



UNIVERSIDAD NACIONAL AUTÓNOMA DE MÉXICO
PROGRAMA DE MAESTRÍA Y DOCTORADO EN CIENCIAS
MATEMÁTICAS Y DE LA ESPECIALIZACIÓN EN
ESTADÍSTICA APLICADA

ESTRUCTURAS GEOMÉTRICAS SOBRE ESPACIOS
CUÁNTICOS

TESIS

QUE PARA OBTENER EL GRADO DE:
DOCTORA

PRESENTA:

PERLA CECILIA LUCIO PEÑA

DIRECTOR DE TESIS:

DR. MICHIO ĐURĐEVICH
INSTITUTO DE MATEMÁTICAS, UNAM

MIEMBROS DEL COMITÉ TUTOR:

DR. OSCAR ADOLFO SÁNCHEZ VALENZUELA
CIMAT, MÉRIDA

DR. ZBIGNIEW OZIEWICZ KWASS †
FES-CUAUTITLÁN, UNAM

CIUDAD DE MÉXICO, MAYO DEL 2024



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.

Dedicatoria

Quisiera expresar mi más profundo agradecimiento a mis padres Ma. Nieves Peña Rodríguez y Marcelo Lucio Aguilar, por su amor incondicional y apoyo moral. Han sido uno de los pilares de este logro, muchas gracias por toda su ayuda y por siempre estar. Gracias a Jonathán Rivera y a mis perritas Eipi y Galoa por que me ayudaron notablemente a sobrellevar el estrés y a superar los momentos difíciles, gracias por su amor, apoyo y acompañamiento en este viaje académico.

Gracias infinitas a mi asesor de tesis, el Dr. Micho Đurđevich, ya que con su apoyo y paciencia contribuyó de manera muy significativa en mi desarrollo y camino hacia la investigación. Gracias por todas esas tardes donde me aclaró dudas sobre las maravillosas teorías de cálculos diferenciales cuánticos y haces principales cuánticos. Sin duda, su guía constante y su fe inquebrantable en mis habilidades me han ayudado a entender el fantástico mundo cuántico.

También quiero expresar mi agradecimiento a los miembros de mi comité tutorial, Dr. Pierre Michel Bayard, Dr. Raymundo Bautista, Dra. María del Carmen González Vallarte, Dr. Oscar Adolfo Sánchez Valenzuela, Dr. Zbigniew Oziewicz Kwass y Dr. Gregor Weingart por tomarse el tiempo de leer la tesis y enviarme sus valiosos comentarios, sugerencias y observaciones.

Un agradecimiento a la Universidad Nacional Autónoma de México por abrirme las puertas y brindarme la oportunidad de avanzar en mi carrera profesional, en particular al Instituto de Matemáticas y al Posgrado en Ciencias Matemáticas por su apoyo constante.

Contenido

1	Introducción	3
2	Antecedentes	11
2.1	Física cuántica	11
2.2	Geometría cuántica	14
3	Teoría de álgebras C^*	21
3.1	Álgebras C^*	21
3.2	Homomorfismos e ideales	24
3.3	Unitalización	26
3.4	El espectro de un elemento	27
3.5	Espectro de álgebras de Banach conmutativas	33
3.6	Teorema de Gelfand-Naimark	35
3.7	Estados y representaciones	40
3.8	Simetrías y grupos cuánticos	46
4	Cálculos diferenciales	55
4.1	Cálculo diferencial sobre espacios cuánticos	56
4.2	Cálculo diferencial sobre grupos cuánticos	61
4.3	Álgebra tensorial	83
4.4	Álgebra exterior	87
4.5	Cálculo diferencial de más alto orden	90
4.5.1	Cálculo exterior trenzado	90
4.5.2	Cálculo envolvente universal diferencial	96
5	Haces principales	107
5.1	Haces principales clásicos y cuánticos	109
5.2	Conexiones cuánticas	116
5.3	Derivada covariante y curvatura	122

6	Haz hiperbólico cuántico de Hopf	131
6.1	Grupo cuántico $SU_\mu(1, 1)$	131
6.1.1	Cálculo diferencial cuántico de primer orden sobre $\mathfrak{B} = SU_\mu(1, 1)$	133
6.1.2	Envolvente universal diferencial sobre $SU_\mu(1, 1)$	135
6.1.3	Cálculo diferencial sobre el círculo $U(1)$	138
6.2	Haz principal cuántico hiperbólico	139
6.2.1	Conexiones, derivada covariante y curvatura	143
6.3	Modelo cuántico de Poincaré	145
6.4	Planos hiperbólicos cuánticos	159
6.5	Representaciones en espacios de Hilbert	162
6.5.1	Kernel reproductor y métrica inducida en el modelo cuántico hiperbólico	167
6.6	Cálculo diferencial	171
6.6.1	Conexiones, derivada covariante y curvatura	174
6.6.2	Solución universal	176
6.7	Conclusiones	179

Capítulo 1

Introducción

Uno de los más grandes logros conceptuales de la humanidad en el siglo XX, es la teoría cuántica. Con ella se establece un nuevo paradigma para entender el universo y nuestro papel en él. ¿Cómo es el mundo en pequeñas escalas? ¿Se pueden separar las cosas? ¿Qué es un átomo? ¿Pueden ocurrir fenómenos sin que tengan una causa? ¿Cómo es la relación entre la parte y el todo? Estas preguntas encuentran nuevas respuestas en la teoría cuántica, y la visión del mundo es una visión holística, donde todo puede ser entrelazado con todo, y donde la diferencia entre aquí y allá es superficial.

Y es que, la mera existencia de una longitud universal para la mecánica cuántica, la longitud de Planck, como que nos sugiere que para describir los espacios en estas escalas, debemos considerar una nueva geometría donde se pueda definir una longitud absoluta. En esta geometría, no se deben incluir a las homotecias como parte de las simetrías del espacio, pues se requiere distinguir lo grande de lo pequeño, y para evitar dificultades de divergencias de algunas cantidades físicas, como lo es la energía, quizá debemos suponer que el espacio no se comporta como una variedad clásica en las pequeñas escalas.

La geometría cuántica también conocida como geometría no conmutativa, propone romper con la suposición de que en los espacios se pueden definir coordenadas y nos invita a pensarlos, como espacios formados de tejidos que no se pueden fragmentar infinitamente. Los espacios cuánticos están descritos por $*$ -álgebras no conmutativas, cuyos elementos son “funciones” de cierta naturaleza (por ejemplo continuas) sobre el espacio cuántico asociado. Y cuando las álgebras son conmutativas, estamos de regreso en la geometría clásica.

Entonces podemos decir que la geometría cuántica es una especie de hermana platónica de la física cuántica, donde en lugar de estudiar átomos,

moléculas, partículas elementales, objetos grandes y pequeños y sus fenómenos, estudiamos los espacios. Estos espacios, son diferentes a los estudiados en la geometría estándar. Allí, el espacio siempre se ve como una *colección de puntos* equipada con la estructura apropiada adicional (ya sea, topología, atlas, métrica, sistema de rectas, planos y círculos, etc.) dependiendo del nivel geométrico de nuestras consideraciones. Sin embargo, los espacios en la geometría cuántica por lo general, no tienen ni puntos ni partes.

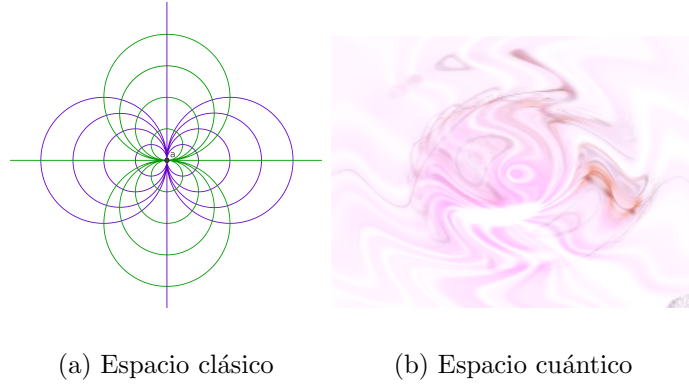


Figure 1.1: Espacios

Para describir los espacios cuánticos, usamos como principales componentes al análisis funcional, a las $*$ -álgebras no conmutativas y a las simetrías. Cuando las álgebras son C^* , éstas se pueden ver como álgebras de operadores en un espacio de Hilbert, encontrando de esta manera, enlaces directos con la teoría cuántica. Podemos pensar a los elementos hermitianos del álgebra como observables de un sistema, los vectores del espacio de Hilbert como estados del sistema y los valores propios correspondientes, como los valores que toman las observables.

Existe una biyección entre cada álgebra C^* conmutativa y el álgebra de funciones continuas que desaparecen en el infinito, sobre un espacio topológico localmente compacto naturalmente asociado al álgebra dada. De manera que cada propiedad geométrica del espacio se puede traducir a su propiedad algebraica análoga y viceversa, creando así un diccionario que relaciona dichas propiedades. Por ejemplo, decimos que el espacio asociado a una álgebra C^* conmutativa es compacto si el álgebra es unital, o podemos estudiar las simetrías del espacio a través de los automorfismos del álgebra, los puntos del espacio a través de los caracteres del álgebra, las medidas de probabilidad a través de los estados del álgebra, etc. Ampliamos esta correspondencia, permitiendo álgebras no conmutativas y dando lugar a

los espacios cuánticos. En esta ampliación podemos encontrar ahora que algunas propiedades algebraicas tales como caracter del álgebra, estado puro y clase de equivalencia de representaciones irreducibles corresponden a una misma definición del concepto de *punto*, pero lejos de hacer ambigua la teoría, solamente puede ayudar a enriquecer nuestros conocimientos del espacio asociado. Incluso la mayoría de los teoremas de geometría clásica, se pueden demostrar de una forma más sencilla y elegante usando conceptos de geometría cuántica.

Los grupos juegan un papel muy importante en el desarrollo de la geometría de los espacios. Para el caso de los espacios cuánticos, diremos que las simetrías de éstos están dadas por los automorfismos de su álgebra asociada. Aquí podemos observar y resaltar que en todas las álgebras no conmutativas siempre existen automorfismos internos, es decir, los espacios cuánticos siempre tienen simetrías. Por otro lado, a la traducción de la estructura de grupo en términos de un álgebra de funciones sobre el grupo, se le conoce como álgebra de Hopf y es a lo que llamamos grupo cuántico. Esta manera de introducir grupos cuánticos a través de las álgebras de Hopf se debe a Stanislaw Lech Woronowicz, y por nuestra parte haremos uso de estas álgebras como los análogos cuánticos de los grupos. En esta definición cuántica de grupo siempre existe al menos un punto clásico que corresponde al elemento unidad. Existen otras formulaciones de grupo que permiten la introducción de grupos cuánticos sin ningún punto. Una de ellas es el lenguaje diagramático profundizado por Micho Đurđevich. Éste se puede ver como una categoría, donde los objetos son los números naturales y los morfismos son ciertos diagramas que si los realizamos en un universo de espacios cuánticos da lugar al concepto de grupo cuántico.

Una forma de estudiar la geometría de los espacios es a través del estudio de las propiedades que no cambian cuando aplicamos un grupo de simetrías. En esta tesis, estudiamos la geometría a través de la teoría de haces principales cuánticos que podemos encontrar en los artículos de Đurđevich llamados *Geometría de haces principales cuánticos parte I, II y III* [4, 5, 6] y de la teoría de grupos cuánticos introducida por Woronowicz en sus artículos *Pseudogrupos de matrices compactas, Cálculos diferenciales en pseudogrupos de matrices compactos y Un ejemplo de cálculo diferencial no conmutativo sobre el grupo cuántico $SU(2)$* [22, 23, 24]. Con estas teorías, los objetos estudiados, como los planos cuánticos, se ven de manera natural, como los invariantes bajo la acción de un grupo que actúa sobre un haz principal cuántico. También podemos decir que al espacio base de un haz principal, le podemos asociar un cálculo diferencial que precisamente se ve como el espacio de formas horizontales invariantes.

La unificación de todo lo mencionado anteriormente da lugar a resultados

interesantes sobre los planos hiperbólicos cuánticos. La tesis se organiza de la siguiente manera: en el Capítulo 2, presentamos una breve introducción sobre la física y la geometría cuántica. A manera de motivación, hablaremos de algunos experimentos que dieron lugar a la necesidad de crear una nueva teoría para explicarlos: la mecánica cuántica. Y que el lenguaje para describirlos en su totalidad son los operadores en espacios de Hilbert, los cuales no siempre conmutan. De la mano con esto, surge la geometría cuántica que permite la introducción de nuevos espacios totalmente diferentes a los espacios clásicos, los espacios cuánticos. Podremos encontrar en este capítulo las equivalencias entre los lenguajes de álgebra y geometría. También usaremos como herramienta complementaria al entorno general no conmutativo, a los espacios de Hilbert de funciones holomorfas. Con éstos, reproducimos espacios cuánticos considerando el álgebra no conmutativa generada por los operadores de multiplicación z y su adjunto z^* . Estos espacios de Hilbert de funciones holomorfas tienen asociadas ciertas funciones núcleos reproductores, que unifican el lenguaje de puntos y ondas y que nos permite, entre otras cosas, calcular la métrica de estos espacios.

La teoría de álgebras C^* conmutativas la encontraremos en el Capítulo 3. Aquí establecemos la equivalencia entre álgebras C^* conmutativas con unidad y álgebras de funciones continuas sobre espacios topológicos compactos. Mostramos cómo algunos de los conceptos geométricos del espacio, se pueden reformular en términos de conceptos algebraicos. Por ejemplo, se muestra que los puntos del espacio se corresponden con los siguientes conceptos: con ideales maximales, caracteres, estados puros y clases de equivalencia de representaciones irreducibles del álgebra asociada, coincidiendo todos ellos en este caso conmutativo. Se muestra además por qué unitalizar un álgebra es equivalente a compactificar el espacio asociado, encontrar proyectores es equivalente a que el espacio sea desconexo, que los estados del álgebra se corresponden con medidas de probabilidad del espacio, los $*$ -automorfismos con simetrías del espacio y que las álgebras de Hopf son el análogo de la estructura de grupo.

Es aquí donde desarrollamos uno de los ejemplos más sencillos no conmutativos: el álgebra de matrices de tamaño 2×2 y coeficientes complejos. Para dicha álgebra mostramos que las diferentes generalizaciones de punto no son equivalentes y que a pesar de no contener puntos, su grupo de simetrías está determinado por el grupo de rotaciones en el espacio tridimensional real, lo que sugiere que el espacio geométrico asociado a dicha álgebra es una esfera cuántica. También, a manera de ejemplo, exploramos la estructura de álgebra de Hopf sobre el círculo unitario, que nos servirá y usaremos más adelante en las exploraciones de esta tesis.

Los Capítulos 4 y 5 están dedicados a exponer las contribuciones de

Woronowicz y Đurđevich sobre los cálculos diferenciales cuánticos y los haces principales cuánticos, que forman la base de los resultados originales de esta tesis que serán presentados en el Capítulo 6. En el Capítulo 4 veremos que un grupo cuántico puede ser dotado de diferentes cálculos diferenciales, contrario al caso clásico donde hay un único cálculo diferencial natural. Cabe mencionar, que también podemos equipar a los espacios clásicos, como el círculo o grupos finitos, de cálculos diferenciales cuánticos no triviales. El concepto de cálculos diferenciales covariantes es especialmente importante, pues se refiere a la construcción de cálculos que permiten la extensibilidad de las acciones al nivel del cálculo. Para cada cálculo covariante sobre un álgebra podemos encontrar un ideal derecho contenido en el kernel de la counidad del álgebra que reproduce el cálculo diferencial completo, si el ideal es pequeño entonces el cálculo se vuelve complicado desde el punto de vista operacional y si el ideal es grande puede acercarse a lo trivial. Cuando el álgebra es conmutativa, podemos recuperar el cálculo diferencial clásico considerando el cuadrado del núcleo de la counidad.

También podemos encontrar en este capítulo las construcciones principales de cálculo diferencial trenzado y de cálculo envolvente universal. El cálculo diferencial trenzado o exterior, directamente generaliza la construcción de De Rham en el caso clásico. Por otro lado el cálculo envolvente universal se puede obtener como el álgebra cociente del álgebra tensorial y el ideal cuadrático generado por los mínimos requerimientos de consistencia interna. Se puede demostrar que si se construye un cálculo diferencial de más alto orden sobre un grupo cuántico tal que el coproducto se extiende a un morfismo de álgebras diferenciales graduadas, entonces el cálculo universal se proyecta sobre dicho cálculo y este último a su vez se proyecta al cálculo exterior. De esta manera, estos dos cálculos representan las construcciones máximas y mínimas de cálculos diferenciales no triviales. Además en el caso de que el álgebra de Hopf sea clásica y el cálculo diferencial de primer orden también lo sea, entonces las dos construcciones de cálculos diferenciales completos coinciden.

En el Capítulo 5, abordamos la teoría de haces principales en su versión cuántica [4], [5] y [6]. Para ver el contraste, entre la teoría clásica y la cuántica presentamos algunos de los conceptos en las dos versiones. Aquí encontraremos la definición de haz principal cuántico, así como también los conceptos de conexión, curvatura y derivada covariante. Este es el lenguaje que utilizamos a lo largo de la tesis para hablar de las estructuras geométricas sobre espacios cuánticos.

Por último, en el Capítulo 6 mostramos las aportaciones originales de esta tesis. Usamos la teoría de Woronowicz para construir un cálculo diferencial

sobre el grupo cuántico $SU_\mu(1, 1)$ cuya geometría codificada, corresponde a una geometría hiperbólica. Además, dado que la parte clásica de este grupo es el círculo, podemos obtener de manera natural, un cálculo diferencial cuántico sobre el círculo donde el generador principal y su diferencial no conmutan.

Podemos ver al grupo $SU_\mu(1, 1)$ como un haz principal cuántico, donde el grupo estructural está dado por el círculo y la base del haz corresponde geoméricamente a un hiperboloide cuántico, motivo por el cual llamamos a este haz, haz hiperbólico. Haciendo un cambio de coordenadas y permitiendo la extensión de $SU_\mu(1, 1)$ por medio de la introducción de las inversas y funciones apropiadas de los generadores α y α^* , podemos ver a este grupo cuántico de una forma más sencilla: estará generado por dos “coordenadas” z y w , donde la primera corresponde clásicamente a la coordenada compleja del plano hiperbólico en el modelo del disco unitario y la segunda, a una transformación unitaria que corresponderá a la rotación alrededor del origen y que cumple la relación

$$z^*z = \psi(zz^*) = f(zz^*),$$

donde ψ es un automorfismo interno del álgebra dado por $\psi(a) = w^*aw$ y f un homeomorfismo en el intervalo $[0, 1]$.

Desde estas nuevas coordenadas podemos ver a $SU_\mu(1, 1)$ como un haz principal cuántico con grupo estructural dado por el círculo y base dada por la $*$ -álgebra generada por el elemento z . Si consideramos el mismo cálculo diferencial de Woronowicz, obtenemos una fórmula muy interesante de la métrica en la variedad base de este haz:

$$ds^2 = \mu^{-2} \frac{1}{(1 - \psi^{-1}(zz^*))(1 - zz^*)} dz dz^*.$$

En el caso clásico cuando $\mu = 1$ y en el caso $\mu = -1$, el automorfismo ψ se trivializa y obtenemos de la fórmula anterior la métrica hiperbólica, confirmando el hecho de que se trata de una generalización cuántica del modelo del disco de Poincaré.

En este capítulo se desarrolla la teoría para encontrar las funciones f que satisfacen la relación principal que define al álgebra y que nos dan versiones cuánticas de los planos hiperbólicos y que además retienen las simetrías clásicas del grupo $SU(1, 1)$. También consideramos representaciones en espacios de Hilbert cuyos elementos son funciones holomorfas sobre cierto dominio común. Resulta que si queremos que se retengan las simetrías clásicas el dominio común es todo el disco unitario, y si se retienen las simetrías cuánticas, el dominio común es el disco unitario sin el cero.

Inspirados en los resultados obtenidos para el haz hiperbólico dado por $SU_\mu(1, 1)$, generalizamos la manera de obtener planos hiperbólicos cuánticos considerando álgebras que se pueden ver como producto cruzado entre un álgebra $\mathcal{S}$ generada por un elemento positivo y equipada por un automorfismo y un álgebra asociada a dicho automorfismo. Aquí hay diferentes maneras de proceder a la hora de definir un cálculo diferencial sobre estas álgebras. La primera de ellas es considerar un cálculo diferencial clásico en $\mathcal{S}$ de manera que la diferencial y el automorfismo conmuten y la segunda, es no poner condiciones sobre la conmutatividad de la diferencial y el automorfismo. Abordamos estos dos caminos y calculamos las principales componentes geométricas como la conexión, la curvatura y la derivada covariante.

Capítulo 2

Antecedentes: Física y Geometría cuántica

2.1 Física cuántica

El universo está lleno de grandes misterios y entre ellos se encuentra el surgimiento de ésta nuestra tierra y el de nuestra existencia. Como parte del desarrollo de intentos para desenvolver tales misterios, los científicos se han dado a la tarea de entender cómo es la materia en sus niveles más profundos. Se sabe que la materia está formada por unidades diminutas entre las que se encuentran los átomos, moléculas, electrones, protones, neutrones, quarks, fotones, etc. A todas estas unidades diminutas podemos llamarlas *cuantos*. Y es la mecánica cuántica, la encargada de describir y explicar todos los fenómenos relacionados con los cuantos. Uno de éstos fenómenos que además dan lugar al nacimiento de ésta es la luz y el misterio de la estabilidad de la materia en algunas de sus configuraciones básicas.

Según la física clásica, la luz es una manifestación del campo electromagnético cuyos fenómenos se describen a través de las ecuaciones de Maxwell. En particular, según esta teoría, la luz exhibe *propiedades ondulatorias*. Uno de los fenómenos que distinguen a las ondas de las partículas es que las primeras manifiestan propiedades de difracción e interferencia. La difracción es la desviación que ocurre cuando una onda se ve interrumpida por un obstáculo o barrera y la interferencia ocurre cuando dos o más ondas se encuentran en un punto en el espacio y éstas dan lugar a una nueva onda que surge de la superposición de ellas.

Por ejemplo, encontramos como una bella ilustración del fenómeno de interferencia el *experimento de la rendija* efectuado por primera vez por el

físico inglés Thomas Young en 1801, donde se observan señales de que la luz interfiere.

Las partículas reflejan la idea de lo localizado; una idealización extrema es el concepto de *punto material*: el punto geométrico, equipado con propiedades físicas como por ejemplo, energía, velocidad y masa. Y por otro lado, las ondas son *extendidas*; su existencia está siempre vinculada a una región del espacio y tiempo. Una estructura física subyacente proporciona la geometría física para las ondas.

Podemos decir que en la física clásica, por lo general, los fenómenos se describen con base en los conceptos de partículas y ondas posiblemente interactuando en un espacio tiempo tetradimensional.

La visión ondulatoria de la luz que nos da la física clásica cambió profundamente a inicios del siglo *XX* cuando Max Planck, estudiando teóricamente la *radiación del cuerpo negro* se dio cuenta de la necesidad lógica de algo completamente diferente de la imagen clásica: que la energía de radiación de un cuerpo negro se puede recibir y emitir solamente en cantidades discretas de cuantos de energía para cada frecuencia de radiación dada. Estos cuantos de energía tienen cantidad mínima directamente proporcional a la frecuencia de oscilaciones a través de una nueva constante fundamental de la naturaleza, h , ahora conocida como constante de Planck. Así este cuanto de la energía de radiación del cuerpo negro se puede calcular con la fórmula:

$$E = h\nu$$

donde ν es la frecuencia de la radiación. La energía total emitida o absorbida en cualquier situación para una frecuencia dada ν es siempre un múltiplo entero de la cantidad elemental.

Fue así como empieza a surgir la teoría cuántica, donde aparecía la hipótesis cuántica de la energía de Planck en la que se afirma que la energía de la luz tiene valores discretos y la posterior hipótesis cuántica de la luz de Einstein que afirma que la luz puede ser tratada como una colección de partículas, donde la partícula de luz en la frecuencia ν tiene su energía dada por esta fórmula de Planck. Esto permitía a los físicos explicar cosas sobre la luz que la teoría clásica no podía esclarecer y la cual describe un mundo que es invisible a simple vista.

Planck además se dio cuenta de algo muy interesante: que era posible combinar las dos constantes universales de la naturaleza; la constante de gravitación G y la velocidad de la luz c , con su constante encontrada, y obtener una unidad de longitud, conocida como *longitud de Planck*:

$$\ell_p = \sqrt{\frac{hG}{c^3}}.$$

Recordemos que Einstein pensaba que la geometría era el lenguaje para describir a la naturaleza. Por lo tanto, la mera existencia de una unidad fundamental de longitud absoluta, como que nos indica que la geometría de la naturaleza es una geometría no-euclidea pues la geometría euclidea no sabe distinguir lo grande de lo pequeño. Si observamos el tamaño de la longitud de Planck $\approx 10^{-33}cm$, ésta es algo extremadamente pequeño, la misma naturaleza nos indica, que en esta escala espacio-temporal de lo exorbitantemente pequeño en términos humanos, algo completamente nuevo pasará con la geometría.

Asociado a ésto, también se puede introducir el *tiempo de Planck* que es el tiempo que tarda la luz en atravesar una distancia igual a la longitud de Planck y está dado por la fórmula

$$\tau_p = \sqrt{\frac{hG}{c^5}},$$

que mide aproximadamente 10^{-43} segundos. Así, es natural preguntarnos, ¿qué estructura geométrica debemos cambiar si nos interesa algo tan pequeño?

Desde tiempos de Euclides sabemos que *cuanto-átomo* de geometría se interpreta naturalmente como punto y que el espacio geométrico es una estructurada colección de sus puntos. Lo cual nos indica que el cambio que la geometría cuántica promoverá tendrá que ver con la idea de punto y todas las nociones basadas en ella.

La materia quiere estabilidad, en particular el átomo de hidrógeno es estable, sin embargo este fenómeno tan simple fue imposible de explicar con los métodos de la física clásica. Además fenómenos de absorción y emisión de luz por parte de los átomos resultaban inexplicables ya que los experimentos mostraban *spectros discretos de luz*.

Niels Bohr resuelve ésto, y así nace el primer modelo cuántico del átomo, postulando que solamente existen órbitas discretas estables para el electrón en el átomo, incluyendo la órbita estable de la energía mínima base. Y que pasando de una a otra de tales órbitas, el electrón emite o absorbe energía en múltiplos enteros del elemental cuanto de Planck.

En la escuela de Copenhague dirigida por Bohr, se cristalizó la estructura matemática de esta nueva teoría que rompe con los paradigmas de la física clásica y que unifica conceptos que la física clásica consideraba muy diferentes. Una manifestación matemática de este profundo cambio es la *no conmutatividad* de observables físicas. En particular, para el sistema físico más simple de una partícula unidimensional, su coordenada q y su momento p cumplen la siguiente relación de Heisenberg:

$$pq - qp = \frac{h}{2\pi i}. \quad (2.1)$$

Otro cambio profundo es el rompimiento con el determinismo y con el concepto de estado del sistema individual. Viendo-desarrollando esta física no conmutativa a fondo, resulta que lo nuevo de la mecánica cuántica como tal, surge de esta primordial no conmutatividad de observables físicas.

2.2 Geometría cuántica

La no conmutatividad de observables físicas como posición y momento dadas por la relación de Heisenberg (2.1), nos invita a considerar una nueva geometría que describirá los espacios que emergen de estas estructuras no-conmutativas. Es claro que para describir dichos espacios, las nociones geométricas fundamentales, incluyendo la de punto del espacio, deben ser profundizadas, generalizadas y reinterpretadas, obteniendo así estos nuevos espacios que llamaremos *espacios cuánticos*.

Para describir cómo son los espacios cuánticos, fue necesario desarrollar un nuevo vocabulario geométrico que dio lugar a la ahora conocida *geometría cuántica*. En ésta se unifican las ideas y conceptos de otras áreas como el análisis funcional, la geometría diferencial, la topología, y el álgebra. El principal vocabulario matemático usado, es el mismo que usa la física cuántica: básicamente la teoría de operadores en espacios de Hilbert. Pero en lugar de estudiar sistemas físicos se estudian los espacios.

Dentro de las estructuras algebraicas de la geometría cuántica se encuentran las álgebras C^* . Cuando estas álgebras son conmutativas siempre se pueden ver como las de funciones continuas sobre un espacio localmente compacto (desapareciendo en la infinitud topológica del espacio).

Y en otros contextos clásico-geométricos – como de variedades suaves, algebraicas o complejas, o bien, de espacios medibles – siempre es posible traducir y completamente expresar las propiedades geométricas del espacio en términos de sus correspondientes funciones (como suaves, polinomiales, holomorfas o medibles). Y estas funciones siempre generan un sistema algebraico conmutativo. En resonancia con esto, existe una correspondencia uno a uno entre las propiedades geométricas del espacio y las propiedades algebraicas del álgebra.

En geometría cuántica ampliamos esta correspondencia permitiendo las álgebras no conmutativas desarrollando de tal forma el concepto de espacio cuántico. Álgebras no-conmutativas siempre llevan un enlace con álgebras C^* , y resulta que estas álgebras se pueden representar por medio de operadores en un espacio de Hilbert, regresándonos así al primordial entorno matemático de la física cuántica.

Los espacios de Hilbert hechos de funciones holomorfas son particularmente interesantes y nos muestran una manera sencilla y elegante de desarrollar espacios cuánticos con sus álgebras no-conmutativas de operadores.

Existe una biyección entre estos espacios de Hilbert H y ciertas funciones *núcleos reproductores* $K: \Omega \times \Omega \rightarrow \mathbb{C}$ donde Ω es el dominio común de las funciones de H . La función $K(w, z)$ es holomorfa en z y antiholomorfa en w , y para w fijo, se manifiesta como una función $w(z)$ del espacio H . Resulta que estas *funciones de onda de puntos* generan todo el espacio H y cumplen con la propiedad básica del núcleo reproductor:

$$\langle w(z) | \psi(z) \rangle = \psi(w) \quad (2.2)$$

donde $\langle | \rangle$ es el producto escalar en H .

Observemos que ésta última igualdad, de forma similar con lo que ocurre en la fórmula de De Broglie, unifica en una sola expresión el lenguaje de puntos y ondas.

El núcleo reproductor K también debe satisfacer una serie de condiciones de positividad, para poder reproducir la estructura del espacio de Hilbert.

Para construir el espacio cuántico asociado al espacio de Hilbert holomorfo H , podemos considerar z (o alguna otra función apropiada) como el operador de multiplicación dentro de H y con su adjunto z^* generar el álgebra, que siempre resulta ser no-conmutativa.

Estos espacios proporcionan un interesante entorno transformacional, donde podemos comenzar con un álgebra, modificándola para que sea interpretable como una de sus asociadas álgebras no-conmutativas.

Las funciones de onda de puntos $w(z)$ permiten una interpretación interesante en términos de puntos estocásticos de la formulación de Prugovečki, como aproximaciones de ondas de lo posible en algo localizado alrededor del punto clásico dado. Esto por su parte, establece otro vínculo natural con el concepto de estado coherente. En particular, la fórmula

$$v(u) = K(v, u) = \langle u(z) | v(z) \rangle \quad (2.3)$$

se puede ver en esta nueva luz como parte de la geometría de los puntos estocásticos, dándonos la amplitud de probabilidad de transición de un punto a otro. Con esta información, entre otras cosas, podemos calcular la métrica clásica.

Un ejemplo donde directamente se utiliza esta teoría es para el *plano cuántico* dado por la relación de Heisenberg (2.1). En lugar de utilizar

coordenadas p y q podemos considerar coordenadas complejas z y z^* definidas de la siguiente manera

$$z^* = \sqrt{\frac{m\omega}{2\hbar}} \left(q + i \frac{p}{m\omega} \right), \quad z = \sqrt{\frac{m\omega}{2\hbar}} \left(q - i \frac{p}{m\omega} \right),$$

donde los parámetros m y ω tienen la interpretación física de la masa y velocidad angular del oscilador armónico simple unidimensional.

Así la relación (2.1) ahora se convierte en la siguiente relación entre las coordenadas z y z^* :

$$[z^*, z] = z^* z - z z^* = 1. \quad (2.4)$$

Podemos considerar H como el espacio de Hilbert de todas las funciones enteras tal que la norma

$$\|f\|^2 = \frac{1}{\pi} \int_{\mathbb{C}} e^{-|z|^2} |f(z)|^2 dz$$

es finita y pensar a z como el operador de multiplicación y la adjunta z^* como el operador de derivación por z en dicho espacio de Hilbert. Así las funciones de onda de punto $w(z)$ para cada $w \in \mathbb{C}$ están dadas de la siguiente manera:

$$w(z) = \exp \bar{w} z.$$

Se puede demostrar que se satisface la propiedad básica de núcleos reproductores. Obsérvese que si definimos $|u\rangle$ y $|v\rangle$ en H como las funciones de onda normalizadas dadas por

$$|u\rangle = \exp(-|u|^2/2) \exp(\bar{u}z), \quad |v\rangle = \exp(-|v|^2/2) \exp(\bar{v}z),$$

entonces

$$\begin{aligned} \langle u|v\rangle &= \exp(-(|u|^2 + |v|^2)/2) \langle \exp(\bar{u}z) | \exp(\bar{v}z) \rangle = \\ &= \exp(-(|u|^2 + |v|^2)/2) \exp(u\bar{v}) = \exp\left(-\frac{|u-v|^2}{2} + i \operatorname{Im} u\bar{v}\right). \end{aligned}$$

Y por lo tanto la probabilidad de transición del punto estocástico u al punto estocástico v , dada por el cuadrado del módulo de su producto escalar, resulta ser

$$|\langle u|v\rangle|^2 = \exp(-|u-v|^2).$$

Así la probabilidad de no-transición nos ayuda a recuperar la métrica euclídeana considerando puntos u y v suficientemente cercanos:

$$1 - |\langle u|v\rangle|^2 \approx |u-v|^2,$$

lo que nos da una de las razones por las que llamamos al espacio asociado al álgebra generada por estas dos coordenadas z y z^* con la relación (2.4)—*plano euclideo cuántico*.

En la física cuántica, como mencionamos en la sección anterior, se puede definir de manera natural una unidad fundamental de longitud, por lo que es de esperarse, que la geometría que describa a los sistemas físicos cuánticos, y en este caso al oscilador armónico simple, debe tener una unidad de medida que pueda ser construida con pura geometría. Es fácil calcular que esta álgebra tiene como simetrías a las siguientes transformaciones:

$$T(z) = az + b,$$

donde $a, b \in \mathbb{C}$ y $|a| = 1$. Es decir, sólo mantiene las simetrías euclidianas que le son necesarias para poder hablar de una escala absoluta de distancia y una orientación natural: rotaciones y traslaciones, excluyendo las homotecias y las reflexiones. Podemos decir que este espacio cuántico es más rico en la geometría que su contraparte clásica. Y la métrica y orientación en este caso, surgen de la geometría del espacio.

También podemos decir que la riqueza de la estructura geométrica que encontramos en los espacios cuánticos, es un fenómeno general.

Existe una variedad de ejemplos de espacios cuánticos que surgen de este entorno holomorfo. En particular, como veremos más adelante, el álgebra generada por los operadores z y z^* con la única relación

$$z^*z = f(zz^*), \quad (2.5)$$

nos da toda una variedad de ejemplos de planos hiperbólicos cuánticos.

En el caso mas simple cuando $f \equiv 1$, obtenemos el álgebra de Toeplitz generada por el “operador de shift” (u “operador de corrimiento”). El espacio cuántico asociado a esta álgebra, no puede ser un plano euclideo cuántico, pues las transformaciones de la forma $T(z) = z + b$, para algún $b \in \mathbb{C} \setminus \{0\}$, no son simetrías del álgebra. Sin embargo, cálculos directos nos muestran que este espacio tiene como simetrías al grupo de transformaciones hiperbólicas $SU(1, 1)$, reforzando con esto la afirmación de que se trata de un espacio hiperbólico.

Otro ejemplo muy importante que abunda en propiedades cuánticas, es basado en el álgebra C^* no-conmutativa más simple—la de las matrices complejas 2×2 . Aquí se destaca la base dada por la matriz identidad y las matrices de Pauli:

$$\sigma_1 = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \quad \sigma_2 = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}, \quad \sigma_3 = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}. \quad (2.6)$$

Estas matrices satisfacen las siguientes relaciones, que determinan completamente al álgebra

$$\begin{aligned} \sigma_1 \sigma_2 &= -\sigma_2 \sigma_1 = i\sigma_3, & \sigma_1^* &= \sigma_1, & \sigma_2^* &= \sigma_2, & \sigma_3^* &= \sigma_3, \\ \sigma_2 \sigma_3 &= -\sigma_3 \sigma_2 = i\sigma_1, & \sigma_1^2 &= \sigma_2^2 = \sigma_3^2 &= \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}. \\ \sigma_3 \sigma_1 &= -\sigma_1 \sigma_3 = i\sigma_2, \end{aligned} \quad (2.7)$$

En concordancia con lo que mencionamos en los ejemplos anteriores, para desarrollar la intuición geométrica podemos investigar cuáles son las simetrías naturales. Como veremos en detalle más adelante, los automorfismos de esta álgebra forman el grupo de Moebius $M(2)$ sin considerar la estructura $*$, y el grupo $SO(3)$ si incluimos la preservación de dicha estructura. En ambos casos reconocemos las mismas simetrías fundamentales de la esfera clásica: tanto de nivel holomorfo como métrico. Entonces, se trata de un objeto cuántico esférico—una verdadera *esfera cuántica* podemos decir.¹

Como parte de la riqueza de este mundo cuántico, podemos mencionar que un mismo concepto clásico se puede ver profundizado y multiplicado en diferentes versiones cuánticas no-equivalentes. Es este el caso del concepto de punto en el espacio. Para la geometría clásica podemos expresar el concepto de punto en términos algebraicos como *caracter* del álgebra, *estado puro* ó como clase de equivalencia de *representaciones irreducibles*, todas ellas equivalentes en este contexto.

Todos estos conceptos, de alguna forma reflejan “punto”. Sin embargo, la geometría cuántica es una nueva tierra en la que debemos desarrollar un lenguaje para que pueda ser explorada. De manera similar como se hizo para la física cuántica usando las ideas y conceptos de la física clásica, desarrollamos para tal lenguaje: ideas, dibujos y metáforas de la tierra ya conocida, la geometría clásica. La misma problemática de cómo explorar nuevas tierras, también se encuentra en las discusiones de los fundadores de la mecánica cuántica [15].

Así, en la nueva tierra: geometría cuántica, podemos conocer las estructuras geométricas de estos espacios, a través de las de la vieja tierra: la geometría clásica. Por ejemplo, si pensamos a los puntos del espacio en términos de caracteres, podremos encontrar espacios cuánticos sin puntos. El álgebra que describe al oscilador armónico y el álgebra de matrices complejas de 2×2 no contienen ningún caracter, sin embargo ésta última, tiene etiquetados a los estados puros con los puntos de la 2-esfera.

Por otro lado, el álgebra de Toeplitz tiene como caracteres a todos los puntos del círculo unitario, podemos pensar así, con lo que ya dijimos antes

¹Un Bebé Cuántico, como la llamamos en el Seminario de Geometría Cuántica en el Instituto de Matemáticas de la UNAM.

de su grupo de simetrías, que se trata de la versión cuántica del modelo del disco de Poincaré.

Relacionado con los estados ρ , podemos mencionar la construcción de Gelfand-Naimark-Segal, que es fundamental para el universo de álgebras C^* . Ésta establece una biyección entre el conjunto de estados del álgebra A y las ternas (H, π, ξ) que consisten de un espacio de Hilbert, una representación del álgebra en H y un vector cíclico de norma uno para la representación, tal que se cumple la siguiente relación

$$\rho(a) = \langle \xi | \pi(a) \xi \rangle,$$

para cada $a \in A$.

Como veremos más adelante, los estados en álgebras C^* generalizan a las medidas de probabilidad, así que esta construcción GNS entre otras cosas establece el puente entre probabilidad y representaciones.

Un corolario importante de esta construcción es que toda álgebra C^* la podemos ver como un álgebra de operadores en un espacio de Hilbert. En el siguiente capítulo, veremos en más detalle este resultado y mostraremos cómo algunos conceptos geométricos se traducen en conceptos algebraicos.

En este trabajo, para desarrollar y estudiar las estructuras geométricas de los espacios cuánticos, seguiremos el camino clásico de simetrías, como lenguaje unificante, tal y como lo establece Felix Klein en su programa de Erlangen en 1872. En éste, afirma que la geometría es el estudio de las propiedades que no cambian bajo la acción de un grupo de transformaciones.

Por lo tanto, para seguir este camino, necesitamos establecer el concepto cuántico de *grupo*. Otro pilar de la geometría clásica, muy en resonancia con el Programa de Erlangen, es la teoría de haces principales. Esta teoría integra simetrías al nivel fundamental de geometría del espacio, y establece la visión del espacio geométrico no solo como un espacio de puntos sino que también como una *estructura viva* en el que los puntos se ven enriquecidos con elementos de la estructura geométrica, y son estos elementos que constituyen el haz principal, y se transforman a lo largo de las fibras, por medio del grupo de simetrías.

Así que la teoría de haces principales toma un papel fundamental e importante para el desarrollo de la geometría clásica, como lo podemos apreciar en el libro de Shoshichi Kobayashi y Katsumi Nomizu [16].

Otro entorno teórico donde aparecen naturalmente haces principales es en las teorías de norma, o también conocidas como teorías de Yang Mills. Esto incluye varios modelos de partículas elementales y sus interacciones, en particular el Modelo Estándar. Aquí, en cierta forma, la idea de espacio

vivo asume dimensiones físicas. Los elementos estructurales que enriquecen puntos del espacio-tiempo físico, se relacionan con las simetrías locales que se pueden realizar.

Esta idea de espacio vivo reaparece en la geometría cuántica si tomamos en cuenta el marco conceptual de espacios de Hilbert de funciones holomorfas mencionado anteriormente, donde los puntos clásicos se ven como funciones de onda. En esta tesis usaremos este lenguaje y las contrapartes cuánticas de la teoría de grupos y haces principales para estudiar la geometría de los espacios cuánticos.

Por lo ya mencionado, de cómo se incorporan los conceptos matemáticos en el desarrollo de la geometría cuántica, nos damos cuenta que es a través de esta teoría que podemos unificar los grandes pilares de las matemáticas como el análisis, la geometría, el álgebra, la probabilidad, la topología y la teoría de grupos, entre otras, con la teoría física que describe a los sistemas cuánticos. Esto de nuevo nos indica en su propia forma que las matemáticas y la física van de la mano [21].

Capítulo 3

Teoría de álgebras C^*

El estudio de las álgebras C^* comienza con los trabajos de Gelfand y Naimark quienes obtuvieron el resultado fundamental de que todas las álgebras C^* conmutativas son isomorfas a ciertas álgebras de funciones continuas complejas sobre los espacios topológicos localmente compactos.

Usando este resultado, podemos decir que a cada álgebra C^* conmutativa le corresponde un espacio topológico localmente compacto. Cada propiedad geométrica del espacio se corresponde con una propiedad algebraica, y viceversa. A lo largo de este capítulo, mostraremos algunas formulaciones algebraicas que se corresponden con conceptos como punto, parte, medida, compactificación por un punto y simetrías de un espacio.

Para los fines de la geometría cuántica, y con analogía a cómo se desarrolló la física cuántica, ampliamos esta correspondencia permitiendo que el álgebra C^* sea no conmutativa, y es entonces, que conoceremos a los espacios cuánticos a través de lo ya conocido en los espacios clásicos. Por lo tanto, si el álgebra es no conmutativa, ésta debe corresponderse con un espacio que llamamos cuántico y que puede ser estudiado desde los ojos clásicos.

3.1 Álgebras C^*

Recordamos que un álgebra asociativa A , es un espacio vectorial sobre $\mathbb{C}$, dotado de una operación multiplicación tal que se satisfacen las siguientes igualdades:

$$\begin{aligned}(a + b)c &= ac + bc, & a(b + c) &= ab + ac, & a(bc) &= (ab)c, \\ \lambda(ab) &= (\lambda a)b = a(\lambda b),\end{aligned}\tag{3.1}$$

para todo $a, b, c \in A$ y $\lambda \in \mathbb{C}$.

Decimos que A es unital si existe un elemento 1 (necesariamente único) en dicha álgebra tal que

$$1a = a1 = a,$$

para todo $a \in A$.

Una *estructura* $*$ en un álgebra A es una transformación $*$: $A \rightarrow A$, denotada también por $a \mapsto a^*$, que es un antiautomorfismo antilineal involutivo, es decir, que cumple con las siguientes propiedades:

$$(\lambda a + b)^* = \bar{\lambda}a^* + b^*, \quad (ab)^* = b^*a^*, \quad (a^*)^* = a, \quad (3.2)$$

para todo $a, b \in A$ y $\lambda \in \mathbb{C}$.

Decimos también que A es una $*$ -álgebra, si A está equipada con una estructura $*$.

Un álgebra normada A es un álgebra equipada con una norma (como espacio vectorial complejo) tal que la norma es submultiplicativa:

$$\|ab\| \leq \|a\|\|b\|,$$

para todo $a, b \in A$.

Un álgebra de Banach es un álgebra normada que es un espacio de Banach, en otras palabras completa en su norma.

Definición 3.1.1. Sea A un álgebra de Banach equipada con una estructura $*$. Decimos que A es una álgebra C^* si se satisface la propiedad C^* :

$$\|a\|^2 = \|a^*a\|, \quad (3.3)$$

para cada $a \in A$.

Más adelante veremos que cuando A es una álgebra C^* hay una conexión intrínseca entre la $*$ y la norma: la estructura $*$ determina completamente la norma, y viceversa; dada una norma en una álgebra C^* , ésta determina por completo la $*$.

Un punto interesante que cabe resaltar es que cuando el álgebra es no conmutativa entonces podemos definir múltiples estrellas; $*$, y por lo tanto múltiples normas. Y en el caso conmutativo que marca la frontera de los espacios clásicos, ambas estructuras la $*$ y la norma, están completamente determinadas por el álgebra compleja.

Ejemplos. Los ejemplos más sencillos de álgebras C^* conmutativas con unidad son los espacios $\mathbb{C}^n$ para cada $n \in \mathbb{N}$. En estos espacios definimos la suma y

la multiplicación por coordenadas, la involución de un elemento de $\mathbb{C}^n$ está dada por la conjugación compleja:

$$(a_1, \dots, a_n)^* = (\bar{a}_1, \dots, \bar{a}_n),$$

y la norma con la que se vuelve C^* está dada por la norma del máximo de las intensidades de las coordenadas:

$$\|(a_1, \dots, a_n)\| = \max\{|a_j| : j = 1, \dots, n\}.$$

De hecho, todas las álgebras conmutativas y de dimensión finita tienen esta forma.

Para el siguiente ejemplo consideramos un espacio localmente compacto¹ X y sea $C_0(X)$ el álgebra de funciones continuas en X que se anulan al infinito, es decir,

$$C_0(X) = \{f : X \rightarrow \mathbb{C} : f \text{ continua y } \lim_{x \rightarrow \infty} |f(x)| = 0\}.$$

En esta álgebra, definimos las operaciones de suma y producto puntualmente y consideramos la norma del supremo:

$$\|f\| = \sup_{x \in X} |f(x)|.$$

Definiendo la involución para cada función continua f como $f^*(x) = \overline{f(x)}$, tenemos que $C_0(X)$ es una álgebra C^* conmutativa. De hecho, mostraremos que toda álgebra conmutativa es de esta forma. El álgebra $C_0(X)$ tendrá unidad si y sólo si X es compacto, y en tal caso coincide con $C(X)$, el álgebra de todas las funciones continuas sobre X .

Como último ejemplo, sea $\mathcal{H}$ un espacio de Hilbert, y $B(\mathcal{H})$ el álgebra de operadores lineales acotados en $\mathcal{H}$, con la suma usual y el producto dado por la composición de operadores. Si además definimos la $*$ como el operador adjunto y la norma como la norma uniforme de operadores:

$$\|T\| = \sup_{x \neq 0} \frac{\|T(x)\|}{\|x\|},$$

entonces tenemos una álgebra C^* con unidad. En el caso finito-dimensional donde podemos poner $\mathcal{H} = \mathbb{C}^n$, tendremos una estructura C^* natural en las matrices complejas $M_n(\mathbb{C})$ de tamaño $n \times n$, ya que podemos identificar $M_n(\mathbb{C}) = B(\mathcal{H})$. Recordemos que en espacios de Hilbert de dimensión finita, cada mapeo lineal es continuo. Aquí la involución $*$ está dada por la transpuesta conjugada y la norma es dada por la norma uniforme de operadores-matrices.

¹Para nosotros la noción de compacto es la combinación de la condición de subcubierta abierta y el segundo axioma de separabilidad de Hausdorff.

Observación 3.1.1. Si A es un álgebra normada, entonces la multiplicación

$$\cdot : A \times A \rightarrow A$$

es continua en ambas entradas, y si A es una álgebra C^* entonces

$$\|a\| = \|a^*\|,$$

y por tanto la operación $*$ es una función continua.

3.2 Homomorfismos e ideales

Definición 3.2.1. Sean A y B álgebras, decimos que un mapeo lineal $\pi : A \rightarrow B$ es un homomorfismo si

$$\pi(ab) = \pi(a)\pi(b),$$

donde $a, b \in A$. Si además A y B son $*$ -álgebras, decimos que el homomorfismo π es un $*$ -homomorfismo si cumple que

$$\pi(a^*) = \pi(a)^*,$$

para todo $a \in A$. Y si las álgebras A y B son unitales, decimos que π es unital si $\pi(1_A) = 1_B$.

Si B es el álgebra de operadores acotados en un espacio de Hilbert $\mathcal{H}$, diremos que π es una representación del álgebra A en el espacio de Hilbert $\mathcal{H}$, si π es un $*$ -homomorfismo tal que se cumple la siguiente condición de regularidad

$$\overline{\{\pi(A)\mathcal{H}\}} = \mathcal{H}.$$

Esta condición implica que en caso de que A sea unital, el homomorfismo π debe ser unital.

En el caso particular cuando A es un álgebra de Banach y $B = \mathbb{C}$, decimos que $\kappa : A \rightarrow \mathbb{C}$ es un caracter de A , si κ es un homomorfismo no cero.

Observación 3.2.1. Si κ es caracter y el álgebra es unital, entonces κ es unital también.

Definición 3.2.2. Un espacio vectorial $I \subseteq A$, es un ideal izquierdo del álgebra A si y sólo si

$$an \in I,$$

para todo $a \in A$ y $n \in I$. Similarmente, decimos que I es un ideal derecho de A si y sólo si

$$na \in I,$$

para todo $a \in A$ y $n \in I$. Finalmente, decimos que I es un ideal bilateral, o simplemente ideal, si es izquierdo y derecho. Además, I es $*$ -ideal si es $*$ -invariante.

Obsérvese que si I es un ideal unilateral (izquierdo o derecho) y además $*$ -invariante es necesariamente bilateral. Esto es la consecuencia directa de la antimultiplicatividad de $*$. Cada ideal es una subálgebra.

Se puede demostrar que si $\pi : A \rightarrow B$ es un homomorfismo, entonces el $\ker(\pi)$ siempre es un ideal de A . Más aún, si π es $*$ -homomorfismo, entonces $\ker(\pi)$ es un $*$ -ideal. En particular, si A tiene pocos ideales, entonces permitirá pocos homomorfismos.

Definición 3.2.3. Sea A un álgebra y sea I un ideal de A . Decimos que I es propio si es diferente de $\{0\}$ y de A . Éstos dos los llamaremos ideales triviales.

Definición 3.2.4. Sea A un álgebra y I un ideal de A . Decimos que I es maximal si es ideal propio y no existe otro ideal propio J tal que $I \subseteq J$ y $I \neq J$.

Obsérvese que si I es un ideal de cualquier tipo introducido arriba, y A un álgebra de Banach entonces su cerradura $\bar{I}$ es también un ideal del mismo tipo.

Obsérvese también que en caso de álgebras conmutativas los conceptos de ideal derecho, izquierdo o bilateral, coinciden.

Más adelante veremos que todo ideal maximal en un álgebra de Banach conmutativa es cerrado. Además, si A tiene ideales propios, entonces como consecuencia del lema de Zorn existen ideales maximales.

Ejemplo. Los ideales maximales del álgebra de funciones continuas sobre un compacto X son de la forma

$$I_x = \{f \in C(X) \mid f(x) = 0\},$$

para cada $x \in X$. Además si $x \neq y$ se tiene que $I_x \neq I_y$. Por lo tanto existe una correspondencia uno a uno entre puntos del espacio X e ideales maximales del álgebra $C(X)$.

Los ideales bilaterales nos permiten pasar al álgebra cociente $A \rightsquigarrow A/I$. Si A es conmutativa entonces A/I es conmutativa también. Si A es unital e $I \subset A$, entonces el álgebra cociente es unital. Si A es $*$ -álgebra y I un $*$ -ideal, la estructura $*$ se proyecta naturalmente a A/I .

Si A es un álgebra de Banach y $I = \bar{I}$, entonces A/I es un álgebra de Banach con la norma

$$\|[a]\| = \inf\{\|a + n\| \mid n \in I\}.$$

Proposición 3.2.1. En todas las álgebras C^ los ideales cerrados siempre son $*$ -ideales.* $\square$

Además si A es una álgebra C^* e I es un ideal cerrado de A entonces el álgebra cociente equipada con la estructura proyectada, es una álgebra C^* .

3.3 Unitalización

En general, podemos encontrar álgebras C^* sin unidad, como es el caso de $C_0(X)$ cuando X es no compacto (y localmente compacto). O bien, la subálgebra C^* de operadores compactos del álgebra $B(\mathcal{H})$ para un espacio de Hilbert de dimensión infinita separable.

En estos casos, siempre podemos extender el álgebra llegando a un álgebra unital. Presentaremos ahora la construcción de la extensión mínima.

Sea A una álgebra C^* y sea $A^\dagger = A \oplus \mathbb{C}$. Definimos la suma, el producto por escalar, multiplicación y $*$ en $A^\dagger$ de la siguiente manera:

$$\begin{aligned} (a, \lambda) + (b, \beta) &= (a + b, \lambda + \beta), & \mu(a, \lambda) &= (\mu a, \mu \lambda), \\ (a, \lambda)(b, \beta) &= (ab + \lambda b + \beta a, \lambda \beta), & (a, \lambda)^* &= (a^*, \bar{\lambda}), \end{aligned}$$

para todo $a, b \in A$ y $\lambda, \mu \in \mathbb{C}$. Entonces $A^\dagger$ es una $*$ -álgebra y el elemento $(0, 1)$ es la unidad en $A^\dagger$.

Definimos la norma en $A^\dagger$ como

$$\|(a, \lambda)\| = \sup\{\|ab + \lambda b\|_A : \|b\|_A = 1\},$$

donde $\|\cdot\|_A$ es la norma en el álgebra A . Se puede demostrar que $A^\dagger$ es una álgebra C^* con dicha norma y que el $*$ -homomorfismo

$$\varphi : A \rightarrow A^\dagger, \quad a \mapsto (a, 0)$$

es isométrico. Por otra parte, el $*$ -homomorfismo

$$\psi : A^\dagger \rightarrow \mathbb{C}, \quad (a, \lambda) \mapsto \lambda$$

tiene como kernel al álgebra A , de manera que A es un ideal cerrado de $A^\dagger$ y $A^\dagger/A \cong \mathbb{C}$.

De esta manera si A es una álgebra C^* no unital, definimos la unitalización de A como $A^\dagger$, la cual es la mínima álgebra C^* unital que contiene a A como su ideal.

3.4 El espectro de un elemento

A continuación mencionaremos algunos elementos especiales en las álgebras, éstos de cierta forma reflejan el tipo de operador en el que se convierten cuando consideramos álgebras de operadores en un espacio de Hilbert.

Definición 3.4.1. Sea A una $*$ -álgebra y $a \in A$. Entonces a es

- *autoadjunto o hermitiano* si $a = a^*$;
- *normal* si $aa^* = a^*a$;
- *idempotente* si $a = a^2$;
- *proyector* si $a = a^* = a^2$;
- *isometría parcial* si $aa^*a = a \iff a^*aa^* = a^*$.

Si A es unital entonces a es

- *invertible* si existe $a^{-1} \in A$ tal que $aa^{-1} = a^{-1}a = 1$;
- *unitario* si $aa^* = a^*a = 1$.

El siguiente teorema es uno de los más importantes sobre la existencia de elementos invertibles en un álgebra.

Teorema 3.4.1. Sea A un álgebra de Banach unital y $a \in A$ tal que $\|a\| < 1$. Entonces el elemento $1 - a$ es invertible y

$$\frac{1}{1 - a} = \sum_{n \geq 0} a^n.$$

Proposición 3.4.2. El conjunto $\text{GL}(A)$ de los elementos invertibles en el álgebra A , forma un conjunto abierto y además es un grupo respecto al producto del álgebra.

Demostración. Sea $a \in \text{GL}(A)$ y $\epsilon = \|a^{-1}\|^{-1}$. Consideremos b en la vecindad ϵ de a y observemos que $b = a(1 - a^{-1}(a - b))$. Por un lado, por la definición de ϵ tenemos que

$$\|a^{-1}\| \|a - b\| < 1, \quad (3.4)$$

y por el otro, dado que A es un álgebra de Banach obtenemos que

$$\|a^{-1}(a - b)\| \leq \|a^{-1}\| \|a - b\|.$$

Por lo tanto $\|a^{-1}(a - b)\| < 1$ y usando el Teorema (3.4.1) concluimos que $1 - a^{-1}(a - b) \in \text{GL}(A)$. Así, como $a \in \text{GL}(A)$ y $1 - a^{-1}(a - b) \in \text{GL}(A)$ entonces $b \in \text{GL}(A)$. Por lo tanto $\text{GL}(A)$ es abierto.

Finalmente, si a y b son invertibles, es claro que ab es invertible con $(ab)^{-1} = b^{-1}a^{-1}$. $\square$

Observación 3.4.1. Observemos que en el caso de álgebras de Banach conmutativas unitales, si $ab \in GL(A)$ entonces a y b también son invertibles. De hecho, esto es una consecuencia de un resultado más general que establece que en álgebras unitales, si el producto de dos elementos que conmutan entre sí es invertible, entonces tales elementos son invertibles.

Dada un álgebra A de Banach unital, definimos los autormorfismos internos de A , como aquellos de la forma

$$F_g : A \rightarrow A, \quad F_g(a) = gag^{-1}, \quad (3.5)$$

para todo $g \in GL(A)$. Observemos que tenemos un homomorfismo natural entre el grupo $GL(A)$ de los elementos invertibles de A y el grupo $\text{Aut}(A)$ de automorfismos de A . La existencia de estas simetrías es un fenómeno completamente cuántico: si A es conmutativa, el automorfismo F_g se trivializa.

En la mecánica cuántica describimos los sistemas físicos en términos de sus *propiedades*. Cada propiedad se representa por medio de un proyector ortogonal en el espacio de Hilbert de funciones de onda asociado al sistema físico. Estos proyectores se pueden construir por medio de diadas de Dirac a través de vectores unitarios en dicho espacio. Por otro lado, en la geometría cuántica es a través de los proyectores que describimos importantes aspectos estructurales de los espacios cuánticos, como conexidad, haces vectoriales e invariantes topológicos de K -teoría.

Proposición 3.4.3. Sea A una álgebra C^* y $u \in A$. Entonces las siguientes afirmaciones son equivalentes:

1. u^*u es proyector;
2. uu^* es proyector;
3. $u = uu^*u$;
4. $u^* = u^*uu^*$.

□

Por lo tanto, en álgebras C^* las isometrías parciales pueden ser definidas por cualquiera de estas 4 condiciones equivalentes.

Los elementos unitarios representan *simetrías* en las álgebras C^* . En este caso, los automorfismos internos F_u de A , cuando u es un elemento unitario, en realidad son $*$ -automorfismos.

Proposición 3.4.4. Sea A una $*$ -álgebra unital. El conjunto $U(A)$ de los elementos unitarios de A forman un grupo. □

Observemos que hay un homomorfismo natural entre el grupo de elementos unitarios de A y el grupo de $*$ -automorfismos de A .

Proposición 3.4.5. Sea A un álgebra de Banach y sea I un ideal maximal de A , entonces I es cerrado.

Demostración. Los ideales $I \neq A$ siempre contienen elementos no-invertibles. Y estos elementos, forman un conjunto cerrado, siendo el complemento de $\text{GL}(A)$ que es abierto (en caso de A unital, si A no es unital pasamos a $A^\dagger$). Por lo tanto, la cerradura de un tal ideal no puede ser todo A . $\square$

Lema 3.4.6. Sea A álgebra de Banach unital. Entonces el mapeo de la inversa $\text{GL}(A) \rightarrow \text{GL}(A)$ dado por $a \mapsto a^{-1}$ es continuo.

Demostración. Primero veamos que es continuo en $1 \in \text{GL}(A)$ y luego que es continuo en cualquier $a \in \text{GL}(A)$.

Sea $\{a_n\} \subset \text{GL}(A)$ una sucesión de elementos invertibles de A tal que $\lim a_n = 1$. Observemos que $b_n = 1 - a_n$ tiene límite cero, y por tanto, sin perder generalidad podemos suponer que $\|b_n\| = q_n < 1$. Por el Teorema (3.4.1) se puede concluir que $a_n = 1 - b_n \in \text{GL}(A)$ y $a_n^{-1} = \sum_{k=0}^{\infty} b_n^k$. Entonces

$$\lim_{n \rightarrow \infty} \|a_n^{-1} - 1\| = \lim_{n \rightarrow \infty} \left\| \sum_{k=0}^{\infty} b_n^k - 1 \right\| = \lim_{n \rightarrow \infty} \left\| \sum_{k=1}^{\infty} b_n^k \right\| \leq \lim_{n \rightarrow \infty} \frac{q_n}{1 - q_n} = 0.$$

Por lo tanto, $\lim a_n^{-1} = 1$. Así el mapeo de la inversa es continuo en 1.

Sea ahora $\{a_n\} \subset \text{GL}(A)$ una sucesión de elementos invertibles de A tal que $\lim a_n = a$. Observemos que $\lim a_n a_n^{-1} = 1$, de manera que aplicando la primera parte de la demostración tenemos que

$$\lim_{n \rightarrow \infty} a a_n^{-1} = \lim_{n \rightarrow \infty} (a_n a^{-1})^{-1} = 1.$$

Ahora como la multiplicación es continua

$$\lim_{n \rightarrow \infty} a a_n^{-1} = \lim_{n \rightarrow \infty} a \lim_{n \rightarrow \infty} a_n^{-1} = a \lim_{n \rightarrow \infty} a_n^{-1} = 1,$$

por lo que

$$\lim_{n \rightarrow \infty} a_n^{-1} = a^{-1}$$

lo que completa la demostración. $\square$

Definición 3.4.2. Sea A un álgebra de Banach unital, para todo $a \in A$ definimos el espectro del elemento a como

$$\sigma(a) = \{\lambda \in \mathbb{C} : (\lambda - a) \notin \text{GL}(A)\}.$$

Decimos que $\rho(a) = \mathbb{C} \setminus \sigma(a)$ es el conjunto resolvente de a y $R_\lambda(a) = (\lambda - a)^{-1}$ es la resolvente de a . En donde $R_\lambda(a)$ la interpretamos como una función definida sobre $\rho(a)$ con valores en A .

Ejemplos. Consideremos el álgebra $M_n(\mathbb{C})$ de matrices complejas de tamaño $n \times n$, para cada matriz a tenemos

$$\begin{aligned}\sigma(a) &= \{\lambda \in \mathbb{C} \mid \lambda I - a \notin \text{GL}(M_n(\mathbb{C}))\} = \{\lambda \in \mathbb{C} \mid \det(\lambda I - a) = 0\} \\ &= \{\lambda \in \mathbb{C} \mid \lambda \text{ es valor propio de } a\}.\end{aligned}$$

Es decir, el espectro de una matriz es igual al conjunto de todos sus valores propios.

Sea X un espacio topológico compacto y $C(X)$ el álgebra de funciones continuas sobre X . Entonces para cada función continua $f : X \rightarrow \mathbb{C}$ tenemos

$$\begin{aligned}\sigma(f) &= \{\lambda \in \mathbb{C} \mid \lambda - f \notin \text{GL}(C(X))\} = \{\lambda \in \mathbb{C} \mid \exists x \in X : (\lambda - f)(x) = 0\} \\ &= \{\lambda \in \mathbb{C} \mid \exists x \in X : \lambda = f(x)\} = \{f(x) \mid x \in X\} = f(X).\end{aligned}$$

Por lo tanto, podemos concluir que el espectro de una función continua del tipo mencionado arriba, es igual a la imagen o rango de la función.

Pensamos al espectro de un elemento en un álgebra como la generalización del rango de una función y de los valores propios de una matriz cuadrada de dimensión finita.

Proposición 3.4.7. Sea A un álgebra de Banach unital y $a \in A$. Si $\lambda > \|a\|$ entonces $\lambda \in \rho(a)$. Es decir, el espectro $\sigma(a)$, está acotado por $\|a\|$.

Demostración. Si $\lambda > \|a\|$ entonces $\lambda - a = \lambda(1 - a/\lambda) \in \text{GL}(A)$ pues $\|a/\lambda\| < 1$. Lo que nos dice que $\lambda \in \rho(a)$. $\square$

Proposición 3.4.8. Sea A un álgebra de Banach con unidad y $a \in A$. Entonces el espectro $\sigma(a)$, es cerrado.

Demostración. Sea $\lambda \in \rho(a)$, entonces $\lambda - a \in \text{GL}(A)$, pero como $\text{GL}(A)$ es abierto existe un $\epsilon > 0$ tal que la vecindad de $\lambda - a$ de radio ϵ está contenida en $\text{GL}(A)$. Si $\mu > 0$ y $|\lambda - \mu| < \epsilon$ entonces

$$\|(\lambda - a) - (\mu - a)\| = |\lambda - \mu| < \epsilon,$$

por lo que $\mu - a \in \text{GL}(A)$, es decir, $\mu \in \rho(a)$. Así que $\rho(a)$ es abierto, y por lo tanto $\sigma(a)$ es cerrado. $\square$

Observemos que las dos proposiciones anteriores implican que el espectro de un elemento es un conjunto compacto, siempre que A sea un álgebra de Banach con unidad.

Proposición 3.4.9. Sea A un álgebra de Banach unital. La función resolvente $R_\lambda(a) : \rho(a) \rightarrow \text{GL}(A)$ es holomorfa. $\square$

Teorema 3.4.10 (Gelfand). *Sea A un álgebra de Banach unital y $a \in A$. Entonces el espectro $\sigma(a)$, es no vacío.*

Demostración. Supongamos que $\sigma(a) = \emptyset$ y para cada f en el dual A^* de A , definimos $F : \rho(a) \rightarrow \mathbb{C}$, tal que $F(\lambda) = f(R_\lambda(a))$. Como f es lineal y acotada se tiene que

$$\lim_{\lambda \rightarrow \beta} \frac{F(\lambda) - F(\beta)}{\lambda - \beta} = f \left(\lim_{\lambda \rightarrow \beta} \frac{R_\lambda(a) - R_\beta(a)}{\lambda - \beta} \right).$$

Dicho límite existe pues la resolvente es diferenciable. Por lo tanto F es una función entera.

Además observando que

$$\lim_{\lambda \rightarrow \infty} \|F(\lambda)\| = 0,$$

obtenemos por el Teorema de Liouville que F es constante y $F = 0$. Lo cual es una contradicción al Teorema de Hanh-Banach que afirma que los puntos del álgebra se pueden separar por funciones del dual. Por lo tanto, $\sigma(a) \neq \emptyset$. $\square$

Proposición 3.4.11. (Gelfand-Mazur) *Sea A un álgebra de Banach unital tal que $\text{GL}(A) = A \setminus \{0\}$. Entonces $A \cong \mathbb{C}$.* $\square$

Proposición 3.4.12. *Sean A un álgebra de Banach unital y $a \in A$. Para todo polinomio en A ,*

$$p(a) = c_n a^n + c_{n-1} a^{n-1} + \cdots + c_0,$$

se cumple que

$$\sigma(p(a)) = p(\sigma(a)),$$

en otras palabras, los operadores de espectro y de polinomio conmutan. $\square$

Definición 3.4.3. *Sea A un álgebra de Banach unital, para todo $a \in A$ definimos el radio espectral de a , $r(a)$, por*

$$r(a) = \max\{|\lambda| : \lambda \in \sigma(a)\}.$$

Observación 3.4.2. Dado que $\sigma(a)$ está acotado por la norma de a , se tiene entonces que $r(a) \leq \|a\|$. Es particularmente interesante considerar los casos extremos $r(a) = 0$ y $r(a) = \|a\|$.

