



Rôle(s) du champ de fond antisymétrique en théorie des cordes.

Stéphane Fidanza

► To cite this version:

Stéphane Fidanza. Rôle(s) du champ de fond antisymétrique en théorie des cordes.. Physique [physics]. Ecole Polytechnique X, 2003. Français. NNT : . pastel-00000709

HAL Id: pastel-00000709

<https://pastel.hal.science/pastel-00000709>

Submitted on 21 Jul 2010

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

CENTRE DE PHYSIQUE THÉORIQUE – ECOLE POLYTECHNIQUE

THÈSE DE DOCTORAT

Spécialité : **Physique Théorique**

présentée par

Stéphane FIDANZA

pour obtenir le grade de

Docteur de l'Ecole Polytechnique

Sujet :

***Rôle(s) du champ de fond antisymétrique
en théorie des cordes***

soutenue le 19 novembre 2003 devant le jury composé de :

MM. Costas Bachas,
Branislav Jurčo, rapporteur,
Elias Kiritsis,
Ruben Minasian, directeur de thèse,
Daniel Waldram, rapporteur ,
& Paul Windey.

Remerciements

Cette thèse est l'aboutissement de trois années au centre de physique théorique de l'Ecole Polytechnique, même si elle n'est qu'une étape dans ma carrière professionnelle. Au-delà de l'enrichissement personnel que j'en retire (malheureusement uniquement immatériel), j'y ai pris beaucoup de plaisir, que ce soit dans les activités scientifiques, mon action au sein de l'association X'Doc ou tout ce qui a étoffé cette expérience, lui donnant la consistance de la vie. Je suis conscient que cela n'a été possible que grâce à la qualité et la sensibilité des personnes avec qui j'ai travaillé et vécu pendant trois ans. Je voudrais ici les en remercier, et leur dire que j'en conserverai toujours un souvenir chaleureux.

J'adresse tout spécialement mes remerciements à Ruben Minasian, qui m'a offert l'occasion d'apprendre la théorie des cordes et de montrer ce que je pouvais y faire. Merci de m'avoir guidé, tout en me laissant l'autonomie de choisir mon chemin et mes méthodes. Je remercie aussi chaleureusement les membres de mon jury, et tout particulièrement mes deux rapporteurs, qui ont accepté de lire ma thèse malgré la barrière de la langue.

Je pense aussi à l'ensemble des doctorants de physique et de mathématique, avec qui j'ai eu de si nombreuses discussions toujours enrichissantes. Je suis très fier de nos séminaires communs, et j'espère qu'ils sauront conserver leur pertinence. Je garderai bien sûr un souvenir particulier de tous ceux avec qui j'ai partagé un bureau, et de nos réflexions engagées sur la science, l'enseignement, l'informatique... A tous ceux qui restent, je souhaite de continuer avec la même motivation et le même dynamisme.

Bien entendu, je remercie l'ensemble du laboratoire pour m'avoir accueilli avec amitié et simplicité. Patrick Mora qui dirige la maison avec sérénité; les membres du secrétariat, pour leur disponibilité, et pour m'avoir simplifié le quotidien; les informaticiens, qui se débattent avec un système disparate et parfois archaïque; les membres du groupe de théorie des cordes qui m'ont fait profiter de leur expérience.

Convaincu que mon expérience de thèse ne se limite pas à ma seule activité scientifique, je voudrais associer ici Jean-Charles Bolomey et Arnaud Bournel, avec qui j'ai eu le plaisir d'enseigner, ainsi que Marie-Christine Henriot qui m'a permis de collaborer à des sujets enthousiasmants. J'espère que mes étudiants sauront faire honneur au métier complexe d'ingénieur auquel nous avons voulu les préparer.

Lorsque plutôt par hasard je me suis engagé avec X'Doc, pour faire du centre de recherche un lieu à la fois plus professionnel et plus convivial pour les doctorants, j'ai fait sans le savoir un choix déterminant pour mon avenir. Je ne pensais pas alors

prendre autant de plaisir dans la diversité d'activités que j'ai pu aborder. Outre les nombreuses relations amicales que j'ai pu lier, j'ai profité du soutien répété de toute l'Ecole Doctorale, et je les en remercie vivement. Je pense aussi à Marie-Dominique Pujol avec qui j'ai eu le bonheur de partager plusieurs formations. Merci donc à X'Doc, pour m'avoir donné l'occasion de m'épanouir en dehors de mon activité scientifique. J'espère avoir été un bon président, et je souhaite beaucoup de réussite à mes successeurs.

Je voudrais aussi remercier tous ceux avec qui j'ai passé des moments exceptionnels, à courir derrière un ballon. Ces rendez-vous du jeudi soir ont été pendant plus de deux ans une véritable bouffée d'oxygène hebdomadaire. Je suis conscient de la chance que m'a offert l'Ecole, avec ses nombreuses installations sportives, de pouvoir ainsi continuer à jouer au football pendant ma thèse.

Enfin, merci à ma famille. C'est grâce à elle que j'ai pu suivre la voie qui a été la mienne, et la poursuivre aujourd'hui encore dans les meilleures conditions. Merci de votre écoute, de votre compréhension, de votre affection.

Table des matières

Introduction	1
1 D-branes non-commutatives	5
1.1 Les théories de jauge	5
1.1.1 Présentation du physicien	5
1.1.2 L'éclairage du mathématicien	6
1.1.3 Le(s) point(s) sur les i	21
1.2 La transformation de Seiberg-Witten (1)	22
1.2.1 Branes et action de Born-Infeld	22
1.2.2 La transformation des champs de fond	24
1.3 Le star-produit	26
1.3.1 La formule de Moyal-Weyl	26
1.3.2 Star-produits et structures de Poisson	26
1.3.3 La transformation de formalité	28
1.4 La transformation de Seiberg-Witten (2)	30
1.4.1 Les théories de jauge non-commutatives	30
1.4.2 La transformation du secteur de jauge	31
1.4.3 Liberté de jauge	32
1.5 Les différentes approches	32
1.5.1 Couplage aux potentiels Ramond-Ramond	32
1.5.2 Développement ordre par ordre	33
1.5.3 Equivalence de star-produits	36
1.6 Le calcul récursif	39
1.6.1 Résoudre l'équation récursive	39
1.6.2 Expression récursive des termes semi-classiques	40
1.6.3 Expression compacte des termes semi-classiques	42
1.6.4 Une récursivité en puissance de A	44
1.7 L'approche algébrique	49
1.7.1 La résolution théorique	49
1.7.2 Les expressions concrètes	54
1.8 Et les théories non-abéliennes ?	57
1.8.1 Quelques calculs à l'ordre 2	57
1.8.2 Le cas d'une pure jauge	59
2 M5-branes non-abéliennes	63
2.1 Les multiplets supersymétriques	64
2.1.1 Anomalies et polynômes de Chern	64

2.1.2	Les multiplets supersymétriques (2,0) en dimension 6	65
2.1.3	Peut-on reproduire l'anomalie attendue?	66
2.2	Théories de jauge non-abéliennes de 2-formes	67
2.2.1	Complexes de Hochschild et algèbres de Gerstenhaber	67
2.2.2	Théorie de jauge de 2-formes	69
2.2.3	2-groupes de Lie, 2-algèbres de Lie et modules croisés	77
2.3	L'approche géométrique	81
2.3.1	Cordes de type (p,q) et réseaux de cordes	81
2.3.2	Membranes M2 tendues entre plusieurs M5	84
3	Généraliser la symétrie-miroir	87
3.1	Variétés de Calabi-Yau	87
3.1.1	Variétés complexes	88
3.1.2	Groupes d'holonomie	90
3.1.3	Structure de Kähler	94
3.1.4	Variétés de Calabi-Yau	95
3.2	Variétés à structure $SU(3)$	96
3.2.1	Forme de Kähler et forme volume holomorphe	96
3.2.2	$SU(3)$ comme groupe de structure	96
3.2.3	La torsion intrinsèque et ses composantes	97
3.3	Calculs explicites pour une fibration toroïdale	98
3.3.1	Définitions et notations	98
3.3.2	Les composantes de la torsion intrinsèque	100
3.3.3	Les composantes du champ de fond H	104
3.3.4	L'action de la T-dualité	106
3.4	Le spineur invariant	110
3.4.1	Décomposition de $D_\mu \epsilon$ en représentations	110
3.4.2	Relations avec les tenseurs	112
3.4.3	Couplage au champ de fond	116
	Conclusion et perspectives	119
	A Towards an explicit expression of the Seiberg-Witten map	121
	B Mirror symmetric $SU(3)$-structure manifolds with NS fluxes	131
	C Proceedings des Journées Jeunes Chercheurs 2001	143
	Références Bibliographiques	149

Introduction

L'imagination est plus importante que la connaissance. La connaissance est limitée. L'imagination, elle, dépasse toutes les frontières.

A. Einstein

Ce mémoire de thèse s'adresse en premier lieu à un public averti. Je l'ai voulu néanmoins aussi pédagogique que possible. Les travaux que j'y présente, et qui sont pour la plupart originaux, sont comme ils se doivent en prise avec la recherche actuelle. Alors bien sûr certains développements mathématiques pourront sembler insondables au profane, certaines formules pourront faire sursauter. Mais une thèse, c'est avant tout deux choses : du désir et de l'effort. Désir de découvrir, qui nous anime parce que nous avons eu la chance de connaître le plaisir qu'il y a à comprendre quelque chose de nouveau. C'est un goût pour l'esthétisme, pour la connaissance, le désir de participer à la création du savoir. Et surtout l'effort sans lequel on ne peut rien créer, l'effort parce qu'avant de comprendre, et que tout s'éclaire, on travaille dans l'obscurité. La réalité, ce sont plusieurs mois d'incertitude et de doute, à vouloir pénétrer un territoire inconnu, celui des résultats à découvrir, sans être sûr que nos outils et notre connaissance puissent y suffire.

J'ai donc voulu essayer de transmettre dans ce mémoire un peu de ce désir qui m'a animé, et éclairer le chemin que j'ai suivi pour que mes efforts puissent profiter à d'autres. En particulier, j'ai inclus deux sections plus didactiques pour ouvrir le premier et le troisième chapitre. La première est une présentation des théories de jauge classiques avec le regard du mathématicien. Son objectif est de définir proprement ces objets communs au physicien du champ, tout en permettant une compréhension globale. La seconde, au début du troisième chapitre, permet d'introduire les variétés de Calabi-Yau, outils de base de la compactification, avant d'aborder mon travail sur les variétés plus générales à structure $SU(3)$.

Par contre, je ne souhaite pas ici faire une introduction détaillée à la théorie des cordes. Pour ceux qui voudraient davantage de précisions techniques sur sa construction, il existe en effet déjà beaucoup de très bons cours introductifs¹. Je me contenterai donc d'en brosser un portrait rapide, lisible je l'espère par un public plus

¹ On pourra se référer par exemple à [1, 2, 3], que l'on trouvera sur www.arxiv.org aux références indiquées (hep-th. . .). D'autre part, on pourra lire avec bonheur le livre de Brian Greene, *L'univers élégant* [4], ouvert à un plus large public mais qui touche à certains aspects essentiels de la théorie, en particulier la compactification. Ce livre a dernièrement été adapté en documentaire d'animation aux Etats-Unis.

large. Cela me permettra de situer mes travaux de thèse, que je décris dans les trois chapitres de ce mémoire.

Pourquoi les cordes sont-elles si attachantes ?

Comme son nom l'indique, la théorie des cordes contient des *cordes* (fig.1). Celles-ci peuvent être ouvertes comme un bout de ficelle, ou fermées comme un élastique. Mais comme son nom ne l'indique pas, elle contient aussi des *branes* (fig.2), qui sont des surfaces à plusieurs dimensions². Ce sont ces cordes et ces branes qui sont les objets fondamentaux de l'univers. Ce sont eux qui créent les particules et les interactions.



FIG. 1 – Cordes ouvertes et cordes fermées sont les ingrédients de base de la théorie des cordes.

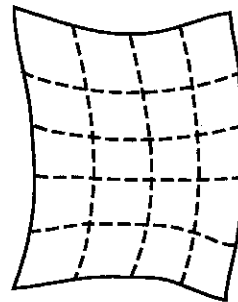


FIG. 2 – Les branes “apparaissent” naturellement dans la théorie, et sont essentielles à sa compréhension.

La relation entre les cordes et les particules observables est la même que celle entre les cordes d'un violon et les notes que l'on entend : ce sont les *vibrations* des cordes qui créent les notes de musique.

Ici, les modes de vibrations des cordes sont les *champs*, c'est-à-dire les particules. Bien sûr, les cordes sont tellement petites que l'on ne peut pas les voir directement, mais on observe leur mode de vibration fondamental (de masse nulle). En particulier, on y trouve :

- ▶ pour les cordes fermées : $g_{\mu\nu}$, $B_{\mu\nu}$, ϕ
- ▶ pour les cordes ouvertes (fig.3) : A_i — le groupe de jauge est abélien ou non-abélien selon le nombre de branes

On retrouve ainsi la *gravitation* ($g_{\mu\nu}$) et les *théories de jauge* (A_i). Le *dilaton* ϕ intervient entre autre dans les modèles cosmologiques, où il permet de recréer l'inflation. Quant au champ de fond antisymétrique $B_{\mu\nu}$, élucider son rôle est la problématique générale de cette thèse.

En étant plus précis, et c'est le premier aspect de mon travail, on peut se demander ici comment décrire le couplage du champ $B_{\mu\nu}$ aux D-branes. En effet, si l'on connaît bien les D-branes plongées dans un espace plat, et sans champ B , on ne sait pas, ou mal, écrire les interactions des champs sur la brane (A_i pour le secteur bosonique) à des champs de fonds non triviaux. Nous verrons ainsi que le couplage au champ

² Le terme vient de membrane, surface à deux dimensions, comme une feuille de papier.

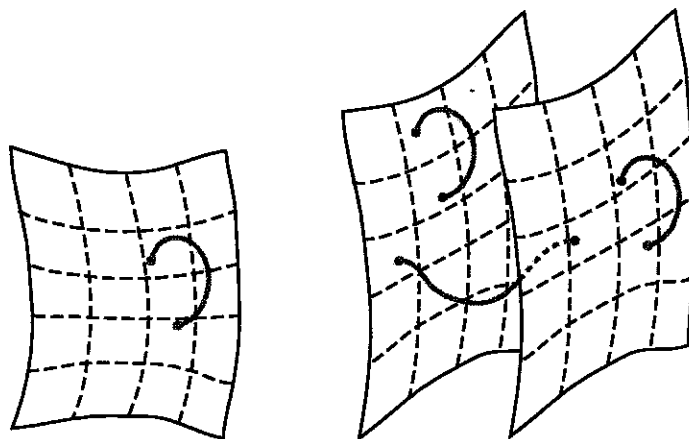


FIG. 3 – Les extrémités des cordes ouvertes s’attachent sur les branes, et créent un champ de jauge. Le groupe de jauge est $U(1)$ pour une brane isolée, et $U(N)$ pour un paquet de N branes coïncidentes.

B peut être décrit par l’intermédiaire d’une théorie de jauge non-abélienne sur la brane, grâce à la transformation de Seiberg-Witten, que nous étudierons dans le premier chapitre.

D’autre part, si l’on croit que les cordes sont les objets fondamentaux de notre monde, alors nous vivons dans un espace-temps à 10 dimensions (9 dimensions d’espace plus le temps). Il faut donc faire “disparaître” 6 dimensions spatiales, ou plus précisément les *compactifier*. Un tuyau d’arrosage est un bon exemple pour illustrer cette idée : vu de loin, il semble unidimensionnel puisque son diamètre est très fin par rapport à sa longueur. Mais s’il n’avait pas deux dimensions, l’eau ne pourrait pas s’écouler à l’intérieur. Ici, ce sont donc 6 dimensions spatiales qui doivent être enroulées très serrées pour disparaître à nos yeux. Cependant, la manière dont elles s’enroulent aura des conséquences sur notre monde à grande échelle. Cette “manière” s’appelle techniquement la variété de compactification. En l’absence de champs de fond, ce doit être une variété de Calabi-Yau. Nous verrons dans le chapitre 3 que la présence d’un champ $B_{\mu\nu}$ oblige à se placer dans un cas plus général, où l’on parle de variétés à structure $SU(3)$. Nous essaierons alors d’élargir la notion de symétrie-miroir, concept essentiel qui relie deux variétés conduisant à la même physique, grâce à la T-dualité.

Enfin, et c’est un comble pour une théorie qui se prétend *la* théorie ultime de l’univers, il y a 5 théories des cordes consistantes. Elles sont différentes, mais partagent des propriétés essentielles, qui ont conduit les théoriciens à comprendre qu’il ne s’agit en fait que de 5 facettes d’une même théorie, appelée *M-théorie* (fig.4). Un peu comme si l’on regardait une pièce à travers 5 trous de serrure : c’est toujours la même pièce, mais chacun n’en voit qu’un morceau, et chaque fois différent.

Cette M-théorie est encore M-ystérieuse ; on sait néanmoins qu’elle a 11 dimensions (10+1) et qu’elle contient deux types d’objets : M2 et M5. La membrane à deux dimensions M2 joue le rôle des cordes. En effet, pour retrouver les théories des cordes à partir de la M-théorie, il faut enrouler une dimension. Et si vous roulez très serrée

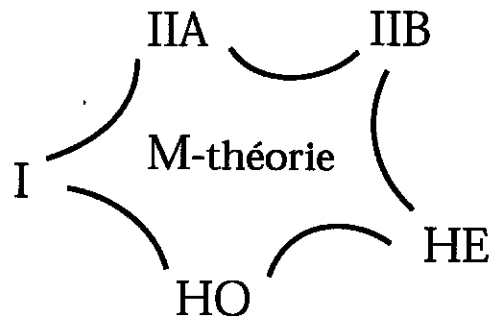


FIG. 4 – Les 5 théories des cordes consistantes sont reliées par un réseau de dualités. Elles sont 5 secteurs perturbatifs de la M-théorie.

une membrane, il vous reste une corde à une dimension. Quant à M5, c'est une brane à 5 dimensions, sur laquelle peut venir s'accrocher M2. C'est elle qui joue le rôle des branes de la figure 3. Souvenons-nous alors que pour les D-branes et les cordes, nous savions décrire non seulement une brane unique (grâce à une théorie de jauge abélienne), mais aussi un paquet de branes coïncidentes (la théorie de jauge devient non abélienne). Sur une brane M5 par contre, si l'on connaît la théorie abélienne associée, on ignore tout de sa version non-abélienne. On ne dispose même pas de candidat bien formulé. Dans le chapitre 2, nous aborderons cette question sous différents aspects.

Ces travaux ont fait l'objet de deux articles originaux, ainsi que de plusieurs séminaires parfois assortis de comptes-rendus. Ces deux articles sont reproduits en annexe. Le premier [36] est basé sur les travaux décrits dans le chapitre 1, le second [71] sur ceux du chapitre 3.

Chapitre 1

D-branes non-commutatives

Objectif A long terme, nous voulons comprendre le couplage du champ de fond antisymétrique $B_{\mu\nu}$ aux D-branes non-abéliennes. Pour cela, nous tenterons d'expliquer la transformation de Seiberg-Witten sur le secteur de jauge, puisque les branes non-commutatives permettent d'interpoler entre la situation abélienne et les théories non-abéliennes.

Plan du chapitre Je vais tout d'abord présenter la transformation de Seiberg-Witten, ainsi que le formalisme utilisé, pour me permettre de bien poser le problème, à savoir trouver une expression explicite de cette transformation sur les champs de jauge. Je décrirai ensuite différentes approches utilisées pour le résoudre, dont la résolution récursive que j'ai développée au cours de ma thèse. J'expliquerai comment elle se rattache aux travaux de B. Jurčo, P. Schupp et J. Wess [30], ainsi qu'à ceux de Kontsevich [37]. Enfin, nous verrons comment aborder le cas non-abélien de la transformation de Seiberg-Witten, puisque tous les calculs explicites ne seront faits que dans le cas abélien. Bien sûr, je traiterai le cas général dans toute la partie de présentation, et aussi longtemps que possible.

1.1 Les théories de jauge

J'aimerais décrire ici un peu plus en détails ce qu'est une théorie de jauge, et en particulier les structures mathématiques sous-jacentes. Je précise que ce ne sont pas les actions qui m'intéresseront, mais la description des champs et des transformations de jauge. D'autre part, les champs seront toujours ici des champs classiques, et je ne m'intéresserai pas à leur quantification. Commençons donc par rappeler les formules bien connues, et nous verrons ensuite qu'elle est leur justification mathématique¹.

1.1.1 Présentation du physicien

Dans le cas le plus simple, c'est-à-dire une théorie abélienne de groupe de jauge $U(1)$, le potentiel de jauge (ou connexion) $A = A_\mu dx^\mu$ est une 1-forme et a pour

¹ L'approche que j'ai choisie ici est tirée du livre de John Baez et Javier Muniain [5].

courbure la 2-forme $F = \frac{1}{2}F_{\mu\nu}dx^\mu \wedge dx^\nu$:

$$F = dA, \quad F_{\mu\nu} = \partial_\mu A_\nu - \partial_\nu A_\mu \quad (1.1)$$

La transformation de jauge δ s'écrit, avec λ une 0-forme quelconque

$$\delta A = d\lambda, \quad \delta A_\mu = \partial_\mu \lambda \quad (1.2)$$

La courbure est donc invariante de jauge, puisque $d^2 = 0$.

Dans le cas plus général où le groupe de jauge est non-abélien, la connexion A , tout comme F et λ , est à valeur dans l'algèbre de Lie du groupe de jauge. Par exemple, lorsque le groupe de jauge est $SU(n)$, ce sont des matrices $n \times n$ hermitiennes² de trace nulle. Dans les formules précédentes apparaissent alors des commutateurs d'algèbre de Lie, qui étaient nuls dans le cas de $U(1)$. La courbure s'écrit

$$F_{\mu\nu} = \partial_\mu A_\nu - \partial_\nu A_\mu - i[A_\mu, A_\nu] \quad (1.3)$$

Cette fois, elle n'est plus invariante de jauge, mais covariante³. En effet, les transformations de jauge prennent la forme

$$\delta A_\mu = \partial_\mu \lambda - i[A_\mu, \lambda] = D_\mu \lambda, \quad \delta F_{\mu\nu} = i[\lambda, F_{\mu\nu}] \quad (1.4)$$

1.1.2 L'éclairage du mathématicien

Quittons maintenant la peau du physicien pour nous glisser dans celle du mathématicien. Et reprenons.

Une théorie de jauge, c'est tout d'abord un fibré vectoriel sur la variété d'espace-temps.

Qu'est-ce qu'un fibré ?

Pour faire simple, un fibré, c'est la réunion d'un espace de base (une variété différentielle M), et d'une fibre (un espace vectoriel V) que l'on colle en chaque point de M . Un peu comme un champ de blé, où la terre joue le rôle de la base, et chaque tige est une fibre, chacune plantée en un point précis de la base. Pour anticiper un peu avec cette image, l'un des problèmes qui se posera par la suite sera de savoir comment une puce se déplace d'un épi à l'autre. Vous trouverez peut-être que tout cela ne fait pas très mathématicien, contrairement à ce que je viens d'annoncer, alors voilà la version moins simple... et plus rigoureuse :

² En général, les physiciens paramétrisent l'algèbre de Lie par $U = e^{i\lambda}$, où $U \in SU(n)$. λ est donc une matrice hermitienne. Tandis que les mathématiciens préfèrent écrire $U = e^\lambda$, faisant de λ une matrice anti-hermitienne. Cela bien sûr ne change pas la dimension de l'algèbre.

³ On dit que la transformation de f est covariante lorsque $f\psi$ se transforme comme ψ , où ψ est un champ de matière. A ce titre, la transformation de jauge $U = e^{i\lambda}$ agit comme

$$\psi \longrightarrow U\psi, \quad \delta\psi = i\lambda\psi$$

Donc, pour que $f\psi$ se transforme de la même manière, U doit agir sur f par

$$f \longrightarrow UfU^{-1}, \quad \delta f = i[\lambda, f]$$

Définition 1.1.1 *Un fibré (E, π, M) est une variété E (l'espace total ou fibré), muni d'une application surjective π (la projection) sur la variété M (la base). En tout point p de M , l'espace*

$$E_p = \{q \in E / \pi(q) = p\}$$

*est appelé **fibre** en p . Si ces fibres sont des espaces vectoriels, alors elles ont toutes la même dimension, et E est un fibré vectoriel.*

Un bon exemple de fibré vectoriel, que nous allons utiliser, est le fibré tangent. L'espace tangent en p de la variété M , qui sera pour nous l'espace-temps (ou en tout cas la zone sur laquelle est définie la théorie de jauge), est noté $T_p M$. Il a bien sûr la même dimension que M . On peut s'en faire l'image géométrique intuitive suggérée par son nom. Sur une surface par exemple, plongée dans l'espace à trois dimensions, l'espace tangent en p est le plan tangent à la surface en p . Une autre formulation, plus abstraite, nous sera plus utile : lorsque l'on a défini des coordonnées x^i au voisinage de p , $T_p M$ est l'espace engendré par les dérivées partielles ∂_i . Le fibré tangent est la réunion de toutes ces fibres, et est noté TM .

Cela nous permet d'introduire un deuxième fibré classique, dual du fibré tangent : le fibré cotangent. Chaque fibre est cette fois engendrée par les différentielles dx^i des fonctions coordonnées. Il est dual du premier au sens où l'espace cotangent en p , $T_p^* M$ est dual de $T_p M$ comme espace vectoriel. Le fibré cotangent est noté $T^* M$.

Qu'est-ce qu'une section ?

Comme son nom l'indique, si le fibré est un champ de blé, alors une section est une coupe de ce champ : imaginez que vous fauchiez ce champ avec une grande faux, coupant chaque brin à la hauteur que vous choisissez. En termes plus précis,

Définition 1.1.2 *Une section s est une fonction de M dans E telle que*

$$\forall p \in M, s(p) \in E_p$$

On note $\Gamma(E)$ l'ensemble des sections de E .

Cela signifie qu'une section associe à tout point p de la base M un vecteur de sa fibre associée E_p . Ainsi, si s et s' sont deux sections, on peut facilement construire la section $s + s'$:

$$(s + s')(p) = s(p) + s'(p)$$

De même, si f est une fonction C^∞ sur M , alors fs est la section :

$$(fs)(p) = f(p)s(p)$$

Ces deux relations ne sont rien d'autre que l'affirmation, en termes techniques,

Proposition 1.1.3 *$\Gamma(E)$ est un $C^\infty(M)$ -module.*

Pour un physicien (qui fait de la théorie des champs), les sections de fibrés sont des objets très naturels. En effet, les objets de base en théorie des champs sont des sections :

- Les champs vectoriels sont des sections du fibré tangent TM . A chaque point de l'espace-temps M , ils associent un vecteur de son espace tangent.
- Les 1-formes sont des sections du fibré cotangent T^*M . En électromagnétisme, par exemple, le champ $A = A_\mu dx^\mu$ associe à chaque point x^ν de l'espace-temps la forme $A(x^\nu)$.
- De manière générale, les champs tensoriels d'ordre (p, q) , p fois contravariants et q fois covariants, sont des sections de $(TM)^{\otimes p} \otimes (T^*M)^{\otimes q}$.
- Les p -formes sont des sections de $(T^*M)^{\wedge p} = \wedge^p(T^*M)$.

Le groupe de jauge G fait de E un G -fibré. Les transformations de jauge sont alors des sections de $\text{End}(E)$.

Qu'est-ce qu'un G -fibré ?

Soit $\{U_\alpha\}$ un recouvrement de M , V un espace vectoriel et ρ une action de G sur V . On peut définir le fibré E de telle sorte que G agisse naturellement sur E , autrement dit, que E soit un G -fibré : $E = \cup_\alpha (U_\alpha \times V)$.

De plus, si $p \in (U_\alpha \cap U_\beta)$, alors sa fibre E_p est définie dans $U_\alpha \times V$ (notons v son élément générique) et dans $U_\beta \times V$ (notons v' son élément générique). Il faut identifier les v aux v' , ce que nous faisons par l'intermédiaire des fonctions de transitions $g_{\alpha\beta}$

$$g_{\alpha\beta} : U_\alpha \cap U_\beta \longrightarrow G$$

par la relation : $v = \rho[g_{\alpha\beta}(p)] v'$.

Pour que cette identification définisse bien un fibré, les fonctions de transitions doivent vérifier deux conditions assez intuitives :

- $g_{\alpha\alpha} = 1$ sur U_α
- $g_{\alpha\beta} g_{\beta\gamma} g_{\gamma\alpha} = 1$ sur $U_\alpha \cap U_\beta \cap U_\gamma$ (condition de cocycle)

La première condition traduit simplement le fait que nous ne voulons pas “mélanger” une fibre. Les fonctions de transitions servent de dictionnaire lorsqu’une fibre possède deux définitions différentes (la définition α et la définition β), mais ne doivent pas déformer une définition donnée.

La deuxième condition, appelée condition de cocycle, vient du même genre de remarque : sur une intersection triple $U_\alpha \cap U_\beta \cap U_\gamma$, passer d’une définition à l’autre circulairement (de α à β , puis à γ et retour à α) doit faire revenir au point de départ. La définition α ne doit pas être déformée après ce petit voyage.

On pourrait en dire de même pour les intersections à n ouverts, mais ces deux conditions suffisent pour assurer toutes les autres. Elles suffisent à faire de E un fibré vectoriel, ce que nous ne démontrerons pas ici.

Pour un physicien, un tel groupe G est un groupe de jauge.

Qu'est-ce que $\text{End}(E)$?

Lorsque V est un espace vectoriel, l'ensemble de ses endomorphismes $\text{End}(V)$ est lui aussi un espace vectoriel, isomorphe à $V \otimes V^*$. De la même manière, nous *définissons* $\text{End}(E)$, le fibré d'endomorphismes de E , comme étant le fibré $E \otimes E^*$. Cette définition assure que $\text{End}(E)$ est bien un fibré. D'autre part, les fibres de $\text{End}(E)$ sont exactement les espaces d'endomorphismes des fibres de E (ce qui justifie la notation) :

$$\text{End}(E)_p = \text{End}(E_p)$$

Ainsi, si T est une section de $\text{End}(E)$, alors en tout point p de M , $T(p)$ est un endomorphisme de l'espace vectoriel E_p . A ce titre, T agit sur les sections de E par

$$\begin{aligned} T : \Gamma(E) &\longrightarrow \Gamma(E) \\ s &\longrightarrow Ts \quad \text{tel que } (Ts)(p) = T(p) s(p) \end{aligned}$$

Mieux encore, cette action est $C^\infty(M)$ -linéaire, c'est-à-dire qu'en plus d'être linéaire, pour tout $f \in C^\infty(M)$,

$$T(fs) = fT(s)$$

Réciproquement, on pourrait montrer que toute application $C^\infty(M)$ -linéaire de $\Gamma(E)$ dans lui-même définit une section de $\text{End}(E)$.

Qu'est-ce qu'une transformation de jauge ?

Supposons que E a pour groupe de jauge G , c'est-à-dire que E est un G -fibré. Soit T une section de $\text{End}(E)$. En tout point p de M , $T(p)$ est donc un endomorphisme de E_p . Mais parmi ces endomorphismes, certains sont spéciaux : tous ceux qui correspondent à l'action d'un élément de G . Plus exactement, cette action est définie par la représentation $\rho : G \longrightarrow \text{End}(E_p)$. On dit alors que l'endomorphisme $T(p)$ *vit dans* G lorsque $T(p) \in \text{Im}(\rho)$.

Définition 1.1.4 *On dit qu'une section T de $\text{End}(E)$ est une transformation de jauge si et seulement si $T(p)$ vit dans G pour tout $p \in M$. On note $\mathcal{G}$ l'ensemble des transformations de jauge.*

Si $\mathfrak{g}$ désigne l'algèbre de Lie de G , alors $\mathfrak{g}$ agit naturellement sur E_p par la représentation $d\rho$, différentielle de ρ :

$$\forall X \in \mathfrak{g}, \quad d\rho(X) = \left. \frac{d}{dt} \right|_{t=0} \rho(e^{tX})$$

$$\begin{array}{ccc} G & \xrightarrow{\rho} & GL(V) \\ \uparrow e^{tX} & & \downarrow \left. \frac{d}{dt} \right|_{t=0} \\ \mathfrak{g} & \xrightarrow{d\rho} & \mathfrak{gl}(V) \end{array}$$

Ainsi, de même que précédemment, on dira que l'endomorphisme $T(p)$ *vit dans* $\mathfrak{g}$ lorsque $T(p) \in \text{Im}(d\rho)$. Et que la section T de $\text{End}(E)$ *vit dans* $\mathfrak{g}$ lorsque $T(p)$ vit dans $\mathfrak{g}$ pour tout $p \in M$.

Le potentiel de jauge A_μ est (le potentiel-vecteur d')une G -connexion.

Qu'est-ce qu'une connexion ?

Parler d'une connexion, c'est d'abord parler d'une dérivée. En effet, une connexion, c'est ce qui sert à dériver les sections. Et ce n'est pas si évident, en particulier parce qu'il y a plusieurs manières de le faire. Lorsqu'on dérive une fonction d'une variable, on prend la limite du taux d'accroissement (c'est en tout cas la définition que j'ai apprise au lycée) :

$$f'(x) = \lim_{\epsilon \rightarrow 0} \frac{f(x + \epsilon) - f(x)}{\epsilon}$$

Cela signifie qu'il faut calculer la différence entre les valeurs de f en deux points voisins. Mais voilà, les valeurs d'une section en deux points voisins sont deux vecteurs qui ne vivent pas dans le même espace ! Il n'y a donc pas de manière unique d'en faire la différence, ni de manière meilleure que les autres. Bien sûr, lorsque l'on a défini un système de coordonnées on peut calculer les dérivées partielles $\partial_\mu s$ de la section s . Mais en passant d'une fibre à l'autre, comme nous l'avons déjà remarqué en parlant des G -fibrés, on peut s'autoriser une transformation (c'est-à-dire un endomorphisme). Globalement, cela donne un terme de la forme $A_\mu s$, où A_μ est une section de $\text{End}(E)$. Ainsi, une connexion D_μ s'écrira

$$D_\mu = \partial_\mu + A_\mu$$

ce qui laisse beaucoup de possibilités.

Maintenant que nous avons allègrement écrit la forme générale d'une connexion, peut-être faudrait-il la définir :

Définition 1.1.5 Une connexion D sur M est une application $\mathcal{C}^\infty(M)$ -linéaire de $\Gamma(TM)$ dans l'espace $\text{Der}(\Gamma(E))$ des dérivations de $\Gamma(E)$.

Puisque tous ceux qui peuvent comprendre au premier regard une telle expression ne liront certainement jamais cette section, accordons-nous quelques explications. Soit donc v un champ de vecteurs sur M (donc une section de TM). La connexion D lui associe une fonction D_v . Souvenons-nous alors que nous savons additionner les sections ($v + w$) et les multiplier par des fonctions $\mathcal{C}^\infty$ sur M (fv). D est une application $\mathcal{C}^\infty(M)$ -linéaire signifie alors tout simplement :

$$D_{v+w} = D_v + D_w \quad \text{et} \quad D_{fv} = fD_v$$

Ainsi, en écrivant $v = v^\mu \partial_\mu$ dans la base ∂_μ associée au coordonnées x^μ sur un ouvert de M , on peut développer

$$D_v = D_{v^\mu \partial_\mu} = v^\mu D_\mu$$

avec la notation $D_\mu = D_{\partial_\mu}$.

Grâce à cette fonction D_v , nous allons ensuite dériver les sections de E . Cela signifie que si $s, t \in \Gamma(E)$, et $f \in \mathcal{C}^\infty(M)$,

$$D_v(s + t) = D_v(s) + D_v(t) \quad \text{et} \quad D_v(fs) = v(f)s + fD_v(s)$$

où $v(f) = v^\mu \partial_\mu f$. Cette deuxième relation n'est rien d'autre que la règle de Leibniz, qui indique comment dériver un produit. En composantes, elle s'écrit en effet

$$D_\mu(fs) = (\partial_\mu f)s + fD_\mu s$$

Revenons alors à la proposition par laquelle nous avons commencé : toute expression $D = \partial + A$, où $A = A_\mu dx^\mu$ est une section de $\text{End}(E) \otimes T^*M$, est une connexion. Réciproquement, toute connexion s'écrit sous cette forme⁴.

A s'appelle le **potentiel-vecteur**, mais est couramment, et abusivement, appelé connexion. C'est une 1-forme à valeur dans $\text{End}(E)$, qui sera d'ici peu le potentiel de jauge du physicien.

$D_v(s)$ est la **dérivée covariante** de s dans la direction v . Nous venons d'expliquer le terme *dérivée*. Il nous reste le terme *covariante*, que nous allons voir en précisant :

Qu'est-ce qu'une G -connexion ?

Puisque notre fibré est un G -fibré, il nous faut expliquer comment les connexions s'articulent avec la structure de jauge, et en particulier, comment les transformations de jauge agissent sur ces connexions. Rappelons que g est une transformation de jauge ($g \in \mathcal{G}$) si et seulement si $g(p)$ vit dans G pour tout $p \in M$. Comme nous allons le vérifier, l'action suivante de g sur la connexion D définit une nouvelle connexion D' :

$$D'_v(s) = gD_v(g^{-1}s) \quad \text{c'est-à-dire} \quad D'_v = gD_vg^{-1} \quad (1.5)$$

Avant toute chose, comme expliqué dans la note 3 (p.6), la dérivée D_v se transforme donc de manière covariante, d'où son nom. D'autre part, laissant le soin au lecteur consciencieux de vérifier que D' est $C^\infty(M)$ -linéaire en v , et linéaire en s , nous nous contenterons ici de vérifier qu'elle suit la règle de Leibniz. Si f est une fonction C^∞ sur M :

$$\begin{aligned} D'_v(fs) &= gD_v(g^{-1}fs) \\ &= gD_v(fg^{-1}s) \\ &= gv(f)g^{-1}s + gfD_v(g^{-1}s) \\ &= v(f)s + fD'_v(s) \end{aligned}$$

où, entre la première à la deuxième ligne, on se souviendra que l'action g^{-1} sur s est $C^\infty(M)$ -linéaire, comme celle de toute section de $\text{End}(E)$ ⁵.

Dans un système de coordonnées locales, nous pouvons donc en déduire la transformation de jauge du potentiel-vecteur A (en écrivant bien sûr $D_\mu = \partial_\mu + A_\mu$ et $D'_\mu = \partial_\mu + A'_\mu$) :

$$\begin{aligned} (\partial_\mu + A'_\mu)s &= D'_\mu = gD_\mu(g^{-1}s) \\ &= g(\partial_\mu + A_\mu)(g^{-1}s) \\ &= gA_\mu g^{-1}s + g(\partial_\mu g^{-1})s + gg^{-1}\partial_\mu s \\ &= \partial_\mu s + (gA_\mu g^{-1} + g\partial_\mu g^{-1})s \end{aligned}$$

⁴ Bien sûr, en toute rigueur, ces expressions sont locales. Elles s'entendent une fois choisi des coordonnées x^μ sur un ouvert U et une trivialisat on locale du fibr  E sur U . De mani re g n rale, il faudrait remplacer ∂ par une connexion de r f rence D^0 , qui peut  tre quelconque. Cela illustre le fait qu'il n'existe pas de connexion meilleure que les autres dans l'absolu. De m me, le potentiel A est en r alit  une section de $\text{End}(E|_U) \otimes T^*U$.

⁵En clair, $(g^{-1}fs)(p) = g^{-1}(p)(fs)(p) = g^{-1}(p)f(p)s(p) = f(p)(g^{-1}s)(p) = (fg^{-1}s)(p)$.

ce qui donne finalement :

$$A'_\mu = gA_\mu g^{-1} + g\partial_\mu g^{-1} \quad (1.6)$$

On peut tout d'abord remarquer qu'avec $g \in \mathcal{G}$, $g\partial_\mu g^{-1}$ vit dans l'algèbre de Lie $\mathfrak{g}$. Donc, pour que l'expression ait un sens, il est naturel de demander que A_μ vive dans $\mathfrak{g}$. En effet, c'est à cette condition qu'il en sera de même pour $gA_\mu g^{-1}$, et donc pour A'_μ .

Définition 1.1.6 *On dit que $D = \partial + A$ est une G -connexion si les composantes A_μ dans un système de coordonnées locales vivent dans $\mathfrak{g}$.*

Cette définition peut sembler maladroite, puisqu'elle privilégie un système de coordonnées locales. Mais si elle est vérifiée dans un système de coordonnées, la propriété *vivre dans $\mathfrak{g}$* l'est dans tous les autres.

Avant de parler de la courbure, j'aimerais expliciter les formules que nous venons de voir dans le cas où G est le groupe $U(1)$. $U(1)$ étant un sous-groupe d'isomorphismes de $\mathbb{C}$, considérons pour simplifier le cas le plus simple : soit E le fibré en droites trivial $M \times \mathbb{C}$. Puisque les endomorphismes de $\mathbb{C}$ sont isomorphes à $\mathbb{C}$, le fibré d'endomorphismes $\text{End}(E)$ est alors isomorphe à E . Le potentiel-vecteur A est donc une 1-forme à valeurs complexes. Ou plus exactement, en imposant qu'elle respecte la structure de jauge, à valeurs dans l'algèbre de Lie de $U(1)$:

$$\mathfrak{u}(1) = \{ix / x \in \mathbb{R}\}$$

Une transformation de jauge $g \in \mathcal{G}$ s'écrit alors $g = e^{-\lambda}$, où λ est une fonction⁶ de M dans $\mathfrak{u}(1)$ ⁷, et agit donc sur A par

$$\begin{aligned} A'_\mu &= e^{-\lambda} A_\mu e^\lambda + e^{-\lambda} \partial_\mu (e^\lambda) \\ &= A_\mu + \partial_\mu \lambda \end{aligned}$$

ce qui revient à ajouter à la forme A la forme exacte $d\lambda$

$$A' = A + d\lambda \quad (1.7)$$

C'est bien la transformation de jauge attendue pour un potentiel de jauge abélien. Lorsque le groupe de jauge est $U(n)$ (ou de manière générale un groupe non-abélien compact), le même raisonnement conduit à la transformation infinitésimale suivante

⁶ λ est ici une simple fonction parce que le fibré E est globalement trivial. C'est bien sûr une section dans le cas général.

⁷ L'application exponentielle de $\mathfrak{g}$ dans G est surjective lorsque G est compact et connexe. Si le groupe est connexe mais non compact, alors tout élément X de G s'écrit comme produit fini de telles exponentielles : $X = e^{x_1} \dots e^{x_p}$, $x_i \in \mathfrak{g}$. Enfin, si G n'est pas connexe, l'exponentielle est à valeurs sur la composante connexe de l'identité.

D'autre part, même dans le cas d'un groupe compact et connexe, l'exponentielle n'est en général pas injective, ce qui autorise des groupes différents ($SU(2)$ et $SO(3)$ par exemple) à avoir la même algèbre de Lie (ils sont alors dits localement isomorphes). Cependant, parmi tous les groupes localement isomorphes à G , il existe un et un seul groupe connexe et simplement connexe : c'est son groupe de revêtement universel G_0 , et $G \simeq G_0/\pi_1(G)$ (dans l'exemple, $SO(3) \simeq SU(2)/\{\pm I\}$).

(au premier ordre en λ), où cette fois A et λ appartiennent à l'algèbre $\mathfrak{u}(n)$:

$$\begin{aligned} A'_\mu &= e^{-\lambda} A_\mu e^\lambda + e^{-\lambda} \partial_\mu (e^\lambda) \\ &= (1 - \lambda) A_\mu (1 + \lambda) + e^{-\lambda} \partial_\mu \lambda (e^\lambda) + o(\lambda) \\ &= A_\mu + A_\mu \lambda - \lambda A_\mu + \partial_\mu \lambda \end{aligned}$$

soit, en appelant δ cette transformation infinitésimale

$$\delta A_\mu \equiv A'_\mu - A_\mu = \partial_\mu \lambda + [A_\mu, \lambda] \quad (1.8)$$

Bien sûr, nous travaillons ici avec des notations de mathématiciens. Nous y reviendrons à la fin de la section.

Le champ de jauge $F_{\mu\nu}$ est la courbure de la G -connexion A_μ .

Qu'est-ce que la courbure ?

La courbure mesure la non-commutativité de la connexion dans les directions de deux vecteurs v et w . En effet, D_v permet de dériver dans la direction de v , c'est-à-dire de faire un déplacement infinitésimal le long de v . Plus exactement, elle indique comment il faut bouger sur la fibre lorsque l'on fait un tel déplacement sur la base. Mais lorsque l'on se déplace d'abord selon v puis w , arrive-t-on au même endroit sur la fibre que si l'on fait l'inverse ? En général, non. Et c'est ce que mesure la courbure.

Définition 1.1.7 La courbure F de la connexion D est la section de $\text{End}(E) \otimes \wedge^2 T^*M$ définie, pour tout $v, w \in TM$, par

$$F(v, w) \equiv D_v D_w - D_w D_v - D_{[v, w]} \quad (1.9)$$

où le crochet $[v, w]$ est le crochet de Lie de TM

$$[v, w] = (v^i \partial_i w^j - w^i \partial_i v^j) \partial_j \quad (1.10)$$

Et la présence de ce troisième terme dans la définition de F n'est pas anodine. Tout d'abord, il permet à la connexion triviale $D^0 = \partial$ d'être **plate**, c'est-à-dire d'avoir une courbure nulle. En effet, on a alors $D^0 v = v^\mu \partial_\mu = v$, et donc

$$F^0(v, w) = vw - wv - [v, w] = 0 \quad (1.11)$$

D'autre part, il n'est pas du tout évident a priori que $F(v, w)$, d'après la définition 1.1.7, soit une section de $\text{End}(E)$. Nous avons vu que ces sections sont exactement les fonctions $C^\infty(M)$ -linéaires de $\Gamma(E)$ dans $\Gamma(E)$. Or, si D_v est bien un endomorphisme de $\Gamma(E)$, il n'est pas $C^\infty(M)$ -linéaire, puisque pour toute section s de E et fonction $f \in C^\infty(M)$:

$$D_v(fs) = fD_v(s) + v(f)s \neq fD_v(s)$$

Cependant, comme nous allons le vérifier, les termes parasites s'éliminent dans l'expression de F

$$\begin{aligned}
F(v, w)(fs) &= D_v(fD_w(s) + w(f)s) - D_w(fD_v(s) + v(f)s) \\
&\quad - fD_{[v, w]}(s) - [v, w](f)s \\
&= fD_vD_ws + v(f)D_ws + w(f)D_vs + v(w(f))s \\
&\quad - fD_wD_vs - w(f)D_vs - v(f)D_ws - w(v(f))s \\
&\quad - fD_{[v, w]}s - [v, w](f)s \\
&= f(D_vD_w - D_wD_v - D_{[v, w]})s \\
&= fF(v, w)s
\end{aligned}$$

On pourrait vérifier de même que F est $\mathcal{C}^\infty(M)$ -linéaire dans ses deux premiers arguments v et w . De sorte que si l'on définit⁸

$$F_{\mu\nu} \equiv F(\partial_\mu, \partial_\nu) = [D_\mu, D_\nu] \quad (1.12)$$

on peut écrire pour tout $v, w : F(v, w) = v^\mu w^\nu F_{\mu\nu}$.

Ces composantes $F_{\mu\nu}$ nous permettent aussi d'écrire l'expression locale de la 2-forme F à laquelle les physiciens sont habitués :

$$F = \frac{1}{2} F_{\mu\nu} dx^\mu \wedge dx^\nu \quad (1.13)$$

Puisque la courbure est définie en fonction de la connexion, qui est covariante sous l'action d'une transformation de jauge, il est très facile de vérifier que la courbure est elle aussi covariante. Ainsi, sous l'action de $g \in \mathcal{G}$, agissant comme élément de $\text{End}(E)$,

$$F' = gFg^{-1} \quad (1.14)$$

D'autre part, on peut calculer l'expression de $F_{\mu\nu}$, lorsque $D_\mu = \partial_\mu + A_\mu$:

$$\begin{aligned}
F_{\mu\nu} &= [\partial_\mu + A_\mu, \partial_\nu + A_\nu] \\
&= [\partial_\mu, \partial_\nu] + [\partial_\mu, A_\nu] + [A_\mu, \partial_\nu] + [A_\mu, A_\nu] \\
&= \partial_\mu A_\nu - \partial_\nu A_\mu + [A_\mu, A_\nu]
\end{aligned} \quad (1.15)$$

Cette expression est la même que (1.3) à des facteurs i près, qui viennent des définitions du potentiel de jauge et de la courbure du physicien. Encore une fois, nous y reviendrons à la fin de cette section pour faire le point.

Enfin, la relation (1.15) reliant les composantes des formes A et F dans un système de coordonnées locales, on peut légitimement se demander comment écrire l'identité sur les formes elles-mêmes. Malheureusement, nous n'avons pas défini F directement à partir de la connexion, mais seulement $F(v, w)$. La forme elle-même a été seulement reconstruite. Pourquoi cela ? Parce que nous ne savons pas écrire la connexion D "comme une forme", c'est-à-dire comme $D_\mu dx^\mu$. Ou pour utiliser le langage approprié, nous ne disposons pas de la différentielle associée à D . Dès que nous l'aurons défini, nous verrons que (1.15) s'écrit effectivement de manière très élégante en terme de formes. Et nous pourrons aussi exprimer simplement l'identité de Bianchi et l'équation de Yang-Mills. Mais pour cela, il nous faut faire...

⁸On se souviendra que $[\partial_\mu, \partial_\nu] = 0$.

Un peu de calcul différentiel

Nous allons maintenant formaliser un peu le calcul différentiel à valeur dans un fibré. En effet, le potentiel de jauge A est une 1-forme à valeurs dans $\text{End}(E)$, tandis que la courbure F est, comme nous venons de le voir, une 2-forme à valeurs dans $\text{End}(E)$. De manière générale, nous allons donc utiliser des p -formes à valeurs dans E et $\text{End}(E)$, c'est-à-dire les sections respectivement de $E \otimes \bigwedge T^*M$ et $\text{End}(E) \otimes \bigwedge T^*M$.

La forme des p -formes Tout d'abord, ces p -formes s'écrivent comme sommes des p -formes primitives $s \otimes \omega$ et $T \otimes \omega$, où ω est une p -forme usuelle sur M (s et T étant bien sûr des sections, respectivement, de E et $\text{End}(E)$). Dans un système de coordonnées locales, on peut donc les exprimer dans la base des dx^I , où I est un multi-indices : $I = (i_1, i_2, \dots, i_p)$.

- Une forme à valeurs dans E s'écrit localement : $s_I \otimes dx^I$.
- Une forme à valeurs dans $\text{End}(E)$ s'écrit localement : $T_I \otimes dx^I$.

Produit(s) extérieur(s) On peut définir le produit extérieur de différentes p -formes, mais attention, on ne peut pas pour autant s'essayer à n'importe quelle combinaison. Tout d'abord, on peut multiplier les formes à valeurs dans E et $\text{End}(E)$ par une forme usuelle $\alpha \in \bigwedge T^*M$, par

$$(s \otimes \omega) \wedge \alpha = s \otimes (\omega \wedge \alpha), \quad (T \otimes \omega) \wedge \alpha = T \otimes (\omega \wedge \alpha) \quad (1.16)$$

Bien sûr, cette formule, comme les suivantes, s'étend aux formes non primitives par $C^\infty(M)$ -linéarité. Remarquons qu'en considérant la section s comme une 0-forme sur E , cette définition nous permet d'écrire $s \wedge \omega$ la p -forme $s \otimes \omega$. Nous n'utiliserons pas ici cette facilité, pour bien marquer la différence entre la partie *section* et la partie *forme* de ces objets.

Dans le cas des formes sur $\text{End}(E)$, on dispose de plus d'une multiplication sur les sections de $\text{End}(E)$, la multiplication des endomorphismes, ce qui nous permet de définir le produit extérieur⁹

$$(S \otimes \alpha) \wedge (T \otimes \beta) = ST \otimes (\alpha \wedge \beta) \quad (1.17)$$

alors que cela est impossible entre deux formes sur E , puisqu'on ne sait pas multiplier deux sections de E . Par contre, les sections de $\text{End}(E)$ agissent sur les sections de E , définissant donc naturellement le produit extérieur :

$$(T \otimes \alpha) \wedge (s \otimes \beta) = T(s) \otimes (\alpha \wedge \beta) \quad (1.18)$$

entre $T \otimes \alpha \in \text{End}(E) \otimes \bigwedge T^*M$ et $s \otimes \beta \in E \otimes \bigwedge T^*M$.

Revenons sur le produit extérieur entre deux formes à valeurs dans $\text{End}(E)$. C'est le seul produit, parmi ceux que nous avons défini ici, qui puisse s'écrire dans les deux sens $S_p \wedge T_q$ et $T_q \wedge S_p$, où S_p et T_q désignent respectivement une p -forme et une

⁹ Notons qu'en général le produit ST n'a rien à voir avec le produit TS , et ce produit extérieur n'est donc pas antisymétrique (même de manière graduée).

q -forme sur $\text{End}(E)$. Comme nous l'avons déjà remarqué, ces expressions ne sont pas antisymétriques, et nous pouvons donc définir le commutateur gradué :

$$[S_p, T_q] = S_p \wedge T_q - (-1)^{pq} T_q \wedge S_p \quad (1.19)$$

Ce commutateur serait identiquement nul sur des formes à valeurs réelles. Ici, on peut le relier au commutateur entre endomorphismes, en utilisant l'antisymétrie graduée de la base locale¹⁰ :

$$\begin{aligned} [S_I \otimes dx^I, T_J \otimes dx^J] &= S_I T_J \otimes dx^I \wedge dx^J - (-1)^{pq} T_J S_I \otimes dx^J \wedge dx^I \\ &= (S_I T_J - T_J S_I) \otimes dx^I \wedge dx^J \\ &= [S_I, T_J] \otimes dx^I \wedge dx^J \end{aligned}$$

Ce commutateur vérifie la propriété d'antisymétrie graduée :

$$[S_p, T_q] = -(-1)^{pq} [T_q, S_p]$$

Différentielle extérieure covariante sur E On veut généraliser la notion de différentielle extérieure (d) pour la rendre compatible avec une connexion D quelconque. En effet, si l'on considère la connexion triviale $D_v^0 = v$ (qui s'écrit simplement en coordonnées locale $D_\mu^0 = \partial_\mu$), elle est naturellement reliée à la différentielle extérieure par la formule

$$D_v^0(f) = v(f) = v^\mu \partial_\mu f = df(v) \quad (1.20)$$

pour toute fonction f . Dans un fibré, ce sont les sections qui jouent le rôle des fonctions, et on définit donc la différentielle extérieure covariante d_D associée à la connexion D par

$$d_D s(v) = D_v(s), \quad \text{soit localement} \quad d_D s = D_\mu s \otimes dx^\mu \quad (1.21)$$

Les sections de E étant les 0-formes à valeurs dans E , cette définition s'étend naturellement sur une p -forme quelconque. d_D est donc l'opérateur de degré 1 sur l'algèbre extérieure à valeurs dans E :¹¹

$$d_D(s_I \otimes dx^I) = D_\mu s_I \otimes dx^\mu \wedge dx^I \quad (1.22)$$

On pourra noter au passage la forme élégante de cet opérateur sur une forme primitive, où l'on voit bien apparaître la règle de Leibniz (qui disparaissait dans la formule précédente puisque $d(dx^I) = 0$) :

$$d_D(s \otimes \omega) = d_D s \wedge \omega + s \otimes d\omega \quad (1.23)$$

¹⁰ Lorsque dx^I est une p -forme et dx^J une q -forme,

$$dx^I \wedge dx^J = (-1)^{pq} dx^J \wedge dx^I$$

¹¹ On peut remarquer que cette formule généralise simplement celle de la différentielle d'une forme usuelle :

$$d(\omega_I dx^I) = \partial_\mu \omega_I dx^\mu \wedge dx^I$$

Précisons aussi qu'il n'y a aucune raison pour que la différentielle d_D soit de carré nul. En fait, ce n'est le cas que si la connexion D est plate. Plus précisément :

$$\begin{aligned} d_D^2(s_I \otimes dx^I) &= D_\mu D_\nu s_I \otimes dx^\mu \wedge dx^\nu \wedge dx^I \\ &= \frac{1}{2} [D_\mu, D_\nu] s_I \otimes dx^\mu \wedge dx^\nu \wedge dx^I \\ &= F \wedge (s_I \otimes dx^I) \end{aligned} \quad (1.24)$$

Différentielle extérieure covariante sur $\text{End}(E)$ Pour pouvoir faire le même travail sur $\text{End}(E)$, il nous faut d'abord y définir une connexion. Plus exactement, puisque nous disposons de la connexion D sur E , nous aimerions disposer sur $\text{End}(E)$ d'une connexion "compatible". Et il existe effectivement une (unique) connexion $\tilde{D}$, dérivée de D au sens où

Définition 1.1.8 La connexion $\tilde{D}$ sur $\text{End}(E)$ est définie par

$$(\tilde{D}_v T)(s) = D_v(Ts) - TD_v(s) = [D_v, T]s \quad (1.25)$$

pour tout champ vectoriel v sur M , section T de $\text{End}(E)$ et section s de E .

Le commutateur est celui des endomorphismes de $\Gamma(E)$ (au sens *linéaire* et pas $C^\infty(M)$ -linéaire), que sont D_v et T . Il n'est pas évident a priori que le résultat soit $C^\infty(M)$ -linéaire, et donc section de $\text{End}(E)$, mais c'est bien le cas, comme on pourra le vérifier facilement. Ainsi, pour faire plus court que la définition, on peut écrire localement

$$\tilde{D}_\mu T = [D_\mu, T] \quad (1.26)$$

Muni de $\tilde{D}$, nous pouvons maintenant dérouler la mécanique de la différentielle extérieure covariante, mais cette fois sur $\text{End}(E)$. Pour ne pas alourdir inutilement les notations, nous continuerons à noter d_D cette différentielle (et pas $d_{\tilde{D}}$), qui est de toute façon l'unique généralisation aux formes à valeurs dans $\text{End}(E)$ de la différentielle d_D initiale. Ainsi, comme précédemment :

– d_D est définie sur $\text{End}(E)$ par

$$(d_D T)(v) = \tilde{D}_v T \quad \text{i.e.} \quad d_D T = \tilde{D}_\mu T \otimes dx^\mu \quad (1.27)$$

– d_D est définie sur $\text{End}(E) \otimes \bigwedge T^*M$ par

$$d_D(T_I \otimes dx^I) = \tilde{D}_\mu T_I \otimes dx^\mu \wedge dx^I \quad (1.28)$$

On pourra noter que ces définitions conduisent à la formule de Leibniz graduée (qui généralise la formule (1.23))

$$d_D(\eta \wedge \chi) = d_D \eta \wedge \chi + (-1)^p \eta \wedge d_D \chi \quad (1.29)$$

où p est le degré de la forme η . Cette expression est valable que les formes η et χ soit des formes usuelles, ou à valeurs dans E ou $\text{End}(E)$ (c'est-à-dire toute combinaison compatible avec l'expression $\eta \wedge \chi$).

Enfin, comme pour les formes à valeurs dans E , le carré de la différentielle d_D appliqué à une forme $\eta = T_I \otimes dx^I$ à valeurs dans $\text{End}(E)$ fait intervenir la courbure :

$$\begin{aligned}
 d_D^2 \eta &= [D_\mu, [D_\nu, T_I]] \otimes dx^\mu \wedge dx^\nu \wedge dx^I \\
 &= \frac{1}{2} ([D_\mu, [D_\nu, T_I]] - [D_\nu, [D_\mu, T_I]]) \otimes dx^\mu \wedge dx^\nu \wedge dx^I \\
 &= \frac{1}{2} [[D_\mu, D_\nu], T_I] \otimes dx^\mu \wedge dx^\nu \wedge dx^I \\
 &= [F, \eta]
 \end{aligned} \tag{1.30}$$

Le potentiel de jauge et la courbure — Acte II : les formes

Si nous décomposons la connexion sous la forme $D = D^0 + A$, où D^0 est la connexion plate standard $D_\mu^0 = \partial_\mu$, nous pouvons expliciter l'action de la différentielle extérieure covariante d_D sur les formes à valeurs dans E et $\text{End}(E)$. On notera que la différentielle extérieure associée à D^0 est la différentielle usuelle : $d_{D^0} = d$. En particulier, $d^2 = 0$.

Sur la forme $\chi = s_I \otimes dx^I$ à valeurs dans E ,

$$\begin{aligned}
 d_D \chi &= D_\mu s_I \otimes dx^\mu \wedge dx^I \\
 &= (D_\mu^0 + A_\mu) s_I \otimes dx^\mu \wedge dx^I \\
 &= d\chi + A \wedge \chi
 \end{aligned} \tag{1.31}$$

De même, sur la forme $\eta = T_I \otimes dx^I$ à valeurs dans $\text{End}(E)$,

$$\begin{aligned}
 d_D \eta &= [D_\mu, T_I] \otimes dx^\mu \wedge dx^I \\
 &= [D_\mu^0, T_I] \otimes dx^\mu \wedge dx^I + [A_\mu, T_I] \otimes dx^\mu \wedge dx^I \\
 &= d\eta + [A, \eta]
 \end{aligned} \tag{1.32}$$

Grâce à ces expressions, nous pouvons écrire la relation entre les formes F et A . En effet, pour toute forme χ à valeurs dans E , et d'après (1.24) :

$$\begin{aligned}
 F \wedge \chi &= d_D^2 \chi = d_D(d\chi + A \wedge \chi) \\
 &= d^2 \chi + dA \wedge \chi - A \wedge d\chi + A \wedge d\chi + A \wedge A \wedge \chi \\
 &= (dA + A \wedge A) \wedge \chi
 \end{aligned}$$

ce qui conduit à la formule

$$F = dA + A \wedge A = dA + \frac{1}{2} [A, A] \tag{1.33}$$

dont on peut vérifier qu'elle redonne effectivement l'expression (1.15) pour les composantes $F_{\mu\nu}$.

L'identité de Bianchi

J'ai dit en introduction de cette section que je ne m'intéresserai qu'aux structures de jauge, et pas aux actions, mais je ne peux résister à la tentation d'écrire les équations de Yang-Mills, tant elles sont élégantes. En toute rigueur, les équations ne sont

pas l'action, donc accordons-nous cette faveur. De plus, nous allons commencer par l'identité de Bianchi, qui est une identité topologique comme nous allons le voir, et donc pas une équation dynamique.

Ecrites dans le formalisme des formes différentielles, les équations de Maxwell se résument à¹²

$$dF = 0, \quad \star d\star F = J \quad (1.34)$$

F étant défini comme étant la différentielle du potentiel de jauge ($F = dA$), la première équation est en fait une identité, puisque $d^2 = 0$. Toute la dynamique se trouve donc dans la deuxième équation.

Dans le cas général, que nous traitons ici, cette situation se reproduit. La première relation s'appelle l'**identité de Bianchi**, et s'écrit :

$$d_D F = 0 \quad (1.35)$$

Nous pouvons facilement la démontrer, grâce au calcul différentiel que nous venons de développer, et en remarquant que

$$[A, A \wedge A] = A \wedge A \wedge A - A \wedge A \wedge A = 0$$

On a alors

$$\begin{aligned} d_D F &= dF + [A, F] \\ &= d(dA + A \wedge A) + [A, dA + A \wedge A] \\ &= dA \wedge A - A \wedge dA + [A, dA] + [A, A \wedge A] \\ &= [dA, A] + [A, dA] \\ &= 0 \end{aligned}$$

D'autre part, remarquons qu'elle s'écrit aussi, en composantes :

$$[D_\mu, [D_\nu, D_\rho]] + [D_\nu, [D_\rho, D_\mu]] + [D_\rho, [D_\mu, D_\nu]] = 0 \quad (1.36)$$

où elle apparaît simplement comme une forme particulière de l'identité de Jacobi.

L'équation de Yang-Mills

Comme dans le cas d'un groupe de jauge $U(1)$ et de l'équation de Maxwell, la dynamique d'une théorie de jauge quelconque est donc entièrement contenue dans la deuxième équation. Dans sa version non-abélienne, elle porte le nom d'**équation de Yang-Mills**, et s'écrit :

$$\star d_D \star F = J \quad (1.37)$$

¹² J est la 1-forme qui désigne le courant électromagnétique. L'opérateur $\star$ est l'étoile de Hodge, qui échange les p -formes et les $(n - p)$ -formes (où n est la dimension de la variété M), selon

$$\star(dx^{\mu_1} \wedge \dots \wedge dx^{\mu_p}) = \frac{1}{(n-p)!} \epsilon_{\nu_1 \dots \nu_{n-p}}^{\mu_1 \dots \mu_p} dx^{\nu_1} \wedge \dots \wedge dx^{\nu_{n-p}}$$

où J est une 1-forme à valeurs dans $\text{End}(E)$, appelée le **courant**, qui joue le rôle de source. Nous avons de plus défini l'**étoile de Hodge** sur les formes à valeurs dans $\text{End}(E)$ à partir de l'étoile de Hodge usuelle par la relation

$$\star (T \otimes \omega) = T \otimes \star \omega \quad (1.38)$$

On peut vérifier que l'équation de Yang-Mills se réduit bien à l'équation de Maxwell dans le cas d'un groupe de jauge $U(1)$. En effet, dans ce cas, les sections de $\text{End}(E) \otimes \wedge T^*M$ sont simplement des formes à valeurs complexes. C'est vrai en particulier pour A et F . Cela signifie que d_D agit sur ces formes comme la différentielle ordinaire, puisque le commutateur gradué est alors identiquement nul. Par exemple, $d_D F = dF$. Attention cependant, puisque d_D reste non trivial sur les formes à valeurs dans E . En effet, il n'y a aucune raison pour que le terme en $A \wedge \cdot$ disparaisse.

Etant une équation sur le champ F , il est évident a priori que l'équation de Maxwell est, comme il se doit, invariante de jauge. L'évidence est moins flagrante pour l'équation de Yang-Mills, puisque dans le cas général, F est covariant et pas invariant. Néanmoins, comme on peut le vérifier assez facilement, (1.37) est bien invariante de jauge. En effet, $g \in \mathcal{G}$ agit sur $\tilde{D}$ comme

$$\begin{aligned} \tilde{D}'_\mu T &= [D'_\mu, T] = [g D_\mu g^{-1}, T] \\ &= g [D_\mu, g^{-1} T g] g^{-1} \\ &= \left(\text{Ad}(g) \tilde{D}_\mu \text{Ad}(g^{-1}) \right) T \end{aligned} \quad (1.39)$$

où $\text{Ad}(g)T = gTg^{-1}$ est l'action adjointe de g . On en déduit

$$d_{D'} = \text{Ad}(g) d_D \text{Ad}(g^{-1}) \quad (1.40)$$

et puisque d'autre part, $F' = gFg^{-1} = \text{Ad}(g)F$ et que l'étoile de Hodge est complètement transparente, cela donne

$$\star d_{D'} \star F' = \text{Ad}(g) \star d_D \star \text{Ad}(g^{-1}) \text{Ad}(g)F = g(\star d_D \star F)g^{-1} = J' \quad (1.41)$$

où $J' = gJg^{-1}$ est bien entendu le transformé de jauge de la 1-forme J à valeurs dans $\text{End}(E)$.

Pour conclure sur l'équation de Yang-Mills et l'identité de Bianchi, nous allons les réécrire sur l'espace de Minkowski, en séparant les parties magnétiques (B) et électriques (E), par analogie avec l'électromagnétisme. L'espace-temps se compose donc de la dimension temporelle (t) et de l'espace à trois dimensions (x^i). On décompose F et J en formes sur l'espace (en dx^i uniquement) :

$$F = B + E \wedge dt, \quad J = j - \rho dt \quad (1.42)$$

De plus, en travaillant dans la jauge statique ($A_t = 0$), la connexion s'écrit

$$D_t = \partial_t, \quad D_i = \partial_i + A_i \quad (1.43)$$

En remplaçant dans la définition de F , on identifie alors les composantes électriques et magnétiques :

$$E = E_i dx^i, \quad B = \frac{1}{2} \epsilon_{ijk} B^i dx^j \wedge dx^k \quad (1.44)$$

avec les expressions

$$E_i = -\partial_t A_i, \quad B^i = \epsilon^{ijk} (\partial_j A_k - \partial_k A_j + [A_j, A_k]) \quad (1.45)$$

On peut alors réécrire l'identité de Bianchi,

$$\partial_i B^i + [A_i, B^i] = 0, \quad \partial_t B^i + \epsilon^{ijk} (\partial_j E_k + [A_j, E_k]) = 0 \quad (1.46)$$

et l'équation de Yang-Mills,

$$\partial_i E^i + [A_i, E^i] = \rho, \quad -\partial_t E^i + \epsilon^{ijk} (\partial_j B_k + [A_j, B_k]) = j^i \quad (1.47)$$

Ces quatre équations sont similaires aux quatre équations de Maxwell, auxquelles se sont ajoutés bien sûr les commutateurs non-abéliens, qui sont responsables de la non-linéarité de la théorie de jauge (le champ de jauge non-abélien est autocouplé : les gluons interagissent entre eux, alors que les photons s'ignorent).

1.1.3 Le(s) point(s) sur les i

Faisons d'abord le point sur les transformations de jauge, et plus exactement les transformations infinitésimales. Pour les écrire, nous devons paramétrer l'élément g du groupe de jauge (plus précisément, $g \in \mathcal{G}$), sous la forme $g = e^{-\lambda}$. Nous l'avons déjà fait, de façon un peu cavalière, pour A , avec néanmoins le résultat correct (1.8). En faisant de même pour F à partir de l'expression (1.14), on trouve donc :

$$\delta A_\mu = \partial_\mu \lambda + [A_\mu, \lambda], \quad \delta F = [F, \lambda] \quad (1.48)$$

Pour être moins cavalier, il aurait fallu définir proprement λ , ce que nous pouvons faire désormais. Tout d'abord, il suffit de supposer que nous travaillons avec un groupe de jauge compact et connexe (voir note 7 page 12) pour s'assurer de la surjectivité de l'application exponentielle définie sur $\mathfrak{g}$. D'autre part, puisque le fibré n'est pas trivial, λ est une section de $\text{End}(E)$, et bien sûr $\lambda(p) \in \mathfrak{g}$ en tout point p de M . Autrement dit, dans notre nouveau langage différentiel, λ est une 0-forme à valeurs dans $\text{End}(E)$ qui vit dans $\mathfrak{g}$. Et d'après (1.32) et (1.30), la transformation infinitésimale δ s'écrit alors simplement :

$$\delta A = d_D \lambda, \quad \delta F = d_D^2 \lambda \quad (1.49)$$

ce qui est une généralisation élégante des transformations de jauge abéliennes

$$\delta A = d\lambda, \quad \delta F = 0 \quad (= d^2 \lambda)$$

Enfin, concluons cette section en faisant le lien entre la notation des mathématiciens et celle des physiciens. Que ce soit pour le potentiel de jauge A , la courbure F ou

le paramètre de jauge λ , la grandeur du physicien Q_p et celle du mathématicien Q_m diffèrent du facteur $-i$:

$$Q_m = -i Q_p \quad (1.50)$$

Ainsi, pour un physicien, les expressions (1.33) et (1.15) pour la courbure s'écriront

$$F = dA - i A \wedge A \quad \text{soit} \quad F_{\mu\nu} = \partial_\mu A_\nu - \partial_\nu A_\mu - i [A_\mu, A_\nu] \quad (1.51)$$

tandis que les transformations de jauge infinitésimales (1.48) deviennent identiques aux transformations (1.4) que nous avons présentées au départ :

$$\delta A_\mu = \partial_\mu \lambda - i [A_\mu, \lambda] = D_\mu \lambda, \quad \delta F_{\mu\nu} = i [\lambda, F_{\mu\nu}] \quad (1.52)$$

1.2 La transformation de Seiberg-Witten (1)

1.2.1 Branes et action de Born-Infeld

La théorie des cordes contient la gravitation (à travers la supergravité, version supersymétrique de la Relativité Générale), et les théories de jauge. Comme nous l'avons dit dans l'introduction, la métrique de la Relativité Générale est l'un des modes de masse nulle des cordes fermées. Mais en théorie des cordes, elle est naturellement accompagnée par le champ antisymétrique $B_{\mu\nu}$, qui joue le rôle d'un champ électromagnétique de fond. D'autre part, les cordes ouvertes, dont les extrémités sont attachées sur des branes, reproduisent les théories de jauge : une brane unique pour une théorie $U(1)$, un paquet de n branes empilées pour une théorie $U(n)$.

