k tendremos entonces a t_k , y sobre U_{k+1} a t_{k+1} ; dado que $U_k \cap U_{k+1} \neq \emptyset$, entonces t_k y t_{k+1} satisfacen la misma condición en cuanto a orientabilidad. En efecto:

Como las fibras en $(U_k \cap U_{k+1}) \times \mathbb{R}^n$ tienen la orientación estándar $\Rightarrow \det(t_k \circ t_{k+1}^{-1}) > 0 \Rightarrow t_k$ y t_{k+1} satisfacen la misma condición.

En consecuencia, t_k y t_{k+1} preservan o revierten la orientación en $U_k \cup U_{k+1}$, como estos elementos de C son cualesquiera, entonces se sigue la validez para todo B . Por lo tanto, hemos dado una orientación sobre todas las fibras de B (i.e.: hemos orientado a E con ξ_n orientaciones).

Si el haz E tiene una orientación $\xi = \{\xi_p\}$, entonces tiene también una orientación $-\xi = \{-\xi_p\}$; sin embargo, no todo haz tiene orientación.

Diremos que un haz es orientable si este tiene una orientación, en contraparte, un haz sin orientación será llamado haz no orientable. Un haz orientado es la pareja dada por (E, ξ) . La definición anterior es aplicable al haz tangente TM de una variedad $M \in C^\infty$. Se dice que M es orientable (no orientable) si TM es orientable (no orientable): una orientación de TM es también llamada una orientación de M , una variedad orientada es la pareja (M, ξ) .

§A.2 Partición de unidad.

Lema A.1 Sea $C \subset U \subset M$, con C compacto y U abierto. Entonces existe una función $f: M \rightarrow [0, 1]$ (C^∞) tal que $f = 1$ sobre C y la cerradura de $\text{sop}(f) = \{p | f(p) \neq 0\}$ (soporte de f) está contenida en U .

Demostración: La demostración consiste en exhibir a tal función f .

Para cada $p \in C$ elegimos un sistema coordenado alrededor de p : (X, V) ; adicionalmente pediremos que $\bar{V}$ este contenida en U y que $X(p) = 0$. Entonces $(-\epsilon, \epsilon) \times \dots \times (-\epsilon, \epsilon) \subset X(V)$ para alguna $\epsilon > 0$. Sea $j: \mathbb{R} \rightarrow \mathbb{R}$ ($j \in C^\infty$) definida por:

$$j(x) = \begin{cases} e^{-(x-1)^{-2}} e^{-(x+1)^{-2}} & x \in (-1, 1) \\ 0 & x \notin (-1, 1) \end{cases}$$

Definamos ahora a $g: \mathbb{R}^n \rightarrow \mathbb{R}$ como sigue: $g(x) = j(x^1/\epsilon)j(x^2/\epsilon)\dots j(x^n/\epsilon)$, esta función es C^∞ , positiva sobre $(-\epsilon, \epsilon) \times \dots \times (-\epsilon, \epsilon)$ y nula fuera de tal región.

La función $g \circ X$ es entonces C^∞ sobre V . Podemos extender a la función fuera de V , preservando el carácter C^∞ , diciendo simplemente que su valor es nulo. Sea f_p la función extendida. Para todo p en C es posible llevar a cabo el mismo procedimiento y construir la f_p correspondiente, esta función es positiva en una vecindad de p cuya cerradura esta contenida en U . Dado que C es compacto podemos cubrirlo con una cantidad finita de dichas vecindades, a saber las que corresponden a los n puntos $\{p(1), \dots, p(n)\}$; el soporte de la suma $f_{p(1)} + \dots + f_{p(n)}$ esta contenido en U ya que el soporte de cada una de las funciones esta contenido en U . Dado que C es cubierto con vecindades tales que en su cerradura $f_p > 0$, entonces sobre C la suma $f_{p(1)} + \dots + f_{p(n)}$ es positiva. Digamos que esta suma es mayor o igual que δ para alguna $\delta > 0$.

Sea $l: \mathbb{R} \rightarrow \mathbb{R}$ definida como:

$$l(x) = \int_0^x k / \int_0^\delta k$$

con $k: \mathbb{R} \rightarrow \mathbb{R}$ positiva sobre $(0, \delta)$ y nula fuera de tal intervalo.

