



Prélude à l'analyse des données du détecteur Virgo: De l'étalonnage à la recherche de coalescences binaires

Fabrice Beauville

► To cite this version:

Fabrice Beauville. Prélude à l'analyse des données du détecteur Virgo: De l'étalonnage à la recherche de coalescences binaires. Astrophysique [astro-ph]. Université de Savoie, 2005. Français. NNT : . tel-00010297

HAL Id: tel-00010297

<https://theses.hal.science/tel-00010297>

Submitted on 26 Sep 2005

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.

Laboratoire d'Annecy-le-Vieux de Physique des particules

Thèse

présentée à

l'Université de Savoie

pour obtenir le titre de

Docteur en sciences,

Spécialité: Physique des particules

par

Fabrice BEAUVILLE

Prélude à l'analyse des données

du détecteur Virgo :

De l'étalonnage à la recherche de coalescences binaires

Soutenue le 18 juillet 2005 devant le jury composé de

A.	BARRAU	Rapporteur
J.	COLAS	Président du jury
F.	MARION	Directeur de thèse
F.	MONTANET	
B.	MOURS	Directeur de thèse
A.	VICERE	Rapporteur

Table des matières

Introduction	7
1 Les ondes gravitationnelles	9
1.1 Existence théorique	9
1.1.1 Le principe d'équivalence	9
1.1.2 L'équation d'Einstein	11
1.1.3 Propagation des ondes gravitationnelles	13
1.1.4 Émission et énergie rayonnée	14
1.1.5 Effets sur des masses-tests	16
1.2 Réalité expérimentale	16
1.2.1 Mise en évidence de l'émission	16
1.2.2 Les espoirs de détection	17
1.2.3 Les coalescences binaires	20
1.2.4 Les autres sources	23
1.2.5 Conclusion	24
2 L'expérience Virgo	25
2.1 Les masses-tests	25
2.1.1 Les masses-tests et leurs suspensions	25
2.1.2 Les super-atténuateurs	26
2.1.3 Agir sur les masses-tests: les actionneurs	28
2.2 L'interféromètre	29
2.2.1 Principe	29
2.2.2 Configuration optique	34
2.2.3 Le contrôle global des longueurs	36
2.3 Objectifs et progression	38
2.3.1 Bruits et sensibilité nominale	38
2.3.2 Le <i>commissioning</i> et les <i>Mock Data Challenges</i>	39
3 Étalonnage du détecteur	43
3.1 Introduction: but de l'étalonnage	43
3.2 Les actionneurs comme système d'étalonnage	44
3.2.1 Description du problème	44

3.2.2	Réponse de l'ensemble DAC - bobines	45
3.2.3	Stabilité de l'ensemble au cours du temps	49
3.2.4	Mesure du gain absolu en m/volts	52
3.3	Réponse du détecteur	61
3.3.1	Description	61
3.3.2	Mesure	61
3.3.3	Analyse	62
3.4	Sensibilité	66
4	Reconstruction du signal gravitationnel $h(t)$	69
4.1	Introduction	69
4.2	Méthode utilisant le signal de frange noire	70
4.2.1	Principe	70
4.2.2	Essai sur des données simulées	71
4.3	Méthode utilisant le signal de frange noire et les signaux de contrôles	73
4.3.1	Principe	74
4.3.2	L'algorithme	75
4.3.3	Effet sur le signal	77
4.3.4	Effet de suppression de bruit	78
4.4	Suivi de la réponse du détecteur	83
4.4.1	Suivi des lignes	83
4.4.2	Cas particuliers et vetos	84
4.4.3	Application aux prises de données C4 et C5	85
4.5	Soustraction des harmoniques de 50 Hz	88
4.5.1	Principe	88
4.5.2	Test et résultats	91
5	Recherche de coalescences binaires: Méthode	99
5.1	Le signal de coalescences binaires	99
5.1.1	Forme du signal détecté	99
5.1.2	Signal Newtonien	101
5.1.3	Corrections d'ordres supérieures	102
5.2	Détection de la forme d'onde dans les données : méthode standard	104
5.2.1	Filtrage adapté	104
5.2.2	Grilles de filtres	108
5.3	Recherche multi-bandes	110
5.3.1	Principe	110
5.3.2	Implémentation	112
5.3.3	Regroupement de candidats	117
5.3.4	Perspectives	119

6	Validation de la recherche multi-bandes	121
6.1	Recherche avec le <i>template</i> exact	121
6.1.1	Fausses alarmes	121
6.1.2	Détection des événements	125
6.1.3	Résolution en temps et en distance effective	129
6.2	Recherche avec une grille de <i>templates</i>	129
6.2.1	Fausses alarmes	129
6.2.2	Détection des événements	131
6.2.3	Estimation des masses de la source	133
6.2.4	Résolution en temps et en distance effective	136
6.3	Conclusion	138
7	Application aux données actuelles de l'interféromètre	139
7.1	Recherche d'injections	140
7.1.1	Les injections	140
7.1.2	Résultats	141
7.2	Bruit du détecteur: Faux événements	143
7.2.1	<i>Runs</i> complets	143
7.2.2	Période silencieuse	145
7.3	Stabilité du détecteur	147
7.3.1	Distance maximale de détection	147
7.3.2	Effets sur les grilles de <i>templates</i>	148
7.4	Perspectives: Vetos	155
8	Comparaison avec la chaîne d'analyse de LIGO	159
8.1	Le premier MDC LIGO-Virgo	160
8.1.1	Les deux jeux de données simulés	160
8.1.2	Les deux chaînes d'analyse	161
8.1.3	La production des candidats	162
8.2	Comparaison des résultats	163
8.2.1	Injections détectées	163
8.2.2	Rapport signal sur bruit	163
8.2.3	Distance	165
8.2.4	Temps d'arrivée	165
8.2.5	Masses de la source	165
8.3	Conclusion	172
	Conclusion	173
	Bibliographie	174

Introduction

L'existence des ondes gravitationnelles a été prédite par A.Einstein en 1916, un an après la publication de sa théorie de la relativité générale, qui décrivait la gravitation comme l'effet des déformations de l'espace-temps induites par les masses de l'univers. En mettant en évidence, dans sa théorie, la possibilité pour une perturbation de la métrique de l'espace-temps de se propager à la vitesse de la lumière, Einstein désignait alors un équivalent gravitationnel du rayonnement électromagnétique: les ondes gravitationnelles.

Produites par des mouvements de matière, ces perturbations de la métrique peuvent se manifester par des variations relatives de distance entre un ensemble de masses libres. Elles n'ont pourtant jamais pu être observées directement: leur couplage à la matière est si faible, que seuls des mouvements violents de sources astrophysiques extrêmement massives et compactes, peuvent engendrer des perturbations suffisamment fortes pour envisager une détection par les instruments les plus sensibles. A ce jour, la seule preuve expérimentale de leur existence est donc purement indirecte: l'observation par Taylor et Weisberg des pertes d'énergie orbitale d'un pulsar et de son étoile compagnon (PSR1913+16).

Depuis les années 60, l'espoir d'une observation directe suscite néanmoins la mise au point de détecteurs toujours plus sensibles, encouragé par la meilleure connaissance des sources astrophysiques potentielles. Après les barres résonnantes, qui cherchaient à recueillir l'énergie de l'onde pour la restituer en vibrant à leur propre fréquence de résonance, la détection par interférométrie des variations de distances séparant des masses-tests est envisagée, permettant d'élargir la recherche à de larges plages de fréquence.

Pour pouvoir espérer une détection, la collaboration franco-italienne Virgo, comme LIGO aux états-unis, a construit un détecteur de plusieurs kilomètres. La construction de Virgo s'est achevée en 2003, pour laisser place à la mise en route de l'instrument, toujours en cours à ce jour.

Inscrit dans le contexte de cette mise en route, ce travail de thèse porte sur l'étalonnage et la reconstruction des données de Virgo, ainsi que sur leur analyse pour une recherche d'ondes gravitationnelles émises par la coalescence de systèmes binaires d'astres compacts.

Après deux chapitres introductifs, qui visent à résumer les notions essentielles relatives aux ondes gravitationnelles et au fonctionnement de Virgo, la première

partie de ce travail sera abordée. Elle porte sur la réponse global du détecteur à un signal gravitationnel, en particulier les effets de filtrage optique et ceux du contrôle en temps réel du point de fonctionnement de l'interféromètre.

Le chapitre 3 présente la mise en oeuvre d'une procédure de mesure de cette réponse. Il décrit d'abord la caractérisation du système qui permet l'injection de signaux s'apparentant à des ondes gravitationnelles dans le détecteur. Viens ensuite la mesure et son application à la détermination de la sensibilité du détecteur, indispensable au suivi des progrès tout au long de sa mise en route.

Le chapitre suivant présente la reconstruction des données, c'est à dire la prise en compte de cette réponse au moyen d'un processus de mise en forme des données. Celui-ci doit produire un signal corrigé des déformations instrumentales et nettoyé de certain bruits, directement comparable aux signaux gravitationnels prédits par la relativité générale.

La deuxième partie de ce travail, la recherche de coalescence d'astres compacts, sera abordée ensuite. Ces coalescences sont les sources les plus prometteuses pour Virgo, et leur signal gravitationnel a le mérite rare d'être en grande partie prédictible avec précision. Ce sont aussi les sources qui occupent le plus largement la plage de fréquence observée par Virgo, et par conséquent les plus exigeantes quand à la reconstruction des données.

Le signal attendu est très faible; il faut donc l'extraire du bruit par un filtrage optimal. La méthode standard pour cette extraction sera présentée au cinquième chapitre, suivi d'une description détaillée de la recherche multi-bande, une implémentation alternative motivée par les limitations liées aux moyens de calculs de la méthode habituelle.

Pour résumer le travail de mise au point sur des données simulées de cette méthode alternative, qui s'est déroulé en parallèle avec le début de la mise en route de Virgo, le chapitre 6 présente la validation en simulation de la recherche multi-bande dans sa version actuelle.

On passera ensuite aux applications de la méthode validée, avec la description au chapitre 7 des premières analyses des données actuelles de Virgo, obtenues durant les prises de données régulières effectués au cours de sa mise en route.

Enfin, la préparation de l'analyse de données en collaboration avec l'expérience LIGO sera abordée dans le dernier chapitre à travers la description d'une étude - dans le cadre du groupe d'analyse jointe LIGO-Virgo - de comparaison de la recherche multi-bande et de la chaîne d'analyse pour la recherche de coalescences de LIGO. On prépare ainsi l'échange de données au sein du réseau de détecteurs, pour augmenter les chances de détection et permettre la pleine exploitation astrophysique des signaux détectés.

Chapitre 1

Les ondes gravitationnelles

Prédites par A. Einstein en 1916, les ondes gravitationnelles sont l'une des nouveautés majeures de sa théorie de la relativité générale.

Cette théorie interprète la gravitation comme l'effet de la courbure de l'espace-temps; celui-ci s'apparente donc à un objet "vivant" de la physique, dont la structure est déformée par la présence de matière. Il peut alors jouer le rôle de milieu élastique dans lequel des petits plis se propagent: les ondes gravitationnelles.

Longuement considérées comme des artefacts mathématiques, ces ondes ont progressivement gagné leur statut d'objets physiques, au même titre que les ondes électromagnétiques, grâce aux progrès théoriques réalisés entre la fin des années 60 et le début des années 80. A ces confirmations de leur existence théorique, s'est ajoutée à la même époque une première preuve expérimentale indirecte de leur émission, grâce à près de vingt ans d'observation astrophysique d'un système émetteur.

Toutefois, près de quatre-vingt-dix ans après la première prédiction, les ondes gravitationnelles n'ont jamais pu être détectées, et ce à cause de la faiblesse de leurs effets. Réussir une première détection du rayonnement émis par les événements astrophysiques les plus violents est donc l'objectif d'expériences comme Virgo.

Ce premier chapitre expose brièvement¹ les notions théorique essentielles à la compréhension du principe de détection de Virgo. Il présente ensuite un rapide état des lieux de la réalité expérimentale: la preuve indirecte dors et déjà acquise de leur existence, et les espoirs de détection.

1.1 Existence théorique

1.1.1 Le principe d'équivalence

Après l'avènement de la théorie de la relativité restreinte énoncée par Einstein en 1905, il apparaît indispensable de re-formuler les lois de la gravitation.

1. en se basant sur les cours plus détaillés [1, 2, 3, 4].

La gravitation newtonienne décrit en effet une force découlant d'un potentiel Φ , qui vit dans un espace universel à trois dimensions, et se propage instantanément à distance. Au contraire, la relativité restreinte adopte pour cadre un espace-temps à quatre dimensions: les grandeurs "géométriques" - indépendantes du repère inertiel choisi - ne sont plus la distance et le temps, mais l'intervalle d'espace-temps séparant deux événements:

$$ds^2 = -c^2 dt^2 + dx^2 + dy^2 + dz^2. \quad (1.1)$$

Elle renonce ainsi à l'idée d'un temps universel absolu et donc enlève tout sens au mot "instantanée"; si la gravitation semble se propager instantanément d'un point à un autre dans un repère inertiel donné, dans d'autres repères inertiels elle pourra sembler se propager vers le passé ($dt < 0$) ou le futur ($dt > 0$), sans que l'intervalle ds^2 séparant les deux points ne varie.

Pour résoudre cet incompatibilité, la solution la plus simple consisterait à modifier l'équation du potentiel Newtonien:

$$\nabla^2 \Phi = 4\pi G \rho \quad (1.2)$$

pour la rendre invariante - comme l'est ds^2 - par changement de repère inertiel. Mais dès 1907, Einstein supposait qu'une nouvelle théorie de la gravitation impliquerait une réforme plus profonde de la relativité, basé sur le principe d'équivalence.

Ce principe repose sur l'égalité expérimentale de la masse grave et de la masse inerte, c'est à dire l'universalité de l'accélération des corps en chute libre, connue depuis Galilée. D'après ce principe, un observateur en chute libre - par exemple en orbite dans une station spatiale - ne voit aucune force de gravitation agir sur lui ou les objets qui l'entourent, tous subissant la même accélération. Il constitue donc localement un repère inertiel, dans lequel la gravité est inexistante et les lois non-gravitationnelles - de la relativité restreinte - sont vérifiées. Cette équivalence entre repère inertiel et repère en chute libre est facilement détruite par la non-uniformité du champ de gravitation: deux objets proches subissant une gravité différente afficheront une accélération relative, dévoilant la présence du champ à l'observateur. Aussi les repères en chute libres ne constituent des repères inertiels que localement, sur des régions d'espaces-temps suffisamment petites pour que le champ de gravitation y soit uniforme.

C'est l'idée clé de la relativité générale: abandonner les restrictions aux repères Galiléens des lois de l'ancienne relativité, et généraliser cette dernière à tous les repères en chute libre, qui deviennent des repères *localement inertiels*. On peut donc trouver en tout point de l'espace-temps un repère localement inertiel, dans lequel les lois de la relativité restreinte sont vérifiés localement, notamment l'expression (1.1) de la métrique.

Dans ce contexte, que devient la gravitation? En tant que force elle n'existe plus, puisque elle est imperceptible dans les repères en chute libre. Elle continue

néanmoins de se manifester à travers le caractère strictement local des repères inertiels en chute libre: en l'absence de gravitation, les observateurs en chute libre par exemple autour de la terre pourrait former un réseau qui constituerait un repère inertiel globale. Par "gravitation", on entend donc à présent l'*impossibilité de joindre bout-à-bout les repères inertiels locaux pour former un repère inertiel global*. En géométrie, une telle contrainte est attribuée à une courbure de l'espace considéré; ces espaces courbes généralisent l'exemple de la surface terrestre, sur laquelle des repères Euclidiens peuvent être adoptés de façon locale - quand la surface semble plane - mais ne peuvent être mis bout-à bout pour former des repères Euclidiens décrivent toute la surface courbe. C'est pourquoi la gravitation, responsable du caractère local des repères inertiels, est considéré en relativité générale comme l'effet d'une *courbure de l'espace-temps*.

Une telle courbure se manifeste concrètement par la déformation des géodésiques de l'espace-temps, comme en témoigne les trajectoires des astres en chute libres. De façon plus intuitive, elle se retrouve dans la métrique de l'espace-temps $g_{\mu\nu}$, qui généralise (1.1):

$$ds^2 = g_{\mu\nu}(x)dx^\mu dx^\nu. \quad (1.3)$$

La courbure n'empêche en aucun cas - ce serait contraire au principe d'équivalence - de trouver *localement* un système de coordonnées dans lequel la métrique se ramène à la métrique $\eta_{\mu\nu}$ de la relativité restreinte (1.1). Elle exclue par contre l'existence d'un système dans lequel $g_{\mu\nu}(x) = \eta_{\mu\nu}$ *en tout point* de l'espace-temps. Ce sont donc bien les variations de la métrique d'un point à un autre, pour un système de coordonnées donné, qui mettent en évidence la courbure.

1.1.2 L'équation d'Einstein

En tant que nouvelle forme de la gravitation, la courbure de l'espace-temps doit être engendré par la présence de matière dans l'univers. L'équation d'Einstein décrit cette relation entre la courbure et la matière qui l'engendre:

$$G^{\mu\nu} = \frac{8\pi G}{c^4} T^{\mu\nu} \quad (1.4)$$

Le terme de droite - coté matière - est le tenseur symétrique impulsion-énergie. C'est la généralisation relativiste immédiate de la densité de masse ρ , c'est à dire du terme de droite de l'équation du potentiel Newtonien (eq. 1.2). Il représente la densité de courant d'énergie-impulsion, comme le montre ses composantes dans un systèmes de coordonnées quelconque:

$$\left(\begin{array}{c|c} T^{00} & T^{0i} \\ \hline T^{i0} & T^{ij} \end{array} \right) = \left(\begin{array}{c|c} \text{densité d'énergie} & \text{flux d'énergie} \\ \hline \text{densité d'impulsion} & \text{flux d'impulsion} \end{array} \right) \quad (1.5)$$

avec $i, j = 1, 2, 3$ et $c = 1$.

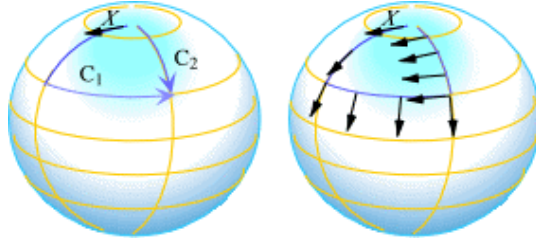


FIG. 1.1 – Illustration des effets sur un vecteur quelconque de son transport parallèle le long d'une courbe fermée, effets décrits localement par le tenseur de Riemann et qui caractérisent la courbure.

Le terme de gauche - coté courbure - est appelé tenseur d'Einstein et dépend de la métrique. Il s'interprète en grande partie par analogie avec la gravitation Newtonienne. La recherche des géodésiques d'un espace-temps courbe à la limite Newtonienne - particules lentes et faibles perturbations statiques à la métrique plate - met en évidence une correspondance directe entre la métrique et le potentiel Newtonien. Dans le cadre de cet analogie, le tenseur des forces Newtonienne de marées $E = \nabla\nabla\Phi$, qui témoigne de la non-uniformité du champ de gravitation, se généralise au tenseur de Riemann $R^\alpha_{\mu\beta\nu}$, défini à partir des dérivées secondes (et premières) de la métrique. Du point de vue géométrique, le tenseur de Riemann décrit localement les effets sur un vecteur quelconque de son transport parallèle le long d'une courbe fermée; ces effets, inexistantes en l'absence de courbure, caractérisent complètement celle-ci.

Le tenseur de Riemann est de rang 4, trop élevé pour intervenir directement dans l'équation d'Einstein. Donc, tout comme l'équation du potentiel Newtonien (eq.1.2) fait intervenir la trace du tenseur E , l'équation d'Einstein doit utiliser le tenseur de courbure de Ricci obtenu par contraction du tenseur de Riemann:

$$R_{\mu\nu} = R^\alpha_{\mu\alpha\nu} \quad (1.6)$$

et de rang 2 comme $T_{\mu\nu}$. En fait, le tenseur d'Einstein est une adaptation du tenseur de Ricci:

$$G^{\mu\nu} \equiv R^{\mu\nu} - \frac{1}{2}g^{\mu\nu}R^\alpha_\alpha \quad (1.7)$$

modifié pour supprimer la divergence non-nulle de ce dernier ($\nabla^\mu R_{\mu\nu} \neq 0$), en contradiction potentielle avec la conservation de l'énergie-impulsion ($\nabla^\mu T_{\mu\nu} = 0$).

Une fois obtenu le tenseur d'Einstein, le facteur de proportionnalité est obtenu par identification avec l'équation (1.2) à la limite Newtonienne. La découverte par Einstein de ce tenseur (1915) - et indépendamment par Hilbert (1915) - marque donc l'achèvement de la théorie de la relativité générale.

1.1.3 Propagation des ondes gravitationnelles

L'équation d'Einstein détermine la structure de l'espace-temps, en présence tout comme en l'absence de matière. Les ondes gravitationnelles, étant des perturbations de cette structure, leur équation de propagation repose donc logiquement sur l'équation d'Einstein. Cette dernière étant non-linéaire vis-à-vis de la métrique, elle doit être linéarisée en considérant les ondes gravitationnelles comme des perturbations $h_{\mu\nu}$ à la métrique de l'espace-temps plat² $\eta_{\mu\nu}$:

$$g_{\mu\nu} = \eta_{\mu\nu} + h_{\mu\nu} \quad , \quad |h_{\mu\nu}| \ll 1 \quad (1.8)$$

Cette décomposition (espace plat + perturbation) de la métrique n'est pas unique, puisqu'elle peut être obtenue dans différents systèmes de coordonnées; la liberté de choix du système de coordonnées est analogue à l'invariance de l'électromagnétisme par changement de jauge $A_\mu \rightarrow A_\mu + \partial_\mu \lambda$. Comme pour les ondes électromagnétiques, un choix de jauge particulier, dit "harmonique", et vérifiant:

$$\square x_\mu = 0 \quad (1.9)$$

donne à l'équation d'Einstein dans le vide ($T_{\mu\nu} = 0$) et linéarisée - ramenée au premier ordre en $h_{\mu\nu}$ - la forme standard d'une équation de propagation relativiste:

$$\square \tilde{h}_{\mu\nu} = 0. \quad (1.10)$$

Les solutions d'une cette équation sont alors des ondes se propageant à la vitesse de la lumière:

$$h_{\mu\nu} = h_{\mu\nu}(ct - z) \quad (1.11)$$

(si l'axe z coïncide avec l'axe de propagation).

Sur les dix composantes a priori indépendantes du tenseur symétrique $h_{\mu\nu}$, les quatre conditions de jauge (1.9) n'en laisse que six. Mais ces conditions ne contraignent toujours pas complètement le système de coordonnées. Il nous reste en effet la liberté de choix liée aux changements de la forme:

$$x_\mu \rightarrow x_\mu + \epsilon_\mu \quad , \quad \square \epsilon_\mu = 0. \quad (1.12)$$

qui n'affectent pas les conditions (1.9). Le choix³ des quatre ϵ_μ ramène donc le nombre de composantes indépendantes du tenseur $h_{\mu\nu}$ de six à deux.

En choisissant ainsi un système de coordonnées adéquat appelée jauge TT (*Transverse Traceless*), on peut annuler sa trace et supprimer ses composantes

2. ou, plus généralement, dont le rayon de courbure et l'échelle de longueur à laquelle la courbure varie sont grands devant la longueur d'onde du rayonnement.

3. en choisissant par exemple $\epsilon_\mu = \varepsilon_\mu \exp(ik^\alpha x_\alpha)$, k^α est imposée par (1.12) et il reste à choisir les quatre ε_μ .

suivant la direction de propagation:

$$h^{TT} = \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & h_+ & h_\times & 0 \\ 0 & h_\times & -h_+ & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} \quad (1.13)$$

et le décomposer suivant deux polarisations indépendantes (+) et (×):

$$(h_+ e_+ + h_\times e_\times) e^{i(kz - \omega t)} \quad (1.14)$$

où e_+ et $e_\times$ sont les deux directions de polarisation, formant un angle de 45° .

$$e_+ = \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} \text{ et } e_\times = \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} \quad (1.15)$$

En résumé, la Relativité Générale prédit l'existence d'ondes gravitationnelles, se propageant à la vitesse de la lumière, et admettant deux états de polarisation transverse.

1.1.4 Émission et énergie rayonnée

En présence de matière, l'équation d'onde (1.10) s'écrit:

$$\square h_{\mu\nu} = -\frac{16\pi G}{c^4} S_{\mu\nu} \quad (1.16)$$

où $S_{\mu\nu} = T_{\mu\nu} - \frac{1}{2} g_{\mu\nu} g_{\alpha\beta} T^{\alpha\beta}$ représente la matière source d'ondes gravitationnelles.

Pour une source de petite dimension devant les longueurs d'ondes du rayonnement émis, la métrique admet la solution de type "potentiel retardé":

$$h_{\mu\nu}(t, \vec{x}) = \frac{4G}{c^4} \int_{source} \frac{S_{\mu\nu}(t - \frac{|\vec{x} - \vec{x}'|}{c}, \vec{x}')}{|\vec{x} - \vec{x}'|} d^3x' \quad (1.17)$$

En supposant de plus la source petite par rapport à la distance r de l'observateur, cette expression peut être développée en termes multipolaires. Le premier terme non-nul est le terme quadrupolaire:

$$h_{ij}(t) = \frac{2G}{r c^4} \ddot{I}_{ij}(t - \frac{r}{c}) \quad (1.18)$$

qui décrit une émission proportionnel à la dérivée seconde du quadrupôle - projeté dans la jauge TT - de la densité d'énergie:

$$I_{ij}(t) = \frac{1}{c^2} \int_{source} T^{00}(t, \vec{x}) (x_i x_j - \frac{1}{3} \delta_{ij} |\vec{x}|^2) d^3x. \quad (1.19)$$

Il n'y a donc pas, contrairement au cas de l'électromagnétisme, de rayonnement dipolaire. Ce qui se justifie physiquement: les variations au cours du temps du dipôle de la densité d'énergie représente le déplacement du centre de gravité de la source; par conservation de l'impulsion, ce déplacement ne peut être accéléré. Le quadrupôle, lui, décrit la forme interne - les asymétries - de la source, et son accélération n'est contrainte par les lois de conservation.

Autre implication importante de la formule de l'émission (1.18), le rayonnement est en $\frac{1}{r}$. On peut le comprendre intuitivement par analogie avec les ondes électromagnétique: l'amplitude serait en $\frac{1}{r}$ pour que la puissance rayonnée soit en $\frac{1}{r^2}$ et conserve ainsi l'énergie émise sur 4π .

Mais associer une énergie aux ondes gravitationnelles s'avère extrêmement délicat⁴, ne serait-ce que parce qu'on l'attend *au second ordre* en $h_{\mu\nu}$, alors même que l'existence de ces ondes est établie dans le cadre d'une version *linéarisée* de la Relativité Générale. Depuis les années 60 de nombreuses approches ont été proposés, qui sont en général équivalente en terme de prédictions d'énergie rayonné. On peut par exemple [?] développer l'équation d'Einstein au second ordre pour définir un pseudo-tenseur énergie-impulsion:

$$T_{\mu\nu}^{GW} = -\frac{c^4}{8\pi G} G_{\mu\nu}^{(2)} \quad (1.20)$$

où $G_{\mu\nu}^{(2)}$ désigne le terme d'ordre 2 du tenseur d'Einstein. Dans un repère inertiel, le flux d'énergie associé à la propagation de l'onde est alors donné par:

$$T_{0z}^{GW} = \frac{c^2}{16\pi G} \langle \dot{h}_+^2 + \dot{h}_\times^2 \rangle \quad (1.21)$$

où la moyenne $\langle \dots \rangle$ se fait sur de nombreux cycles, cette énergie n'ayant localement aucun sens puisque associée à une perturbation de la métrique. De retour au formalisme quadrupolaire (eq. 1.18), la puissance rayonnée sous forme d'ondes gravitationnelles par une source s'écrit alors:

$$\frac{dE}{dt} = \frac{G}{c^5} \sum_{j,k} \langle \ddot{I}_{jk}^2 \rangle. \quad (1.22)$$

Pour l'estimer grossièrement on considère une source de masse M , de dimension R , de vitesse de rotation Ω , et de facteur d'asymétrie a tel que (1.22) puisse s'écrire approximativement:

$$\frac{dE}{dt} \sim \frac{G}{c^5} (aMR^2\Omega^3)^2, \quad (1.23)$$

4. Cette difficulté a longtemps constitué la principale objection à la réalité physique des ondes gravitationnelles.

que l'on re-exprime:

$$\frac{dE}{dt} \sim \frac{c^5}{G} a^2 \left(\frac{v}{c}\right)^6 \left(\frac{R_s}{R}\right)^2 \quad (1.24)$$

en fonction de la vitesse caractéristique des éléments de la source $v \sim R\Omega$, et de son rayon de Schwarzschild $R_s = 2\frac{GM}{c^2}$.

La luminosité d'une source augmente donc avec ses vitesses, son asymétrie, et sa compacité. C'est pourquoi les sources les plus puissantes sont constituées d'objets astrophysiques compacts comme il n'en existe pas sur terre: trous noirs, étoiles à neutrons ...

1.1.5 Effets sur des masses-tests

De même que la détection ou réception des ondes électromagnétiques passe par le mouvement des "charges-tests" d'une antenne, les ondes gravitationnelles se manifestent à travers leurs effets sur des masses-tests.

On considère des masses-tests en chute libre, séparées par des distances l_i faibles devant la longueur d'onde du rayonnement. Au passage d'une onde gravitationnelle polarisée (+), la métrique dans la jauge TT prend la forme:

$$ds^2 = -c^2 dt^2 + (1 + h_+) dx^2 + (1 - h_+) dy^2 + dz^2 \quad (1.25)$$

Les distances propres mesurées physiquement entre deux de ces masses-tests sont alors modifiées par l'onde:

$$l_x \rightarrow l_x + dl_x = l_x \sqrt{1 + h_+} = l_x \left(1 + \frac{1}{2} h_+\right) \text{ selon la direction } x \quad (1.26)$$

$$l_y \rightarrow l_y + dl_y = l_y \sqrt{1 - h_+} = l_y \left(1 - \frac{1}{2} h_+\right) \text{ selon la direction } y \quad (1.27)$$

$$(1.28)$$

Dans le cas d'une onde ($\times$), les effets sont identiques après rotation de 45° . Les effets du passage d'une onde sur un anneau de masses-tests situé dans un plan orthogonal à la direction de propagation de l'onde sont illustrés sur la figure 1.2 pour les deux polarisations.

C'est sur cet effet que se basent les tentatives de détection.

1.2 Réalité expérimentale

1.2.1 Mise en évidence de l'émission

Les ondes gravitationnelles constituent également une réalité expérimentale, et ce depuis la mise en évidence de leur émission par le pulsar PSR1913+16.

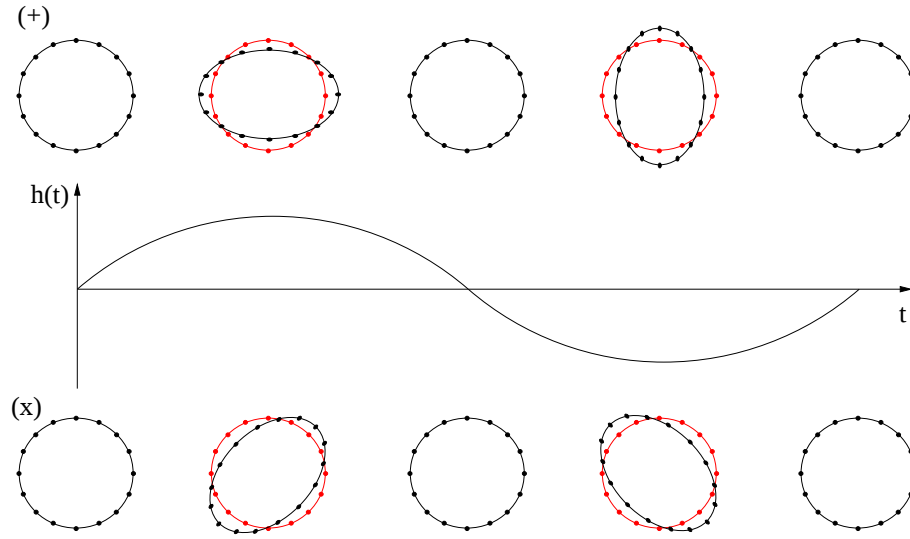


FIG. 1.2 – Déformation d'un anneau de masses libres sous les effets d'une onde gravitationnelle de polarisations (+) ou ($\times$).