Lorsqu'une D-brane est plongée dans un champ de fond $g_{\mu\nu} + B_{\mu\nu}$ quelconque, sa dynamique est décrite par l'action de Born-Infeld¹³

$$S_{BI} = \frac{1}{g_s} \int dx^\mu \sqrt{\det(g + B + 2\pi\alpha' F)} \quad (1.53)$$

Cette action est exacte si l'on néglige les dérivées des champs (en particulier, elle est exacte en α' à cette approximation près, c'est-à-dire que les seules corrections sont des corrections dérivatives, contenant bien sûr autant de α' que nécessaire). Au premier ordre en α' , elle redonne l'action de Maxwell sur la brane, comme on peut s'y attendre. C'est par contre un problème ouvert que de déterminer les différents types de correction dérivatives (en dérivées de g , B , F ou tout mélange) de cette action.

Lorsque plusieurs branes sont coïncidentes, la théorie de jauge devient non-abélienne. Au premier ordre en α' , on doit cette fois retrouver une théorie de Yang-Mills. Et la forme elle-même de l'action de Born-Infeld doit être généralisée. En particulier, en effet, le champ F est alors à valeurs matricielles, et la combinaison $B + F$ n'a plus de sens. Cela oblige à prendre la trace (matricielle) de F . Mais plus grave encore,

¹³ C'est l'action sur le secteur de Neveu-Schwarz, ici en l'absence de dilaton. Pour l'action totale, il faut lui ajouter l'action de Chern-Simons pour le secteur de Ramond-Ramond.

le développement du déterminant (sur les indices d'espace-temps) est indifférent à l'ordre des facteurs, qui est pourtant important. Prenons un exemple simple,

$$\det \begin{pmatrix} a & b \\ c & d \end{pmatrix} = ad - bc = da - cb$$

ce qui pose problème lorsque les coefficients a, b, c, d sont par ailleurs des matrices, et ne commutent pas. Cette fois, résoudre le problème signifie définir l'ordre des différents coefficients (qui sont ici les composantes spatio-temporelles de F) avant d'en prendre la trace. Le plus simple consiste bien sûr à considérer la trace symétrisée, c'est-à-dire en moyennant sur tous les ordres possibles [15], mais cette solution ne suffit pas, elle doit être corrigée à partir de l'ordre F^6 [16]. Cependant, la situation est complexe, puisque pour une théorie de jauge non-abélienne, les dérivées de F sont très peu discernables des commutateurs de F :

$$[D_i, D_j]F_{kl} = [F_{ij}, F_{kl}]$$

Or ces commutateurs sont justement la différence entre ordres différents des facteurs F_{ij} . Trouver le bon ordre dans le développement de l'action de Born-Infeld non-abélienne revient donc à en calculer les corrections dérivatives en F . En anticipant sur les théories de jauge non-commutatives (et en particulier sur les star-produits), la situation est la même si l'on veut écrire l'action de Born-Infeld pour un champ F non-commutatif, même de groupe de jauge $U(1)$. En effet, cette fois la différence entre deux expressions ne différant que par l'ordre des facteurs est du type

$$a \star b - b \star a = [a \star, b] = i\theta^{ij} \partial_i a \partial_j b + \dots$$

et met comme précédemment en jeu des expressions dérivatives en a et b . Cette similarité entre la situation non-commutative et la situation non-abélienne "ordinaire" (commutative) n'est pas surprenante. Elle n'est que le reflet de l'équivalence de Morita, qui relie des théories définies sur des espaces présentant des paramètres de non-commutativité différents, au prix d'un changement du rang du groupe de jauge.

Une caractéristique importante de l'action de Born-Infeld est qu'elle est invariante sous la transformation de Seiberg-Witten qui nous intéresse ici. Cette transformation associe à chaque champ $g, B, F, \dots$ un champ $\hat{g}, \hat{B}, \hat{F}, \dots$ défini sur la version non-commutative de la brane initiale. L'action de Born-Infeld sur cette brane s'écrit sous la même forme, en fonction des nouveaux champs¹⁴ :

$$S_{BI}^\theta = \frac{1}{\hat{g}_s} \int dx^\mu \sqrt{\det(\hat{g} + \hat{B} + 2\pi\alpha' \hat{F})} \quad (1.54)$$

Encore mieux, elle est l'unique action invariante sous cette transformation, et bien que nous ne nous intéressions pas davantage ici à ses corrections dérivatives, c'est l'un des moyens envisagés pour les calculer [17].

¹⁴ Comme nous l'avons déjà remarqué, cette expression n'a de sens qu'à l'approximation des champs lentement variables. Trouver ses corrections est équivalent trouver les corrections de l'action de Born-Infeld initiale.

1.2.2 La transformation des champs de fond

Nous allons maintenant nous pencher sur cette transformation, dans un premier temps en détaillant son action sur les champs de fond g et B . Cela nous permettra en particulier de nous familiariser avec la description d'un espace non-commutatif, et son paramètre θ . Nous aborderons à la section 1.4 l'action de la transformation de Seiberg-Witten sur le secteur de jauge ; en trouver la forme explicite est d'ailleurs l'objectif que nous poursuivons tout au long de ce chapitre.

Une brane non-commutative est tout simplement une brane sur laquelle les coordonnées, que nous noterons $\hat{x}^\mu$, ne commutent plus :

$$[\hat{x}^\mu, \hat{x}^\nu] = i\theta^{\mu\nu} \quad (1.55)$$

où $\theta^{\mu\nu}$ est une matrice (antisymétrique), constante pour l'instant, appelée paramètre de non-commutativité. La situation "ordinaire" est alors simplement obtenue lorsque $\theta = 0$. Ainsi, nous venons simplement d'ajouter des degrés de liberté aux théories envisageables (en supposant bien sûr que l'on sache écrire des théories consistantes sur un tel espace). Nous pouvons bien sûr plonger cette brane dans le champ de fond $\hat{g}_{\mu\nu} + \hat{B}_{\mu\nu}$, créé par un rayonnement de cordes fermées.

Seiberg et Witten ont montré qu'une telle description n'était finalement pas si étrange, puisqu'elle était équivalente à une brane ordinaire plongée dans le champ $g + B$, avec les relations suivantes entre les champs¹⁵ :

$$\frac{1}{g + B} = \theta + \frac{1}{\hat{g} + \hat{B}} \quad (1.56)$$

C'est-à-dire que la quantité $\theta + (\hat{g} + \hat{B})^{-1}$ est conservée lorsque le paramètre θ varie.

Inversement, on peut dire qu'une brane ordinaire plongée dans un champ B antisymétrique non nul est équivalente à une certaine brane non-commutative, où le champ antisymétrique a été "remplacé" par le paramètre θ . Nous reviendrons bien sûr sur ce calcul, mais d'un point de vue conceptuel, cela ressemble fortement à l'essence de la Relativité Générale : au lieu de décrire la théorie par un champ complexe sur un espace simple (plat), toutes les subtilités sont cachées "sous le tapis", dans l'espace-temps qui acquiert une courbure tandis que le champ complexe disparaît. Ici, plutôt que de décrire le couplage de la brane avec un champ de fond électromagnétique, on peut donc encoder ce champ dans la non-commutativité de l'espace, et comprendre comment fonctionne la théorie sur cet espace plus subtil. Il est d'autant plus remarquable qu'alors que la courbure (la "non-platitudo") est liée à la partie symétrique $g_{\mu\nu}$, la non-commutativité est liée à la partie antisymétrique $B_{\mu\nu}$.

Signalons d'autre part que le champ B comporte, à l'instar d'un champ électromagnétique usuel, une partie électrique et une partie magnétique. Cependant, lorsque l'on efface le champ B dans la description non-commutative équivalente, la partie magnétique est responsable de la non-commutativité espace/espace tandis que la partie électrique est responsable de la non-commutativité espace/temps. Au niveau de la théorie des champs, ce deuxième type de cas est non causal, même si la théorie des cordes semble réguler ce problème [11, 12]. Cela explique pourquoi les travaux

¹⁵En posant $\alpha' = 1$

décrivant les théories non-commutatives se limitent en général à des composantes magnétiques uniquement. En ce qui nous concerne, nous ne nous intéresserons qu'à la formulation de la transformation de Seiberg-Witten, permettant de passer d'une description à l'autre, et pas aux caractéristiques physiques des théories de jauge non-commutatives. De telles précautions seront donc inutiles.

Revenons enfin à l'équivalence elle-même, pour étudier de plus près ces formules. Puisque cette équivalence est valable quel que soit θ , c'est une infinité de descriptions qui sont ainsi reliées par la transformation de Seiberg-Witten. On peut par exemple fixer le paramètre θ , et en déduire $\hat{g}$ et $\hat{B}$, parties symétriques et antisymétriques de l'expression. On retrouve bien sûr les champs initiaux g et B lorsque $\theta = 0$, ce qui est bien le cas d'une brane "ordinaire" dont les coordonnées commutent :

$$\begin{aligned}
 2\hat{g} &= (\hat{g} + \hat{B}) + {}^t(\hat{g} + \hat{B}) \\
 &= \frac{1}{\frac{1}{g+B} - \theta} + \frac{1}{\frac{1}{g-B} + \theta} \\
 &= \frac{1}{\frac{1}{g+B} - \theta} \left[\frac{1}{g-B} + \theta + \frac{1}{g+B} - \theta \right] \frac{1}{\frac{1}{g-B} + \theta} \\
 &= \frac{1}{\frac{1}{g+B} - \theta} \frac{1}{g+B} 2g \frac{1}{g-B} \frac{1}{\frac{1}{g-B} + \theta} \\
 \hat{g} &= \frac{1}{1 - (g+B)\theta} g \frac{1}{1 - \theta(g-B)} \tag{1.57}
 \end{aligned}$$

De la même manière, on trouve la partie antisymétrique

$$\hat{B} = \frac{1}{\frac{1}{g+B} - \theta} \left[\theta + \frac{1}{g+B} B \frac{1}{g-B} \right] \frac{1}{\frac{1}{g-B} + \theta} \tag{1.58}$$

Inversement, on peut choisir la valeur de $\hat{B}$, par exemple nulle, et on en déduit $\hat{g}$ et θ . Dans ce cas de figure, on trouve

$$\hat{g} = g - Bg^{-1}B, \quad \theta = -\frac{1}{g+B} B \frac{1}{g-B} \tag{1.59}$$

Supposons alors que B soit inversible. θ l'est alors aussi et

$$\theta^{-1} = B - gB^{-1}g \tag{1.60}$$

Dernier exemple, toujours si B est inversible : fixons $\theta = B^{-1}$, alors

$$\begin{aligned}
 \hat{g} + \hat{B} &= \left[\frac{1}{g+B} - \frac{1}{B} \right]^{-1} \\
 &= \left[\frac{1}{g+B} (B - g - B) \frac{1}{B} \right]^{-1} \\
 &= -Bg^{-1}(g+B) \\
 &= -B - Bg^{-1}B \tag{1.61}
 \end{aligned}$$

c'est-à-dire $\hat{B} = -B$ et $\hat{g} = -Bg^{-1}B$.

1.3 Le star-produit

1.3.1 La formule de Moyal-Weyl

Pour travailler dans un espace non-commutatif, nous allons commencer par déformer le produit sur la variété de base, en définissant un star-produit. Commençons par l'exemple le plus simple que nous utiliserons d'ailleurs par la suite : la formule de Moyal-Weyl.

$$\begin{aligned} f \star g &= f e^{\frac{i}{2} \overleftarrow{\partial}_\rho \theta^{\rho\sigma} \overrightarrow{\partial}_\sigma} g \\ &= fg + \frac{i}{2} (\partial_\rho f) \theta^{\rho\sigma} (\partial_\sigma g) - \frac{1}{8} (\partial_\rho \partial_\kappa f) \theta^{\rho\sigma} \theta^{\kappa\tau} (\partial_\sigma \partial_\tau g) + O(\theta^3) \end{aligned} \quad (1.62)$$

où θ est un tenseur antisymétrique d'ordre deux, constant, qui joue le rôle de paramètre de déformation et donc de non-commutativité. En effet, si l'on applique cette formule aux fonctions coordonnées :

$$[x^\mu \star x^\nu] = x^\mu \star x^\nu - x^\nu \star x^\mu = i\theta^{\mu\nu} \quad (1.63)$$

et les coordonnées ne commutent plus, comme nous le souhaitions. Ainsi, en déformant le produit ($f \star g$ est égal à fg plus des corrections en puissance de θ), nous avons bien déformé l'espace.

Remarquons que la formule de Moyal-Weyl est valide seulement lorsque le paramètre θ est constant. La réciproque est vraie aussi, mais pour le voir, il nous faut d'abord définir le star-produit d'une manière un peu plus générale.

1.3.2 Star-produits et structures de Poisson

Un star-produit est un produit *associatif* qui s'écrit sous la forme

$$f \star g = fg + \sum_{n=1}^{\infty} (i\hbar)^n B_n(f, g) \quad (1.64)$$

où $\hbar$ est un paramètre formel et les B_n sont des applications bilinéaires.

On peut montrer qu'un star-produit est toujours la déformation d'une structure de Poisson :

$$[f \star g] = i\hbar \{\partial f, \partial g\} + O(\hbar^2), \text{ avec } \{\partial f, \partial g\} = \partial_i f \theta^{ij}(x) \partial_j g \quad (1.65)$$

La déformation au sens où tout autre déformation de la même structure de Poisson donne un star-produit équivalent. On peut en effet définir un star-produit $\tilde{\star}$ à l'aide d'un opérateur inversible D par :

$$D(f \tilde{\star} g) = Df \star Dg, \text{ avec } Df \equiv f + \sum_{n=1}^{\infty} \hbar^n D_n(f) \quad (1.66)$$

où les D_n sont des opérateurs linéaires. Deux tels star-produits sont alors dits équivalents.

Si l'on explicite cette transformation au premier ordre en $\hbar$, on trouve l'expression de $\tilde{B}_1$:

$$\tilde{B}_1(f, g) = B_1(f, g) + iD_1(fg) - iD_1(f)g - ifD_1(g) \quad (1.67)$$

de telle sorte que si l'on sépare B_1 en ses parties symétrique B_1^+ et antisymétrique B_1^- , cette transformation ne modifie que la partie symétrique : $(\tilde{B}_1)^- = B_1^-$.

Or c'est la partie antisymétrique qui définit la structure de Poisson, puisque :

$$\{\partial f^\theta, \partial g\} = 2B_1^-(f, g)$$

Nous venons donc de vérifier que deux star-produits équivalents définissent bien la même structure de Poisson. Il n'est pas du tout évident qu'inversement, étant donné une structure de Poisson, on puisse la déformer en star-produit, mais c'est néanmoins vrai. Ce résultat est dû à Kontsevich [37], la procédure s'appelant transformation de formalité. Elle est définie à des changements de variables près, qui conduisent à autant de star-produits équivalents. Comme nous le verrons, c'est le fait que θ soit une structure de Poisson qui assure l'associativité du star-produit. Il existe donc une identification naturelle entre les structures de Poisson et les classes d'équivalences de star-produit.

Avant d'aborder cette transformation de formalité, nous allons voir un exemple concret de cette relation. En effet, θ définit une structure de Poisson lorsque son crochet $\{.^\theta, .\}$ vérifie l'identité de Jacobi. Cela signifie, pour θ^{ij} antisymétrique :

$$\theta^{il}\partial_l\theta^{jk} + \theta^{jl}\partial_l\theta^{ki} + \theta^{kl}\partial_l\theta^{ij} = 0 \quad (1.68)$$

Pour relier cette identité est reliée à l'associativité du star-produit, écrivons explicitement cette associativité sur un exemple. La formule de Moyal-Weyl est, cette fois, trop simple, puisque la relation d'associativité est triviale à tous les ordres¹⁶. La première correction lorsque θ n'est plus constant intervient à l'ordre 2, et s'écrit par exemple [7] :

$$f \star g = f \star_0 g - \frac{1}{12} \theta^{il}\partial_l\theta^{jk}(\partial_i\partial_j f \partial_k g - \partial_j f \partial_i\partial_k g) + O(\theta^3) \quad (1.69)$$

où $\star_0$ désigne la formule de Moyal-Weyl. De plus, comme au début de cette section, le paramètre formel $\hbar$ a été absorbé par θ . Il est alors facile de vérifier que l'associativité de $\star$ s'écrit :

$$(f \star g) \star h - f \star (g \star h) = \frac{1}{6} (\theta^{il}\partial_l\theta^{jk} + \theta^{jl}\partial_l\theta^{ki} + \theta^{kl}\partial_l\theta^{ij}) \partial_i f \partial_j g \partial_k h + O(\theta^3) \quad (1.70)$$

C'est-à-dire que $\star$ est associatif si et seulement si θ est de Poisson.

J'en profite pour faire ici un commentaire supplémentaire. Supposons que θ soit inversible et appelons ω son inverse. Alors, $\partial_l\theta^{jk} = \theta^{ja}\theta^{kb}\partial_l\omega_{ab}$. Et la condition (1.68) pour que θ définisse une structure de Poisson devient

$$\partial_a\omega_{bc} + \partial_b\omega_{ca} + \partial_c\omega_{ab} = 0 \quad \text{i.e.} \quad H = d\omega = 0 \quad (1.71)$$

¹⁶ On pourrait bien sûr le vérifier. Cela vient du fait que c'est cette relation qui fixe ordre par ordre les coefficients de la formule. On pouvait d'ailleurs s'attendre à ce qu'elle ne contienne rien de plus, puisque θ est constant, et donc nécessairement de Poisson.

De même, l'associateur du star-produit associé s'écrit alors

$$(f \star g) \star h - f \star (g \star h) = \frac{1}{6} \theta^{ia} \theta^{jb} \theta^{kc} H_{abc} \partial_i f \partial_j g \partial_k h + O(\theta^3) \quad (1.72)$$

1.3.3 La transformation de formalité

Cette transformation est due à Kontsevich [37], et est en réalité un ensemble $\mathcal{U}$ de transformations multilinéaires. U_n associe à n champs polyvectoriels (d'ordres k_i) un opérateur m -différentiel, où m est déterminé par la condition

$$m = 2 - 2n + \sum_{k=1}^n k_i \quad (1.73)$$

En particulier, U_0 est un opérateur à deux variables : par définition, c'est le produit usuel des fonctions

$$U_0(f, g) = fg \quad (1.74)$$

Pour $n \geq 1$, U_n vérifie la condition de formalité, qui le relie à tous les U_i précédents. Mais avant de pouvoir l'évoquer un peu plus précisément, il nous faut introduire un peu d'algèbre.

Crochet de Schouten-Nijenhuis Commençons par définir le commutateur de Schouten-Nijenhuis sur les champs polyvectoriels,

$$\begin{aligned} [\xi_1 \wedge \cdots \wedge \xi_p, \eta_1 \wedge \cdots \wedge \eta_q]_S &= \sum_{i,j} (-1)^{i+j} [\xi_i, \eta_j] \wedge \xi_1 \wedge \cdots \wedge \hat{\xi}_i \wedge \cdots \wedge \hat{\eta}_j \wedge \cdots \wedge \eta_q \\ [\xi_1 \wedge \cdots \wedge \xi_p, f]_S &= \sum_i (-1)^i \xi_i(f) \xi_1 \wedge \cdots \wedge \hat{\xi}_i \wedge \cdots \wedge \xi_p \end{aligned} \quad (1.75)$$

où les ξ_i et η_j sont des champs vectoriels, et f une fonction — la notation $\hat{\xi}_i$ indique que le champ ξ_i est omis. Ce crochet est un opérateur de degré -1 sur l'ordre polyvectoriel : le crochet d'un champ d'ordre p et d'un champ d'ordre q est un champ d'ordre $(p + q - 1)$. En particulier, un champ bivectoriel $\theta = \frac{1}{2} \theta^{ij} \partial_i \wedge \partial_j$ définit l'opérateur de cobord

$$\mathbf{d}_\theta = -[\cdot, \theta]_S \quad (1.76)$$

Ainsi, θ définit une structure de Poisson si et seulement si il vérifie l'une des trois relations équivalentes

$$[\theta, \theta]_S = 0 \quad \Leftrightarrow \quad \mathbf{d}_\theta \theta = 0 \quad \Leftrightarrow \quad \mathbf{d}_\theta^2 = 0 \quad (1.77)$$

Crochet de Gerstenhaber D'autre part, on définit le crochet de Gerstenhaber sur les opérateurs multidifférentiels

$$[\phi_1, \phi_2]_G = \phi_1 \circ \phi_2 - (-1)^{(p_1-1)(p_2-1)} \phi_2 \circ \phi_1 \quad (1.78)$$

où ϕ_1 est d'ordre p_1 et ϕ_2 d'ordre p_2 . Ce crochet un opérateur de degré -1 , tout comme la composition $\circ$, qui est définie, pour toutes fonctions $f_1, \dots, f_{p_1+p_2-1}$, par :

$$\begin{aligned} & (\phi_1 \circ \phi_2)(f_1, \dots, f_{p_1+p_2-1}) \\ &= \sum_{i=1}^{p_1} (-1)^{(i-1)(p_2-1)} \phi_1(f_1, \dots, \phi_2(f_i, \dots, f_{i+p_2-1}), \dots, f_{p_1+p_2-1}) \end{aligned} \quad (1.79)$$

Comme on peut s'en douter d'après ces formules, ϕ_1 et ϕ_2 n'ont pas réellement besoin d'être des opérateurs différentiels pour définir une telle algèbre de Gerstenhaber. Ces définitions sont valables plus généralement pour des homomorphismes d'algèbres. Nous y reviendrons d'ailleurs dans le chapitre 2, où les opérateurs différentiels seront remplacés par le complexe de Hochschild $\bigoplus_n \text{Hom}(A^{\otimes n}, A)$ d'une algèbre A . Etant donné un produit $\star$, on peut encore une fois définir l'opérateur de cobord associé,

$$\mathbf{d}_\star = -[\cdot, \star]_G \quad (1.80)$$

Ainsi, $\star$ est associatif, c'est-à-dire définit un star-produit si et seulement si il vérifie l'une des trois relations équivalentes

$$[\star, \star]_G = 0 \quad \Leftrightarrow \quad \mathbf{d}_\star \star = 0 \quad \Leftrightarrow \quad \mathbf{d}_\star^2 = 0 \quad (1.81)$$

Condition de formalité Cette condition impose des restrictions sur les transformations U_n . Je n'en donnerai qu'un aperçu ici, au sens où je laisserai indéfinis les signes $\pm$. On pourra se référer à [37, 38] pour plus de précision. Pour tout $n \geq 1$,

$$\begin{aligned} & \frac{1}{2} \sum_{I+J=(1,\dots,n)} \pm [U_{|I|}(\alpha_I), U_{|J|}(\alpha_J)]_G \\ &= \sum_{i < j} \pm U_{n-1}([\alpha_i, \alpha_j]_S, \alpha_1, \dots, \hat{\alpha}_i, \dots, \hat{\alpha}_j, \dots, \alpha_n) \end{aligned} \quad (1.82)$$

où I et J sont des multi-indices, éventuellement vides, d'intersection vide. $I + J$ désigne leur union, et $|I|$ le cardinal de I . De plus, comme précédemment, les quantités $\hat{\alpha}_i$ sont omises.

Star-produit et transformation de formalité D'après la condition (1.73), pour tout $n > 0$ et tout champ polyvectoriel α d'ordre p , $U_n(\alpha, \theta, \dots, \theta)$ est un opérateur p -différentiel (où bien sûr, θ est un opérateur bidifférentiel). En particulier, tout bivecteur de Poisson θ définit un star-produit

$$\star = \sum_{n=0}^{\infty} \frac{(i\hbar)^n}{n!} U_n(\theta, \dots, \theta) \quad (1.83)$$

c'est-à-dire, pour deux fonctions f et g ,

$$f \star g \equiv \sum_{n=0}^{\infty} \frac{(i\hbar)^n}{n!} U_n(\theta, \dots, \theta)(f, g) = fg + \frac{i\hbar}{2} \partial_i f \theta^{ij} \partial_j g + \dots \quad (1.84)$$

Si θ est une structure de Poisson (et donc $[\theta, \theta]_S = 0$), alors la condition de formalité implique que $\star$ est associatif, et est donc un star-produit. Ce que nous avons plus haut vérifié au premier ordre sur un cas particulier est donc toujours valable.

Nous avons vu précédemment qu'un star-produit définit une structure de Poisson, en voici donc la réciproque : on peut recréer tout un star-produit à partir simplement de sa structure de Poisson associée. Plus précisément, c'est tout une classe d'équivalence de star-produits qui définit la même structure de Poisson, et même si nous ne l'avons pas démontré, cette réciproque-là est vraie aussi. En effet, les U_n ne sont pas uniques, mais les star-produits définis par deux séries d'opérateurs différents, à partir du même θ , sont bien équivalents.

Revenons à la formalité, qui nous permet de définir bien d'autres objets utiles. Ainsi, la transformation Φ ,

$$\Phi(\alpha) = \sum_{n=0}^{\infty} \frac{(i\hbar)^n}{n!} U_{n+1}(\alpha, \theta, \dots, \theta) \quad (1.85)$$

définit un opérateur p -différentiel à partir du champ α d'ordre p . En particulier, $\Phi(f)$ est une fonction si f est une fonction ; $\Phi(\xi)$ est un opérateur différentiel linéaire si ξ est un champ vectoriel. En appliquant la condition de formalité, si θ est une structure de Poisson, on trouve alors

$$d_\star \Phi(\alpha) = i\hbar \Phi(d_\theta \alpha) \quad (1.86)$$

où $d_\star$ et d_θ sont les opérateurs de cobord que nous avons définis plus haut, respectivement pour les structures de Gerstenhaber et de Schouten-Nijenhuis.

1.4 La transformation de Seiberg-Witten (2)

Maintenant que nous disposons de l'outil du star-produit pour calculer efficacement sur l'espace non-commutatif, nous allons pouvoir revenir à l'action de la transformation de Seiberg-Witten sur le secteur de jauge. Elle associe la théorie ordinaire à la théorie non-commutative correspondante (de même groupe de jauge¹⁷), que nous allons donc tout d'abord présenter.

1.4.1 Les théories de jauge non-commutatives

Sur l'espace non-commutatif, on peut définir une théorie de jauge comme précédemment, simplement en remplaçant toutes les multiplications par des star-produits¹⁸.

¹⁷Pour autant qu'on sache ce qu'est la version non-commutative d'un groupe. Disons plutôt de même algèbre de jauge, puisque de toute façon tous les champs avec lesquels nous travaillons vivent dans (l'algèbre enveloppante de) l'algèbre de jauge.

¹⁸On absorbera ici le facteur $\hbar$ dans θ , comme on l'avait déjà fait pour la formule de Moyal-Weyl.

On notera avec un chapeau toutes les quantités de cette théorie ($\hat{\lambda}$, $\hat{A}_\mu$, $\hat{F}_{\mu\nu}$). Ainsi, il y a deux raisons pour lesquelles les produits de $\hat{\lambda}$, $\hat{A}$, $\hat{F}$... ne commutent plus. Non seulement ce sont des produits matriciels (dans l'algèbre de Lie $\mathfrak{g}$ du groupe de jauge G), mais leurs composantes sont aussi multipliées par un star-produit. La courbure est définie par

$$\hat{F}_{\mu\nu} = \partial_\mu \hat{A}_\nu - \partial_\nu \hat{A}_\mu - i[\hat{A}_\mu * \hat{A}_\nu] \quad (1.87)$$

et la transformation de jauge $\hat{\delta}$ de paramètre $\hat{\lambda}$ s'écrit

$$\hat{\delta} \hat{A}_\mu = \partial_\mu \hat{\lambda} + i[\hat{\lambda} * \hat{A}_\mu] = \hat{D}_\mu \hat{\lambda}, \quad \hat{\delta} \hat{F}_{\mu\nu} = i[\hat{\lambda} * \hat{F}_{\mu\nu}] \quad (1.88)$$

Une première conséquence évidente de ces formules, est que même avec un groupe de jauge $U(1)$ la théorie ne se réduit pas à la théorie $U(1)$ ordinaire. En particulier, le champ de jauge est autocouplé, et ressemble fortement à un champ de jauge non-abélien.

Une deuxième conséquence est que les champs vivent en réalité, non pas dans l'algèbre $\mathfrak{g}$, mais dans son algèbre enveloppante. C'est-à-dire l'algèbre formée par les générateurs, leurs commutateurs et leurs anticommutateurs. En effet, si l'on développe l'expression du star-commutateur dans le cas non-abélien, on obtient¹⁹

$$[f * g] = [f, g] + \frac{i}{2} \{\partial f^\theta, \partial g\} - \frac{1}{8} [\partial \partial f^{\theta\theta}, \partial \partial g] + \dots \quad (1.89)$$

Un terme sur deux est un commutateur d'algèbre de Lie, l'autre étant un anticommutateur. Dans le cas de $SU(n)$, par exemple, l'algèbre de jauge non-commutative est donc en réalité $\mathfrak{u}(n)$, son algèbre enveloppante.

1.4.2 La transformation du secteur de jauge

La transformation de Seiberg-Witten transforme une théorie de jauge commutative en sa contrepartie non-commutative en respectant les transformations de jauge (ou de manière plus générale deux théories vivant dans des espaces à paramètres θ différents). C'est-à-dire qu'elle transporte toute la structure du fibré d'un espace à l'autre. Autrement dit, et c'est ce qui va nous permettre de la calculer, elle transforme les orbites de jauge en orbites de jauge.

Ainsi, si l'on écrit les grandeurs non-commutatives comme fonctions des grandeurs commutatives

$$\hat{A}_\mu = \hat{A}_\mu(A_\mu; \theta), \quad \hat{F}_{\mu\nu} = \hat{F}_{\mu\nu}(A_\mu; \theta), \quad \hat{\lambda} = \hat{\lambda}(\lambda, A_\mu; \theta) \quad (1.90)$$

¹⁹ Avec les notations suivantes, où l'opérande de gauche porte le(s) premier(s) indices de θ , de sorte que ces symboles sont tous deux antisymétriques :

$$\begin{aligned} \{\partial f^\theta, \partial g\} &= \theta^{ij} \{\partial_i f, \partial_j g\} = \theta^{ij} (\partial_i f \partial_j g + \partial_j g \partial_i f) \\ [\partial \partial f^{\theta\theta}, \partial \partial g] &= \theta^{ij} \theta^{kl} [\partial_i \partial_k f, \partial_j \partial_l g] \end{aligned}$$

la transformation se résume au diagramme suivant :

$$\begin{array}{ccc}
 A_\mu & \xrightarrow{\delta} & A_\mu + \delta A_\mu \\
 \downarrow & & \downarrow \\
 \hat{A}_\mu & \xrightarrow{\hat{\delta}} & \hat{A}_\mu + \hat{\delta} \hat{A}_\mu
 \end{array}
 \quad \Leftrightarrow \quad \widehat{A_\mu + \delta A_\mu} = \hat{A}_\mu + \hat{\delta} \hat{A}_\mu \quad \Leftrightarrow \quad \delta \hat{A}_\mu = \hat{\delta} \hat{A}_\mu$$

Il équivaut de transformer par δ la *fonction* $\hat{A}_\mu$ et de transformer par $\hat{\delta}$ la *connexion* non-commutative $\hat{A}_\mu$. C'est ce que nous appellerons l'équation de Seiberg-Witten :

$$\delta \hat{A}_\mu - \partial_\mu \hat{\lambda} = i[\hat{\lambda} \star \hat{A}_\mu] \quad (1.91)$$

1.4.3 Liberté de jauge

Par définition, la transformation de Seiberg-Witten agit sur les orbites de jauge. Cela signifie qu'en cherchant une expression explicite pour un potentiel de jauge nous trouverons nécessairement une infinité de solutions toutes bien entendues reliées par une transformation de jauge. Ainsi, trouver une expression explicite de cette transformation sur les champs de jauge est donc tout autant un problème de choix d'une "bonne" solution. Malheureusement, nous ne disposons a priori pas de critères pour savoir quelle solution serait "meilleure" qu'une autre, d'autant que cela peut dépendre de la méthode employée pour la chercher. Nous reviendrons sur cette difficulté par la suite, en décrivant la résolution récursive de l'équation (1.91) que j'ai mise en oeuvre pendant cette thèse.

1.5 Les différentes approches

Pour trouver la relation explicite entre le champ de jauge ordinaire et sa contrepartie non-commutative, on peut essentiellement utiliser trois méthodes :

- exprimer les courants qui couplent aux potentiels de Ramond-Ramond, ²⁰
- développer ordre par ordre l'équation de Seiberg-Witten, ²¹
- ou utiliser le formalisme des star-produits. ²²

Je ne m'attarderai pas sur la première méthode. Par contre nous allons mettre en oeuvre les deux autres, pour pouvoir en aborder la résolution dans les sections 1.6 et 1.7 respectivement.

1.5.1 Couplage aux potentiels Ramond-Ramond

Dans la partie Chern-Simons de l'action d'une p -brane, le champ $F_{\mu\nu}$ sur la brane joue le rôle de courant pour le potentiel de Ramond-Ramond C^{p-1} . Le principe de cette première méthode est donc de trouver une expression de ce courant en terme

²⁰On pourra se référer à [20, 21, 22, 23, 24, 25] pour plus de détails.

²¹Voir [26, 36].

²²Voir [29, 30, 31, 32].

des grandeurs caractéristiques de la théorie non-commutative. On disposera alors du champ électromagnétique commutatif $F_{\mu\nu} = F_{\mu\nu}(\hat{A}_\mu; \theta)$.

Je ne décrirai pas davantage le calcul ici, je me contenterai d'en donner les résultats.

Dans un espace non-commutatif de paramètre θ , une particule d'impulsion k dans une direction subit un étirement spatial dans les directions perpendiculaires, de longueur $l = k\theta$. Plus précisément, si le vecteur impulsion est k_j alors l'étirement spatial du quantum de champ vaut $l_i = k_j \theta^{ji}$. Notons $\hat{f}$ la valeur moyennée du champ autour d'une particule située en x et ainsi étirée

$$\hat{f}_{ij} = \int_0^1 \hat{F}_{ij}(x + \tau l) d\tau$$

Notons d'autre part W la ligne de Wilson calculée sur cette configuration (P indique que les produits dans la série exponentielle sont ordonnés le long de la ligne) :

$$W = P \exp(i l^i \int_0^1 \hat{A}_i(x + \tau l) d\tau)$$

On peut alors exprimer la valeur du quanta de champ "ordinaire" d'impulsion k en fonction du champ de la théorie non-commutative

$$F_{ij}(k) = \int dx \star [e^{ikx} \sqrt{\det(1 - \hat{f}\theta)} \frac{1}{1 - \hat{f}\theta} \hat{f} W]$$

où $\star[\cdot]$ signifie que tous les produits effectués dans les crochets sont des star-produits.

1.5.2 Développement ordre par ordre

En considérant les grandeurs non-commutatives comme des fonctions des grandeurs commutatives (cf 1.90) écrites comme série entière de θ , on peut résoudre l'équation (1.91) ordre par ordre en θ :

$$\begin{aligned} \hat{A}_\mu &= A_\mu + A_\mu^{(1)} + A_\mu^{(2)} + \dots \\ \hat{F}_{\mu\nu} &= F_{\mu\nu} + F_{\mu\nu}^{(1)} + F_{\mu\nu}^{(2)} + \dots \\ \hat{\lambda} &= \lambda + \lambda^{(1)} + \lambda^{(2)} + \dots \end{aligned}$$

où $A_\mu^{(n)}$ est d'ordre θ^n . L'ordre 0 est fixé par la valeur pour $\theta = 0$, c'est-à-dire lorsque l'espace est commutatif. On retrouve donc les valeurs commutatives A_μ , $F_{\mu\nu}$ et λ .

Dans l'équation (1.91), le membre de gauche est alors, à l'ordre n , $\delta A_\mu^{(n)} - \partial_\mu \lambda^{(n)}$, tandis que le membre de droite ne contient que des termes d'ordre inférieur ou égal à n (strictement inférieur pour une théorie U(1)). Voilà pourquoi il est possible de chercher une solution récursivement, ordre par ordre en θ .

On voit ici explicitement apparaître l'infinité de solutions due à la liberté de jauge. Il est important de noter que cette ambiguïté se retrouve à tous les ordres, et que le choix d'une solution à un ordre donné va influencer sur les équations aux ordres suivants. Pour être concrets, voyons tout d'abord le cas de la théorie abélienne. Déjà à l'ordre 0, l'équation est

$$\delta A_\mu^{(0)} - \partial_\mu \lambda^{(0)} = 0 \quad (1.92)$$

qui admet bien sûr la solution que nous avons implicitement choisie plus haut

$$A_\mu^{(0)} = A_\mu, \quad \lambda^{(0)} = \lambda \quad (1.93)$$

Mais tout autre représentant de l'orbite de jauge est aussi valable. Nous disposons cependant à cet ordre du critère supplémentaire qui veut que $\hat{A}_\mu$ coïncide avec A_μ lorsque $\theta = 0$. Par contre, à l'ordre n , il n'y a a priori aucun critère pour choisir une solution plutôt qu'une autre. Et toute solution "homogène" $(A_\mu^{(n)}; \Lambda^{(n)})$, c'est-à-dire telle que

$$\delta A_\mu^{(n)} - \partial_\mu \Lambda^{(n)} = 0 \quad (1.94)$$

peut être librement ajoutée à une solution pour en former une nouvelle. Ces solutions homogènes ne changent pas la physique, et peuvent se réécrire sous la forme de transformations de jauge [28] (quand ce ne sont pas de simples redéfinitions du potentiel A_μ). Bien sûr, faire une transformation de jauge non-commutative modifie tous les ordres en θ , en mélangeant potentiel de jauge et paramètre λ , ce qui explique pourquoi choisir une solution à l'ordre n influence les équations à tous les ordres supérieurs.

Remarquons que les deux opérateurs δ et ∂_μ commutent, nous permettant ainsi d'écrire les solutions homogènes correspondant à des transformations de jauge

$$A_\mu^{(n)} = \partial_\mu \alpha, \quad \Lambda^{(n)} = \delta \alpha \quad (1.95)$$

où α est une expression quelconque d'ordre θ^n , fonction de A_μ .

Voyons maintenant ce qu'il en est des théories non-abéliennes. Bien sûr, tout ce qui précède est toujours valable, à part (1.94) et (1.95) qui doivent être généralisées. C'est ce que nous allons faire ici. Pour cela, il nous suffit de trouver les opérateurs qui remplacent δ et ∂_μ .

À l'ordre θ , l'équation de Seiberg-Witten pour une théorie non-abélienne s'écrit

$$\delta A_\mu^{(1)} - i[\lambda, A_\mu^{(1)}] - \partial_\mu \lambda^{(1)} + i[A_\mu, \lambda^{(1)}] = -\frac{1}{2} \{ \partial \lambda, \partial A_\mu \} \quad (1.96)$$

L'équation homogène associée est donc

$$\Delta A_\mu^{(1)} - D_\mu \lambda^{(1)} = 0 \quad (1.97)$$

où D_μ est la dérivée covariante, et Δ reproduit la transformation de jauge abélienne²³ sur les champs de jauge non-abéliens

$$\Delta = \delta - i[\lambda, \cdot], \quad D_\mu = \partial_\mu - i[A_\mu, \cdot] \quad (1.98)$$

On se convaincra facilement que l'équation homogène à l'ordre n est de même

$$\Delta A_\mu^{(n)} - D_\mu \lambda^{(n)} = 0 \quad (1.99)$$

²³ En effet, en utilisant les transformations de jauge (1.4), on obtient

$$\Delta A_\mu = \partial_\mu \lambda, \quad \Delta F_{\mu\nu} = 0$$

où les commutateurs que l'on a ajoutés par rapport à l'équation homogène abélienne correspondent à ceux du star-commutateur du membre de droite de l'équation de Seiberg-Witten, qui ont une contribution à l'ordre n dans le cas non-abélien. Nous nous sommes alors replacés, comme précédemment, dans la situation où les deux opérateurs commutent. En effet,

$$\begin{aligned}
 [\Delta, D_\mu]\alpha &= [\delta, \partial_\mu]\alpha - i[\delta, [A_\mu, \cdot]]\alpha - i[[\lambda, \cdot], \partial_\mu]\alpha - [[\lambda, \cdot], [A_\mu, \cdot]]\alpha \\
 &= -i\delta[A_\mu, \alpha] + i[A_\mu, \delta\alpha] - i\partial_\mu[\lambda, \alpha] + i\partial_\mu[\lambda, \alpha] \\
 &\quad - [\lambda, [A_\mu, \alpha]] + [A_\mu, [\lambda, \alpha]] \\
 &= 0
 \end{aligned}$$

en utilisant l'expression (1.4) pour la transformation de jauge non-abélienne, et l'identité de Jacobi. Ainsi, nous pouvons exprimer la généralisation non-abélienne de la formule (1.95) : les expressions

$$\mathbb{A}_\mu^{(n)} = D_\mu\alpha, \quad \Lambda^{(n)} = \Delta\alpha \quad (1.100)$$

forment une solution homogène de l'équation de Seiberg-Witten non-abélienne à l'ordre n .

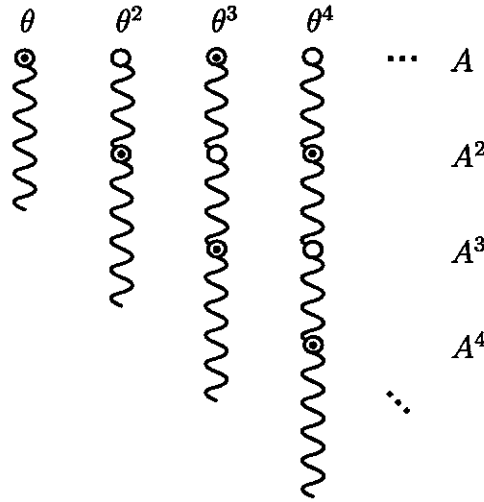


FIG. 1.1 – Les cercles représentent les termes non-abéliens, qui s'étalent sur plusieurs ordres en A , ce qui est signifié par les queues ondulées. Lorsque le groupe est abélien, de nombreux termes disparaissent, ne laissant que les points, dont l'ordre est cette fois bien identifiable.

Je n'ai jusqu'ici parlé que d'un développement en puissance de θ , mais il est aussi intéressant de développer les expressions en puissance de A . Tout d'abord parce que les résultats, suggérés par cette méthode ou par d'autres, semblent indiquer que c'est ce développement qui revêt en réalité un sens plus profond (je fais référence en particulier aux $\star_n$ -produits). Nous en reparlerons bien sûr dans la section suivante. D'autre part, d'un point de vue très algébrique et calculatoire, les solutions

cherchées sont des expressions contenant A , θ et des dérivées. Donc si l'on spécifie, pour un terme du développement, le nombre de θ (disons n , ce qui fait $2n$ indices contravariants) et le nombre de A (disons m , soit m indices covariants), alors le nombre de dérivées est fixé (et vaut naturellement $2n - m$). La structure algébrique du terme est donc partiellement définie. Restent à choisir les contractions des indices et l'action des dérivées.

Signalons enfin que l'ordre en A , s'il est parfaitement défini pour la théorie abélienne, est a priori assez flou pour une théorie non-abélienne. En effet, la courbure F , la dérivée covariante et la transformation de jauge contiennent tous trois une partie quadratique en A . Il faut donc être particulièrement précautionneux pour identifier l'ordre de ces termes, si l'on espère voir apparaître une structure. La figure 1.1 est une représentation schématique de ce double développement, et de sa complexité dans le cas non-abélien par rapport au cas abélien.

1.5.3 Equivalence de star-produits

Sur un espace non-commutatif, caractérisé par un star-produit $\star$, on peut introduire la théorie de jauge par son action sur les champs de matière ψ ,

$$\psi \longrightarrow g \star \psi$$

sous la transformation induite par l'élément g . Mais si f est une fonction sur cet espace, f est donc invariante, et le produit $f \star \psi$ ne se comporte plus comme un champ de matière, puisqu'en général

$$g \star f \star \psi \neq f \star g \star \psi$$

Pour cela, il faut que la fonction se transforme de manière covariante, et nous appellerons donc "covarianteur" l'opérateur $\mathcal{D}$ tel que $\mathcal{D}f$ soit covariant :

$$\mathcal{D}f \longrightarrow g \star \mathcal{D}f \star g^{-1}$$

On peut écrire $\mathcal{D}f = f + f_A$, où f_A est introduit pour compenser l'invariance de f par transformation de jauge. En allant plus loin, on peut introduire un potentiel de jauge abstrait $\mathcal{A}$, défini par

$$\mathcal{D} = \text{id} + \mathcal{A} \tag{1.101}$$

ce qui revient à écrire $\mathcal{D}f = f + \mathcal{A}(f)$. On déduit la transformation de ce potentiel de jauge de la covariance de $\mathcal{D}$:

$$\begin{aligned} \mathcal{A} &\longrightarrow g \star \mathcal{D} \star g^{-1} - \text{id} \\ &= g \star \text{id} \star g^{-1} + g \star \mathcal{A} \star g^{-1} - \text{id} \\ &= g \star \mathcal{A} \star g^{-1} + g \star (\text{id} \star g^{-1} - g^{-1} \star \text{id}) \\ &= g \star \mathcal{A} \star g^{-1} + g \star \mathbf{d}_\star g^{-1} \end{aligned} \tag{1.102}$$

en remplaçant id par $g \star g^{-1} \star \text{id}$ entre la deuxième et la troisième ligne. L'opérateur de cobord $\mathbf{d}_\star$ a été défini en (1.80). D'autre part, $\mathcal{A}$ est un élément du complexe de

Hochschild de l'algèbre de fonctions $\mathcal{A}_x$ sur l'espace non-commutatif (plus simplement $\mathcal{A}$ est un opérateur linéaire sur cette algèbre). L'opérateur $\star$, qui prolonge le star-produit sur ce complexe, est en fait le cup-produit

$$(C_1 \star C_2)(f_1, \dots, f_{p_1+p_2}) = C_1(f_1, \dots, f_{p_1}) \star C_2(f_{p_1+1}, \dots, f_{p_1+p_2}) \quad (1.103)$$

où C_i est un opérateur p_i -linéaire.

Sous la forme infinitésimale (en remplaçant g par $1 + i\lambda$), cette transformation devient :

$$\delta \mathcal{A} = -i \mathbf{d}_\star \lambda + i \lambda \star \mathcal{A} - i \mathcal{A} \star \lambda \quad (1.104)$$

La courbure $\mathcal{F}$ est l'opérateur bilinéaire *alterné* défini par la formule usuelle,

$$\mathcal{F} = \mathbf{d}_\star \mathcal{A} + \mathcal{A} \wedge \mathcal{A} \quad (1.105)$$

$\mathcal{F}$ appartient à la cohomologie de Chevalley de $\mathcal{A}_x$, qui est le complexe formé des applications multilinéaires alternées $\text{Hom}(\mathcal{A}_x^{\wedge n}, \mathcal{A}_x)$. En toute rigueur, $\mathcal{A}$ appartenait lui aussi à ce complexe, plutôt qu'au complexe de Hochschild, mais les deux sont confondus jusqu'à l'ordre 1. L'opérateur $\mathbf{d}_\star$ s'étend naturellement sur ce complexe, tout comme le cup-produit, que nous avons ici renommé $\wedge$. Ainsi, en appliquant la courbure aux deux fonctions f et g , en se souvenant que $\mathcal{A} = \mathcal{D} - \text{id}$,

$$\begin{aligned} \mathcal{F}(f \wedge g) &= -[\mathcal{A}, \star]_{\mathcal{G}}(f \wedge g) + \mathcal{A} \wedge \mathcal{A}(f \wedge g) \\ &= -\mathcal{A}([f \star g]) + [\mathcal{A}(f) \star g] + [f \star \mathcal{A}(g)] + [\mathcal{A}(f) \star \mathcal{A}(g)] \\ &= -\mathcal{D}([f \star g]) + [f \star g] + [\mathcal{D}(f) \star g] - [f \star g] + [f \star \mathcal{D}(g)] - [f \star g] \\ &\quad + [\mathcal{D}(f) \star \mathcal{D}(g)] - [\mathcal{D}(f) \star g] - [f \star \mathcal{D}(g)] + [f \star g] \\ &= [\mathcal{D}(f) \star \mathcal{D}(g)] - \mathcal{D}([f \star g]) \end{aligned} \quad (1.106)$$

Dans un système de coordonnées x^i , nous pouvons maintenant introduire les coordonnées covariantes

$$X^i = \mathcal{D}x^i = x^i + \mathcal{A}(x^i) = x^i + \theta^{ij} \hat{A}_j \quad (1.107)$$

où, pour la dernière égalité, on suppose que le star-produit est associé à une structure de Poisson θ constante et inversible. Sous cette forme plus concrète, on peut recalculer la transformation de jauge infinitésimale de $\hat{A}$:

$$\begin{aligned} \theta^{ij} \delta \hat{A}_j &= i [\lambda \star x^i + \theta^{ij} \hat{A}_j] \\ &= \theta^{ij} \partial_j \lambda + i \theta^{ij} [\lambda \star \hat{A}_j] \\ \delta \hat{A}_j &= \partial_j \lambda + i [\lambda \star \hat{A}_j] \end{aligned}$$

Ainsi, connaître le covarianteur $\mathcal{D}$ est équivalent à connaître le potentiel de jauge $\hat{A}$, ou plus généralement $\mathcal{A}$. Pour résoudre le problème qui nous préoccupe ici, à savoir calculer l'expression de $\hat{A}$ en fonction du paramètre θ , il nous suffit donc de trouver l'expression du covarianteur. D'autre part, on peut remarquer la similarité entre l'expression (1.66) pour les déformations des star-produits et la relation $\mathcal{D}f = f + f_{\mathcal{A}}$ pour le covarianteur. Cela suggère que si l'on connaît l'opérateur d'entrelacement

entre deux star-produits, alors celui-ci pourra peut-être aussi jouer le rôle de covarianteur. Bien sûr cela reste très vague tant qu'il n'y a pas de théorie de jauge quelque part, mais nous allons le préciser. Avant cela, remarquons que puisque le covarianteur $\mathcal{D}$ est inversible, il définit un star-produit $\star'$ par

$$f \star' g \equiv \mathcal{D}^{-1}(\mathcal{D}f \star \mathcal{D}g)$$

En effet, l'opérateur $\star'$ ainsi défini est bien associatif :

$$\begin{aligned} (f \star' g) \star' h &= \mathcal{D}^{-1}[\mathcal{D}\mathcal{D}^{-1}(\mathcal{D}f \star \mathcal{D}g) \star \mathcal{D}h] \\ &= \mathcal{D}^{-1}[(\mathcal{D}f \star \mathcal{D}g) \star \mathcal{D}h] \\ &= \mathcal{D}^{-1}[\mathcal{D}f \star (\mathcal{D}g \star \mathcal{D}h)] \\ &= f \star' (g \star' h) \end{aligned}$$

Notre remarque, bien qu'imprécise, fonctionne donc au moins dans un sens : si $\mathcal{D}$ est un covarianteur, alors il est un opérateur d'entrelacement entre deux star-produits. Nous aimerions ici plutôt faire le contraire. Précisons aussi que ces deux star-produits ne sont pas nécessairement équivalents au sens de Kontsevich (la correction f_A n'est pas nécessairement d'ordre $\hbar$).

Nous avons vu qu'une structure de Poisson θ peut être quantifiée en star-produit. Ainsi, si on se place dans le cas où B est inversible et $\theta = B^{-1}$, les deux structures de Poisson²⁴ suivantes θ et θ' définissent deux star-produits $\star$ et $\star'$, avec chaque fois une structure interpolante θ_t et $\star_t$:

$$\begin{array}{lll} \theta & = & B^{-1} \longrightarrow \star \\ \theta_t & \equiv & (B + tF)^{-1} \longrightarrow \star_t \\ \theta' & \equiv & (B + F)^{-1} \longrightarrow \star' \end{array}$$

où t est un paramètre variant entre 0 et 1, et bien sûr $\theta_0 = \theta$ et $\theta_1 = \theta'$. On peut calculer la relation entre θ et θ_t ou θ' , en éliminant B

$$\begin{aligned} \theta_t &= \frac{1}{B + tF} = \frac{1}{\theta^{-1} + tF} = \theta \frac{1}{1 + tF\theta} \\ \theta' &= \theta \frac{1}{1 + F\theta} \end{aligned} \tag{1.108}$$

Les deux structures introduites sont chacune une déformation de la première, et $\star$ et $\star'$ sont reliées par un opérateur d'entrelacement $\mathcal{D}_a$, que nous appellerons flux quantique,

$$\mathcal{D}_a(f \star' g) = \mathcal{D}_a f \star \mathcal{D}_a g \tag{1.109}$$

La stratégie sera alors de calculer $\mathcal{D}_a$, en déduire le potentiel de jauge associé, et vérifier qu'il s'agit bien du potentiel de jauge non-commutatif, comme nous l'avons anticipé.

²⁴ On rappelle que dans ce cas, θ définit une structure de Poisson si et seulement si $dB = 0$. En particulier, dans le cas abélien, $F = dA$, et donc $d(B + F) = dB = 0$. $\theta' \equiv (B + F)^{-1}$ définit donc aussi une structure de Poisson.

L'esprit de la méthode consiste donc à construire un star-produit qui encode la théorie de jauge, puis à l'extraire grâce au covarianteur. En effet, alors que $\star$ ne "connaît" rien de la structure de jauge sur la variété, $\star_t$, comme $\star'$, contient encodé un certain nombre d'informations, à l'image de sa variation par rapport à t , qui fait apparaître, comme nous le verrons, la courbure du champ de jauge. Ainsi, en reliant par le flux quantique $\mathcal{D}_a$ les deux star-produits, on pourra extraire cette information et la rendre explicite, sous la forme du potentiel de jauge non-commutatif.

1.6 Le calcul récursif

Je vais maintenant présenter plus en détail les résultats que j'ai obtenus au cours de ma thèse, en utilisant la méthode présentée en 1.5.2. Nous nous placerons, dans cette section comme dans la suivante, dans le cas d'une théorie abélienne. Nous évoquerons le cas non-abélien à la section 1.8.

Nous commencerons par les termes de la première diagonale de la figure 1.1, c'est-à-dire les termes d'ordre (A^n, θ^n) que nous appellerons "semi-classiques", pour reprendre la terminologie de [30]. Nous allons voir qu'il existe une solution qui s'exprime récursivement, aussi bien sur le potentiel de jauge que la courbure. J'expliquerai ensuite comment on pourrait généraliser cette formulation récursive à l'ensemble des termes, pour peu que l'on puisse réguler les ambiguïtés qui apparaissent.

1.6.1 Résoudre l'équation récursive

Les termes d'ordre (A^n, θ^n) forment un sous-ensemble complet, au sens où ils suffisent à écrire les ordres correspondants de l'équation de Seiberg-Witten. En voici les premiers ordres :

$$\delta A_\mu^{(1,1)} - \partial_\mu \lambda^{(1,1)} = -\partial \lambda \theta \partial A_\mu \quad (1.110)$$

$$\delta A_\mu^{(2,2)} - \partial_\mu \lambda^{(2,2)} = -\partial \lambda^{(1,1)} \theta \partial A_\mu - \partial \lambda \theta \partial A_\mu^{(1,1)} \quad (1.111)$$

$$\delta A_\mu^{(3,3)} - \partial_\mu \lambda^{(3,3)} = -\partial \lambda^{(2,2)} \theta \partial A_\mu - \partial \lambda^{(1,1)} \theta \partial A_\mu^{(1,1)} - \partial \lambda \theta \partial A_\mu^{(2,2)} \quad (1.112)$$

Nous allons résoudre l'ordre 1, c'est-à-dire trouver des solutions pour $A_\mu^{(1,1)}$ et $\lambda^{(1,1)}$. Dans la première équation, le terme de droite contient A_μ . Il est donc impossible qu'il vienne de $\partial_\mu \lambda^{(1,1)}$ (l'indice μ est porté par une dérivée), et $A_\mu^{(1,1)}$ doit donc contenir $A \theta \partial A_\mu$. Mais la variation de ce terme va aussi créer $A \theta \partial \partial_\mu \lambda$, qui doit à son tour être compensé par un terme de $\lambda^{(1,1)}$. Celui-ci contient donc autant de $A \theta \partial \lambda$. Cela donne donc jusqu'ici

$$\begin{aligned} \lambda^{(1,1)} &= \alpha A \theta \partial \lambda + \dots \\ A_\mu^{(1,1)} &= \alpha A \theta \partial A_\mu + \dots \end{aligned}$$

et en remplaçant dans l'équation (1.110)

$$\begin{aligned}
0 &= \delta A_\mu^{(1,1)} - \partial_\mu \lambda^{(1,1)} + \partial \lambda \theta \partial A_\mu \\
&= \alpha \partial \lambda \theta \partial A_\mu + \alpha A \theta \partial \partial_\mu \lambda - \alpha \partial_\mu A \theta \partial \lambda - \alpha A \theta \partial \partial_\mu \lambda + \partial \lambda \theta \partial A_\mu + \dots \\
&= (1 + \alpha) \partial \lambda \theta \partial A_\mu + \alpha \partial \lambda \theta \partial_\mu A + \dots
\end{aligned}$$

Le terme $\partial \lambda \theta \partial_\mu A$ vient de $\lambda^{(1,1)}$. Il faut donc l'annuler par un terme de $A_\mu^{(1,1)}$, soit $A \theta \partial_\mu A$. Mais la variation de ce terme va redonner du $A \theta \partial \partial_\mu \lambda$, dont on s'était déjà débarrassé. Pour ne pas tourner en rond, il faut donc remplacer $\partial_\mu A$ par un objet invariant de jauge : F_μ . Cela donne alors

$$A_\mu^{(1,1)} = \alpha A \theta \partial A_\mu + \alpha A \theta F_\mu$$

et l'équation devient

$$(1 + 2\alpha) \partial \lambda \theta \partial A_\mu = 0$$

soit $(1 + 2\alpha) = 0$, c'est-à-dire $\alpha = -\frac{1}{2}$.

On a ainsi trouvé la solution à l'ordre 1.

1.6.2 Expression récursive des termes semi-classiques

Voilà les solutions que je propose aux premiers ordres. Les ordres (2, 2) et (3, 3) ont été calculés de la même manière que le premier ordre que nous venons de faire. Comme nous l'avons vu, la difficulté vient surtout du fait de trouver une solution "esthétique", présentant une structure, parmi toutes les solutions possibles.

$$\begin{cases} \lambda^{(1,1)} &= -\frac{1}{2} (A \theta \partial) \lambda \\ A_\mu^{(1,1)} &= -\frac{1}{2} (A \theta \partial) A_\mu - \frac{1}{2} A \theta F_\mu \end{cases} \quad (1.113)$$

$$\begin{cases} \lambda^{(2,2)} &= \frac{1}{6} [(A \theta \partial) (A \theta \partial) \lambda + A \theta F \theta \partial \lambda] \\ A_\mu^{(2,2)} &= \frac{1}{6} [(A \theta \partial) (A \theta \partial) A_\mu + A \theta F \theta \partial A_\mu + 2(A \theta \partial) (A \theta F_\mu) + 2A \theta F \theta F_\mu] \end{cases} \quad (1.114)$$

$$\begin{cases} \lambda^{(3,3)} &= -\frac{1}{24} [(A \theta \partial) (A \theta \partial) (A \theta \partial) \lambda + (A \theta F \theta \partial) (A \theta \partial) \lambda \\ &\quad + 2(A \theta \partial) (A \theta F \theta \partial) \lambda + 2A \theta F \theta F \theta \partial \lambda] \\ A_\mu^{(3,3)} &= -\frac{1}{24} [(A \theta \partial) (A \theta \partial) (A \theta \partial) A_\mu + (A \theta F \theta \partial) (A \theta \partial) A_\mu \\ &\quad + 2(A \theta \partial) (A \theta F \theta \partial) A_\mu + 2A \theta F \theta F \theta \partial A_\mu \\ &\quad + 3(A \theta \partial) (A \theta \partial) (A \theta F_\mu) + 3(A \theta F \theta \partial) (A \theta F_\mu) \\ &\quad + 6(A \theta \partial) (A \theta F \theta F_\mu) + 6A \theta F \theta F \theta F_\mu] \end{cases} \quad (1.115)$$

Inversement, une fois l'expression écrite, on peut vérifier explicitement que ces expressions sont solutions des équations de Seiberg-Witten. En particulier, une fois la structure identifiée, on peut proposer une expression pour tous les ordres suivants.

De plus, ces solutions, que j'ai proposée dans [36], s'écrivent récursivement sous la forme

$$\begin{cases} \lambda^{(n,n)} &= -\frac{1}{n+1} A^{(n-1,n-1)} \theta \partial \lambda \\ A_\mu^{(n,n)} &= -\frac{1}{n+1} A^{(n-1,n-1)} \theta (\partial A_\mu + n F_\mu) \end{cases} \quad (1.116)$$

où $A^{(n-1,n-1)}$ agit comme un opérateur sur le reste du terme, c'est-à-dire que ses dérivées explicites (par opposition à celles contenues dans F) agissent aussi sur ce qui est à droite. On peut se faire une idée concrète de ce que cela signifie sur les premiers ordres déjà écrits. Par exemple, c'est la raison pour laquelle à l'ordre 2

$$A^{(1,1)} \theta \partial \lambda \longrightarrow (A \theta \partial)(A \theta \partial) \lambda$$

Pour vérifier les ordres $n = 4, 5, 6$, j'ai utilisé un logiciel que j'ai écrit spécialement dans ce but. Il calcule et simplifie l'équation à l'ordre considéré, en remplaçant tous les termes par leur valeur. Cela évite de longues et fastidieuses manipulations.

Cela dit, il serait encore plus simple de trouver une preuve générale. Nous verrons dans la prochaine section que ces expressions peuvent s'exprimer encore différemment, et nous aurons en même temps la justification voulue. La présentation des résultats sous cette forme me semble cependant intéressante. D'une part parce que c'est une forme explicite, qui peut être réutilisée dans n'importe quel calcul, contrairement à d'autres expressions, certes plus concises, mais en général aussi moins concrètes. Et d'autre part parce qu'on voit clairement apparaître la structure "esthétique" qui m'a conduit à les choisir parmi la multitude de solutions, et qui n'a pas livré tous ses secrets. Nous verrons en particulier comment tous cela s'écrit dans le cas d'une pure jauge quand nous nous attaquerons au cas non-abélien.

Avant tout cela, écrivons la courbure aux premiers ordres

$$F_{\mu\nu}^{(1,1)} = -(A \theta \partial) F_{\mu\nu} - F_\mu \theta F_\nu \quad (1.117)$$

$$F_{\mu\nu}^{(2,2)} = -(A^{(1,1)} \theta \partial) F_{\mu\nu} - F_\mu^{(1,1)} \theta F_\nu \quad (1.118)$$

$$F_{\mu\nu}^{(3,3)} = -(A^{(2,2)} \theta \partial) F_{\mu\nu} - F_\mu^{(2,2)} \theta F_\nu \quad (1.119)$$

Le même genre de structure récursive apparaît donc sur F . Signalons que cette structure est vraiment liée au choix particulier de solution qui est fait à chaque ordre, et profitons-en pour comparer ces résultats aux formules que l'on peut trouver dans d'autres références [10, 20, 31]. Les ordres 1 pour A et λ sont partout identiques, et les différences ne commencent qu'à l'ordre 2 (c'est-à-dire (A^2, θ^2)). Il y a là en effet trois solutions, différant l'une de l'autre par une transformation de jauge. Si l'on définit

$$\alpha = A \theta (\partial A) \theta A \quad (1.120)$$

où les contractions sont faites dans l'ordre naturel de lecture, alors on obtient une solution homogène par les formules (1.95). Ma solution diffère de celle [31] de la transformation de jauge associée à $\frac{1}{6}\alpha$, et de celle de [20, 26, 27] par $\frac{5}{12}\alpha$. Bien entendu, dans tous les cas, la courbure est identique. Cela signifie en particulier que les deux premières relations (1.117) et (1.118) sont déjà valables pour tous ces articles. Par contre, la relation à l'ordre 3 (1.119), continuation naturelle des deux premières, n'est vraie qu'avec les formules proposées ici.

1.6.3 Expression compacte des termes semi-classiques

Essayons maintenant de trouver une expression plus compacte pour ces formules. Il est clair que l'un des objets qui semble indispensable pour un tel programme est le bloc $A\theta\partial$, autrement dit $A_i\theta^{ij}\partial_j$. Maintenant, pour espérer générer tous les termes, reste à trouver un moyen de transformer ce bloc en $A\theta F\theta\partial$, et le reste devrait suivre. En anticipant un peu sur la suite, introduisons naïvement une ou deux notations bien choisies :

- Remplaçons θ par θ_t , fonction du paramètre formel t . Décidons aussi que l'on retrouve la structure initiale en $t = 0$, c'est-à-dire $\theta_0 = \theta$.
- On définit la variation de θ_t par rapport à t par

$$\partial_t \theta_t = -\theta_t F \theta_t \quad \text{i.e.} \quad \partial_t \theta_t^{ij} = -\theta_t^{ik} F_{kl} \theta_t^{lj} \quad (1.121)$$

- Enfin, on notera a_{θ_t} l'opérateur différentiel d'ordre 1 fondamental, ou plus exactement

$$a_{\theta_t} = -A\theta_t\partial \quad (1.122)$$

Voyons si grâce à tout cela, on peut exprimer les formules (1.113) et suivantes plus simplement. Nous nous intéresserons tout d'abord à λ , ce qui fera en même temps la moitié du travail pour A . En remplaçant dans (1.113), (1.114) et (1.115)

$$\lambda^{(1,1)} = \frac{1}{2} a_{\theta_t} \lambda|_{t=0} = \frac{1}{2} (a_{\theta_t} + \partial_t) \lambda|_{t=0} \quad (1.123)$$

$$\lambda^{(2,2)} = \frac{1}{6} (a_{\theta_t} + \partial_t)^2 \lambda|_{t=0} \quad (1.124)$$

$$\lambda^{(3,3)} = \frac{1}{24} (a_{\theta_t} + \partial_t)^3 \lambda|_{t=0} \quad (1.125)$$

où l'on notera bien sûr que $\partial_t \lambda = 0$. De plus, c'est sur la troisième formule que commencent à apparaître des coefficients non triviaux. On remarquera que cette expression compacte génère exactement le bon nombre de termes de chaque sorte. Cela conduit à la formule générale

$$\lambda^{(n,n)} = \frac{1}{(n+1)!} (a_{\theta_t} + \partial_t)^n \lambda|_{t=0} \quad (1.126)$$

Pour trouver l'expression de A , réécrivons cette formule et comparons-la à (1.116).

$$\begin{aligned} \lambda^{(n,n)} &= \frac{(a_{\theta_t} + \partial_t)^{n-1}}{(n+1)!} a_{\theta_t} \lambda|_{t=0} \\ &= -\frac{1}{n+1} \left[\frac{(a_{\theta_t} + \partial_t)^{n-1}}{n!} (A\theta_t) \right]_{t=0} \partial \lambda \\ &= -\frac{1}{n+1} A^{(n-1,n-1)} \theta \partial \lambda \end{aligned}$$

ce qui conduit donc à l'hypothèse

$$A^{(n-1,n-1)} \theta = \frac{(a_{\theta_t} + \partial_t)^{n-1}}{n!} (A\theta_t) \Big|_{t=0} \quad (1.127)$$

Démontrons maintenant toutes ces formules. On a déjà vérifié λ aux premiers ordres. Reste à le faire pour A . Puis nous vérifierons la récurrence (1.116). En utilisant (1.127), on trouve bien sûr $A^{(0)} = A$, mais aussi

$$\begin{aligned} A_i^{(1,1)} \theta^{ij} &= \frac{1}{2} (a_{\theta_t} + \partial_t) (A_i \theta_t^{ij})|_{t=0} \\ &= -\frac{1}{2} [(A\theta\partial) A_i \theta^{ij} + A_i \theta^{ik} F_{kl} \theta^{lj}] \\ &= -\frac{1}{2} [(A\theta\partial) A_i + A_i \theta^{lk} F_{ki}] \theta^{ij} \end{aligned}$$

où l'on a renommé les indices du terme en F à la troisième ligne. Cette expression est bien celle de $A^{(1,1)}$ donnée en (1.113). De même à l'ordre suivant

$$\begin{aligned} A_i^{(2,2)} \theta^{ij} &= \frac{1}{6} (a_{\theta_t} + \partial_t)^2 (A_i \theta_t^{ij})|_{t=0} \\ &= -\frac{1}{6} (a_{\theta_t} + \partial_t) [(A\theta_t\partial) (A_i \theta_t^{ij}) + A_i \theta_t^{lk} F_{ki} \theta_t^{ij}]|_{t=0} \\ &= \frac{1}{6} [(A\theta\partial)(A\theta\partial) A_i + (A\theta\partial)(A\theta F_i) + (A\theta F\theta\partial) A_i \\ &\quad + (A\theta\partial)(A_i \theta^{lk} F_{ki}) + (A\theta F\theta F_i) + A_i \theta^{lk} F_{kp} \theta^{pq} F_{qi}] \theta^{ij} \end{aligned}$$

qui, encore une fois, est bien la valeur attendue. Plutôt qu'à l'ordre 3, attaquons-nous maintenant à la récurrence, dont le calcul ressemble fort au précédent. Armé de l'expression (1.127) et de la formule de récurrence (1.116), il nous faut donc démontrer que (1.127) est aussi valable à l'ordre n . Pour cela, nous allons utiliser un résultat que nous avons vu dans les deux calculs précédents

$$A\theta_t F_i \theta_t^{ij} = -\partial_t (A_i \theta_t^{ij}) \quad (1.128)$$

Cela donne donc pour la récurrence

$$\begin{aligned} A_i^{(n,n)} \theta^{ij} &= -\frac{1}{n+1} A^{(n-1,n-1)} \theta (\partial A_i + n F_i) \theta^{ij} \\ &= -\frac{1}{n+1} \frac{(a_{\theta_t} + \partial_t)^{n-1}}{n!} (A\theta_t) \Big|_{t=0} (\partial A_i + n F_i) \theta^{ij} \\ &= -\frac{(a_{\theta_t} + \partial_t)^{n-1}}{(n+1)!} [A\theta_t \partial (A_i \theta_t^{ij}) + A\theta_t F_i \theta_t^{ij}] \Big|_{t=0} \\ &= \frac{(a_{\theta_t} + \partial_t)^{n-1}}{(n+1)!} (a_{\theta_t} + \partial_t) (A_i \theta_t^{ij}) \Big|_{t=0} \\ &= \frac{(a_{\theta_t} + \partial_t)^n}{(n+1)!} (A_i \theta_t^{ij}) \Big|_{t=0} \end{aligned}$$

On remarque que la prescription " $A^{(n-1,n-1)}$ agit comme un opérateur" se traduit ici naturellement par le fait que les $(n-1)$ premiers facteurs $(a_{\theta_t} + \partial_t)$ agissent tous sur le dernier. De plus, on pourra se convaincre que cette expression reproduit effectivement le facteur n devant F_i , comme elle l'avait fait à l'ordre 2, au passage de la deuxième à la troisième ligne.