Sea $f = l \circ (f_{p(1)} + \dots + f_{p(n)})$; es decir:

$$f = \int_0^{f_{p(1)} + \dots + f_{p(n)}} k / \int_0^\delta k$$

Como en C resulta que $f_{p(1)} + \dots + f_{p(n)} \geq \delta$, entonces para todo p en C tenemos que:

$$f = \int_0^{\infty} k / \int_0^{\delta} k = 1$$

La función f decae suavemente en la región compuesta por puntos fuera de C pero tal que se cumple que $0 < f_{p(1)} + \dots + f_{p(n)}$. Por otro lado, si $f_{p(1)} + \dots + f_{p(n)} = 0$, entonces $f = 0$.

De tal manera tenemos que $f : M \rightarrow [0, 1]$, es C^∞ dado que es la composición de funciones C^∞ , $f = 1$ sobre C y el soporte de f esta contenido en U ya que $\text{sop}(f)$ es la unión de los soportes de $f_{p(1)}, \dots, f_{p(n)}$ que están contenidos en U . $\square$

Teorema A.1 Sea O una cubierta localmente finita de una variedad M . Entonces hay una colección de funciones $h_U : M \rightarrow [0, 1]$ (C^∞), una para cada U en O , tales que:

- (1).- $\text{sop}(h_U) \subset U$ para cada U .
- (2).- $\sum_U h_U(p) = 1 \quad \forall p \in M$.

Demostración:

Caso 1: Cada U en O tiene cerradura compacta.

Para mostrar este primer caso haremos uso del siguiente resultado: Sea O una cubierta localmente finita de una variedad M . Entonces es posible elegir, para cada U en O , un conjunto abierto U' con $\bar{U}'$ contenida en U , de tal forma que la colección de todos los U' es también una cubierta abierta de M [18]. Elijamos U' como en el teorema. Apliquemos el lema anterior con $\bar{U}' \subset U \subset M$ para obtener una función $A_U : M \rightarrow [0, 1]$ (C^∞) que es igual a la unidad sobre U' y tiene soporte contenido en U .

Dado que U' cubre a M , entonces es evidente que $\sum_{U \in O} A_U > 0$ para todo p en M . Definamos la siguiente cantidad:

$$h_U = \frac{A_U}{\sum_{U \in O} A_U}$$

$\text{sop}(h_U) \subset U$ ya que $\text{sop}(h_U) = \text{sop}(A_U)$, $h_U \in C^\infty$ dado que $A_U \in C^\infty$ y $\sum_{U \in O} A_U \neq 0$. Notemos que:

$$\sum_{U \in O} h_U = \sum_{U \in O} \frac{A_U}{\sum_{U \in O} A_U} = \frac{\sum_{U \in O} A_U}{\sum_{U \in O} A_U} = 1$$

lo cual muestra que $\sum_U h_U = 1$ para todo p en M . Adicionalmente es claro que $h_U : M \rightarrow [0, 1]$.

Caso 2: Caso general.

Estamos de acuerdo en que la cerradura de U no tiene porque ser un compacto, ello implica que la cerradura de U' no necesariamente es un compacto (U' es cerrado). En consecuencia, para demostrar de manera general el teorema, habrá de probarse que el lema anterior es cierto también cuando C (U') no es compacto. Al demostrar tal afirmación, entonces podremos dar las funciones h_U como en el caso 1 sin mayores obstáculos:

Lema A.2 Sea $C \subset K \subset M$ con C cerrado (pero no necesariamente compacto) y K abierto. Entonces existe una función $f : M \rightarrow [0, 1]$ (C^∞) tal que $f = 1$ sobre C y la cerradura de $\text{sop}(f)$ está contenida en K .

Demostración: Exhibamos a tal f .

Para todo p en C demos un $U_p \subset K$ abierto tal que $\bar{U}_p \subset K$ sea compacto. Cubramos al complemento de C con V_α abiertos contenidos en el complemento de C tales que $\bar{V}_\alpha$ sean compactos.

Recurramos al siguiente teorema que dice: Si C es una cubierta abierta de M , entonces existe una cubierta C' de M tal que C' es localmente finita y refina a C [18].