Découvert en 1974 à 7 kpc de la terre par Joseph Taylor et Russell Hulse [5], ce pulsar de $1.44 M_{\odot}$ vit en orbite avec un compagnon de $1.39 M_{\odot}$ révélé par la modulation Doppler qu'il fait subir au signal radio de son grand frère. La période de révolution des deux objets - environ huit heures - est mesurée régulièrement depuis 1975.

La perte d'énergie par rayonnement gravitationnel de ce système binaire peut être ainsi observée à travers la période orbitale de ce système, de plus en plus courte à mesure que les deux étoiles se rapprochent. Les mesures de cette baisse depuis 1975 - de près de 15 s de période orbitale perdues en trois décennies - coïncident à 0.4% avec les prédictions de la relativité générale (figure 1.3).

Cette première preuve expérimentale de l'existence des ondes gravitationnelles, confirmant au passage la validité - en tant qu'approximation - de la formule du quadrupôle (1.22), valut le prix Nobel de physique à R.A. Hulse et J.H. Taylor en 1993.

1.2.2 Les espoirs de détection

Après la mise en évidence de l'émission, l'étape suivante est la détection. Une "antenne" réceptive aux ondes gravitationnelles permettrait bien sûr de confirmer leur existence, mais aussi leurs propriétés: propagation à la vitesse de la lumière, polarisations h_+ et $h_{\times}$. Ces propriétés sont des pistes pour une quantification de la gravitation, puisqu'elles suggèrent un boson de jauge de masse nulle et d'hélicité 2. Fierz et Pauli [7] ont en effet montré, en étudiant les théories de champs de spin quelconque qui sont quantifiables de façon canonique, que celle d'un champ de

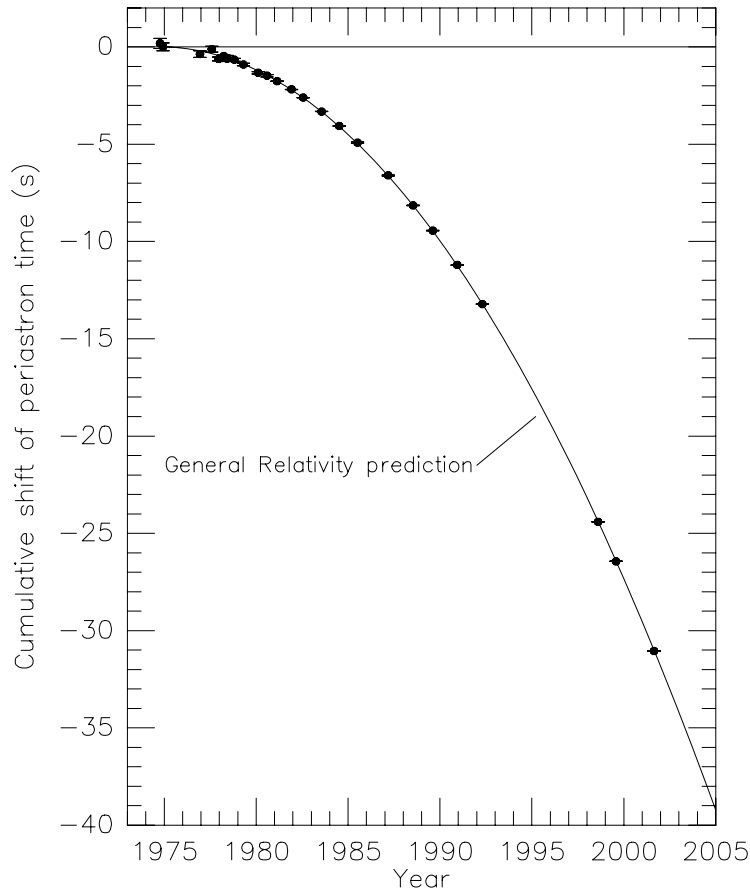


FIG. 1.3 – *Diminution de la période orbitale du pulsar PSR1913+16 au cours du temps, et comparaison avec la prédiction de la relativité générale*[6].

spin 2 est identique à la relativité générale sans ses effets non-linéaires, et que les ondes planes de cette théorie ont ces propriétés des ondes décrites jusqu'à présent.

A cet intérêt pour la nature de ces ondes s'ajoute l'intérêt pour leurs sources; la détection des ondes gravitationnelles devient alors un nouveau moyen d'observation des objets célestes. Grâce à la faiblesse de leur couplage à la matière, l'image gravitationnelle de l'univers que transportent ces ondes est à peine affectée par ce qu'elles rencontrent sur leur parcours. De plus cette image, contrairement à celle que nous offrent les ondes électromagnétiques, est une image des masses de l'univers. Il est de plus en plus clair depuis la deuxième moitié du XX^e siècle que la lumière du ciel n'apporte pas l'information suffisante à la compréhension du contenu de l'univers. La détection des ondes gravitationnelles serait donc un pas en avant dans la recherche à long terme de nouvelles fenêtres d'observation.

Les détecteurs résonnants

Les premières tentatives de détection datent des années 60. Elles reposaient alors sur l'utilisation de barres métalliques qui doivent révéler les ondes gravitationnelles les traversant par une excitation de leur résonance mécanique. Il s'agit donc de détecteur à bande passante resserrée autour de la fréquence de résonance.

Ces prototypes étaient limités par leur bruit thermique; aujourd'hui des détecteurs fondés sur le même principe utilisent les techniques cryogéniques pour lutter contre ce bruit: ALLEGRO aux États-Unis, AURIGA et NAUTILUS en Italie, EXPLORER au CERN, et NIOBE en Australie.

Leurs fréquences de résonance - et donc de détection - se situent autour du kHz.

Les barres sont qualifiées de "directionnelles", au sens où la détection n'est possible que si l'un des axes de polarisation de l'onde coïncide avec l'axe de la barre. Des détecteurs en projet remplacent la barre par une sphère, permettant une détection omni-directionnelle: MiniGRAIL, TIGA, SFERA, GRAVITON, ou encore OMEGA.

Les détecteurs interférométriques

Les détecteurs interférométriques, dont fait partie Virgo, visent une meilleure sensibilité ainsi qu'un élargissement de la bande passante. Le principe général est de mesurer par interférométrie les variations de distances entre des masses-tests en chute libre, ces variations étant liées à celle de la métrique par les équations (1.26). On s'affranchit ainsi du problème de la fréquence de résonance des barres.

La caractéristique frappante de ces détecteurs est leur grande taille, indispensable pour offrir des chances de détection; Virgo [8], qui sera décrit plus précisément au chapitre suivant, possède des bras 3 km de long. Sur terre, on rencontre plusieurs détecteurs du même type que Virgo en développement ou en opération:

- LIGO (Laser Interferometer Gravitational-Wave Observatory) [9] est constitué de 3 détecteurs dont deux (2 km et 4 km) sont situés à Hanford, le troisième (4 km) étant lui situé à Livingston. Ce sont avec Virgo les seuls détecteurs kilométriques à l'heure actuelle.
- GEO600 [10] (600 m), issue d'une collaboration germano-britannique, construit près de Hambourg.
- TAMA300 [11] (300 m), construit à Tokyo, a quasiment atteint sa sensibilité de référence. Il doit servir de prototype à un détecteur kilométrique souterrain et cryogénique.
- AIGO, situé à proximité de Perth en Australie, est le seul détecteur interférométrique de l'hémisphère Sud. Son éloignement des autres détecteurs existant devrait permettre d'augmenter la précision quand à la localisation des sources détectées. Il se limite pour l'instant à des bras de 80 m mais pourrait être étendu à 4 km.

Ces détecteurs sont suspendus pour libérer les masses de la contrainte terrestre. Leur sensibilité au mouvement des masses-tests engendré par le bruit sismique leur interdit toute détection au dessous de quelques Hz ou plus, suivant les efforts mis en oeuvre pour abaisser ce "mur" sismique.

Pour explorer les plus basses fréquences, le projet de détecteur interférométrique spatial en orbite solaire LISA (Laser Interferometer Space Antenna) est lancé conjointement par l'ESA et la NASA; le lancement est prévu pour 2013. Les bras de LISA - en triangle équilatéral et non en L - font 5 millions de km, pour une bande passante pouvant aller de 10^{-5} Hz à 1 Hz.

Il reste alors à savoir quelles sont les sources qui pourront être observée par ces détecteurs, et tout particulièrement par Virgo. Si le principe théorique de l'émission des ondes gravitationnelles est compris et confirmé par la mise en évidence de l'émission du pulsar binaire PSR1913+16 (§1.2.1), il est délicat de prédire les sources détectable, et ce pour deux raisons:

- L'émission étant - au minimum - quadrupolaire, elle dépend de l'asymétrie de la source, qui, en dehors du cas simple des systèmes binaires, peut être complètement inconnue.
- L'amplitude détectée est en $\frac{1}{r}$ (eq. 1.18); même si l'amplitude du rayonnement d'un système donné est connu, il reste à déterminer sa probabilité d'être présent dans la fraction d'univers vue par le détecteur.

C'est pourquoi les principaux candidats à la détection qui vont être évoqué à présent doivent être simplement considérés comme les plus prometteurs, ceux vers lesquels l'analyse des données des détecteurs se tourne en priorité.

1.2.3 Les coalescences binaires

Deux objets compacts (étoiles à neutrons ou trous noirs) en orbite l'un autour de l'autre - comme PSR1913+16 - perdent de l'énergie orbitale par émission d'ondes gravitationnelles. Ils voient donc leur distance relative diminuer et leur fréquence orbitale augmenter. Parallèlement, les ondes émises voient leur amplitude et leur fréquence - double de la fréquence orbitale - augmenter.

Ces systèmes binaires peuvent être détectés par LISA lorsqu'ils émettent dans sa bande de fréquence; leurs signaux sont alors des quasi-sinusoïdes dont la fréquence et l'amplitude augmentent lentement.

A l'approche de la coalescence - la mort du système binaire - ils entrent dans la bande de détection des détecteurs terrestres ($>\sim$ Hz). Les amplitudes et fréquences émises augmentent alors de plus en plus rapidement, les deux corps spirant l'un vers l'autre jusqu'à la fusion. Celle-ci survient dans les quelques secondes - ou quelques dizaines de secondes, suivant les masses - qui suivent.

Le signal entendu par les détecteurs terrestre - le *chirp* (figure 1.4)- durant la phase spirante a la particularité d'être prédictible, comme on le verra plus précisément au cinquième chapitre. Il peut donc être plus facilement extrait du

bruit des détecteurs qu'un signal inconnu; c'est en recherchant de tels signaux que l'on espère détecter des coalescences binaires. Son autre particularité est de couvrir une très large portion de la bande passante des détecteurs terrestres (figure 1.5)

Taux d'événements attendus

Le taux de tels événements dans la région audible de l'univers est en revanche mal déterminé. Il y a essentiellement deux façons distinctes de l'estimer.

La première est basée sur les observations de systèmes existants; elle ne s'applique donc pas aux systèmes formés d'un trou noir et d'une étoile à neutron ou de deux trous noirs, de tels systèmes n'ayant jamais été observés. Grâce aux recherches de pulsars, six systèmes binaires d'étoiles à neutrons sont par contre connus dans la galaxie. Deux d'entre eux ne peuvent contribuer aux prédictions, car le temps qui leur reste à vivre avant la coalescence est plus grand que le temps d'Hubble. Un troisième est exclu par son appartenance à l'amas globulaire M15; la densité spatiale de tels amas rend négligeable leur contribution au taux de coalescences total [15]. Il reste donc trois systèmes, dont le pulsar de Hulse et Taylor PSR1913+16, et J0737-3039 découvert en 2003 [16].

En considérant ces trois systèmes comme représentatifs de la population galactique, on peut estimer, sachant que ces trois là ont pu être vus, le nombre de systèmes identiques existant probablement dans la galaxie mais non détectés. Dans la méthode la plus récente, décrite dans [17], des populations galactiques de systèmes binaires de pulsars similaire au trois détectés sont générés, suivant différents modèles de distribution spatiale et de luminosité; les effets de sélection inhérent à l'ensemble des recherches de pulsars recensées sont ensuite simulés sur ces populations, et une analyse Bayésienne permet de remonter à la densité de probabilité du nombre de systèmes présent dans la galaxie.

En combinant ces résultats avec les temps de vie estimés des trois pulsars, on obtient cette même densité de probabilité pour le taux de coalescences galactique.

L'extrapolation de ce taux à la région audible par LIGO ou Virgo dépend alors la distribution spatiale des galaxies, estimée sur la base de la luminosité dans la bande B [19], et du volume de cette région, facile à estimer puisque le signal est connu.

Après inclusion du pulsar J0737-3039, le taux de coalescences pour LIGO - ou Virgo - a augmenté d'un facteur 6-7 par rapport aux précédentes estimations. Il est à présent donné, dans un intervalle de confiance à 95%, à $75^{+200}_{-6} \times 10^{-3}$ détections par an [18] pour les modèles standards de distribution spatiale et de luminosité. La large barre d'erreur témoigne du manque de statistique et n'inclue par les erreurs possibles sur les modèles: le résultat dépend principalement du modèle de luminosité, pouvant aller pour le modèle le plus (ou le moins) favorable, jusqu'à $188^{+495}_{-151} \times 10^{-3}$ (ou $4^{+11}_{-3} \times 10^{-3}$) détections par an.

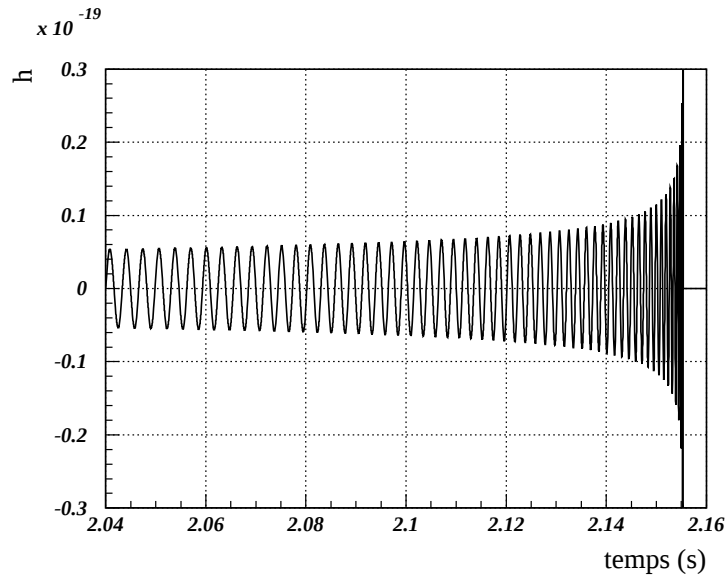


FIG. 1.4 – *Fin du signal théoriquement émis lors de la phase spiralante d'une coalescence binaire (couples d'étoiles de masses $m_1 = m_2 = 1.4 M_\odot$ à une distance $R = 10 \text{ Mpc}$ du détecteur).*

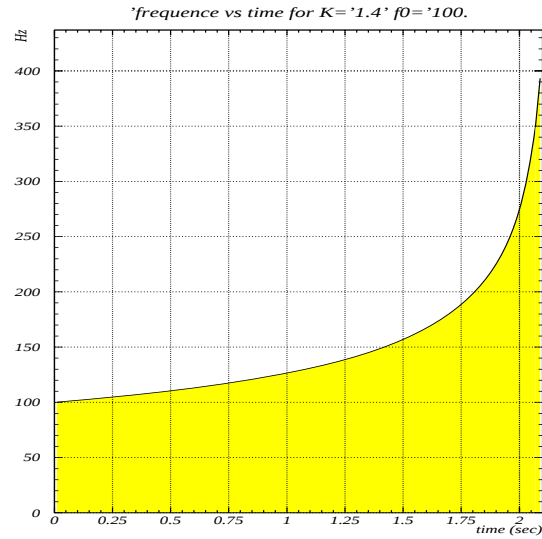


FIG. 1.5 – *Évolution de la fréquence du signal - à partir de 100 Hz - au cours du temps.*

La deuxième méthode, consiste à modéliser la formation de systèmes binaires. Des étoiles binaires sont générés, leurs paramètres obéissant à des distribution soit issues d’observations astrophysiques soit purement théoriques. Leur évolution est ensuite prédites par des méthodes numériques, incluant la transformation successives des deux astres en étoiles à neutrons ou en trou noir. Cette méthode - ou ce type de méthode - est plus diversifiée et peut se décliner suivant de nombreux scénarios d’évolution; on en trouve une revue dans [20]. Elle est aussi très dépendante des paramètres du modèles, en particulier de la vitesse donnée à l’étoile à neutron ou le trou noir naissant par la supernova qui l’engendre - la *kick velocity*; cette vitesse doit être suffisamment faible pour ne pas séparer le système binaire, mais elle peut quand le système binaire survit à l’explosion jouer un rôle positif en rapprochant le système de la coalescence.

Comme dans la méthode précédente, les prédictions sont d’abord des taux galactiques de coalescences, extrapolé ensuite à des taux de détection. Les prédictions de ces modèles sont alors systématiquement favorable aux trous noirs qui sont détectés plus loin. On trouve par exemple dans [21] un taux attendu de coalescences de trous noirs détectables de 8×10^{-1} par an pour un modèle standard, mais variant de 0 à 2 sur une gamme de modèle. Malgré toutes ces incertitudes, ces prédictions font donc des coalescences de trous noirs les candidats les plus prometteurs.

1.2.4 Les autres sources

Les supernovae

L’effondrement d’une étoile de masse suffisante pour former une étoile à neutrons ou un trou noir peut produire des ondes gravitationnelles s’il se produit de façon asymétrique.

En tant qu’objets astrophysiques, ces supernovae sont bien connues, parce qu’elles sont avant tout de fortes sources lumineuses. Contrairement aux systèmes binaires, l’estimation du taux d’événements attendus dans une région donnée de l’espace ne pose donc pas de problème: environ [14] trois par siècle dans la Galaxie, ou quelques uns par an si l’on atteint l’amas de la vierge (15 Mpc).

A cette distance, il n’est ni garanti ni exclu que les signaux émis soit audibles. En effet, ces phénomènes complexes et violents ne peuvent être modélisées aussi simplement que les systèmes binaires; leur asymétrie et donc leur émission d’ondes gravitationnelles sont par conséquent mal connues. On s’attend d’après les modèles existant à des signaux haute fréquence, de l’ordre du kHz, et ne durant pas plus de quelques millisecondes. Leur amplitude varie de plusieurs ordres de grandeurs d’un modèle à l’autre, si bien qu’il est à nouveau impossible de garantir une détection par les détecteurs de première génération.

Les pulsars

Les pulsars sont des étoiles à neutrons en rotation, connus pour émettre des ondes radio (*pulses*) qui permettent leur détection. Un pulsar asymétrique émet une onde gravitationnelle périodique à la fréquence double de sa fréquence de rotation, qui peut être contenue d'après des pulsars connus entre quelques Hz et 1 kHz. Plusieurs mécanismes sont proposés pour justifier la possibilité d'une telle asymétrie, comme par exemple [13]:

- Les déviations de l'axisymétrie acquises par la croûte du pulsar lors de sa cristallisation.
- Un désalignement entre l'axe de rotation du pulsar et son champ magnétique, la pression magnétique engendrant alors une asymétrie de l'étoile.
- Un désalignement entre l'axe de rotation du pulsar et l'axe principale de son moment d'inertie, la précession résultante du pulsar entraînant l'émission.
- Pour un pulsar binaire, son éventuelle déformation causée par l'accrétion de matière provenant du compagnon.

Si cette asymétrie et donc l'amplitude de ces émissions est presque complètement imprévisible, des limites supérieures sont obtenus en supposant la perte d'énergie des pulsars connus entièrement due au rayonnement gravitationnel. Ces limites montrent qu'une détection n'est envisageable qu'en faisant usage de la périodicité du signal, en l'intégrant sur des longues durées (~ 1 an).

1.2.5 Conclusion

Il ressort de cette revue des candidats les plus prometteurs qu'aucune détection n'est garantie (ni exclue) pour les détecteurs terrestres de première génération. On peut néanmoins remarquer que comme le signal est en $\frac{1}{r}$, le volume accessible à un détecteur- et en première approximation le nombre d'événement - augmente d'un facteur 1000 quand sa sensibilité est amélioré d'un seul ordre de grandeur.

Les détecteurs comme Virgo, qui sera présenté au chapitre suivant, peuvent espérer être les outils de la première détection directe d'onde gravitationnelle de l'histoire; ils sont de toute façon assurés de servir de fondation aux détecteurs avancés qui suivront, pour lesquels une détection sera fortement probable.

Chapitre 2

L'expérience Virgo

La mise en oeuvre d'un détecteur terrestre d'ondes gravitationnelles est l'objectif de la collaboration franco-italienne Virgo [8]. Construit à proximité de la ville de Pise, ce détecteur est entré dans sa phase de mise en route (*commissioning*) depuis l'automne 2003.

Virgo est un détecteur interférométrique terrestre, et repose donc sur:

- la réalisation de masse-tests "en chute libre" ou se comportant comme telles, dont les distances relatives révèlent par leurs variations le passage d'une onde gravitationnelle.
- la formation d'un interféromètre avec ces masses-tests - qui sont alors des miroirs - permettant précisément de mesurer ces variations de distances.

Afin de décrire le fonctionnement de base de Virgo, ces deux aspects sont abordés dans les deux premiers sous-chapitres. Les objectifs de Virgo en termes de sensibilité aux ondes gravitationnelles sont décrits ensuite. Enfin, la dernière partie résume la progression récente de Virgo en direction de la prise de données astrophysiques et de leur analyse.

2.1 Les masses-tests

2.1.1 Les masses-tests et leurs suspensions

La détection d'ondes gravitationnelles par Virgo repose sur leur interaction avec des masses-tests supposées libres; la distance entre deux masses-tests éloignées d'une distance L varie (§1.1.5) de:

$$\delta L = \frac{1}{2}hL \quad (2.1)$$

pour une onde h dont l'un des axes de polarisation coïncide avec l'axe $\vec{L}$ défini par les deux masses.

En pratique travailler avec des masses en chute libre reste délicat pour un détecteur terrestre, et on travaille avec des masses suspendues. Le premier rôle des suspensions est donc de compenser la pesanteur terrestre tout en laissant les masses-tests se comporter comme des masses libres dans le plan horizontal.

Pour ce faire, la suspension d'une masse-test doit se comporter comme un ensemble d'oscillateurs harmoniques dont la fréquence de résonance f_0 selon l'axe $\vec{L}$ est bien inférieure aux fréquences des ondes gravitationnelles que l'on cherche à détecter. A ces fréquences, la réponse mécanique (en N/m) du pendule à une force donnée est en $\frac{f_0^2}{f^2}$, tout comme le serait la réponse d'une masse libre.

Au delà de la "libération" des masses-tests se pose le problème de la propagation des vibrations sismiques du sol à travers les suspensions, se confondant avec les ondes gravitationnelles. Or, ce bruit sismique est, à 10 Hz, environ dix ordres de grandeurs au-dessus de la sensibilité espérée. Les suspensions doivent donc aussi servir de filtre sismique.

2.1.2 Les super-atténuateurs

Dans le cas d'un oscillateur simple, la fonction de transfert entre les vibrations du point d'attache et celle transmise à la masse-test est donné au dessus de la résonance f_0 par:

$$G(f) \simeq \frac{f_0^2}{f^2} \quad (2.2)$$

Pour un pendule de longueur réalisable, une telle atténuation n'est pas suffisante pour regagner dix ordres de grandeurs à 10 Hz.

Les suspensions utilisées dans Virgo - les *super-atténuateurs* (figure 2.1) - visent un isolement bien plus efficace en cumulant les effets de plusieurs pendules montés en cascade, filtrant successivement les vibrations sismiques.

Cette chaîne de filtres commence avec le pendule inversé qui contient l'ensemble de la suspension dont la hauteur est proche de 6 m. La fréquence propre de ce "ressort" est ramenée à quelques dizaines de mHz par la pesanteur, qui s'oppose à sa raideur.

Les filtres suspendus qui suivent sont des cylindres d'acier d'une centaine de kg qui atténuent:

- les mouvements horizontaux, par effet pendule ($f_0 \simeq 0.5Hz$);
- les mouvements verticaux¹, grâce aux lames qui relient le cylindre au fil de suspension, jouant ainsi le rôle de ressort vertical dont la fréquence propre est réduite à $f_0 \simeq 1.5Hz$ par des anti-ressorts magnétiques.

1. Ces mouvements sont couplés aux distances entre masses-tests et donc à la détection d'ondes gravitationnelles par la courbure de la terre et les imperfections de la suspensions (découplage imparfait des degrés de libertés).

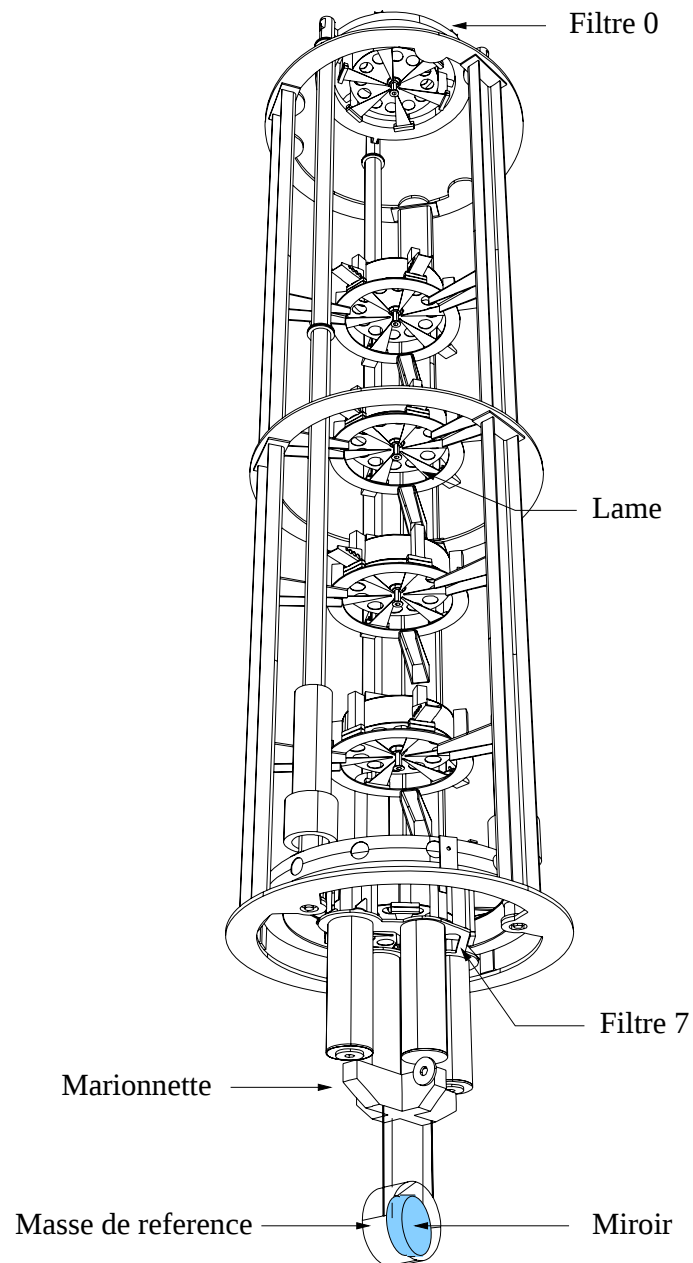


FIG. 2.1 – Les super-atténuateurs de Virgo.

Cette succession de filtres permet de ramener le bruit sismique sous les autres bruits du détecteur, aux fréquences dépassant quelques Hz.

Enfin, suspendu au dernier étage de la suspension - le "filtre 7" - on trouve la marionnette, à laquelle la masse-test (miroir) est elle-même suspendue, ainsi qu'une masse de référence qui permet comme on va le voir d'agir simplement sur la masse-test.

2.1.3 Agir sur les masses-tests: les actionneurs

Les super-atténuateurs sont donc essentiellement des suspensions passives. Il est néanmoins nécessaire d'agir sur les masses-tests et leur suspensions, pour limiter les mouvements de grandes amplitudes aux fréquences de résonance des différents modes de la suspension (à basse fréquence) et assurer l'alignement de l'interféromètre ainsi que les conditions d'interférence optimales.

Les super-atténuateurs sont contrôlés en trois points:

1. Au niveau du "filtre 0", on trouve un amortissement actif agissant bien en dessous des fréquences de détection. Des accéléromètres - qui lui valent le qualificatif d'amortissement inertiel - lui permettent de réduire les variations de vitesse du pendule inversé au dessus d'environ 30 mHz, sans utiliser de repère au sol donc sans ré-injection du bruit sismique. A plus basse fréquence, l'amortissement n'est plus inertiel: il utilise des mesures directes de mouvement ("LVDT"), relatives à un repère externe. Le filtre 0 est ainsi plus fortement ancré au sol, éliminant les mouvements résiduels lents.
2. Quatre bobines fixé sur le "filtre 7" agissent sur des aimants placé sur la marionnette et autorise un contrôle longitudinal - selon l'axe du faisceau - ou angulaire de cette dernière. Les bobines sont pilotées en tension à partir de signaux issues de DSPs , convertis de numérique en analogique par des DACs situés, comme les DSPs, dans les châssis électroniques des suspensions.
3. Quatre bobines similaires sont placées sur la masse de référence, qui comme le miroir, est suspendue à la marionnette. Elles agissent sur des aimants collés aux miroirs. La masse de référence a une masse comparable à celle du miroir, l'action de ces bobines permet donc de déplacer le miroir sans déplacer le centre de gravité de l'ensemble miroir - masse de référence, et donc sans risquer d'exciter le reste de la suspension.

Le contrôle longitudinal - suivant l'axe formé par les masses-tests - a ici un statut particulier, puisqu'il est directement couplé à la détection des ondes gravitationnelles.

Ceci implique d'une part qu'il doit donc être pris en compte à l'analyse des données, ce qui constituera un aspect important des chapitre 3 et 4.

D'autre part, l'amplification des signaux pilotant les bobines - pour obtenir les forces voulues - est alors limitée par le risque de polluer le mouvement longitudinal

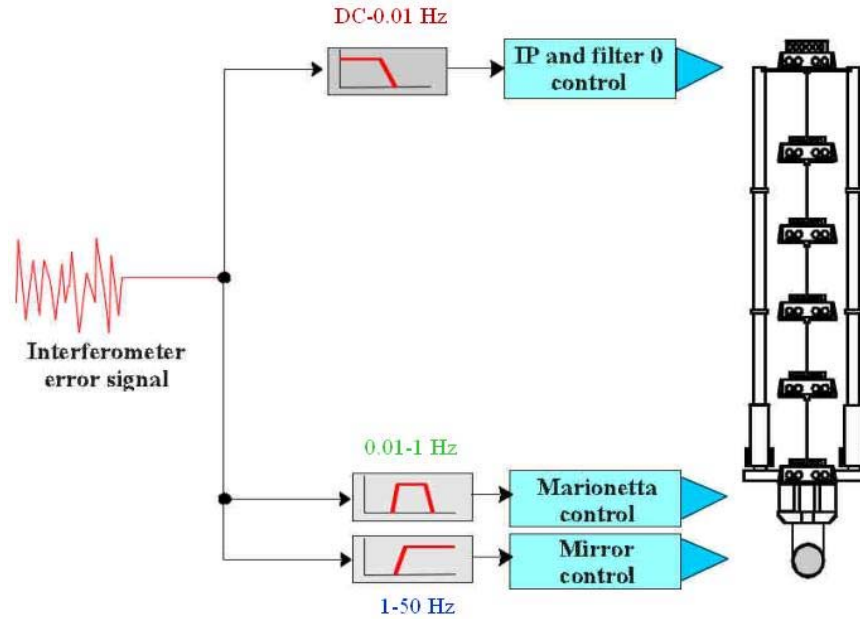


FIG. 2.2 – Répartition du contrôle longitudinal sur l'ensemble de la suspension.

des miroirs par du bruit de l'électronique de pilotage des bobines. Pour obtenir malgré tout ces forces, deux solutions complémentaires sont utilisées:

- L'électronique de pilotage des bobines de la masse de référence est dédoublée: une première configuration bruyante mais puissante produit les forces importantes permettant d'amener Virgo à son point de fonctionnement, relayée ensuite par la configuration à faible bruit et faibles forces qui permet seulement de l'y maintenir.
- Une fois le point de fonctionnement atteint, les composantes basses fréquences - les plus fortes en amplitude - de ce contrôles sont re-dirigées vers les deux autres points d'action de la suspension (figure 2.2).