Maintenant que l'on dispose des expressions des $A^{(n,n)}$ et $\lambda^{(n,n)}$, on peut sommer tous ces termes, pour avoir une version partielle du potentiel et du paramètre de jauge non-commutatif. Comme je l'ai déjà signalé, nous appellerons "semi-classiques" ces grandeurs, en suivant la dénomination de [30], et nous les noterons A^{sc} et λ^{sc}

$$\lambda^{sc} \equiv \sum_{n=0}^{\infty} \lambda^{(n,n)} = \sum_{n=0}^{\infty} \frac{(a_{\theta_t} + \partial_t)^n}{(n+1)!} \lambda \Big|_{t=0} \quad (1.129)$$

$$A^{sc} \theta \equiv \sum_{n=0}^{\infty} A^{(n,n)} \theta = \sum_{n=0}^{\infty} \frac{(a_{\theta_t} + \partial_t)^n}{(n+1)!} (A \theta_t) \Big|_{t=0} \quad (1.130)$$

$$(1.131)$$

Avec encore un peu plus d'anticipation, on peut exprimer la formule de A^{sc} plus simplement. On reconnaît en effet le développement de l'exponentielle. Mais autant il y a un décalage d'indice pour λ , qui ne nous ferait rien gagner à l'exprimer ainsi, autant dans le cas de A , la présence du θ_t , jusque-là intrigante, va nous aider. En effet, remarquons que

$$A_i \theta_t^{ij} = A_i \theta_t^{ik} \partial_k x^j = -a_{\theta_t} x^j$$

Cela rajoute le facteur manquant, puisque d'autre part $\partial_t x^j = 0$. Reste à ajouter le premier terme de la série, pour donner

$$\begin{aligned} A_i^{sc} \theta^{ij} &= x^j - x^j - \sum_{n=1}^{\infty} \frac{(a_{\theta_t} + \partial_t)^n}{n!} x^j \Big|_{t=0} \\ &= x^j - e^{(a_{\theta_t} + \partial_t)} \Big|_{t=0} x^j \end{aligned}$$

Nous nous arrêterons à cette formule, que nous retrouvons plus loin, avec d'autres développements. Concluons en l'écrivant sous une forme encore plus sympathique, avant de reprendre notre récurrence où nous l'avions laissée.

$$e^{(a_{\theta_t} + \partial_t)} \Big|_{t=0} x^i = x^i + \theta^{ij} A_j^{sc} \quad (1.132)$$

1.6.4 Une récursivité en puissance de A

Maintenant que nous avons trouvé les termes de la première diagonale, intéressons-nous au reste du triangle (fig. 1.1). Puisque nous sommes toujours dans le cas abélien, un terme sur deux est nul, ne laissant que les termes dénotés par des points sur la figure. Bien sûr, nous pourrions continuer dans la direction entamée à la section précédente, et essayer de généraliser l'opérateur a_{θ_t} , et c'est ce que nous verrons par la suite. Mais pour l'instant, soyons fidèle à la méthode, et regardons les expressions explicites ordre par ordre.

A l'ordre $(1, 3)$, c'est-à-dire (A, θ^3) , l'équation s'écrit

$$\delta A_{\mu}^{(1,3)} - \partial_{\mu} \lambda^{(1,3)} = \frac{1}{24} \partial \partial \partial \lambda \theta \theta \theta \partial \partial \partial A_{\mu} \quad (1.133)$$

Pour la résoudre, on raisonne comme pour les ordres précédents : le terme de droite ne peut venir que de la variation de $A_{\mu}^{(1,3)}$ (puisque sinon, l'indice μ serait porté par

une dérivée). On commence donc avec le terme $\partial\partial A\theta\theta\theta\partial\partial\partial A_\mu$, et donc nécessairement le terme correspondant dans $\lambda^{(1,3)}$ (pour éliminer la variation du A_μ final, en $\partial_\mu\lambda$). Et on trouve alors le terme manquant. Ici, cela donne

$$\lambda^{(1,3)} = \frac{1}{48} \partial\partial A \theta\theta\theta \partial\partial\partial\lambda \quad (1.134)$$

$$A_\mu^{(1,3)} = \frac{1}{48} \partial\partial A \theta\theta\theta \partial\partial\partial A_\mu + \frac{1}{48} \partial\partial_i A \theta\theta\theta \partial\partial F_{i\mu} \quad (1.135)$$

$$F_{\mu\nu}^{(1,3)} = \frac{1}{24} \partial\partial A \theta\theta\theta \partial\partial\partial F_{\mu\nu} + \frac{1}{24} \partial\partial_i F_{\mu j} \theta\theta\theta \partial\partial_j F_{i\nu} \quad (1.136)$$

où les indices notés i (∂_i et $F_{i\mu}$) sont contractés ensemble (c'est-à-dire sur le même θ)²⁵. Remarquons que cette contraction n'est pas la seule solution. En fait, le premier indice de F peut être contracté sur n'importe lequel des trois indices $\partial\partial A$. Cependant, cette solution s'avèrera "meilleure" que les autres, parce qu'elle nous permettra de retrouver (à un détail près que nous expliquerons) la récurrence commencée sur la première diagonale.

Ce calcul était donc assez simple, et il en est de même pour toute la première ligne (ordre 1 en A). On peut remarquer aussi que $A_\mu^{(1,3)}$ s'exprime en fonction de $(\partial A_\mu + F_{\mu})$, et cela reste vrai pour les autres termes de la ligne. Alors que sur la première diagonale, $A_\mu^{(n,n)}$ s'exprimait en fonction de $(\partial A_\mu + nF_{\mu})$. C'est là un bon argument pour regrouper les termes par lignes, en puissance de A , et nous utiliserons la notation $(1, \infty)$ pour désigner la somme de tous les termes d'ordre A

$$\begin{aligned} A_\mu^{(1,\infty)} &= A_\mu^{(1,1)} + A_\mu^{(1,3)} + \dots \\ &= \frac{1}{2} \left[-A\theta(\partial A_\mu + F_{\mu}) + \frac{1}{24} \partial\partial_i A \theta\theta\theta \partial\partial(\partial_i A_\mu + F_{i\mu}) + \dots \right] \end{aligned} \quad (1.137)$$

Cette expression a déjà été calculée par ailleurs. En particulier, on trouve une formule équivalente, à une transformation de jauge près, par la méthode expliquée en 1.5.1. Ce résultat est connu sous la forme d'un produit nommé $\star_2$ [20, 26]. Cette notation est trompeuse, puisque ce produit n'est pas associatif, et n'est donc pas un star-produit, mais elle consacrée par l'usage :

$$A_\mu^{(1,\infty)} = -\frac{1}{2} \theta^{ij} A_i \star_2 (\partial_j A_\mu + F_{j\mu}) \quad (1.138)$$

où $\star_2$ est défini par

$$f(x) \star_2 g(x) = \frac{\sin(\partial_1\theta\partial_2/2)}{\partial_1\theta\partial_2/2} f(x_1)g(x_2) \Big|_{x_i=x} \quad (1.139)$$

Comme je l'ai signalé, on peut voir que dans cette formule, à l'ordre $(1, 3)$, l'indice libre du F est contracté avec le A du $\partial\partial A$, contrairement à (1.135).

Puisque les expressions semblent vouloir se regrouper ligne par ligne, tentons d'élargir la récurrence dans ce sens. Rappelons-nous que la formule récursive était la même que celle de l'ordre $(1, 1)$, où l'on avait simplement remplacé

²⁵ A contrario, les marques $\cdot$ ne font que symboliser la position de l'indice. C'est important puisque F est antisymétrique. Deux telles marques (et a fortiori davantage) dans un même terme ne sont donc pas nécessairement contractées.

- le A de l'expression $(A\theta\partial)$ par $A^{(n-1,n-1)}$
- le $\frac{\partial A_\mu + F_\mu}{2}$ à la fin du terme $A_\mu^{(1,1)}$ par $\frac{\partial A_\mu + nF_\mu}{n+1}$

Faisons donc de même pour la ligne entière, ce qui donne pour λ :

$$\begin{aligned}
 \lambda^{(1,\infty)} &= \frac{1}{2} \left[-A\theta\partial\lambda + \frac{1}{24} \partial\partial A \theta\theta\theta \partial\partial\partial\lambda + \dots \right] \\
 \lambda^{(2,\infty)} &= \frac{1}{3} \left[-A^{(1,\infty)}\theta\partial\lambda + \frac{1}{24} \partial\partial A^{(1,\infty)} \theta\theta\theta \partial\partial\partial\lambda + \dots \right] \\
 &= \underbrace{-\frac{1}{3} A^{(1,1)}\theta\partial\lambda}_{\lambda^{(2,2)}} - \underbrace{\frac{1}{3} A^{(1,3)}\theta\partial\lambda + \frac{1}{72} \partial\partial A^{(1,1)} \theta\theta\theta \partial\partial\partial\lambda + \dots}_{\lambda^{(2,4)}} \quad (1.140)
 \end{aligned}$$

Comme on le souhaite, cette prescription redonne pour $\lambda^{(2,2)}$ la formule de récurrence déjà trouvée. On remarque aussi que $\lambda^{(2,4)}$ apparaît comme une somme de deux termes, que nous allons regarder de plus près. En particulier, souvenons-nous que $A^{(1,3)}$ et $A^{(1,1)}$ doivent dans cette formule agir comme des opérateurs. Si le premier ne pose pas de problème, parce qu'on a déjà rencontré ce genre de termes, le deuxième est nouveau, et source d'ambiguïtés. En effet, si l'on écrit mieux et complètement ce morceau

$$\partial\partial A^{(1,1)} \theta\theta\theta \partial\partial\partial\lambda \propto (A\theta\overrightarrow{\partial})A(\overleftarrow{\partial}\partial\theta\theta\theta\overrightarrow{\partial}\partial\partial\lambda) \quad (1.141)$$

on voit que le $\overrightarrow{\partial}$ inclus dans $A^{(1,1)}$ doit agir sur tout ce qui est à droite, tandis que $\overleftarrow{\partial}\partial$ agit sur ce qui est à gauche. On pourrait dire, bien sûr, que c'est se compliquer la vie d'écrire les choses comme cela, et qu'on pourrait par défaut les écrire plus simplement. C'est vrai, mais en l'occurrence "par défaut", ça ne marche pas : sans cette subtilité, les expressions générées ne sont pas solutions de l'équation. Autrement dit, c'est comme ça, malheureusement, et c'est à nous de trier un peu les choses. D'autre part, comme nous l'avons déjà fait pour les termes semi-classiques, nous verrons plus loin une expression compacte pour ces formules. Mais encore une fois, je pense qu'il est utile de les manipuler un peu, parce qu'elles sont explicites et que dans le pire des cas, elles sont la seule chose à laquelle on peut se raccrocher. Bien sûr, il est fort peu probable qu'un jour quelqu'un est besoin de l'ordre (4, 7) de ces expressions, mais si c'est le cas, je pense que cette méthode sera la plus facile pour les trouver. Revenons donc à notre terme pour l'instant mal défini. Il nous faut une prescription pour pouvoir calculer l'expression, puis vérifier qu'elle est bien solution, et bien sûr espérer que cette prescription fonctionne toujours aux ordres supérieurs. La longueur des formules se prêtant assez peu à un exercice de devinettes, je vais indiquer la règle que je propose. Bien sûr, pour la trouver il a fallu faire des allers et retours entre des solutions à l'ordre (2, 4) et l'expression générée par la formule (1.140). L'objectif étant de trouver des formules suffisamment concordantes pour espérer avoir choisie la bonne solution, et pouvoir donc faire la bonne prescription. En voilà donc le résultat, précédé de quelques définitions :

Sur la deuxième diagonale, l'ambiguïté est toujours de la forme

$$(\overrightarrow{B_1}\overrightarrow{B_2}\dots\overrightarrow{B_p})T_0(\overleftarrow{\partial}\partial\theta\theta\theta\overrightarrow{\partial}\partial\partial)$$

où les B_k sont des blocs différentiels d'ordre 1, autrement dit $A\theta\partial$, $A\theta F\theta\partial$, etc... Ces blocs s'écrivent $B_k = T_k\theta\partial$, où T_k est une expression de la forme $A(\theta F)^m$ (θF répété m fois). T_0 est aussi une telle expression.

Le premier exemple est celui que l'on vient de rencontrer, où $T_0 = A$ et le seul bloc est $B_1 = A\theta\partial$. Aux ordres suivants, les T_k seront des expressions plus complexes. Ces blocs sont bien entendu ceux qui composent les formules des $A^{(n,n)}$.

Règle :

$$(\overrightarrow{B_1 B_2 \dots B_p}) T_0 (\overleftarrow{\partial \partial \theta \theta \theta \partial \partial \partial}) \equiv (B_1 B_2 \dots B_p) (\partial \partial T_0 \theta \theta \theta \partial \partial \partial) \quad (1.142)$$

$$+ n B_1 \dots B_{p-1} \partial \partial_i T_p \theta F_i (\theta F)^{m_0} \theta \theta \theta \partial \partial \partial$$

où n est l'ordre en A de l'ensemble des blocs B_k (sans oublier que F est d'ordre 1 en A). De plus, $F_i (\theta F)^{m_0}$ signifie que l'on a remplacé dans T_0 le A initial par F_i .

A l'ordre (2, 4), les choses sont simples. Le bloc B_1 est d'ordre 1 en A , donc $n = 1$, et on a :

$$(A\theta\partial)A(\overleftarrow{\partial \partial \theta \theta \theta \partial \partial \partial}) \equiv (A\theta\partial)(\partial \partial A \theta \theta \theta \partial \partial \partial) + \partial \partial_i A \theta F_i \theta \theta \theta \partial \partial \partial \quad (1.143)$$

Nous pouvons donc écrire les termes qui vont composer $\lambda^{(2,4)}$. Sa première partie ne contient aucune ambiguïté, et nous venons d'écrire celle de la deuxième partie :

$$A^{(1,3)}\theta\partial\lambda \propto \partial\partial A \theta \theta \theta \partial \partial \partial (A\theta\partial\lambda) + \partial\partial_i A \theta \theta \theta \partial \partial (F_i \theta \partial \lambda)$$

$$\partial\partial A^{(1,1)} \theta \theta \theta \partial \partial \partial \lambda \propto \partial\partial (A_i \theta F_i) \theta \theta \theta \partial \partial \partial \lambda$$

$$+ (A\theta\partial)(\partial\partial A \theta \theta \theta \partial \partial \partial \lambda) + \partial\partial_i A \theta F_i \theta \theta \theta \partial \partial \partial \lambda$$

Nous voilà donc presque prêts à brandir notre candidat-solution à l'ordre (2, 4). "Presque" parce que bien sûr, il faut encore lui adjoindre $A_\mu^{(2,4)}$, qui s'écrit de la même manière, en remplaçant "simplement" $\partial\lambda$ par $(\partial A_\mu + 2F_\mu)$. Il reste cependant une dernière question à régler : lorsque l'on rencontre non pas $\partial\lambda$, mais $\partial\partial\partial\lambda$ par exemple, quel dérivée faut-il remplacer par F ? Car si les trois indices de $\partial\partial\partial$ sont symétriques, ceux de $\partial\partial F_\mu$ ne le sont plus, et il y a donc au pire trois possibilités. En effet, ces trois indices sont contractés avec trois autres (par l'intermédiaire des θ), et il faut choisir celui qui sera contracté avec F . Exemples :

– $\partial\partial A \theta \theta \theta \partial \partial \partial A_\mu$ sera accompagné de $\partial\partial_i A \theta \theta \theta \partial \partial F_{i\mu}$ (et non pas de $\partial\partial A_i \dots$).

Cet exemple est exactement le choix que nous avons fait pour $A_\mu^{(1,3)}$.

– $\partial\partial_i A \theta F_i \theta \theta \theta \partial \partial \partial A_\mu$ sera accompagné de $\partial_j \partial_i A \theta F_i \theta \theta \theta \partial \partial F_{j\mu}$.

Pour faire simple, la règle est donc : *toujours choisir la dérivée*.

Munis de ces deux règles, nous pouvons donc finalement écrire les expressions à l'ordre (2, 4). Je ferai grâce de la vérification (d'autant qu'il est beaucoup plus rapide de rentrer ces expressions dans le logiciel que j'ai conçu spécialement pour cela, et de lui laisser dix secondes pour faire les calculs), mais elles sont effectivement solutions

de l'équation de Seiberg-Witten à cet ordre :

$$\begin{aligned} \lambda^{(2,4)} = \frac{-1}{144} [& (A\theta\partial)(\partial\partial A \theta\theta\theta \partial\partial\partial\lambda) + \partial\partial A \theta\theta\theta \partial\partial\partial(A\theta\partial\lambda) \\ & + \partial\partial(A\theta F) \theta\theta\theta \partial\partial\partial\lambda + \partial\partial_i A \theta\theta\theta \partial\partial(F_i \theta\partial\lambda) \\ & + \partial\partial_i A \theta F_i \theta\theta\theta \partial\partial\partial\lambda] \end{aligned} \quad (1.144)$$

$$\begin{aligned} A_\mu^{(2,4)} = \frac{-1}{144} [& (A\theta\partial)(\partial\partial A \theta\theta\theta \partial\partial\partial A_\mu) + \partial\partial A \theta\theta\theta \partial\partial\partial(A\theta\partial A_\mu) \\ & + \partial\partial(A\theta F) \theta\theta\theta \partial\partial\partial A_\mu + \partial\partial_i A \theta\theta\theta \partial\partial(F_i \theta\partial A_\mu) \\ & + \partial\partial_i A \theta F_i \theta\theta\theta \partial\partial\partial A_\mu \\ & + 2(A\theta\partial)(\partial\partial_i A \theta\theta\theta \partial\partial F_{i\mu}) + 2\partial\partial A \theta\theta\theta \partial\partial\partial(A\theta F_\mu) \\ & + 2\partial\partial_i(A\theta F) \theta\theta\theta \partial\partial F_{i\mu} + 2\partial\partial_i A \theta\theta\theta \partial\partial(F_i \theta F_\mu) \\ & + 2\partial_i \partial_j A \theta F_j \theta\theta\theta \partial\partial F_{i\mu}] \end{aligned} \quad (1.145)$$

Pour le plaisir des yeux, et pour vérifier que les coefficients de la règle énoncée sont bien corrects, osons nous pencher rapidement sur l'ordre suivant. En effet, à l'ordre (2, 4), il n'y avait qu'un bloc à commuter, et T_0 était réduit à sa plus simple expression. À l'ordre (3, 5), par contre,

$$\begin{aligned} \lambda^{(3,\infty)} &= \frac{1}{4} \left[-A^{(2,\infty)} \theta\partial\lambda + \frac{1}{24} \partial\partial A^{(2,\infty)} \theta\theta\theta \partial\partial\partial\lambda + \dots \right] \\ &= \underbrace{-\frac{1}{4} A^{(2,2)} \theta\partial\lambda}_{\lambda^{(3,3)}} - \underbrace{\frac{1}{4} A^{(2,4)} \theta\partial\lambda + \frac{1}{96} \partial\partial A^{(2,2)} \theta\theta\theta \partial\partial\partial\lambda + \dots}_{\lambda^{(3,5)}} \end{aligned} \quad (1.146)$$

Ainsi, si le premier terme de $\lambda^{(3,5)}$ est trivial, le second ne l'est pas :

$$\begin{aligned} A^{(2,2)} \overleftarrow{\partial\partial} \theta\theta\theta \partial\partial\partial\lambda &\propto 2\partial\partial(A\theta F\theta F) \theta\theta\theta \partial\partial\partial\lambda + 2(A\theta\partial)(A\theta F) \overleftarrow{\partial\partial} \theta\theta\theta \partial\partial\partial\lambda \\ &+ (A\theta F\theta\partial)A \overleftarrow{\partial\partial} \theta\theta\theta \partial\partial\partial\lambda + (A\theta\partial)(A\theta\partial)A \overleftarrow{\partial\partial} \theta\theta\theta \partial\partial\partial\lambda \end{aligned}$$

Ainsi, dans le deuxième terme de cette expression, T_0 n'est pas trivial ($T_0 = A\theta F$). Dans le suivant, c'est le bloc T_1 qui prend cette valeur. Quant au dernier, il présente deux blocs B_1 et B_2 . La règle proposée plus haut donne alors les résultats suivants :

$$\begin{aligned} (A\theta\partial)(A\theta F) \overleftarrow{\partial\partial} \theta\theta\theta \partial\partial\partial\lambda &= (A\theta\partial)(\partial\partial(A\theta F) \theta\theta\theta \partial\partial\partial\lambda) \\ &+ \partial\partial_i A \theta F_{ij} \theta F_j \theta\theta\theta \partial\partial\partial\lambda \\ (A\theta F\theta\partial)A \overleftarrow{\partial\partial} \theta\theta\theta \partial\partial\partial\lambda &= (A\theta F\theta\partial)(\partial\partial A \theta\theta\theta \partial\partial\partial\lambda) \\ &+ 2\partial\partial_i(A\theta F) \theta F_i \theta\theta\theta \partial\partial\partial\lambda \\ (A\theta\partial)(A\theta\partial)A \overleftarrow{\partial\partial} \theta\theta\theta \partial\partial\partial\lambda &= (A\theta\partial)(A\theta\partial)(\partial\partial A \theta\theta\theta \partial\partial\partial\lambda) \\ &+ 2(A\theta\partial)(\partial\partial_i A \theta F_i \theta\theta\theta \partial\partial\partial\lambda) \end{aligned}$$

Regroupant tous ces morceaux, on en déduit l'expression de $\lambda^{(3,5)}$ et, en étant attentifs aux contractions du F_μ final, celle de $A_\mu^{(3,5)}$. Encore une fois, j'ai bien sûr vérifié

que ces expressions prolongent les solutions précédentes à cet ordre.

$$\begin{aligned}
 \lambda^{(3,5)} = \frac{1}{576} [& (\partial\partial A \theta\theta\theta \partial\partial\partial)(A\theta\partial)(A\theta\partial)\lambda + (A\theta\partial)(\partial\partial A \theta\theta\theta \partial\partial\partial)(A\theta\partial)\lambda \\
 & + (A\theta\partial)(A\theta\partial)(\partial\partial A \theta\theta\theta \partial\partial\partial)\lambda + 2 \partial\partial A \theta\theta\theta \partial\partial\partial(A\theta F\theta\partial\lambda) \\
 & + 2 (A\theta\partial)(\partial\partial(A\theta F) \theta\theta\theta \partial\partial\partial\lambda) + 2 (A\theta\partial)(\partial\partial_i A \theta F_i \theta\theta\theta \partial\partial\partial\lambda) \\
 & + 2 (A\theta\partial)(\partial\partial_i A \theta\theta\theta \partial\partial(F_i \theta\partial\lambda)) + (A\theta F\theta\partial)(\partial\partial A \theta\theta\theta \partial\partial\partial\lambda) \\
 & + \partial\partial(A\theta F) \theta\theta\theta \partial\partial\partial(A\theta\partial\lambda) + \partial\partial_i A \theta F_i \theta\theta\theta \partial\partial\partial(A\theta\partial\lambda) \\
 & + \partial\partial_i A \theta\theta\theta \partial\partial(F_i \theta\partial(A\theta\partial\lambda)) + 2 \partial\partial(A\theta F\theta F) \theta\theta\theta \partial\partial\partial\lambda \\
 & + 2 \partial\partial_i A \theta F_{ij} \theta F_j \theta\theta\theta \partial\partial\partial\lambda + 2 \partial\partial_i A \theta\theta\theta \partial\partial(F_i \theta F\theta\partial\lambda) \\
 & + 2 \partial\partial_i(A\theta F) \theta F_i \theta\theta\theta \partial\partial\partial\lambda + 2 \partial\partial_i(A\theta F) \theta\theta\theta \partial\partial(F_i \theta\partial\lambda) \\
 & + 2 \partial_i \partial_j A \theta F_j \theta\theta\theta \partial\partial(F_i \theta\partial\lambda)] \quad (1.147)
 \end{aligned}$$

$$\begin{aligned}
 A_\mu^{(3,5)} = \frac{1}{576} [& (\partial\partial A \theta\theta\theta \partial\partial\partial)(A\theta\partial)(A\theta\partial)A_\mu + (A\theta\partial)(\partial\partial A \theta\theta\theta \partial\partial\partial)(A\theta\partial)A_\mu \\
 & + (A\theta\partial)(A\theta\partial)(\partial\partial A \theta\theta\theta \partial\partial\partial)A_\mu + 2 \partial\partial A \theta\theta\theta \partial\partial\partial(A\theta F\theta\partial A_\mu) \\
 & + 2 (A\theta\partial)(\partial\partial(A\theta F) \theta\theta\theta \partial\partial\partial A_\mu) + 2 (A\theta\partial)(\partial\partial_i A \theta F_i \theta\theta\theta \partial\partial\partial A_\mu) \\
 & + 2 (A\theta\partial)(\partial\partial_i A \theta\theta\theta \partial\partial(F_i \theta\partial A_\mu)) + (A\theta F\theta\partial)(\partial\partial A \theta\theta\theta \partial\partial\partial A_\mu) \\
 & + \partial\partial(A\theta F) \theta\theta\theta \partial\partial\partial(A\theta\partial A_\mu) + \partial\partial_i A \theta F_i \theta\theta\theta \partial\partial\partial(A\theta\partial A_\mu) \\
 & + \partial\partial_i A \theta\theta\theta \partial\partial(F_i \theta\partial(A\theta\partial A_\mu)) + 2 \partial\partial(A\theta F\theta F) \theta\theta\theta \partial\partial\partial A_\mu \\
 & + 2 \partial\partial_i A \theta F_{ij} \theta F_j \theta\theta\theta \partial\partial\partial A_\mu + 2 \partial\partial_i A \theta\theta\theta \partial\partial(F_i \theta F\theta\partial A_\mu) \\
 & + 2 \partial\partial_i(A\theta F) \theta F_i \theta\theta\theta \partial\partial\partial A_\mu + 2 \partial\partial_i(A\theta F) \theta\theta\theta \partial\partial(F_i \theta\partial A_\mu) \\
 & + 2 \partial_i \partial_j A \theta F_j \theta\theta\theta \partial\partial(F_i \theta\partial A_\mu) \\
 & + 3 (\partial\partial A \theta\theta\theta \partial\partial\partial)(A\theta\partial)(A\theta F_{,\mu}) + 3 (A\theta\partial)(\partial\partial A \theta\theta\theta \partial\partial\partial)(A\theta F_{,\mu}) \\
 & + 3 (A\theta\partial)(A\theta\partial)(\partial\partial_i A \theta\theta\theta \partial\partial F_{i\mu}) + 6 \partial\partial A \theta\theta\theta \partial\partial\partial(A\theta F\theta F_{,\mu}) \\
 & + 6 (A\theta\partial)(\partial\partial_i(A\theta F) \theta\theta\theta \partial\partial F_{i\mu}) + 6 (A\theta\partial)(\partial_i \partial_j A \theta F_j \theta\theta\theta \partial\partial F_{i\mu}) \\
 & + 6 (A\theta\partial)(\partial\partial_i A \theta\theta\theta \partial\partial(F_i \theta F_{,\mu})) + 3 (A\theta F\theta\partial)(\partial\partial_i A \theta\theta\theta \partial\partial F_{i\mu}) \\
 & + 3 \partial\partial(A\theta F) \theta\theta\theta \partial\partial\partial(A\theta F_{,\mu}) + 3 \partial\partial_i A \theta F_i \theta\theta\theta \partial\partial\partial(A\theta F_{,\mu}) \\
 & + 3 \partial\partial_i A \theta\theta\theta \partial\partial(F_i \theta\partial(A\theta F_{,\mu})) + 6 \partial\partial_i(A\theta F\theta F) \theta\theta\theta \partial\partial F_{i\mu} \\
 & + 6 \partial_i \partial_j A \theta F_{jk} \theta F_k \theta\theta\theta \partial\partial F_{i\mu} + 6 \partial\partial_i A \theta\theta\theta \partial\partial(F_i \theta F\theta F_{,\mu}) \\
 & + 6 \partial_i \partial_j(A\theta F) \theta F_j \theta\theta\theta \partial\partial F_{i\mu} + 6 \partial\partial_i(A\theta F) \theta\theta\theta \partial\partial(F_i \theta F_{,\mu}) \\
 & + 6 \partial_i \partial_j A \theta F_j \theta\theta\theta \partial\partial(F_i \theta F_{,\mu})] \quad (1.148)
 \end{aligned}$$

1.7 L'approche algébrique

1.7.1 La résolution théorique

Nous allons maintenant reprendre la méthode expliquée en 1.5.3, et calculer l'opérateur $\mathcal{D}$, qui réalise l'équivalence entre les star-produits $\star$ et $\star'$. Pour cela, nous avons besoin d'un résultat préliminaire :

Lemme : Si f est une fonction d'un paramètre t qui vérifie

$$(\partial_t + A(t)) f = 0 \quad (1.149)$$

où $A(t)$ est un opérateur quelconque. Alors les valeurs $f(0)$ et $f(1)$ sont reliées par

$$e^{(\partial_t + A(t))} e^{-\partial_t} \Big|_{t=0} f(1) = f(0) \quad (1.150)$$

Démonstration

La formule de Taylor pour f s'écrit

$$\begin{aligned} f(t + t_0) &= \sum_{k=0}^{\infty} \frac{t_0^k}{k!} f^{(k)}(t) = \sum_{k=0}^{\infty} \frac{t_0^k}{k!} \partial_t^k (f)(t) \\ &= \left[\sum_{k=0}^{\infty} \frac{(t_0 \partial_t)^k}{k!} \right] (f)(t) \\ &= [e^{t_0 \partial_t} f](t) \end{aligned}$$

En particulier, avec $t_0 = 1$

$$f(t + 1) = [e^{\partial_t} f](t) = [e^{\partial_t} e^{-(\partial_t + A(t))} f](t) \quad (1.151)$$

puisque f vérifie (1.149), et donc $e^{-(\partial_t + A(t))} f = f$. Il n'y a plus maintenant qu'à réarranger cette expression en passant les opérateurs dans le membre de gauche, et l'évaluer en $t = 0$, pour trouver la formule (1.150). $\square$

Ainsi, pour relier $\star$ et $\star'$ par un tel opérateur, il nous faut d'abord introduire un paramètre de déformation t , un star-produit $\star_t$, fonction de t qui interpole entre $\star$ et $\star'$, et qui vérifie une équation d'évolution similaire à (1.149). Pour cela, reprenons la relation (1.108) entre θ et θ' , et définissons θ_t par

$$\theta_t \equiv \theta \frac{1}{1 + tF\theta} \quad (1.152)$$

On a alors $\theta_0 = \theta$ et $\theta_1 = \theta'$, et la relation

$$\partial_t \theta_t = -\theta \frac{1}{1 + tF\theta} F\theta \frac{1}{1 + tF\theta} = -\theta_t F\theta_t \quad (1.153)$$

Il nous faut vérifier que θ_t est bien une structure de Poisson, à partir de laquelle on peut définir un star-produit. Nous allons réutiliser l'arsenal algébrique défini en 1.3.3. Introduisons le champ vectoriel

$$\mathbf{a}_\theta = A_i \mathbf{d}_\theta x^i = -A_i \theta^{ij} \partial_j \quad (1.154)$$

où $\mathbf{d}_\theta$ est l'opérateur défini en (1.76). On définit aussi $\mathbf{f}_\theta$ par

$$\mathbf{f}_\theta = \mathbf{d}_\theta \mathbf{a}_\theta = -\frac{1}{2} \theta^{ik} F_{kl} \theta^{lj} \partial_i \wedge \partial_j \quad (1.155)$$

On peut alors réécrire la variation de θ_t comme

$$\partial_t \theta_t = \mathbf{f}_{\theta_t} = \mathbf{d}_{\theta_t} \mathbf{a}_{\theta_t} = -[\mathbf{a}_{\theta_t}, \theta_t]_S \quad (1.156)$$

De plus, on peut maintenant justifier que θ_t est une structure de Poisson, parce que $[\theta_t, \theta_t]_S = 0$ en $t = 0$, et

$$\partial_t [\theta_t, \theta_t]_S = -2\mathbf{d}_{\theta_t} \mathbf{f}_{\theta_t} \propto [\theta_t, \theta_t]_S \quad (1.157)$$

Avant de revenir aux star-produits, remarquons juste que l'opérateur $\mathbf{a}_{\theta_t}$ est exactement celui que nous avons introduit en (1.122). De même que (1.153) est équivalente à (1.121). Si l'on utilise le lemme précédent sur la structure de Poisson θ_t , qui suit l'équation d'évolution (1.156), on obtient l'opérateur réalisant le transport entre les structures θ et θ'

$$\rho_a^* = e^{(\mathbf{a}_{\theta_t} + \partial_t)} e^{-\partial_t} |_{t=0} \quad (1.158)$$

qui est une version partielle de l'opérateur quantique total $\mathcal{D}_a$ que nous cherchons. Ainsi, son action sur la coordonnée x^i définit une version partielle du potentiel de jauge non-commutatif, à savoir sa partie semi-classique (écrite ici lorsque θ est constant)

$$\rho_a^* x^i = x^i + \theta^{ij} A_j^{sc} \quad (1.159)$$

On remarque en effet que le potentiel A^{sc} défini ici est exactement le même que celui de la section 1.6 : la formule (1.159) se réduit à (1.132) puisque $e^{-\partial_t} x^i = x^i$.

Quantifions maintenant tout ce petit monde, en définissant le star-produit $\star_t$ par la transformation de formalité

$$g \star_t h \equiv \sum_{n=0}^{\infty} \frac{(i\hbar)^n}{n!} U_n(\theta_t, \dots, \theta_t)(g, h) \quad (1.160)$$

$\star_t$ est une déformation du star-produit initial, qui interpole entre $\star = \star_0$ et $\star' = \star_1$. De même, définissons les opérateurs quantiques associés aux opérateurs semi-classiques précédents

$$\mathbf{a}_{\star_t} = \Phi_t(\mathbf{a}_{\theta_t}) = \sum_{n=0}^{\infty} \frac{(i\hbar)^n}{n!} U_{n+1}(\mathbf{a}_{\theta_t}, \theta_t, \dots, \theta_t) \quad (1.161)$$

$$\mathbf{f}_{\star_t} = \mathbf{d}_{\star_t} \mathbf{a}_{\star_t} = i\hbar \Phi_t(\mathbf{d}_{\theta_t} \mathbf{a}_{\theta_t}) = \sum_{n=0}^{\infty} \frac{(i\hbar)^{n+1}}{n!} U_{n+1}(f_{\theta_t}, \theta_t, \dots, \theta_t) \quad (1.162)$$

où $\mathbf{d}_{\star_t}$ est l'opérateur de cobord défini en (1.80), et où l'on a utilisé (1.86). On peut alors calculer la variation de $\star_t$ par rapport au paramètre de déformation t , en utilisant la multi-linéarité de U_n

$$\partial_t(g \star_t h) = \sum_{n=1}^{\infty} \frac{(i\hbar)^n}{(n-1)!} U_n(f_{\theta_t}, \theta_t, \dots, \theta_t)(g, h) = \mathbf{f}_{\star_t}(g, h) \quad (1.163)$$

ce qui conduit à l'équation opératorielle vérifiée par $\star_t$

$$\partial_t \star_t = \mathbf{f}_{\star_t} = \mathbf{d}_{\star_t} \mathbf{a}_{\star_t} = -[\mathbf{a}_{\star_t}, \star_t]_G \quad (1.164)$$

Comme précédemment avec la structure semi-classique, le lemme nous donne donc l'expression du flux qui réalise l'entrelacement des star-produits $\star$ et $\star'$

$$\mathcal{D}_a = e^{(\mathbf{a}_{\star_t} + \partial_t)} e^{-\partial_t} |_{t=0} \quad (1.165)$$

Comme nous allons le vérifier, ce flux est aussi l'opérateur qui covariantise les fonctions par rapport à la théorie de jauge non-commutative, et qui définit donc le potentiel de jauge non-commutatif

$$\mathcal{A}_a = \mathcal{D}_a - \text{id} \quad (1.166)$$

Pour vérifier cette affirmation, et donc justifier la méthode, il nous reste à calculer la variation de $\mathcal{A}_a$ associée à une variation $d\lambda$ du potentiel commutatif. En d'autres termes, vérifier que $\mathcal{A}_{a+d\lambda} - \mathcal{A}_a$ a bien la valeur voulue.

Tout d'abord, la transformation de jauge $A \mapsto A + d\lambda$ agit sur le champ vectoriel $\mathbf{a}_\theta$ par

$$\mathbf{a}_\theta \mapsto -A_i \theta^{ij} \partial_j - \partial_i \lambda \theta^{ij} \partial_j = \mathbf{a}_\theta + \mathbf{d}_\theta \lambda \quad (1.167)$$

où $\mathbf{d}_\theta \lambda$ est le champ vectoriel hamiltonien

$$\mathbf{d}_\theta \lambda = -[\lambda, \theta]_S = \{ \cdot, \theta; \lambda \} = -\partial_i \lambda \theta^{ij} \partial_j \quad (1.168)$$

Ensuite, cette variation est transportée par la formalité sur $\mathbf{a}_\star = \Phi(\mathbf{a}_\theta)$. En utilisant l'identité (1.86), cette variation s'écrit

$$\mathbf{a}_\star \mapsto \Phi(\mathbf{a}_\theta + \mathbf{d}_\theta \lambda) = \mathbf{a}_\star + \frac{1}{i\hbar} \mathbf{d}_\star \Phi(\lambda) \quad (1.169)$$

Nous pouvons maintenant évaluer la variation du flux $\mathcal{D}_a$

$$\mathcal{D}_{a+d\lambda} - \mathcal{D}_a = \left[e^{(\mathbf{a}_{\star t} + \frac{1}{i\hbar} \mathbf{d}_{\star t} \Phi_t(\lambda) + \partial_t)} - e^{(\mathbf{a}_{\star t} + \partial_t)} \right] e^{-\partial_t} \Big|_{t=0} \quad (1.170)$$

Cette expression est de la forme

$$e^{A+\epsilon B} - e^A = \epsilon \sum_{n=0}^{\infty} \frac{B_n}{(n+1)!} e^A + o(\epsilon^2) \quad (1.171)$$

où ϵ symbolise la valeur infinitésimale de λ , et

$$B_0 \equiv B, \quad B_{n+1} \equiv [A, B_n] \quad (1.172)$$

où dans notre cas bien sûr

$$A = \mathbf{a}_{\star t} + \partial_t, \quad B = \frac{1}{i\hbar} \mathbf{d}_{\star t} \Phi_t(\lambda) \quad (1.173)$$

et où le crochet est le crochet de Gerstenhaber (1.78). On peut alors remarquer que²⁶, quelle que soit la fonction U_t dépendant de t

$$\begin{aligned} [\mathbf{a}_{\star t} + \partial_t, \mathbf{d}_{\star t} U_t]_G(g) &= \mathbf{a}_{\star t}(\mathbf{d}_{\star t} U_t)(g) + \partial_t(\mathbf{d}_{\star t} U_t)(g) - (\mathbf{d}_{\star t} U_t)(\mathbf{a}_{\star t} + \partial_t)(g) \\ &= \mathbf{a}_{\star t}([g \star U_t]) + \partial_t[g \star U_t] - [\mathbf{a}_{\star t} g \star U_t] \\ &= [g \star \mathbf{a}_{\star t} U_t] + [g \star \partial_t U_t] \\ &= \mathbf{d}_{\star t}[(\mathbf{a}_{\star t} + \partial_t)U_t](g) \end{aligned}$$

²⁶ En utilisant à plusieurs reprises l'expression de $\mathbf{d}_\star U$, par définition de l'opérateur $\mathbf{d}_\star$ et du crochet de Gerstenhaber

$$\mathbf{d}_\star U = -[U, \star]_G = [\cdot \star U] \quad (1.174)$$

De plus, $\mathbf{f}_\star$ mesure la variation à la règle de Leibniz de $\mathbf{a}_\star$ en tant que dérivée

$$\mathbf{f}_\star(g, U) = -[\mathbf{a}_\star, \star]_G(g, U) = -\mathbf{a}_\star(g \star U) + \mathbf{a}_\star(g) \star U + g \star \mathbf{a}_\star(U) \quad (1.175)$$

mais aussi la variation du star-produit $\star_t$ en fonction de t : $\partial_t(g \star_t U_t) = \mathbf{f}_{\star_t}(g, U_t) + g \star_t \partial_t U_t$. Entre la deuxième et la troisième ligne, ces deux relations se compensent partiellement.

En utilisant cette relation récursivement, on trouve donc l'expression de B_n

$$B_n = \frac{1}{i\hbar} \mathbf{d}_{\star_t} [(\mathbf{a}_{\star_t} + \partial_t)^n \Phi_t(\lambda)] \quad (1.176)$$

et en remplaçant dans (1.170), au premier ordre en λ

$$\begin{aligned} \mathcal{D}_{a+d\lambda} - \mathcal{D}_a &= \left[\frac{1}{i\hbar} \sum_{n=0}^{\infty} \frac{\mathbf{d}_{\star_t} [(\mathbf{a}_{\star_t} + \partial_t)^n \Phi_t(\lambda)]}{(n+1)!} e^{(\mathbf{a}_{\star_t} + \partial_t)} \right] e^{-\partial_t} \Big|_{t=0} \\ &= \frac{1}{i\hbar} (\mathbf{d}_{\star} \hat{\lambda}) \mathcal{D}_a \end{aligned} \quad (1.177)$$

avec l'expression de $\hat{\lambda}$:

$$\hat{\lambda} = \sum_{n=0}^{\infty} \frac{(\mathbf{a}_{\star_t} + \partial_t)^n}{(n+1)!} \Phi_t(\lambda) \Big|_{t=0} \quad (1.178)$$

Si maintenant on remplace dans (1.177) $\mathcal{D}_a$ par $(\text{id} + \mathcal{A}_a)$, on obtient la transformation de jauge du potentiel non-commutatif abstrait $\mathcal{A}$

$$\mathcal{A}_{a+d\lambda} - \mathcal{A}_a = \frac{1}{i\hbar} [\mathbf{d}_{\star} \hat{\lambda} + (\mathbf{d}_{\star} \hat{\lambda}) \mathcal{A}_a] \quad (1.179)$$

ce qui est bien la formule attendue. En effet, en se souvenant que le potentiel $\hat{A}$ est défini à partir du potentiel abstrait $\mathcal{A}$ par la formule

$$\mathcal{A}(x^i) = \theta^{ij} \hat{A}_j \quad (1.180)$$

et avec

$$(\mathbf{d}_{\star} \hat{\lambda})(x^i) = [x^i \star \hat{\lambda}] = i\hbar \theta^{ij} \partial_j \hat{\lambda} \quad (1.181)$$

on retrouve bien²⁷ la transformation de jauge de $\hat{A}$

$$\begin{aligned} (\mathcal{A}_{a+d\lambda} - \mathcal{A}_a)(x^i) &= \frac{1}{i\hbar} [x^i \star \hat{\lambda}] - \frac{1}{i\hbar} [\hat{\lambda} \star \theta^{ij} \hat{A}_j] \\ \theta^{ij} \delta \hat{A}_j &= \theta^{ij} \left(\partial_j \hat{\lambda} - \frac{i}{\hbar} [\hat{A}_j \star \hat{\lambda}] \right) \end{aligned} \quad (1.182)$$

Ainsi, en résumé, nous disposons des expressions exactes du potentiel de jauge et du paramètre de jauge non-commutatifs :

$$\theta^{ij} \hat{A}_j = \sum_{n=0}^{\infty} \frac{(\mathbf{a}_{\star_t} + \partial_t)^n}{(n+1)!} \mathbf{a}_{\star_t} x^i \Big|_{t=0} = \sum_{n=0}^{\infty} \theta^{ij} A_j^{(n,\infty)} \quad (1.183)$$

$$\hat{\lambda} = \sum_{n=0}^{\infty} \frac{(\mathbf{a}_{\star_t} + \partial_t)^n}{(n+1)!} \lambda \Big|_{t=0} = \sum_{n=0}^{\infty} \theta^{ij} \lambda^{(n,\infty)} \quad (1.184)$$

²⁷en comparant (1.182) et (1.88), on se souviendra que le star-produit utilisé dans (1.88) ne contenait pas de $\hbar$ dans sa définition.

1.7.2 Les expressions concrètes

Pour expliciter les expressions de cette solution, nous allons devoir calculer l'opérateur $\mathbf{a}_{\star t}$, et développer les formules (1.183) et (1.184). Nous allons voir que les résultats seront exactement ceux de la section 1.6, non seulement pour les termes semi-classiques, comme le suggéraient (1.129) et (1.130), mais aussi pour les autres.

$\mathbf{a}_{\star t}$ est une série en puissance de θ_t , dont tous les termes d'ordre pair sont nuls, à cause du caractère abélien de la théorie de jauge. Souvenons-nous que c'est cela aussi qui annule un ordre sur deux dans l'équation de Seiberg-Witten. En notant $a_t^{(n)}$ l'ordre n de ce développement

$$\mathbf{a}_{\star t} = a_t^{(1)} + a_t^{(3)} + a_t^{(5)} + \dots \quad (1.185)$$

où nous savons déjà que $a_t^{(1)} = \mathbf{a}_{\theta t} = -A\theta_t\partial$.

Ecrivons alors les premiers ordres pour λ

$$\lambda^{(0,\infty)} = \lambda \quad (1.186)$$

$$\lambda^{(1,\infty)} = \underbrace{\frac{1}{2} a_t^{(1)} \lambda|_{t=0}}_{\lambda^{(1,1)}} + \underbrace{\frac{1}{2} a_t^{(3)} \lambda|_{t=0}}_{\lambda^{(1,3)}} + \dots \quad (1.187)$$

$$\begin{aligned} \lambda^{(2,\infty)} = & \underbrace{\frac{1}{6} \left[a_t^{(1)} a_t^{(1)} + \partial_t a_t^{(1)} \right]_{t=0}}_{\lambda^{(2,2)}} \lambda \\ & + \underbrace{\frac{1}{6} \left[a_t^{(1)} a_t^{(3)} + \partial_t a_t^{(3)} + a_t^{(3)} a_t^{(1)} \right]_{t=0}}_{\lambda^{(2,4)}} \lambda + \dots \end{aligned} \quad (1.188)$$

et pour la connexion A

$$\theta^{ij} A_j^{(0,\infty)} = a_t^{(1)} x^i|_{t=0} = \theta^{ij} A_j \quad (1.189)$$

$$\begin{aligned} \theta^{ij} A_j^{(1,\infty)} = & \underbrace{\frac{1}{2} \left[a_t^{(1)} a_t^{(1)} + \partial_t a_t^{(1)} \right]_{t=0} x^i}_{\theta^{ij} A_j^{(1,1)}} \\ & + \underbrace{\frac{1}{2} \left[a_t^{(1)} a_t^{(3)} + \partial_t a_t^{(3)} + a_t^{(3)} a_t^{(1)} \right]_{t=0} x^i}_{\theta^{ij} A_j^{(1,3)}} + \dots \end{aligned} \quad (1.190)$$

$$\begin{aligned} \theta^{ij} A_j^{(2,\infty)} = & \underbrace{\frac{1}{6} \left[a_t^{(1)} a_t^{(1)} a_t^{(1)} + (\partial_t a_t^{(1)}) a_t^{(1)} + 2a_t^{(1)} \partial_t a_t^{(1)} + \partial_t \partial_t a_t^{(1)} \right]_{t=0} x^i}_{\theta^{ij} A_j^{(2,2)}} \\ & + \underbrace{\frac{1}{6} \left[\begin{aligned} & a_t^{(3)} a_t^{(1)} a_t^{(1)} + a_t^{(1)} a_t^{(3)} a_t^{(1)} + a_t^{(1)} a_t^{(1)} a_t^{(3)} \\ & + \left(\partial_t a_t^{(3)} \right) a_t^{(1)} + 2a_t^{(3)} \partial_t a_t^{(1)} \\ & + \left(\partial_t a_t^{(1)} \right) a_t^{(3)} + 2a_t^{(1)} \partial_t a_t^{(3)} + \partial_t \partial_t a_t^{(3)} \end{aligned} \right]_{t=0} x^i}_{\theta^{ij} A_j^{(2,4)}} \\ & + \dots \end{aligned} \quad (1.191)$$

Les termes d'ordre (n, n) , semi-classiques, sont donc bien ceux de la section 1.6, comme on peut s'en convaincre en remplaçant $a_t^{(1)}$ et ses dérivées par leur valeur :

$$\begin{aligned} a_t^{(1)} &= -A\theta_t\partial & a_t^{(1)}x^i|_{t=0} &= \theta^{i\mu}A_\mu \\ \partial_t a_t^{(1)} &= A\theta_t F\theta_t\partial & \partial_t a_t^{(1)}x^i|_{t=0} &= -\theta^{i\mu}A\theta F_\mu \\ \partial_t\partial_t a_t^{(1)} &= -2A\theta_t F\theta_t F\theta_t\partial & \partial_t\partial_t a_t^{(1)}x^i|_{t=0} &= 2\theta^{i\mu}A\theta F\theta F_\mu \end{aligned}$$

Pour calculer les termes d'ordre $(n, n+2)$, il nous faut l'expression explicite de $a_t^{(3)}$:

$$a_t^{(3)} = -\frac{\hbar^2}{2}U_3(-A\theta_t\partial, \theta_t, \theta_t) \quad (1.192)$$

Pour calculer U_3 , nous allons utiliser, jusqu'à un certain point, la méthode diagrammatique décrite dans [37, 39]. A chacun des trois opérateurs auxquels s'applique U_3 on associe un point, d'où partent autant de flèches que leur degré différentiel. En clair, $-A\theta_t\partial$ contient une dérivée : une flèche. θ_t est associée à deux dérivées (une pour chacun de ses indices) : deux flèches. De plus, le résultat est un opérateur agissant sur une variable, donc on ajoute un point correspondant à une fonction-test f . Cela donne la figure 1.2. Ensuite, les règles sont simples. Les flèches se terminent sur l'un quelconque des points, excepté leur point de départ. Et puisque θ_t est antisymétrique, ses deux flèches doivent terminer sur deux points différents. Pour traduire les diagrammes, il n'y a plus qu'à remplacer les flèches par les dérivées qu'elles représentent.

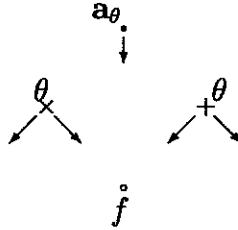


FIG. 1.2 – Le diagramme de base.

La figure 1.3 montre tous les diagrammes non nuls contenus dans $a_t^{(3)}$. Reste à calculer les coefficients de tous ces termes. Pour cela, la procédure voudrait que nous calculions des intégrales sur les espaces de configuration (la position des différents points dans le plan), mais nous allons faire autrement. Nous allons injecter tous ces termes dans l'expression de λ et A à l'ordre $(1, 3)$, puis à l'ordre $(2, 4)$, avec chacun un paramètre inconnu pour coefficient. Puis, nous écrirons l'équation de Seiberg-Witten, qui nous donnera un système d'équations sur ces paramètres. Avec un peu de nez, nous pourrions même choisir un jeu de paramètres, et il ne restera qu'à vérifier l'équation. Bien sûr, si j'ai choisi cette méthode, c'est que je disposais du programme que j'ai déjà évoqué, et qu'il a donc fait l'essentiel du calcul à ma place.

Dans l'expression de $\lambda^{(1,3)}$ les termes sont à utiliser tels quels (sans faire agir ∂_t), et tous ceux contenant des dérivées de θ_t s'annulent donc, n'en laissant qu'un seul. Dans l'expression de $A_\mu^{(1,3)}$, il y a en plus un terme venant de $\partial_t a_t^{(3)}|_{t=0}x^i$ (le dernier

sur la troisième ligne de la figure 1.3). L'équation à cet ordre fixe alors très simplement les deux coefficients correspondants (ils valent $\frac{1}{24}$) et redonnant les solutions (1.134) et (1.135).

Les douze paramètres restants sont fixés par l'ordre suivant, et redonnent les solutions (1.144) et (1.145). Il y a alors deux autres coefficients de diagrammes égaux à $\frac{1}{24}$, tous les autres étant nuls (il s'agit du premier de la ligne 2 et du troisième de la ligne 3). En les additionnant, on obtient alors la valeur de $a_t^{(3)}$:

$$a_t^{(3)} = \frac{1}{24} \partial \partial_\bullet (A \theta_t) \theta_t \partial \partial (\dot{\theta}_t \partial_\bullet) \quad (1.193)$$

où les deux $\partial_\bullet$ sont contractés sur $\dot{\theta}_t$.

1.8 Et les théories non-abéliennes ?

1.8.1 Quelques calculs à l'ordre 2

Nous allons maintenant essayer de décrire la transformation de Seiberg-Witten pour une théorie de jauge non-abélienne dans l'esprit de ce que nous avons fait jusqu'à présent.

Précisons que les expressions dans le cas non-abélien sont bien plus complexes pour différentes raisons. Tout d'abord, le développement du star-commutateur présente des termes à tous les ordres, puisque le commutateur sur l'algèbre de jauge n'est pas nul. Ainsi, par exemple, alors que l'équation à l'ordre $(1, 2)$ abélien était identiquement nulle, elle ne le sera pas ici. D'autre part, il est très difficile d'identifier l'ordre en A d'un terme, puisque la courbure, comme la dérivée covariante, contiennent des termes quadratiques. Malgré tout, nous pouvons tenter d'écrire explicitement les solutions aux premiers ordres (en θ), quitte à faire le tri par la suite pour faire émerger une structure récursive généralisant le cas abélien.

A l'ordre θ — que nous noterons $(\infty, 1)$ comme nous avons précédemment noté $(1, \infty)$ l'ordre A , l'équation de Seiberg-Witten s'écrit

$$\delta A_\mu^{(\infty, 1)} - i [\lambda, A_\mu^{(\infty, 1)}] - \partial_\mu \lambda^{(\infty, 1)} + i [A_\mu, \lambda^{(\infty, 1)}] = -\frac{1}{2} \{ \partial \lambda^\theta, \partial A_\mu \} \quad (1.194)$$

où $\{^\theta\}$ est le symbole défini à la note 19 page 31. Une solution de cette équation est

$$\begin{aligned} \lambda^{(\infty, 1)} &= -\frac{1}{4} \{ A^\theta, \partial \lambda \} \\ A_\mu^{(\infty, 1)} &= -\frac{1}{4} \{ A^\theta, \partial A_\mu + F_{\cdot\mu} \} \\ F_{\mu\nu}^{(\infty, 1)} &= -\frac{1}{2} \{ F_{\mu\cdot}^\theta, F_{\cdot\nu} \} - \frac{1}{2} \{ A^\theta, \frac{\partial + D}{2} F_{\mu\nu} \} \end{aligned} \quad (1.195)$$

où $D = \partial - i[A; \cdot]$ est la dérivée covariante.

Alors que dans le cas abélien, la solution à cet ordre était unique, il existe ici des

solutions homogènes non triviales. Choisir cette solution précise revient à choisir de généraliser le produit abélien en anticommutateur non-abélien (au lieu du produit, ou de toute autre combinaison linéaire du produit et du commutateur). En particulier, l'expression abélienne $A\theta\partial$ est ici remplacée par $-\frac{1}{2}\{A^\theta, \partial\lambda\}$.

D'autre part, on voit concrètement sur cette expression la difficulté d'identifier l'ordre des termes. En effet, dans le cas abélien, tous ces termes sont d'ordre A (c'est-à-dire plus exactement $A\lambda$ ou AA_μ), alors qu'ici il y a de l'ordre A^2 et même A^3 . Faut-il tout de même les regrouper, comme semble le suggérer l'écriture compacte avec F et D ?

Pour répondre à cette question, nous ne pouvons que tenter d'écrire "la" solution à l'ordre θ^2 . L'équation à cet ordre,

$$\begin{aligned} \delta A_\mu^{(\infty,2)} - i[\lambda, A_\mu^{(\infty,2)}] - \partial_\mu \lambda^{(\infty,2)} + i[A_\mu, \lambda^{(\infty,2)}] \\ = i[\lambda^{(\infty,1)}, A_\mu^{(\infty,1)}] - \frac{1}{2}\{\partial\lambda^{(\infty,1)}, \partial A_\mu\} - \frac{1}{2}\{\partial\lambda^\theta, \partial A_\mu^{(\infty,1)}\} - \frac{i}{8}[\partial\partial\lambda^{\theta\theta}, \partial\partial A_\mu] \end{aligned} \quad (1.196)$$

a déjà été résolue de différentes manières [31, 33, 34]. Mais aucune des solutions proposées ne généralise les solutions abéliennes que nous avons développées ici. A cette fin, je voudrais proposer une solution qui se rapproche de la forme attendue, même si elle n'éclaire pas encore l'éventuelle structure récursive :

$$\begin{aligned} \lambda^{(\infty,2)} &= -\frac{i}{16}[\partial A^{\theta\theta}, \partial\partial\lambda] \\ &\quad + \frac{1}{24}\left[\frac{1}{2}\{A^\theta, \{\partial A_i^\theta, \partial_i\lambda\}\} + \{A^\theta, \{A^\theta, \partial\partial\lambda\}\} + \{\{A^\theta, F\}^\theta, \partial\lambda\}\right] \\ &\quad + \frac{i}{96}[\{\{A^\theta, [A, A_i]\}^\theta, \partial_i\lambda\} + \{\{A^\theta, [A, \partial_i\lambda]\}^\theta, A_i\}] \\ A_\mu^{(\infty,2)} &= -\frac{i}{16}[\partial A^{\theta\theta}, \partial\partial A_\mu] - \frac{i}{16}[\partial_i A^{\theta\theta}, \partial F_{i\mu}] \\ &\quad + \frac{1}{24}\left[\frac{1}{2}\{A^\theta, \{\partial A_i^\theta, \partial_i A_\mu\}\} + \{A^\theta, \{A^\theta, \partial\partial A_\mu\}\} + \{\{A^\theta, F\}^\theta, \partial A_\mu\}\right] \\ &\quad + \frac{1}{24}\left[\{A^\theta, \{\partial A_i^\theta, F_{i\mu}\}\} + 2\{A^\theta, \{A_i^\theta, \partial F_{i\mu}\}\} + 3\{A^\theta, \{F^\theta, F_\mu\}\}\right] \\ &\quad + \frac{1}{24}\left[\{\{A^\theta, \partial A_i\}^\theta, F_{i\mu}\} + \{\{A^\theta, F\}^\theta, F_\mu\}\right] \\ &\quad + \frac{i}{96}[\{\{A^\theta, [A, A_i]\}^\theta, \partial_i A_\mu\} + \{\{A^\theta, [A, \partial_i A_\mu]\}^\theta, A_i\}] \\ &\quad + \frac{i}{96}[-\{\{A^\theta, [A, A_i]\}^\theta, F_{i\mu}\} + 3\{\{A^\theta, [A, F_{i\mu}]\}^\theta, A_i\}] \end{aligned} \quad (1.197)$$

Quant à la courbure, on aura pu remarquer jusqu'ici qu'elle se compose de deux types de termes : ceux qui font explicitement apparaître $F_{\mu\nu}$ (les deux indices sont sur le même objet), et ceux pour lesquels μ et ν sont séparés. L'expression $F_{\mu\nu}^{(2,2)}$ correspondant à la solution précédente présente la moitié de la structure récursive :

$$F_{\mu\nu}^{(2,2)} = -\frac{1}{2}\{F_\mu^{(1,1)}, F_\nu\} + \dots \quad (1.198)$$

où les termes restants sont tous du premier type (avec $F_{\mu\nu}$ explicite).

1.8.2 Le cas d'une pure jauge

Une pure jauge en physique désigne un potentiel de jauge associé à une connexion plate, c'est-à-dire dont la courbure F est nulle. Dans ce cas, la solution à l'équation de Seiberg-Witten est triviale : il s'agit de l'orbite de la pure jauge non-commutative. En effet l'un des représentants de l'orbite de la pure jauge commutative est $A = 0$, qui a pour image $\hat{A} = 0$, et les deux orbites sont donc liées. La question que nous souhaitons soulever ici est de trouver une paramétrisation de ces orbites, qui puisse s'appliquer sans distinction à tous les cas : abélien ou non-abélien, commutatif ou non-commutatif.

Dans le cas abélien, cette orbite est évidemment paramétrée par

$$A_\rho = \partial_\rho \alpha$$

et la transformation de jauge de paramètre λ agit comme

$$\delta \alpha = \lambda$$

Nous pouvons alors réécrire la solution à l'ordre $(1, \infty)$,

$$\begin{aligned} \lambda^{(1, \infty)} &= -\frac{1}{2} \partial \alpha \theta \partial \lambda + \frac{1}{48} \partial \partial \partial \alpha \theta \theta \theta \partial \partial \partial \lambda + \dots = \frac{i}{2} [\alpha * \lambda] \\ A_\mu^{(1, \infty)} &= -\frac{1}{2} \partial \alpha \theta \partial A_\mu + \frac{1}{48} \partial \partial \partial \alpha \theta \theta \theta \partial \partial \partial A_\mu + \dots = \frac{i}{2} [\alpha * A_\mu] \end{aligned}$$

Si nous considérons les seuls termes semi-classiques, les expressions se simplifient considérablement, puisque $F = 0$,

$$\lambda^{(n, n)} = \frac{1}{(n+1)!} (A \theta \partial)^n \lambda = \frac{1}{n+1} (A \theta \partial) \lambda^{(n-1, n-1)}$$

et de même pour A . Mais $A \theta \partial$ s'écrit aussi $\partial \alpha \theta \partial$, et est le premier terme du développement du star-commutateur $[\alpha * \cdot]$. Faisons alors l'hypothèse de récurrence que, pour $k \leq n$,

$$\begin{aligned} A_\mu^{(k)} &= \frac{i}{k+1} [\alpha * A_\mu^{(k-1)}] \\ \lambda^{(k)} &= \frac{i}{k+1} [\alpha * \lambda^{(k-1)}] \end{aligned}$$

sont solutions de l'équation de Seiberg-Witten correspondante. En espérant ne pas créer d'ambiguïtés, nous sous-entendons ici le symbole ∞ dans la notation (n, ∞) , ainsi,

$$A_\mu^{(k)} \equiv A_\mu^{(k, \infty)}, \quad \lambda^{(k)} \equiv \lambda^{(k, \infty)}$$

Nous avons vu plus haut que l'hypothèse de récurrence était vraie pour $n = 1$, en notant bien sûr $A_\mu^{(0)} = A_\mu$ et $\lambda^{(0)} = \lambda$. Pour démontrer qu'elle est vraie au rang $n + 1$, où l'équation de Seiberg-Witten s'écrit

$$\delta A_\mu^{(n+1)} - \partial_\mu \lambda^{(n+1)} = \sum_{p=0}^n i [\lambda^{(p)} * A_\mu^{(n-p)}] \quad (1.199)$$

il nous faut vérifier que notre formule de récurrence est bien solution. Nous avons,

$$\begin{aligned}
& \delta A_\mu^{(n+1)} - \partial_\mu \lambda^{(n+1)} \\
&= \frac{i}{n+2} [\delta \alpha \star A_\mu^{(n)}] - \frac{i}{n+2} [\partial_\mu \alpha \star \lambda^{(n)}] + \frac{i}{n+2} [\alpha \star \delta A_\mu^{(n)} - \partial_\mu \lambda^{(n)}] \\
&= \frac{i}{n+2} [\lambda \star A_\mu^{(n)}] + \frac{i}{n+2} [\lambda^{(n)} \star A_\mu] + \frac{i}{n+2} [\alpha \star \sum_{p=0}^{n-1} i[\lambda^{(p)} \star A_\mu^{(n-1-p)}]] \\
&= \frac{i}{n+2} [\lambda \star A_\mu^{(n)}] + \frac{i}{n+2} [\lambda^{(n)} \star A_\mu] \\
&\quad + \frac{i}{n+2} \sum_{p=0}^{n-1} ([\lambda^{(p)} \star i[\alpha \star A_\mu^{(n-1-p)}]] + [i[\alpha \star \lambda^{(p)}] \star A_\mu^{(n-1-p)}]) \\
&= \frac{i}{n+2} [\lambda \star A_\mu^{(n)}] + \frac{i}{n+2} [\lambda^{(n)} \star A_\mu] + \frac{i}{n+2} \sum_{p=0}^{n-1} (n+1-p)[\lambda^{(p)} \star A_\mu^{(n-p)}] \\
&\quad + \frac{i}{n+2} \sum_{p=0}^{n-1} (p+2)[\lambda^{(p+1)} \star A_\mu^{(n-1-p)}] \\
&= [\lambda \star A_\mu^{(n)}] + [\lambda^{(n)} \star A_\mu] + \frac{i}{n+2} \sum_{p=1}^{n-1} (n+1-p)[\lambda^{(p)} \star A_\mu^{(n-p)}] \\
&\quad + \frac{i}{n+2} \sum_{p=1}^{n-1} (p+1)[\lambda^{(p)} \star A_\mu^{(n-p)}] \\
&= \sum_{p=0}^n i[\lambda^{(p)} \star A_\mu^{(n-p)}]
\end{aligned}$$

ce qui démontre la récurrence. Ainsi, la solution s'écrit, dans le cas particulier d'une pure jauge,

$$\begin{aligned}
A_\mu^{(n+1)} &= \frac{i}{n+2} [\alpha \star A_\mu^{(n)}] \\
\lambda^{(n+1)} &= \frac{i}{n+2} [\alpha \star \lambda^{(n)}]
\end{aligned} \tag{1.200}$$

Nous pouvons nous convaincre que cette expression est bien une pure jauge non-commutative en vérifiant que la courbure $\hat{F}$ est bien nulle. Encore une fois, récur-

vement en puissance de A ,

$$\begin{aligned}
F_{\mu\nu}^{(0)} &\equiv F_{\mu\nu} = 0 \\
F_{\mu\nu}^{(1)} &\equiv \partial_\mu A_\nu^{(1)} - \partial_\nu A_\mu^{(1)} - i[A_\mu, A_\nu] \\
&= \frac{i}{2}[A_\mu, A_\nu] + \frac{i}{2}[\alpha, \partial_\mu A_\nu] - \frac{i}{2}[A_\nu, A_\mu] - \frac{i}{2}[\alpha, \partial_\nu A_\mu] - i[A_\mu, A_\nu] \\
&= \frac{i}{2}[\alpha, F_{\mu\nu}] = 0 \\
F_{\mu\nu}^{(n+1)} &\equiv \partial_\mu A_\nu^{(n+1)} - \partial_\nu A_\mu^{(n+1)} - \sum_{k=0}^n i[A_\mu^{(k)}, A_\nu^{(n-k)}] \\
&= \frac{i}{n+2}[A_\mu, A_\nu^{(n)}] - \frac{i}{n+2}[A_\nu, A_\mu^{(n)}] + \frac{i}{n+2}[\alpha, \partial_\mu A_\nu^{(n)} - \partial_\nu A_\mu^{(n)}] \\
&\quad - \sum_{k=0}^n i[A_\mu^{(k)}, A_\nu^{(n-k)}]
\end{aligned}$$

or, par définition de $F_{\mu\nu}^{(n)}$,

$$\partial_\mu A_\nu^{(n)} - \partial_\nu A_\mu^{(n)} = F_{\mu\nu}^{(n)} + \sum_{k=0}^{n-1} i[A_\mu^{(k)}, A_\nu^{(n-1-k)}]$$

et grâce à l'identité de Jacobi

$$\begin{aligned}
[\alpha, i[A_\mu^{(k)}, A_\nu^{(n-1-k)}]] &= [i[\alpha, A_\mu^{(k)}], A_\nu^{(n-1-k)}] + [A_\mu^{(k)}, i[\alpha, A_\nu^{(n-1-k)}]] \\
&= (k+2)[A_\mu^{(k+1)}, A_\nu^{(n-1-k)}] + (n-k+1)[A_\mu^{(k)}, A_\nu^{(n-k)}]
\end{aligned}$$

la plupart des termes s'éliminent dans $F_{\mu\nu}^{(n+1)}$, laissant seulement

$$F_{\mu\nu}^{(n+1)} = \frac{i}{n+2}[\alpha, F_{\mu\nu}^{(n)}] = 0 \quad (1.201)$$

Cela nous permet donc de paramétrer l'orbite de la pure jauge non-commutative. Pour faire le lien avec la section précédente, on peut remarquer qu'avec $F = 0$ les structures de Poisson θ , θ' et θ_t sont confondues, de même que les star-produits correspondants. De plus, on pourra se convaincre aisément que l'opérateur $\mathbf{a}_*$ a alors pour expression

$$\mathbf{a}_* = i[\alpha, \cdot] \quad (1.202)$$

ce qui permet de resommer les séries $\hat{A}$ et $\hat{\lambda}$ sous la forme

$$\hat{A}_\mu = \sum_{n=0}^{\infty} \frac{\mathbf{a}_*^n}{(n+1)!} A_\mu = \frac{e^{\mathbf{a}_*} - 1}{\mathbf{a}_*} A_\mu, \quad \hat{\lambda} = \frac{e^{\mathbf{a}_*} - 1}{\mathbf{a}_*} \lambda \quad (1.203)$$

qui sont exactement les formules (1.183) et (1.184) dans ce cas.