Entonces, la cubierta abierta $\{U_p, V_\alpha\}$ tiene una cubierta abierta O que la refina y es localmente finita. El refinamiento implica que para todo U_p existe el conjunto $O' = \{U \in O \mid U \subset U_p \text{ para alguna } p\}$. Como U_p es compacto entonces $\bar{O}'$ es compacto. Del caso 1 se sigue la construcción de h_U para cada $U \in O'$.

Sea $f = \sum_{U \in O'} h_U$. Esta expresión es C^∞ puesto que es la suma finita (O es localmente finita) de funciones C^∞ .

Dado que $\sum_U h_U(p) = 1$ para todo $p \in U \subset O'$ y $h_U(p) = 0$ cuando $p \in U \subset V_\alpha$, entonces $f(p) = 1 \forall p \in C$. Como O es localmente finita, la suma es finita en una vecindad de cada punto; por construcción la cerradura del soporte no abarca mas allá de $\bigcup \bar{U}_p \subset K \Rightarrow$ la cerradura de $\text{sop } f \subset K$. $\square$

Corolario A.1 Si O es cualquier cubierta abierta de una variedad M , entonces hay una colección de funciones C^∞ $h_i : M \rightarrow [0, 1]$ tales que:

- (1).- La colección de conjuntos $\{p \mid h_i(p) \neq 0\}$ es localmente finito.
- (2).- $\sum_i h_i(p) = 1 \forall p \in M$.
- (3).- Para cada i hay una $U \in O$ tal que $\text{sop}(h_i) \subset U$.

Una colección $\{h_i : M \rightarrow [0, 1]\}$ que cumpla (1) y (2) es llamada una partición de unidad; si además se satisface (3) entonces se dice que la partición esta subordinada a O [18].

Apéndice B

Densidad de Routh y ecuaciones “principales” de Einstein axisimétricas estacionarias.

§B.1 Densidad de Routh.

Comencemos considerando una densidad lagrangiana $\mathcal{L}$ que depende de las coordenadas generalizadas q y s y de sus primeras derivadas $\dot{q}$ y $\dot{s}$ [10]. Al llevar a cabo la transformación $(q, s, \dot{q}, \dot{s}) \mapsto (q, s, p, \dot{s})$, donde p es el momento generalizado correspondiente a q , observamos que la diferencial de la densidad lagrangiana está dada por:

$$\begin{aligned}d\mathcal{L} &= \mathcal{L}_{,q} dq + \mathcal{L}_{,q} d\dot{q} + \mathcal{L}_{,s} ds + \mathcal{L}_{,s} d\dot{s} \\ &= \dot{p}dq + p d\dot{q} + \mathcal{L}_{,s} ds + \mathcal{L}_{,s} d\dot{s}\end{aligned}$$

como $p d\dot{q} = dp\dot{q} - \dot{q}dp$, entonces:

$$d(\mathcal{L} - p\dot{q}) = \dot{p}dq - \dot{q}dp + \mathcal{L}_{,s} ds + \mathcal{L}_{,s} d\dot{s}$$

Definimos a la densidad de Routh como $\mathcal{R}(q, p, s, \dot{s}) \equiv p\dot{q} - \mathcal{L}$, donde $p = \mathcal{L}_{,\dot{q}}$. En consecuencia:

$$d\mathcal{R} = -\dot{p}dq + \dot{q}dp - \mathcal{L}_{,s} ds - \mathcal{L}_{,s} d\dot{s}$$

y por lo tanto tenemos que:

$$\begin{aligned} \dot{q} &= \mathcal{R}_{,p} \quad , \quad \dot{p} = -\mathcal{R}_{,q} \\ \mathcal{L}_{,s} &= -\mathcal{R}_{,s} \quad , \quad \mathcal{L}_{,t} = -\mathcal{R}_{,t} \end{aligned} \quad (\text{B.1})$$

Puesto que $\frac{d}{dt}(\mathcal{L}_{,s}) = \mathcal{L}_{,st}$, entonces las ecuaciones que se derivan de la densidad Routhiana quedan finalmente como sigue:

$$\begin{aligned} \dot{q} &= \mathcal{R}_{,p} \quad , \quad \dot{p} = -\mathcal{R}_{,q} \\ \frac{d}{dt} \left(\frac{\partial \mathcal{R}}{\partial s} \right) &= \frac{\partial \mathcal{R}}{\partial t} \end{aligned} \quad (\text{B.2})$$

La densidad Routhiana es útil cuando existen coordenadas cíclicas, pues el momento asociado será una cantidad conservada. De tal manera, si la transformación es $(s, \dot{q}, \dot{s}) \rightarrow (s, \Pi_q, \dot{s})$, sabemos que Π_q (momento generalizado correspondiente a la coordenada cíclica q) es una cantidad conservada y las ecuaciones resultantes de la densidad Routhiana son las de Euler-Lagrange para la(s) coordenada(s) no cíclica(s) y las de Hamilton para la(s) cíclica(s).