Lema 3.4.13 (Fórmula del radio espectral). *Sea A un álgebra de Banach unital. Se tiene que*

$$r(a) = \lim_{n \rightarrow \infty} \|a^n\|^{1/n},$$

para todo $a \in A$. $\square$

Corolario 3.4.14. Sea A una álgebra C^* unital y $a \in A$ un elemento normal, entonces

$$r(a) = \|a\|.$$

Demostración. Como a es normal tenemos:

$$\|a\|^4 = \|a^*a\|^2 = \|(a^*a)^2\| = \|(a^2)^*a^2\| = \|a^2\|^2.$$

Iterando la expresión anterior tenemos que

$$\|a\| = \|a^{2^n}\|^{2^{-n}},$$

que es una subsucesión de $\|a^n\|^{1/n}$. Por lo tanto

$$\|a\| = \lim_{n \rightarrow \infty} \|a^{2^n}\|^{2^{-n}} = \lim_{n \rightarrow \infty} \|a^n\|^{1/n} = r(a),$$

y nuestra fórmula queda demostrada. $\square$

Corolario 3.4.15. Sea A una álgebra C^* unital. Se cumple que

$$\|a\| = \sqrt{r(a^*a)}, \quad (3.6)$$

para todo $a \in A$. $\square$

Observación 3.4.3. Este simple pero muy importante corolario indica que la norma en una álgebra C^* con unidad está completamente definida por la estructura algebraica y por lo tanto es única. Obsérvese que la norma se escribe en términos del espectro como

$$\|a\| = \max\{|\lambda|^{1/2} : (\lambda - a^*a) \notin \text{GL}(A)\}.$$

Más adelante veremos que el fenómeno es simétrico, es decir que la norma también determina por completo a la $*$.

Obsérvese que si $\pi : A \rightarrow B$ es un homomorfismo entre dos álgebras de Banach unitales, entonces si $\lambda - a$ es invertible para cualquier $a \in A$, entonces $\lambda - \pi(a)$ es invertible, por lo que

$$\sigma(\pi(a)) \subseteq \sigma(a). \quad (3.7)$$

En particular, si A es un álgebra de Banach unital y $\kappa : A \rightarrow \mathbb{C}$ es un caracter entonces

$$|\kappa(a)| \leq r(a) \leq \|a\|,$$

lo cual demuestra que κ es continuo y $\|\kappa\| \leq 1$. Más aún, dado que A es unital tenemos que $\|\kappa\| = 1$. Por lo tanto el caracter κ pertenece al dual A^* de A y

$$\{\kappa(a) : \kappa \in A^* \text{ y } \kappa \text{ caracter de } A\} \subseteq \sigma(a).$$

Por otro lado, observemos que si $\pi : A \rightarrow B$ es un $*$ -homomorfismo entre álgebras C^* unitales entonces usando la fórmula (3.6) y la contención (3.7) tenemos que para todo $a \in A$ se satisface

$$\|\pi(a)\| = \sqrt{r(\pi(a)^*\pi(a))} = \sqrt{r(\pi(a^*a))} \leq \sqrt{r(a^*a)} = \|a\|,$$

es decir el $*$ -homomorfismo π es automáticamente continuo. En particular cada $*$ -isomorfismo es isométrico en el mundo de álgebras C^* .

3.5 Espectro de álgebras de Banach conmutativas

Lema 3.5.1. Sea A un álgebra de Banach unital y $\kappa : A \rightarrow \mathbb{C}$ un caracter, entonces $\ker(\kappa)$ es un ideal maximal.

Demostración. Es claro que $\ker(\kappa)$ es un ideal. La conclusión se obtiene al observar que $\ker(\kappa)$ tiene codimensión uno. $\square$

Lema 3.5.2. Sea A un álgebra de Banach conmutativa unital y sea $I \subseteq A$ un ideal maximal entonces $A/I \cong \mathbb{C}$.

Demostración. Sea I un ideal maximal, entonces por la Proposición (3.4.5) tenemos que I es cerrado y por lo tanto A/I es un álgebra de Banach unital.

Supongamos que $A/I \not\cong \mathbb{C}$, entonces por la Proposición (3.4.11) existe un elemento $[a] \in A/I$ tal que $[a] \neq [0]$ y $[a] \notin \text{GL}(A/I)$. Sea $J = aA$ entonces J es un ideal de A tal que $I \subset J$ y $J \neq A$. Lo cual contradice que I es un ideal maximal. Por lo tanto $A/I \cong \mathbb{C}$. $\square$

Estos dos últimos lemas nos muestran que si A es un álgebra de Banach conmutativa unital entonces existe una biyección natural entre el conjunto de ideales maximales de A y el conjunto de caracteres $\kappa : A \rightarrow \mathbb{C}$. En más detalle, el mapeo entre caracteres e ideales maximales dado por

$$\kappa \mapsto \ker(\kappa),$$

tiene como inverso al mapeo entre ideales maximales y caracteres:

$$I \mapsto \kappa,$$

donde $\kappa : A \rightarrow \mathbb{C}$ es el único caracter cuyo núcleo es I .

De esta manera, para toda álgebra A de Banach conmutativa unital, definimos el espectro del álgebra A , denotado por $\text{Spec}(A)$, como el conjunto de sus caracteres, es decir,

$$\text{Spec}(A) = \{\kappa : A \rightarrow \mathbb{C} : \kappa \text{ es caracter de } A\}. \quad (3.8)$$

Y por la biyección anterior tenemos que

$$\text{Spec}(A) = \{I \mid I \text{ es ideal maximal de } A\}. \quad (3.9)$$

Observemos que cuando A es un álgebra de Banach unital el conjunto de caracteres es un subconjunto de la bola cerrada con centro en cero y radio uno en el dual de A , la cual por el teorema de Alaoglu es compacta con la topología $*$ -débil. Se puede demostrar que $\text{Spec}(A)$ es un subconjunto cerrado de dicha bola con la topología $*$ -débil y por lo tanto es compacto.

Proposición 3.5.3. Sea A una álgebra C^* unital y sea $\kappa : A \rightarrow \mathbb{C}$ un caracter. Entonces para todo $a \in A$, $a = a^*$, se tiene que $\kappa(a) \in \mathbb{R}$. $\square$

Obsérvese que si a es un elemento en una álgebra C^* unital entonces podemos escribir al elemento a como $a = b + ic$ donde b y c son elementos hermitianos. En más detalle tenemos que $b = (a + a^*)/2$ y $c = (a - a^*)/2i$. Por lo tanto, si κ es un caracter entonces

$$\kappa(a^*) = \kappa(b - ic) = \kappa(b) - i\kappa(c) = \overline{\kappa(b) + i\kappa(c)} = \kappa(a)^*,$$

lo que demuestra que κ es automáticamente un $*$ -homomorfismo.

Proposición 3.5.4. Sea A un álgebra de Banach conmutativa con unidad. Entonces el espectro de un elemento $a \in A$ está dado por

$$\sigma(a) = \{\kappa(a) : \kappa \text{ es un caracter de } A\}, \quad (3.10)$$

para todo $a \in A$. $\square$

Corolario 3.5.5. Sea A una álgebra C^* conmutativa con unidad. Si I es un ideal maximal de A entonces $I = I^*$.

Demostración. Por la correspondencia entre ideales maximales y caracteres tenemos que I es el kernel de un caracter $\kappa : A \rightarrow \mathbb{C}$, pero al ser A una álgebra C^* unital, automáticamente κ es un $*$ -homomorfismo, lo que muestra que I es un $*$ -ideal. $\square$

3.6 Teorema de Gelfand-Naimark

Definición 3.6.1 (Transformada de Gelfand). Sea A una álgebra C^ conmutativa con unidad. Llamamos a la función*

$$\begin{aligned}\wedge : A &\longrightarrow C(\text{Spec}(A)) \\ a &\mapsto \widehat{a}(\kappa) = \kappa(a)\end{aligned}$$

transformada de Gelfand.

De la definición de la topología $*$ -débil tenemos que $\widehat{a}$ es un funcional continuo para todo $a \in A$. Además si $a, b \in A$ entonces

$$[\widehat{ab}](\kappa) = \kappa(ab) = \kappa(a)\kappa(b) = \widehat{a}(\kappa)\widehat{b}(\kappa) = [\widehat{a}\widehat{b}](\kappa).$$

De modo que $\wedge$ es homomorfismo de álgebras. Por otro lado, si κ y ω pertenecen al $\text{Spec}(A)$ y son distintos, es porque existe algún $a \in A$ tal que $\kappa(a) \neq \omega(a)$ es decir, $\widehat{a}(\kappa) \neq \widehat{a}(\omega)$. Por lo tanto $\text{im}(\wedge)$ separa los puntos de $\text{Spec}(A)$, por lo que aplicando el teorema de Stone Weierstrass se tiene que la imagen es toda el álgebra de las funciones continuas sobre el $\text{Spec}(A)$.

Además

$$\|\widehat{a}\| = \sup\{|\widehat{a}(\kappa)| : \kappa \in \text{Spec}(A)\} = \sup\{|\lambda| : \lambda \in \sigma(a)\} = r(a). \quad (3.11)$$

Teorema 3.6.1. (Gelfand-Naimark) Sea A una álgebra C^ conmutativa con unidad. La transformada de Gelfand definida sobre A es un $*$ -isomorfismo isométrico.*

Demostración. Para cada elemento a de A , tomando en cuenta la igualdad dada en (3.11) se tiene que

$$\|a\|^2 = r(a^*a) = \|\widehat{a^*a}\| = \|\widehat{a}\widehat{a}\|^2 = \|\widehat{a}\|^2.$$

Por lo tanto es una isometría. Por los comentarios previos a este teorema ya se ha visto que $\wedge$ es un homomorfismo de álgebras, sólo falta ver que es un $*$ -homomorfismo, lo que es consecuencia de que los caracteres $\kappa \in \text{Spec}(A)$ son $*$ -homomorfismos:

$$\widehat{a^*}(\kappa) = \kappa(a^*) = \kappa(a)^* = (\widehat{a})^*(\kappa).$$

Por lo tanto la transformada de Gelfand es un $*$ -isomorfismo isométrico. $\square$

Este teorema nos muestra que toda álgebra C^* conmutativa unital se puede ver como el espacio $C(X)$ de funciones continuas sobre un espacio compacto.

Ahora, sea X un espacio topológico localmente compacto y $X^\dagger = X \cup \{\infty\}$ su compactificación por un punto. Entonces la unitalización de $C_0(X)$ y $C(X^\dagger)$ son isomorfas: el $*$ -homomorfismo

$$\psi : C_0(X)^\dagger \rightarrow C(X^\dagger), \quad \psi[(f, \lambda)] = \begin{cases} f(x) + \lambda & \text{si } x \in X \\ \lambda & \text{si } x = \infty \end{cases}$$

es biyectivo. De esta manera vemos que compactificar el espacio es equivalente a unitalizar el álgebra, y que en general la unitalización mínima corresponde a “compactificar” espacios cuánticos por un punto.

Teorema 3.6.2. Las categorías de álgebras C^ conmutativas con unidad, C^* -álg, y de espacios compactos, $Comp$, son duales. Esto es, existe un funtor contravariante de la categoría $Comp$ en la categoría C^* -álg.* $\square$

En otras palabras, cada $*$ -homomorfismo unital $\pi : A \rightarrow B$ entre álgebras C^* conmutativas y uniales se corresponde con una función $f : Y \rightarrow X$ continua, donde $A \cong C(X)$ y $B \cong C(Y)$. Si π es inyectivo, entonces f es sobre, si π es sobre entonces f es inyectivo y si π es $*$ -automorfismo entonces f es homeomorfismo.

Corolario 3.6.3. Sea A una álgebra C^ unital y $B \subseteq A$ una subálgebra C^* unital. Si $b \in B$ y $b \in GL(A)$ entonces $b \in GL(B)$, es decir*

$$\sigma_A(b) = \sigma_B(b).$$

En otras palabras el espectro de un elemento no depende del álgebra que lo contiene. $\square$

Proposición 3.6.4. Sean A una álgebra C^ unital y $a \in A$ un elemento normal de A . Entonces la álgebra C^* generada por a y a^* y denotada por $C^*[a, a^*]$, es isomorfa al álgebra de las funciones continuas definidas sobre $\sigma(a)$. Es decir,*

$$C^*[a, a^*] \cong C(\sigma(a)).$$

Demostración. La función $\psi : C(\sigma(a)) \rightarrow C^*[a, a^*]$ tal que $\psi(z) = a$, donde $z(\lambda) = \lambda$ para todo $\lambda \in \sigma(a)$, es un $*$ -homomorfismo isométrico biyectivo. $\square$

Ejemplo. Sea A una álgebra C^* conmutativa con unidad y $u \in A$. Entonces u es unitario si y solamente si $\sigma(u) \subseteq S^1 := \{z \in \mathbb{C} \mid |z| = 1\}$.

En efecto, si u es unitario entonces por la transformada de Gelfand se tiene que $r(u) = \|u\|$. Como $u^*u = 1$, entonces

$$\|u\|^2 = \|u^*u\| = \|1\| = 1.$$

Por lo tanto $r(u) = 1$. De manera análoga $r(u^*) = 1$. De esta manera tenemos que si $\lambda \in \sigma(u)$ entonces $|\lambda| \leq 1$. Por otro lado, también se tiene que $\lambda^{-1} \in \sigma(u^*)$ de manera que $|1/\lambda| \leq 1$. Por lo tanto $|\lambda| = 1$ lo que implica que $\sigma(u) \subseteq S^1$.

Inversamente, si $\sigma(u) \subseteq S^1$. Por la transformada de Gelfand tenemos que $C^*[u, u^*] \cong C(\sigma(u))$. El $*$ -isomorfismo está dado por $\psi : C^*[u, u^*] \rightarrow C(\sigma(u))$, $\psi(u) = z$, donde $z(\lambda) = \lambda$ para todo $\lambda \in \sigma(u) \subseteq S^1$. También se tiene que $z^*(\lambda) = \bar{\lambda}$, para todo $\lambda \in \sigma(u) \subseteq S^1$.

Por lo tanto, $zz^*(\lambda) = |\lambda|^2 = 1 = z^*z(\lambda)$, por ser ψ isomorfismo isométrico se debe cumplir

$$uu^* = 1 = u^*u.$$

Así tenemos que u es unitario.

Definición 3.6.2. Sea A una álgebra C^* unital y $a \in A$ tal que $a = a^*$. Decimos que a es positivo, y escribimos $a \geq 0$, si $\sigma(a) \subseteq [0, \infty)$.

Observación 3.6.1. Observemos que en una álgebra C^* conmutativa si $a \in A$ es positivo, entonces existe un único elemento $\sqrt{a}$ tal que $\sqrt{a}^2 = a$. También podemos afirmar que si $a = a^*$ entonces existen únicos $a_+ \geq 0$ y $a_- \geq 0$ tal que $a = a_+ - a_-$ y $a_+a_- = 0 = a_-a_+$. En ambos casos estos elementos adicionales $\sqrt{a}, a_{\pm}$ pertenecen al álgebra $C^*[a]$.

Proposición 3.6.5. Los elementos positivos en una álgebra C^* se pueden definir equivalentemente como aquellos de la forma a^*a donde a es arbitrario. $\square$

Proposición 3.6.6. Sea A una álgebra C^* unital. Entonces cada $a \in A$ es una combinación lineal de no más de 4 elementos unitarios.

Demostración. Sea $a = a^*$ tal que $\|a\| \leq 1$, entonces $a^2 \geq 0$ y $1 - a^2 \geq 0$. Por lo tanto está bien definida la raíz cuadrada de $1 - a^2$. Definimos los elementos unitarios en A como

$$u = a + i\sqrt{1 - a^2}, \quad u^* = a - i\sqrt{1 - a^2}.$$

De esta forma se tiene que

$$a = \frac{u + u^*}{2}.$$

Ahora supongamos que $a \neq a^*$ y que se cumple que $\|(a + a^*)/2\| \leq \alpha$ y $\|(a - a^*)/(2i)\| \leq \beta$. Entonces

$$a = \alpha \left(\frac{a + a^*}{2\alpha} \right) + i\beta \left(\frac{a - a^*}{2i\beta} \right).$$

Aplicamos la primera parte de la demostración a cada una de las componentes de la suma anterior de manera que podemos escribir a cualquier elemento en el álgebra como combinación lineal de no más de 4 operadores unitarios. $\square$

Observación 3.6.2. Se puede demostrar que para álgebras C^* en general, podemos definir que un operador u es unitario si y sólo si, u es invertible y se cumple que

$$\|u\| = 1 = \|u^{-1}\|. \quad (3.12)$$

Esto nos enseña que ser unitario en un álgebra C^* depende exclusivamente de la norma. Y por otro lado, la $*$ actúa sobre los unitarios transformándolos en sus respectivos inversos. De esta manera, tomando en cuenta la proposición anterior podemos definir la estructura $*$ exclusivamente a través de la norma. Conjuntamente con la observación 3.4.3, podemos decir que la norma y la $*$ se determinan mutuamente.

Sea A una álgebra C^* conmutativa unital y $p \in A$ un proyector. Entonces sabemos que $\sigma(p) = \{0, 1\}$. Sean

$$\begin{aligned} X &= \{\kappa \in \text{Spec}(A) \mid \kappa(p) = 1\} = \widehat{p}^{-1}(\{1\}), \\ Y &= \{\kappa \in \text{Spec}(A) \mid \kappa(p) = 0\} = \widehat{p}^{-1}(\{0\}), \end{aligned}$$

donde $\widehat{\cdot}$ denota la transformada de Gelfand. Como $\widehat{p}$ es continua, entonces X y Y son conjuntos cerrados tales que $X \cap Y = \emptyset$ y $\text{Spec}(A) = X \cup Y$. Además como $X = \text{Spec}(A) \setminus Y$ y $Y = \text{Spec}(A) \setminus X$ entonces son abiertos. Lo que nos muestra que el espacio asociado al álgebra es desconexo.

De esta forma, los proyectores a nivel topológico clásico corresponden a las manifestaciones de la desconexidad del espacio. Por otro lado, si el espacio geométrico X es desconexo, entonces para los abiertos U y V que forman la desconexión del espacio, definimos la función $q : X \rightarrow \mathbb{C}$ como la característica de uno de ellos, digamos V . Es decir $q(U) = \{0\}$ y $q(V) = \{1\}$. Esta función es continua sobre X y $q = q^2 = q^*$. Esto es, si el espacio es desconexo, entonces el álgebra asociada contiene proyectores.

Observación 3.6.3. En álgebras C^* conmutativas cada idempotente es proyector.

Los ideales cerrados de un álgebra C^* conmutativa unital A nos definen en un sentido geométrico las partes del espacio asociado. Podemos considerar la proyección $p : A \rightarrow A/I$ entre el álgebra A y el álgebra cociente A/I , donde $A \cong C(X)$ y $A/I \cong C(Y)$. Esto a nivel dual, nos define una función inyectiva $f : Y \rightarrow X$ entre los espacios topológicos asociados, de manera que podemos ver al espacio Y como subconjunto cerrado de X , y en un nivel intuitivo, pensar a Y como *parte* de X . Por tanto, entre más ideales tenga un álgebra, más partes podemos encontrar en el espacio cuántico asociado. En particular, si el álgebra es simple, entonces el espacio asociado no tiene ninguna parte.

Ejemplo. Vamos a demostrar que el álgebra $M_2(\mathbb{C})$ no tiene ideales propios. Usando la notación de Dirac definimos los vectores

$$|1\rangle = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad |2\rangle = \begin{pmatrix} 0 \\ 1 \end{pmatrix},$$

de manera que el conjunto

$$\{|1\rangle\langle 1|, |1\rangle\langle 2|, |2\rangle\langle 1|, |2\rangle\langle 2|\},$$

es la base canónica de $M_2(\mathbb{C})$. Notemos que la matriz identidad I , se escribe como $I = |1\rangle\langle 1| + |2\rangle\langle 2|$.

Suponamos que $J \neq \{0\}$ es un ideal de $M_2(\mathbb{C})$ y $T : \mathbb{C}^2 \rightarrow \mathbb{C}^2$ un operador no cero en J . En particular si $a, b \in \mathbb{C}^2 \setminus \{0\}$ y $Ta \neq 0$ tenemos que la diada $|Ta\rangle\langle b| = T|a\rangle\langle b| \in J$. Por lo tanto,

$$\|Ta\|^2 \|b\|^2 |i\rangle\langle i| = |i\rangle\langle Ta|(|Ta\rangle\langle b|)|b\rangle\langle i| \in J,$$

donde $i = 1, 2$. Observando que J es en particular un espacio vectorial, obtenemos que la matriz identidad pertenece al ideal. Por lo tanto $J = M_2(\mathbb{C})$. Esto demuestra que el álgebra $M_2(\mathbb{C})$ no tiene ideales propios.

El mismo argumento funciona para demostrar que el álgebra $M_n(\mathbb{C})$ no tiene ideales propios, es decir, que es un álgebra simple. Más aún, si $\mathcal{H}$ es un espacio de Hilbert de dimensión infinita separable, las sumas de diadas son operadores de rango finito y éstas forman el mínimo ideal no-trivial. Y el máximo ideal propio en $B(\mathcal{H})$ es el ideal de operadores compactos $K(\mathcal{H})$.

Proposición 3.6.7. Sea A una álgebra C^* con unidad y $a \in A$ un elemento hermitiano. Entonces $a \geq 0$ si y sólo si $a = b^2$, con $b \in A$ tal que $b = b^*$.

También, si existe un $t > 0$ tal que $\|t - a\| \leq t$ entonces a es positivo. En tal caso podemos siempre poner $t \geq \|a\|$.

Demostración. Supongamos que se cumple que $a \geq 0$, entonces existe una única raíz cuadrada positiva $b \geq 0$ tal que $a = b^2$.

Inversamente si $b = b^*$ y $a = b^2$, entonces el espectro del elemento b está contenido en el intervalo $[-\|b\|, \|b\|]$ y $\sigma(a) = \sigma(b^2) = \sigma(b)^2 \subset [0, \|b\|^2]$, y por lo tanto $a \geq 0$.

Observemos que para $a = a^*$ se tiene que $C^*[1, a] \cong C(\sigma(a))$ (por la transformada de Gelfand), $\sigma(a) \subseteq [-\|a\|, \|a\|]$ y

$$\|t - a\| = \|t - z\|_\infty = \sup_{\lambda \in \sigma(a)} |t - z(\lambda)| = \sup_{\lambda \in \sigma(a)} |t - \lambda|,$$

donde $z(\lambda) = \lambda$ para todo $\lambda \in \sigma(a)$. Sea $a \geq 0$, si $t \geq \|a\|$ entonces $\sigma(a) \subseteq [0, t]$ y por tanto

$$\|t - a\| = \sup_{\lambda \in \sigma(a)} |t - \lambda| \leq \sup_{\lambda \in [0, t]} |t - \lambda| \leq t.$$

Ahora supongamos que $\|t - a\| \leq t$ para algún $t \geq \|a\|$. Supongamos que $\sigma(a) \cap (-\infty, 0) \neq \emptyset$, entonces existe $\mu < 0$ tal que $\mu \in \sigma(a)$. Por lo tanto

$$\|t - a\| = \sup_{\lambda \in \sigma(a)} |t - \lambda| \geq |t - \mu| = t + |\mu| > t,$$

lo cual contradice lo supuesto. Luego $\sigma(a) \subseteq [0, \|a\|]$. Por lo tanto $a \geq 0$, lo que completa la demostración. $\square$

3.7 Estados y representaciones

Definición 3.7.1. Sea A una álgebra C^* unital. Un mapeo $\rho : A \longrightarrow \mathbb{C}$ es un estado si ρ es lineal, positivo y normalizado, es decir, para todo $a, b \in A$ y $\lambda \in \mathbb{C}$ se satisface:

1. $\rho(a + \lambda b) = \rho(a) + \lambda \rho(b)$,
2. $\rho(a) \geq 0$ para todo $a \geq 0$,
3. $\rho(1) = 1$.

Denotamos por $S(A)$ al conjunto de todos los estados sobre A .

Proposición 3.7.1. Sea A una álgebra C^* unital y $f : A \rightarrow \mathbb{C}$ un funcional lineal tal que $f(1) = 1$. Entonces f es un estado si y sólo si f es acotado y $\|f\| = 1$.

Demostración. Supongamos f es un estado, es decir $f(a^*a) \geq 0$, para todo $a \in A$. Entonces tomando en cuenta que $a^*a \leq \|a^*a\|$ se tiene

$$|f(a)|^2 \leq f(1)f(a^*a) = f(a^*a) \leq \|a^*a\| = \|a\|^2.$$

Es decir, f es acotado y $\|f\| \leq 1$, pero como $f(1) = 1$, entonces $\|f\| = 1$.

Inversamente, si $\|f\| = 1 = f(1)$ y $a \geq 0$, escribimos $f(a) = \alpha + i\beta$, donde $\alpha, \beta \in \mathbb{R}$. Queremos ver que $\beta = 0$ y que $\alpha \geq 0$. Sea $b = a - \alpha + it\beta \in A$, $t \in \mathbb{R}$. Entonces

$$|f(b)|^2 = |f(a) - \alpha + it\beta|^2 = |i\beta(1+t)|^2 = \beta^2(1+t)^2,$$

y dado que la norma de f es igual a uno, tenemos

$$|f(b)|^2 \leq \|b\|^2 = \|b^*b\| = \|(a - \alpha)^2 + t^2\beta^2\| \leq \|a - \alpha\|^2 + t^2\beta^2.$$

Juntando las dos últimas expresiones tenemos que

$$\beta^2(1+t)^2 \leq \|a - \alpha\|^2 + t^2\beta^2, \quad \forall t \in \mathbb{R}.$$

Es decir

$$\beta^2 + 2t\beta^2 \leq \|a - \alpha\|^2, \quad \forall t \in \mathbb{R}.$$

Observando además que t puede ser tan grande como deseamos, se debe tener que $\beta = 0$.

Ahora bien, como $a \geq 0$ entonces $\sigma(a) \subseteq [0, \|a\|]$, lo que nos lleva a que $\sigma(1 - a/\|a\|) \subseteq [0, 1]$. De esta manera $\|1 - a/\|a\|\| \leq 1$ y por lo tanto

$$\left| f\left(1 - \frac{a}{\|a\|}\right) \right| \leq \|1 - \frac{a}{\|a\|}\| \leq 1.$$

Además como $f(1 - a/\|a\|) = 1 - \alpha/\|a\|$, insertándolo en la última ecuación se tiene que

$$\left| 1 - \frac{\alpha}{\|a\|} \right| \leq 1,$$

lo que muestra que $\alpha \geq 0$, por lo que f es positivo. $\square$

Definición 3.7.2. Un conjunto S de un espacio afín es convexo si para cualesquiera $a, b \in S$ todo el segmento que conecta a y b también pertenece al conjunto S .

Proposición 3.7.2. Sea A una álgebra C^* unital, el conjunto $S(A)$ de estados de A forman un conjunto convexo no vacío en el espacio dual A^* que además es compacto con la topología $*$ -débil. $\square$

Definición 3.7.3. Sea S un conjunto convexo y $z \in S$. Decimos que z es un punto extremal de S , si para todo $a, b \in S$ tales que $z = ta + (1-t)b$, $t \in (0, 1)$ entonces $z = a$ o $z = b$.

Denotamos por $\text{ext}(S)$ al conjunto de todos los puntos extremales de S .

Definición 3.7.4. Sea A una álgebra C^* unital, llamamos estados puros a los elementos del conjunto $\text{ext}(S(A))$.

Sea X un espacio afín topológico localmente convexo y $S \subseteq X$. Denotamos por $\text{co}(S)$ a la cáscara convexa de S , es decir, a la intersección de todos los conjuntos convexos que contienen a S . Los elementos de $\text{co}(S)$ son todas las combinaciones convexas finitas de los elementos de S .

Teorema 3.7.3 (Krein-Milman). Sea X un espacio afín topológico localmente convexo y $S \subseteq X$ un subconjunto compacto y convexo no vacío de X , entonces

$$S = \overline{\text{co}}(\text{ext}(S)),$$

es decir, S es la cerradura de la cáscara convexa de sus puntos extremales. $\square$

Observación 3.7.1. Como consecuencia del teorema anterior se tiene que el conjunto $\text{ext}(S) \neq \emptyset$. En particular, en las álgebras C^* , siempre existen los estados puros y ellos son suficientes para reconstruir el espacio completo de todos los estados del álgebra. En el caso que A es una álgebra C^* conmutativa unital se puede demostrar que el conjunto de estados puros del álgebra coincide con el conjunto de caracteres, que como habíamos mencionado anteriormente, corresponden a los puntos del espacio. Así, estado puro y caracter son dos formas equivalentes de describir los puntos del espacio clásico asociado.

Ejemplo (Estados en $M_2(\mathbb{C})$). Existe una correspondencia biunívoca entre elementos ρ del espacio dual de $M_2(\mathbb{C})$ y matrices $\hat{\rho}$ tal que

$$\rho(b) = \text{Tr}(\hat{\rho}b).$$

En particular, ρ es un estado si y solo si la matriz $\hat{\rho}$ es positiva y de traza 1.

Consideremos la base en $M_2(\mathbb{C})$ dada por las matrices de Pauli y la matriz identidad. Como las matrices de Pauli tienen traza cero, entonces

$$\hat{\rho} = \frac{1}{2} (I + x\sigma_1 + y\sigma_2 + z\sigma_3) = \frac{1}{2} \begin{pmatrix} 1+z & x-iy \\ x+iy & 1-z \end{pmatrix}.$$

Como $\hat{\rho} = \hat{\rho}^*$ entonces $x, y, z \in \mathbb{R}$. Además como $\hat{\rho}$ es positiva entonces sus valores propios son positivos y por lo tanto $\det(\hat{\rho}) = 1 - x^2 - y^2 - z^2 \geq 0$. Es

decir, cada estado en las matrices queda determinado por un punto de la bola unitaria en $\mathbb{R}^3$. Además podemos observar que el conjunto de estados puros están en correspondencia con los puntos $(x, y, z) \in \mathbb{R}^3$ tales que

$$x^2 + y^2 + z^2 = 1.$$

De esta forma se tiene que los estados puros de $M_2(\mathbb{C})$ forman una esfera. Por lo tanto, podemos imaginar que el álgebra de matrices $M_2(\mathbb{C})$ tiene asociado un espacio cuántico esférico.

Teorema 3.7.4 (Teorema de representación para funcionales lineales). *Sea X un espacio localmente compacto. Para cada funcional lineal positivo y continuo ϕ sobre $C_0(X)$ existe una única medida de Borel regular μ sobre X tal que*

$$\phi(f) = \int_X f(x)\mu(x),$$

para todo $f \in C_0(X)$. □

Observación 3.7.2. Como consecuencia de este teorema existe una correspondencia uno a uno entre los estados de una álgebra C^* conmutativa unital y las medidas de probabilidad en el espacio subyacente. Más aún, los estados puros serán interpretados como aquellas medidas de probabilidad que no pueden escribirse como mezcla de otras. Es decir, las medidas donde no hay incertidumbre. Estas medidas cuando toman el valor uno, nos dicen el evento ocurre con certeza y cuando toma el valor cero corresponde a la situación que nunca sucede. El conocimiento completo del sistema es equivalente a la especificación de su estado físico, que es un punto del espacio X . El evento $\Lambda \subseteq X$ ocurre con certeza si y solo si $\Lambda \ni x$ y no ocurre si y solo si $x \notin \Lambda$. Por lo tanto, podemos decir que los puntos del espacio están en correspondencia biunívoca con los estados puros (las medidas de Dirac) del álgebra, como ya habíamos mencionado.

Diremos que una representación $\pi : A \rightarrow \mathcal{B}(\mathcal{H})$ es irreducible si el conjunto $\pi(A)$ es irreducible, esto es, los únicos subespacios lineales cerrados de $\mathcal{H}$ que son invariantes bajo todos los operadores de $\pi(A)$ son los subespacios triviales $\{0\}$ y $\mathcal{H}$.

En particular, cada $*$ -representación $\pi : A \rightarrow \mathcal{B}(\mathcal{H})$ donde $\dim(\mathcal{H}) = 1$ es irreducible. Estas representaciones se pueden ver como caracteres del álgebra y por lo tanto podemos decir que corresponden a *puntos clásicos* del espacio.

Una representación $\pi : A \rightarrow \mathcal{B}(\mathcal{H})$ es cíclica, si existe un vector $\xi \in \mathcal{H}$ tal que

$$\overline{\pi(A)\xi} = \mathcal{H}.$$

Al vector ξ se le llama vector cíclico para la representación π .

Teorema 3.7.5. *Sea $\pi : A \rightarrow \mathcal{B}(\mathcal{H})$ una $*$ -representación no trivial. Las siguientes afirmaciones son equivalentes:*

1. π es irreducible;
2. Cada vector no cero de $\mathcal{H}$ es un vector cíclico para π ;
3. Los únicos operadores acotados que conmutan con todos los operadores de $\pi(A)$, son los operadores escalares, es decir de la forma λI , con $\lambda \in \mathbb{C}$ e I el operador identidad. $\square$

Por último, mencionaremos un teorema que establece el puente entre estados y representaciones del álgebra. Cuando los estados son puros, es decir, que corresponden a puntos del espacio clásico, las representaciones son irreducibles. En el contexto conmutativo, como hemos visto a lo largo de este capítulo, existen diferentes formulaciones equivalentes para punto de un espacio compacto: caracter, ideal maximal, estado puro y como veremos ahora, los estados puros tienen una interesante conexión con las representaciones irreducibles.

Teorema 3.7.6 (Gelfand-Naimark-Segal). *Sea A una álgebra C^* unital y sea $\rho : A \rightarrow \mathbb{C}$ un estado en A . Entonces existe una $*$ -representación cíclica $\pi : A \rightarrow \mathcal{B}(\mathcal{H})$ de A sobre un espacio de Hilbert $\mathcal{H}$ y un vector cíclico $\xi \in \mathcal{H}$ tal que $\|\xi\| = 1$ y*

$$\rho(a) = \langle \xi | \pi(a) \xi \rangle,$$

para todo $a \in A$. La representación es única salvo equivalencia unitaria.

Demostración. Para todo $a, b \in A$ definimos un producto escalar en A como $\langle a | b \rangle = \rho(a^*b)$. Sin embargo, notemos que puede pasar que $\langle a | a \rangle = 0$ y $a \neq 0$. Para tener un producto escalar no degenerado, sea

$$I = \{a \in A : \rho(a^*a) = 0\} = \{a \in A : \rho(a^*b) = 0, \forall b \in A\}.$$

Observemos que I es un ideal izquierdo de A y por lo tanto podemos considerar el espacio cociente $L = A/I$. Este espacio vectorial es un espacio pre-Hilbert con el producto escalar inducido:

$$\langle a + I | b + I \rangle = \rho(a^*b).$$

Sea $\pi : A \rightarrow \mathcal{B}(L)$ el mapeo lineal de A en los operadores acotados sobre L dado por $\pi(a) = D_a$ y $D_a(b + I) = ab + I$.

Para cada $a \in A$ y $b \notin I$ se cumple que

$$\frac{\|D_a(b + I)\|^2}{\|b + I\|^2} = \frac{\|ab + I\|^2}{\|b + I\|^2} = \frac{\rho(b^*a^*ab)}{\rho(b^*b)} \leq \|a\|^2,$$

lo que nos indica que $\|D_a\| \leq \|a\|$. Por lo tanto D_a es un operador acotado sobre L .

Obsérvese que π es un $*$ -homomorfismo pues para todo $a, b, c \in A$ se tiene que

$$\begin{aligned}\pi(ab)(c + I) &= D_{ab}(c + I) = (ab)c + I = a(bc) + I \\ &= D_a D_b(c + I) = \pi(a)\pi(b)(c + I),\end{aligned}$$

es decir, es multiplicativo. Además también tenemos que

$$\begin{aligned}\langle D_a(b + I)|c + I \rangle &= \langle ab + I|c + I \rangle = \rho((ab)^*c) = \rho((b^*a^*)c) = \rho(b^*(a^*c)) \\ &= \langle b + I|a^*c + I \rangle = \langle b + I|D_{a^*}(c + I) \rangle,\end{aligned}$$

es decir $D_a^* = D_{a^*}$ y por lo tanto π respeta la $*$. Finalmente $\pi(1) = I$, donde I es el operador identidad en $\mathcal{B}(L)$. De esta manera hemos demostrado que π es un $*$ -homomorfismo.

Sea $\mathcal{H}$ la completación de L a un espacio de Hilbert, y sea $\mathcal{D}_a$ la única extensión de D_a a un operador lineal acotado en $\mathcal{H}$. Podemos ver inmediatamente que si $\xi = 1 + I$ entonces $\|\xi\| = 1$ y para todo a en A se tiene que

$$\rho(a) = \langle 1 + I|a + I \rangle = \langle \xi|\mathcal{D}_a(\xi) \rangle = \langle \xi|\pi(a)\xi \rangle.$$

La representación π es cíclica con vector cíclico ξ pues

$$\overline{\{\mathcal{D}_a(\xi) : a \in A\}} = \overline{\{a + I : a \in A\}} = \mathcal{H}.$$

Para demostrar la unicidad de la representación salvo operadores unitarios, suponemos que existen dos representaciones cíclicas $\pi_1 : A \rightarrow \mathcal{B}(\mathcal{H}_1)$ y $\pi_2 : A \rightarrow \mathcal{B}(\mathcal{H}_2)$ con vectores cíclicos ξ_1 y ξ_2 , respectivamente, tal que

$$\langle \xi_2|\pi_2(a)\xi_2 \rangle = \rho(a) = \langle \xi_1|\pi_1(a)\xi_1 \rangle, \quad \forall a \in A.$$

Entonces para todo $a, b \in A$ se satisface

$$\langle \pi_2(a)\xi_2|\pi_2(b)\xi_2 \rangle = \langle \xi_2|\pi_2(a^*b)\xi_2 \rangle = \langle \xi_1|\pi_1(a^*b)\xi_1 \rangle = \langle \pi_1(a)\xi_1|\pi_1(b)\xi_1 \rangle.$$

De esta forma, si a es igual a b en la ecuación anterior, se tiene que

$$\|\pi_2(a)\xi_2\| = \|\pi_1(a)\xi_1\|,$$

para todo $a \in A$.

Sea $u_0 : \pi_1(A)\xi_1 \rightarrow \pi_2(A)\xi_2$ tal que $\pi_1(a)\xi_1 \mapsto \pi_2(a)\xi_2$. Notemos que $\|u_0(\pi_1(a)\xi_1)\| = \|\pi_1(a)\xi_1\|$, es decir, u_0 es un operador unitario de un lineal

denso de $\mathcal{H}_1$ a un lineal denso de $\mathcal{H}_2$, y por lo tanto se puede extender a un operador unitario $u : \mathcal{H}_1 \rightarrow \mathcal{H}_2$. Este operador verifica que

$$\begin{aligned} u\pi_1(a)u^{-1}(\pi_2(b)\xi_2) &= u\pi_1(a)(\pi_1(b)\xi_1) = u\pi_1(ab)\xi_1 = \pi_2(ab)\xi_2 \\ &= \pi_2(a)(\pi_2(b)\xi_2). \end{aligned}$$

Por lo tanto $u\pi_1(a)u^{-1} = \pi_2(a)$, lo que demuestra que la representación es única salvo equivalencia unitaria. $\square$

Así, con este teorema podemos establecer una correspondencia entre estados que representan medidas de probabilidad y la teoría de representaciones sobre espacios de Hilbert. Las medidas de probabilidad que corresponden a estados puros, las podemos identificar como representaciones irreducibles como lo mencionamos en la siguiente proposición.

Proposición 3.7.7. Sean A una álgebra C^ unital, $\rho : A \rightarrow \mathbb{C}$ un estado y π_ρ la $*$ -representación del Teorema GNS anterior. Entonces ρ es puro si y solamente si π_ρ es irreducible.* $\square$

3.8 Simetrías y grupos cuánticos

Los grupos de simetría juegan un papel fundamental para describir geométricamente a los espacios clásicos: es a través de los invariantes bajo un grupo de simetrías que se estudia la geometría del espacio. Así que es de vital importancia tener una traducción cuántica de estas estructuras para describir la geometría de los espacios cuánticos.

Ya hemos visto que existe una correspondencia entre álgebras C^* conmutativas (unidades) y los espacios topológicos localmente compactos (compactos). En términos de esta correspondencia los elementos del álgebra se interpretan como las funciones continuas sobre el espacio correspondiente (que desaparecen en la infinidad topológica del espacio en el caso no-compacto) topológico compacto.

Las transformaciones continuas propias (es decir cuyas pre-imágenes de compactos son compactos) entre los espacios se corresponden, en forma contravariante, con $*$ -homomorfismos de sus álgebras, de manera que las simetrías del espacio se corresponden con $*$ -automorfismos de su álgebra.

Así, para los espacios cuánticos, diremos que las simetrías de éstos están dadas por los automorfismos del álgebra. Y obsérvese que en este caso, cuando las álgebras no son conmutativas, siempre existen automorfismos internos. Tenemos entonces, ejemplos de espacios cuánticos sin puntos pero con muchas simetrías, un ejemplo de este tipo de espacio, es el dado por el álgebra de matrices 2×2 .

Ejemplo (Simetrías de conjuntos finitos). Supongamos que X es un conjunto de n elementos. Este conjunto es compacto con la topología discreta. Entonces por la teoría de Gelfand-Naimark, estudiar este espacio es equivalente a estudiar el álgebra de funciones continuas sobre dicho espacio. Dado que X es finito, el álgebra $C(X)$ es isomorfa a $\mathbb{C}^n$ y por lo tanto la dimensión de $C(X)$ es igual al número de elementos de X .

De hecho, los espacios clásicos finitos así son caracterizados con la dimensionalidad finita de su álgebra C^* conmutativa.

Ahora bien el grupo de simetrías de X consiste de todos sus homeomorfismos. Observemos que todas las funciones en la topología discreta son continuas y por lo tanto, el grupo de simetrías de X está dado por el grupo de permutaciones de sus n elementos.

Ejemplo (Simetrías de $M_2(\mathbb{C})$). Estudiemos ahora las simetrías del espacio asociado al álgebra C^* más simple no conmutativa: $M_2(\mathbb{C})$; el álgebra de matrices con coeficientes complejos de tamaño 2×2 .

Como primer paso veremos que todos los automorfismos de $M_2(\mathbb{C})$ son internos. En efecto, sea $F : M_2(\mathbb{C}) \rightarrow M_2(\mathbb{C})$ un automorfismo entre las matrices cuadradas de tamaño 2×2 y sea $w \in \mathbb{C}^2$ un elemento tal que $\|w\| = 1$. Observemos que $F(|w\rangle\langle w|) \neq 0$ pues $|w\rangle\langle w|$ es idempotente. Fijemos $z \in \mathbb{C}^2 \setminus \{0\}$ tal que $F(|w\rangle\langle w|)|z\rangle \neq 0$, y definamos el operador lineal $g : \mathbb{C}^2 \rightarrow \mathbb{C}^2$, $g(|u\rangle) = F(|u\rangle\langle w|)|z\rangle$. Como $g(|w\rangle) \neq 0$ entonces el operador g no es el operador cero y además como F es automorfismo, dado un vector $|\psi\rangle \in \mathbb{C}^2$ siempre existe una matriz a tal que $F(a)g|w\rangle = |\psi\rangle$. Por lo tanto tenemos que $g(|aw\rangle) = F(|aw\rangle\langle w|)|z\rangle = F(a)g(|w\rangle) = |\psi\rangle$, lo que nos muestra que el operador lineal es sobre y por lo tanto g debe ser inyectivo. Es decir, g es invertible.

Así que, para todo $a \in M_2(\mathbb{C})$ se tiene que

$$ga|u\rangle = g|au\rangle = F(|au\rangle\langle w|)|z\rangle = F(a)F(|u\rangle\langle w|)|z\rangle = F(a)g|u\rangle.$$

Por lo que $ga = F(a)g$, o bien, $F(a) = gag^{-1}$, para todo $a \in M_2(\mathbb{C})$ y $g \in \text{GL}(2, \mathbb{C})$, es decir, F es un automorfismo interno.

Más aún, si $F : M_2(\mathbb{C}) \rightarrow M_2(\mathbb{C})$ es un $*$ -automorfismo de $M_2(\mathbb{C})$, veamos que $F(a) = uau^*$ donde $u^* = u^{-1}$. En efecto, por la parte anterior tenemos que existe $g \in \text{GL}(M_2(\mathbb{C}))$ tal que F es el automorfismo interno dado por g . Usando que F es $*$ -automorfismo tenemos que

$$g^{-1*}a^*g^* = ga^*g^{-1},$$

para todo a en $M_2(\mathbb{C})$. Por lo tanto g^*g pertenece al centro del álgebra de las matrices, es decir, existe $\lambda > 0$ tal que $g^*g = \lambda I$. Tomando $u = g/\sqrt{\lambda}$

tenemos que u es unitario y que

$$uau^* = \frac{g}{\sqrt{\lambda}} a \frac{g^*}{\sqrt{\lambda}} = \frac{1}{\lambda} ga(\lambda g^{-1}) = F(a).$$

Por lo tanto los $*$ -automorfismos de $M_2(\mathbb{C})$ son de la forma $F(a) = uau^*$, donde u es un elemento unitario.

Veamos que si no consideramos la estructura $*$ entonces el grupo de automorfismos de $M_2(\mathbb{C})$ es isomorfo al grupo cociente $\text{GL}(2, \mathbb{C})/\mathbb{C}_*$, donde $\mathbb{C}_* = \mathbb{C} \setminus \{0\}$, el cual a su vez es isomorfo al grupo de transformaciones de Möbius $M(2)$.

En efecto, el homomorfismo de grupos $\text{GL}(2, \mathbb{C}) \rightarrow \text{Aut}(M_2(\mathbb{C}))$ dado por $g \mapsto F$, es sobre y el kernel está dado por el conjunto $\{\lambda I : \lambda \in \mathbb{C} \setminus \{0\}\} \cong \mathbb{C} \setminus \{0\} = \mathbb{C}_*$. Sumarizando este análisis, se tiene que

$$\text{GL}(2, \mathbb{C})/\mathbb{C}_* \cong \text{SL}(2, \mathbb{C})/\{I, -I\} \cong M(2) \cong \text{Aut}(M_2(\mathbb{C})).$$

Similarmente, si restringimos el homomorfismo anterior al grupo $\text{SU}(2)$ tenemos que el kernel consta de las matrices unitarias de determinante uno que conmutan con todas las matrices, es decir, el kernel está constituido por la matriz identidad y menos la matriz identidad. De esta forma, por el primer teorema del isomorfismo tendremos que el grupo de $*$ -automorfismos de $M_2(\mathbb{C})$ es isomorfo al grupo especial unitario factorizado por las matrices I y $-I$, esto es,

$$\text{Aut}(M_2(\mathbb{C}), *) \cong \text{SU}(2)/\{I, -I\}. \quad (3.13)$$

Además, podemos ver al campo no-conmutativo de cuaterniones $\mathbb{H}$, como una subálgebra real de $M_2(\mathbb{C})$, identificando a los imaginarios i, j, k con las matrices $-i\sigma_1, -i\sigma_2, -i\sigma_3$, respectivamente. Y diremos que un cuaternión ψ es imaginario si y sólo si $\psi^* = -\psi$ y los identificamos con los vectores de $\mathbb{R}^3$.

Notemos que si $\psi = \delta + \alpha i + \beta j + \gamma k$ es un cuaternión de norma uno, entonces se corresponde con la matriz

$$\sigma_\psi = \begin{pmatrix} \delta - i\gamma & -\beta - i\alpha \\ \beta - i\alpha & \delta + i\gamma \end{pmatrix},$$

unitaria de determinante uno, es decir, $\sigma_\psi \in \text{SU}(2)$. De esta manera, podemos ver a los elementos de $\text{SU}(2)$ como el conjunto de cuaterniones de norma uno, estableciendo también la identificación topológica de $\text{SU}(2)$ y la 3-esfera S^3 .

Consideremos el morfismo de grupos $F : \text{SU}(2) \rightarrow \text{SO}(3)$ dado por $\sigma \mapsto F_\sigma$, donde $F_\sigma : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ es la transformación dada por $F_\sigma(\psi) = \sigma\psi\sigma^*$. Primero que nada, veamos que F_σ está bien definida: en efecto, si $\psi \in \mathbb{R}^3$, es decir,

$\psi^* = -\psi$, entonces $(\sigma\psi\sigma^*)^* = -\sigma\psi\sigma^*$, lo que nos lleva a concluir que $\sigma\psi\sigma^* \in \mathbb{R}^3$, y por tanto F_σ está bien definida. Además

$$\|F_\sigma(\psi)\|^2 = (\sigma\psi\sigma^*)^*(\sigma\psi\sigma^*) = \|\psi\|^2,$$

implica que $F_\sigma \in O(3)$. También tenemos por la continuidad de la construcción que el determinante de $F_\sigma = 1$, es decir, $F_\sigma \in SO(3)$.

Veamos que F es sobre: sea $f \in SO(3)$ una rotación de $\mathbb{R}^3$ por un ángulo θ alrededor de la recta generada por el vector unitario η . Definimos el cuaternión unitario $\sigma = \cos(\theta/2) + \eta \sin(\theta/2)$. Si ψ pertenece a la recta generada por η , es decir $\psi = \lambda\eta$ para algún $\lambda \in \mathbb{R}$, entonces

$$F_\sigma(\psi) = (\cos(\theta/2) + \eta \sin(\theta/2)) \lambda \eta (\cos(\theta/2) - \eta \sin(\theta/2)) = \lambda \eta = \psi.$$

Por otro lado, si ψ es perpendicular a η entonces

$$\begin{aligned} F_\sigma(\psi) &= (\cos(\theta/2) + \eta \sin(\theta/2)) \psi (\cos(\theta/2) - \eta \sin(\theta/2)) \\ &= \psi \cos^2(\theta/2) + 2 \sin(\theta/2) \cos(\theta/2) (\eta \times \psi) - \psi \sin^2(\theta/2) \\ &= \psi \cos(\theta) + (\eta \times \psi) \sin(\theta), \end{aligned}$$

lo que nos lleva a concluir que $F_\sigma = f$ y por lo tanto F es sobre. Notando que el kernel de F está dado por las matrices $\sigma \in SU(2)$ tales que $\sigma\psi\sigma^* = \psi$, es decir, en términos de cuaterniones, son los cuaterniones unitarios que conmutan con todos los imaginarios, y estos son los cuaterniones que tienen parte imaginaria igual a cero. Por lo tanto $\ker(F) = \{I, -I\}$, lo que nos lleva usando el primer teorema de isomorfismo de grupos que

$$SU(2)/\{I, -I\} \cong SO(3). \quad (3.14)$$

Uniendo (3.13) y (3.14) obtenemos

$$\text{Aut}(M_2(\mathbb{C}), *) \cong SO(3).$$

Resumiendo, hemos visto que el grupo de automorfismos de $M_2(\mathbb{C})$ está dado por el grupo cociente $GL(2, \mathbb{C})/\mathbb{C}_*$ el cual es isomorfo al grupo de transformaciones de Möbius. Y si consideramos al álgebra de matrices con su estructura $*$ se tiene que el grupo de $*$ -automorfismos es isomorfo a $SU(2)/\{I, -I\}$ el cual a su vez es isomorfo al grupo $SO(3)$. Notemos que ambos grupos en geometría clásica corresponden a las simetrías de la 2-esfera, sea como superficie de Riemann (simetrías holomorfas) o bien equipada con su métrica Euclídeana. Notemos también, que aunque el álgebra es de dimensión finita – invitándonos de pensar intuitivamente sobre un espacio cuántico finito subyacente – su grupo de simetrías considerando o no la $*$ no lo es. Esto se puede decir que es un fenómeno particular de los espacios cuánticos.

Ya hemos visto que podemos pensar a las simetrías de un espacio como los automorfismos de su álgebra asociada. Sin embargo, no hemos mencionado aún la generalización cuántica de todo el grupo de simetrías de un espacio, donde lo *colectivo* se cuantiza y las simetrías individuales dejan de existir como tales. ¿Qué estructura algebraica se corresponde con la estructura de grupo? Los grupos de simetrías son de vital importancia para estudiar la geometría de los espacios clásicos. De hecho, según Felix Klein y su programa de Erlangen, cada geometría es el estudio de las propiedades que no cambian bajo un grupo de simetrías, y de la mano con esto, aparecen naturalmente los haces principales como fuente natural para describir la geometría de los espacios.

Así, para desarrollar un nuevo lenguaje que describa la geometría de los espacios cuánticos, necesitamos introducir los *grupos cuánticos*. Resulta que para la formulación cuántica del grupo de simetrías, se destaca por su simplicidad y generalidad el lenguaje diagramático [11].

Este lenguaje, puede ser formalizado como una categoría donde los objetos son los números naturales y los morfismos son diagramas que introduciremos a continuación. En general, el grupo cuántico se puede definir como realización de diagramas en un universo de espacios cuánticos. O bien, en términos formales, como un funtor de la categoría diagramática en la categoría de los espacios cuánticos.

Los principales elementos estructurales del lenguaje diagramático son

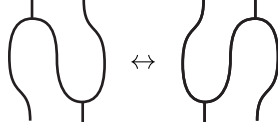


Los diagramas se leen de arriba hacia abajo y la transformación identidad se representa por una línea vertical. Morfismos-diagramas arbitrarios son las combinaciones finitas de los diagramas de producto y coproducto.

Como un eco de la idea de asociatividad, entonces se piden las siguientes igualdades diagramáticas:



Estas identidades, las complementamos con el postulado de que los siguientes diagramas mutuamente simétricos que mezclan 2 símbolos primarios, sean invertibles:



Por último, vamos a suponer que los diagramas cumplen la *ley de cancelación*: si dos diagramas encapsulados entre líneas verticales resultan así iguales, entonces la igualdad sigue válida después de quitar dichas líneas verticales. Esta propiedad es básica para el concepto de identidad en nuestra categoría diagramática.

Con estas simples propiedades podemos hablar de grupo en cualquier contexto matemático, sea clásico o cuántico. En particular, se puede deducir la existencia universal del ‘elemento neutro’ y ‘la inversa’.

Una importante realización de este programa es introducir grupos cuánticos a través de álgebras de Hopf. Aquí, el punto de partida es saber cuál es la estructura sobre un álgebra de funciones que refleja la estructura de grupo.

Para un grupo clásico G podemos escribir la estructura del grupo por medio de los mapeos de multiplicación $m : G \times G \rightarrow G$, dado por $m(x, y) = xy$, de los inversos $s : G \rightarrow G$, dado por $s(g) = g^{-1}$ y el mapeo que corresponde al neutro $e : \{1\} \rightarrow G$ definido como $e(1) = e$, donde $e \in G$ es el elemento neutro. Aquí nos interesa principalmente la *estructura geométrica* de G , que puede ser por ejemplo medible, topológica, suave (de Lie) o algebraica. En cada uno de estos contextos geométricos, como ya habíamos explicado, la estructura se especifica con el álgebra $\mathcal{A}$ de funciones apropiadas sobre G .

Si consideramos los ‘pullback’ de los mapeos que definen la estructura del grupo obtenemos sus contrapartes duales

$$\phi : \mathcal{A} \rightarrow \mathcal{A} \otimes \mathcal{A}, \quad \kappa : \mathcal{A} \rightarrow \mathcal{A}, \quad \epsilon : \mathcal{A} \rightarrow \mathbb{C},$$

definidas de la siguiente manera:

$$\phi(f)(x, y) = f(xy), \quad \kappa(f)(x) = f(x^{-1}), \quad \epsilon(f)(1) = f(e).$$

Observemos que los mapeos ϕ , κ y ϵ son homomorfismos de álgebras.

A lo largo de estas notas, usaremos la notación de Sweedler para simplificar la notación de $\phi(f)$ para todo $f \in \mathcal{A}$. Es decir, si $\phi(f) = \sum_{i=1}^n g_i \otimes h_i$, para algunos $g_i, h_i \in \mathcal{A}$ entonces escribimos lo mismo como $\phi(f) = f^{(1)} \otimes f^{(2)}$. Notemos que la propiedad de asociatividad en un grupo, es equivalente a la identidad $f((xy)z) = f(x(yz))$ para cada $f \in \mathcal{A}$ y $x, y, z \in G$. Esto se escribe en términos del homomorfismo ϕ de la siguiente manera:

$$\begin{aligned} f((xy)z) &= \phi(f)(xy, z) = f^{(1)}(xy) \otimes f^{(2)}(z) = \phi(f^{(1)})(x, y) \otimes f^{(2)}(z) \\ &= (\phi \otimes \text{id})\phi(f)(x, y, z), \end{aligned}$$

y por otro lado:

$$\begin{aligned} f(x(yz)) &= \phi(f)(x, yz) = f^{(1)}(x) \otimes f^{(2)}(yz) = f^{(1)}(x) \otimes \phi(f^{(2)})(y, z) \\ &= (\text{id} \otimes \phi)\phi(f)(x, y, z). \end{aligned}$$

Así la propiedad de asociatividad en G es equivalente a la siguiente propiedad de co-asociatividad en el álgebra $\mathcal{A}$:

$$(\phi \otimes \text{id})\phi = (\text{id} \otimes \phi)\phi. \quad (3.15)$$

De manera análoga, la propiedad del neutro e en el grupo G nos lleva a la siguiente relación entre los homomorfismos ϵ y ϕ :

$$(\epsilon \otimes \text{id})\phi = \text{id} = (\text{id} \otimes \epsilon)\phi, \quad (3.16)$$

y por último, la relación que se obtiene de la propiedad de los inversos se traduce a la siguiente relación:

$$m(\text{id} \otimes \kappa)\phi = 1_{\mathcal{A}}\epsilon = m(\kappa \otimes \text{id})\phi, \quad (3.17)$$

donde $m : \mathcal{A} \otimes \mathcal{A} \rightarrow \mathcal{A}$ ahora es la multiplicación (dada por $m(a \otimes b) = ab$) en el álgebra de funciones sobre el grupo y $1_{\mathcal{A}}$ es la función constante 1, es decir la unidad, en $\mathcal{A}$. Las tres ecuaciones (3.15), (3.16), (3.17), traducen la estructura de grupo a nivel de funciones sobre el grupo, y las álgebras que las satisfacen se les llama álgebras de Hopf. Así, esta es una manera de traducir la estructura de grupo en términos de un álgebra de funciones sobre el grupo.

Definición 3.8.1. Una biálgebra es un álgebra asociativa $\mathcal{A}$ equipada con el homomorfismo $\phi : \mathcal{A} \rightarrow \mathcal{A} \otimes \mathcal{A}$ llamado co-producto, tal que

$$(\phi \otimes \text{id})\phi = (\text{id} \otimes \phi)\phi. \quad (3.18)$$

Un álgebra de Hopf es una biálgebra con unidad tal que existen mapeos lineales $\epsilon : \mathcal{A} \rightarrow \mathbb{C}$ llamado co-unidad y $\kappa : \mathcal{A} \rightarrow \mathcal{A}$ llamado antípoda, cumpliendo con las siguientes relaciones:

$$(\epsilon \otimes \text{id})\phi = \text{id} = (\text{id} \otimes \epsilon)\phi, \quad (3.19)$$

$$m(\kappa \otimes \text{id})\phi = 1_{\mathcal{A}}\epsilon = m(\text{id} \otimes \kappa)\phi. \quad (3.20)$$

Observación 3.8.1. Para cada $a \in \mathcal{A}$, sea $\phi(a) = a^{(1)} \otimes a^{(2)}$. Usando la notación de Sweedler podemos escribir la propiedad (3.18) de la definición anterior, conocida como co-asociatividad, de la siguiente manera:

$$(a^{(1)} \otimes a^{(2)}) \otimes a^{(3)} = a^{(1)} \otimes (a^{(2)} \otimes a^{(3)}).$$

Esta notación nos dice que la ecuación (3.19) de la definición anterior se reescribe como $\epsilon(a^{(1)}) \otimes a^{(2)} = a = a^{(1)} \otimes \epsilon(a^{(2)})$, la cual usando el hecho de que $\mathbb{C} \otimes \mathcal{A} \cong \mathcal{A} \cong \mathcal{A} \otimes \mathbb{C}$ la podemos reescribir simplemente como

$$\epsilon(a^{(1)})a^{(2)} = a = a^{(1)}\epsilon(a^{(2)}).$$

Y la propiedad (3.20) en esta notación se verá como

$$\kappa(a^{(1)})a^{(2)} = \epsilon(a)1_{\mathcal{A}} = a^{(1)}\kappa(a^{(2)}).$$

Para los principales ejemplos de esta tesis, consideraremos precisamente a las álgebras de Hopf como los análogos cuánticos de los grupos. En estos contextos cuánticos $\mathcal{A}$ es un álgebra no conmutativa. En lo que sigue, escribiremos simplemente 1 para $1_{\mathcal{A}}$, si el contexto descarta posible ambigüedad.

Proposición 3.8.1. Sea $\mathcal{A}$ un álgebra de Hopf. Entonces los mapeos co-unidad y antípoda son únicamente determinados y se cumple lo siguiente:

$$\epsilon\kappa = \epsilon, \quad \kappa(1) = 1, \quad (3.21)$$

$$\phi(1) = 1 \otimes 1, \quad \epsilon(ab) = \epsilon(a)\epsilon(b), \quad (3.22)$$

$$\kappa(ab) = \kappa(b)\kappa(a), \quad \phi\kappa(a) = \kappa(a^{(2)}) \otimes \kappa(a^{(1)}), \quad (3.23)$$

para todo $a, b \in \mathcal{A}$. □

Observemos que el mapeo $\epsilon : \mathcal{A} \rightarrow \mathbb{C}$, es entonces un caracter del álgebra de Hopf $\mathcal{A}$, es decir, el grupo cuántico G correspondiente a $\mathcal{A}$ tiene al menos un punto clásico. Además, el conjunto de todos los caracteres de $\mathcal{A}$ forman un grupo que geométricamente corresponde a la parte clásica G_{cl} de G . La estructura de grupo clásico es dada por

$$gh = (g \otimes h)\phi, \quad h^{-1} = h\kappa, \quad (3.24)$$

y elemento neutro es la counidad ϵ .

Definición 3.8.2. Sea $\mathcal{A}$ un álgebra de Hopf. Decimos que $\mathcal{A}$ es una $$ -álgebra de Hopf si $\mathcal{A}$ tiene una involución $*$: $\mathcal{A} \rightarrow \mathcal{A}$ tal que para todo $a \in \mathcal{A}$ se cumple que*

$$\phi(a^*) = \phi(a)^* = a^{(1)*} \otimes a^{(2)*}, \quad (3.25)$$

es decir, el co-producto es un $$ -homomorfismo.*

Proposición 3.8.2. Sea $\mathcal{A}$ una $$ -álgebra de Hopf. Entonces $\epsilon(a^*) = \overline{\epsilon(a)}$. Además κ es invertible y $\kappa^{-1} = *\kappa*$. □*

Observación 3.8.2. La propiedad anterior para la inversa de κ se puede reescribir como $\kappa * \kappa * = \text{id}$ y es conocida como condición de Woronowicz.

El siguiente ejemplo de álgebra de Hopf, es el que utilizaremos extensivamente a lo largo de esta tesis. Esta álgebra corresponde al círculo unitario.

Ejemplo (Estructura de álgebra de Hopf sobre el círculo unitario). Sea S^1 el círculo unitario de números complejos y $u : S^1 \hookrightarrow \mathbb{C}$ la inclusión canónica. Es claro que $u^* = u^{-1}$. Consideremos la $*$ -álgebra $\mathcal{A}$ generada por u y u^* . Entonces $\mathcal{A}$ es una $*$ -álgebra de Hopf definiendo los mapeos de co-producto, co-unidad y antípoda de la siguiente manera:

$$\begin{aligned} \phi(u) &= u \otimes u, & \phi(u^*) &= u^* \otimes u^*, & \epsilon(u) &= 1 = \epsilon(u^*), \\ \kappa(u) &= u^*, & \kappa(u^*) &= u. \end{aligned} \tag{3.26}$$

De hecho, $\mathcal{A}$ es una $*$ -álgebra conmutativa unital y como vimos antes, naturalmente se completa en el álgebra C^* dada por $A \cong C(S^1)$ donde $\sigma(u) = S^1$. Así el grupo clásico asociado a dicha álgebra de Hopf no es otra cosa que el grupo de números complejos unitarios $U(1)$.

Capítulo 4

Cálculos diferenciales

En este capítulo abordaremos una generalización cuántica del cálculo integral y diferencial.

En el cálculo diferencial 1-dimensional real, uno de los teoremas más importantes es el Teorema Fundamental del Cálculo el cuál entrelaza la derivada y la integral como un tipo de operadores opuestos. En este cálculo diferencial, la diferencial de una función suave y cualquier función suave conmutan, es decir, el cálculo unidimensional es siempre conmutativo en el caso clásico. Para obtener un teorema fundamental del cálculo en dimensiones más altas necesitamos definir un integrando que nos de un número al integrar sobre dominios orientados.

Las k -formas son los integrandos correctos para integrar sobre dominios orientados k -dimensionales. Podemos pensar a las k -formas en $\mathbb{R}^n$ como funciones que toman k vectores y regresan un número $\varphi(v_1, \dots, v_k)$ tal que φ es multilineal y antisimétrica como una función de los vectores. O más intuitivamente, las k -formas toman k vectores, juegan con ellos para obtener una matriz de tamaño $k \times k$ y formar un paralelogramo de tal forma que arroja como resultado el área con signo de dicho paralelogramo.

Y una k -forma diferencial en un subconjunto U de $\mathbb{R}^n$ la pensamos como un mapeo φ que toma paralelogramos anclados en los puntos x de U y regresa números de acuerdo a la siguiente fórmula:

$$\varphi(P_x(v_1, \dots, v_k)) = \varphi(x)(v_1, \dots, v_k),$$

donde $\varphi(x)$ es una k -forma.

También podemos escribir a las formas diferenciales en U como expresiones de la forma

$$\varphi = \sum_{1 \leq i_1 < \dots < i_k \leq n} a_{i_1, \dots, i_k}(x) dx_{i_1} \wedge \dots \wedge dx_{i_k},$$

donde las $a_{i_1, \dots, i_k}(x)$ son funciones de U en los reales. De esta manera, las formas diferenciales son nuevos objetos que incluyen a las diferenciales del cálculo 1-dimensional.

El espacio de formas diferenciales forma un álgebra graduada

$$\Omega = \bigoplus_{k=0}^{\infty} \Omega^k,$$

donde Ω^0 contiene a todos los elementos de grado cero los cuales son funciones suaves en las coordenadas locales $x_1, \dots, x_n$ de la variedad; en Ω^1 encontramos todos los elementos de grado 1 que son todos los polinomios en las diferenciales con coeficientes en Ω^0 , y así sucesivamente y se cumple que si $\omega, \eta \in \Omega$ son homogéneas, entonces

$$\omega\eta = (-1)^{\partial\omega\partial\eta}\eta\omega,$$

donde ∂ es el operador de grado de formas diferenciales homogéneas.

Resumiendo lo anterior, se define la diferencial exterior como un operador lineal $d : \Omega \rightarrow \Omega$ que satisface las siguientes propiedades:

1. Aumenta el grado por 1.
2. Actúa en Ω^0 como la diferencial.
3. Satisface la regla graduada de Leibnitz, esto es, para todo $\omega, \eta \in \Omega$ se cumple

$$d(\omega\eta) = d(\omega)\eta + (-1)^{\partial\omega}\omega d(\eta).$$

4. Su cuadrado es cero:

$$d^2 = 0.$$

En lo que sigue de este capítulo veremos una realización cuántica de este programa.

4.1 Cálculo diferencial sobre espacios cuánticos

Dada una variedad suave M , podemos considerar el álgebra $\mathcal{A} = C^\infty(M)$ de funciones suaves sobre la variedad. La estructura completa del álgebra $\Omega(M)$ de formas diferenciales sobre M , se puede reconstruir conociendo las *uno formas* $\Gamma = \Omega^1(M)$, las cuales forman un bimódulo sobre $\mathcal{A}$ equipado con la diferencial $d : \mathcal{A} \rightarrow \Gamma$.