Mais on peut faire mieux, et reconnaître le développement de l'expression d'une pure jauge non-abélienne sous la forme infinitésimale :

$$A_\mu = ig^{-1}dg = \partial_\mu \alpha + \frac{1}{2}[i\alpha, \partial_\mu \alpha] + \frac{1}{6}[i\alpha, [i\alpha, \partial_\mu \alpha]] + \dots \quad (1.204)$$

lorsque g est un élément du groupe de jauge que l'on écrit $g = e^{-i\alpha}$. De même,

$$\lambda = ig^{-1}\delta g = \delta\alpha + \frac{1}{2} [i\alpha; \delta\alpha] + \frac{1}{6} [i\alpha; [i\alpha; \delta\alpha]] + \dots \quad (1.205)$$

Notons que cette deuxième relation n'est rien d'autre que la transformation de jauge de g , comme on le voit bien dans le cas abélien

$$\delta g = -ie^{-i\alpha}\delta\alpha = -ig\lambda, \quad \text{et donc } \lambda = ig^{-1}\delta g$$

Ces expressions sont la généralisation non-abélienne de $A_\mu = \partial_\mu\alpha$ et $\lambda = \delta\alpha$. En plus de ces développements explicites, nous avons en effet montré qu'elles étaient compatibles avec $F = 0$ et la transformation de jauge non-abélienne. Ou plutôt, nous l'avons fait dans le cas non-commutatif, mais le calcul est algébriquement le même. Cela nous permet ainsi d'étendre notre paramétrisation de l'orbite des pures jagues au cas général, non-abélien et non-commutatif, sous la forme

$$\begin{aligned} \hat{A} &= i\hat{g}^{-1} \star d\hat{g} = d\alpha + \frac{1}{2} [i\alpha \star d\alpha] + \frac{1}{6} [i\alpha \star [i\alpha \star d\alpha]] + \dots \\ \hat{\lambda} &= i\hat{g}^{-1} \star \delta\hat{g} = \delta\alpha + \frac{1}{2} [i\alpha \star \delta\alpha] + \frac{1}{6} [i\alpha \star [i\alpha \star \delta\alpha]] + \dots \end{aligned} \quad (1.206)$$

où $\hat{g} = (\star e)^{-i\alpha} = 1 - i\alpha - \frac{\alpha \star \alpha}{2} + \dots$ est toujours construit à partir de l'élément de l'algèbre de jauge α , mais en utilisant cette fois des star-produits. Il semble donc que cela nous donne en même temps une paramétrisation du groupe de jauge non-commutatif en fonction de l'algèbre de jauge (en tout cas de la composante connexe de l'identité, sans présager d'éventuelles déformations plus globales).

Chapitre 2

M5-branes non-abéliennes

Objectif Nous souhaitons trouver un contenu en champs décrivant la théorie de basse énergie d'un système de N M5-branes coïncidentes. Une telle théorie devra reproduire l'anomalie calculée dans [42]. Elle formera aussi une généralisation non-abélienne des théories de jauge de 2-formes.

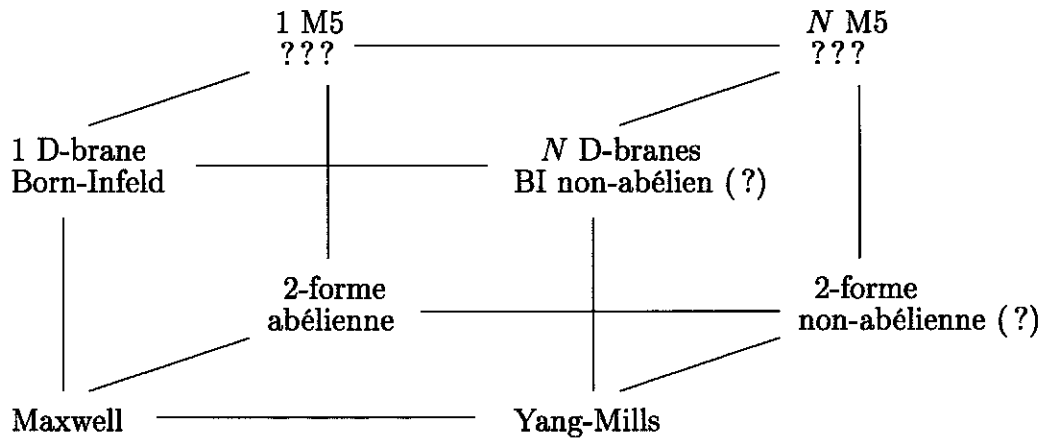


FIG. 2.1 – Les théories de jauge de 2-formes non-abéliennes sont des généralisations des théories de Yang-Mills. De même, l'action décrivant un paquet de M5-branes généralise l'action de Born-Infeld non-abélienne, déjà mal connue.

Plan du chapitre Nous allons décrire les différentes méthodes envisagées pour trouver les champs décrivant un paquet de M5-branes. Tout d'abord, nous allons tenter de reproduire l'anomalie attendue (cf [42]) à partir des multiplets supersymétriques de type (2,0) en dimension 6. En effet, dans le cas d'une M5-brane unique, les champs de basse énergie forment un multiplet tensoriel de cette supersymétrie de type (2,0).

D'autre part, on s'attend à ce que le secteur de jauge du paquet de M5-branes décrive une théorie de 2-formes non-abéliennes. En effet, dans les théories de cordes, la théorie de jauge correspondant à un paquet de D-branes est la version non-abélienne

de la théorie de jauge sur une brane isolée. Or, le secteur de jauge d'une M5-brane isolée est un champ tensoriel B_2 d'ordre 2, qui est de plus anti-autodual¹, et dont on ne connaît pas de version non-abélienne (voir fig. 2.1).

Enfin, nous essaierons de trouver quelle image géométrique pourrait constituer une réalisation intuitive de ces états, de manière similaire aux cordes ouvertes tendues entre les différentes branes en théories des cordes. En effet, dans le cas des D-branes, la situation est bien comprise : chaque D-brane porte un champ de jauge $U(1)$, qui est le mode de masse nulle des cordes ouvertes qui sont attachées à cette D-brane. Lorsque N D-branes coïncident, la symétrie $U(1)^N$ est élargie en $U(N)$, correspondant à l'élargissement du spectre de masse nulle. La masse (des états de basse énergie) des cordes ouvertes reliant deux branes différentes dépend de la distance entre ces branes, et s'annule lorsque les branes deviennent coïncidentes. Au total, ces cordes ouvertes décrivent donc les N^2 degrés de liberté du paquet de N D-branes. Nous essaierons d'adapter un tel comptage pour un paquet de N M5-branes, en remplaçant bien sûr les cordes ouvertes par des membranes M2.

2.1 Les multiplets supersymétriques

2.1.1 Anomalies et polynômes de Chern

A 6 dimensions, un champ chiral, un vecteur chiral et un tenseur antisymétrique présentent respectivement l'anomalie (cf [43])

$$\begin{aligned} I_{1/2} &= 1 - \frac{1}{24} p_1(TW) + \frac{7p_1(TW)^2 - 4p_2(TW)}{5760} \\ I_{3/2} &= 5 + \frac{23}{24} p_1(TW) + \frac{55}{24 \times 48} p_1(TW)^2 - \frac{49}{6 \times 48} p_2(TW) \\ I_B &= -\frac{1}{8} - \frac{1}{24} p_1(TW) + \frac{p_1(TW)^2 - 7p_2(TW)}{360} \end{aligned}$$

D'autre part, lorsque le fibré normal N a $SO(5)$ comme groupe de structure,

$$\begin{aligned} \text{ch}(N) &= 5 + p_1(N) + \frac{1}{12} p_1(N)^2 - \frac{1}{6} p_2(N) \\ \text{ch}(S(N)) &= 4 + \frac{1}{2} p_1(N) + \frac{1}{96} p_1(N)^2 + \frac{1}{24} p_2(N) \end{aligned}$$

En effet, de manière générale,

$$\text{ch}(N) = \sum_i e^{\lambda_i}$$

où les λ_i sont les racines de Chern de N . Si de plus p_1 et p_2 désignent les deux premières classes de Pontryagin de N ,

$$p_1(N) = \sum_i \lambda_i^2, \quad p_2(N) = \sum_{i < j} \lambda_i^2 \lambda_j^2$$

on peut exprimer $\text{ch}(N)$ en fonction de $p_1(N)$ et $p_2(N)$:

¹ Sa courbure est la 3-forme $H_3 = dB_2$. Elle est anti-autoduale au sens où $*H_3 = -H_3$, en désignant par $*$ l'étoile de Hodge sur le volume d'univers à six dimensions de la M5-brane.

– pour $\text{SO}(2n)$, les racines sont $\pm\lambda_1, \dots, \pm\lambda_n$, et

$$\begin{aligned}\text{ch}(\text{SO}(2n)) &= 2n + (\lambda_1^2 + \dots + \lambda_n^2) + \frac{1}{12}(\lambda_1^4 + \dots + \lambda_n^4) + \dots \\ &= 2n + p_1 + \frac{1}{12}p_1^2 - \frac{1}{6}p_2 + \dots\end{aligned}$$

– pour $\text{SO}(2n+1)$, les racines sont $\pm\lambda_1, \dots, \pm\lambda_n, 0$, et le même calcul donne donc

$$\text{ch}(\text{SO}(2n+1)) = 2n+1 + p_1 + \frac{1}{12}p_1^2 - \frac{1}{6}p_2 + \dots$$

Ainsi, l'expression est la même pour $\text{SO}(d)$, que d soit pair ou impair. D'autre part, calculons ce même polynôme pour la représentation spinorielle de $\text{SO}(5)$:

$\text{SO}(5)$ a pour racine $\pm\lambda_1, \pm\lambda_2, 0$, ce qui donne pour $\text{S}(\text{SO}(5))$ les racines $\pm\frac{\lambda_1 \pm \lambda_2}{2}$, et donc

$$\begin{aligned}\text{chS}(\text{SO}(5)) &= \sum_{\epsilon_1, \epsilon_2 = \pm 1} e^{\frac{\epsilon_1 \lambda_1 + \epsilon_2 \lambda_2}{2}} \\ &= 4 + \frac{1}{2}(\lambda_1^2 + \lambda_2^2) + \frac{1}{2 \times 96}[(\lambda_1 + \lambda_2)^4 + (\lambda_1 - \lambda_2)^4] \\ &= 4 + \frac{1}{2}p_1 + \frac{1}{96}p_1^2 + \frac{1}{24}p_2\end{aligned}$$

2.1.2 Les multiplets supersymétriques (2,0) en dimension 6

Les modes de masse nulle d'une brane M5 constituent un multiplet tensoriel d'une théorie supersymétrique de type (2,0) à six dimensions. Considérons donc les autres multiplets de type (2,0) à six dimensions [44] :

$$\begin{aligned}M_1 &= (3, 1; 1) + (1, 1; 5) + (2, 1; 4) \\ M_2 &= (3, 2; 1) + (1, 2; 5) + (2, 2; 4) \\ M_3 &= (3, 3; 1) + (1, 3; 5) + (2, 3; 4)\end{aligned}$$

Les triplets désignent une représentation² de $\text{SO}(4) \times \text{USp}(4)$. La représentation de $\text{SO}(4)$ est donnée par le premier doublet, sous la forme $\text{SU}(2) \times \text{SU}(2)$. Le troisième nombre est donc la représentation sous l'action du groupe de jauge $\text{SO}(5) = \text{USp}(4)$ du fibré normal N .

Le premier multiplet est le multiplet tensoriel, composé d'un tenseur antisymétrique (3,1) singlet de $\text{SO}(5)$, de vecteurs non chiraux (1,1) dans la représentation spinorielle (4) de $\text{SO}(5)$, et de spineurs (2,1) dans la représentation vectorielle (5) de $\text{SO}(5)$. On peut de même analyser les autres multiplets. Ils comprennent entre autres les champs de spin 3/2 (3,2;1) et (2,3;4), et un graviton (3,3;1).

On peut alors calculer l'anomalie totale des trois multiplets.

² On rappelle que les champs de masse nulle en dimension d (ou plus exactement $(d-1, 1)$, avec une dimension de genre temps) sont des représentations de $\text{SO}(d-2)$ — et les champs massifs des représentations de $\text{SO}(d-1)$. Le volume d'univers de la M5-brane compte 6 dimensions, dont la dimension temporelle, d'où le groupe $\text{SO}(4)$. D'autre part, il reste au fibré normal les 5 dimensions transverses dans l'espace total à 11 dimensions, d'où $\text{SO}(5)$ pour le groupe de jauge transverse.

– $M_1 = (3, 1; 1) + (1, 1; 5) + (2, 1; 4)$ a pour anomalie

$$I_B + \frac{1}{2} I_{1/2} \text{ch}(S(N)) = \frac{1}{48} \left(\frac{p_1(TW)^2}{4} - p_2(TW) \right) - \frac{1}{2 \times 48} p_1(TW) p_1(N) \\ + \frac{1}{4 \times 48} p_1(N)^2 + \frac{1}{48} p_2(N)$$

– $M_2 = (3, 2; 1) + (1, 2; 5) + (2, 2; 4)$ a pour anomalie

$$\frac{1}{2} I_{3/2} - \frac{1}{2} I_{1/2} \text{ch}(N) = \frac{4}{48} \left(\frac{p_1(TW)^2}{4} - p_2(TW) \right) + \frac{1}{48} p_1(TW) p_1(N) \\ - \frac{2}{48} p_1(N)^2 + \frac{4}{48} p_2(N)$$

– $M_3 = (3, 3; 1) + (1, 3; 5) + (2, 3; 4)$ a pour anomalie (les coefficients m_i permettent de prendre en compte un polynôme de couplage différent de $\text{ch}(N)$, puisqu'il n'est pas clair qu'un tenseur antisymétrique soit couplé ainsi)

$$-I_B * (5 + m_1 p_1(N) + m_2 p_1(N)^2 + m_3 p_2(N)) - \frac{1}{2} I_{3/2} \text{ch}(S(N)) \\ = -\frac{21}{48} \left(\frac{p_1(TW)^2}{4} - p_2(TW) \right) - \frac{1}{2 \times 48} (23 - 4m_1) p_1(TW) p_1(N) \\ - \frac{1}{4 \times 48} (5 - 24m_2) p_1(N)^2 - \frac{1}{48} (5 - 6m_3) p_2(N)$$

2.1.3 Peut-on reproduire l'anomalie attendue ?

L'anomalie attendue pour un paquet de Q M5-branes est (cf [42])

$$\mathcal{A} = \frac{Q}{48} \left(\frac{p_1(TW)^2}{4} - p_2(TW) \right) - \frac{Q}{2 \times 48} p_1(TW) p_1(N) \\ + \frac{Q}{4 \times 48} p_1(N)^2 + \frac{2Q^3 - Q}{48} p_2(N)$$

On veut reproduire cette anomalie avec n_1 copies du multiplet M_1 , n_2 copies du multiplet M_2 et n_3 copies du multiplet M_3 . Cela donne le système d'équations

$$\begin{cases} n_1 + 4n_2 - 21n_3 = Q & (L_1) \\ n_1 - 2n_2 + (23 - 4m_1)n_3 = Q & (L_2) \\ n_1 - 8n_2 - (5 - 24m_2)n_3 = Q & (L_3) \\ n_1 + 4n_2 - (5 - 6m_3)n_3 = 2Q^3 - Q & (L_4) \end{cases}$$

Le sous-système (L_1, L_2, L_3) est équivalent à

$$(9 - m_1 - 3m_2)n_3 = 0$$

Il est donc dégénéré si et seulement si $m_1 + 3m_2 = 9$. Si ce n'est pas le cas, il a pour unique solution $n_1 = Q$, $n_2 = n_3 = 0$, qui est incompatible avec L_4 . Supposons donc

qu'il soit dégénéré. Le système complet est alors équivalent à

$$\begin{aligned}(8 + 3m_3)n_3 &= Q^3 - Q \\ n_2 &= \frac{2}{3}(11 - m_1)n_3 \\ n_1 &= Q + \frac{1}{3}(8m_1 - 25)n_3\end{aligned}$$

Ainsi, non seulement ce système n'a pas de solutions pour toute valeur de Q , mais selon les paramètres m_i , il arrive même qu'il n'en ait aucune. Il n'est donc pas clair que cette approche permette de reproduire l'anomalie attendue, en tout cas sans subtilités supplémentaires. De plus, même si cette méthode fournit un résultat, elle ne peut rien dire sur la structuration des différents multiplets. Par exemple au sein d'une théorie de jauge non abélienne de 2-formes pour le secteur bosonique, que nous allons étudier maintenant.

2.2 Théories de jauge non-abéliennes de 2-formes

2.2.1 Complexes de Hochschild et algèbres de Gerstenhaber

Considérons tout d'abord l'ensemble $\text{Hom}(A^{\otimes n}, A)$ des homomorphismes (au sens d'applications n -linéaires) d'une algèbre A . La réunion de ces espaces forme le complexe de Hochschild de A . Comme nous allons le voir plus précisément les opérations algébriques sur A , ainsi que les différentielles, appartiennent à ce complexe, puisqu'elles sont toutes multilinéaires. En outre, les déformations de la structure algébrique sont elles aussi naturellement des éléments du complexe de Hochschild.

Commençons par définir l'action d'une application n -linéaire ϕ sur un m -uplet d'élément de A , c'est-à-dire³ un élément de $A^{\otimes m}$:

$$\phi(a_1, \dots, a_m) = \sum_{k=1}^{m-n+1} (-1)^{(k-1)(n-1)} (a_1, \dots, \phi(a_k, \dots, a_{k+n-1}), \dots, a_m) \quad (2.1)$$

Cette action permet aux applications n -linéaires d'agir non plus seulement sur $A^{\otimes n}$, par leur action naturelle, mais sur toute l'algèbre tensorielle $\mathcal{T}A = \bigoplus_k A^{\otimes k}$. Cela est nécessaire pour pouvoir manipuler les homomorphismes au sein du complexe de Hochschild, par exemple pour les additionner. Ou pour composer deux homomorphismes. En effet, pour $\phi_1 \in \text{Hom}(A^{\otimes n_1}, A)$ et $\phi_2 \in \text{Hom}(A^{\otimes n_2}, A)$, la composée $\phi_1 \circ \phi_2 \in \text{Hom}(A^{\otimes n_1+n_2-1}, A)$ est naturellement définie par

$$\begin{aligned}\phi_1 \circ \phi_2(a_1, \dots, a_{n_1+n_2-1}) \\ = \sum_{k=1}^{n_1} (-1)^{(k-1)(n_2-1)} \phi_1(a_1, \dots, \phi_2(a_k, \dots, a_{k+n_2-1}), \dots, a_{n_1+n_2-1})\end{aligned} \quad (2.2)$$

³ On se permettra de confondre les applications multilinéaires et leur application linéaire associée sur l'espace tensoriel : à ϕ m -linéaire sur A est associée de manière universelle $\tilde{\phi}$, application linéaire de $A^{\otimes m}$ dans A , telle que,

$$\phi(a_1, \dots, a_m) = \tilde{\phi}(a_1 \otimes \dots \otimes a_m)$$

ce qui n'est rien d'autre que l'action de ϕ_2 sur l'élément de $A^{\otimes(n_1+n_2-1)}$, dont le résultat est un élément de $A^{\otimes n_1}$, suivie de l'action naturelle de ϕ_1 . Cela fait de cette loi, comme toute loi de composition, une loi associative. De plus, son commutateur est alors exactement le crochet de Gerstenhaber, que nous avons déjà rencontré au chapitre précédent — équation (1.78) — et s'écrit

$$[\phi_1, \phi_2]_G = \phi_1 \circ \phi_2 - (-1)^{(n_1-1)(n_2-1)} \phi_2 \circ \phi_1, \quad \phi_i \in \text{Hom}(A^{\otimes n_i}, A) \quad (2.3)$$

Voyons maintenant les objets qui vivent dans ce complexe, et plus particulièrement deux exemples : la différentielle et le produit.

- Une différentielle sur A est un endomorphisme Q de carré nul, c'est-à-dire

$$Q \circ Q = 0$$

Bien que tentés, il faut nous retenir de l'écrire en termes du crochet $[Q, Q]_G$, qui est par définition identiquement nul pour tout opérateur d'ordre 1.

- Un produit est une application bilinéaire $m \in \text{Hom}(A^{\otimes 2}, A)$. Son associativité est exprimée par la condition

$$m \circ m = \frac{1}{2} [m, m]_G = 0$$

En effet, pour $a_1, a_2, a_3 \in A$, en notant $m(a, b) = a \star b$:

$$\begin{aligned} (m \circ m)(a_1, a_2, a_3) &= m(m(a_1, a_2), a_3) - m(a_1, m(a_2, a_3)) \\ &= (a_1 \star a_2) \star a_3 - a_1 \star (a_2 \star a_3) \end{aligned}$$

- Le produit m vérifie la règle de Leibniz par rapport à la différentielle Q — ou plus simplement Q dérive m — si et seulement si

$$[Q, m]_G = Q \circ m - m \circ Q = 0$$

En effet, en notant toujours $m(a, b) = a \star b$:

$$(Q \circ m)(a, b) = Q(a \star b), \quad (m \circ Q)(a, b) = Q(a) \star b + a \star Q(b)$$

Il est intéressant de noter que l'on peut regrouper ces trois conditions en une seule : m un produit associatif et Q est une différentielle dérivant m si et seulement si

$$(Q + m) \circ (Q - m) = 0 \quad (2.4)$$

En effet, cette relation se décompose en opérateurs de différents ordres, pour redonner les trois expressions voulues. Cette structure définit une **algèbre associative différentielle**. On peut être un peu plus général, et utiliser une infinité d'applications $m_n \in \text{Hom}(A^{\otimes n}, A)$, en généralisant l'objet $Q + m$ que nous venons de construire en $m = \sum m_n$. D'autre part, définissons l'involution

$$\bar{\phi} = (-1)^{n-1} \phi, \quad \text{pour } \phi \in \text{Hom}(A^{\otimes n}, A) \quad (2.5)$$

autrement dit, cette involution inverse le signe des applications d'ordre pair. Ainsi,

$$m = m_1 + m_2 + m_3 + \dots, \quad \bar{m} = \bar{m}_1 - \bar{m}_2 + \bar{m}_3 - \dots \quad (2.6)$$

On appelle **algèbre** A_∞ , ou **algèbre associative à homotopie près**, une algèbre munie d'un tel opérateur m vérifiant

$$m \circ \bar{m} = 0 \quad (2.7)$$

Notons qu'une algèbre différentielle associative est donc une algèbre A_∞ pour laquelle $m_n = 0$ pour $n \geq 3$. Mais dans le cas le plus général, le produit m_2 n'est pas associatif, la déviation étant due à m_3 , ce qu'indique le nom de l'algèbre. En effet, à l'ordre 3, on trouve

$$m_2 \circ m_2 = m_3 \circ m_1 + m_1 \circ m_3 \quad (2.8)$$

où comme précédemment, $m_2 \circ m_2$ est l'associateur du produit m_2 .

En espérant ne pas causer de confusion, revenons à un produit bilinéaire simple, que nous noterons à nouveau m . Il nous permet de définir deux nouveaux objets sur le complexe de Hochschild. Tout d'abord, l'opérateur de cobord, dit différentielle de Hochschild

$$\delta = [m, \cdot]_G : \text{Hom}(A^{\otimes n}, A) \longrightarrow \text{Hom}(A^{\otimes n+1}, A) \quad (2.9)$$

En toute rigueur, δ est une différentielle si et seulement si m est associatif :

$$[m, m]_G = 0 \quad \Leftrightarrow \quad m \text{ est associative} \quad \Leftrightarrow \quad \delta^2 = 0 \quad (2.10)$$

D'autre part, on définit le cup-produit $\cup_m$ de ϕ_1 et ϕ_2 par

$$(\phi_1 \cup_m \phi_2)(a_1, \dots, a_{n_1+n_2}) = m(\phi_1(a_1, \dots, a_{n_1}), \phi_2(a_{n_1+1}, \dots, a_{n_1+n_2})) \quad (2.11)$$

Muni du crochet, de la différentielle de Hochschild et du cup-produit, le complexe de Hochschild est une **algèbre de Gerstenhaber différentielle**.

2.2.2 Théorie de jauge de 2-formes

Les champs que nous allons introduire sont à la fois des formes différentielles sur la variété M et des éléments du complexe de Hochschild d'une algèbre $\mathfrak{g}$ sur $\mathbb{R}$ (ou $\mathbb{C}$). Ils ont donc un double indice : un degré différentiel d en tant que forme de $\bigwedge^d T^*M$ et un degré de Hochschild p en tant qu'application linéaire de $\text{Hom}(\mathfrak{g}^{\otimes p}, \mathfrak{g})$. Nous noterons

$$\Omega^d(M, \text{Hom}(\mathfrak{g}^{\otimes p}, \mathfrak{g})) = \bigwedge^d T^*M \otimes \text{Hom}(\mathfrak{g}^{\otimes p}, \mathfrak{g})$$

les sous-espaces de ce complexe bigradué Ω^* . En particulier, notons tout de suite que

$$\text{Hom}(\mathfrak{g}, \mathfrak{g}) = \text{End}(\mathfrak{g}) \quad \text{pour } p = 1, \quad \text{Hom}(\mathbb{R}, \mathfrak{g}) = \mathfrak{g} \quad \text{pour } p = 0$$

Nous appellerons degré (ou degré total) la somme $p + d$. Cela signifie aussi que par rapport aux formules de la section précédente, il faut tenir compte du degré différentiel dans la graduation totale. Ainsi, le crochet de Gerstenhaber entre ϕ_1 et ϕ_2 s'écrit

$$[\phi_1, \phi_2] = \phi_1 \circ \phi_2 - (-1)^{(p_1-1)(p_2-1)+d_1d_2} \phi_2 \circ \phi_1 \quad (2.12)$$

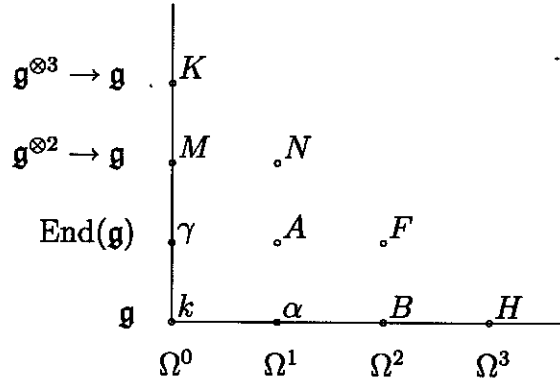


FIG. 2.2 – Les composantes de la courbure (H, F, N, K), du champ de jauge (B, A, M), du paramètre de jauge (α, γ) et du paramètre de changement de jauge (k).

La théorie de jauge en composantes

Commençons par présenter les composantes du champ de jauge et de la courbure, comme elles apparaissent dans [47]. Nous chercherons ensuite une manière plus compacte d'écrire la théorie. Comme on peut le voir sur la figure 2.2, le potentiel de jauge se compose en réalité des trois champs B, A, M :

- B est une 2-forme à valeurs dans $\mathfrak{g}$,
- A est une 1-forme à valeurs dans les endomorphismes de $\mathfrak{g}$,
- M est une 0-forme à valeurs dans les applications bilinéaires sur $\mathfrak{g}$.

A eux trois, ils forment un triplet de degré total 2. La courbure est de même un quadruplet, de degré total 3, avec pour composantes H, F, N, K , comme indiqué sur la figure 2.2. Ces champs s'expriment en fonction du potentiel par les formules suivantes :

$$\begin{aligned}
 H &= dB + A(B) \\
 F &= dA + A^2 + [M, B] \\
 N &= dM + [A, M] \\
 K &= \frac{1}{2} [M, M]
 \end{aligned} \tag{2.13}$$

D'autre part, comme dans une théorie de jauge de 1-forme, la courbure vérifie l'identité de Bianchi, qui s'écrit en composantes :

$$\begin{aligned}
 dH + [A, H] + [B, F] &= 0 \\
 dF + [A, F] - [B, N] - [M, H] &= 0 \\
 dN + [A, N] + [B, K] + [M, F] &= 0 \\
 dK + [A, K] - [M, N] &= 0 \\
 [M, K] &= 0
 \end{aligned} \tag{2.14}$$

Les transformations de jauge dépendent d'un paramètre de degré total 1, composé des champs α et γ . Elles agissent sur le champ de jauge selon

$$\begin{aligned} B' &= \gamma^{-1}(B + d\alpha + A(\alpha) - M(\alpha, \alpha)) \\ A' &= \gamma^{-1}(d\gamma + A\gamma - [M, \alpha]\gamma) \\ M' &= \gamma^{-1}M\gamma^{\otimes 2} \end{aligned} \quad (2.15)$$

et induisent les transformations suivantes de la courbure

$$\begin{aligned} H' &= \gamma^{-1}(H + F(\alpha) - N(\alpha, \alpha) - K(\alpha, \alpha, \alpha)) \\ F' &= \gamma^{-1}(F - [N, \alpha] - \frac{1}{2} [[K, \alpha], \alpha])\gamma \\ N' &= \gamma^{-1}(N + [K, \alpha])\gamma^{\otimes 2} \\ K' &= \gamma^{-1}K\gamma^{\otimes 3} \end{aligned} \quad (2.16)$$

où l'on a noté $A(\alpha) = A \circ \alpha = [A, \alpha]$, puisque A est un endomorphisme agissant sur l'élément α de l'algèbre. La même remarque vaut pour M , F , N et K , avec par exemple $M(\alpha, \alpha) = M \circ \alpha \circ \alpha = \frac{1}{2} [[M, \alpha], \alpha]$.

Enfin, comme on le constate sur la figure 2.2, il reste de la place pour un champ scalaire à valeurs dans $\mathfrak{g}$, que nous noterons k . Il permet de transformer les paramètres précédents α et γ , dans ce que nous appellerons des *transformations de jauge de la jauge*, qui prennent la forme infinitésimale

$$\begin{aligned} \delta_k \alpha &= dk + A(k) + [[M, k], \alpha] \\ \delta_k \gamma &= [M, k]\gamma \end{aligned} \quad (2.17)$$

Ces transformations ne laissent pas le champ de jauge invariant. En effet,

$$\delta_k B = [F, k], \quad \delta_k A = [N, k], \quad \delta_k M = [K, k] \quad (2.18)$$

Après cette liste de formules, examinons la structure du groupe de jauge. Nous oublierons pour l'instant les transformations (2.17), auxquelles nous reviendrons. Appelons $U_{\gamma, \alpha}$ la transformation de jauge (2.15) de paramètres (α, γ) , et remarquons tout d'abord que

$$U_{\gamma, \alpha} = U_{\gamma, 0} U_{1, \alpha} = U_{1, \gamma^{-1} \alpha} U_{\gamma, 0} \quad (2.19)$$

Si la première relation est explicite dans les formules (2.15), la deuxième est moins évidente. Etudions-la rapidement :

$$(B, A, M) \xrightarrow{U_{\gamma, 0}} (B', A', M') \xrightarrow{U_{1, \gamma^{-1} \alpha}} (B'', A'', M'')$$

$$\begin{aligned} U_{\gamma, 0} &\begin{cases} B' = \gamma^{-1} B \\ A' = \gamma^{-1} (d\gamma + A\gamma) \\ M' = \gamma^{-1} M \gamma^{\otimes 2} \end{cases} \\ U_{1, \gamma^{-1} \alpha} &\begin{cases} B'' = B' + d(\gamma^{-1} \alpha) + A'(\gamma^{-1} \alpha) - M'(\gamma^{-1} \alpha, \gamma^{-1} \alpha) \\ A'' = A' - [M', \gamma^{-1} \alpha] \\ M'' = M' \end{cases} \end{aligned}$$

Il ne reste qu'à remplacer les expressions et vérifier que l'on retrouve bien (2.15). Pour cela, notons que

$$\begin{aligned}
 [M', \gamma^{-1}\alpha] &= (\gamma^{-1}M\gamma^{\otimes 2}) \circ (\gamma^{-1}\alpha) \\
 &= \gamma^{-1}(M \circ \alpha)\gamma \\
 &= \gamma^{-1}[M, \alpha]\gamma \\
 d(\gamma^{-1}\alpha) + A'(\gamma^{-1}\alpha) &= \gamma^{-1}d\alpha - \gamma^{-1}d\gamma\gamma^{-1}\alpha + \gamma^{-1}(d\gamma + A\gamma)\gamma^{-1}\alpha \\
 &= \gamma^{-1}d\alpha + \gamma^{-1}A(\alpha)
 \end{aligned}$$

D'autre part, on vérifie aisément que $U_{\gamma,0}U_{\gamma',0} = U_{\gamma'\gamma,0}$ — attention à l'inversion $\gamma'\gamma$ qui vient de l'ordre de la composition — et $U_{1,\alpha}U_{1,\alpha'} = U_{1,\alpha+\alpha'}$. Les premiers se comportent comme des rotations, les seconds comme des translations. On peut alors calculer le produit de deux transformations de jauge quelconques :

$$\begin{aligned}
 U_{\gamma,\alpha}U_{\gamma',\alpha'} &= U_{\gamma,0}U_{1,\alpha}U_{\gamma',0}U_{1,\alpha'} \\
 &= U_{\gamma,0}U_{\gamma',0}U_{1,\gamma'\alpha}U_{1,\alpha'} \\
 &= U_{\gamma'\gamma,0}U_{1,\gamma'\alpha+\alpha'} \\
 &= U_{\gamma'\gamma,\gamma'\alpha+\alpha'}
 \end{aligned} \tag{2.20}$$

Le groupe de jauge a donc la structure du produit semi-direct

$$\Omega^0(M, \text{End}(\mathfrak{g})) \ltimes \Omega^1(M, \mathfrak{g}) \tag{2.21}$$

La théorie de jauge sur le complexe Ω^*

Nous allons maintenant essayer d'écrire toutes ces formules d'une manière plus élégante et plus familière (voir la section 1.1). Pour cela, nous allons utiliser les objets "complets", par opposition aux composantes, dont nous avons déjà parlé, mais sans les définir. Ainsi, définissons le potentiel de jauge $\mathcal{B}$ et la courbure $\mathcal{F}$ sur le complexe par :

$$\mathcal{B} = B + A - M \tag{2.22}$$

$$\mathcal{F} = H + F - N - K \tag{2.23}$$

Ces sommes sont réellement à comprendre sur le complexe total, puisque les composantes vivent dans des sous-espaces différents. C'est d'ailleurs grâce à cela que l'on va pouvoir identifier terme à terme les expressions qui suivent avec celles que j'ai déjà présentées. D'autre part, souvenons-nous que nous avons défini une involution à la section précédente — formule (2.5) — qui inverse le signe des composantes ayant un degré de Hochschild pair. Nous disposons donc aussi des conjugués $\bar{\mathcal{B}}$ et $\bar{\mathcal{F}}$

$$\bar{\mathcal{B}} = -B + A + M \tag{2.24}$$

$$\bar{\mathcal{F}} = -H + F + N - K \tag{2.25}$$

On peut alors exprimer la courbure en fonction du potentiel par la relation

$$\mathcal{F} = d\mathcal{B} + \frac{1}{2} [\bar{\mathcal{B}}, \mathcal{B}] \tag{2.26}$$

très similaire à (1.33) pour une théorie de jauge de 1-forme. On pourra vérifier que les différentes composantes redonnent bien les expressions (2.13), en utilisant les identités suivantes sur les commutateurs :

$$[A, B] = -[B, A], \quad [M, B] = [B, M], \quad [A, M] = -[M, A]$$

On peut aussi définir la différentielle extérieure covariante d_B

$$d_B = d + [B, \cdot] \quad (2.27)$$

grâce à laquelle on réécrit l'(es) identité(s) de Bianchi (2.14) :

$$d_B \mathcal{F} = d\mathcal{F} + [B, \mathcal{F}] = 0 \quad (2.28)$$

Pour s'attaquer aux transformations de jauge, il faut connaître leur forme infinitésimale, qui devrait s'exprimer simplement grâce à d_B . Pour cela, on écrit $\gamma = 1 + u$, et on ne conserve que le premier ordre en u et α . Cela donne, en remarquant entre autres que $[B, u] = -u(B)$,

$$\begin{aligned} \delta_{u,\alpha} \mathcal{B} = & \quad d\alpha + [B, u] + [A, \alpha] \\ & + du + [A, u] - [M, \alpha] \\ & \quad - [M, u] \end{aligned}$$

ce qui s'écrit aussi en fonction du paramètre de jauge $\Lambda = u + \alpha$

$$\delta_\Lambda \mathcal{B} = d_B \Lambda = d\Lambda + [B, \Lambda] \quad (2.29)$$

Avant de continuer, il va nous falloir démontrer quelques résultats intermédiaires pour pouvoir manipuler les expressions. Je noterai ici, autant que possible, un champ sur le complexe Ω^* avec une lettre grecque majuscule, et ses composantes sur les sous-espaces avec des lettres grecques minuscules. De plus, on appellera Φ *champ homogène de degré n* si toutes ses composantes sont de degré total n :

$$\Phi = \sum_{i=0}^n \phi_{i,n-i}$$

où $\phi_{i,n-i}$ est une composante de degré de Hochschild i et de degré différentiel $n - i$. Nous noterons aussi s l'involution

$$s\Phi = -\bar{\Phi} = \sum_i (-1)^i \phi_{i,n-i} \quad (2.30)$$

Tout d'abord, la différentielle agit sur le crochet par

$$d[\Phi, \cdot] = [d\Phi, \cdot] + (-1)^n [s\Phi, d\cdot] \quad (2.31)$$

où, Φ est homogène de degré n . En effet,

$$\begin{aligned} d[\Phi, \cdot] &= \sum_i d[\phi_{i,n-i}, \cdot] \\ &= \sum_i ([d\phi_{i,n-i}, \cdot] + (-1)^{n-i} [\phi_{i,n-i}, d\cdot]) \\ &= [d\Phi, \cdot] + (-1)^n [s\Phi, d\cdot] \end{aligned}$$

En particulier, $\mathcal{B}$ étant homogène de degré 2,

$$d[\mathcal{B}, \cdot] = [d\mathcal{B}, \cdot] - [\bar{\mathcal{B}}, d\cdot] \quad (2.32)$$

Ensuite, si Φ_1 et Φ_2 sont homogènes de degré respectifs n_1 et n_2 , alors⁴

$$[\Phi_1, \Phi_2] = (-1)^{n_1 n_2} [s^{n_1-1} \Phi_2, s^{n_2-1} \Phi_1] \quad (2.33)$$

ce qui permet de permuter les crochet (presque) plus facilement que pour les composantes seules. De même, trois champs homogènes Φ_1, Φ_2, Φ_3 vérifient l'identité de Jacobi graduée suivante :

$$[\Phi_1, [\Phi_2, \Phi_3]] = [[\Phi_1, \Phi_2], \Phi_3] - (-1)^{n_1 n_2} [s^{n_1-1} \Phi_2, [s^{n_2-1} \Phi_1, \Phi_3]] \quad (2.34)$$

La démonstration est très similaire à celle de la formule précédente. Encore une fois, on en déduit en particulier :

$$[\bar{\mathcal{B}}, [\mathcal{B}, \cdot]] = \frac{1}{2} [[\bar{\mathcal{B}}, \mathcal{B}], \cdot] \quad (2.35)$$

Grâce à (2.32) et (2.35), nous pouvons maintenant démontrer

$$d_{\bar{\mathcal{B}}} d_{\mathcal{B}} = [\mathcal{F}, \cdot] \quad (2.36)$$

où $d_{\bar{\mathcal{B}}} = d + [\bar{\mathcal{B}}, \cdot]$ est la différentielle extérieure covariante associée à $\bar{\mathcal{B}}$. En effet,

$$\begin{aligned} d_{\bar{\mathcal{B}}} d_{\mathcal{B}} &= (d + [\bar{\mathcal{B}}, \cdot]) (d + [\mathcal{B}, \cdot]) \\ &= d^2 + [d\mathcal{B}, \cdot] - [\bar{\mathcal{B}}, d\cdot] + [\bar{\mathcal{B}}, d\cdot] + [\bar{\mathcal{B}}, [\mathcal{B}, \cdot]] \\ &= [d\mathcal{B}, \cdot] + \frac{1}{2} [[\bar{\mathcal{B}}, \mathcal{B}], \cdot] = [\mathcal{F}, \cdot] \end{aligned}$$

Revenons maintenant sur les transformations de jauge. Nous nous sommes arrêtés à l'action (2.29) de Λ sur le champ de jauge. Pour la courbure, nous pouvons désormais réécrire la version infinitésimale de (2.16) — avec u et α — sous la même forme. Ainsi,

$$\begin{aligned} \delta_{\Lambda} \mathcal{B} &= d\Lambda + [\mathcal{B}, \Lambda] = d_{\mathcal{B}} \Lambda \\ \delta_{\Lambda} \mathcal{F} &= [\mathcal{F}, \Lambda] = d_{\bar{\mathcal{B}}} d_{\mathcal{B}} \Lambda \end{aligned} \quad (2.37)$$

En ce qui concerne les transformations de jauge de la jauge, elles se réécrivent elles aussi élégamment

$$\begin{aligned} \delta_k \Lambda &= dk + [\bar{\mathcal{B}}, k] = d_{\bar{\mathcal{B}}} k \\ \delta_k \mathcal{B} &= [\bar{\mathcal{F}}, k] = d_{\mathcal{B}} d_{\bar{\mathcal{B}}} k \\ \delta_k \mathcal{F} &= [\mathcal{F}, d_{\bar{\mathcal{B}}} k] = d_{\bar{\mathcal{B}}} d_{\mathcal{B}} d_{\bar{\mathcal{B}}} k \end{aligned} \quad (2.38)$$

⁴ Pour le montrer, on commence par le commutateur de Φ avec $\varphi_{p,d}$, composante de degré (p, d) avec $p + d = m$:

$$\begin{aligned} [\Phi, \varphi_{p,d}] &= \sum_i [\phi_{i,n-i}, \varphi_{p,d}] = \sum_i -(-1)^{(i-1)(p-1)+(m-p)(n-i)} [\varphi_{p,d}, \phi_{i,n-i}] \\ &= (-1)^{mn} [(-1)^{p(n-1)} \varphi_{p,d}, \sum_i (-1)^{i(m-1)} \phi_{i,n-i}] \\ &= (-1)^{mn} [(-1)^{p(n-1)} \varphi_{p,d}, s^{m-1} \Phi] \end{aligned}$$

Et si maintenant les $\varphi_{p,d}$ sont les composantes de Φ_2 , il n'y a plus qu'à sommer pour trouver la formule proposée.

On remarque que ces expressions sont compatibles avec (2.37), au sens où, en notant $\Lambda_k = \delta_k \Lambda$, on a

$$\delta_k \mathcal{B} = d_B \Lambda_k, \quad \delta_k \mathcal{F} = d_{\mathcal{B}} d_B \Lambda_k$$

On en trouvera l'explication dans la note 5 : δ_k agit comme δ_{Λ_k} .

Les relations (2.38) ne sont pas particulièrement difficiles à vérifier sur les composantes, si l'on est bien attentif aux signes en manipulant les différents commutateurs. A titre d'exemple, et pour se convaincre, voyons quelques uns de ces calculs. Tout d'abord⁵, on écrit les transformations de α et u , à l'ordre 1 en α, u, k , tirées bien sûr de (2.17) :

$$\begin{aligned} \delta_k \alpha &= dk + [A, k] \\ \delta_k u &= [M, k] \end{aligned}$$

On en déduit la transformation induite sur A , par exemple⁶,

$$\begin{aligned} \delta_{u,\alpha} A &= du + [A, u] - [M, \alpha] \\ \delta_k A &= d[M, k] + [A, [M, k]] - [M, dk + [A, k]] \\ &= [dM, k] + [M, dk] + [[A, M], k] + [M, [A, k]] - [M, dk] - [M, [A, k]] \\ &= [N, k] \end{aligned}$$

Enfin, on calcule la variation des composantes de la courbure. Par exemple, pour N ,

$$\begin{aligned} N + \delta_k N &= d(M + \delta_k M) + [A + \delta_k A, M + \delta_k M] \\ \delta_k N &= d[K, k] + [A, [K, k]] - [M, [N, k]] \\ &= [dK, k] + [K, dk] + [[A, K], k] + [K, [A, k]] - [[M, N], k] + [N, [M, k]] \\ &= [K, dk + [A, k]] + [N, [M, k]] \end{aligned}$$

où l'on a utilisé l'identité de Bianchi (2.14, ligne 4) entre l'avant-dernière et la dernière ligne.

⁵ On peut aussi considérer la transformation $U_{\gamma,\alpha}$, suivie de son inverse $U_{\gamma^{-1},-\gamma^{-1}\alpha}$. Le produit est bien sûr l'identité :

$$U_{\gamma^{-1},-\gamma^{-1}\alpha} U_{\gamma,\alpha} = U_{1,\alpha-\alpha} = U_{1,0}$$

Maintenant, si on change la première transformation, en utilisant $\gamma + \delta_k \gamma$ et $\alpha + \delta_k \alpha$, on obtiendra au lieu de l'identité la transformation induite par k sur le champ de jauge :

$$U_k = U_{\gamma^{-1},-\gamma^{-1}\alpha} U_{\gamma+\delta_k \gamma, \alpha+\delta_k \alpha} = U_{1+\delta_k \gamma \gamma^{-1}, -\delta_k \gamma \gamma^{-1} \alpha + \delta_k \alpha}$$

Autrement dit, c'est la transformation de jauge paramétrée par $(\gamma_k = 1 + u_k)$

$$\begin{aligned} u_k &= \delta_k \gamma \gamma^{-1} = [M, k] \\ \alpha_k &= -\delta_k \gamma \gamma^{-1} \alpha + \delta_k \alpha = dk + A(k) \end{aligned}$$

⁶ On peut noter que $\delta_k A$ et $\delta_k M$ ne sont pas nulles, contrairement à ce qui est affirmé dans [47].

Réduction du complexe

Nous allons réduire les degrés de liberté de cette théorie en réduisant le complexe de Hochschild, de manière consistante avec la structure de jauge, bien sûr. Pour cela, nous allons prendre pour champ M la multiplication naturelle de l'algèbre associative $\mathfrak{g}$:

$$\forall g_1, g_2 \in \mathfrak{g}, \quad M(g_1, g_2) = g_1 g_2 \quad (2.39)$$

Nous appellerons alors δ_M la différentielle de Hochschild (2.9) associée à M ,

$$\delta_M = [M, \cdot] \quad (2.40)$$

et puisque M est associative, $\delta_M^2 = 0$, et en particulier $K = 0$. D'autre part, soit $\mathfrak{h}$ l'espace des dérivations sur $\mathfrak{g}$. Les dérivations sont en particulier linéaires sur l'algèbre, ce qui fait de $\mathfrak{h}$ un sous-espace de $\text{End}(\mathfrak{g})$, l'espace de degré 1 du complexe de Hochschild. Nous imposons alors au champ A d'appartenir à $\mathfrak{h}$. Enfin, nous limiterons les transformations de jauge infinitésimales $\Lambda = u + \alpha$ à celles pour lesquelles u est une dérivation.

Concrètement, "être une dérivation" signifie vérifier la règle de Leibniz vis-à-vis du produit naturel de $\mathfrak{g}$, c'est-à-dire M : pour tout $g_1, g_2 \in \mathfrak{g}$

$$\begin{aligned} u(g_1 g_2) &= u(g_1) g_2 + g_1 u(g_2) \\ (u \circ M)(g_1, g_2) &= (M \circ u)(g_1, g_2) \\ 0 &= [M, u](g_1, g_2) \end{aligned}$$

Autrement dit, u est une dérivation si et seulement si $\delta_M(u) = 0$. Ce qui signifie aussi : $\mathfrak{h} = \text{Ker } \delta_M|_{\text{End}(\mathfrak{g})}$.

En résumé, la réduction du complexe de Hochschild est réalisée en imposant :

- ▷ M est la multiplication (associative) sur $\mathfrak{g}$ (et donc $K = 0$),
- ▷ $A \in \mathfrak{h} = \text{Der}(\mathfrak{g})$ (et donc $N = dM + [A, M] = 0$),
- ▷ les transformations de jauge admissibles sont celles pour lesquelles $u \in \mathfrak{h}$.

Pour que la théorie ainsi réduite soit consistante, il nous faut alors vérifier que :

- les transformations admissibles forment un sous-groupe du groupe de jauge initial,
 - M est invariant par transformation de jauge, et reste ainsi fixé à la valeur choisie.
- Le premier point est évident si l'on se souvient de la formule de composition des transformations de jauge (2.20) : $\gamma' \gamma = (1 + u')(1 + u) = 1 + u + u'$ au premier ordre, et les dérivations sont stables par addition. De même, si $\gamma = 1 + u$, alors $\gamma^{-1} = 1 - u$, et les transformations de jauge admissibles sont donc stables aussi par inversion. Cela fait de ces transformations un sous-groupe du groupe de jauge initial.

Vérifier la deuxième condition est encore plus simple, puisque la transformation de M s'écrit :

$$\delta_A M = [M, u] = \delta_M(u) = 0$$

La théorie réduite est donc bien consistante. Comme nous l'avons vu plus haut, la courbure se réduit aux champs H et F , puisque $K = N = 0$:

$$\begin{aligned} H &= d_A B = dB + [A, B] \\ F &= dA + \frac{1}{2} [A, A] + \delta_M B \end{aligned} \quad (2.41)$$

où $d_A = d + [A, \cdot]$ est la différentielle extérieure covariante associée à A . Les identités de Bianchi s'écrivent :

$$\begin{aligned} d_A H &= [F, B] \\ d_A F &= \delta_M H \\ 0 &= \delta_M F \end{aligned} \tag{2.42}$$

2.2.3 2-groupes de Lie, 2-algèbres de Lie et modules croisés

L'introduction de ces concepts est motivé par [48].

2-groupes de Lie comme modules croisés de Lie

La théorie des catégories permet d'introduire le concept de 2-groupe de Lie, qui décrit de manière abstraite, comme nous allons le voir, la structure que nous avons développée à la section précédente. Commençons donc par présenter ce qu'est une catégorie :

une catégorie C consiste en un ensemble C_0 d'objets et un ensemble C_1 de morphismes, ainsi que des fonctions s, t, i et $\circ$. Notons qu'un morphisme g de C_1 est un couple (x, f) , où $x \in C_0$ et f est une application de C_0 dans C_0 .

– s et t sont les fonctions *source* et *destination* ("target")

$$s, t : C_1 \rightarrow C_0, \quad s(x, f) = x, \quad t(x, f) = f(x),$$

– $i : C_0 \rightarrow C_1$ associe à chaque objet son morphisme identité : $i(x) = (x, id)$,

– $\circ$ est la composition des morphismes :

$$\circ : C_1 \times_{C_0} C_1 \longrightarrow C_1$$

où $C_1 \times_{C_0} C_1 = \{(g_1, g_2) \in C_1 \mid t(g_1) = s(g_2)\}$ est l'ensemble des morphismes composables.

Cela nous permet de définir ce qu'est un 2-groupe de Lie, essentiellement en remplaçant les mots "ensemble" et "fonction" par "groupe de Lie" et "homomorphisme".

Définition 2.2.1 *Un 2-groupe de Lie est une catégorie C dont l'ensemble C_0 des objets et l'ensemble C_1 des morphismes sont des groupes de Lie, l'ensemble des morphismes composables $C_1 \times_{C_0} C_1$ est un sous-groupe de Lie de $C_1 \times C_1$, et les fonctions s, t, i et $\circ$ sont des homomorphismes (continus).*

On peut remarquer qu'un 2-groupe de Lie dont tous les morphismes de C_1 sont des identités n'est rien d'autre que le groupe de Lie C_0 . C'est en ce sens que les 2-groupes de Lie méritent leur appellation et généralisent la notion de groupe de Lie.

Dans la situation générale, si G et H sont les groupes de Lie

$$G = C_0, \quad H = \text{Ker } s \subset C_1$$

alors la fonction destination t se réduit à un homomorphisme de H dans G . De plus, l'inclusion i définit une action α de G sur H , par

$$\alpha_g(h) = i_g h i_{g^{-1}}$$

Comme on pourrait le vérifier, cela fait alors de ce 2-groupe de Lie un module croisé de Lie. Cette structure est définie comme suit :

Définition 2.2.2 *Un module croisé de Lie (G, H, t, α) est formé des groupes de Lie G et H , d'un homomorphisme t de H dans G et d'une action α de G sur H (c'est-à-dire un homomorphisme $\alpha : G \rightarrow \text{Aut}(H)$) tels que*

$$t(\alpha_g(h)) = g t(h) g^{-1}, \quad \text{et} \quad \alpha_{t(h)}(h') = h h' h^{-1}$$

pour tout $g \in G$ et $h, h' \in H$.

Inversement, étant donné un module croisé de Lie (G, H, t, α) , on peut construire un 2-groupe de Lie avec $C_0 = G$ et $C_1 = G \ltimes H$, où C_1 est muni de la multiplication du produit semi-direct,

$$(g, h)(g', h') = (gg', h \alpha_g(h'))$$

Ainsi,

Proposition 2.2.3 *La catégorie des 2-groupes de Lie est équivalente à la catégorie des modules croisés de Lie.*

2-algèbres de Lie comme modules croisés différentiels

Après les 2-groupes de Lie, présentons les 2-algèbres de Lie, dont la construction est similaire.

Définition 2.2.4 *Une 2-algèbre de Lie est une catégorie $\mathfrak{c}$ dont l'ensemble $\mathfrak{c}_0$ des objets et l'ensemble $\mathfrak{c}_1$ des morphismes sont des algèbres de Lie, et les fonctions $s, t : \mathfrak{c}_1 \rightarrow \mathfrak{c}_0$, $i : \mathfrak{c}_0 \rightarrow \mathfrak{c}_1$ et $\circ : \mathfrak{c}_1 \times_{\mathfrak{c}_0} \mathfrak{c}_1 \rightarrow \mathfrak{c}_1$ sont des homomorphismes d'algèbres de Lie.*

De même que les groupes de Lie admettent une algèbre de Lie, les 2-groupes de Lie admettent des 2-algèbres de Lie. Bien sûr, ces deux notions coïncident si le 2-groupe de Lie est aussi un groupe de Lie.

Proposition 2.2.5 *A tout 2-groupe de Lie C est associé la 2-algèbre de Lie $\mathfrak{c}$, pour laquelle $\mathfrak{c}_0$ et $\mathfrak{c}_1$ sont les algèbres de Lie respectives de C_0 et C_1 , et les fonctions $s, t, i, \circ$ sont les différentielles de celles de C .*

On peut ici encore écrire une 2-algèbre comme un module croisé, mettant en jeu deux algèbres de Lie $\mathfrak{g}$ et $\mathfrak{h}$, dit module croisé différentiel. En effet, l'action α de $\mathfrak{g}$ sur $\mathfrak{h}$ est cette fois-ci à valeurs dans l'ensemble $\text{Der}(\mathfrak{h})$ des dérivations de $\mathfrak{h}$, qui est l'algèbre de Lie formée des applications linéaires $f : \mathfrak{h} \rightarrow \mathfrak{h}$ vérifiant

$$f([y, y']) = [f(y), y'] + [y, f(y')]$$

Définition 2.2.6 Un module croisé différentiel $(\mathfrak{g}, \mathfrak{h}, t, \alpha)$ est formé des algèbres de Lie $\mathfrak{g}$ et $\mathfrak{h}$, d'un homomorphisme t de $\mathfrak{h}$ dans $\mathfrak{g}$ et d'un homomorphisme α faisant agir $\mathfrak{g}$ comme dérivations de $\mathfrak{h}$ (c'est-à-dire $\alpha : \mathfrak{g} \rightarrow \text{Der}(\mathfrak{h})$) tels que

$$t(\alpha_x(y)) = [x, t(y)]_{\mathfrak{g}}, \quad \text{et} \quad \alpha_{t(y)}(y') = [y, y']_{\mathfrak{h}}$$

pour tout $x \in \mathfrak{g}$ et $y, y' \in \mathfrak{h}$.

Les formules de cette définition deviennent plus naturelles encore si l'on écrit sous forme de crochet l'action α , puisque α_x agit comme une dérivée : $\alpha_x(y) = [x, y]_{\alpha}$. On a alors,

$$t([x, y]_{\alpha}) = [x, t(y)]_{\mathfrak{g}}, \quad [t(y), y']_{\alpha} = [y, y']_{\mathfrak{h}}$$

De plus, $\alpha_x \in \text{Der}(\mathfrak{h})$ se traduit par les identités de Jacobi :

$$\begin{aligned} [[x, x']_{\mathfrak{g}}, y]_{\alpha} &= [x, [x', y]_{\alpha}]_{\alpha} - [x', [x, y]_{\alpha}]_{\alpha}, \\ \text{et} \quad [x, [y, y']_{\mathfrak{h}}]_{\alpha} &= [[x, y]_{\alpha}, y']_{\mathfrak{h}} - [[x, y']_{\alpha}, y]_{\mathfrak{h}} \end{aligned}$$

pour tout $x, x' \in \mathfrak{g}$ et $y, y' \in \mathfrak{h}$.

Proposition 2.2.7 La catégorie des 2-algèbres de Lie est équivalente à la catégorie des modules croisés différentiels.

Il en découle que l'association d'une 2-algèbre de Lie à un 2-groupe de Lie peut aussi s'écrire avec le langage des modules croisés.

Proposition 2.2.8 A tout module croisé de Lie (G, H, t, α) est associé le module croisé différentiel $(\mathfrak{g}, \mathfrak{h}, dt, d\alpha)$, où $\mathfrak{g}$ et $\mathfrak{h}$ sont les algèbres de Lie respectives des groupes de Lie G et H , $dt : \mathfrak{h} \rightarrow \mathfrak{g}$ est la différentielle de $t : H \rightarrow G$, et l'action $d\alpha : \mathfrak{g} \rightarrow \text{Der}(\mathfrak{h})$ est la différentielle de l'action $\alpha : G \rightarrow \text{Aut}(H)$.

Retour sur le complexe Ω^* réduit

De la même manière qu'un groupe de Lie peut servir de groupe de jauge à un fibré, un 2-groupe de Lie C peut définir une structure de jauge dans un 2-fibré sur une variété M . Permettons-nous de noter (G, H, t, α) un tel 2-groupe de Lie, et $(\mathfrak{g}, \mathfrak{h}, dt, d\alpha)$ sa 2-algèbre de Lie associée. La connexion et la courbure de la théorie de jauge peuvent alors faire l'objet de la définition suivante :

Définition 2.2.9 Une connexion (A, B) sur M est constituée d'une 1-forme A sur M à valeurs dans $\mathfrak{g}$ et d'une 2-forme B sur M à valeurs dans $\mathfrak{h}$. Elle a pour courbure la paire (F, H) , où F est la 2-forme à valeurs dans $\mathfrak{g}$ et H la 3-forme à valeurs dans $\mathfrak{h}$:

$$\begin{aligned} F &= dA + \frac{1}{2} [A, A] + \delta(B) \\ H &= dB + [A, B] \end{aligned} \tag{2.43}$$

Dans cette définition, on a utilisé les crochets des formes à valeurs dans $\mathfrak{g}$ ou $\mathfrak{h}$ n'est rien d'autre que le crochet sur l'algèbre correspondante tensorisé par le produit extérieur des formes. Ainsi, si $X = x \otimes \mu$ et $X' = x' \otimes \mu'$ sont des formes différentielles à valeurs dans $\mathfrak{g}$,

$$[X, X'] = [x, x']_{\mathfrak{g}} \otimes \mu \wedge \mu'$$

De même, si $Y = y \otimes \nu$ est une forme différentielle à valeurs dans $\mathfrak{h}$,

$$[X, Y] = [x, y]_{\alpha} \otimes \mu \wedge \nu$$

De plus, δ est l'opérateur remontant les formes Y à valeurs dans $\mathfrak{h}$ en formes à valeurs dans $\mathfrak{g}$,

$$\delta(Y) = -dt(y) \otimes \nu$$

Introduisons la notation d_A pour désigner la différentielle extérieur covariante associée à A ,

$$d_A X = dX + [A, X], \quad d_A Y = dY + [A, Y]$$

où les crochets sont ceux que nous venons de définir. Un calcul direct montre alors que les identités de Bianchi de cette théorie s'écrivent simplement :

$$d_A F = \delta(H), \quad d_A H = [F, B] \quad (2.44)$$

Remarquons que pour démontrer la première identité, on trouve $d_A F = d_A \delta(B)$, puis le résultat énoncé en commutant les deux opérateurs d_A et δ . Cette commutation est justifiée par la définition du module croisé différentiel : en oubliant les parties formes différentielles qui commutent trivialement, il reste en effet

$$[A, dt(B)]_{\mathfrak{g}} = dt([A, B]_{\alpha})$$

qui est la première relation demandée dans la définition 2.2.6.

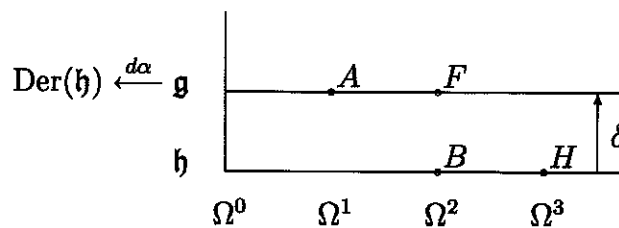


FIG. 2.3 – Les composantes de la courbure (H, F) et de la connexion (B, A) pour une 2-algèbre de jauge $(\mathfrak{g}, \mathfrak{h}, dt, d\alpha)$.

Sous cette forme, et en comparant les figures 2.2 et 2.3, on voit maintenant le lien avec les formules (2.41) et (2.42), écrites sur le complexe Ω^* réduit. En effet, si l'on choisit pour G le groupe $\text{Aut}(H)$, l'algèbre $\mathfrak{g}$ devient alors $\text{Der}(\mathfrak{h})$, faisant des actions α et $d\alpha$ des identités. Les différents crochets que nous avons introduits ici sont en fait les différentes combinaisons du crochet de Gerstenhaber. Enfin, l'opérateur δ s'identifie à la différentielle de Hochschild δ_M .

2.3 L'approche géométrique

Pour conclure ce chapitre, abordons maintenant le problème sous un angle plus intuitif et géométrique. Nous allons d'abord revenir sur les jonctions de cordes, qui sont obtenues par compactification du diagramme "en pantalon" de la membrane M2, et qui constituent dans certains cas des degrés de liberté dynamiques. Cela suggère d'inclure ces configurations de M2 dans la description des paquets de M5-branes.

2.3.1 Cordes de type (p,q) et réseaux de cordes

Plaçons-nous dans la théorie des cordes de type IIB. Les champs de Ramond-Ramond qui survivent sont ceux de degré pairs, associés aux D p -branes pour lesquelles p est impair. En particulier, le secteur de Ramond-Ramond contient la 2-forme C_2 , associée aux D1-branes. Par ailleurs, dans le secteur de Neveu-Schwarz, on trouve toujours la 2-forme B . La présence de ces objets est la raison de l'existence d'une symétrie $SL(2, \mathbb{Z})$ supplémentaire, agissant sur le doublet. En quelques mots, la théorie est incapable de décider laquelle de ces 2-formes, et de leurs combinaisons, est la forme fondamentale, rôle joué par B usuellement. Par conséquent, en plus de la corde fondamentale notée $(1,0)$, et de la D1-brane notée $(0,1)$, on trouve des cordes de type (p,q) avec p et q premiers entre eux, associées à chacune des versions différentes du champ de Ramond-Ramond. Elles sont obtenues à partir de la corde fondamentale par l'action de $SL(2, \mathbb{Z})$

$$A \begin{pmatrix} 1 \\ 0 \end{pmatrix} = \begin{pmatrix} p \\ q \end{pmatrix}, \quad \text{donc } A = \begin{pmatrix} p & v \\ q & u \end{pmatrix}$$

et $A \in SL(2, \mathbb{Z}) \Leftrightarrow up - vq = 1 \Leftrightarrow p \wedge q = 1$ (p et q premiers entre eux).

Mieux encore, ces différentes cordes peuvent s'attacher les unes sur les autres, en formant des vertex à trois branches (fig.2.4a). Ces objets sont appelés "jonctions à trois cordes". Cependant, toute combinaison (p_i, q_i) n'est pas valide. D'une part, la conservation de la charge sous $SL(2, \mathbb{Z})$ impose

$$\sum_{i=1}^3 p_i = \sum_{i=1}^3 q_i = 0 \quad (2.45)$$

D'autre part, le système n'est stable que si la somme vectorielle des tensions s'annule. Assez remarquablement, cela correspond à l'équilibre classique d'une telle configuration de trois cordes. La tension d'une corde de type (p,q) est donnée par l'expression

$$T_{p,q} = \frac{\sqrt{g_s}}{2\pi\alpha'} |p + q\tau|, \quad \text{avec } \tau = \frac{i}{g_s} + \frac{\chi}{2\pi} \quad (2.46)$$

où χ est la valeur dans le vide du champ scalaire de Ramond-Ramond C_0 . En particulier, la jonction est plane, et on peut utiliser des coordonnées complexes dans ce plan.

Prenons pour exemple la jonction la plus simple, formée de la corde fondamentale $(1,0)$ et de la corde $(0,1)$ — qui est aussi la D1-brane puisqu'elle est naturellement

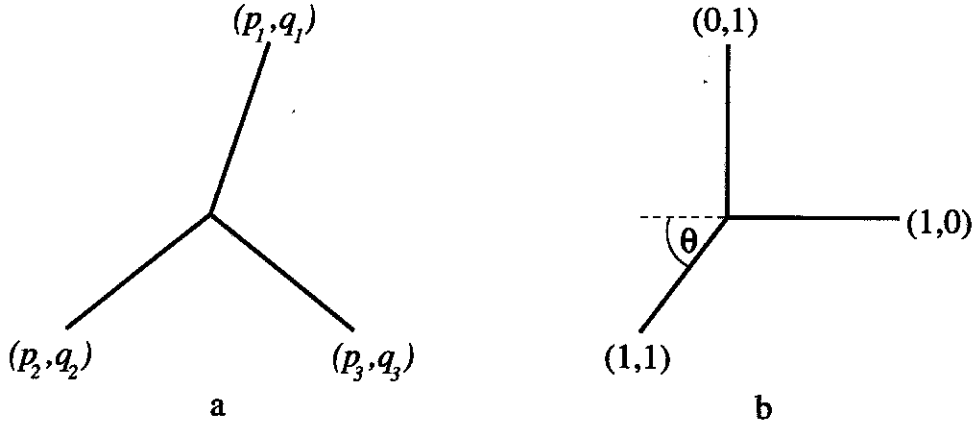


FIG. 2.4 – (a) La jonction à trois cordes n'est possible que si la charge totale s'annule. (b) Pour être stable, elle doit de plus prendre une géométrie particulière. L'angle θ vaut $\varphi_3 - \pi$, dans la formule (2.47)

associée au champ C_2 . Le troisième morceau est alors nécessairement une corde $(1,1)$ par conservation de la charge. En situant la corde de type $(1,0)$ sur le demi-axe $[Ox)$ (voir fig.2.4b), on obtient alors l'équation de stabilité

$$1 + e^{i\varphi_2}|\tau| + e^{i\varphi_3}|1 + \tau| = 0$$

où $e^{i\varphi_2}$ est la direction de la corde $(0,1)$ et $e^{i\varphi_3}$ celle de la corde $(1,1)$. Cette équation complexe est équivalente aux deux équations réelles

$$\begin{cases} 1 + |\tau| \cos \varphi_2 + |1 + \tau| \cos \varphi_3 = 0 \\ |\tau| \sin \varphi_2 + |1 + \tau| \sin \varphi_3 = 0 \end{cases}$$

D'autre part, si l'on prend le module carré de l'équation complexe, on obtient

$$1 + |\tau|^2 + |1 + \tau|^2 + 2|\tau| \cos \varphi_2 + 2|1 + \tau| \cos \varphi_3 + 2|\tau||1 + \tau| \cos(\varphi_2 - \varphi_3) = 0$$

On peut alors développer $\cos(\varphi_2 - \varphi_3)$ et remplacer φ_3 par φ_2 grâce aux deux équations réelles. De plus, en écrivant $\tau = |\tau|e^{i\phi}$, on a

$$|1 + \tau|^2 = 1 + |\tau|^2 + 2|\tau| \cos \phi$$

Après simplification, l'équation précédente devient alors simplement

$$\cos \varphi_2 = \cos \phi$$

Cela laisse deux solutions ($\varphi_2 = \pm\phi$), ce à quoi on pouvait s'attendre à cause de la symétrie d'axe (Ox) dans le plan. Avec par exemple la solution $\varphi_2 = \phi$, le troisième morceau a pour direction

$$e^{i\varphi_3} = -\frac{1 + \tau}{|1 + \tau|}, \quad \text{donc } \tan \varphi_3 = \frac{\text{Im}\tau}{1 + \text{Re}\tau} = \frac{\frac{1}{g_s}}{1 + \frac{\chi}{2\pi}}$$

En particulier, lorsque $\chi = 0$, les cordes $(1,0)$ et $(0,1)$ forment un angle droit ($\tau = \frac{i}{g_s}$, donc $\phi = \frac{\pi}{2}$), et

$$\tan \varphi_3 = \frac{1}{g_s} \tag{2.47}$$

Maintenant que l'on connaît sa géométrie, voyons si l'on peut attacher cette jonction à plusieurs D-branes, tout comme on peut attacher les cordes fondamentales. Les D3-branes étant invariantes sous $SL(2, \mathbb{Z})$, elles peuvent recevoir n'importe quel type de cordes. Considérons donc un système de trois D3-branes parallèles, étendues dans une direction transverse au plan de la jonction que nous venons de décrire. Ce système est paramétré par les distances entre les branes, ou par une distance et les angles α, β, γ entre les côtés du triangle (fig.2.5). Nous allons déterminer dans quels cas la jonction liant ces trois D3-branes est stable, c'est-à-dire ne se désintègre pas en deux cordes séparées.

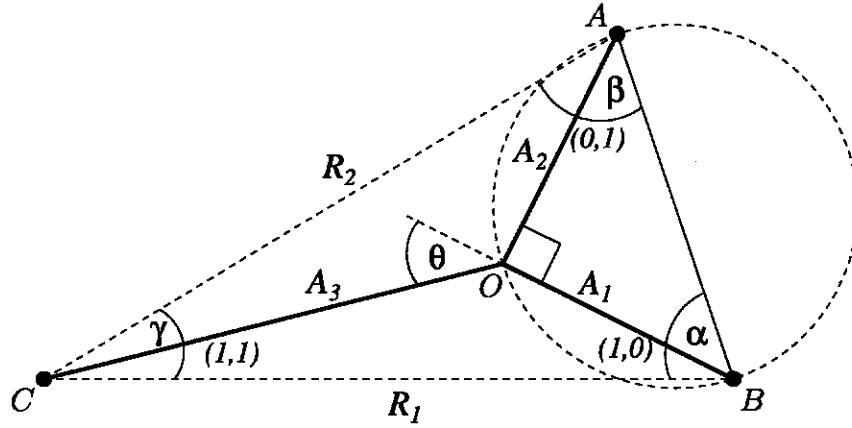


FIG. 2.5 – Les trois D3-branes (aux extrémités du triangle) peuvent être reliées par une corde $(1,0)$ sur le côté inférieur de longueur R_1 et une corde $(0,1)$ sur R_2 , ou par la jonction de la figure. La configuration la plus stable est déterminée à partir des paramètres géométriques indiqués.