Ahora bien, si la densidad lagrangiana es como en (5.23):

$$\mathcal{L}' = 2D\rho D\gamma + \frac{e^{4\psi}}{2\rho} (D\omega)^2 - 2\rho (D\psi)^2$$

entonces $\mathcal{L}' = \mathcal{L}'(\psi, D\gamma, D\omega, D\psi)$; puesto que γ y ω son cíclicas:

$$\mathcal{R} = \Pi_\gamma D\gamma + \Pi_\omega D\omega - \mathcal{L}'$$

donde $\Pi_\kappa = (\mathcal{L}'),_{D\kappa}$ son los momentos asociados a las "coordenadas" $\kappa = \{\gamma, \omega\}$. Entonces:

$$\mathcal{R} = 2\rho (D\psi)^2 + \frac{\rho}{2} e^{-4\psi} \Pi_\omega^2$$

donde evidentemente $\mathcal{R} = \mathcal{R}(\psi, \Pi_\omega, D\psi)$. Por (B.1)se sigue que las ecuaciones son:

$$D\gamma = \frac{\partial \mathcal{R}}{\partial \Pi_\gamma} \quad , \quad D\Pi_\gamma = -\frac{\partial \mathcal{R}}{\partial \gamma}$$

$$D\omega = \frac{\partial \mathcal{R}}{\partial \Pi_\omega} \quad , \quad D\Pi_\omega = -\frac{\partial \mathcal{R}}{\partial \omega}$$

$$(\mathcal{L}'),_{\psi} = -(\mathcal{R}),_{\psi} \quad , \quad (\mathcal{L}'),_{D\psi} = -(\mathcal{R}),_{D\psi}$$

puesto que la ecuación de Lagrange es $D((\mathcal{L}), D_X) = (\mathcal{L})_{,X}$ [9], entonces:

$$D\gamma = \frac{\partial \mathcal{L}}{\partial \Pi_\gamma}, \quad D\Pi_\gamma = -\frac{\partial \mathcal{L}}{\partial \gamma} = 0$$

$$D\omega = \frac{\partial \mathcal{L}}{\partial \Pi_\omega}, \quad D\Pi_\omega = -\frac{\partial \mathcal{L}}{\partial \omega} = 0$$

$$D((\mathcal{R}), D_\psi) - (\mathcal{R})_{,\psi} = 0$$

Introduciendo la función $\Omega(\rho, z)$ como $\Pi_\omega = \bar{D}\Omega$, con $\bar{D} = (-\partial_z, \partial_\rho)$, la densidad Routhiana y las ecuaciones correspondientes son las siguientes:

$$\mathcal{R} = 2\rho(D\psi)^2 + \frac{\rho}{2}e^{-4\psi}(\bar{D}\Omega)^2 \quad (\text{B.3})$$

y

$$D((\mathcal{R}), D_\psi) - (\mathcal{R})_{,\psi} = 0 \quad (\text{B.4})$$

$$D\omega = (\mathcal{R})_{,\bar{D}\Omega}, \quad D(\bar{D}\Omega) = 0$$

§B.2 Ecuaciones “principales” de Einstein axisimétricas estacionarias.

Las ecuaciones de Einstein en el vacío están dadas por el siguiente conjunto de ecuaciones diferenciales de segundo orden para la métrica g [11]:

$$R_{ik} - \frac{1}{2}g_{ik}R = 0 \quad (\text{B.5})$$

donde R_{ik} y g_{ik} son las componentes covariantes del tensor de Ricci y de la métrica g respectivamente, en tanto que R es el escalar de curvatura.