En el caso cuántico, formulamos de manera análoga la noción de cálculo diferencial. Consideraremos un álgebra compleja $\mathcal{A}$, en general no-conmutativa y la pensaremos como un álgebra de funciones apropiadas sobre una variedad cuántica. Consideraremos un bimódulo sobre esta álgebra equipado con la diferencial $d: \mathcal{A} \rightarrow \Gamma$.

Definición 4.1.1. Sean $\mathcal{A}$ un álgebra unital, Γ un bimódulo unital sobre el álgebra $\mathcal{A}$ y $d: \mathcal{A} \rightarrow \Gamma$ un mapeo lineal. Decimos que (Γ, d) es un cálculo diferencial de primer orden sobre $\mathcal{A}$ si

1. El operador d es una derivación de $\mathcal{A}$ con valores en Γ . Es decir

$$d(ab) = (da)b + adb,$$

para todo $a, b \in \mathcal{A}$.

2. Para cada elemento $\rho \in \Gamma$, existe una cantidad finita de elementos $a_1, b_1, \dots, a_n, b_n \in \mathcal{A}$ tales que

$$\rho = \sum_{k=1}^n a_k db_k.$$

De aquí en adelante, cuando el contexto lo permita usaremos la notación $\sum adb$ para referirnos a las sumas finitas de la forma $\sum_{k=1}^n a_k db_k$ para algunos $a_1, b_1, \dots, a_n, b_n \in \mathcal{A}$ y también escribiremos $\sum a \otimes b$ en lugar de la suma finita $\sum_{k=1}^n a_k \otimes b_k$, para algunos $a_1, b_1, \dots, a_n, b_n \in \mathcal{A}$.

Observemos que en la definición anterior hemos omitido la condición de que las ‘funciones suaves’ conmuten con d , es decir los elementos de $\mathcal{A}$, conmuten con las diferenciales. La primera propiedad es básicamente la regla de Leibnitz y la segunda nos dice que las diferenciales generan algebraicamente a todo el bimódulo Γ .

Observación 4.1.1. Aplicando la regla de Leibnitz al elemento 1 en el álgebra $\mathcal{A}$, tenemos que necesariamente $d(1) = 0$. Por lo tanto, dado que d es lineal, tendremos que

$$d(\lambda) = 0,$$

para todo $\lambda \in \mathbb{C}$.

Definición 4.1.2. Dos cálculos diferenciales (Γ_1, d_1) y (Γ_2, d_2) sobre $\mathcal{A}$ son iguales si existe un isomorfismo $\iota: \Gamma_1 \rightarrow \Gamma_2$ de bimódulos tal que $\iota(d_1 a) = d_2 a$ para todo $a \in \mathcal{A}$.

Presentaremos a continuación la construcción dada en [23] de cálculos diferenciales de primer orden sobre un álgebra unital $\mathcal{A}$. Sea m el mapeo multiplicación de $\mathcal{A}$. Definimos el subespacio vectorial $\mathcal{A}^2$ del álgebra $\mathcal{A} \otimes \mathcal{A}$ como

$$\mathcal{A}^2 = \{q \in \mathcal{A} \otimes \mathcal{A} \mid m(q) = 0\}. \quad (4.1)$$

El subespacio $\mathcal{A}^2$ tiene estructura de $\mathcal{A}$ -bimódulo definiendo para todo $c \in \mathcal{A}$ las acciones izquierda y derecha como:

$$c \left(\sum a \otimes b \right) = \sum ca \otimes b, \quad \left(\sum a \otimes b \right) c = \sum a \otimes bc. \quad (4.2)$$

El operador $D : \mathcal{A} \rightarrow \mathcal{A}^2$ dado por

$$D(a) = 1 \otimes a - a \otimes 1, \quad (4.3)$$

para todo $a \in \mathcal{A}$, es un operador diferencial pues si $a, b \in \mathcal{A}$, entonces

$$\begin{aligned} D(a)b + aD(b) &= (1 \otimes a - a \otimes 1)b + a(1 \otimes b - b \otimes 1) = 1 \otimes ab - a \otimes b + \\ &\quad a \otimes b - ab \otimes 1 = 1 \otimes ab - ab \otimes 1 = D(ab). \end{aligned}$$

Es decir, D satisface la regla de Leibnitz. Por otro lado, si $q = \sum a \otimes b \in \mathcal{A}^2$, entonces

$$\sum aD(b) = \sum a(1 \otimes b - b \otimes 1) = \sum a \otimes b - \sum ab \otimes 1 = q.$$

Por lo tanto podemos concluir que $(\mathcal{A}^2, D)$ es un cálculo diferencial de primer orden sobre $\mathcal{A}$ conocido como el *cálculo universal* sobre $\mathcal{A}$.

El siguiente teorema nos muestra que todo cálculo diferencial sobre un álgebra, se puede obtener como el cálculo cociente del cálculo universal.

Teorema 4.1.1. Sea $\mathcal{N}$ un $\mathcal{A}$ sub-bimódulo de $\mathcal{A}^2$ y $\Gamma = \mathcal{A}^2/\mathcal{N}$ el $\mathcal{A}$ bimódulo cociente. Si $\Pi : \mathcal{A}^2 \rightarrow \Gamma$ es el epimorfismo canónico y $\delta = \Pi \circ D$, entonces (Γ, δ) es un cálculo diferencial de primer orden sobre $\mathcal{A}$. Más aún, cualquier cálculo diferencial de primer orden sobre $\mathcal{A}$ se puede obtener de esta manera.

Demostración. Es obvio que el operador δ satisface la regla de Leibniz pues Π es un morfismo de $\mathcal{A}$ -módulos. Además como Π es epimorfismo, para todo elemento $\gamma \in \Gamma$ existe $q \in \mathcal{A}^2$, $q = \sum a \otimes b$ tal que $\gamma = \Pi(q)$. Por la construcción anterior tenemos que $q = \sum aD(b)$, lo que nos lleva a que

$$\gamma = \Pi(q) = \sum a\delta(b).$$

Por lo tanto, (Γ, δ) es un cálculo diferencial de primer orden sobre $\mathcal{A}$. Por otra parte, si (Γ, d) es otro cálculo diferencial de primer orden sobre $\mathcal{A}$, definimos el mapeo lineal $\tilde{\Pi} : \mathcal{A}^2 \rightarrow \Gamma$, dado por

$$\tilde{\Pi} \left[\sum a \otimes b \right] = \sum ad(b).$$

Entonces $\tilde{\Pi}$ es un morfismo de $\mathcal{A}$ -bimódulos. En efecto, para todo $c \in \mathcal{A}$ tenemos:

$$\tilde{\Pi} \left[c \left(\sum a \otimes b \right) \right] = \tilde{\Pi} \left(\sum ca \otimes b \right) = c \sum ad(b) = c \tilde{\Pi} \left[\sum a \otimes b \right],$$

es decir, $\tilde{\Pi}$ es un morfismo de $\mathcal{A}$ -módulo izquierdo. Por el otro lado, también se tiene que

$$\begin{aligned} \tilde{\Pi} \left[\left(\sum a \otimes b \right) c \right] &= \tilde{\Pi} \left[\sum a \otimes bc \right] = \sum ad(bc) = \left[\sum ad(b) \right] c \\ &= \tilde{\Pi} \left[\sum a \otimes b \right] c, \end{aligned}$$

es decir, $\tilde{\Pi}$ es un morfismo de $\mathcal{A}$ -módulo derecho. Por lo tanto $\tilde{\Pi}$ es un morfismo de $\mathcal{A}$ -bimódulos.

Notemos además que si $\rho = \sum ad(b) \in \Gamma$ entonces el elemento $q = \sum a \otimes b - \sum ab \otimes 1 \in \mathcal{A}^2$ satisface que $\tilde{\Pi}(q) = \rho$, lo que demuestra que $\tilde{\Pi}$ es sobre.

Así, dado que $\tilde{\Pi}$ es morfismo de $\mathcal{A}$ -bimódulo entonces $\mathcal{N} = \ker(\tilde{\Pi})$ es un $\mathcal{A}$ -sub-bimódulo de $\mathcal{A}^2$ y por el primer teorema de isomorfismo se tiene que $\Gamma \cong \tilde{\Gamma} = \mathcal{A}^2 / \mathcal{N}$.

Observemos que $\tilde{\Pi}$ induce un isomorfismo $\bar{\Pi} : \tilde{\Gamma} \rightarrow \Gamma$ de $\mathcal{A}$ -bimódulos y además observemos también que $\tilde{\Gamma}$ es un cálculo diferencial sobre $\mathcal{A}$ con mapeo diferencial dado por $\delta = \Pi \circ D$, donde $\Pi : \mathcal{A}^2 \rightarrow \tilde{\Gamma}$ es la proyección canónica y $D : \mathcal{A} \rightarrow \mathcal{A}^2$ es la diferencial dada por la fórmula (4.3). Veamos que los cálculos $(\tilde{\Gamma}, \delta)$ y (Γ, d) son iguales. En efecto,

$$\bar{\Pi}\delta(b) = \bar{\Pi}(\Pi(D(b))) = \tilde{\Pi}(1 \otimes b - b \otimes 1) = 1d(b) - bd(1) = d(b),$$

lo que termina la demostración. $\square$

En general, en un álgebra podemos definir más de un cálculo diferencial sobre ésta. En el caso de que $\mathcal{N} = \mathcal{A}^2$ obtenemos el cálculo diferencial trivial y si $\mathcal{N} = \{0\}$ entonces tenemos el cálculo inicial universal. Denotaremos por δ a la diferencial del factor cálculo $\mathcal{A}^2 / \mathcal{N}$, es decir, $\delta = \Pi \circ D$.

Definición 4.1.3. Sean $\mathcal{A}$ una $*$ -álgebra y (Γ, d) un cálculo diferencial de primer orden sobre $\mathcal{A}$. Decimos que (Γ, d) es un $*$ -cálculo si se cumple que

$$\sum adb = 0, \quad \text{entonces} \quad \sum d(b^*)a^* = 0.$$

En tal caso, la fórmula

$$(adb)^* = (db^*)a^* \tag{4.4}$$

determina de manera consistente y única un operador antilineal $*$: $\Gamma \rightarrow \Gamma$ que se puede pensar como $*$ -estructura inducida de $\mathcal{A}$ a Γ .

Observemos que si (Γ, d) es un $*$ -cálculo, entonces como consecuencia de la última igualdad en la definición anterior, se tendrá que d y las $*$ conmutan y además podemos ver fácilmente que se satisfacen las siguientes igualdades

$$(a\varrho)^* = \varrho^*a^* \quad (\varrho^*)^* = \varrho \quad (\varrho a)^* = a^*\varrho^*$$

de la definición del operador $*$ en Γ .

Definición 4.1.4. Sea Ψ un bimódulo sobre una $*$ -álgebra $\mathcal{A}$. Decimos que Ψ es un $*$ -bimódulo sobre $\mathcal{A}$, si existe una operación $*$: $\Psi \rightarrow \Psi$, definida como $*(\varrho) = \varrho^*$ y que satisface las siguientes propiedades:

$$(\lambda\varrho + \theta)^* = \bar{\lambda}\varrho^* + \theta^*, \quad (\varrho^*)^* = \varrho, \quad (a\varrho)^* = \varrho^*a^*,$$

para todo $a \in \mathcal{A}$, $\varrho, \theta \in \Psi$ y $\lambda \in \mathbb{C}$.

Observación 4.1.2. Si (Γ, d) es un $*$ -cálculo, entonces Γ es automáticamente un $*$ -bimódulo tal que la diferencial d intercambia la estructura $*$ de $\mathcal{A}$ y Γ . Por la definición de $*$ -cálculo, como ya mencionamos, la operación $*$ en Γ queda únicamente determinada.

Observación 4.1.3. Si $\mathcal{A}$ es una $*$ -álgebra entonces el $\mathcal{A}$ -bimódulo $\mathcal{A}^2$ es un $*$ -bimódulo donde la operación $*$: $\mathcal{A}^2 \rightarrow \mathcal{A}^2$ está dada por

$$(\sum a \otimes b)^* = - \sum b^* \otimes a^*.$$

Observemos además que D y $*$ conmutan:

$$D(a^*) = (1 \otimes a^* - a^* \otimes 1) = (1 \otimes a)^* - (a \otimes 1)^* = (1 \otimes a - a \otimes 1)^* = D(a)^*,$$

para todo $a \in \mathcal{A}$. Es decir, el cálculo universal $(\mathcal{A}^2, D)$ es un $*$ -cálculo con la estructura $*$ dada arriba.

Proposición 4.1.2. Sea $\mathcal{A}$ una $*$ -álgebra y $\mathcal{N}$ un $\mathcal{A}$ -sub-bimódulo de $\mathcal{A}^2$, entonces $(\mathcal{A}^2/\mathcal{N}, \delta)$ es un $*$ -cálculo si y sólo si $\mathcal{N}^* = \mathcal{N}$.

Demostración. Supongamos que $(\mathcal{A}^2/\mathcal{N}, \delta)$ es un $*$ -cálculo, y consideremos $\Pi: \mathcal{A}^2 \rightarrow \mathcal{A}^2/\mathcal{N}$ el $*$ -epimorfismo canónico, entonces para todo $q \in \mathcal{N}$, se satisface que $\Pi(q^*) = \Pi(q)^* = 0$, es decir $q^* \in \mathcal{N}$. Por lo tanto $\mathcal{N}$ es invariante bajo la estructura $*$.

Inversamente, si $\mathcal{N}^* = \mathcal{N}$ y $\sum a\delta(b) = 0$, para $a, b \in \mathcal{A}$, entonces observando que

$$0 = \sum a\delta(b) = \sum a\Pi(D(b)) = \Pi(\sum aD(b)),$$

tenemos que $\sum aD(b) \in \mathcal{N}$, lo cual, usando que $\mathcal{N}$ es invariante bajo $*$, implica que

$$(\sum aD(b))^* = \sum D(b^*)a^* \in \mathcal{N}.$$

Por lo tanto

$$\sum \delta(b^*)a^* = \Pi(\sum D(b^*)a^*) = 0,$$

es decir, $\mathcal{A}^2/\mathcal{N}$ es un $*$ -cálculo. $\square$

4.2 Cálculo diferencial sobre grupos cuánticos

En esta sección veremos cómo son los cálculos diferenciales sobre las álgebras de Hopf. Como hemos mencionado, estas álgebras nos proporcionan estructuras básicas para definir el concepto de grupo cuántico.

Así que la problemática natural con los cálculos diferenciales sobre un álgebra de Hopf (en general no conmutativa) es investigar en qué forma la estructura de grupo cuántico (dada por el coproducto) se relaciona con la del cálculo.

Definiremos primeramente los conceptos de cálculo diferencial covariante. Este concepto se divide en izquierda, derecha y bicovariante. Estos nos dan versiones cuánticas de la idea de extensibilidad de las acciones izquierda y derecha del grupo al nivel del cálculo. Obtendremos también una variedad de fórmulas explícitas para conocer el cálculo diferencial con el que comenzamos.

Definición 4.2.1. Sean $\mathcal{A}$ un álgebra de Hopf y (Γ, d) un cálculo diferencial de primer orden sobre $\mathcal{A}$. Decimos que el cálculo (Γ, d) es izquierda covariante si se cumple que cada vez que $\sum ad(b) = 0$ en Γ , entonces

$$\sum \phi(a)(\text{id} \otimes d)\phi(b) = \sum a^{(1)}b^{(1)} \otimes a^{(2)}d(b^{(2)}) = 0.$$

Similarmente el cálculo (Γ, d) es derecha covariante si cada que $\sum ad(b) = 0$ en Γ , entonces

$$\sum \phi(a)(d \otimes \text{id})\phi(b) = \sum a^{(1)}d(b^{(1)}) \otimes a^{(2)}b^{(2)} = 0.$$

Finalmente, el cálculo (Γ, d) es bicovariante si es izquierda y derecha covariante.

Al igual que la multiplicación del grupo tiene su mapeo dual ϕ , podemos pensar en la versión dual de una acción.

Definición 4.2.2. Sea $\mathcal{A}$ un álgebra de Hopf y Ψ un bimódulo sobre $\mathcal{A}$. Decimos que Ψ es un bimódulo izquierda covariante si está equipado con un mapeo lineal $\ell : \Psi \rightarrow \mathcal{A} \otimes \Psi$ que es una co-acción izquierda de $\mathcal{A}$ en Ψ , es decir los diagramas

$$\begin{array}{ccc} \Psi & \xrightarrow{\ell} & \mathcal{A} \otimes \Psi \\ \ell \downarrow & & \downarrow \phi \otimes \text{id} \\ \mathcal{A} \otimes \Psi & \xrightarrow{\text{id} \otimes \ell} & \mathcal{A} \otimes \mathcal{A} \otimes \Psi \end{array} \quad \begin{array}{ccc} \Psi & \xrightarrow{\ell} & \mathcal{A} \otimes \Psi \\ \text{id} \downarrow & & \downarrow \epsilon \otimes \text{id} \\ \Psi & \xrightarrow{\cong} & \mathbb{C} \otimes \Psi \end{array}$$

son conmutativos, y además se cumple

$$\ell(a\rho) = \phi(a)\ell(\rho), \quad \ell(\rho a) = \ell(\rho)\phi(a),$$

para cada $a \in \mathcal{A}$ y $\rho \in \Psi$.

Notemos que si usamos súper índices negativos para indicar los elementos del álgebra y súper índice cero para los elementos del bimódulo, obtendremos que si $\ell(\rho) = \rho^{(-1)} \otimes \rho^{(0)}$, entonces el diagrama del lado izquierdo de la definición anterior nos dice que

$$\rho^{(-2)} \otimes (\rho^{(-1)} \otimes \rho^{(0)}) = (\rho^{(-2)} \otimes \rho^{(-1)}) \otimes \rho^{(0)},$$

Por otro lado, el diagrama del lado derecho de la definición anterior nos dice que

$$\epsilon(\rho^{(-1)})\rho^{(0)} = \rho. \quad (4.5)$$

Lema 4.2.1. Sea $\mathcal{A}$ un álgebra de Hopf y $\mathcal{A}^2 = \ker(m)$ el bimódulo definido anteriormente. Entonces $\mathcal{A}^2$ es un bimódulo izquierda covariante con co-acción izquierda definida como

$$\ell\left(\sum a \otimes b\right) = a^{(1)}b^{(1)} \otimes a^{(2)} \otimes b^{(2)}, \quad (4.6)$$

para todo $a, b \in \mathcal{A}$. □

De manera análoga, podemos definir $\mathcal{A}$ -bimódulo derecha covariante:

Definición 4.2.3. Sea $\mathcal{A}$ un álgebra de Hopf y Ψ un bimódulo sobre $\mathcal{A}$. Decimos que Ψ es un bimódulo derecha covariante si está equipado con un mapeo lineal $\wp : \Psi \rightarrow \Psi \otimes \mathcal{A}$ que es una co-acción derecha de $\mathcal{A}$ en Ψ , es decir los diagramas

$$\begin{array}{ccc} \Psi & \xrightarrow{\wp} & \Psi \otimes \mathcal{A} \\ \wp \downarrow & & \downarrow \text{id} \otimes \phi \\ \Psi \otimes \mathcal{A} & \xrightarrow{\wp \otimes \text{id}} & \Psi \otimes \mathcal{A} \otimes \mathcal{A} \end{array} \qquad \begin{array}{ccc} \Psi & \xrightarrow{\wp} & \Psi \otimes \mathcal{A} \\ \text{id} \downarrow & & \downarrow \text{id} \otimes \epsilon \\ \Psi & \xrightarrow{\cong} & \Psi \otimes \mathbb{C} \end{array}$$

son conmutativos, y además se cumple:

$$\wp(a\varrho) = \phi(a)\wp(\varrho), \qquad \wp(\varrho a) = \wp(\varrho)\phi(a),$$

para todo $\varrho \in \Psi$ y $a \in \mathcal{A}$.

Si ahora usamos el súper índice cero para indicar los elementos del bimódulo, y súper índices positivos para indicar a los elementos del álgebra, se tiene que para $\wp(\varrho) = \varrho^{(0)} \otimes \varrho^{(1)}$, entonces el diagrama del lado izquierdo de la definición anterior nos dice que

$$(\varrho^{(0)} \otimes \varrho^{(1)}) \otimes \varrho^{(2)} = \varrho^{(0)} \otimes (\varrho^{(1)} \otimes \varrho^{(2)}). \quad (4.7)$$

Por otro lado, el diagrama del lado derecho de la definición anterior es equivalente a decir

$$\varrho^{(0)}\epsilon(\varrho^{(1)}) = \varrho. \quad (4.8)$$

Lema 4.2.2. Sea $\mathcal{A}$ un álgebra de Hopf y $\mathcal{A}^2 = \ker(m)$ definido en (4.1). Entonces $\mathcal{A}^2$ es un bimódulo derecha covariante con co-acción derecha dada por la fórmula:

$$\wp\left(\sum a \otimes b\right) = \sum a^{(1)} \otimes b^{(1)} \otimes a^{(2)}b^{(2)}, \quad (4.9)$$

para todo $a, b \in \mathcal{A}$. □

Definición 4.2.4. Sea $\mathcal{A}$ un álgebra de Hopf y Ψ un bimódulo sobre $\mathcal{A}$. Decimos que Ψ es un bimódulo bicovariante si está equipado con una co-acción izquierda ℓ y una co-acción derecha $\wp$ con las que se convierte en bimódulo izquierda y derecha covariante respectivamente y además el siguiente diagrama

$$\begin{array}{ccc} \Psi & \xrightarrow{\wp} & \Psi \otimes \mathcal{A} \\ \ell \downarrow & & \downarrow \ell \otimes \text{id} \\ \mathcal{A} \otimes \Psi & \xrightarrow{\text{id} \otimes \wp} & \mathcal{A} \otimes \Psi \otimes \mathcal{A} \end{array} \quad (4.10)$$

es conmutativo.

Lema 4.2.3. Sea $\mathcal{A}$ un álgebra de Hopf y $\mathcal{A}^2 = \ker(m)$ el bimódulo definido en (4.1). Entonces $\mathcal{A}^2$ es un bimódulo bicovariante con co-acciones izquierda y derecha dadas por las fórmulas (4.6) y (4.9), respectivamente. $\square$

Definición 4.2.5. Sean $\mathcal{A}$ un álgebra de Hopf, Ψ un bimódulo izquierda covariante sobre $\mathcal{A}$ y $\ell: \Psi \rightarrow \mathcal{A} \otimes \Psi$ la co-acción izquierda. Decimos que un elemento $\varrho \in \Psi$ es izquierda invariante si

$$\ell(\varrho) = 1 \otimes \varrho.$$

Denotamos por Ψ_{inv} al espacio vectorial de todos los elementos izquierda invariantes de Ψ .

De aquí en adelante escribiremos (Ψ, ℓ) para referirnos al bimódulo Ψ izquierda covariante con su co-acción izquierda ℓ . De la misma manera, escribiremos (Ψ, φ) para referirnos al bimódulo Ψ derecha covariante con su co-acción derecha φ y escribiremos (Ψ, ℓ, φ) para referirnos al bimódulo bicovariante Ψ con sus co-acciones izquierda y derecha ℓ y φ , respectivamente. También escribiremos simplemente Ψ , si el contexto lo permite.

Cada bimódulo Ψ se corresponde geoméricamente con un haz vectorial sobre el grupo, equipado con una acción del grupo, es decir, con un haz covariante. En este caso, Ψ_{inv} será el espacio de todas las secciones del haz que son invariantes bajo la acción izquierda del grupo.

Observación 4.2.1. Obsérvese que estamos usando bimódulos y no módulos izquierdos o derechos como se usa en K -teoría algebraica o en la formulación de Alain Connes. Los bimódulos como análogos de haces vectoriales, también aparecen naturalmente en la teoría de haces principales cuánticos [8] como sus haces vectoriales cuánticos asociados.

De esta manera, cuando $\mathcal{A}$ es conmutativa, le corresponde un grupo clásico, y la formulación cuántica incluye como caso particular toda esta teoría clásica.

Proposición 4.2.4. Sea $\mathcal{A}$ un álgebra de Hopf y Ψ un bimódulo izquierda covariante sobre $\mathcal{A}$. Entonces existe una única proyección $p: \Psi \rightarrow \Psi_{\text{inv}}$ tal que

$$p(b\varrho) = \epsilon(b)p(\varrho), \quad (4.11)$$

para todo $b \in \mathcal{A}$ y $\varrho \in \Psi$. Más aún, se cumple que si $\ell(\varrho) = \varrho^{(-1)} \otimes \varrho^{(0)}$ entonces

$$\varrho = \varrho^{(-1)} p(\varrho^{(0)}), \quad (4.12)$$

para todo $\varrho \in \Psi$.

Demostración. Definimos el mapeo lineal $p : \Psi \longrightarrow \Psi_{\text{inv}}$ como

$$p(\varrho) = \kappa(\varrho^{(-1)})\varrho^{(0)}.$$

Primeramente veamos que la imagen de p está en los elementos izquierda invariantes. En efecto, para todo $\varrho \in \Psi$ se tiene que

$$\begin{aligned} \ell(p(\varrho)) &= \ell(\kappa(\varrho^{(-1)})\varrho^{(0)}) = \phi(\kappa(\varrho^{(-1)}))\ell(\varrho^{(0)}) \\ &= (\kappa(\varrho^{(-2)}) \otimes \kappa(\varrho^{(-3)})) (\varrho^{(-1)} \otimes \varrho^{(0)}) \\ &= \kappa(\varrho^{(-2)})\varrho^{(-1)} \otimes \kappa(\varrho^{(-3)})\varrho^{(0)} \\ &= \epsilon(\varrho^{(-1)}) \otimes \kappa(\varrho^{(-2)})\varrho^{(0)} \\ &= 1 \otimes \kappa(\varrho^{(-1)})\varrho^{(0)} = 1 \otimes p(\varrho), \end{aligned}$$

es decir, $p(\varrho) \in \Psi_{\text{inv}}$. Por otro lado, observemos que $\ell(b\varrho) = b^{(1)}\varrho^{(-1)} \otimes b^{(2)}\varrho^{(0)}$ implica que:

$$p(b\varrho) = \kappa(b^{(1)}\varrho^{(-1)})b^{(2)}\varrho^{(0)} = \kappa(\varrho^{(-1)})\kappa(b^{(1)})b^{(2)}\varrho^{(0)} = \epsilon(b)p(\varrho),$$

y por lo tanto p satisface (4.11). Además,

$$\varrho^{(-1)}p(\varrho^{(0)}) = \varrho^{(-2)}\kappa(\varrho^{(-1)})\varrho^{(0)} = \epsilon(\varrho^{(-1)})\varrho^{(0)} = \varrho,$$

lo que muestra que p también satisface (4.12).

Finalmente debido a que $\ell(p(\varrho)) = 1 \otimes p(\varrho)$, se tiene entonces que $p^2(\varrho) = \kappa(1)p(\varrho) = p(\varrho)$ y si $\varrho \in \Psi_{\text{inv}}$ entonces $p(\varrho) = \kappa(1)\varrho = \varrho$. Lo que demuestra que p es una proyección.

Para demostrar la unicidad, sea $q : \Psi \rightarrow \Psi_{\text{inv}}$ otra proyección con las propiedades de la proposición. Entonces tenemos que

$$\varrho = \varrho^{(-1)}p(\varrho^{(0)}) = \varrho^{(-1)}q(\varrho^{(0)}),$$

y de esta forma se tiene que

$$q(\varrho) = q(\varrho^{(-1)}q(\varrho^{(0)})) = \epsilon(\varrho^{(-1)})q(\varrho^{(0)}) = \epsilon(\varrho^{(-1)})p(\varrho^{(0)}) = p(\epsilon(\varrho^{(-1)})\varrho^{(0)}) = p(\varrho),$$

lo que demuestra la unicidad. $\square$

Proposición 4.2.5. Sea $\mathcal{A}$ un álgebra de Hopf y Ψ un $\mathcal{A}$ -bimódulo izquierda covariante. El espacio Ψ_{inv} posee una estructura natural de $\mathcal{A}$ -módulo derecho definiendo la acción derecha $\circ : \Psi_{\text{inv}} \times \mathcal{A} \rightarrow \Psi_{\text{inv}}$ como

$$\varrho \circ b = \kappa(b^{(1)})\varrho b^{(2)}. \quad (4.13)$$

Demostración. Observemos que el elemento $\kappa(b^{(1)})\varrho b^{(2)}$ efectivamente es izquierda invariante:

$$\begin{aligned}\ell(\kappa(b^{(1)})\varrho b^{(2)}) &= \phi(\kappa(b^{(1)}))\ell(\varrho)\phi(b^{(2)}) = \\ &= (\kappa(b^{(2)}) \otimes \kappa(b^{(1)}))(1 \otimes \varrho)(b^{(3)} \otimes b^{(4)}) \\ &= \kappa(b^{(2)})b^{(3)} \otimes \kappa(b^{(1)})\varrho b^{(4)} \\ &= \epsilon(b^{(2)}) \otimes \kappa(b^{(1)})\varrho b^{(3)} \\ &= 1 \otimes \kappa(b^{(1)})\varrho b^{(2)}.\end{aligned}$$

Y veamos ahora que $\circ$ define una acción derecha. Por la definición se tiene que $\varrho \circ 1 = \kappa(1)\varrho 1 = \varrho$, para todo $\varrho \in \Psi_{\text{inv}}$. Además,

$$\begin{aligned}(\varrho \circ a) \circ b &= (\kappa(a^{(1)})\varrho a^{(2)}) \circ b = \kappa(b^{(1)}) (\kappa(a^{(1)})\varrho a^{(2)}) b^{(2)} \\ &= \kappa(a^{(1)}b^{(1)})\varrho a^{(2)}b^{(2)} = \varrho \circ (ab),\end{aligned}$$

para todo $a, b \in \mathcal{A}$ y $\varrho \in \Psi_{\text{inv}}$. Lo que muestra que efectivamente $\circ$ es una acción derecha de $\mathcal{A}$ en Ψ_{inv} . $\square$

Si además es el caso que Ψ es un bimódulo bicovariante entonces se satisface que

$$\wp(\Psi_{\text{inv}}) \subseteq \Psi_{\text{inv}} \otimes \mathcal{A}.$$

Es decir, el espacio Ψ_{inv} además de ser un módulo derecho sobre $\mathcal{A}$ también está equipado con una estructura de comódulo derecho.

Proposición 4.2.6. Las estructuras de módulo y de comódulo en Ψ_{inv} satisfacen la siguiente identidad de compatibilidad mutua:

$$\wp(\theta \circ a) = (\theta^{(0)} \circ a^{(2)}) \otimes \kappa(a^{(1)})\theta^{(1)}a^{(3)}. \quad (4.14)$$

Demostración. El cálculo directo nos lleva a

$$\begin{aligned}\wp(\theta \circ a) &= \wp(\kappa(a^{(1)})\theta a^{(2)}) = \phi(\kappa(a^{(1)}))\wp(\theta)\phi(a^{(2)}) \\ &= (\kappa(a^{(2)}) \otimes \kappa(a^{(1)})) (\theta^{(0)} \otimes \theta^{(1)}) (a^{(3)} \otimes a^{(4)}) \\ &= \kappa(a^{(2)})\theta^{(0)}a^{(3)} \otimes \kappa(a^{(1)})\theta^{(1)}a^{(4)} \\ &= (\theta^{(0)} \circ a^{(2)}) \otimes \kappa(a^{(1)})\theta^{(1)}a^{(3)},\end{aligned}$$

por lo tanto se cumple la identidad de compatibilidad. $\square$

Proposición 4.2.7. Sea $\mathcal{A}$ un álgebra de Hopf y Ψ un bimódulo sobre $\mathcal{A}$ izquierda covariante. Entonces $\mathcal{A} \otimes \Psi_{\text{inv}}$ posee una estructura de $\mathcal{A}$ -bimódulo naturalmente isomorfo a Ψ .

Demostración. Para todo $a, b \in \mathcal{A}$ y $\varrho \in \Psi_{\text{inv}}$ definimos la acción izquierda de $\mathcal{A}$ en $\mathcal{A} \otimes \Psi_{\text{inv}}$ de la manera natural, es decir

$$a(b \otimes \varrho) = ab \otimes \varrho.$$

La acción derecha de $\mathcal{A}$ en $\mathcal{A} \otimes \Psi_{\text{inv}}$ la definimos como:

$$(a \otimes \varrho)b = ab^{(1)} \otimes (\varrho \circ b^{(2)}), \quad (4.15)$$

para todo $\varrho \in \Psi_{\text{inv}}$ y $a, b \in \mathcal{A}$. Con las acciones así definidas $\mathcal{A} \otimes \Psi_{\text{inv}}$ se convierte en bimódulo sobre $\mathcal{A}$.

Además, $\tilde{\ell} : \mathcal{A} \otimes \Psi_{\text{inv}} \rightarrow \mathcal{A} \otimes \mathcal{A} \otimes \Psi_{\text{inv}}$ definida como

$$\tilde{\ell}(b \otimes \varrho) = \phi(b) \otimes \varrho = b^{(1)} \otimes b^{(2)} \otimes \varrho, \quad (4.16)$$

define una co-acción izquierda de $\mathcal{A}$ en $\mathcal{A} \otimes \Psi_{\text{inv}}$. En efecto,

$$\begin{aligned} (\phi \otimes \text{id})\tilde{\ell}(b \otimes \varrho) &= (\phi \otimes \text{id})(\phi(b) \otimes \varrho) = \phi(b^{(1)}) \otimes b^{(2)} \otimes \varrho \\ &= (\phi \otimes \text{id})\phi(b) \otimes \varrho = (\text{id} \otimes \phi)\phi(b) \otimes \varrho \\ &= b^{(1)} \otimes \phi(b^{(2)}) \otimes \varrho = b^{(1)} \otimes \tilde{\ell}(b^{(2)} \otimes \varrho) \\ &= (\text{id} \otimes \tilde{\ell})(\phi(b) \otimes \varrho) = (\text{id} \otimes \tilde{\ell})\tilde{\ell}(b \otimes \varrho), \end{aligned}$$

para todo $b \in \mathcal{A}$ y $\varrho \in \Psi_{\text{inv}}$. Nótese que el mapeo id es el mapeo identidad en $\mathcal{A} \otimes \Psi_{\text{inv}}$. Por otra parte también se tiene que

$$(\epsilon \otimes \text{id})\tilde{\ell}(b \otimes \varrho) = (\epsilon \otimes \text{id})(\phi(b) \otimes \varrho) = \epsilon(b^{(1)}) \otimes b^{(2)} \otimes \varrho = b \otimes \varrho,$$

para todo $b \in \mathcal{A}$ y $\varrho \in \Psi_{\text{inv}}$. Por lo tanto, $\tilde{\ell}$ define una co-acción izquierda.

Además por la definición se tiene que

$$\tilde{\ell}(a(b \otimes \varrho)) = \tilde{\ell}(ab \otimes \varrho) = \phi(ab) \otimes \varrho = \phi(a)\phi(b) \otimes \varrho = \phi(a)\tilde{\ell}(b \otimes \varrho).$$

Y por otra parte también tenemos que

$$\begin{aligned} \tilde{\ell}((a \otimes \rho)b) &= \tilde{\ell}(ab^{(1)} \otimes [\varrho \circ b^{(2)}]) = \phi(ab^{(1)}) \otimes [\varrho \circ b^{(2)}] \\ &= a^{(1)}b^{(1)} \otimes a^{(2)}b^{(2)} \otimes [\varrho \circ b^{(3)}] = a^{(1)}b^{(1)} \otimes (a^{(2)} \otimes \varrho)b^{(2)} \\ &= [a^{(1)} \otimes (a^{(2)} \otimes \varrho)]\phi(b) = \tilde{\ell}(a \otimes \varrho)\phi(b). \end{aligned}$$

Lo que demuestra que el bimódulo $\mathcal{A} \otimes \Psi_{\text{inv}}$ es izquierda covariante.

Finalmente, si denotamos por ℓ a la co-acción izquierda de $\mathcal{A}$ en Ψ , entonces el mapeo $\iota : \Psi \rightarrow \mathcal{A} \otimes \Psi_{\text{inv}}$ definido como

$$\iota = (\text{id} \otimes p)\ell,$$

es un isomorfismo de bimódulos. Aquí id es el mapeo identidad definido en $\mathcal{A}$ y p es la proyección sobre Ψ_{inv} dada en la proposición 4.2.4.

Más explícitamente tenemos

$$\iota(\theta) = \theta^{(-1)} \otimes p(\theta^{(0)}),$$

donde $\ell(\theta) = \theta^{(-1)} \otimes \theta^{(0)}$. Veamos primeramente que ι es un morfismo de bimódulos. Para todo $a \in \mathcal{A}$ y $\theta \in \Psi$ se tiene que

$$\ell(a\theta) = \phi(a)\ell(\theta) = a^{(1)}\theta^{(-1)} \otimes a^{(2)}\theta^{(0)},$$

pues (Ψ, ℓ) es un bimódulo izquierda covariante. Por lo tanto,

$$\begin{aligned} \iota(a\theta) &= a^{(1)}\theta^{(-1)} \otimes p(a^{(2)}\theta^{(0)}) = a^{(1)}\theta^{(-1)} \otimes \epsilon(a^{(2)})p(\theta^{(0)}) \\ &= a\theta^{(-1)} \otimes p(\theta^{(0)}) = a\iota(\theta), \end{aligned}$$

es decir, ι es un morfismo de módulo izquierdo. Por otro lado

$$\begin{aligned} \iota(\theta a) &= \theta^{(-1)}a^{(1)} \otimes p(\theta^{(0)}a^{(2)}) = \theta^{(-2)}a^{(1)} \otimes \kappa(\theta^{(-1)}a^{(2)})\theta^{(0)}a^{(3)} \\ &= \theta^{(-2)}a^{(1)} \otimes \kappa(a^{(2)})\kappa(\theta^{(-1)})\theta^{(0)}a^{(3)} = \theta^{(-2)}a^{(1)} \otimes (\kappa(\theta^{(-1)})\theta^{(0)} \circ a^{(2)}) \\ &= (\theta^{(-2)} \otimes \kappa(\theta^{(-1)})\theta^{(0)})a = (\theta^{(-1)} \otimes p(\theta^{(0)}))a = \iota(\theta)a, \end{aligned}$$

donde hemos usado la derecha ϕ -linealidad de ℓ y también que el mapeo proyección se define como $p(\theta^{(0)}) = \kappa(\theta^{(-1)})\theta^{(0)}$. Así que tenemos que ι es un morfismo de $\mathcal{A}$ -bimódulos.

Finalmente, ι es biyectivo pues el mapeo multiplicación $m : \mathcal{A} \otimes \Psi_{\text{inv}} \rightarrow \Psi$ es su mapeo inverso. Por lo tanto los bimódulos izquierda covariantes Ψ y $\mathcal{A} \otimes \Psi_{\text{inv}}$ son isomorfos. $\square$

Lema 4.2.8. Sea $\mathcal{A}$ un álgebra de Hopf y $(\mathcal{A}^2, D)$ el cálculo universal con co-acciones definidas por las ecuaciones (4.6) y (4.9). Entonces, $(\mathcal{A}^2, D)$ es un cálculo bicovariante. $\square$

Proposición 4.2.9. Sea $\mathcal{A}$ un álgebra de Hopf y (Γ, d) un cálculo diferencial de primer orden sobre $\mathcal{A}$ izquierda covariante. Entonces Γ es un bimódulo izquierda covariante equipado con una única co-acción izquierda $\ell : \Gamma \rightarrow \mathcal{A} \otimes \Gamma$, tal que

$$\ell(d(a)) = a^{(1)} \otimes d(a^{(2)}),$$

para todo $a \in \mathcal{A}$.

Demostración. Sea $\ell : \Gamma \rightarrow \mathcal{A} \otimes \Gamma$ definida como

$$\ell\left(\sum ad(b)\right) = \sum a^{(1)}b^{(1)} \otimes a^{(2)}d(b^{(2)}). \quad (4.17)$$

Entonces ℓ es una co-acción izquierda que satisface la condición de la proposición. En efecto, primeramente observemos que ℓ está bien definida pues el cálculo (Γ, d) es izquierda covariante.

El mapeo lineal ℓ satisface lo siguiente:

$$\begin{aligned} (\phi \otimes \text{id})\ell\left(\sum ad(b)\right) &= \sum \phi(a^{(1)}b^{(1)}) \otimes a^{(2)}d(b^{(2)}) \\ &= \sum a^{(1)}b^{(1)} \otimes a^{(2)}b^{(2)} \otimes a^{(3)}d(b^{(3)}) \\ &= \sum a^{(1)}b^{(1)} \otimes \ell(a^{(2)}d(b^{(2)})) \\ &= (\text{id} \otimes \ell)\left(\sum a^{(1)}b^{(1)} \otimes a^{(2)}d(b^{(2)})\right) \\ &= (\text{id} \otimes \ell)\ell\left(\sum ad(b)\right). \end{aligned}$$

Así mismo, también se cumple que

$$\begin{aligned} (\epsilon \otimes \text{id})\ell\left(\sum ad(b)\right) &= \sum \epsilon(a^{(1)}b^{(1)})a^{(2)}d(b^{(2)}) \\ &= \sum \epsilon(a^{(1)})\epsilon(b^{(1)})a^{(2)}d(b^{(2)}) \\ &= \sum (\epsilon(a^{(1)})a^{(2)})d(\epsilon(b^{(1)})b^{(2)}) \\ &= \sum ad(b). \end{aligned}$$

Es decir ℓ es una co-acción izquierda. Además,

$$\begin{aligned} \ell\left(\sum ad(b)\right) &= \sum a^{(1)}b^{(1)} \otimes a^{(2)}d(b^{(2)}) = \sum (a^{(1)} \otimes a^{(2)})(b^{(1)} \otimes d(b^{(2)})) \\ &= \sum \phi(a)\ell(d(b)), \end{aligned}$$

y también

$$\begin{aligned} \ell\left(\sum d(b)a\right) &= \sum \ell(d(ba)) - \ell(bd(a)) \\ &= \sum b^{(1)}a^{(1)} \otimes d(b^{(2)}a^{(2)}) - b^{(1)}a^{(1)} \otimes b^{(2)}d(a^{(2)}) \\ &= \sum b^{(1)}a^{(1)} \otimes (d(b^{(2)}a^{(2)}) - b^{(2)}d(a^{(2)})) \\ &= \sum b^{(1)}a^{(1)} \otimes d(b^{(2)})a^{(2)} = \sum \ell(d(b))\phi(a). \end{aligned}$$

Por lo tanto Γ es un bimódulo izquierda covariante tal que

$$\ell(d(a)) = a^{(1)} \otimes d(a^{(2)}).$$

Para demostrar la unicidad, sea $\tilde{\ell}$ otra co-acción izquierda tal que se cumple la misma propiedad de covarianza

$$\tilde{\ell}(d(a)) = a^{(1)} \otimes d(a^{(2)}).$$

Entonces como todo elemento de Γ se escribe como suma finita de elementos de la forma $ad(b)$ se tiene que

$$\ell\left(\sum ad(b)\right) = \sum \phi(a)\ell(d(b)) = \sum \phi(a)\tilde{\ell}(d(b)) = \tilde{\ell}\left(\sum ad(b)\right),$$

es decir, $\ell = \tilde{\ell}$. Lo que termina la demostración. $\square$

Observación 4.2.2. Como consecuencia de esta proposición tenemos que la co-acción izquierda de $\mathcal{A}$ en el bimódulo $\mathcal{A}^2$ dada por la fórmula (4.6), es la única co-acción izquierda tal que

$$\ell(D(b)) = b^{(1)} \otimes D(b^{(2)}),$$

para cada $b \in \mathcal{A}$.

De forma análoga, si el cálculo diferencial (Γ, d) es derecha covariante, entonces Γ es un bimódulo derecha covariante tal que la co-acción derecha, $\wp$ de $\mathcal{A}$ en Γ satisface que

$$\wp(d(a)) = d(a^{(1)}) \otimes a^{(2)},$$

para todo $a \in \mathcal{A}$. La co-acción derecha es única y está definida por la fórmula

$$\wp\left(\sum ad(b)\right) = \sum a^{(1)}d(b^{(1)}) \otimes a^{(2)}b^{(2)}, \quad (4.18)$$

para todo $a, b \in \mathcal{A}$.

Como consecuencia, también tendremos que la fórmula (4.9) para la co-acción derecha de $\mathcal{A}$ en el bimódulo $\mathcal{A}^2$, define una única co-acción derecha tal que se satisface la correspondiente propiedad de covarianza.

Corolario 4.2.10. Sea $\mathcal{A}$ un álgebra de Hopf y (Γ, d) un cálculo diferencial de primer orden sobre $\mathcal{A}$ bicovariante. Sean ℓ y $\wp$ las co-acciones izquierda y derecha, definidas en (4.17) y (4.18) respectivamente. Entonces

$$(\ell \otimes \text{id})\wp = (\text{id} \otimes \wp)\ell.$$

Es decir, Γ es un bimódulo bicovariante. $\square$

Definición 4.2.6. Sean $\mathcal{A}$ un álgebra de Hopf y (Γ, d) un cálculo diferencial de primer orden sobre $\mathcal{A}$. Definimos el mapeo de gérmenes $\pi : \mathcal{A} \rightarrow \Gamma$ como $\pi(a) = \kappa(a^{(1)})da^{(2)}$.

Observación 4.2.3. Observemos que el mapeo de gérmenes π se puede definir para cualquier cálculo diferencial de primer orden. En el caso particular cuando el cálculo (Γ, d) es izquierda covariante, tenemos que el mapeo π es la composición de la diferencial d y la proyección p dada en la Proposición 4.2.4. En particular se tiene que los gérmenes $\pi(a)$ están en Γ_{inv} .

Proposición 4.2.11. Sea $\mathcal{A}$ un álgebra de Hopf y (Γ, d) un cálculo diferencial sobre $\mathcal{A}$ izquierda covariante. Entonces el mapeo de gérmenes π es sobre Γ_{inv} .

Demostración. Sean $\varrho \in \Gamma_{\text{inv}}$ y ℓ la co-acción izquierda de $\mathcal{A}$ en Γ . Como Γ es cálculo diferencial de primer orden sobre $\mathcal{A}$, entonces todo elemento en Γ se escribe como combinación lineal finita de elementos de la forma $ad(b)$, con $a, b \in \mathcal{A}$. Observemos que

$$\sum \ell(ad(b)) = \sum a^{(1)}b^{(1)} \otimes a^{(2)}d(b^{(2)}) = 1 \otimes \varrho.$$

Aplicando $m(\kappa \otimes \text{id})$ en ambos lados de la última ecuación tendremos que

$$\begin{aligned} \varrho &= \sum \kappa(a^{(1)}b^{(1)})a^{(2)}d(b^{(2)}) = \sum \kappa(b^{(1)})\kappa(a^{(1)})a^{(2)}d(b^{(2)}) \\ &= \sum \kappa(b^{(1)})\epsilon(a)d(b^{(2)}) = \sum \epsilon(a)\pi(b), \end{aligned}$$

lo que nos muestra que π es un mapeo sobre. $\square$

Observación 4.2.4. Dado un cálculo diferencial (Γ, d) de primer orden sobre un álgebra de Hopf $\mathcal{A}$, podemos en general definir el espacio Γ_{inv} como la imagen del mapeo de gérmenes π y la estructura de $\circ$ en Γ_{inv} está bien definida.

Lema 4.2.12. Sean $\mathcal{A}$ un álgebra de Hopf, (Γ, d) un cálculo diferencial de primer orden sobre $\mathcal{A}$ y π el mapeo de gérmenes. La estructura de $\mathcal{A}$ -módulo derecho en Γ_{inv} se relaciona con π de una forma particularmente simple:

$$\pi(ab) = \pi(a) \circ b + \epsilon(a)\pi(b).$$

Esto lo podemos interpretar como la regla de Leibniz localizada en el ‘elemento neutro’ del grupo.

Demostración. Por la definición del mapeo de gérmenes se tiene que

$$\begin{aligned} \pi(ab) &= \kappa(a^{(1)}b^{(1)})d(a^{(2)}b^{(2)}) = \kappa(b^{(1)})\kappa(a^{(1)}) (d(a^{(2)})b^{(2)} + a^{(2)}d(b^{(2)})) \\ &= \pi(a) \circ b + \epsilon(a)\pi(b), \end{aligned}$$

para todo $a, b \in \mathcal{A}$. $\square$

Observación 4.2.5. Sea $\mathcal{R} = \ker(\epsilon) \cap \ker(\pi)$, entonces como consecuencia directa del lema anterior tenemos que para todo $a \in \mathcal{R}$ y $b \in \ker(\epsilon)$ se cumple que

$$\pi(ab) = \pi(a) \circ b = 0 \circ b = 0, \quad \epsilon(ab) = \epsilon(a)\epsilon(b) = \epsilon(a)0 = 0,$$

es decir, $ab \in \mathcal{R}$ y por lo tanto $\mathcal{R}$ es un ideal derecho de $\ker(\epsilon)$. Además como $\mathcal{A} = \ker(\epsilon) \oplus \mathbb{C}$ y $\pi(\mathbb{C}) = 0$, entonces

$$\Gamma_{\text{inv}} = \text{im}(\pi) = \ker(\epsilon)/\mathcal{R}. \quad (4.19)$$

Dada un álgebra de Hopf $\mathcal{A}$ y un cálculo diferencial (Γ, d) izquierda covariante sobre $\mathcal{A}$, ¿podemos definir en el bimódulo $\mathcal{A} \otimes \Gamma_{\text{inv}}$ una estructura diferenciable sobre $\mathcal{A}$ que defina el mismo cálculo diferencial sobre $\mathcal{A}$ dado en Γ ? En lo que sigue veremos que efectivamente, podemos definir una tal estructura. El producto tensorial aquí, se debe entender como $\mathcal{A}$ -módulo izquierdo libre.

Proposición 4.2.13. Sea (Γ, d) un cálculo diferencial de primer orden sobre $\mathcal{A}$. Entonces $\mathcal{A} \otimes \Gamma_{\text{inv}}$ tiene una estructura diferencial sobre $\mathcal{A}$ definida como

$$d(a) = a^{(1)} \otimes \pi(a^{(2)}).$$

Más aún, $(\mathcal{A} \otimes \Gamma_{\text{inv}}, d)$ es un cálculo diferencial de primer orden sobre $\mathcal{A}$ izquierda covariante.

Demostración. En la proposición (4.2.7) ya hemos visto que $\mathcal{A} \otimes \Gamma_{\text{inv}}$ es un bimódulo sobre $\mathcal{A}$, con acción izquierda dada por

$$a(b \otimes \varrho) = ab \otimes \varrho,$$

y acción derecha

$$(b \otimes \varrho)a = ba^{(1)} \otimes (\varrho \circ a^{(2)}).$$

Es claro que d es lineal pues es la composición de los mapeos lineales $(1 \otimes \pi)$ y ϕ . Además

$$\begin{aligned} d(ab) &= a^{(1)}b^{(1)} \otimes \pi(a^{(2)}b^{(2)}) \\ &= a^{(1)}b^{(1)} \otimes (\pi(a^{(2)}) \circ b^{(2)}) + a^{(1)}b^{(1)} \otimes \epsilon(a^{(2)})\pi(b^{(2)}) \\ &= (a^{(1)} \otimes \pi(a^{(2)}))b + a(b^{(1)} \otimes \pi(b^{(2)})) \\ &= d(a)b + ad(b), \end{aligned}$$

por lo tanto se cumple la regla de Leibnitz. Además, como π es un mapeo sobre, entonces todo elemento $\rho \in \Gamma_{\text{inv}}$ se puede escribir como $\rho = \pi(c)$, para algún $c \in \mathcal{A}$. Por lo tanto, para todo $a \in \mathcal{A}$ y $\rho \in \Gamma_{\text{inv}}$ tenemos

$$\begin{aligned} (a\kappa(c^{(1)})) dc^{(2)} &= a\kappa(c^{(1)}) (c^{(2)} \otimes \pi(c^{(3)})) = a (\kappa(c^{(1)})c^{(2)}) \otimes \pi(c^{(3)}) \\ &= a\epsilon(c^{(1)}) \otimes \pi(c^{(2)}) = a \otimes \pi(c) = a \otimes \rho. \end{aligned}$$

Por lo tanto $(\mathcal{A} \otimes \Gamma_{\text{inv}}, d)$ es un cálculo diferencial de primer orden sobre el grupo cuántico.

Similarmente como en 4.2.7 definimos la co-acción izquierda como

$$\ell(a \otimes \varrho) = \phi(a) \otimes \varrho.$$

Entonces $\mathcal{A} \otimes \Gamma_{\text{inv}}$ es un cálculo diferencial de primer orden sobre $\mathcal{A}$ izquierda covariante. Hay que ver que si $\sum adb = 0$ entonces $\sum a^{(1)}b^{(1)} \otimes a^{(2)}db^{(2)} = 0$. Sea entonces $\sum adb = 0$ es decir, $\sum ab^{(1)} \otimes \pi(b^{(2)}) = 0$ entonces

$$\begin{aligned} \sum a^{(1)}b^{(1)} \otimes a^{(2)}db^{(2)} &= \sum a^{(1)}b^{(1)} \otimes (a^{(2)}b^{(2)} \otimes \pi(b^{(3)})) = \\ \sum \phi(a)\phi(b^{(1)}) \otimes \pi(b^{(2)}) &= \sum \phi(ab^{(1)}) \otimes \pi(b^{(2)}) = \ell(\sum ab^{(1)} \otimes \pi(b^{(2)})) = 0. \end{aligned}$$

Por lo tanto el cálculo $(\mathcal{A} \otimes \Gamma_{\text{inv}}, d)$ es izquierda covariante. $\square$

Al cálculo diferencial $(\mathcal{A} \otimes \Gamma_{\text{inv}}, d)$ izquierda covariante mencionado en la proposición anterior, lo podemos pensar como una *envolvente covariante* del cálculo diferencial (Γ, d) . De esta manera, siempre podemos asociar un cálculo covariante a un cálculo diferencial dado.

Teorema 4.2.14. Sea (Γ, d) un cálculo diferencial de primer orden sobre $\mathcal{A}$ izquierda covariante. Entonces los cálculos diferenciales $(\mathcal{A} \otimes \Gamma_{\text{inv}}, d)$ y (Γ, d) de primer orden sobre $\mathcal{A}$ son iguales.

Demostración. Por la proposición (4.2.9) el bimódulo Γ es un bimódulo izquierda covariante y por la proposición (4.2.7) $\mathcal{A} \otimes \Gamma_{\text{inv}}$ tiene una estructura de bimódulo isomorfo a Γ .

El isomorfismo de bimódulos $\iota : \Gamma \rightarrow \mathcal{A} \otimes \Gamma_{\text{inv}}$ dado por $\iota(\theta) = \theta^{(1)} \otimes p(\theta^{(0)})$ cumple que:

$$\iota(da) = a^{(1)} \otimes p(d(a^{(2)})) = a^{(1)} \otimes \pi(a^{(2)}) = d(a),$$

donde usamos que la co-acción izquierda en Γ cumple $\ell(da) = a^{(1)} \otimes d(a^{(2)})$. Por lo tanto los cálculos son iguales. $\square$

Observación 4.2.6. Observemos que el mapeo multiplicación $m : \mathcal{A} \otimes \Gamma_{\text{inv}} \rightarrow \Gamma$ dado por $m(b \otimes \varrho) = b\varrho$ es el mapeo inverso del isomorfismo ι de la demostración anterior.

Así que tenemos que si (Γ, d) es izquierda covariante, entonces el morfismo m (ó bien ι) es isomorfismo.

Por otro lado, si (Γ, d) es un cálculo diferencial de primer orden sobre $\mathcal{A}$ y m es un isomorfismo entonces

$$m(d(a)) = m(a^{(1)} \otimes \pi(a^{(2)})) = a^{(1)}\pi(a^{(2)}) = a^{(1)}\kappa(a^{(2)})d(a^{(3)}) = da.$$

Es decir, los cálculos (Γ, d) y $(\mathcal{A} \otimes \Gamma_{\text{inv}}, d)$ son iguales.

Ahora, como $(\mathcal{A} \otimes \Gamma_{\text{inv}}, d)$ es un cálculo diferencial izquierda covariante (véase proposición 4.2.13) entonces (Γ, d) es izquierda covariante también.

Por lo tanto, el cálculo (Γ, d) es izquierda covariante si y solo si el mapeo multiplicación $m : \mathcal{A} \otimes \Gamma_{\text{inv}} \rightarrow \Gamma$ es isomorfismo.

Ya hemos visto que cualquier cálculo diferencial (Γ, d) de primer orden sobre $\mathcal{A}$ se obtiene como factor cálculo del cálculo universal $(\mathcal{A}^2, D)$, es decir, $(\Gamma, d) \cong (\mathcal{A}^2/\mathcal{N}, \delta)$ donde $\mathcal{N}$ es un sub-bimódulo de $\mathcal{A}^2$, $\delta = \Pi \circ D$ y Π la proyección de $\mathcal{A}^2$ sobre $\mathcal{A}^2/\mathcal{N}$.

Además tenemos que el cálculo universal es un cálculo diferencial bicovariante. ¿Qué condiciones debemos poner sobre el sub-bimódulo $\mathcal{N}$ para que el factor cálculo sea bicovariante también?

Proposición 4.2.15. Sea $\mathcal{N}$ un sub-bimódulo de $\mathcal{A}^2$ y $\delta : \mathcal{A} \rightarrow \mathcal{A}^2/\mathcal{N}$ la diferencial del teorema 4.1.1. Entonces $(\mathcal{A}^2/\mathcal{N}, \delta)$ es un cálculo diferencial de primer orden sobre $\mathcal{A}$ izquierda covariante si y solamente si $\mathcal{N}$ es izquierda invariante, es decir, $\ell(\mathcal{N}) \subseteq \mathcal{A} \otimes \mathcal{N}$.

Demostración. Supongamos que el cálculo $(\mathcal{A}^2/\mathcal{N}, \delta)$ es izquierda covariante. Sean ℓ la única co-acción izquierda de $\mathcal{A}$ en $\mathcal{A}^2/\mathcal{N}$ dada en la proposición 4.2.9 y $\Pi : \mathcal{A}^2 \rightarrow \mathcal{A}^2/\mathcal{N}$ la proyección canónica. Entonces para todo elemento $q \in \mathcal{A}^2/\mathcal{N}$ se cumple lo siguiente

$$\begin{aligned} (\text{id} \otimes \Pi)\ell(q) &= (\text{id} \otimes \Pi)(q^{(1)} \otimes q^{(0)}) = q^{(1)} \otimes \Pi(q^{(0)}) = \Pi(q^{(1)} \otimes q^{(0)}) \\ &= \Pi(\ell(q)) = \ell(\Pi(q)). \end{aligned}$$

En particular, si $q \in \mathcal{N}$, $\ell(q) \in \ker(\text{id} \otimes \Pi) \in \mathcal{A} \otimes \mathcal{N}$, es decir, $\mathcal{N}$ es izquierda invariante.

Inversamente, supongamos que el sub-bimódulo $\mathcal{N}$ es izquierda invariante y que $\sum a\delta(b) = 0$. Entonces como $\Pi(\sum aD(b)) = \sum a\delta(b)$ tenemos que

$\sum aD(b) \in \mathcal{N}$. Así, aplicando nuestra hipótesis de que $\mathcal{N}$ es izquierda invariante tendremos que $\ell(\sum aD(b)) \in \mathcal{A} \otimes \mathcal{N}$.

Por lo tanto,

$$\begin{aligned} 0 &= (\text{id} \otimes \Pi)\ell(\sum aD(b)) = (\text{id} \otimes \Pi)(\sum a^{(1)}b^{(1)} \otimes a^{(2)}D(b^{(2)})) \\ &= \sum a^{(1)}b^{(1)} \otimes a^{(2)}\delta(b^{(2)}). \end{aligned}$$

Es decir, $(\mathcal{A}^2/\mathcal{N}, \delta)$ es un cálculo diferencial de primer orden sobre $\mathcal{A}$ izquierda covariante. $\square$

Proposición 4.2.16. Sea $\mathcal{N}$ un sub-bimódulo de $\mathcal{A}^2$ y $\delta : \mathcal{A} \rightarrow \mathcal{A}^2/\mathcal{N}$ la diferencial del teorema 4.1.1. Entonces $(\mathcal{A}^2/\mathcal{N}, \delta)$ es un cálculo diferencial de primer orden sobre $\mathcal{A}$ derecha covariante si y solamente si $\mathcal{N}$ es derecha invariante, es decir, $\wp(\mathcal{N}) \subseteq \mathcal{N} \otimes \mathcal{A}$. $\square$

Corolario 4.2.17. Sea $\mathcal{N}$ un sub-bimódulo de $\mathcal{A}^2$ y $\delta : \mathcal{A} \rightarrow \mathcal{A}^2/\mathcal{N}$ la diferencial del teorema 4.1.1. El factor cálculo $(\mathcal{A}^2/\mathcal{N}, \delta)$ es un cálculo diferencial de primer orden sobre $\mathcal{A}$ bicovariante si y solamente si $\mathcal{N}$ es izquierda y derecha invariante. $\square$

Como veremos en seguida, el cálculo diferencial sobre un álgebra de Hopf no es único. En este contexto, ocurre un fenómeno completamente cuántico, existe una multitud de cálculos diferenciales y cuando el álgebra de Hopf es conmutativa este fenómeno sigue válido. Algunos de estos cálculos se proyectan al cálculo diferencial clásico, y otros en general no.

Observación 4.2.7. Así ejemplos de espacios triviales a primera vista dados por los grupos finitos, por ejemplo, donde el cálculo clásico se trivializa a 0, esconden una nueva naturaleza que sólo se hace visible en el contexto cuántico.

Sea $\mathcal{A}$ un álgebra de Hopf, y sean $r : \mathcal{A} \otimes \mathcal{A} \rightarrow \mathcal{A} \otimes \mathcal{A}$ y $s : \mathcal{A} \otimes \mathcal{A} \rightarrow \mathcal{A} \otimes \mathcal{A}$ mapeos lineales dados por las siguientes fórmulas:

$$r(a \otimes b) = (a \otimes 1)\phi(b) = ab^{(1)} \otimes b^{(2)}, \quad (4.20)$$

$$s(a \otimes b) = (1 \otimes a)\phi(b) = b^{(1)} \otimes ab^{(2)}. \quad (4.21)$$

Se puede ver que ambos mapeos son biyectivos y la inversa de cada uno está dada por los mapeos

$$r^{-1}(a \otimes b) = a\kappa(b^{(1)}) \otimes b^{(2)}, \quad (4.22)$$

$$s^{-1}(a \otimes b) = b\kappa^{-1}(a^{(2)}) \otimes a^{(1)}, \quad (4.23)$$

respectivamente.

Podemos definir en el espacio vectorial $\mathcal{A} \otimes \ker(\epsilon)$ una estructura de $\mathcal{A}$ -bimódulo izquierda covariante de la siguiente manera: las acciones izquierda y derecha de $\mathcal{A}$ en $\mathcal{A} \otimes \ker(\epsilon)$ están dadas por

$$c(a \otimes b) = ca \otimes b, \quad (a \otimes b)c = ac^{(1)} \otimes bc^{(2)},$$

para todo $a, c \in \mathcal{A}$ y $b \in \ker(\epsilon)$. La co-acción izquierda está dada por

$$\ell(a \otimes b) = \phi(a) \otimes b, \quad (4.24)$$

para todo $a \in \mathcal{A}$ y $b \in \ker(\epsilon)$.

Observación 4.2.8. Todos nuestros cálculos para este caso, se pueden ver como caso particular de los anteriores realizados para cálculos diferenciales Γ izquierda covariantes arbitrarios.

Los elementos izquierda invariantes de $\mathcal{A} \otimes \ker(\epsilon)$ son todos los elementos del espacio $1 \otimes \ker(\epsilon) \cong \ker(\epsilon)$.

En efecto, $\ell(1 \otimes b) = \phi(1) \otimes b = 1 \otimes 1 \otimes b$, para todo $b \in \ker(\epsilon)$. Es decir, $1 \otimes b$ es izquierda invariante.

Por otro lado, si suponemos que $\sum a \otimes b$ es izquierda invariante, para una suma finita de elementos $a \in \mathcal{A}$ y $b \in \ker(\epsilon)$, entonces

$$\ell(\sum a \otimes b) = \sum \phi(a) \otimes b = 1 \otimes (\sum a \otimes b) = \sum 1 \otimes a \otimes b.$$

Obsérvese que

$$(\text{id} \otimes \epsilon \otimes \text{id})(\sum 1 \otimes a \otimes b) = \sum 1 \otimes \epsilon(a) \otimes b = \sum 1 \otimes \epsilon(a)b,$$

y también

$$(\text{id} \otimes \epsilon \otimes \text{id})(\sum \phi(a) \otimes b) = \sum (\text{id} \otimes \epsilon)\phi(a) \otimes b = \sum a \otimes b.$$

Por lo tanto

$$\sum a \otimes b = \sum 1 \otimes \epsilon(a)b.$$

Así, como cada $b \in \ker(\epsilon)$ entonces $\sum a \otimes b \in 1 \otimes \ker(\epsilon)$.

Observación 4.2.9. Sea $\mathcal{A}$ un álgebra de Hopf y $\mathcal{A}^2 = \ker(m)$. Entonces $r(\mathcal{A}^2) = \mathcal{A} \otimes \ker(\epsilon)$. En efecto, primeramente observemos que se satisface la siguiente relación:

$$mr^{-1} = (\text{id} \otimes \epsilon).$$

Por lo tanto, para todo $q \in \mathcal{A}^2$ tenemos que

$$(\text{id} \otimes \epsilon)r(q) = mr^{-1}r(q) = m(q) = 0,$$

lo cual implica que $r(\mathcal{A}^2) \subseteq \mathcal{A} \otimes \ker(\epsilon)$.

Por otro lado, dado que r ya es biyectivo, entonces su restricción a $\mathcal{A}^2$ es inyectiva. También si $a \in \mathcal{A}$ y $b \in \ker(\epsilon)$ entonces

$$mr^{-1}(a \otimes b) = (\text{id} \otimes \epsilon)(a \otimes b) = 0,$$

es decir, $r^{-1}(a \otimes b) \in \mathcal{A}^2$. Por lo tanto $r(\mathcal{A}^2) = \mathcal{A} \otimes \ker(\epsilon)$.

Observación 4.2.10. La biyección r restringida a $\mathcal{A}^2 = \ker(m)$ es un isomorfismo de $\mathcal{A}$ -bimódulos y el siguiente diagrama

$$\begin{array}{ccc} \mathcal{A}^2 & \xrightarrow{r} & \mathcal{A} \otimes \ker(\epsilon) \\ \ell \downarrow & & \downarrow \phi \otimes \text{id} \\ \mathcal{A} \otimes \mathcal{A}^2 & \xrightarrow[\text{id} \otimes r]{} & \mathcal{A} \otimes \mathcal{A} \otimes \ker(\epsilon) \end{array} \quad (4.25)$$

es conmutativo.

Corolario 4.2.18. Un elemento $q \in \mathcal{A}^2$ es izquierda invariante si y sólo si $q = r^{-1}(1 \otimes b)$ para algún $b \in \ker(\epsilon)$. $\square$

Observación 4.2.11. El isomorfismo r mueve la parte invariante de $\mathcal{A}^2$ en la parte invariante de $\mathcal{A} \otimes \ker(\epsilon)$. Por lo tanto, $\mathcal{A}_{\text{inv}}^2 \cong 1 \otimes \ker(\epsilon) \cong \ker(\epsilon)$.

Observación 4.2.12. Sea $\mathcal{A}$ un álgebra de Hopf y $(\mathcal{A}^2, D)$ el cálculo universal. Aplicando el teorema 4.2.14, tenemos que éste es igual al cálculo $(\mathcal{A} \otimes \mathcal{A}_{\text{inv}}^2, d) = (\mathcal{A} \otimes \ker(\epsilon), d)$, donde $d(a) = a^{(1)} \otimes \pi(a^{(2)})$, para todo $a \in \mathcal{A}$. Más explícitamente, usando que el mapeo de gérmenes está dado por la fórmula $\pi(a) = a^{(1)}D(a^{(2)})$ se tiene lo siguiente:

$$\begin{aligned} d(a) &= a^{(1)} \otimes \pi(a^{(2)}) = a^{(1)} \otimes \kappa(a^{(2)})D(a^{(3)}) \\ &= a^{(1)} \otimes \kappa(a^{(2)})(1 \otimes a^{(3)} - a^{(3)} \otimes 1) \\ &= a^{(1)} \otimes \kappa(a^{(2)}) \otimes a^{(3)} - a^{(1)} \otimes \epsilon(a^{(2)}) \otimes 1 \\ &= a^{(1)} \otimes \kappa(a^{(2)}) \otimes a^{(3)} - a \otimes 1 \otimes 1 \\ &= a^{(1)} \otimes r^{-1}(1 \otimes a^{(2)}) - a \otimes r^{-1}(1 \otimes 1) \\ &\cong a^{(1)} \otimes a^{(2)} - a \otimes 1, \end{aligned}$$

para todo $a \in \mathcal{A}$.