Pour cela, nous allons comparer la masse de la jonction avec celle des deux cordes de longueurs respectives R_1 et R_2 , lorsque les trois angles sont inférieurs ou égaux à 90° . Les deux cordes séparées, de type $(1,0)$ et $(0,1)$, ont une masse totale égale à

$$M_{2\text{ cordes}}^2 = \left(\frac{\sqrt{g_s}}{2\pi\alpha'} R_1 + \frac{\sqrt{g_s}}{2\pi\alpha'} \frac{1}{g_s} R_2 \right)^2 = \left(\frac{1}{2\pi\alpha'} \right)^2 \left(g_s R_1^2 + \frac{1}{g_s} R_2^2 + 2R_1 R_2 \right)$$

D'autre part, la masse de la jonction est la somme des masses des trois cordes la composant (il n'y a pas d'énergie de liaison [51]) :

$$\begin{aligned} M_{\text{jonction}}^2 &= \frac{g_s}{(2\pi\alpha')^2} \left(A_1 + \frac{1}{g_s} A_2 + \sqrt{1 + \frac{1}{g_s^2}} A_3 \right)^2 \\ &= \frac{1}{(2\pi\alpha')^2} \left(g_s A_1^2 + \frac{1}{g_s} A_2^2 + \left(g_s + \frac{1}{g_s} \right) A_3^2 + 2A_1 A_2 \right. \\ &\quad \left. + 2g_s \sqrt{1 + \frac{1}{g_s^2}} A_1 A_3 + 2\sqrt{1 + \frac{1}{g_s^2}} A_2 A_3 \right) \end{aligned}$$

En utilisant les relations géométriques dans le triangle ABC de la figure 2.5, on peut

relier les longueurs A_i aux côtés R_1 et R_2 du triangle

$$\begin{aligned} R_1^2 &= A_1^2 + A_3^2 + 2A_1A_3 \cos \theta \\ R_2^2 &= A_2^2 + A_3^2 + 2A_2A_3 \sin \theta \\ A_1^2 + A_2^2 &= R_1^2 + R_2^2 - 2R_1R_2 \cos \gamma \end{aligned}$$

où $\theta = \varphi_3 - \pi$ est compris entre 0 et $\frac{\pi}{2}$. On a alors

$$\frac{1}{g_s} = \tan \theta, \quad \sqrt{1 + \frac{1}{g_s^2}} = \frac{1}{\cos \theta}$$

Cela permet de réécrire M_{jonction}^2 sous la forme

$$\begin{aligned} M_{\text{jonction}}^2 &= \frac{1}{(2\pi\alpha')^2} \left(g_s R_1^2 - 2g_s A_1 A_3 \cos \theta + \frac{1}{g_s} R_2^2 - \frac{2}{g_s} A_2 A_3 \sin \theta \right. \\ &\quad \left. + 2A_1 A_2 + \frac{2}{\sin \theta} A_1 A_3 + \frac{2}{\cos \theta} A_2 A_3 \right) \\ &= \frac{1}{(2\pi\alpha')^2} \left[g_s R_1^2 + \frac{1}{g_s} R_2^2 + 2(A_1 A_2 + A_1 A_3 \sin \theta + A_2 A_3 \cos \theta) \right] \end{aligned}$$

De plus, l'expression entre parenthèse s'exprime grâce au produit vectoriel

$$\begin{aligned} R_1 R_2 \sin \gamma &= \vec{CA} \wedge \vec{CB} = (\vec{CO} + \vec{OA}) \wedge (\vec{CO} + \vec{OB}) \\ &= \vec{CO} \wedge \vec{OB} + \vec{OA} \wedge \vec{CO} + \vec{OA} \wedge \vec{OB} \\ &= A_2 A_3 \sin(\pi - \theta - \frac{\pi}{2}) + A_1 A_3 \sin \theta + A_1 A_2 \\ &= A_1 A_2 + A_1 A_3 \sin \theta + A_2 A_3 \cos \theta \end{aligned}$$

Ainsi, on peut finalement écrire la masse au carré de la jonction sous une forme comparable à la masse au carré des deux cordes séparées

$$M_{\text{jonction}}^2 = \left(\frac{1}{2\pi\alpha'} \right)^2 [g_s R_1^2 + \frac{1}{g_s} R_2^2 + 2R_1 R_2 \sin \gamma] \quad (2.48)$$

ce qui montre que cette jonction est stable tant que γ est strictement inférieur à 90° , cette valeur étant une courbe de stabilité marginale. En effet, c'est la valeur pour laquelle le vertex se situe sur la D3-brane C (la corde CO est de longueur nulle), et les deux autres cordes sont alors libres de se séparer. Si l'angle γ augmente davantage, il est clair que c'est l'état non-lié qui demande le moins d'énergie, ce qui fournit un processus dynamique pour créer et désassembler cette jonction à trois cordes (fig.2.6) [52].

2.3.2 Membranes M2 tendues entre plusieurs M5

On obtient la théorie des cordes de type IIB en compactifiant la M-théorie sur un tore, dans la limite où la taille du tore tend vers 0. En particulier, les cordes sont les membranes M2 dont une direction est enroulée sur un cycle du tore. Le type (p, q)

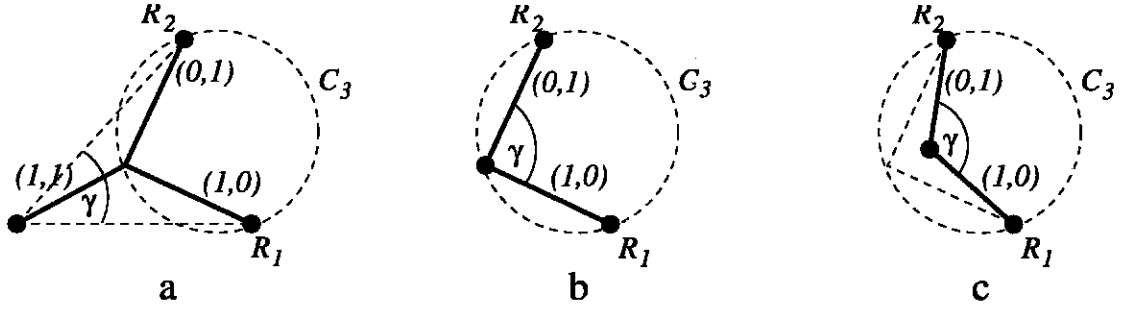


FIG. 2.6 – Les différentes configurations stables en fonction de l'angle γ .

de la corde est déterminé par la classe d'homologie du cycle choisi. Ainsi, la jonction à trois cordes est une compactification particulière du diagramme "en pantalon" de la membrane M2 [50, 51], c'est-à-dire une configuration à trois pattes ou tripode. Puisqu'une telle configuration semble décrire, au moins dans un certain régime, les degrés de liberté des théories de jauge vivant sur les branes dans le cas de la théorie IIB, pourquoi ne pas imaginer qu'il en soit de même en M-théorie. Avant de préciser, commençons par (tenter de) décrire les degrés de liberté de la théorie de jauge vivant sur un paquet de branes M5 de manière géométrique.

En théorie des cordes, les excitations d'une brane unique sont décrites par les cordes ouvertes dont les deux extrémités sont attachées à son volume d'univers. En présence de N branes, les cordes ouvertes peuvent alors s'étirer entre deux branes différentes, formant au total N^2 états différents, en non N si l'on s'était limité aux cordes fixées sur une seule brane.

En M-théorie, on peut faire jouer le rôle de la corde à une M2-brane cylindrique. Elle possède alors deux extrémités, que l'on peut relier au choix à une M5-brane unique ou à deux M5 différentes. Ainsi, en présence d'un paquet de N M5-branes, on dispose toujours de N^2 états. Mais ici, M2 peut aussi prendre d'autres configurations, classifiées par leur nombre d'extrémités :

- A part la sphère ne présentant aucune extrémité (et donc inutile ici), la plus simple est la demi-sphère. Une membrane M2 sous forme de demi-sphère peut s'accrocher à une M5, mais il n'est pas évident que cela fournisse un degré de liberté dynamique de la M5. Un tel système ressemble plutôt topologiquement à une déformation du volume d'univers de la M5. Supposons donc que nous laissons de côté la demi-sphère comme une déformation non-dynamique.
- La configuration suivante est le cylindre (fig.2.7), qui, comme nous l'avons déjà dit, joue exactement le rôle de la corde, avec ses deux extrémités.
- On trouve ensuite le tripode (fig.2.7), ou diagramme en pantalon. Le fait que les jonctions qui en sont dérivées en théorie IIB décrivent des états de la théorie suggère que ces tripodes sont aussi acceptables que les cylindres. Inversement, nous n'avons aucune raison de les éliminer a priori. Ils permettent de relier trois M5-branes en une seule fois, et créent donc N^3 états supplémentaires.

- Qu'en est-il alors des quadrupodes, configuration à quatre extrémités ? Pour plusieurs raisons, nous aimerions qu'ils ne contribuent pas à la dynamique comme états supplémentaires. D'abord parce que le calcul de l'anomalie du paquet de N M5-branes suggère N^3 degrés de liberté, et que ces quadrupodes en fourniraient N^4 . Ensuite parce que s'il existe des jonctions à trois cordes, ce n'est pas le cas de jonctions à quatre cordes (au sens de vertex à quatre pattes). Bien sûr, ces deux explications tiennent plutôt du domaine des conséquences indésirables que des causes a priori. Mais la deuxième réflexion nous ouvre peut-être une voie. En effet, topologiquement, le quadrupode n'est rien d'autre que deux tripodes accolés, et les jonctions à quatre cordes existent donc sous la forme de deux jonctions à trois cordes successives. Voilà qui pourrait être une raison de ne pas compter ces états.
- De même, nous pouvons considérer que toutes les configurations ultérieures ne forment pas de nouveaux états, puisqu'on peut toutes les considérer comme des assemblages de tripodes.

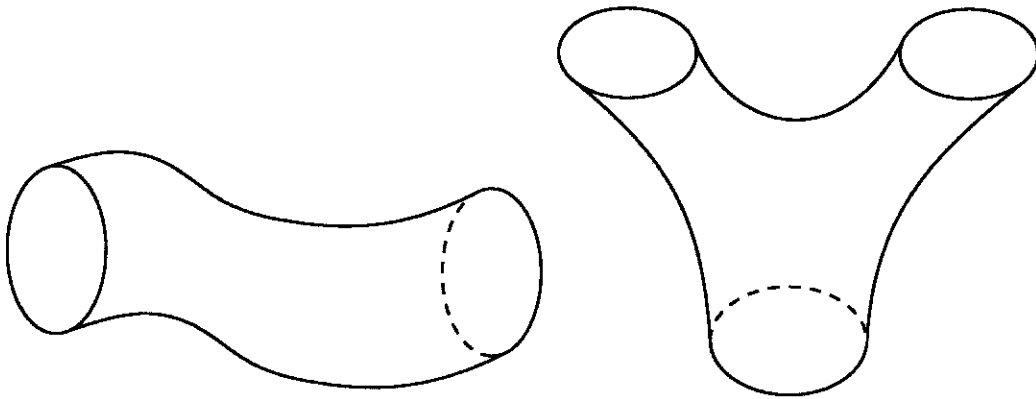


FIG. 2.7 – Le cylindre et le tripode sont les deux configurations “élémentaires” de M2.

Ainsi, il ne reste que deux objets “élémentaires”, le cylindre et le tripode. Remarquons que le deuxième ne peut pas s’obtenir à partir du premier, puisqu’il faudrait pour cela commencer par déchirer le cylindre en un point. Alors que les assemblages qui nous font considérer les autres configurations comme composées ne sont que des recollements sans déchirures. Ce dénombrement des états “élémentaires” nous permet d’arriver à un total de degrés de liberté dynamiques d’ordre N^3 pour un paquet de N M5-branes.

Nous pouvons reformuler cela de manière plus formalisée. Appelons “configuration de niveau n ” la configuration à n extrémités. L’ensemble des configurations est muni de la loi de composition par recollement : deux configurations peuvent être recollées selon l’une de leurs extrémités respectives pour en donner une troisième. Ainsi, deux configurations de niveau respectifs p et q se recollent pour donner une configuration de niveau $p + q - 2$. En particulier, en laissant de côté les sphères (niveau 0) et les demi-sphères (niveau 1), l’ensemble minimal pour générer toutes les configurations est constitué des configurations de niveau 2 (les cylindres) et 3 (les tripodes).

Chapitre 3

Généraliser la symétrie-miroir

Objectif La symétrie-miroir est une symétrie discrète sur les variétés de Calabi-Yau, échangeant les déformations de la structure complexe $\mathcal{J}$ et de la structure de Kähler J . Elle s'interprète comme une T-dualité le long d'une fibration toroïdale de la variété ; c'est la conjecture SYZ (Strominger-Yau-Zaslow [57]). Si l'on souhaite considérer des configurations où le champ de fond antisymétrique $B_{\mu\nu}$ est non-trivial, alors il faut travailler dans des variétés plus générales : les variétés à structure $SU(3)$. Il n'y a alors plus de symétrie-miroir à proprement parler, mais on peut toujours faire une T-dualité le long de la fibration toroïdale. Nous souhaitons ici comprendre cette transformation en termes géométriques sur la variété, et éclaircir en particulier le lien entre géométrie et champ de fond antisymétrique.

Plan du chapitre Nous commencerons par décrire les objets mathématiques dont nous aurons besoin par la suite, à travers la présentation de la structure de Kähler, des variétés de Calabi-Yau et des variétés à structure $SU(3)$. Nous calculerons alors explicitement l'action de la T-dualité sur les composantes W_i de la torsion intrinsèque et du champ $B_{\mu\nu}$, et nous montrerons quelle structure on peut en faire émerger.

3.1 Variétés de Calabi-Yau

Puisque les différentes théories de cordes consistantes sont définies à dix dimensions, cela fait six dimensions (d'espace) en trop par rapport à l'espace-temps que nous pouvons observer autour de nous. On fait alors l'hypothèse que la variété d'espace-temps dans laquelle vit la théorie s'écrit localement $\mathbb{R}^{1,3} \times M$, où $\mathbb{R}^{1,3}$ est l'espace de Minkowski et M est une variété à six dimensions, de taille caractéristique si petite qu'elle n'a pas encore pu être observée. On l'appelle variété de compactification. Ses propriétés définissent bien sûr les paramètres effectifs de la théorie à quatre dimensions, ce qui permet de réduire la liberté de choix de M à ce que les mathématiciens et les cordistes appellent *variétés de Calabi-Yau*. Etude de la bête.

3.1.1 Variétés complexes

Commençons par la définition la plus "naturelle". En effet, une variété *topologique* est définie par un atlas de cartes sur $\mathbb{R}^n$, dont les fonctions de transition sont *continues*. Pour une variété *différentielle*, on demande cette fois un atlas *différentiable*, c'est-à-dire aux fonctions de transition C^∞ . De la même manière, une variété *complexe* de dimension n est définie par un atlas de cartes sur $\mathbb{C}^n$, avec des fonctions de transitions *holomorphes*.

Nous allons maintenant voir comment on définit une variété complexe à partir d'une variété différentielle réelle, et introduire ainsi les outils dont nous nous servirons dans la suite. Soit donc M une variété réelle de dimension $2n$.

Définition 3.1.1 Une structure presque complexe $\mathcal{J}$ sur M est un tenseur $(\mathcal{J}_a^b) \in (TM \otimes T^*M)$ tel que

$$\mathcal{J}_a^b \mathcal{J}_b^c = -\delta_a^c \quad (3.1)$$

Cette structure presque complexe définit un endomorphisme du fibré tangent TM . En effet, si v est un champ de vecteurs sur M (c'est-à-dire une section de TM), $\mathcal{J}v$ est le champ de vecteurs défini par

$$(\mathcal{J}v)^b = \mathcal{J}_a^b v^a \quad (3.2)$$

Ainsi $\mathcal{J}^2 = -1$, faisant de chaque espace tangent $T_p M$ un espace vectoriel complexe. Si cet objet est appelé structure *presque* complexe, c'est qu'il manque quelque chose pour que la variété M , munie de $\mathcal{J}$ soit complexe. En effet, comme souvent avec les variétés, la stratégie consiste à

- définir un objet localement, où tout se passe comme dans $\mathbb{R}^{2n}$,
- recoller les morceaux sur un atlas pour que les propriétés soient réellement globales.

Or ici, justement, il y a une obstruction à la construction d'un atlas holomorphe. Si les dérivées de $\mathcal{J}$ ne vérifient pas une certaine identité, alors il est impossible de transporter la structure presque complexe d'une carte à l'autre. Plus précisément, son tenseur de Nijenhuis associé doit s'annuler :

Définition 3.1.2 Le tenseur de Nijenhuis N d'une structure presque complexe est défini par

$$N_{bc}^a = \mathcal{J}_d^a \partial_b \mathcal{J}_c^d - \mathcal{J}_d^a \partial_c \mathcal{J}_b^d + \mathcal{J}_c^d \partial_d \mathcal{J}_b^a - \mathcal{J}_b^d \partial_d \mathcal{J}_c^a \quad (3.3)$$

Une structure complexe $\mathcal{J}$ est une structure presque complexe dont le tenseur de Nijenhuis est nul.

Proposition 3.1.3 Une variété complexe $(M, \mathcal{J})$ est une variété réelle M munie d'une structure complexe $\mathcal{J}$.

Revenons quelques instants sur les variétés presque complexes, qui permettent déjà de définir les objets essentiels. Tout d'abord, même si elle n'est pas complexe, la structure presque complexe a bien entendu pour valeurs propres $\pm i$, la limitation

étant qu'il n'existe pas nécessairement de base de coordonnées locales qui la diagonalise. Cependant, cela suffit à définir les projecteurs orthogonaux sur les deux sous-espaces propres :

$$(P^\pm)_b^a = \frac{1}{2} (\delta_b^a \mp \mathcal{J}_b^a) \quad (3.4)$$

Grâce à eux, on peut projeter toute forme de degré réel d , et définir ses composantes de degré (p, q) sur la variété (presque) complexe (avec $p + q = d$). De manière générale,

$$\omega^{(d)} = \sum_{p+q=d} \omega^{(p,q)} \quad (3.5)$$

On parle alors pour $\omega^{(p,q)}$ de forme p fois holomorphe et q fois anti-holomorphe. Comme le suggère avec raison la base complexe "idéale" sur l'espace des formes $(dz_i, d\bar{z}_i)$ (qui encore une fois n'existe pas nécessairement sur une variété presque complexe), une forme non nulle sur une variété de dimension réelle $2n$ aura ses deux degrés p et q inférieurs ou égaux à n . Jusque-là, tout se passe donc de la même manière que si la variété était complexe. Par contre, lorsque l'on différencie une (p, q) -forme, les premières différences apparaissent. En effet, sur une variété complexe, on obtient comme on pourrait le deviner naïvement, deux composantes, de degré $(p+1, q)$ et $(p, q+1)$. Par contre, si la variété est seulement presque complexe, on ne peut diagonaliser la structure complexe sur un voisinage, et sa variation n'est donc pas nulle et contribue à la différentielle totale :

$$d\omega^{(p,q)} = (d\omega)^{(p-1,q+2)} + (d\omega)^{(p,q+1)} + (d\omega)^{(p+1,q)} + (d\omega)^{(p+2,q-1)} \quad (3.6)$$

D'autre part, supposons que l'on dispose d'une métrique riemannienne g_{ab} sur la variété. On dit que g est une **métrique hermitienne** si elle vérifie

$$g_{ab} = \mathcal{J}_c^a \mathcal{J}_d^b g_{cd} \quad (3.7)$$

Une variété presque complexe munie d'une métrique hermitienne est dite **variété presque hermitienne**. Lorsque c'est le cas, on définit la forme hermitienne, ou forme fondamentale

$$J_{ac} = \mathcal{J}_b^a g_{bc} \quad (3.8)$$

J est une 2-forme non-dégénérée, dont la donnée est équivalente à celle de la métrique au sens où l'on peut reconstruire g à partir de J :

$$J_{ac} \mathcal{J}_d^c = g_{ad} \quad (3.9)$$

La métrique hermitienne comme la forme fondamentale sont de type $(1, 1)$, c'est-à-dire se décomposent en une composante holomorphe et une composante anti-holomorphe. En particulier, la métrique est utilisée pour monter et descendre les indices (comme elle l'a fait entre $\mathcal{J}$ et J), et ce faisant elle échange les indices holomorphes et anti-holomorphes. De plus, puisque c'est elle qui contracte les formes, toute contraction entre deux indices de même nature est nulle. Nous verrons plus loin sur les expressions explicites l'importance que cela peut avoir.

Enfin, à la métrique g est associée la connexion de Levi-Civita ∇ . Je rappelle qu'il s'agit de l'unique connexion sans torsion qui préserve la métrique ($\nabla g = 0$). En composantes, la connexion de Levi-Civita s'écrit à l'aide des symboles de Christoffel Γ_{jk}^i :

$$\Gamma_{jk}^i = \frac{1}{2} g^{il} (\partial_j g_{kl} + \partial_k g_{lj} - \partial_l g_{jk}) \quad (3.10)$$

Le fait que la connexion soit sans torsion se retrouve dans la symétrie des indices covariants j, k . La dérivée covariante de la forme fondamentale, par exemple, s'écrit alors

$$\nabla_i \mathcal{J}_b^a = \partial_i \mathcal{J}_b^a + \Gamma_{ij}^a \mathcal{J}_b^j - \Gamma_{ib}^j \mathcal{J}_j^a \quad (3.11)$$

On peut alors voir que dans le cas d'une variété presque hermitienne, le tenseur de Nijenhuis s'écrit aussi en remplaçant les dérivées partielles par des dérivées covariantes, puisque les termes supplémentaires introduits s'annulent globalement :

$$N_{bc}^a = \mathcal{J}_d^a \nabla_b \mathcal{J}_c^d - \mathcal{J}_d^a \nabla_c \mathcal{J}_b^d + \mathcal{J}_c^d \nabla_d \mathcal{J}_b^a - \mathcal{J}_b^d \nabla_d \mathcal{J}_c^a \quad (3.12)$$

Pour conclure, comme on pouvait s'en douter, une variété complexe possédant une métrique hermitienne est dite **variété hermitienne**.

3.1.2 Groupes d'holonomie

Avant d'aborder les variétés de Kähler et de Calabi-Yau, parlons un peu des groupes d'holonomie, ce qui nous permettra d'en présenter la classification de Berger. Nous verrons alors apparaître naturellement les variétés de Kähler et de Calabi-Yau, ainsi que les variétés G_2 à sept dimensions, utilisées pour compactifier la M-Théorie. Nous apercevrons aussi très rapidement le lien intrigant qui les lie aux algèbres de division $\mathbb{R}$, $\mathbb{C}$, $\mathbb{H}$ et $\mathbb{O}$.

Transport parallèle et holonomie

Soit E un fibré vectoriel sur une variété M de dimension n , muni d'une connexion ∇^E . Soit $\gamma : [0, 1] \rightarrow M$ un chemin continu (par morceaux) sur la variété, reliant les points $x = \gamma(0)$ et $y = \gamma(1)$. On peut "déplacer" un vecteur de la fibre E_x à la fibre E_y le long de γ , grâce à l'application de **transport parallèle** P_γ :

Proposition 3.1.4 *Soit $\gamma^*(E)$ le fibré sur $[0, 1]$ ayant $E_{\gamma(t)}$ pour fibre en t . ∇^E induit une connexion sur $\gamma^*(E)$, notée elle aussi ∇^E . Pour tout v dans E_x , il existe une unique section continue s de $\gamma^*(E)$ telle que $s(0) = v$ et*

$$\forall t \in [0, 1], \quad \nabla_{\dot{\gamma}(t)}^E s(t) = 0$$

Par définition, $P_\gamma(v) = s(1)$.

Plus simplement, le transport parallèle "déplace" le vecteur le long de la courbe γ de telle sorte qu'il soit toujours "le plus naturellement" (c'est là que la connexion

intervient) tangent à la variété¹. Si l'on transporte un vecteur de E_x le long d'une boucle fermée, on lui associe un autre vecteur de E_x , et P_γ est alors un automorphisme de la fibre (P_γ est inversible puisqu'on obtient son inverse en parcourant la boucle en sens inverse). Ainsi, on définit le **groupe d'holonomie** de la connexion au point x ,

$$\text{Hol}_x(\nabla^E) = \{P_\gamma / \gamma \text{ boucle en } x\}$$

Ce groupe est un sous-groupe de Lie de $\text{GL}(E_x)$, qui est de plus indépendant de x au sens où, si γ est un chemin continu de x à y , $\text{Hol}_x(\nabla^E) = P_\gamma \text{Hol}_y(\nabla^E) P_\gamma^{-1}$. D'autre part, $\text{Hol}_x(\nabla^E)$ est connexe lorsque M est simplement connexe².

Holonomie riemannienne

Supposons que M soit munie d'une métrique riemannienne g . On note $\text{Hol}_x(g)$ le groupe d'holonomie $\text{Hol}_x(\nabla)$ de la connexion de Levi-Civita ∇ de g sur le fibré tangent TM . Par définition de ∇ , $\nabla g = 0$, et le tenseur deux fois covariant $g|_x$ est donc invariant sous l'action³ du groupe d'holonomie en x . Cela signifie que ce groupe d'holonomie est isomorphe à un sous-groupe du groupe orthogonal $O(n)$. On notera dans la suite $\text{Hol}(g)$ ce groupe, défini à une conjugaison près.

Variétés riemanniennes irréductibles

Si (P, g) et (Q, h) sont deux variétés riemanniennes, alors on peut construire leur **produit riemannien** $(P \times Q, g \times h)$. L'espace tangent $T_{(p,q)}(P \times Q)$ est isomorphe à $T_p P \oplus T_q Q$, ce qui permet de définir la métrique-produit sur cette variété au point (p, q) par

$$g \times h|_{(p,q)} = g|_p + h|_q$$

Le groupe d'holonomie de la métrique-produit est alors le produit des groupes d'holonomie de g et h :

$$\text{Hol}(g \times h) = \text{Hol}(g) \times \text{Hol}(h)$$

On dit qu'une variété riemannienne (M, g) est (localement) réductible si elle est (localement) un produit riemannien. Si ce n'est pas le cas, elle est dite **irréductible**. En particulier, en notant n la dimension de M , la représentation de $\text{Hol}(g)$ sur $\mathbb{R}^n$ est alors irréductible.

¹ Prenons un vecteur sur la terre par exemple. Initialement tangent au pôle nord, imaginons qu'il pointe vers le méridien de Greenwich. Transportons-le alors parallèlement à la sphère selon ce méridien jusqu'au pôle sud. Puis remontons au pôle nord selon un méridien à 90°. Il est alors revenu à son point de départ, mais il est dirigé vers la direction opposée. En choisissant d'autres boucles, on peut le faire pointer dans n'importe quelle direction du plan horizontal : le groupe d'holonomie de la sphère S^2 est $\text{SO}(2)$.

² Toute boucle γ se ramène alors continûment au point x , et cette contraction définit un arc dans l'espace des automorphismes de E_x reliant P_γ à l'identité.

³ Cette action est la généralisation aux tenseurs de l'action de $\text{Hol}_x(g)$ sur les vecteurs, qui a servi à le définir.

Variétés riemanniennes non-symétriques

Une variété riemannienne est dite **(localement) symétrique** si, pour tout point p de M , il existe une isométrie (locale) involutive dont p soit un point fixe isolé.

Par exemple, les symétries centrales de $\mathbb{R}^n$ en font des espaces symétriques. Ou les rotations d'angle π autour d'un diamètre pour les sphères S^n . En fait, les espaces localement symétriques sont exactement ceux dont le tenseur de Riemann est constant :

Proposition 3.1.5 *Soit (M, g) une variété riemannienne, ∇ sa connexion de Levi-Civita et R sa courbure de Riemann. (M, g) est localement symétrique si et seulement si $\nabla R = 0$.*

Comme on pourrait le deviner, une variété qui n'est pas localement symétrique est dite non-symétrique.

Classification de Berger des groupes d'holonomie

Nous avons maintenant tous les éléments pour présenter cette classification. Nous avons déjà remarqué qu'un groupe d'holonomie riemannien était un sous-groupe de $\mathrm{SO}(n)$. En voici toutes les possibilités :

Théorème 3.1.6 (Berger) *Soit (M, g) une variété riemannienne de dimension n simplement connexe, irréductible et non-symétrique. Il n'y a alors que sept cas possibles pour son groupe d'holonomie :*

- (i) $\mathrm{Hol}(g) = \mathrm{SO}(n)$,
- (ii) $n = 2m$ avec $m \geq 2$, et $\mathrm{Hol}(g) = \mathrm{U}(m)$,
- (iii) $n = 2m$ avec $m \geq 2$, et $\mathrm{Hol}(g) = \mathrm{SU}(m)$,
- (iv) $n = 4m$ avec $m \geq 2$, et $\mathrm{Hol}(g) = \mathrm{Sp}(m)$,
- (v) $n = 4m$ avec $m \geq 2$, et $\mathrm{Hol}(g) = \mathrm{Sp}(m)\mathrm{Sp}(1)$,
- (vi) $n = 7$ et $\mathrm{Hol}(g) = G_2$,
- (vii) $n = 8$ et $\mathrm{Hol}(g) = \mathrm{Spin}(7)$.

Il est remarquable que tous ces cas soient effectivement réalisables, et présentent un intérêt en physique. Ainsi,

- Les métriques riemanniennes d'holonomie $\mathrm{U}(m)$ sont dites *kähleriennes*, et définissent les variétés de Kähler dont nous allons parler.
- Les métriques d'holonomie $\mathrm{SU}(m)$ définissent les *variétés de Calabi-Yau*, que nous aborderons juste après. Elles sont bien entendu aussi kähleriennes, et leur tenseur de Ricci est nul.
- Les métriques d'holonomie $\mathrm{Sp}(m)$ sont dites *hyperkähleriennes*. En tant que sous-ensemble des métriques de Calabi-Yau ($\mathrm{Sp}(m) \subset \mathrm{SU}(2m) \subset \mathrm{U}(2m)$), elles sont aussi kähleriennes et Ricci-plates.
- Les métriques d'holonomie $\mathrm{Sp}(m)\mathrm{Sp}(1)$ sont dites *quaternion-kähleriennes*, assez bizarrement d'ailleurs puisqu'elles ne sont pas kähleriennes.
- Enfin, les groupes d'holonomie exceptionnels G_2 et $\mathrm{Spin}(7)$ sont associés à des métriques Ricci-plates.

$\mathbb{R}$, $\mathbb{C}$, $\mathbb{H}$, $\mathbb{O}$, les briques du réel...

Les nombres réels $\mathbb{R}$, les nombres complexes $\mathbb{C}$, les quaternions $\mathbb{H}$ et les octonions $\mathbb{O}$ (ou octaves de Cayley) sont les quatre algèbres intègres normées. Une algèbre intègre est une algèbre dans laquelle il n'y a pas de diviseurs de zéro⁴. Une algèbre intègre normée est une algèbre munie d'une norme (au sens de son espace vectoriel sous-jacent) telle que $\|ab\| = \|a\|\|b\|$. Il s'avère que ces algèbres sont exactement au nombre de quatre. Avant d'évoquer leurs liens avec les groupes d'holonomie, je ne résiste pas au plaisir de citer la présentation qu'en fait John Baez dans [56] :

"There are exactly four normed division algebras : the real numbers ($\mathbb{R}$), complex numbers ($\mathbb{C}$), quaternions ($\mathbb{H}$), and octonions ($\mathbb{O}$). The real numbers are the dependable breadwinner of the family, the complete ordered field we all rely on. The complex numbers are a slightly flashier but still respectable younger brother : not ordered, but algebraically complete. The quaternions, being noncommutative, are the eccentric cousin who is shunned at important family gatherings. But the octonions are the crazy old uncle nobody lets out of the attic : they are *nonassociative*."

Bien que non-associative, l'algèbre des octonions est tout de même *alternative*. Il y a en effet trois niveaux d'associativité imbriqués, que l'on peut caractériser grâce à l'associateur $[a, b, c] = (ab)c - a(bc)$. L'algèbre A est :

- associative ssi $[a, b, c] = 0$ pour tout a, b, c dans A ,
- alternative ssi $[a, a, b] = [b, a, a] = 0$ pour tout a, b dans A ,
- associative en puissance ssi $[a, a, a] = 0$ pour tout a dans A .

Revenons aux quatre algèbres $\mathbb{R}$, $\mathbb{C}$, $\mathbb{H}$ et $\mathbb{O}$. Tout d'abord, il est intéressant de noter que tous les groupes d'holonomie précédents sont des groupes d'automorphismes de l'une de ces algèbres :

- $\mathrm{SO}(n)$ est un groupe d'automorphismes de $\mathbb{R}^n$.
- $\mathrm{U}(m)$ et $\mathrm{SU}(m)$ sont des groupes d'automorphismes de $\mathbb{C}^m$.
- $\mathrm{Sp}(m)$ et $\mathrm{Sp}(m)\mathrm{Sp}(1)$ sont des groupes d'automorphismes de $\mathbb{H}^m$.
- G_2 et $\mathrm{Spin}(7)$ sont des groupes d'automorphismes de $\mathbb{O}$.

Plus remarquable encore est leur lien avec les algèbres de Lie simples. En effet, les séries infinies $\mathfrak{so}(n+1)$, $\mathfrak{su}(n+1)$ et $\mathfrak{sp}(n+1)$ sont les algèbres des isométries des espaces projectifs réels, complexes et quaternioniques respectivement. En quelques mots, l'espace projectif réel $\mathbb{RP}^n$ est l'ensemble des lignes passant par l'origine de $\mathbb{R}^{n+1}$. Algébriquement, c'est $\mathbb{R}^{n+1} \setminus \{0\}$ quotienté par la relation d'équivalence :

$$(x_1, \dots, x_{n+1}) \sim (\lambda x_1, \dots, \lambda x_{n+1}), \quad \text{pour tout } \lambda \in \mathbb{R} \setminus \{0\}$$

Si $x_1 \neq 0$, on peut prendre le représentant avec $\lambda = \frac{1}{x_1}$. Avec les n autres cartes de ce type, on obtient un recouvrement de la variété $\mathbb{RP}^n$. La même construction

⁴ C'est-à-dire, pour tout élément a, b de l'algèbre, $ab = 0$ si et seulement si $a = 0$ ou $b = 0$.

La construction de Cayley-Dickson permet de construire une $*$ -algèbre (une algèbre munie d'une conjugaison) de dimension $2n$ à partir d'une $*$ -algèbre de dimension n . En particulier, en partant des réels, c'est ainsi qu'on obtient les nombres complexes, puis les quaternions, puis les octonions. En poursuivant, on peut définir une algèbre de dimension 16, les *sédénions*, mais cette algèbre n'est pas intègre, comme toutes les suivantes. Pour l'anecdote, le sous-espace de ses diviseurs de zéro de norme 1 est isomorphe à G_2 .

définit les espaces projectifs complexes $\mathbb{CP}^n$ et quaternioniques $\mathbb{HP}^n$.

C'est moins évident pour les octonions, puisqu'ils ne sont pas associatifs. Or pour définir une variété par des cartes, il faut que les changements de cartes soient compatibles (changer de système de coordonnées sur l'intersection de deux cartes, puis revenir, doit être une identité). Cela dit, l'alternativité de $\mathbb{O}$ suffit à définir $\mathbb{OP}^1$ et $\mathbb{OP}^2$, mais la série s'arrête là. Ce fait remarquable est lié à l'existence des cinq algèbres exceptionnelles $\mathfrak{g}_2$, $\mathfrak{f}_4$, $\mathfrak{e}_6$, $\mathfrak{e}_7$ et $\mathfrak{e}_8$. Si $\mathfrak{g}_2$ est l'algèbre des dérivations de $\mathbb{O}$, les quatre autres sont les algèbres des isométries de plans projectifs octonionique, bioctonionique, quateroctonionique et octooctonionique. Pour faire simple, il s'agit des plans projectifs⁵ sur les algèbres suivantes :

- $\mathbb{R} \otimes \mathbb{O} = \mathbb{O}$, c'est-à-dire simplement les octonions,
- $\mathbb{C} \otimes \mathbb{O}$, les bioctonions,
- $\mathbb{H} \otimes \mathbb{O}$, les quateroctonions,
- $\mathbb{O} \otimes \mathbb{O}$, les octooctonions.

Ainsi, en résumé, les algèbres de Lie de la classification de Cartan⁶ s'écrivent :

$$\begin{array}{ll}
 \mathfrak{so}(n+1) &= \text{isom}(\mathbb{RP}^n) & \mathfrak{f}_4 &= \text{isom}(\mathbb{OP}^2) \\
 \mathfrak{su}(n+1) &= \text{isom}(\mathbb{CP}^n) & \mathfrak{e}_6 &= \text{isom}((\mathbb{C} \otimes \mathbb{O})\mathbb{P}^2) \\
 \mathfrak{sp}(n+1) &= \text{isom}(\mathbb{HP}^n) & \mathfrak{e}_7 &= \text{isom}((\mathbb{H} \otimes \mathbb{O})\mathbb{P}^2) \\
 \mathfrak{g}_2 &= \mathfrak{der}(\mathbb{O}) & \mathfrak{e}_8 &= \text{isom}((\mathbb{O} \otimes \mathbb{O})\mathbb{P}^2)
 \end{array}$$

3.1.3 Structure de Kähler

Soit $(M, \mathcal{J})$ une variété complexe munie d'une métrique hermitienne g . Si l'une des trois conditions équivalentes suivantes est vérifiée, alors on dit que g est une **métrique kählerienne** :

- (i) $dJ = 0$,
- (ii) $\nabla J = 0$,
- (iii) $\nabla \mathcal{J} = 0$,

où ∇ est la connexion de Levi-Civita associée à g .

Définition 3.1.7 Une variété de Kähler $(M, \mathcal{J}, g)$ est une variété complexe munie d'une métrique kählerienne. La forme fondamentale J est alors appelée **forme de Kähler**.

On peut remarquer que la condition $\nabla \mathcal{J} = 0$ implique l'annulation du tenseur de Nijenhuis, et une variété kählerienne est donc nécessairement complexe.

D'autre part, comme nous l'avions signalé, une variété complexe de dimension n est une variété de Kähler si et seulement si son groupe d'holonomie est inclus dans $U(n)$. On peut rendre cette remarque plus précise :

⁵A part $\mathbb{OP}^2$, et encore, la définition de ces plans projectifs est loin d'être évidente. On se référera à [56] pour plus de renseignements.

⁶En plus des cinq algèbres exceptionnelles, elle comprend les quatre séries infinies $A_n = \mathfrak{su}(n+1)$, $B_n = \mathfrak{so}(2n+1)$, $C_n = \mathfrak{sp}(n)$ et $D_n = \mathfrak{so}(2n)$.

Proposition 3.1.8 *Soit g une métrique sur M . Il existe une structure complexe $\mathcal{J}$ sur M pour laquelle g est kählerienne si et seulement $\text{Hol}(g) \subset U(n)$.*

Enfin, la forme de Kähler est une $(1, 1)$ -forme réelle fermée, qui peut donc être écrite **localement** sous la forme

$$J = i\partial\bar{\partial}\phi$$

où ϕ est une fonction à valeurs réelles, appelée *potentiel de Kähler*. Cependant, puisque J n'est pas exacte, cette écriture n'est que locale. La classe $[J] \in H^2(M, \mathbb{R})$ de J dans la cohomologie de De Rham n'est donc pas triviale. Elle s'écrit bien sûr $[J] = J + i\partial\bar{\partial}\phi$.

3.1.4 Variétés de Calabi-Yau

Définition 3.1.9 *Une variété de Calabi-Yau de dimension n est une variété de Kähler $(M, \mathcal{J}, g)$ compacte, de dimension n et d'holonomie $SU(n)$.*

Sur une telle variété, il existe une $(n, 0)$ -forme non-nulle Ω constante (en particulier $d\Omega = 0$), unique à une phase près, telle que

$$\frac{J^n}{n!} = (-1)^{\frac{n(n-1)}{2}} \left(\frac{i}{2}\right)^n \Omega \wedge \bar{\Omega} \quad (3.13)$$

Ω est la **forme volume holomorphe**. De plus, puisque la forme de Kähler est une $(1, 1)$ -forme, leur produit extérieur est nécessairement nul :

$$J \wedge \Omega = 0 \quad (3.14)$$

D'autre part, une variété de Calabi-Yau est Ricci-plate, et sa première classe de Chern est donc nulle. L'implication est assez directe, contrairement à la réciproque, conjecturée par Calabi en 1954 et prouvée par Yau en 1976 seulement. Appelons ρ la forme de Ricci

$$\rho_{ac} = \mathcal{J}_a^b R_{bc}, \quad R_{ab} = \rho_{ac} \mathcal{J}_b^c$$

où R_{ab} est la courbure de Ricci. ρ est une $(1, 1)$ -forme réelle fermée. Sa classe dans la cohomologie de De Rham $[\rho] \in H^2(M, \mathbb{R})$ est proportionnelle à la première de Chern de M

$$[\rho] = 2\pi c_1(M)$$

En effet, la k -ième classe de Chern $c_k(M)$ est un élément de $H^{2k}(M, \mathbb{R})$, défini par le polynôme

$$c(M) = 1 + \sum_k c_k(M) = \det(1 + \mathcal{R}) = 1 + \text{tr}\mathcal{R} + \text{tr}(\mathcal{R} \wedge \mathcal{R} - 2\text{tr}(\mathcal{R})^2) + \dots$$

où $\mathcal{R}$ est la $(1, 1)$ -forme de courbure à valeurs matricielles

$$\mathcal{R} = R_{i\bar{j}}^k dz^i \wedge d\bar{z}^{\bar{j}}$$

La réciproque, à savoir que l'annulation de la première classe de Chern implique l'existence d'une métrique Ricci-plate sur une variété de Kähler, fait l'objet du théorème suivant :

Théorème 3.1.10 (Yau) *Soit $(M, \mathcal{J})$ une variété complexe compacte dont la première classe de Chern est nulle. Alors chaque classe de Kähler de M contient une unique métrique kählérienne Ricci-plate.*

Si g est l'une de ces métriques kählérienne Ricci-plate, alors bien sûr $(M, \mathcal{J}, g)$ une variété de Calabi-Yau.

3.2 Variétés à structure $SU(3)$

3.2.1 Forme de Kähler et forme volume holomorphe

Nous allons maintenant nous placer dans une variété à trois dimensions complexes, ou plutôt presque complexes. En effet, nous allons élargir un peu le cadre des variétés de Calabi-Yau, en relaxant la condition d'intégrabilité de la structure complexe. Autrement dit, nous imposons seulement à la variété d'être presque complexe. Cela ne nous empêche pas de définir la forme fondamentale J , à partir de la structure presque complexe $\mathcal{J}$ et de la métrique g presque hermitienne. Par abus de langage, nous appellerons tout de même forme de Kähler cette forme fondamentale, bien qu'elle ne soit pas nécessairement fermée. D'autre part, la forme volume holomorphe Ω est toujours bien définie. Il s'agit ici d'une $(3,0)$ -forme, qui n'est a priori pas fermée elle non plus. Cependant, les deux relations (3.13) et (3.14) sont toujours vérifiées (ici avec $n = 3$)

$$J \wedge J \wedge J = \frac{3i}{4} \Omega \wedge \bar{\Omega} \quad (3.15)$$

$$J \wedge \Omega = 0 \quad (3.16)$$

Par contre, ce qui change par rapport aux variétés de Calabi-Yau, c'est que dJ et $d\Omega$ ne sont pas forcément nuls. On peut les décomposer en représentations de $SU(3)$, en faisant apparaître les composantes W_i de la torsion intrinsèque

$$dJ = -\frac{3}{2} \text{Im}(W_1 \bar{\Omega}) + W_4 \wedge J + W_3 \quad (3.17)$$

$$d\Omega = W_1 J^2 + W_2 \wedge J + \bar{W}_5 \wedge \Omega \quad (3.18)$$

Afin de pouvoir revenir sur ces égalités en termes de représentations, donnons la définition d'une variété à structure $SU(3)$.

3.2.2 $SU(3)$ comme groupe de structure

Pour toute variété riemannienne, on définit une base mobile (ou vielbein) comme étant une base de vecteurs e_a tels que $e_a^i e_b^j g_{ij} = \delta_{ab}$. L'ensemble de toutes ces bases mobiles définit le fibré en bases de la variété. Le groupe d'holonomie de la métrique riemannienne g , dont nous avons déjà dit qu'il était inclus dans $SO(6)$ (pour une variété orientable de dimension 6), agit naturellement sur ce fibré comme groupe de structure. Cependant, lorsque la variété est presque complexe et que g est presque hermitienne, on peut toujours se ramener à des bases complexes, et le groupe de structure est réduit à $U(3)$. Cela implique l'existence de la 2-forme J invariante

sous $U(3)$. Ici, nous voulons deux formes invariantes J et Ω , ce qui se traduit par la réduction du groupe de structure à $SU(3)$. “Invariante” signifiant singlet de $SU(3)$, nous allons voir que nous pouvons retrouver ces formes, ainsi que les équations (3.15) et (3.16) par la théorie des représentations.

En effet, en dimension 6, les 1-formes vivent dans la représentation **6** de $SO(6)$, les 2-formes dans la représentation **15** (2 indices antisymétriques à choisir : $C_6^2 = 15$), les 3-formes dans la représentation **20** ($C_6^3 = 20$), ... Ces représentations de $SO(6)$ se décomposent en représentations de $SU(3)$, selon

$$\begin{array}{llll}
0 - \text{forme} : & \mathbf{1} & = & \mathbf{1} \\
1 - \text{forme} : & \mathbf{6} & = & \mathbf{3} \oplus \bar{\mathbf{3}} \\
2 - \text{forme} : & \mathbf{15} & = & \mathbf{1} \oplus \mathbf{8} \oplus \mathbf{3} \oplus \bar{\mathbf{3}} \quad \Rightarrow J \\
3 - \text{forme} : & \mathbf{20} & = & \mathbf{1} \oplus \mathbf{1} \oplus \mathbf{3} \oplus \bar{\mathbf{3}} \oplus \mathbf{6} \oplus \bar{\mathbf{6}} \quad \Rightarrow \Omega = \Omega^+ + i\Omega^- \\
4 - \text{forme} : & \mathbf{15} & = & \mathbf{1} \oplus \mathbf{8} \oplus \mathbf{3} \oplus \bar{\mathbf{3}} \quad \Rightarrow J^2 \\
5 - \text{forme} : & \mathbf{6} & = & \mathbf{3} \oplus \bar{\mathbf{3}} \quad \Rightarrow J \wedge \Omega = 0 \\
6 - \text{forme} : & \mathbf{1} & = & \mathbf{1} \quad \Rightarrow J^3 \propto \Omega \wedge \bar{\Omega}
\end{array}$$

Les singlets **1** de $SU(3)$ montrent l'existence de formes invariantes. Inversement, lorsqu'il n'y a pas de singlet, cela signifie qu'il ne peut y avoir de forme invariante de ce degré. Par exemple, il n'y a pas de 5-forme invariante, donc $J \wedge \Omega$ ne peut être que nul. De même, les 6-formes J^3 et $\Omega \wedge \bar{\Omega}$ ne peuvent qu'être proportionnelles.

3.2.3 La torsion intrinsèque et ses composantes

On trouvera une bonne présentation de la torsion intrinsèque dans [62, annexe C.3]. Nous allons ici seulement faire le point sur la nature des W_i , qui sont la décomposition de la torsion intrinsèque en représentation de $SU(3)$:

- W_1 est une 0-forme complexe dans $\mathbf{1} \oplus \mathbf{1}$,
- W_2 est une 2-forme complexe dans $\mathbf{8} \oplus \mathbf{8}$,
- W_3 est une 3-forme réelle dans $\mathbf{6} \oplus \bar{\mathbf{6}}$,
- W_4 est une 1-forme réelle dans $\mathbf{3} \oplus \bar{\mathbf{3}}$,
- W_5 est une 1-forme complexe holomorphe dans $\mathbf{3} \oplus \bar{\mathbf{3}}$.

On peut déjà remarquer que le comptage des représentations correspond aux formules (3.17) et (3.18)⁷. En étant plus précis, on peut identifier les composantes holomorphes et anti-holomorphes. Ainsi :

- W_2 est une $(1, 1)$ -forme primitive⁸,
- W_3 est une $(2, 1) \oplus (1, 2)$ -forme réelle,
- W_4 est une $(1, 0) \oplus (0, 1)$ -forme réelle,
- et comme déjà signalé, W_5 est une $(1, 0)$ -forme.

Cela permet de relire les équations (3.17) et (3.18) en termes d'holomorphicité,

$$dJ = \underbrace{-\frac{3}{2} \operatorname{Im}(W_1 \bar{\Omega})}_{(3,0) \oplus (0,3)} + \underbrace{W_4 \wedge J + W_3}_{(2,1) \oplus (1,2)}, \quad d\Omega = \underbrace{W_1 J^2 + W_2 \wedge J}_{(2,2)} + \underbrace{\bar{W}_5 \wedge \Omega}_{(3,1)}$$

⁷ On notera que Ω est complexe, donc $d\Omega \in \mathbf{15} \oplus \mathbf{15}$. Cependant, elle est holomorphe, ce qui traduit par l'holomorphicité de W_5 , et explique la disparition d'un des deux termes $\mathbf{3} \oplus \bar{\mathbf{3}}$.

⁸Sans contraction sur J , nous y reviendrons.

Comme nous l'avions vu à la formule (3.6), la différentielle dans une structure presque complexe fait apparaître des termes supplémentaires par rapport à une structure complexe. On remarquera ici qu'il s'agit des termes contenant W_1 et W_2 . On pourrait vérifier qu'ils permettent d'exprimer le tenseur de Nijenhuis, et que par conséquent la variété est complexe si et seulement si $W_1 = W_2 = 0$.

Signalons d'autre part que nous aurions pu définir ces composantes en passant par les parties réelles et imaginaires de $\Omega = \psi - i(*\psi)$:

$$\begin{aligned} d\psi &= W_1^+ J^2 + W_2^+ \wedge J + 2W_5^+ \wedge \psi \\ d(*\psi) &= W_1^- J^2 + W_2^- \wedge J + 2W_5^- \wedge (*\psi) \end{aligned}$$

où $W_i = W_i^+ - iW_i^-$. Dès que nous aurons défini les contractions ($\lrcorner$) nécessaires à séparer les termes (cf 3.3.2), il ne sera pas très difficile de vérifier que

$$-2W_5^+ = 4\psi \lrcorner d\psi = 4(*\psi) \lrcorner d(*\psi) = \Omega \lrcorner d\bar{\Omega} + \bar{\Omega} \lrcorner d\Omega$$

Ainsi, d'une part la quantité $2W_5^+$ est bien la même sur les deux lignes, d'autre part elle vaut bien exactement deux fois la partie réelle de W_5 . Inversement, W_5 en est la partie holomorphe.

3.3 Calculs explicites pour une fibration toroïdale

Nous allons maintenant calculer explicitement les composantes de la torsion intrinsèque, ainsi que l'action de la T-dualité, pour tenter d'en trouver l'interprétation géométrique. Pour cela, nous allons définir une variété à structure $SU(3)$ dont la seule restriction par rapport au cas général sera de présenter une fibration toroïdale. Concrètement, notre système de coordonnées sera constitué de deux triplets de dimensions (la base y^i et la fibre x^α) ayant chacune leur métrique (respectivement g_{ij} et $h_{\alpha\beta}$), reliées par une connexion λ , sur laquelle nous ne faisons pas d'hypothèses particulières.

Outre le fait de nous permettre de mener à bien des calculs explicites, cette restriction est liée au cas des variétés de Calabi-Yau. En effet, la conjecture de Strominger-Yau-Zaslow [57] affirme que toute variété de Calabi-Yau admet une fibration toroïdale. Puisque nous voulons ici étendre le domaine d'application de la symétrie-miroir, il semble donc naturel, voire prudent, de se restreindre aux variétés à structure $SU(3)$ qui possèdent cette propriété de fibration toroïdale. Nous verrons dans la section suivante comment envisager de dépasser cette restriction.

3.3.1 Définitions et notations

Nous noterons $(y^1, y^2, y^3, x^1, x^2, x^3)$ les six dimensions de la variété de compactification. Les trois dimensions y désignent la base et les trois dimensions x la fibre de la fibration toroïdale. Tous les champs ne dépendront que des coordonnées y , faisant des trois directions x des vecteurs de Killing⁹, selon lesquels nous pourrions appliquer

⁹ x^α est un vecteur de Killing de G si et seulement si la dérivée de Lie de G par rapport à x^α s'annule : $\mathcal{L}_{x^\alpha} G = 0$

la transformation de T-dualité. De plus, nous utiliserons la convention suivante pour les indices :

- $i, j, k, \dots$ pour le sous-espace à 3 dimensions des coordonnées y ,
- $\alpha, \beta, \gamma, \dots$ pour le sous-espace à 3 dimensions des coordonnées x ,
- $I, J, K, \dots$ dans l'espace total à 6 dimensions : $dy^I = (dy^i, dx^\alpha)$
- $A, B, C, \dots$ dans l'espace tangent total à 3 dimensions complexes,
- $a, b, c, \dots$ et $a', b', c', \dots$ dans les espaces tangents réels à 3 dimensions des y et des x respectivement (nous oublierons les primes très rapidement).

Nous écrirons la métrique g et le champ B_2 antisymétrique sous la forme générale¹⁰

$$ds^2 = g_{ij} dy^i dy^j + h_{\alpha\beta} (dx^\alpha + \lambda^\alpha)(dx^\beta + \lambda^\beta) = G_{IJ} dy^I dy^J \quad (3.19)$$

$$\begin{aligned} B_2 &= \frac{1}{2} B_{ij} dy^i \wedge dy^j + B_\alpha \wedge (dx^\alpha + \frac{1}{2} \lambda^\alpha) \\ &\quad + \frac{1}{2} B_{\alpha\beta} (dx^\alpha + \lambda^\alpha) \wedge (dx^\beta + \lambda^\beta) \end{aligned} \quad (3.20)$$

avec

$$\lambda^\alpha = \lambda_i^\alpha dy^i, \quad B_\alpha = B_{i\alpha} dy^i, \quad B_{\alpha\beta} = -B_{\beta\alpha}, \quad B_{ij} = -B_{ji}$$

Sur cette variété riemannienne, nous allons définir une structure presque complexe. Tout d'abord, nous complexifions le vielbein

$$E^A = ie_i^a dy^i + V_\alpha^{a'} (dx^\alpha + \lambda^\alpha) \quad (3.21)$$

où $A = a = a'$ prend les valeurs 1, 2, 3, et où nous avons utilisé les vielbeine e_i^a et $V_\alpha^{a'}$ de g et h respectivement :

$$\begin{aligned} \delta_{ab} e_i^a e_j^b &= g_{ij} & \delta_{a'b'} V_\alpha^{a'} V_\beta^{b'} &= h_{\alpha\beta} \\ g^{ij} e_i^a e_j^b &= \delta^{ab} & h^{\alpha\beta} V_\alpha^{a'} V_\beta^{b'} &= \delta^{a'b'} \end{aligned}$$

De même, la définition de E^A permet d'établir les relations correspondantes,

$$\delta_{AB} E_I^A E_J^B = G_{IJ}, \quad G^{IJ} E_I^A E_J^B = \delta^{AB}$$

A cette structure holomorphe correspond bien sûr la structure anti-holomorphe $\overline{E^A} = \overline{E^A}$, aux propriétés similaires.

Notons qu'il n'y a aucune raison pour que E^A soit intégrable ($E^A \neq dz^A$). C'est pour cela que la structure est seulement presque complexe. Cela dit, nous allons tout de même définir, par abus de notation, la quantité " dz^j ", que l'on ne confondra pas avec la différentielle d'une hypothétique coordonnée z^j ,

$$dz^j \equiv dy^j - i V_\alpha^j e^\alpha = -i e_a^j E^A \quad (3.22)$$

où l'on a introduit

$$e^\alpha = dx^\alpha + \lambda^\alpha, \quad V_{i\alpha} = \delta_{aa'} e_i^a V_\alpha^{a'} = e_i^a V_{a\alpha}$$

¹⁰ avec la convention $\omega_1 \wedge \omega_2 = \omega_1 \otimes \omega_2 - \omega_2 \otimes \omega_1$

La “forme de Kähler” est alors définie par

$$J = \frac{i}{2} \delta_{AB} E^A \wedge E^{\bar{B}} = \frac{i}{2} g_{ij} dz^i \wedge d\bar{z}^j = -V_{i\alpha} dy^i \wedge e^\alpha \quad (3.23)$$

tandis que la 3-forme holomorphe s'écrit

$$\Omega = E^1 \wedge E^2 \wedge E^3 = \frac{1}{6} \epsilon_{ABC} E^A \wedge E^B \wedge E^C = -\frac{i}{6} \epsilon_{ijk} dz^i \wedge dz^j \wedge dz^k \quad (3.24)$$

où $\epsilon^{ijk} = \epsilon_{abc} e^{ai} e^{bj} e^{ck}$.

Définissons aussi les projecteurs holomorphe P et antiholomorphe $\bar{P}$, dont nous servirons par la suite,

$$P = \frac{1}{2} (1 - iJ), \quad \bar{P} = \frac{1}{2} (1 + iJ) \quad (3.25)$$

Ainsi, J et Ω vérifient les identités (3.15) et (3.16). Leurs différentielles ne sont pas nulles, mais se décomposent sous la forme (3.17) et (3.18) grâce aux composantes de la torsion intrinsèque, auxquelles nous allons maintenant nous intéresser.

3.3.2 Les composantes de la torsion intrinsèque

Pour calculer les composantes W_i à partir des formules (3.17) et (3.18), il nous faut introduire la contraction $\lrcorner$ sur les formes. Elle est définie sur la base complexe par

$$E^A \lrcorner E^{\bar{B}} = \delta^{AB} \quad (3.26)$$

les autres combinaisons étant nulles. Cela s'écrit aussi

$$dz^i \lrcorner d\bar{z}^j = g^{ij} \quad (3.27)$$

En ce qui concerne les composantes réelles, la base naturellement adaptée aux contractions est la base (dy^i, e^α) . En effet, on pourra vérifier que

$$e^\alpha \lrcorner e^\beta = \frac{1}{2} h^{\alpha\beta}, \quad e^\alpha \lrcorner dy^j = dy^i \lrcorner e^\beta = 0, \quad dy^i \lrcorner dy^j = \frac{1}{2} g^{ij}$$

De plus, les contractions sont linéaires, et pour toutes formes A, B, C (telles que l'expression suivante ait un sens)

$$(A \wedge B) \lrcorner C = B \lrcorner (A \lrcorner C)$$

Enfin, si ω_1 est une 1-forme

$$\omega_1 \lrcorner (A \wedge B) = (\omega_1 \lrcorner A) \wedge B + (-1)^{|A|} A \wedge (\omega_1 \lrcorner B)$$

Quelques exemples utiles

Puisque les seules contractions non nulles sont celles entre un indice holomorphe et un indice anti-holomorphe, la contraction de la $(1, 1)$ -forme J et de la $(3, 0)$ -forme Ω , par exemple, est trivialement nulle.

D'autre part, nous pouvons remarquer quelques identités intéressantes :

$$J \lrcorner J = J^2 \lrcorner J^2 = \frac{3}{4}, \quad J \lrcorner J^2 = J, \quad \Omega \lrcorner \bar{\Omega} = 1, \quad J^2 \lrcorner d\Omega = -i \Omega \lrcorner dJ$$

Nous nous contenterons ici de démontrer la dernière égalité, qui nous donnera deux manières de calculer W_1 . Tout d'abord,

$$\begin{aligned} i \Omega \lrcorner dJ &= \left(\frac{i}{6} \epsilon_{ABC} E^A \wedge E^B \wedge E^C \right) \lrcorner \left(\frac{i}{2} dE^D \wedge E^{\bar{D}} - \frac{i}{2} E^D \wedge dE^{\bar{D}} \right) \\ &= -\frac{3}{12} \epsilon_{ABC} [(E^A \wedge E^B) \lrcorner dE^D] \delta^{C\bar{D}} + 0 \\ &= -\frac{1}{4} \epsilon_{ABC} (E^A \wedge E^B) \lrcorner dE^C \end{aligned}$$

Et de même pour la deuxième expression

$$\begin{aligned} J^2 \lrcorner d\Omega &= \left(-\frac{1}{4} E^D \wedge E^{\bar{D}} \wedge E^F \wedge E^{\bar{F}} \right) \lrcorner \left(\frac{1}{2} \epsilon_{ABC} dE^A \wedge E^B \wedge E^C \right) \\ &= \frac{1}{8} \epsilon_{ABC} [(E^D \wedge E^F) \lrcorner dE^A] [(E^{\bar{D}} \wedge E^{\bar{F}}) \lrcorner (E^B \wedge E^C)] \\ &= \frac{2}{8} \epsilon_{ABC} \delta^{B\bar{D}} \delta^{C\bar{F}} [(E^D \wedge E^F) \lrcorner dE^A] \\ &= \frac{1}{4} \epsilon_{ABC} (E^A \wedge E^B) \lrcorner dE^C \end{aligned}$$

D'autre part, on peut trouver une deuxième série d'identités, lorsque J et Ω sont mélangées à d'autres formes. En particulier,

$$\begin{aligned} J \lrcorner (Q^1 \wedge J) &= \frac{1}{2} Q^1 & \Omega \lrcorner (Q^1 \wedge \bar{\Omega}) &= -PQ^1 \\ J^2 \lrcorner (Q^2 \wedge J) &= (J \lrcorner Q^2) & \bar{\Omega} \lrcorner (Q^1 \wedge \Omega) &= -\bar{P}Q^1 \\ J^2 \lrcorner (Q^3 \wedge J) &= \frac{1}{2} (J \lrcorner Q^3) \end{aligned}$$

où Q^p est une p -forme.

Voyons pour l'exemple la démonstration de la première expression. Pour cela, séparons Q^1 en ses parties holomorphes et anti-holomorphes, $Q^1 = Q_A^1 E^A + Q_{\bar{A}}^1 E^{\bar{A}}$:

$$\begin{aligned} J \lrcorner (Q^1 \wedge J) &= -\frac{1}{4} (E^C \wedge E^{\bar{C}}) \lrcorner (Q_A^1 E^A \wedge E^D \wedge E^{\bar{D}} + Q_{\bar{A}}^1 E^{\bar{A}} \wedge E^D \wedge E^{\bar{D}}) \\ &= -\frac{1}{4} \left[Q_A^1 \delta^{C\bar{D}} (\delta^{\bar{C}A} E^D - \delta^{\bar{C}D} E^A) + Q_{\bar{A}}^1 \delta^{\bar{C}D} (\delta^{CA} E^{\bar{D}} - \delta^{C\bar{D}} E^{\bar{A}}) \right] \\ &= -\frac{1}{4} \left[Q_A^1 (E^A - 3E^A) + Q_{\bar{A}}^1 (E^{\bar{A}} - 3E^{\bar{A}}) \right] \\ &= \frac{1}{2} Q^1 \end{aligned}$$

Enfin, en plus de ces identités, nous utiliserons les propriétés de primitivité de W_2 et W_3

$$J \lrcorner W_2 = 0, \quad J \lrcorner W_3 = 0, \quad \Omega \lrcorner W_3 = 0$$

La composante W_1

W_1 est un scalaire complexe, que nous pouvons isoler à partir de dJ ou de $d\Omega$. Par exemple, le premier calcul, très court, s'écrit :

$$\Omega \lrcorner dJ = \Omega \lrcorner \left[\frac{3}{4i} (W_1 \bar{\Omega} - \bar{W}_1 \Omega) + W_4 \wedge J + W_3 \right] = \frac{3}{4i} W_1$$

En remplaçant J et Ω par leur expression, on trouve alors celle de W_1 , que nous noterons w_1 , par homogénéité avec la suite.¹¹

$$W_1 = w_1 = -\frac{i}{12} \epsilon^{ijk} V_{i\alpha} [d(V - i\lambda)]_{jk}^\alpha \quad (3.28)$$

La composante W_4

W_4 est une 1-forme réelle. Elle s'écrit donc dans la base complexe

$$W_4 = w_k^4 dz^k + \overline{w_k^4} d\bar{z}^k$$

c'est-à-dire que $w_k^4 = \overline{w_k^4}$. De plus, d'après les identités citées plus haut,

$$W_4 = 2J \lrcorner dJ = -\frac{1}{2} V^{j\alpha} [dV_\alpha]_{jk} dy^k$$

D'où l'expression de sa partie holomorphe

$$w_k^4 = -\frac{1}{4} V^{j\alpha} [dV_\alpha]_{jk} \quad (3.29)$$

La composante W_5

W_5 est une 1-forme holomorphe, et s'écrit donc $W_5 = w_k^5 dz^k$.

Encore une fois, on peut l'isoler grâce aux identités précédentes :

$$\begin{aligned} W_5 &= -\Omega \lrcorner d\bar{\Omega} \\ &= -\frac{1}{4} \{ V_\alpha^j [d(V + i\lambda)]_{jk}^\alpha - h^{\alpha\beta} \partial_k h_{\alpha\beta} \} dz^k \end{aligned}$$

D'où la valeur de w_k^5

$$w_k^5 = -\frac{1}{4} V_\alpha^j [d(V + i\lambda)]_{jk}^\alpha + \frac{1}{4} h^{\alpha\beta} \partial_k h_{\alpha\beta} \quad (3.30)$$

¹¹ En notant que $[d\lambda^\alpha]_{jk} = \partial_j \lambda_k^\alpha - \partial_k \lambda_j^\alpha$, puisque l'on définit la 2-forme $d\lambda$ comme

$$d\lambda = \frac{1}{2} [d\lambda^\alpha]_{jk} dy^j \wedge dy^k = \partial_j \lambda_k^\alpha dy^j \wedge dy^k$$

La composante W_2

W_2 est une $(1, 1)$ -forme¹² : $W_2 = w_{ij}^2 dz^i \wedge d\bar{z}^j$.

Connaissant W_1 et W_5 , il ne reste qu'à soustraire leur contribution de $d\Omega$ pour isoler W_2 .

$$\begin{aligned} W_2 &= 4J \lrcorner [d\Omega - W_1 J^2 - \bar{W}_5 \wedge \Omega] \\ &= \frac{1}{12} \epsilon^{pqk} [d(V - i\lambda)]_{pq}^\alpha (V_{k\alpha} g_{ij} - 3V_{j\alpha} g_{ik}) dz^i \wedge d\bar{z}^j \end{aligned}$$

Cette fois, nous ne arrêterons pas là. w_{ij}^2 peut en effet être encore décomposée. Elle possède une partie symétrique $w_{\{ij\}}^2$, naturellement sans trace, et une partie antisymétrique, qui se dualise en un vecteur à trois dimensions $w_2^k = \frac{1}{2} \epsilon^{ijk} w_{ij}^2$.

$$w_{ij}^2 = w_{\{ij\}}^2 + w_{[ij]}^2 = w_{\{ij\}}^2 + \epsilon_{ijk} w_2^k$$

On en déduit les expressions

$$w_{\{ij\}}^2 = \frac{1}{24} \epsilon^{pqk} [d(V - i\lambda)]_{pq}^\alpha [2V_{k\alpha} g_{ij} - 3V_{j\alpha} g_{ik} - 3V_{i\alpha} g_{jk}] \quad (3.31)$$

$$w_k^2 = -\frac{1}{4} V_\alpha^j [d(V - i\lambda)]_{jk}^\alpha \quad (3.32)$$

La composante W_3

W_3 est une $(2, 1) \oplus (1, 2)$ -forme réelle primitive. Les deux composantes sont donc complexes conjuguées :

$$W_3 = \frac{1}{2} w_{ijk}^3 dz^i \wedge d\bar{z}^j \wedge d\bar{z}^k + \text{c. c.}$$

Comme pour W_2 , nous notons ici w_{ijk}^3 ce qui est en réalité w_{ijk}^3 . On obtient

$$\begin{aligned} W_3 &= dJ + \frac{3}{2} \text{Im}(W_1 \bar{\Omega}) - W_4 \wedge J \\ &= \frac{3}{8} V_{i\alpha} [d\lambda^\alpha]_{jk} dy^i \wedge dy^j \wedge dy^k \\ &\quad - \frac{1}{4} [dV_\alpha]_{ik} \left[\frac{3}{2} \delta_j^\alpha \delta_\beta^\alpha + V_j^\alpha V_\beta^k - 2V^{k\alpha} V_{j\beta} \right] dy^i \wedge dy^j \wedge e^\beta \\ &\quad + \frac{1}{4} [d\lambda^\alpha]_{jk} \left[\frac{1}{2} V_{i\alpha} V_\beta^j V_\gamma^k - \delta_i^\alpha h_{\alpha\beta} V_\gamma^k \right] dy^i \wedge e^\beta \wedge e^\gamma \\ &\quad - \frac{1}{8} V_\beta^i V_\gamma^j [dV_\alpha]_{ij} e^\alpha \wedge e^\beta \wedge e^\gamma \\ &= \frac{1}{16} \{ -i[dV_\alpha]_{jk} V_i^\alpha + i[dV_\alpha]_{ij} V_k^\alpha + i[dV_\alpha]_{ki} V_j^\alpha \\ &\quad + i[dV_\alpha]_{jl} V^{l\alpha} g_{ik} - i[dV_\alpha]_{kl} V^{l\alpha} g_{ij} \\ &\quad + [d\lambda^\alpha]_{jk} V_{i\alpha} + [d\lambda^\alpha]_{ij} V_{k\alpha} + [d\lambda^\alpha]_{ki} V_{j\alpha} \} dz^i \wedge d\bar{z}^j \wedge d\bar{z}^k + \text{c. c.} \end{aligned}$$

¹² En toute rigueur, nous devrions noter cette composante w_{ij}^2 . Les deux autres composantes, $w_{\bar{i}\bar{j}}^2$ et la vraie w_{ij}^2 , sont nulles. Cependant, ce ne sont pas les propriétés holomorphes qui nous intéressent par la suite, mais simplement la symétrie $\text{SO}(3)$ sur la base. Notons cependant que les indices i et j ne sont pas antisymétriques, comme cette notation simplifiée pourrait le faire croire.

Heureusement, encore une fois, nous ne nous arrêterons pas à cet objet. En effet, la représentation **6** de $SU(3)$, ici exprimée comme (1,2)-forme primitive, peut aussi l'être par un tenseur symétrique à deux indices holomorphes, que nous noterons w_{ij}^3 . On l'obtient en contractant les deux indices anti-holomorphes de w_{ipq}^3 (donc p et q) sur Ω . Comme attendu, le résultat est naturellement symétrique. Nous séparerons cependant sa trace sur la base tridimensionnelle (bien sûr, la trace est nulle à six dimensions, puisqu'elle est égale à la contraction complète de W_3 et Ω , et W_3 est primitive). En résumé,

$$w_{ij}^3 \equiv w_{ipq}^3 \Omega^{pq}{}_j = w_{\{ij\}^0}^3 + \frac{1}{3} w_t^3 g_{ij}$$

Le résultat se simplifie alors quelque peu,

$$w_{\{ij\}^0}^3 = \frac{1}{24} \epsilon^{pqk} [dV_\alpha]_{pq} [2V_k^\alpha g_{ij} - 3V_j^\alpha g_{ik} - 3V_i^\alpha g_{jk}] \quad (3.33)$$

$$w_t^3 = \frac{1}{8} \epsilon^{ijk} V_{i\alpha} [d(V - 3i\lambda)]_{jk}^\alpha \quad (3.34)$$

En résumé

Nous avons donc, à partir des cinq composantes W_i , obtenu les expressions de deux scalaires (w_1, w_t^3), trois vecteurs (w_k^2, w_k^4, w_k^5) et deux tenseurs symétriques de trace nulle ($w_{\{ij\}^0}^2, w_{\{ij\}^0}^3$) à trois dimensions. Ces expressions sont regroupées dans la table 3.1, p.107.