Notemos que estas ecuaciones se reducen simplemente a $R_{ik} = 0$. En efecto:

$$\begin{aligned} R_{ik} - \frac{1}{2}g_{ik}R = 0 &\Rightarrow g^{ik}R_{ik} - \frac{1}{2}g^{ik}g_{ik}R = 0 \\ &\Rightarrow R - 2R = 0 \Rightarrow R = 0 \end{aligned}$$

puesto que $g^{ik}R_{ik} = R$ y $g^{ik}g_{ik} = 4$; en consecuencia, las ecuaciones de Einstein están dadas por:

$$R_{ik} = 0 \quad (\text{B.6})$$

Las ecuaciones de campo que de (B.6) se derivan para el elemento de línea (5.11) son las siguientes:

$$\begin{aligned} R_{00} &= -\frac{e^{4\psi-2\gamma}}{\rho} \{ \rho(\psi_{\rho\rho} + \psi_{zz}) + \psi_{\rho} + \frac{e^{4\psi}}{2\rho} (\omega_{\rho}^2 + \omega_z^2) \} = 0 \\ R_{03} + \omega^{-1} R_{33} &= -\frac{e^{4\psi-2\gamma}}{2} \{ 4(\psi_{\rho}\omega_{\rho} + \psi_z\omega_z) + (\omega_{\rho\rho} + \omega_{zz}) - \rho^{-1}\omega_{\rho} \} \\ &\quad - \frac{\rho e^{-2\gamma}}{\omega} \{ \rho(\psi_{\rho\rho} + \psi_{zz}) + \psi_{\rho} + \frac{e^{4\psi}}{2\rho} (\omega_{\rho}^2 + \omega_z^2) \} = 0 \\ R_{12} &= \frac{1}{\rho} \{ 2\rho\psi_{\rho}\psi_z - \frac{e^{4\psi}}{2\rho} (\omega_{\rho}\omega_z - \gamma_z) \} = 0 \\ R_{11} - R_{22} &= \frac{2}{\rho} \{ \rho(\psi_{\rho}^2 - \psi_z^2) - \gamma_{\rho} + \frac{e^{4\psi}}{4\rho} (\omega_z^2 - \omega_{\rho}^2) \} = 0 \end{aligned}$$

y por lo tanto las ecuaciones de Einstein axisimétricas estacionarias en el vacío son:

$$\psi_{\rho} + \rho(\psi_{\rho\rho} + \psi_{zz}) + \frac{e^{4\psi}}{2\rho} (\omega_{\rho}^2 + \omega_z^2) = 0 \quad (\text{B.7})$$

$$4(\psi_{\rho}\omega_{\rho} + \psi_z\omega_z) + (\omega_{\rho\rho} + \omega_{zz}) - \rho^{-1}\omega_{\rho} = 0$$

y

$$\begin{aligned} \gamma_z &= 2\rho\psi_{\rho}\psi_z - \frac{e^{4\psi}}{2\rho}\omega_{\rho}\omega_z \\ \gamma_{\rho} &= \rho(\psi_{\rho}^2 - \psi_z^2) + \frac{e^{4\psi}}{4\rho} (\omega_z^2 - \omega_{\rho}^2) \end{aligned} \quad (\text{B.8})$$

Nótese que al resolver el sistema (B.7) de ecuaciones diferenciales de segundo orden acopladas, la solución al sistema (B.8) se reduce simplemente a un problema de integración. Es por esta razón que a las ecuaciones (B.7) habremos de llamarlas *ecuaciones principales de Einstein axisimétricas estacionarias*.

Obsérvese que el sistema (B.7) es equivalente al sistema (B.4). En efecto:

$$D\omega = \frac{\partial \mathcal{R}}{\partial \tilde{D}\Omega}$$

$$D\omega = \frac{\partial \mathcal{R}}{\partial \tilde{D}\Omega} = \rho e^{-4\psi} \tilde{D}\Omega$$

por lo tanto:

$$\Omega_\rho = \rho^{-1} e^{4\psi} \omega_z, \quad \Omega_z = -\rho^{-1} e^{4\psi} \omega_\rho$$