D'une façon plus générale, la réduction des bruits introduits par le contrôle des masses-tests dans la sensibilité de Virgo constitue l'un des éléments essentiels de sa mise en route.

2.2 L'interféromètre

2.2.1 Principe

Les miroirs se comportant à présent comme des masses-tests libres et isolées du bruit sismique terrestre, on doit détecter les variations des distances qui les

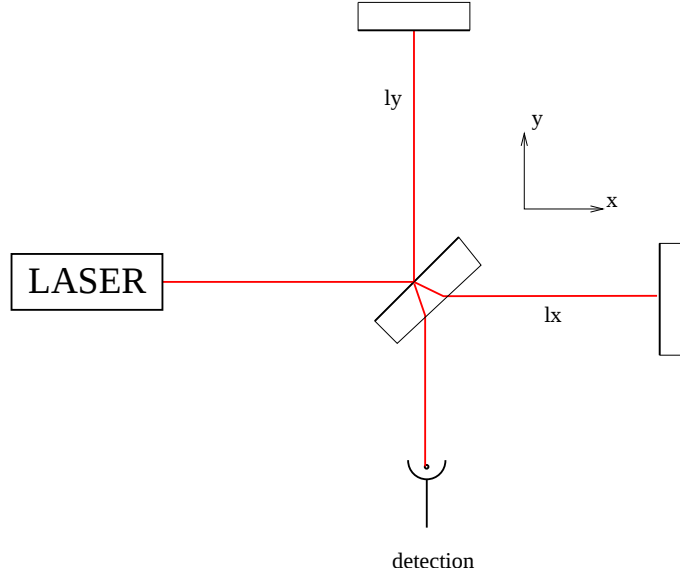


FIG. 2.3 – *Interféromètre de Michelson*

séparent. Le principe de cette détection dans Virgo - décrit en détails dans [22] - est celui d'un interféromètre de Michelson "amélioré".

L'interféromètre de Michelson

Dans cet instrument (figure 2.3), le faisceau lumineux émis par une source cohérente (laser) de puissance P_0 et de longueur d'onde λ est divisé par la lame séparatrice en deux faisceaux de puissances égales se propageant respectivement suivant les direction perpendiculaires Ox et Oy ; après réflexion sur les miroirs de bout de bras, les deux faisceaux sont recombinaés sur la séparatrice. La puissance P alors transmise par l'interféromètre vers la détection dépend du déphasage mutuel des deux faisceaux ϕ :

$$P = \frac{P_0}{2}(1 - C \cos \phi) \quad (2.3)$$

où C désigne le contraste de l'interféromètre, qui mesure la qualité de l'interférence: il vaut 1 si l'interféromètre est parfaitement symétrique (les deux faisceaux recombinaés sont identiques à l'exception du déphasage), et diminue en présence d'asymétries (différence de réflectivité global des miroirs de bout de bras, ou plus certainement déformation différente de la géométrie du faisceau).

Un interféromètre de Michelson de longueur de bras $l_x \simeq l_y \simeq L$, initialement réglé sur un déphasage ϕ_0 , voit le passage d'une onde gravitationnelle h comme une variation de la différence de chemin optique des deux parcours et donc comme

un déphasage² des deux faisceaux:

$$\delta\phi = \frac{4\pi}{\lambda}(\delta l_x - \delta l_y) = \frac{4\pi}{\lambda}hL \quad (2.4)$$

qui entraîne une variation de la puissance détectée:

$$\delta P_h = \frac{dP}{d\phi}\delta\phi = -CP_0 \sin(\phi_0) \frac{4\pi}{\lambda}hL \quad (2.5)$$

permettant finalement la détection de l'onde.

Bruit de photons et frange noire

A ce principe de détection interférométrique est associé un bruit de mesure fondamental lié à la nature quantique de la lumière. Mesurer la puissance transmise P revient en effet à compter le nombre n de photons incidents (pendant une durée δT):

$$P\delta T = n\hbar\omega \quad (2.6)$$

Le *bruit de photon* s'écrit alors:

$$\Delta n = \sqrt{n} \quad (2.7)$$

$$\Delta P_{sn} = \frac{\sqrt{n}}{\delta T}\hbar\omega = \sqrt{\frac{P\hbar\omega}{\delta T}} \quad (2.8)$$

L'onde h n'est détectée que si la variation de puissance qu'elle provoque dépasse les fluctuations de puissance en son absence, soit d'après les équations 2.3 et 2.5:

$$CP_0 \sin(\phi_0) \frac{4\pi}{\lambda}hL > \sqrt{\frac{\frac{P_0}{2}(1 - C \cos \phi_0)\hbar\omega}{\delta T}} \quad (2.9)$$

La plus petite onde h détectable dépend donc du déphasage initial ou *point de fonctionnement* de l'interféromètre ϕ_0 , la meilleure valeur étant celle qui minimise le rapport (figure 2.4):

$$\frac{\sqrt{1 - C \cos \phi_0}}{|\sin \phi_0|}$$

Pour un interféromètre sans défaut de contraste ($C = 1$), la meilleure valeur est $\phi_0 = \pi \pmod{2\pi}$. La puissance transmise est alors nulle (eq. 2.3), c'est pourquoi on l'appelle la frange noire. Lorsque le défaut de contraste ($1 - C$) apparaît, le bruit de photons non-nul même à la "frange noire" se fait sentir et le point de fonctionnement optimal s'en éloigne légèrement, pour mieux amplifier le signal.

2. A priori l'interféromètre ne dépasse pas quelques km, on considère donc h constant sur la durée de l'aller-retour ($\lambda_{og} \gg L$)

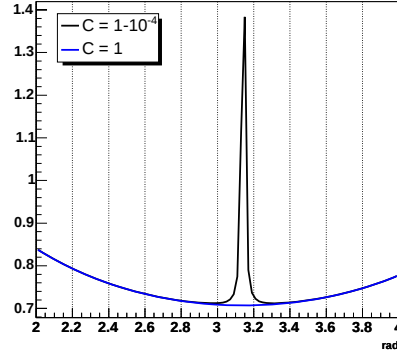


FIG. 2.4 – Illustration du choix du point de fonctionnement: rapport $\frac{\sqrt{1-C \cos \phi_0}}{|\sin \phi_0|}$ en fonction de ϕ_0 pour un défaut de contraste nul ou de 10^{-4} .

Modulation frontale

En pratique ce n'est pas la puissance transmise en détection P qui sert de signal de mesure, et ce pour plusieurs raisons:

- L'équation 2.5 montre que ce signal est toujours positif, et répond symétriquement à une élongation d'un bras ou de l'autre. Un signal impair et signé est plus commode d'utilisation, notamment pour amener et maintenir l'interféromètre à la frange noire.
- Le défaut de contraste éloigne le point de fonctionnement optimal de la frange noire et introduit dans la puissance transmise une composante continue proportionnel à la puissance du laser; or les fluctuations de cette puissance sont importantes dans la bande de fréquence de détection de Virgo.

La technique de modulation frontale employée dans Virgo consiste à moduler en phase le faisceau du laser à l'aide d'une cellule de Pockels. Cette opération fait apparaître, en plus de l'onde à la pulsation ω_0 du laser qu'on appelle la porteuse, des bandes latérales à $\omega_0 \pm n\Omega$ ($2\pi\Omega$ étant la fréquence de modulation). La puissance transmise contient alors un terme à la fréquence de modulation, issue du battement entre les premières bandes latérales et la porteuse.

On choisit des bras asymétriques assurant la transmission des bandes latérales lorsque la porteuse est réglée sur la frange noire. Le signal à la fréquence de modulation n'apparaît alors plus, sauf quand la porteuse est légèrement transmise (s'écarte de la frange noire) à cause d'un déphasage $\delta\phi$. En démodulant la puissance transmise à la fréquence $2\pi\Omega$, on obtient donc un signal proportionnel au déphasage:

$$P_\Omega = P_0 C J_1(2m) \delta\phi \quad (2.10)$$

où J_1 est la fonction de Bessel d'ordre 1 et m la profondeur de modulation.

L'utilisation du signal démodulé comme "signal de frange noire" répond aux problèmes évoqués plus haut:

- Le signal démodulé répond linéairement au déphasage et change de signe lorsque l'interféromètre traverse la frange noire.
- La mesure est effectuée à haute fréquence (démodulation à $2\pi\Omega$); au dessus du MHz le bruit en puissance du laser est plus bas³ que le bruit de photons.

Cavités

La sensibilité de Virgo peut être améliorée en augmentant la réponse de l'interféromètre au déplacement d'une masse-test. On réduit ainsi comparativement tout les bruits autres que des bruits de déplacement, tels que le bruit de photons ou le bruit de l'électronique de détection.

Une première amélioration consiste à remplacer les bras de l'interféromètre par des cavités Fabry-Perot (FP), chacune formée du miroir de bout de bras du Michelson et d'un miroir d'entrée semi-transparent.

Le point de fonctionnement optimal d'une cavité FP correspond intuitivement au cas où la lumière y fait de nombreux allers-retours, c'est à dire lorsque la fréquence du laser résonne avec la longueur L de la cavité ($L = q\frac{\lambda}{2}$).

La cavité se comporte alors comme un miroir, dont le déphasage à la réflexion dépend du petit écart δL qui éloigne la cavité de cette condition de résonance:

$$\phi_{FP} = \pi + \frac{2\mathcal{F}}{\pi} \frac{4\pi}{\lambda} \delta L \quad (2.11)$$

où $\mathcal{F}$ est la finesse de la cavité.

Avec des cavités occupant pratiquement toute la longueur des bras, le déphasage mesuré par le Michelson après recombinaison des deux faisceaux réfléchis par chaque cavité augmente donc d'un facteur $\frac{2\mathcal{F}}{\pi}$ par rapport au Michelson simple.

Cependant, ce facteur diminue pour des ondes gravitationnelles $h(t)$ qui varient à des échelles de temps comparables au temps moyen de stockage de la lumière dans la cavité $\tau = \frac{2L}{c} \frac{\mathcal{F}}{\pi}$, l'effet de l'onde étant intégré sur cette durée. Cette atténuation s'assimile dans la bande passante de Virgo à la réponse d'un filtre passe-bas du premier ordre à la fréquence de coupure f_c :

$$f_c = \frac{1}{2\pi\tau} = \frac{c}{4\mathcal{F}L} \quad (2.12)$$

L'autre moyen utilisé dans Virgo pour améliorer la sensibilité est d'augmenter la puissance incidente P_0 , à laquelle la réponse de l'interféromètre est proportionnelle (eq. 2.10). En augmentant cette puissance d'un facteur R , le rapport du

3. Le bruit en puissance du laser est toujours introduit par la présence de P_0 dans la relation 2.10. Mais cette re-introduction est réduite par un bon contrôle de la condition de frange noire $\delta\phi \simeq 0$.

signal au bruit s'améliore d'un facteur $\sqrt{R}$ pour le bruit de photons qui est en $\sqrt{P_0}$.

L'interféromètre étant réglé sur la frange noire, il réfléchit l'essentiel de cette puissance vers le laser, se comportant ainsi comme un miroir. L'idée est alors d'utiliser un miroir semi-transparent situé entre le laser et l'interféromètre pour ré-injecter cette puissance vers la séparatrice. On forme ainsi une cavité entre ce miroir de recyclage et l'interféromètre, dans laquelle le laser doit être résonnant (la lumière ré-injectée doit être en phase avec la lumière injectée).

Le gain de recyclage R dépend des propriétés de réflexion et transmission du miroir de recyclage; il est cependant limité par les pertes dans l'interféromètre. Ainsi pour des pertes de 2%, il peut atteindre une valeur optimale d'environ 50.

Virgo utilise donc un Michelson sur sa frange noire, amélioré de deux cavités Fabry-Perot et d'une cavité de recyclage, toutes trois en résonance avec le laser. Le principe de la modulation frontale y reste le même, pourvu que les bandes latérales soient résonnantes dans la cavité de recyclage - pour être transmise en détection - et anti-résonnantes dans les cavités FP - pour ne pas annuler le déphasage des cavités ϕ_{FP} dans le battement avec la porteuse (signal démodulé).

2.2.2 Configuration optique

Les principaux éléments de l'interféromètre sont représentés sur la figure 2.5.

Injection

Le système d'injection de Virgo utilise un laser NdYAG de 20 W et de longueur d'onde $1.064 \mu m$, dont le faisceau est filtré par une cavité optique triangulaire suspendue (*mode-cleaner*) de longueur $l_{mc} = 145 m$ - imposée par la nécessité de transmettre la à la fois la porteuse et les bandes latérales. Cette cavité de finesse 1000 est maintenue en résonance avec le mode TEM00 du laser pour éliminer les défauts géométriques du faisceau et les bruits de position et d'angle du laser, servant par la même occasion à la pré-stabilisation en fréquence du laser. Le système d'injection comprend également un télescope chargé d'adapter la largeur du faisceau à l'interféromètre, et dont la dernière lentille n'est autre que le miroir de recyclage (PR). Ce système est en partie provisoire, puisqu'un nouveau banc d'injection doit être installé au cours de l'été 2005.

Les miroirs

Les principaux miroirs suspendus de l'interféromètre sont au nombre de six:

- le miroir de recyclage (PR),
- la séparatrice (BS),
- les miroirs d'entrées des cavités FP nord (NI) et ouest (WI),

cavité de recyclage et quasiment anti-résonnante dans les cavités FP, est de 6.27 MHz.

L'asymétrie du petit Michelson ($l_1 - l_2$) choisie pour la transmission des bandes latérales en détection, est alors de 0.8 m.

Détection

A l'image de l'injection, le système de détection comprend successivement:

- Deux télescopes, en amont et en aval du *mode-cleaner* de sortie, ramenant la taille du faisceau à l'échelle des photodiodes (1 mm).
- le *mode-cleaner* de sortie, cavité optique monolithique de 2.5 cm de longueur asservie - par contrôle thermique - à la résonance avec le mode TEM00, pour améliorer le contraste de l'interféromètre.
- un réseau de 16 photodiodes InGaAs; ses photodiode sont linéaires jusqu'à 100 mW, et sont donc en nombre suffisant pour répartir une puissance transmise pouvant atteindre environ 1W.
- l'électronique de détection, incluant l'amplification, la démodulation, la compression (réduction de la composante basse fréquence du signal), le filtrage anti-repliement à 10 kHz et enfin la numérisation des signaux mesurés par des ADC 16 bits.

Les mouvements du *mode-cleaner* par rapport au faisceau transmis ainsi que ses vibrations peuvent générer un signal parasite, le télescope et le *mode-cleaner* sont donc suspendus, tandis que les photodiodes sont situées sur un banc externe.

Le système de détection ne sert pas uniquement à fournir le signal de frange noire, mais aussi des signaux auxiliaires utiles pour le contrôle de l'interféromètre; il y a ailleurs dans l'interféromètre d'autres bancs externes de détection similaires (figure 2.6).

2.2.3 Le contrôle global des longueurs

Le principe de fonctionnement de l'interféromètre prévoit en effet le contrôle de quatre longueurs, puisque la porteuse doit être sur la frange noire et résonner avec les trois cavités. La précision nécessaire, de $\sim 10^{-4}$ à $\sim 10^{-5}\lambda$ suivant les longueurs, n'est pas compatible avec le mouvement naturel des masses suspendues autour de la fréquence propre des suspensions, et un système de contrôle local utilisant pour chaque masse-test des capteurs optiques ne permet qu'une précision d'une fraction de λ . La précision requise pour les différentes longueurs ne peut donc être fournie que par l'interféromètre lui-même (contrôle global).

Si l'on sait agir sur les masses-tests grâce aux actionneurs (§2.1.3), quatre signaux d'erreurs sont nécessaires pour contrôler en permanence les quatre modes, définis de façon standard comme:

- DARM: La différence de longueur des cavités FP ($L1 - L2$), ou "canal

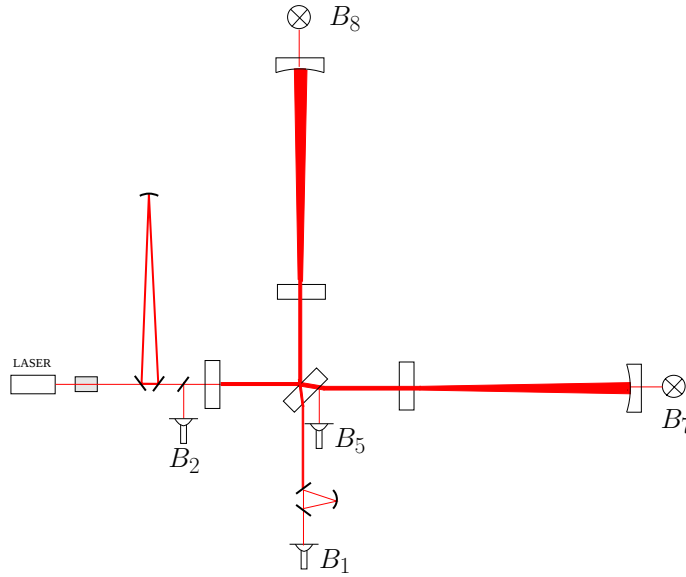


FIG. 2.6 – Positions des photodiodes servant au contrôle de l'interféromètre.

gravitationnel" puisqu'il s'agit du mode sensible aux ondes;

- CARM: La longueur totale des cavités FP ($L1 + L2$), sensible aux bruits communs aux deux bras, comme les fluctuations de fréquence du laser;
- MICH: La frange noire du petit Michelson ($l1 - l2$);
- PRCL: La longueur de la cavité de recyclage ($l_0 + (l_1 + l_2)/2$).

Ces signaux d'erreurs sont à choisir parmi les signaux démodulés - en phase ou en quadrature - des photodiodes B1, B2, B5 (figure 2.6). La relation entre ces différents signaux possibles et les longueurs des 4 modes constitue la matrice optique, dont on peut trouver des prédictions théoriques pour des cas idéaux dans [23] et [24]. On note que:

- Cette matrice permet d'associer un signal de photodiode au mode correspondant mais n'est pas pour autant diagonale, aussi s'attend on à rencontrer des couplages entre boucles de contrôles, comme on le verra aux chapitres 3 et 4.
- le signal de frange noire (B1) n'est surtout sensible qu'au deux modes différentiels DARM et MICH qu'il voit tous deux filtrés par les cavités FP, le premier dominant le deuxième de près d'un facteur 25.

Une fois les 4 modes asservis, un contrôle similaire de la position des miroirs dans leurs degrés de liberté angulaires est engagé, reposant sur l'utilisation de signaux de photodiodes à quadrants, qui reconnaissent les modes d'ordres supérieurs dont la présence renseigne sur le désalignement.

2.3 Objectifs et progression

2.3.1 Bruits et sensibilité nominale

Les bruits limitant la détection des ondes gravitationnelles [25] peuvent venir du déplacement des masses-tests comme du bruit de mesure des distances, autrement dit de l'interféromètre lui-même.

Du côté des masses-tests, on peut citer:

- Le bruit sismique, raison d'être des super-atténuateurs (§2.1.2);
- Le bruit thermique des suspensions ou des miroirs. Pour combattre ce bruit on s'efforce d'obtenir les meilleurs facteurs de qualité possibles pour chaque mode propre d'un miroir ou de sa suspension: par le choix des matériaux, ainsi qu'en plaçant l'ensemble dans le vide pour éviter la friction de l'air.
- L'effet du bruit acoustique sur les masses-tests, lui aussi éliminé par l'isolation de Virgo dans une enceinte à vide.

Les bruits amenés par l'interféromètre sont par exemple:

- Le bruit de photons (§2.2.1), limite fondamentale dont la réduction est à l'origine du schéma optique de l'interféromètre.
- Les fluctuations d'indices de l'air simulant une variation du chemin optique, sont éliminées par le vide des bras de Virgo.

Prenant en compte les différentes sources de bruit et les moyens mis en oeuvre pour les réduire, la courbe de sensibilité théorique de Virgo (figure 2.7) est finalement dominée par trois sources principales de bruits fondamentaux:

- A très basse fréquence ($f < 4Hz$), on trouve la "barrière" sismique excluant toute recherche dans cette région.
- Entre 4 Hz et 500 Hz, la sensibilité est dominée successivement par le bruit thermique des suspensions puis des miroirs.
- Au dessus de 500 Hz, on atteint le bruit de photons, auquel s'ajoutent les pics des modes violons des fils des suspensions. Le bruit de photon augmente alors linéairement avec la fréquence mais reste faible jusqu'à 10 kHz, fréquence où s'interrompt l'acquisition - échantillonnée à 20 kHz - car il n'y a pas de source attendue au-dessus.

Mais cette courbe n'inclut pas les bruits techniques que rencontre Virgo au cours de sa mise en route et qui peuvent a priori être éliminés progressivement par l'optimisation de l'instrument, comme le bruit électronique de détection ou la lumière diffusée sur les photodiodes, qui s'introduisent sur le signal de frange noire ou sur les photodiodes auxiliaires, les effets de la ré-injection de ces bruits dans les boucles du contrôle longitudinal et angulaire, le bruit en fréquence du laser, le bruit de l'électronique de pilotage des actionneurs...

La sensibilité évolue donc au cours de la mise en route de Virgo, la courbe 2.7 constituant l'objectif à atteindre.

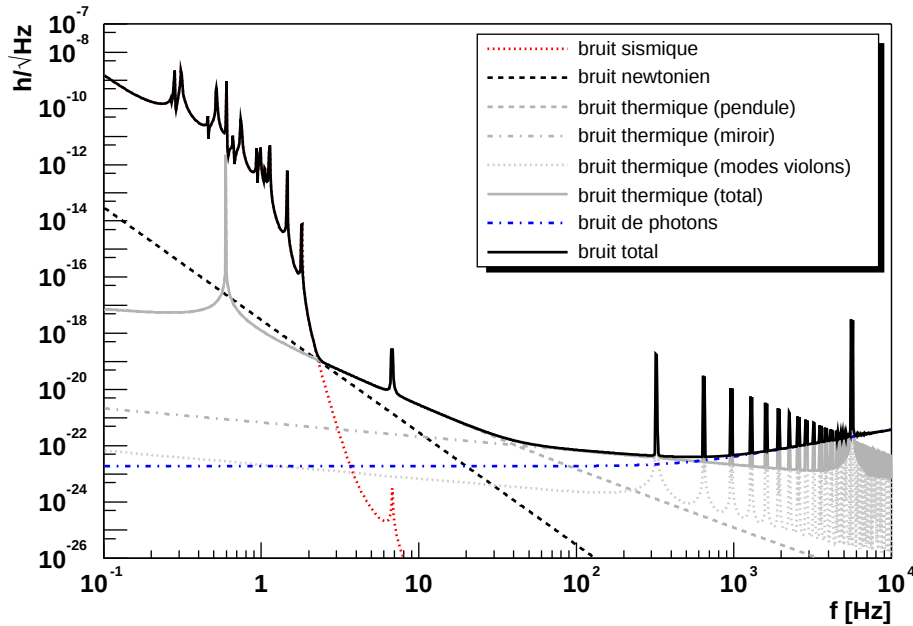


FIG. 2.7 – Contribution des principaux bruits à la sensibilité du détecteur Virgo en fonction de la fréquence.

2.3.2 Le *commissioning* et les *Mock Data Challenges*

Mise en route de l'instrument: les *commissioning runs*

La validation des techniques employées dans Virgo a commencé avec l'étude de l'interféromètre central (CITF), qui s'est terminé en juillet 2002 ([26]). L'année qui a suivi a vu la fin de l'assemblage de Virgo. Sa mise en route (*commissioning*) a commencé en Octobre 2003 avec l'acquisition du contrôle d'une cavité FP. Depuis, la configuration de Virgo ne cesse d'évoluer, pour se rapprocher de la configuration finale. Cette évolution est bien représentée par des prises de données en continu de quelques jours appelés *commissioning runs*, effectuées régulièrement au cours de la progression pour permettre la caractérisation du détecteur à un instant donné, d'effectuer des études de bruit et de préparer l'analyse des données. Cinq *runs* ont eu lieu à ce jour:

- **C1**: 4 jours en novembre 2003, avec la cavité FP du bras nord asservi à la résonance.
- **C2**: 3 jours en février 2004, avec la cavité FP du bras nord asservie à la résonance, et le contrôle angulaire global; 1 jour avec la cavité FP ouest asservi à la résonance.
- **C3**: 4 jours en avril 2004, avec la stabilisation de la fréquence du laser sur

la cavité nord (SSFS⁴); 1 jour avec l'interféromètre "recombiné": les deux cavités asservies indépendamment et la recombinaison des faisceaux sur la séparatrice pour former une frange noire contrôlée avec B1.

- **C4**: 6 jours en juin 2004, avec l'interféromètre recombinaison: contrôle du mode différentiel des cavités (DARM) avec le signal de frange noire, stabilisation de la fréquence de l'injection sur le mode commun des cavités (CARM), contrôle de la frange noire MICH par le signal B2 en réflexion de l'interféromètre, contrôle angulaire global des cavités.
- **C5**: 4 jours en Novembre 2004 dans la configuration C4 avec l'automatisation des opérations de l'interféromètre, contrôle longitudinal hiérarchique complet et pilotage des bobines en mode "bas bruit" (§2.1.3); 2 jours avec l'interféromètre "recyclé": interféromètre recombinaison avec asservissement de la longueur de la cavité recyclage (PRCL), autrement dit l'interféromètre dans sa configuration optique finale.

Les performances obtenues au cours de ces *runs* seront décrites dans les chapitres suivants.

Chaîne d'analyse: les *Mock Data Challenges*

Parallèlement à la mise en route de l'instrument, a commencé la mise en oeuvre de la chaîne de recherche de signaux astrophysiques dans les données, à la fois pour préparer l'analyse des futures données de Virgo et pour permettre une meilleure caractérisation du bruit du détecteur - de sa capacité à simuler un signal astrophysique - au cours de sa progression.

Cette chaîne doit, au moins partiellement, fonctionner "en ligne", c'est à dire intégrée à la suite de l'acquisition de données (DAQ), analysant en temps réel les données directement issues de celle-ci. Elle est constituée d'algorithmes spécifiques correspondant aux différents types de signaux recherchés, précédés d'une partie commune de conditionnement des données, incluant notamment la reconstruction du signal gravitationnel h à partir des données de Virgo.

L'intégration et la validation de la chaîne d'analyse s'est d'abord faite sur des données simulées. Pour jouer un rôle équivalent à celui des *runs*, des *Mock Data Challenges* (MDCs) ont été organisés au sein de la collaboration. Il s'agit d'épreuves d'environ trois jours, durant lesquels un jeu de données, produit par le code de simulation SIESTA [27], doit être analysé par les différentes méthodes existantes, pour retrouver des signaux astrophysiques d'événements impulsifs ou de coalescences binaires injectés lors de la simulation.

Six MDCs ont été organisés dans Virgo entre Mars 2003 et Juin 2004:

- **MDC1** (mars 2003): Ce MDC concernait uniquement les recherches de coalescences binaires, assez lourdes du point de vue informatique. Il s'agissait essentiellement d'un MDC d'intégration "software" dont l'objectif était

4. Second Stage of Frequency Stabilisation

l'analyse complète du jeu de six heures de donnés. Le bruit des données était simulé à l'aide d'un processus ARMA reproduisant la sensibilité du *run* E4 (dernier *run* effectué sur le CITF).

- **MDC2** (juin 2003): Sur des données similaires à celles du MDC1 mais contenant également des injections d'événements impulsifs (*bursts*), le but était de tester l'efficacité des algorithmes en vérifiant la détection des événements injectés.
- **MDC3** (octobre 2003): Il s'agit d'un MDC presque identique au MDC2, mais des non-stationnarités sont introduites dans le bruit sous la forme d'un facteur multiplicatif ou de variation de position de pics spectraux variant au cours du temps.
- **MDC4** (janvier 2004): Pour tester les algorithmes sur des bandes passantes plus larges, on simule cette fois-ci des données reproduisant la sensibilité nominale de Virgo. On introduit également des harmoniques de 50 Hz (fréquence du secteur) et des pics de résonances thermiques pour tester les algorithmes de réduction de ces bruits.
- **MDC5** (mars 2004): Pour tester la chaîne dans son intégralité, on simule entièrement la cavité nord dans sa configuration du *run* C1. On doit ainsi tester la reconstruction de h , avant d'utiliser les algorithmes sur les données reconstruites.
- **MDC6** (Juin 2004): On teste à nouveau la chaîne complète, mais sur une simulation d'interféromètre recombinaison (configurations du *run* C3). De plus la chaîne tourne "en ligne", les données étant lues sur fichier par un processus simulant l'acquisition et fournissant une seconde de données par seconde.

Après ces six MDCs la chaîne a pu être appliquée aux données des *run* C4 et C5, marquant le passage des données simulées aux vraies données.

Par ailleurs un MDC sur données simulées en commun avec l'expérience LIGO s'est déroulé en septembre et octobre 2004 dans le but de comparer les méthodes d'analyses, première étape de la collaboration qui démarre entre les deux expériences. Il fera l'objet du dernier chapitre.

Chapitre 3

Étalonnage du détecteur

3.1 Introduction: but de l'étalonnage

L'étalonnage du détecteur consiste à mesurer sa sensibilité aux ondes gravitationnelles. Cette étape intermédiaire entre la mise en route et l'analyse des données doit répondre à deux questions:

- Que voit-on dans le signal de sortie fourni par le système de détection lorsqu'une onde gravitationnelle atteint Virgo?
- Quelle doit être l'intensité minimale de cette onde pour être vue au dessus du bruit naturel, déjà présent dans le signal de sortie en l'absence de toute onde?

L'interaction d'une onde gravitationnelle h avec Virgo s'interprète comme l'apparition d'une différence de longueur des bras nord et ouest:

$$\Delta L_{OG} = h \cdot 3km \quad (3.1)$$

Cette différence de longueur éloigne l'interférence de la frange noire et génère un signal S à la sortie de l'interféromètre. La relation entre ce signal et le déplacement ΔL_{OG} constitue la *réponse* de l'interféromètre. A proximité de la frange noire elle est linéaire et peut-être représentée par une fonction de transfert R exprimée en *watts/m*:

$$S(f) = R(f) \cdot \Delta L_{OG}(f) \quad (3.2)$$

En l'absence d'onde gravitationnelle ($\Delta L_{OG} = 0$), le signal de sortie S ne contient que le bruit naturel de l'interféromètre, caractérisé par son spectre $\bar{S}_n(f)$ (densité spectrale d'amplitude ou *ASD*, en $\text{watts}/\sqrt{\text{Hz}}$). Le plus petit signal d'onde gravitationnelle ΔL_{OG}^{min} dominant ce bruit est donc celui pour lequel:

$$\bar{S}_n(f) = R(f) \cdot \Delta L_{OG}^{min}(f) \quad (3.3)$$

d'où l'on déduit la *sensibilité* en $\text{m}/\sqrt{\text{Hz}}$ du détecteur:

$$\Delta L_{OG}^{min}(f) = \bar{S}_n(f) \cdot R(f)^{-1} \quad (3.4)$$

et sa sensibilité "en $h/\sqrt{Hz}$ ":

$$h_{OG}^{min}(f) = \frac{\Delta L_{OG}^{min}(f)}{3km} \quad (3.5)$$

où R^{-1} , la réponse inverse, joue donc le rôle de "fonction de calibration" permettant de déduire, du niveau de bruit mesuré en sortie de l'interféromètre, sa sensibilité.