3.3.3 Les composantes du champ de fond H

En l'absence du champ de Ramond-Ramond B_2 , la supersymétrie impose à la variété de compactification d'être une variété de Calabi-Yau. C'est donc la présence d'un champ de fond antisymétrique non trivial qui autorise une déviation de la géométrie, ce qui se matérialise par le fait que les quantités dJ et $d\Omega$ sont non nulles. En quelque sorte, c'est donc $H = dB_2$ qui autorise la présence des composantes W_i . Il est alors naturel de s'attendre à ce que la T-dualité, que nous verrons dans la section suivante, mélange ces composantes avec leurs équivalents pour H , c'est-à-dire les projections de H sur les représentations de $SU(3)$.

$$\begin{aligned} H &= dB_2 \\ &= \frac{1}{2} \partial_k B_{\alpha\beta} dy^k \wedge e^\alpha \wedge e^\beta + [\partial_k B_{i\alpha} - B_{\alpha\beta} \partial_k \lambda_i^\beta] dy^k \wedge dy^i \wedge e^\alpha \\ &\quad + \frac{1}{2} [\partial_k B_{ij} - \partial_k B_{i\alpha} \lambda_j^\alpha + B_{i\alpha} \partial_k \lambda_j^\alpha] dy^i \wedge dy^j \wedge dy^k \end{aligned} \quad (3.35)$$

Nous pouvons ensuite projeter H sur les représentations de $SU(3)$ de la même manière que la 3-forme dJ :

$$H = -\frac{3}{2} \text{Im}(H_1 \bar{\Omega}) + H_4 \wedge J + H_3 \quad (3.36)$$

où les composantes H_1, H_3 et H_4 sont numérotées comme les W_i correspondants. On peut alors calculer ces composantes avec les mêmes contractions que pour W_1, W_3 et W_4 . Puis les exprimer sur la base holomorphe, pour en tirer h_1, h_3 et h_4 .

La composante H_1

H_1 est le scalaire complexe de la représentation $1 \oplus 1$,

$$\begin{aligned} h_1 &= H_1 = -\frac{4i}{3} \Omega_{\perp} H \\ &= \frac{1}{12} \epsilon^{ijk} V_i^{\alpha} V_j^{\beta} \partial_k B_{\alpha\beta} + \frac{i}{12} \epsilon^{ijk} V_i^{\alpha} [dB_{\alpha} - B_{\alpha\beta} d\lambda^{\beta}]_{jk} \\ &\quad - \frac{1}{12} \epsilon^{ijk} [\partial_k B_{ij} - \partial_k B_{i\alpha} \lambda_j^{\alpha} + B_{i\alpha} \partial_k \lambda_j^{\alpha}] \end{aligned} \quad (3.37)$$

On remarque tout de suite que cette expression est beaucoup plus compliquée que w_1 , à qui on voudrait la comparer. Cependant, dans le cas où $B_{\alpha\beta} = 0$, elle se simplifie et la ressemblance apparaît plus clairement. Bien sûr, nous reviendrons plus complètement là-dessus dès que nous aurons défini les transformations de T-dualité, puisque ce sont elles qui se cachent derrière cette “similarité” nécessairement imprécise pour l’instant.

La composante H_4

H_4 étant la $(1, 0) \oplus (0, 1)$ -forme réelle de la représentation $3 \oplus \bar{3}$, elle s’écrit

$$H_4 = h_k^4 dz^k + \text{c. c.}$$

De plus, on trouve son expression de la même manière que pour W_4 précédemment :

$$\begin{aligned} H_4 &= 2J_{\perp} H \\ &= -\frac{1}{2} V^{k\alpha} [dB_{\alpha} - B_{\alpha\beta} d\lambda^{\beta}]_{jk} dy^j - \frac{1}{2} V^{k\alpha} \partial_k B_{\alpha\beta} e^{\beta} \end{aligned}$$

ce qui nous permet d’identifier le terme h_k^4 qui va nous intéresser,

$$h_k^4 = \frac{1}{4} V^{j\alpha} \left\{ [dB_{\alpha} - B_{\alpha\beta} d\lambda^{\beta}]_{jk} - i V_k^{\beta} \partial_j B_{\alpha\beta} \right\} \quad (3.38)$$

La composante H_3

H_3 est la $(2, 1) \oplus (1, 2)$ -forme réelle primitive de la représentation $6 \oplus \bar{6}$,

$$\begin{aligned} H_3 &= H + \frac{3}{2} \text{Im}(H_1 \bar{\Omega}) - H_4 \wedge J \\ &= \frac{1}{4} V^{k\beta} V^{i\gamma} [\partial_k B_{i\alpha} - B_{\alpha\mu} \partial_k \lambda_i^{\mu}] e^{\alpha} \wedge e^{\beta} \wedge e^{\gamma} \\ &\quad + \frac{1}{2} \left[\frac{5}{4} \partial_k B_{\alpha\beta} - V^{j\gamma} V_{k\alpha} \partial_j B_{\gamma\beta} - \frac{1}{2} V_k^{\gamma} V_{\alpha}^j \partial_j B_{\gamma\beta} \right] dy^k \wedge e^{\alpha} \wedge e^{\beta} \\ &\quad + \frac{1}{8} [\partial_k B_{ij} - \partial_i B_{j\gamma} \lambda_k^{\gamma} - \partial_i \lambda_j^{\gamma} B_{k\gamma}] [V_{\beta}^k V_{\alpha}^j dy^i - V_{\beta}^k V_{\alpha}^i dy^j + V_{\alpha}^i V_{\beta}^j dy^k] \wedge e^{\alpha} \wedge e^{\beta} \\ &\quad + \frac{1}{4} [(dB_{\alpha})_{ik} - B_{\alpha\beta} (d\lambda^{\beta})_{ik}] \left[\frac{3}{2} \delta_j^k \delta_{\gamma}^{\alpha} + V_j^{\alpha} V_{\gamma}^k - 2V^{k\alpha} V_{j\gamma} \right] dy^i \wedge dy^j \wedge e^{\gamma} \\ &\quad + \frac{3}{8} [\partial_k B_{ij} - \partial_i B_{j\gamma} \lambda_k^{\gamma} - \partial_i \lambda_j^{\gamma} B_{k\gamma}] dy^i \wedge dy^j \wedge dy^k \\ &\quad - \frac{1}{8} V_j^{\alpha} V_k^{\beta} \partial_i B_{\alpha\beta} dy^i \wedge dy^j \wedge dy^k \\ &= \frac{1}{2} h_{ijk}^3 dz^i \wedge dz^j \wedge dz^k + \text{c. c.} \end{aligned}$$

Et de même que pour W_3 , nous pouvons définir h_{ij}^3 , tenseur symétrique à deux indices holomorphes,

$$h_{ij}^3 = h_{ipq}^3 \Omega^{pq}{}_j = h_{\{ij\}^0}^3 + \frac{1}{3} h_i^3 g_{ij}$$

et encore une fois, les formules sont alors quelque peu plus simples :

$$h_{\{ij\}^0}^3 = -\frac{1}{24} \epsilon^{pqk} [dB_\alpha]_{pq} [2V_k^\alpha g_{ij} - 3V_j^\alpha g_{ik} - 3V_i^\alpha g_{jk}] \quad (3.39)$$

$$h_i^3 = -\frac{1}{8} \epsilon^{ijk} V_i^\alpha [dB_\alpha]_{jk} - \frac{3i}{8} \epsilon^{ijk} [\partial_k B_{ij} - \partial_k B_{i\alpha} \lambda_j^\alpha + B_{i\alpha} \partial_k \lambda_j^\alpha] \quad (3.40)$$

En résumé

Nous avons fait subir à H , champ de fond non trivial, le même traitement qu'à dJ et $d\Omega$, qui définissent la géométrie de la configuration. Nous disposons donc de deux scalaires supplémentaires (h_1 , h_i^3), ainsi que d'un vecteur (h_k^4) et d'un tenseur symétrique de trace nulle ($h_{\{ij\}^0}^3$). Ces expressions sont elles aussi regroupées dans la table 3.1, p.107.

3.3.4 L'action de la T-dualité

Voyons maintenant comment agit la T-dualité sur ces champs, pour pouvoir écrire la version T-duale de la géométrie et du champ de fond. Et éventuellement l'identifier en fonction des champs initiaux.

En désignant par E la matrice complète $E = G + B_2$, somme de la métrique et du champ antisymétrique, la T-dualité T_α dans les directions x^α agit comme

$$E = \left[\begin{array}{c|c} E_{ij} & E_{i\beta} \\ \hline E_{\alpha j} & E_{\alpha\beta} \end{array} \right] \xrightarrow{T_\alpha} \tilde{E} = \left[\begin{array}{c|c} E_{ij} - E_{i\alpha} E^{\alpha\beta} E_{\beta j} & E_{i\alpha} E^{\alpha\beta} \\ \hline -E^{\alpha\beta} E_{\beta j} & E^{\alpha\beta} \end{array} \right]$$

où $E^{\alpha\beta}$ est l'inverse du bloc $E_{\alpha\beta}$: $E^{\alpha\beta} E_{\beta\gamma} = \delta_\gamma^\alpha$.

Nous avons donc besoin en premier lieu de l'inverse de $(G + B_2)_{\alpha\beta} = (h + B)_{\alpha\beta}$, que nous pouvons écrire

$$\left(\frac{1}{h+B} \right)^{\alpha\beta} = \hat{h}^{\alpha\beta} + \hat{B}^{\alpha\beta}, \quad \text{avec} \quad \begin{cases} \hat{h} = \frac{1}{h+B} h \frac{1}{h-B} \\ \hat{B} = \frac{1}{h+B} (-B) \frac{1}{h-B} \end{cases} \quad (3.41)$$

De plus, $\hat{h}$ et $\hat{B}$ vérifient les relations

$$\hat{B} \hat{h}^{-1} = -h^{-1} B, \quad \hat{h}^{-1} \hat{B} = -B h^{-1}, \quad \hat{B} \hat{h}^{-1} \hat{B} = \hat{h} - h^{-1}$$

En écrivant $\tilde{E} = \tilde{G} + \tilde{B}_2$, on peut alors identifier la métrique T-duale,

$$d\tilde{s}^2 = g_{ij} dy^i dy^j + \hat{h}^{\alpha\beta} (dx^\alpha + B_\alpha)(dx^\beta + B_\beta) \quad (3.42)$$

et le champ antisymétrique $\tilde{B}_2$,

$$\tilde{B}_2 = \frac{1}{2} B_{ij} dy^i \wedge dy^j + \lambda^\alpha \wedge (dx^\alpha + \frac{1}{2} B_\alpha) + \frac{1}{2} \hat{B}^{\alpha\beta} (dx^\alpha + B_\alpha) \wedge (dx^\beta + B_\beta) \quad (3.43)$$

$$\begin{aligned}
w_1 &= -\frac{i}{12} \epsilon^{ijk} V_{i\alpha} [d(V - i\lambda)]_{jk}^\alpha \\
w_t^3 &= \frac{1}{8} \epsilon^{ijk} V_{i\alpha} [d(V - 3i\lambda)]_{jk}^\alpha \\
h_1 &= \frac{1}{12} \epsilon^{ijk} V_i^\alpha V_j^\beta \partial_k B_{\alpha\beta} + \frac{i}{12} \epsilon^{ijk} V_i^\alpha [dB_\alpha - B_{\alpha\beta} d\lambda^\beta]_{jk} \\
&\quad - \frac{1}{12} \epsilon^{ijk} [\partial_k B_{ij} - \partial_k B_{i\alpha} \lambda_j^\alpha + B_{i\alpha} \partial_k \lambda_j^\alpha] \\
h_t^3 &= -\frac{1}{8} \epsilon^{ijk} V_i^\alpha [dB_\alpha]_{jk} - \frac{3i}{8} \epsilon^{ijk} [\partial_k B_{ij} - \partial_k B_{i\alpha} \lambda_j^\alpha + B_{i\alpha} \partial_k \lambda_j^\alpha] \\
\\
w_k^2 &= -\frac{1}{4} V_\alpha^j [d(V - i\lambda)]_{jk}^\alpha \\
w_k^4 &= -\frac{1}{4} V^{j\alpha} [dV_\alpha]_{jk} \\
w_k^5 &= -\frac{1}{4} V_\alpha^j [d(V + i\lambda)]_{jk}^\alpha + \frac{1}{4} h^{\alpha\beta} \partial_k h_{\alpha\beta} \\
h_k^4 &= \frac{1}{4} V^{j\alpha} [dB_\alpha - B_{\alpha\beta} d\lambda^\beta]_{jk} - \frac{i}{4} V^{j\alpha} V_k^\beta \partial_j B_{\alpha\beta} \\
\\
w_{\{ij\}^0}^2 &= \frac{1}{24} \epsilon^{pqk} [d(V - i\lambda)]_{pq}^\alpha [2V_{k\alpha} g_{ij} - 3V_{j\alpha} g_{ik} - 3V_{i\alpha} g_{jk}] \\
w_{\{ij\}^0}^3 &= \frac{1}{24} \epsilon^{pqk} [dV_\alpha]_{pq} [2V_k^\alpha g_{ij} - 3V_j^\alpha g_{ik} - 3V_i^\alpha g_{jk}] \\
h_{\{ij\}^0}^3 &= -\frac{1}{24} \epsilon^{pqk} [dB_\alpha]_{pq} [2V_k^\alpha g_{ij} - 3V_j^\alpha g_{ik} - 3V_i^\alpha g_{jk}]
\end{aligned}$$

TAB. 3.1 – Les composantes de la torsion intrinsèque et du champ de fond H dans la base holomorphe

Ainsi, on peut résumer cette T-dualité par l'échange des trois champs

$$\boxed{h_{\alpha\beta} \leftrightarrow \hat{h}^{\alpha\beta}, \quad B_{\alpha\beta} \leftrightarrow \hat{B}^{\alpha\beta}, \quad B_\alpha \leftrightarrow \lambda^\alpha} \quad (3.44)$$

Naturellement, le vielbein V_α^a se transforme lui aussi. Son image est le vielbein $\hat{V}^{a\alpha}$ de la métrique T-duale $\hat{h}^{\alpha\beta}$, défini par $\hat{V}^{a\alpha}\hat{V}^{a\beta} = \hat{h}^{\alpha\beta}$.

Leur relation est immédiate d'après les formules (3.41),

$$\hat{V}^{a\alpha} = \left(\frac{1}{h+B} \right)^{\alpha\beta} V_\beta^a = V_\beta^a \left(\frac{1}{h-B} \right)^{\beta\alpha} \quad (3.45)$$

Son inverse $\hat{V}_\alpha^a \equiv \hat{h}_{\alpha\beta}\hat{V}^{a\beta}$ tel que $\hat{V}_\alpha^a\hat{V}_\beta^a = \hat{h}_{\alpha\beta}$ s'écrit de même

$$\hat{V}_\alpha^a = (h-B)_{\alpha\beta}V^{a\beta} = V^{a\beta}(h+B)_{\beta\alpha} \quad (3.46)$$

Ces relations s'expriment bien sûr aussi sur les formes $V^\alpha = V_i^\alpha dy^i, \dots$ C'est ainsi qu'elles nous seront utiles par la suite :

$$\hat{V}_\alpha = V_\alpha - B_{\alpha\beta}V^\beta, \quad V^\beta = \hat{V}^\beta + \hat{V}_\alpha\hat{B}^{\alpha\beta} \quad (3.47)$$

Et l'action de la T-dualité s'écrit,

$$\boxed{V_\alpha \leftrightarrow \hat{V}^\alpha, \quad V^\alpha \leftrightarrow \hat{V}_\alpha} \quad (3.48)$$

Cas particulier : $B_{\alpha\beta} = 0$

Dans ce cas, les formules (3.41) se réduisent à $\hat{h} = h^{-1}$ et $\hat{B} = 0$. Remarquons que la condition est conservée par T-dualité. C'est ce qui lui donne toute sa pertinence. De plus, puisque h ne fait que s'inverser par T-dualité, il en est de même de son vielbein : $\hat{V}_\alpha = V_\alpha$ et $\hat{V}^\alpha = V^\alpha$.

Ainsi, mis à part l'échange de B_α et λ^α , la T-dualité ne consiste qu'à monter et descendre les indices le long de la fibre (c'est-à-dire $\alpha, \beta, \dots$).

Une simple observation de la table 3.1 nous offre alors les relations suivantes,

$$\begin{aligned} w_{\{ij\}^0}^2 &\longleftrightarrow (w_3 + ih_3)_{\{ij\}^0} \\ w_k^5 - \frac{1}{4} h^{\alpha\beta} \partial_k h_{\alpha\beta} &= \overline{w_k^2} \longleftrightarrow (w_4 - ih_4)_k \end{aligned}$$

Ainsi, même si la T-dualité échange le fibré tangent et le fibré cotangent le long de la fibre, les expressions dans la base complexe dz^i "ne le voient pas". Au moins dans ce cas particulier.

En ce qui concerne les champs scalaires, on peut facilement calculer la relation entre w_1 et h_1 . Tous deux complexes, faisons apparaître leurs parties réelles et imaginaires :

$$w_1 = w_1^+ - iw_1^-, \quad h_1 = h_1^+ - ih_1^-$$

Encore une fois, un simple calcul donne alors les transformations

$$\begin{aligned} \tilde{w}_1^+ &= h_1^- \\ \tilde{w}_1^- &= w_1^- \\ \tilde{h}_1^+ &= h_1^+ \\ \tilde{h}_1^- &= w_1^+ \end{aligned}$$

Cas général

Plus intéressant encore, si elles ne sont plus complètement découplées dans le cas général, deux combinaisons linéaires survivent :

$$\begin{aligned}\tilde{w}_1^+ - \tilde{h}_1^- &= -(w_1^+ - h_1^-) \\ \tilde{w}_1^- + \tilde{h}_1^+ &= w_1^- + h_1^+\end{aligned}$$

Si l'on définit la 0-forme "complexifiée"¹³ x_1 par

$$x_1 = w_1 - ih_1 = x_1^+ - ix_1^- \quad (3.49)$$

alors les deux transformations précédentes ne sont rien d'autre que

$$\begin{cases} \tilde{x}_1^+ &= -x_1^+ \\ \tilde{x}_1^- &= x_1^- \end{cases} \Leftrightarrow \boxed{\tilde{x}_1 = -\bar{x}_1} \quad (3.50)$$

Démonstration

Par définition, $x_1^+ = w_1^+ - h_1^-$ et $x_1^- = w_1^- + h_1^+$. Donc,

$$\begin{aligned}x_1^+ &= -\frac{1}{12} \epsilon^{ijk} \{V_{i\alpha} [d\lambda^\alpha]_{jk} - V_i^\alpha [dB_\alpha - B_{\alpha\beta} d\lambda^\beta]_{jk}\} \\ x_1^- &= \frac{1}{12} \epsilon^{ijk} \{V_{i\alpha} [dV^\alpha]_{jk} + V_i^\alpha V_j^\beta \partial_k B_{\alpha\beta} - [\partial_k B_{ij} - \partial_k B_{i\alpha} \lambda_j^\alpha + B_{i\alpha} \partial_k \lambda_j^\alpha]\}\end{aligned}$$

Ce qui est plus clair en utilisant une notation de formes (en remplaçant ϵ^{ijk} par $dy^i \wedge dy^j \wedge dy^k$),

$$\begin{aligned}x_1^+ &\propto V_\alpha \wedge d\lambda^\alpha - V^\alpha \wedge dB_\alpha + V^\alpha \wedge B_{\alpha\beta} d\lambda^\beta \\ &\propto [V_\alpha - B_{\alpha\beta} V^\beta] \wedge d\lambda^\alpha - V^\alpha \wedge dB_\alpha \\ &\propto \hat{V}_\alpha \wedge d\lambda^\alpha - V^\alpha \wedge dB_\alpha \\ x_1^- &\propto \underbrace{2V_\alpha \wedge dV^\alpha - V^\alpha \wedge dB_{\alpha\beta} \wedge V^\beta}_Q - [2dB - dB_\alpha \wedge \lambda^\alpha - d\lambda^\alpha \wedge B_\alpha]\end{aligned}$$

où l'on a utilisé (3.47) pour réécrire x_1^+ , dont la transformation est alors immédiate. De même, le dernier terme de x_1^- est invariant. Seule la quantité Q nécessite un peu plus de calcul. Notons d'abord que, par symétrie de $h_{\alpha\beta}$,

$$V^\alpha B_{\alpha\beta} \wedge V^\beta = V^\alpha (h + B)_{\alpha\beta} \wedge V^\beta = \hat{V}_\beta \wedge V^\beta$$

puis que le développement de la différentielle s'écrit

$$\begin{aligned}d[V^\alpha B_{\alpha\beta} \wedge V^\beta] &= dV^\alpha B_{\alpha\beta} \wedge V^\beta - V^\alpha \wedge dB_{\alpha\beta} \wedge V^\beta - V^\alpha B_{\alpha\beta} \wedge dV^\beta \\ &= -V^\alpha \wedge dB_{\alpha\beta} \wedge V^\beta - 2V^\alpha B_{\alpha\beta} \wedge dV^\beta\end{aligned}$$

¹³ x_1 est la composante sur $\tilde{\Omega}$ de la forme de Kähler complexifiée $J + iB_2$.

Cela nous permet de réécrire Q , toujours en utilisant (3.47),

$$\begin{aligned}
 Q &= 2V_\alpha \wedge dV^\alpha + d[V^\alpha B_{\alpha\beta} \wedge V^\beta] + 2V^\alpha B_{\alpha\beta} \wedge dV^\beta \\
 &= 2[V_\beta + V^\alpha B_{\alpha\beta}] \wedge dV^\beta + d[\hat{V}_\beta \wedge V^\beta] \\
 &= 2\hat{V}_\beta \wedge dV^\beta + [d\hat{V}_\beta \wedge V^\beta - \hat{V}_\beta \wedge dV^\beta] \\
 &= \hat{V}_\beta \wedge dV^\beta + V^\beta \wedge d\hat{V}_\beta
 \end{aligned}$$

Sous cette forme, il est alors clair que Q est invariant, comme anticipé.

3.4 Le spineur invariant

Jusqu'ici, nous avons caractérisé la géométrie, à savoir la structure $SU(3)$, par l'existence de deux tenseurs invariants, J et Ω . Mais elle peut aussi être caractérisée de manière équivalente par l'existence d'un spineur invariant, que nous noterons ϵ . Il se décompose sur les états propres de chiralité selon $\epsilon_+ + \epsilon_-$.

Nous nous sommes intéressés aux tenseurs parce qu'ils mesuraient l'écart à la géométrie de Calabi-Yau. En effet, $dJ = d\Omega = 0$ pour une variété de Calabi-Yau. Pour une structure $SU(3)$, cette différentielle — et donc les composantes W_i — est alors une mesure de la déviation créée par la présence du champ de fond H . Mais l'histoire se répète pour le spineur ϵ . Dans une variété de Calabi-Yau, sa dérivée covariante est nulle : $D_\mu \epsilon = 0$. Ce n'est pas le cas ici, et c'est cette quantité qui va remplacer les différentielles dJ et $d\Omega$.

Nous voulons donc maintenant relier ces différentielles à $D_\mu \epsilon$, pour pouvoir exprimer les composantes de celle-ci sur les représentations de $SU(3)$ en fonction des composantes W_i . Nous verrons alors comment exprimer plus naturellement les transformations de T-dualité sur ces champs.

Notons que dans cette section, nous utiliserons des indices grecs ($\mu, \nu, \rho, \dots$) pour les coordonnées réelles à 6 dimensions (au lieu de $I, J, K, \dots$). α et β sont par contre conservées pour les coordonnées sur la fibre comme auparavant.

3.4.1 Décomposition de $D_\mu \epsilon$ en représentations

Commençons par décomposer $D_\mu \epsilon$ sur la base des spineurs $(\epsilon, \gamma\epsilon, \gamma^\nu \epsilon)$:

$$D_\mu \epsilon = q_\mu \epsilon + iq'_\mu \gamma\epsilon + iq_{\mu\nu} \gamma^\nu \epsilon \quad (3.51)$$

où les champs q_μ , q'_μ et $q_{\mu\nu}$ sont réels.

q_μ et q'_μ sont dans la représentation $\mathbf{3} \oplus \bar{\mathbf{3}}$. $q_{\mu\nu}$ est un tenseur à deux indices, dans la représentation réductible

$$\mathbf{36} = \mathbf{1} \oplus \mathbf{1} \oplus \mathbf{3} \oplus \bar{\mathbf{3}} \oplus \mathbf{6} \oplus \bar{\mathbf{6}} \oplus \mathbf{8} \oplus \mathbf{8}$$

On le décomposera donc sous la forme

$$q_{\mu\nu} = q_+ G_{\mu\nu} + q_- J_{\mu\nu} + \psi_{\mu\nu\rho} q_0^\rho + s_{\mu\nu}^0 + b_{\mu\nu} + a_{\mu\nu} \quad (3.52)$$

où q_+ , q_- , q_0^ρ , $s_{\mu\nu}^0$, $b_{\mu\nu}$ et $a_{\mu\nu}$ sont réels, et

- q_+ et q_- sont les deux scalaires,
- q_0^ρ est une 1-forme dans $\mathbf{3} \oplus \bar{\mathbf{3}}$,
- $a_{\mu\nu}$ est antisymétrique dans $\mathbf{8}$, donc une $(1, 1)$ -forme primitive

$$J^{\mu\nu} a_{\mu\nu} = 0, \quad \Omega^{\mu\nu\rho} a_{\mu\nu} = \bar{\Omega}^{\mu\nu\rho} a_{\mu\nu} = 0$$

- $s_{\mu\nu} = s_{\mu\nu}^0 + b_{\mu\nu}$ est un tenseur symétrique de trace nulle ($s_{\mu\nu} G^{\mu\nu} = 0$), dont $s_{\mu\nu}^0$ est la partie $(2, 0) \oplus (0, 2)$ de $\mathbf{6} \oplus \bar{\mathbf{6}}$, et $b_{\mu\nu}$ la partie $(1, 1)$ de $\mathbf{8}$.

Avant de poursuivre, voyons quelques résultats utiles sur ces composantes. Rappelons que P , projecteur holomorphe, et $\bar{P}$ son projecteur conjugué antiholomorphe, ont été définis en (3.25). On peut remarquer en particulier que

$$P_\mu{}^\nu = \frac{1}{2} (\delta_\mu{}^\nu - iJ_\mu{}^\nu) = \bar{P}^\nu{}_\mu$$

De plus, ces projecteurs sont très liés à la forme holomorphe Ω , puisque

$$\Omega_{\mu\nu\sigma} \bar{\Omega}^{\sigma\rho\lambda} = 8[P_\mu{}^\rho P_\nu{}^\lambda - P_\mu{}^\lambda P_\nu{}^\rho] \quad (3.53)$$

$$\Omega_{\mu\nu\rho} \bar{\Omega}^{\nu\rho\sigma} = 16P_\mu{}^\sigma \quad (3.54)$$

On a bien sûr d'autres relations plus évidentes, sur des tenseurs dont le degré holomorphe est identifié :

- Ω est une $(3, 0)$ -forme, donc $P_\mu{}^\sigma \Omega_{\sigma\nu\rho} = 0$, i.e. $J_\mu{}^\sigma \Omega_{\sigma\nu\rho} = i\Omega_{\mu\nu\rho}$.
- De même pour $\bar{\Omega}$, $J_\mu{}^\sigma \bar{\Omega}_{\sigma\nu\rho} = -i\bar{\Omega}_{\mu\nu\rho}$.
- $a_{\mu\nu}$ est une $(1, 1)$ -forme, donc $P_\mu{}^\rho P_\nu{}^\sigma a_{\rho\sigma} = 0$, ce qui implique que $J_\mu{}^\sigma a_{\sigma\nu}$ est un tenseur symétrique $(1, 1)$, de trace nulle.
- $b_{\mu\nu}$ est un tenseur symétrique $(1, 1)$, donc $P_\mu{}^\rho P_\nu{}^\sigma b_{\rho\sigma} = 0$, ce qui implique que $J_\mu{}^\sigma b_{\sigma\nu}$ est une $(1, 1)$ -forme.
- $s_{\mu\nu}^0$ est un tenseur symétrique $(2, 0) \oplus (0, 2)$, donc $P_\mu{}^\rho \bar{P}_\nu{}^\sigma s_{\rho\sigma}^0 = 0$, ce qui implique que $2P_\mu{}^\sigma s_{\sigma\nu}^0 = s_{\mu\nu}^0 - iJ_\mu{}^\sigma s_{\sigma\nu}^0$ est un tenseur symétrique $(2, 0) \oplus (0, 2)$, par ailleurs de trace nulle.

Il nous faut maintenant exprimer ces composantes en fonction des W_i . On pourra remarquer qu'un simple comptage des degrés de liberté nous donne 42 pour les W_i , tandis qu'on arrive ici à 48. La faute à q_μ et q'_μ , qui semblent indépendants, mais qui forment en fait un champ complexe holomorphe

$$q_\mu - iq'_\mu = 2P_\mu{}^\nu q_\nu$$

Cela réduit au 42 attendu le nombre de degrés de liberté dans cette nouvelle description.

De plus, les remarques précédentes permettent d'anticiper sur les identifications entre les deux séries de composantes. On voit en effet que les différents objets des deux séries s'appartiennent par nature (scalaires, vecteurs...) :

La série W	La série q
Un scalaire complexe W_1	Deux scalaires réels q_+ et q_-
Deux vecteurs W_4 et W_5	Deux vecteurs q et q_0
Une $(1, 1)$ -forme complexe W_2	Deux $(1, 1)$ -formes réelles $a_{\mu\nu}$ et $J_\mu^\sigma b_{\sigma\nu}$
Une $(2, 1) \oplus (1, 2)$ -forme réelle W_3 ou sous la forme w_{ij}^3	Un tenseur symétrique $(2, 0) \oplus (0, 2)$ réel s^0
un tenseur symétrique $(2, 0) \oplus (0, 2)$	

3.4.2 Relations avec les tenseurs

Les tenseurs J et Ω s'obtiennent à partir du spineur ϵ grâce aux bilinéaires

$$\epsilon^\dagger \gamma_{\mu\nu} \gamma \epsilon = iJ_{\mu\nu}, \quad \epsilon^\dagger \gamma_{\mu\nu\rho} \epsilon = i\psi_{\mu\nu\rho}, \quad \epsilon^\dagger \gamma_{\mu\nu\rho} \gamma \epsilon = *\psi_{\mu\nu\rho} \quad (3.55)$$

où ψ et $*\psi$ sont les parties réelles et imaginaires de Ω : $i\Omega = i\psi + *\psi$.

On peut déduire de ces bilinéaires les expressions des différentielles dJ et $d\Omega$,

$$dJ = 2q_\mu e^\mu \wedge J + \text{Im}(q_\mu^\sigma \Omega_{\sigma\nu\rho}) e^\mu \wedge e^\nu \wedge e^\sigma \quad (3.56)$$

$$d\Omega = 2(q_\mu + iq'_\mu) e^\mu \wedge \Omega + 4iP_\mu^\sigma q_{\nu\sigma} e^\mu \wedge e^\nu \wedge J \quad (3.57)$$

Démonstration

Pour la première formule, d'après la décomposition (3.51) de $D_\mu \epsilon$,

$$\begin{aligned} i\partial_\mu J_{\nu\rho} |^{[\mu\nu\rho]} &= D_\mu [\epsilon^\dagger \gamma_{\mu\nu} \gamma \epsilon] |^{[\mu\nu\rho]} \\ &= q_\mu \epsilon^\dagger \gamma_{\nu\rho} \gamma \epsilon - iq'_\mu \epsilon^\dagger \gamma \gamma_{\nu\rho} \gamma \epsilon - iq_{\mu\sigma} \epsilon^\dagger \gamma^\sigma \gamma_{\nu\rho} \gamma \epsilon \\ &\quad + q_\mu \epsilon^\dagger \gamma_{\nu\rho} \gamma \epsilon + iq'_\mu \epsilon^\dagger \gamma_{\nu\rho} \gamma \gamma \epsilon + iq_{\mu\sigma} \epsilon^\dagger \gamma_{\nu\rho} \gamma \gamma^\sigma \epsilon |^{[\mu\nu\rho]} \\ &= 2iq_\mu J_{\nu\rho} - 2iq_{\mu\sigma} \epsilon^\dagger \{\gamma_{\nu\rho}, \gamma^\sigma\} \gamma \epsilon |^{[\mu\nu\rho]} \\ &= 2iq_\mu J_{\nu\rho} - 2iq_\mu^\sigma (*\psi)_{\nu\rho\sigma} |^{[\mu\nu\rho]} \end{aligned}$$

$$\text{avec } \begin{aligned} \gamma^\dagger &= \gamma, & \gamma\gamma_\mu &= -\gamma_\mu\gamma, & \gamma\gamma &= \gamma_\mu\gamma^\mu = 1, & \{\gamma_{\nu\rho}, \gamma^\sigma\} &= 2\gamma_{\nu\rho}^\sigma, \\ \gamma_\mu^\dagger &= \gamma_\mu, & \gamma\gamma_{\mu\nu} &= \gamma_{\mu\nu}\gamma. \end{aligned}$$

Il ne reste alors plus qu'à revenir à la différentielle,

$$\begin{aligned} dJ &= \frac{1}{2} \partial_\mu J_{\nu\rho} e^\mu \wedge e^\nu \wedge e^\rho \\ &= [q_\mu J_{\nu\rho} - q_\mu^\sigma (*\psi)_{\sigma\nu\rho}] e^\mu \wedge e^\nu \wedge e^\rho \end{aligned}$$

d'où la formule (3.56). On fait de même pour la formule (3.57).

Puis, en remplaçant $q_{\mu\nu}$ par son développement en composantes, avec les notations

$q = q_\mu e^\mu$ et $q_0 = q_{0\mu} e^\mu$, on identifie les W_i :

$$dJ = -\frac{3}{2} \underbrace{\text{Im}[4(q_+ - iq_-)]}_{W_1} \bar{\Omega} + \underbrace{[2q - 4q_0]}_{W_4} \wedge J + \underbrace{s_\mu^{0\sigma} (\text{Im } \Omega)_{\sigma\nu\rho} e^\mu \wedge e^\nu \wedge e^\rho}_{W_3} \quad (3.58)$$

$$d\Omega = \underbrace{4(q_+ - iq_-)}_{W_1} J^2 + \underbrace{2(J_\mu^\sigma b_{\sigma\nu} - ia_{\mu\nu}) e^\mu \wedge e^\nu \wedge J}_{W_2} + \underbrace{4\bar{P}(q - q_0)}_{W_5} \wedge \Omega \quad (3.59)$$

Démonstration

Pour changer, voyons cette fois la deuxième formule. Dans (3.57), il nous faut expliciter le deuxième terme. Tout d'abord,

$$P_\mu^\sigma q_{\nu\sigma} = \frac{1}{2} (\delta_\mu^\sigma - iJ_\mu^\sigma) (q_+ G_{\nu\sigma} + q_- J_{\nu\sigma} + s_{\nu\sigma}^0 + b_{\nu\sigma} + a_{\nu\sigma}) + \frac{1}{2} \Omega_{\nu\mu\rho} q_0^\rho$$

Ensuite, en contractant sur $e^\mu \wedge e^\nu$, on ne conserve que les termes antisymétriques. Ainsi, G , s^0 , b , Js^0 et Ja disparaissent. Parmi les termes restants, un seul nécessite encore un peu de travail : $\Omega_{\nu\mu\rho} q_0^\rho e^\mu \wedge e^\nu \wedge J$ est une (3,1)-forme et contribue donc exclusivement à W_5 . Pour l'identifier, nous allons la réécrire $\bar{Q} \wedge \Omega$, où Q est une (1,0)-forme.

$$\begin{aligned} \bar{Q} &= -\bar{\Omega}_\perp (\bar{Q} \wedge \Omega) = -\bar{\Omega}_\perp [\Omega_{\nu\mu\rho} q_0^\rho e^\mu \wedge e^\nu \wedge J] \\ &= -\frac{1}{8} \bar{\Omega}^{\mu\nu\kappa} \Omega_{\nu\mu\rho} q_0^\rho J_{\kappa\lambda} e^\lambda \\ &= -\frac{i}{8} \bar{\Omega}_{\mu\nu\lambda} \Omega^{\nu\mu\rho} q_{0\rho} e^\lambda \\ &= \frac{16i}{8} \bar{P}_\lambda{}^\rho q_{0\rho} e^\lambda = 2i\bar{P}q_0 \end{aligned}$$

Pour pouvoir conclure en regroupant les composantes de $q_{\mu\nu}$ dans sa décomposition, nous devons encore inverser les relations que nous venons d'obtenir.

La composante W_1

Si l'inversion est évidente, nous allons tout de même écrire, en anticipant un peu, ce que vaut $q_+ G_{\mu\nu} + q_- J_{\mu\nu}$, qui intervient dans l'expression de $q_{\mu\nu}$. Pour cela, remarquons que

$$(q_+ - iq_-) \bar{P}_{\mu\nu} = \frac{1}{2} [q_+ G_{\mu\nu} + iq_+ J_{\mu\nu} - q_- G_{\mu\nu} + q_- J_{\mu\nu}]$$

Donc

$$q_+ G_{\mu\nu} + q_- J_{\mu\nu} = 2\text{Re} [(q_+ - iq_-) \bar{P}_{\mu\nu}] = \text{Re} \left[\frac{1}{2} w_1 \bar{P}_{\mu\nu} \right] \quad (3.60)$$

Les composantes W_4 et W_5

Notons tout d'abord que $W_5 = W_5^+ - iW_5^- = 2PW_5^+$. Puisque q et q_0 sont réels, on peut donc identifier

$$\begin{cases} W_4 = 2q - 4q_0 \\ W_5^+ = 2q - 2q_0 \end{cases} \Leftrightarrow \begin{cases} q = \frac{1}{2} (2W_5^+ - W_4) \\ q_0 = \frac{1}{2} (W_5^+ - W_4) \end{cases}$$

En composantes holomorphes, souvenons-nous que $w_5 = W_5$ et $w_4 = PW_4$:

$$Pq = \frac{1}{2} (w_5 - w_4), \quad Pq_0 = \frac{1}{4} (w_5 - 2w_4) \quad (3.61)$$

La composante W_2

W_2 s'exprime en fonction de $a_{\mu\nu}$ et $b_{\mu\nu}$,

$$W_2 = 2(J_\mu^\sigma b_{\sigma\nu} - ia_{\mu\nu}) e^\mu \wedge e^\nu = \frac{1}{2} W_{2\mu\nu} e^\mu \wedge e^\nu$$

On en déduit leur expression en séparant partie réelle et partie imaginaire

$$b_{\mu\nu} = -\frac{1}{4} J_\mu^\sigma W_{2\sigma\nu}^+, \quad a_{\mu\nu} = \frac{1}{4} W_{2\mu\nu}^- \quad (3.62)$$

De plus, remarquons que

$$\bar{P}_\mu^\sigma W_{2\sigma\nu} = \frac{1}{2} [W_{2\mu\nu}^+ - iW_{2\mu\nu}^- + iJ_\mu^\sigma W_{2\sigma\nu}^+ + J_\mu^\sigma W_{2\sigma\nu}^-]$$

donc la somme $a + b$, qui apparaît dans $q_{\mu\nu}$, peut s'écrire en regroupant W_2^+ et W_2^- ,

$$b_{\mu\nu} + a_{\mu\nu} = -\frac{1}{2} \text{Im}(\bar{P}_\mu^\sigma W_{2\sigma\nu}) = \text{Re} \left[\frac{i}{2} \bar{P}_\mu^\sigma w_{\sigma\nu}^2 \right]$$

La composante W_3

On a trouvé pour W_3

$$W_3 = s_\mu^{0\sigma} (\text{Im } \Omega)_{\sigma\nu\rho} e^\mu \wedge e^\nu \wedge e^\rho = \frac{1}{6} W_{3\mu\nu\rho} e^\mu \wedge e^\nu \wedge e^\rho$$

Ce qui donne, en tenant compte de l'antisymétrisation,

$$W_{3\mu\nu\rho} = 2[s_\mu^{0\sigma} (\text{Im } \Omega)_{\sigma\nu\rho} + s_\nu^{0\sigma} (\text{Im } \Omega)_{\sigma\rho\mu} + s_\rho^{0\sigma} (\text{Im } \Omega)_{\sigma\mu\nu}]$$

Pour isoler s^0 , on contracte avec $(\text{Im } \Omega)_{\nu\rho\kappa}$, en se souvenant pour développer le produit des parties imaginaires que Ω ne se contracte que sur $\bar{\Omega}$, et vice versa.

$$\begin{aligned} W_{3\mu\nu\rho} (\text{Im } \Omega)^{\nu\rho\kappa} &= \frac{1}{2} [s_\mu^{0\sigma} (\Omega_{\sigma\nu\rho} \bar{\Omega}^{\nu\rho\kappa} + \text{c. c.}) + 2s_\rho^{0\sigma} (\Omega_{\sigma\mu\nu} \bar{\Omega}^{\nu\rho\kappa} + \text{c. c.})] \\ &= \frac{1}{2} [s_\mu^{0\sigma} (16P_\sigma^\kappa + 16\bar{P}_\sigma^\kappa) + 16s_\rho^{0\sigma} (P_\sigma^\rho P_\mu^\kappa - P_\sigma^\kappa P_\mu^\rho + \text{c. c.})] \\ &= 8s_\mu^{0\kappa} \end{aligned}$$

les autres termes sont nuls parce que $s^0 \in (2, 0) \oplus (0, 2)$, et est de trace nulle sur G et J . Finalement, on a donc

$$s_{\mu\nu}^0 = \frac{1}{8} W_{3\mu\rho\sigma} (\text{Im } \Omega)^{\rho\sigma}{}_{\nu} = \text{Re} \left[-\frac{i}{2} w_{\mu\nu}^3 \right] \quad (3.63)$$

Le facteur final $\frac{i}{2}$ mérite explication. En effet, nous avons défini w_{ij}^3 comme la contraction $w_{ipq}^3 \Omega^{pq}{}_j$, où les indices p et q sont les deux indices antiholomorphes de w_{ipq}^3 ; la somme va donc de 1 à 3. Or ici la contraction se fait sur les indices réels ρ et σ ; la somme va de 1 à 6. Cependant, l'holomorphie de Ω assure que le résultat est le même, on a simplement compté 4 fois trop de termes, d'où le facteur 4.

L'expression de $q_{\mu\nu}$

Pour finir cette section, il ne reste qu'à rassembler les morceaux que nous venons d'énumérer :

$$q_{\mu\nu} = \text{Re} \left[\frac{1}{2} w_1 \bar{P}_{\mu\nu} + \frac{1}{4} \Omega_{\mu\nu\rho} (\bar{w}_5 - 2\bar{w}_4)^\rho - \frac{i}{2} w_{\mu\nu}^3 + \frac{i}{2} \bar{P}_\mu{}^\sigma w_{\sigma\nu}^2 \right]$$

Remarquons qu'à l'intérieur de cette partie réelle, l'indice ν est toujours holomorphe : sur $\bar{P}_{\mu\nu}$, tenseur $(1, 1)$ qui projette sur l'indice antiholomorphe μ , sur Ω évidemment, comme sur w_3 , tenseur $(2, 0)$, et enfin sur w_2 , tenseur $(1, 1)$ dont le premier indice est projeté par $\bar{P}$ sur μ antiholomorphe. Cela signifie deux choses. D'abord qu'appliquer $P_n{}^\nu$ à $q_{\mu\nu}$, où n représente explicitement un indice holomorphe, fait disparaître la partie réelle :

$$2P_n{}^\nu q_{\mu\nu} = \frac{1}{2} w_1 \bar{P}_{\mu n} + \frac{1}{4} \Omega_{\mu n\rho} (\bar{w}_5 - 2\bar{w}_4)^\rho - \frac{i}{2} w_{\mu n}^3 + \frac{i}{2} \bar{P}_\mu{}^\sigma w_{\sigma n}^2$$

Et ensuite que le premier indice de chaque terme étant soit holomorphe, soit antiholomorphe, appliquer le deuxième projecteur pour séparer les composantes $(1, 1)$ et $(2, 0)$ est immédiat :

$$q_{mn} = \frac{1}{4} \Omega_{mnp} (\bar{w}_5 - 2\bar{w}_4)^p - \frac{i}{4} w_{mn}^3, \quad q_{\bar{m}\bar{n}} = \frac{1}{4} w_1 \bar{P}_{\bar{m}\bar{n}} + \frac{i}{4} w_{\bar{m}\bar{n}}^2$$

On notera que comme pour w_3 dans le paragraphe précédent, en remplaçant la somme sur ρ , indice réel (de 1 à 6), en somme sur p , indice holomorphe (de 1 à 3), dans la contraction de Ω , on a fait apparaître un facteur 2. De plus, on pourra remarquer que $\Omega_{mnp} = -i\epsilon_{mnp}$ et $P_{m\bar{n}} = g_{m\bar{n}}$. Ainsi, q_{mn} et $q_{\bar{m}\bar{n}} = \overline{q_{mn}}$ s'écrivent finalement

$$q_{mn} = -\frac{i}{4} \epsilon_{mnp} (\bar{w}_5 - 2\bar{w}_4)^p - \frac{i}{4} w_{mn}^3, \quad q_{\bar{m}\bar{n}} = \frac{1}{4} \bar{w}_1 g_{\bar{m}\bar{n}} - \frac{i}{4} w_{\bar{m}\bar{n}}^2 \quad (3.64)$$

Quant à la dernière composante, elle a pour valeur holomorphe

$$2P_m{}^\mu q_\mu = q_m - iq'_m = (w_5 - w_4)_m \quad (3.65)$$

3.4.3 Couplage au champ de fond

Attaquons nous maintenant au couplage du champ de fond H , ce que nous n'avions pas pu faire jusqu'ici, faute d'angle d'attaque. Dans la description spinorielle, ce couplage intervient naturellement : il se glisse dans la contribution de H à la dérivée covariante. Ainsi, dans notre expression de départ $D_\mu \epsilon$, les composantes H_i vont s'ajouter au W_i que nous venons d'identifier.

Définissons

$$D_\mu^H = D_\mu + \frac{1}{8} H_{\mu\nu\rho} \gamma^{\nu\rho} \quad (3.66)$$

En notant que

$$\gamma^{\nu\rho} \epsilon = i J^{\nu\rho} \gamma \epsilon + i \psi^{\nu\rho}{}_\sigma \gamma^\sigma \epsilon$$

nous pouvons calculer les contributions de ce nouveau terme aux composantes q , renommées Q pour l'occasion

$$D_\mu^H \epsilon = Q_\mu \epsilon + i Q'_\mu \gamma \epsilon + i Q_{\mu\nu} \gamma^\nu \epsilon \quad (3.67)$$

où bien sûr

$$\begin{aligned} Q_\mu &= q_\mu \\ Q'_\mu &= q'_\mu + \frac{1}{8} H_{\mu\rho\sigma} J^{\rho\sigma} \\ Q_{\mu\nu} &= q_{\mu\nu} + \frac{1}{8} H_{\mu\rho\sigma} \psi^{\rho\sigma}{}_\nu \end{aligned}$$

Or H se décompose en représentations selon (3.36)

$$H = -\frac{3}{2} \text{Im}(H_1 \bar{\Omega}) + H_4 \wedge J + H_3$$

et les contractions précédentes s'écrivent donc

$$\begin{aligned} H_{\mu\rho\sigma} J^{\rho\sigma} &= 8(J \lrcorner H)_\mu = 8J \lrcorner (H_4 \wedge J)_\mu = 4H_{4\mu} \\ H_{\mu\rho\sigma} \Omega^{\rho\sigma}{}_\nu &= -\frac{3}{4i} H_1 \bar{\Omega}_{\mu\rho\sigma} \Omega^{\rho\sigma}{}_\nu + [H_4 \wedge J]_{\mu\rho\sigma} \Omega^{\rho\sigma}{}_\nu + H_{3\mu\rho\sigma} \Omega^{\rho\sigma}{}_\nu \\ &= \frac{3i}{4} H_1 \cdot 16 \bar{P}_{\mu\nu} + [H_{4\mu} J_{\rho\sigma} + H_{4\rho} J_{\sigma\mu} + H_{4\sigma} J_{\mu\rho}] \Omega^{\rho\sigma}{}_\nu + 4h_{\mu\nu}^3 \\ &= 12ih_1 \bar{P}_{\mu\nu} - 2i\Omega_{\mu\nu\sigma} h_4^\sigma + 4h_{\mu\nu}^3 \end{aligned}$$

Dans l'expression de $Q_{\mu\nu}$, c'est la partie réelle du deuxième terme qui apparaît, ce qui donne

$$\begin{aligned} Q_{\mu\nu} &= \text{Re} \left[\frac{1}{2} (w_1 + 3ih_1) \bar{P}_{\mu\nu} + \frac{1}{4} \Omega_{\mu\nu\rho} (\bar{w}_5 - 2\bar{w}_4 - ih_4)^\rho \right. \\ &\quad \left. - \frac{i}{2} (w_3 - ih_3)_{\mu\nu} + \frac{i}{2} \bar{P}_\mu{}^\sigma w_{\sigma\nu}^2 \right] \end{aligned}$$

Comme précédemment, on en déduit la valeur des composantes Q_{mn} et $Q_{m\bar{n}}$

$$Q_{mn} = -\frac{i}{4} \epsilon_{mnp} (\bar{w}_5 - 2\bar{w}_4 - ih_4)^p - \frac{i}{4} (w_3 - ih_3)_{mn} \quad (3.68)$$

$$Q_{m\bar{n}} = \frac{1}{4} (\bar{w}_1 - 3i\bar{h}_1) g_{m\bar{n}} - \frac{i}{4} w_{m\bar{n}}^2 \quad (3.69)$$

De même pour le vecteur :

$$Q_\mu - iQ'_\mu = q_\mu - iq'_\mu - \frac{i}{8} H_{\mu\rho\sigma} J^{\rho\sigma} = (w_5 - w_4)_\mu - \frac{i}{2} H_{4\mu}$$

D'où l'expression de la composante holomorphe

$$Q_m - iQ'_m = P_m{}^\mu (Q_\mu - iQ'_\mu) = (w_5 - w_4 - \frac{i}{2} h_4)_m \quad (3.70)$$

Conclusion et perspectives

Dans le premier chapitre, nous avons vu que le champ de fond B antisymétrique peut être codé dans la non-commutativité des branes, et déformer les théories de jauge qu'elles portent. On sait écrire cette transformation dans le cas abélien, grâce à la transformation de formalité de Kontsevitch. De plus, l'utilisation conjointe de la formalité et d'une formulation récursive de la transformation de Seiberg-Witten permet d'en calculer plus rapidement les termes explicites (cf (1.193)).

Par contre, la situation est moins claire dans le cas non-abélien. Le calcul explicite direct est plus difficile, bien qu'il semble fonctionner pour partie au moins avec le même genre de récursivité que dans le cas abélien (cf 1.8). Peut-être est-il possible d'inclure l'algèbre de jauge comme des dimensions supplémentaires au sein d'une théorie abélienne plus large [30]. Peut-être pourrait-on utiliser l'équivalence de Morita pour capitaliser les formules déjà obtenues. Cette équivalence regroupe en effet les théories de jauge sur un tore en orbites de $SL(2, \mathbb{Z})$. Elle agit à la fois sur le rang du groupe de jauge et sur le paramètre de non-commutativité, reliant ainsi les théories non-abéliennes ordinaires et les théories abéliennes non-commutatives ayant un paramètre rationnel.

Nous nous sommes demandés dans le chapitre 2 quels étaient les champs d'un paquet de N M5-branes. L'approche directe, en utilisant les multiplets supersymétriques à 6 dimensions (cf 2.1), n'a rien donné pour une raison encore inconnue. Plus exactement, nous n'avons pas été capables de reproduire l'anomalie attendue en N^3 . Par contre, j'ai présenté en 2.2 un formalisme qui permet de décrire de manière unifiée les différentes théories de jauge (1-formes et 2-formes, abéliennes et non-abéliennes, voir fig.2.1). En particulier les théories de 2-formes non-abéliennes dont on s'attend à ce qu'elles soient les théories de basse énergie du paquet de M5-branes. Enfin, j'ai montré dans la section 2.3 comment on pouvait obtenir, à partir des configurations géométriques de membranes M2, un comptage des degrés de liberté d'ordre N^3 .

Malgré tout, nous n'avons pas répondu à la question initiale. Et mon travail appelle de nombreux prolongements. En particulier, peut-on relier les états géométriques en cylindre et en tripode aux différents champs qui apparaissent dans la théorie de jauge de 2-formes? C'est le cas pour les D-branes, puisque les composantes du champ A dans l'algèbre de jauge sont associées aux cordes qui relient les branes. On pourrait alors penser que les cylindres de M2, identifiables aux cordes, correspondent à la 1-forme A , tandis que les tripodes correspondent à la 2-forme B . Encore faut-il que le comptage des degrés de liberté puisse s'accommoder de cette identification (cf fig.2.3) : il y aurait N^3 objets dans l'algèbre $\mathfrak{h}$ et seulement N^2 dans l'algèbre $\mathfrak{g}$, plongé dans $\text{Der}(\mathfrak{h})$ par l'action $d\alpha \dots$

Finalement, nous avons vu dans le troisième chapitre que si la présence d'un flux de Neveu-Schwarz non trivial oblige la géométrie de la variété de compactification à dévier du cas Calabi-Yau, on peut toujours définir une symétrie-miroir à la SYZ, qui mélange géométrie et champs de fond. Les expressions dans la base holomorphe semblent les plus adaptées pour écrire explicitement ces transformations, comme si l'échange des fibrés tangent et cotangent n'avait pas d'effet sur elles. De plus, les formules (3.68), (3.69) et (3.70) auxquelles nous sommes finalement parvenus suggèrent fortement que les deux vides supersymétriques correspondant à ϵ_+ et ϵ_- sont échangés par la T-dualité. Ainsi, T-dualité et supersymétrie seraient naturellement compatibles. Cette simplicité apparente ne doit cependant pas masquer qu'elle s'accompagne de la transformation $T + H \longleftrightarrow -(T + H)$, où le mélange de la torsion et du champ de Neveu-Schwarz généralise l'échange des modules de la structure complexe et de la structure de Kähler d'une symétrie-miroir classique.

Ainsi formulée, nous pouvons conjecturer que cette T-dualité est toujours valide dans le cas d'une variété à structure $SU(3)$ générale. En effet, nous avons travaillé ici avec une fibration T^3 explicite. D'autre part, les résultats présentés concernent en majeure partie le cas $B_{\alpha\beta} = 0$, pour lequel les transformations sont simplifiées. Cependant la structure spinorielle est elle aussi modifiée si cette composante est non nulle, apportant une contribution supplémentaire. Nous nous attendons à ce que le résultat final ne soit pas modifié, comme le suggère la transformation des champs scalaires qui reste inchangée. Enfin, il serait intéressant de pouvoir inclure les champs de Ramond-Ramond dans cette description. Dans une configuration supersymétrique, leur contribution à la dérivée covariante devrait être l'opposée de celle de la torsion et du champ de Neveu-Schwarz, puisque la somme totale est alors nulle. Cela suggère qu'elle se transforme de manière similaire à ce que nous avons trouvé jusqu'à présent.

Annexe A

Towards an explicit expression of the Seiberg-Witten map

Cet article a été publié par la revue *Journal of High Energy Physics*, sous la référence JHEP **0206** (2002) 016. Il est aussi référencé sur www.arXiv.org comme hep-th/0112027.

Towards an explicit expression of the Seiberg-Witten map at all orders

Stéphane Fidanza
fidanza@cphpt.polytechnique.fr

*Centre de Physique Théorique, École Polytechnique¹
91123 Palaiseau Cedex, France*

Abstract

The Seiberg-Witten map links noncommutative gauge theories to ordinary gauge theories, and allows to express the noncommutative variables in terms of the commutative ones. Its explicit form can be found order by order in the noncommutative parameter θ and the gauge potential A by the requirement that gauge orbits are mapped on gauge orbits. This of course leaves ambiguities, corresponding to gauge transformations, and there is an infinity of solutions. Is there one better, clearer than the others ? In the abelian case, we were able to find a solution, linked by a gauge transformation to already known formulas, which has the property of admitting a recursive formulation, uncovering some pattern in the map. In the special case of a pure gauge, both abelian and non-abelian, these expressions can be summed up, and the transformation is expressed using the parametrisation in terms of the gauge group.

¹Unité mixte du CNRS et de l'EP, UMR 7644

1 Introduction

Noncommutativity has been studied extensively since it has become clear that noncommutative gauge theories can describe the low energy effective action of a brane in a constant magnetic background (see [1, 2] for reviews). In particular in [3], a correspondence, known as the Seiberg-Witten map, between gauge fields living on D-brane worldvolume in the background of a non vanishing constant electromagnetic field and noncommutative gauge theory on a space with coordinate x^μ satisfying

$$[x^\mu, x^\nu] = i\theta^{\mu\nu} \quad (1.1)$$

has been established.

The simplest example of noncommutative gauge theory is the abelian $U(1)$, which is not a free theory, and is more similar to a $U(N)$ gauge theory. In fact, at least on a torus, there exists a transformation, the Morita equivalence, that allows to change the noncommutativity parameter of the space and the rank of the gauge group at the same time. For example on a 2-torus, $U(1)$ with noncommutativity parameter $\theta = \frac{1}{N}$ is equivalent to an ordinary $U(N)$ theory. In this context, Morita equivalence can be seen from rewriting the objects in a matrix language [4, 5].

The action of Seiberg-Witten (SW) map and Morita equivalence could be combined, for example, to link the usual abelian Born-Infeld action to its non-abelian version [6], for which the ambiguities have not yet been solved [7, 8, 9, 15]. In this process, Morita equivalence brings all these ordering ambiguities back into the noncommutative world, where we know that the DBI lagrangian is the lagrangian invariant under SW map, in the slowly varying fields limit. Fixing the ambiguities requires going beyond this limit, which means also having a better knowledge and understanding of the SW map.

Another related issue is about the dynamics of the B field. For instance, if we take two snapshots of a $U(1)$ theory, with two different B fields, they would correspond to two different noncommutativity parameters. Let us say that the first one is linked to an ordinary $U(5)$ gauge group by Morita equivalence, while the second one is related to $U(3)$. Would that mean that the rank of the gauge group itself has become dynamic ? We believe that there are still things to be understood in the case of a constant θ , that could help dealing with a dynamic background.

In this paper, we would like to explore further the order by order formulation of the Seiberg-Witten map, and the structures associated with it. There already exists a formula, known as Liu's conjecture [10], recently proven in [11, 12, 13, 14], which expresses the commutative field strength for a $U(1)$ gauge group in terms of the noncommutative variables non perturbatively. This formula was found by considering the couplings of the noncommutative brane to Ramond-Ramond potentials, and its order by order expansion involves the $\star_n$ products [16, 17].

Our approach here is completely different. It is based on solving order by order the Seiberg-Witten equation, and is closer in spirit to [19]. Using the freedom in the possible solutions

This allows us to write the non trivial star commutator²

$$i[f * g] = i[f, g] - \frac{1}{2} \{ \partial f \circ \partial g \} - \frac{i}{8} \{ \partial \partial f \circ \partial \partial g \} + \frac{1}{48} \{ \partial \partial \partial f \circ \partial \partial \partial g \} + \dots \quad (2.2)$$

which, in the abelian case, reduces to

$$i[f * g] = -\partial f \partial g + \frac{1}{24} \partial \partial f \partial \partial \partial g + \dots \quad (2.3)$$

In particular, we have the noncommutativity of the coordinates: $[x^\mu * x^\nu] = i\theta^{\mu\nu}$

On this space lives the noncommutative gauge theory with gauge potential $\hat{A}_\mu$ and field strength

$$\hat{F}_{\mu\nu} = \partial_\mu \hat{A}_\nu - \partial_\nu \hat{A}_\mu - i[\hat{A}_\mu * \hat{A}_\nu]$$

The gauge transformation δ has gauge parameter $\hat{\lambda}$ and

$$\delta \hat{A}_\mu = \partial_\mu \hat{\lambda} + i[\hat{\lambda} * \hat{A}_\mu] = \hat{D}_\mu \hat{\lambda} \quad \delta \hat{F}_{\mu\nu} = i[\hat{\lambda} * \hat{F}_{\mu\nu}]$$

2.1 The Seiberg-Witten map

The Seiberg-Witten map is a map between the commutative and the noncommutative gauge theory, which is compatible with gauge transformations. In other words, it maps gauge orbits into gauge orbits.

We can write noncommutative objects as functions of the commutative ones

$$\hat{A}_\mu = \hat{A}_\mu(A_\mu; \theta) \quad ; \quad \hat{F}_{\mu\nu} = \hat{F}_{\mu\nu}(A_\mu; \theta) \quad ; \quad \hat{\lambda} = \hat{\lambda}(\lambda, A_\mu; \theta)$$

and we have the following diagram, which tells that gauge transforming $\hat{A}_\mu$ as the noncommutative gauge potential or as a function of the commutative one is the same:

$$\begin{array}{ccc} A_\mu & \xrightarrow{\delta} & A_\mu + \delta A_\mu \\ \downarrow & & \downarrow \\ \hat{A}_\mu & \xrightarrow{\delta} & \hat{A}_\mu + \delta \hat{A}_\mu \end{array} \quad \Leftrightarrow \quad \begin{array}{ccc} & & \\ & \xrightarrow{\delta} & \\ & & \end{array} \quad A_\mu + \delta A_\mu = \hat{A}_\mu + \delta \hat{A}_\mu \quad \Leftrightarrow \quad \delta \hat{A}_\mu = \delta \hat{A}_\mu \quad (2.4)$$

This is the Seiberg-Witten equation, which will allow us to find the map explicitly:

$$\delta \hat{A}_\mu - \partial_\mu \hat{\lambda} = i[\hat{\lambda} * \hat{A}_\mu] \quad (2.4)$$

²where we have introduced the notations $\{\partial f \circ \partial g\} = \theta^{\mu\nu} \{\partial_\mu f, \partial_\nu g\}$,

$\{\partial \partial f \circ \partial \partial g\} = \theta^{\mu\nu} \theta^{\rho\sigma} \{\partial_\mu \partial_\nu f, \partial_\rho \partial_\sigma g\}$, $\{\partial \partial \partial f \circ \partial \partial \partial g\} = \theta^{\mu\nu} \theta^{\rho\sigma} \theta^{\alpha\beta} \{\partial_\mu \partial_\nu \partial_\alpha f, \partial_\rho \partial_\sigma \partial_\beta g\}$.

The first operand carries the first index of θ , making these symbols antisymmetric. Generally speaking, we will not write the indices when contractions are done in a natural way.

(see [18]) we have also chosen expressions that differ from those of [17] by a gauge transformation. The hope was to express the map in a simpler way, by extending the recursive formulation that appears in the usual first orders of the field strength (see 3.1.5). Our guiding criterion for choosing a solution is to preserve its structure at all orders. Hence, our expressions, if still order by order, can be expressed recursively and explicitly written at any order with little work, at least in the abelian case. Hopefully, these recursive expressions will lead to the full non perturbative formula, as in the case of the pure gauge, where we were able to express the solution in a particularly suggestive form. Of course, what is interesting in the pure gauge case is not finding a solution but rather a clue on the particular form of the general solution.

The structure of the paper is as follows: definitions and notations will be found in Section 2, as well as some remarks about our method. Section 3 contains the results for an abelian gauge theory, and the order by order solutions are expressed recursively. This is the main perturbative feature of the paper, and has been explicitly checked up to order 6 (cf 3.1). Work still remains to be done to extend the pattern to all orders in A_μ (cf 3.2). A non perturbative explanation for the existence of this structure has been found only in the case of the pure gauge. Then indeed, the solution can be compactly expressed using the parametrisation in terms of the gauge group, be it abelian (cf 3.3) or not (cf 4.1). The same kind of parametrisation may be possible in the general case, but has not yet been achieved. Some remarks about solving the non-abelian equations are made in 4.2.

2 Discovering the landscape

Let us consider a commutative gauge theory. The gauge group may be general, and will not be explicit. The gauge potential is A_μ , with field strength

$$F_{\mu\nu} = \partial_\mu A_\nu - \partial_\nu A_\mu - i[A_\mu, A_\nu]$$

The gauge transformation δ , with gauge parameter λ , will be acting as

$$\delta A_\mu = \partial_\mu \lambda + i[\lambda; A_\mu] = D_\mu \lambda \quad \delta F_{\mu\nu} = i[\lambda; F_{\mu\nu}]$$

We will first study the abelian theory, so that all commutators vanish and the expressions simplify. When working with a non-abelian theory, the gauge structure will then be encoded in the commutators and anticommutators.

On the noncommutative side, usual multiplication of functions is replaced by star product, which here is the Moyal-Weyl formula for a constant noncommutativity parameter $\theta^{\mu\nu}$:

$$\begin{aligned} f \star g &= f e^{\frac{i}{2} \theta^{\mu\nu} \partial_\mu \partial_\nu} g \\ &= f g + \frac{i}{2} (\partial_\mu f) \theta^{\mu\nu} (\partial_\nu g) - \frac{1}{8} (\partial_\rho \partial_\sigma f) \theta^{\rho\sigma} \theta^{\alpha\beta} (\partial_\alpha \partial_\beta g) + O(\theta^3) \end{aligned} \quad (2.1)$$

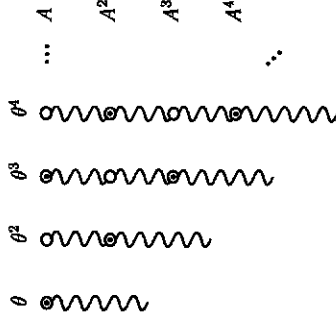


Figure 1: Circles are non-abelian terms, that spread on higher orders in A , hence the tails. When the group is abelian, many terms drop out, leaving only the dots, which do not spread anymore.

2.2 Freedom in the solution

To solve equation (2.4), we develop the noncommutative variables in powers of θ :

$$\begin{aligned}\hat{A}_\mu &= A_\mu + A_\mu^{(1)} + A_\mu^{(2)} + \dots \\ \hat{F}_{\mu\nu} &= F_{\mu\nu} + F_{\mu\nu}^{(1)} + F_{\mu\nu}^{(2)} + \dots \\ \hat{\lambda} &= \lambda + \lambda^{(1)} + \lambda^{(2)} + \dots\end{aligned}$$

Then the left hand side of equation (2.4) is, at order n , exactly $\delta A_\mu^{(n)} - \partial_\mu \lambda^{(n)}$, while development of the star commutator on the right hand side only involves terms of lower (or same) order. This is the key why we can find explicit solution to this equation, order by order in θ . Of course, since the only requirement was mapping gauge orbits into gauge orbits, there is some freedom in the solution. In fact, any "homogeneous solution", meaning $(\hat{A}_\mu^{(n)}; \hat{\lambda}^{(n)})$ so that $\delta \hat{A}_\mu^{(n)} - \partial_\mu \hat{\lambda}^{(n)} = 0$, can be freely added to a particular solution of the complete equation. This "homogeneous solution" does not change the physics, and can be rewritten as a gauge transformation [18]. We will see an example of it in section 3.1.2, where we relate our solution with the already known solutions of [10, 12].

In the following, there will be another type of development: order by order in powers of A (fig.1). In fact, each term in a solution can only involve, algebraically, A , θ and derivatives. Specifying the number of A and the number of θ then characterizes a class of terms (freedom remains in the contraction of indices, and the action of derivatives). We will denote a term of order $(A^n; \theta^m)$ in the development of $\hat{\lambda}$ or $\hat{A}_\mu$ by $\lambda^{(n,m)}$ or $A_\mu^{(n,m)}$.

Let us stress that this order in A , whereas well defined in the abelian theory, becomes rather unclear in a non-abelian theory. The field strength, the covariant derivative and the gauge transformation contain a quadratic part in A . This does not mean that order in A does not make sense, but that we will have to be very careful in the identification of the terms. Misidentification would spoil the picture.

We will also see that it may be more efficient to sum all the terms of a given order in A , and not in θ , as advertised. Thus, we will use the notation $\lambda^{(n,\infty)}$, $A_\mu^{(n,\infty)}$ meaning order n in A

$$\lambda^{(n,\infty)} = \sum_{m=0}^{\infty} \lambda^{(n,m)} \quad ; \quad A_\mu^{(n,\infty)} = \sum_{m=0}^{\infty} A_\mu^{(n,m)}$$

3 The abelian theory

In the abelian theory, different orders in A do not mix together. We will start with the first diagonal of the diagram: the terms of order (A^n, θ^n) . We will see that we can write solutions so that each point is given from the previous one. The expressions can then also be written easily without recursivity. The emphasized expressions have been checked up to and including $n = 6$. For $n = 4, 5, 6$, this has been done using a computer program specially written on this purpose.

3.1 The first diagonal: orders (A^n, θ^n)

3.1.1 First order

The first order equation is reduced to $\delta A_\mu^{(1,1)} - \partial_\mu \lambda^{(1,1)} = -\partial \lambda \theta \partial A_\mu$ and admits the usual solution [3]

$$\begin{aligned}\lambda^{(1,1)} &= -\frac{1}{2} \theta^{\sigma\tau} A_\sigma \partial_\tau \lambda \\ A_\mu^{(1,1)} &= -\frac{1}{2} \theta^{\sigma\tau} A_\sigma (\partial_\tau A_\mu + F_{\sigma\mu})\end{aligned}$$

We can rewrite this solution to make its structure more apparent, so that we can recognize the form of next orders

$$\begin{aligned}\lambda^{(1,1)} &= -\frac{1}{2} (A \theta \partial) \lambda \\ A_\mu^{(1,1)} &= -\frac{1}{2} (A \theta \partial) A_\mu - \frac{1}{2} (A \theta F)_\mu\end{aligned}\quad (3.1)$$

3.1.2 Second order

The equation at order (A^2, θ^2) is $\delta A_\mu^{(2,2)} - \partial_\mu \lambda^{(2,2)} = -\partial \lambda^{(1,1)} \theta \partial A_\mu - \partial \lambda \theta \partial A_\mu^{(1,1)}$. Solving with the previous (1,1) expressions, we get the solution

$$\begin{aligned}\lambda^{(2,2)} &= \frac{1}{6} [(A \theta \partial)(A \theta \partial) \lambda + A \theta F \theta \partial \lambda] \\ A_\mu^{(2,2)} &= \frac{1}{6} [(A \theta \partial)(A \theta \partial) A_\mu + A \theta F \theta \partial A_\mu + 2(A \theta \partial)(A \theta F)_\mu + 2A \theta F \theta F_\mu]\end{aligned}\quad (3.2)$$

Recalling the first order solution

$$\begin{aligned}\lambda^{(1,1)} &= -\frac{1}{2} A\theta\partial\lambda \\ A_\mu^{(1,1)} &= -\frac{1}{2} A\theta(\partial A_\mu + F_\mu)\end{aligned}$$

and comparing to (3.2) and (3.3), we can write the second and third orders in a similar way, with $A^{(1,1)}$ or $A^{(2,2)}$ acting as operators, under the form given above. That means explicit derivatives (contrary to those hidden in F s) also act on what is on the right.

$$\begin{aligned}\lambda^{(2,2)} &= -\frac{1}{3} A^{(1,1)}\theta\partial\lambda & \lambda^{(3,3)} &= -\frac{1}{4} A^{(2,2)}\theta\partial\lambda \\ A_\mu^{(2,2)} &= -\frac{1}{3} A^{(1,1)}\theta(\partial A_\mu + 2F_\mu) & A_\mu^{(3,3)} &= -\frac{1}{4} A^{(2,2)}\theta(\partial A_\mu + 3F_\mu)\end{aligned}$$

This leads to the conjecture that the same iterative relation gives a solution for general n

$$\begin{aligned}\lambda^{(n,n)} &= -\frac{1}{n+1} A^{(n-1,n-1)}\theta\partial\lambda \\ A_\mu^{(n,n)} &= -\frac{1}{n+1} A^{(n-1,n-1)}\theta(\partial A_\mu + nF_\mu)\end{aligned}\quad (3.4)$$

It is also possible to express an algebraically factorized form without recursivity. Let us denote as d and f the two basic blocks that are building all the terms

$$d \equiv A\theta\partial \quad f \equiv A\theta F A^{-1}$$

with $A_p^{-1}A_\sigma = \delta_{p\sigma}$ and the rule $A^{-1}\lambda = 0$ (or $f\lambda = 0$). Then

$$\begin{aligned}\lambda^{(n,n)} &= \frac{(-1)^n}{(n+1)!} (d+f)(d+2f)\dots(d+nf)\lambda \\ A_\mu^{(n,n)} &= \frac{(-1)^n}{(n+1)!} (d+f)(d+2f)\dots(d+nf)A_\mu\end{aligned}$$

The only effect of the $f\lambda = 0$ rule is that when developing the last factor, the formulas for A have twice the number of terms than those for λ , as has been seen previously. It might be interesting to learn more about the significance and properties of these f and d , even if they can be thought of as the first order of more general objects. Indeed, we will see the first correction in section 3.2.