Cabe observar que éste cálculo para $d(a)$ se puede obtener de forma más directa notando que para el cálculo diferencial universal el ideal $\mathcal{R} = \{0\}$ y $\pi(a) = a - \epsilon(a)$, para cada $a \in \mathcal{A}$, es decir, los cálculos anteriores se reducen al siguiente:

$$\begin{aligned} d(a) &= a^{(1)} \otimes \pi(a^{(2)}) = a^{(1)} \otimes (a^{(2)} - \epsilon(a^{(2)})) = a^{(1)} \otimes a^{(2)} - a^{(1)} \epsilon(a^{(2)}) \otimes 1 \\ &= \phi(a) - a \otimes 1. \end{aligned}$$

Por lo tanto el cálculo universal es igual al cálculo $(\mathcal{A} \otimes \ker(\epsilon), d)$ donde $d(a) = \phi(a) - a \otimes 1$.

Podemos traducir el teorema 4.1.1 para clasificar todos los cálculos diferenciales de primer orden sobre $\mathcal{A}$ en términos del cálculo $(\mathcal{A} \otimes \ker(\epsilon), d)$. Ya hemos visto que los cálculos cocientes serán izquierda (o derecha) covariantes cuando el sub-bimódulo por el que se factoriza es izquierda (o derecha) invariante.

Notemos que si consideramos $\mathcal{R} \subseteq \ker(\epsilon)$ un ideal derecho arbitrario del $\ker(\epsilon)$, entonces $\mathcal{A} \otimes \mathcal{R}$ es un sub-bimódulo de $\mathcal{A} \otimes \ker(\epsilon)$. Recordemos que la co-acción izquierda ℓ de $\mathcal{A}$ en $\mathcal{A} \otimes \ker(\epsilon)$ está dada por la fórmula $\ell(a \otimes b) = \phi(a) \otimes b$, de manera que

$$\ell(\mathcal{A} \otimes \mathcal{R}) = \phi(\mathcal{A}) \otimes \mathcal{R} \subseteq \mathcal{A} \otimes \mathcal{A} \otimes \mathcal{R},$$

es decir, $\mathcal{A} \otimes \mathcal{R}$ es izquierda invariante.

Por lo tanto el cálculo diferencial que se obtiene de $\mathcal{A} \otimes \ker(\epsilon)$ factorizando por el sub-bimódulo $\mathcal{A} \otimes \mathcal{R}$ es izquierda covariante.

De hecho, como anunciaremos a continuación, todos los cálculos diferenciales izquierda covariantes sobre un álgebra de Hopf son de este tipo. Esta clasificación la realiza por primera vez Woronowicz en su artículo [23].

Teorema 4.2.19. Sea $\mathcal{A}$ un álgebra de Hopf y $\mathcal{R}$ un ideal derecho de $\mathcal{A}$ contenido en $\ker \epsilon$. Entonces $\mathcal{N} = \mathcal{A} \otimes \mathcal{R}$ es un sub-bimódulo de $\mathcal{A} \otimes \ker(\epsilon)$ y el cálculo $(\mathcal{A} \otimes \ker(\epsilon)/\mathcal{N}, \delta)$, donde $\delta = \Pi \circ d$, $d(a) = \phi(a) - a \otimes 1$ y Π es la proyección canónica $\Pi : \mathcal{A} \otimes \ker(\epsilon) \rightarrow \mathcal{A} \otimes \ker(\epsilon)/\mathcal{N}$, es izquierda covariante.

Más aún, cualquier cálculo diferencial de primer orden sobre $\mathcal{A}$ izquierda covariante es de la forma $(\mathcal{A} \otimes \ker(\epsilon)/(\mathcal{A} \otimes \mathcal{R}), \delta)$, para algún ideal derecho $\mathcal{R} \subseteq \ker \epsilon$. $\square$

Resumiendo, si (Γ, d) es un cálculo diferencial izquierda covariante de primer orden sobre un álgebra de Hopf $\mathcal{A}$, entonces por el teorema 4.2.14, tenemos la igualdad de este cálculo con el cálculo diferencial $(\mathcal{A} \otimes \Gamma_{\text{inv}}, d)$. Ahora por el teorema 4.2.19, existe un ideal derecho $\mathcal{R} \subseteq \ker(\epsilon)$ tal que

$$(\mathcal{A} \otimes \Gamma_{\text{inv}}, d) = (\mathcal{A} \otimes \ker(\epsilon)/(\mathcal{A} \otimes \mathcal{R}), \delta).$$

Por lo tanto, como ya habíamos mencionado (ver observación 4.2.5) llegamos por otro camino a la misma conclusión, es decir, para cada cálculo diferencial izquierda covariante, existe un ideal derecho $\mathcal{R}$ del $\ker(\epsilon)$ tal que

$$\Gamma_{\text{inv}} \cong \ker \epsilon / \mathcal{R}, \quad (4.26)$$

y el cálculo es igual al cálculo diferencial $(\mathcal{A} \otimes \Gamma_{\text{inv}}, d)$.

Así cada cálculo izquierda covariante está en correspondencia con un ideal derecho $\mathcal{R}$ contenido en el kernel de la co-unidad. Si el ideal es muy pequeño, entonces el cálculo correspondiente se vuelve más complicado desde el punto de vista operacional. Por otro lado, si el ideal es grande, el cálculo resultante se puede acercar a lo trivial.

Hay un teorema análogo que clasifica a los cálculos diferenciales de primer orden derecha covariante, no lo mencionaremos en esta tesis, sin embargo se puede ver en [23].

Definición 4.2.7. Sea $\mathcal{A}$ un álgebra de Hopf, definimos la acción adjunta $\text{ad} : \mathcal{A} \rightarrow \mathcal{A} \otimes \mathcal{A}$ como el mapeo lineal

$$\text{ad}(a) = a^{(2)} \otimes \kappa(a^{(1)})a^{(3)}, \quad (4.27)$$

para todo $a \in \mathcal{A}$.

Observación 4.2.13. Consideremos los isomorfismos r y s definidos en las ecuaciones (4.20) y (4.21), entonces

$$\begin{aligned} sr^{-1}(1 \otimes a) &= s(\kappa(a^{(1)}) \otimes a^{(2)}) = (1 \otimes \kappa(a^{(1)}))\phi(a^{(2)}) = a^{(2)} \otimes \kappa(a^{(1)})a^{(3)} \\ &= \text{ad}(a), \end{aligned}$$

es decir, $\text{ad} = sr^{-1}$.

Proposición 4.2.20. El mapeo ad es una co-acción derecha de $\mathcal{A}$ en sí mismo. En otras palabras

$$\begin{aligned} (\text{id} \otimes \phi)\text{ad} &= (\text{ad} \otimes \text{id})\text{ad}, \\ (\text{id} \otimes \epsilon)\text{ad}(a) &= a, \end{aligned}$$

para cada $a \in \mathcal{A}$.

Demostración. En efecto, para todo $a \in \mathcal{A}$ se cumple lo siguiente:

$$\begin{aligned} (\text{id} \otimes \phi)\text{ad}(a) &= a^{(2)} \otimes \phi(\kappa(a^{(1)})a^{(3)}) = a^{(3)} \otimes \phi(\kappa(a^{(1)}))\phi(a^{(3)}) \\ &= a^{(3)} \otimes (\kappa(a^{(2)})a^{(4)} \otimes \kappa(a^{(1)})a^{(5)}) \\ &= (a^{(3)} \otimes \kappa(a^{(2)})a^{(4)}) \otimes \kappa(a^{(1)})a^{(5)} \\ &= \text{ad}(a^{(2)}) \otimes \kappa(a^{(1)})a^{(3)} = (\text{ad} \otimes \text{id})\text{ad}(a). \end{aligned}$$

Además

$$(\text{id} \otimes \epsilon)\text{ad}(a) = a^{(2)} \otimes \epsilon(\kappa(a^{(1)})a^{(3)}) = a^{(2)} \otimes \epsilon(a^{(1)}) = a.$$

Así ad es una co-acción derecha de $\mathcal{A}$ en sí mismo. $\square$

Definición 4.2.8. Sea V un subespacio lineal de $\mathcal{A}$. Decimos que V es *ad invariante* si

$$\text{ad}(V) \subseteq V \otimes \mathcal{A}.$$

Observemos que si V es un subespacio ad -invariante entonces la restricción $\text{ad} : V \rightarrow V \otimes \mathcal{A}$ es una co-acción derecha.

Proposición 4.2.21. Sea $\mathcal{A}$ un álgebra de Hopf, entonces el ideal $\ker(\epsilon)$ es *ad invariante*.

Demostración. Para todo $a \in \mathcal{A}$ se satisface lo siguiente:

$$(\epsilon \otimes \text{id})\text{ad}(a) = \epsilon(a^{(2)}) \otimes \kappa(a^{(1)})a^{(3)} = \kappa(a^{(1)})a^{(2)} = \epsilon(a).$$

Es decir, ϵ es ad -invariante. Y en particular el ideal $\ker(\epsilon)$ es ad invariante. $\square$

Lema 4.2.22. Supongamos que (Γ, d) es un cálculo diferencial bicovariante con una co-acción derecha $\wp$ de $\mathcal{A}$ en Γ . Entonces el siguiente diagrama es conmutativo:

$$\begin{array}{ccc} \mathcal{A} & \xrightarrow{\text{ad}} & \mathcal{A} \otimes \mathcal{A} \\ \pi \downarrow & & \downarrow \pi \otimes \text{id} \\ \Gamma_{\text{inv}} & \xrightarrow[\wp]{} & \Gamma_{\text{inv}} \otimes \mathcal{A} \end{array} \quad (4.28)$$

Demostración. La co-acción derecha actúa en el espacio Γ_{inv} de la siguiente manera:

$$\begin{aligned} \wp(\pi(a)) &= \wp(\kappa(a^{(1)})da^{(2)}) = (\kappa(a^{(2)}) \otimes \kappa(a^{(1)})) (da^{(3)} \otimes a^{(4)}) \\ &= \kappa(a^{(2)})da^{(3)} \otimes \kappa(a^{(1)})a^{(4)} = \pi(a^{(2)}) \otimes \kappa(a^{(1)})a^{(3)}, \end{aligned}$$

es decir, el diagrama 4.28 conmuta. Por lo tanto $\wp$ coincide con la acción proyectada adjunta ad en el espacio Γ_{inv} . $\square$

Proposición 4.2.23. Sea $\mathcal{A}$ un álgebra de Hopf y $\Gamma_{\text{inv}} = \ker(\epsilon)/\mathcal{R}$ para un ideal derecho $\mathcal{R}$ de $\ker(\epsilon)$. Sea $(\Gamma = \mathcal{A} \otimes \Gamma_{\text{inv}}, d)$ el cálculo diferencial izquierda covariante de la proposición 4.2.13. Entonces (Γ, d) es bicovariante si y sólo si $\mathcal{R}$ es ad invariante. $\square$

Proposición 4.2.24. Sea $\mathcal{A}$ un álgebra de Hopf y $\Gamma_{\text{inv}} = \ker(\epsilon)/\mathcal{R}$ para un ideal derecho $\mathcal{R}$ de $\ker(\epsilon)$. Sea $(\Gamma = \mathcal{A} \otimes \Gamma_{\text{inv}}, d)$ el cálculo diferencial izquierda covariante de la proposición 4.2.13. Entonces (Γ, d) es $*$ -cálculo si y sólo si $\kappa(a)^* \in \mathcal{R}$ para todo $a \in \mathcal{R}$. En este caso se cumple

$$\pi(a)^* = -\pi[\kappa(a)^*], \quad (4.29)$$

para todos $a \in \mathcal{A}$.

Demostración. Notemos que si (Γ, d) es $*$ -cálculo entonces se cumple lo siguiente

$$\begin{aligned} \pi(a)^* &= (\kappa(a^{(1)})d(a^{(2)}))^* = d(a^{(2)*})\kappa(a^{(1)*}) \\ &= d(a^{(2)*}\kappa(a^{(1)*}) - a^{(2)*}d(\kappa(a^{(1)*})) \\ &= d(\epsilon(a)^*) - a^{(2)*}d(\kappa(a^{(1)*})) \\ &= -a^{(2)*}d(\kappa(a^{(1)*})) = -\kappa[\kappa(a^{(2)*})]d(\kappa(a^{(1)*})) = -\pi[\kappa(a)^*]. \end{aligned}$$

En particular si $a \in \mathcal{R}$ se debe satisfacer que $\kappa(a)^* \in \mathcal{R}$. Es decir, $\mathcal{R} = \kappa(\mathcal{R})^*$.

Inversamente, supongamos que $\mathcal{R} = \kappa(\mathcal{R})^*$. Por la observación 4.1.2, la $*$ queda definida de manera única en el bimódulo Γ de tal manera que d y $*$ conmutan. Así que si encontramos una estructura $*$ en Γ habremos terminado. Definamos primero la estructura $*$ en el espacio de los izquierda invariantes por la fórmula

$$\pi(a)^* = -\pi(\kappa(a)^*).$$

Obsérvese que si $\pi(a) = \pi(b)$ para elementos de $\ker(\epsilon)$ entonces $a - b \in \mathcal{R}$, y por la hipótesis entonces $\kappa(a - b)^* = \kappa(a)^* - \kappa(b)^* \in \mathcal{R}$. Así se tiene que $\pi(a)^* = \pi(b)^*$, es decir la fórmula para la estrella en los elementos izquierda invariantes está bien definida.

El bimódulo $\mathcal{A} \otimes \Gamma_{\text{inv}}$ es isomorfo al bimódulo Γ por medio de la identificación $a \otimes \theta \rightsquigarrow a\theta$. La estructura $*$ en Γ se define como

$$(a\theta)^* = \theta^*a^*,$$

para todo $a \in \mathcal{A}$ y $\theta \in \Gamma_{\text{inv}}$.

Observemos que la estructura de $\mathcal{A}$ -módulo derecho en Γ se ve como

$$\pi(a)b = b^{(1)}\pi[(a - \epsilon(a))b^{(2)}],$$

de manera que

$$\begin{aligned} \pi[\kappa(b^{(2)*})]b^{(1)*} &= b^{(1)*}\pi[(\kappa(b^{(3)*}) - \epsilon(\kappa(b^{(3)*})))b^{(2)*}] \\ &= b^{(1)*}\pi[\epsilon(b^{(2)*})] - b^{(1)*}\pi(b^{(2)*}) \\ &= -b^{(1)*}\pi(b^{(2)*}) = -d(b^*), \end{aligned}$$

para todo $b \in \mathcal{A}$.

Así que

$$[ad(b)]^* = [ab^{(1)}\pi(b^{(2)})]^* = \pi(b^{(2)})^*(ab^{(1)})^* = -\pi[\kappa(b^{(2)})^*]b^{(1)*}a^* = d(b^*)a^*,$$

para todo $a, b \in \mathcal{A}$. Por lo tanto (Γ, d) es un $*$ -cálculo. $\square$

Como se puede ver en el artículo [23], se puede definir el operador de intercambio σ entre los espacios de elementos izquierda invariantes, para cada bimódulo bicovariante (y en particular para los cálculos bicovariantes de primer orden). Y en [4] encontramos la fórmula explícita

$$\sigma(\eta \otimes \theta) = \theta^{(0)} \otimes (\eta \circ \theta^{(1)}), \quad (4.30)$$

donde $\text{ad}(\theta) = \theta^{(0)} \otimes \theta^{(1)}$ es la acción derecha en bimódulos.

Observación 4.2.14. En el caso clásico se cumple que

$$\theta \circ a = \epsilon(a)\theta,$$

para todo $\theta \in \Gamma_{\text{inv}}$ y $a \in \mathcal{A}$, de manera que el operador σ definido anteriormente cumple

$$\sigma(\eta \otimes \theta) = \theta^{(0)}\epsilon(\theta^{(1)}) \otimes \eta = \theta \otimes \eta,$$

es decir, el operador σ en el caso clásico coincide con la transposición estándar.

Proposición 4.2.25. El operador σ satisface la ecuación de trenza, es decir,

$$(\sigma \otimes \text{id})(\text{id} \otimes \sigma)(\sigma \otimes \text{id}) = (\text{id} \otimes \sigma)(\sigma \otimes \text{id})(\text{id} \otimes \sigma). \quad (4.31)$$

Demostración. Sean $\eta, \theta, \varrho \in \Gamma_{\text{inv}}$, $\text{ad}(\theta) = \theta^{(0)} \otimes \theta^{(1)}$ y $\text{ad}(\varrho) = \varrho^{(0)} \otimes \varrho^{(1)}$ entonces

$$\begin{aligned} (\sigma \otimes \text{id})(\text{id} \otimes \sigma)(\sigma \otimes \text{id})(\eta \otimes \theta \otimes \varrho) &= (\sigma \otimes \text{id})(\text{id} \otimes \sigma)(\theta^{(0)} \otimes (\eta \circ \theta^{(1)}) \otimes \varrho) \\ &= (\sigma \otimes \text{id})(\theta^{(0)} \otimes \varrho^{(0)} \otimes ((\eta \circ \theta^{(1)}) \circ \varrho^{(1)})) = (\varrho^{(0)} \otimes (\theta^{(0)} \circ \varrho^{(1)}) \otimes (\eta \circ (\theta^{(1)} \varrho^{(2)}))), \end{aligned}$$

por otro lado tenemos:

$$\begin{aligned} (\text{id} \otimes \sigma)(\sigma \otimes \text{id})(\text{id} \otimes \sigma)(\eta \otimes \theta \otimes \varrho) &= (\text{id} \otimes \sigma)(\sigma \otimes \text{id})(\eta \otimes \varrho^{(0)} \otimes \theta \circ \varrho^{(1)}) \\ &= (\text{id} \otimes \sigma)(\varrho^{(0)} \otimes (\eta \circ \varrho^{(1)}) \otimes (\theta \circ \varrho^{(2)})) \\ &= (\varrho^{(0)} \otimes (\theta^{(0)} \circ \varrho^{(3)}) \otimes (\eta \circ \varrho^{(1)}) \circ (\kappa(\varrho^{(2)})\theta^{(1)}\varrho^{(4)})) \\ &= \varrho^{(0)} \otimes (\theta^{(0)} \circ \varrho^{(2)}) \otimes (\eta \circ (\epsilon(\varrho^{(1)})\theta^{(1)}\varrho^{(3)})) = \varrho^{(0)} \otimes (\theta^{(0)} \circ \varrho^{(1)}) \otimes (\eta \circ (\theta^{(1)}\varrho^{(2)})), \end{aligned}$$

por lo tanto se cumple la ecuación de trenza. $\square$

4.3 Álgebra tensorial

Definición 4.3.1. Sean $\mathfrak{B}$ un álgebra con unidad y $\mathcal{A}$ un álgebra de Hopf. Decimos que $\mathfrak{B}$ es un álgebra bicovariante sobre $\mathcal{A}$ si $\mathfrak{B}$ es un bimódulo bicovariante tal que las co-acciones de $\mathcal{A}$ en $\mathfrak{B}$ son homomorfismos de álgebras.

Si $\mathfrak{B}$ es un álgebra bicovariante sobre $\mathcal{A}$, decimos que $\mathfrak{B}$ es un álgebra graduada bicovariante sobre $\mathcal{A}$, si

$$\mathfrak{B} = \sum_{n=0}^{\infty} \oplus \mathfrak{B}^n$$

es un álgebra graduada y las co-acciones preservan el grado, es decir

$$\ell(\mathfrak{B}^n) \subseteq \mathcal{A} \otimes \mathfrak{B}^n, \quad \wp(\mathfrak{B}^n) \subseteq \mathfrak{B}^n \otimes \mathcal{A},$$

para todo $n \in \mathbb{N}$. Denotamos en este caso por $\ell_{\mathfrak{B}^n}$, y $\wp_{\mathfrak{B}^n}$ a la restricción de ℓ y $\wp$ a $\mathfrak{B}^n$, respectivamente.

Observemos que toda álgebra de Hopf es un álgebra bicovariante sobre sí misma, donde las co-acciones izquierda y derecha están dadas por el coproducto.

Definición 4.3.2. Sea $(\Psi, \ell_{\Psi}, \wp_{\Psi})$ un bimódulo bicovariante sobre un álgebra de Hopf $\mathcal{A}$. Decimos que un álgebra graduada bicovariante $(\mathfrak{B}, \ell, \wp)$ está construida sobre Ψ si

1. $(\mathfrak{B}^0, \ell_{\mathfrak{B}^0}, \wp_{\mathfrak{B}^0}) = (\mathcal{A}, \phi, \phi)$;
2. El bimódulo bicovariante $(\mathfrak{B}^1, \ell_{\mathfrak{B}^1}, \wp_{\mathfrak{B}^1})$ coincide con $(\Psi, \ell_{\Psi}, \wp_{\Psi})$;
3. El álgebra $\mathfrak{B}$ está generada por elementos de grado uno, es decir, si $\tau \in \mathfrak{B}^n$, $n = 2, 3, \dots$, entonces

$$\tau = \sum \tau_i, \text{ donde } \tau_i \text{ es producto de } n \text{ elementos de } \Psi.$$

Dado un bimódulo Ψ sobre $\mathcal{A}$, denotamos por $\Psi^{\otimes n}$ al producto tensorial sobre el álgebra $\mathcal{A}$ de Ψ consigo mismo n veces, $n = 2, 3, \dots$ es decir,

$$\Psi^{\otimes n} = \Psi \otimes_{\mathcal{A}} \Psi \otimes_{\mathcal{A}} \cdots \otimes_{\mathcal{A}} \Psi.$$

De aquí en adelante, si el contexto lo permite, escribiremos simplemente $a_1 \otimes a_2 \otimes \cdots \otimes a_n$ en lugar de la expresión $a_1 \otimes_{\mathcal{A}} a_2 \otimes_{\mathcal{A}} \cdots \otimes_{\mathcal{A}} a_n$ que representa a los elementos del producto tensorial $\Psi^{\otimes n}$, es decir, omitiremos los subíndices que acompañan al producto tensorial.

Proposición 4.3.1. Sea Ψ un bimódulo sobre un álgebra de Hopf $\mathcal{A}$. Entonces $\Psi^{\otimes n}$ es un bimódulo sobre $\mathcal{A}$ con acciones dadas por las siguientes fórmulas

$$\begin{aligned} a(\varrho_1 \otimes \varrho_2 \otimes \cdots \otimes \varrho_n) &= (a\varrho_1) \otimes \varrho_2 \otimes \cdots \otimes \varrho_n, \\ (\varrho_1 \otimes \varrho_2 \otimes \cdots \otimes \varrho_n)a &= \varrho_1 \otimes \varrho_2 \otimes \cdots \otimes (\varrho_n a), \end{aligned}$$

para todo $a \in \mathcal{A}$ y $\varrho_1, \dots, \varrho_n \in \Psi$. □

Proposición 4.3.2. Sea Ψ un bimódulo bicovariante sobre un álgebra de Hopf $\mathcal{A}$. El bimódulo $\Psi^{\otimes n}$ es bicovariante con las co-acciones izquierda y derecha dadas por las siguientes fórmulas:

$$\begin{aligned} \ell^{\otimes n}(\varrho_1 \otimes \varrho_2 \otimes \cdots \otimes \varrho_n) &= a_1^{(1)} a_2^{(1)} \cdots a_n^{(1)} \otimes \varrho_1^{(0)} \otimes \varrho_2^{(0)} \otimes \cdots \otimes \varrho_n^{(0)}, \\ \wp^{\otimes n}(\varrho_1 \otimes \varrho_2 \otimes \cdots \otimes \varrho_n) &= \varrho_1^{(0)} \otimes \varrho_2^{(0)} \otimes \cdots \otimes \varrho_n^{(0)} \otimes b_1^{(1)} b_2^{(1)} \cdots b_n^{(1)}, \end{aligned}$$

donde $\ell(\varrho_i) = a_i^{(1)} \otimes \varrho_i^{(0)}$ y $\wp(\varrho_i) = \varrho_i^{(0)} \otimes b_i^{(1)}$, $i = 1, \dots, n$. □

Definimos $(\Psi^{\otimes 0}, \ell^{\otimes 0}, \wp^{\otimes 0}) = (\mathcal{A}, \phi, \phi)$ y $(\Psi^{\otimes 1}, \ell^{\otimes 1}, \wp^{\otimes 1}) = (\Psi, \ell, \wp)$. Sean

$$\Psi^{\otimes} = \sum_{n=0}^{\infty} \oplus \Psi^{\otimes n}, \quad \ell^{\otimes} = \sum_{n=0}^{\infty} \oplus \ell^{\otimes n}, \quad \wp^{\otimes} = \sum_{n=0}^{\infty} \oplus \wp^{\otimes n}.$$

Observemos que $\Psi^{\otimes}$ es un álgebra graduada que contiene a $\mathcal{A}$ como subálgebra de elementos de grado cero y a Ψ como subespacio de todos los elementos de grado uno. Los mapeos $\ell^{\otimes}$ y $\wp^{\otimes}$ son mapeos lineales multiplicativos actuando de $\Psi^{\otimes}$ a $\mathcal{A} \otimes \Psi^{\otimes}$ y a $\Psi^{\otimes} \otimes \mathcal{A}$, respectivamente. Y por lo tanto obtenemos el siguiente resultado:

Proposición 4.3.3. Sea Ψ un bimódulo bicovariante sobre un álgebra de Hopf $\mathcal{A}$, entonces $(\Psi^{\otimes}, \ell^{\otimes}, \wp^{\otimes})$ es un álgebra graduada bicovariante construida sobre $(\Psi, \ell, \wp)$. □

Proposición 4.3.4. Sea $\mathcal{A}$ un álgebra de Hopf y Ψ un bimódulo bicovariante sobre $\mathcal{A}$, entonces los mapeos definidos como

$$\phi(\theta) = \ell(\theta) \oplus \wp(\theta), \quad \epsilon(\theta) = 0, \quad \kappa(\theta) = -\theta^{(0)} \kappa(\theta^{(1)}),$$

para todo $\theta \in \Psi_{inv}$ se extienden naturalmente al álgebra tensorial $\Psi^{\otimes}$ dándole una estructura de álgebra de Hopf. □

Observación 4.3.1. Consideremos $\mathcal{A}$ un álgebra de Hopf y Ψ un bimódulo sobre $\mathcal{A}$. En la proposición anterior podemos considerar el coproducto ϕ

definido sobre el álgebra tensorial $\Psi^\otimes$ al álgebra tensorial $\Psi^\otimes \widehat{\otimes} \Psi^\otimes$ donde $\widehat{\otimes}$ es el producto tensorial graduado. Es decir si $\alpha \otimes \beta$ y $a \otimes b$ son elementos de $\Psi^\otimes \widehat{\otimes} \Psi^\otimes$ entonces

$$(\alpha \otimes \beta)(a \otimes b) = (-1)^{\partial\beta\partial a} \alpha a \otimes \beta b. \quad (4.32)$$

De esta manera, $\Psi^\otimes$ se convierte en una *super álgebra de Hopf*, donde el coproducto en los elementos de grado uno es la suma de la coacción izquierda y la coacción derecha, así como también la counidad y la coinversa están dadas como en la proposición anterior, módulo el simple cambio de la coinversa inducido por la graduación.

Proposición 4.3.5. Sean Ψ un bimódulo izquierda covariante, con co-acción izquierda ℓ y $\Psi_{\text{inv}}^{\otimes n}$, $n \in \mathbb{N}$ el espacio de elementos izquierda invariantes de $\Psi^{\otimes n}$ con respecto a la co-acción izquierda $\ell^{\otimes n}$. Entonces

$$\Psi_{\text{inv}}^{\otimes n} = \underbrace{\Psi_{\text{inv}} \otimes \cdots \otimes \Psi_{\text{inv}}}_n. \quad (4.33)$$

Podemos extender la co-acción derecha $\text{ad} : \Psi_{\text{inv}} \rightarrow \Psi_{\text{inv}} \otimes \mathcal{A}$, a una co-acción derecha en $\Psi_{\text{inv}} \otimes \Psi_{\text{inv}}$ sobre $\mathcal{A}$, de tal forma que el siguiente diagrama

$$\begin{array}{ccc} \Psi_{\text{inv}} \otimes \Psi_{\text{inv}} & \xrightarrow{\text{ad}} & \Psi_{\text{inv}} \otimes \Psi_{\text{inv}} \otimes \mathcal{A} \\ \text{ad} \otimes \text{ad} \downarrow & & \uparrow \text{id} \otimes \text{id} \otimes m \\ \Psi_{\text{inv}} \otimes \mathcal{A} \otimes \Psi_{\text{inv}} \otimes \mathcal{A} & \longrightarrow & \Psi_{\text{inv}} \otimes \Psi_{\text{inv}} \otimes \mathcal{A} \otimes \mathcal{A} \end{array}$$

sea conmutativo. Donde hemos usado la transposición estándar entre los elementos de posición 2 y 3 en la flecha inferior. Es decir, si η y θ son elementos izquierda invariantes, entonces el diagrama nos dice que

$$\text{ad}(\eta \otimes \theta) = \eta^{(0)} \otimes \theta^{(0)} \otimes \eta^{(1)} \theta^{(1)}. \quad (4.34)$$

De igual forma, también extendemos la acción derecha $\circ$ de $\mathcal{A}$ en Ψ_{inv} , a una acción derecha de $\mathcal{A}$ en $\Psi_{\text{inv}} \otimes \Psi_{\text{inv}}$ de manera que el siguiente diagrama

$$\begin{array}{ccc} \Psi_{\text{inv}} \otimes \Psi_{\text{inv}} \otimes \mathcal{A} & \xrightarrow{\circ} & \Psi_{\text{inv}} \otimes \Psi_{\text{inv}} \\ \text{id} \otimes \text{id} \otimes \phi \downarrow & & \uparrow \circ \otimes \circ \\ \Psi_{\text{inv}} \otimes \Psi_{\text{inv}} \otimes \mathcal{A} \otimes \mathcal{A} & \longrightarrow & \Psi_{\text{inv}} \otimes \mathcal{A} \otimes \Psi_{\text{inv}} \otimes \mathcal{A} \end{array}$$

sea conmutativo. En la flecha inferior, usamos la transposición estándar entre elementos de posición 2 y 3. Es decir,

$$(\eta \otimes \theta) \circ a = (\eta \circ a^{(1)}) \otimes (\theta \circ a^{(2)}), \quad (4.35)$$

para todo $\eta, \theta \in \Psi_{\text{inv}}$ y $a \in \mathcal{A}$.

Proposición 4.3.6. El operador σ definido en la ecuación (4.30) es un automorfismo de módulo derecho.

Demostración. Sea $a \in \mathcal{A}$ y $\eta \otimes \theta \in \Psi_{\text{inv}} \otimes \Psi_{\text{inv}}$ entonces

$$\begin{aligned}
 \sigma((\eta \otimes \theta) \circ a) &= \sigma((\eta \circ a^{(1)}) \otimes (\theta \circ a^{(2)})) \\
 &= (\theta^{(0)} \circ a^{(3)}) \otimes ((\eta \circ a^{(1)}) \circ (\kappa(a^{(2)})\theta^{(1)}a^{(4)})) \\
 &= (\theta^{(0)} \circ a^{(1)}) \otimes (\eta \circ (\theta^{(1)}a^{(2)})) \\
 &= (\theta^{(0)} \circ a^{(1)}) \otimes ((\eta \circ \theta^{(1)}) \circ a^{(2)}) \\
 &= (\theta^{(0)} \otimes (\eta \circ \theta^{(1)})) \circ a \\
 &= \sigma(\eta \otimes \theta) \circ a.
 \end{aligned}$$

Por lo tanto σ es un homomorfismo de módulo derecho. Y además como también es invertible, se tiene que es automorfismo. $\square$

El siguiente diagrama nos muestra que el automorfismo σ conmuta con la acción adjunta ad .

Proposición 4.3.7. El siguiente diagrama

$$\begin{array}{ccc}
 \Psi_{\text{inv}} \otimes \Psi_{\text{inv}} & \xrightarrow{\text{ad}} & \Psi_{\text{inv}} \otimes \Psi_{\text{inv}} \otimes \mathcal{A} \\
 \sigma \downarrow & & \downarrow \sigma \otimes \text{id} \\
 \Psi_{\text{inv}} \otimes \Psi_{\text{inv}} & \xrightarrow{\text{ad}} & \Psi_{\text{inv}} \otimes \Psi_{\text{inv}} \otimes \mathcal{A},
 \end{array} \tag{4.36}$$

es conmutativo.

Demostración. Sean $\theta, \eta \in \Psi_{\text{inv}}$ y $\text{ad}(\theta \otimes \eta) = \theta^{(0)} \otimes \eta^{(0)} \otimes \theta^{(1)}\eta^{(1)}$, entonces

$$\begin{aligned}
 (\sigma \otimes \text{id})\text{ad}(\theta \otimes \eta) &= (\sigma \otimes \text{id})(\theta^{(0)} \otimes \eta^{(0)} \otimes \theta^{(1)}\eta^{(1)}) \\
 &= \eta^{(0)} \otimes (\theta^{(0)} \circ \eta^{(1)}) \otimes \theta^{(1)}\eta^{(2)},
 \end{aligned}$$

y por otro lado se tiene que

$$\begin{aligned}
 \text{ad}\sigma(\theta \otimes \eta) &= \text{ad}(\eta^{(0)} \otimes (\theta \circ \eta^{(1)})) \\
 &= \eta^{(0)} \otimes (\theta^{(0)} \circ \eta^{(3)}) \otimes \eta^{(1)}\kappa(\eta^{(2)})\theta^{(1)}\eta^{(4)} \\
 &= \eta^{(0)} \otimes (\theta^{(0)} \circ \eta^{(2)}) \otimes \epsilon(\eta^{(1)})\theta^{(1)}\eta^{(3)} \\
 &= \eta^{(0)} \otimes (\theta^{(0)} \circ \eta^{(1)}) \otimes \theta^{(1)}\eta^{(2)},
 \end{aligned}$$

por lo que el diagrama (4.36) conmuta. $\square$

4.4 Álgebra exterior

En esta sección nos basamos en las notas del libro Principal Bundles, the Quantum Case [20]. Para cada entero $j > 0$ tal que $1 \leq j < n$ definimos la permutación $\tau_j \in S_n$, en el grupo de permutaciones de n elementos, como el 2-ciclo que intercambia j y $j + 1$, es decir, $\tau_j = (j, j + 1)$. Estos satisfacen las siguientes relaciones para todo $1 \leq i < n$ y $1 \leq j < n$:

$$\tau_i \tau_j = \tau_j \tau_i \quad \text{si } |i - j| > 2, \quad \tau_i \tau_j \tau_i = \tau_j \tau_i \tau_j \quad \text{si } |i - j| = 1, \quad \tau_i^2 = e.$$

Sabemos que el grupo de permutaciones de n elementos, S_n , es isomorfo al grupo abstracto generado por $n - 1$ símbolos, $\tau_1, \dots, \tau_{n-1}$ que satisfacen las relaciones antes mencionadas.

El grupo trenzado B_n para $n \geq 2$ se define como el grupo generado por $n - 1$ símbolos, $\sigma_1, \dots, \sigma_{n-1}$ que satisface solamente las relaciones de conmutación y de trenza, y no las de involutividad. Es decir,

$$\sigma_i \sigma_j = \sigma_j \sigma_i \quad \text{si } |i - j| > 2, \quad \sigma_i \sigma_j \sigma_i = \sigma_j \sigma_i \sigma_j \quad \text{si } |i - j| = 1.$$

Existe un único homomorfismo de grupos $p_n : B_n \rightarrow S_n$ sobreyectivo, para todo $n \geq 2$, definido como

$$p_n(\sigma_i) = \tau_i, \quad i = 1, \dots, n - 1.$$

El grupo trenzado B_n , se puede representar naturalmente como un grupo de automorfismos de $\Psi^{\otimes n}$, haciendo

$$\sigma_i \mapsto \text{id} \otimes \dots \otimes \text{id} \otimes \sigma \otimes \dots \otimes \text{id}, \quad (4.37)$$

donde σ es el operador de intercambio que aparece en la i -ésima entrada y la función id ocurre $n - 2$ veces.

Sea $T = \{\tau_1, \dots, \tau_{n-1}\}$ el conjunto de transposiciones de S_n antes mencionadas. Denotamos por $I(\tau)$ al mínimo número de transposiciones en T en las que se descompone la permutación τ .

Por definición diremos que $I(e) = 0$ para e la permutación identidad. Si $\tau \in S_2$ se tiene que su descomposición con elementos de T es única, sin embargo la descomposición de cada permutación en S_n , para $n \geq 3$, ya no lo es.

Supongamos que $\tau \in S_n$ se descompone con elementos de T como $\tau = \tau_{i_1} \tau_{i_2} \dots \tau_{i_m}$ donde $I(\tau) = m$. Definimos la función $\pi_n : S_n \rightarrow B_n$ como

$$\pi_n(\tau) = \sigma_{i_1} \sigma_{i_2} \dots \sigma_{i_m}.$$

Por definición tendremos que $\pi_n(e) = e$, y la función π_n no depende de la descomposición de τ en transposiciones con elementos de T pues los elementos σ_{i_k} cumplen la ecuación de trenza.

Sin embargo π_n no es un morfismo de grupos, ya que al elemento neutro en S_n se puede escribir como el cuadrado de una transposición, pero el cuadrado de la imagen de la transposición no da el elemento neutro de B_n . Observemos que $p_n \circ \pi_n(\tau) = \tau$, por lo tanto la función π_n es inyectiva.

Proposición 4.4.1. Si $\alpha, \beta \in S_n$ son dos permutaciones de n elementos tales que $I(\alpha\beta) = I(\alpha) + I(\beta)$, entonces $\pi_n(\alpha\beta) = \pi_n(\alpha)\pi_n(\beta)$. $\square$

Definición 4.4.1. Para cada $n \geq 2$ definimos el operador antisimetrizador, $A_n : \Psi^{\otimes n} \rightarrow \Psi^{\otimes n}$, como

$$A_n = \sum_{\alpha \in S_n} \text{sign}(\alpha) \pi_n(\alpha),$$

donde $\text{sign}(\alpha) = (-1)^{I(\alpha)}$ y $\pi_n(\alpha) \in B_n$ actúa naturalmente en $\Psi^{\otimes n}$ como en (4.37). Además, definimos $A_0 = \text{id}_A$ y $A_1 = \text{id}_\Gamma$.

Sean h un entero $0 \leq h \leq n$ y

$$SH_{n,h} = \left\{ \tau \in S_k \mid \tau(i) < \tau(j) \text{ si } (1 \leq i < j \leq h \text{ ó } h < i < j \leq n) \right\}.$$

Proposición 4.4.2. Sea $\tau \in S_n$, entonces para todo $0 \leq h \leq n$, existen permutaciones únicas α , β y ν de n elementos tal que α deja fijos a los números $h+1, h+2, \dots, n$, β deja fijos a los primeros h números y $\nu \in SH_{n,h}$ tal que

$$\tau = \nu\alpha\beta. \quad (4.38)$$

Definimos el h -operador de antisimetrizador $A_{n,h} : \Psi^{\otimes n} \rightarrow \Psi^{\otimes n}$ como

$$A_{n,h} = \sum_{\alpha \in SH_{n,h}} \text{sign}(\alpha) \pi_n(\alpha).$$

Como consecuencia directa de la proposición anterior, para todo $0 \leq h \leq n$ se satisface

$$A_n = A_{n,h}(A_h \otimes A_{n-h}),$$

por lo tanto podemos concluir:

Proposición 4.4.3. Si $K_n := \ker A_n \subset \Psi^{\otimes n}$, entonces

$$K = \sum_{n=0}^{\infty} \oplus K_n$$

es un ideal bilateral graduado de $\Psi^{\otimes}$. $\square$

Definimos el álgebra exterior trenzada $\Psi^\vee$ asociada al bimódulo Ψ , como el cociente del álgebra tensorial y el ideal K , es decir,

$$\Psi^\otimes / K \cong \sum_{n=0}^{\infty} \oplus \Psi^{\otimes n} / K_n = \sum_{n=0}^{\infty} \oplus \Psi^{\vee n}. \quad (4.39)$$

Esta álgebra, es un álgebra graduada con el producto cociente tal que $\Psi^{\vee 0} = \mathcal{A}$ y $\Psi^{\vee 1} = \Psi$. Si $\omega_1 \in \Psi^{\vee k}$ y $\omega_2 \in \Psi^{\vee l}$, denotamos su producto como

$$\omega_1 \wedge \omega_2 \in \Psi^{\vee k+l},$$

el cual es llamado producto cuña (trenzado) ó bien, producto exterior (trenzado).

El mapeo A_n es un morfismo de $\Psi^{\otimes n}$ de $\mathcal{A}$ co-módulo derecho e izquierdo, y el ideal K_n es invariante bajo las co-acciones izquierdas y derechas $\ell_{\Psi^{\otimes n}}$ y $\wp_{\Psi^{\otimes n}}$. Denotamos por

$$\ell^\vee := \sum_{n=0}^{\infty} \oplus \ell_n^\vee, \quad \wp^\vee := \sum_{n=0}^{\infty} \oplus \wp_n^\vee,$$

a las co-acción izquierda y derecha, respectivamente, de $\mathcal{A}$ en $\Psi^\vee$.

Teorema 4.4.4. El álgebra $(\Psi^\vee, \ell^\vee, \wp^\vee)$ es un álgebra graduada bicovariante construida sobre el bimódulo bicovariante $(\Psi, \ell, \wp)$.

Las álgebras bicovariantes sobre un álgebra $\mathcal{A}$ se pueden ver como bimódulos bicovariantes sobre $\mathcal{A}$, los cuales forman una categoría. Las flechas entre Ψ y Φ son las transformaciones $\mathcal{A}$ -lineales $F: \Psi \rightarrow \Phi$ tales que los siguientes diagramas

$$\begin{array}{ccc} \Psi & \xrightarrow{\wp_\Psi} & \Psi \otimes \mathcal{A} \\ F \downarrow & & \downarrow F \otimes \text{id} \\ \Phi & \xrightarrow[\wp_\Phi]{} & \Phi \otimes \mathcal{A} \end{array} \quad \begin{array}{ccc} \Psi & \xrightarrow{\ell_\Psi} & \mathcal{A} \otimes \Psi \\ F \downarrow & & \downarrow \text{id} \otimes F \\ \Phi & \xrightarrow[\ell_\Phi]{} & \mathcal{A} \otimes \Phi \end{array}$$

son conmutativos. Dentro de esta categoría, las álgebras bicovariantes forman una subcategoría. En efecto, si $F: \Psi \rightarrow \Phi$ es un morfismo en la categoría de bimódulos se satisface que los operadores de intercambio correspondientes y $F^{\otimes 2}$ conmutan, es decir,

$$\begin{array}{ccc} \Psi \otimes_{\mathcal{A}} \Psi & \xrightarrow{F^{\otimes 2}} & \Phi \otimes_{\mathcal{A}} \Phi \\ \sigma_\Psi \downarrow & & \downarrow \sigma_\Phi \\ \Psi \otimes_{\mathcal{A}} \Psi & \xrightarrow{F^{\otimes 2}} & \Phi \otimes_{\mathcal{A}} \Phi \end{array}$$

y por lo tanto también el diagrama

$$\begin{array}{ccc} \Psi^\otimes & \xrightarrow{F^\otimes} & \Phi^\otimes \\ A_\Psi \downarrow & & \downarrow A_\Phi \\ \Psi^\otimes & \xrightarrow{F^\otimes} & \Phi^\otimes \end{array}$$

es conmutativo, donde $A = \sum^\oplus A_n$ y el subíndice Ψ y Φ indican que es el antisimetrizador del álgebra tensorial correspondiente. Por lo tanto $K = \sum^\oplus K_n$ es invariante bajo $F^\otimes$, lo que implica que pasa al cociente $\Psi^\vee$ y por lo tanto tenemos un morfismo $F^\vee : \Psi^\vee \rightarrow \Phi^\vee$ de álgebras graduadas bicovariantes.

Teorema 4.4.5. El álgebra exterior trenzada es un funtor covariante entre la categoría de bimódulos bicovariantes sobre $\mathcal{A}$ y la categoría de álgebras bicovariantes sobre $\mathcal{A}$.

Woronowicz prueba en su artículo [23], que el álgebra exterior construida sobre el bimódulo bicovariante tiene la siguiente propiedad:

Teorema 4.4.6. Sean Ψ y Φ bimódulos sobre $\mathcal{A}$ tales que Φ es sub-bimódulo de Ψ entonces la inyección $\Phi \hookrightarrow \Psi$ induce un morfismo inyectivo entre las álgebras exteriores $\Phi^\vee$ y $\Psi^\vee$.

De hecho, el álgebra $\Psi^\vee$ de Woronowicz es una super álgebra de Hopf que extiende a $\mathcal{A}$, donde el coproducto en el primer orden esta dado por la suma de acciones izquierda y derecha, como lo vimos en el caso del álgebra tensorial (ver observación (4.3.1)). De hecho se puede demostrar que el ideal $\ker(A)$ es el coideal en el álgebra tensorial.

4.5 Cálculo diferencial de más alto orden

4.5.1 Cálculo exterior trenzado

En esta sección veremos la construcción del cálculo diferencial exterior presentada por Woronowicz en el artículo [23], el cual consiste de diferenciales en el álgebra exterior trenzada que son análogos de las diferenciales de De Rham en el caso clásico, para la construcción usaremos el método de bimódulos extendidos presentado en [24].

Consideremos (Γ, d) un cálculo diferencial de primer orden sobre $\mathcal{A}$. Como ya hemos notado antes, $\mathcal{A}$ se puede ver como bimódulo bicovariante sobre sí mismo, con las co-acciones dadas por el co-producto. Consideramos la suma directa de $\mathcal{A}$ -módulos izquierdos

$$\tilde{\Gamma} := \mathcal{A} \oplus \Gamma. \quad (4.40)$$

Notemos que $\tilde{\Gamma}$ es naturalmente un $\mathcal{A}$ -módulo izquierdo. Además, si denotamos por $X := (1, 0) \in \tilde{\Gamma}$, al elemento de $\tilde{\Gamma}$, entonces cualquier elemento $\tilde{\omega}$ en $\tilde{\Gamma}$, se puede escribir de manera única como

$$\tilde{\omega} = aX + \omega, \quad a \in \mathcal{A} \text{ y } \omega \in \Gamma. \quad (4.41)$$

Notemos que si en la ecuación (4.41) hacemos $a = 0$, tenemos naturalmente incluido a Γ en $\tilde{\Gamma}$.

Proposición 4.5.1. *El módulo izquierdo $\tilde{\Gamma}$ es un bimódulo con acción derecha dada por*

$$\tilde{\omega}b = abX + ad(b) + \omega b, \quad (4.42)$$

para todo $b \in \mathcal{A}$ y $\tilde{\omega} = (aX + \omega) \in \tilde{\Gamma}$.

Demostración. En efecto, si $\tilde{\omega} = aX + \omega$ entonces

$$\tilde{\omega}1 = (a1)X + ad(1) + \omega 1 = aX + \omega = \tilde{\omega},$$

pues $d(1) = 0$. Además si $b_1, b_2 \in \mathcal{A}$, entonces usando la asociatividad de $\mathcal{A}$ se tiene que

$$\begin{aligned} (\tilde{\omega}b_1)b_2 &= (ab_1X + ad(b_1) + \omega b_1)b_2 \\ &= (ab_1)b_2X + (ab_1)d(b_2) + (ad(b_1))b_2 + (\omega b_1)b_2 \\ &= a(b_1b_2)X + ad(b_1b_2) + \omega(b_1b_2) \\ &= (aX + \omega)(b_1b_2) = \tilde{\omega}(b_1b_2), \end{aligned}$$

así que tenemos una acción derecha de $\mathcal{A}$ en $\tilde{\Gamma}$.

Ahora veamos que $\tilde{\Gamma}$ es un bimódulo sobre $\mathcal{A}$. Sean entonces $b, c \in \mathcal{A}$ y $\tilde{\omega} \in \tilde{\Gamma}$, entonces

$$\begin{aligned} b(\tilde{\omega}c) &= b(acX + ad(c) + \omega c) = b(ac)X + b(ad(c)) + b(\omega c) \\ &= (ba)cX + (ba)d(c) + (b\omega)c = (baX + b\omega)c = (b\tilde{\omega})c, \end{aligned}$$

por lo tanto las acciones son compatibles y $\tilde{\Gamma}$ es un bimódulo sobre $\mathcal{A}$. $\square$

Observación 4.5.1. Por definición de la acción derecha de $\mathcal{A}$ sobre $\tilde{\Gamma}$, tenemos que $Xa = aX + 1d(a) = aX + da$, por lo que

$$Xa - aX = da, \quad (4.43)$$

es decir la diferencial de Γ se ve como el conmutador con X de $\tilde{\Gamma}$.

Para obtener un álgebra exterior sobre $\tilde{\Gamma}$ debemos darle estructura de bimódulo bicovariante sobre $\mathcal{A}$.

Proposición 4.5.2. Sea $(\Gamma, \ell, \wp)$ un bimódulo bicovariante sobre $\mathcal{A}$. Entonces $(\tilde{\Gamma}, \tilde{\ell}, \tilde{\wp})$ es un bimódulo bicovariante sobre $\mathcal{A}$, donde las co-acciones izquierda y derecha $\tilde{\ell} : \tilde{\Gamma} \rightarrow \mathcal{A} \otimes \tilde{\Gamma}$, $\tilde{\wp} : \tilde{\Gamma} \rightarrow \tilde{\Gamma} \otimes \mathcal{A}$, están dadas por

$$\begin{aligned} \tilde{\ell}(\tilde{\omega}) &= a^{(1)} \otimes a^{(2)}X + \ell(\omega), \\ \tilde{\wp}(\tilde{\omega}) &= a^{(1)}X \otimes a^{(2)} + \wp(\omega), \end{aligned}$$

para todo $\tilde{\omega} = aX + \omega$.

Demostración. Veamos que $\tilde{\ell}$ es una co-acción derecha de $\mathcal{A}$ en $\tilde{\Gamma}$. Sea $\tilde{\omega} = aX + \omega$, notemos que $\tilde{\ell}(\omega) = \ell(\omega)$, así

$$\begin{aligned} (\phi \otimes \text{id})\tilde{\ell}(\tilde{\omega}) &= (\phi \otimes \text{id})(a^{(1)} \otimes a^{(2)}X) + (\phi \otimes \text{id})\ell(\omega) \\ &= (\phi(a^{(1)}) \otimes a^{(2)}X) + (\text{id} \otimes \ell)\ell(\omega) \\ &= ((a^{(1)} \otimes a^{(2)}) \otimes a^{(3)}X) + (\text{id} \otimes \ell)\ell(\omega) \\ &= (a^{(1)} \otimes (a^{(2)} \otimes a^{(3)}X)) + (\text{id} \otimes \ell)\ell(\omega) \\ &= a^{(1)} \otimes \tilde{\ell}(a^{(2)}X) + (\text{id} \otimes \tilde{\ell})\ell(\omega) \\ &= (\text{id} \otimes \tilde{\ell})(a^{(1)} \otimes a^{(2)}X + \ell(\omega)) \\ &= (\text{id} \otimes \tilde{\ell})\tilde{\ell}(\tilde{\omega}), \end{aligned}$$

por lo tanto, $(\phi \otimes \text{id})\tilde{\ell} = (\text{id} \otimes \tilde{\ell})\tilde{\ell}$.

Por otro lado,

$$\begin{aligned} (\epsilon \otimes \text{id})\tilde{\ell}(\tilde{\omega}) &= (\epsilon \otimes \text{id})(a^{(1)} \otimes a^{(2)}X) + (\epsilon \otimes \text{id})\ell(\omega) \\ &= (\epsilon(a^{(1)}) \otimes a^{(2)}X) + (\epsilon \otimes \text{id})\ell(\omega) \\ &= 1 \otimes (\epsilon(a^{(1)})a^{(2)}X) + (1 \otimes \omega) \\ &= 1 \otimes (aX + \omega) = \tilde{\omega}. \end{aligned}$$

Además, para todo $b \in \mathcal{A}$ tenemos

$$\begin{aligned} \tilde{\ell}(b\tilde{\omega}) &= \phi(b)(a^{(1)} \otimes a^{(2)}X) + \phi(b)\ell(\omega) \\ &= \phi(b)\tilde{\ell}(\tilde{\omega}), \end{aligned}$$

donde usamos que ℓ es una co-acción izquierda de $\mathcal{A}$ en Γ . Por otra parte:

$$\begin{aligned}
 \tilde{\ell}(\tilde{\omega}b) &= \tilde{\ell}((aX + \omega)b) = \tilde{\ell}(abX + adb + \omega b) \\
 &= (a^{(1)}b^{(1)} \otimes a^{(2)}b^{(2)}X) + a^{(1)}b^{(1)} \otimes a^{(2)}db^{(2)} + \ell(\omega)\phi(b) \\
 &= \phi(a) \left((b^{(1)} \otimes b^{(2)}X + b^{(1)} \otimes db^{(2)}) + \ell(\omega)\phi(b) \right) \\
 &= \phi(a)(b^{(1)} \otimes Xb^{(2)}) + \ell(\omega)\phi(b) \\
 &= \tilde{\ell}(\tilde{\omega})\phi(b).
 \end{aligned}$$

Por lo tanto $\tilde{\ell}$ es una co-acción izquierda de $\mathcal{A}$ en $\tilde{\Gamma}$. De manera análoga se prueba que $\tilde{\wp}$ es una co-acción derecha de $\mathcal{A}$ en $\tilde{\Gamma}$.

Por último veamos que $(\tilde{\Gamma}, \tilde{\ell}, \tilde{\wp})$ es un bimódulo bi-covariante, es decir,

$$(\text{id} \otimes \tilde{\wp})\tilde{\ell}(\tilde{\omega}) = (\tilde{\ell} \otimes \text{id})\tilde{\wp}(\tilde{\omega}),$$

para todo $\tilde{\omega} \in \tilde{\Gamma}$. En efecto:

$$\begin{aligned}
 (\text{id} \otimes \tilde{\wp})\tilde{\ell}(\tilde{\omega}) &= (\text{id} \otimes \tilde{\wp})(a^{(1)} \otimes a^{(2)}X) + (\text{id} \otimes \tilde{\wp})\ell(\omega) \\
 &= a^{(1)} \otimes \tilde{\wp}(a^{(2)}X) + (\text{id} \otimes \tilde{\wp})\ell(\omega) \\
 &= a^{(1)} \otimes (a^{(2)}X \otimes a^{(3)}) + (\text{id} \otimes \tilde{\wp})\ell(\omega) \\
 &= \tilde{\ell}(a^{(1)}X) \otimes a^{(2)} + (\ell \otimes \text{id})\wp(\omega) \\
 &= (\tilde{\ell} \otimes \text{id})\tilde{\wp}(\tilde{\omega}).
 \end{aligned}$$

Por lo tanto, $(\tilde{\Gamma}, \tilde{\ell}, \tilde{\wp})$ es un bimódulo bicovariante sobre $\mathcal{A}$. $\square$

Observación 4.5.2. Observemos que en el bimódulo $\tilde{\Gamma}$ el elemento X satisface

$$\tilde{\ell}(X) = 1 \otimes X, \quad \tilde{\wp}(X) = X \otimes 1,$$

es decir X es un elemento invariante por la izquierda y por la derecha.

Proposición 4.5.3. Sea Γ un bimódulo bicovariante sobre un álgebra de Hopf $\mathcal{A}$, entonces la estructura de $\mathcal{A}$ -módulo derecho en Γ_{inv} se extiende a $\tilde{\Gamma}_{\text{inv}}$ como

$$X \circ a = \pi(a) + \epsilon(a)X, \tag{4.44}$$

para todo $a \in \mathcal{A}$.

Demostración. Sea $X \in \tilde{\Gamma}$ entonces por la definición de $\circ$ se tiene que

$$X \circ a = \kappa(a^{(1)})Xa^{(2)} = \kappa(a^{(1)}) \left(d(a^{(2)}) + a^{(2)}X \right) = \pi(a) + \epsilon(a)X,$$

para todo $a \in \mathcal{A}$. $\square$

Definición 4.5.1. Sea $\tilde{\Gamma}$ el bimódulo construido anteriormente. Definimos el conmutador graduado en $\tilde{\Gamma}^\vee$ de sus dos elementos homogéneos como

$$[\varrho, \theta] = \varrho \wedge \theta - (-1)^{\partial\varrho\partial\theta} \theta \wedge \varrho. \quad (4.45)$$

Teorema 4.5.4. Sean $(\Gamma, \ell, \wp)$ un cálculo diferencial bicovariante de primer orden sobre $\mathcal{A}$ y $(\Gamma^\vee, \ell^\vee, \wp^\vee)$ el álgebra exterior trenzada construida sobre $(\Gamma, \ell, \wp)$. Entonces existe un único mapeo lineal $d : \Gamma^\vee \rightarrow \Gamma^\vee$ de grado uno tal que

1. La restricción $d : \mathcal{A} = \Gamma^{\vee 0} \rightarrow \Gamma = \Gamma^{\vee 1}$ coincide con la diferencial inicial d ;
2. El mapeo d es derivada graduada:

$$d(\varrho \wedge \theta) = d\varrho \wedge \theta + (-1)^{\partial\varrho} \varrho \wedge d\theta;$$

3. Su cuadrado se anula, es decir, $d^2 = 0$.

Además, se cumple que d es bicovariante, es decir,

$$\begin{aligned} \ell^\vee d &= (\text{id} \otimes d) \ell^\vee, \\ \wp^\vee d &= (d \otimes \text{id}) \wp^\vee. \end{aligned}$$

Demostración. Consideremos $\tilde{\Gamma}$ el bimódulo bicovariante definido en (4.40), y sea $(\tilde{\Gamma}^\vee, \tilde{\ell}^\vee, \tilde{\wp}^\vee)$ el álgebra exterior construida sobre $(\tilde{\Gamma}, \tilde{\ell}, \tilde{\wp})$.

Definimos la diferencial de θ , como el conmutador graduado de X con θ , se tiene que,

$$d\theta := X \wedge \theta - (-1)^{\partial\theta} \theta \wedge X = [X, \theta]. \quad (4.46)$$

Por la observación 4.5.2 y por la definición del operador de trenza $\tilde{\sigma}$ se tiene que

$$\tilde{\sigma}(X \otimes X) = X \otimes (X \circ 1) = X \otimes X.$$

Además, como el grupo de permutaciones de dos elementos se compone de la permutación identidad e y la permutación τ que transpone a dichos elementos, se tiene que

$$\begin{aligned} A_2(X \otimes X) &= \pi_2(e)(X \otimes X) - \pi_2(\tau)(X \otimes X) \\ &= e(X \otimes X) - \tilde{\sigma}(X \otimes X) = 0, \end{aligned}$$

es decir, $X \otimes X \in \ker(A_2)$, por lo que $X \wedge X = 0$.

Veamos que las condiciones del teorema se satisfacen para d . Observemos que por la definición de d y debido a que X tiene grado 1, se tiene que d

incrementa el grado por uno y por la ecuación (4.43) se tiene que d coincide con la derivada en elementos de grado cero.

Por la construcción se cumple que d es una derivada graduada. En efecto, si ϱ y θ son elementos homogéneos en $\tilde{\Gamma}$ entonces

$$\begin{aligned} d\varrho \wedge \theta + (-1)^{\partial\varrho} \varrho \wedge d\theta &= [X, \varrho] \wedge \theta + (-1)^{\partial\varrho} \varrho \wedge [X, \theta] \\ X \wedge \varrho \wedge \theta - (-1)^{\partial\varrho} \varrho \wedge X \wedge \theta + (-1)^{\partial\varrho} \varrho \wedge X \wedge \theta - (-1)^{\partial\varrho + \partial\theta} \varrho \wedge \theta \wedge X \\ X \wedge \varrho \wedge \theta - (-1)^{\partial\varrho + \partial\theta} \varrho \wedge \theta \wedge X &= d(\varrho \wedge \theta), \end{aligned}$$

y por lo tanto d es una derivada graduada (como siempre lo son los operadores dados por los conmutadores graduados con primer argumento fijo).

Además,

$$\begin{aligned} d(d\theta) &= X \wedge d\theta - (-1)^{\partial\theta + 1} d\theta \wedge X = X \wedge (X \wedge \theta - (-1)^{\partial\theta} \theta \wedge X) \\ &\quad - (-1)^{\partial\theta + 1} (X \wedge \theta \wedge X - (-1)^{\partial\theta} \theta \wedge X \wedge X) \\ &= (-1)^{\partial\theta + 1} X \wedge \theta \wedge X - (-1)^{\partial\theta + 1} X \wedge \theta \wedge X = 0. \end{aligned}$$

Ahora veamos que d es bicovariante. Por la definición de $\tilde{\ell}^\vee$ tenemos que

$$\begin{aligned} \tilde{\ell}^\vee(\theta) &= \theta^{(-1)} \otimes \theta^{(0)} \\ \tilde{\ell}^\vee(X \wedge \theta) &= \theta^{(-1)} \otimes (X \wedge \theta^{(0)}) \\ \tilde{\ell}^\vee(\theta \wedge X) &= \theta^{(-1)} \otimes (\theta^{(0)} \wedge X). \end{aligned}$$

Por lo que

$$\begin{aligned} \tilde{\ell}^\vee(d(\theta)) &= \tilde{\ell}^\vee([X, \theta]) = \tilde{\ell}^\vee(X \wedge \theta - (-1)^{\partial\theta} \theta \wedge X) \\ &= \theta^{(-1)} \otimes (X \wedge \theta^{(0)}) - (-1)^{\partial\theta} \theta^{(-1)} \otimes (\theta^{(0)} \wedge X) \\ &= \theta^{(-1)} \otimes [X, \theta^{(0)}] = (\text{id} \otimes d)\tilde{\ell}^\vee(\theta), \end{aligned}$$

es decir, d es covariante a la izquierda.

De forma similar, la co-acción derecha $\tilde{\wp}^\vee$ satisface que

$$\begin{aligned} \tilde{\wp}^\vee(\theta) &= \theta^{(0)} \otimes \theta^{(1)} \\ \tilde{\wp}^\vee(X \wedge \theta) &= (X \wedge \theta^{(0)}) \otimes \theta^{(1)} \\ \tilde{\wp}^\vee(\theta \wedge X) &= (\theta^{(0)} \wedge X) \otimes \theta^{(1)}, \end{aligned}$$

y por lo tanto

$$\begin{aligned} \tilde{\wp}^\vee(d\theta) &= \tilde{\wp}^\vee([X, \theta]) = \tilde{\wp}^\vee(X \wedge \theta - (-1)^{\partial\theta} \theta \wedge X) \\ &= (X \wedge \theta^{(0)}) \otimes \theta^{(1)} - (-1)^{\partial\theta} (\theta^{(0)} \wedge X) \otimes \theta^{(1)} \\ &= [X, \theta^{(0)}] \otimes \theta^{(1)} = (d \otimes \text{id})\tilde{\wp}^\vee(\theta), \end{aligned}$$

así tenemos que d es derecha covariante, y por lo tanto d es bicovariante.

Hasta el momento hemos construido d sobre el bimódulo extendido $\tilde{\Gamma}$. Lo que procede ahora es ver que $\Gamma^\vee$ se queda invariante bajo d , es decir, si $\theta \in \Gamma^\vee$ entonces $d\theta \in \Gamma^\vee$. Sea $\theta \in \Gamma^\vee$ dado por

$$\theta = \sum a_0 da_1 \wedge \cdots \wedge da_k.$$

Entonces es inmediato usando las propiedades de d enunciadas en este teorema que

$$d\theta = \sum da_0 \wedge da_1 \wedge \cdots \wedge da_k,$$

por lo que $d(\Gamma^\vee) \subset \Gamma^\vee$ y automáticamente se prueba la unicidad de d , pues éste queda completamente determinado por sus propiedades. $\square$

Observación 4.5.3. Como mencionamos antes, $\Gamma^\vee$ también se puede ver como una super álgebra de Hopf y por este último teorema además diferencial. Por lo tanto el coproducto ϕ en dicha álgebra se extiende a un homomorfismo $\phi^\vee : \Gamma^\vee \rightarrow \Gamma^\vee \hat{\otimes} \Gamma^\vee$ de super álgebras diferenciales, es decir, entrelaza las diferenciales correspondientes.

Observación 4.5.4. En el caso clásico, cuando G es un grupo de Lie compacto, el cálculo diferencial $\Omega(G) = \Gamma^\vee$ sobre el cálculo de primer orden clásico nos da el álgebra de formas diferenciales clásicas.

Además la identificación natural de $\Omega(G \times G)$ con $\Omega(G) \hat{\otimes} \Omega(G)$ nos muestra que el mapeo $\phi^\vee$, es precisamente el pullback a nivel de formas diferenciales del mapeo producto $m : G \times G \rightarrow G$ y que a nivel cero corresponde al coproducto $\phi : \mathcal{A} \rightarrow \mathcal{A} \otimes \mathcal{A}$.

4.5.2 Cálculo envolvente universal diferencial

Consideremos para esta sección que $\mathcal{A}$ es un álgebra asociativa con unidad y (Γ, d) un cálculo diferencial de primer orden sobre $\mathcal{A}$. Sea además $S^\wedge$ el ideal graduado de $\Gamma^\otimes$ generado por el conjunto

$$\left\{ Q = \sum da \otimes db \mid \sum adb = 0 \right\}. \quad (4.47)$$

Definimos el cálculo diferencial universal como el álgebra cociente

$$\Gamma^\wedge = \Gamma^\otimes / S^\wedge.$$

Notemos que $S^0 = S^1 = 0$ lo que implica que $\Gamma^{\wedge 0} = \mathcal{A}$ y $\Gamma^{\wedge 1} = \Gamma$. La multiplicación del álgebra cociente la indicaremos por yuxtaposición.

Proposición 4.5.5. Sea Γ un cálculo diferencial de primer orden sobre $\mathcal{A}$. Entonces existe un único mapeo lineal $d : \Gamma^\wedge \rightarrow \Gamma^\wedge$ de grado uno que satisface las siguientes propiedades:

1. La restricción $d : \mathcal{A} = \Gamma^{\wedge 0} \rightarrow \Gamma = \Gamma^{\wedge 1}$ coincide con la diferencial inicial d ;
2. El mapeo d es una derivada graduada:

$$d(\nu\eta) = d(\nu)\eta + (-1)^{\partial\nu}\nu d(\eta),$$

para todo $\nu, \eta \in \Gamma^\wedge$;

3. El cuadrado de d se anula, es decir, $d^2 = 0$.

Demostración. Sea $d : \Gamma \rightarrow \Gamma$ definida como

$$d\left(\sum adb\right) = \sum dadb. \quad (4.48)$$

Notemos que si $\sum adb = 0$ entonces $\sum da \otimes db \in S^\wedge$ lo que implica que en el álgebra cociente $\Gamma^\wedge$ se tiene que $\sum dadb = 0$ y por lo tanto d es un mapeo lineal bien definido.

Veamos que dicho d es una derivada graduada en los elementos donde está definida. En efecto sea $\nu \in \Gamma$ tal que $\nu = \sum bdc$ y $a \in \mathcal{A}$, entonces

$$d(a\nu) = \sum d(ab)dc = \sum (da)b(dc) + \sum adbdc = da\nu + ad\nu,$$

y por otro lado

$$\begin{aligned} d(\nu a) &= \sum d(bd(ca) - bcda) = \sum dbd(ca) - \sum d(bc)da \\ &= \sum dbdca + \sum dbcd a - \sum d(bc)da \\ &= \sum dbdca - \sum b(dc)da = d\nu a - \nu da, \end{aligned}$$

además $d(ab) = (da)b + a(db)$ para todo $a, b \in \mathcal{A}$ pues (Γ, d) es un cálculo diferencial, y por lo que d definida en (4.48) es una derivación graduada.

También se tiene que para cada $a \in \mathcal{A}$ se cumple que

$$d(da) = d(1da) = d1da = 0, \quad (4.49)$$

es decir $d^2 = 0$.