Ce chapitre présente l'ensemble des méthodes étudiées et mises en oeuvre pour obtenir les courbes de sensibilité, avec comme principaux objectifs:

- de s'adapter à l'évolution permanente du détecteur, d'une simple cavité à l'interféromètre recyclé.
- de mesurer la sensibilité avec une précision de l'ordre de 10% à 20%. Cet objectif, très largement suffisant pour décrire l'évolution de la sensibilité au cours de la première phase du commissioning - alors qu'elle progresse par ordres de grandeur - pourra être rehaussé lorsque elle approchera la courbe nominale.

Il se divise en trois parties:

1. L'étude préalable des actionneurs du système de contrôle, afin de les utiliser pour l'injection artificielle de signaux de calibration connus dans l'interféromètre.
2. L'extraction de la fonction de calibration, par l'analyse des résultats de telles injections.
3. Les mesure de sensibilités.

3.2 Les actionneurs comme système d'étalonnage

3.2.1 Description du problème

L'étude de la réponse du détecteur repose sur l'injection de signaux de calibration connus imitant le passage d'une onde gravitationnelle, soit une variation du mode différentiel ΔL . Pour cela on déplace un miroir de bout de bras - le plus souvent du bras nord (NE) - à l'aide des actionneurs du système de contrôle. Un signal de calibration S_{calib} en volts (bruit blanc ou coloré, sinusoïde) est d'abord créé en temps réel par le programme de calibration CaRT puis transmis au DSP de la suspension. Le signal rejoint le signal de contrôle qui peut alors être présent. Il est converti en signal analogique via un DAC, incluant un filtre passe-bas

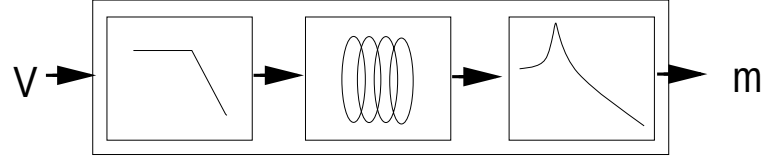


FIG. 3.1 – La réponse des actionneurs: DAC, bobine, réponse mécanique.

supprimant la composante haute-fréquence introduite par la conversion. Le signal analogique est ensuite transmis aux pilotes des 2 bobines (haute et basse) de la masse de référence qui sont utilisées pour le contrôle longitudinal du miroir, et converti par les bobines en force magnétique agissant sur les aimants du miroir. Cette force engendre finalement le déplacement du miroir souhaité ΔL_{calib} .

On ne peut utiliser les actionneurs comme système de calibration que si cette variation ΔL_{calib} est connue. Elle se déduit du signal de calibration S_{calib} envoyé aux actionneurs:

$$\Delta L_{calib}(f) = A(f) \cdot S_{calib}(f) \quad (3.6)$$

où $A(f)$ est la fonction de transfert du système en $m/volts$, et f la fréquence. $A(f)$ représente les différentes déformations subies par le signal de calibration. Elle se décompose en quelques éléments (figure 3.1):

- la réponse de l'ensemble DAC-bobine et électronique associée, soit le passage d'un signal de calibration à l'intensité correspondante du courant circulant dans les bobines (en $A/volts$);
- la conversion du courant en force magnétique, supposée se limiter à un simple gain (en N/A);
- la réponse mécanique $M(f)$ entre force et déplacement du miroir (en N/m).

La forme de la réponse mécanique est déjà connue par des mesures antérieures [28]. Grâce à la masse de référence et à la simplicité du dernier étage de suspensions du miroir, elle se résume à celle d'un pendule résonnant à 600 mHz:

$$M(f) \propto \left[-\left(\frac{f}{600mHz} \right)^2 + \frac{i}{Q} \frac{f}{600mHz} + 1 \right]^{-1} \quad (3.7)$$

dont le facteur de qualité Q est au moins supérieur à 500.

L'étude de la réponse des actionneurs s'est donc concentrée sur les maillons manquants: la réponse DAC-bobine et le gain absolu de la réponse en $m/volts$ ¹.

3.2.2 Réponse de l'ensemble DAC - bobines

L'ensemble électronique constitué par le DAC et les bobines forme la première partie de l'actionneur, et aussi la plus complexe en terme de réponse. Malgré l'iso-

1. dans lequel se fondent le gain de la réponse mécanique 3.7, le gain de conversion en N/A et le gain en $A/volts$ de la partie DAC-bobine et électronique associée.

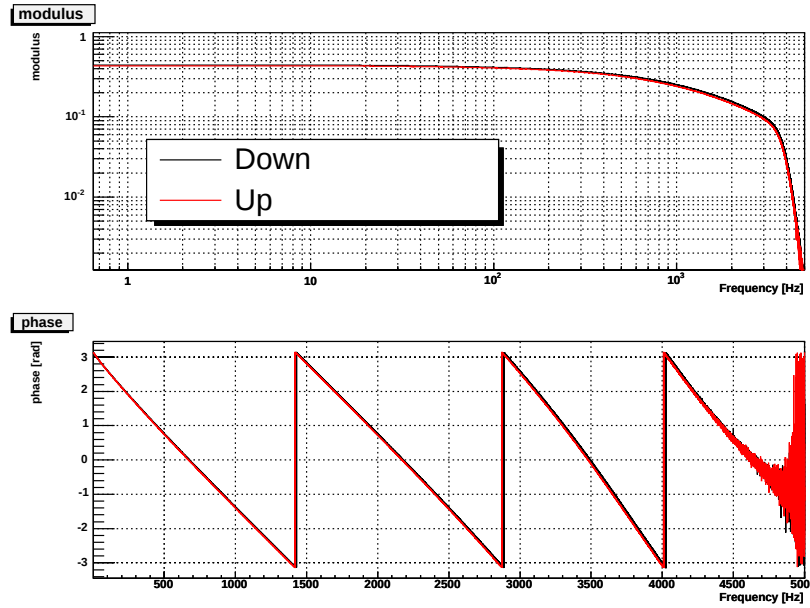


FIG. 3.2 – Fonctions de transfert DAC-bobine-ADC mesurées pour les deux bobines de l'actionneur NE (en haut: module, en bas: phase).

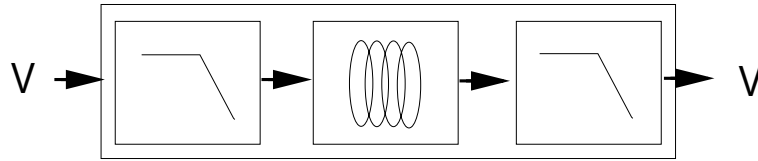


FIG. 3.3 – Réponse de l'ensemble DAC-Bobine + ADC de lecture.

lation des bobines dans l'enceinte à vide, son étude peut se faire indépendamment du reste puisqu'une mesure du courant circulant dans chaque bobine est réalisée en sortie de l'amplificateur qui les pilote, puis enregistrée dans les données. On peut donc, en injectant un bruit blanc de calibration dans les actionneurs, mesurer la fonction de transfert entre ce bruit et le courant ainsi engendré dans les bobines (réponse de l'ensemble DAC - bobine).

La figure 3.2 montre une telle mesure dans les cas superposés des bobines UP et DOWN de la masse de référence du bout de bras Nord (NE); les deux fonctions de transfert sont identiques au premier ordre, l'écart entre les deux ne dépassant pas 2%, sur le module ou sur la phase.

Si la très bonne cohérence (~ 1) des mesures fournit des fonctions de transfert très peu bruitées, celle-ci ne peuvent être utilisées telles quelles en raison des effets de lecture qu'elles incluent; le dispositif de mesure du courant comprend en effet un ADC qui contient un filtre passe-bas, empêchant le repliement du spectre

autour de la fréquence de Nyquist, partiellement responsable de la coupure haute fréquence observée. On doit donc déduire la réponse DAC-bobine de celle mesurée (DAC-bobine-ADC, figure 3.3), ce qui oblige à identifier les différentes contribution présente dans cette dernière. Pour cela on effectue un ajustement de la fonction de transfert par un modèle² incluant:

- Le filtre passe-bas du DAC, dont les pôles et zéros sont donnés par les valeurs nominales des résistance et capacités utilisés, connus à $\pm 5\%$;
- Les bobines, représentées par un simple pôle laissé en paramètre libre, attendu entre quelques centaines de Hz et 1 kHz;
- Le filtre passe-bas (anti-repliement) de l'ADC de lecture, de conception identique à celui du DAC;
- Le retard et le gain DC, laissés en paramètres libres.

L'ajustement est effectué entre 10 Hz et 3 kHz³. Il résulte de la minimisation par MINUIT d'une fonction de χ^2 incluant simultanément les écarts sur le module et la phase. La précision de la mesure étant largement au-dessus de nos exigences pour cette modélisation, on fixe arbitrairement les erreurs de mesure, qui interviennent dans la normalisation du χ^2 , à 0,5% sur toute la bande de fréquence - à comparer aux 10 à 20% attendu sur la calibration.

Les ajustements pour les bobines UP et DOWN se superposent aux fonctions mesurées (figure 3.4), en moyenne à $\sim 1\%$ près soit un χ^2 légèrement inférieur à deux fois le nombre de degrés de liberté; l'ajustement place le pôle de la bobine à 788 Hz pour la UP et à 821 Hz pour la DOWN, soit moins de 4% d'écart. Le modèle peut donc être utilisé pour la calibration, après retrait des contributions associées à la lecture du courant. On conserve alors, comme fonction de transfert DAC-bobines de l'actionneur NE, un modèle faisant intervenir une seule fois les pôles et zéros du filtre passe-bas du DAC, ainsi qu'un pôle à 805 Hz représentant les deux bobines (moyenne des deux).

Des mesures identiques (figure 3.5) ont été effectuées sur l'actionneur du bout de bras ouest (WE). On observe à partir de 1 kHz une différence de comportement des bobines UP et DOWN, atteignant 20 % d'écart sur le module à 2 kHz. L'utilisation de deux fonctions de transfert distinctes pour les deux bobines est pourtant exclue par l'impossibilité de les distinguer dans leur action magnétique sur les aimants du miroir; on doit donc se contraindre à utiliser une fonction de transfert unique. En appliquant la procédure d'ajustement utilisée pour NE sur la bobine DOWN, on obtient un χ^2 de l'ordre de 20 fois le nombre de degré de liberté, montrant l'inadéquation du modèle. Dans le cas de la bobine UP, l'ajustement se comporte comme pour les bobines de NE et place le pôle de la bobine à 920 Hz.

2. Un tel modèle pôles-zéros sera par ailleurs utile pour la reconstruction (*cf* chapitre 4).

3. On verra dans la deuxième partie du chapitre qu'il est de toute façon impossible d'injecter des signaux d'étalonnage au delà, et l'analyse permettant de compenser cette contrainte sera exposée

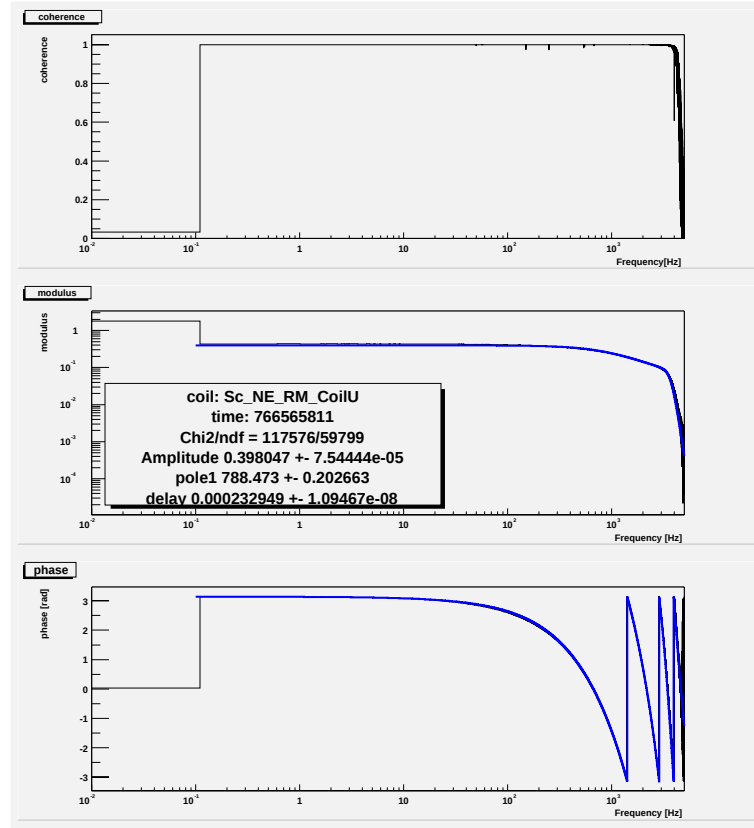


FIG. 3.4 – Ajustement (en bleu) de la fonction de transfert DAC-bobines-ADC pour la bobine UP de l'actionneur NE (en haut: cohérence de la mesure, au milieu: module, en bas: phase).

Ce modèle a été retenu tout en limitant la bande de fréquence de confiance à 1 kHz.

Il faut noter que cet asymétrie des bobines n'est pas choquante dans la mesure où les actionneurs sont conçus pour le système de contrôle, dont les signaux de correction ne dépassent pas 500 Hz. On rencontre ici les difficultés liées à l'utilisation des actionneurs dans un but différent de leur objet principal. Par la suite, on a tout naturellement préféré assurer l'essentiel de la calibration avec l'actionneur NE.

Accessoirement, des mesures similaires ont été conduites sur l'actionneur de la séparatrice (BS). La raison principale était de savoir si les disparités observées sur WE s'étendaient à un plus grand nombre de bobines; de plus cette fonction de transfert sera utile pour la reconstruction (*cf* chapitre 4). Il s'avère que les 4 bobines utilisées pour le contrôle longitudinal de BS ne présentent pas d'écart entre elles comme celles de WE (figure 3.6). Les ajustements des fonctions de transfert par le modèle habituel produisent des valeurs de χ^2 proches de quatre

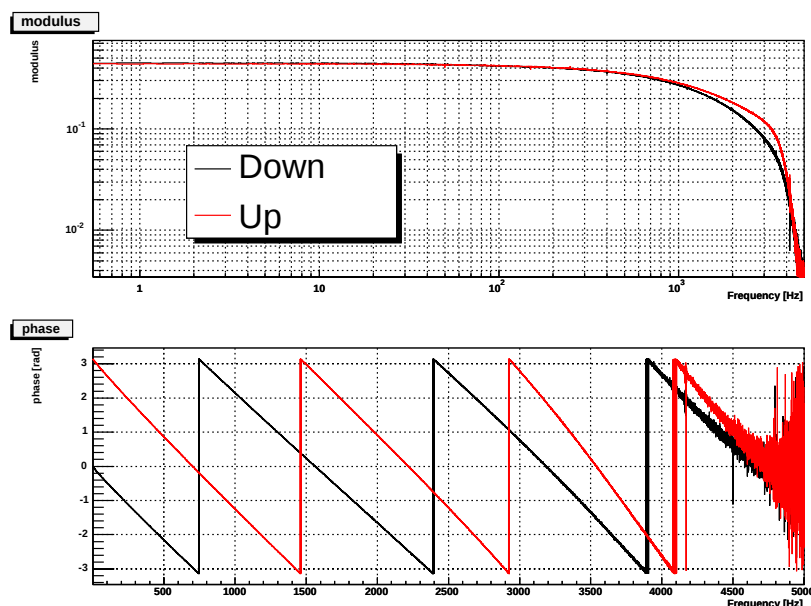


FIG. 3.5 – Fonctions de transfert DAC-bobine-ADC mesurées pour les deux bobines de l'actionneur WE (en haut: module, en bas: phase).

fois le nombre de degrés de liberté, pour des pôles situés autour de 425 Hz ($\pm 2\%$).

3.2.3 Stabilité de l'ensemble au cours du temps

Les résultats précédents doivent être utilisés pour toute injection de signaux de calibration et n'ont d'intérêt que s'ils restent valides au cours du temps. On s'intéresse donc à toute évolution de la réponse des actionneurs, à court terme -durée d'un segment de *lock* ou d'un *run*- comme à long terme -tout au long du *commissioning*.

Variations avec la température des bobines

Des problèmes liés à la température des bobines ont été observés lors du *commissioning run* C2 (20 - 23 février 2004). Durant ce *run* utilisant une simple cavité, quatre lignes de calibration (sinusoïdes à différentes fréquences) étaient injectées en permanence sur NE, ainsi que une sur NI. Le suivi des lignes de calibration dans le signal de frange noire sur une période de *lock* de plus de deux heures (figure 3.7) fait apparaître une contradiction entre la dérive des quatre lignes de NE et la stabilité de celle sur NI. La cavité ne distingue pas un déplacement de l'un ou l'autre des miroirs: l'asymétrie entre NE et NI se trouve donc au niveau des actionneurs. La même dérive est observée sur le suivi des lignes

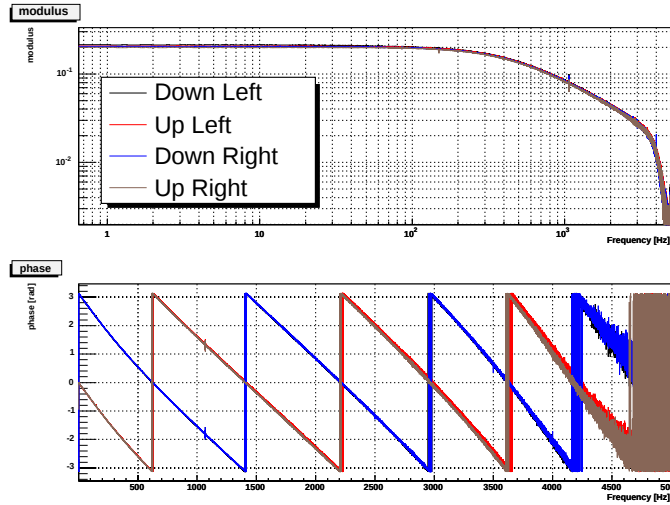


FIG. 3.6 – *Fonctions de transfert DAC-bobine-ADC mesurées pour les quatre bobines de l'actionneur NE (en haut: module, en bas: phase).*

de NE dans la mesure de courant des bobines (figure 3.8), montrant que c'est la réponse DAC-bobine de NE qui varie au cours du temps. Cette dérive s'explique par l'utilisation de l'actionneur NE par le système de contrôle, combiné à la grande amplitude des marées terrestres durant le *run*. Celles-ci font en effet varier la longueur du bras, obligeant le système de contrôle à appliquer un signal de correction toujours plus intense au miroir (figure 3.9). La bobine s'échauffe alors et sa réponse - sa résistivité- diminue. Les figures 3.8 et 3.9 donnent les ordres de grandeurs du phénomène : le signal de correction doit atteindre quelques volts pour engendrer une variation conséquente de la réponse des actionneurs. Depuis avril 2004 la composante basse fréquence des corrections (DC - 10 mHz) contenant les effets de marées est appliqué à la partie haute de la suspension, et les signaux de correction au niveau du miroir ne dépassent plus $\pm 1 \text{ volts}$; les problèmes de variation de la réponse des actionneurs n'ont donc plus été observés.

Modifications des actionneurs

Comme de nombreuses parties de l'interféromètre les actionneurs évoluent avec le commissioning. Les actionneurs présentés ont été utilisés pendant toute la première année du commissioning. Ils sont néanmoins trop bruyants pour permettre d'atteindre la sensibilité finale de Virgo, et ont donc été remplacés en septembre 2004 par un système moins bruyant. On attend en conséquence une modification de la fonction de transfert DAC-bobines, en particulier la disparition du pôle de la bobine, repoussé à plus haute fréquence. Par ailleurs les ADC de lecture du courant des bobines ont été munis de filtres de compression pour améliorer la

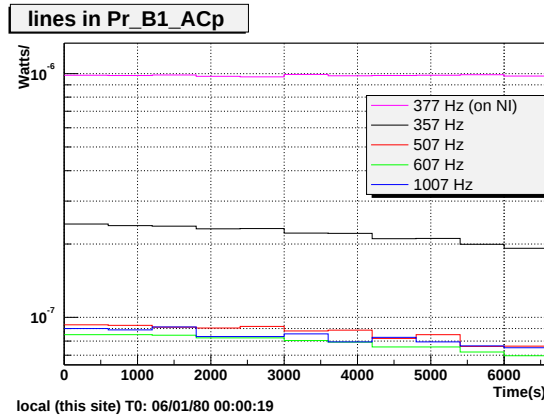


FIG. 3.7 – Amplitude des lignes de calibration (injectées sur NE et NI) dans le signal de frange noire, au cours d'un segment de lock (C2).

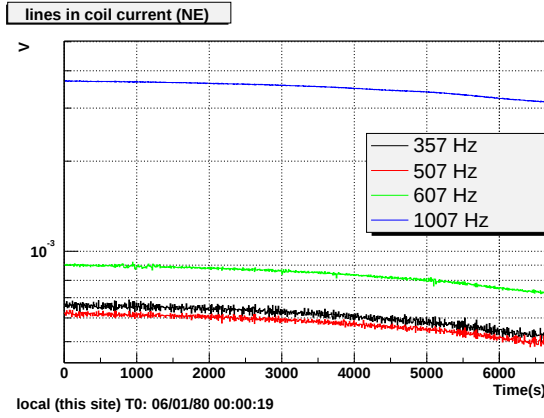


FIG. 3.8 – Amplitude des lignes de calibration injectées sur NE, mesurée dans le signal de courant des bobines de l'actionneur, sur le même segment que la figure 3.7

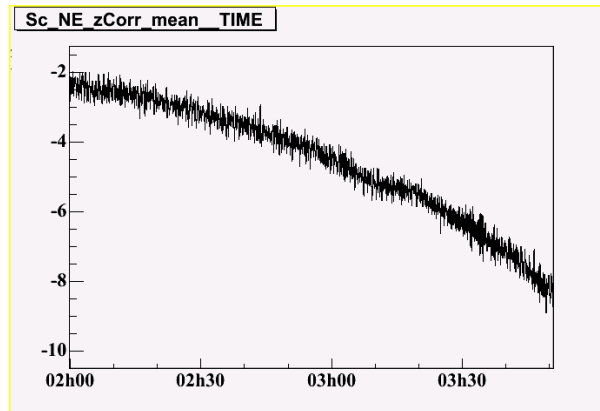


FIG. 3.9 – Évolution du signal de contrôle envoyé à l'actionneur NE, sur le même segment que la figure 3.7).

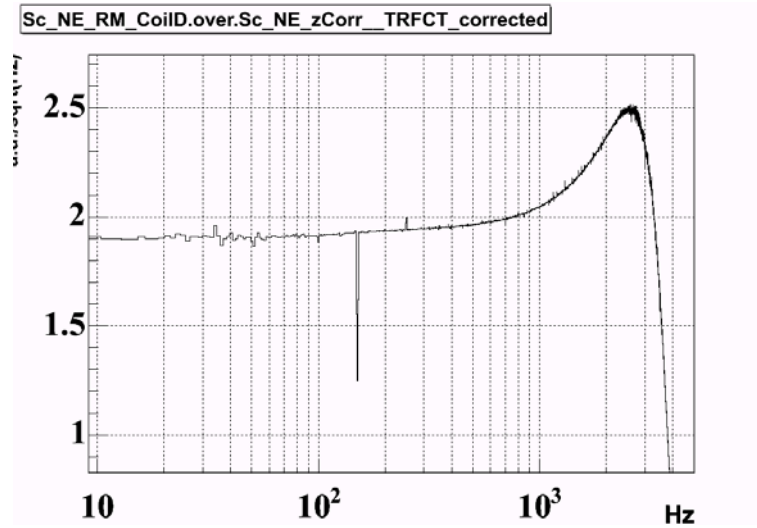


FIG. 3.10 – Fonction de transfert DAC-bobines-ADC mesurée après la modification du système (l'effet des filtres de compression est préalablement déduit en utilisant les valeurs nominales des composants de ces filtres).

mesure.

Une mesure de fonction de transfert DAC-bobines-ADC similaires au précédentes a pu être effectuée après cette modification (figure 3.10). Si la fonction de transfert est essentiellement plate comme prévu, la bosse après 1 kHz et l'abaissement de la fréquence de coupure - en référence à la coupure des filtres passe-bas à 3.7 kHz- ne sont pas compris, on ne peut donc savoir s'ils font partie de la réponse DAC-bobine ou s'ils proviennent de l'ADC de lecture. L'étude de la réponse DAC-bobines devra donc être répétée une fois le nouveau système finalisé.

3.2.4 Mesure du gain absolu en $m/volts$

Revenons à la réponse complète $A(f)$ du système d'étalonnage. Sa forme se déduit des modèles de réponse DAC-bobines adoptés, en incluant le terme mécanique (3.7). Il reste donc à déterminer son gain absolu.

Mesurer ce gain en $m/volts$ implique de pouvoir mesurer les déplacements de miroir générés par les actionneurs. Dans la bande de fréquence d'intérêt (au dessus de la résonance), ces déplacements ne peuvent dépasser $10^{-7}m$. Un dispositif de mesure suffisamment sensible est obtenu en formant, avec la séparatrice et deux miroirs de Virgo, un interféromètre de Michelson non-asservi. Cette méthode a été utilisée avec succès lors du *commissioning* de l'interféromètre central [29]: le déplacement différentiel des miroirs du Michelson peut être reconstruit à partir du signal d'interférence par un algorithme développé à l'époque. Cette reconstruction utilise la puissance non-démodulée et l'une des quadratures du signal de sortie de l'interféromètre, pour estimer son contraste et la puissance moyenne - intégrés sur

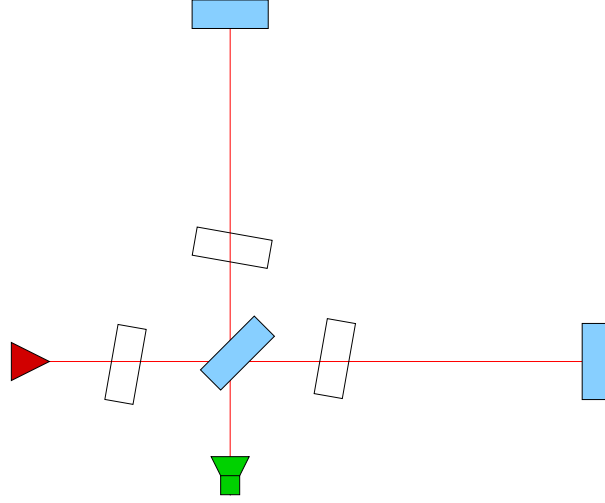


FIG. 3.11 – Configuration du "long Michelson" formé des miroirs de bout de bras NE et WE.

quelques secondes - puis reconstruire le déphasage $\Delta\phi$ entre les faisceaux nord et ouest. Le déplacement est alors donné par:

$$\Delta L = \frac{\lambda}{4\pi} \Delta\phi \quad (3.8)$$

la longueur d'onde du faisceau $\lambda = 1.064 \mu m$ servant d'étalon de distance.

Mesures en "long Michelson"

On commence donc, pour mesurer les gains des actionneurs de bout de bras, par aligner l'interféromètre dit "long Michelson" constitué de la séparatrice et des miroirs NE et WE (Figure 3.11). On vérifie le bon fonctionnement de la reconstruction on comparant le signal reconstruit aux franges d'interférences à la sortie de l'interféromètre:

- à cours terme (figure 3.12): pour 6 franges correspondant à un déplacement $\Delta L = 6\lambda = 3.19\mu m$ la reconstruction donne $\Delta L = 3.17\mu m$. L'écart observé, inférieur à 1%, est négligeable en comparaison des incertitudes qui vont être rencontrés par la suite.
- à plus long terme, on peut compter le nombre d'oscillations du miroir en comptant les éclaircissements visibles entre les paquets de franges. Ces éclaircissements correspondent aux instants où le miroir s'arrête pour faire demi-tour, chaque paquet de frange correspond donc à une demi-période de l'oscillation du miroir (figure 3.13). On a pu vérifier, pour 96 paquets de franges, que l'on comptait 96 demi-périodes du signal reconstruit δL .

Trois lignes de calibration sont injectées sur chaque miroir entre 2 et 10 Hz. Les lignes apparaissent sur le spectre du signal reconstruit δL (figure 3.14) au

3.2. Les actionneurs comme système d'étalonnage

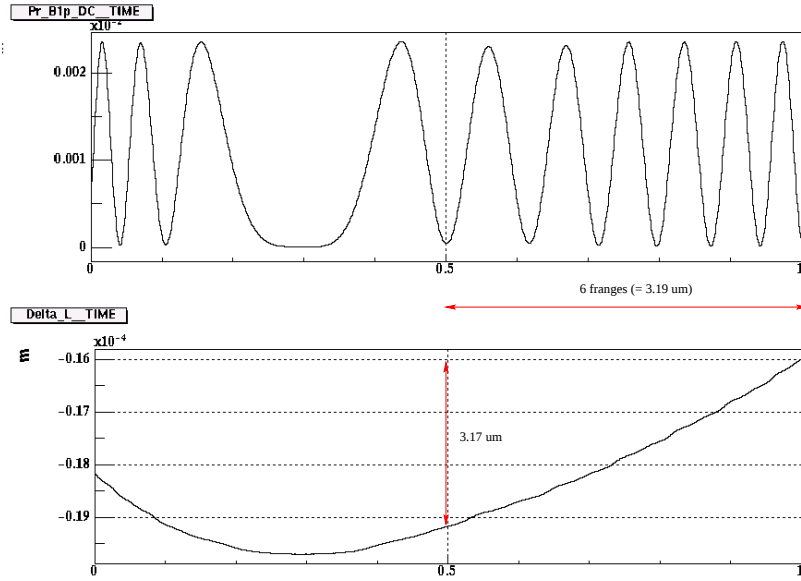


FIG. 3.12 – Puissance mesurée en sortie du Michelson (en haut) et signal reconstitué (en bas). Pour 6 franges comptées dans la puissance, correspondant à un déplacement $\Delta L = 6\lambda = 3.19\mu\text{m}$, le signal reconstruction donne $\Delta L = 3.17\mu\text{m}$.

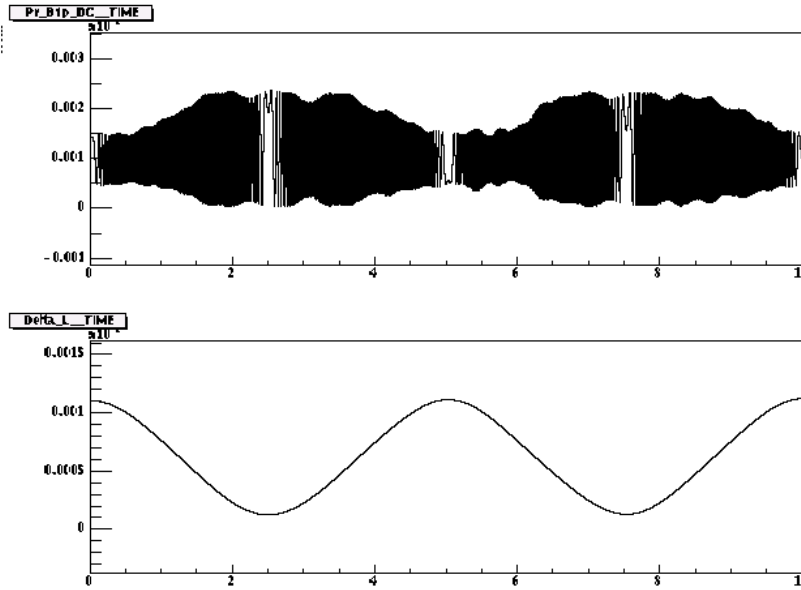


FIG. 3.13 – Puissance mesurée en sortie du Michelson (en haut) et signal reconstitué (en bas). Les éclaircissement dans les franges correspondent aux instants où le miroir s'arrête pour faire demi-tour, chaque paquet de frange correspond donc à une demi-période de l'oscillation du miroir. On vérifie donc que les demi-périodes du signal reconstitué coïncident avec les paquets de franges.