But there are other possibilities. In [19] another solution was found

$$\begin{aligned}\lambda^{(2,2)}_{\text{JMSW}} &= \frac{1}{2} \theta^{\rho\sigma} \theta^{\tau\kappa} A_\rho \partial_\sigma A_\kappa \partial_\tau \lambda \\ A_\mu^{(2,2)}_{\text{JMSW}} &= \frac{1}{2} \theta^{\rho\sigma} \theta^{\tau\kappa} A_\rho [\partial_\sigma A_\kappa \partial_\tau A_\mu + A_\kappa \partial_\sigma F_{\tau\mu} + F_{\sigma\kappa} F_{\tau\mu}]\end{aligned}$$

The difference between the two is the “homogeneous solution” $\frac{1}{6}H$, where H is

$$\begin{aligned}\Lambda^{(2,2)} &= \theta^{\rho\sigma} \theta^{\tau\kappa} A_\rho [\partial_\sigma A_\kappa + \partial_\sigma A_\kappa - A_\kappa \partial_\sigma] \partial_\tau \lambda = \delta[A\theta(\partial A)\theta A] \\ A_\mu^{(2,2)} &= \theta^{\rho\sigma} \theta^{\tau\kappa} A_\rho [\partial_\sigma A_\kappa + \partial_\sigma A_\kappa - A_\kappa \partial_\sigma] \partial_\tau A_\mu = \partial_\mu [A\theta(\partial A)\theta A]\end{aligned}$$

In particular, it is a gauge transformation on $\hat{A}_\mu$, with parameter $\frac{1}{6}A\theta(\partial A)\theta A$. The solution (3.2) is also different from the order 2 expansion of Liu’s formula [10, 16, 17]. This time, the difference is $\frac{5}{12}H$, and it is again a gauge transformation. Of course, in both cases, the field strength is the same.

3.1.3 Third order

The third order begins to get very messy. The equation is

$$\delta A_\mu^{(3,3)} - \partial_\mu \lambda^{(3,3)} = -\partial \lambda^{(2,2)} \theta \partial A_\mu - \partial \lambda^{(1,1)} \theta \partial A_\mu^{(1,1)} - \partial \lambda \theta \partial A^{(2,2)}_\mu$$

Expanding it with the actual expressions (3.1) and (3.2) gives a great number of terms. But we can build an expression on the model arising from (3.1) and (3.2). And indeed, we explicitly checked that the following is a solution to this equation

$$\begin{aligned}\lambda^{(3,3)} &= -\frac{1}{24} [(A\theta\partial)(A\theta\partial)(A\theta\partial)\lambda + (A\theta F\theta\partial)(A\theta\partial)\lambda \\ &\quad + 2(A\theta\partial)(A\theta F\theta\partial)\lambda + 2A\theta F\theta\partial\theta\lambda] \\ A_\mu^{(3,3)} &= -\frac{1}{24} [(A\theta\partial)(A\theta\partial)(A\theta\partial)A_\mu + (A\theta F\theta\partial)(A\theta\partial)A_\mu \\ &\quad + 2(A\theta\partial)(A\theta F\theta\partial)A_\mu + 2A\theta F\theta\partial\theta A_\mu \\ &\quad + 3(A\theta\partial)(A\theta F)(A\theta F)_\mu + 3(A\theta F\theta\partial)(A\theta F)_\mu \\ &\quad + 6(A\theta\partial)(A\theta F\theta F)_\mu + 6A\theta F\theta F\theta F_\mu]\end{aligned}\quad (3.3)$$

The interest of having this explicit solution is that we see a structure emerging. The operators look very similar, the coefficients seem to follow some rule. The first two orders were not really enough to draw conclusions, but order 3 strongly suggests that the form can be generalized to a solution of order n . Furthermore, we developed a computer code which allowed us to check that these expressions are also solutions at order 4, 5 and 6.

3.1.4 Recursive formulation

Better than describing some prescriptions leading to the general form of order n or rather (n, n) , we can express the results obtained so far in a recursive manner.

$\partial A_\mu + nF_{\mu\nu}$, as can be seen on (3.4). We have calculated a solution at order (2,4), the second term on the second line, and indeed it happens to depend on $\partial A_\mu + 2F_{\mu\nu}$. Then, one may wonder about the extended recursive rules. Recall that the expression at order 2 was the same as the solution for order 1, with

- the starting A replaced by $A^{(1,1)}$
- $\frac{\partial A_\mu + F_{\mu\nu}}{2}$ changed to $\frac{\partial A_\mu + 2F_{\mu\nu}}{3}$

If we do the same for all the line, we get

$$\begin{aligned}\lambda^{(1,\infty)} &= \frac{1}{2} \left[-A\partial\lambda + \frac{1}{24} \partial\partial A \partial\partial\partial\lambda + \dots \right] \\ \lambda^{(2,\infty)} &= \frac{1}{3} \left[-A^{(1,\infty)}\partial\partial\lambda + \frac{1}{24} \partial\partial A^{(1,\infty)} \partial\partial\partial\lambda + \dots \right] \\ &= \underbrace{-\frac{1}{3} A^{(1,1)}\partial\partial\lambda - \frac{1}{3} A^{(1,3)}\partial\partial\lambda + \frac{1}{72} \partial\partial A^{(1,1)} \partial\partial\partial\partial\lambda + \dots}_{\lambda^{(2,4)}}\end{aligned}$$

And the same for $A_\mu^{(2,\infty)}$, replacing $\partial\lambda$ by $(\partial A_\mu + 2F_{\mu\nu})$. Do not forget that the $A^{(1,m)}$ have to act as operators (the rest of the term is "glued" inside the action of their derivatives). This prescription leads to ambiguities³, up to which this proposal for order (2,4) reproduces an actual solution (out of a dozen, all but one terms are reproduced unambiguously).

The same should also hold for the field strength. Generalizing the formulas (3.5) would give $F_{\mu\nu}^{(2,\infty)}$ from $F_{\mu\nu}^{(1,\infty)}$, with

$$F_{\mu\nu}^{(1,\infty)} = \begin{bmatrix} -A\partial\partial F_{\mu\nu} + \frac{1}{24} \partial\partial A \partial\partial\partial F_{\mu\nu} + \dots \\ -F_{\mu\nu}\partial F_{\nu\rho} + \frac{1}{24} \partial\partial F_{\mu\nu} \partial\partial\partial F_{\nu\rho} + \dots \end{bmatrix}$$

3.3 Pure gauge case

We consider here the case when the gauge potential A is a pure gauge. For an abelian theory, this can be done with $A_\rho = \partial_\rho\alpha$. As expected, the field strength vanishes, and the gauge transformation of α should be $\delta\alpha = \lambda$.

The objective here is not to find the solutions: it is the orbit of the noncommutative pure gauge. Indeed, we expect from the map that zero is mapped to zero, and this statement can

³ We shall try to comment on this here. In the term

$$\partial\partial A^{(1,1)} \partial\partial\partial\lambda \propto (A\partial\partial)A (\partial\partial\partial\lambda)$$

$A^{(1,1)}$ should act as an operator on the right, while $\partial\partial$ is acting on the left. Note that this is the first instance where such a situation occurs. One will keep meeting similar ambiguities at higher orders.

3.1.5 The field strength

We calculated the noncommutative field strength corresponding to the previous expressions for A_μ

$$\begin{aligned}F_{\mu\nu}^{(1,1)} &= \partial_\mu A_\nu^{(1,1)} - \partial_\nu A_\mu^{(1,1)} + \partial_\mu A_\nu \partial\partial A_\nu - (A\partial\partial)F_{\mu\nu} - F_{\mu\nu}\partial F_{\nu\rho} \\ F_{\mu\nu}^{(2,2)} &= \partial_\mu A_\nu^{(2,2)} - \partial_\nu A_\mu^{(2,2)} + \partial A_\mu^{(1,1)}\partial\partial A_\nu + \partial A_\nu\partial\partial A_\mu^{(1,1)} \\ &= \frac{1}{2}(A\partial\partial)(A\partial\partial)F_{\mu\nu} + \frac{1}{2}(A\partial F\partial\partial)F_{\mu\nu} + (A\partial\partial)(F_{\mu\nu}\partial F_{\nu\rho}) + F_{\mu\nu}\partial F\partial F_{\nu\rho} \\ F_{\mu\nu}^{(3,3)} &= \partial_\mu A_\nu^{(3,3)} - \partial_\nu A_\mu^{(3,3)} + \partial A_\mu^{(2,2)}\partial\partial A_\nu + \partial A_\nu\partial\partial A_\mu^{(2,2)} + \partial A_\mu^{(1,1)}\partial\partial A_\nu^{(1,1)} \\ &= -\frac{1}{6} \left[(A\partial\partial)(A\partial\partial)(A\partial\partial)F_{\mu\nu} + 2(A\partial\partial)(A\partial F\partial\partial)F_{\mu\nu} \right] \\ &\quad + \frac{1}{2} \left[(A\partial F\partial\partial)(A\partial\partial)F_{\mu\nu} + 2(A\partial F\partial F\partial\partial)F_{\mu\nu} \right] \\ &= -\frac{1}{2} \left[(A\partial\partial)(A\partial\partial)(F_{\mu\nu}\partial F_{\nu\rho}) \right] - \left[(A\partial\partial)(F_{\mu\nu}\partial F\partial F_{\nu\rho}) \right] \\ &\quad + \frac{1}{2} \left[(A\partial F\partial\partial)(F_{\mu\nu}\partial F_{\nu\rho}) \right] - \left[(A\partial F\partial F\partial\partial)F_{\mu\nu} \right]\end{aligned}\tag{3.5}$$

And again we identify in these expressions the same recursive pattern

$$\begin{aligned}F_{\mu\nu}^{(1,1)} &= -(A\partial\partial)F_{\mu\nu} - F_{\mu\nu}\partial F_{\nu\rho} \\ F_{\mu\nu}^{(2,2)} &= -(A^{(1,1)}\partial\partial)F_{\mu\nu} - F_{\mu\nu}^{(1,1)}\partial F_{\nu\rho} \\ F_{\mu\nu}^{(3,3)} &= -(A^{(2,2)}\partial\partial)F_{\mu\nu} - F_{\mu\nu}^{(2,2)}\partial F_{\nu\rho}\end{aligned}$$

Let us stress that the second line only involves usual formulas: $F_{\mu\nu}^{(1,1)}$, $F_{\mu\nu}^{(2,2)}$ and $A_\mu^{(1,1)}$ are the same here as in [10, 19]. The particular choice that we have made for $A_\mu^{(2,2)}$ and others then allowed the same pattern for the third line.

3.2 A line by line relation ?

Now that we have some expression for the first diagonal, it still remains to fill the entire triangle on fig.1 to get all the terms. In fact, the whole first line (order 1 in A) is really easy to find, and the solution is essentially unique. It is known as a $\star_2$ or $\star'$ product expression [16]:

$$A_\mu^{(1,\infty)} = -\frac{1}{2}\theta^{\mu\nu}A_\nu \star_2 (\partial_\mu A_\nu + F_{\mu\nu}) \quad \text{with} \quad f(x) \star_2 g(x) = \frac{\sin(\frac{\partial_\mu \partial_\nu}{2})}{\frac{\partial_\mu \partial_\nu}{2}} f(x_1)g(x_2)|_{x_1=x_2}$$

which have expansion in θ

$$\begin{aligned}A_\mu^{(1,\infty)} &= A_\mu^{(1,1)} + A_\mu^{(1,3)} + \dots \\ A_\mu^{(1,\infty)} &= \frac{1}{2} \left[-A\partial(\partial A_\mu + F_{\mu\nu}) + \frac{1}{24} \partial\partial A \partial\partial\partial(\partial A_\mu + F_{\mu\nu}) + \dots \right]\end{aligned}$$

We want to use this to extend the previous recursive relation, holding for the first diagonal, to a line by line relation. The basic reason to try to sum by lines (see fig.1), is that all terms of the first line depend on $\partial A_\mu + F_{\mu\nu}$, while terms of order (n, n) are expected to depend on

only hold for the whole orbits. What we are looking for is an expression that, in some sense, is nicer than the others. The solutions emphasized up to now will indeed greatly simplify in the pure gauge case. That will allow two different improvements, which are still out of reach for a general potential: first, the solutions have a natural non-abelian generalisation, and second, they will sum up to a non-perturbative form (as we will see in 4.1). This suggests that the same kind of expressions could hold for general potentials, once they are parametrized by elements of the gauge group.

3.3.1 The solution for A and λ

We can rewrite the solution at order $(1, \infty)$

$$A_\mu^{(1, \infty)} = -\frac{1}{2} \partial \alpha \partial \partial A_\mu + \frac{1}{48} \partial \partial \partial \alpha \partial \partial \partial A_\mu + \dots$$

and recognize the development of a shorter formula (the same holds for $\lambda^{(1, \infty)}$)

$$\begin{aligned} A_\mu^{(1, \infty)} &= \frac{i}{2} [\alpha \star A_\mu] \\ \lambda^{(1, \infty)} &= \frac{i}{2} [\alpha \star \lambda] \end{aligned}$$

The Seiberg-Witten equation, written at each order in A , gives⁴

$$\begin{aligned} \text{Order } A : \quad & \delta A_\mu^{(1)} - \partial_\mu \lambda^{(1)} = i[\lambda \star A_\mu] \\ \text{Order } A^2 : \quad & \delta A_\mu^{(2)} - \partial_\mu \lambda^{(2)} = i[\lambda \star A_\mu^{(1)}] + i[\lambda^{(1)} \star A_\mu] \\ \text{Order } A^n : \quad & \delta A_\mu^{(n)} - \partial_\mu \lambda^{(n)} = \sum_{p=0}^{n-1} i[\lambda^{(p)} \star A_\mu^{(n-1-p)}] \end{aligned}$$

Now we can check directly that the expression is indeed solution at order A

$$\begin{aligned} \delta A_\mu^{(1)} - \partial_\mu \lambda^{(1)} &= \frac{1}{2} i[\lambda \star A_\mu] + \frac{1}{2} i[\alpha \star \partial_\mu \lambda] - \frac{1}{2} i[\partial_\mu \alpha \star \lambda] - \frac{1}{2} i[\alpha \star \partial_\mu \lambda] \\ &= i[\lambda \star A_\mu] \equiv i[\lambda^{(0)} \star A_\mu^{(0)}] \end{aligned}$$

This formula is still valid to all orders, as we shall now prove recursively. Suppose we have the solution for $k \leq n$

$$\begin{aligned} A_\mu^{(k)} &= \frac{i}{k+1} [\alpha \star A_\mu^{(k-1)}] \\ \lambda^{(k)} &= \frac{i}{k+1} [\alpha \star \lambda^{(k-1)}] \end{aligned}$$

and define order $(n+1)$ by the same recursive formula. Then

$$\delta A_\mu^{(n+1)} - \partial_\mu \lambda^{(n+1)} = \frac{i}{n+2} [\lambda \star A_\mu^{(n)}] + \frac{i}{n+2} [\lambda^{(n)} \star A_\mu] + \frac{i}{n+2} [\alpha \star \delta A_\mu^{(n)} - \partial_\mu \lambda^{(n)}]$$

⁴ From now on, we drop the ∞ since no confusion can arise: $A_\mu^{(n, \infty)}$ means $A_\mu^{(n, \infty)}$

and with the help of the Jacobi identity

$$\begin{aligned} [\alpha \star \delta A_\mu^{(n)} - \partial_\mu \lambda^{(n)}] &= [\alpha \star \sum_{p=0}^{n-1} i[\lambda^{(p)} \star A_\mu^{(n-1-p)}]] \\ &= \sum_{p=0}^{n-1} (\lambda^{(p)} \star i[\alpha \star A_\mu^{(n-1-p)}]) - [A_\mu^{(n-1-p)} \star i[\alpha \star \lambda^{(p)}]] \\ &= \sum_{p=0}^{n-1} ((n+1-p)[\lambda^{(p)} \star A_\mu^{(n-p)}] - (p+2)[A_\mu^{(n-1-p)} \star \lambda^{(p+1)}]) \end{aligned}$$

which is exactly what is needed to finally give

$$\delta A_\mu^{(n+1)} - \partial_\mu \lambda^{(n+1)} = \sum_{p=0}^n i[\lambda^{(p)} \star A_\mu^{(n-p)}]$$

So we arrive at a solution to all orders for A_μ and λ , in the form

$$\begin{aligned} A_\mu^{(n+1)} &= \frac{i}{n+2} [\alpha \star A_\mu^{(n)}] \\ \lambda^{(n+1)} &= \frac{i}{n+2} [\alpha \star \lambda^{(n)}] \end{aligned} \quad (3.6)$$

$A_\mu^{(0)}$ and $\lambda^{(0)}$ being A_μ and λ .

3.3.2 The field strength

Here we compute the field strength corresponding to the solution (3.6) for the noncommutative connection, and show recursively that it indeed vanishes. $\tilde{F}_{\mu\nu}$ can be developed in powers of A , giving

$$\begin{aligned} F_{\mu\nu}^{(0)} &\equiv F_{\mu\nu} = 0 \\ F_{\mu\nu}^{(1)} &\equiv \partial_\mu A_\nu^{(1)} - \partial_\nu A_\mu^{(1)} - i[A_\mu \star A_\nu] \\ &= \frac{i}{2} [A_\mu \star A_\nu] + \frac{i}{2} [\alpha \star \partial_\mu A_\nu] - \frac{i}{2} [A_\nu \star A_\mu] - \frac{i}{2} [\alpha \star \partial_\nu A_\mu] - i[A_\mu \star A_\nu] \\ &= \frac{i}{2} [\alpha \star F_{\mu\nu}] = 0 \\ F_{\mu\nu}^{(n+1)} &\equiv \partial_\mu A_\nu^{(n+1)} - \partial_\nu A_\mu^{(n+1)} - \sum_{k=0}^n i[A_\mu^{(k)} \star A_\nu^{(n-k)}] \\ &= \frac{i}{n+2} [A_\mu \star A_\nu^{(n)}] - \frac{i}{n+2} [A_\nu \star A_\mu^{(n)}] + \frac{i}{n+2} [\alpha \star \partial_\mu A_\nu^{(n)} - \partial_\nu A_\mu^{(n)}] \\ &\quad - \sum_{k=0}^n i[A_\mu^{(k)} \star A_\nu^{(n-k)}] \end{aligned}$$

where the commutators take care at the same time of the noncommutativity and of the non-abelian gauge group. The relation between the above A_μ and $\hat{A}_\mu$ is precisely the SW map.

In particular, we get A_μ and λ for $\theta = 0$, as expected. So that the expressions of the noncommutative variables in terms of the commutative ones, order by order in θ , start like

$$\begin{aligned}\hat{A}_\mu &= A_\mu + A_\mu^{(1)} + \dots \\ \hat{\lambda} &= \lambda + \lambda^{(1)} + \dots\end{aligned}$$

To understand why such a parametrisation exists, and what it means, let us see that we can go beyond this perturbative development. If A is a pure gauge, then it can be written in terms of an element $g = e^{-i\alpha}$ of the gauge group

$$A = ig^{-1}dg$$

Developed in powers of α , this expression is nothing but (4.1). The non-perturbative formula for λ is $\lambda = ig^{-1}\delta g$ (which is nothing but the gauge transformation of g). Of course, it can be checked directly that F is zero or that the right gauge transformation holds.

On the noncommutative side, the formulas are the same. We use the corresponding element $\hat{g} = (\star e)^{-i\alpha} = 1 - i\alpha - \frac{\alpha \star \alpha}{2} + \dots$, and replace ordinary multiplications by star products

$$\begin{aligned}\hat{A} &= i\hat{g}^{-1} \star d\hat{g} = d\alpha + \frac{i}{2}[\alpha \star d\alpha] + \frac{i}{6}[\alpha \star i[\alpha \star d\alpha]] + \dots \\ \hat{\lambda} &= i\hat{g}^{-1} \star \delta \hat{g} = \delta\alpha + \frac{i}{2}[\alpha \star \delta\alpha] + \frac{i}{6}[\alpha \star i[\alpha \star \delta\alpha]] + \dots\end{aligned}\quad (4.3)$$

4.2 The general case

Is it possible to generalize the recursive picture from abelian theory to non-abelian one? To do this, we first have to find a "nice" solution at order θ^2 , which means that it should at least reduce to the abelian solution that we found in 3.1.2. But with the appearance of gauge commutators, there is a lot of freedom in generalizing, and we face some new problems:

- one term out of two in fig.1 was a vanishing commutator, now it does not vanish,
- the derivative could be replaced by the covariant derivative or stay as an ordinary derivative,
- the order in A (the lines) is not well defined, now that the field strength, the gauge transformation and the covariant derivative have a quadratic part.

The last point is the reason why we have to calculate the columns first (see fig.1), and sort the terms thereafter, identifying which term belongs to which line.

Let us start with order θ . The Seiberg-Witten equation is

$$\delta A_\mu^{(\infty,1)} - i[\lambda; A_\mu^{(\infty,1)}] - \partial_\mu \lambda^{(\infty,1)} + i[A_\mu; \lambda^{(\infty,1)}] = -\frac{1}{2}\{\partial\lambda; \partial A_\mu\}$$

Using the definition of $F_{\mu\nu}^{(n)}$

$$\partial_\mu A_\nu^{(n)} - \partial_\nu A_\mu^{(n)} = F_{\mu\nu}^{(n)} + \sum_{k=0}^{n-1} i[A_\mu^{(k)} \star A_\nu^{(n-1-k)}],$$

the Jacobi identity

$$\begin{aligned}[\alpha \star i[A_\mu^{(k)} \star A_\nu^{(n-1-k)}]] &= [A_\mu^{(k)} \star i[\alpha \star A_\nu^{(n-1-k)}]] - [A_\nu^{(n-1-k)} \star i[\alpha \star A_\mu^{(k)}]] \\ &= (n-k+1)[A_\mu^{(k)} \star A_\nu^{(n-k)}] + (k+2)[A_\mu^{(k+1)} \star A_\nu^{(n-1-k)}],\end{aligned}$$

and replacing in $F_{\mu\nu}^{(n+1)}$, most of the terms cancel, leaving

$$F_{\mu\nu}^{(n+1)} = \frac{i}{n+2}[\alpha \star F_{\mu\nu}^{(n)}] = 0 \quad (3.7)$$

4 Non-abelian theory

4.1 The pure gauge

To go on with the pure gauge in the non-abelian theory, we have to find some parametrisation compatible with the gauge transformation

$$\delta A_\mu = \partial_\mu \lambda + i[\lambda; A_\mu]$$

Indeed, the crucial point in the abelian theory was the existence of α , such that $A_\mu = \partial_\mu \alpha$ and $\lambda = \delta\alpha$. And we already know a generalization of this, allowing the non-abelian gauge transformation:

$$\begin{aligned}A_\mu &= \partial_\mu \alpha + \frac{i}{2}[\alpha; \partial_\mu \alpha] + \frac{i}{6}[\alpha; i[\alpha; \partial_\mu \alpha]] + \dots \\ \lambda &= \delta\alpha + \frac{i}{2}[\alpha; \delta\alpha] + \frac{i}{6}[\alpha; i[\alpha; \delta\alpha]] + \dots\end{aligned}\quad (4.1)$$

These are exactly the formulas (3.6) for $\hat{A}_\mu$ and $\hat{\lambda}$, the commutator being now a gauge commutator instead of a star commutator. Applying the results, without the hats and the stars, we see that the gauge transformation of a non-abelian A_μ is the expected one, and that the field strength vanishes. We can also point out that all gauge commutators vanish in the abelian case, and that A_μ and λ then reduce to the shorter expected formulas.

On the noncommutative side, we still have the expressions

$$\begin{aligned}\hat{A}_\mu &= \partial_\mu \alpha + \frac{i}{2}[\alpha \star \partial_\mu \alpha] + \frac{i}{6}[\alpha \star i[\alpha \star \partial_\mu \alpha]] + \dots \\ \hat{\lambda} &= \delta\alpha + \frac{i}{2}[\alpha \star \delta\alpha] + \frac{i}{6}[\alpha \star i[\alpha \star \delta\alpha]] + \dots\end{aligned}\quad (4.2)$$

- [4] K. Saraikin, *Comments on the Morita equivalence*, [hep-th/0005138].
- [5] B. Pioline and A. Schwarz, *Morita equivalence and T-Duality (or B versus Θ)*, [hep-th/9908019].
- [6] A. A. Tseytlin, *On non-abelian generalisation of the Born-Infeld action in string theory*, Nucl. Phys. **B501** (1997) 41, [hep-th/9701125].
- [7] L. Cornalba, *On the general structure of the Non-Abelian Born-Infeld Action*, [hep-th/0006018].
- [8] L. Cornalba, *Some computations with Seiberg-Witten invariant actions*, Mod. Phys. Lett. **A16** (2001) 227-234, [hep-th/0101015].
- [9] N. Wyllard, *Derivative corrections to the D-brane Born-Infeld action: non-geodesic embeddings and the Seiberg-Witten map*, [hep-th/0107185].
- [10] H. Liu, **-Trek II: $\star_n$ Operations, Open Wilson Lines and the Seiberg-Witten Map*, [hep-th/0011125].
- [11] H. Liu and J. Michelson, *Ramond-Ramond Couplings of Noncommutative D-branes*, Phys. Lett. **B518** (2001) 143-152, [hep-th/0104139].
- [12] Y. Okawa and H. Ooguri, *An exact solution to Seiberg-Witten Equation of Noncommutative Gauge Theory*, [hep-th/0104036].
- [13] S. R. Das and S. P. Trivedi, *Supergravity couplings to Noncommutative Branes, Open Wilson lines and Generalized Star Products*, JHEP **0102** (2001) 046, [hep-th/0011131].
- [14] S. Mukhi and N. V. Suryanarayana, *Gauge Invariant Couplings of Noncommutative Branes to Ramond-Ramond Backgrounds*, JHEP **0105** (2001) 023, [hep-th/0104045].
- [15] S. R. Das, S. Mukhi and N. V. Suryanarayana, *Derivative Corrections from Noncommutativity*, JHEP **0108** (2001) 039, [hep-th/0104045].
- [16] T. Mehen and M. B. Wise, *Generalized $\star$ -Products, Wilson Lines and the Solution of the Seiberg-Witten Equations*, JHEP **0012** (2000) 008, [hep-th/0010204].
- [17] S. S. Pal, *Derivative corrections to Dirac-Born-Infeld and Chern-Simons actions from Non-commutativity*, [hep-th/0108104].
- [18] T. Asakawa and I. Kishimoto, *Comments on Gauge Equivalence in Noncommutative Geometry*, JHEP **9911** (1999) 024, [hep-th/9909139].
- [19] B. Jurčo, L. Möller, S. Schraml, P. Schupp and J. Wess, *Construction of non-Abelian gauge theories on noncommutative spaces*, [hep-th/0104153].
- [20] P. Schupp, *Non-Abelian gauge theory on noncommutative spaces*, [hep-th/0111038].

We can point out that here δ is non-abelian, but its non-abelian part is partly removed by the second term. The two other terms build up the covariant derivative of $\lambda^{(\infty,1)}$. It has the solution

$$\begin{aligned}\lambda^{(\infty,1)} &= -\frac{1}{4}\{A \sharp \partial \lambda\} \\ A_\mu^{(\infty,1)} &= -\frac{1}{4}\{A \sharp \partial A_\mu + F_\mu\} \\ F_{\mu\nu}^{(\infty,1)} &= -\frac{1}{2}\{F_\mu \sharp F_\nu\} - \frac{1}{2}\{A \sharp \frac{\partial + D}{2} F_{\mu\nu}\}\end{aligned}\quad (4.4)$$

where D is the gauge covariant derivative: $D = \partial - i[A, \cdot]$. As already emphasized, in the abelian case all these terms were of order A (meaning λA or AA_μ). Here they go to order A^3 . The question is whether they all belong to the first “line” anyway, higher orders in A being hidden in D and F , or not.

At order θ^2 , the equation is

$$\begin{aligned}\delta A_\mu^{(\infty,2)} - i[\lambda; A_\mu^{(\infty,2)}] - \partial_\mu \lambda^{(\infty,2)} + i[A_\mu; \lambda^{(\infty,2)}] = \\ i[\lambda^{(\infty,1)}; A_\mu^{(\infty,1)}] - \frac{1}{2}\{\partial \lambda^{(\infty,1)} \sharp \partial A_\mu\} - \frac{1}{2}\{\partial \lambda \sharp \partial A_\mu^{(\infty,1)}\} - \frac{i}{8}\{\partial \partial \lambda \sharp \partial \partial A_\mu\}\end{aligned}$$

for which solutions are known [19, 23, 24]. We have found another one, getting closer to the more or less expected form, but only fitting it partially. If we try to guess a recursive formula, we again face the problem of identifying order (1,1) and (2,2) among their respective columns. For example, the conjectured generalization for $F_{\mu\nu}^{(2,2)}$ could be

$$F_{\mu\nu}^{(2,2)} \stackrel{??}{=} -\frac{1}{2}\{F_{\mu\nu}^{(1,1)} \sharp F_{\nu\mu}\} - \frac{1}{2}\{A^{(1,1)} \sharp \frac{\partial + 2D}{3} F_{\mu\nu}\} \quad (4.5)$$

Paradoxically enough, the first part, which was harder to check in the abelian case, gives exactly the right result, while there are still problems with the second part.

ACKNOWLEDGMENTS

I would like to thank A. Armoni and R. Minasian for useful discussions and advice.

References

- [1] M. R. Douglas and N. Nekrasov, *Noncommutative Field Theory*, [hep-th/0106048].
- [2] R. J. Szabo, *Quantum Field Theory on Noncommutative Spaces*, [hep-th/0109162].
- [3] N. Seiberg and E. Witten, *String theory and noncommutative geometry*, JHEP **9909** (1999) 032, [hep-th/9908142].

- [21] B. Jurčo, P. Schupp and J. Wess, *Nonabelian noncommutative gauge fields and Seiberg-Witten map*, Mod. Phys. Lett. **A16** (2001) 343-348, [hep-th/0012225].
- [22] J. Madore, S. Schraml, P. Schupp and J. Wess, *Gauge Theory on Noncommutative Spaces*, Eur. Phys. J. **C16** (2000) 161-167, [hep-th/0001203].
- [23] S. Goto and H. Hata, *Noncommutative Monopole at the Second Order in θ* , Phys. Rev. **D62** (2000) 085, [hep-th/0005101].
- [24] D. Brace, B. L. Cerchiai, A. F. Pasqua, U. Varadarajan and B. Zumino, *A Cohomological Approach to the Non-Abelian Seiberg-Witten Map*, JHEP **0106** (2001) 047, [hep-th/0105192].

Annexe B

Mirror symmetric SU(3)-structure manifolds with NS fluxes

Cet article a été soumis pour publication à la revue *Communications in Mathematical Physics*. Il est aussi référencé sur www.arXiv.org comme hep-th/0311122.

La première version, reproduite ici, sera modifiée dès que possible. En effet, sans présager d'autres modifications avant sa publication effective, la formule (5.5) de cet article est incorrecte, et doit être remplacée par la formule (3.64) de ce mémoire. De même, dans la valeur finale indiquée pour W_3 sur la base holomorphe, en annexe, il faut lire trois fois $[d\lambda]$ et non $i[d\lambda]$.

1 Introduction

Mirror symmetry is a pairing between different compactifications which give rise to the same four-dimensional effective theory. For Calabi-Yau compactifications it is well-understood and has played an important role, becoming arguably the most interesting mathematical application of string theory. More general compactifications with fluxes on manifolds which are not Ricci-flat have become focus of much attention recently, and it would be important to extend to these at least partially the machinery which proved so useful for Calabi-Yaus.

If we had to consider only supersymmetric vacua, our search would be premature. The conditions on fluxes and warping to compensate non-Ricci-flatness and preserve supersymmetry are well-understood for some types of fluxes. To some extent, as we review later, these conditions are even translated into mathematical requirements: the manifold has to have $SU(3)$ structure and fall into a certain class in the mathematical classification of these objects. But Bianchi identity becomes an equation for which there is no existence theorem in the literature, unlike the famous Yau's theorem for Calabi-Yaus (not even the analogue of Calabi conjecture seems to have been formulated: this might be a task for string theory). If there is no singularity in the internal compact manifold, and the higher derivative terms are not taken into account, one can actually show even *non-existence* theorems.

Fortunately mirror symmetry as a more general equivalence of effective theories, and not only of vacua, still makes sense. As emphasized in [1], to have supersymmetry of the effective action $SU(3)$ structure is enough, without the extra requirements mentioned above, which ensure we are actually in a supersymmetric vacuum. Not only looking for mirror symmetric $SU(3)$ manifolds makes sense, but it is sensible to expect that a formal advance in this direction might help to understand the still elusive problem of compactifications with fluxes.

Much in this spirit, [1] (building on a comment in [2]) considered a particular case. Namely, an H flux on Calabi-Yau manifolds (without back-reaction: we are not dealing with a vacuum) get mapped to so-called half-flat manifolds, a particular class of $SU(3)$ structure manifolds, without any H flux. The amount by which these half-flat manifolds fail to be Ricci-flat is measured by a certain quantity called *intrinsic torsion*. We can thus say that, in this example, H flux gets exchanged by mirror symmetry with components of the intrinsic torsion associated with lack of integrability of the complex structure.

It is natural to wonder what happens in more general cases, when on both sides one has both H and intrinsic torsion. (As mentioned above, this for example is necessary in order to have supersymmetric vacua.) In the Calabi-Yau case, a concrete approach to mirror symmetry is Strominger-Yau-Zaslow (SYZ) [3] conjecture. This states that $\imath$ every Calabi-Yau is a T^3 special lagrangian fibration over a three-dimensional base, and $\Re$ mirror symmetry is T-duality along the three circles of the T^3 . It is natural to try and generalize this method to the present problem. Part $\imath$ of the conjecture came originally from considering moduli spaces of D-branes on Calabi-Yaus; generalizing this to background with fluxes does seem premature, and in any case we do not attempt it here, although later we will comment more on it. So we simply assume the manifold and flux we start with have this property, of admitting three Killing vectors. The idea is that the mirror transformations found in this class of examples will generalize to some extent to the most general case.

Having assumed this, we perform T-duality along the three isometries at once. T-duality will preserve four-dimensional effective theories, but since eventually we hope this procedure could be extended to more general situations by including singular fibers as in SYZ, we want to show why this should be called mirror symmetry – for that matter, indeed, why is there any mirror symmetry at all. A good framework for answering this is Hitchin's method based on Clifford(6,6) spinors [4]. As we review later in more detail, these are simply formal sums of forms on the manifold. Existence on a manifold of a Clifford(6,6) spinor without zeroes which is also *pure* (annihilated by half of the gamma matrices) is the same as saying that there is a $SU(3,3)$ structure on the manifold. (If the spinor is also closed, Hitchin calls these manifolds *generalized Calabi-Yaus*.) For a $SU(3)$ structure,

hep-th/0311122
CPHT-RR 106.1103

Mirror symmetric $SU(3)$ -structure manifolds with NS fluxes

Stéphane Fidanze, Ruben Minasian and Alessandro Tomasiello

Centre de Physique Théorique, Ecole Polytechnique
Unité mixte du CNRS et de l'EP, UMR 7644
91128 Palaiseau Cedex, France

Abstract

When string theory is compactified on a six-dimensional manifold with a nontrivial NS flux turned on, mirror symmetry exchanges the flux with a purely geometrical composite NS form associated with lack of integrability of the complex structure on the mirror side. Considering a general class of T^3 -fibered geometries admitting $SU(3)$ structure, we find an exchange of pure spinors ($e^{A,J}$ and Ω) in dual geometries under fiberwise T-duality, and study the transformations of the NS flux and the components of intrinsic torsion. A complementary study of action of twisted covariant derivatives on invariant spinors allows to extend our results to generic geometries and formulate a proposal for mirror symmetry in compactifications with NS flux.

November 17, 2003

where we denoted by H_1 and H_2 the components in the representations 3 and 6 of $SU(3)$ respectively, in analogy with the notation for torsions (see also the appendix A). So the next natural step would be to try and include more general classes, among which maybe supersymmetric vacua¹. In order to do so, it is natural to wonder to what extent the transformation rules can be put in a nicer form, and in particular be covariantized. It turns out that it is convenient to use spinors. Although also the previously mentioned Clifford(6,6) spinors can be used, here we mean a more conventional Clifford(6) spinor without zeros. One such spinor, call it ϵ , always exists on any $SU(3)$ structure manifold and can actually be used to define it. It turns out that using ϵ a different basis for W 's can be defined, which is diagonal under T-duality: elements of the basis transform picking a sign.

The idea of this different spinorial basis for W 's is roughly speaking the following. Usual W 's are defined, as we review later, from dJ and $d\Omega$. Now, not only ϵ is equivalent to the pair J, Ω , but the information contained in dJ and $d\Omega$ can also be completely extracted from $D_M \epsilon$. Using $SU(3)$ structure, this can be decomposed as $D_M \epsilon = (q_M + i\tilde{q}_M \gamma + i\tilde{q}_{MN} \gamma^N) \epsilon$, where γ is the chiral gamma in six dimensions and the group representations inside the quantities $q_M, \tilde{q}_M, q_{MN}$ are in one-to-one correspondence with the W 's. Switching to the spinorial basis accomplishes two things. First, it allows to capture the exchange of the pure spinors e^J and Ω and the exchange of their integrability properties simultaneously. More importantly, it allows to conjecture the six-dimensional covariantization of the mirror transformation (1.2), written in terms of the forms pulled back to the base of the T^3 fibration. Details can be found in section 5.

For the purposes of studying mirror symmetry/T-duality we will need first to introduce the covariant derivative twisted by the NS flux:

$$D_M^H \epsilon = (Q_M + i\tilde{Q}_M \gamma + iQ_{MN} \gamma^N) \epsilon \quad (1.4)$$

where, as we will see in detail, Q 's are obtained from q 's by complexifying certain components of the intrinsic torsion by the matching components of the flux (as in (1.2)). We will show that their restrictions to the base (denoted by hatted quantities) transform as

$$\hat{Q}_{ij} \longrightarrow -\hat{\tilde{Q}}_{ij}, \quad \hat{Q}_i \longrightarrow -\hat{\tilde{Q}}_i. \quad (1.5)$$

We will then argue that this simplification is due to the simple transformation of the ten-dimensional spinors under T-duality.

Finally we will try, in section 5.1, to collect these several points of view to argue that in general a rule like

$$Q_{mn} \longleftrightarrow -Q_{m\tilde{n}}, \quad Q_m \longleftrightarrow -\tilde{Q}_m$$

should hold. This rule is consistent with what we found in the T^3 fibered case, and with the principle that supersymmetric vacua should map in supersymmetric vacua (not necessarily the same). There are however more checks that could be done if one understood better examples; we discuss this in section 6. For example, in the case we mentioned above of compactifications with H only described by (1.3), one should understand moduli spaces and then check that a kind of exchange of complex and Kähler moduli (although, as we will argue, this has to be taken with a grain of salt) should happen. This might be interesting for the problem of fixing moduli. We end with a section on open problems.

2 Geometric setting

We start with an introductory section on T-duality, mainly to fix the notations.

The six-dimensional manifold will be taken to be a T^3 fibration over a base B . Coordinates on the base will be denoted by (y^1, y^2, y^3) , and on the fiber by (x^1, x^2, x^3) . All the quantities will only depend on the y coordinates, so that the x directions are Killing vectors. We will use conventions for indices as follows:

¹There may actually be supersymmetric vacua involving T^3 fibrations, if other fields are included.

there are two pure spinors which are orthogonal and of unit norm. From this point of view it is natural to conjecture that mirror symmetry between two $SU(3)$ structure manifolds exchanges these two pure spinors. We can be more explicit if we compare this Clifford(6,6) spinor definition of $SU(3)$ structure with the more usual one, existence of a two-form J and three-form Ω obeying $J \wedge \Omega = 0$ and $\Omega \wedge \bar{\Omega} = (2J)^3/3!$. In these terms the two pure spinors are e^J and Ω . We can actually multiply first spinor by e^B leaving it pure [4]. So what we are claiming is

$$e^{B+J} \longleftrightarrow \Omega. \quad (1.1)$$

The arrows here will be made precise in section 3. In the Calabi-Yau case, this exchange is implicit in many applications of mirror symmetry. For example, the even periods and the D-brane charge can be written using e^{B+J} , and its exchange with Ω was used in mapping [5] string-corrected DUY equations [6] to the special lagrangian condition; e^{B+J} was also used in formulating the concept of Π -stability [7].

With this in mind, we check that T-duality along T^3 (when it is possible) realizes the exchange (1.1), and for this reason we call it mirror symmetry. In this sense we have generalized part (i) of SYZ. However as it stands, (1.1) is hardly useful in predicting the mirror background starting from a particular six-manifold and NS flux.

After having discussed and justified the method, we can schematically describe here our results. The usual quantities which measure non-Ricci-flatness of the $SU(3)$ manifold are the five components of the intrinsic torsion (mentioned above) labeled as W_i , $i = 1 \dots 5$, in the representations $1 \oplus 1, 8 \oplus 8, 6 \oplus \bar{6}, 3 \oplus \bar{3}, 3 \oplus \bar{3}$ respectively. What is puzzling at first is that one does not see many ways of mirror pairing these representations, except for W_4 and W_5 which are two vectors. The answer is that the two mirrors have indeed two different $SU(3)$ structures: the two $SU(3)$ are differently embedded into $\text{Spin}(6,6)$, because the fiber directions change from tangent bundle to cotangent bundle, roughly speaking. As a result, representations get actually mixed. What is preserved is the representations that these objects have once pulled back to the base manifold, which is untouched by T-duality. W_2 and W_3 get then split as $W_2 = w_2^1 + w_2^2$ ($8 \rightarrow 5 \oplus 3$) and $W_3 = w_3^1 + w_3^2$ ($6 \rightarrow 5 \oplus 1$), and we get

$$\begin{aligned} W_1 - iH_1 &\longleftrightarrow -(\bar{W}_1 - i\bar{H}_1), \\ w_2^1 &\longleftrightarrow w_3^1 - i\bar{h}_3^1, \\ w_5, w_2^2 &\longleftrightarrow w_4 - i\bar{h}_4. \end{aligned} \quad (1.2)$$

A more detailed discussion of these equations can be found in section 4 (see in particular (4.11) for the precise statement). In (1.2) one can see that W_1, W_3, W_4 get naturally complexified by the components of H in the corresponding representation. This is no surprise as these torsions appear, as we review later, in dJ , and the natural object in string theory is always $B + iJ$. In the present context, this arises rather due to the usual combination $E = g + B$ of T-duality. As we will specify in section 2, we mostly work with a purely base-fiber type B -field, which is not the most general form allowed by T^3 invariance. However, we will see that this is just a simplifying technical assumption, and may eventually be relaxed. Note also that (1.2) complements (1.1) in an essential way by specifying in a more practical fashion the data of the mirror background (the metric and the NS flux). In particular, it quantifies the exchange of components of the flux and the intrinsic torsion on mirror sides.

We have also indicated in (1.2) that some of the components of the intrinsic torsion we begin with are actually related; so T^3 fibrations are not the most general $SU(3)$ manifold. This in a way answers in the negative the question about generalizing part (i) of SYZ: we are getting mirror symmetry only for a subclass of manifolds. In particular, supersymmetric vacua with only the three-form switched on are outside this class: indeed the conditions for these are [8] (reinterpreted in terms of torsions for example in [9] and [10])

$$W_1 = W_2 = 0, \quad W_3 = *H_3, \quad W_4 = d\phi = iH_4, \quad W_5 = 2d\phi \quad (1.3)$$

one can show that the T-dual metric and B field can be obtained by the original ones (2.1), (2.2) by the substitutions

$$h_{\alpha\beta} \longleftrightarrow \hat{h}^{\alpha\beta}; \quad B_{\alpha\beta} \longleftrightarrow \hat{B}^{\alpha\beta}; \quad B_\alpha \longleftrightarrow \lambda^\alpha. \quad (2.6)$$

and leaving the g_{ij} and B_{ij} in variant. Notice that last equation in (2.6) means that the twisting of each of the three S^1 bundles gets exchanged with B field. This fact played for example a role in a number of applications and was recently formalized in mathematical terms in [11].

We can also find the vielbein $\hat{V}^{\alpha\beta}$ of the T-dual metric $\hat{h}^{\alpha\beta}$, that satisfies $\hat{V}^{\alpha\beta} \hat{V}_{\alpha\beta} = \hat{h}^{\alpha\beta}$.

$$\hat{V}^{\alpha\beta} = \left(\frac{1}{h+B} \right)^{\alpha\beta} V_\beta^\alpha = V_\beta^\alpha \left(\frac{1}{h-B} \right)^{\beta\alpha} \quad (2.7)$$

whose inverse is

$$\hat{V}_\alpha^\beta \equiv \hat{h}^{\alpha\beta} \hat{V}^{\alpha\beta} = (h-B)_{\alpha\beta} V^{\alpha\beta} = V^{\alpha\beta} (h+B)_{\beta\alpha}. \quad (2.8)$$

The T-duality transformations of the vielbeine then are:

$$V_\alpha^\beta \longleftrightarrow \hat{V}^{\alpha\beta}; \quad V^{\alpha\beta} \longleftrightarrow \hat{V}_\alpha^\beta. \quad (2.9)$$

We will mostly work in the case when the B -field is purely of base-fiber type in frame indices. Transformation (2.6) shows that this condition is conserved by T-duality, while (2.5) reduces to $\hat{h}^{\alpha\beta} = h^{\alpha\beta}$. Consequently, $\hat{V}^{\alpha\beta} = V^{\alpha\beta}$ and $\hat{V}_\alpha^\beta = V_\alpha^\beta$. T-duality then only amounts to moving fiber indices up and down (still exchanging B_α and λ^α though).

For later use, we also define here the tensors defining the $SU(3)$ structure. These would be a priori only a two-form J and a three-form Ω satisfying $J \wedge \Omega = 0$ and $\Omega \wedge \Omega = (2J)^3/3!$, but here we define the structure in a more conventional way starting from an almost complex structure. The latter is defined by giving the (1,0) vielbein

$$E^A = ie_1^\alpha dy^1 + V_\alpha^{a'} e^a \quad (2.10)$$

where $A = a = a'$ goes from 1 to 3. The corresponding (0,1) vielbein is $E^{\bar{B}} = \bar{E}^{\bar{B}}$. This almost complex structure is in general not integrable, as (even after rescaling) it is not expressible as d of a complex coordinate, $E^A \neq \alpha^A dz^A$. However, with an abuse of language we will use the quantity

$$dz^j \equiv dy^j - iV_j^{\bar{a}} e^{\bar{a}} = -ie_j^{\bar{a}} E^{\bar{a}}, \quad (2.11)$$

keeping in mind that there is no reason for an actual coordinate z^j to exist. We also used in this expression the notation

$$V_{\alpha\alpha} \equiv \delta_{\alpha\alpha'} e_1^\alpha V_{\alpha'}^{\alpha'} = e_1^\alpha V_{\alpha\alpha}$$

The two-form J (sometimes called *fundamental form*) is defined by

$$J = \frac{i}{2} \delta_{AB} E^A \wedge E^{\bar{B}} = \frac{i}{2} g_{ij} dz^i \wedge d\bar{z}^j = -V_{i\alpha} dy^i \wedge e^\alpha \quad (2.12)$$

The holomorphic 3-form reads instead

$$\Omega = E^1 \wedge E^2 \wedge E^3 = \frac{1}{6} \epsilon_{ABC} E^A \wedge E^B \wedge E^C = -\frac{i}{6} \epsilon_{ijk} dz^i \wedge dz^j \wedge dz^k, \quad (2.13)$$

where $\epsilon^{ijk} = \epsilon_{abc} e^{a1} e^{b2} e^{c3}$.

The choices we are making for the $SU(3)$ structure are inspired by the SYZ approach. As we stressed above, these choices reduce the structure group further and are thus not to be expected to be as general (not even locally) as the T^3 fibration structure was in the SYZ approach. In particular some unphysical features will arise later in the dual of the complex coordinates. Anyway, in section 5.1 we will try to amend to this loss of generality.

- $i, j, k, \dots$ are used in the 3d y subspace,
- $\alpha, \beta, \gamma, \dots$ are used in the 3d x subspace,
- $M, N, \dots$ are used in the total 6d space for real coordinates: $dy^M = (dy^i, dx^{\alpha'})$,
- $m, n, \dots$ are used for holomorphic/antiholomorphic indices,
- $A, B, C, \dots$ are indices in the total 3d complex frame space,
- $a, b, c, \dots$ and $a', b', c', \dots$ are used in the 3d real y and x frame spaces. Primes will be dropped quickly.

We write then the most general metric and B field as ²

$$ds^2 = g_{ij} dy^i dy^j + h_{\alpha\beta} e^\alpha e^\beta = G_{MN} dy^M dy^N \quad (2.1)$$

$$B_2 = \frac{1}{2} B_{ij} dy^i \wedge dy^j + B_\alpha \wedge (dx^\alpha + \frac{1}{2} \lambda^\alpha) + \frac{1}{2} B_{\alpha\beta} e^\alpha \wedge e^\beta \quad (2.2)$$

where $\lambda^\alpha = \lambda_\alpha^\alpha dy^\alpha$, $B_\alpha = B_{\alpha\alpha} dy^\alpha$ and we have defined

$$e^\alpha \equiv dx^\alpha + \lambda^\alpha.$$

Of course the vielbein reads $(e_1^\alpha dy^\alpha, V_\alpha^{a'} e^a)$, where

$$\begin{aligned} \delta_{ab} e_1^\alpha e_1^\beta &= g_{ij}, & \delta_{a'b'} V_\alpha^{a'} V_\beta^{b'} &= h_{\alpha\beta}, \\ g^{ij} e_1^\alpha e_1^\beta &= g^{\alpha\beta}, & h^{\alpha\beta} V_\alpha^{a'} V_\beta^{b'} &= \delta^{a'b'}; \end{aligned}$$

we also record that the inverse vielbein has instead the form

$$e_1^\alpha \left(\frac{\partial}{\partial y^i} - \lambda_i^\alpha \frac{\partial}{\partial x^\alpha} \right), \quad V_\alpha^{a'} \frac{\partial}{\partial x^a}. \quad (2.3)$$

T-duality along the three x^α directions can be expressed conveniently in terms of the quantity $E = g + B$:

$$\begin{aligned} E_{ij} dy^i dy^j + E_{\alpha\alpha} dy^\alpha dx^\alpha + E_{\alpha\alpha} dx^\alpha dy^\alpha + E_{\alpha\beta} dx^\alpha dx^\beta \\ \mapsto E_{ij} dy^i dy^j + E^{\alpha\beta} (dx^\alpha + \lambda^\alpha) (dx^\beta + \lambda^\beta) + E_{\beta\gamma} dx^\beta dx^\gamma; \end{aligned} \quad (2.4)$$

notice that in this expression all the (implicit) tensor products are neither symmetrized nor antisymmetrized, for example $dy^i dy^j = dy^i \otimes dy^j$. Also remark that in this expression we used dy^i, dx^α basis instead of dy^i, e^α as virtually everywhere else. $E^{\alpha\beta}$ is the inverse of $E_{\alpha\beta}$ and can be decomposed in symmetric and antisymmetric part:

$$E^{\alpha\beta} = \left(\frac{1}{h+B} \right)^{\alpha\beta} = \hat{h}^{\alpha\beta} + \hat{B}^{\alpha\beta} \quad \text{where} \quad \begin{cases} \hat{h} = \frac{1}{h+B} \frac{1}{h-B} \\ \hat{B} = \frac{1}{h+B} (-B) \frac{1}{h-B} \end{cases} \quad (2.5)$$

The objects $\hat{h}$ and $\hat{B}$ would also be called in other contexts H and Θ . (In this paper H denotes instead the three-form field.)

Using the relations

$$\hat{B} \hat{h}^{-1} = -\hat{h}^{-1} B; \quad \hat{h}^{-1} \hat{B} = -B \hat{h}^{-1}; \quad \hat{B} \hat{h}^{-1} \hat{B} = \hat{h} - \hat{h}^{-1}$$

² with the convention that $\omega_1 \wedge \omega_2 = \omega_1 \otimes \omega_2 - \omega_2 \otimes \omega_1$

3 Mirror symmetry as T-duality

We start in this section showing, as promised in the introduction, how e^{B+J} and Ω get exchanged by T-duality.

First we do the easier case, in which there is no B field and λ twisting of the T^3 bundle. The basic idea is that Ω can be written in a sense as an exponential of the almost complex structure J_M^N applied to a degenerate three-form $e_{ijk}dy^i/dy^j/dy^k$, that can be thought of as the holomorphic three-form in the large complex structure limit. A way to be more explicit is the following. Expand Ω from (2.13) using the expression for the holomorphic vielbein in (2.10). One obtains four terms, with dy^3 , dy^2e , and so on. Define now the operation $V_{-J}^{\perp}(\cdot)$ by

$$V_{-J}^{\perp}(e^{e_1} \dots e^{e_k}) = \frac{1}{(3-k)!} e^{e_1} \dots e^{e_k} e_{e_{k+1}} \dots e_{e_3}, \quad k = 0 \dots 3.$$

This is essentially a Hodge star on the fiber, except it sends a k -form in the fiber into a $3-k$ -vector (a section of $\Lambda^{3-k}T$). Lower e_a are indeed vectors $\partial_a \equiv \partial/\partial x^a$. This operation is very similar to the T-duality transformation of spinors to be discussed shortly. Using this, on every component of the expansion in dy and e of Ω , we get a sum of (k, k) tensors, namely k indices up and k down: those down are along the fiber. The sum can be expressed as an exponent of $V_{-J}^{\perp}e_a dy^a$, which is the complex structure. T-duality is now easy to perform. According to (2.9) its action is simply to raise and lower α index: the tangent bundle (in the fiber direction) of the starting manifold is equal to the cotangent bundle (again in the fiber direction) of the T-dual manifold. As a result the complex structure gets now mapped to $V_{\alpha} e^a dy^a$, the fundamental two-form J . So we have gotten

$$T(V_{-J}^{\perp}\Omega) = \frac{i}{3!} e^{e^J}. \quad (3.1)$$

The case with B -field and λ is less trivial. Although this is not strictly required here, we find already at this point helpful to think about this in terms of Clifford(d, d) spinors. So we make a brief intermezzo explaining these and then we get back to our computation. Much of this material is taken from [4].

3.1 Clifford(d, d) spinors

Clifford algebra is usually defined on the tangent bundle (or cotangent) of a manifold using the metric. In physical notation this amounts to defining d gamma matrices which satisfy $\{\gamma^M, \gamma^N\} = 2g^{MN}$, where g^{MN} is the metric on the cotangent bundle of the manifold. On $SU(3)$ manifold there is moreover a well-known representation of this Clifford algebra, on $\Omega^{0,p}$ forms. If we on the contrary forget about the metric (thus about the $SO(d)$ structure), this algebra cannot be defined.

If we consider, however, both the tangent and the cotangent bundles of the manifold at the same time, there is a natural pairing between them (namely contraction between a vector and a form, $(dy^M, \partial_N) = \delta^M_N$), in which the metric does not enter. This "metric" on $T \oplus T^*$ is block-off-diagonal

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

and thus of signature (d, d) . Concretely, what this means is that one has to define $2d$ independent gamma matrices, γ^M, γ_M , that satisfy $\{\gamma^M, \gamma^N\} = 0 = \{\gamma_M, \gamma_N\}$ and $\{\gamma^M, \gamma_M\} = \delta^M_M$. Even though the Clifford structure has been defined on $T \oplus T^*$, fortunately the algebra still has a representation in terms of the forms on the manifold. Only now we have twice the number of creators and annihilators, and instead of using simply $(0, p)$ forms as before, we have to use forms of all possible degrees. On this space $\oplus_{p=1}^d \Lambda^p T^*$, an explicit representation is

$$\gamma^M = dy^M \wedge, \quad \gamma_M = \iota_{\partial_M}. \quad (3.2)$$

In all this we stress again that we have to consider γ_M and γ^M as independent: we cannot raise and lower indices using the metric. In this Clifford(d, d) algebra, however, the usual Clifford(d) is embedded: indeed a combination of wedge and contraction in (3.2) is the more conventional Clifford product, and if we use that we can raise and lower indices.

As stated in the introduction, a pure spinor is one which is annihilated exactly by half of the gamma matrices. If we come back at the application we have in mind, both e^{e^J} and $\tilde{\Omega}$ are pure:

$$\begin{aligned} (\gamma^M - iJ^MN\gamma_N)e^{e^J} &= 0, & (\gamma^M - iJ^MN\gamma^N)\tilde{\Omega} &= 0, \\ (\gamma_M - iJ_M^N\gamma_N)e^{e^J} &= 0, & (\gamma_M - iJ_M^N\gamma^N)\tilde{\Omega} &= 0. \end{aligned} \quad (3.3)$$

The gammas that annihilate the pure spinor $\tilde{\Omega}$ are more familiar if one expresses them in holomorphic/antiholomorphic indices: $\gamma^{\alpha}\tilde{\Omega} = \gamma_{\alpha}\tilde{\Omega} = 0$. Indeed $\tilde{\Omega}$ is one of the Clifford vacua for the Clifford(d) representation mentioned above (this is why we wrote the relations for $\tilde{\Omega}$ rather than Ω). Let us also notice that the annihilators of the two Clifford(d, d) spinors in (3.3) become the same when we allow ourselves to raise and lower indices on gammas, that is, when we descend to Clifford(d): e^{e^J} becomes then an alternative expression for a Clifford vacuum of Clifford(d).

As already mentioned, this dual way of realizing Clifford(d) from Clifford(d, d) is obviously in the center of mirror symmetry - exchange of the the Kähler form and the holomorphic three-form (or their non-integrable generalizations) is seen as different choices of Clifford vacuum.

3.2 Back to mirror symmetry

In this section, the only parts of the above theory that we actually use here are the formulas for the annihilators (3.3), which of course could have been derived independently. This insight gives however a useful rule of thumb, in particular when dealing with e^{e^J} , where we can save ourselves expanding the exponential as we did above. What we will do in the following will be to consider e^{e^J} and Ω as Clifford(6, 6) spinors, and other forms acting on them as combinations of gamma matrices. This is of course not the only possibility. One might have included B in the definition of the pure spinor. Due to technical details in how T-duality works we preferred this way. Also, we will work here in the case $B_{\alpha\beta} = 0$.

Let us consider for example the expression $e^{\beta\Omega}$. Due to $\gamma^{\alpha}\Omega = i\gamma^{\alpha}\gamma^{\beta}\Omega$, this equals $e^{\beta\alpha}\gamma^{\alpha}\Omega$. If we act on this with the operator $V_{-J}^{\perp}$ defined above, the prefactor can be taken out (it does not contain any e^{β}). On Ω we get $V_{-J}^{\perp}(\Omega) = e^{-iV^{\alpha}e_{\alpha}}$ as above; the only thing to notice is that e_{α} is simply ∂_{α} , as seen on (2.3). If we finally apply T-duality, $V^{\alpha}e_{\alpha} \mapsto V_{\alpha}e^{\alpha}$, putting it together with the inert factor $e^{iB_{\alpha}\wedge V^{\alpha}} = e^{iB^{\alpha}\wedge V_{\alpha}}$, we have shown

$$\frac{i}{3!} T(e^{e^J}) = V_{-J}^{\perp}(e^{\beta\Omega}) e^{-B_{\alpha}\wedge V^{\alpha}}. \quad (3.4)$$

It is a little surprising that the B field has to be subtracted on right hand side rather than being already present on left hand side. In the same way we can also prove the more reassuring

$$T(\tilde{\Omega}) = \frac{i}{3!} V_{-J}^{\perp}(e^{\beta}e^{e^J}) e^{B_{\alpha}\wedge V^{\alpha}}. \quad (3.5)$$

The exchange $e^{B+J} \longleftrightarrow \Omega$ as presented in (3.4) and (3.5) is not very aesthetically pleasing, however the exponents involving the T-duality anti-invariant $B_{\alpha}\wedge V^{\alpha}$ are easy to explain going back to (3.3). The condition of purity e.g. on Ω is essentially $d\alpha\wedge\Omega = 0$, and the holomorphic coordinates change under T-duality. The reason of this is that the dx^i which we have defined above as $dy^i - iV^i e^{\gamma}$ has a λ hidden inside e^{γ} . Since λ gets exchanged with B due to (2.6), dx on the original manifold does not map exactly to dx , but $dx \rightarrow dx - i(B_{\alpha}V^{\alpha} - \lambda^{\alpha}V_{\alpha})$ shifting by another T-duality anti-invariant. Thus the role of $e^{\pm B_{\alpha}\wedge V^{\alpha}}$ is to compensate for this change, preserving the condition for purity.

Before we start computing, notice that from this classical definition it would be not obvious to guess transformation laws for W 's, other than some qualitative features. There are two vectors, but the 8 and the 6 are different representations. If one thinks already at this stage about decomposing in base representations, guessing becomes easier, but one feels rapidly the need for a more solid ground. One way, which we pursue in this section, is to compute blindly. The other way is to put J and Ω on a more symmetrical basis, using the formalism of Clifford(d, d), or, which is another manifestation of the same idea, to actually use the $SU(3)$ invariant spinor directly. We do this in next section.

4.2 Computations of torsions in the T^3 fibered case

We can now compute W 's from the expressions (2.12) and (2.13). This is done by doing contractions, partial or total, appropriate to isolate the component of interest. For example W_4 is computed contracting $J \cdot dJ$.³ First we give W_1, W_4, W_5 , expressed in the holomorphic basis.⁴ Note that W_1 is real and W_5 holomorphic, so that $W_4 = w_4^i dz^i + c$ and $W_5 = w_5^j dz^j$, while $W_1 = w_1$ is a scalar.⁵ These components read:

$$w_1 = -\frac{i}{12} \epsilon^{ijk} V_{\alpha} [d(V - i\lambda)]_{jk}^{\alpha} \quad (4.2)$$

$$w_4 = -\frac{1}{4} V^{\alpha\beta} [dV_{\alpha}]_{\beta} \quad (4.3)$$

$$w_5 = -\frac{1}{4} \{V_{\alpha}^j [d(V + i\lambda)]_{jk}^{\alpha} - h^{\alpha\beta} \partial_k h_{\alpha\beta}\} \quad (4.4)$$

where $[d(\cdot)]_{ij} = 2\partial_k(\cdot)_{ij}$.

We now pass to W_2 and W_3 . W_2 is a (1,1)-form, and W_3 is a real (2,1) $\oplus$ (1,2), and are written as

$$W_2 = w_2^i dz^i \wedge d\bar{z}^j, \quad W_3 = \frac{1}{2} w_3^{\alpha\beta} dz^{\alpha} \wedge d\bar{z}^{\beta} + c. \quad (4.5)$$

However, since the representation 6 can be expressed not only as a primitive (2,1) form, but also as a symmetric tensor with two holomorphic indices, we will give this latter expression for W_3 . The way to pass from one to another is $w_3^{\alpha\beta} = w_3^{\alpha\beta} \Omega^{\alpha\beta}$. This is already a little in the spirit of the different basis for intrinsic torsion that we will give later. Furthermore, these two matrices with indices ij can actually be further decomposed in representation theory of the $SO(3)$ of the base. w_3^i has a symmetric and an antisymmetric part; the symmetric part does not drop out, it only contributes to $dJ \wedge e$ part; the antisymmetric part can be dualized to a three dimensional vector $w_3^j = \frac{1}{2} \epsilon^{ijk} w_3^{\alpha\beta}$. As for W_3 , $w_3^j = w_3^{\alpha\beta} \Omega^{\alpha\beta} + \frac{1}{3} w_3^j g_{ij}$ is already symmetric but has a trace part w_3^j on the three-dimensional base (of course, it is traceless in six dimensions).

$$w_2^i = \frac{1}{24} \epsilon^{ijk} [d(V - i\lambda)]_{jk}^{\alpha} [2V_{\alpha\beta} g_{ik} - 3V_{\alpha} g_{ik} - 3V_{\beta} g_{ik}] \quad (4.6)$$

$$w_3^j = -\frac{1}{4} V_{\alpha}^j [d(V - i\lambda)]_{jk}^{\alpha} \quad (4.7)$$

³Complete expressions for all five components of the intrinsic torsion for a metric of the form (2.1) can be found in appendix A.

⁴In what follows, we will denote W 's in complex coordinates as lower case w 's. For example, w_2^i is W_2^i , even if we did not explicitly mark i and j as holomorphic and antiholomorphic indices in this expression. This is also true for the other components.

⁵As already emphasized, one has to bear in mind that the almost complex structure is in general not integrable, so that $d\bar{z}^i$ is not to be understood as the differential of a hypothetical coordinate $\bar{z}^i$.

The combinations ϵ^{IJ} and Ω allow, as we have commented on in the introduction and as we will see further later on, to treat J and Ω more symmetrically. The most symmetrical object one might imagine is actually the $SU(3)$ invariant spinor ϵ itself. Given also the role that we anticipated it will have in torsions, one might wonder at this point if it is more convenient to use T -duality transformation of ϵ and forget all the rest. The problem is, so to say, that the spinor is too symmetric. The transformation rule of the ten-dimensional spinors are known: in the case without $B_{\alpha\beta}$, we simply have $\psi_+ \rightarrow \psi_+$, $\psi_- \rightarrow \gamma_f \psi_-$, where γ_f is the product of the three gammas in the fiber directions [12]. However, when we express $\psi_{\pm}$ in terms of the chirality projected $\epsilon_{\pm}$ of the six-dimensional spinor, $\gamma_f \epsilon_{\pm}$ is actually $\epsilon_{\pm}$ and all the information we get is that a IIA compactification has been exchanged with a IIB one. This means that the spinor is essentially on both the original and the T -dual manifold the pull-back of a spinor in the base. Still, using the familiar bilinear definitions for J and Ω (5.2) and $\gamma_f \epsilon_{\pm} = \epsilon_{\pm}$, one can show the identities above in a different way.

4 Intrinsic torsions and their duals

This section is the technical core of the paper. Here we define and compute intrinsic torsions for our T^3 fibered manifolds. As stressed in the introduction, these are not the most general $SU(3)$ structure manifolds. Performing T -duality along the T^3 is then easy using (2.6) and (2.9).

4.1 Conventional definition of torsions

We do not aim here at reviewing intrinsic torsions on manifolds with G -structures as discussions already exist in the literature, see for example [13] and among recent physics papers [1, 9]. Here we give a good working definition. It is familiar that, if we are on a $SU(3)$ holonomy manifold, not only J and Ω are well defined, but also they are closed: $dJ = 0 = d\Omega$. If they are not, dJ and $d\Omega$ give a good measure of how far the manifold is from having $SU(3)$ holonomy. The usual definitions require to split them in $SU(3)$ representations:

$$dJ = -\frac{3}{2} \text{Im}(W_1 \bar{\Omega}) + W_4 \wedge J + W_5 \quad (4.1)$$

$$d\Omega = W_1 J^2 + W_2 \wedge J + \bar{W}_2 \wedge \Omega$$

where the representations of the W_i are as follows:

- W_1 is a complex zero-form in $1 \oplus 1$;
- W_2 is a complex primitive two-form, so it lies in $8 \oplus 8$;
- W_3 is a real primitive (2,1) $\oplus$ (1,2) form, so it lies in $6 \oplus \bar{6}$;
- W_4 is a real one-form in $3 \oplus \bar{3}$;
- W_5 is a complex (1,0)-form (notice that in (4.1) the (0,1) part drops out), so its degrees of freedom are again $3 \oplus \bar{3}$.

These W_i allow to classify quickly any $SU(3)$ manifold. We will later define them in an alternative way using directly the spinor; that definition will be more natural for T -duality, but the W 's are often better to analyze the type of the manifold. For example, notice that in (4.1) the exterior derivative d does not satisfy the usual rule $d : \Omega^{p,q} \rightarrow \Omega^{p+1,q} \oplus \Omega^{p,q+1}$. For an almost complex manifold as we have here, there are also $(p+2, q-1)$ and $(p-1, q+2)$ contributions. Hence in (4.1) the $(3,0) \oplus (0,3)$ part of dJ , namely $\text{Im}(W_1 \bar{\Omega})$, and the $(2,2)$ part of $d\Omega$, which reads $W_1 J^2 + W_2 \wedge J$. So we know actually that $W_1 = W_2 = 0$ iff the manifold is complex. One can check indeed that the Nijenhuis tensor can be expressed in terms of W_1 and W_2 . Other examples of the use of these W 's abound in the literature. Notice also that the information of dJ and $d\Omega$ is a little redundant, as W_1 appears in both.

that appear in $D_M \epsilon$ transform better. Second, they might be useful in future occasions to analyze the geometry behind a given supersymmetry transformation without even having to bother to construct bilinears. In particular, we can find from this approach immediately the conditions (1.3) for supersymmetric vacua with H .