El mapeo d admite una única extensión a $\Gamma^\otimes$ tal que si $\vartheta, \varphi \in \Gamma^\otimes$ entonces

$$d(\vartheta\varphi) = d(\vartheta)\Pi(\varphi) + (-1)^{\partial\vartheta}\Pi(\vartheta)d(\varphi), \quad (4.50)$$

donde $\Pi : \Gamma^\otimes \rightarrow \Gamma^\wedge$ es el mapeo proyección.

Observemos además que

$$d(\sum da \otimes db) = \sum d(da)\Pi(db) - \sum \Pi(da)d(db) = 0.$$

Por lo tanto $S^\wedge \subset \ker(d)$, y de esta manera existe un único mapeo $d : \Gamma^\wedge \rightarrow \Gamma^\wedge$ definido como la factorización del d definido anteriormente a través de Π , dicho d satisface las propiedades de la proposición. $\square$

Este cálculo diferencial tiene las siguientes propiedades universales que pueden ser vistas en [4] o en el libro [20]:

Proposición 4.5.6 (Propiedad universal uno). Sean (Ω, d_Ω) un álgebra diferencial y (Γ, d) un cálculo diferencial de primer orden sobre $\mathcal{A}$. Sean también $\varphi^0 : \mathcal{A} = \Gamma^{\wedge 0} \rightarrow \Omega$ un morfismo de álgebras y $\varphi^1 : \Gamma = \Gamma^{\wedge 1} \rightarrow \Omega$ un mapeo lineal tal que

$$\varphi^1(ad b) = \varphi^0(a)d_\Omega(\varphi^0(b)),$$

para todo $a, b \in \mathcal{A}$. Entonces existe un único mapeo lineal $\varphi^n : \Gamma^{\wedge n} \rightarrow \Omega$ para cada $n \geq 2$ tal que

$$\varphi := \sum_{n \geq 0}^{\oplus} \varphi^n : \Gamma^\wedge \rightarrow \Omega$$

es un morfismo de álgebras diferenciales.

Observación 4.5.5. Si consideramos $\Gamma^\vee$ el álgebra exterior trenzada asociada al cálculo diferencial de primer orden (Γ, d) , entonces existe un único monomorfismo de álgebras $\varphi : \Gamma^\wedge \rightarrow \Gamma^\vee$ que es la identidad en los elementos de grado cero y uno.

Proposición 4.5.7 (Propiedad universal dos). Sean (Ω, d_Ω) un álgebra diferencial y (Γ, d) un cálculo diferencial de primer orden sobre $\mathcal{A}$. Sean $\varphi^0 : \mathcal{A} = \Gamma^{\wedge 0} \rightarrow \Omega$ un morfismo antimultiplicativo de espacios vectoriales y $\varphi^1 : \Gamma = \Gamma^{\wedge 1} \rightarrow \Omega$ un mapeo lineal tal que

$$\varphi^1(ad b) = d_\Omega(\varphi^0(b))\varphi^0(a),$$

para todo $a, b \in \mathcal{A}$. Entonces existe un único mapeo lineal $\varphi : \Gamma^\wedge \rightarrow \Omega$ que es morfismo de espacios vectoriales tal que

1. El morfismo φ y d conmutan, es decir, $\varphi d = d_\Omega \varphi$;
2. El morfismo φ es graduadamente antimultiplicativo, es decir, $\varphi(\theta\eta) = (-1)^{\partial\theta\partial\eta}\varphi(\eta)\varphi(\theta)$, para todo $\theta, \eta \in \Gamma^\wedge$.

Teorema 4.5.8. Sea $\mathcal{A}$ un álgebra de Hopf con coproducto $\phi : \mathcal{A} \rightarrow \mathcal{A} \otimes \mathcal{A}$ y (Γ, d) un cálculo diferencial de primer orden sobre $\mathcal{A}$ bicovariante. Entonces existe un único morfismo de álgebras diferenciales graduadas

$$\hat{\phi} : \Gamma^\wedge \rightarrow \Gamma^\wedge \hat{\otimes} \Gamma^\wedge$$

que extiende al coproducto y tal que

$$\hat{\phi}(\theta) = \ell(\theta) + \wp(\theta), \quad (4.51)$$

para todo $\theta \in \Gamma$. En particular si $\omega \in \Gamma_{\text{inv}}$ entonces $\hat{\phi}(\omega) = 1 \otimes \omega + \text{ad}(\omega)$.

Además, el mapeo $\hat{\phi}$ es homomorfismo de álgebras diferenciales y cumple

$$(\text{id} \otimes \hat{\phi})\hat{\phi} = (\hat{\phi} \otimes \text{id})\hat{\phi}. \quad (4.52)$$

Demostración. Notemos que Γ está generado por los elementos de la forma $\sum adb$. Además, como $\hat{\phi}$ debe ser morfismo de álgebras diferenciales y que además extiende al coproducto, entonces necesariamente se debe satisfacer

$$\hat{\phi}(\sum adb) = \sum \phi(a) \hat{\phi}(db) = \sum \phi(a) d\phi(b),$$

para los $a, b \in \mathcal{A}$. Por lo tanto, usando la notación de Sweedler, $\hat{\phi}$ queda definido en los elementos de Γ como:

$$\begin{aligned} \hat{\phi}(\sum adb) &= \sum (a^{(1)} \otimes a^{(2)}) d(b^{(1)} \otimes b^{(2)}) \\ &= \sum (a^{(1)} \otimes a^{(2)}) (db^{(1)} \otimes b^{(2)}) + \sum (a^{(1)} \otimes a^{(2)}) (b^{(1)} \otimes db^{(2)}) \\ &= \sum a^{(1)} db^{(1)} \otimes a^{(2)} b^{(2)} + \sum a^{(1)} b^{(1)} \otimes a^{(2)} db^{(2)} \\ &= \wp(\sum adb) + \ell(\sum adb). \end{aligned}$$

En general, para un elemento $\theta = \sum a_0 da_1 \cdots da_n$, se define

$$\hat{\phi}(\theta) = \sum \phi(a_0) (d \otimes \text{id} + \text{id} \otimes d) \phi(a_1) \cdots (d \otimes \text{id} + \text{id} \otimes d) \phi(a_n).$$

En particular,

$$\hat{\phi}(da) = (da^{(1)} \otimes a^{(2)}) + (a^{(1)} \otimes da^{(2)}).$$

Lo que, conjuntamente con la multiplicatividad de $\hat{\phi}$ y del hecho de que $\Gamma^\wedge$ es generada como álgebra diferencial por $\mathcal{A}$, implica que este mapeo entrelaza

las diferenciales correspondientes. Por la construcción $\widehat{\phi}$ es un morfismo de álgebras diferenciales graduadas. La unicidad se da pues las condiciones del teorema determinan por completo a la extensión. Ahora consideremos $\omega = \pi(b)$, $b \in \mathcal{A}$. Entonces

$$\begin{aligned}\ell(\pi(b)) &= \ell(\kappa(b^{(1)})db^{(2)}) = (\kappa(b^{(2)}) \otimes \kappa(b^{(1)}))(b^{(3)} \otimes db^{(4)}) \\ &= \epsilon(b^{(2)}) \otimes \kappa(b^{(1)})db^{(3)} = 1 \otimes \kappa(b^{(1)})db^{(2)} = 1 \otimes \pi(b).\end{aligned}$$

Y también se tiene para la co-acción derecha

$$\begin{aligned}\wp(\pi(b)) &= \wp(\kappa(b^{(1)})db^{(2)}) = (\kappa(b^{(2)}) \otimes \kappa(b^{(1)}))(db^{(3)} \otimes b^{(4)}) \\ &= \pi(b^{(2)}) \otimes \kappa(b^{(1)})b^{(3)} = \text{ad}(\pi(b)).\end{aligned}$$

Por lo que $\widehat{\phi}(\omega) = 1 \otimes \omega + \text{ad}(\omega)$, para todo $\omega \in \Gamma_{\text{inv}}$.

Y para demostrar la coasociatividad de $\widehat{\phi}$, basta observar que ambos lados de la identidad son homomorfismos entre álgebras graduadas diferenciales $\Gamma^\wedge$ y $\Gamma^\wedge \widehat{\otimes} \Gamma^\wedge \widehat{\otimes} \Gamma^\wedge$, y que $\phi: \mathcal{A} \rightarrow \mathcal{A} \otimes \mathcal{A}$ es coasociativo. $\square$

Proposición 4.5.9. Sea $\mathcal{A}$ un álgebra de Hopf y (Γ, d) un cálculo diferencial de primer orden sobre $\mathcal{A}$ izquierda covariante. Entonces se cumple que

$$d(\pi(a)) = -\pi(a^{(1)})\pi(a^{(2)}),$$

en $\Gamma^\wedge$ (y por la universalidad, en cualquier cálculo diferencial extendiendo el de primer orden $-\Gamma$) para todo $a \in \mathcal{A}$.

Demostración. Sea $a \in \mathcal{A}$ entonces,

$$\begin{aligned}d(\pi(a)) &= d(\kappa(a^{(1)})d(a^{(2)})) = d(\kappa(a^{(1)}))d(a^{(2)}) \\ &= \kappa(a^{(1)})a^{(2)}d(\kappa(a^{(3)}))d(a^{(4)}) - \kappa(a^{(1)})d(a^{(2)})\kappa(a^{(3)})d(a^{(4)}) \\ &= -\kappa(a^{(1)})d(a^{(2)})\kappa(a^{(3)})d(a^{(4)}) = -\pi(a^{(1)})\pi(a^{(2)})\end{aligned}$$

y la proposición queda demostrada. $\square$

Definición 4.5.2. Sean $\mathcal{A}$ un Álgebra de Hopf y Γ un cálculo diferencial de primer orden sobre $\mathcal{A}$. Entonces el conmutador transpuesto de Lie, denotado por $c^T: \Gamma_{\text{inv}} \rightarrow \Gamma_{\text{inv}} \otimes \Gamma_{\text{inv}}$, se define como el mapeo

$$c^T(\theta) = (\text{id} \otimes \pi)\text{ad}(\theta) = \theta^{(0)} \otimes \pi(\theta^{(1)}),$$

para todo $\theta \in \Gamma_{\text{inv}}$.

Definición 4.5.3. Sea Γ un cálculo diferencial de primer orden sobre un álgebra de Hopf $\mathcal{A}$. Definimos un encaje diferencial $\delta : \Gamma_{\text{inv}} \rightarrow \Gamma_{\text{inv}} \otimes \Gamma_{\text{inv}}$ como el mapeo

$$\delta(v) = -(\pi \otimes \pi)\phi(a) = -\pi(a^{(1)}) \otimes \pi(a^{(2)}),$$

donde $v = \pi(a)$ y $a \in \mathcal{L}$, con $\mathcal{L}$ siendo un apropiado complemento covariante de $\mathcal{R}$ en $\ker(\epsilon)$ tal y como está explicado en [4]. Tal complementación, entre otras cosas nos permite identificar Γ_{inv} con $\mathcal{L}$.

Proposición 4.5.10. Sea Γ un cálculo diferencial de primer orden sobre $\mathcal{A}$, c^T y δ el conmutador transpuesto y el encaje diferencial, respectivamente. Entonces los siguientes diagramas

$$\begin{array}{ccc} \Gamma_{\text{inv}} & \xrightarrow{\text{ad}} & \Gamma_{\text{inv}} \otimes \mathcal{A} \\ \delta \downarrow & & \downarrow \delta \otimes \text{id} \\ \Gamma_{\text{inv}} \otimes \Gamma_{\text{inv}} & \xrightarrow{\text{ad}} & \Gamma_{\text{inv}} \otimes \Gamma_{\text{inv}} \otimes \mathcal{A} \end{array} \quad \begin{array}{ccc} \Gamma_{\text{inv}} & \xrightarrow{\text{ad}} & \Gamma_{\text{inv}} \otimes \mathcal{A} \\ c^T \downarrow & & \downarrow c^T \otimes \text{id} \\ \Gamma_{\text{inv}} \otimes \Gamma_{\text{inv}} & \xrightarrow{\text{ad}} & \Gamma_{\text{inv}} \otimes \Gamma_{\text{inv}} \otimes \mathcal{A} \end{array}$$

son conmutativos. Es decir, δ y c^T son covariantes respecto a la acción adjunta.

Demostración. Sea $a \in \mathcal{L}$, entonces

$$\begin{aligned} (\delta \otimes \text{id})\text{ad}(\pi(a)) &= (\delta \otimes \text{id})(\pi(a^{(2)}) \otimes \kappa(a^{(1)})a^{(3)}) \\ &= -\pi(a^{(2)}) \otimes \pi(a^{(3)}) \otimes \kappa(a^{(1)})a^{(4)}. \end{aligned}$$

Y por otro lado,

$$\begin{aligned} \text{ad}\delta(\pi(a)) &= \text{ad}(-\pi(a^{(1)}) \otimes \pi(a^{(2)})) \\ &= -\pi(a^{(2)}) \otimes \pi(a^{(5)}) \otimes \kappa(a^{(1)})a^{(3)}\kappa(a^{(4)})a^{(6)} \\ &= -\pi(a^{(2)}) \otimes \pi(a^{(4)}) \otimes \kappa(a^{(1)})\epsilon(a^{(3)})a^{(5)} \\ &= -\pi(a^{(2)}) \otimes \pi(a^{(3)}) \otimes \kappa(a^{(1)})a^{(4)}, \end{aligned}$$

por lo tanto, el primer diagrama es conmutativo.

Por otro lado, si $\theta \in \Gamma_{\text{inv}}$ entonces

$$\begin{aligned} \text{ad}(c^T(\theta)) &= \text{ad}(\theta^{(0)} \otimes \pi(\theta^{(1)})) = \theta^{(0)} \otimes \pi(\theta^{(3)}) \otimes \theta^{(1)}\kappa(\theta^{(2)})\theta^{(4)} \\ &= \theta^{(0)} \otimes \pi(\theta^{(2)}) \otimes \epsilon(\theta^{(1)})\theta^{(3)} = \theta^{(0)} \otimes \pi(\theta^{(1)}) \otimes \theta^{(2)} \\ &= (c^T \otimes \text{id})\text{ad}(\theta), \end{aligned}$$

lo que termina la demostración. $\square$

Proposición 4.5.11. Sea Γ un cálculo diferencial de primer orden sobre $\mathcal{A}$ izquierda covariante. Sea σ el operador de intercambio definido como en (4.30), δ un encaje diferencial y c^T el conmutador transpuesto. Entonces se satisface la siguiente identidad

$$c^T = \sigma\delta - \delta. \quad (4.53)$$

Demostración. Sea $a \in \ker(\epsilon)$ tal que $\theta = \pi(a)$, entonces

$$\begin{aligned} -\sigma(\delta(\theta)) &= \sigma(\pi(a^{(1)}) \otimes \pi(a^{(2)})) = \pi(a^{(3)}) \otimes [\pi(a^{(1)}) \circ (\kappa(a^{(2)})a^{(4)})] \\ &= \pi(a^{(3)}) \otimes [\pi(a^{(1)})\kappa(a^{(2)})a^{(4)} - \epsilon(a^{(1)})\pi(\kappa(a^{(2)})a^{(4)})] \\ &= \pi(a^{(2)}) \otimes \pi(\epsilon(a^{(1)})a^{(3)}) - \pi(a^{(2)}) \otimes \pi(\kappa(a^{(1)})a^{(3)}) \\ &= \pi(a^{(1)}) \otimes \pi(a^{(2)}) - c^T(\pi(a)) = -\delta(\theta) - c^T(\theta), \end{aligned}$$

como se quería demostrar. $\square$

Observación 4.5.6. En el caso clásico, si $\theta, \eta \in \Gamma_{\text{inv}}$ se tiene que $\sigma(\theta \otimes \eta) = \eta \otimes \theta$ y además pasando del álgebra tensorial al álgebra cociente $\Gamma^\wedge$ se tiene que $\delta(\pi(a))$ se proyecta a la diferencial en $\pi(a)$ y $\sigma(\delta(\pi(a)))$ se proyecta a $-d\pi(a)$ por lo que

$$d = -\frac{1}{2}c^T.$$

Es decir, la fórmula (4.53) corresponde en el caso clásico a la ecuación de Maurer-Cartan.

Regresemos por un momento a los bimódulos extendidos y consideremos el operador de trenza $\tilde{\sigma} : \tilde{\Gamma} \rightarrow \tilde{\Gamma}$. Si $\theta \in \Gamma$ entonces por una parte como X es derecha invariante se tiene que

$$\tilde{\sigma}(\theta \otimes X) = X \otimes (\theta \circ 1) = X \otimes \theta. \quad (4.54)$$

Y por otro lado usando la ecuación (4.44) tenemos que

$$\begin{aligned} \tilde{\sigma}(X \otimes \theta) &= \theta^{(0)} \otimes (X \circ \theta^{(1)}) = \theta^{(0)} \otimes (\pi(\theta^{(1)}) + \epsilon(\theta^{(1)})X) = \\ &= c^T(\theta) + \theta \otimes X. \end{aligned} \quad (4.55)$$

El operador $\tilde{\sigma}$ satisface la ecuación de trenza y por lo tanto obsérvese que si $\theta \in \Gamma_{\text{inv}}$ entonces

$$\begin{aligned} (\tilde{\sigma} \otimes \text{id})(\text{id} \otimes \tilde{\sigma})(\tilde{\sigma} \otimes \text{id})(X \otimes X \otimes \theta) &= (\tilde{\sigma} \otimes \text{id})(\text{id} \otimes \tilde{\sigma})(X \otimes X \otimes \theta) \\ &= (\tilde{\sigma} \otimes \text{id})(X \otimes c^T(\theta) + X \otimes \theta \otimes X) \\ &= (c^T \otimes \text{id})c^T(\theta) + \theta^{(0)} \otimes X \otimes \pi(\theta^{(1)}) + c^T(\theta) \otimes X + \theta \otimes X \otimes X. \end{aligned}$$

Y por otro lado,

$$\begin{aligned}
(\text{id} \otimes \tilde{\sigma})(\tilde{\sigma} \otimes \text{id})(\tilde{\sigma} \otimes \text{id})(X \otimes X \otimes \theta) &= (\text{id} \otimes \tilde{\sigma})(\tilde{\sigma} \otimes \text{id})(X \otimes c^T(\theta) \\
&+ X \otimes \theta \otimes X) = (\text{id} \otimes \tilde{\sigma})(c^T(\theta^{(0)}) \otimes \pi(\theta^{(1)}) + \theta^{(0)} \otimes X \otimes \pi(\theta^{(1)}) \\
&+ c^T(\theta) \otimes X + \theta \otimes X \otimes X) = (\text{id} \otimes \sigma)(c^T \otimes \text{id})c^T(\theta) + (\text{id} \otimes c^T)c^T(\theta) \\
&+ c^T(\theta) \otimes X + \theta^{(0)} \otimes X \otimes \pi(\theta^{(1)}) + \theta \otimes X \otimes X.
\end{aligned}$$

Por lo que igualando ambos resultados obtenemos la ecuación

$$(c^T \otimes \text{id})c^T = (\text{id} \otimes \sigma)(c^T \otimes \text{id})c^T + (\text{id} \otimes c^T)c^T, \quad (4.56)$$

que corresponde en caso clásico a la identidad de Jacobi.

Por otro lado, con cálculos directos se tiene que la ecuación de trenza aplicada a los elementos $X \otimes \theta \otimes X$, $\theta \otimes X \otimes X$ y $\theta \otimes \eta \otimes X$ no genera nuevas relaciones, para $\theta, \eta \in \Gamma_{\text{inv}}$.

Calculemos ahora el valor de la ecuación de trenza en los elementos de la forma $X \otimes \theta \otimes \eta$, donde $\theta, \eta \in \Gamma_{\text{inv}}$:

$$\begin{aligned}
(\tilde{\sigma} \otimes \text{id})(\text{id} \otimes \tilde{\sigma})(\tilde{\sigma} \otimes \text{id})(X \otimes \theta \otimes \eta) &= (\tilde{\sigma} \otimes \text{id})(\text{id} \otimes \tilde{\sigma})(c^T(\theta) \otimes \eta + \theta \otimes X \otimes \eta) \\
&= (\tilde{\sigma} \otimes \text{id})((\text{id} \otimes \sigma)(c^T \otimes \text{id})(\theta \otimes \eta) + \theta \otimes c^T(\eta) + \theta \otimes \eta \otimes X) \\
&= (\sigma \otimes \text{id})(\text{id} \otimes \sigma)(c^T \otimes \text{id})(\theta \otimes \eta) + (\sigma \otimes \text{id})(\text{id} \otimes c^T)(\theta \otimes \eta) + \sigma(\theta \otimes \eta) \otimes X
\end{aligned}$$

y por otro lado

$$\begin{aligned}
(\text{id} \otimes \tilde{\sigma})(\tilde{\sigma} \otimes \text{id})(\tilde{\sigma} \otimes \text{id})(X \otimes \theta \otimes \eta) &= (\text{id} \otimes \tilde{\sigma})(\tilde{\sigma} \otimes \text{id})(X \otimes \sigma(\theta \otimes \eta)) \\
&= (\text{id} \otimes \tilde{\sigma})((c^T \otimes \text{id})\sigma(\theta \otimes \eta) + \eta^{(0)} \otimes X \otimes (\theta \circ \eta^{(1)})) \\
&= (\text{id} \otimes \sigma)(c^T \otimes \text{id})\sigma(\theta \otimes \eta) + (\text{id} \otimes c^T)\sigma(\theta \otimes \eta) + \sigma(\theta \otimes \eta) \otimes X,
\end{aligned}$$

lo que nos lleva igualando ambas ecuaciones a la siguiente relación

$$\begin{aligned}
(\sigma \otimes \text{id})(\text{id} \otimes \sigma)(c^T \otimes \text{id}) + (\sigma \otimes \text{id})(\text{id} \otimes c^T) &= \\
&= (\text{id} \otimes \sigma)(c^T \otimes \text{id})\sigma + (\text{id} \otimes c^T)\sigma. \quad (4.57)
\end{aligned}$$

Calculando la ecuación de trenza en los elementos de la forma $\theta \otimes X \otimes \eta$ obtenemos

$$\begin{aligned}
(\tilde{\sigma} \otimes \text{id})(\text{id} \otimes \tilde{\sigma})(\tilde{\sigma} \otimes \text{id})(\theta \otimes X \otimes \eta) &= (\tilde{\sigma} \otimes \text{id})(\text{id} \otimes \tilde{\sigma})(X \otimes \theta \otimes \eta) \\
&= (\tilde{\sigma} \otimes \text{id})(X \otimes \sigma(\theta \otimes \eta)) \\
&= (c^T \otimes \text{id})\sigma(\theta \otimes \eta) + \eta^{(0)} \otimes X \otimes (\theta \circ \eta^{(1)}),
\end{aligned}$$

y por otro lado

$$\begin{aligned} (\text{id} \otimes \tilde{\sigma})(\tilde{\sigma} \otimes \text{id})(\tilde{\sigma} \otimes \text{id})(\theta \otimes X \otimes \eta) &= (\text{id} \otimes \tilde{\sigma})(\tilde{\sigma} \otimes \text{id})(\theta \otimes c^T(\eta) \\ &+ \theta \otimes \eta \otimes X) = (\text{id} \otimes \tilde{\sigma})((\sigma \otimes \text{id})(\text{id} \otimes c^T)(\theta \otimes \eta) + \sigma(\theta \otimes \eta) \otimes X) \\ &= (\text{id} \otimes \sigma)(\sigma \otimes \text{id})(\text{id} \otimes c^T)(\theta \otimes \eta) + \eta^{(0)} \otimes X \otimes (\theta \circ \eta^{(1)}), \end{aligned}$$

lo que nos lleva a una nueva ecuación:

$$(c^T \otimes \text{id})\sigma = (\text{id} \otimes \sigma)(\sigma \otimes \text{id})(\text{id} \otimes c^T). \quad (4.58)$$

Observación 4.5.7. En el caso clásico cuando el operador de trenza es trivial las ecuaciones (4.57) y (4.58) siempre se cumplen.

Podemos definir entonces el concepto de álgebra de Lie cuántica, dualizando esta construcción y extrayendo el contenido algebraico abstracto. En otras palabras, consideraremos un espacio vectorial $\mathfrak{g}$ equipado con un operador de trenza $\sigma: \mathfrak{g} \otimes \mathfrak{g} \rightarrow \mathfrak{g} \otimes \mathfrak{g}$ y un conmutador $[\cdot, \cdot]: \mathfrak{g} \otimes \mathfrak{g} \rightarrow \mathfrak{g}$ tal que las siguientes ecuaciones se satisfacen:

$$\begin{aligned} [\cdot, \cdot](\text{id} \otimes [\cdot, \cdot]) &= [\cdot, \cdot](\text{id} \otimes [\cdot, \cdot])(\sigma \otimes \text{id}) + [\cdot, \cdot]([\cdot, \cdot] \otimes \text{id}); \\ \sigma(\text{id} \otimes [\cdot, \cdot]) &= ([\cdot, \cdot] \otimes \text{id})(\text{id} \otimes \sigma)(\sigma \otimes \text{id}); \\ (\text{id} \otimes [\cdot, \cdot])(\sigma \otimes \text{id})(\text{id} \otimes \sigma) &+ ([\cdot, \cdot] \otimes \text{id})(\text{id} \otimes \sigma) = \sigma(\text{id} \otimes [\cdot, \cdot])(\sigma \otimes \text{id}) + \sigma([\cdot, \cdot] \otimes \text{id}). \end{aligned}$$

Observación 4.5.8. Esto todavía no incluye la propiedad

$$\sum \sigma(x \otimes y) = \sum x \otimes y \quad \Rightarrow \quad \sum [x, y] = 0$$

que es la versión cuántica de la antisimetricidad del corchete de Lie.

Proposición 4.5.12. El espacio $S_{\text{inv}}^{\wedge 2} \subseteq \Gamma_{\text{inv}} \otimes \Gamma_{\text{inv}}$ está formado por todos los elementos de la forma

$$q = \pi(a^{(1)}) \otimes \pi(a^{(2)}),$$

donde $a \in \mathcal{R}$.

Demostración. El espacio $S_{\text{inv}}^{\wedge 2}$ consiste de las proyecciones izquierda invariantes de los elementos $Q = \sum da \otimes db$, donde $\sum adb = 0$.

Identificando $\Gamma^{\otimes}$ con $\mathcal{A} \otimes \Gamma_{\text{inv}}^{\otimes}$ se tiene que

$$Q = \sum a^{(1)}b^{(1)} \otimes \{(\pi(a^{(2)}) \circ b^{(2)}) \otimes \pi(b^{(3)})\}.$$

Así que

$$\begin{aligned}
 (\epsilon \otimes \text{id})(Q) &= \sum \epsilon(a^{(1)}b^{(1)}) \otimes \{(\pi(a^{(2)}) \circ b^{(2)}) \otimes \pi(b^{(3)})\} \\
 &= \sum 1 \otimes (\pi(a) \circ b^{(1)}) \otimes \pi(b^{(2)}) \\
 &\cong \sum \pi(ab^{(1)}) \otimes \pi(b^{(2)}) - \sum \epsilon(a)\pi(b^{(1)}) \otimes \pi(b^{(2)}) \\
 &= - \sum \epsilon(a)\pi(b^{(1)}) \otimes \pi(b^{(2)}) \in S_{\text{inv}}^{\wedge 2},
 \end{aligned}$$

donde hemos usado que $\sum adb = \sum ab^{(1)} \otimes \pi(b^{(2)}) = 0$. Finalmente observemos que los elementos $\sum \epsilon(a)b$ cubren todo el espacio $\ker(\pi) = \mathbb{C} \oplus \mathcal{R}$, de manera que obtenemos el resultado de la proposición. $\square$

Observación 4.5.9. El espacio $S_{\text{inv}}^{\wedge 2}$ genera a todo el ideal $S_{\text{inv}}^{\wedge}$ en $\Gamma_{\text{inv}}^{\otimes}$, es decir $\Gamma_{\text{inv}}^{\wedge}$ es un álgebra cuadrática.

Observación 4.5.10. En vista de los elementos que generan al ideal $S_{\text{inv}}^{\wedge}$, observemos que podemos considerar el mapeo $(\pi \otimes \pi)\phi$ definido en todo $\mathcal{A}$ y se cumple que

$$\begin{aligned}
 (\pi \otimes \pi)\phi(ab) &= \pi(a^{(1)}b^{(1)}) \otimes \pi(a^{(2)}b^{(2)}) \\
 &= (\pi(a^{(1)}) \circ b^{(1)}) \otimes (\pi(a^{(2)}) \circ b^{(2)}) + \epsilon(a^{(1)})\pi(b^{(1)}) \otimes (\pi(a^{(2)}) \circ b^{(2)}) \\
 &\quad + (\pi(a^{(1)}) \circ b^{(1)}) \otimes \epsilon(a^{(2)})\pi(b^{(2)}) + \epsilon(a^{(1)})\pi(b^{(1)}) \otimes \epsilon(a^{(2)})\pi(b^{(2)}) \\
 &= (\pi(a^{(1)}) \circ b^{(1)}) \otimes (\pi(a^{(2)}) \circ b^{(2)}) + \pi(b^{(1)}) \otimes (\pi(a) \circ b^{(2)}) \\
 &\quad + (\pi(a) \circ b^{(1)}) \otimes \pi(b^{(2)}) + \epsilon(a)\pi(b^{(1)}) \otimes \pi(b^{(2)}),
 \end{aligned}$$

para todo $a, b \in \mathcal{A}$.

En particular, si $a \in \mathcal{R}$ entonces

$$(\pi \otimes \pi)\phi(ab) = (\pi(a^{(1)}) \circ b^{(1)}) \otimes (\pi(a^{(2)}) \circ b^{(2)}),$$

para todo $b \in \mathcal{A}$.

Proposición 4.5.13. Si V es un espacio lineal de generadores del ideal derecho $\mathcal{R}$, entonces los elementos

$$q = (\pi(a^{(1)}) \circ b^{(1)}) \otimes (\pi(a^{(2)}) \circ b^{(2)}), \quad (4.59)$$

generan al espacio $S_{\text{inv}}^{\wedge 2}$, para todo $a \in V$ y $b \in \mathcal{A}$. $\square$

Observación 4.5.11. Sea $\Gamma^{\vee}$ el álgebra exterior trenzada construida sobre Γ . En vista de la universalidad de $\Gamma^{\wedge}$ existe un único morfismo $P : \Gamma^{\wedge} \rightarrow \Gamma^{\vee}$

de álgebras graduadas diferenciales que se reduce a la identidad tanto en Γ como en $\mathcal{A}$. En particular

$$S^{\wedge 2} \subseteq \ker(I - \sigma).$$

El mapeo P es sobre, pero generalmente no es inyectivo. Más aún, el álgebra $\Gamma^\vee$ no siempre es cuadrática.

Proposición 4.5.14. Sea Γ un cálculo diferencial de primer orden sobre un álgebra de Hopf $\mathcal{A}$ bicovariante. Sea Ω un cálculo diferencial sobre $\mathcal{A}$ de más alto orden tal que $\phi : \mathcal{A} \rightarrow \mathcal{A} \otimes \mathcal{A}$ se extiende a un morfismo de álgebras diferenciales graduadas. Entonces el cálculo universal $\Gamma^\wedge$ se proyecta sobre Ω y Ω se proyecta sobre el cálculo exterior $\Gamma^\vee$.

Además, en el caso cuando el álgebra de Hopf es clásica y el cálculo diferencial Γ también lo es, entonces $\Gamma^\wedge = \Gamma^\vee$. $\square$

Observación 4.5.12. Para que un cálculo diferencial de primer orden se pueda extender a un cálculo de más alto orden, tal que el coproducto ϕ se extiende a nivel de tal cálculo diferencial, necesariamente debe ser bicovariante, pues si el coproducto se extiende a un homomorfismo entre cálculos diferenciales $\hat{\phi}$, entonces

$$\begin{aligned} \hat{\phi}(d(a)) &= d(\phi(a)) = d(a^{(1)} \otimes a^{(2)}) = d(a^{(1)}) \otimes a^{(2)} + a^{(1)} \otimes d(a^{(2)}) \\ &= \wp(da) + \ell(da). \end{aligned}$$

Observando que el cálculo que extiende al de primer orden está generado por éste, se tiene que la extensión del coproducto es igual a la suma de las co-acciones izquierda y derecha.

Capítulo 5

Haces principales

Según el programa de Erlangen de Felix Klein, la geometría es el estudio de las propiedades de un espacio que son invariantes bajo un grupo de transformaciones. Por ejemplo la geometría Euclidiana estudia las propiedades de los cuerpos que no cambian cuando se les aplica una transformación de desplazamiento, rotación ó reflexión.

El estudio de geometrías no Euclidianas surge a principios del siglo *XIX* negando el *V* postulado de Euclides sobre las paralelas y que da lugar a nuevas geometrías como lo son la elíptica y la hiperbólica. En la geometría elíptica se tiene que por un punto exterior a una recta no existe paralela alguna y en la geometría hiperbólica se especifica que por un punto exterior a una recta hay más de una paralela.

Los conceptos de la mecánica clásica son invariantes bajo simetrías Euclidianas del espacio y tiempo. Por lo que podemos pensar que la geometría que le corresponde debe ser una geometría Euclidiana.

Sin embargo la existencia de una longitud universal para la mecánica cuántica: la longitud de Planck, acoplada con la visión de Einstein que física en su naturaleza más profunda es geometría, sugiere que debemos considerar una nueva geometría que describa a los espacios en estas escalas y que con puros conceptos geométricos podamos definir una longitud absoluta.

En la geometría Euclidiana, no es posible definir una longitud universal pues entre sus simetrías se encuentran las homotecias las cuales no distinguen lo grande de lo pequeño. Sin embargo en las geometrías no Euclidianas como lo son la elíptica y la hiperbólica se puede definir de forma geométrica una unidad de longitud. ¿Estas geometrías pueden describir completamente los fenómenos cuánticos?

En la teoría cuántica, existen algunas dificultades profundas sobre cantidades físicas donde aparecen divergencias de las cantidades físicas claves, como

energía por ejemplo. La inspección más detallada revela que estas dificultades están asociadas a la suposición de que el espacio y tiempo se comportan como una variedad clásica suave en todas las escalas, incluyendo las arbitrariamente pequeñas. Y que precisamente en las escalas de lo ultra-pequeño de Planck, debemos considerar la introducción de una nueva geometría.

Un posible y prometedor camino es romper con la suposición de que se pueden definir coordenadas y pensar que el espacio está formado de un tejido que no es posible fragmentar infinitamente. Es decir, suponer que el espacio no está formado de puntos como tales. Así que las geometrías no Euclideanas, ni las otras aún más generales que proporciona el entorno de geometría diferencial, no son aptas para modelar lo que pasa en estas escalas.

Es así como surge, a principios del siglo *XX* una nueva geometría, la cuántica, con el fin de dar vida a nuevos espacios, los *espacios cuánticos*. Esta nueva geometría recurre a los conocimientos clásicos, los transforma y profundiza, y lo que propone es explorarlos desde una óptica distinta permitiendo estos espacios cuánticos mucho más elaborados y generales, que no se pueden reducir a partes, incluyendo puntos como átomos de geometría.

La geometría diferencial tiene una larga historia. Hay diversos enfoques fundacionales. Nosotros tomaremos el enfoque de Kobayashi y Nomizu de describirla y desarrollarla a través de *haces principales*.

Los haces principales expresan una bella simbiosis entre la filosofía del programa de Erlangen, y la visión de la estructura geométrica como algo que surge del espacio, internamente, localmente. Es pensar al espacio como un ‘espacio vivo’, donde los elementos de la estructura geométrica se conectan y mueven a través de la acción del grupo.

Es muy importante también mencionar, que haces principales tienen un papel fundamental en física. En particular, nos permiten describir de una manera natural los conceptos básicos de física teórica, especialmente los que aparecen cuando consideramos teorías con simetrías locales, como las teorías calibracionales (conocidas también como las de Yang-Mills). Otro ejemplo de la inexplicable efectividad de matemáticas para la descripción de la naturaleza [21].

Usaremos la versión cuántica de haces principales desarrollada en [4, 5, 6]. Con esta teoría, conocemos cómo son las simetrías del espacio alrededor de cada ‘punto’ que ni siquiera existe. Es decir, conocemos localmente la geometría del espacio, además de describir las propiedades de éste en términos de cálculo, conexiones, curvatura y otras estructuras.

A lo largo de este capítulo, desarrollaremos a fondo el formalismo cuántico de haces principales comparándolo cuando sea apropiado, con su análogo clásico. Para tal desarrollo, mencionaremos las definiciones de geometría

clásica y después explicaremos como se traduce cada una de ellas en la geometría cuántica.

5.1 Haces principales clásicos y cuánticos

Definición 5.1.1. Sea M una variedad y G un grupo de Lie. Un haz principal $P(M, G, \Pi)$ sobre M y con grupo estructural G consiste de una variedad P y una acción derecha $P \times G \rightarrow P$ de G en P , que satisface las siguientes condiciones:

1. El grupo G actúa libremente en P , es decir, para cada $p \in P$ la transformación $G \rightarrow P$, dada por $g \mapsto pg$ es inyectiva. En otras palabras, el estabilizador en cada punto respecto de la acción derecha dada, es trivial.
2. El espacio M es el cociente de P por la relación de equivalencia inducida por la acción (donde las clases de equivalencia son las órbitas), y la proyección canónica $\Pi : P \rightarrow M$ es diferenciable.
3. El haz P es localmente trivial, es decir, existe una cubierta abierta $\mathcal{U}$ de M tal que para cada $U \in \mathcal{U}$ existe un difeomorfismo ψ_U tal que el siguiente diagrama

$$\begin{array}{ccc} P \supseteq \Pi^{-1}(U) & \xrightarrow{\psi_U} & U \times G \\ & \searrow \Pi \quad \swarrow & \\ & U \subseteq M & \end{array}$$

es conmutativo. Donde la flecha diagonal derecha es la proyección canónica. En particular podemos ver que el difeomorfismo ψ ‘preserva’ los puntos de U .

Observación 5.1.1. Las condiciones 1 y 2 de la definición anterior se pueden deducir de la condición 3.

En la versión cuántica de haz principal, consideramos en lugar de las variedades P y M apropiadas $*$ -álgebras unitarias (asociativas), y en lugar de grupo de Lie, consideramos un grupo cuántico G dado al nivel primario algebraico por un álgebra de Hopf $\mathcal{A}$ tal y como hemos explicado antes.

Definición 5.1.2. Un G -haz principal cuántico sobre un espacio cuántico M descrito por una $*$ -álgebra unital $\mathcal{V}$, es una terna $P = (\mathfrak{B}, i, F)$, donde $\mathfrak{B}$ es una $*$ -álgebra unital y $i : \mathcal{V} \rightarrow \mathfrak{B}$ y $F : \mathfrak{B} \rightarrow \mathfrak{B} \otimes \mathcal{A}$, son $*$ -homomorfismos unitales, tales que

1. El mapeo $i : \mathcal{V} \longrightarrow \mathfrak{B}$ es inyectivo y además $b \in i(\mathcal{V})$ si y sólo si, $F(b) = b \otimes 1$ para todo $b \in \mathfrak{B}$.
2. Se cumple que F es una co-acción de $\mathcal{A}$ en $\mathfrak{B}$, es decir, las siguientes igualdades se satisfacen:

$$\begin{aligned} \text{id} &= (\text{id} \otimes \epsilon)F, \\ (\text{id} \otimes \phi)F &= (F \otimes \text{id})F. \end{aligned}$$

3. El mapeo lineal $X : \mathfrak{B} \otimes \mathfrak{B} \longrightarrow \mathfrak{B} \otimes \mathcal{A}$ definido por

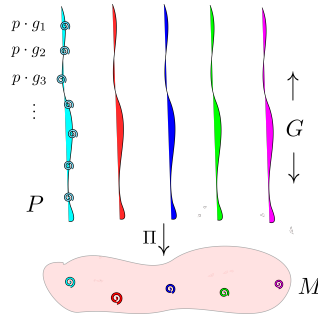
$$X(a \otimes b) = aF(b)$$

es sobre.

En la definición anterior podemos notar que en la condición (1), el homomorfismo i es la versión dualizada del mapeo Π y la condición que le continúa corresponde a la propiedad clásica de que M se puede identificar con el espacio de las órbitas de la acción derecha.

La propiedad (2) es la versión dualizada de que G actúa por la derecha en P , y por último la condición (3) nos dice que la acción es libre. Notemos que no hay una condición análoga a la de localidad trivial pues las álgebras que representan a las variedades cuánticas que constituyen el haz por lo general no poseen caracteres, es decir, los espacios no necesariamente tienen ‘puntos’ ni ‘partes’ sobre los cuales algo se puede trivializar.

Si P es un haz principal clásico, para cada punto $p \in P$, podemos definir un mapeo $\iota_p : G \rightarrow P$ como $\iota_p(g) = p \cdot g$, para todo $g \in G$. Podemos pensar que el grupo empuja por medio de la acción a los elementos de P a lo largo de las fibras y que cada fibra $\iota_p(G)$ es un conjunto vertical arriba del punto $x \in \Pi(p) \in M$.



Debido a que $\iota_p(e) = p$, entonces tenemos un mapeo inducido entre los espacios tangentes

$$(\iota_p)_* : \mathfrak{g} = T_e G \rightarrow T_p P. \quad (5.1)$$

Por lo que podemos definir un vector vertical en $T_p P$ como aquellos vectores que caen en la imagen de $(\iota_p)_*$. Denotamos al espacio de vectores verticales en p como $\text{Vert}(T_p P)$. Esto es simplemente el espacio tangente en p a la órbita de la acción derecha. Por lo que vimos

$$\text{Vert}(T_p P) = \text{Im}(\iota_p)_* = (\iota_p)_*(\mathfrak{g}).$$

Por la construcción, se satisface que $\dim[\text{Vert}(T_p P)] = \dim(G)$ y el mapeo $(\iota_p)_*$ es un isomorfismo de espacios vectoriales.

De forma equivalente se puede definir también el espacio de vectores verticales como

$$\text{Vert}(T_p P) = \ker[(\Pi_*)_p],$$

y las direcciones en el espacio base M corresponden a direcciones horizontales.

Para cada $A \in \mathfrak{g}$ se denota el vector vertical correspondiente en el punto p por

$$A^\sharp(p) := (\iota_p)_*(A) \in \text{Vert}(T_p P),$$

de manera que podemos definir el *campo vectorial fundamental en P asociado a A* como el mapeo

$$A^\sharp : P \rightarrow TP, \quad p \mapsto A^\sharp(p).$$

Definición 5.1.3. Una conexión en un haz principal clásico P es una asignación suave que a cada punto $p \in P$ le asigna un subespacio vectorial $\text{Hor}(T_p P)$ del espacio tangente $T_p P$, tal que

$$1. \text{Vert}(T_p P) \oplus \text{Hor}(T_p P) = T_p P.$$

2. La condición de G -invarianza se cumple:

$$(R_g)_*(\text{Hor}(T_p P)) = \text{Hor}(T_{pg} P),$$

para todo $g \in G$, $p \in P$ y $R_g : P \rightarrow P$, definido como $R_g(p) = pg$.

A los elementos de $\text{Hor}(T_p P)$ se les llama vectores horizontales de P en p con respecto a la conexión, la dimensión de $\text{Hor}(T_p P)$ es igual a la dimensión de M para cada punto $p \in P$ y la diferencial del mapeo Π nos da un isomorfismo de espacios vectoriales

$$\text{Hor}(T_p P) \rightarrow T_x M,$$

donde $p \in P$ y $\Pi(p) = x$.

Cada conexión de un haz principal nos define de manera natural una proyección $\text{pr}_p : T_p P \rightarrow \text{Vert}(T_p P)$. Además como el espacio de vectores verticales $\text{Vert}(T_p P)$ es isomorfo al álgebra de Lie $\mathfrak{g}$ podemos definir de manera natural un mapeo lineal

$$\omega_p : T_p P \rightarrow \mathfrak{g},$$

donde $\omega_p = (\iota_p)_*^{-1} \circ \text{pr}_p$. Éste nos define una 1-forma $\mathfrak{g}$ -valuada

$$\omega : TP \rightarrow \mathfrak{g},$$

tal que

- $\omega|_{T_p P} = \omega_p$,
- $\omega_p(A^\sharp(p)) = A$, para todo $A \in \mathfrak{g}$ y $p \in P$,
- $\ker(\omega_p) = \text{Hor}(T_p P)$.

De esta manera podemos definir conexión de Ehresmann en términos de 1-formas en un haz principal.

Definición 5.1.4. Sea P un haz principal clásico. Decimos que una conexión de Ehresmann en este haz es una forma $\mathfrak{g}$ -valuada, $\omega : TP \rightarrow \mathfrak{g}$, tal que para todo $p \in P$ se satisfacen las siguientes condiciones:

1. $\omega_p(A^\sharp(p)) = A$, para todo $A \in \mathfrak{g}$,
2. $\omega \circ (R_g)_* = \text{Ad}_{g^{-1}} \circ \omega$.

En la definición anterior $\text{Ad}_{g^{-1}} : \mathfrak{g} \rightarrow \mathfrak{g}$ es la diferencial en el elemento neutro de la acción adjunta

$$\text{ad}_{g^{-1}}(h) = g^{-1}hg, \quad g, h \in G.$$

Esta caracterización de conexión basada en covectores tangentes es justo la representación dual de la definición 5.1.3 basada en vectores tangentes. Como nos podemos dar cuenta, la teoría de conexiones clásica involucra la parte diferencial de la variedad y del grupo. Para extender esto al nivel no conmutativo, es necesario introducir la noción de cálculo diferencial sobre el haz cuántico.

Definición 5.1.5. Una $$ -álgebra diferencial graduada $\Omega(P)$ es un cálculo diferencial sobre un haz principal cuántico $P = (\mathfrak{B}, i, F)$, si se satisfacen las siguientes propiedades:*

1. Como álgebra diferencial, $\Omega(P)$ está generada por $\mathfrak{B} = \Omega^0(P)$.
2. El homomorfismo $F : \mathfrak{B} \longrightarrow \mathfrak{B} \otimes \mathcal{A}$ se extiende a un homomorfismo

$$\widehat{F} : \Omega(P) \longrightarrow \Omega(P) \widehat{\otimes} \Gamma^\wedge$$

de álgebras diferenciales graduadas.

Proposición 5.1.1. El homomorfismo $\widehat{F}$ satisface las siguientes dos propiedades:

1. $(\widehat{F} \otimes \text{id})\widehat{F} = (\text{id} \otimes \widehat{\phi})\widehat{F}$,
2. $\widehat{F}* = (* \otimes *)\widehat{F}$.

Demostración. Sea $\varphi \in \Omega^n(P)$, tal que

$$\varphi = \sum \varphi_0 d\varphi_1 \cdots d\varphi_n.$$

Por la definición de $\widehat{F}$ se tiene que

$$\begin{aligned} (\widehat{F} \otimes \text{id})\widehat{F}(\varphi) &= (\widehat{F} \otimes \text{id})\left(\sum F(\varphi_0)dF(\varphi_1) \cdots dF(\varphi_n)\right) \\ &= \sum [(F \otimes \text{id})F(\varphi_0)] d[(F \otimes \text{id})F(\varphi_1)] \cdots d[(F \otimes \text{id})F(\varphi_n)] \\ &= \sum [(\text{id} \otimes \phi)F(\varphi_0)] d[(\text{id} \otimes \phi)F(\varphi_1)] \cdots d[(\text{id} \otimes \phi)F(\varphi_n)] \\ &= \sum (\text{id} \otimes \widehat{\phi}) [F(\varphi_0)dF(\varphi_1) \cdots dF(\varphi_n)] \\ &= (\text{id} \otimes \widehat{\phi})\widehat{F}(\varphi). \end{aligned}$$

También tenemos que

$$\begin{aligned} \widehat{F}(\varphi^*) &= \widehat{F}\left(\sum (-1)^{n(n-1)/2} d\varphi_n^* \cdots d\varphi_1^* \varphi_0^*\right) \\ &= \sum (-1)^{n(n-1)/2} dF(\varphi_n^*) \cdots dF(\varphi_1^*) F(\varphi_0^*) \\ &= \sum (-1)^{n(n-1)/2} d[(*) \otimes (*)F(\varphi_n)] \cdots d[(*) \otimes (*)F(\varphi_1)] [(*) \otimes (*)F(\varphi_0)] \\ &= (* \otimes *) \left[\sum F(\varphi_0)dF(\varphi_1) \cdots dF(\varphi_n) \right] \\ &= (* \otimes *)\widehat{F}(\varphi). \end{aligned}$$

Lo que completa la demostración. $\square$

Al igual que en el caso clásico, en el caso cuántico podemos hablar del espacio horizontal y vertical del haz. Podemos pensar a los elementos del espacio vertical $\mathbf{vert}(P)$, como formas diferenciales verticalizadas en el haz.

Sea $P = (\mathfrak{B}, i, F)$ un G -haz principal cuántico sobre M . Sea Γ un $*$ -cálculo diferencial de primer orden sobre $\mathcal{A}$ bicovariante. Definimos el espacio vertical en el haz cuántico P , como el espacio vectorial graduado

$$\mathbf{ver}(P) = \mathfrak{B} \otimes \Gamma_{\text{inv}}^\wedge.$$

Este espacio tiene una estructura de $*$ -álgebra diferencial graduada, tal que está generada por $\mathfrak{B} = \mathbf{ver}^0(P)$.

Lema 5.1.2. Las siguientes fórmulas determinan en $\mathbf{ver}(P)$ la estructura de $$ -álgebra diferencial graduada:*

$$(q \otimes \eta)(b \otimes \vartheta) = qb^{(0)} \otimes (\eta \circ b^{(1)})\vartheta, \quad (5.2)$$

$$(b \otimes \vartheta)^* = b^{(0)*} \otimes (\vartheta^* \circ b^{(1)*}), \quad (5.3)$$

$$d_v(b \otimes \vartheta) = b \otimes d\vartheta + b^{(0)} \otimes \pi(b^{(1)})\vartheta, \quad (5.4)$$

donde $F(b) = b^{(0)} \otimes b^{(1)}$. □

Proposición 5.1.3. Como álgebra diferencial $\mathbf{ver}(P)$ está generada por $\mathfrak{B} = \mathbf{ver}^0(P)$. □

Proposición 5.1.4. Existe un único homomorfismo

$$\widehat{F}_v : \mathbf{ver}(P) \rightarrow \mathbf{ver}(P) \widehat{\otimes} \Gamma^\wedge$$

de álgebras diferenciales graduadas que extiende a la co-acción F y tal que

$$\begin{aligned} (\widehat{F}_v \otimes \text{id})\widehat{F}_v &= (\text{id} \otimes \widehat{\phi})\widehat{F}_v, \\ \widehat{F}_v * &= (* \otimes *)\widehat{F}_v. \end{aligned}$$

Demostración. Como $\mathbf{ver}(P)$ está generada por $\mathfrak{B}$, entonces si existe $\widehat{F}_v$ éste debe ser único.

Definimos entonces el mapeo lineal $\widehat{F}_v : \mathbf{ver}(P) \rightarrow \mathbf{ver}(P) \widehat{\otimes} \Gamma^\wedge$ como

$$\widehat{F}_v(b \otimes \vartheta) = b^{(0)} \otimes \vartheta^{(0)} \otimes b^{(1)}\vartheta^{(1)}, \quad (5.5)$$

donde $F(b) = b^{(0)} \otimes b^{(1)}$ y $\text{ad}(\vartheta) = \vartheta^{(0)} \otimes \vartheta^{(1)}$.

Es fácil ahora ver que $\widehat{F}_v$ es un homomorfismo de álgebras diferenciales que satisface las condiciones de la proposición. □

También podemos definir el álgebra que representa a las formas horizontales cuánticas y los podemos caracterizar como aquellos elementos de $\Omega(P)$ que tienen propiedades diferenciales triviales sobre las fibras verticales.

Definición 5.1.6. Supongamos $\Omega(P)$ es un cálculo diferencial sobre $\mathfrak{B}$ y $P = (\mathfrak{B}, i, F)$ un haz principal cuántico. Entonces la $*$ -subálgebra graduada

$$\mathfrak{hor}(P) = \hat{F}^{-1}(\Omega(P) \otimes \mathcal{A}),$$

es llamada el álgebra de formas horizontales. Además se cumple $\mathfrak{B} = \mathfrak{hor}^0(P)$.

Recordemos que en el caso clásico el espacio horizontal quedaba invariante por la acción del grupo. En este nuevo contexto, una vez definido cálculo diferencial sobre el haz cuántico, tenemos una propiedad análoga.

Proposición 5.1.5. La restricción de $\hat{F}$ a $\mathfrak{hor}(P) \subseteq \Omega(P)$ satisface que

$$\hat{F}(\mathfrak{hor}(P)) \subseteq \mathfrak{hor}(P) \otimes \mathcal{A}.$$

En otras palabras, la restricción de $\hat{F}$ sobre $\mathfrak{hor}(P)$ coincide con la acción $F^\wedge$, donde $F^\wedge = (\text{id} \otimes p_0)\hat{F}$ y $p_0 : \Gamma^\wedge \rightarrow \mathcal{A}$ es la proyección.

Observemos que el mapeo $F^\wedge : \Omega(P) \rightarrow \Omega(P) \otimes \mathcal{A}$ extiende a la coacción F y es un $*$ -homomorfismo que satisface

$$\begin{aligned} (F^\wedge \otimes \text{id})F^\wedge &= (\text{id} \otimes \phi)F^\wedge, \\ (\text{id} \otimes \epsilon)F^\wedge &= \text{id}, \\ (d \otimes \text{id})F^\wedge &= F^\wedge d, \end{aligned}$$

es decir, $F^\wedge$ es una coacción derecha de $\mathcal{A}$ en el espacio de formas $\Omega(P)$.

Los cálculos diferenciales sobre los haces principales cuánticos, nos dan una manera natural de equipar al espacio base de un cálculo diferencial. Éste consiste de una $*$ -subálgebra graduada diferencial $\Omega(M) \subseteq \Omega(P)$ de formas horizontales invariantes por la derecha, es decir,

$$\Omega(M) = \{\varphi \in \Omega(P) : \hat{F}(\varphi) = \varphi \otimes 1\}. \quad (5.6)$$

Definición 5.1.7. Sea P un haz principal cuántico con cálculo diferencial $\Omega(P)$. Definimos el espacio de formas pseudotensoriales en P , $\psi(P)$, como el espacio de todos los mapeos lineales $\zeta : \Gamma_{\text{inv}} \rightarrow \Omega(P)$ tal que el diagrama

$$\begin{array}{ccc} \Gamma_{\text{inv}} & \xrightarrow{\zeta} & \Omega(P) \\ \text{ad} \downarrow & & \downarrow F^\wedge \\ \Gamma_{\text{inv}} \otimes \mathcal{A} & \xrightarrow[\zeta \otimes \text{id}]{} & \Omega(P) \otimes \mathcal{A} \end{array} \quad (5.7)$$

es conmutativo.

Decimos que una forma $\zeta \in \psi(P)$ es una forma tensorial si $\zeta(\vartheta) \in \mathfrak{hor}(P)$. Al espacio de formas tensoriales en P lo denotamos por $\tau(P)$.

Observación 5.1.2. Los espacios $\psi(P)$ y $\tau(P)$ son naturalmente espacios graduados:

$$\psi(P) = \sum_{\kappa \geq 0} \psi^k(P), \quad \tau(P) = \sum_{\kappa \geq 0} \tau^k(P),$$

donde $\psi^k(P)$ consiste de mapeos lineales $\zeta : \Gamma_{\text{inv}} \rightarrow \Omega^k(P)$ y $\tau^k(P)$ de mapeos lineales $\eta : \Gamma_{\text{inv}} \rightarrow \text{hor}^k(P)$. Además cada $\tau^k(P)$ es un subespacio vectorial de $\psi^k(P)$.

Proposición 5.1.6. El espacio de formas pseudotensoriales $\psi(P)$ en P es naturalmente un bimódulo sobre el álgebra de formas invariantes por la derecha.

Demostración. Sea $\zeta \in \psi(P)$ y $\varphi \in \Omega(P)$ tal que $F^\wedge(\varphi) = \varphi \otimes 1$ entonces se cumple

$$\begin{aligned} F^\wedge(\varphi\zeta(\theta)) &= (\varphi \otimes 1)((\zeta \otimes \text{id})\text{ad}(\theta)) = (\varphi \otimes 1)(\zeta(\theta^{(0)}) \otimes \theta^{(1)}) \\ &= \varphi\zeta(\theta^{(0)}) \otimes \theta^{(1)} = (\varphi\zeta \otimes \text{id})\text{ad}(\theta), \end{aligned}$$

es decir, $\varphi\zeta \in \psi(P)$. Análogamente

$$\begin{aligned} F^\wedge(\zeta(\theta)\varphi) &= ((\zeta \otimes \text{id})\text{ad}(\theta))(\varphi \otimes 1) = (\zeta(\theta^{(0)}) \otimes \theta^{(1)})(\varphi \otimes 1) \\ &= \zeta(\theta^{(0)})\varphi \otimes \theta^{(1)} = (\zeta\varphi \otimes \text{id})\text{ad}(\theta), \end{aligned}$$

es decir, $\zeta\varphi \in \psi(P)$. Por lo tanto, $\psi(P)$ es un bimódulo sobre las formas invariantes por la derecha. $\square$

Observación 5.1.3. Dado que $F^\wedge$ y $\widehat{F}$ coinciden en $\text{hor}(P)$ se tiene también que $\psi(P)$ es un bimódulo sobre $\Omega(M)$.

5.2 Conexiones cuánticas

Definición 5.2.1. Sean Γ un cálculo diferencial de primer orden sobre $\mathcal{A}$ y $\Omega(P)$ un cálculo diferencial graduado sobre el haz $(\mathfrak{B}, i, F)$. Decimos que un mapeo lineal de primer orden $\omega : \Gamma_{\text{inv}} \rightarrow \Omega(P)$ es una conexión cuántica del haz P si y sólo si se satisfacen las siguientes dos condiciones:

1. La conexión es hermitiana, es decir, $\omega(\vartheta^*) = \omega(\vartheta)^*$;
2. Se cumple que $\widehat{F}\omega(\vartheta) = (\omega \otimes \text{id})\text{ad}(\vartheta) + 1 \otimes \vartheta$,

donde $\text{ad} : \Gamma_{\text{inv}} \rightarrow \Gamma_{\text{inv}} \otimes \mathcal{A}$ es la co-acción derecha adjunta de $\mathcal{A}$, introducida anteriormente como $\text{ad}\pi(a) = \pi(a^{(2)}) \otimes \kappa(a^{(1)})a^{(3)}$.

Denotamos por $\text{cn}(P)$ al conjunto de todas las conexiones en P .

Definición 5.2.2. Sean ω y $\tilde{\omega}$ dos conexiones cuánticas en el haz P . Denotamos a la diferencia de éstas por

$$\Delta\omega := \omega - \tilde{\omega},$$

y lo llamamos el desplazamiento (de las conexiones).

Denotamos por $\mathfrak{qcd}$ al espacio vectorial formado por todos los desplazamientos.

Observación 5.2.1. Por la definición del desplazamiento es inmediato observar que se cumple la siguiente igualdad

$$\widehat{F}[\Delta\omega(\vartheta)] = (\Delta\omega \otimes \text{id})\text{ad}(\vartheta). \quad (5.8)$$

Observación 5.2.2. Si ω es una conexión cuántica en el haz, cualquier otra conexión cuántica $\tilde{\omega}$ se puede escribir como la suma de ω con un desplazamiento $\Delta\omega$:

$$\tilde{\omega} = \omega + \Delta\omega. \quad (5.9)$$

En diversos ejemplos importantes, tanto clásicos como cuánticos, la estructura geométrica del haz promueve una conexión especial, como por ejemplo, la conexión trivial en haces triviales y la conexión de Levi-Civita en Geometría Riemanniana. En tales casos los desplazamientos se ven como una manera natural de parametrizar todas las conexiones.

Proposición 5.2.1. Si $\Delta\omega$ es un desplazamiento en el haz P , entonces

$$\Delta\omega(\vartheta) \in \mathfrak{hor}(P).$$

Demostración. Como $\Delta\omega$ es un desplazamiento, entonces se cumple

$$\widehat{F}[\Delta\omega(\vartheta)] = (\Delta\omega \otimes \text{id})\text{ad}(\vartheta),$$

para todo $\vartheta \in \Gamma_{\text{inv}}$. Notando que la imagen de la acción adjunta es un subconjunto de $\Gamma_{\text{inv}} \otimes \mathcal{A}$ se tiene que

$$\widehat{F}[\Delta\omega(\vartheta)] \in \Omega(P) \otimes \mathcal{A}.$$

Por lo tanto $\Delta\omega(\vartheta) \in \mathfrak{hor}(P)$. □

Observación 5.2.3. Por la ecuación (5.8) y la proposición 5.2.1 se tiene la conmutatividad del siguiente diagrama

$$\begin{array}{ccc} \Gamma_{\text{inv}} & \xrightarrow{\Delta\omega} & \mathfrak{hor}(P) \\ \text{ad} \downarrow & & \downarrow F^\wedge \\ \Gamma_{\text{inv}} \otimes \mathcal{A} & \xrightarrow{\Delta\omega \otimes \text{id}} & \mathfrak{hor}(P) \otimes \mathcal{A} \end{array}$$

Es decir, $\Delta\omega$ es una forma tensorial.

Proposición 5.2.2. El espacio $\mathfrak{qc}\mathfrak{d}$ de todos los desplazamientos es naturalmente un $*$ -bimódulo sobre $\mathcal{V}$.

Demostración. Sea $f \in \mathcal{V}$ y $\vartheta \in \Gamma_{\text{inv}}$, entonces $\widehat{F}(f) = f \otimes 1$ y además escribiendo $\text{ad}(\vartheta) = \vartheta^{(0)} \otimes \vartheta^{(1)}$ se tiene que

$$\begin{aligned} \widehat{F}[f\Delta\omega(\vartheta)] &= (f \otimes 1) (\Delta\omega(\vartheta^{(0)}) \otimes \vartheta^{(1)}) = f\Delta\omega(\vartheta^{(0)}) \otimes \vartheta^{(1)} \\ &= (f\Delta\omega \otimes \text{id})\text{ad}(\vartheta). \end{aligned}$$

De manera análoga tenemos que

$$\begin{aligned} \widehat{F}[\Delta\omega(\vartheta)f] &= (\Delta\omega(\vartheta^{(0)}) \otimes \vartheta^{(1)}) (f \otimes 1) = \Delta\omega(\vartheta^{(0)})f \otimes \vartheta^{(1)} \\ &= (\Delta\omega f \otimes \text{id})\text{ad}(\vartheta). \end{aligned}$$

Por lo tanto $f\Delta\omega$ y $\Delta\omega f$ también son desplazamientos.

Si $\Delta\omega$ es un desplazamiento, definimos

$$\Delta\omega^*(\vartheta) = \Delta\omega(\vartheta^*)^*. \quad (5.10)$$

Entonces

$$(f\Delta\omega)^*(\vartheta) = (f\Delta\omega(\vartheta^*))^* = \Delta\omega(\vartheta^*)^* f^* = \Delta\omega^*(\vartheta) f^* = \Delta\omega^* f^*(\vartheta),$$

es decir, $\mathfrak{qc}\mathfrak{d}$ es un $*$ -bimódulo sobre $\mathcal{V}$. □

Definición 5.2.3. La conexión cuántica ω es multiplicativa si y solamente si, para todo $a \in \mathcal{R}$

$$\omega\pi(a^{(1)})\omega\pi(a^{(2)}) = 0.$$

Cuando la conexión es multiplicativa entonces existe una única extensión, $\omega^\wedge : \Gamma_{\text{inv}}^\wedge \rightarrow \Omega(P)$, unital y multiplicativa sobre el cálculo diferencial de más alto orden del grupo.

Observación 5.2.4. En la teoría clásica, donde el cálculo diferencial tanto en el haz como en el grupo estructural es basado en formas diferenciales clásicas, todas las conexiones son multiplicativas.

Observación 5.2.5. Si ω es una conexión en el haz P , entonces el mapeo $r_\omega : \mathcal{R} \rightarrow \Omega(P)$ definido como

$$r_\omega(a) = \omega\pi(a^{(1)})\omega\pi(a^{(2)}), \quad (5.11)$$

mide la falta de multiplicatividad de la conexión.

Proposición 5.2.3. Sea ω una conexión cuántica en el haz P . El mapeo r_ω satisface las siguientes igualdades.

$$r_\omega(\kappa(a)^*) = -r_\omega(a)^*,$$

$$\widehat{F}r_\omega(a) = F^\wedge r_\omega(a) = (r_\omega \otimes \text{id})\text{ad}(a).$$

En particular, $r_\omega(a) \in \text{hor}^2(P)$, para cada $a \in \mathcal{R}$.

Demostración. Se $a \in \mathcal{R}$, entonces la primera igualdad se obtiene directamente:

$$r_\omega(a)^* = -\omega\pi(a^{(2)})^*\omega\pi(a^{(1)})^* = -\omega\pi(\kappa(a^{(2)})^*)\omega\pi(\kappa(a^{(1)})^*) = -r_\omega(\kappa(a)^*).$$

Para demostrar la segunda igualdad, sea $a \in \mathcal{R}$ entonces $\pi(a) = 0$ y además observando que se cumplen las siguientes igualdades:

$$d(\pi(b)) = -\pi(b^{(1)})\pi(b^{(2)}), \quad d(b) = b^{(1)}\pi(b^{(2)}), \quad \pi(b) = \kappa(b^{(1)})d(b^{(2)}),$$

para todo $b \in \mathcal{A}$ se tiene lo siguiente:

$$\begin{aligned} \widehat{F}r_\omega(a) &= \widehat{F}[\omega(\pi(a^{(1)}))\omega(\pi(a^{(2)}))] \\ &= [(\omega \otimes \text{id})\text{ad}(\pi(a^{(1)}))][(\omega \otimes \text{id})\text{ad}(\pi(a^{(2)}))] \\ &\quad + [1 \otimes \pi(a^{(1)})][1 \otimes \pi(a^{(2)})] + [1 \otimes \pi(a^{(1)})][(\omega \otimes \text{id})\text{ad}(\pi(a^{(2)}))] \\ &\quad + [(\omega \otimes \text{id})\text{ad}(\pi(a^{(1)}))][1 \otimes \pi(a^{(2)})] \\ &= [\omega(\pi(a^{(2)})) \otimes \kappa(a^{(1)})a^{(3)}][\omega(\pi(a^{(5)})) \otimes \kappa(a^{(4)})a^{(6)}] \\ &\quad + 1 \otimes \pi(a^{(1)})\pi(a^{(2)}) + [1 \otimes \pi(a^{(1)})][\omega(\pi(a^{(3)})) \otimes \kappa(a^{(2)})a^{(4)}] \\ &\quad + [\omega(\pi(a^{(2)})) \otimes \kappa(a^{(1)})a^{(3)}][1 \otimes \pi(a^{(4)})] \\ &= \omega(\pi(a^{(2)}))\omega(\pi(a^{(3)})) \otimes \kappa(a^{(1)})a^{(4)} - 1 \otimes d\pi(a) \\ &\quad - \omega(\pi(a^{(3)})) \otimes \pi(a^{(1)})\kappa(a^{(2)})a^{(4)} + \omega(\pi(a^{(2)})) \otimes \kappa(a^{(1)})a^{(3)}\pi(a^{(4)}) \\ &= r_\omega(a^{(2)}) \otimes \kappa(a^{(1)})a^{(4)} - \omega(\pi(a^{(4)})) \otimes \kappa(a^{(3)})\pi[a^{(1)}\kappa(a^{(2)})]a^{(5)} \\ &\quad + \omega(\pi(a^{(4)})) \otimes \kappa(a^{(3)})\pi[\epsilon(a^{(1)})\kappa(a^{(2)})]a^{(5)} + \omega(\pi(a^{(2)})) \otimes \kappa(a^{(1)})d\pi(a^{(3)}) \\ &= (r_\omega \otimes \text{id})\text{ad}(a) + \omega(\pi(a^{(3)})) \otimes \kappa(a^{(2)})\pi(\kappa(a^{(1)}))a^{(4)} \\ &\quad + \omega(\pi(a^{(2)})) \otimes d(\kappa(a^{(1)})a^{(3)}) - \omega(\pi(a^{(2)})) \otimes d(\kappa(a^{(1)}))a^{(3)} \\ &= (r_\omega \otimes \text{id})\text{ad}(a) + \omega(\pi(a^{(4)})) \otimes \kappa(a^{(3)})\kappa(\kappa(a^{(2)}))d(\kappa(a^{(1)}))a^{(5)} \\ &\quad + (\omega \otimes \text{id})(\text{id} \otimes d)\text{ad}(\pi(a)) - \omega(\pi(a^{(2)})) \otimes d(\kappa(a^{(1)}))a^{(3)} \\ &= (r_\omega \otimes \text{id})\text{ad}(a) + \omega(\pi(a^{(3)})) \otimes \epsilon(\kappa(a^{(2)}))d(\kappa(a^{(1)}))a^{(4)} \\ &\quad - \omega(\pi(a^{(2)})) \otimes d(\kappa(a^{(1)}))a^{(3)} = (r_\omega \otimes \text{id})\text{ad}(a). \end{aligned}$$

Por lo tanto, $r_\omega(a) \in \text{hor}^2(P)$, para cada $a \in \mathcal{R}$, como se quería demostrar. $\square$

En esta teoría cuántica-geométrica, las conexiones son interpretadas como formas verticales. La proposición anterior nos muestra un fenómeno totalmente cuántico, que una combinación cuadrática de formas verticales nos da como resultado una forma horizontal.