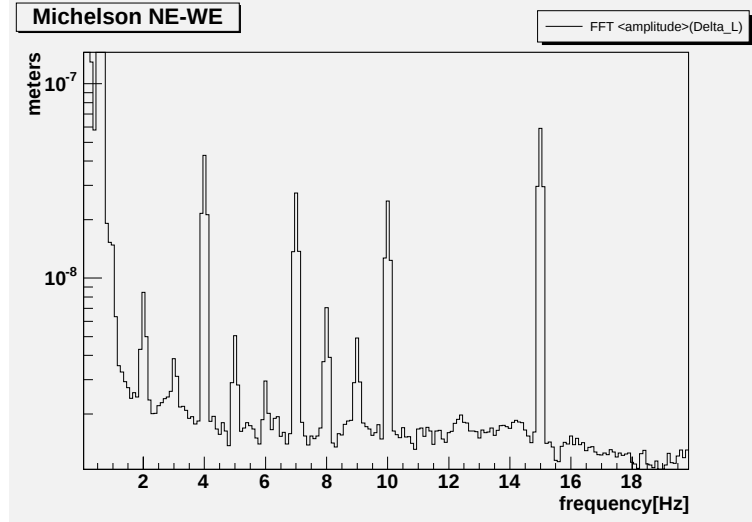


FIG. 3.14 – Mesures en "long Michelson" (21/09/04): spectre du signal reconstruit ΔL , pour un temps d'intégration de 10s. On reconnaît les lignes injectées sur NE et WE (cf tableau 3.1) pour la mesure des gains absolus, ainsi que des lignes supplémentaires plus forte injectées sur BS.

dessus du plancher de bruit estimé à :

$$\sim 2.10^{-9} \frac{m}{\sqrt{Hz}}.$$

Pour mesurer l'amplitude des lignes on choisit un temps d'intégration $T = 5$ min, compromis entre la nécessité d'intégrer longtemps pour augmenter le rapport signal sur bruit et les contraintes de stabilité du Michelson non-asservi. Le bruit sur l'amplitude d'une ligne est alors :

$$\Delta A \simeq \frac{2.10^{-9}}{\sqrt{T}} m \lesssim 2 * 10^{-10} m.$$

Les amplitudes mesurées pour chaque ligne sont comparées aux amplitudes injectées en volts. On en déduit la réponse de l'actionneur $|A(f)|$ pour cette fréquence, puis connaissant la forme de cette réponse on remonte au gain DC. Les valeurs obtenues pour chaque ligne sont présentées dans le tableau 3.1; pour les deux actionneurs la dispersion des mesures associées aux trois lignes est du même ordre que le bruit de mesure.

Après la modification des actionneurs des mesures identiques, mais avec des lignes à plus haute fréquence ont été effectuées. Le rapport signal sur bruit du "long Michelson" (figure 3.17) était alors un peu moins bon que lors des premières mesures et les résultats 3.2 montrent une plus grande dispersion ainsi qu'un biais par rapport aux valeurs précédentes, qui n'était pas attendu à priori même après la modification.

	fréquence (Hz)	amplitude (volts)	$\Delta L(nm)$	Gain $g(\mu m/V)$
NE	2	0.01	14.1	14.2
NE	5	0.04	8.2	14.0
NE	8	0.15	11.8	13.9
WE	3	0.01	5.0	12.1
WE	6	0.04	4.6	11.3
WE	9	0.15	7.9	11.7

TAB. 3.1 – Lignes de calibration et mesures de gain en configuration "long Michelson" (21/09/04)

	fréquence (Hz)	amplitude (V)	$\Delta L(nm)$	Gain $g(\mu m/volts)$
NE	7	0.1	9.01	12.2
NE	27	0.5	2.89	11.7
NE	107	3	1.05	11.3
WE	9	0.1	4.98	11.2
WE	29	0.5	2.59	12.1
WE	109	3	0.93	10.4

TAB. 3.2 – Lignes de calibration et mesures de gain en configuration "long Michelson" (23/11/04)

D'autres configurations de mesures ont été envisagés, dans un double but:

- multiplier les mesures afin de tester leur reproductibilité et estimer l'importance des erreurs systématiques;
- repérer les configurations les plus favorables en vue de mesures plus intensives qui devront être envisagées lorsque Virgo approchera sa sensibilité nominale.

Mesures en "Michelson asymétrique"

La difficulté de cette mesure de gain, en comparaison du cas de l'interféromètre central, s'explique au moins en partie par l'utilisation des miroirs de bout de bras; d'une part à cause des 3 km qui les séparent du point d'interférence et rendent la mesure très sensible à l'alignement, d'autre part parce que la traversée obligatoire des miroirs d'entrée (NI et WI) réduit la puissance des faisceaux interférant. On pense alors résoudre la moitié du problème en formant un Michelson asymétrique avec un miroir de bout de bras (celui dont on souhaite mesurer le gain) et le miroir d'entrée de l'autre bras; cette configuration permet la mesure successive des gains WE - avec un Michelson utilisant WE et NI - et NE - avec un Michelson utilisant NE et WI.

Cette configuration a pu être testée avec un Michelson WE - NI (figure 3.15).

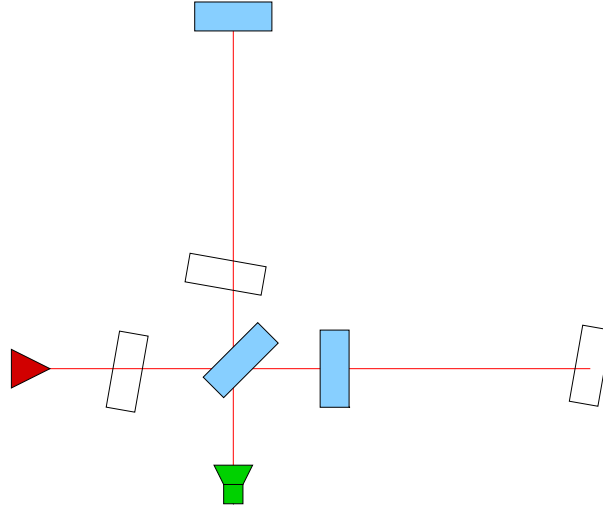


FIG. 3.15 – Configuration du "Michelson asymétrique" formé de NI et WE

	fréquence (Hz)	amplitude (volts)	$\Delta L(nm)$	Gain $g(\mu m/volts)$
WE	9	0.1	5.36	12.0
WE	29	0.5	2.33	10.9
WE	109	3	0.93	10.4

TAB. 3.3 – Lignes de calibration et mesures de gain en configuration "Michelson asymétrique" (23/11/04).

On observe effectivement une amélioration du rapport signal sur bruit des lignes de calibration (figure 3.17), qui permettent la mesure du gain WE (tableau 3.3).

Mesures en "petit Michelson et cavité"

La troisième méthode de mesure envisagée consiste à étalonner un actionneur de bout de bras, par exemple NE, sur son "voisin de cavité" NI, lui-même étant étaloné sur la longueur d'onde du laser. On utilise pour cela les cavités, et en particulier la symétrie de leur réponse à un déplacement du miroir de bout de bras ou d'entrée.

Le première étape consiste à mesurer le gain absolu de NI g_{NI} . Pour cela on injecte une ligne sur NI en configuration "petit Michelson" (constitué des miroirs NI et WI (figure 3.16). Comme il n'y a pas de mesure de courant des bobines sur NI (ou WI), on choisit de faire toutes les mesures à des fréquences pour lesquelles la réponse DAC-bobines peut raisonnablement être laissée de côté ($f < 100Hz$). Comme précédemment, on utilise l'amplitude de la ligne dans le signal reconstruit ΔL pour mesurer la réponse $|A_{NI}(f)|$ en $m/volts$ à la fréquence de la ligne, dont on déduit le gain g_{NI} en corrigeant simplement de l'atténuation mécanique (eq. 3.7). Les résultats obtenus pour g_{NI} et g_{WI} sont présentés dans le tableau 3.4.

3.2. Les actionneurs comme système d'étalonnage

	fréquence (Hz)	amplitude (volts)	$\Delta L(nm)$	Gain $g(\mu m/volts)$
NI	7	0.1	9.83	13.3
NI	27	0.5	3.14	12.7
WI	9	0.1	5.86	13.1
WI	29	0.1	2.74	12.8

TAB. 3.4 – Lignes de calibration et mesures de gain en configuration "petit Michelson" (23/11/04)

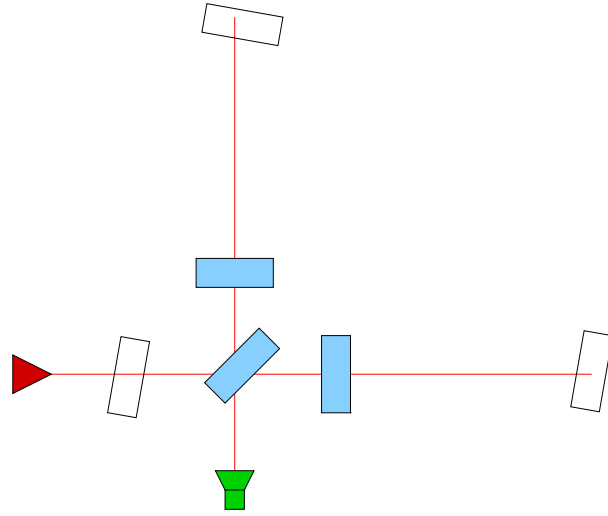


FIG. 3.16 – Configuration du "petit Michelson", formé des miroirs d'entrées NI et WI.

Pour la deuxième étape, on utilise la cavité nord asservie à la frange noire, sur laquelle on injecte simultanément:

1. une ligne de calibration de fréquence f_{NI} sur NI. On mesure alors le gain G_{NI} en *watts/volts* de la ligne entre le signal de calibration et le signal de frange noire de la cavité.
2. une ligne de calibration de fréquence f_{NE} sur NE. On mesure alors le gain G_{NE} en *watts/volts* de la ligne entre le signal de calibration et le signal de frange noire de la cavité.

La cavité ne distinguant pas un déplacement sur NE d'un déplacement sur NI, répond identiquement aux deux. On a donc

$$\begin{aligned}
 G_{NI} &= |A_{NI}(f_{NI})| \cdot |R_{cavity}(f_{NI})| \\
 G_{NE} &= |A_{NE}(f_{NE})| \cdot |R_{cavity}(f_{NE})|
 \end{aligned}$$

où R_{cavity} est la réponse - inconnue mais unique - de la cavité asservie. Dans le cas où $f_{NI} \simeq f_{NE}$ les différents termes se compensent et l'on obtient l'étalonnage

	$f_{NI}(Hz)$	$f_{NE}(Hz)$	Gain $g(\mu m/volts)$
NE	9	7	11.2
NE	29	27	11.4
	$f_{WI}(Hz)$	$f_{WE}(Hz)$	Gain $g(\mu m/volts)$
WE	9	7	10.6
WE	29	27	11.5

TAB. 3.5 – *Lignes de calibration et mesures de gain avec les cavités combinées au "petit Michelson" (23/11/04).*

de NE sur NI:

$$\frac{g_{NE}(f)}{g_{NI}(f)} = \frac{G_{NE}}{G_{NI}}$$

que l'on applique aux mesures de g_{NI} en "petit Michelson" pour déduire le gain absolu g_{NE} .

Les mesures en cavité ont été effectuées pour les mêmes fréquences que les mesures de g_{NI} en "petit Michelson", puis répétées dans le cas symétrique de WE (tableau 3.5).

Bilan des mesures de gain

La dispersion des mesures présentées ci-dessus par chaque actionneur donne une estimation de la précision globale de la mesure, de l'ordre de $\pm 1\mu m/volts$. Le plus grand écart est observé dans le cas de NE pour lequel la valeur mesurée du gain absolu chute de 20% entre les mesures de septembre 2004 et celles effectuées deux mois plus tard. Bien que les actionneurs aient été modifiés entre temps (cf 3.2.3), il n'existe aucune indication indépendante d'une variation du gain à cette occasion; on doit donc inclure cet écart dans les erreurs de mesures.

Une sous-estimation du gain de l'actionneur utilisé pour l'étalonnage - principalement NE - revient à sous-estimer les signaux d'étalonnage injectés dans l'interféromètre. La procédure d'étalonnage conduira alors à le croire plus sensible que ce qu'il est réellement. On préfère au contraire adopter les valeurs les plus élevées :

$$g_{NE} = 14\mu m/volts \quad (3.9)$$

$$g_{WE} = 12\mu m/volts \quad (3.10)$$

et retenir que la sensibilité pourrait être jusqu'à 20% *meilleure* que celle fournie par la procédure d'étalonnage.

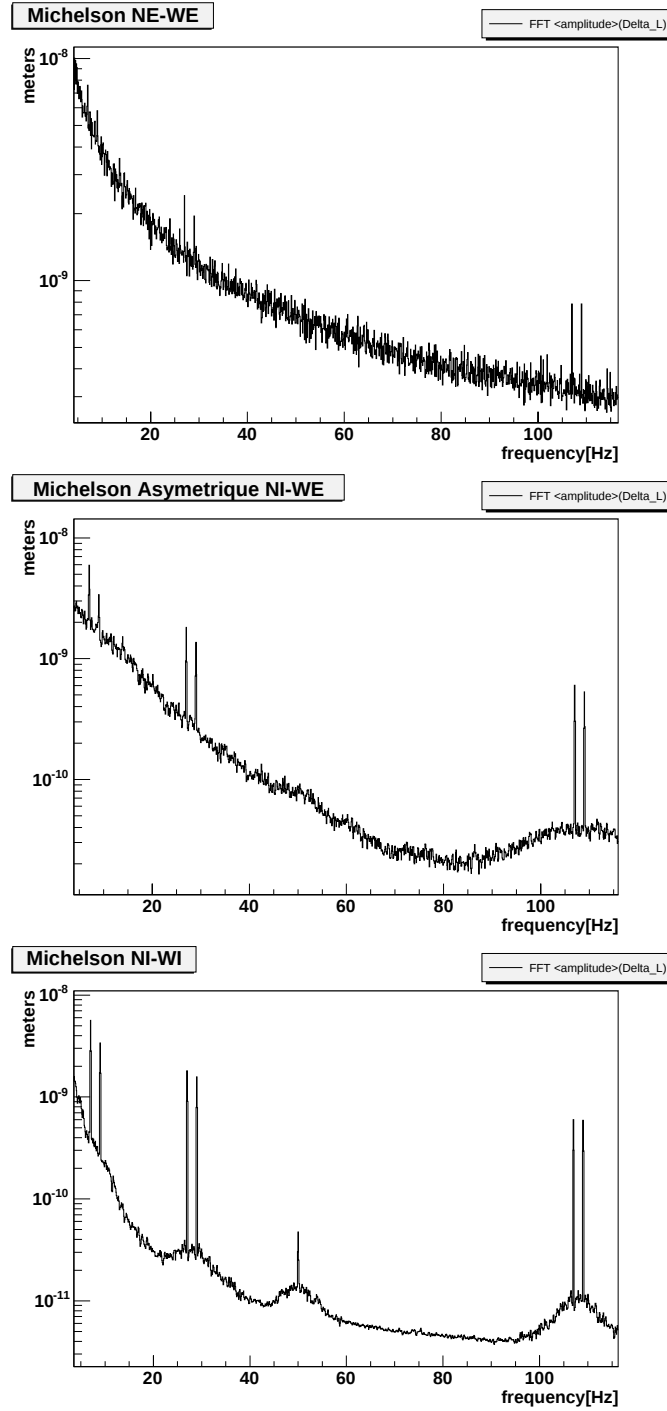


FIG. 3.17 – Spectre du signal reconstruit ΔL , pour un temps d'intégration de 10s, pour les différentes configuration de mesure (23/11/04). En haut: le "long Michelson", au milieu: le "Michelson asymétrique", en bas: le "petit Michelson". Les amplitudes en volts des lignes injectées sont les mêmes d'une configuration à l'autre (cf tableaux 3.2, 3.3 et 3.4). Le meilleur rapport signal sur bruit est offert par le "petit Michelson".

3.3 Réponse du détecteur

3.3.1 Description

La réponse $R(f)$ inclut les deux effets bien distincts que Virgo fait "subir" au signal ΔL_{OG} :

- Les effets des contrôles:

Le signal est atténué à basse fréquence par le système de contrôle de Virgo, qui applique des forces aux miroirs de bout de bras pour maintenir l'interférence à la frange noire. La différence de longueur réelle s'écrit donc:

$$\Delta L_{reel}(f) = H(f) \cdot \Delta L_{OG}(f) \quad (3.11)$$

où H est la fonction de transfert - en boucle fermée - des asservissements. Ceux-ci n'agissant qu'aux basses fréquences, jusqu'à environ $300Hz$. Au delà on a $H \simeq 1$

- La réponse optique:

C'est la réponse (notée O) entre le déplacement réel et le signal détecté:

$$S(f) = O(f) \cdot \Delta L_{reel}(f) \quad (3.12)$$

Elle inclue l'effet d'intégration du signal sur le temps de stockage $\tau = \frac{2L}{c} \frac{\mathcal{F}}{\pi}$ de la lumière dans les cavités Fabry-Pérot des bras de Virgo, qui se manifeste comme une coupure à la fréquence:

$$f_c = \frac{1}{2\pi\tau_s} \quad (3.13)$$

A très haute fréquence ($\gtrsim 10kHz$), elle inclue également la coupure des filtres anti-repliement utilisés pour la détection du signal de frange noire.

Finalement la réponse $R(f)$ est la combinaison de ces deux fonctions de transfert:

$$R(f) = H(f) \cdot O(f) \quad (3.14)$$

3.3.2 Mesure

Une fois la fonction de transfert des actionneurs connue, ceux-ci peuvent être utilisés pour injecter des signaux de calibration dans l'interféromètre. On peut alors mesurer la réponse du détecteur, ou plus précisément la réponse inverse permettant de traduire le signal de frange noire en un déplacement du mode différentiel. On procède pour cela en deux étapes:

La prise de données:

On injecte pendant 5 min du bruit blanc - ou coloré pour mieux s'adapter au bruit naturel du détecteur - *via* l'actionneur NE, qui se retrouve dans le signal de frange noire et domine le bruit naturel du détecteur dans la bande de fréquence la plus large possible. Le niveau d'injection est limité par la dynamique des DACs des actionneurs, celle des ADC de la détection du signal de frange noire, et la stabilité de l'asservissement global de l'interféromètre qui ne doit pas quitter le régime linéaire sous l'effet du déplacement de calibration. La réponse des actionneurs $A(f)$ joue de plus son rôle de filtre passe bas; la coupure des DAC à $3.7kHz$ interdit toute injection au dessus de cette fréquence, et l'atténuation de la réponse mécanique réduit plus encore la bande de fréquence accessible. En pratique on dépasse difficilement $1kHz$.

Le calcul de la fonction de transfert:

On calcule la fonction de transfert TF en *Watt/volts* entre le signal de calibration et le signal de frange noire, sur des segments consécutifs de 10s - pour une résolution en fréquence de $0.1Hz$ - que l'on moyenne sur les 5mn de données. La réponse inverse R^{-1} en *m/watts* s'obtient alors par simple inversion, après déduction de la réponse des actionneurs:

$$R^{-1}(f) = \left[\frac{TF(f)}{A(f)} \right]^{-1} \quad (3.15)$$

3.3.3 Analyse

La validité de R^{-1} étant limitée à la bande de fréquence de mesure ($< 1kHz$), on doit extrapoler ce résultat expérimental jusqu'à 10 kHz. L'idée générale est alors d'utiliser un modèle de la réponse inverse, que l'on ajuste à la mesure là ou celle-ci est disponible.

Ajustement par un modèle

De façon générale, les ajustements sont effectués simultanément sur le module et la phase - comme pour les fonctions de transfert DAC-bobines. Les erreurs utilisées pour la normalisation du χ^2 sont calculées à partir de la cohérence C entre le signal de calibration et le signal de frange noire, suivant:

$$\frac{dM}{M} = \alpha \sqrt{\frac{(1-C)}{C}} \frac{1}{\sqrt{n}} \quad \text{avec} \quad \alpha = 0.85 \quad (3.16)$$

$$dP = \beta \sqrt{\frac{(1-C)}{C}} \frac{1}{\sqrt{n}} \quad \text{avec} \quad \beta = 0.88 \quad (3.17)$$

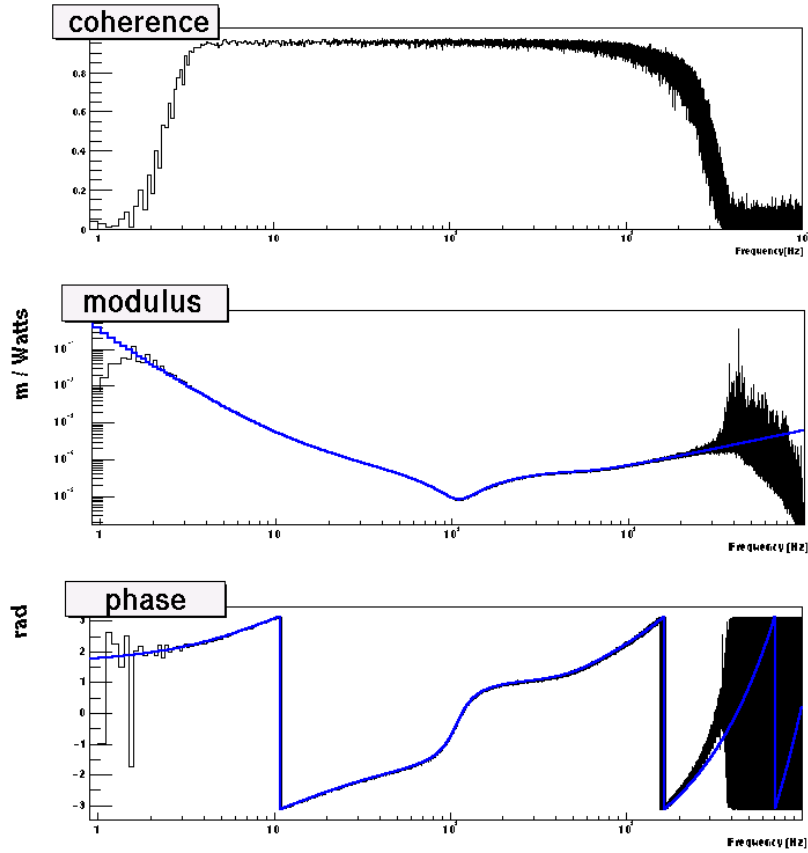


FIG. 3.18 – Ajustement (en bleu) de la réponse inverse (en noir) d'une cavité asservie simulé par SIESTA.

où n est le nombre de moyennes effectuées pour calculer la fonction de transfert. Le choix du modèle est délicat; cet aspect de la procédure de calibration a beaucoup évolué au cours du commissioning, notamment face à la difficulté de prendre en compte les effets des asservissements dans le modèle (les plus difficiles à modéliser, notamment en présence de couplage entre plusieurs boucles de contrôle), alors même que ceux-ci évoluent très rapidement. L'idée initiale d'une procédure utilisant un modèle complet de la réponse inverse a pu être mise en oeuvre dans certains cas. La figure 3.18 montre ainsi l'ajustement de la réponse inverse d'une cavité asservie simulée par SIESTA; la figure 3.19 montre l'ajustement dans le cas de la cavité nord de Virgo (C1), limitée à la partie "effet des contrôles".

Cette stratégie lourde à mettre en oeuvre et à mettre à jour s'est progressivement effacée au profit de la procédure semi-empirique décrite ici, qui ne dépend ni de la stratégie de contrôle ni de la configuration du détecteur - d'une simple cavité à l'interféromètre recyclé.

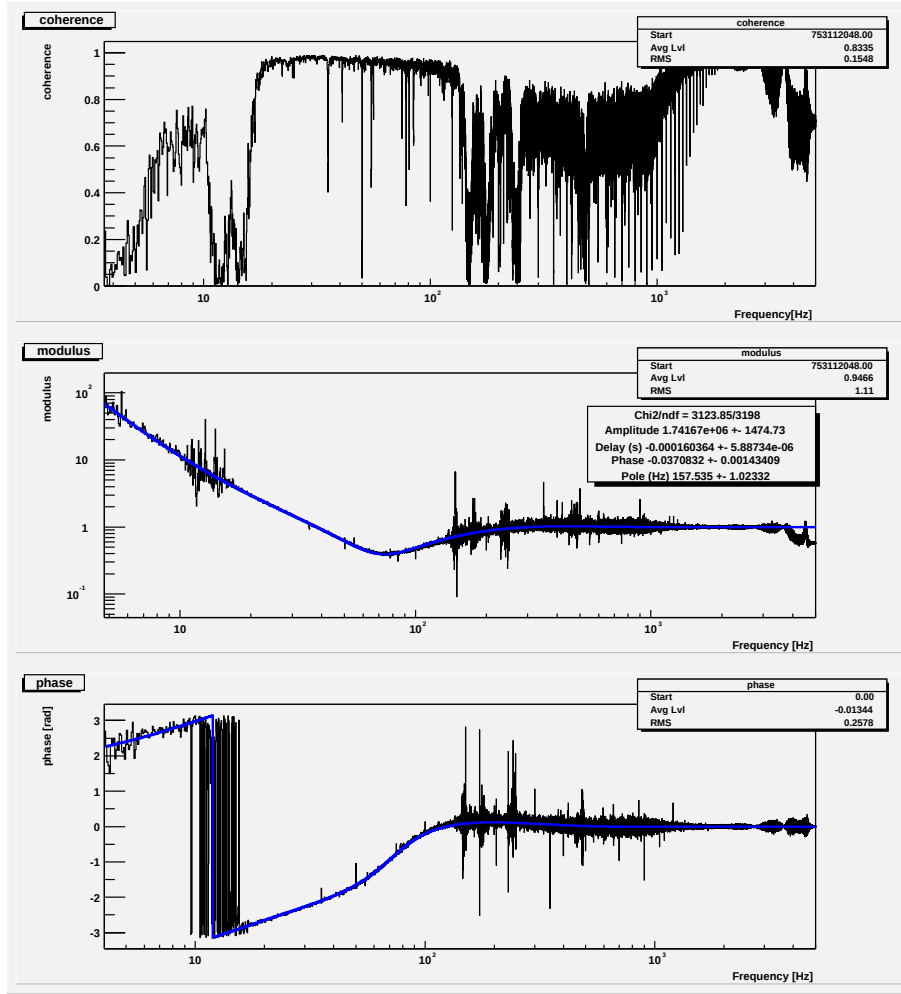


FIG. 3.19 – Ajustement (en bleu) de la fonction de transfert des contrôles (en noir) de la cavité nord (C1).

Ajustement par un modèle partiel

Cette procédure plus souple repose sur quelques hypothèses simples:

1. R^{-1} est correctement mesurée entre $f_l \sim 10Hz$ et $f_u \sim 1kHz$
2. les effets des contrôles sont complètement négligeables au dessus de $f_c \sim 500Hz$, la réponse du détecteur se limitant alors à la réponse optique.

Le principe est alors de reconstituer R^{-1} par morceaux en utilisant:

- la mesure entre f_l et f_c ,
- un ajustement entre f_c et $10KHz$

L'ajustement peut alors utiliser un modèle tenant compte de la seule réponse optique, soit :

- un gain et un retard (paramètres libres de l'ajustement),

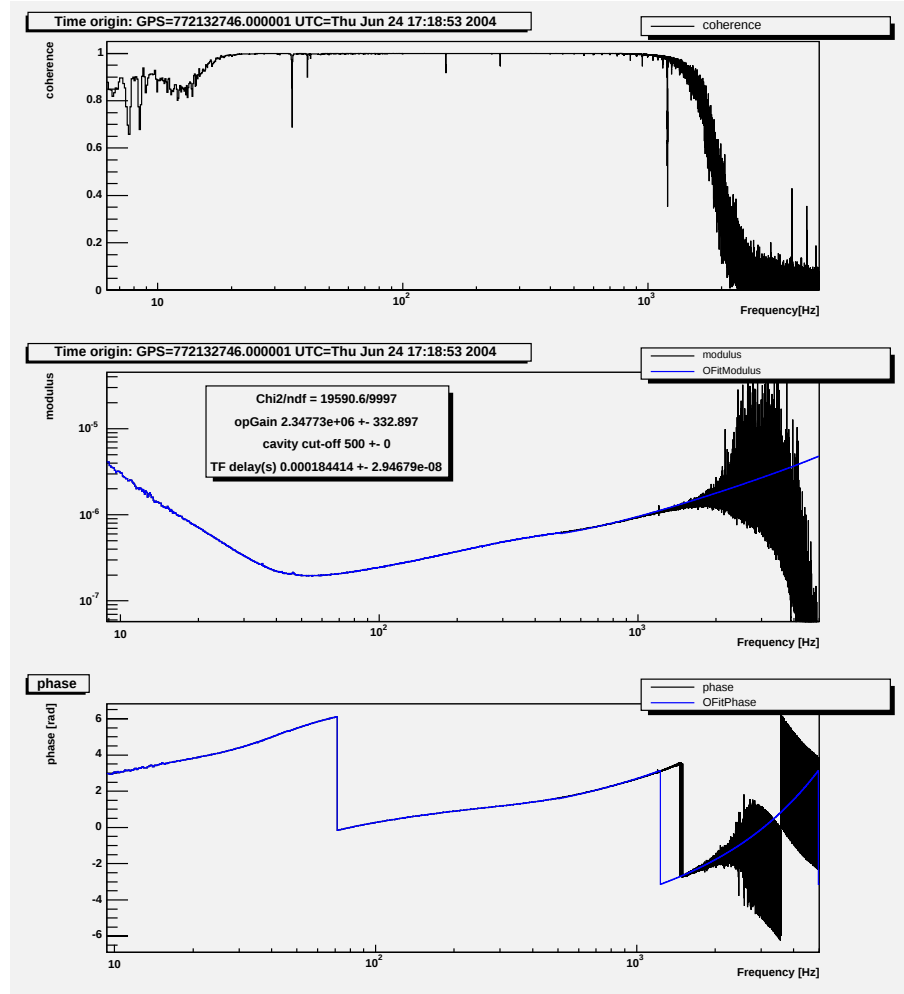


FIG. 3.20 – Application de la procédure générique à la calibration de l'interféromètre recombiné

- un pôle à 500 Hz représentant la coupure des cavités,
- les pôles et zéros des filtres anti-repliement appliqués à la détection du signal de frange noire.

La figure 3.20 montre un exemple d'application de cette procédure à la calibration de l'interféromètre recombiné.

Les cas particuliers

La procédure décrite ci-dessus est apparue avec l'interféromètre recombiné. Les méthodes utilisées au cours des *runs* précédents illustrent bien les cas particuliers et l'évolution de la procédure de calibration au cours du temps :

- **C1 et C2** : La position du pôle des bobines est complètement inconnue -

et il est difficile d'injecter des signaux de calibration à haute fréquence - par conséquent la mesure de la réponse optique est impossible. On effectue un ajustement par un modèle sur la fonction de transfert des contrôles uniquement (figure 3.19). Pour compléter la fonction de transfert on utilise le modèle réponse optique sans ajustement, une simple ligne de calibration à $\sim 350Hz$ servant à déterminer le gain optique.

- **C3 bras nord** : La stabilisation en fréquence du laser agit exceptionnellement sur le mode différentiel (qui se confond avec le mode commun dans le cas d'une simple cavité). Comme il s'agit d'une boucle rapide les effets dans la réponse se font sentir jusqu'à très haute fréquence, et l'extrapolation par un modèle simple de la réponse optique n'a plus de sens. On parvient cependant à utiliser cette boucle rapide pour injecter des signaux de calibration, sous la forme de perturbations du signal d'erreur. La mesure est ainsi étendue de $1kHz$ à $4.5kHz$.