Por construcción sabemos que:

$$D(\tilde{D}\Omega) = 0 \Rightarrow \Omega_{\rho z} - \Omega_{z\rho} = 0$$

$$\Rightarrow \rho^{-1} e^{4\psi} (4(\psi_\rho \omega_\rho + \psi_z \omega_z) + (\omega_{\rho\rho} + \omega_{zz}) - \rho^{-1} \omega_\rho) = 0$$

y por lo tanto:

$$4(\psi_\rho \omega_\rho + \psi_z \omega_z) + (\omega_{\rho\rho} + \omega_{zz}) - \rho^{-1} \omega_\rho = 0$$

Por otra parte, la primera de las ecuaciones de (B.7) es justamente la ecuación lagrangiana correspondiente a la densidad Routhiana (B.3):

$$D\left(\frac{\partial \mathcal{R}}{\partial D\psi}\right) - \frac{\partial \mathcal{R}}{\partial \psi} = 0$$

$$\Rightarrow 4D\rho D\psi + 4\rho D^2\psi + 2\rho e^{-4\psi} (D\Omega)^2 = 0$$

donde se ha tomado en cuenta que $(D\Omega)^2 = (\tilde{D}\Omega)^2$. Aplicando el operador D y escribiendo las derivadas de Ω en función de las correspondientes en ω obtenemos que:

$$4\psi_\rho + 4\rho(\psi_{\rho\rho} + \psi_{zz}) + \frac{2e^{4\psi}}{\rho}(\omega_\rho^2 + \omega_z^2) = 0$$

$$\Rightarrow \psi_\rho + \rho(\psi_{\rho\rho} + \psi_{zz}) + \frac{e^{4\psi}}{2\rho}(\omega_\rho^2 + \omega_z^2) = 0$$

Con lo cual queda probado que las ecuaciones principales de Einstein axisimétricas estacionarias para el vacío son las que de la densidad Routhiana axisimétrica estacionaria se derivan.

Bibliografía

- [1] D. Maison, Phys. Rev. Lett. 41, 521 (1978)
- [2] John L. Kelley. General Topology. Springer-Verlag, Graduate Texts in Mathematics 27, 298 pp.
- [3] Yvonne Choquet-Bruhat, Cécile DeWitt-Morette, Margaret Dillard-Bleick. Analysis, Manifolds and Physics. North-Holland 1982, 630 pp.
- [4] Schutz. Bernard Schutz. Geometrical Methods of Mathematical Physics. Cambridge University Press 1980, 250 pp.
- [5] Charles Nash, Siddhartha Sen. Topology and Geometry for Physicists. Academic Press, Inc. 1983, 312 pp.
- [6] M.Göckeler, T.Schücker. Differential Geometry, Gauge Theories, and Gravity. Cambridge University Press, 230 pp.
- [7] Steven Weinberg. Gravitation and Cosmology: Principles and Applications of The General Theory of Relativity. John Wiley and Sons 1972, 657 pp.
- [8] E.C.G. Sudarshan, N.Mukunda. Classical Dynamics: A Modern Perspective John Wiley and Sons, 1974. 615 pp.
- [9] Herbert Goldstein. Classical Mechanics Addison-Wesley, Second Edition, 672 pp.
- [10] Landau y Lifshitz. Mecánica Vol 1, Reverté, S.A., Segunda Edición, 200 pp.
- [11] Landau y Lifshitz. Teoría Clásica de Campos Vol 2, Reverté, S.A., Segunda Edición, 525 pp.

- [12] Robert M. Wald. General Relativity The University of Chicago Press 1984, 491 pp.
- [13] Jamal Nazrul Islam. Rotating Fields in General Relativity. Cambridge University Press 1985, 122pp.
- [14] A.P.Balachandran, G.Marmo, B.S.Skagerstam, A.Stern.
Classical Topology and Quantum States. World Scientific, 358 pp.
- [15] H.Nicolai. Two-Dimensional Gravities and Supergravities as Integrable Systems. Lectures presented at the 30.Schladming, February 1991.
- [16] Lars V. Ahlfors. Complex Analysis. Mc Graw-Hill International Editions, Mathematics Series. Third Edition, 331 pp.
- [17] Darío Núñez, Hernando Quevedo, Alberto Sánchez.
Einstein's Equations as Functional Geodesics.
- [18] Michael Spivak. A Comprehensive Introduction to Differential Geometry. Volume I. Publish or Perish, Inc. Berkeley 1979, 668 pp.