One proceeds in the following way. What we call ϵ in what follows is the $SU(3)$ invariant spinor, which can be furthermore decomposed by chirality as $\epsilon_+ + \epsilon_-$. Again, if we were on a manifold of $SU(3)$ holonomy, we would have a covariantly constant spinor, $D_M \epsilon = 0$. This is not the case, but still decomposing $D_M \epsilon$ into representations will give us a measure of how far we are from $SU(3)$ holonomy. The way of decomposing $D_M \epsilon$ into representations is again implicit in the literature. On a $SU(3)$ invariant manifold, a basis for spinors is given by ϵ_+ and $\gamma_M \epsilon_+$ (or alternatively we can trade ϵ_+ with ϵ and $\gamma \epsilon$). So, for example, anything else in Clifford algebra acting on ϵ , say $\gamma_{M_1 \dots M_n}$, can be reexpressed in terms of this basis. Explicit formulas for this are known (see for example [6]; in [14] a complete set of these equations are provided, along with the simple group theoretical description of how to get them, for the case of seven-manifolds with G_2 structure). We will not however need them here, it is enough to know that this decomposition can be done. Actually, with one exception: the relation $\gamma^M \gamma^N \epsilon = i J^M{}_N \gamma^N \epsilon$ can be used to eliminate one possible term. So we can write in general

$$D_M \epsilon = (q_M + i \tilde{q}_M \gamma + i q_{MN} \gamma^N) \epsilon. \quad (5.1)$$

The real q 's that we have defined in this equation are just another definition of intrinsic torsion. To see that they can be compared with the W 's above, it suffices to use group theory. q_M and $\tilde{q}_M$ are vectors, $3 \oplus \bar{3}$; as to q_{MN} , it can be decomposed into $(3 \oplus \bar{3}) \otimes (3 \oplus \bar{3}) \otimes (8 \oplus 1) \oplus (8 \oplus 1)$. We see that all the representations of the W 's are present. There is one redundancy, since we get three vectors (q_M , $\tilde{q}_M$ and one from q_{MN}). The objects we get in this way are the same as the W 's up to factors. Qualitatively we could stop here; in the present context we are actually interested in getting the factors, as they are important for being able to express q 's in terms of W 's explicitly. This is done as follows. After having decomposed q_{MN} as above, we can define J and Ω as bilinears

$$\epsilon^\dagger \gamma_{MN} \gamma \epsilon = i J_{MN}, \quad -i \epsilon^\dagger \gamma_{MN} P (\gamma + 1) \epsilon = \Omega_{MN} P. \quad (5.2)$$

One can now compute their exterior derivative using (5.1). Comparing the result with (4.1) gives the desired coefficients. The result is

$$q_{MN} = \frac{1}{4} (W_1^+ G_{MN} + W_1^- J_{MN}) + \frac{1}{8} (\Omega_{MN} P (\bar{W}_5 - 2 P W_4) + c.c.) \\ + \frac{1}{4} \left(-J_M P W_{PN}^2 + W_{MN}^2 \right) + \frac{1}{8} \text{Im}(W_{MN}^3) \quad (5.3)$$

$$q_M + i \tilde{q}_M = \left[\frac{1}{2} W_1 \bar{P}_{MN} + \frac{1}{4} \Omega_{MN} P (\bar{W}_5 - 2 W_4) + \frac{1}{2} \bar{P}_M P W_{PN}^2 - \frac{i}{8} W_{MN}^3 \right] \\ = [\bar{W}_5 - \bar{P} W_4]_M \quad (5.4)$$

where $\bar{W}_4 = W_4^+ - i W_4^-$ as usual in the literature, and we have defined $W_{MN}^3 = W_{MNP} \Omega^{PQ}{}_N$ and used a holomorphic projector $P = \frac{1}{2}(1 - iJ)$. We had observed in the previous section that the split of W_2 in two parts upon restriction to the base is crucial in the mirror transformation. Here we can see that the split nature of W_2 reveals itself in covariant six-dimensional expressions: W_2^+ and W_2^- enter respectively into the symmetric and antisymmetric parts of q_{MN} .

It is worth recording the same expression in holomorphic/antiholomorphic basis:

$$q_{mn} = -\frac{i}{16} w_{mn}^3 + \frac{1}{8} \Omega_{mnp} (w_5 - 2 w_4)^p, \quad q_{m\bar{n}} = -\frac{i}{4} w_{m\bar{n}}^2 + \frac{1}{4} \bar{w}_1 q_{m\bar{n}}. \quad (5.5)$$

And for the remaining vector:

$$q_m - i \tilde{q}_m = (w_5 - w_4)_m. \quad (5.6)$$

$$w_{(ij)0}^3 = \frac{1}{24} e^{\rho \eta k} [dV_a]_{pq} [2V_a^{\alpha} g_{ij} - 3V_j^{\alpha} g_{ik} - 3V_i^{\alpha} g_{jk}] \quad (4.8)$$

$$w_{ij}^3 = \frac{1}{8} e^{\rho \eta k} [d(V - 3t\lambda)]_{pq} V_{ka} \quad (4.9)$$

Before turning to the T-duality transformations of components of the intrinsic torsion and the flux, we observe that the conditions for a supersymmetric vacuum with H only (1.3) are not compatible in a nontrivial way with the expressions above, as we anticipated in the introduction. For example, demanding $W_1 = W_2 = 0$ sets λ and V to constants.

4.3 T-duality

It is now easy to see what the transformation rules of the W 's are. Decomposing in base representations says essentially where to look. One sees immediately that the various three-dimensional vectors and symmetric matrices are all similar. Before spelling this out, one should however stress that the full six-dimensional quantities have a more complicated transformation rule. As explained in section 3.2, due to presence of λ in e^T , $dx^i = dy^i - i V_a^i e^T$ on the original manifold does not map exactly to dx on the mirror side.

With this important caveat in mind, let us proceed to give T-duality transformations. As we said, many of the expressions we have for W 's are similar (see appendix A). The differences are mainly because of V_a versus V^a . This is already good, as these quantities are exchanged by T-duality (2.9). One also sees that some of the quantities contain λ , that after T-duality become B as we just recalled. So, we are led naturally to complexify some of the torsions adding dB projected in the appropriate representation. As this projections are verbatim those we did for dJ in previous subsection, this step is trivial. Thus defining components for H as for other forms⁶

$$H = -\frac{3}{2} \text{Im}(H_1 \bar{\Omega}) + H_4 \wedge J + H_5 \quad (4.10)$$

we find the transformations:

$$w_1 - i \tilde{w}_1 \longleftrightarrow -(\bar{w}_1 - i \tilde{\bar{w}}_1), \\ w_{(ij)}^2 \longleftrightarrow (w_3 + i \tilde{w}_3) (i \tilde{w}_j)^0, \\ w_k^2 - \frac{1}{4} h^{\alpha\beta} \partial_k h_{\alpha\beta} = \bar{w}_k^2 \longleftrightarrow (w_4 - i \tilde{w}_4)_k. \quad (4.11)$$

describing the mixing of the components of the flux and of the intrinsic torsion under mirror symmetry.

The central role in the mirror/T-duality transformation (4.11) is obviously played by W_2 (a component of the torsion associated with the non-integrability of the complex structure). It splits in two different pieces upon restriction to the base and the respective mixing of the two parts of W_2 with complexified H_3 and H_4 is an essential ingredient of the mirror map.

We will now try to rederive and generalize to generic geometries these results from a different point of view, using spinors rather than differential forms.

5 Spinorial basis

The idea is that the same information we have in dJ and $d\Omega$ are contained in $D_M \epsilon$. Doing the effort of reexpressing torsions in these terms pays off for several reasons. First of all, the combinations

⁶The explicit expressions for components of H in the T^3 -fibered geometry can be found in appendix A. We labeled these components so that they match the corresponding ones in dJ . H_1 is then the 1 $\oplus$ 1 complex scalar, H_3 the $6 \oplus \bar{6}$ real 3-form and H_4 the $3 \oplus \bar{3}$ real 1-form.

with the expressions

$$\tilde{Q}_i = Q_i + i\tilde{Q}_i = (\tilde{W}_5 - \frac{1}{2}(W_4 - iH_4))_i \quad (5.12)$$

$$\begin{aligned} \tilde{Q}_{ij} &= Q_{ij} - iQ_{ia}V_j^a = 2\tilde{P}^M Q_{iM} \\ &= \frac{1}{4} \left[\tilde{W}_1 + 3i\tilde{H}_{(1)} + \frac{i}{12} \tilde{P}^{iM} (W_5 - iH_5)_{Mj} \right] g_{ij} - \frac{i}{4} \left[\tilde{W}_2^2 + \frac{1}{4} \tilde{P}^M (W_5 - iH_5)_{Mj} \right] \\ &\quad + \frac{i}{2} \epsilon_{ijk} \left[W_5 - W_4 - iH_4 - \frac{1}{2} \tilde{W}_2 \right] \end{aligned} \quad (5.13)$$

This means that at an effective level the rule tells us $\epsilon_+ \leftrightarrow \epsilon_-$.

We should remark that working so far with a finite-size T^3 -fiber, we have extra (nowhere vanishing) vector fields, and thus reduces structure. This may in particular allow to locally preserve supersymmetry even when conditions (1.3) are violated. Since when fibers degenerate this restricted structure no longer exists, we avoided making explicit use of it, even though doing restrictions to the base manifolds implicitly uses the existence of a restricted structure. It is reasonable to expect that the results based on representations are valid over the entire moduli space, and thus next we turn to the six-dimensional covariantization of mirror transformation (4.11).

5.1 Approaches to the general case

At this point it is natural to wonder if we have enough information to simply guess what mirror symmetry should be in the general case. We have a precise set of transformation rules in the case of T^3 fibrations, and we also know that supersymmetric vacua should be sent to supersymmetric vacua. As we remarked above, T-duality is induced by an exchange of ϵ_+ with ϵ_- . Since we also have $\gamma^{\tilde{M}}\epsilon_- = 0$, these two facts together would suggest following proposal naturally generalizing (5.11):

$$Q_{mn} \longleftrightarrow -Q_{m\tilde{n}}, \quad Q_m \longleftrightarrow -\tilde{Q}_m. \quad (5.14)$$

We noticed above that representations of W 's do not match in such a way as to suggest immediately a transformation law. In the T-duality approach above this was solved by decomposing further in representations of the $SO(3)$ of the base. The proposal (5.14), on the contrary, gets around this problem collecting together $SU(3)$ representations rather than decomposing them further: qualitatively, $6 \oplus \bar{3} \leftrightarrow 8 \oplus 1$.

Let us now check that this proposal for mirror symmetry agrees with T-duality and with supersymmetry, as we just required. First of all, (5.14) agrees with the exchange (5.11). Indeed we have

$$\tilde{Q}_{M\tilde{A}} = P_m^P \tilde{Q}_{P\tilde{A}} + \tilde{P}^P \tilde{Q}_{P\tilde{A}} = 2Q_{m\tilde{A}} + 2Q_{m\tilde{A}}; \quad (5.15)$$

similarly one can consider the transformation of $Q_m = \tilde{Q}_m$.

Turning now to supersymmetry, the two transformations in (5.14) induce simply

$$D_m^H \epsilon_+ \longrightarrow -D_m^H \epsilon_- \quad (5.16)$$

So if only H is present we are sending $D_m^H \epsilon_+ = 0$ to $-D_m^H \epsilon_- = 0$; in the latter case supersymmetry is of course still preserved. In this form the duality might seem a little tautological, in the sense that it sends a supersymmetric vacuum in another one in an obvious way. Compare however with the usual mirror symmetry: a Calabi-Yau is sent to another Calabi-Yau, and the nontriviality lies in the exchange of Kähler and complex structure moduli. This should be happening for vacua with H only as well, and in a sense this would be yet another check to do; we will comment on this in next section.

Coming back to checking compatibility with supersymmetry, the situation becomes more complicated with RR fluxes, because the latter also transform, and one would have to check that they

The quantities we have defined so far would not be expected to behave nicely under T-duality, for the following simple reason. The transformation laws we have computed in (4.11) have, as one would expect also from the arguments in [1] and from (1.1), the feature of exchanging some torsions with H . Therefore we have to add a dependence on H to the covariant derivative in (5.1). Then also the q defined in (5.1) will change and (5.1) will become

$$D_M^H \epsilon = (Q_M + i\tilde{Q}_M \gamma + iQ_{MN} \gamma^N) \epsilon. \quad (5.7)$$

We have defined D^H (and as a consequence the Q 's) in such a way as to find good T-duality transformation properties afterwards. Not too surprisingly, we have found that the best definition is exactly the same as the one which appears in supergravity supersymmetry transformations: $D_M^H \equiv (D_M + \frac{i}{8} H_{MNP} \gamma^{NP})$. We find then

$$\begin{aligned} Q_{MN} &= \text{Re} \left[\frac{1}{2} (W_1 + 3iH_1) \tilde{F}_{MN} + \frac{1}{4} \Omega_{MNP} (\tilde{W}_5 - 2(W_4 + iH_4))^P \right. \\ &\quad \left. + \frac{i}{2} \tilde{P}^P W_{PN}^2 - \frac{i}{8} (W^3 + iH^3)_{MN} \right]. \end{aligned} \quad (5.8)$$

So, adding H as $D \rightarrow D^H$ complexifies W as $W + iH$, though at the end the Re in (5.8) makes the Q 's real. It should also be possible to write directly a formula for the (con)torion, as an alternative to formulas for the q 's that we have given. The fact that H appears as $H_{MNP} \gamma^{NP}$ tells us already that this formula will have a piece $K_{MNP} = dJ_{MNP} + \dots$ that will combine with H . As we will not need it here, we do not pursue this. Notice also that the G_2 analogue of what we just did for H is discussed in detail in [14] for the G_2 case.

The fact that the natural combination for T-duality and for supersymmetry is the same will be useful later, when we will try to extend our results to the general case. Then this is also a good place to see that of course the conditions for supersymmetry in the case with H only (1.3) can be recovered from the spinor equation. To have supersymmetry it is enough that one chirality, say ϵ_+ , is annihilated by D^H . We have the expressions

$$\begin{aligned} D_m^H \epsilon_+ &= (Q_m + i\tilde{Q}_m) \epsilon_+ + iQ_{mn} \gamma^n \epsilon_- & D_m^H \epsilon_- &= (Q_m - i\tilde{Q}_m) \epsilon_- - iQ_{mn} \gamma^n \epsilon_+ \\ D_m^H \epsilon_+ &= (Q_m + i\tilde{Q}_m) \epsilon_+ + iQ_{mn} \gamma^n \epsilon_- & D_m^H \epsilon_- &= (Q_m - i\tilde{Q}_m) \epsilon_- - iQ_{mn} \gamma^n \epsilon_+ \end{aligned} \quad (5.9)$$

Notice that $Q_{m\tilde{A}}$ and $Q_{\tilde{m}A}$ have disappeared from $D^H \epsilon_+$, because ϵ_- , being a Clifford vacuum, is annihilated by $\gamma^{\tilde{A}}$. From this one obtains directly that the complexified Q_{mn} and $Q_{\tilde{m}A}$ have to vanish. These will say that the complexified W_3 has to be purely antiholomorphic, which in more usual terms means of type (1,2) (this is the condition $W_3 = {}^* H_3$) and that W_2 has to vanish. The vectors require a little more care because usually the dilaton is rescaled in the metric (as a warping) and in the spinor itself. More generally it is clear that one can use gamma matrices identities mentioned above to reduce the expression to a form like (5.1), and then use (5.4) or (5.5).

For us the main advantage of having computed these quantities is to compare with T-duality transformations given in previous sections, although we will see shortly how these supersymmetry considerations can play a role in understanding the general case (without T^3 fibration structure). We can restrict the free index in (5.1) to be on the base, $M = i$, and furthermore apply a chirality projector

$$D_i \epsilon_+ = \tilde{Q}_i \epsilon_+ + i\tilde{Q}_{ij} \gamma^j \epsilon_-, \quad D_i \epsilon_- = \tilde{Q}_i \epsilon_- + i\tilde{Q}_{ij} \gamma^j \epsilon_+. \quad (5.10)$$

having introduced hatted quantities for restrictions to the base. The quantities $\tilde{Q}_i$ and $\tilde{Q}_{ij}$ in these expressions turn out to transform neatly under T-duality:

$$\tilde{Q}_i \longrightarrow -\tilde{Q}_i, \quad \tilde{Q}_{ij} \longrightarrow -\tilde{Q}_{ij}, \quad (5.11)$$

isomorphic to it (see the comments in previous section). Less trivial is the statement that complex and Kähler moduli are exchanged. To check this one would have first of course to know by what groups moduli spaces are computed.

This check we will not be able to perform here, and we limit ourselves to some comments. First of all, in general moduli spaces of solutions with fluxes are likely not to be simply factorized in Kähler and complex part. This is because, unlike the Calabi-Yau case, the conditions are no longer $dJ = 0 = d\bar{J}$, but something involving torsions; and the definition of torsions (4.1) mixes Ω and J . Also, in the Calabi-Yau case the fact that the conditions were of simple closure allowed to reduce the counting to a cohomology problem. In general, here, we are dealing with conditions involving projections $P_{\text{rep}} dJ$ and $P_{\text{rep}} d\bar{J}$, where P_{rep} is a projector on a certain representation. These conditions mean roughly that a form is closed “up to” a contribution from the other form, schematically $dJ = \text{operator}(\Omega)$. In general it should be possible to restate this as the cohomology of a double complex. Coming back to the case with H only switched on, a preliminary analysis of moduli spaces was sketched in [18], following ideas in [19]. Indeed the H -twisted cohomology groups proposed there are total cohomologies of a double complex with $\bar{\partial}$ and $H^{2,1}\wedge$ as differentials. We will unfortunately not say more on this here, but plan to come back on the issue in the future. For now we just observe that, in known examples, fluxes fix complex structure moduli. These considerations tell us that in a mirror picture Kähler moduli will be fixed.

Type B solution for IIB strings presented in [20] provides with another related case flux compactifications with the back-reaction taken into account. The metric now is conformally CY, and RR-fluxes are turned on as well. In addition, supersymmetry conservation imposes restrictions on H -flux, which now turns out to be primitive. We will not attempt here to present a complete analysis of the mirror transformation and will ignore the RR sector (which mirror symmetry maps to RR fields in Type IIA theory). Thus our starting data include $W_4 \sim W_5$ and H_4 . Note that this means in particular that we have $Q_{mn} = 0$, and thus we need $\bar{Q}_{mn} = 0$ on the mirror side. The two previous examples have this feature: we could either take $\bar{H} = 0$ and $\bar{W}_3 = \bar{W}_4 - \bar{W}_5 = 0$ or have $\bar{H}$ with imaginary selfdual primitive part and geometry given by $\bar{W}_5 = *H_3$ and $2\bar{W}_4 = \bar{W}_5 = 2d\bar{\phi} = 2iH_4$ as in [8]. However differently from previous cases $Q_{mn} \neq 0$ and this results in an additional non-integrability of the complex structure on the mirror side (in particular, $\bar{W}_2$ cannot be zero now). Of course, explicit constructions of such IIA string backgrounds would be of some interest.

The last application we will discuss here concerns the possibility of lifting the $SU(3)$ mirror symmetry picture to the G_2 -structure case. We could start from IIA string theory in a monopole background and lift it to M-theory, using the explicit relations between the components of intrinsic torsion for $SU(3)$ and G_2 for $U(1)$ -fibered manifolds. The components of the torsion for the representations 1, 7 and 27 get complexified by the corresponding representations of the G_4 -flux. The analogy with the $SU(3)$ case is rather close. There as well there was a number of components of the intrinsic torsion that get complexified by the H -flux; mirror symmetry then mixed these with the components corresponding to representations that are not contained in the flux (essentially $8 \oplus 8$ in that case, with some extra subtleties having to do with $3 \oplus \bar{3}$ appearing twice). In the G_2 geometry, 14 is such a representation, and the corresponding component of the torsion is the lifting of W_2^- [17, 21], the component of $SU(3)$ -torsion central in the exchange with the NS flux. Once more one would be hoping that going to spinorial basis and writing for the invariant spinor the twisted covariant derivative will lead to a covariant expression for a mirror transformation for the G_2 geometry. Indeed, as in (5.1) the torsion for the G_2 -structure manifolds is also encoded in a covariant derivative $D_M \epsilon = (q_M + iQ_{MN}\gamma^N)\epsilon$, where q 's are real. Then the eleven-dimensional supersymmetry transformations restricted to seven-dimensions twist the covariant derivative by a term $\frac{1}{3}G_{MNP} + G_{MN} + 2G_{NM}\gamma^N\epsilon$, where we have defined $G \equiv \frac{1}{3}G_{MNPQ}(*\Phi)_{MNPQ}$, $G_M \equiv \frac{1}{3}G_{MNPQ}\Phi_{NPQ}$, $G_{MN} \equiv \frac{1}{3}G_{MPQR}(*\Phi)^{PQR}$, using the associative form Φ . Putting all together we arrive at the twisted operator

$$D_M^G \epsilon = (Q_M + iQ_{MN}\gamma^N)\epsilon$$

which we can now use to extend the $SU(3)$ mirror symmetry proposal. Indeed, the G_2 analogue of

do it in a way compatible with the one we are giving for geometry and H . This can be elaborated as follows. Just as the entire NS contribution to the covariant derivative of the invariant spinor got summarized in Q 's (see (5.8)), the RR contribution can be accounted by introduction of similar objects, R_M , $\bar{R}_M$ and R_{MN} with a group decomposition matching that of Q 's. On supersymmetric backgrounds, the total action of the covariant derivative of the invariant spinor should be zero and thus $R = -Q$. Thus from this point of view the mirror transformation of the RR sector can also be brought to the form (5.14). From other side, in the T^3 fibered case, one could use the known transformation rules of RR fields. From the above, it is clear that the natural way to do this check in general would be to consider RR fields not as sums of forms but as bispinors, expressing for example in terms of the latter also supersymmetry transformations.

Even after all these motivations, the proposal (5.14) stands as a conjecture, and there would be other possible checks to be made. One possibility is to use again the formalism of Clifford(6,6) spinors. One can give an alternative definition of torsions, that we have not mentioned so far, using the Clifford(6,6) spinors e^{IJ} and Ω . Schematically one gets

$$D_M e^{IJ} = q_M e^{IJ} + \text{Im}(q_M^{(2)} \cdot \Omega), \quad D_M \Omega = (q_M + i\bar{q}_M)\Omega + q_M^{(2)} \cdot e^{IJ}. \quad (5.17)$$

In these equations, $q_M^{(2)}$ is the Clifford product of q_{MN} using only second index. These formulas seem indeed to be consistent with the general rule (5.14) given above.

6 Applications and examples

In this section we analyze some simple consequences of the mirror symmetry transformation that we have proposed. Apart from the case in which only geometry and B -field are present, the situation will be different from the usual one for Calabi-Yau's in that RR fluxes will transform, and so solutions with some types of fluxes switched on get mapped generically to solutions with other types of fluxes. On top of this we should also have the usual exchange of Kähler and complex structure moduli, in the sense of (1.1). Simple checks of both claims have been listed in previous section; here we take these statements for granted and examine the consequences.

The natural starting point is to check how the picture developed so far reduces to known cases. We start from a brief discussion of an example which has already been mentioned, and involves a CY manifold with B -field turned on. This case was considered in [1] in great detail. Since the intrinsic torsion vanishes on CY, we start from Q_{MN} built purely from components of H . The Q_{mn} gets a single contribution from H_1 . If we follow [1] and look for a purely geometrical mirror, on the mirror side we may have non-zero $\bar{W}_3$ and $\bar{W}_4 - \bar{W}_5 = 0$. Looking at Q_{mn} , we see that the reality of remaining components of the flux ensures that on the mirror side only $\bar{W}_1$ and $\bar{W}_2$ survive. This agrees with [1] up to a conventional $\pm$ exchange. So we recover as a particular case the half-flat geometries and the G_2 lifts discussed in [16, 17]. Note that neither the starting configuration, nor its mirror are vacua but rather domain walls.

The simplest background is when the B -field is turned off and we just deal with Calabi-Yau geometry. This case was also discussed in section 3, where we recover the exchange of the complex structure and (the exponentiated) Kähler form for mirror Calabi-Yau manifolds. An exchange of complex and Kähler moduli for a metric of the form (2.1) with $\lambda = 0$ and the integrability properties of its complex structure were studied in [15]. Here we easily see that the exchange of the e^{IJ} and Ω is accompanied by an exchange of their integrability conditions.

Without turning RR fields on, we can also consider yet another possibility of vacua.⁷ These cases are to obey the conditions given in (1.3). In our language these conditions read $Q_{mn} = 0 = Q_{mn}$. What one gets by the proposal (5.14) is the condition $Q_{mn} = 0 = Q_{mn}$, which is obviously

⁷Here and in the parts with also RR on, the word vacuum should be understood with the usual grain of salt: no-go theorems force us to consider noncompact or singular cases, or to hope (in a less well-defined way) that some of the features analyzed here will survive after taking into account higher-derivatives corrections.

spaces for D-branes wrapping (sub)manifolds. One may hope that eventually developing the picture of D-brane moduli spaces in geometries with NS fluxes may lead to refining the proposal for mirror symmetry presented here.

Acknowledgments. We would like to thank Peter Kaste for participation in early stages of this work, and Mariana Graña and Pawel Haseen for useful conversations. This work is supported in part by EU contract HPRN-CT-2000-00122 and by INTAS contracts 95-1-590 and 00-0334.

A Intrinsic torsion for T^3 -fibered manifolds

The components of the intrinsic torsion are defined by

$$\begin{aligned} dJ &= -\frac{3}{2} \text{Im}(W_1 \bar{\Omega}) + W_4 \wedge J + W_3, \\ d\Omega &= W_1 J^2 + W_2 \wedge J + \bar{W}_6 \wedge \Omega. \end{aligned}$$

They can be computed using contractions $(\cdot)$ with J and Ω :

$$\begin{aligned} W_1 &= \frac{4}{3} J^2 \lrcorner d\Omega = -\frac{4i}{3} \Omega \lrcorner dJ = \frac{1}{3} \epsilon_{ABC} (E^A \wedge E^B) \lrcorner dE^C \\ &= -\frac{i}{12} e^{ijk} V_{ia} [d(V - t\lambda)]_{jk}^\alpha \end{aligned}$$

$$\begin{aligned} W_2 &= 4J \lrcorner [d\Omega - W_1 J^2 - \bar{W}_6 \wedge \Omega] \\ &= \frac{1}{12} e^{ijk} [d(V - t\lambda)]_{jk}^\alpha [\theta_{pq} V_{\alpha} - 3g_{pq} V_{\alpha}] dx^p \wedge dx^q \end{aligned}$$

$$\begin{aligned} W_3 &= dJ + \frac{3}{2} \text{Im}(W_1 \bar{\Omega}) - W_4 \wedge J \\ &= \frac{3}{8} V_{ia} [d\lambda^\alpha]_{jk} dy^i \wedge dy^j \wedge dy^k \\ &\quad - \frac{1}{4} [dV_a]_{ik} [\frac{3}{2} \delta_j^k \delta_2^\alpha + V_j^\alpha V_k^\beta - 2V^{ka} V_{j\beta}] dy^i \wedge dy^j \wedge dy^k \\ &\quad + \frac{1}{4} [d\lambda^\alpha]_{jk} [\frac{1}{2} V_{ia} V_j^\beta V_k^\gamma - \delta_{ia}^j h_{\alpha\beta} V_\gamma^k] dy^i \wedge dy^j \wedge dy^k \\ &\quad - \frac{1}{8} V_j^\beta V_\gamma^j [dV_a]_{ik} e^\alpha \wedge e^\beta \wedge e^\gamma \\ &= \frac{1}{16} \{ -i(dV_a)_{jk} V_{\alpha}^\alpha + i(dV_a)_{ij} V_{\alpha}^\alpha + i(dV_a)_{ik} V_{\alpha}^\alpha + i(dV_a)_{ik} V_{\alpha}^\alpha g_{jk} - i(dV_a)_{ik} V_{\alpha}^\alpha g_{ij} \\ &\quad + i(d\lambda^\alpha)_{jk} V_{\alpha}^\alpha + i(d\lambda^\alpha)_{ij} V_{\alpha}^\alpha + i(d\lambda^\alpha)_{ik} V_{\alpha}^\alpha \} dx^i \wedge dx^j \wedge dx^k + \text{c. c.} \end{aligned}$$

$$\begin{aligned} W_4 &= 2J \lrcorner dJ = \frac{1}{2} V^{\alpha k} [dV_a]_{jk} dy^j \\ &= \frac{1}{2} h^{\alpha\beta} [dh_{\alpha\beta} - \mathcal{L}_g V_\beta] \end{aligned}$$

(5.14) can be written as

$$Q_{[MN]}^+ \longleftrightarrow -Q_{[MN]}^- \quad Q_{\{MN\}} \longrightarrow -Q_{\{MN\}} \quad (6.1)$$

where $\pm$ denote selfdual and anti-selfdual representations respectively. Note that only the former is complexified by G -flux, and (6.1) exchanges 14 with 7+7. In view of this, we may go back to (5.14) and note that there as well, modulo the trace part, mirror symmetry can be thought of as an exchange of selfdual and anti-selfdual matrices (à la Hermitian Yang-Mills).

7 Discussion

We conclude by mentioning some open technical and conceptual problems. Throughout the paper we have worked with a $B_{\alpha\beta} = 0$ case. Obviously, this choice simplifies greatly the T-duality transformation. The reason for this is most clear on the spinorial picture. As shown in [12], the only change in the simple T-duality transformations used above (see section 3.2) occurs when $B_{\alpha\beta}$ component of the B -field is nonzero. In this case we have to use instead

$$\psi_+ \rightarrow \psi_+, \quad \psi_- \rightarrow e^E \gamma_f \psi_-$$

where $E^{\alpha\beta}$ is defined in (2.5). We have here a gamma matrix exponential of $E \equiv \frac{1}{2} E^{\alpha\beta} \gamma_{\alpha\beta}$ which has the same form of the kappa-symmetry Γ operator; in the power series expansion the products of all gamma matrices are antisymmetrized.

Note that without a $B_{\alpha\beta}$ component, there is a certain ambiguity in the choice of T-duality invariants (5.12). The ambiguity is in the of complexification by H in Q 's. He have chosen everywhere the plus sign (and correspondingly T-dual expressions which become complex conjugates) for the following reason. The singlet representation allows a simple calculation even with a non-vanishing $B_{\alpha\beta}$. The result is then the first formula in (4.11), which fixes the ambiguity. For all other components we have chosen the complexification rule consistent with that of W_1 , hence the choice of sign in the definition of the twisted covariant derivative (5.7). The T-duality rule for the spinors given above should allow to lift restrictions from the B -field and verify this explicitly. We would like to emphasize though that this restriction is of technical nature - for a number of applications the B -field is generic enough. First, the H -flux contains all the representations it can. Second, in the holomorphic coordinate basis it is not hard to see that B is of generic type and contains both (1,1) and (2,0) components. The latter is important for several aspects of topological B-branes (see [22] for a recent discussion, in which also Clifford(d,d) spinors appear) and mirror symmetry [23].

Clearly there are two directions in which our results have to be extended. As mentioned many times we have worked with a T^3 fibration with finite-size fibers (and thus had a luxury of having extra vector fields without zeroes) and most of our formulae explicitly involve restrictions to the base of the fibration. At the end we succeeded in finding a basis in which the mirror/T-duality transformations can be covariantized and written over the entire six-dimensional manifold. The final simple rule for the mirror transformation

$$Q_{mn} \longleftrightarrow -Q_{mn}, \quad Q_m \longleftrightarrow -Q_m$$

is of group-theoretical nature, and we conjectured it to be true for general geometries, even without fibration structure at all, not even locally. In particular it should also work when there is a fibration but with singular fibers. From other side, singular T^3 fibers hold the key to SYZ picture, and would be extremely important to understand their fate in any generalization of SYZ.

Finally, one would like to complete the picture by incorporating D-branes. A better understanding of submanifolds in generalized CY manifolds as well as vector bundles on these would be essential preliminaries. Extending the picture developed in [6] for the exchange of branes (a pair of calibration and bundle conditions) and T-duality to generalized CY case would be of great interest. We may recall once more that in SYZ picture both mirror manifolds appear as moduli

$$\begin{aligned}
H_3 &= H + \frac{3}{2} \text{Im}(H_1 \bar{\Omega}) - H_4 \wedge J \\
&= \frac{1}{4} V^{\delta\bar{\delta}} V^{\gamma\bar{\gamma}} [\partial_k B_{\alpha\beta} - B_{\alpha\mu} \partial_k \lambda^{\bar{\mu}}] e^\alpha \wedge e^\beta \wedge e^\gamma \\
&\quad + \frac{1}{2} \left[\frac{5}{4} \partial_k B_{\alpha\beta} - V^{\gamma\bar{\gamma}} V^{\delta\bar{\delta}} \partial_k B_{\gamma\beta} - \frac{1}{2} V^{\gamma\bar{\gamma}} V^{\delta\bar{\delta}} \partial_k B_{\gamma\beta} \right] dy^k \wedge e^\alpha \wedge e^\beta \\
&\quad + \frac{1}{8} [\partial_k B_{ij} - \partial_k B_{j\bar{i}} \lambda_k^{\bar{i}} - \partial_k \lambda_j^{\bar{i}} B_{k\bar{i}}] [V_\beta^{\delta\bar{\delta}} V_\alpha^{\gamma\bar{\gamma}} dy^\beta - V_\beta^{\delta\bar{\delta}} V_\alpha^{\gamma\bar{\gamma}} dy^\beta + V_\beta^{\delta\bar{\delta}} V_\alpha^{\gamma\bar{\gamma}} dy^\beta] \wedge e^\alpha \wedge e^\beta \\
&\quad + \frac{1}{4} [(dB_\alpha)_{ik} - B_{\alpha\beta} (d\lambda^\beta)_{ik}] \left[\frac{3}{2} \delta_j^{\bar{k}} \delta_\alpha^{\bar{\gamma}} + V_j^\alpha V_\alpha^{\bar{\gamma}} - 2V^{\alpha\bar{\gamma}} V_{j\bar{\gamma}} \right] dy^k \wedge dy^j \wedge e^\alpha \wedge e^\gamma \\
&\quad + \frac{3}{8} [\partial_k B_{ij} - \partial_k B_{j\bar{i}} \lambda_k^{\bar{i}} - \partial_k \lambda_j^{\bar{i}} B_{k\bar{i}}] dy^j \wedge dy^i \wedge dy^k \\
&\quad - \frac{1}{8} V_\alpha^{\gamma\bar{\gamma}} V_\beta^{\delta\bar{\delta}} \partial_k B_{\alpha\beta} dy^k \wedge dy^j \wedge dy^i \\
&= \frac{1}{2} h_{ij}^3 dx^i \wedge dx^j \wedge dx^k + \text{c. c.}
\end{aligned} \tag{A.7}$$

In analogy with w 's (see (4.6) - (4.9)) we have introduced h :

$$h_{ij}^3 = h_{ijpq}^3 \Omega^{pq} = h_{\{ij\}^0}^3 + \frac{1}{3} h_{ij}^3 g_{ij} \tag{A.8}$$

$$\begin{aligned}
h_{\{ij\}^0}^3 &= -\frac{1}{24} e^{\mu\bar{\mu}} [dB_\alpha]_{\mu\bar{\mu}} [2V_\alpha^{\delta\bar{\delta}} g_{ij} - 3V_\alpha^{\delta\bar{\delta}} g_{ij} - 3V_\alpha^{\delta\bar{\delta}} g_{ij}] \\
h_{ij}^3 &= -\frac{1}{8} e^{\alpha\bar{\alpha}} [dB_\alpha]_{\mu\bar{\mu}} V_\alpha^{\delta\bar{\delta}} \\
&\quad - \frac{i}{8} e^{\alpha\bar{\alpha}} [dB_\alpha - \frac{1}{2} (dB_\alpha \wedge \lambda^\alpha + d\lambda^\alpha \wedge B_\alpha)]_{\mu\bar{\mu}}
\end{aligned} \tag{A.9}$$

References

- [1] S. Grier, J. Louis, A. Micu and D. Waldram, "Mirror symmetry in generalized Calabi-Yau compactifications," Nucl. Phys. B **654**, 61 (2003) [arXiv:hep-th/0211102].
- [2] C. Vafa, "Superstrings and topological strings at large N ," J. Math. Phys. **42** (2001) 2798 [arXiv:hep-th/0008142].
- [3] A. Strominger, S. T. Yau and E. Zaslow, "Mirror symmetry is T-duality," Nucl. Phys. B **479** (1996) 243 [arXiv:hep-th/9606040].
- [4] N. Hitchin, "Generalized Calabi-Yau manifolds," arXiv:math.dg/0209099.
- [5] N. C. Leung, S. T. Yau and E. Zaslow, "From special Lagrangian to Hermitian-Yang-Mills via Fourier-Mukai transform," Adv. Theor. Math. Phys. **4** (2002) 1319 [arXiv:math.dg/0005118].
- [6] M. Marino, R. Minasian, G. W. Moore and A. Strominger, "Nonlinear instantons from supersymmetric p-branes," JHEP **0001** (2000) 005 [arXiv:hep-th/9911206].
- [7] M. R. Douglas, B. Fiol and C. Romelsberger, "Stability and BPS branes," arXiv:hep-th/0002037.
- [8] A. Strominger, "Superstrings With Torsion," Nucl. Phys. B **274** (1986) 253.
- [9] J. P. Gauntlett, D. Martelli, S. Pakis and D. Waldram, "G-structures and wrapped NS5-branes," arXiv:hep-th/0205050.

$$\begin{aligned}
W_6^{(1,0)} &= -\Omega \wedge d\bar{\Omega} \\
&= \frac{1}{4} \{V_\alpha^{\delta\bar{\delta}} [d(V + i\lambda)]_{\bar{\delta}}^{\delta} + h^{\alpha\bar{\alpha}} \partial_j h_{\alpha\bar{\beta}}\} dx^{\bar{\delta}} \\
&= \frac{1}{4} \{[h_{\alpha\bar{\beta}} \mathcal{L}_g V^\alpha V^{\bar{\beta}}]_{\bar{j}} - iV_\alpha^{\delta\bar{\delta}} [d\lambda]_{\bar{j}}^{\delta}\} dx^{\bar{\delta}}
\end{aligned} \tag{A.1}$$

where $dx^{\bar{j}} = dy^{\bar{j}} - iV_j^{\bar{j}} e^{\bar{j}}$.

In the last two expressions, we have used the Lie derivative $\mathcal{L}$, which is defined by

$$\mathcal{L}_X Y = [X, Y] = X^i \partial_i Y^j - Y^i \partial_i X^j, \quad \mathcal{L}_X \omega = [X^i \partial_i \omega_j + \omega_i \partial_j X^i] dy^j, \tag{A.1}$$

on the vector field Y and the 1-form ω , with respect to the vector field X . We wrote V^α and V_β for the 1-forms $V_j^\alpha dy^j$ and $V_{\bar{j}} \beta dy^{\bar{j}}$, while gV^α and gV_α are the vector fields $V^\alpha \partial_\alpha$ and $V_\alpha \partial_\alpha$.

We also give here the components of the H field

$$\begin{aligned}
H &= dB_2 \\
&= \frac{1}{2} \partial_k B_{\alpha\beta} dy^k \wedge e^\alpha \wedge e^\beta + [\partial_k B_{i\alpha} - B_{\alpha\beta} \partial_k \lambda_i^{\bar{\beta}}] dy^k \wedge dy^i \wedge e^\alpha \\
&\quad + \frac{1}{2} [\partial_k B_{ij} - \partial_k B_{i\bar{j}} \lambda_k^{\bar{j}} + B_{i\alpha} \partial_k \lambda_j^{\bar{\alpha}}] dy^k \wedge dy^j \wedge dy^i
\end{aligned} \tag{A.2}$$

As a 3-form, we project H on representations of SU(3) as we did for dJ :

$$H = -\frac{3}{2} \text{Im}(H_1 \bar{\Omega}) + H_4 \wedge J + H_5 \tag{A.3}$$

These components are computed with the same contractions used for W_7 's:

$$\begin{aligned}
h_1 = H_1 &= -\frac{4i}{3} \Omega \wedge H \\
&= \frac{1}{12} e^{\delta\bar{\delta}} V_\alpha^{\gamma\bar{\gamma}} \partial_k B_{\alpha\beta} \\
&\quad + \frac{i}{12} e^{\delta\bar{\delta}} V_\alpha^{\gamma\bar{\gamma}} [dB_\alpha - B_{\alpha\beta} d\lambda_j^{\bar{\beta}}]_{jk} \\
&\quad - \frac{1}{12} e^{\delta\bar{\delta}} [\partial_k B_{ij} - \partial_k B_{i\bar{j}} \lambda_k^{\bar{j}} + B_{i\alpha} \partial_k \lambda_j^{\bar{\alpha}}]
\end{aligned} \tag{A.4}$$

$$\begin{aligned}
H_4 &= 2J \wedge H \\
&= \frac{1}{2} V^{\delta\bar{\delta}} [dB_\alpha - B_{\alpha\beta} d\lambda_j^{\bar{\beta}}]_{jk} dy^j - \frac{1}{2} V^{\delta\bar{\delta}} \partial_k B_{\alpha\beta} e^\beta \\
&= h_k^4 dx^k + h_k^4 dx^k
\end{aligned} \tag{A.5}$$

$$h_k^4 = \frac{1}{4} \{V^{\delta\bar{\delta}} [dB_\alpha - B_{\alpha\beta} d\lambda_j^{\bar{\beta}}]_{jk} - iV^{\delta\bar{\delta}} \partial_j B_{\alpha\beta} V_k^{\bar{\alpha}}\} \tag{A.6}$$

- [10] G. L. Cardoso, G. Curio, G. Dall'Agata, D. Lust, P. Marousselis and G. Zoupanos, "Non-Kähler string backgrounds and their five torsion classes," Nucl. Phys. B **652**, 5 (2003) [[arXiv:hep-th/0211118](#)].
- [11] P. Bouwknegt, J. Evain and V. Mathai, "T-duality: Topology change from H-flux," [arXiv:hep-th/0306062](#).
- [12] S. F. Hassan, "SO(d,d) transformations of Ramond-Ramond fields and space-time spinors," Nucl. Phys. B **583** (2000) 431 [[arXiv:hep-th/9912236](#)].
- [13] D. D. Joyce, *Compact manifolds with special holonomy*, Oxford, 2000.
- [14] P. Kaste, R. Minasian and A. Tomasiello, "Supersymmetric M-theory compactifications with fluxes on seven-manifolds and G-structures," JHEP **0307** (2003) 004 [[arXiv:hep-th/0303127](#)].
- [15] J. T. Liu and R. Minasian, "U-branes and T**3 fibrations," Nucl. Phys. B **510** (1998) 538 [[arXiv:hep-th/9707125](#)].
- [16] N. Hitchin, "Stable forms and special metrics," in *Global differential geometry: the mathematical legacy of Alfred Gray*, volume 288 of *Contemp. Math.*, pages 70-89. AMS, 2001. [[math.DG/0107101](#)].
- [17] S. Chiossi and S. Salamon, "The intrinsic torsion of SU(3) and G₂ structures," Proc. conf. Differential Geometry Valencia 2001, [math.DG/0202282](#).
- [18] K. Becker, M. Becker, P. S. Green, K. Dasgupta and E. Sharpe, "Compactifications of heterotic strings on non-Kähler complex manifolds. II," [arXiv:hep-th/0310058](#).
- [19] R. Rohm and E. Witten, *Annals Phys.* **170** (1986) 454.
- [20] M. Graña and J. Polchinski, "Gauge / gravity duals with holomorphic dilaton," Phys. Rev. D **65** (2002) 126005 [[arXiv:hep-th/0106014](#)].
- [21] P. Kaste, R. Minasian, M. Petrini and A. Tomasiello, "Nontrivial RR two-form field strength and SU(3)-structure," [arXiv:hep-th/0301063](#).
- [22] A. Kapustin, "Topological strings on noncommutative manifolds," [arXiv:hep-th/0310057](#).
- [23] A. Kapustin and D. Orlov, "Vertex algebras, mirror symmetry, and D-branes: The case of complex tori," Commun. Math. Phys. **233** (2003) 79 [[arXiv:hep-th/0010293](#)].

Annexe C

Proceedings des Journées Jeunes Chercheurs 2001

Cet article a été publié dans les proceedings des Journées Jeunes Chercheurs 2001.

Balade dans un espace-temps à dix dimensions... non-commutatives

Stéphane Fianza
CPHT - Ecole Polytechnique, Palaiseau

Résumé

Partons pour une balade aux pays des supercordes, bien trop courte pour comprendre tous les objets bizarres qui y vivent (qui sont d'ailleurs loin d'avoir livré tous leurs secrets), mais qui nous permettra de les découvrir. Nous irons admirer en particulier des espace-temps exotiques, dont les coordonnées ne commutent pas, et qui sont pourtant des cousins des espace-temps de nos contrées. Suivez le guide et ne donnez pas à manger aux animaux.

1 Introduction

Le monde qui nous entoure nous pose une grande question : de quoi est-il fait ? Quelles sont les particules élémentaires qui composent toute matière, quels sont les grands principes auxquels tout obéit ? Aujourd'hui, on ne sait répondre à cette question que partiellement, parce qu'on a encore besoin de deux morceaux pour tout décrire : la Relativité Générale qui décrit la gravitation, et le modèle standard qui décrit la physique des particules élémentaires. Il nous reste à trouver une théorie unique qui unifie ces deux aspects dans un seul langage. Le candidat le plus prometteur aujourd'hui est la théorie des cordes.

2 Un soupçon de théorie des cordes

2.1 Des cordes pour faire le monde ?

Si vous imaginez une particule, vous imaginez un point, et vous la représentez mathématiquement par un point. Il faut alors une petite "étiquette" pour différencier l'électron du photon ou du quark (sa masse, sa charge...). Maintenant, remplaçons ce point par une petite corde, comme un élastique, ouvert ou fermé. Bien sûr elle est très petite (environ 10^{-35} m), et vue de notre échelle, elle ressemble à un point. Mais mathématiquement, c'est un objet à une dimension. Et on n'a plus besoin d'étiquette, puisque, comme une corde d'instrument, elle peut vibrer en produisant différentes notes, qui nous donnent différentes particules. C'est la base de la théorie des cordes.

Reproduire les particules, c'est bien, mais ça ne veut rien dire si l'on ne reproduit pas les interactions. Et cela, en théorie quantique des champs, signifie reproduire les bosons vecteurs qui la transportent. Ce sont le photon de l'électromagnétisme (champ de jauge $U(1)$), les trois bosons $W^\pm$, Z_0 de l'interaction faible (champs de jauge $SU(2)$) et les huit gluons de l'interaction forte (champs de jauge $SU(3)$). Sans oublier le graviton qui propage la gravitation.

C'est ce graviton qui rend la métrique de la Relativité Générale dynamique, à partir de la métrique plate de Minkowski. Et ils sont bien sûr présents en théorie des cordes : ce sont les modes de basse énergie des cordes fermées. Etant fermées, celles-ci n'ont pas d'extrémités libres pour s'accrocher quelque part, contrairement aux cordes ouvertes comme nous le verrons bientôt, et peuvent baigner tout l'espace-temps. Mais si l'on regarde mieux ces cordes fermées, on s'aperçoit que le mode de basse énergie est en fait une matrice complète : sa partie symétrique $g_{\mu\nu}$ est la métrique de l'espace-temps, et sa partie antisymétrique $B_{\mu\nu}$ joue le rôle d'un champ électromagnétique de fond. Ainsi, la gravitation est naturellement incluse dans les théories de cordes. Nous reparlerons de l'importance du champ $B_{\mu\nu}$ à la section 3.

Avant d'expliquer comment obtenir des théories de jauge avec des cordes, faisons encore deux remarques générales, qui justifient l'intérêt que l'on porte à la théorie des cordes.

La première tient du vocabulaire : il faudrait chaque fois entendre "supercorde" et non "corde", signe que ces objets sont supersymétriques. Le modèle plus simple des cordes bosoniques, non supersymétriques, s'il offre un réel intérêt pédagogique, est entaché de difficultés que l'on ne sait pas résoudre. Ou plutôt, c'est en voulant les résoudre que l'on a introduit la supersymétrie.

Ma deuxième remarque tient aux divergences de théories des champs, qui conduisent à la renormalisation. Ces divergences apparaissent dans les diagrammes de Feynmann principalement à cause de la localité des interactions : les vertex sont les points d'intersection de plusieurs segments, qui conduisent dans le calcul à la multiplication de fonctions δ au même point, ce qui est mathématiquement mal défini. A ce problème, la théorie des cordes apporte une solution élégante, comparable à

ce que font les bosons intermédiaires pour les interactions de Fermi. La zone d'interaction est étalée, tout en conservant la covariance d'espace-temps, puisque ce sont des cordes, et plus des points, qui interagissent. Ainsi, les diagrammes de Feynmann deviennent des tuyauteries de plombier, qui sont lisses partout, et qui se comportent beaucoup mieux mathématiquement. Le développement perturbatif devient par la même occasion un développement topologique.

2.2 La terre des Dragons

Une fois que l'on a remplacé les objets ponctuels par des objets à une dimension, pourquoi s'arrêter ? En fait, il n'y a même pas besoin de rajouter autre chose, la théorie contient déjà des objets à deux, trois... dimensions, comme des petites feuilles, ou des bulles de savon : on les appelle des p -branes (membranes à $p + 1$ dimensions). D'autre part, la théorie a besoin de dix dimensions pour exister, neuf d'espace et une de temps, donc il y a de la place pour ces branes.

Pour être honnête, il ne s'agit pas d'une seule théorie, mais de cinq, avec des noms bizarres (de type I, IIA, IIB, Hétérotiques $E_8 \times E_8$ et $SO(32)$). C'est triste pour la théorie ultime de l'univers d'exister en cinq versions différentes ! Heureusement, on s'est aperçu en 1995, que ces cinq théories étaient en fait cinq morceaux d'un même objet, vu par cinq petits trous différents, cinq secteurs perturbatifs d'une seule théorie : la théorie M. On ne sait pas grand chose sur elle, elle reste donc encore Mystérieuse, mais il semble qu'elle ait besoin de onze dimensions pour être à son aise...

Revenons sur les branes, puisque ces objets sont les objets dynamiques fondamentaux qu'il faut comprendre. D'autant plus qu'elles vont nous permettre de parler des théories de jauge, que nous avons oubliées jusqu'ici. Dire que les branes sont dynamiques, ça veut dire que, comme les cordes dont nous sommes partis, elles peuvent se déplacer et vibrer. Leurs modes de vibration sont décrits par différents champs qui vivent sur sa surface, et qui peuvent s'interpréter comme des degrés de liberté de particules. Par exemple, le premier mode de vibration d'une brane, seule dans l'espace, peut être décrit par des cordes ouvertes, dont les extrémités sont attachées sur la brane. A leur tour, ces cordes ouvertes peuvent être interprétées comme des photons, autrement dit un champ de jauge $U(1)$. Compliquons maintenant un peu l'image, et imaginons N branes empilées, comme un tas de feuilles de papier. Les cordes ouvertes peuvent maintenant fixer leur extrémités sur plusieurs surfaces, et elles ne se privent pas. De nouveaux modes apparaissent donc sous la formes de petites cordes reliant entre elles deux branes différentes. Si ces branes sont éloignées d'une distance d , alors ces modes sont massifs. Mais si les branes coïncident, alors ces modes sont

de masse nulle, comme le photon de tout-à-l'heure, et forment tous ensemble une théorie de jauge $U(N)$ non abélienne.

3 Une pincée de géométrie non-commutative

La question que je me suis posée jusqu'à présent commence comme cela : que se passe-t-il lorsque l'on approche un aimant d'une brane ? Une brane, ça pourrait très bien être notre univers, à 3+1 dimensions. La réponse, c'est que ses coordonnées ne commutent plus : $xy \neq yx$. Déjà qu'il était courbe depuis la relativité générale, l'espace-temps est maintenant devenu non-commutatif ! Cependant, il y a toujours un photon (et plus généralement des théories de jauge), mais beaucoup plus compliqué que le photon ordinaire dont on a l'habitude. Il y a quand même une astuce : ils sont reliés par une transformation, compliquée, la transformation de Seiberg-Witten[2]. Ce que j'essaie de faire, c'est trouver la formule explicite de cette transformation[3].

3.1 Des dimensions non-commutatives

Soyons plus précis. Pour ne pas confondre avec les groupes de jauge non abéliens, où les objets ne commutent pas parce que ce sont des matrices, on parle, lorsque les coordonnées de l'espace-temps ne commutent pas, de non-commutatif. Un moyen simple pour obtenir une telle propriété est de modifier le produit que l'on utilise pour multiplier les fonctions. Si f et g sont deux fonctions de x^μ , on définit leur star-produit par

$$\begin{aligned} f \star g &= f e^{\frac{i}{2} \overleftrightarrow{\partial}_\rho \theta^{\rho\sigma} \overleftrightarrow{\partial}_\sigma} g \\ &= fg + \frac{i}{2} (\partial_\rho f) \theta^{\rho\sigma} (\partial_\sigma g) + O(\theta^2) \end{aligned}$$

où $\theta^{\rho\sigma}$ est constant : c'est le paramètre de non-commutativité. En effet, si l'on prend pour f et g les fonctions coordonnées

$$[x^\mu \star x^\nu] = x^\mu \star x^\nu - x^\nu \star x^\mu = i\theta^{\mu\nu}$$

Ainsi, au lieu de coordonnées bizarres, on a une multiplication bizarre, ce qui est quand même moins choquant.

Revenons alors au problème de la brane plongée dans un champ $B_{\mu\nu}$ (l'aimant). Plus précisément, c'est une brane ordinaire (ses coordonnées commutent : $\theta = 0$), avec le champ de fond $g + B$ (où g est la métrique). La transformation de Seiberg-Witten nous dit alors que cette situation est équivalente à une autre situation, dans laquelle la brane a pour paramètre de non-commutativité θ , et est plongée dans le champ de fond $\hat{g} + \hat{B}$, pourvu que

$$0 + \frac{1}{g + B} = \theta + \frac{1}{\hat{g} + \hat{B}}$$

Si l'on considère que l'on préfère travailler lorsque le champ B est nul, il suffit de choisir $\hat{B} = 0$, mais alors, θ ne sera pas nul, et les coordonnées sur la brane ne commutent plus.

3.2 Le photon non-commutatif

On dispose ainsi d'une infinité de descriptions équivalentes, qui nous permettent de choisir le point de vue que l'on préfère. Mais changer de point de vue, ça ne sert que si l'on sait encore calculer dans la nouvelle description. En particulier, qu'en est-il du champ de jauge $U(1)$, qui vivait sur la brane ordinaire, notre photon A_μ ? Est-il lui aussi transformé? La réponse est oui, et d'une très jolie manière, puisqu'il ressemble désormais beaucoup à une théorie de jauge non abélienne : le photon est devenu une sorte de gluon. En effet, le champ électromagnétique s'écrit maintenant

$$\hat{F}_{\mu\nu} = \partial_\mu \hat{A}_\nu - \partial_\nu \hat{A}_\mu - i[\hat{A}_\mu, \hat{A}_\nu]$$

et de même pour la transformation de jauge

$$\begin{aligned}\delta \hat{A}_\mu &= \partial_\mu \hat{\lambda} + i[\hat{\lambda}, \hat{A}_\mu] \\ \delta \hat{F}_{\mu\nu} &= i[\hat{\lambda}, \hat{F}_{\mu\nu}]\end{aligned}$$

Cette ressemblance va plus loin, puisque dans certains cas, que nous ne précisons pas ici, cette théorie est réellement équivalente à une théorie $U(N)$. C'est ce que l'on appelle l'équivalence de Morita.

Continuons encore un peu avec notre photon non-commutatif, pour préciser quelle est la transformation qui le lie au photon ordinaire, en faisant des cousins qui, s'ils ont l'air éloignés, n'en sont pas moins équivalents. Cette transformation s'appelle toujours transformation de Seiberg-Witten. Au paragraphe précédent, nous avons vu son action sur le champ de fond, tandis que son action sur la théorie de jauge peut être résumée par le diagramme suivant

$$\begin{array}{ccc} A_\mu & \xrightarrow{\delta} & A_\mu + \delta A_\mu \\ \downarrow & & \downarrow \\ \hat{A}_\mu & \xrightarrow{\delta} & \hat{A}_\mu + \delta \hat{A}_\mu \end{array} \quad \Leftrightarrow \quad \delta \hat{A}_\mu = \delta \hat{A}_\mu$$

où le champ $\hat{A}_\mu$ est une fonction du photon ordinaire A_μ (et du paramètre θ), et se transforme donc sous la transformation de jauge ordinaire δ . Cette équation, dite "de Seiberg-Witten", peut être résolue ordre par ordre en θ . Elle signifie que les orbites de jauge des deux théories sont les mêmes, bien que leurs transformations soient très différentes. J'insiste cependant sur le fait que les solutions ne sont pas linéaires, loin de là, ce qui explique qu'une théorie "simple", celle du photon, soit transformée en une théorie bien plus compliquée, celle des gluons.

Pour finir, voyons à quoi peut ressembler la solution de ces équations, c'est-à-dire à quoi ressemble $\hat{A}_\mu$ lorsque l'on connaît A_μ . Je ne donnerai ici que la solution simple lorsque A_μ est une pure jauge ($F_{\mu\nu} = 0$). On peut alors écrire

$$A = ig^{-1} dg$$

où $g = e^{-i\alpha}$ est un élément du groupe de jauge. A cet élément, on peut faire correspondre $\hat{g} = (\star e)^{-i\alpha}$, où tous les produits dans le développement de l'exponentielle sont des star-produits. On a alors simplement

$$\hat{A} = i\hat{g}^{-1} \star d\hat{g}$$

Si les formules, dans un cas plus général, pouvaient ressembler à celles-ci, tout irait pour le mieux dans le meilleur des mondes possibles. D'autant plus que les équations se résolvent ordre par ordre en θ , et rien n'assure que les solutions puissent ensuite être mises sous une telle forme non perturbative. D'ailleurs, en général, elles ne le sont pas. Pire, pouvoir être résolues ne signifie pas que l'on dispose effectivement des solutions. C'est pour toutes ces raisons que j'ai présentées les formules ci-dessus comme simples.

4 Conclusion

Il y aurait de nombreuses choses à dire encore, et il reste beaucoup de questions. En particulier pour comprendre à quoi tout cela peut bien servir. Une des raisons compliquées est que ces théories non-commutatives peuvent parfois être reliées à des théories ordinaires, mais avec un groupe de jauge différent. Les comprendre est une autre manière de comprendre, par exemple, la dynamique d'un paquet de branes. En quelque sorte, plusieurs branes empilées, c'est pareil qu'une certaine brane, toute seule, mais non-commutative.

Je citerai une deuxième raison, peut-être personnelle, mais qui a le mérite d'être plus simple. Comme je l'ai dit, la non-commutativité provient du champ de fond $B_{\mu\nu}$, qui est la partie antisymétrique d'une certaine matrice, dont l'autre partie, symétrique, est bien connue et a de très profondes conséquences : c'est la métrique $g_{\mu\nu}$. Quoi de plus naturel alors que d'étudier sa petite sœur.

Remerciements

Merci à tous les participants d'avoir accepté un théoricien dans un monde composé majoritairement de physiciens des particules. C'est bien sûr un plaisir de remercier les organisateurs et les coordinateurs des différentes sessions. J'adresse enfin mes sincères remerciements aux serveurs et aux cuisiniers du Village Club Khelus de La Hume.

Références

- [1] Si vous voulez en savoir plus sur la théorie des cordes, vous pourrez lire le tout début de l'article suivant, que vous trouverez sur www.arXiv.org : C. V. Johnson, *D-Brane Primer*, [hep-th/0007170].
- [2] N. Seiberg and E. Witten, *String theory and non-commutative geometry*, JHEP **9909** (1999) 032, [hep-th/9908142].
- [3] S. Fidanza, *Towards an explicit expression of the Seiberg-Witten map at all orders*, JHEP **0206** (2002) 016, [hep-th/0112027].

Bibliographie

- [1] C. V. Johnson, *D-brane primer*, [hep-th/0007170].
- [2] J. H. Schwarz, *Introduction to superstring theory*, [hep-ex/0008017].
- [3] E. Kiritsis, *Introduction to superstring theory*, [hep-th/9709062].
- [4] B. Greene, *The elegant universe*, traduit en français sous le titre *L'univers élégant*.
- [5] J. Baez and J. P. Muniain, *Gauge fields, knots and gravity*, published by World Scientific, Singapore (1994).
- [6] M. Nakahara, *Geometry, Topology And Physics*, published by Bristol, UK (1990).
- [7] L. Cornalba and R. Schiappa, *Nonassociative star product deformations for D-brane worldvolumes in curved backgrounds*, Commun. Math. Phys. **225** (2002) 33, [hep-th/0101219].
- [8] M. R. Douglas and N. Nekrasov, *Noncommutative field theory*, Rev. Mod. Phys. **73** (2001) 977, [hep-th/0106048].
- [9] R. J. Szabo, *Quantum field theory on noncommutative spaces*, Phys. Rept. **378** (2003) 207, [hep-th/0109162].
- [10] N. Seiberg and E. Witten, *String theory and noncommutative geometry*, JHEP **9909** (1999) 032, [hep-th/9908142].
- [11] N. Seiberg, L. Susskind and N. Toumbas, *Strings in background electric field, space/time noncommutativity and a new noncritical string theory*, JHEP **0006** (2000) 021, [hep-th/0005040].
- [12] N. Seiberg, L. Susskind and N. Toumbas, *Space/time non-commutativity and causality*, JHEP **0006** (2000) 044, [hep-th/0005015].
- [13] K. Saraikin, *Comments on the Morita equivalence*, J. Exp. Theor. Phys. **91** (2000) 653, [hep-th/0005138].
- [14] B. Pioline and A. Schwarz, *Morita equivalence and T-Duality (or B versus Θ)*, JHEP **9908** (1999) 021, [hep-th/9908019].
- [15] A. A. Tseytlin, *On non-abelian generalisation of the Born-Infeld action in string theory*, Nucl. Phys. **B501** (1997) 41, [hep-th/9701125].
- [16] P. Bain, *On the non-abelian Born-Infeld action*, [hep-th/9909154].
- [17] L. Cornalba, *On the general structure of the non-abelian Born-Infeld action*, Adv. Theor. Math. Phys. **4** (2002) 1259, [hep-th/0006018].
- [18] L. Cornalba, *Some computations with Seiberg-Witten invariant actions*, Mod. Phys. Lett. **A16** (2001) 227-234, [hep-th/0101015].

- [19] N. Wyllard, *Derivative corrections to the D-brane Born-Infeld action : non-geodesic embeddings and the Seiberg-Witten map*, JHEP **0108** (2001) 027, [hep-th/0107185].
- [20] H. Liu, *$\star$ -Trek II : $\star_n$ operations, open Wilson lines and the Seiberg-Witten map*, Nucl. Phys. **B614** (2001) 305, [hep-th/0011125].
- [21] H. Liu and J. Michelson, *Ramond-Ramond couplings of noncommutative D-branes*, Phys. Lett. **B518** (2001) 143-152, [hep-th/0104139].
- [22] Y. Okawa and H. Ooguri, *An exact solution to Seiberg-Witten equation of noncommutative gauge theory*, Phys. Rev. **D64** (2001) 046009, [hep-th/0104036].
- [23] S. R. Das and S. P. Trivedi, *Supergravity couplings to noncommutative branes, open Wilson lines and generalized star products*, JHEP **0102** (2001) 046, [hep-th/0011131].
- [24] S. Mukhi and N. V. Suryanarayana, *Gauge invariant couplings of noncommutative branes to Ramond-Ramond backgrounds*, JHEP **0105** (2001) 023, [hep-th/0104045].
- [25] S. R. Das, S. Mukhi and N. V. Suryanarayana, *Derivative corrections from noncommutativity*, JHEP **0108** (2001) 039, [hep-th/0104045].
- [26] T. Mehen and M. B. Wise, *Generalized $\star$ -products, Wilson lines and the solution of the Seiberg-Witten equations*, JHEP **0012** (2000) 008, [hep-th/0010204].
- [27] S. S. Pal, *Derivative corrections to Dirac-Born-Infeld and Chern-Simons actions from non-commutativity*, Int. J. Mod. Phys. **A17** (2002) 1253, [hep-th/0108104].
- [28] T. Asakawa and I. Kishimoto, *Comments on gauge equivalence in noncommutative geometry*, JHEP **9911** (1999) 024, [hep-th/9909139].
- [29] J. Madore, S. Schraml, P. Schupp and J. Wess, *Gauge theory on noncommutative spaces*, Eur. Phys. J. **C16** (2000) 161-167, [hep-th/0001203].
- [30] B. Jurčo, P. Schupp and J. Wess, *Nonabelian noncommutative gauge theory via noncommutative extra dimensions*, Nucl. Phys. **B604** (2001) 148-180, [hep-th/0102129].
- [31] B. Jurčo, L. Möller, S. Schraml, P. Schupp and J. Wess, *Construction of non-abelian gauge theories on noncommutative spaces*, Eur. Phys. J. **C21** (2001) 383, [hep-th/0104153].
- [32] P. Schupp, *Non-abelian gauge theory on noncommutative spaces*, [hep-th/0111038].
- [33] S. Goto and H. Hata, *Noncommutative monopole at the second order in θ* , Phys. Rev. **D62** (2000) 085022, [hep-th/0005101].
- [34] D. Brace, B. L. Cerchiai, A. F. Pasqua, U. Varadarajan and B. Zumino, *A cohomological approach to the non-abelian Seiberg-Witten map*, JHEP **0106** (2001) 047, [hep-th/0105192].
- [35] B. L. Cerchiai, A. F. Pasqua and B. Zumino, *The Seiberg-Witten map for noncommutative gauge theories*, [hep-th/0206231].
- [36] S. Fidanza, *Towards an explicit expression of the Seiberg-Witten map at all orders*, JHEP **0206** (2002) 016, [hep-th/0112027].

- [37] M. Kontsevich, *Deformation quantization of Poisson manifolds, I*, [q-alg/9709040].
- [38] D. Arnal, D. Manchon et M. Masmoudi, *Choix des signes pour la formalité de M. Kontsevich*, [math.QA/0003003].
- [39] A. S. Cattaneo and G. Felder, *A path integral approach to the Kontsevich quantization formula*, Commun. Math. Phys. **212** (2000) 591, [math.QA/9902090].
- [40] R. Jackiw, S. Y. Pi and A. P. Polychronakos, *Noncommuting gauge fields as a Lagrange fluid*, Annals Phys. **301** (2002) 157, [hep-th/0206014];
B. Bistrovic, R. Jackiw, H. Li, V. P. Nair and S. Y. Pi, *Non-Abelian fluid dynamics in Lagrangian formulation*, Phys. Rev. **D67** (2003) 025013, [hep-th/0210143].
- [41] D. Freed, J. A. Harvey, R. Minasian and G. W. Moore, *Gravitational anomaly cancellation for M-theory fivebranes*, Adv. Theor. Math. Phys. **2** (1998) 601, [hep-th/9803205].
- [42] J. A. Harvey, R. Minasian and G. W. Moore, *Non-abelian tensor-multiplet anomalies*, JHEP **9809** (1998) 004, [hep-th/9808060].
- [43] L. Alvarez-Gaume and E. Witten, *Gravitational Anomalies*, Nucl. Phys. **B234** (1984) 269.
- [44] J. Strathdee, *Extended Poincaré Supersymmetry*, Int. J. Mod. Phys. **A2** (1987) 273.
- [45] B. Keller, *Introduction to A-infinity algebras and modules*, [math.RA/9910179].
- [46] C. Hofman and W. K. Ma, *Deformations of closed strings and topological open membranes*, JHEP **0106** (2001) 033, [hep-th/0102201].
- [47] C. Hofman, *Nonabelian 2-forms*, [hep-th/0207017].
- [48] J. C. Baez, *Higher Yang-Mills theory*, [hep-th/0206130].
- [49] L. Breen and W. Messing, *Differential geometry of gerbes*, [math.AG/0106083].
- [50] J. H. Schwarz, *Lectures on superstring and M theory dualities*, Nucl. Phys. Proc. Suppl. **55B** (1997) 1, [hep-th/9607201].
- [51] M. Krogh and S. M. Lee, *String network from M-theory*, Nucl. Phys. **B516** (1998) 241, [hep-th/9712050].
- [52] O. Bergman, *Three-pronged strings and 1/4 BPS states in $N = 4$ super-Yang-Mills theory*, Nucl. Phys. **B525** (1998) 104, [hep-th/9712211].
- [53] M. R. Gaberdiel, T. Hauer and B. Zwiebach, *Open string-string junction transitions*, Nucl. Phys. **B525** (1998) 117, [hep-th/9801205].
- [54] M. R. Gaberdiel and B. Zwiebach, *Exceptional groups from open strings*, Nucl. Phys. **B518** (1998) 151, [hep-th/9709013].
- [55] O. DeWolfe and B. Zwiebach, *String junctions for arbitrary Lie algebra representations*, Nucl. Phys. **B541** (1999) 509, [hep-th/9804210].
- [56] J. C. Baez, *The Octonions*, [math.RA/0105155].
- [57] A. Strominger, S. T. Yau and E. Zaslow, *Mirror symmetry is T-duality*, Nucl. Phys. **B479** (1996) 243, [hep-th/9606040].

- [58] D. Joyce, *Lectures on Calabi-Yau and special Lagrangian geometry*, [math.DG/0108088].
- [59] B. R. Greene, *String theory on Calabi-Yau manifolds*, [hep-th/9702155].
- [60] N. Hitchin, *Lectures on special lagrangian submanifolds*, [math.DG/9907034].
- [61] J. T. Liu and R. Minasian, *U-branes and T^3 fibrations*, Nucl. Phys. **B510** (1998) 538, [hep-th/9707125].
- [62] S. Gurrieri, J. Louis, A. Micu and D. Waldram, *Mirror symmetry in generalized Calabi-Yau compactifications*, Nucl. Phys. **B654** (2003) 61, [hep-th/0211102].
- [63] C. Vafa, *Superstrings and topological strings at large N* , J. Math. Phys. **42** (2001) 2798, [hep-th/0008142].
- [64] N. Hitchin, *Generalized Calabi-Yau manifolds*, [math.DG/0209099].
- [65] A. Strominger, *Superstrings with torsion*, Nucl. Phys. **B274** (1986) 253.
- [66] J. P. Gauntlett, D. Martelli, S. Pakis and D. Waldram, *G-structures and wrapped NS5-branes*, [hep-th/0205050].
- [67] G. L. Cardoso, G. Curio, G. Dall'Agata, D. Lust, P. Manousselis and G. Zoupanos, *Non-Kähler string backgrounds and their five torsion classes*, Nucl. Phys. **B652** (2003) 5, [hep-th/0211118].
- [68] P. Kaste, R. Minasian and A. Tomasiello, *Supersymmetric M-theory compactifications with fluxes on seven-manifolds and G-structures*, JHEP **0307** (2003) 004, [hep-th/0303127].
- [69] S. Chiossi and S. Salamon, *The intrinsic torsion of $SU(3)$ and G_2 structures*, Proc. Conf. Differential Geometry Valencia 2001, [math.DG/0202282].
- [70] P. Kaste, R. Minasian, M. Petrini and A. Tomasiello, *Nontrivial RR two-form field strength and $SU(3)$ -structure*, Fortsch. Phys. **51** (2003) 764, [hep-th/0301063].
- [71] S. Fidanza, R. Minasian and A. Tomasiello, *Mirror symmetric $SU(3)$ -structure manifolds with NS fluxes*, [hep-th/0311122].