3.4 Sensibilité

Une fois connue la réponse inverse du détecteur de $10Hz$ à $1kHz$, on calcule le spectre du signal de sortie en l'absence de tout signal de calibration $\tilde{S}_n(f)$ en *watt*, sur des segments consécutifs de $10s$ - comme pour la fonction de transfert - que l'on moyenne sur les 5 min de données. Le produit par la fonction de calibration nous donne alors la sensibilité (équations 3.4 ou 3.5).

La figure 3.21 montre les courbes de sensibilité ainsi obtenues pour les *run* C1 à C5, résultat final de l'étape d'étalonnage.

L'incertitude de mesures de ces courbes est dominée par l'erreur sur le gain absolu des actionneurs, soit 20%.

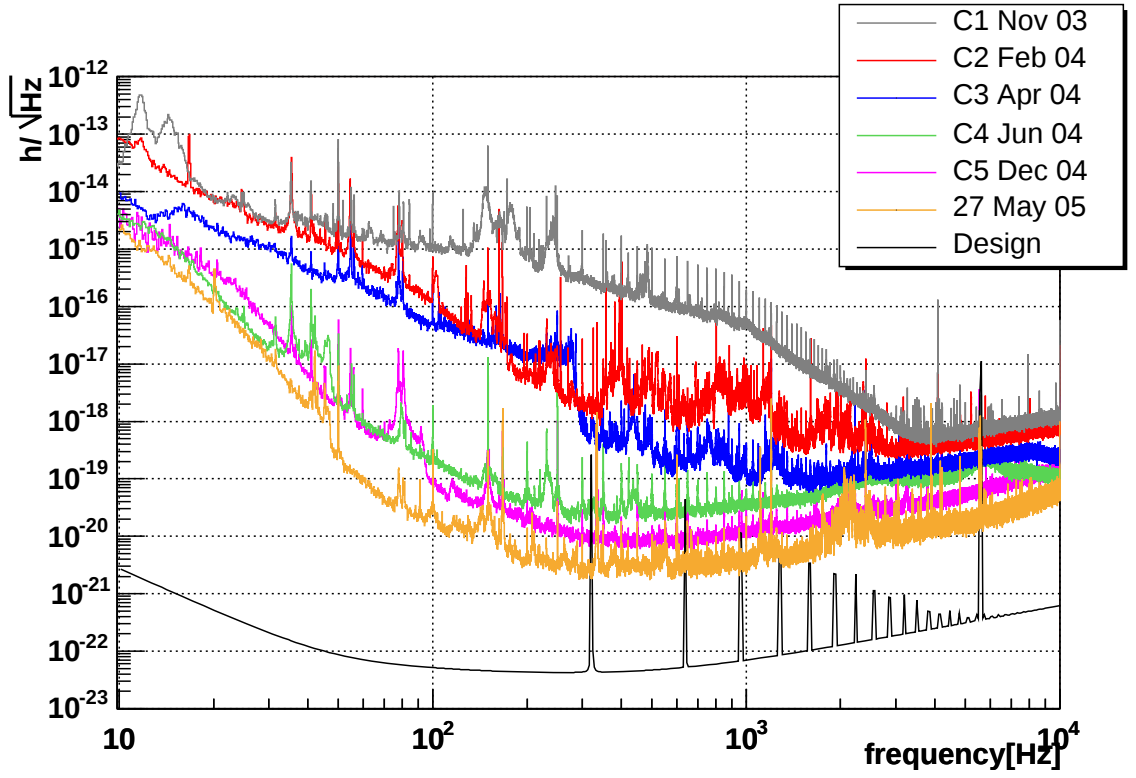


FIG. 3.21 – Courbes de sensibilité en $\frac{h}{\sqrt{\text{Hz}}}$ obtenues lors des différents runs au cours de la mise en route de Virgo (+ courbe du 5 mai 2005).

Chapitre 4

Reconstruction du signal gravitationnel $h(t)$

4.1 Introduction

La reconstruction du signal gravitationnel $h(t)$ détecté par Virgo constitue la première étape de l'analyse de données.

Les algorithmes de recherche de signaux d'ondes gravitationnelles ne peuvent être appliqués directement au signal de sortie de Virgo $s(t)$, car les signaux recherchés sont transformés par la réponse du détecteur, dont la forme complexe¹ a été mise en évidence au chapitre précédent.

Intégrer cette réponse dans chaque algorithme n'est pas trivial dans la mesure où :

- Elle s'exprime dans le domaine fréquentiel, alors que certains algorithmes ne travaillent que dans le domaine temporel.
- Elle peut varier au cours du temps, et doit donc être suivie au moyen de lignes de calibration, injectées en permanence.

Il est donc naturellement préférable de concevoir un processus situé en amont des divers algorithmes, chargé de suivre les variations de la réponse et de reconstruire en permanence le signal $h(t)$ qu'ils analysent ensuite. C'est aussi l'occasion d'améliorer la qualité des données en soustrayant certains bruits du détecteur.

La méthode la plus évidente semble être de traiter le signal de frange noire pour "inverser" l'effet de la réponse de Virgo. Le principe de cette méthode et les difficultés rencontrées lors de sa mise en oeuvre, qui ont conduit à la rejeter, sont décrits dans la première partie de ce chapitre.

La deuxième partie présente la mise au point et la validation d'une méthode

1. La recherche de coalescences binaires, qui explore toute la bande passante du détecteur jusqu'à quelques kHz, et s'appuie sur une bonne connaissance du signal attendu, est particulièrement sensible aux effets de la réponse sur les signaux recherchés.

alternative combinant les différents signaux de contrôle de l'interféromètre au signal de frange noire.

L'adaptation de cette méthode au cas réaliste où la réponse varie au cours du temps et doit être suivie au moyen de lignes permanentes de calibration est décrite ensuite.

Enfin, la dernière partie concerne la soustraction des harmoniques de 50 Hz dans le signal reconstruit, dernière étape avant la recherche de signaux d'ondes gravitationnelles.

4.2 Méthode utilisant le signal de frange noire

4.2.1 Principe

La première méthode envisagée est issue de l'expérience du CITF, pour lequel elle a pu être utilisée avec succès [29] sur les données du *run* E4. Elle consiste à filtrer le signal de frange noire issu de Virgo pour en déduire le signal gravitationnel $h(t)$. Le filtre à implémenter est alors déterminé par sa fonction de transfert, qui n'est autre que la réponse inverse de Virgo:

$$h(f) = \frac{\Delta L(f)}{3km} = \frac{R^{-1}(f)}{3km} \cdot s(f) \quad (4.1)$$

Un tel filtre est facilement réalisé si l'on s'autorise un "aller-retour" dans le domaine fréquentiel par voie de transformées de Fourier rapides (FFT) et de transformées inverse (FFT^{-1}):

$$h(t) = FFT^{-1}\left\{\frac{R^{-1}(f)}{3km} \cdot FFT[s(t)]\right\} \quad (4.2)$$

la reconstruction étant alors effectué par segments de données de durée finie, se chevauchant suffisamment pour éviter les effets de bord.

Si elle paraît simple, cette méthode n'est pas pour autant satisfaisante: le suivi progressif de l'évolution de la réponse est alors impossible, celle-ci ne pouvant être mise à jour que par à-coups (entre deux segments), perturbant la continuité du signal reconstruit.

On préfère donc un filtre dans le domaine temporel, dont les paramètres peuvent être mis à jour en temps réel. Pour construire ce filtre de façon standard on doit utiliser un modèle pôles-zéros de R^{-1} , issu de l'ajustement de la fonction mesurée (§3.3.3). Un tel modèle tiens compte des effets des contrôles aussi bien que de la réponse optique, et repose donc sur trois termes principaux:

- O : la coupure des cavités, représentée par un pôle à 500 Hz,
- A : la réponse des actionneurs, soit le pôle des bobines et le terme du second ordre associé au pendule,

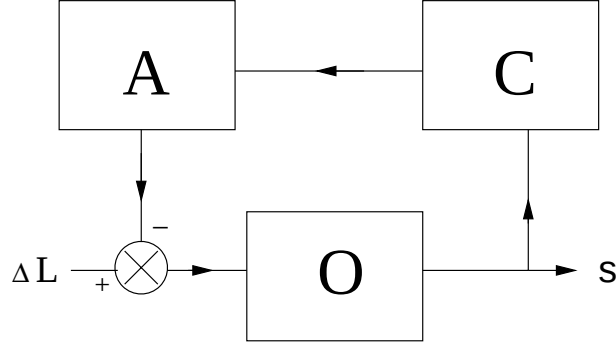


FIG. 4.1 – schéma de la réponse pour un asservissement à un degré de liberté.

- C : le filtre numérique (contrôleur) de l'asservissement longitudinal, défini à la base comme une suite de pôles et de zéros.

Dans le cas où la boucle de contrôle du mode différentiel est seule (cavité simple) ou indépendante des autres boucles, le schéma de principe de cet asservissement est très simple (figure 4.1) et le modèle s'écrit:

$$\Delta L = (\alpha_1 O^{-1} + \alpha_2 AC) \cdot s \quad (4.3)$$

où les paramètres α_k , de la forme:

$$\alpha_k = g_k e^{-i\omega\tau_k} \quad (4.4)$$

sont déterminés par l'ajustement.

En présence de couplage entre les boucles de contrôles, des fonction plus complexes des mêmes termes de base sont introduites pour en tenir compte, comme on le verra dans l'exemple de mise en application qui suit.

Une fois le modèle pôle-zéro disponible, il est convertit en un filtre discret dans le domaine temporelle. Pour cela le modèle est inclus dans la librairie de reconstruction, qui factorise les différents pôle et zéro en un produit de terme du second ordre:

$$R^{-1} = \prod_{k=1}^L \frac{b_{2k}s^2 + b_{1k}s + b_{0k}}{a_{2k}s^2 + a_{1k}s + a_{0k}} \quad (4.5)$$

et les convertit en filtre du second ordre pouvant être successivement appliqué au signal de sortie de Virgo $s(t)$ pour reconstruire le signal gravitationnel $h(t)$.

4.2.2 Essai sur des données simulées

Afin de l'appliquer à Virgo, cette méthode a d'abord été testée en simulation. La configuration choisie, celle du MDC6, met particulièrement en avant les problèmes de couplages entre boucle de contrôles évoqués précédemment. il s'agit en

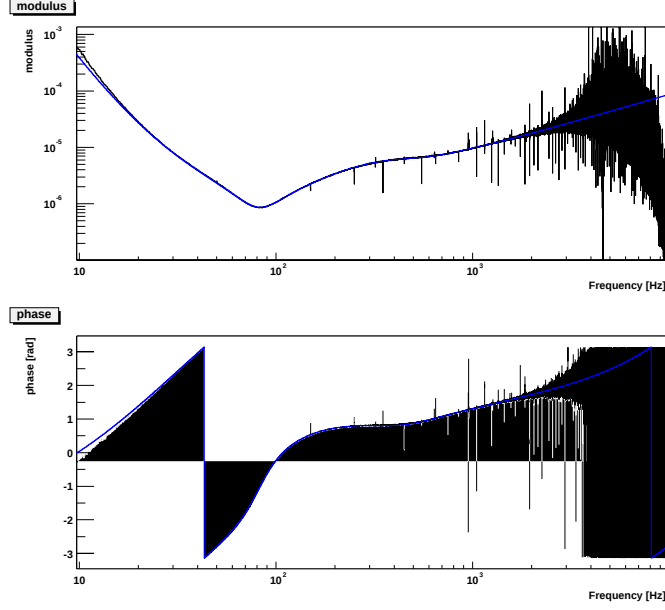


FIG. 4.2 – Ajustement par un modèle (en bleu) de la fonction de transfert mesurée (en noir) d'un interféromètre recombiné simulé. En haut: module, en bas: phase.

effet de l'interféromètre recombiné, avec la stratégie de contrôle utilisée durant C3:

- Les longueurs des deux cavités nord et ouest sont contrôlées indépendamment, respectivement par les photodiodes de bout de bras B7 et B8.
- la frange noire du petit Michelson (MICH) est contrôlée par le signal principal B1, alors que celui-ci est plus sensible au mode différentiel des bras (DARM).

On doit alors inclure dans le modèle de base un terme supplémentaire prenant en compte le couplage entre les modes MICH et DARM:

$$\Delta L = (\alpha_1 O^{-1} + \alpha_2 AC) \cdot (1 - \alpha_3 \frac{AC}{1 + \alpha_4 AC})^{-1} \cdot s \quad (4.6)$$

L'ajustement par ce modèle de la fonction de transfert - mesurée en injectant du bruit de calibration dans la simulation - s'avère satisfaisant malgré un léger écart à basse fréquence (figure 4.2).

Après factorisation des pôles et zéros du modèle, le filtre de reconstruction est constitué d'un enchaînement de 20 filtres du second ordre, auxquels s'ajoute un filtre de compression d'ordre 6 permettant de réduire la dynamique du signal reconstruit.

Ce filtre peut alors être appliqué aux données simulées, en l'absence de signal de calibration, mais la divergence quasi-immédiate du signal reconstruit (Figure

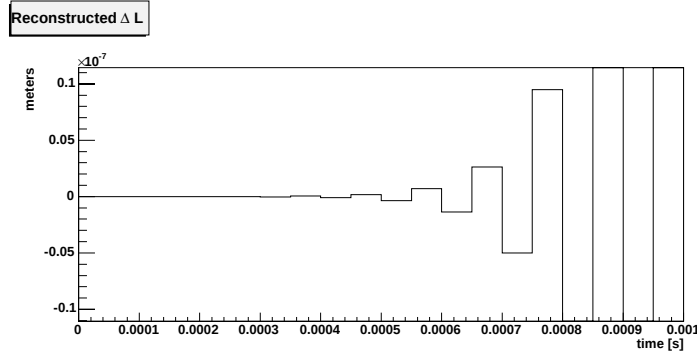


FIG. 4.3 – Les 20 premiers échantillons (1ms) du signal reconstruit à partir du signal de frange noire.

4.3) démontre l'instabilité de ce filtre. Cette instabilité n'est pas surprenante, au regard des nombreuses factorisations appliquées aux termes de la réponse inverse (4.6), factorisations particulièrement sensibles aux effets numériques.

Après vérification, cinq des vingt-deux filtres du deuxième ordre s'avèrent être instable. Si l'on s'attend à pouvoir compenser ce problème, au cas par cas, par l'ajout de zéros à très haute fréquence ou en jouant sur les paramètres de la factorisation du modèle, la complexité du modèle rend cette opération particulièrement délicate.

Cette difficulté s'ajoute à d'autres inconvénients majeurs de cette méthode:

- Elle impose de construire et valider un modèle complet de la réponse systématiquement après chaque changement de configuration, alors même qu'un tel modèle n'est plus nécessaire pour la mesure de sensibilité (§3.3.3).
- Un tel modèle repose entre autres sur la connaissance des pôles et zéros de l'asservissement longitudinal, alors que ceux-ci sont fréquemment modifiés.
- La mise à jour du modèle - pour suivre les variations de la réponse - sur la base des lignes permanentes de calibration n'est pas triviale, car il n'existe pas de relation simple entre les différents paramètres α_k et l'amplitude de ces lignes.

La tentative de mise en oeuvre de cette méthode s'est alors interrompue au profit d'une méthode alternative qui contourne ces difficultés.

4.3 Méthode utilisant le signal de frange noire et les signaux de contrôles

L'instabilité de la cascade de filtres représentant la réponse inverse de Virgo résulte des factorisation compliquées qui doivent être appliquées aux nombreux termes qui la composent (4.6).

Une solution élégante aux problèmes de factorisations de filtres est utilisée dans la reconstruction des données de GEO 600 [30], dans sa configuration "recyclage de puissance", plus proche de celle de Virgo. Pour GEO 600, la boucle du mode gravitationnel n'est pas - ou très peu - couplée aux trois autres boucles longitudinales; la réponse inverse est donc essentiellement de la forme (4.3). Il reste néanmoins à gérer la somme du terme optique ($\alpha_1 O^{-1}$) et du terme lié aux contrôles ($\alpha_2 AC$). Plutôt que de les factoriser, ces deux termes sont appliqués séparément au signal de frange noire; les deux signaux résultant sont ensuite additionnés numériquement, échantillon par échantillon.

Cette méthode donne de bons résultats, et continue actuellement d'être utilisée du moins sur le principe, puisque elle a dû évoluer [31] pour s'adapter à la nouvelle configuration incluant le recyclage du signal. Elle est également développée de façon tout à fait similaire dans LIGO [32].

Dans notre cas elle ne peut s'appliquer directement, toujours à cause de la forme complexe de la réponse inverse (4.6). On peut néanmoins s'en inspirer et retenir la philosophie générale: mieux vaut traiter séparément la partie effet des contrôles de la partie optique. Une nouvelle méthode va donc être mise en oeuvre, qui renonce à la description de la réponse inverse par une unique chaîne de filtres; elle peut être vue comme une extension de la méthode [30] au cas de contrôles plus complexe. Elle visera également à se passer de la connaissance de la stratégie de contrôle et des filtres d'asservissement.

4.3.1 Principe

Les difficultés rencontrées jusqu'ici sont toutes liées à la complexité des effets des contrôles dans la réponse du détecteur. Plutôt que d'inclure ces effets dans la réponse, il est possible de les éliminer au préalable, en les traitant comme des perturbations venant s'ajouter au signal de frange noire, puisqu'une mesure indépendante est fournie en permanence par les *signaux de contrôle* de l'interféromètre appliqués sur les miroirs.

Pour éliminer l'ensemble des contributions du système de contrôle au signal de frange noire, on doit utiliser les signaux $c_i(t)$ de toutes les boucles de contrôle dont les effets sont visibles dans ce signal; au minimum les trois qui agissent sur les modes différentiel DARM et MICH:

- $c_1(t) = c_{NE}(t)$ envoyé à l'actionneur de NE
- $c_2(t) = c_{WE}(t)$ envoyé à l'actionneur de WE
- $c_3(t) = c_{BS}(t)$ envoyé à l'actionneur de BS

auxquels peuvent s'ajouter un nombre illimité de signaux supplémentaires comme les signaux de contrôles angulaires ou du mode commun, en cas de couplages importants entre ces boucles et celles des modes différentiels.

Connaissant les réponses A_i des actionneurs qui ont été mesurées au chapitre précédents, les effets de filtrage optique O_i subi par les déplacements qu'ils génèrent, et les gains optiques (+ retard) α_i en *Watts/m* entre ces déplacements et le signal de frange noire, la relation de base² qui détermine le signal de frange noire s'écrit:

$$s = \alpha O \cdot \Delta L + \sum_i \alpha_i O_i \cdot A_i \cdot c_i \quad (4.7)$$

où³:

$$\alpha O \simeq \alpha_{NE} O_{NE} \simeq \alpha_{WE} O_{WE} \quad (4.8)$$

désigne la réponse optique du mode différentiel.

On en déduit:

$$\Delta L = (\alpha O)^{-1} \cdot (s - \sum_i \alpha_i O_i \cdot A_i \cdot c_i) \quad (4.9)$$

La conversion des termes O_i et A_i en filtres dans le domaine temporel discret permet alors de reconstruire successivement:

1. Les contributions des contrôles:

$$s_i(t) = \alpha_i O_i \cdot A_i \cdot c_i(t)$$

2. Le signal de frange noire qu'aurait donné un interféromètre libre:

$$s_{libre}(t) = s(t) - \sum_i s_i(t)$$

3. Le signal gravitationnel:

$$h(t) = (3km)^{-1} \cdot \Delta L = (3km)^{-1} \cdot (\alpha O)^{-1} \cdot s_{libre}(t)$$

4.3.2 L'algorithme

L'algorithme de reconstruction, développé sur la base de ce nouveau principe, est représenté schématiquement par la figure 4.4. Les données sont traitées par blocs d'une seconde (*frame*), lus dans un fichier ou transmis par l'acquisition dans le cas où l'algorithme tourne en ligne. Pour chaque *frame*, on applique ces quelques étapes:

1. Le signal de frange noire et les signaux de corrections sont lus, convertis en nombres flottants de 64 bits, et sur-échantillonnés si nécessaire à la fréquence de travail de 20 *kHz*. Afin de réduire leur forte dynamique à basse fréquence

2. d'où découlent *tous* les modèles de réponse, quelle que soit la configuration

3. par convention dans la reconstruction les signes des gains des actionneurs sont choisis pour que tous les gains optiques α_i soient positifs.

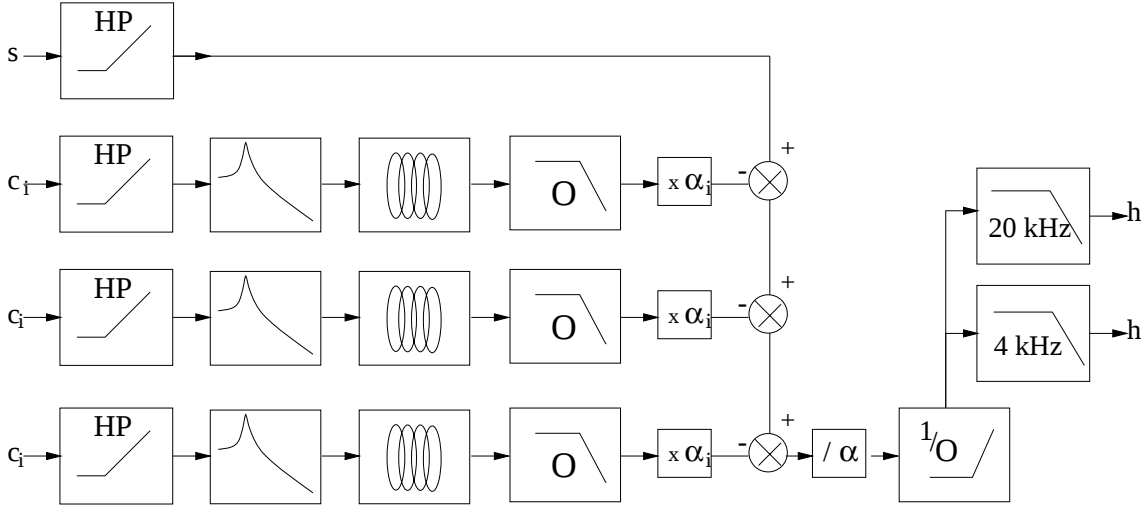


FIG. 4.4 – Schéma général de l'algorithme de reconstruction

ils sont tous filtrés passe-haut; cette précaution est indispensable pour les analyser dans le domaine fréquentiel, notamment pour le suivi des lignes de calibration. Le filtre, identique pour tous les signaux, est configurable suivant la situation mais consiste généralement en un *Butterworth* d'ordre 4, à 10 ou 20 Hz (figure 4.5).

2. Chaque signal de contrôle est spécifiquement préparé pour la recombinaison: filtres, gain et retard associés à la réponse des actionneurs A_i , filtrage optique O_i .
3. les signaux de contrôles sont convertis en *Watts* (multiplication par α_i) puis soustraits du signal de frange noire (eq. 4.9).
4. Le signal résultant $s_{libre}(t)$ est divisé par le gain optique α du mode différentiel calculé comme:

$$\alpha = \frac{1}{2}(\alpha_{NE} + \alpha_{WE}).$$

Il est en effet impossible de reconstruire séparément la contribution d'un bras ou de l'autre au mode différentiel, et les gains optiques doivent donc être supposé égaux: $\alpha \simeq \alpha_{NE} \simeq \alpha_{WE}$. Le signal est ensuite sur-échantillonné à 40 kHz pour éviter les effets de bords (utilisation de filtres discrets). Le filtre inverse O^{-1} des cavités est appliqué pour obtenir $\Delta L(t)$, puis $h(t)$.

5. Le signal $h(t)$ est enfin sous-échantillonné à 20 kHz et 4 kHz pour les différentes analyses, après filtrage anti-repliement (*Butterworth* d'ordre 6).

Il faut noter que les effets des filtres analogiques anti-repliement situés à l'entrée des ADC de détection du signal de frange noire ne sont pour l'instant pas pris en compte.

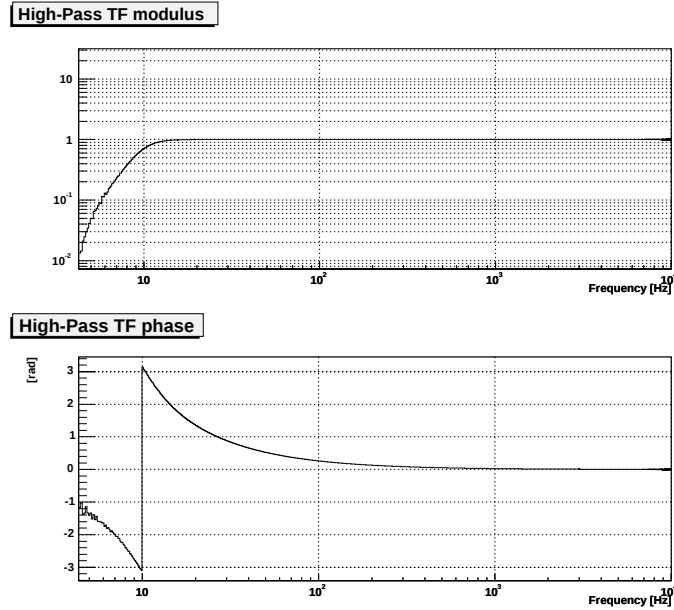


FIG. 4.5 – Fonction de transfert du filtre de réduction de la dynamique appliqué en entrée de la reconstruction (coupure à 20Hz)

4.3.3 Effet sur le signal

Pour valider cette méthode, on l'applique à la simulation d'interféromètre recombinaison "MDC6", déjà utilisée pour tester la méthode utilisant seulement le signal de frange noire. La reconstruction est configurée pour prendre en compte les trois signaux de contrôles principaux (NE, WE, BS), les gains optiques α_i correspondants étant au préalable calculés à la main à l'aide de lignes de calibration injectés sur chacun des trois miroirs (la méthode automatique de détermination des gains sera décrite ultérieurement).

Contrairement à la méthode précédente on constate à présent que, grâce à la simplicité des filtres appliqués, les problèmes d'instabilité ont disparu : le signal reconstruit reste borné même après plusieurs heures de données reconstruites.

On vérifie par ailleurs sur le spectre en fréquence des différentes contributions $s_i(t)$ comparées au signal de frange noire (figure 4.6) que :

- Au dessous du gain unité (100 Hz) des boucles de contrôles, il y a plus de signal dans les signaux de corrections que dans la frange noire.
- la contribution des corrections sur BS n'est pas négligeable et vient en deuxième position après celle de WE, même si le signal de frange noire est environ vingt fois plus sensible au mode différentiel des cavités qu'au petit Michelson.

Pour s'assurer de la validité du signal reconstruit, on injecte dans la simulation une perturbation de la métrique sous la forme de bruit blanc $h_{inj}(t)$. Là où ce signal

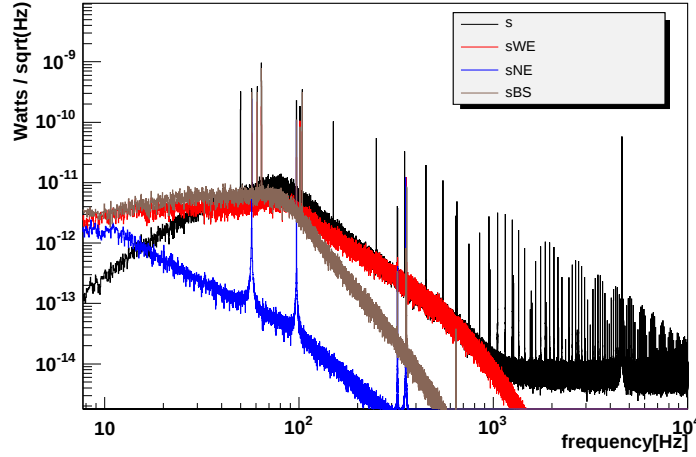


FIG. 4.6 – Spectre du signal de frange noire s et des contributions s_i des différents contrôles (Données simulées "MDC6"; les pics au-dessus de 1 kHz sont des résonances thermique simulés plus fortes que la normale).

domine le bruit du détecteur, on attend du signal reconstruit $h_{rec}(t)$ qu'il vérifie:

$$h_{rec}(t) = h_{inj}(t) \quad (4.10)$$

La qualité de la reconstruction peut donc être estimée par la fonction de transfert $T(f)$ entre h_{inj} et h_{rec} (fonction de transfert "totale" de l'ensemble interféromètre + reconstruction).

La mesure de $T(f)$ obtenue (figure 4.7) est valide au moins jusqu'à 2 kHz, au-delà le bruit électronique (blanc) domine les signaux injectés atténués par l'effet de filtrage des cavités. Le module de $T(f)$ est proche de 1 - à l'erreur de l'estimation des gains optiques près - sauf à basse fréquence; la phase de $T(f)$ présente par contre une forme inattendue. Ces effets sont dus au filtre passe-haut de réduction de la dynamique. Pour s'en assurer, on corrige $T(f)$ de la fonction de transfert de ce filtre (figure 4.8); Le module est alors complètement plat, et la phase (linéaire) se limite au retard de $\simeq 150\mu s$, soit 3 échantillons à 20 kHz. On en conclut que:

- La reconstruction inverse correctement la réponse du détecteur pour retrouver $h(t)$,
- L'effet du filtre passe-haut, caractérisé par sa fonction de transfert, devra néanmoins être pris en compte dans certaines analyses, sensible à la déformation du signal visible sur la phase de $T(f)$.

4.3.4 Effet de suppression de bruit

Pour mesurer l'effet de la reconstruction sur le bruit, on compare le spectre du signal reconstruit en l'absence de tout signal gravitationnel avec la sensibilité

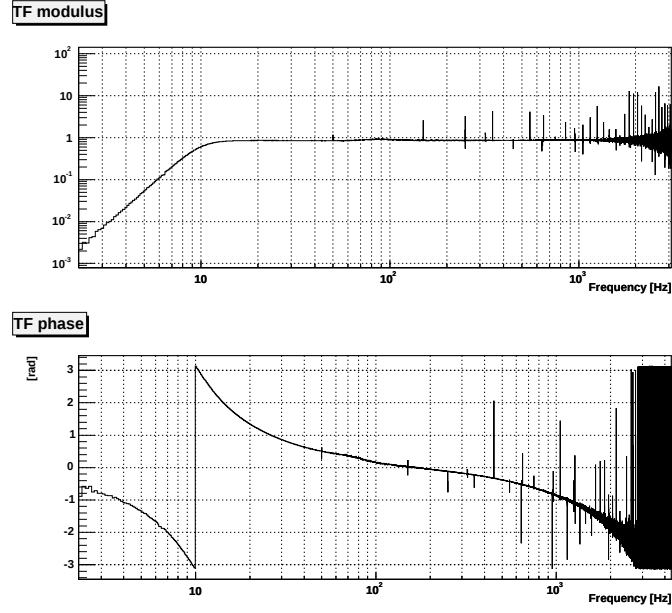


FIG. 4.7 – Fonction de transfert "totale" entre un signal injecté dans la simulation h_{inj} et le signal reconstruit h_{rec} . Au-delà de ~ 2 kHz le bruit électronique domine les signaux injectés

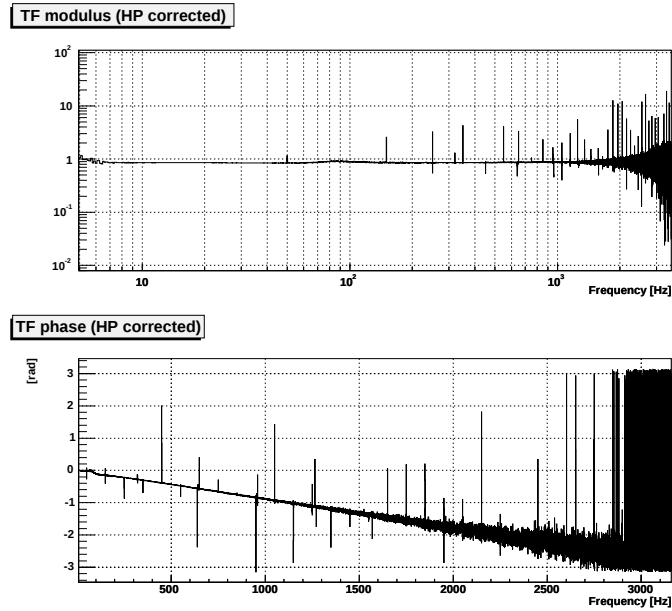


FIG. 4.8 – Fonction de transfert "totale", corrigée de la fonction de transfert du filtre passe-haut de la reconstruction.