Definición 5.2.4. Una conexión cuántica ω es regular si y sólo si para todo $\vartheta \in \Gamma_{\text{inv}}$ y $\varphi \in \mathfrak{hor}(P)$ se satisface

$$\omega(\vartheta)\varphi = (-1)^{\partial\varphi}\varphi^{(0)}\omega(\vartheta \circ \varphi^{(1)}),$$

donde $F^\wedge(\varphi) = \varphi^{(0)} \otimes \varphi^{(1)}$.

Observación 5.2.6. Observemos que si ω es una conexión cuántica regular y φ es una forma diferencial en la variedad base, es decir, $F^\wedge(\varphi) = \varphi \otimes 1$ entonces

$$\omega(\vartheta)\varphi = (-1)^{\partial\varphi}\varphi\omega(\vartheta \circ 1) = (-1)^{\partial\varphi}\varphi\omega(\vartheta).$$

Es decir, las conexiones regulares conmutan graduadamente con las formas de la variedad base $\Omega(M)$.

Observación 5.2.7. En el caso de haces principales clásicos, donde el cálculo diferencial tanto en el haz como en el grupo estructural es clásico también, todas las conexiones son regulares.

Proposición 5.2.4. Sea ω una conexión cuántica regular en el haz P y m_Ω la multiplicación en $\Omega(P)$. Entonces

$$m_\Omega\{\omega \otimes \varphi\} = (-1)^{\partial\varphi}m_\Omega\{\varphi \otimes \omega\}, \quad (5.12)$$

para cada $\varphi \in \tau(P)$.

Demostración. Observemos que si $\eta \in \Gamma_{\text{inv}}$ y $\varphi \in \tau(P)$ entonces

$$\widehat{F}(\varphi(\eta)) = (\varphi \otimes \text{id})\text{ad}(\eta) = \varphi(\eta^{(0)}) \otimes \eta^{(1)}.$$

Por lo tanto si $\theta \in \Gamma_{\text{inv}}$ se tiene que

$$\begin{aligned} m_\Omega\{\varphi \otimes \omega\}\sigma(\theta \otimes \eta) &= m_\Omega\{\varphi \otimes \omega\}(\eta^{(0)} \otimes (\theta \circ \eta^{(1)})) = m_\Omega\{\varphi(\eta^{(0)}) \otimes \omega(\theta \circ \eta^{(1)})\} \\ &= \varphi(\eta^{(0)})\omega(\theta \circ \eta^{(1)}) = (-1)^{\partial\varphi}\omega(\theta)\varphi(\eta) = (-1)^{\partial\varphi}m_\Omega\{\omega \otimes \varphi\}(\theta \otimes \eta), \end{aligned}$$

concluyendo así la demostración. $\square$

Observación 5.2.8. El mapeo $\ell_\omega : \Gamma_{\text{inv}} \times \mathfrak{hor}(P) \rightarrow \Omega(P)$ definido como

$$\ell_\omega(\vartheta, \varphi) = \omega(\vartheta)\varphi - (-1)^{\partial\varphi}\varphi^{(0)}\omega(\vartheta \circ \varphi^{(1)}), \quad (5.13)$$

mide la falta de regularidad de una conexión cuántica ω . Obviamente se tiene que $\ell_\omega = 0$ si y sólo si ω es regular.

Este mapeo entonces, mide que tanto se desvía ω de ser regular.

Proposición 5.2.5. La evaluación del mapeo ℓ_ω de falta de regularidad de una conexión ω es horizontal, es decir, siempre se cumple $\ell_\omega(\vartheta, \varphi) \in \mathfrak{hor}(P)$.

Demostración. Sea $\vartheta \in \Gamma_{\text{inv}}$ y $\varphi \in \mathfrak{hor}(P)$. Escribimos $\text{ad}(\vartheta) = \vartheta^{(0)} \otimes \vartheta^{(1)}$ y $\widehat{F}(\varphi) = \varphi^{(0)} \otimes \varphi^{(1)}$ de manera que

$$\begin{aligned} \widehat{F}\ell_\omega(\vartheta, \varphi) &= \widehat{F}[\omega(\vartheta)\varphi - (-1)^{\partial\varphi}\varphi^{(0)}\omega(\vartheta \circ \varphi^{(1)})] = [(\omega \otimes \text{id})\text{ad}(\vartheta)][\varphi^{(0)} \otimes \varphi^{(1)}] \\ &\quad + [1 \otimes \vartheta][\varphi^{(0)} \otimes \varphi^{(1)}] - (-1)^{\partial\varphi}[\varphi^{(0)} \otimes \varphi^{(1)}][(\omega \otimes \text{id})\text{ad}(\vartheta \circ \varphi^{(1)})] \\ &\quad - (-1)^{\partial\varphi}\varphi^{(0)} \otimes \varphi^{(1)}(\vartheta \circ \varphi^{(1)}) \\ &= \omega(\vartheta^{(0)})\varphi^{(0)} \otimes \vartheta^{(1)}\varphi^{(1)} + (-1)^{\partial\varphi}\varphi^{(0)} \otimes \vartheta\varphi^{(1)} - (-1)^{\partial\varphi}\varphi^{(0)}\omega(\vartheta^{(0)} \circ \varphi^{(3)}) \otimes \\ &\quad \otimes \varphi^{(1)}\kappa(\varphi^{(2)})\vartheta^{(1)}\varphi^{(4)} - (-1)^{\partial\varphi}\varphi^{(0)} \otimes \varphi^{(1)}\kappa(\varphi^{(2)})\vartheta\varphi^{(3)} \\ &= \omega(\vartheta^{(0)})\varphi^{(0)} \otimes \vartheta^{(1)}\varphi^{(1)} + (-1)^{\partial\varphi}\varphi^{(0)} \otimes \vartheta\varphi^{(1)} - (-1)^{\partial\varphi}\varphi^{(0)}\omega(\vartheta^{(0)} \circ \varphi^{(1)}) \otimes \vartheta^{(1)}\varphi^{(2)} \\ &\quad - (-1)^{\partial\varphi}\varphi^{(0)} \otimes \vartheta\varphi^{(1)} = \ell_\omega(\vartheta^{(0)}, \varphi^{(0)}) \otimes \vartheta^{(1)}\varphi^{(1)}, \end{aligned}$$

donde usamos que $\text{ad}(\vartheta \circ \varphi) = (\vartheta^{(0)} \circ \varphi^{(2)}) \otimes \kappa(\varphi^{(1)})\vartheta^{(1)}\varphi^{(3)}$. Y por lo tanto podemos concluir que $\ell_\omega(\vartheta, \varphi) \in \mathfrak{hor}(P)$. $\square$

Observación 5.2.9. Cabe observar que la falta de regularidad de una conexión es un fenómeno exclusivamente cuántico en el que la combinación algebraica de una forma horizontal y una vertical da lugar a una forma horizontal.

Proposición 5.2.6. Sea ω una conexión cuántica en el haz P . Entonces

$$\ell_\omega(\vartheta, \varphi\psi) = \ell_\omega(\vartheta, \varphi)\psi + (-1)^{\partial\varphi}\varphi^{(0)}\ell_\omega(\vartheta \circ \varphi^{(1)}, \psi), \quad (5.14)$$

$$\ell_\omega(\vartheta, \varphi)^* = -\ell_\omega(\vartheta^* \circ \kappa(\varphi^{(1)})^*, \varphi^{(0)*}), \quad (5.15)$$

para todo $\vartheta \in \Gamma_{\text{inv}}$, $\varphi, \psi \in \mathfrak{hor}(P)$.

Demostración. Sean $\vartheta \in \Gamma_{\text{inv}}$ y $\varphi, \psi \in \mathfrak{hor}(P)$, entonces

$$\begin{aligned} \ell_\omega(\vartheta, \varphi)\psi + (-1)^{\partial\varphi}\varphi^{(0)}\ell_\omega(\vartheta \circ \varphi^{(1)}, \psi) &= [\omega(\vartheta)\varphi - (-1)^{\partial\varphi}\varphi^{(0)}\omega(\vartheta \circ \varphi^{(1)})]\psi \\ &\quad + (-1)^{\partial\varphi}\varphi^{(0)}[\omega(\vartheta \circ \varphi^{(1)})\psi - (-1)^{\partial\psi}\psi^{(0)}\omega(\vartheta \circ \varphi^{(1)} \circ \psi^{(1)})] = \omega(\vartheta)\varphi\psi \\ &\quad - (-1)^{\partial\varphi+\partial\psi}\psi^{(0)}\omega(\vartheta \circ (\varphi^{(1)}\psi^{(1)})) = \ell_\omega(\vartheta, \varphi\psi). \end{aligned}$$

Por otro lado, observemos que se satisface lo siguiente

$$\vartheta^* \circ \kappa(\varphi)^* = \kappa(\kappa(\varphi^{(2)})^*)\vartheta^*\kappa(\varphi^{(1)})^* = \varphi^{(2)*}\vartheta^*\kappa(\varphi^{(1)})^* = (\vartheta \circ \varphi)^*,$$

donde hemos usado que se cumple la identidad $\kappa * \kappa * = \text{id}$. Por lo tanto

$$\begin{aligned} \ell_\omega(\vartheta^* \circ \kappa(\varphi^{(1)})^*, \varphi^{(0)*}) &= \omega(\vartheta^* \circ \kappa(\varphi^{(1)})^*) \varphi^{(0)*} \\ &\quad - (-1)^{\partial\varphi} \varphi^{(0)*} \omega(\vartheta^* \circ \kappa(\varphi^{(1)})^* \circ \varphi^{(2)*}) = \omega(\vartheta \circ \varphi^{(1)})^* \varphi^{(0)*} \\ &\quad - (-1)^{\partial\varphi} \varphi^* \omega(\vartheta^*) = (-1)^{\partial\varphi} [\varphi^{(0)} \omega(\vartheta \circ \varphi^{(1)})]^* - [\omega(\vartheta) \varphi]^* \\ &= -\ell_\omega(\vartheta, \varphi)^*, \end{aligned}$$

completando así lo que se quería demostrar. $\square$

Proposición 5.2.7. Sea ω una conexión cuántica regular en el haz P , si $\varphi \in \mathfrak{hor}(P)$ entonces

$$r_\omega(a)\varphi = \varphi^{(0)} r_\omega(a\varphi^{(1)}),$$

para todo $a \in \mathcal{R}$.

Demostración. Sea $\omega \in \mathbf{rcn}$, $\varphi \in \mathfrak{hor}(P)$ y $a \in \mathcal{R}$, entonces

$$\begin{aligned} r_\omega(a)\varphi &= \omega(\pi(a^{(1)}))\omega(\pi(a^{(2)}))\varphi \\ &= (-1)^{\partial\varphi} \omega(\pi(a^{(1)}))\varphi^{(0)}\omega[\pi(a^{(2)}) \circ \varphi^{(2)}] \\ &= \varphi^{(0)}\omega[\pi(a^{(1)}) \circ \varphi^{(1)}]\omega[\pi(a^{(2)}) \circ \varphi^{(2)}] \\ &= \varphi^{(0)}\omega\pi[(a^{(1)} - \epsilon(a^{(1)}))\varphi^{(1)}]\omega\pi[(a^{(2)} - \epsilon(a^{(2)}))\varphi^{(2)}] \\ &= \varphi^{(0)}\omega\pi(a^{(1)}\varphi^{(1)})\omega\pi(a^{(2)}\varphi^{(2)}) - \varphi^{(0)}\omega\pi(\varphi^{(1)})\omega\pi(a\varphi^{(2)}) \\ &\quad - \varphi^{(0)}\omega\pi(a\varphi^{(1)})\omega\pi(\varphi^{(2)}) + \epsilon(a)\omega(\varphi^{(1)})\omega(\varphi^{(2)}) \\ &= \varphi^{(0)}\omega\pi(a^{(1)}\varphi^{(1)})\omega\pi(a^{(2)}\varphi^{(2)}) = \varphi^{(0)}r_\omega(a\varphi^{(1)}), \end{aligned}$$

donde hemos usado que $\mathcal{R} \subseteq \ker(\epsilon)$ es un ideal derecho. $\square$

5.3 Derivada covariante y curvatura

Observación 5.3.1. En el caso clásico, dado un haz principal con una conexión ω podemos definir el mapeo derivada covariante sobre el espacio de formas diferenciales. Para cada punto en el espacio tangente TP queda definido un mapeo proyección que a cada vector tangente en TP lo manda a su parte horizontal en $\text{Hor}(TP)$ asociado a la descomposición $TP = \text{Vert}(TP) \oplus \text{Hor}(TP)$ dada por la conexión ω .

Así se define la derivada covariante asociada a la conexión ω como el mapeo $D : \Omega^k(P) \rightarrow \Omega^{k+1}(P)$ tal que

$$D\theta(v_1, v_2, \dots, v_{k+1}) = d\theta(v_1^H, v_2^H, \dots, v_{k+1}^H),$$

donde $d\theta$ es la derivada exterior de la k -forma $\theta \in \Omega^k(P)$ y v_i^H es la parte horizontal de v_i .

Además por la definición de derivada covariante, y dado que el mapeo que a cada vector v le asocia su parte horizontal v^H es una proyección, se tiene que

$$D\theta(v_1, v_2, \dots, v_{k+1}) = D\theta(v_1^H, v_2^H, \dots, v_{k+1}^H),$$

de manera que es suficiente saber como actúa la derivada covariante sobre los vectores horizontales. Además se puede demostrar que $D\theta$ es una forma horizontal.

Con estas observaciones en mente, se puede definir la derivada covariante de una conexión cuántica, extendiendo la construcción clásica:

Definición 5.3.1. Si ω es una conexión cuántica en el haz y $\varphi \in \mathfrak{hor}(P)$, definimos la derivada covariante asociada a la conexión ω como el operador $D_\omega : \mathfrak{hor}(P) \rightarrow \mathfrak{hor}(P)$ dado por

$$D_\omega(\varphi) = d\varphi - (-1)^{\partial\varphi} \varphi^{(0)} \omega \pi(\varphi^{(1)}),$$

donde $F^\wedge(\varphi) = \varphi^{(0)} \otimes \varphi^{(1)}$.

En la definición anterior, en general la diferencial de una forma horizontal no es horizontal, de manera que el término restante es lo que le falta a la derivada covariante para que el resultado sea una forma horizontal, como verificaremos mediante un cálculo directo. Podemos pensar en el término $(-1)^{\partial\varphi} \varphi^{(0)} \omega \pi(\varphi^{(1)})$ como la ‘parte vertical’ de $d\varphi$.

Proposición 5.3.1. La siguiente igualdad se satisface

$$\widehat{F}D_\omega = (D_\omega \otimes \text{id})F^\wedge. \quad (5.16)$$

En particular, $D_\omega(\varphi) \in \mathfrak{hor}(P)$.

Demostración. Sea $\varphi \in \mathfrak{hor}(P)$ y $F^\wedge(\varphi) = \varphi^{(0)} \otimes \varphi^{(1)}$, entonces

$$\begin{aligned} F^\wedge(D_\omega(\varphi)) &= F^\wedge(d\varphi - (-1)^{\partial\varphi} \varphi^{(0)} \omega(\pi(\varphi^{(1)}))) \\ &= dF^\wedge(\varphi) - (-1)^{\partial\varphi} F^\wedge(\varphi^{(0)}) F^\wedge(\omega(\pi(\varphi^{(1)}))) = d(\varphi^{(0)} \otimes \varphi^{(1)}) \\ &\quad - (-1)^{\partial\varphi} (\varphi^{(0)} \otimes \varphi^{(1)}) (\omega \otimes \text{id}) \text{ad}(\pi(\varphi^{(2)})) - (-1)^{\partial\varphi} (\varphi^{(0)} \otimes \varphi^{(1)}) (1 \otimes \pi(\varphi^{(2)})) \\ &= d\varphi^{(0)} \otimes \varphi^{(1)} + (-1)^{\partial\varphi} \varphi^{(0)} \otimes d\varphi^{(1)} \\ &\quad - (-1)^{\partial\varphi} (\varphi^{(0)} \otimes \varphi^{(1)}) (\omega(\pi(\varphi^{(3)})) \otimes \kappa(\varphi^{(2)}) \varphi^{(4)}) \\ &\quad - (-1)^{\partial\varphi} (\varphi^{(0)} \otimes \varphi^{(1)}) (1 \otimes \pi(\varphi^{(2)})) \end{aligned}$$

$$\begin{aligned}
&= d\varphi^{(0)} \otimes \varphi^{(1)} + (-1)^{\partial\varphi} \varphi^{(0)} \otimes d\varphi^{(1)} - (-1)^{\partial\varphi} \varphi^{(0)} \omega(\pi(\varphi^{(2)})) \otimes \epsilon(\varphi^{(1)}) \varphi^{(3)} \\
&\quad - (-1)^{\partial\varphi} \varphi^{(0)} \otimes d\varphi^{(1)} \\
&= d\varphi^{(0)} \otimes \varphi^{(1)} - (-1)^{\partial\varphi} \varphi^{(0)} \omega(\pi(\varphi^{(1)})) \otimes \varphi^{(2)} = D_\omega(\varphi^{(0)}) \otimes \varphi^{(1)} \\
&= (D_\omega \otimes \text{id}) F^\wedge(\varphi),
\end{aligned}$$

obteniendo así lo que se quería demostrar. $\square$

Proposición 5.3.2. Sea ω una conexión cuántica en el haz P y D_ω la derivada covariante asociada a la conexión ω , entonces

$$D_\omega(\varphi\psi) = D_\omega(\varphi)\psi + (-1)^{\partial\varphi} \varphi D_\omega(\psi) + (-1)^{\partial\varphi} \varphi^{(0)} \ell_\omega(\pi(\varphi^{(1)}), \psi),$$

para todo $\varphi, \psi \in \mathfrak{hor}(P)$.

Demostración. Sean $\varphi, \psi \in \mathfrak{hor}(P)$, entonces

$$\begin{aligned}
D_\omega(\varphi\psi) &= d(\varphi\psi) - (-1)^{\partial\varphi+\partial\psi} \varphi^{(0)} \psi^{(0)} \omega\pi(\varphi^{(1)} \psi^{(1)}) \\
&= d(\varphi)\psi + (-1)^{\partial\varphi} \varphi d(\psi) - (-1)^{\partial\varphi+\partial\psi} \varphi^{(0)} \psi^{(0)} \omega[\pi(\varphi^{(1)}) \circ \psi^{(1)}] \\
&\quad - (-1)^{\partial\varphi+\partial\psi} \varphi \psi^{(0)} \omega\pi(\psi^{(1)}) \\
&= [d(\varphi) - (-1)^{\partial\varphi} \varphi^{(0)} \omega\pi(\varphi^{(1)})] \psi \\
&\quad + (-1)^{\partial\varphi} \varphi [d(\psi) - (-1)^{\partial\psi} \psi^{(0)} \omega(\psi^{(1)})] \\
&\quad + (-1)^{\partial\varphi} \varphi^{(0)} [\omega\pi(\varphi^{(1)})\psi - (-1)^{\partial\psi} \psi^{(0)} \omega[\pi(\varphi^{(1)}) \circ \psi^{(1)}]] \\
&= D_\omega(\varphi)\psi + (-1)^{\partial\varphi} \varphi D_\omega(\psi) + (-1)^{\partial\varphi} \varphi^{(0)} \ell_\omega(\pi(\varphi^{(1)}), \psi),
\end{aligned}$$

demostrando así la ecuación. $\square$

Proposición 5.3.3. Sea ω una conexión cuántica en P , si ω es regular entonces

$$D_\omega(\varphi\psi) = D_\omega(\varphi)\psi + (-1)^{\partial\varphi} \varphi D_\omega(\psi), \quad (5.17)$$

$$D_\omega(\varphi^*) = D_\omega(\varphi)^*, \quad (5.18)$$

para todo $\varphi, \psi \in \mathfrak{hor}(P)$.

Demostración. Sean $\varphi, \psi \in \mathfrak{hor}(P)$ y ω una conexión regular entonces $\ell_\omega = 0$ y por la proposición 5.3.2 obtenemos la ecuación (5.17). Para demostrar la segunda ecuación calculamos directamente:

$$\begin{aligned}
D_\omega(\varphi)^* &= d(\varphi)^* - (-1)^{\partial\varphi} [\varphi^{(0)} \omega\pi(\varphi^{(1)})]^* \\
&= d(\varphi^*) - \omega(\pi(\varphi^{(1)})^*) \varphi^{(0)*} \\
&= d(\varphi^*) + \omega\pi(\kappa(\varphi^{(1)})^*) \varphi^{(0)*} \\
&= d(\varphi^*) + (-1)^{\partial\varphi} \varphi^{(0)*} \omega(\pi(\kappa(\varphi^{(2)*})) \circ \varphi^{(1)*}) \\
&= d(\varphi^*) - (-1)^{\partial\varphi} \varphi^{(0)*} \omega(\pi(\varphi^{(1)*})) = D_\omega(\varphi^*),
\end{aligned}$$

completando así la demostración. $\square$

Proposición 5.3.4. Sea $\omega \in \mathfrak{cn}(P)$, entonces la derivada covariante coincide con la derivada inicial en el espacio de las formas diferenciales sobre la variedad base, es decir, $(d - D_\omega)(\Omega(M)) = \{0\}$.

Demostración. Observemos que si $\varphi \in \Omega(M)$ entonces $\widehat{F}(\varphi) = \varphi \otimes 1$ y por lo tanto

$$D_\omega(\varphi) = d(\varphi) - (-1)^{\partial\varphi} \varphi \omega(\pi(1)) = d(\varphi) - (-1)^{\partial\varphi} \varphi \omega(0) = d(\varphi).$$

Por lo tanto D_ω y d coinciden en $\Omega(M)$. $\square$

Definición 5.3.2. Sean $\delta, c^T : \Gamma_{\text{inv}} \rightarrow \Gamma_{\text{inv}} \otimes \Gamma_{\text{inv}}$ el encaje diferencial y el conmutador transpuesto dados en las definiciones 4.5.3 y 4.5.2, respectivamente. Sean también $\vartheta, \eta : \Gamma_{\text{inv}} \rightarrow \Omega(P)$ dos mapeos lineales. Definimos dos nuevos mapeos lineales $\langle \vartheta | \eta \rangle, [\vartheta, \eta] : \Gamma_{\text{inv}} \rightarrow \Omega(P)$ como

$$\langle \vartheta | \eta \rangle = m_\Omega(\vartheta \otimes \eta) \delta, \quad (5.19)$$

$$[\vartheta, \eta] = m_\Omega(\vartheta \otimes \eta) c^T. \quad (5.20)$$

Definición 5.3.3. Sea ω una conexión cuántica en el Haz Principal P . Definimos la curvatura de la conexión como el mapeo $R_\omega : \Gamma_{\text{inv}} \rightarrow \Omega(P)$ tal que

$$R_\omega = d\omega - \langle \omega | \omega \rangle.$$

Proposición 5.3.5. Se satisface la siguiente ecuación:

$$\widehat{F}R_\omega(\theta) = (R_\omega \otimes \text{id})\text{ad}(\theta),$$

para todo $\theta \in \Gamma_{\text{inv}}$. En particular $R_\omega(\theta) \in \mathfrak{hor}(P)$, para todo $\theta \in \Gamma_{\text{inv}}$, y por lo tanto la curvatura es una 2-forma tensorial. $\square$

Proposición 5.3.6. Sea $\varphi \in \mathfrak{hor}(P)$ y ω una conexión, entonces

$$D_\omega^2(\varphi) = -\varphi^{(0)} (d\omega\pi(\varphi^{(1)}) + \omega\pi(\varphi^{(1)})\omega\pi(\varphi^{(2)})).$$

Demostración. Sea $\varphi \in \mathfrak{hor}(P)$, observemos que

$$\widehat{F}(d(\varphi)) = d(\varphi^{(0)}) \otimes \varphi^{(1)} + \varphi^{(0)} \otimes d(\varphi^{(1)}).$$

También se tiene que

$$\widehat{F}(\varphi^{(0)}\omega(\pi(\varphi^{(1)}))) = \varphi^{(0)}\omega(\pi(\varphi^{(1)})) \otimes \varphi^{(2)} + \varphi^{(0)} \otimes d(\varphi^{(1)}).$$

Por lo tanto, calculando directamente tenemos lo siguiente:

$$\begin{aligned}
D_\omega^2(\varphi) &= D_\omega[d(\varphi) - (-1)^{\partial\varphi}\varphi^{(0)}\omega\pi(\varphi^{(1)})] \\
&= D_\omega(d(\varphi)) - (-1)^{\partial\varphi}D_\omega(\varphi^{(0)}\omega\pi(\varphi^{(1)})) \\
&= (-1)^{\partial\varphi} [d(\varphi^{(0)})\omega(\varphi^{(1)}) + (-1)^{\partial\varphi}\varphi^{(0)}\omega\pi(d(\varphi^{(1)}))] \\
&\quad - (-1)^{\partial\varphi}d(\varphi^{(0)})\omega\pi(\varphi^{(1)}) - \varphi^{(0)}d(\omega\pi(\varphi^{(1)})) \\
&\quad - [\varphi^{(0)}\omega\pi(\varphi^{(1)})\omega\pi(\varphi^{(2)}) + \varphi^{(0)}\omega\pi(d(\varphi^{(1)}))] \\
&= -\varphi^{(0)}d\omega\pi(\varphi^{(1)}) - \varphi^{(0)}\omega(\pi(\varphi^{(1)}))\omega\pi(\varphi^{(2)}) \\
&= -\varphi^{(0)} [d\omega\pi(\varphi^{(1)}) + \omega\pi(\varphi^{(1)})\omega\pi(\varphi^{(2)})],
\end{aligned}$$

lo que completa la demostración. $\square$

Observación 5.3.2. La curvatura depende implícitamente del mapeo δ y esta dependencia desaparece cuando la conexión es multiplicativa, en este caso se tiene que $R_\omega = d\omega$. Por lo tanto, usando la Proposición 5.3.6 cuando la conexión es multiplicativa, se tiene que

$$D_\omega^2(\varphi) = -\varphi^{(0)}R_\omega\pi(\varphi^{(1)}). \quad (5.21)$$

Se puede ampliar el concepto de curvatura, extendiendo el mapeo r_ω definido en la ecuación (5.11) a un mapeo (denotado de la misma forma) $r_\omega : \mathcal{A} \rightarrow \Omega^2(P)$ definido como

$$r_\omega(a) = d\omega\pi(a) + \omega\pi(a^{(1)})\omega\pi(a^{(2)}),$$

para todo $a \in \mathcal{A}$. De esta manera tendremos que

$$D_\omega^2(\varphi) = -\varphi^{(0)}r_\omega(\varphi^{(1)}).$$

Este nuevo mapeo verifica que el siguiente diagrama

$$\begin{array}{ccc}
\mathcal{A} & \xrightarrow{r_\omega} & \text{hor}^2(P) \\
\text{ad} \downarrow & & \downarrow F^\wedge \\
\mathcal{A} \otimes \mathcal{A} & \xrightarrow[r_\omega \otimes \text{id}]{} & \text{hor}^2(P) \otimes \mathcal{A}
\end{array}$$

es conmutativo. En particular r_ω exhibe la propiedad clave de tensorialidad.

Observación 5.3.3. Podemos observar varios puntos respecto a los conceptos de curvatura:

- Como hemos explicado a lo largo de la tesis, un concepto clásico se puede generalizar y enriquecer de diferentes maneras en el contexto cuántico, y la curvatura cuando no es multiplicativa es uno de ellos. La nueva generalización de curvatura r_ω sólo se hace visible en el mundo cuántico y no depende del encaje diferencial.
- Se tiene que la imagen de R_ω forma un subespacio de la imagen de r_ω y cuando la conexión es multiplicativa se relacionan mediante la fórmula

$$r_\omega = R_\omega \pi.$$

- El mapeo que mide la falta de multiplicatividad es en realidad parte de la curvatura y sólo lo podemos percibir en el contexto cuántico.

Proposición 5.3.7. Sea ω una conexión cuántica en el haz P y $r_\omega : \mathcal{A} \rightarrow \text{hor}^2(P)$ la curvatura. Entonces

$$D_\omega r_\omega(a) + \ell_\omega \{ \pi(a^{(1)}), r_\omega(a^{(2)}) \} = 0, \quad (5.22)$$

para todo $a \in \hat{\mathcal{A}}$. Una versión cuántica de la identidad de Bianchi. $\square$

Proposición 5.3.8. Sea ω una conexión cuántica en el haz P . Entonces para todo $\eta \in \tau(P)$ se satisface

$$D_\omega \eta = d\eta - (-1)^{\partial \eta} [\eta, \omega].$$

Demostración. Sea $\theta \in \Gamma_{\text{inv}}$ y $\eta \in \tau(P)$. Entonces por la tensorialidad de η se tiene que $\hat{F}(\eta(\theta)) = (\eta \otimes \text{id})\text{ad}(\theta) = \eta(\theta^{(0)}) \otimes \theta^{(1)}$ y por la definición del conmutador transpuesto también tenemos que $c^T(\theta) = \theta^{(0)} \otimes \pi(\theta^{(1)})$. Por lo tanto

$$\begin{aligned} D_\omega(\eta(\theta)) &= d(\eta(\theta)) - (-1)^{\partial \eta} \eta(\theta^{(0)}) \omega(\pi(\theta^{(1)})) \\ &= d(\eta(\theta)) - (-1)^{\partial \eta} m_\Omega(\eta \otimes \omega)(\theta^{(0)} \otimes \pi(\theta^{(1)})) = d(\eta(\theta)) - (-1)^{\partial \eta} m_\Omega(\eta \otimes \omega) c^T(\theta) \\ &= d(\eta(\theta)) - (-1)^{\partial \eta} [\eta, \omega] \theta, \end{aligned}$$

concluyendo así la demostración. $\square$

Definición 5.3.4. Sea ω una conexión cuántica en el haz P . Definimos el mapeo lineal $q_\omega : \psi(P) \rightarrow \psi(P)$ como el mapeo dado por:

$$q_\omega(\varphi) = \langle \omega | \varphi \rangle - (-1)^{\partial \varphi} \langle \varphi | \omega \rangle - (-1)^{\partial \varphi} [\varphi, \omega]. \quad (5.23)$$

Proposición 5.3.9. El mapeo q_ω satisface la siguiente igualdad

$$\hat{F}q_\omega(\varphi) = (q_\omega \otimes \text{id})F^\wedge\varphi,$$

para cada $\varphi \in \tau(P)$. En particular $q_\omega\tau(P) \subseteq \tau(P)$.

Demostración. Sea $\varphi \in \tau(P)$ y $\theta = \pi(a)$ para un $a \in \ker(\epsilon)$ apropiado. Entonces escribimos

$$\begin{aligned}\delta(\theta) &= -\pi(a^{(1)}) \otimes \pi(a^{(2)}), \\ \text{ad}(\theta) &= \theta^{(0)} \otimes \theta^{(1)} = \pi(a^{(2)}) \otimes \kappa(a^{(1)})a^{(3)}, \\ c^T(\theta) &= \theta^{(0)} \otimes \pi(\theta^{(1)}) = \pi(a^{(2)}) \otimes \pi(\kappa(a^{(1)})a^{(3)}).\end{aligned}$$

Entonces

$$\begin{aligned}\hat{F}(q_\omega(\varphi)(\theta)) &= \hat{F}[\langle \omega | \varphi \rangle \theta - (-1)^{\partial\varphi} \langle \varphi | \omega \rangle \theta - (-1)^{\partial\varphi} [\varphi, \omega] \theta] \\ &= -\hat{F}[\omega \pi(a^{(1)}) \varphi \pi(a^{(2)})] + \\ &\quad (-1)^{\partial\varphi} \hat{F}[\varphi \pi(a^{(1)}) \omega \pi(a^{(2)})] (-1)^{\partial\varphi} \hat{F}[\varphi(\theta^{(0)}) \omega \pi(\theta^{(1)})] \\ &= -[\omega \pi(a^{(2)}) \otimes \kappa(a^{(1)})a^{(3)}][\varphi \pi(a^{(5)}) \otimes \kappa(a^{(4)})a^{(6)}] \\ &\quad - [1 \otimes \pi(a^{(1)})][\varphi \pi(a^{(3)}) \otimes \kappa(a^{(2)})a^{(4)}] \\ &\quad + (-1)^{\partial\varphi} [\varphi \pi(a^{(2)}) \otimes \kappa(a^{(1)})a^{(3)}][\omega \pi(a^{(5)}) \otimes \kappa(a^{(4)})a^{(6)}] \\ &\quad + (-1)^{\partial\varphi} [\varphi \pi(a^{(2)}) \otimes \kappa(a^{(1)})a^{(3)}][1 \otimes \pi(a^{(4)})] \\ &\quad - (-1)^{\partial\varphi} [\varphi(\theta^{(0)}) \otimes \theta^{(1)}][\omega \pi(\theta^{(3)}) \otimes \kappa(\theta^{(2)})\theta^{(4)}] \\ &\quad - (-1)^{\partial\varphi} [\varphi(\theta^{(0)}) \otimes \theta^{(1)}][1 \otimes \pi(\theta^{(2)})] \\ &= -\omega \pi(a^{(2)}) \varphi \pi(a^{(3)}) \otimes \kappa(a^{(1)})a^{(4)} - (-1)^{\partial\varphi} \varphi \pi(a^{(3)}) \otimes \pi(a^{(1)})\kappa(a^{(2)})a^{(4)} \\ &\quad + (-1)^{\partial\varphi} \varphi \pi(a^{(2)}) \omega \pi(a^{(3)}) \otimes \kappa(a^{(1)})a^{(4)} + (-1)^{\partial\varphi} \varphi \pi(a^{(2)}) \otimes \kappa(a^{(1)})d(a^{(3)}) \\ &\quad - (-1)^{\partial\varphi} \varphi(\theta^{(0)}) \omega \pi(\theta^{(1)}) \otimes \theta^{(2)} - (-1)^{\partial\varphi} \varphi(\theta^{(0)}) \otimes d(\theta^{(1)}) \\ &= \{ \langle \omega | \varphi \rangle \pi(a^{(2)}) - (-1)^{\partial\varphi} \langle \varphi | \omega \rangle \pi(a^{(2)}) - (-1)^{\partial\varphi} [\varphi, \omega] \pi(a^{(2)}) \} \otimes \kappa(a^{(1)})a^{(3)} \\ &\quad + (-1)^{\partial\varphi} \varphi \pi(a^{(3)}) \otimes d(\kappa(a^{(1)}))\epsilon(a^{(2)})a^{(4)} + (-1)^{\partial\varphi} \varphi(\pi(a^{(2)})) \otimes \kappa(a^{(1)})d(a^{(3)}) \\ &\quad - (-1)^{\partial\varphi} \varphi(\pi(a^{(2)})) \otimes d(\kappa(a^{(1)})a^{(3)}) = q_\omega(\varphi)(\theta^{(0)}) \otimes \theta^{(1)},\end{aligned}$$

lo que completa la demostración. $\square$

Proposición 5.3.10. Sea ω una conexión regular en el haz P . Entonces se satisface que

$$q_\omega(\varphi) = 0,$$

para todo $\varphi \in \tau(P)$.

Demostración. Sea ω una conexión cuántica regular y $\varphi \in \tau(P)$ entonces usando la igualdad

$$m_\Omega\{\omega \otimes \varphi\} = (-1)^{\partial\varphi} m_\omega(\varphi \otimes \omega)\sigma,$$

donde σ es el operador de trenza, se tiene que si $\vartheta = \pi(a)$ y $a \in \ker(\epsilon)$ entonces

$$\begin{aligned} \langle \omega | \varphi \rangle \vartheta &= -\omega \pi(a^{(1)}) \varphi \pi(a^{(2)}) \\ &= -(-1)^{\partial\varphi} \varphi \pi(a^{(3)}) \omega (\pi(a^{(1)}) \circ \kappa(a^{(2)}) a^{(4)}) \\ &= -(-1)^{\partial\varphi} \varphi \pi(a^{(3)}) \omega (\pi(a^{(1)}) \kappa(a^{(2)}) a^{(4)} - \epsilon(a^{(1)}) \pi(\kappa(a^{(2)}) a^{(4)})) \\ &= -(-1)^{\partial\varphi} \varphi \pi(a^{(1)}) \omega \pi(a^{(2)}) + (-1)^{\partial\varphi} \varphi \pi(a^{(2)}) \omega \pi(\kappa(a^{(1)}) a^{(3)}) \\ &= (-1)^{\partial\varphi} \langle \varphi | \omega \rangle \vartheta + (-1)^{\partial\varphi} [\varphi, \omega] \vartheta, \end{aligned}$$

lo que implica que $q_\omega(\varphi) = 0$. $\square$

Observemos que aplicando la definición de curvatura dada en 5.3.3 y usando 5.3.8, entonces tenemos que para toda conexión cuántica ω en el haz P se satisface

$$\begin{aligned} D_\omega(R_\omega) - q_\omega(R_\omega) &= dR_\omega - [R_\omega, \omega] - \langle \omega | R_\omega \rangle + \langle R_\omega | \omega \rangle + [R_\omega, \omega] \\ &= d(d\omega - \langle \omega | \omega \rangle) - \langle \omega | d\omega \rangle + \langle \omega | \langle \omega | \omega \rangle \rangle + \langle d\omega | \omega \rangle - \langle \langle \omega | \omega \rangle | \omega \rangle \\ &= \langle \omega | \langle \omega | \omega \rangle \rangle - \langle \langle \omega | \omega \rangle | \omega \rangle. \end{aligned}$$

Observación 5.3.4. La identidad

$$(D_\omega - q_\omega)(R_\omega) = \langle \omega | \langle \omega | \omega \rangle \rangle - \langle \langle \omega | \omega \rangle | \omega \rangle, \quad (5.24)$$

es otra contraparte cuántica de la identidad clásica de Bianchi.

Observemos que si la conexión ω es multiplicativa entonces la ecuación (5.24) se convierte en $(D_\omega - q_\omega)(R_\omega) = 0$ y si además la conexión ω es una conexión cuántica regular entonces $q_\omega(R_\omega) = 0$, y por lo tanto se recupera la identidad clásica de Bianchi: $D_\omega(R_\omega) = 0$.

En general, en el caso cuántico, no todas las conexiones regulares son multiplicativas, de modo que si consideramos dichas conexiones la identidad (5.24) nos mostrará que la derivada covariante de la curvatura no se anulará, contrario a lo que ocurre en el caso clásico. Por lo tanto es interesante calcular dichos términos pues nos muestran un fenómeno que sólo puede ocurrir en el contexto cuántico.

Capítulo 6

Haz hiperbólico cuántico de Hopf

En este capítulo desarrollaremos de acuerdo a la teoría de S. L. Woronowicz [23], la construcción de un cálculo diferencial cuántico sobre el grupo $SU_\mu(1, 1)$, análogamente a la ya construida versión de $SU(2)$ presentada en [24]. Como veremos, hay una muy sutil diferencia de signo en las relaciones de conmutación entre los elementos del álgebra, pero que conllevan a grandes diferencias tanto geométricas como algebraicas. La geometría codificada en el grupo cuántico $SU_\mu(2)$ de Woronowicz corresponde a una geometría elíptica mientras que la de $SU_\mu(1, 1)$ a una geometría hiperbólica.

Abordaremos este ejemplo desde la visión de haces principales desarrollada a detalle en el capítulo anterior.

6.1 Grupo cuántico $SU_\mu(1, 1)$

Recordemos que el grupo clásico $U(1, 1)$ está formado por las matrices complejas T de tamaño 2×2 que satisfacen

$$T \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} T^* = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}. \quad (6.1)$$

Esta última ecuación es equivalente a la condición

$$T^* \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} T = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}.$$

El grupo $SU(1, 1)$ es un subgrupo de $U(1, 1)$ que consta de los elementos de determinante igual a 1. Este grupo puede ser descrito también equivalentemente

como

$$\mathrm{SU}(1, 1) = \left\{ \begin{pmatrix} \alpha & \bar{\gamma} \\ \gamma & \bar{\alpha} \end{pmatrix} \in \mathrm{M}_{2 \times 2}(\mathbb{C}) : |\alpha|^2 - |\gamma|^2 = 1 \right\}.$$

Es importante observar que si definimos T como la matriz

$$T = \begin{pmatrix} \alpha & \bar{\gamma} \\ \gamma & \bar{\alpha} \end{pmatrix}.$$

Entonces la condición (6.1) es equivalente a la ecuación de determinante $|\alpha|^2 - |\gamma|^2 = \det(T) = 1$.

Este es el punto inicial para construir $\mathrm{SU}_\mu(1, 1)$, la versión cuántica de $\mathrm{SU}(1, 1)$. Ésta consiste de una $*$ -álgebra $\mathfrak{B}$ generada por las entradas matriciales de

$$u = \begin{pmatrix} \alpha & \mu\gamma^* \\ \gamma & \alpha^* \end{pmatrix} \quad (6.2)$$

donde postulamos que u cumple las relaciones dadas en (6.1). Aquí el número $\mu \in [-1, 1] \setminus \{0\}$. Lo anterior implica las siguientes relaciones de conmutatividad entre los generadores de $\mathfrak{B}$:

$$\begin{aligned} \alpha\alpha^* - \mu^2\gamma\gamma^* &= 1, & \alpha^*\alpha - \gamma^*\gamma &= 1, & \gamma\gamma^* &= \gamma^*\gamma, \\ \alpha\gamma^* &= \mu\gamma^*\alpha, & \alpha^*\gamma &= \mu^{-1}\gamma\alpha^*, & \alpha\gamma &= \mu\gamma\alpha, & \alpha^*\gamma^* &= \mu^{-1}\gamma^*\alpha^*. \end{aligned} \quad (6.3)$$

Observación 6.1.1. Los elementos $\alpha^m\gamma^k\gamma^{*l}$ forman una base para el álgebra $\mathfrak{B}$, donde $m \in \mathbb{Z}$, $k, l \in \mathbb{N}$. Entendemos a las potencias negativas de α como las correspondientes potencias positivas de α^* , de esta manera $\alpha^{-1} = \alpha^*$.

Observación 6.1.2. Los puntos clásicos de $\mathrm{SU}_\mu(1, 1)$ están dados por todos los caracteres de $\mathfrak{B}$ que respetan las relaciones.

Obsérvese que si $\kappa : \mathfrak{B} \rightarrow \mathbb{C}$ es un caracter de $\mathfrak{B}$, entonces la relación

$$\kappa(\alpha)\kappa(\gamma) = \mu\kappa(\gamma)\kappa(\alpha),$$

se debe satisfacer en $\mathbb{C}$, lo cual si $\mu \neq 1$, pasa sólo si $\kappa(\alpha) = 0$ ó $\kappa(\gamma) = 0$.

Por otro lado, cualquiera de las relaciones $\alpha\alpha^* - \mu^2\gamma\gamma^* = 1 = \alpha^*\alpha - \gamma^*\gamma$ independiente del valor de μ implica que $\kappa(\alpha) = \kappa(\alpha^*) \neq 0$.

Por lo tanto $\kappa(\gamma) = 0$ y $|\kappa(\alpha)| = 1$. Esto es, si $\mu \neq 1$ entonces cada caracter está etiquetado por un punto $z \in \mathbb{C}$ del círculo unitario. Así la parte clásica de este grupo cuántico es un círculo.

Por último observemos que si $\mu = 1$, el álgebra $\mathfrak{B}$ es conmutativa y los puntos clásicos coinciden con $\mathrm{SU}(1, 1)$.

Proposición 6.1.1. El álgebra $\mathfrak{B}$ es una $$ -álgebra de Hopf, con coproducto $\phi : \mathfrak{B} \rightarrow \mathfrak{B} \otimes \mathfrak{B}$, co unidad $\epsilon : \mathfrak{B} \rightarrow \mathbb{C}$ y antípoda $\kappa : \mathfrak{B} \rightarrow \mathfrak{B}$ dados por*

$$\begin{aligned}\phi(u_{ij}) &= \sum_{k=1}^2 u_{ik} \otimes u_{kj}, & \epsilon(u_{ij}) &= \delta_{ij}, \\ \kappa(u_{ij}) &= \left\{ \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} u^* \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \right\}_{ij}.\end{aligned}$$

Ó explicitamente

$$\begin{aligned}\phi(\alpha) &= \alpha \otimes \alpha + \mu \gamma^* \otimes \gamma, & \epsilon(\alpha) &= 1, & \kappa(\alpha) &= \alpha^*, \\ \phi(\alpha^*) &= \alpha^* \otimes \alpha^* + \mu \gamma \otimes \gamma^*, & \epsilon(\alpha^*) &= 1, & \kappa(\alpha^*) &= \alpha, \\ \phi(\gamma) &= \gamma \otimes \alpha + \alpha^* \otimes \gamma, & \epsilon(\gamma) &= 0, & \kappa(\gamma) &= -\mu \gamma, \\ \phi(\gamma^*) &= \gamma^* \otimes \alpha^* + \alpha \otimes \gamma^*, & \epsilon(\gamma^*) &= 0, & \mu \kappa(\gamma^*) &= -\gamma^*. \quad \square\end{aligned}$$

6.1.1 Cálculo diferencial cuántico de primer orden sobre $\mathfrak{B} = SU_\mu(1, 1)$

Usando la teoría de Woronowicz expuesta en el capítulo 4, construimos un cálculo diferencial (Ψ, d) de primer orden sobre $\mathfrak{B} = SU_\mu(1, 1)$. Este cálculo diferencial se puede ver como el bimódulo $\mathfrak{B} \otimes \Psi_{\text{inv}}$, donde $\Psi_{\text{inv}} \cong \ker(\epsilon)/\mathcal{J}$, para un ideal derecho $\mathcal{J}$.

Para fines de esta tesis consideraremos el ideal:

$$\begin{aligned}\langle \gamma^2, \gamma^{*2}, \gamma\gamma^*, \alpha\gamma - \gamma, \alpha\gamma^* - \gamma^*, \alpha^*\gamma - \gamma, \alpha^*\gamma^* - \gamma^*, \\ \mu^2\alpha + \alpha^* - (1 + \mu^2)1 \rangle = \mathcal{J}\end{aligned}\quad (6.4)$$

y exploraremos cómo es dicho cálculo diferencial.

Recordemos que el mapeo de gérmenes $\pi : \mathfrak{B} \rightarrow \Psi_{\text{inv}}$ es un mapeo sobre, es decir, recuperamos el espacio Ψ_{inv} mediante la fórmula $\pi(b) = \kappa(b^{(1)})db^{(2)}$, para cada $b \in \mathfrak{B}$. Además, se satisface que $\pi(b) = 0$ para todo $b \in \mathcal{J}$, lo cual nos lleva a las siguientes igualdades en Ψ_{inv} :

$$\begin{aligned}\pi(\gamma^2) &= 0, & \pi(\alpha\gamma) &= \pi(\gamma), & \pi(\alpha^*\gamma^*) &= \pi(\gamma^*), \\ \pi(\gamma^{*2}) &= 0, & \pi(\alpha\gamma^*) &= \pi(\gamma^*), & \pi(\alpha^*) &= -\mu^2\pi(\alpha), \\ \pi(\gamma\gamma^*) &= 0, & \pi(\alpha^*\gamma) &= \pi(\gamma), & \mu^2\pi(\alpha^2) &= (1 + \mu^2)\pi(\alpha).\end{aligned}$$

Proposición 6.1.2. Los elementos

$$\eta_+ = \pi(\gamma), \quad \eta_- = \mu\pi(\gamma^*), \quad \eta_3 = \pi(\alpha - \alpha^*), \quad (6.5)$$

generan el espacio Ψ_{inv} . $\square$

El espacio Ψ_{inv} tiene estructura de módulo derecho definiendo la acción derecha $\circ$ como en (4.13). Observando además que se satisface la ecuación

$$\pi(ab) = \pi(a) \circ b + \epsilon(a)\pi(b),$$

para todo $a, b \in \mathfrak{B}$, tenemos la acción derecha de $\mathfrak{B}$ sobre Ψ_{inv} dada por las siguientes ecuaciones:

$$\begin{aligned} \pi(\alpha) \circ \alpha &= \mu^{-2}\pi(\alpha), & \pi(\alpha) \circ \alpha^* &= \mu^2\pi(\alpha), & \pi(\alpha) \circ \{\gamma, \gamma^*\} &= 0, \\ \pi(\alpha^*) \circ \alpha &= \mu^{-2}\pi(\alpha^*), & \pi(\alpha^*) \circ \alpha^* &= \mu^2\pi(\alpha^*), & \pi(\alpha^*) \circ \{\gamma, \gamma^*\} &= 0, \\ \pi(\gamma) \circ \alpha &= \mu^{-1}\pi(\gamma), & \pi(\gamma) \circ \alpha^* &= \mu\pi(\gamma), & \pi(\gamma) \circ \{\gamma, \gamma^*\} &= 0, \\ \pi(\gamma^*) \circ \alpha &= \mu^{-1}\pi(\gamma^*), & \pi(\gamma^*) \circ \alpha^* &= \mu\pi(\gamma^*), & \pi(\gamma^*) \circ \{\gamma, \gamma^*\} &= 0. \end{aligned}$$

Proposición 6.1.3. La acción $\circ : \Psi_{\text{inv}} \times \mathfrak{B} \rightarrow \Psi_{\text{inv}}$ dada por:

$$\begin{aligned} \eta_{\pm} \circ \alpha &= \mu^{-1}\eta_{\pm}, & \eta_3 \circ \alpha &= \mu^{-2}\eta_3, \\ \eta_{\pm} \circ \alpha^* &= \mu\eta_{\pm}, & \eta_3 \circ \alpha^* &= \mu^2\eta_3, \\ \eta_{\pm} \circ \{\gamma, \gamma^*\} &= 0, & \eta_3 \circ \{\gamma, \gamma^*\} &= 0. \end{aligned} \tag{6.6}$$

es una acción derecha de $\mathfrak{B}$ en el espacio Ψ_{inv} . □

Proposición 6.1.4. Sea (Ψ, d) el cálculo diferencial de primer orden sobre $\mathfrak{B}$ determinado por el ideal $\mathcal{I}$ definido anteriormente. Las siguientes reglas de conmutación se satisfacen

$$\begin{aligned} \eta_3 \alpha &= \mu^{-2}\alpha \eta_3, & \eta_3 \alpha^* &= \mu^2\alpha^* \eta_3, & \eta_3 \gamma &= \mu^{-2}\gamma \eta_3, & \eta_3 \gamma^* &= \mu^2\gamma^* \eta_3, \\ \eta_{\pm} \alpha &= \mu^{-1}\alpha \eta_{\pm}, & \eta_{\pm} \alpha^* &= \mu\alpha^* \eta_{\pm}, & \eta_{\pm} \gamma &= \mu^{-1}\gamma \eta_{\pm}, & \eta_{\pm} \gamma^* &= \mu\gamma^* \eta_{\pm}. \end{aligned}$$

Demostración. El cálculo diferencial (Ψ, d) de primer orden sobre $\mathfrak{B}$ se identifica con el cálculo diferencial $(\mathfrak{B} \otimes \Psi_{\text{inv}}, d)$ mediante el isomorfismo $i : \mathfrak{B} \otimes \Psi_{\text{inv}} \rightarrow \Psi$ dado por $i(b \otimes \theta) = b\theta$, para todo $b \in \mathfrak{B}$ y $\theta \in \Psi_{\text{inv}}$ (ver el Teorema 4.2.14). De esta manera, tenemos que la multiplicación de $\theta \in \Psi$ con un elemento $b \in \mathfrak{B}$ queda determinada por la siguiente ecuación:

$$\theta b = b^{(1)}(\theta \circ b^{(2)}).$$

De esta manera los cálculos directos nos llevan a las reglas de conmutación en Ψ establecidas en la proposición. □

Proposición 6.1.5. El espacio Ψ_{inv} tiene estructura de $*$ -módulo definiendo en los generadores las siguientes igualdades:

$$\eta_3^* = -\eta_3, \quad \eta_{\pm}^* = \eta_{\mp}. \tag{6.7}$$

Demostración. La estructura $*$ queda determinada mediante la fórmula $\pi(a)^* = -\pi[\kappa(a)^*]$ como se puede ver en (4.29). De esta manera los cálculos directos nos llevan a

$$\pi(\alpha)^* = -\pi(\alpha), \quad \pi(\alpha^*)^* = -\pi(\alpha^*), \quad \pi(\gamma)^* = \mu\pi(\gamma^*), \quad \pi(\gamma^*)^* = \mu^{-1}\pi(\gamma).$$

Por la definición de $\eta_\pm$ y η_3 se tiene el resultado. $\square$

Proposición 6.1.6. La diferencial $d : \mathfrak{B} \rightarrow \Psi$ queda determinada mediante la siguientes fórmulas

$$d\alpha = \frac{1}{1+\mu^2}\alpha\eta_3 + \mu\gamma^*\eta_+, \quad d\alpha^* = -\frac{\mu^2}{1+\mu^2}\alpha^*\eta_3 + \gamma\eta_-, \quad (6.8)$$

$$d\gamma = \frac{1}{1+\mu^2}\gamma\eta_3 + \alpha^*\eta_+, \quad d\gamma^* = -\frac{\mu^2}{1+\mu^2}\gamma^*\eta_3 + \mu^{-1}\alpha\eta_-. \quad (6.9)$$

Demostración. Observemos que la diferencial d en $\mathfrak{B} \otimes \Psi_{\text{inv}}$ está dada por la fórmula $db = b^{(1)} \otimes \pi(b^{(2)})$, por lo que $d : \mathfrak{B} \rightarrow \Psi$ queda determinada por la fórmula

$$db = b^{(1)}\pi(b^{(2)}),$$

para todo $b \in \mathfrak{B}$. Esta fórmula directamente nos lleva a los resultados. $\square$

Observación 6.1.3. Podemos escribir a los generadores del espacio Ψ_{inv} en término de las diferenciales de la siguiente manera

$$\begin{aligned} \eta_3 &= \alpha^*d\alpha - \alpha d\alpha^* + \mu^2\gamma d\gamma^* - \gamma^*d\gamma, \\ \eta_+ &= \alpha d\gamma - \mu\gamma d\alpha, \quad \eta_- = \mu\alpha^*d\gamma^* - \gamma^*d\alpha^*. \end{aligned}$$

Observación 6.1.4. Observemos que el ideal $\mathcal{J}$ se obtiene de multiplicar por la derecha a los generadores de $\mathcal{J}$ por todos los elementos de $\mathfrak{B}$ y considerar sumas finitas de éstos. Observemos que si $g \in \mathcal{J}$ es un generador del ideal, y $b \in \mathfrak{B}$ entonces dado que $\mathfrak{B}$ es un álgebra de Hopf, se tiene que la antípoda κ es un mapeo lineal antimultiplicativo y por lo tanto

$$\kappa(gb)^* = (\kappa(b)\kappa(g))^* = \kappa(g)^*\kappa(b)^*.$$

Además como $\kappa(g)^* \in \mathcal{J}$ entonces $\kappa(gb)^* \in \mathcal{J}$. Por lo tanto el cálculo diferencial Ψ es un $*$ -cálculo.

6.1.2 Envolvente universal diferencial sobre $SU_\mu(1, 1)$

Proposición 6.1.7. Las relaciones de conmutación entre los generadores de Ψ_{inv} quedan determinadas de la siguiente manera:

$$\begin{aligned} \eta_+\eta_- &= -\mu^2\eta_-\eta_+, & \eta_+\eta_3 &= -\mu^4\eta_3\eta_+, \\ \eta_3\eta_- &= -\mu^4\eta_-\eta_3, & \eta_+^2 &= \eta_-^2 = \eta_3^2 = 0. \end{aligned}$$

Demostración. Recordemos que el espacio $\Psi_{\text{inv}}^\wedge = \Psi_{\text{inv}}^\otimes / S_{\text{inv}}^{\wedge 2}$, donde $S_{\text{inv}}^{\wedge 2}$ es el ideal cuadrático generado por los elementos de la forma

$$q = \pi(a^{(1)}) \otimes \pi(a^{(2)}),$$

donde $a \in \mathcal{J}$. Además usando la Proposición 4.5.13, el ideal $S_{\text{inv}}^{\wedge 2}$ está generado por todos los elementos de la forma

$$(\pi(a^{(1)}) \circ b^{(1)}) \otimes (\pi(a^{(2)}) \circ b^{(2)}),$$

donde los elementos a son generadores del ideal $\mathcal{J}$ y $b \in \mathfrak{B}$. Observando también que la acción derecha solo genera potencias de μ entonces todas las relaciones en Ψ_{inv} son cubiertas al considerar los elementos q obtenidos a partir de los generadores de $\mathcal{J}$.

Por lo tanto, como $\gamma^2 \in \mathcal{J}$ y

$$\phi(\gamma^2) = \gamma^2 \otimes \alpha^2 + \alpha^* \gamma \otimes \gamma \alpha + \gamma \alpha^* \otimes \alpha \gamma + \alpha^{*2} \otimes \gamma^2,$$

entonces

$$0 = \pi(\alpha^* \gamma) \pi(\gamma \alpha) + \pi(\gamma \alpha^*) \pi(\alpha \gamma) = \mu^{-1} \pi(\gamma)^2 + \mu \pi(\gamma)^2 = (\mu^{-1} + \mu) \eta_+^2.$$

Por lo tanto, $\eta_+^2 = 0$. De manera análoga, usando que $\gamma^{*2} \in \mathcal{J}$ se demuestra que $\eta_-^2 = 0$.

Ahora consideremos $\gamma \gamma^* \in \mathcal{J}$, su coproducto está dado por

$$\phi(\gamma \gamma^*) = \gamma \alpha \otimes \alpha \gamma^* + \alpha^* \alpha \otimes \gamma \gamma^* + \gamma \gamma^* \otimes \alpha \alpha^* + \alpha^* \gamma^* \otimes \gamma \alpha^*.$$

Así que

$$\begin{aligned} 0 &= \pi(\gamma \alpha) \pi(\alpha \gamma^*) + \pi(\alpha^* \gamma^*) \pi(\gamma \alpha^*) = \mu^{-1} \pi(\gamma) \pi(\gamma^*) + \mu \pi(\gamma^*) \pi(\gamma) \\ &= \mu^{-2} \eta_+ \eta_- + \eta_- \eta_+, \end{aligned}$$

es decir $\eta_+ \eta_- = -\mu^2 \eta_- \eta_+$.

Sea $\alpha \gamma - \gamma \in \mathcal{J}$ entonces tenemos el coproducto de dicho elemento dado por la fórmula

$$\phi(\alpha \gamma - \gamma) = \alpha \gamma \otimes \alpha^2 + \mu \gamma^* \gamma \otimes \gamma \alpha + \alpha \alpha^* \otimes \alpha \gamma + \mu \gamma^* \alpha^* \otimes \gamma^2 - \gamma \otimes \alpha - \alpha^* \otimes \gamma.$$

Por lo que

$$\begin{aligned} 0 &= \pi(\alpha \gamma) \pi(\alpha^2) - \pi(\gamma) \pi(\alpha) - \pi(\alpha^*) \pi(\gamma) = \mu^{-2} \eta_+ \eta_3 - \frac{1}{1 + \mu^2} \eta_+ \eta_3 \\ &+ \frac{\mu^2}{1 + \mu^2} \eta_3 \eta_+ = \frac{\mu^{-2}}{1 + \mu^2} \eta_+ \eta_3 + \frac{\mu^2}{1 + \mu^2} \eta_3 \eta_+, \end{aligned}$$

es decir, $\eta_+\eta_3 = -\mu^4\eta_3\eta_+$. Observemos que si aplicamos la $*$ a esta última igualdad se tiene que $\eta_3\eta_- = -\mu^4\eta_-\eta_3$.

Finalmente sea $\mu^2\alpha + \alpha^* - (1 + \mu^2)1 \in \mathcal{J}$, consideremos el coproducto de tal elemento:

$$\phi(\mu^2\alpha + \alpha^* - (1 + \mu^2)1) = \mu^2\alpha \otimes \alpha + \mu^3\gamma^* \otimes \gamma + \alpha^* \otimes \alpha^* + \mu\gamma \otimes \gamma^* - (1 + \mu^2)1 \otimes 1.$$

Entonces

$$\begin{aligned} 0 &= \mu^2\pi(\alpha)^2 + \mu^3\pi(\gamma^*)\pi(\gamma) + \pi(\alpha^*)^2 + \mu\pi(\gamma)\pi(\gamma^*) \\ &= \frac{\mu^2}{(1 + \mu^2)^2}\eta_3^2 + \mu^2\eta_-\eta_+ + \frac{\mu^4}{(1 + \mu^2)^2}\eta_3^2 + \eta_+\eta_- \\ &= \frac{\mu^2 + \mu^4}{(1 + \mu^2)^2}\eta_3^2, \end{aligned}$$

lo que implica que $\eta_3^2 = 0$, como se quería demostrar. $\square$

Proposición 6.1.8. El cálculo diferencial Ψ de primer orden sobre el grupo cuántico $SU_\mu(1, 1)$ se extiende al cálculo envolvente universal $\Psi^\wedge$ definiendo las diferenciales en los generadores de la siguiente manera

$$\begin{aligned} d\eta_3 &= -(1 + \mu^2)\eta_-\eta_+, \\ d\eta_+ &= \mu^2\eta_3\eta_+, \\ d\eta_- &= \mu^2\eta_-\eta_3. \end{aligned}$$

Demostración. La teoría general nos muestra que si el cálculo diferencial se extiende entonces se cumple que $d\pi(b) = -\pi(b^{(1)})\pi(b^{(2)})$, para todo $b \in \mathfrak{B}$. Por lo tanto,

$$d\pi(\alpha) = -[\pi(\alpha)^2 + \mu\pi(\gamma^*)\pi(\gamma)] = -\left[\frac{1}{(1 + \mu^2)^2}\eta_3^2 + \eta_-\eta_+\right] = -\eta_-\eta_+.$$

Análogamente

$$d\pi(\alpha^*) = -[\pi(\alpha^*)^2 + \mu\pi(\gamma)\pi(\gamma^*)] = -\left[\frac{\mu^4}{(1 + \mu^2)^2}\eta_3^2 + \eta_+\eta_-\right] = -\eta_+\eta_-.$$

Por lo tanto

$$d\eta_3 = d\pi(\alpha) - d\pi(\alpha^*) = -\eta_-\eta_+ + \eta_+\eta_- = -(1 + \mu^2)\eta_-\eta_+.$$

Calculemos ahora la diferencial en η_+ :

$$\begin{aligned} d\eta_+ &= d\pi(\gamma) = -[\pi(\gamma)\pi(\alpha) + \pi(\alpha^*)\pi(\gamma)] = \frac{-1}{1 + \mu^2}\eta_+\eta_3 + \frac{\mu^2}{1 + \mu^2}\eta_3\eta_+ \\ &= \frac{\mu^4 + \mu^2}{1 + \mu^2}\eta_3\eta_+ = \mu^2\eta_3\eta_+. \end{aligned}$$

Finalmente obtenemos la última ecuación al aplicar la $*$ a la forma $d\eta_+$ y observar que el cálculo es $*$ -cálculo diferencial graduado. $\square$

6.1.3 Cálculo diferencial sobre el círculo $U(1)$

Dado que $U(1) := \mathcal{A}$ es la parte clásica de $SU_\mu(1, 1)$ cuando $\mu \neq 1$, se puede construir un cálculo diferencial sobre $\mathcal{A}$ de manera natural.

Proposición 6.1.9. Si definimos el co-producto, co-unidad y antípoda como

$$\begin{aligned}\phi_{\mathcal{A}}(U) &= U \otimes U, & \kappa_{\mathcal{A}}(U) &= U^{-1}, & \epsilon_{\mathcal{A}}(U) &= 1 \\ \phi_{\mathcal{A}}(U^{-1}) &= U^{-1} \otimes U^{-1}, & \kappa_{\mathcal{A}}(U^{-1}) &= U, & \epsilon_{\mathcal{A}}(U^{-1}) &= 1.\end{aligned}$$

Entonces el álgebra $\mathcal{A}$ se convierte en un álgebra de Hopf. $\square$

Consideremos la proyección natural $\Pi : \mathfrak{B} \rightarrow \mathcal{A}$, es decir

$$\Pi(\alpha) = U, \quad \Pi(\alpha^*) = U^{-1}, \quad \Pi(\gamma) = 0, \quad \Pi(\gamma^*) = 0,$$

y proyectemos el ideal $\mathcal{J}$ del cálculo sobre $\mathfrak{B}$ al ideal $\mathcal{J}_{\mathcal{A}} = \Pi(\mathcal{J})$. Entonces obtenemos de manera natural un cálculo diferencial Γ sobre $\mathcal{A}$.

Observación 6.1.5. El cálculo diferencial Γ sobre $\mathcal{A} = U(1)$ se identifica con el bimódulo $\mathcal{A} \otimes \Gamma_{\text{inv}}$, donde el espacio Γ_{inv} de los izquierdas invariantes es isomorfo a $\ker(\epsilon)/\mathcal{J}_{\mathcal{A}}$. En este caso, tenemos directamente que

$$\mathcal{J}_{\mathcal{A}} = \{\mu^2 U + U^{-1} - (1 + \mu^2)1\}.$$

Observación 6.1.6. Se cumplen las siguientes relaciones en Γ_{inv} :

$$\mu^2 \pi_{\mathcal{A}}(U^2) = (1 + \mu^2) \pi_{\mathcal{A}}(U), \quad \pi_{\mathcal{A}}(U^{-1}) = -\mu^2 \pi_{\mathcal{A}}(U).$$

Proposición 6.1.10. El elemento

$$\chi = \pi_{\mathcal{A}}(U - U^{-1})$$

genera al espacio Γ_{inv} . $\square$

Observación 6.1.7. El elemento χ^* es antihermitiano, es decir, $\chi^* = -\chi$.

Proposición 6.1.11. La acción derecha de $\mathcal{A}$ en el espacio Γ_{inv} está definida como

$$\chi \circ U = \mu^{-2} \chi, \quad \chi \circ U^{-1} = \mu^2 \chi. \quad (6.10)$$

Demostración. Observemos que se satisfacen las siguientes igualdades

$$\begin{aligned}\pi_{\mathcal{A}}(U) \circ U &= \mu^{-2} \pi_{\mathcal{A}}(U), & \pi_{\mathcal{A}}(U^{-1}) \circ U &= \mu^{-2} \pi_{\mathcal{A}}(U^{-1}), \\ \pi_{\mathcal{A}}(U) \circ U^{-1} &= \mu^2 \pi_{\mathcal{A}}(U), & \pi_{\mathcal{A}}(U^{-1}) \circ U^{-1} &= \mu^2 \pi_{\mathcal{A}}(U^{-1}),\end{aligned}$$

de manera que aplicando la acción derecha de U y U^{-1} en el generador χ obtenemos el resultado. $\square$

Proposición 6.1.12. Sea (Γ, d) el cálculo diferencial de primer orden sobre $\mathcal{A}$ determinado por el ideal $\mathcal{I}_{\mathcal{A}}$. Entonces las siguientes reglas de conmutación se satisfacen:

$$\chi U = \mu^{-2} U \chi, \quad \chi U^{-1} = \mu^2 U^{-1} \chi.$$

Demostración. Las relaciones se obtienen al identificar $\mathcal{A} \otimes \Gamma_{\text{inv}}$ con Γ mediante el isomorfismo $a \otimes \theta \mapsto a\theta$, para todo $a \in \mathcal{A}$ y $\theta \in \Gamma_{\text{inv}}$. Por lo que se cumple la fórmula $\chi \circ a = a^{(1)}(\chi \circ a^{(2)})$, para todo $a \in \mathcal{A}$. $\square$

Proposición 6.1.13. El mapeo lineal $d : \mathcal{A} \rightarrow \Gamma$ definido en los generadores como

$$dU = \frac{1}{1 + \mu^2} U \chi, \quad dU^{-1} = -\frac{\mu^2}{1 + \mu^2} U^{-1} \chi,$$

es la diferencial en el álgebra $\mathcal{A}$. $\square$

Observación 6.1.8. Primeramente observemos que $(dU)^* = dU^{-1}$ pues

$$(dU)^* = \frac{1}{1 + \mu^2} \chi^* U^* = \frac{-1}{1 + \mu^2} \chi U^{-1} = \frac{-\mu^2}{1 + \mu^2} U^{-1} \chi = dU^{-1}.$$

También se puede observar que este cálculo diferencial sobre el círculo es no conmutativo:

$$d(U)U = \mu^{-2} U d(U), \quad d(U)U^{-1} = \mu^2 U^{-1} d(U).$$

Y finalmente, observemos también que la extensión a más alto orden del cálculo Γ es igual a cero.