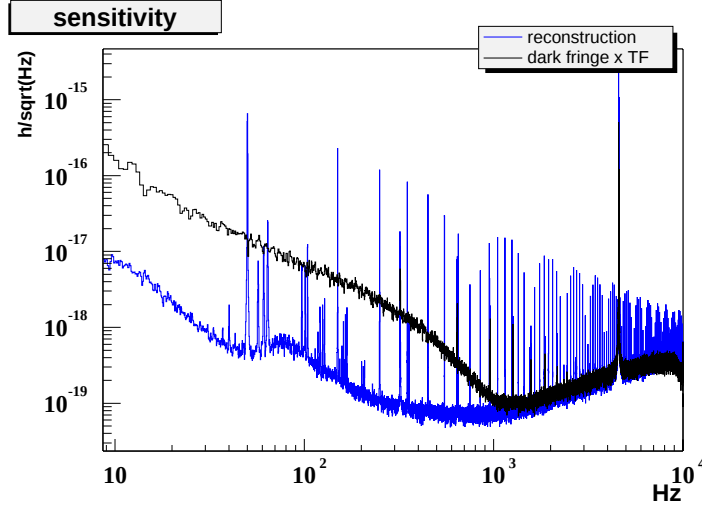


FIG. 4.9 – Spectre de bruit du signal reconstruit (en bleu) et sensibilité calibrée (signal de frange noire corrigé de la fonction de transfert du détecteur, en noir) pour la simulation du recombéné.

estimée suivant la procédure décrite au chapitre précédent : produit du spectre du signal de frange noire par la fonction de transfert du détecteur.

La comparaison (figure 4.9) suggère que la reconstruction réduit le bruit du détecteur en dessous de 1 kHz. Cet effet n'est cependant pas un artefact, il s'interprète comme un effet de suppression des bruit de contrôles lié à la nouvelle méthode de reconstruction. En effet, les contrôles de l'interféromètre n'ont pas pour seul effet d'atténuer le signal de frange noire: il peuvent aussi, lorsque les boucles utilisent des photodiodes auxiliaires, y introduire le bruit de ces photodiodes. La figure 4.6 montre en particulier que la contribution du contrôle sur WE (photodiode B7 dans la configuration C3-MDC6) est conséquente bien au-delà du gain unité des boucles, ce qui suggère l'introduction d'un bruit parasite. En utilisant les signaux de contrôle comme mesure de leurs effets, on s'attend à soustraire ces bruits de contrôle; la reconstruction peut donc améliorer la sensibilité du détecteur.

Pour valider cette hypothèse, on compare le spectre du signal reconstruit à la sensibilité d'un interféromètre simulé, identique au précédent mais dont les photodiodes auxiliaires B7 et B8 sont idéales et n'introduisent aucun bruit dans l'interféromètre. Les deux courbes (figure 4.10) se superposent presque parfaitement, montrant que le bruit du signal reconstruit est bien celui d'un interféromètre sans bruit de contrôle. Le bruit résiduel autour de 100 Hz correspond à 3% de la contribution initiale de WE restant après soustraction, et laisse donc supposer une erreur de 3% sur le gain optique correspondant. Cet effet de suppression de bruit - dans la limite de l'erreur sur les gains optiques - constitue un avantage essentiel de la nouvelle méthode. La configuration de ce test (C3-MDC6) rend

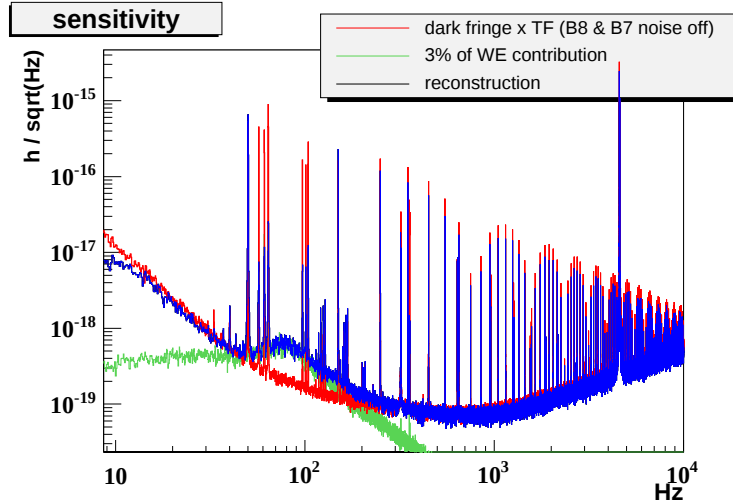


FIG. 4.10 – *Spectre de bruit du signal reconstruit (en bleu), bruit résiduel de contrôle de WE pour 3% d'erreur (en vert), et sensibilité d'un interféromètre simulé sans bruit sur les photodiodes auxiliaires (en rouge).*

cet effet particulièrement spectaculaire du fait de l'utilisation d'une photodiode auxiliaire très bruitée (B7). Il a néanmoins pu être observé à moindre échelle dans des configurations plus standards : lors de la reconstruction des *run* C4 et C5, où le bruit (de lumière diffusée) de la photodiode auxiliaire B2 utilisée pour contrôler BS se retrouvait dans le signal de frange noire, principalement autour de 100 Hz. Après la reconstruction, ce bruit est soustrait de la sensibilité (figure 4.11) et la cohérence avec le signal de contrôle de BS s'annule (figure 4.12). Notons enfin que pour bénéficier au mieux de cet avantage, on peut envisager :

- D'utiliser les signaux de lecture du courant des bobines au lieu des signaux de contrôles ; en plus des bruits de photodiodes auxiliaires on soustrait ainsi le bruit de l'électronique de pilotage des bobines des actionneurs, qui est inclus dans ces signaux. La qualité insuffisante de cette lecture ne permet actuellement pas d'utiliser cette option, mais celle-ci sera améliorée prochainement dans ce but par des filtres de compression, qui réduiront la gamme dynamique du signal mesuré afin de mieux l'adapter à celle des ADCs de lecture.
- D'étendre la soustraction à des signaux de contrôles n'intervenant pas dans la réponse de l'interféromètre pouvant néanmoins introduire du bruit, tels que les contrôles angulaires ou du mode commun.

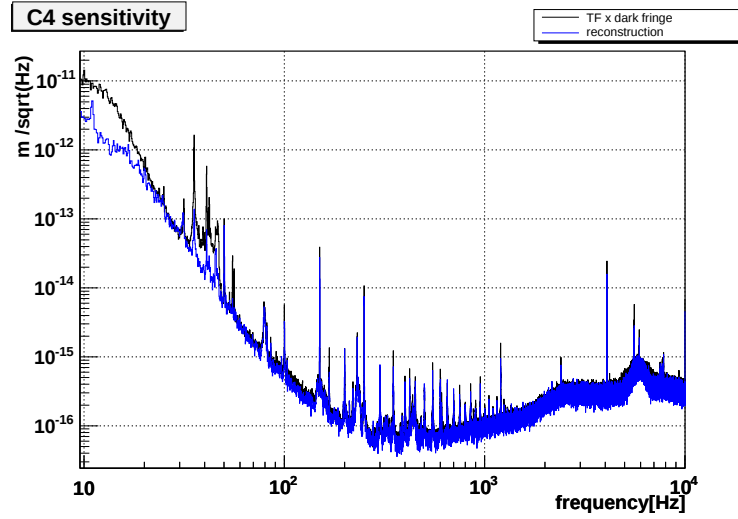


FIG. 4.11 – Spectre de bruit du signal reconstruit (en noir) et sensibilité calibrée (signal de frange noire corrigé de la fonction de transfert du détecteur, en bleu) dans le cas du run C_4 .

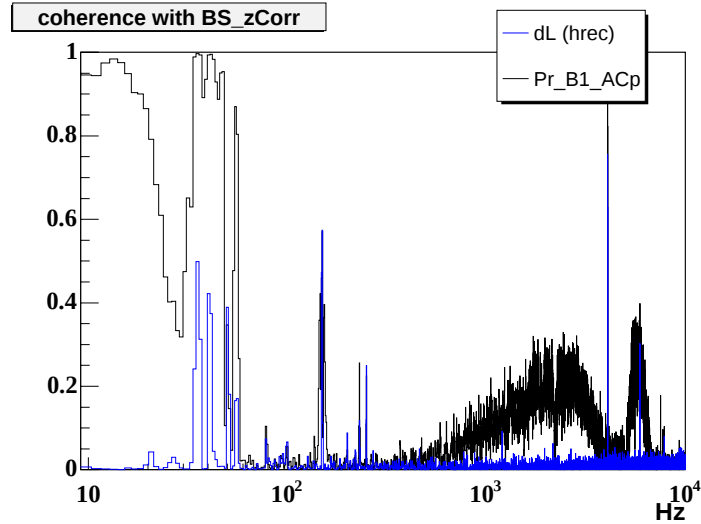


FIG. 4.12 – Cohérence du signal de contrôle de BS avec le signal de frange noire (en noir) et avec le signal reconstruit (en bleu) durant le run C_4 .

4.4 Suivi de la réponse du détecteur

La réponse de l'interféromètre ou plus particulièrement les gains optiques α_i varient au cours du temps. Pour prendre en compte ces variations comme pour faire la mesure initiale, on injecte, durant toute prise de données destinées à être reconstruites, des lignes permanentes de calibration dans chacun des signaux de contrôle. Le code de reconstruction doit alors analyser ces lignes en continu afin de :

- surveiller la qualité de la reconstruction;
- mettre à jour les gains optiques utilisés.

4.4.1 Suivi des lignes

Les lignes permanentes de calibration sont traitées par la nouvelle méthode de reconstruction comme tout bruit venant s'ajouter au signaux de contrôles: elles sont éliminées du signal lors de la soustraction des contributions des contrôles (§4.3.4). Les lignes résiduelles observées fournissent alors une estimation des erreurs relatives sur les gains optiques:

$$\epsilon = \frac{|s_{libre}(f_l)|}{|s(f_l)|} \quad (4.11)$$

pour une ligne à la fréquence f_l .

Ces erreurs s'annulent lorsque que les contributions des contrôles - celui dans lequel la ligne est injectée mais également les autres, en raison des couplages entre les boucles de contrôles - compensent exactement les lignes dans le signal de frange noire:

$$s(f_l) = \sum_i \alpha_i O_i \cdot A_i \cdot c_i(f_l). \quad (4.12)$$

Les valeurs des gains optiques doivent donc vérifier cette équation pour toutes les lignes de calibration.

Le mécanisme de suivi des lignes et de mise à jour des gains - qui se greffe sur l'algorithme statique décrit au chapitre précédent - est basé sur ce principe. Pour chaque *frame* (chaque seconde):

1. les signaux nécessaires sont convertis dans le domaine fréquentiel par transformées de Fourier rapides (*FFT*) calculées à la fin de l'étape 2 de la reconstruction :

$$s(f) = FFT[s] \quad (4.13)$$

$$s'_i(f) = FFT[O_i \cdot A_i \cdot c_i]. \quad (4.14)$$

Les FFTs sont généralement intégrées sur plus d'une seconde, auquel cas les FFTs successives se chevauchent. Le temps d'intégration est un compromis

entre l'augmentation du rapport signal sur bruit des lignes et la limitation du temps d'adaptation de la reconstruction.

2. Pour chaque ligne de calibration (à la fréquence f_i , injectée sur le signal de contrôle c_i), l'erreur est calculée par l'équation (4.11):

$$\epsilon_i(t) = \frac{|s(f_i) - \sum_i \alpha_i(t) s'_i(f_i)|}{|s(f_i)|}. \quad (4.15)$$

3. Les nouvelles valeurs des gains optiques correspondants se déduisent quant à elles de l'équation (4.12):

$$\alpha_i = \frac{1}{|s'_i(f_i)|} \cdot |s(f_i) - \sum_{j \neq i} \alpha_j s'_j(f_i)|. \quad (4.16)$$

Le nombre de gains et de lignes étant *a priori* illimité, ce système est résolu de façon perturbative, en négligeant les variations des autres gains optiques α_j :

$$\alpha_i(t+1) = \frac{1}{|s'_i(f_i)|} \cdot |s(f_i) - \sum_{j \neq i} \alpha_j(t) s'_j(f_i)|. \quad (4.17)$$

4. Si nécessaire pour compenser le bruit de mesure des lignes, les nouvelles valeurs des gains et erreurs sont moyennées avec les N précédentes *via* des moyenne glissantes. Afin d'éviter des sauts dans le signal reconstruit, les gains de la reconstruction sont mis à jour progressivement, une interpolation linéaire faisant passer le gain de l'ancienne à la nouvelle valeur sur la durée d'une *frame* (une seconde). Les gains utilisés à l'instant t sont donc moyennés entre $t - (N + 1) \times 1s$ et t . Ce retard à l'utilisation des gains n'est pas compensé - les variations des gains à des échelles de temps inférieures à $(N + 1) \times 1s$ étant de toute façon négligées - pour permettre de fournir les données "en-ligne".

Dans le cas où plusieurs lignes sont injectées dans le même canal, une seule d'entre elle est utilisée pour la mise à jour, les autres servant à estimer l'erreur uniquement.

4.4.2 Cas particuliers et vetos

Suivre la réponse du détecteur implique également, au-delà du fonctionnement régulier décrit ci-dessus, de détecter et gérer correctement les cas particuliers pour lesquels la reconstruction ne peut être appliquée:

- pertes de l'asservissement de Virgo
- données manquantes
- absence ou insuffisance des lignes de calibration

Les pertes d'asservissement sont détectées à l'entrée de la reconstruction par une combinaison de signaux du détecteur, garantissant que l'interféromètre est asservi dans sa configuration finale qui utilise le mode-cleaner de sortie. Les données ne sont alors pas transmises à la suite du code et l'on est donc ramené au cas des données manquantes. A la reprise de l'asservissement, une durée minimum de cinq secondes est attendu avant de transmettre à nouveau les données.

Le cas des données manquantes se produit fréquemment lorsque la reconstruction tourne en ligne. L'ensemble du code est protégé contre ce cas par une ré-initialisation systématique de tous les éléments "à mémoire" tels que les mémoires tampons des FFTs, les valeurs des gains optiques ou les moyennes glissantes, et ce dès qu'une discontinuité temporelle est détectée dans les données.

La présence des lignes de calibration est vérifiée juste avant leur analyse. Leur rapport signal sur bruit (SNR) est estimée sur la FFT du signal de frange noire, en comparant l'amplitude de la ligne au plancher moyen estimé sur 2 points voisins de la ligne:

$$SNR = \frac{\sqrt{2} \cdot s(f)}{\sqrt{s^2(f - 4df) + s^2(f + 4df)}} \quad (4.18)$$

où df est la résolution de la FFT.

L'analyse des lignes n'a lieu que si $SNR > 3$. A partir de cette valeur, la distribution d'une ligne dans du bruit se rapproche fortement d'une gaussienne, et le SNR peut être augmenté d'un facteur $\sqrt{N}$ en moyennant les gains calculés sur N valeurs. En dessous les lignes sont noyées dans le bruit et ne peuvent être extraites. Lorsque les lignes reviennent, une durée excédant le temps de calcul d'une FFT est attendue avant de reprendre l'analyse des lignes.

Le veto final de la reconstruction, transmis au reste de l'analyse pour garantir la validité des données reconstruites, doit tenir compte non seulement des cas précédents mais aussi des périodes de transition lors du retour à la normale. Ce veto est donc maintenu lorsque la plus grande erreur ϵ , de toutes les estimations fournies par les différentes lignes de calibration, dépasse un seuil donné - dépendant du contexte mais généralement fixé à $\epsilon_{max} = 30\%$.

En résumé, des données étiquetées comme analysables par la reconstruction impliquent:

- la continuité du flux de données en amont de celle-ci,
- l'interféromètre asservi depuis plus de cinq secondes,
- la présence de lignes de calibration suffisamment fortes,
- une erreur de reconstruction suffisamment faible.

4.4.3 Application aux prises de données C4 et C5

Cette extension adaptative du code de reconstruction a essentiellement été testée sur les données des *runs* C4 et C5, pour permettre la reconstruction com-

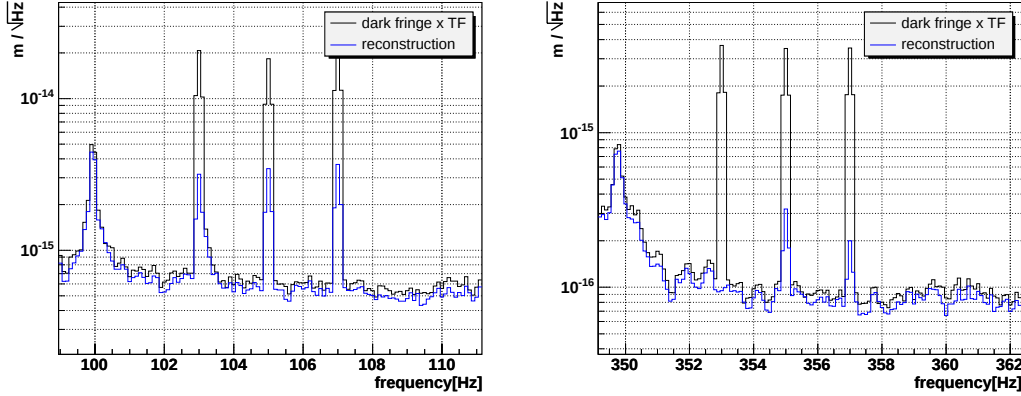


FIG. 4.13 – Les six lignes de calibration de $C4$ dans le dans le spectre du signal reconstruit (en bleu) et dans celui du signal de frange noire, corrigé de la fonction de calibration (en noir). Les lignes au-dessus de 350 Hz sont utilisés pour mettre à jour les gains optiques.

plète des jeux de données. Durant ces deux *runs*, deux lignes permanentes étaient injectées par miroir (NE, WE et BS), les unes autour de 350 Hz servant à mettre à jour les gains et les autres vers 100 Hz - où la réponse est 2 à 3 fois plus faible en module qu'à 350 Hz (figure 3.20) - ne servant qu'aux calculs d'erreurs.

Les deux triplets de lignes sont visibles (figure 4.13) dans le spectre du signal reconstruit comme dans celui du signal de frange noire, corrigé de la fonction de calibration. L'effet de suppression des lignes sur lequel repose l'estimation d'erreur est bien visible et laisse prévoir une erreur maximale de l'ordre de 10% à 15%, qui correspond aux valeurs prises par l'erreur ϵ_{max} annoncée par la reconstruction.

Le bon fonctionnement du calcul des gains à partir des lignes se vérifie en initialisant la reconstruction avec des valeurs de gains fausses d'un facteur 10. La figure 4.14 montre les gains calculés durant les premières secondes de reconstruction. Ils retrouvent leurs vraies valeurs ($\pm \epsilon_{max}$) après 15 s, c'est à dire dès qu'un premier calcul de FFT est disponible. On vérifie ainsi également la robustesse de la méthode.

L'estimation du rapport signal sur bruit des lignes est testée dès le début de la reconstruction puisque celle-ci est lancée avant les lignes de calibration; la reconstruction attend donc de détecter les lignes pour entamer la mise à jour des gains (figure 4.15).

Le mécanisme de mise à jour des gains à partir des valeurs calculées est quand à lui validé par l'absence de saut brutaux en début de *frame* dans le signal reconstruit. De tels sauts ou *glitches* peuvent par contre être observés dans des périodes de veto, lors de l'apparition ou la perte des lignes de calibration (figure 4.16), et sont détectés par la recherche de coalescences binaires ou d'événements impulsifs.

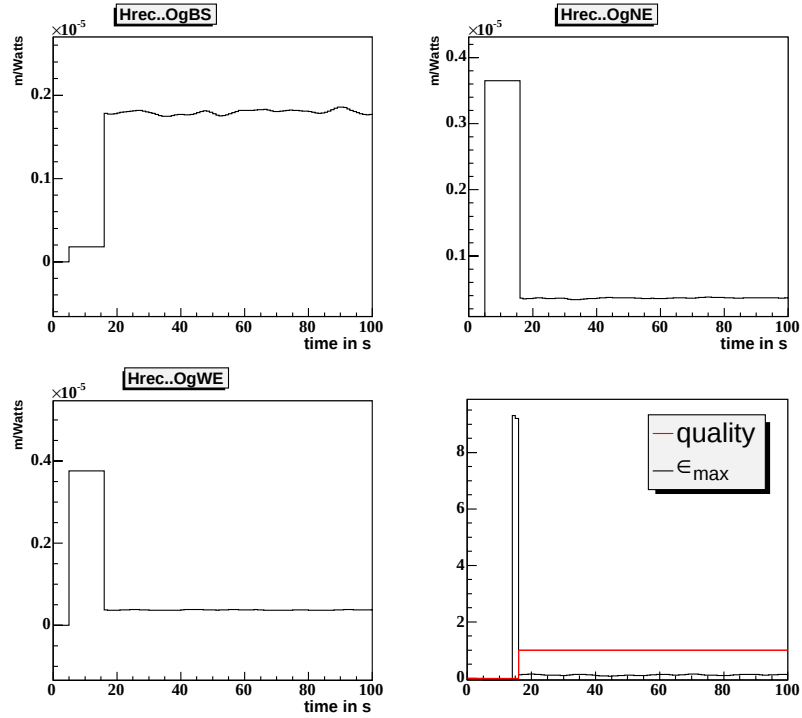


FIG. 4.14 – Lancement de la reconstruction (C_4) avec des gains optiques initiaux faux d'un facteur 10. De gauche à droite puis de haut en bas: Les valeurs des gains (α_i^{-1} , en m/Watts) mesurés par la reconstruction pour BS, NE et WE, l'erreur maximum et l'indicateur de qualité de la reconstruction.

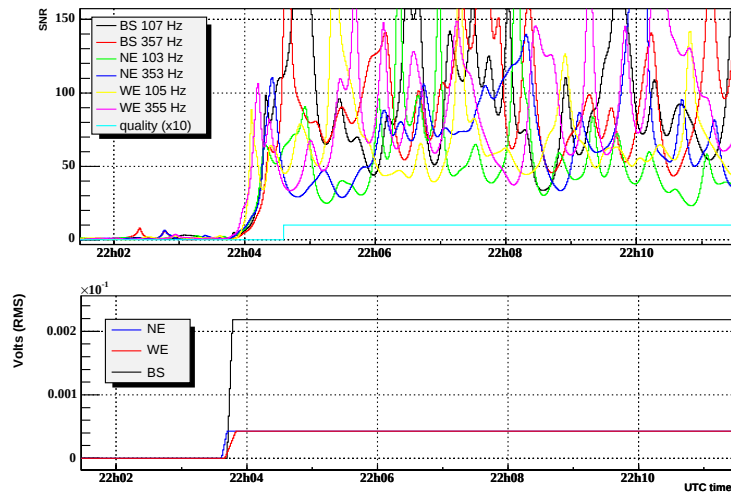


FIG. 4.15 – Mise en route des lignes de calibration durant le run C_4 . En haut: estimations du SNR des lignes par la reconstruction et indicateur de qualité. En bas: RMS signaux d'injection de lignes sur les trois miroirs.

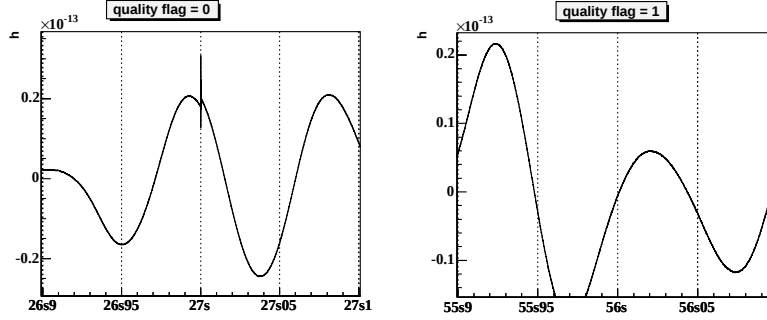


FIG. 4.16 – *Signal reconstruit au cours du temps, autour d'un changement de frame (C5). A gauche durant une période de veto (lignes absentes), à droite en période normale.*

Pour la reconstruction complète des données, le choix du temps de calcul des FFTs T et du nombre de moyenne N pour les gains et les erreurs repose en principe sur un compromis entre la nécessité éventuelle de suivre des variations rapides des gains et les rapports signal sur bruit des lignes de calibration. Ces derniers variant beaucoup au cours des *runs*, les paramètres sont choisis assez largement pour reconstruire un maximum de données, quitte à ne prendre en compte que les dérives à long terme des gains optiques: $T = 40s$ et $N = 100$ - correspondant globalement à 100s de moyenne/integration (cinq FFTs) - pour C4, $N = 10$ pour C5.

Il est alors possible de reconstruire l'essentiel des *run* C4 et C5. Les figures 4.17 et 4.18 représentent l'erreur maximum ϵ_{max} enregistrée au cours du temps pour chaque (*run*). Les bandes vertes indique les zones sans donnés, le plus souvent lorsque l'interféromètre n'est pas asservi; les valeurs négatives (-1) de l'erreur indique que certaines lignes sont trop faible pour être analysé. On observe ainsi, surtout pour C4, de longues périodes où les lignes sont à la limite de détectabilité, l'erreur maximum étant alors disponible par intermittence.

4.5 Soustraction des harmoniques de 50 Hz

Une fois les données reconstruites, une dernière étape subsiste pour fournir un signal $h(t)$ propre à l'analyse: la soustraction des harmoniques de 50 Hz (fréquence du secteur), ces harmoniques constituant une important contribution au bruit et surtout la principale composante non-gaussienne.

4.5.1 Principe

La méthode implémentée s'inspire d'une étude [33] sur la suppression des harmoniques de 50 Hz dans le bruit électronique des photodiodes du banc de dé-

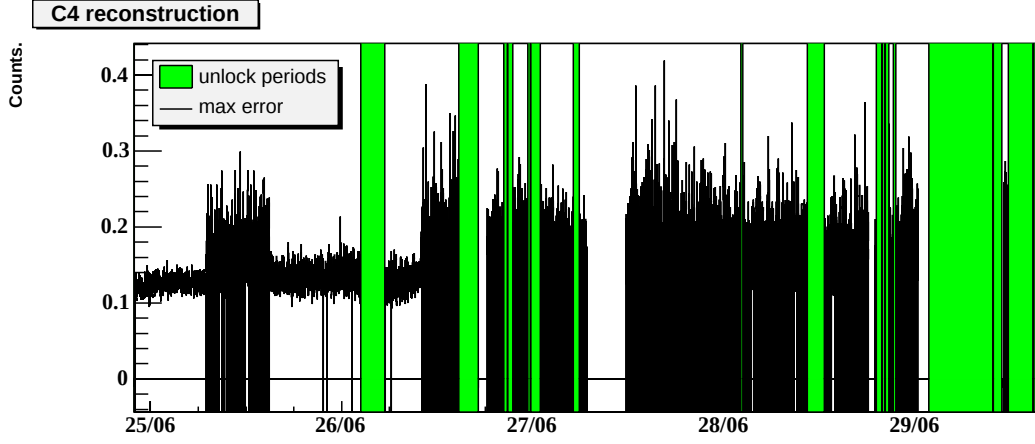


FIG. 4.17 – Erreur maximum de reconstruction $\epsilon_{\max}$ enregistrée au cours du run C4. Les bandes vertes indique les zones sans donnés (interféromètre non asservi); les valeurs négatives (-1) de l'erreur indique que certaines lignes sont trop faible pour être analysé.

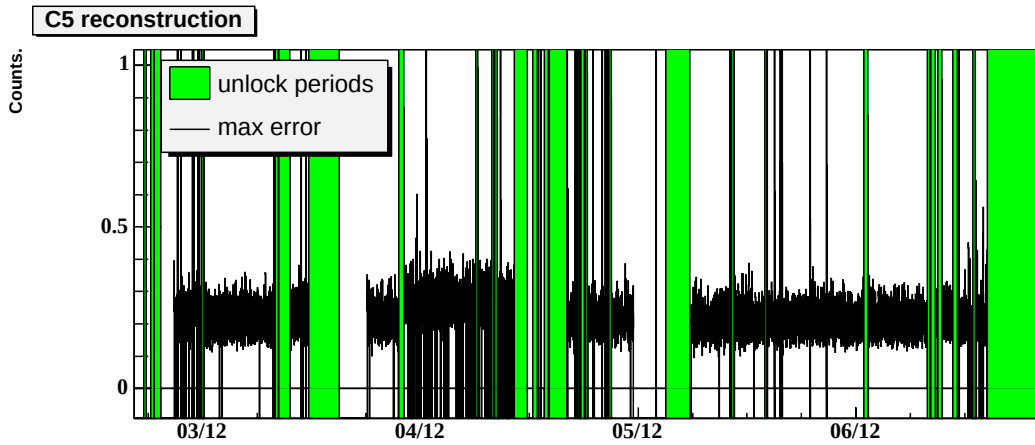


FIG. 4.18 – Erreur maximum de reconstruction $\epsilon_{\max}$ enregistrée au cours du run C5. (cf figure 4.17)

tection. Pour séparer ces harmoniques du reste du signal, on utilise une mesure indépendante du signal à 50 Hz de l'alimentation $s_{50}(t)$, fournie par l'acquisition des données de Virgo. On en déduit la contribution des harmoniques $h_{50}(t)$ que l'on peut alors soustraire pour fournir à l'analyse un signal "nettoyé":

$$h_{no50}(t) = h(t) - h_{50}(t) \quad (4.19)$$

On peut imaginer de déduire directement la contribution $h_{50}(t)$ du signal de mesure $s_{50}(t)$, par le biais d'une fonction de transfert convertie en filtres, comme on l'a fait précédemment pour les contributions des contrôles au signal de frange noire. Au moins deux raisons majeures s'y opposent:

- Le signal de mesure $s_{50}(t)$ contient du bruit blanc de lecture, qui sera introduit dans le signal reconstruit, et sera visible si le SNR des harmoniques dans $s_{50}(t)$ excède leur SNR dans le signal à nettoyer $h(t)$.
- Le processus de contamination des données de Virgo par ce signal est fortement non-linéaire, comme le prouve la création d'harmoniques en plus du mode fondamental, et ne peut donc être correctement reproduit par une fonction de transfert convertie en filtres linéaires.

On contourne cette difficultés en profitant de la simplicité du signal $h_{50}(t)$, qui se modélise comme une somme de signaux harmoniques:

$$h_{50}(t) = \sum_{n=1}^N [c_n(t) \cos(2n\pi f(t) \cdot t) + s_n(t) \sin(2n\pi f(t) \cdot t)] \quad (4.20)$$

où la fréquence fondamentale $f(t)$ est identique à la fréquence du signal d'alimentation mesuré par $s_{50}(t)$, tandis que les coefficient de couplages $c_n(t)$ et $s_n(t)$ se mesurent dans le signal "contaminé" $h(t)$.

Comme pour les étapes précédentes de la reconstruction, les données sont traitées par *frame* (1s):

1. Le signal de mesure $s_{50}(t)$ est démodulé à la fréquence de référence $f_r = 50Hz$ avec une durée d'intégration T_{freq} d'une seconde à priori⁴. La phase du signal $\phi_{50}(t)$ issue de la démodulation est comparé à la phase précédente $\phi_{50}(t - 1s)$, la dérive de phase permettant de corriger la fréquence:

$$\delta\phi = \phi_{50}(t) - \phi_{50}(t - 1s) - 2\pi f_r \pmod{2\pi} \quad (4.21)$$

$$f(t) = f_r + \frac{\delta\phi}{2\pi} \quad (4.22)$$

On obtient ainsi un suivi à 1 Hz de la fréquence du signal d'alimentation.