6.2 Haz principal cuántico hiperbólico

Seguimos analizando al grupo cuántico $\text{SU}_{\mu}(1, 1)$, ahora por medio de la teoría de haces principales, donde el espacio total será $\text{SU}_{\mu}(1, 1)$ representado por el álgebra $\mathfrak{B}$ y el grupo estructural será el círculo unitario $\text{U}(1)$ representado por el álgebra $\mathcal{A}$, introducidas anteriormente. Como veremos este haz tiene como variedad base un hiperboloide cuántico, razón por la cuál lo llamamos haz hiperbólico cuántico (de Hopf). Estudiaremos las principales estructuras geométricas como lo son la conexión, derivada covariante y curvatura.

Proposición 6.2.1. El mapeo lineal $F : \mathfrak{B} \rightarrow \mathfrak{B} \otimes \mathcal{A}$ dado por

$$F(\alpha) = \alpha \otimes U, \quad F(\alpha^*) = \alpha^* \otimes U^*, \quad F(\gamma) = \gamma \otimes U, \quad F(\gamma^*) = \gamma^* \otimes U^*, \quad (6.11)$$

es una co-acción derecha de $\mathcal{A}$ en $\mathfrak{B}$.

Demostración. El mapeo F se define de manera general como $F = (\text{id} \otimes \Pi)\phi$. Cabe mencionar que F es un $*$ -homomorfismo unital, pues la proyección Π , el coproducto ϕ y la identidad son $*$ -homomorfismos.

Por lo tanto basta ver que se cumplen las propiedades de co-acción en los generadores α y γ . En efecto, calcularemos únicamente para α pues la acción en γ se comporta de forma similar. Así que

$$\begin{aligned}(F \otimes \text{id})F(\alpha) &= F(\alpha) \otimes U = \alpha \otimes U \otimes U = \alpha \otimes \phi_{\mathcal{A}}(U) = (\text{id} \otimes \phi_{\mathcal{A}})F(\alpha), \\ (\text{id} \otimes \epsilon_{\mathcal{A}})F(\alpha) &= \alpha \otimes \epsilon_{\mathcal{A}}(U) = \alpha \otimes 1 \cong \alpha.\end{aligned}$$

Por lo tanto el $*$ -homomorfismo F es una co-acción derecha de $\mathcal{A}$ en $\mathfrak{B}$. $\square$

Proposición 6.2.2. La $*$ -álgebra $\mathcal{V} = \{b \in \mathfrak{B} : F(b) = b \otimes 1\}$ está generada por lo elementos

$$\rho = \gamma\gamma^*, \quad \eta = \alpha\gamma^*, \quad \eta^* = \gamma\alpha^*.$$

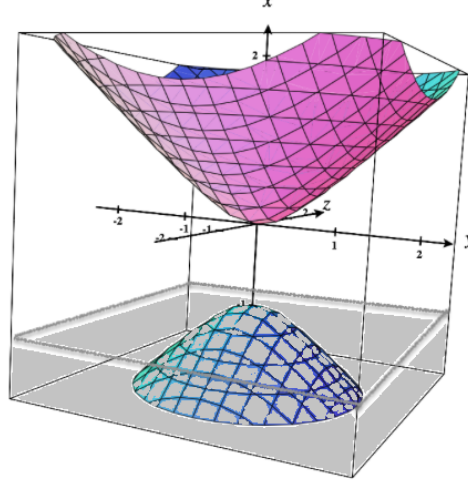
Demostración. Todo elemento $b \in \mathfrak{B}$ se puede escribir como combinación lineal de elementos de la forma $\alpha^k \gamma^m \gamma^{*n}$, donde $m, n \in \mathbb{N}$, $k \in \mathbb{Z}$, y si $k < 0$ entonces $\alpha^k = \alpha^{*-k}$. De esta manera se tiene que $\alpha^k \gamma^m \gamma^{*n} \in \mathcal{V}$ si y sólo si $k + m - n = 0$. Lo anterior se satisface para combinaciones lineales de $\rho = \gamma\gamma^*$, $\eta = \alpha\gamma^*$, $\eta^* = \gamma\alpha^*$. $\square$

Proposición 6.2.3. Las siguientes relaciones algebraicas

$$\begin{aligned}\rho &= \rho^*, & \rho\eta &= \mu^{-2}\eta\rho, & \rho\eta^* &= \mu^2\eta^*\rho, \\ (\mu^2\rho + \frac{1}{2})^2 - \eta\eta^* &= \frac{1}{4}, & (\rho + \frac{1}{2})^2 - \eta^*\eta &= \frac{1}{4}.\end{aligned}$$

se satisfacen en $\mathcal{V}$. $\square$

Observación 6.2.1. Podemos pensar al elemento ρ como una coordenada real, pero dado que $\rho = \gamma\gamma^*$ entonces estamos en el caso de que la coordenada real es positiva, y también pensamos al elemento η como una coordenada compleja, de manera que si $\mu = 1$ ó -1 , la $*$ -álgebra $\mathcal{V}$ geoméricamente corresponderá a la hoja superior del hiperboloide clásico.



Si $\mu \in (-1, 1) \setminus \{0\}$ el álgebra $\mathcal{V}$ es no conmutativa y podemos pensar que de manera geométrica se ve como un hiperboloide cuántico.

Proposición 6.2.4. La terna $P = (\mathfrak{B}, \iota, F)$ construida anteriormente, donde $\iota : \mathcal{V} \rightarrow \mathfrak{B}$ es el mapeo inclusión, es un haz principal cuántico. $\square$

Observación 6.2.2. Por todo lo anterior, observemos que si $\mu = -1$ tenemos un haz principal cuántico, donde resulta que tanto el grupo estructural como la variedad base son clásicos. A este haz principal lo llamamos haz principal hiperbólico cuántico.

Proposición 6.2.5. La co-acción F dada en la proposición 6.2.1 se extiende a un homomorfismo $\widehat{F} : \Omega(P) \rightarrow \Omega(P) \widehat{\otimes} \Gamma^\wedge$ definido como

$$\widehat{F}(\eta_3) = \eta_3 \otimes 1 + 1 \otimes \chi, \quad \widehat{F}(\eta_+) = \eta_+ \otimes U^2, \quad \widehat{F}(\eta_-) = \eta_- \otimes U^{-2}, \quad (6.12)$$

donde $\Omega(P) = \Psi^\wedge$.

Demostración. Observemos que el cálculo diferencial $\Omega(P)$ está generado por los elementos de $\mathfrak{B}$ y todos los elementos $\pi(b)$, $b \in \mathfrak{B}$. Aunque el cálculo $\Omega(P)$ no es bicovariante, se puede demostrar, como consecuencia de la propiedad de extendibilidad de los mapeos definidos sobre la envolvente universal que vimos antes, que la co-acción F se extiende y de manera única a $\widehat{F}$, donde resulta que

$$\widehat{F}(\pi(b)) = 1 \otimes \Pi(\pi(b)) + \pi(b^{(2)}) \otimes \Pi(\kappa(b^{(1)})b^{(3)}),$$

para todo $b \in \mathfrak{B}$. Aplicando directamente esta fórmula para $\widehat{F}$ obtenemos el resultado. $\square$

Observación 6.2.3. El álgebra de formas horizontales $\mathfrak{hor}(P)$ está generado por los elementos de $\mathfrak{B}$, y por los elementos η_+ y η_- .

Proposición 6.2.6. El álgebra de formas diferenciales sobre la base se puede escribir como

$$\Omega(M) = *\text{-}\text{alg}\{\rho, \eta, \eta^*, \eta_+\eta_-, \alpha\gamma\eta_-, \eta_+\gamma^*\alpha^*, \gamma^2\eta_-, \eta_+\gamma^{*2}, \alpha^2\eta_-, \eta_+\alpha^{*2}\}.$$

Demostración. Notemos que $\Omega(M) \subseteq \Omega(P)$ está formada por todos los $\varphi \in \Omega(P)$ tales que $\widehat{F}(\varphi) = \varphi \otimes 1$. Los elementos de la forma $\alpha^k \gamma^m \gamma^{*n} \eta_+^l \eta_-^s$, donde $k \in \mathbb{Z}$, $m, n, l, s \in \mathbb{N}$ pertenecen a $\Omega(M)$ si y sólo si $k + m - n + 2l - 2s = 0$. Así, el espacio de formas diferenciales en la base se escribe como sumas finitas de productos de los siguientes elementos

$$\begin{aligned} \rho &= \gamma\gamma^*, & \eta &= \alpha\gamma^*, & \eta^* &= \gamma\alpha^*, & \eta_+\eta_-, \\ \alpha\gamma\eta_-, & \eta_+\gamma^*\alpha^*, & \gamma^2\eta_-, & \eta_+\gamma^{*2}, & \alpha^2\eta_-, & \eta_+\alpha^{*2}. \end{aligned}$$

Por lo que estas relaciones definen completamente a $\Omega(M)$. $\square$

Proposición 6.2.7. La diferencial $d : \mathcal{V} = \Omega^0(M) \rightarrow \Omega(M)$ sobre la variedad base se define en los generadores de la siguiente manera:

$$\begin{aligned} d\rho &= \mu^{-2}\eta_+\gamma^*\alpha^* + \mu^{-2}\alpha\gamma\eta_-, \\ d\eta &= \eta_+\gamma^{*2} + \mu^{-1}\alpha^2\eta_-, \\ d\eta^* &= \mu^{-1}\eta_+\alpha^{*2} + \gamma^2\eta_-. \end{aligned}$$

Demostración. El resultado se obtiene aplicando a los generadores ρ , η y η^* las diferenciales de los elementos de $\mathfrak{B}$ dados en las ecuaciones (6.8) y (6.9), y aplicando las relaciones de conmutación dadas en (6.3) que verifican los generadores de $\mathfrak{B}$ y la proposición 6.1.4. $\square$

Observación 6.2.4. Los elementos de $\Omega^1(M)$ se escriben en término de los elementos de $\Omega^0(M) = \mathcal{V}$ de la siguiente manera:

$$\begin{aligned} \alpha\gamma\eta_- &= \mu^2(1 + \mu^2\rho)d\rho - \eta d\eta^*, & \eta_+\gamma^*\alpha^* &= \eta d\eta^* - \mu^4\rho d\rho, \\ \gamma^2\eta_- &= \mu^2\eta^*d\rho - \rho d\eta^*, & \eta_+\gamma^{*2} &= \eta d\rho - \mu^2\rho d\eta, \\ \alpha^2\eta_- &= \mu(1 + \mu^2\rho)d\eta - \mu\eta d\rho, & \eta_+\alpha^{*2} &= \mu(1 + \rho)d\eta^* - \mu^3\eta^*d\rho. \end{aligned}$$

Observemos también que el generador de $\Omega^2(M)$ se escribe como combinación lineal de elementos de $\Omega^0(M)$ y sus diferenciales como:

$$\eta_+\eta_- = -\left(1 + (\mu^2 + \mu^4)\rho + (\mu^6 - 1)\rho^2\right)^{-1} d\eta d\eta^*.$$

Por lo tanto, el cálculo diferencial $\Omega(M)$ sobre la variedad base representada por el álgebra $\mathcal{V}$ está generado por todos los elementos de grado cero.

6.2.1 Conexiones, derivada covariante y curvatura

Observemos que la acción adjunta $\text{ad} : \Gamma_{\text{inv}} \rightarrow \Gamma_{\text{inv}} \otimes \mathcal{A}$ en el generador χ es igual a

$$\begin{aligned} \text{ad}(\chi) &= \text{ad}(\pi_{\mathcal{A}}(U - U^{-1})) = \text{ad}(\pi_{\mathcal{A}}(U)) - \text{ad}(\pi_{\mathcal{A}}(U^{-1})) \\ &= \pi_{\mathcal{A}}(U) \otimes \kappa(U)U - \pi_{\mathcal{A}}(U^{-1}) \otimes \kappa(U^{-1})U^{-1} \\ &= \pi_{\mathcal{A}}(U) \otimes U^{-1}U - \pi_{\mathcal{A}}(U^{-1}) \otimes UU^{-1} \\ &= \pi_{\mathcal{A}}(U - U^{-1}) \otimes 1 = \chi \otimes 1. \end{aligned}$$

Por lo tanto una conexión cuántica en P debe satisfacer que

$$\widehat{F}\omega(\chi) = \omega(\chi) \otimes 1 + 1 \otimes \chi.$$

Observando la primera igualdad de (6.12) nos damos cuenta que podemos definir naturalmente una conexión $\omega : \Gamma_{\text{inv}} \rightarrow \Omega(P)$ en el haz P como:

$$\omega(\chi) = \eta_3. \quad (6.13)$$

Proposición 6.2.8. La conexión $\omega : \Gamma_{\text{inv}} \rightarrow \Omega(P)$ dada por $\omega(\chi) = \eta_3$ es una conexión regular.

Demostración. Recordemos que una conexión es regular si se cumple que

$$\omega(\theta)\varphi = (-1)^{\partial\varphi}\varphi^{(0)}\omega(\theta \circ \varphi^{(1)}),$$

para todo $\theta \in \Gamma_{\text{inv}}$ y $\varphi \in \mathfrak{hor}(P)$.

Por la observación 6.2.3 se tiene que el espacio de formas horizontales está generado por los elementos de $\mathfrak{B}$ y por los elementos η_+ y η_- , así como también por la proposición 6.1.10 tenemos que el espacio Γ_{inv} está generado por el elemento χ .

Por lo tanto usando las definiciones de las co-acciones F dada en la proposición 6.2.1 y $\widehat{F}$ en la ecuación (6.12) y que la acción derecha está definida como en (6.10) tenemos mediante cálculos directos que:

$$\begin{aligned} \omega(\chi)\alpha &= \eta_3\alpha = \mu^{-2}\alpha\eta_3 = \alpha\omega(\mu^{-2}\chi) = \alpha\omega(\chi \circ U), \\ \omega(\chi)\gamma &= \eta_3\gamma = \mu^{-2}\gamma\eta_3 = \gamma\omega(\mu^{-2}\chi) = \gamma\omega(\chi \circ U), \\ \omega(\chi)\gamma^* &= \eta_3\gamma^* = \mu^2\gamma^*\eta_3 = \gamma^*\omega(\mu^2\chi) = \gamma^*\omega(\chi \circ U^*), \\ \omega(\chi)\eta_+ &= \eta_3\eta_+ = -\mu^{-4}\eta_+\eta_3 = -\eta_+\omega(\mu^{-4}\chi) = -\eta_+\omega(\chi \circ U^2), \\ \omega(\chi)\eta_- &= \eta_3\eta_- = -\mu^4\eta_-\eta_3 = -\eta_-\omega(\mu^4\chi) = -\eta_-\omega(\chi \circ U^{*2}). \end{aligned}$$

Es decir, la conexión es regular. $\square$

Proposición 6.2.9. La conexión $\omega : \Gamma_{\text{inv}} \rightarrow \Omega(P)$, dada por $\omega(\chi) = \eta_3$ es una conexión multiplicativa.

Demostración. El ideal $\mathcal{J}_{\mathcal{A}}$ está generado por el elemento $\mu^2 U + U^{-1} + (1 - \mu^2)1$, por lo que basta ver que el mapeo $r_\omega : \mathcal{J} \rightarrow \Omega(P)$, dado por $r_\omega(a) = \omega\pi_{\mathcal{A}}(a^{(1)})\omega\pi_{\mathcal{A}}(a^{(2)})$, se anula en dicho generador. En efecto, usando que $\pi_{\mathcal{A}}(U) = 1/(1 + \mu^2)\chi$ y que $\pi_{\mathcal{A}}(U^{-1}) = -\mu^2/(1 + \mu^2)\chi$ se tiene que

$$\mu^2 (\omega\pi_{\mathcal{A}}(U))^2 + (\omega\pi_{\mathcal{A}}(U^{-1}))^2 = \frac{\mu^2}{(1 + \mu^2)^2} \eta_3^2 + \frac{\mu^4}{(1 + \mu^2)^2} \eta_3^2 = 0,$$

es decir, la conexión ω es multiplicativa. $\square$

Observación 6.2.5. Por la proposición 5.3.3 la derivada covariante cumple la regla graduada de Leibniz y $D_\omega(\varphi^*) = D_\omega(\varphi)^*$.

Proposición 6.2.10. La derivada covariante $D_\omega : \mathfrak{hor}(P) \rightarrow \mathfrak{hor}(P)$ se escribe en los generadores de la siguiente manera:

$$\begin{aligned} D_\omega(\alpha) &= \mu\gamma^*\eta_+, & D_\omega(\gamma) &= \alpha^*\eta_+, & D_\omega(\eta_+) &= 0, \\ D_\omega(\alpha^*) &= \gamma\eta_-, & D_\omega(\gamma^*) &= \mu^{-1}\alpha\eta_-, & D_\omega(\eta_-) &= 0. \end{aligned}$$

Demostración. Tenemos que $D_\omega(\varphi) = d\varphi - (-1)^{\partial\varphi}\varphi^{(0)}\omega\pi_{\mathcal{A}}(\varphi^{(1)})$ se cumple por definición para todo $\varphi \in \mathfrak{hor}(P)$. Por lo tanto, usando las ecuaciones (6.8), (6.9) y la diferencial definida en la proposición 6.1.8 se tiene que

$$\begin{aligned} D_\omega(\alpha) &= d\alpha - \alpha\omega\pi_{\mathcal{A}}(U) = d\alpha - \frac{1}{1 + \mu^2}\alpha\eta_3 = \mu\gamma^*\eta_+, \\ D_\omega(\gamma) &= d\gamma - \gamma\omega\pi_{\mathcal{A}}(U) = d\gamma - \frac{1}{1 + \mu^2}\gamma\eta_3 = \alpha^*\eta_+, \\ D_\omega(\eta_+) &= d\eta_+ + \eta_+\omega\pi_{\mathcal{A}}(U^2) = d\eta_+ + \frac{1}{\mu^2}\eta_+\eta_3 = 0. \end{aligned}$$

Aplicando la $*$ a las igualdades anteriores se completa la demostración. $\square$

Proposición 6.2.11. La curvatura $R_\omega : \Gamma_{\text{inv}} \rightarrow \text{hor}^2(P)$ se define en el generador χ como

$$R_\omega(\chi) = -(1 + \mu^2)\eta_-\eta_+.$$

Más aún, la curvatura extendida $r_\omega : \mathcal{A} \rightarrow \text{hor}^2(P)$ en el generador U^n está dada por la expresión

$$r_\omega(U^n) = -\frac{1 - \mu^{2n}}{\mu^{2(n-1)}(1 - \mu^2)}\eta_-\eta_+.$$

Demostración. Por definición $R_\omega(\chi) = d\omega(\chi) - \langle \omega | \omega \rangle \chi$. Sin embargo tenemos que

$$\langle \omega | \omega \rangle \chi = (1 + \mu^2) m_\omega(\omega \otimes \omega) \delta(\pi_{\mathcal{A}}(U)) = -(1 + \mu^2) (\omega \pi_{\mathcal{A}}(U))^2 = -\eta_3^2 = 0,$$

lo que implica que

$$R_\omega(\chi) = d\eta_3 = -(1 + \mu^2) \eta_- \eta_+.$$

Además, como la conexión ω es multiplicativa se tiene que $r_\omega = R_\omega \pi_{\mathcal{A}}$ y por lo tanto

$$r_\omega(U) = R_\omega \pi_{\mathcal{A}}(U) = \frac{1}{1 + \mu^2} R_\omega(\chi) = -\eta_- \eta_+.$$

Para demostrar la expresión general en el generador U^n obsérvese que

$$\pi_{\mathcal{A}}(U^n) = \frac{1 - \mu^{2n}}{\mu^{2(n-1)}(1 - \mu^2)} \pi_{\mathcal{A}}(U),$$

de manera que

$$r_\omega(U^n) = d\omega \pi_{\mathcal{A}}(U^n) = \frac{1 - \mu^{2n}}{\mu^{2(n-1)}(1 - \mu^2)} d\omega \pi_{\mathcal{A}}(U) = -\frac{1 - \mu^{2n}}{\mu^{2(n-1)}(1 - \mu^2)} \eta_- \eta_+,$$

como se quería demostrar. $\square$

Observación 6.2.6. Observemos que como la derivada covariante D_ω satisface la regla graduada de Leibniz se tiene que

$$\begin{aligned} D_\omega(R_\omega(\chi)) &= -(1 + \mu^2) D_\omega(\eta_- \eta_+) = -(1 + \mu^2) (D_\omega(\eta_-) \eta_+ - \eta_- D_\omega(\eta_+)) \\ &= 0, \end{aligned}$$

es decir, obtenemos la identidad clásica de Bianchi.

Observación 6.2.7. En el caso clásico, la conexión ω definida en esta sección corresponde a la conexión de Levi-Civita, esta conexión es caracterizada como la única conexión compatible con la métrica que además tiene torsión idénticamente cero.

6.3 Modelo cuántico de Poincaré

La geometría hiperbólica está muy estrechamente relacionada con la teoría de la relatividad y resulta ser una importante herramienta en el estudio de las 3-variedades, nudos, superficies de Riemann, entre otras. De hecho, el

plano hiperbólico proporciona un marco de referencia para la construcción de todas las superficies de Riemann compactas de curvatura constante negativa. Resulta que todas ellas se pueden obtener factorizando este plano por una acción libre de un grupo discreto infinito.

En esta tesis extendemos la teoría clásica a una nueva teoría hiperbólica cuántica, generalizando el modelo de Poincaré. Como veremos, muchas cosas de este modelo podrán ser reincorporadas naturalmente en el contexto cuántico, preservando el espíritu de las construcciones clásicas. Sin embargo, es natural esperar que aparecerán diversos fenómenos completamente ausentes en la teoría clásica.

Ya vimos que el grupo cuántico $SU_\mu(1, 1)$ representado por el álgebra $\mathfrak{B}$ se puede ver como un haz principal cuántico donde el grupo estructural es el círculo y la base corresponde a un espacio hiperbólico. Podemos dar un nuevo giro, y ver el álgebra $\mathfrak{B}$ como un producto cruzado entre el plano hiperbólico y el círculo. Para esto introduciremos nuevas coordenadas, que corresponderán clásicamente a las coordenadas del plano hiperbólico complejo en el modelo del disco unitario, y consideraremos una transformación unitaria que corresponderá a la rotación alrededor del origen.

Como vimos anteriormente, en geometría cuántica es crucial considerar representaciones (aunque posiblemente no acotadas) en un espacio $\mathcal{H}$ de Hilbert. Y para nuestra álgebra $\mathfrak{B}$, se puede demostrar que existen representaciones no acotadas en un espacio de Hilbert. Notemos que el elemento α en $\mathfrak{B}$ se puede ver como un operador de norma mayor o igual que uno, pues si $\psi \in \mathcal{H}$, se cumple que

$$\|\alpha\psi\|^2 = \langle \alpha\psi | \alpha\psi \rangle = \langle \psi | \alpha^* \alpha \psi \rangle = \langle \psi | (1 + \gamma^* \gamma) \psi \rangle = \|\psi\|^2 + \|\gamma\psi\|^2,$$

es decir, $\|\alpha\psi\| \geq \|\psi\|$, donde hemos usado las relaciones dadas en (6.3).

Esto último implica que el elemento α visto como operador es invertible y cuya inversa tiene norma menor o igual a uno.

Comentario. De aquí en adelante supondremos que nuestra álgebra $\mathfrak{B}$ se puede extender incluyendo al elemento α^{-1} y a apropiadas funciones holomorfas de $\alpha\alpha^*$ y de $\alpha^*\alpha$.

Recordemos que el álgebra $\mathfrak{B}$ está generada por los elementos dados en las entradas matriciales de (6.2). Dichos elementos cumplen las relaciones de conmutatividad que surgen al considerar la propiedad de determinante igual a uno. Por lo tanto la matriz u es invertible y podemos considerar su descomposición polar dada por $u = |u|W$, donde $|u| = \sqrt{uu^*}$ y W es una matriz unitaria.

Observación 6.3.1. Sea u la matriz dada en (6.2) y u^* su matriz adjunta, entonces la matriz uu^* se puede escribir como

$$\begin{aligned} uu^* &= \begin{pmatrix} \alpha & \mu\gamma^* \\ \gamma & \alpha^* \end{pmatrix} \begin{pmatrix} \alpha^* & \gamma^* \\ \mu\gamma & \alpha \end{pmatrix} = \begin{pmatrix} \alpha\alpha^* + \mu^2\gamma^*\gamma & \alpha\gamma^* + \mu\gamma^*\alpha \\ \gamma\alpha^* + \mu\alpha^*\gamma & \alpha^*\alpha + \gamma\gamma^* \end{pmatrix} \\ &= \begin{pmatrix} 2\alpha\alpha^* - 1 & 2\alpha\gamma^* \\ 2\gamma\alpha^* & 2\alpha^*\alpha - 1 \end{pmatrix} \end{aligned}$$

donde hemos usado las relaciones de conmutación dadas en (6.3) entre los generadores de $\mathfrak{B}$.

Proposición 6.3.1. Es posible introducir una coordenada compleja z tal que

$$uu^* = \begin{pmatrix} \frac{1 + zz^*}{1 - zz^*} & \frac{2}{1 - zz^*}z \\ z^*\frac{2}{1 - zz^*} & \frac{1 + z^*z}{1 - z^*z} \end{pmatrix}. \quad (6.14)$$

Más aún, zz^* y z^*z se relacionan mediante las fórmulas

$$z^*z = \frac{zz^*}{\mu^2 + (1 - \mu^2)zz^*}, \quad zz^* = \frac{z^*z}{\mu^{-2} + (1 - \mu^{-2})z^*z}, \quad (6.15)$$

y las siguientes reglas de conmutación se satisfacen

$$\frac{1}{1 - zz^*}z = z\frac{1 + (\mu^{-2} - 1)zz^*}{1 - zz^*}, \quad z^*\frac{1}{1 - zz^*} = \frac{1 + (\mu^{-2} - 1)zz^*}{1 - zz^*}z^*.$$

Demostración. Consideremos el siguiente cambio de coordenadas

$$z = \alpha^{*-1}\gamma^*, \quad z^* = \gamma\alpha^{-1},$$

entonces

$$\begin{aligned} 1 - z^*z &= 1 - \gamma\alpha^{-1}\alpha^{*-1}\gamma^* = 1 - \mu^2\alpha^{-1}\gamma\gamma^*\alpha^{*-1} = 1 - \alpha^{-1}(\alpha\alpha^* - 1)\alpha^{*-1} \\ &= (\alpha^*\alpha)^{-1}, \end{aligned}$$

es decir, $\alpha^*\alpha = 1/(1 - z^*z)$. Similarmente

$$1 - zz^* = 1 - \alpha^{*-1}\gamma^*\gamma\alpha^{-1} = 1 - \alpha^{*-1}(\alpha^*\alpha - 1)\alpha^{-1} = (\alpha\alpha^*)^{-1},$$

por lo que $\alpha\alpha^* = 1/(1 - zz^*)$. Por otro lado también se tiene que

$$\frac{1}{1 - zz^*}z = \alpha\alpha^*\alpha^{*-1}\gamma^* = \alpha\gamma^*.$$

De esta manera la matriz uu^* se escribe como en (6.14).

Mediante cálculos directos, las dos relaciones $\alpha^*\alpha = (1 - \mu^{-2}) + \mu^{-2}\alpha\alpha^*$ y $\alpha\alpha^* = (1 - \mu^2) + \mu^2\alpha^*\alpha$ nos llevan a las ecuaciones (6.15).

Finalmente observemos que

$$\begin{aligned} z^* \frac{1}{1 - zz^*} &= \gamma\alpha^{-1}\alpha\alpha^* = \gamma\alpha^{-1}(1 + \mu^2\gamma\gamma^*) = (1 + \gamma\gamma^*)\gamma\alpha^{-1}\alpha^*\alpha\gamma\alpha^{-1} \\ &= ((1 - \mu^{-2}) + \mu^{-2}\alpha\alpha^*)\gamma\alpha^{-1} = ((1 - \mu^{-2}) + \mu^{-2}\frac{1}{1 - zz^*})z^* \\ &= \frac{1 + (\mu^{-2} - 1)zz^*}{1 - zz^*}z^*. \end{aligned}$$

Aplicando la $*$ a esta última relación terminamos la demostración. $\square$

Proposición 6.3.2. La matriz

$$P = \begin{pmatrix} \frac{1}{\sqrt{1 - zz^*}} & \frac{1}{\sqrt{1 - zz^*}}z \\ z^* \frac{1}{\sqrt{1 - zz^*}} & \frac{\sqrt{1 + (\mu^{-2} - 1)zz^*}}{\sqrt{1 - zz^*}} \end{pmatrix}$$

es la raíz cuadrada de la matriz uu^ dada en (6.14), es decir $P = |u|$.* $\square$

Observación 6.3.2. Continuando con la descomposición polar de la matriz u tenemos que la matriz unitaria W está dada por

$$W = \begin{pmatrix} w & 0 \\ 0 & w^* \end{pmatrix}$$

donde $w = (1/\sqrt{\alpha\alpha^*})\alpha$, es obvio que el elemento w es unitario. Por lo tanto la matriz u se escribe en coordenadas z y w como

$$u = \begin{pmatrix} \frac{1}{\sqrt{1 - zz^*}}w & \frac{1}{\sqrt{1 - zz^*}}zw^* \\ z^* \frac{1}{\sqrt{1 - zz^*}}w & \frac{\sqrt{1 + (\mu^{-2} - 1)zz^*}}{\sqrt{1 - zz^*}}w^* \end{pmatrix}.$$

Así por la construcción se tiene que las antiguas coordenadas de la matriz u ahora se escriben en nuevas coordenadas z y w de la siguiente manera

$$\begin{aligned} \alpha &= \frac{1}{\sqrt{1 - zz^*}}w, & \alpha^* &= \frac{\sqrt{1 + (\mu^{-2} - 1)zz^*}}{\sqrt{1 - zz^*}}w^*, \\ \gamma &= z^* \frac{1}{\sqrt{1 - zz^*}}w, & \mu\gamma^* &= \frac{1}{\sqrt{1 - zz^*}}zw^*. \end{aligned}$$

Cabe observar que la identidad (6.15) implica la siguiente relación

$$\frac{1}{\sqrt{1 - z^*z}} = \frac{\sqrt{1 + (\mu^{-2} - 1)zz^*}}{\sqrt{1 - zz^*}}. \quad (6.16)$$

Por lo tanto las relaciones de conmutación dadas en la proposición 6.3.1 se pueden reescribir de la siguiente manera:

$$\frac{1}{1 - zz^*}z = z\frac{1}{1 - z^*z}, \quad z^*\frac{1}{1 - zz^*} = \frac{1}{1 - z^*z}z^*. \quad (6.17)$$

Y la igualdad entre las matrices en coordenadas α, γ y coordenadas z, w nos da las siguientes reglas de conmutación:

$$\frac{1}{\sqrt{1 - zz^*}}w = w\frac{1}{\sqrt{1 - z^*z}}, \quad \frac{1}{\sqrt{1 - z^*z}}z^*w = \mu^{-1}wz^*\frac{1}{\sqrt{1 - zz^*}}, \quad (6.18)$$

$$w^*\frac{1}{\sqrt{1 - zz^*}} = \frac{1}{\sqrt{1 - z^*z}}w^*, \quad w^*z\frac{1}{\sqrt{1 - z^*z}} = \mu^{-1}\frac{1}{\sqrt{1 - zz^*}}zw^*. \quad (6.19)$$

Definimos $\psi : \mathfrak{B} \rightarrow \mathfrak{B}$, como el homomorfismo dado por conjugación, es decir, $\psi(b) = w^*bw$ para todo $b \in \mathfrak{B}$. Con las relaciones anteriores podemos ver que éste evaluado en z o en z^* es igual al producto de una función de zz^* , por z o z^* , respectivamente.

Proposición 6.3.3. El homomorfismo $\psi : \mathfrak{B} \rightarrow \mathfrak{B}$ en las coordenadas z y w , lo podemos ver como

$$\begin{aligned} \psi(z) &= \mu^{-1} \sqrt{\frac{1 - z^*z}{1 - zz^*}} z, & \psi(z^*) &= \mu^{-1} z^* \sqrt{\frac{1 - z^*z}{1 - zz^*}}, \\ \psi(w) &= w, & \psi(w^*) &= w^*. \end{aligned} \quad (6.20)$$

Más aún, dicho homomorfismo es un automorfismo del álgebra, con inversa $\psi^{-1}(b) = wbw^*$ para todo $b \in \mathfrak{B}$, y que se define en las coordenadas z y w como:

$$\begin{aligned} \psi^{-1}(z) &= \mu z \sqrt{\frac{1 - zz^*}{1 - z^*z}}, & \psi^{-1}(z^*) &= \mu \sqrt{\frac{1 - zz^*}{1 - z^*z}} z^*, \\ \psi^{-1}(w) &= w, & \psi^{-1}(w^*) &= w^*. \end{aligned} \quad (6.21)$$

Demostración. Observemos que si consideramos la segunda ecuación de (6.18) se tiene que

$$\mu^{-1}z^*\frac{1}{\sqrt{1 - zz^*}} = w^*\frac{1}{\sqrt{1 - z^*z}}z^*w = w^*z^*w\frac{1}{\sqrt{1 - z^*z}},$$

es decir,

$$w^* z^* w = \mu^{-1} z^* \sqrt{\frac{1 - z^* z}{1 - z z^*}} = \psi(z^*).$$

Ahora, usando la segunda ecuación de (6.19), conmutando z con la función de $z z^*$ del lado derecho de la igualdad y usando la identidad (6.16) llegamos a

$$\begin{aligned} \mu^{-1} z \frac{1}{\sqrt{1 - z^* z}} &= w^* z \frac{1}{\sqrt{1 - z^* z}} w = w^* z \frac{\sqrt{1 + (\mu^{-2} - 1) z z^*}}{\sqrt{1 - z z^*}} w \\ &= w^* z w \frac{\sqrt{1 + (\mu^{-2} - 1) z z^*}}{\sqrt{1 - z^* z}}, \end{aligned}$$

es decir,

$$\begin{aligned} w^* z w &= \mu^{-1} z \frac{1}{\sqrt{1 + (\mu^{-2} - 1) z z^*}} = \mu^{-1} \frac{1}{\sqrt{1 + (\mu^{-2} - 1) z z^*}} z \\ &= \mu^{-1} \sqrt{\frac{1 - z^* z}{1 - z z^*}} z = \psi(z). \end{aligned}$$

De igual manera de las ecuaciones (6.18) y (6.19) se tiene que

$$w z^* w^* = \mu \sqrt{\frac{1 - z z^*}{1 - z^* z}} z^*, \quad w z w^* = \mu z \sqrt{\frac{1 - z z^*}{1 - z^* z}},$$

teniendo de esta manera que el homomorfismo $\psi^{-1}(b) = w b w^*$, $b \in \mathfrak{B}$ es la inversa de ψ . Por lo tanto ψ es automorfismo de $\mathfrak{B}$. $\square$

Proposición 6.3.4. El automorfismo $\psi : \mathfrak{B} \rightarrow \mathfrak{B}$ dado por $\psi(b) = w^* b w$, para todo $b \in \mathfrak{B}$ y definido en la proposición 6.3.3 satisface la siguiente relación

$$z^* z = \psi(z z^*).$$

Demostración. Usando la relación (6.16) podemos reescribir el automorfismo ψ en los generadores z y z^* de la siguiente manera

$$\psi(z) = \mu^{-1} \frac{1}{\sqrt{1 + (\mu^{-2} - 1) z z^*}} z, \quad \psi(z^*) = \mu^{-1} z^* \frac{1}{\sqrt{1 + (\mu^{-2} - 1) z z^*}},$$

y por lo tanto

$$\begin{aligned} \psi(z z^*) &= \mu^{-2} \frac{1}{\sqrt{1 + (\mu^{-2} - 1) z z^*}} z z^* \frac{1}{\sqrt{1 + (\mu^{-2} - 1) z z^*}} \\ &= \mu^{-2} \frac{z z^*}{1 + (\mu^{-2} - 1) z z^*} = \frac{z z^*}{\mu^2 + (1 - \mu^2) z z^*} = z^* z, \end{aligned}$$

donde en el último renglón hemos usado que se satisface (6.15). $\square$

Observación 6.3.3. Notemos que la igualdad

$$\psi\left(\frac{1}{1-zz^*}\right) = \mu^{-2}\frac{1}{1-zz^*} + (1-\mu^{-2})$$

se satisface, de manera que la función $zz^*/(1-zz^*)$ se ve como un vector propio del automorfismo ψ , con valor propio μ^{-2} . Por lo tanto podemos concluir que todas las potencias iteradas de dicho vector forman una base de vectores propios y el espectro del operador está dado por todas las potencias pares de μ^{-1} .

Observación 6.3.4. Hasta este momento, hemos reescrito el álgebra $\mathfrak{B}$ en nuevas coordenadas z y w donde se satisfacen las siguientes reglas de conmutación entre los generadores

$$z^*z = \psi(zz^*), \quad zw = w\psi(z), \quad w^* = w^{-1}, \quad (6.22)$$

donde ψ es el automorfismo interno dado en la proposición 6.3.3.

Observación 6.3.5. Notemos que si $\mu = 1$ se satisface que $\psi(b) = b$, para todo $b \in \mathfrak{B}$ y por lo tanto recuperamos el caso conmutativo.

Definición 6.3.1. Denotamos por $\mathcal{S}$ la subálgebra conmutativa de $\mathcal{V}$ generada por el elemento zz^*

Observación 6.3.6. Equivalentemente podemos decir que $\mathcal{S}$ está generada por el elemento z^*z pues se satisface la relación (6.15). Además por la proposición 6.3.4 también se tiene que $\mathcal{S}$ es ψ -invariante.

Como el elemento zz^* es autoadjunto, podemos pensarlo como una coordenada real y con espectro dentro del intervalo $[0, 1]$. Definimos entonces la función $f : [0, 1] \rightarrow [0, 1]$ como

$$f(t) := \frac{t}{\mu^2 + (1 - \mu^2)t}. \quad (6.23)$$

Notemos que por definición se tiene la relación

$$z^*z = f(zz^*) = \psi(zz^*),$$

es decir la función f y el automorfismo ψ coinciden en zz^* , dándonos el producto opuesto z^*z .

Proposición 6.3.5. Consideremos el grupo cuántico $\mathrm{SU}_\mu(1, 1)$ representado por el álgebra $\mathfrak{B}$ en coordenadas z, w y el círculo unitario representado por el álgebra $\mathcal{A}$. El $*$ -homomorfismo $T : \mathfrak{B} \rightarrow \mathfrak{B} \otimes \mathcal{A}$ definido como

$$T(z) = z \otimes 1, \quad T(z^*) = z^* \otimes 1, \quad T(w) = w \otimes U, \quad T(w^*) = w^* \otimes U^*,$$

es una co-acción derecha de $\mathcal{A}$ en $\mathfrak{B}$.

Demostración. Es fácil verificar que el $*$ -homomorfismo T es compatible con las relaciones dadas en (6.22). Ahora veamos que para los generadores z y w se cumplen las propiedades de co-acción derecha:

$$\begin{aligned} (\text{id} \otimes \epsilon_{\mathcal{A}})T(z) &= z \otimes \epsilon_{\mathcal{A}}(1) = z \otimes 1 \cong z, \\ (\text{id} \otimes \epsilon_{\mathcal{A}})T(w) &= w \otimes \epsilon_{\mathcal{A}}(U) = w \otimes 1 \cong w, \end{aligned}$$

por lo tanto la relación $(\text{id} \otimes \epsilon)T = \text{id}$ se satisface en $\mathfrak{B}$. Por otro lado:

$$\begin{aligned} (\text{id} \otimes \phi_{\mathcal{A}})T(z) &= z \otimes \phi_{\mathcal{A}}(1) = z \otimes 1 \otimes 1 = T(z) \otimes 1 = (T \otimes \text{id})T(z), \\ (\text{id} \otimes \phi_{\mathcal{A}})T(w) &= w \otimes \phi_{\mathcal{A}}(U) = w \otimes U \otimes U = T(w) \otimes U = (T \otimes \text{id})T(w), \end{aligned}$$

lo que muestra que la relación $(\text{id} \otimes \phi_{\mathcal{A}})T = (T \otimes \text{id})T$ se satisface. Por lo tanto T es una co-acción derecha de $\mathcal{A}$ en $\mathfrak{B}$. $\square$

Observación 6.3.7. El álgebra $\mathcal{W}$ de todos los elementos de $\mathfrak{B}$ que quedan invariantes bajo la co-acción T dada en la proposición anterior esta generada por los elementos $1, z$ y z^* . Esto es,

$$\mathcal{W} = *\text{-}\text{álgebra}\{1, z, z^*\}. \quad (6.24)$$

Proposición 6.3.6. Sea $\iota : \mathcal{W} \rightarrow \mathfrak{B}$ el mapeo inclusión del álgebra $\mathcal{W}$ definida en la ecuación (6.24) en el álgebra $\mathfrak{B}$ que representa al grupo cuántico $\text{SU}_{\mu}(1, 1)$ en coordenadas z, w . Entonces, la terna $(\mathfrak{B}, T, \iota)$ es un haz principal cuántico, con grupo estructural dado por el álgebra del círculo $\mathcal{A}$ y donde la co-acción T está dada como en la proposición 6.3.5.

Demostración. Observemos que los elementos z y w generan al álgebra $\mathfrak{B}$ y la relación de conmutatividad en $\mathfrak{B}$ dada en (6.22), nos dice que los elementos de $\mathfrak{B}$ van a ser sumas finitas de elementos $z^l \psi^m (zz^*)^n w^k$, donde $k, l, m \in \mathbb{Z}$ y $n \in \mathbb{N}$. También, dado que el conjunto $\{U^n : n \in \mathbb{Z}\}$ genera al álgebra $\mathcal{A}$ tenemos que cualquier elemento de $\mathfrak{B} \otimes \mathcal{A}$ se escribe como sumas finitas de elementos de la forma $z^l \psi^m (zz^*)^n w^k \otimes U^p$.

Así el mapeo lineal $X : \mathfrak{B} \otimes \mathfrak{B} \rightarrow \mathfrak{B} \otimes \mathcal{A}$ dado por $X(a \otimes b) = aT(b)$ es sobre, pues

$$X(z^l \psi^m (zz^*)^n w^{k-p} \otimes w^p) = z^l \psi^m (zz^*)^n w^{k-p} w^p \otimes U^p = z^l \psi^m (zz^*)^n w^k \otimes U^p$$

De esta forma concluimos que $(\mathfrak{B}, T, \iota)$ es un haz principal cuántico. $\square$

Observación 6.3.8. Si $\mu = 1$ ó -1 el álgebra $\mathcal{W}$ es conmutativa pues $zz^* = z^*z$. En este último caso de $\mu = -1$ tenemos como ya habíamos observado en la sección anterior, un haz principal cuántico cuyo grupo estructural y variedad base son clásicos.

Proposición 6.3.7. Sea Ψ el cálculo diferencial construido sobre $\mathfrak{B}$ dado por el ideal $\mathcal{I}$ definido en la ecuación (6.4). Entonces se cumplen las siguientes relaciones de conmutación entre los generadores del cálculo y los generadores del álgebra:

$$\eta_3 z = z \eta_3, \quad \eta_{\pm} z = z \eta_{\pm}, \quad \eta_3 w = \mu^{-2} w \eta_3, \quad \eta_{\pm} w = \mu^{-1} w \eta_{\pm}.$$

Demostración. En efecto, cálculos directos considerando las reglas de conmutación dadas en la proposición (6.1.4) nos llevan a los resultados deseados. $\square$

Proposición 6.3.8. El cálculo diferencial $\Omega^1(M)$ sobre $\mathcal{W}$ está dado mediante las siguientes fórmulas

$$dz = \sqrt{1 - zz^*} w \eta_- w \sqrt{1 - z^* z} = \mu^{-1} (\sqrt{1 - zz^*} w)^2 \eta_-, \quad (6.25)$$

$$dz^* = \sqrt{1 - z^* z} w^* \eta_+ w^* \sqrt{1 - zz^*} = \mu (\sqrt{1 - z^* z} w^*)^2 \eta_+. \quad (6.26)$$

Demostración. Sustituyendo los valores de α , α^* , γ y γ^* en términos de las coordenadas z, w (ver página 149) en las ecuaciones (6.8) y (6.9) tenemos las siguientes igualdades

$$d\left(\frac{1}{\sqrt{1 - zz^*}} w\right) = \frac{1}{1 + \mu^2} \frac{1}{\sqrt{1 - zz^*}} w \eta_3 + \frac{1}{\sqrt{1 - zz^*}} z w^* \eta_+, \quad (6.27)$$

$$d\left(z^* \frac{1}{\sqrt{1 - zz^*}} w\right) = \frac{1}{1 + \mu^2} z^* \frac{1}{\sqrt{1 - zz^*}} w \eta_3 + \frac{1}{\sqrt{1 - z^* z}} w^* \eta_+, \quad (6.28)$$

$$d\left(\frac{1}{\sqrt{1 - z^* z}} w^*\right) = -\frac{\mu^2}{1 + \mu^2} \frac{1}{\sqrt{1 - z^* z}} w^* \eta_3 + \mu^{-1} w z^* \frac{1}{\sqrt{1 - zz^*}} \eta_-, \quad (6.29)$$

$$d\left(\frac{1}{\sqrt{1 - z^* z}} w^* z\right) = -\frac{\mu^2}{1 + \mu^2} \frac{1}{\sqrt{1 - z^* z}} w^* z \eta_3 + \mu^{-1} \frac{1}{\sqrt{1 - zz^*}} w \eta_-. \quad (6.30)$$

Multiplicando la ecuación (6.27) por la izquierda con el elemento z^* y comparándolo con la ecuación (6.28) tenemos que

$$\begin{aligned} d(z^*) \frac{1}{\sqrt{1 - zz^*}} w &= \frac{1}{\sqrt{1 - z^* z}} w^* \eta_+ - z^* \frac{1}{\sqrt{1 - zz^*}} z w^* \eta_+ \\ &= \frac{1}{\sqrt{1 - z^* z}} w^* \eta_+ - \frac{1}{\sqrt{1 - z^* z}} z^* z w^* \eta_+ \\ &= \sqrt{1 - z^* z} w^* \eta_+, \end{aligned}$$

es decir,

$$dz^* = \sqrt{1 - z^* z} w^* \eta_+ w^* \sqrt{1 - zz^*}.$$

Además, usando las reglas de conmutación dadas en la ecuación (6.19) podemos reescribir lo anterior como

$$dz^* = \mu(\sqrt{1 - z^*zw^*})^2 \eta_+,$$

como se quería demostrar.

Por otro lado, multiplicando la ecuación (6.29) por la derecha por el elemento z y comparándola con la ecuación (6.30) tenemos

$$\begin{aligned} \frac{1}{\sqrt{1 - z^*z}} w^* dz &= \mu^{-1} \frac{1}{\sqrt{1 - zz^*}} w \eta_- - \mu^{-1} w z^* \frac{1}{\sqrt{1 - zz^*}} z \eta_- \\ &= \mu^{-1} \frac{1}{\sqrt{1 - zz^*}} w \eta_- - \mu^{-1} \frac{1}{\sqrt{1 - zz^*}} w z^* z \eta_- \\ &= \mu^{-1} \frac{1}{\sqrt{1 - zz^*}} (1 - z z^*) w \eta_- \\ &= \mu^{-1} \sqrt{1 - zz^*} w \eta_-. \end{aligned}$$

Por lo tanto,

$$dz = \mu^{-1} w \sqrt{1 - z^*z} \sqrt{1 - zz^*} w \eta_- = \mu^{-1} (\sqrt{1 - zz^*} w)^2 \eta_-,$$

lo que concluye la demostración. $\square$

Observación 6.3.9. Podemos escribir a los elementos η_- y η_+ en término de las diferenciales de z y z^* respectivamente, como

$$\begin{aligned} \eta_- &= \mu \left(w^* \frac{1}{\sqrt{1 - zz^*}} \right)^2 dz, \\ \eta_+ &= \mu^{-1} \left(w \frac{1}{\sqrt{1 - z^*z}} \right)^2 dz^*. \end{aligned}$$

Y por la proposición (6.3.7) se tiene que $w^* \eta_{\pm} w = \mu^{-1} \eta_{\pm}$. Por lo tanto

$$\mu^{-1} \eta_- = \left(w^* \frac{1}{\sqrt{1 - zz^*}} \right)^2 dz.$$

Y por otra parte

$$\begin{aligned} w^* \eta_- w &= \mu w^* \left(w^* \frac{1}{\sqrt{1 - zz^*}} \right)^2 dz w \\ &= \mu \left(w^* \frac{1}{\sqrt{1 - z^*z}} \right)^2 w^* dz w. \end{aligned}$$

Igualando estas dos últimas expresiones llegamos a que

$$w^* dz w = \mu^{-1} \sqrt{\frac{1 - z^* z}{1 - \psi^{-2}(z^* z)}} dz.$$

Similarmente

$$\mu^{-1} \eta_+ = \mu^{-2} \left(w \frac{1}{\sqrt{1 - z^* z}} \right)^2 dz^*,$$

y también

$$\begin{aligned} w^* \eta_+ w &= \mu^{-1} w^* \left(w \frac{1}{\sqrt{1 - z^* z}} \right)^2 dz^* w \\ &= \mu^{-1} \left(\frac{1}{\sqrt{1 - z^* z}} w \right)^2 w^* dz^* w. \end{aligned}$$

De esta manera, igualando ambas expresiones tenemos

$$w^* dz^* w = \mu^{-1} \sqrt{\frac{1 - \psi^2(z^* z)}{1 - z^* z}} dz^*.$$

Observación 6.3.10. La multiplicación de las ecuaciones (6.25) y (6.26) nos lleva a lo siguiente:

$$\begin{aligned} dz \, dz^* &= (\sqrt{1 - z z^*} w)^2 \eta_- (\sqrt{1 - z^* z} w^*)^2 \eta_+ \\ &= (w \sqrt{1 - z^* z})^2 \eta_- (w \sqrt{1 - z^* z})^{*2} \eta_+ \\ &= \mu^2 (1 - z z^*) (1 - \psi^{-1}(z z^*)) \eta_- \eta_+. \end{aligned}$$

Por lo tanto tenemos que

$$\eta_- \eta_+ = \mu^{-2} \frac{1}{(1 - \psi^{-1}(z z^*)) (1 - z z^*)} dz \, dz^*. \quad (6.31)$$

De esta manera, es claro que si $\mu = \pm 1$, se tiene que $z^* z = z z^* = |z|^2$, y por lo tanto el automorfismo $\psi = \text{id}$ y

$$\eta_- \eta_+ = \frac{dz \, dz^*}{(1 - |z|^2)^2} = \frac{1}{4} \left[\frac{4 \, dz \, dz^*}{(1 - |z|^2)^2} \right],$$

es decir, obtenemos la métrica hiperbólica en la variedad base del haz cuántico. Se trata entonces, de una generalización del modelo cuántico del disco de Poincaré.

Proposición 6.3.9. Se satisfacen las siguientes reglas de conmutación entre los elementos del cálculo $\Omega^1(M)$ y del álgebra $\mathcal{W}$

$$\begin{aligned} dz^* z^* &= \mu^{-2} z^* dz^*, & dz z &= \mu^2 z dz, \\ dz^* z &= \frac{\mu^2}{\mu^4 + (1 - \mu^4)zz^*} z dz^*, & dz z^* &= \frac{\mu^{-2}}{\mu^{-4} + (1 - \mu^{-4})z^*z} z^* dz. \end{aligned}$$

Demostración. En efecto, por la definición de dz y usando la definición de $\psi^{-1}(z)$ dada en la demostración de la proposición 6.3.3 se tiene lo siguiente

$$\begin{aligned} dz z &= \mu^{-1}(\sqrt{1 - zz^*w})^2 \eta_- z = \mu^{-1}(\sqrt{1 - zz^*w})^2 z \eta_- \\ &= \sqrt{1 - zz^*w} \sqrt{1 - zz^*z} \sqrt{\frac{1 - zz^*}{1 - z^*z}} w \eta_- \\ &= \sqrt{1 - zz^*w} z \sqrt{1 - z^*z} \sqrt{\frac{1 - zz^*}{1 - z^*z}} w \eta_- \\ &= \mu \sqrt{1 - zz^*z} \sqrt{\frac{1 - zz^*}{1 - z^*z}} w \sqrt{1 - zz^*w} \eta_- \\ &= \mu z (\sqrt{1 - zz^*w})^2 \eta_- = \mu^2 z dz. \end{aligned}$$

Por otra parte, usando la definición de dz^* y que la igualdad

$$\left(\frac{1}{1 + (\mu^{-2} - 1)zz^*} \right) \left(\frac{1}{1 + (\mu^{-2} - 1)z^*z} \right) = \frac{1}{1 + (\mu^{-4} - 1)zz^*}$$

se satisface, tendremos que

$$\begin{aligned} dz^* z &= \mu(\sqrt{1 - z^*zw^*})^2 \eta_+ z = \sqrt{1 - z^*zw^*} \sqrt{1 - z^*z} \sqrt{\frac{1 - z^*z}{1 - zz^*}} zw^* \eta_+ \\ &= \sqrt{1 - z^*zw^*} \sqrt{1 - z^*z} \frac{1}{\sqrt{1 + (\mu^{-2} - 1)zz^*}} zw^* \eta_+ \\ &= \sqrt{1 - z^*zw^*} \sqrt{1 - z^*zz} \frac{1}{\sqrt{1 + (\mu^{-2} - 1)z^*z}} w^* \eta_+ \\ &= \sqrt{1 - z^*zw^*} \sqrt{\frac{1 - zz^*}{1 + (\mu^{-2} - 1)zz^*}} z \frac{1}{\sqrt{1 + (\mu^{-2} - 1)z^*z}} w^* \eta_+ \\ &= \sqrt{1 - z^*zw^*} z \frac{\sqrt{1 - z^*z}}{1 + (\mu^{-2} - 1)z^*z} w^* \eta_+ \end{aligned}$$

$$\begin{aligned}
&= \mu^{-1} \sqrt{1 - z^* z} \sqrt{\frac{1 - z^* z}{1 - z z^*}} z w^* \frac{\sqrt{1 - z^* z}}{1 + (\mu^{-2} - 1) z^* z} w^* \eta_+ \\
&= \mu^{-1} (1 - z^* z) z \frac{1}{\sqrt{1 - z^* z}} w^* \frac{\sqrt{1 - z^* z}}{1 + (\mu^{-2} - 1) z^* z} w^* \eta_+ \\
&= \mu^{-1} \frac{1 - z z^*}{1 + (\mu^{-2} - 1) z z^*} z \frac{1}{\sqrt{1 - z^* z}} w^* \frac{\sqrt{1 - z^* z}}{1 + (\mu^{-2} - 1) z^* z} w^* \eta_+ \\
&= \mu^{-1} z \frac{1 - z^* z}{1 + (\mu^{-2} - 1) z^* z} \frac{1}{\sqrt{1 - z^* z}} w^* \frac{\sqrt{1 - z^* z}}{1 + (\mu^{-2} - 1) z^* z} w^* \eta_+ \\
&= \mu^{-1} z \frac{\sqrt{1 - z^* z}}{1 + (\mu^{-2} - 1) z^* z} w^* \frac{\sqrt{1 - z^* z}}{1 + (\mu^{-2} - 1) z^* z} w^* \eta_+ \\
&= \mu^{-1} \frac{1}{1 + (\mu^{-2} - 1) z z^*} z \sqrt{1 - z^* z} w^* \frac{\sqrt{1 - z^* z}}{1 + (\mu^{-2} - 1) \psi^{-1}(z z^*)} w^* \eta_+ \\
&= \mu^{-1} \frac{1}{1 + (\mu^{-2} - 1) z z^*} z \frac{\sqrt{1 - z^* z}}{1 + (\mu^{-2} - 1) \psi^{-1}(z z^*)} w^* \sqrt{1 - z^* z} w^* \eta_+ \\
&= \mu^{-1} \frac{1}{1 + (\mu^{-2} - 1) z z^*} \frac{1}{1 + (\mu^{-2} - 1) z^* z} z (\sqrt{1 - z^* z} w^*)^2 \eta_+ \\
&= \mu^{-2} \frac{1}{1 + (\mu^{-4} - 1) z z^*} z d z^* = \frac{\mu^2}{\mu^4 + (1 - \mu^4) z z^*} z d z^*.
\end{aligned}$$

Aplicando la $*$ a las 2 igualdades obtenidas se completa la demostración. $\square$

Observación 6.3.11. Por la Observación 6.3.10 se tiene que

$$dz dz^* = \mu^2 (1 - z z^*) \psi^{-1} (1 - z z^*) \eta_- \eta_+.$$

Y por otro lado

$$\begin{aligned}
dz^* dz &= (\sqrt{1 - z^* z} w^*)^2 \eta_+ (\sqrt{1 - z z^*} w)^2 \eta_- \\
&= \mu^{-2} (\sqrt{1 - z^* z} w^*)^2 (w \sqrt{1 - z^* z})^2 \eta_+ \eta_- \\
&= \mu^{-2} (1 - z^* z) \psi (1 - z^* z) \eta_+ \eta_- \\
&= -(1 - z^* z) \psi (1 - z^* z) \eta_- \eta_+.
\end{aligned}$$

De manera que se satisface la siguiente relación

$$dz^* dz = -\mu^{-2} \frac{\psi(1 - z z^*) \psi^2(1 - z z^*)}{(1 - z z^*) \psi^{-1}(1 - z z^*)} dz dz^* \quad (6.32)$$

o equivalentemente,

$$dz^* dz = -\mu^{-2} \frac{1 + (\mu^2 - 1) z z^*}{(1 + (\mu^{-2} - 1) z z^*) (1 + (\mu^{-4} - 1) z z^*)} dz dz^*.$$

Observación 6.3.12. Observemos que si $\mu = 1$ en la ecuación anterior se tiene el caso clásico como debe de ser.

Observación 6.3.13. Cálculos directos nos muestran que la siguientes igualdades se satisfacen

$$\begin{aligned}\frac{1 - \psi(z^*z)}{1 - \psi^{-1}(z^*z)} &= \frac{1 - \psi^2(zz^*)}{1 - zz^*} = \frac{\mu^4}{\mu^4 + (1 - \mu^4)zz^*}, \\ \frac{1 - \psi^{-1}(zz^*)}{1 - \psi(zz^*)} &= \frac{1 - \psi^{-2}(z^*z)}{1 - z^*z} = \frac{\mu^{-4}}{\mu^{-4} + (1 - \mu^{-4})z^*z}.\end{aligned}$$

Por lo que usando estas igualdades en las relaciones dadas en la proposición 6.3.9 tenemos que

$$dz^* z = \mu^{-2} \frac{1 - \psi(z^*z)}{1 - \psi^{-1}(z^*z)} z dz^*, \quad dz z^* = \mu^2 \frac{1 - \psi^{-1}(zz^*)}{1 - \psi(zz^*)} z^* dz.$$

Observación 6.3.14. El cálculo $\Omega(M)$ en la variedad base es el cálculo envolvente universal del cálculo de primer orden sobre $\mathcal{W}$, pues derivando las relaciones de conmutación dadas en (6.3.9), obtenemos la única relación dada en (6.32).

Proposición 6.3.10. Las siguientes relaciones de conmutación se satisfacen entre los elementos del cálculo $\Omega^1(M)$ y los elementos de $\mathcal{S}$:

$$dz (zz^*) = \frac{zz^*}{\mu^{-4} + (1 - \mu^{-4})zz^*} dz, \quad dz^* (zz^*) = \frac{zz^*}{\mu^4 + (1 - \mu^4)zz^*} dz^*.$$

Demostración. En efecto, los resultados se obtienen al aplicar las reglas de conmutación anteriores. $\square$

Observación 6.3.15. Por la definición del automorfismo ψ y usando la primera relación dada en (6.15) tenemos que

$$\begin{aligned}\psi^2(zz^*) &= \psi\left(\frac{zz^*}{\mu^2 + (1 - \mu^2)zz^*}\right) = w^* \frac{zz^*}{\mu^2 + (1 - \mu^2)zz^*} w \\ &= w^* w \frac{z^*z}{\mu^2 + (1 - \mu^2)z^*z} = \frac{z^*z}{\mu^4 + (1 - \mu^4)z^*z}.\end{aligned}$$

Por lo tanto, es claro que las relaciones de conmutación entre los elementos de $\Omega^1(M)$ y los elementos de $\mathcal{S}$ están dados por las siguientes fórmulas:

$$dz h = \psi^{-2}(h) dz, \quad dz^* h = \psi^2(h) dz^*,$$

para todo $h \in \mathcal{S}$.

6.4 Planos hiperbólicos cuánticos

En esta sección, daremos una generalización de planos hiperbólicos cuánticos inspirados en el ejemplo presentado en la sección anterior. Para este fin, consideraremos una $*$ -álgebra conmutativa $\mathcal{S}$ (denotada de igual manera que antes para seguir la analogía) generada por un único elemento positivo $\varrho = \varrho^*$ y equipada con un $*$ -automorfismo $\psi: \mathcal{S} \rightarrow \mathcal{S}$.

Con el fin de poder generalizar la construcción de la sección anterior será fundamental suponer que $\mathcal{S}$ admite un cálculo funcional y en particular, interpretaremos a los elementos de $\mathcal{S}$ como funciones de ϱ definidas sobre su espectro.

Definición 6.4.1. Denotemos por $\mathcal{U}$ al álgebra que surge al considerar el producto cruzado de $\mathcal{S}$ con el álgebra del círculo asociada a la transformación ψ . Es decir, construimos $\mathcal{U}$ introduciendo un elemento unitario abstracto ϖ , tal que

$$\varpi h \varpi^* = \psi(h), \quad (6.33)$$

para cada $h \in \mathcal{U}$.

Cabe observar que el automorfismo ψ de esta sección es el mismo que el de la sección anterior haciendo $\varpi = w^*$, o bien, igual al automorfismo inverso. En la definición, también podemos postular que se cumple la ecuación (6.33) para el generador $h = \varrho$, y después extendemos ψ a todo $\mathcal{U}$.

Con el fin de tener un análogo cuántico del modelo del disco de Poincaré, definiremos ‘coordenadas complejas’ en términos de los elementos ϱ y ϖ anteriores.

Definición 6.4.2. Definimos la ‘coordenada compleja’ $\mathfrak{z}$ y su estrella $\mathfrak{z}^*$ como los elementos del álgebra $\mathcal{U}$ dados por

$$\mathfrak{z} = \varrho \varpi, \quad \mathfrak{z}^* = \varpi^* \varrho. \quad (6.34)$$

Observación 6.4.1. Pensaremos al elemento ϱ como una coordenada radial abstracta y supondremos que su espectro $\sigma(\varrho)$ está contenido en el intervalo $[0, 1]$.

Entonces en la definición anterior, podemos imaginar que el elemento $\mathfrak{z}$ es una coordenada compleja que pertenece al disco unitario.

Proposición 6.4.1. Sea ψ el $*$ -automorfismo de $\mathcal{U}$ dado en la definición 6.4.1. Entonces se satisface la siguiente relación

$$\mathfrak{z}\mathfrak{z}^* = \varrho^2 = \psi(\mathfrak{z}^*\mathfrak{z}). \quad (6.35)$$

Demostración. En efecto, dado que $\psi(\mathfrak{z}) = \varpi \mathfrak{z} \varpi^*$ y $\psi(\mathfrak{z}^*) = \varpi \mathfrak{z}^* \varpi^*$ se tiene que

$$\psi(\mathfrak{z}) = \varpi \varrho \varpi \varpi^* = \varpi \varrho \quad \psi(\mathfrak{z}^*) = \varpi \varpi^* \varrho \varpi^* = \varrho \varpi^*.$$

Por lo tanto,

$$\psi(\mathfrak{z}^* \mathfrak{z}) = \psi(\mathfrak{z}^*) \psi(\mathfrak{z}) = \varrho \varpi^* \varpi \varrho = \varrho^2 = \mathfrak{z} \mathfrak{z}^*,$$

como se quería demostrar. $\square$

En todo lo que continúa, vamos a suponer que existe un homeomorfismo f del intervalo $[0, 1]$ que deja fijos al cero y al uno y tal que $f(\mathfrak{z}^* \mathfrak{z}) = \mathfrak{z} \mathfrak{z}^*$, es decir, se satisface la siguiente relación

$$\mathfrak{z} \mathfrak{z}^* = \varrho^2 = \psi(\mathfrak{z}^* \mathfrak{z}) = f(\mathfrak{z}^* \mathfrak{z}). \quad (6.36)$$

Y también que existe una única función continua positiva $q: [0, 1] \rightarrow \mathbb{R}$ tal que

$$f(t) = tq(t)^2, \quad \forall t \in [0, 1]. \quad (6.37)$$

Observemos que lo anterior se cumple si f es una función analítica o más generalmente diferenciable en $t = 0$.

Definición 6.4.3. Denotamos por $\mathcal{V}$ a la $*$ -subálgebra unital de $\mathcal{U}$ generada por los elementos $\mathfrak{z}$ y $\mathfrak{z}^*$. Supongamos que al igual que $\mathcal{S}$ esta subálgebra permite un cálculo funcional tal que $\mathcal{S} \subseteq \mathcal{V}$.

Observación 6.4.2. De manera equivalente, podemos definir $\mathcal{V}$ como la mínima subálgebra de $\mathcal{U}$ que contiene a $\mathcal{S}$ y a las coordenadas $\mathfrak{z}$ y $\mathfrak{z}^*$.

Lema 6.4.2. La subálgebra $\mathcal{V}$ es ψ -invariante, es decir $\psi(\mathcal{V}) = \mathcal{V}$. Más aún, se satisfacen las siguientes relaciones

$$\psi(\mathfrak{z}) = q(\mathfrak{z} \mathfrak{z}^*) \mathfrak{z} = \mathfrak{z} q(\mathfrak{z}^* \mathfrak{z}), \quad \psi(\mathfrak{z}^*) = q(\mathfrak{z}^* \mathfrak{z}) \mathfrak{z}^* = \mathfrak{z}^* q(\mathfrak{z} \mathfrak{z}^*). \quad (6.38)$$

Demostración. En efecto, usando (6.36), la definición de la función f dada en (6.37) y que $\mathfrak{z} = \varrho \varpi$ se tiene que

$$\begin{aligned} \varrho^2 = f(\mathfrak{z}^* \mathfrak{z}) &= \mathfrak{z}^* \mathfrak{z} q(\mathfrak{z}^* \mathfrak{z})^2 = \mathfrak{z}^* q(\mathfrak{z} \mathfrak{z}^*)^2 \mathfrak{z} = \varpi^* \varrho q(\varrho^2)^2 \varrho \varpi = \varpi^* \varrho^2 q(\varrho^2)^2 \varpi \\ &= \psi^{-1}[f(\varrho^2)], \end{aligned}$$

de manera que si aplicamos ψ y raíz cuadrada tendremos que $\psi(\varrho) = \varrho q(\varrho^2)$. Por otra parte, por la definición de la coordenada $\mathfrak{z}$, se tiene que

$$\psi(\mathfrak{z}) = \psi(\varrho \varpi) = \psi(\varrho) \varpi.$$

Por lo tanto $\psi(\mathfrak{z}) = q(\varrho^2) \varrho \varpi = q(\mathfrak{z} \mathfrak{z}^*) \mathfrak{z}$. Finalmente, usando que ψ es un $*$ -automorfismo obtenemos la relación faltante. $\square$

Comentario. Podemos construir todo un grupo de simetrías generalizando el procedimiento anterior y considerando homeomorfismos $g: [0, 1] \rightarrow [0, 1]$ que conmutan con el f inicial.

Por definición, cada uno de ellos induce un $*$ -automorfismo ψ_g of $\mathcal{U}$ que actúa trivialmente en ϖ . Si asumimos la misma condición de regularidad: $g(t) = ts(t)^2$ para g , obtenemos $\psi_g(\varrho) = (\psi_g(\varrho^2))^{1/2} = (g(\varrho^2))^{1/2} = \varrho s(\varrho^2)$ y por lo tanto que $\psi_g(\mathfrak{z}) = \psi_g(\varrho)\varpi = s(\varrho^2)\mathfrak{z} = s(\mathfrak{z}\mathfrak{z}^*)\mathfrak{z}$.

Pensaremos al elemento ϖ como el generador del álgebra $\mathcal{A}$ que representa al círculo. Esta álgebra es conmutativa y la estructura de grupo en el círculo es descrita a través del coproducto $\phi_{\mathcal{A}}$ definido como $\phi_{\mathcal{A}}(\varpi) = \varpi \otimes \varpi$ ($\phi_{\mathcal{A}}$ es el pull-back del producto en el círculo).

Proposición 6.4.3. Consideremos el álgebra $\mathcal{A}$ que corresponde al círculo y con generador ϖ . El coproducto $\phi_{\mathcal{A}} : \mathcal{A} \rightarrow \mathcal{A} \otimes \mathcal{A}$ se extiende a un único $*$ -homomorfismo $F : \mathcal{U} \rightarrow \mathcal{U} \otimes \mathcal{A}$ tal que

$$F(\mathfrak{z}) = \mathfrak{z} \otimes 1.$$

Es decir, $\phi_{\mathcal{A}}$ se extiende a una única co-acción derecha de $\mathcal{A}$ en $\mathcal{U}$.

Demostración. En efecto, veamos que el $*$ -homomorfismo F es una co-acción derecha de $\mathcal{A}$ en $\mathcal{U}$. Observemos que $\epsilon_{\mathcal{A}}(\varpi) = 1$ y por lo tanto

$$\begin{aligned} (\text{id} \otimes \epsilon_{\mathcal{A}})F(\mathfrak{z}) &= \mathfrak{z} \otimes \epsilon_{\mathcal{A}}(1) = \mathfrak{z} \otimes 1 \cong \mathfrak{z}, \\ (\text{id} \otimes \epsilon_{\mathcal{A}})F(\varpi) &= \varpi \otimes \epsilon_{\mathcal{A}}(\varpi) = \varpi \otimes 1 \cong \varpi, \end{aligned}$$

y como $\mathcal{U}$ está generado por $\mathfrak{z}$ y ϖ entonces las igualdades anteriores se satisfacen para todo $\mathcal{U}$. Por otro lado se tiene que

$$\begin{aligned} (\text{id} \otimes \phi_{\mathcal{A}})F(\mathfrak{z}) &= \mathfrak{z} \otimes \phi_{\mathcal{A}}(1) = \mathfrak{z} \otimes 1 \otimes 1 = F(\mathfrak{z}) \otimes 1 = (F \otimes \text{id})F(\mathfrak{z}), \\ (\text{id} \otimes \phi_{\mathcal{A}})F(\varpi) &= \varpi \otimes \phi_{\mathcal{A}}(\varpi) = \varpi \otimes \varpi \otimes \varpi = F(\varpi) \otimes \varpi = (F \otimes \text{id})F(\varpi). \end{aligned}$$

Por lo tanto se satisface la ecuación $(\text{id} \otimes \phi_{\mathcal{A}})F = (F \otimes \text{id})F$ en $\mathcal{U}$. Por lo tanto F es una co-acción derecha de $\mathcal{A}$ en $\mathcal{U}$.

La unicidad de la co-acción la determina por completo la definición en el generador $\mathfrak{z}$. $\square$

Observación 6.4.3. Cálculos directos nos llevan a que

$$F(\psi(\mathfrak{z})) = (\psi \otimes \text{id})F(\mathfrak{z}).$$

Observación 6.4.4. El álgebra $\mathcal{V}$ consiste de todos los elementos derecha invariantes bajo la co-acción F construida anteriormente.

Proposición 6.4.4. La terna $(\mathcal{U}, F, \iota)$ es un haz principal cuántico con base representada por el álgebra $\mathcal{V}$, el grupo estructural representado por el álgebra $\mathcal{A}$, y donde la co-acción F está definida en la proposición 6.4.3 y $\iota : \mathcal{V} \rightarrow \mathcal{U}$ es el mapeo inclusión. $\square$

6.5 Representaciones en espacios de Hilbert

Consideramos la $*$ -álgebra $\mathcal{V}$ generada por $\mathfrak{z}$ y $\mathfrak{z}^*$ definida anteriormente. Y sea $f : [0, 1] \rightarrow [0, 1]$ una biyección analítica tal que $f(0) = 0$, $f(1) = 1$ y tal que la relación $f(\mathfrak{z}\mathfrak{z}^*) = \mathfrak{z}^*\mathfrak{z}$ se satisface.

Observación 6.5.1. Para nuestro ejemplo principal de la fibración de Hopf cuántica hiperbólica de $SU_\mu(1, 1)$ sobre el plano hiperbólico cuántico tenemos que la función f es una transformación de Möbius que deja fijos a los puntos 0 y 1 y que preserva el intervalo $(0, 1)$. En este caso, la función se define como

$$f(t) = \frac{t}{q + (1 - q)t}, \quad q = \mu^2 \quad (6.39)$$

la cual se puede reescribir de la siguiente manera

$$\frac{1}{1 - f(t)} = \frac{1}{q} \frac{1}{1 - t} + 1 - \frac{1}{q}.$$

Es decir, f se puede ver como una conjugación de una transformación lineal L por el elemento g , esto es, $f = g^{-1}Lg$, donde tenemos $g(t) = 1/(1 - t)$ y $L(t) = (1/q)t + (1 - 1/q)$.

Una primera clase de representaciones surge al considerar $\ell^2(\mathbb{Z})$ con su base canónica ortonormal $\{\cdots | -2\rangle, | -1\rangle, |0\rangle, |1\rangle, |2\rangle \cdots\}$ y postulando la acción de la forma

$$\mathfrak{z}|n\rangle = \lambda_n^{1/2}|n+1\rangle, \quad \mathfrak{z}^*|n\rangle = \lambda_{n-1}^{1/2}|n-1\rangle, \quad (6.40)$$

tal que se satisface la relación funcional dada. Además vamos a suponer que todos los λ son no negativos. Nuestros operadores se pueden ver como caso especial de los operadores bilaterales de shift con peso.

Notemos que

$$\mathfrak{z}^*\mathfrak{z}|n\rangle = \lambda_n|n\rangle, \quad \mathfrak{z}\mathfrak{z}^*|n\rangle = \lambda_{n-1}|n\rangle, \quad (6.41)$$

se satisface para cada $n \in \mathbb{Z}$ y dado que se cumple la relación funcional se tiene inmediatamente que

$$f(\lambda_n) = \lambda_{n+1}. \quad (6.42)$$

Es decir, los números $\{\cdots \lambda_{-2}, \lambda_{-1}, \lambda_0, \lambda_1, \lambda_2 \cdots\}$ están en la misma órbita bajo la acción de f en el intervalo $[0, 1]$. Además, como el 0 y el 1 son puntos fijos para f podemos excluirlos y asumir que $\lambda_n \in (0, 1)$ para cada $n \in \mathbb{Z}$.

El mapeo f podría tener algunos puntos fijos dentro del intervalo $(0, 1)$. Sin embargo, si $t \in (0, 1)$ no es un punto fijo entonces todos los elementos

$f^n(t)$ para $n \in \mathbb{Z}$ son mutuamente diferentes y forman una sucesión monótona. Es decir, si $t < f(t)$ entonces $f^n(t) < f^{n+1}(t)$ y la sucesión será creciente. Similarmente, si $t > f(t)$ entonces $f^n(t) > f^{n+1}(t)$ y la sucesión será decreciente. En ambos casos, los puntos límite para $f^n(t)$ cuando $n \rightarrow \pm\infty$ serán puntos fijos para f . Intercambiando f por f^{-1} cambiaremos entre los modos decrecientes y crecientes. Y por lo tanto, no existe órbita para f que pueda brincar entre los puntos fijos. En particular, cuando los únicos puntos fijos para f son el cero y el uno, entonces cada órbita en el intervalo $(0, 1)$ es infinita con cero y uno como sus únicos puntos límite. Más aún, todas las órbitas siempre son crecientes o decrecientes.

Si fijamos $|0\rangle$ como un vector de referencia, entonces las acciones sucesivas de $\mathfrak{z}$ y $\mathfrak{z}^*$ generan la entera base

$$|k\rangle = \frac{\mathfrak{z}^k |0\rangle}{\sqrt{\lambda_0 \lambda_1 \cdots \lambda_{k-1}}}, \quad |-k\rangle = \frac{\mathfrak{z}^{*k} |0\rangle}{\sqrt{\lambda_{-1} \cdots \lambda_{-k}}}, \quad (6.43)$$

del espacio de Hilbert $\ell^2(\mathbb{Z})$, donde $k \in \mathbb{N}$.

De esta manera, cada órbita infinita de f , da lugar, de manera natural, a una representación infinito dimensional del álgebra del plano hiperbólico cuántico.

Proposición 6.5.1. Para cada órbita infinita, la representación infinito dimensional construida es irreducible. $\square$

Nuestro siguiente paso es construir una realización de las representaciones en espacio de Hilbert por funciones holomorfas. Vamos a empezar con el escenario más simple, cuando la transformación f es tal que sus únicos puntos fijos son el 0 y el 1, y las órbitas son crecientes. En este caso, veremos que el dominio común para las funciones holomorfas correspondientes es el interior del disco unitario sin el 0.

Proposición 6.5.2. Las identificaciones

$$|k\rangle \leftrightarrow \frac{z^k}{\sqrt{\lambda_0 \lambda_1 \cdots \lambda_{k-1}}}, \quad |-k\rangle \leftrightarrow \frac{\sqrt{\lambda_{-1} \cdots \lambda_{-k}}}{z^k}, \quad (6.44)$$

para cada $k \in \mathbb{N}$, induce la realización del espacio de Hilbert por funciones holomorfas en $\mathbb{D} \setminus \{0\}$. En términos de esta realización, la variable compleja cuántica $\mathfrak{z}$ actúa simplemente como el operador de multiplicación por la coordenada z .

Demostración. En la proposición estamos identificando $|0\rangle$ con 1 y la variable cuántica compleja $\mathfrak{z}$ con la coordenada compleja z . Y por lo tanto, dado que

$\mathfrak{z}|k\rangle = \sqrt{\lambda_k}|k+1\rangle$ tenemos que $\mathfrak{z}|-1\rangle = \sqrt{\lambda_{-1}}|0\rangle$, es decir, $|-1\rangle \leftrightarrow \sqrt{\lambda_{-1}}/z$. Siguiendo inductivamente identificamos $|-k\rangle$ con $\sqrt{\lambda_{-1} \cdots \lambda_{-k}}/z^k$.

Así tenemos que los elementos en el espacio de Hilbert se escriben como

$$c_0 + \sum_{k>0} \left\{ \frac{c_k z^k}{\sqrt{\lambda_0 \lambda_1 \cdots \lambda_{k-1}}} + c_{-k} \sqrt{\lambda_{-1} \cdots \lambda_{-k}} \frac{1}{z^k} \right\}, \quad \sum_{k \in \mathbb{Z}} |c_k|^2 < \infty. \quad (6.45)$$

La fórmula de Cauchy-Hadamard nos da el anillo de convergencia de esta serie de Laurent

$$\frac{1}{R} = \overline{\lim}_{k \rightarrow \infty} \left(\frac{|c_k|}{\sqrt{\prod_{j=0}^{k-1} \lambda_j}} \right)^{1/k} = \overline{\lim}_{k \rightarrow \infty} \frac{|c_k|^{1/k}}{\left(\sqrt{\prod_{j=0}^{k-1} \lambda_j} \right)^{1/k}} = \overline{\lim}_{k \rightarrow \infty} |c_k|^{1/k} \leq 1,$$

por lo que el radio exterior de convergencia R es mayor o igual a 1. Por otro lado,

$$r = \overline{\lim}_{k \rightarrow \infty} \left(|c_{-k}| \sqrt{\lambda_{-1} \cdots \lambda_{-k}} \right)^{1/k} = \overline{\lim}_{k \rightarrow \infty} |c_{-k}|^{1/k} \lim_{k \rightarrow \infty} \left(\sqrt{\lambda_{-1} \cdots \lambda_{-k}} \right)^{1/k} = 0.$$

Así, el anillo de convergencia común en estas series de Laurent en el espacio de Hilbert está dado por $\mathbb{D} \setminus \{0\}$. En este anillo la convergencia es normal. $\square$

Vamos ahora a definir los números

$$\Lambda_n = \prod_{k \geq n} \lambda_k.$$

Por lo tanto las identificaciones dadas en la proposición 6.5.2 ahora se simplifican como

$$|n\rangle \leftrightarrow \sqrt{\frac{\Lambda_n}{\Lambda_0}} z^n.$$

Estas expresiones ahora son válidas para todo $n \in \mathbb{Z}$. Los números Λ_n forman una sucesión monótona en $\mathbb{Z}$ con 0 y 1 como puntos límites de la sucesión. Claramente $\lambda_n = \Lambda_n / \Lambda_{n+1}$. Las potencia de z son mutuamente ortogonales con

$$\langle z^n | z^n \rangle = \frac{\Lambda_0}{\Lambda_n}.$$

Por lo tanto se tiene que el kernel reproductor para este espacio de Hilbert está dado por

$$K(\bar{w}, z) = \frac{1}{\Lambda_0} \sum_{k \in \mathbb{Z}} \Lambda_k (\bar{w} z)^k. \quad (6.46)$$

Observación 6.5.2. Este es un caso especial de la situación cuando el kernel se obtiene al componer el producto $\bar{w}z$ con una función holomorfa de un solo argumento:

$$K(\bar{w}, z) = \mathbb{K}(\bar{w}z).$$

Llamamos a la función $\mathbb{K}$ kernel generador. A la función K de dos argumentos también se le conoce como núcleo o kernel generador.

Si fijamos w y hacemos correr la variable z , obtenemos la función de onda $|w\rangle$ asociada al punto w , lo cual podemos escribir en término de los kets iniciales como

$$|w\rangle = \frac{1}{\sqrt{\Lambda_0}} \sum_{k \in \mathbb{Z}} \sqrt{\Lambda_k} \bar{w}^k |k\rangle. \quad (6.47)$$

Estos vectores generan un espacio lineal denso en el espacio de Hilbert $\mathcal{H}$ y satisfacen que

$$\mathfrak{z}^*|w\rangle = \bar{w}|w\rangle.$$

En particular vemos que tanto el espectro de $\mathfrak{z}$ como el de $\mathfrak{z}^*$ cubren el entero disco $\mathbb{D}$, como debería de ser.

Vale la pena recordar que la identidad

$$\langle u|v\rangle = K(\bar{v}, u)$$

siempre se cumple en espacios de Hilbert de funciones holomorfas. También recordemos que cuando tenemos espacios de Hilbert de funciones holomorfas sobre un dominio Ω , siempre es posible asociar de manera natural un estado coherente a cada punto de Ω .

$$\omega \leftrightarrow \frac{|\omega\rangle\langle\omega|}{\langle\omega|\omega\rangle}.$$

Los estados coherentes son vistos como ‘puntos estocásticos cuánticos’ en esta interpretación. La métrica de los operadores de Hilbert-Schmidt en $\mathcal{H}$ induce una métrica en Ω .

Proposición 6.5.3. La métrica inducida en Ω está dada por

$$ds^2 = \left\{ \frac{\partial^2}{\partial \bar{w} \partial w} \log K(\bar{w}, w) \right\} \delta \bar{w} \delta w. \quad (6.48)$$

Demostración. La distancia, como forma cuadrática sobre el desplazamiento

infinitesimal está dado por

$$\begin{aligned}
& \frac{1}{2} \text{Tr} \left\{ \left(\frac{|w + \delta w\rangle \langle w + \delta w|}{\langle w + \delta w | w + \delta w \rangle} - \frac{|w\rangle \langle w|}{\langle w | w \rangle} \right)^2 \right\} \sim ds^2 = \\
& = \frac{1}{2} \left(\text{Tr} \left(\frac{|w + \delta w\rangle \langle w + \delta w|}{\langle w + \delta w | w + \delta w \rangle} \right) + \text{Tr} \left(\frac{|w\rangle \langle w|}{\langle w | w \rangle} \right) - \frac{2 |\langle w + \delta w | w \rangle|^2}{\langle w + \delta w | w + \delta w \rangle \langle w | w \rangle} \right) \\
& = \frac{1}{2} \left(\frac{K(\overline{w + \delta w}, w + \delta w) K(\bar{w}, w) - |K(\bar{w}, w + \delta w)|^2}{K(\overline{w + \delta w}, w + \delta w) K(\bar{w}, w)} \right).
\end{aligned}$$

Desarrollando en series de potencias la función $K(\overline{w + \delta w}, w + \delta w)$ alrededor del punto $(\bar{w}, w)$, y además viendo la función $K(\bar{w}, w + \delta w)$ como función que solo dependen de la segunda variable podemos desarrollar alrededor del punto w de modo que los cálculos anteriores se transforman de la siguiente manera

$$\begin{aligned}
ds^2 &= \frac{1}{K(\bar{w}, w)^2} \left(K(\bar{w}, w) \frac{\partial^2}{\partial \bar{w} \partial w} K(\bar{w}, w) \delta w \delta \bar{w} - \right. \\
&\quad \left. - \frac{\partial}{\partial w} K(\bar{w}, w) \frac{\partial}{\partial \bar{w}} K(\bar{w}, w) \delta w \delta \bar{w} + \dots \right) \\
&\quad \left(1 + \frac{\partial}{\partial w} \log K(\bar{w}, w) \delta w + \frac{\partial}{\partial \bar{w}} \log K(\bar{w}, w) \delta \bar{w} + \dots \right)^{-1} \\
&\approx \frac{2}{K(\bar{w}, w)^2} \left(K(\bar{w}, w) \frac{\partial^2}{\partial \bar{w} \partial w} K(\bar{w}, w) \delta w \delta \bar{w} - \right. \\
&\quad \left. - \frac{\partial}{\partial \bar{w}} K(\bar{w}, w) \frac{\partial}{\partial w} K(\bar{w}, w) \delta w \delta \bar{w} \right) = \left\{ \frac{\partial^2}{\partial \bar{w} \partial w} \log K(\bar{w}, w) \right\} \delta w \delta \bar{w},
\end{aligned}$$

donde hemos descartado los términos de más alto orden en δw y $\delta \bar{w}$ concluyendo así con la demostración. $\square$

Observación 6.5.3. Esta es una fórmula universal válida para espacios de Hilbert de funciones holomorfas. La interpretación de la mecánica cuántica es que el cuadrado del elemento de la distancia es proporcional a la probabilidad de no transición de un estado a otro. Es decir, si $\varphi, \xi \in \mathcal{H}$ son vectores unitarios entonces la fórmula

$$\text{Tr} \{ (|\varphi\rangle \langle \varphi| - |\xi\rangle \langle \xi|)^2 \} = 2(1 - |\langle \varphi | \xi \rangle|^2)$$

se satisface.

6.5.1 Kernel reproductor y métrica inducida en el modelo cuántico hiperbólico

Vamos a ilustrar todo esto considerando la función dada en (6.39) del modelo cuántico hiperbólico dado por $SU_\mu(1, 1)$ ya desarrollado. Calcularemos el kernel reproductor y la métrica cuántica inducida. Sin embargo, antes de continuar vamos a definir algunas q -expresiones que aparecerán a lo largo de esta sección.

Definimos los siguientes q -números

$$(k)_q = \frac{1 - q^k}{1 - q} = \begin{cases} 1 + q + \cdots + q^{k-1} & k > 0 \\ 0 & k = 0 \\ -q^{-1} - \cdots - q^k & k < 0 \end{cases} \quad (6.49)$$

donde suponemos que q es central en el álgebra e invertible.

Observación 6.5.4. Por la definición de los q -números podemos ver que se satisfacen las siguientes dos fórmulas:

$$(k + j)_q = (k)_q + q^k(j)_q, \quad q^k = (k + 1)_q - (k)_q. \quad (6.50)$$

Definimos también los q -símbolos como

$$(z \mid q)_k = \prod_{j=0}^{k-1} (1 - zq^j) \quad (6.51)$$

y su versión con productos infinitos

$$(z \mid q)_\infty = \prod_{j=0}^{\infty} (1 - zq^j) \quad (6.52)$$

que para que esté bien definido, se requieren algunas condiciones apropiadas de convergencia. Por ejemplo si q es un número complejo dentro del disco unitario $\mathbb{D}$, entonces el producto converge absolutamente y normalmente para todo z en el plano complejo.

Observación 6.5.5. Notemos que los símbolos finitos se pueden escribir en términos de los infinitos de la siguiente manera

$$(z \mid q)_n = (z \mid q)_\infty / (zq^n \mid q)_\infty.$$

Podemos extender la definición de los símbolos de tal manera que queden definidos para todo los números $n \in \mathbb{Z}$ y si $q \in (0, 1)$ para todos los números complejos.

Observación 6.5.6. Notemos que siempre se cumple lo siguiente

$$(0 \mid q)_\alpha = (z \mid q)_0 = 1.$$

Además por la definición escribimos $(z \mid q)_{-\alpha} = (z \mid q)_\infty / (zq^{-\alpha} \mid q)_\infty$ y también $(q^{-\alpha}z \mid q)_\alpha = (q^{-\alpha}z \mid q)_\infty / (q^{-\alpha}zq^\alpha \mid q)_\infty$, de tal manera que la siguiente fórmula de reflexión se satisface

$$(z \mid q)_{-\alpha} (q^{-\alpha}z \mid q)_\alpha = 1.$$

Para calcular el kernel reproductor del espacio de Hilbert $\mathcal{H}$ cuando la función está dada por (6.39) fijamos λ_0 como un punto origen de tal forma que su órbita está dada por

$$\lambda_k = \frac{\lambda_0}{(1 - q^k)\lambda_0 + q^k} = \frac{1}{1 - q^k \varkappa}, \quad \varkappa = 1 - \frac{1}{\lambda_0}$$

donde $k \in \mathbb{Z}$. De esta manera tenemos que los productos que figuran en la construcción del espacio de Hilbert se escriben como

$$\begin{aligned} \prod_{k=0}^{n-1} \frac{1}{\lambda_k} &= \prod_{k=0}^{n-1} (1 - q^k \varkappa) = (\varkappa \mid q)_n, \\ \prod_{k=1}^n \lambda_{-k} &= \prod_{k=1}^n \frac{1}{1 - q^{-k} \varkappa} = (\varkappa \mid q)_{-n}. \end{aligned}$$

Por lo que la serie de Laurent para el kernel generador está dada por

$$\mathbb{K}(z) = \sum_{n \in \mathbb{Z}} (\varkappa \mid q)_n z^n. \quad (6.53)$$

Recordemos la suma de Ramanujan para series infinitas dobles:

$$\sum_{n \in \mathbb{Z}} \frac{(a \mid q)_n}{(b \mid q)_n} z^n = \frac{(az \mid q)_\infty (q/az \mid q)_\infty (q \mid q)_\infty (b/a \mid q)_\infty}{(z \mid q)_\infty (b/az \mid q)_\infty (b \mid q)_\infty (q/a \mid q)_\infty}, \quad (6.54)$$

la cual converge para el anillo

$$|b/a| < |z| < 1. \quad (6.55)$$

Haciendo $b = 0$ en la suma anterior tenemos el caso particular que nos interesa:

$$\sum_{n \in \mathbb{Z}} (a \mid q)_n z^n = \frac{(az \mid q)_\infty (q/az \mid q)_\infty (q \mid q)_\infty}{(z \mid q)_\infty (q/a \mid q)_\infty}. \quad (6.56)$$

Proposición 6.5.4. La serie de Laurent para el kernel generador se suma en lo siguiente

$$\mathbb{K}(z) = \frac{(\varkappa z \mid q)_\infty (q/\varkappa z \mid q)_\infty (q \mid q)_\infty}{(z \mid q)_\infty (q/\varkappa \mid q)_\infty}. \quad (6.57)$$

Demostración. En la fórmula (6.56) sustituimos a por $\varkappa$ de manera que la suma dada en (6.53) se convierte en la expresión de la proposición. $\square$

Proposición 6.5.5. La métrica cuántica inducida en $\mathbb{D} \setminus \{0\}$ está dada por

$$ds^2 = \delta w \delta \bar{w} \left\{ \sum_{k \geq 0} \frac{q^k}{(1 - q^k |w|^2)^2} - \varkappa \sum_{n \in \mathbb{Z}} \frac{q^n}{(1 - \varkappa q^n |w|^2)^2} \right\}. \quad (6.58)$$

Demostración. Aplicamos la fórmula (6.48) al kernel reproductor dado por el kernel generador calculado en (6.57), de manera que

$$\begin{aligned} ds^2 &= \delta w \delta \bar{w} \left\{ \frac{\partial^2}{\partial w \partial \bar{w}} \log \frac{(\varkappa |w|^2 \mid q)_\infty (q/\varkappa |w|^2 \mid q)_\infty (q \mid q)_\infty}{(|w|^2 \mid q)_\infty (q/\varkappa \mid q)_\infty} \right\} \\ &= \delta w \delta \bar{w} \frac{\partial^2}{\partial w \partial \bar{w}} \{ \log(\varkappa |w|^2 \mid q)_\infty + \log(q/\varkappa |w|^2 \mid q)_\infty + \log(q \mid q)_\infty \\ &\quad - \log(|w|^2 \mid q)_\infty - \log(q/\varkappa \mid q)_\infty \} \\ &= \delta w \delta \bar{w} \frac{\partial^2}{\partial w \partial \bar{w}} \{ \log(\varkappa |w|^2 \mid q)_\infty + \log(q/\varkappa |w|^2 \mid q)_\infty - \log(|w|^2 \mid q)_\infty \} \\ &= \delta w \delta \bar{w} \frac{\partial}{\partial \bar{w}} \sum_{k \geq 0} \left(\frac{-\varkappa \bar{w} q^k}{1 - \varkappa |w|^2 q^k} + \frac{q^{1+k}}{\varkappa |w|^2 w - q^{1+k} w} + \frac{\bar{w} q^k}{1 - |w|^2 q^k} \right) \\ &= \delta w \delta \bar{w} \sum_{k \geq 0} \left(\frac{-\varkappa q^k}{(1 - \varkappa q^k |w|^2)^2} - \frac{\varkappa q^{1+k}}{(\varkappa |w|^2 - q^{1+k})^2} + \frac{q^k}{(1 - |w|^2 q^k)^2} \right) \\ &= \delta w \delta \bar{w} \left(\sum_{k \geq 0} \frac{q^k}{(1 - |w|^2 q^k)^2} - \varkappa \sum_{n \in \mathbb{Z}} \frac{q^n}{(1 - \varkappa |w|^2 q^n)^2} \right), \end{aligned}$$

lo que termina la demostración. $\square$

Observación 6.5.7. La métrica exhibe dos divergencias importantes. La primera de ellas se presenta en el horizonte cuando $|w| = 1$, tal como ocurre en el modelo clásico del plano hiperbólico de Poincaré, relacionado con la infinidad del espacio. la segunda divergencia es totalmente cuántica, ésta ocurre cuando $w = 0$ y es promovida por la parte negativa de la segunda suma, el espacio se vuelve más y más abundante cuando nos acercamos al único punto clásico, comportandose como una fuente de energía infinita.

En el artículo [12] exploramos el caso cuando el álgebra $\mathcal{V} = * - \text{alg}\{1, z\}$ satisface la relación

$$\mathfrak{z}^* \mathfrak{z} = f(\mathfrak{z} \mathfrak{z}^*),$$

donde $f : [0, 1] \rightarrow [0, 1]$ es una función inyectiva no sobreyectiva tal que $f(1) = 1$ y $f(0) \in (0, 1]$. Para este caso se tiene una realización del espacio de Hilbert $\mathcal{H}$ por funciones holomorfas cuyo dominio es el entero disco unitario $\mathbb{D}$ y donde la variable compleja cuántica $\mathfrak{z}$ actúa simplemente como el operador multiplicación por la coordenada z .

También en [13] podemos encontrar a detalle la presentación de todas las funciones en este modelo cuya geometría retiene las mismas simetrías que en el caso clásico. Estas funciones f resultan ser transformaciones de Möbius parabólicas tal que

$$\frac{1}{1-f(t)} = \frac{1}{1-t} + \nu, \quad \nu = \frac{f(0)}{1-f(0)}. \quad (6.59)$$

Más explícitamente, la función f es de la forma

$$f(t) = \frac{(1-\nu)t + \nu}{1 + \nu - \nu t} \quad \longleftrightarrow \quad \begin{pmatrix} 1-\nu & \nu \\ -\nu & 1+\nu \end{pmatrix}. \quad (6.60)$$

Si ν toma cualquier valor real se tiene un grupo uniparamétrico de transformaciones de Möbius parametrizadas por ν . Además f se convierte en la transformación identidad en el límite $\nu = 0$, lo cual nos da un álgebra conmutativa y el modelo clásico de Poincaré del plano hiperbólico.

En el caso no conmutativo $\nu > 0$ el kernel reproductor para el espacio de Hilbert resultante está dado por

$$K(\bar{w}, z) = (1 - \bar{w}z)^{-\lambda}, \quad \lambda = 1 + \frac{1}{\nu}. \quad (6.61)$$

Y la métrica inducida en el espacio de Hilbert $\mathcal{H}$ es de la forma

$$ds^2 = \frac{\lambda \delta w \delta \bar{w}}{(1 - |w|^2)^2}, \quad (6.62)$$

la cual, salvo por un escalar, es la métrica hiperbólica clásica.

Observación 6.5.8. La extensión de Toeplitz la obtenemos tomando el límite $\nu \rightarrow \infty$.

6.6 Cálculo diferencial

¿Cómo construir el cálculo diferencial en $\mathcal{U}$? Presentaremos los principales resultados dados en el artículo [12]. Para construir un tal cálculo seguiremos el siguiente procedimiento.

Primeramente, observemos que $\mathcal{S}$ es un álgebra conmutativa, de manera que le corresponde una variedad clásica, y por lo tanto la podemos equipar con un cálculo diferencial clásico el cual denotaremos por $\mathcal{S}_\wedge$. Observemos que éste queda definido por una única relación dada por $\varrho d(\varrho) = d(\varrho)\varrho$. Nuestro automorfismo ψ se extiende a un único $*$ -automorfismo $\psi: \mathcal{S}_\wedge \rightarrow \mathcal{S}_\wedge$ que conmuta con la diferencial d , es decir,

$$d\psi(\rho) = \varpi d(\rho)\varpi^* = \psi(d\rho).$$

Podemos construir el cálculo diferencial en el haz P como el cálculo diferencial graduado $\mathcal{U}_\wedge = \Omega(P)$ dado por las siguientes relaciones:

$$\begin{aligned} \varpi\lambda &= \psi(\lambda)\varpi, & d(\varpi)\lambda &= (-)^{\partial\lambda}\psi(\lambda)d(\varpi), \\ \varpi d(\varpi) &= d(\varpi)\varpi, & d(\varpi)\varpi^* &= -\varpi d(\varpi^*). \end{aligned} \quad (6.63)$$

donde $\lambda \in \mathcal{S}_\wedge$. Restringimos el cálculo diferencial de $\mathcal{U}$ al álgebra $\mathcal{A}$ para de esta manera equiparla de un cálculo que denotaremos como $\bigcirc_\wedge$.

Ahora podemos definir $\Omega(M)$ como la $*$ -subálgebra diferencial generada por z y z^* . Observemos que

$$\begin{aligned} \psi(dz) &= \varpi(d(\rho)\varpi + \rho d\varpi)\varpi^* = \varpi d\rho + \varpi\rho d(\varpi)\varpi^* = \psi(d\rho)\varpi + \psi(\rho d\varpi) \\ &= \psi(d\rho)\varpi + \psi(\rho)d\varpi = d(\psi(\rho)\varpi) = d\psi(z), \end{aligned}$$

es decir,

$$d\psi(z) = \varpi dz\varpi^* = \psi(dz).$$

Por definición $\mathcal{V} = \Omega^0(M)$. Denotamos por $d_M: \Omega(M) \rightarrow \Omega(M)$ al cálculo diferencial correspondiente.

Lema 6.6.1. La co-acción $F: \mathcal{U} \rightarrow \mathcal{U} \otimes \mathcal{A}$ se extiende a un único homomorfismo diferencial $\widehat{F}: \Omega(P) \rightarrow \Omega(P) \widehat{\otimes} \bigcirc_\wedge$.

Demostración. Chequemos que $\widehat{F}$ es compatible con las relaciones dadas en (6.63) que definen a $\Omega(P)$.

$$\begin{aligned} \widehat{F}(\varpi d\varpi - d\varpi\varpi) &= (\varpi \otimes \varpi)(d\varpi \otimes \varpi + \varpi \otimes d\varpi) \\ &\quad - (d\varpi \otimes \varpi + \varpi \otimes d\varpi)(\varpi \otimes \varpi) \\ &= (\varpi d\varpi \otimes \varpi^2 + \varpi^2 \otimes \varpi d\varpi) - (d\varpi\varpi \otimes \varpi^2 + \varpi^2 \otimes d\varpi\varpi) \\ &= (\varpi d\varpi - d\varpi\varpi) \otimes \varpi^2 + \varpi^2 \otimes (\varpi d\varpi - d\varpi\varpi) = 0, \end{aligned}$$

por lo que $\widehat{F}$ es compatible con la relación $\varpi d\varpi - d\varpi\varpi = 0$. También tenemos que

$$\begin{aligned}\widehat{F}(\varpi\lambda - \psi(\lambda)\varpi) &= (\varpi \otimes \varpi)(\lambda \otimes 1) - (\psi(\lambda) \otimes 1)(\varpi \otimes \varpi) \\ &= \varpi\lambda \otimes \varpi - \psi(\lambda)\varpi \otimes \varpi = (\varpi\lambda - \psi(\lambda)\varpi) \otimes \varpi = 0,\end{aligned}$$

de manera que la relación $\varpi\lambda = \psi(\lambda)\varpi$, para todo $\lambda \in \Omega(M)$ es respetada por $\widehat{F}$. Y por último

$$\begin{aligned}\widehat{F}(d\varpi\lambda - (-)^{\partial\lambda}\psi(\lambda)d\varpi) &= (d\varpi \otimes \varpi + \varpi \otimes d\varpi)(\lambda \otimes 1) \\ &\quad - (-)^{\partial\lambda}(\psi(\lambda) \otimes 1)(d\varpi \otimes \varpi + \varpi \otimes d\varpi) \\ &= (d\varpi\lambda \otimes \varpi + (-1)^{\partial\lambda}\varpi\lambda \otimes d\varpi) - (-)^{\partial\lambda}(\psi(\lambda)d\varpi \otimes \varpi + \psi(\lambda)\varpi \otimes d\varpi) \\ &= (d\varpi\lambda - (-)^{\partial\lambda}\psi(\lambda)d\varpi) \otimes \varpi + (-)^{\partial\lambda}(\varpi\lambda - \psi(\lambda)\varpi) \otimes d\varpi = 0,\end{aligned}$$

donde $\lambda \in \Omega(M)$. Por lo tanto, las relaciones que definen a $\Omega(P)$ son completamente compatibles. $\square$

Observación 6.6.1. El homomorfismo $\widehat{F}$ es $*$ y preserva grados: en efecto, si $\sum a_0 da_1 \cdots da_n \in \Omega(P)$ entonces

$$\begin{aligned}\widehat{F}[(\sum a_0 da_1 \cdots da_n)^*] &= \sum (-1)^{n(n-1)/2} \widehat{F}[da_n^* \cdots da_1^* a_0^*] \\ &= \sum (-1)^{n(n-1)/2} dF(a_n)^* \cdots dF(a_1)^* F(a_0)^* \\ &= \sum [F(a_0) dF(a_1) \cdots dF(a_n)]^* \\ &= \widehat{F}(\sum a_0 da_1 \cdots da_n)^*.\end{aligned}$$

Dado que los elementos $\sum a_0 da_1 \cdots da_n$ cubren $\Omega^n(P)$ se concluye el resultado. Es claro que $\widehat{F}$ preserva grados.

Lo anterior se cumple más generalmente siempre que el cálculo diferencial esté generado por el álgebra inicial.

De esta manera tenemos un cálculo diferencial en el haz P definido por las relaciones (6.63).

Observación 6.6.2. El automorfismo ψ se extiende al cálculo diferencial sobre el haz y se cumple que

$$\widehat{F}\psi(\lambda) = (\psi \otimes \text{id})\widehat{F}(\lambda),$$

para todo $\lambda \in \Omega(P)$.

Observación 6.6.3. El álgebra diferencial $\Omega(M)$ es preservada por el automorfismo ψ , pues si $\lambda \in \Omega(M)$ entonces $\widehat{F}(\lambda) = \lambda \otimes 1$ y

$$\widehat{F}(\psi(\lambda)) = (\psi \otimes \text{id})(\lambda \otimes 1) = \psi(\lambda) \otimes 1.$$

Lema 6.6.2. Las siguientes reglas de conmutación se satisfacen:

$$\frac{1}{z}d_M(z) = \psi^{-1}\left\{d_M(z)\frac{1}{z}\right\}, \quad z^*d_M(z) = \psi^{-1}\{d_M(z)z^*\}, \quad (6.64)$$

$$\frac{1}{z^*}d_M(z^*) = \psi\left\{d_M(z^*)\frac{1}{z^*}\right\}, \quad zd_M(z^*) = \psi\{d_M(z^*)z\}. \quad (6.65)$$

Demostración. Haciendo $\lambda = \rho$ en la ecuación central de (6.63) se tiene que $d\varpi\rho = \psi(\rho)d\varpi$ y por lo tanto

$$\rho\varpi^*d\varpi = \varpi^*\psi(\rho)d\varpi = \varpi^*d\varpi\rho,$$

es decir, ρ conmuta con la expresión $\varpi^*d\varpi$. Así se tiene que

$$\begin{aligned} \frac{1}{z}d_M(z) &= \varpi^*\frac{1}{\varrho}(d\varrho\varpi + \varrho d\varpi) = \psi^{-1}\left(\frac{1}{\varrho}d\varrho\right) + \varpi^*d\varpi = \psi^{-1}\left(\frac{1}{\varrho}d\varrho + \varpi^*d\varpi\right) \\ &= \psi^{-1}\left(d\varrho\frac{1}{\varrho} + \varrho\varpi^*d\varpi\frac{1}{\varrho}\right) = \psi^{-1}\left((d\varrho\varpi + \varrho d\varpi)\varpi^*\frac{1}{\varrho}\right) = \psi^{-1}\left(d_M(z)\frac{1}{z}\right). \end{aligned}$$

Similarmente,

$$\begin{aligned} z^*d_M(z) &= \varpi^*\varrho(d\varrho\varpi + \varrho d\varpi) = \psi^{-1}(\varrho d\varrho) + \varpi^*\varrho^2d\varpi = \psi^{-1}(\varrho d\varrho + \varrho d\varpi\varpi^*\varrho) \\ &= \psi^{-1}((d\varrho\varpi + \varrho d\varpi)\varpi^*\varrho) = \psi^{-1}(d_M(z)z^*). \end{aligned}$$

Esto prueba el primer par de relaciones. Las segundas son simplemente las conjugadas de las primeras. $\square$

Observemos que el álgebra $\Omega(P)$ tiene dos componentes:

$$\Omega(P) = \{\Omega(M)\mathcal{A}\} \oplus \{\Omega(M)\mathcal{A}\}\varpi^*d\varpi, \quad (6.66)$$

por lo que las formas horizontales del haz están dadas por

$$\mathfrak{hor}(P) = \Omega(M)\mathcal{A}.$$

Veremos en la siguiente sección que la expresión $\varpi^*d\varpi$ define una conexión canónica en el haz P .

6.6.1 Conexiones, derivada covariante y curvatura

Vamos a considerar la expresión $\omega = \varpi^* d\varpi$. Este elemento es precisamente el mapeo de germen en el generador ϖ , es decir $\omega = \pi_{\mathcal{A}}(\varpi) = \kappa_{\mathcal{A}}(\varpi) d\varpi$.

Observación 6.6.4. El elemento ω es antihermitiano, es decir,

$$\omega^* = (d\varpi^*)\varpi = -\varpi^* d\varpi = -\omega.$$

Por otro lado, viendo a ω como elemento de $\Omega(P)$ se tiene que

$$\widehat{F}(\omega) = (\varpi^* \otimes \varpi^*)(d\varpi \otimes \varpi + \varpi \otimes d\varpi) = \omega \otimes 1 + 1 \otimes (\varpi^* d\varpi).$$

Por lo tanto, se puede definir naturalmente una conexión $\omega : \bigcirc_{\text{inv}} \rightarrow \Omega(P)$ en el haz P . Ésta asigna al germen $\pi_{\mathcal{A}}(\varpi)$ su representación en el cálculo diferencial del haz.

En efecto, observemos que aunque ω es antihermitiano, como mapeo se tiene que

$$\omega(\pi_{\mathcal{A}}(\varpi)^*) = -\omega(\pi_{\mathcal{A}}(\kappa_{\mathcal{A}}(\varpi)^*)) = -\omega(\pi_{\mathcal{A}}(\varpi)) = -\omega = \omega^* = \omega(\pi_{\mathcal{A}}(\varpi))^*.$$

Observación 6.6.5. La conexión ω superconmuta con todos los elementos de $\Omega(P)$. Pues para $\lambda \in \Omega(M)$ y $\varpi^n \in \mathcal{A}$, $n \in \mathbb{Z}$ se tiene que

$$(\lambda \varpi^n)(\varpi^* d\varpi) = \lambda \varpi^* d\varpi \varpi^n = \varpi^* \psi(\lambda) d\varpi \varpi^n = (-1)^{\partial \lambda} (\varpi^* d\varpi)(\lambda \varpi^n).$$

Y por lo tanto usando (6.66) se concluye lo deseado.

Observación 6.6.6. Dado que el álgebra $\mathcal{A}$ generada por el elemento ϖ es conmutativa se tiene que

$$\pi_{\mathcal{A}}(\varpi) \circ \varpi = \kappa_{\mathcal{A}}(\varpi) \pi_{\mathcal{A}}(\varpi) \varpi = \varpi^* \pi_{\mathcal{A}}(\varpi) \varpi = \pi_{\mathcal{A}}(\varpi).$$

Más aún, $\pi_{\mathcal{A}}(\varpi) \circ \varpi^n = \pi_{\mathcal{A}}(\varpi)$, para todo $n \in \mathbb{N}$.

Proposición 6.6.3. La conexión ω en el haz P es regular.

Demostración. Para cada $\varphi \in \mathfrak{hor}(P)$ existe una cantidad finita de $\lambda_k \in \Omega(M)$ tales que

$$\varphi = \sum_k \lambda_k \varpi^k.$$

Por lo que

$$\widehat{F}(\varphi) = \sum_k \widehat{F}(\lambda_k) F(\varpi)^k = \sum_k (\lambda_k \otimes 1) (\varpi^k \otimes \varpi^k) = \sum_k \lambda_k \varpi^k \otimes \varpi^k.$$

Por lo tanto, usando que la expresión ω conmuta gradualmente con todos los elementos de $\Omega(P)$ se tiene

$$\omega\varphi = (-1)^{\partial\varphi}\varphi\omega = (-1)^{\partial\varphi}\left(\sum_k \lambda_k \varpi^k\right)\omega = (-1)^{\partial\varphi}\sum_k \lambda_k \varpi^k \omega(\pi_{\mathcal{A}}(\varpi) \circ \varpi^k),$$

es decir, la conexión es regular. $\square$

Observación 6.6.7. Como comentamos en la Observación 5.2.6, es una propiedad universal para conexiones regulares que éstas superconmutan con las formas de $\Omega(M)$. Es decir,

$$\omega(\theta)\lambda = (-1)^{\partial\lambda}\lambda\omega(\theta), \quad (6.67)$$

para todo $\lambda \in \Omega(M)$.

Observación 6.6.8. Si la estructura dada por $\circ$ en el espacio $\bigcirc_{\text{inv}}$ de uno-formas izquierda invariantes en el grupo cuántico estructural $\bigcirc$ es trivial en el sentido de que se satisface la relación $\vartheta \circ a = \epsilon(a)\vartheta$, para todo $a \in \bigcirc$ y $\vartheta \in \bigcirc_{\text{inv}}$, entonces la regularidad para una conexión $\omega : \bigcirc_{\text{inv}} \rightarrow \Omega(P)$ es equivalente a que $\omega(\vartheta)$ superconmuta con las formas horizontales.

Observemos que lo anterior se satisface para grupos estructurales clásicos equipados con cálculos diferenciales clásicos.

Ahora analicemos como son todas las posibles conexiones regulares. Cada tal conexión debe ser de la forma $\omega + \xi$ donde ξ es un elemento antihermitiano de $\Omega^1(M)$, el ‘tensor de desplazamiento’ superconmuta con todas las formas horizontales. En otras palabras

$$\xi\varphi = (-1)^{\partial\varphi}\varphi\xi,$$

donde $\varphi \in \mathfrak{hor}(P)$. En particular, se tiene que ξ conmuta con ϖ , lo cual es equivalente a decir que ξ es invariante bajo el automorfismo ψ : $\xi = \psi(\xi)$. Y además ξ debe superconmutar con todos los elementos de $\Omega(M)$. Claramente estas dos condiciones son equivalentes a la superconmutatividad con $\mathfrak{hor}(P)$, y por la construcción del cálculo en el haz, se tiene la superconmutatividad con todo $\Omega(P)$. En resumen:

Lema 6.6.4. Existe una correspondencia afín natural entre las conexiones en el haz P y los elementos antihermitianos de $\Omega^1(M)$.

La conexión es regular, sí y solamente si, el desplazamiento ξ es ψ -invariante y supercentral en $\Omega(M)$. $\square$

Ahora calculemos la curvatura de tales conexiones. Primeramente observemos que si la conexión ω es regular entonces el tensor de desplazamiento ξ

superconmuta con todas las formas horizontales, en particular consigo mismo, lo cual nos lleva a que $\xi^2 = 0$. También por la observación (6.6.5) se tiene que $\omega\xi + \xi\omega = 0$. Por lo tanto, usando la definición 5.3.3 tenemos que

$$r_{\omega+\xi} = d(\omega + \xi) + (\omega + \xi)^2 = d(\xi) + \xi^2 + \xi\omega + \omega\xi = d(\xi), \quad (6.68)$$

donde hemos usado que $d\omega = d\pi(\varpi) = -\pi(\varpi)^2 = -\omega^2$. Observemos que dado que el grupo es abeliano, y la ecuación

$$\hat{F}r_\omega(\vartheta) = (r_\omega \otimes \text{id})\text{ad}(\vartheta)$$

se satisface, entonces la curvatura es de $\Omega^2(M)$. Sea $D = D_{\omega+\xi}$ la derivada covariante correspondiente. Ésta extiende a la diferencial $d_M: \Omega(M) \rightarrow \Omega(M)$ a una antiderivación hermitiana covariante de primer orden de $\mathfrak{hor}(P)$.

Como tal, es completamente determinada por su valor en el generador ϖ . De tal forma que

$$D(\varpi) = d\varpi - \varpi(\omega + \xi) = d\varpi - \varpi\varpi^*d\varpi - \varpi\xi = -\varpi\xi. \quad (6.69)$$

En este contexto particular comparamos el cuadrado de la derivada covariante con la curvatura:

$$D^2(\varpi) = -D(\varpi)\xi - \varpi d(\xi) = -\varpi r_{\omega+\xi}, \quad (6.70)$$

en concordancia con la teoría general.

6.6.2 Solución universal

En esta sección consideraremos el caso general donde no se pide ninguna propiedad de compatibilidad entre el automorfismo ψ y la diferencial d_M . Y donde admitimos conexiones regulares. Supondremos que $\Omega(M)$ es una $*$ -álgebra graduada diferencial y que $\psi: \Omega(M) \rightarrow \Omega(M)$ es un $*$ -automorfismo graduado. Construiremos el álgebra producto cruzado $\mathfrak{hor}(P)$ de la misma manera como antes, es decir, pidiendo que

$$\varpi\lambda = \psi(\lambda)\varpi,$$

para cada $\lambda \in \Omega(M)$.

La existencia de conexiones regulares está intrínsecamente relacionada con la posibilidad de extender la diferencial d_M a un mapeo derivada covariante $D: \mathfrak{hor}(P) \rightarrow \mathfrak{hor}(P)$, la cual es una antiderivación hermitiana de primer orden y que es covariante; en el sentido de que el siguiente diagrama es conmutativo:

$$\begin{array}{ccc}
\mathfrak{hor}(P) & \xrightarrow{F^\wedge} & \mathfrak{hor}(P) \otimes \mathcal{A} \\
D \downarrow & & \downarrow D \otimes \text{id} \\
\mathfrak{hor}(P) & \xrightarrow{F^\wedge} & \mathfrak{hor}(P) \otimes \mathcal{A}
\end{array} \tag{6.71}$$

Cada tal mapeo, está completamente determinado por su valor en ϖ , y debido a la covarianza, éste debe ser de la forma

$$D(\varpi) = -\varpi\xi, \tag{6.72}$$

donde $\xi \in \Omega^1(M)$. Como D es hermitiano, tenemos que

$$0 = D(1) = D(\varpi^*)\varpi + \varpi^*D(\varpi) = -(\xi^* + \xi),$$

es decir, ξ debe ser antihermitiano.

Proposición 6.6.5. Se satisface la siguiente igualdad:

$$\xi\lambda - (-)^{\partial\lambda}\lambda\xi = d_M(\lambda) - \psi^{-1}d_M\psi(\lambda), \tag{6.73}$$

para cada $\lambda \in \Omega(M)$.

Inversamente, si ξ es un elemento antihermitiano de primer orden en $\Omega(M)$ que satisface dicha igualdad, entonces la fórmula (6.72) define consistentemente una única antiderivación hermitiana de primer orden covariante $D: \mathfrak{hor}(P) \rightarrow \mathfrak{hor}(P)$ que extiende a la diferencial $d_M: \Omega(M) \rightarrow \Omega(M)$.

Demostración. Aplicando D a la relación del producto cruzado obtenemos

$$\begin{aligned}
0 &= -\varpi^* \{ D(\varpi)\lambda + \varpi d_M(\lambda) - d_M\psi(\lambda)\varpi - (-)^{\partial\lambda}\psi(\lambda)D(\varpi) \} = \\
&= \xi\lambda - d_M(\lambda) + \psi^{-1}d_M\psi(\lambda) - (-)^{\partial\lambda}\lambda\xi.
\end{aligned}$$

En otras palabras se satisface (6.73). Como el álgebra de formas horizontales está definida por la única relación de producto cruzado, una condición necesaria y suficiente para la consistencia y extendibilidad a D de manera única como una antiderivación, es la compatibilidad con esta única relación. $\square$

Observación 6.6.9. En la sección anterior consideramos como caso especial, cuando ψ preserva d_M y vimos que era equivalente a que ξ superconmutara con todos los elementos de $\Omega(M)$. Ahora vemos, que la situación general no es muy diferente: aunque ψ y d_M no conmutan, la diferencia entre las diferenciales d_M y $\psi^{-1}d_M\psi$ es, de cierta manera pequeña; es la derivación interior asociada a ξ .

Cabe recordar que las superderivaciones interiores siempre forman un ideal de la superálgebra de Lie de todas las superderivaciones. La fórmula de la proposición 6.73 fija completamente la extensión del automorfismo ψ , de $\mathcal{V} = \Omega^0(M)$ a todo $\Omega(M)$, todo esto gracias a que estamos suponiendo que $\Omega(M)$ está generada como álgebra diferencial por $\mathcal{V}$.

Ahora calculemos el operador de curvatura para tales conexiones. De acuerdo con la teoría general, la curvatura es un mapeo $\varrho_D: \mathcal{A} \rightarrow \mathfrak{hor}(P)$ que se define de manera única por la fórmula

$$D^2(\varphi) = -\varphi^{(0)}\varrho_D(\varphi^{(1)}),$$

donde $F^\wedge(\varphi) = \varphi^{(0)} \otimes \varphi^{(1)}$. En nuestro contexto, como el grupo del círculo es abeliano, la curvatura debe ser $\Omega(M)$ -valuada. Tenemos entonces

$$\varrho_D(\varpi) = -\varpi^* D^2(\varpi) = \varpi^* D(\varpi\xi) = \varpi^*(D(\varpi)\xi + \varpi d_M(\xi)) = d_M(\xi) - \xi^2.$$

Esto se extiende fácilmente a todas las potencias de ϖ :

Proposición 6.6.6. Las siguientes igualdades se satisfacen:

$$\varrho_D(\varpi^n) = \sum_{k=0}^{n-1} \psi^{-k} [d_M(\xi) - \xi^2], \quad \varrho_D(\varpi^{-n}) = -\sum_{k=1}^n \psi^k [d_M(\xi) - \xi^2], \quad (6.74)$$

para cada $n \in \mathbb{N}$.

Demostración. El cuadrado D^2 es una derivación. Entonces,

$$\begin{aligned} \varrho_D(\varpi^n) &= -\varpi^{-n} D^2(\varpi^n) = -\varpi^{-n} \sum_{k=0}^{n-1} \varpi^{n-k-1} D^2(\varpi) \varpi^k = \\ &= \varpi^{-n} \sum_{k=0}^{n-1} \varpi^{n-k-1} \varpi [d_M(\xi) - \xi^2] \varpi^k = \sum_{k=0}^{n-1} \psi^{-k} [d_M(\xi) - \xi^2]. \end{aligned}$$

Y similarmente

$$\begin{aligned} \varrho_D(\varpi^{-n}) &= -\varpi^n D^2(\varpi^{-n}) = -\varpi^n \sum_{k=0}^{n-1} \varpi^{1+k-n} D^2(\varpi^*) \varpi^{-k} = \\ &= \varpi^n \sum_{k=0}^{n-1} \varpi^{1+k-n} [\xi^2 - d_M(\xi)] \varpi^* \varpi^{-k} = \sum_{k=1}^n \psi^k [\xi^2 - d_M(\xi)], \end{aligned}$$

donde hemos usado el hecho de que D conmuta con $*$, y que la expresión $d_M(\xi) - \xi^2$ es antihermitiana. $\square$

De la teoría general, sabemos que la derivada covariante del operador curvatura es cero. Esto es la identidad cuántica de Bianchi para conexiones regulares. En nuestro caso, como la curvatura toma valores de $\Omega(M)$, esto significa que todos los $\varrho_D(\varpi^m)$, para $m \in \mathbb{Z}$, deben ser cerrados. Equivalentemente

$$d_M \psi^m [d_M(\xi) - \xi^2] = 0, \quad (6.75)$$

para cada $m \in \mathbb{Z}$. La más simple de estas identidades de Bianchi es $d_M(\xi^2) = 0$. La curvatura también conmuta con todos los elementos de $\Omega(M)$. En particular tenemos que

$$d_M(\xi)\xi = \xi d_M(\xi), \quad (6.76)$$

lo que es equivalente como mencionamos antes, a que ξ^2 es cerrada.

Observación 6.6.10. El operador curvatura define de manera natural, el mínimo cálculo diferencial compatible en el grupo estructural. Este cálculo esta basado en el ideal derecho $\mathcal{R}$ en el kernel de la counidad $\epsilon: \mathcal{A} \rightarrow \mathbb{C}$ que consiste precisamente de todos los elementos que son anulados por la curvatura.

También podemos considerar, en lugar de una conexión particular preferida, el espacio afín de todas ellas (incluyendo todas las posibles realizaciones via derivadas covariantes). Como vimos en las expresiones anteriores para la curvatura, en general, dicho cálculo podría ser arbitrariamente diferente al cálculo clásico. Incluyendo al cálculo universal sobre el álgebra del círculo (correspondiente al kernel trivial $\mathcal{R} = \{0\}$).

6.7 Conclusiones

En el estudio del grupo $SU_\mu(1, 1)$ a través de la teoría de haces principales cuánticos, pudimos ver que se trata de un modelo hiperbólico cuántico, cuya versión clásica es el haz hiperbólico de Hopf. El grupo estructural de este haz cuántico está dado por el círculo y la variedad base es un hiperboloide cuántico. Además, con la teoría de cálculos diferenciales cuánticos, hemos estudiado en detalle un simple y natural cálculo diferencial 3-dimensional sobre este grupo cuántico y sobre su parte clásica $U(1)$, donde el principal generador U del círculo y su diferencial dU , no conmutan.

Permitiendo que el álgebra que representa al grupo cuántico $SU_\mu(1, 1)$ se pueda extender, incluyendo las inversas y funciones apropiadas de los generadores α y α^* , pudimos hacer una generalización cuántica del modelo de Poincaré del plano hiperbólico clásico. Resultó que dicha álgebra se puede

escribir como producto cruzado entre el plano hiperbólico y el círculo, donde se cumple la única relación

$$z^*z = \psi(zz^*) = f(zz^*).$$

Aquí ψ es un automorfismo interno del álgebra asociado al elemento unitario y f es un homeomorfismo en el intervalo $[0, 1]$. Para el ejemplo estudiado del haz hiperbólico cuántico, f está dado por

$$f(zz^*) = \frac{zz^*}{\mu^2 + (1 - \mu^2)zz^*}.$$

El cálculo diferencial sobre el haz respalda la interpretación de que se trata de un modelo cuántico de Poincaré. Pues para los casos $\mu = \pm 1$ se recupera la métrica hiperbólica. Para el resto de los casos se tiene una pequeña modificación de la métrica que involucra al automorfismo ψ que como hemos visto representa la clase de las criaturas puramente cuánticas:

$$\eta_- \eta_+ = \mu^{-2} \frac{dz dz^*}{(1 - \psi^{-1}(zz^*))(1 - zz^*)}.$$

Una primer interrogante que nos surgió fue: ¿Para qué funciones f la relación principal que define al álgebra nos da versiones cuánticas de planos hiperbólicos y tal que su geometría retiene las simetrías clásicas de $SU(1, 1)$?

La respuesta fue que lo anterior se cumple precisamente para aquellas transformaciones de Möbius dadas por

$$f(t) = \frac{(1 - \nu)t + \nu}{-\nu t + (1 - \nu)}, \quad \nu \in (0, \infty].$$

Incluimos el caso límite cuando $\nu = \infty$. Para este caso, la función f no es homeomorfismo. Sin embargo, también describe una geometría hiperbólica representada por la extensión de Toeplitz. Toda esta familia de espacios se describe a detalle en el artículo [12].

Otro camino que abordamos en esta tesis fue considerar representaciones en espacios de Hilbert cuyos elementos son funciones holomorfas sobre cierto dominio común, para el caso de las álgebras que retienen las simetrías clásicas el dominio donde éstas están definidas corresponde al disco unitario. En el caso del álgebra que trabajamos, las de simetría cuántica, el espacio de Hilbert consta de funciones que se escriben como series de Laurent cuyo dominio común de todas ellas es el disco unitario sin el cero. Esta singularidad en cero muestra un fenómeno completamente cuántico en la métrica inducida.

Entonces este espacio cuántico asociado al álgebra $SU_\mu(1,1)$ lo podemos pensar como un cilindro holomorfo.

La métrica inducida muestra dos divergencias, una de ellas aparece en el horizonte, caso cuando consideramos puntos del círculo unitario, y que al igual que en el caso clásico se relaciona con el horizonte de la infinidad del espacio. Otra divergencia ocurre cuando estamos en el punto clásico asociado a 0, y en ese caso, el espacio se comporta como una fuente de energía ilimitada que se vuelve cada vez más abundante conforme nos acercamos a este punto.

En el modelo hiperbólico del disco de Poincaré estudiado, consideramos el cálculo arriba mencionado 3D de Woronowicz, ahí pudimos calcular las principales componentes geométricas como lo son la conexión, la derivada covariante y la curvatura. Un tema de interés separado, que va mas allá de los cálculos desarrollados aquí, es ver si el automorfismo ψ que define al álgebra se extiende al cálculo diferencial de más alto orden de manera que ψ y la diferencial d conmuten. Tal extensión, si existe, siempre es única.

De manera más general, consideramos ya no solo el modelo cuántico dado por el álgebra asociada a $SU_\mu(1,1)$, sino que consideramos cualquier álgebra que se pueda ver como producto cruzado de un álgebra $\mathcal{S}$ generada por un elemento positivo y equipada con un automorfismo ψ , donde la parte radial correspondiente al elemento positivo, se mueve con la apropiada función de zz^* o z^*z .

Luego construimos el haz hiperbólico como producto cruzado de $\mathcal{S}$ y el círculo representado por el álgebra $\mathcal{A}$ y cuyo generador es el elemento unitario asociado al automorfismo ψ . La base de este haz corresponde al espacio generado por la coordenada compleja z . Una manera directa de definir el cálculo diferencial en dicha álgebra, fue considerar un cálculo diferencial clásico en $\mathcal{S}$ tal que el automorfismo interno y la diferencial conmutan. Para este caso, el álgebra diferencial sobre la base se preserva por el automorfismo ψ y el cálculo diferencial en la base tiene relaciones de conmutación condicionadas por el automorfismo.

En estos espacios hay una conexión regular natural que superconmuta con todos los elementos del cálculo diferencial en la base.

Además, hemos explorado la condición de que el automorfismo interno del haz y la diferencial no conmuten. Aquí la diferencia entre la diferencial en la base y $\psi^{-1}d\psi$ es en cierto sentido pequeña.

En el artículo [12], hemos presentado en detalle todos los cálculos diferenciales que se le pueden asociar al círculo. Y se ha estudiado de manera más profunda el caso cuando el álgebra que representa al haz se puede ver como producto cruzado entre el álgebra conmutativa generada por una coordenada

real y el álgebra de Toeplitz. En este caso, ψ no es automorfismo sino homomorfismo y el generador U del álgebra de Toeplitz verifica que $U^*U = 1$ y $UU^* = \psi(1)$ es proyector ortogonal. Ahí se responde, cuál de tales cálculos diferenciales sobre el grupo estructural del haz, recupera el cálculo diferencial de $SU_\mu(1, 1)$.

Esta serie de modelos cuánticos hiperbólicos proporciona una base que nos permite construir y conocer las versiones cuánticas de todas las superficies de Riemann compactas de curvatura constante negativa. Para tal construcción será necesario factorizar dicho plano cuántico por la acción libre del grupo discreto infinito (isomorfo al grupo fundamental de la superficie). Aquí la construcción cuántica sigue la filosofía clásica, de ver estas superficies de Riemann como base de un haz principal, cuyo espacio total es el plano hiperbólico y el grupo estructural es dado por el grupo discreto correspondiente. En una publicación separada [14] se presentará esta construcción, que incluirá enlaces con las superficies cuánticas desarrolladas en [9, 10].

Otro interesante camino de exploración, es determinar cómo es la geometría del haz basado en el cálculo $4D$ de Woronowicz, o el cálculo mínimo admisible en las álgebras desarrolladas y exploradas en esta tesis. Aquí vale la pena mencionar que el cálculo mínimo admisible se puede definir para cualquier grupo cuántico compacto [4]. Es el mínimo cálculo bicovariante compatible con todas las funciones de transición para haces localmente triviales sobre variedades clásicas, cuyo grupo estructural es dado grupo cuántico G . También se puede definir como el mínimo cálculo bicovariante sobre el grupo G , que se proyecta al cálculo clásico sobre la parte clásica de G . En general, la parte clásica de un espacio cuántico, es una importante invariante geométrica del espacio. Como tal, se preserva por todas las simetrías clásicas del espacio cuántico (automorfismos del álgebra correspondiente). También, al nivel algebraico, se puede asociar naturalmente al *conmutante* del álgebra, es decir el ideal generado por todos los conmutadores—la parte clásica le corresponde el factor álgebra que por la construcción es el máximo factor álgebra conmutativo.

Para el grupo $SU_\mu(1, 1)$ su cálculo mínimo admisible es de dimensión infinita. Su estructura base será presentada en la versión extendida de [12].

Bibliografía

- [1] A. Connes, *Non-Commutative Differential Geometry*. I.H.E.S., Ext. Publ. Math. 62, (1986).
- [2] V. Drinfeld, *Quantum Groups*. Proc. Int. Cong. Math., Berkeley (1986) 793–820.
- [3] M. Đurđevich, *Quantum Geometry and New Concept of Space*. Artículo en línea: <http://www.matem.unam.mx/~micho>
- [4] M. Đurđevich, *Geometry of Quantum Principal Bundles I*. Commun. Math. Phys. 175 (3) (1996) 457–521.
- [5] M. Đurđevich, *Geometry of Quantum Principal Bundles II – Extended Version*. Rev. Math. Phys. 9 (5) (1997) 531–607.
- [6] M. Đurđevich, *Geometry of Quantum Principal Bundles III – Structure of Calculi and Around*. Algebras, Groups and Geometries, 27 (2010) 247–336.
- [7] M. Đurđevich, *Differential Structures on Quantum Principal Bundles*. Rep Math Phys 41 (1) 91–115 (1998).
- [8] M. Đurđevich, *Quantum Principal Bundles & Tannaka-Krein Duality Theory*. Rep Math Phys 38 (3) 313–324 (1996).
- [9] M. Đurđevich, *Quantum Riemann Surfaces I*. Algebras, Groups and Geometries. Vol 33, 323–368. (2016).
- [10] M. Đurđevich, *Quantum Riemann Surfaces II*. QAI Publication. (2024).
- [11] M. Đurđevich, *Quantum Groups as Realizations of Diagrammatic Symmetry Category*. Algebras, Groups and Geometries (2016).
- [12] M. Đurđevich, P. C. Lucio Peña, *Geometrical Structures On Quantum Hyperbolic Planes*. Varias versiones desarrolladas (2020–2024).

- [13] M. Đurdevich, P. C. Lucio Peña, *Generating Function for Quantum Hyperbolic Planes*. Scientific Legacy of Professor Zbigniew Oziewicz: Selected Papers from the International Conference “Applied Category Theory Graph-Operad-Logic”. Chapter 33, páginas 695-720. Series on Knots and Everything: Volume 75. World Scientific (Octubre de 2023).
- [14] M. Đurdevich, P. C. Lucio Peña, *Quantum Riemann Surfaces of Classical Symmetry*. En preparación (2024).
- [15] W. Heisenberg, *Der teil und Ganze/ La Parte y el Todo*. Traducción del Alemán por Rocio da Riva Muñoz. El Lago Ediciones (2004).
- [16] S. Kobayashi and K. Nomizu, *Foundations of Differential Geometry*. Interscience Publishers, New York, London (1963).
- [17] P. C. Lucio Peña, *Integración Invariante Sobre un Hiperboloide Cuántico*. Tesis de Maestría, UMSNH (2010).
- [18] J-P. Serre, *Lie Algebras and Lie Groups*. Benjamin (1965).
- [19] S. B. Sontz, *Principal Bundles, The Classical Case*. Universitext, Springer. (2015). ISBN 978-3-319-14764-2.
- [20] S. B. Sontz, *Principal Bundles: The Quantum Case*. Universitext, Springer. (2015). ISBN 978-3-319-15828-0.
- [21] E. Wigner, *The Unreasonable Effectiveness of Mathematics in the Natural Sciences*. Reprinted from Communications in Pure and Applied Mathematics, Vol. 13, No. I (February 1960). New York: John Wiley & Sons, Inc. Copyright 1960 by John Wiley & Sons, Inc.
- [22] S. L. Woronowicz, *Compact Matrix Pseudogroups*. Commun. Math. Phys. 111 (1987) 613–665.
- [23] S. L. Woronowicz, *Differential Calculus on Compact Matrix Pseudogroups (Quantum Groups)*. Commun. Math. Phys. 122 (1989) 125–170.
- [24] S. L. Woronowicz, *Twisted $SU(2)$ Group. An Example of a Non-Commutative Differential Calculus*. Publ. RIMS, Kyoto Univ., Vol. 23, 1 (1987) 117–23 (1), 117–181.
- [25] Transnational College of Lex, *¿Qué es la Mecánica Cuántica? Una Aventura en la Física*. Primera Edición: 2 de Julio de 2010. D.R. Universidad Nacional Autónoma de México, Ciudad Universitaria, Ciudad de México. ISBN: 978-607-02-1511-7.