4. Dans le cas d'un mauvais SNR du signal de mesure on peut choisir ($T > 1s$) en ré-utilisant des données des *frames* précédentes

2. Les harmoniques sont reconstruites en phase ($\cos(2n\pi f(t) \cdot t + \theta)$) et quadrature ($\sin(2n\pi f(t) \cdot t + \theta)$). La phase θ assure la continuité avec les harmoniques reconstruites de la *frame* précédente.
3. Le signal reconstruit $h(t)$ doit être à son tour démodulé par les harmoniques reconstruites pour estimer les coefficients de couplages:

$$c_n = \frac{1}{T} \int_{t-T}^t h(t) \cos(2n\pi f(t) \cdot t) dt \quad (4.23)$$

$$s_n = \frac{1}{T} \int_{t-T}^t h(t) \sin(2n\pi f(t) \cdot t) dt \quad (4.24)$$

avec une durée d'intégration $T \gg \frac{1}{50\text{Hz}}$. Les intégrales glissantes sont dans la pratique remplacées par un filtre passe-bas du second ordre ($Q = \frac{\sqrt{2}}{2}$) à la fréquence $f = (2\pi T)^{-1}$. La contribution $h_{50}(t)$ au signal $h(t)$ est alors reconstruite suivant (4.20) puis soustraite.

La durée d'intégration T constitue un paramètre essentiel de cette méthode. elle doit être:

- suffisamment faible pour permettre un suivi précis des coefficients de couplages;
- suffisamment élevée pour ne pas affecter le reste des données et en particulier ne pas risquer de soustraire un éventuel signal d'onde gravitationnelle.

4.5.2 Test et résultats

Simulation

Pour s'assurer du bon fonctionnement de l'algorithme décrit plus haut, on l'applique à des données simulées incluant des harmoniques de 50 Hz et un canal de mesure $s_{50}(t)$. Un tel test a été réalisé pour la première fois lors du MDC4. On simule directement des données reconstruites à l'aide d'un processus ARMA reproduisant approximativement la sensibilité nominale de Virgo, auquel s'ajoute quelques résonances thermiques, et les harmoniques de 50 Hz simulées par la procédure suivante:

- La fréquence fondamentale $f(t)$ est donnée par:

$$f(t) = 50\text{Hz} + \delta f(t) \quad (4.25)$$

où $\delta f(t)$ est un bruit blanc gaussien de largeur σ généré à 20kHz auquel on applique un filtre passe-bas coupant à 1Hz . On en déduit la ligne fondamentale à 50 Hz ($\cos(f(t))$). La largeur des variations de fréquence simulées ainsi est prise à $\sigma = 0.1\text{ Hz}$, correspondant approximativement aux fluctuations visibles dans le canal de mesure $s_{50}(t)$ de Virgo.

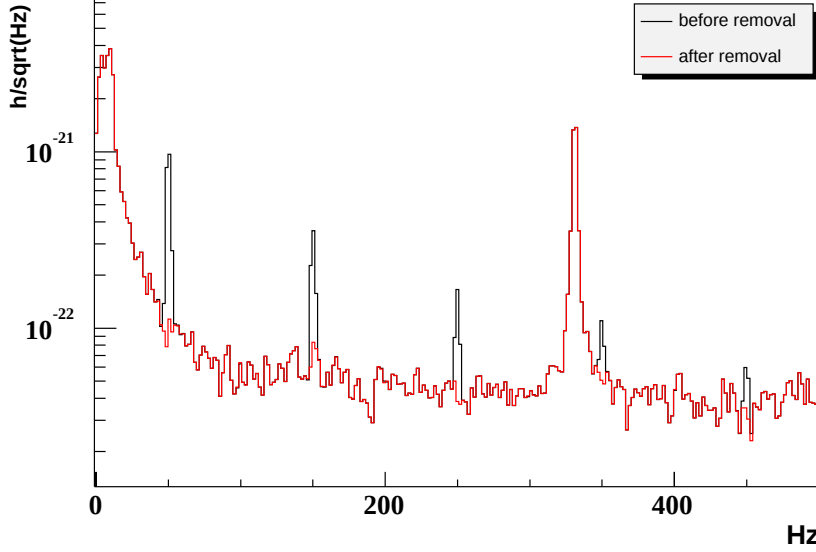


FIG. 4.19 – *MDC4*: Spectre des données simulées avant (en noir, en-dessous) et après (en rouge, au-dessus) la soustraction des harmoniques de 50 Hz - calculés sur $\simeq 5s$ de données prises après 45s de soustraction.

- Pour simuler le processus non-linéaire à l'origine des harmoniques d'ordres supérieures, on simule la lecture de ce signal par un ADC saturé. Le signal ainsi obtenu est ensuite mélangé aux données, avec un couplage constant.
- Pour simuler le canal de mesure $s_{50}(t)$ on utilise la ligne fondamentale, à laquelle on ajoute du bruit blanc gaussien (bruit de lecture).

Le spectre des données avant la soustraction apparaît sur la figure 4.19. On note que seul des harmoniques impaires sont générés par la simulation, et qu'ils sont visibles - sur des spectres intégrés quelques secondes - jusqu'à environ 450 Hz.

La soustraction est donc appliquée pour les 10 premières harmoniques de 50 Hz. Le couplage des harmoniques aux données étant constant dans la simulation, la durée d'intégration T peut-être prise aussi grande qu'on le souhaite. En pratique pour éviter un temps d'initialisation trop important (temps de chauffe des filtres) on la limite à 15 s. Après soustraction on vérifie que les lignes sont supprimées tandis que le "plancher" de bruit n'est lui pas affecté (figure 4.19).

On note néanmoins que ce n'est le cas que après 45 s de données: après 15 s les lignes ne sont que partiellement soustraites (figure 4.20), ce qui est le principal inconvénient d'utiliser des filtres récurrents au lieu des intégrales (4.23).

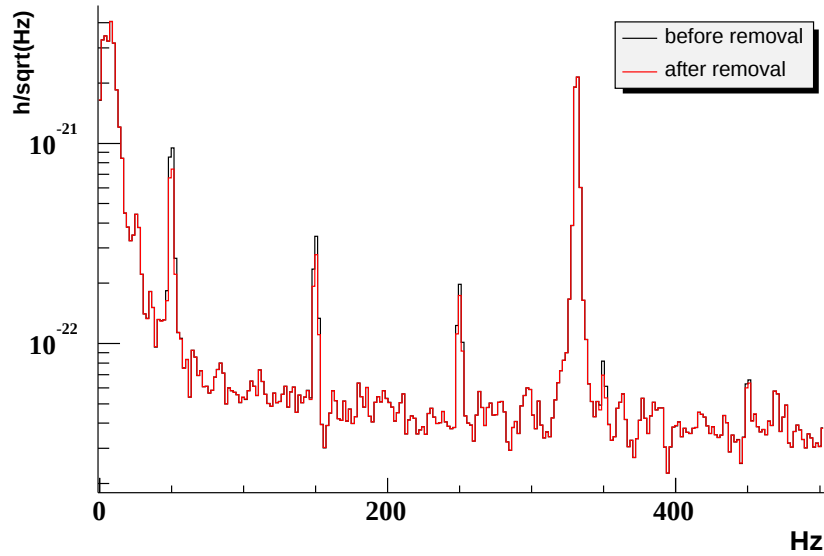


FIG. 4.20 – *MDC4*: Spectres de la figure 4.19, calculés après seulement 15s de soustraction.

Données réelles

L'algorithme se comportant comme attendu en simulation, on l'applique à des vrais données issue de *run* techniques, confrontant ainsi le modèle sur lequel il repose à la réalité.

L'exemple présenté ici est celui des données reconstruites de C5. Comme pour la simulation, les harmoniques impaires y sont dominantes et visibles jusqu' à environ 350Hz, bien que dans ce cas-ci des harmoniques paires soient également présentes - plus faiblement - notamment à 200 et 300 Hz.

La configuration de l'algorithme est identique à celle testé en simulation: soustraction des 10 premières harmoniques, avec une durée d'intégration de 15 secondes. La figure 4.21 présente le spectre des donnés reconstruites avant et après la soustraction. On vérifie que les harmoniques sont soit complètement supprimées comme en simulation, ou au moins ramenées plus bas que les autres lignes avoisinantes; l'harmonique le moins bien soustrait est celui à 100 Hz, réduit d'un facteur 1,5. Un tel cas de faible efficacité peut généralement s'expliquer par:

- une variation du couplage sur des échelles de temps inférieures à 15s, d'où découlerait une mauvaise estimation de ce couplage, donc une soustraction soit trop forte soit trop faible et dans les deux cas une ligne résiduelle
- une structure plus complexe du bruit autour de la fréquence concernée, ajoutant d'autres contributions à la ligne que l'harmonique du 50Hz de l'alimentation.

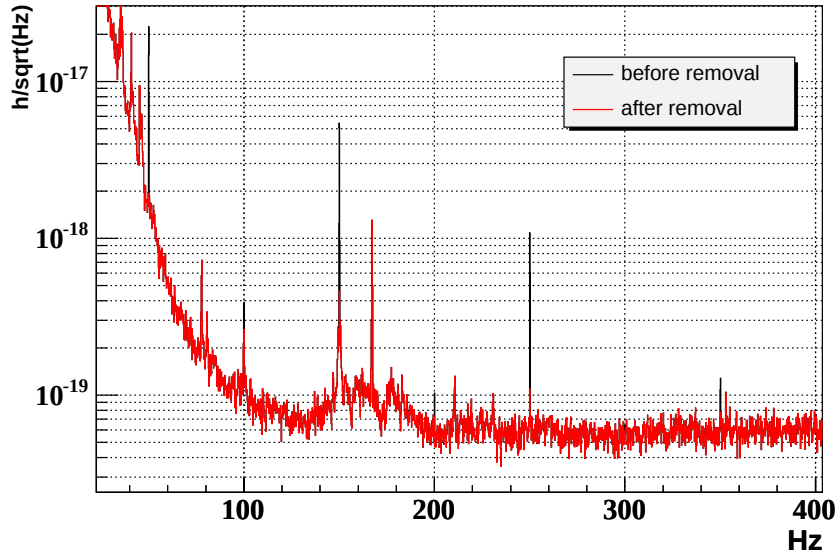


FIG. 4.21 – C5: Spectre des données reconstruites avant (en noir, en-dessous) et après (en rouge, au-dessus) la soustraction des harmoniques de 50 Hz avec une durée d'intégration $T = 15s$.

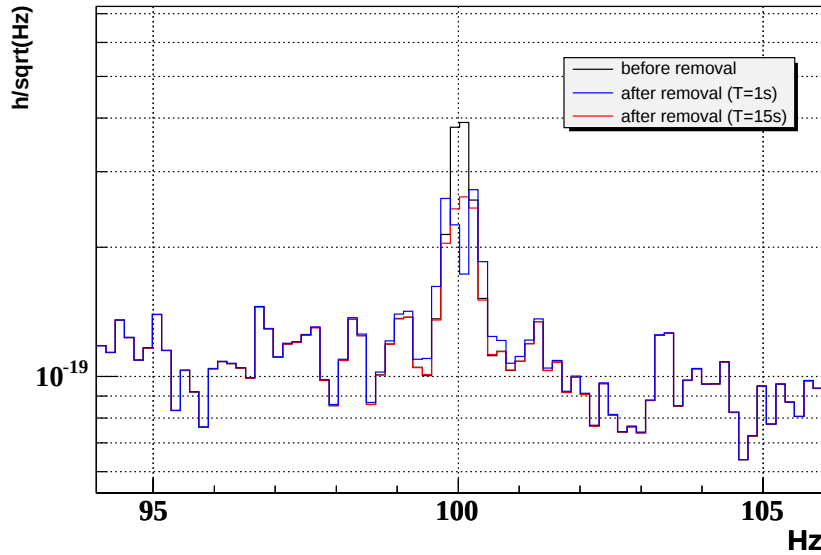


FIG. 4.22 – C5: Spectre des données reconstruites autour de 100 Hz avant la soustraction des harmoniques de 50 Hz (en noir), après soustraction avec une durée d'intégration $T = 15s$ (en rouge), et après soustraction avec une durée d'intégration $T = 1s$ (en bleu).

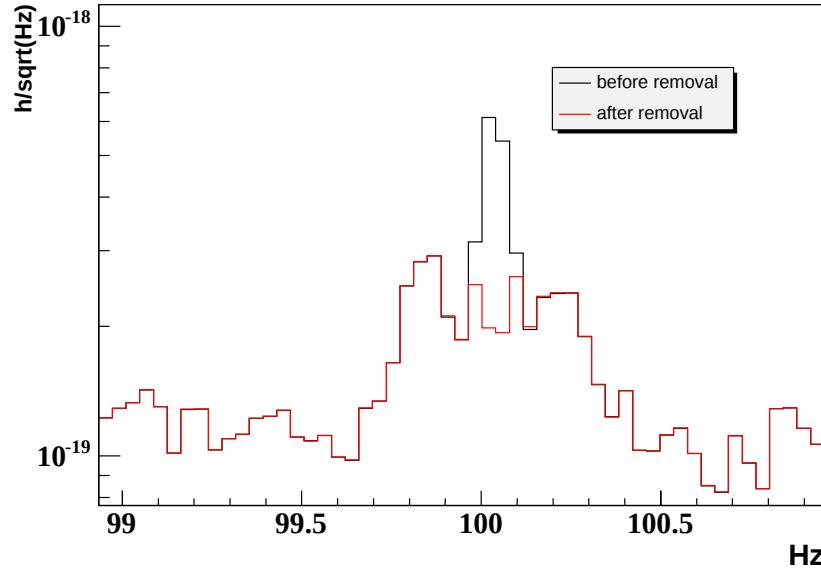


FIG. 4.23 – C5: Spectre de la figure 4.21 autour de 100 Hz, calculés sur des FFTs de $\sim 25s$.

La première possibilité est éliminée en réduisant la durée d'intégration à 1 s au lieu de 15 s - en-dessous d'1 s la méthode n'a plus de sens, la fréquence $f(t)$ étant suivie à 1 Hz. On constate l'absence d'amélioration significative et une plus grande perturbation des données aux fréquences voisines (figure 4.22).

La deuxième est par contre confirmée en augmentant la résolution en fréquence des spectres (FFTs de $\sim 25s$, figure 4.23): on vérifie que la ligne à 100 Hz résulte de la confusion de trois pics, dont l'un seulement est soustrait.

Le bilan de ces tests est donc que l'algorithme peut-être appliqué aux données reconstruites, en conservant la durée d'intégration de 15 s.

Justification de l'utilisation d'un canal de mesure

Le fait que cette méthode parvienne à soustraire les harmoniques en simulation comme sur des données reconstruites réelles ne la justifie pas nécessairement: le signal étant quasi-périodique on peut imaginer soustraire les harmoniques en leur attribuant une fréquence constante et donc sans avoir à analyser constamment le signal de mesure $s_{50}(t)$.

Pour valider ce choix qui complique la méthode, on teste l'efficacité d'une soustraction restreinte à l'étape 3 de l'algorithme, en utilisant comme harmoniques reconstruites des sinusoïdes pures aux fréquence multiples de 50 Hz. L'amplitude de la ligne la mieux soustraites - le fondamental - ne diminue alors de pas plus de 20% (figure 4.24).

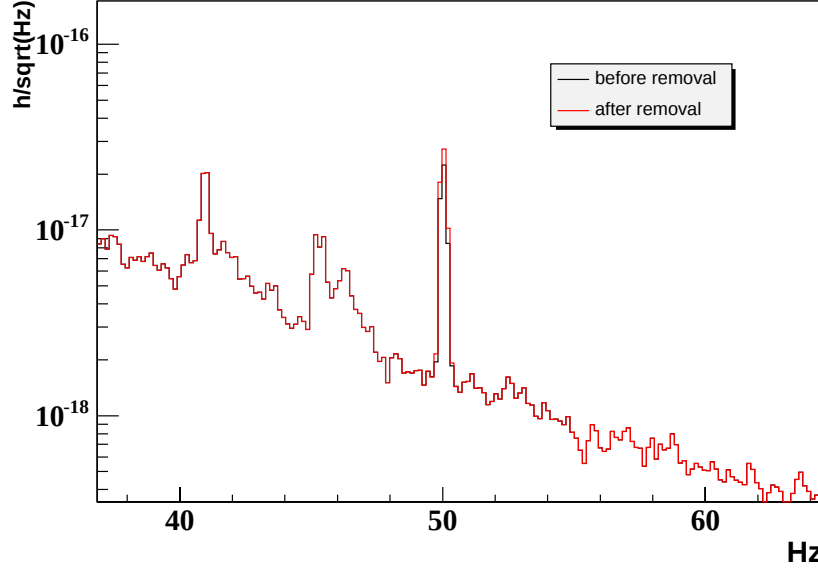


FIG. 4.24 – C5: Soustraction des harmoniques de 50 Hz avec une durée d'intégration $T = 15s$, sans utiliser le canal de mesure du 50Hz.

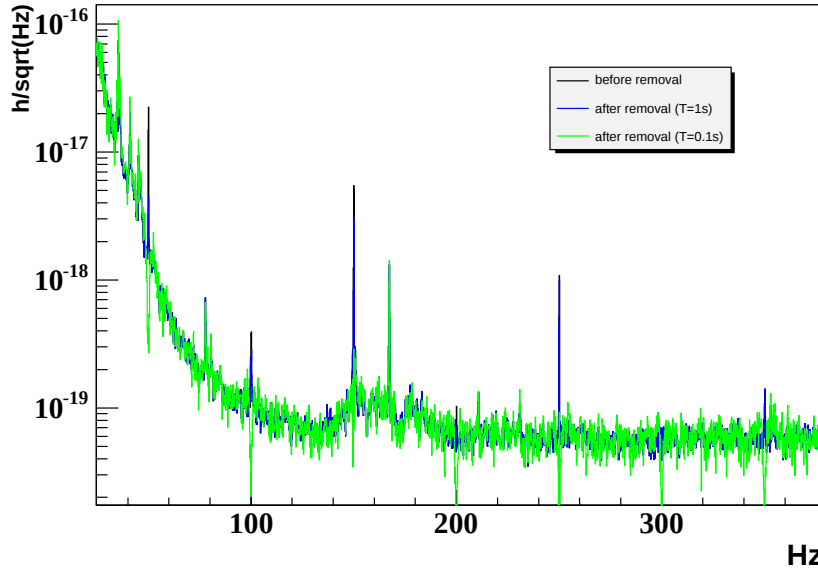


FIG. 4.25 – Soustraction des harmoniques de 50 Hz avec une durée d'intégration $T = 1s$ (en bleu), et $T = 0.1s$ (en vert), sans utiliser le canal de mesure du 50Hz.

En réduisant la durée d'intégration on peut rendre la soustraction moins "sélective". La figure 4.25 montre les résultats obtenus pour toutes les harmoniques avec une durée d'intégration d'une seconde. Si l'efficacité s'améliore elle ne se compare toujours pas à celle obtenu avec le suivi du canal de mesure. En réduisant d'un ordre de grandeur encore cette durée ($T = 0.1s$, même figure), on supprime entièrement les lignes au prix d'une dégradation totale des données. Ces derniers tests confirment que seul le suivi du canal de mesure permet de soustraire les harmoniques de 50 Hz distinctement du reste du signal.

Chapitre 5

Recherche de coalescences binaires: Méthode

Les coalescences binaires ont été brièvement décrites au premier chapitre: elles marquent la fin de vie d'un système binaire au terme de l'épuisement complet de son énergie orbitale, dissipée par rayonnement gravitationnel.

Ce chapitre décrit le principe et la méthode de la recherche de signaux de coalescences binaires dans les données de Virgo.

Les principales caractéristiques [34] de ces signaux sont d'abord abordées telles qu'elle sont prédites par la relativité générale.

La méthode générale [35] d'analyse du signal expérimental, utilisée dans tous les détecteurs interférométriques, est le filtrage adapté. Cette méthode et les problèmes qu'elle pose en terme de moyens de calculs sont abordés ensuite.

Enfin, viendra la description de son implémentation dans Virgo sous une forme alternative - la recherche multi-bande - pour justement contourner ces problèmes de moyens de calculs.

5.1 Le signal de coalescences binaires

5.1.1 Forme du signal détecté

Du point de vue du signal à détecter, ces sources englobent en fait trois phases :

- La phase spirale, qui précède la coalescence proprement dite, durant laquelle la progression régulière de la fréquence et de l'amplitude du signal accélère violemment, amenant le signal à balayer rapidement la bande passante de Virgo jusqu'à quelques kHz. Le système spiralant est relativement simple (deux points matériels), et ce signal - le *chirp* - est, comme on va le voir, prédictible avec précision par les lois de la gravitation.
- La fusion du système en un seul objet, le plus souvent en trou noir. Cette phase commence dès que la structure interne des deux astres entre en jeu :

d'abord à travers les forces de marées qui s'exercent entre eux, puis lors de la collision. Elle est donc difficilement prédictible, et dépend de l'équation d'état de la matière nucléaire dans le cas où il y a une étoile à neutrons parmi les astres.

- Le retour à l'équilibre du trou noir souvent formé, qui se fait en partie par rayonnement d'énergie sous forme d'ondes gravitationnelles suivant les modes propres de vibration du trou noir. Le signal s'approche donc d'une superposition de sinusoïdes amorties.

La recherche de coalescences binaires s'appuie essentiellement sur la première phase, tirant parti de la connaissance poussée du *chirp* qu'offrent les équations de la relativité générale.

Celui-ci est obtenu en considérant des effets relativistes faibles : on utilise un développement dit "post-newtonien" (PN) en $\frac{v}{c}$ et en $\sqrt{\frac{M}{r}}$. L'excentricité de l'orbite est également négligée, on considère en effet qu'elle doit être quasi-circulaire à l'approche de la coalescence.

La formule du quadrupôle (eq.1.22) donne alors accès à la dissipation d'énergie par rayonnement gravitationnel, d'où dérive l'évolution décroissante du rayon orbitale et la forme d'onde émise au fur et à mesure de cette décroissance:

$$\begin{aligned} h_+(t) &= h_c(t) \frac{1 + \cos^2 i}{2} \cos(\phi(t) + \phi_0) \\ h_\times(t) &= h_c(t) \cos i \sin(\phi(t) + \phi_0) \end{aligned}$$

essentiellement décrite par son amplitude croissante $h(t)$ et sa fréquence instantanée $F(t) = \frac{\dot{\phi}(t)}{2\pi}$, double de la fréquence orbitale et donc croissante aussi. L'angle i décrit l'inclinaison entre l'axe de rotation du système binaire et la direction de l'observateur. Il s'agit enfin d'une onde de polarisation circulaire, caractérisée par sa phase initiale ϕ_0 .

Le signal vu par l'interféromètre s'écrit alors comme une combinaison de h_+ et $h_\times$:

$$h(t) = F_+(t)h_+(t) + F_\times(t)h_\times(t)$$

où F_+ et $F_\times$, les fonctions d'antennes, représentent la réponse du détecteur en fonction de la direction de la source et de la polarisation de l'onde incidente.

La dépendance en temps des fonctions d'antennes étant négligeable sur la durée du signal pour les détecteurs terrestres, on peut ré-écrire h sous la forme:

$$h(t) = \Theta h_c(t) \cos(\phi(t) + \phi_0 - \Phi)$$

avec

$$\begin{aligned} \Theta &= \sqrt{\left(F_+ \frac{1 + \cos^2 i}{2}\right)^2 + (F_\times \cos i)^2} \\ \Phi &= \arctan\left(\frac{2 \cos i F_\times}{1 + \cos^2 i F_\times}\right) \end{aligned}$$

expression qui montre que - grâce à la polarisation circulaire de l'onde - l'influence de la direction et la polarisation de la source se limite à un facteur d'atténuation et un déphasage du signal, mais n'affecte pas sa forme qui est complètement déterminée par $h_c(t)$ et $\phi(t)$.

5.1.2 Signal Newtonien

L'ordre 0 du développement PN correspond à un calcul Newtonien de l'orbite du système. L'amplitude $h_c(t)$ du signal est alors donnée par:

$$h_c(t) = \frac{4KG^{5/3}}{rc^4}(\pi F(t))^{2/3} \quad (5.1)$$

où r est la distance de la source et K une fonction de la masse totale et de la masse réduite du système:

$$\begin{aligned} K &= \mu M^{2/3} \\ M &= m_1 + m_2 \\ \mu &= \frac{m_1 m_2}{M} \end{aligned}$$

La fréquence instantanée est donnée par:

$$F(t) = F_0 \left(1 - \frac{t - t_0}{\tau_0} \right)^{-3/8}$$

où F_0 est la fréquence du signal à l'instant initial arbitraire¹ $t = t_0$, et τ_0 est la durée restante jusqu'à la coalescence à partir de cet instant $t = t_0$:

$$\tau_0 = \frac{5}{256} \frac{c^5}{KG^{5/3}} (\pi F_0)^{-8/3} \quad (5.2)$$

On en déduit la phase $\Phi(t)$ par intégration:

$$\Phi(t) = \frac{16\pi\nu_0\tau(\nu_0)}{11} \left[1 - \left(1 - \frac{t - t_0}{\tau(\nu_0)} \right)^{5/8} \right]$$

qui ne dépend donc des masses du système binaires qu'à travers l'unique paramètre τ_0 .

Le *chirp* Newtonien fournit les caractéristiques principales du signal attendu (figure 1.4). Son spectre s'estime rapidement à l'aide de l'approximation de phase stationnaire (SPA) qui suppose que les variations relatives de l'amplitude $h_c(t)$

1. en pratique il est commode de définir t_0 de sorte que f_0 soit la fréquence à laquelle commence l'analyse.

et les variations de la fréquence instantanée $F(t)$ sont lentes devant celles de la phase $\phi(t)$. On obtient ainsi :

$$|h(f)| \propto f^{-7/6} \quad (5.3)$$

qui montre que le signal se concentre à basse fréquence, en dépit de la montée en amplitude au cours du temps. Cette caractéristique se justifie par le temps que passe le signal à basse fréquence, où l'amélioration de la sensibilité des détecteurs est donc de première importance.

5.1.3 Corrections d'ordres supérieures

Pour juger de la présence du signal dans le bruit de Virgo, on peut s'attendre à devoir évaluer la corrélation entre le signal expérimental et le signal théorique. Autrement dit il s'agira d'intégrer le déphasage entre signal expérimental et théorique sur la durée de ce dernier ($\sim \tau_0$).

Une bonne connaissance de la phase $\phi(t)$ est donc requise; il est nécessaire de prendre en compte les corrections post-Newtoniennes qui modifient légèrement l'expression de cette phase. Par exemple à l'ordre 2PN le signal s'écrit:

$$h(t) = \Theta h_c(t) \cos(\Phi(F(t)))$$

La nouvelle phase est donné par:

$$\begin{aligned} \Phi(F) = \Phi_l + \frac{16\pi}{5} \tau_0 F_0 \left[\left(1 - \left(\frac{F}{F_0}\right)^{-5/3}\right) + \frac{5}{4} \frac{\tau_1}{\tau_0} \left(1 - \left(\frac{F}{F_0}\right)^{-1}\right) \right. \\ \left. - \frac{25}{16} \frac{\tau_{1.5}}{\tau_0} \left(1 - \left(\frac{F}{F_0}\right)^{-2/3}\right) + \frac{5}{2} \frac{\tau_2}{\tau_0} \left(1 - \left(\frac{F}{F_0}\right)^{-1/3}\right) \right] \end{aligned}$$

où la fréquence instantanée F est donnée par l'intégration et l'inversion de l'équation:

$$\frac{dF}{dt} = \frac{3F_0}{8} \tau_0^{-1} \left(\frac{F}{F_0}\right)^{11/3} \left[1 - \frac{3}{4} \frac{\tau_1}{\tau_0} \left(\frac{F}{F_0}\right)^{2/3} + \frac{5}{8} \frac{\tau_{1.5}}{\tau_0} \frac{F}{F_0} - \frac{1}{2} \left(\frac{\tau_2}{\tau_0} - \frac{9}{8} \left(\frac{\tau_1}{\tau_0}\right)^2 \right) \left(\frac{F}{F_0}\right)^{4/3} \right]$$

Du point de vue de l'analyse, la principale nouveauté par rapport au signal Newtonien est l'apparition dans l'expression de la phase, en plus de la durée Newtonienne τ_0 , des durée post-Newtoniennes τ_1 , $\tau_{1.5}$ et τ_2 . Elles sont toutes fonctions

des masses du système binaire:

$$\tau_0 \equiv \frac{5}{256\pi} F_0^{-1} (\pi M F_0)^{-5/3} \eta^{-1} \quad (5.4)$$

$$\tau_1 \equiv \frac{5}{192\pi} F_0^{-1} (\pi M F_0)^{-1} \left(\frac{743}{336} + \frac{11}{4} \eta \right) \eta^{-1} \quad (5.5)$$

$$\tau_{1.5} \equiv \frac{1}{8} F_0^{-1} (\pi M F_0)^{-2/3} \eta^{-1} \quad (5.6)$$

$$\tau_2 \equiv \frac{5}{128\pi} F_0^{-1} (\pi M F_0)^{-1/3} \eta^{-1} \left(\frac{3058673}{1016064} + \frac{5429}{1008} \eta + \frac{617}{144} \eta^2 \right) \quad (5.7)$$

où η est le rapport de la masse réduite μ sur la masse total M . Le signal dépend donc à présent des deux masses du système et non d'un unique paramètre de masse comme le signal Newtonien.

Quant au temps restant jusqu'à la coalescence à partir $t = t_0$, il n'est plus simplement donné par τ_0 mais par:

$$t_c - t_0 = \tau_0 + \tau_1 - \tau_{1.5} + \tau_2$$

À l'heure actuelle le développement PN existe jusqu'à l'ordre 3.5, et les dernières contributions n'ont qu'un effet limité sur la phase, en terme du nombre de cycles accumulés sur la durée du signal dans la bande passante des détecteurs terrestre. Ce développement est donc satisfaisant pour l'analyse.

Par contre le développement PN est limité intrinsèquement à la phase où le mouvement du système se comporte comme une succession d'orbites; lorsque les astres plongent l'un vers l'autre la notion d'orbite perd son sens et l'approximation fait de même. C'est pourquoi, on doit restreindre le signal décrit jusqu'ici aux fréquences qui ne dépassent pas la fréquence F_{LSO} de dernière orbite stable (Last Stable Orbit) décrite par le développement, donnée à l'ordre 2PN par :

$$F_{LSO} = \frac{1}{\pi M} \left[\frac{2 \left(\sqrt{1539 - 1008\eta + 19\eta^2} - 9 - \eta \right)}{3(81 - 57\eta + \eta^2)} \right]^{3/2} \quad (5.8)$$

Cette fréquence limite est plus basse pour les hautes masses: elle est proche de 450 Hz - bien à l'intérieur de la bande de détection - pour un couple de trous noirs de $10M_\odot$ chacun, au lieu d'environ 3 kHz pour un couple d'étoiles à neutrons de $1.4M_\odot$. La part de rapport signal sur bruit contenue dans la partie au-dessus, mal décrite par le développement PN, est donc plus importante pour les couples de trous noirs. C'est pourquoi d'autres méthodes de calculs viennent s'ajouter à ce développement dans le cas des coalescences de trous noirs, comme par exemple les méthodes effectives à un corps (EOB) [?], qui permettent de compléter le développement PN pour obtenir la suite du signal et augmenter les chances de détection.

5.2 Détection de la forme d'onde dans les données : méthode standard

5.2.1 Filtrage adapté

Dans les données de Virgo, un signal tel que celui décrit à l'instant sera toujours mélangé au bruit instrumental. Le signal reconstruit de Virgo $h(t)$ s'écrira donc :

$$h(t) = s(t) + n(t)$$

où $s(t)$ désigne le signal éventuellement présent et $n(t)$ est le bruit du détecteur.

On considère $n(t)$ comme un processus gaussien et stationnaire, caractérisé par sa densité spectrale de puissance $S_n(f)$. La méthode optimale pour la détection d'un signal connu au sein d'un tel bruit porte le nom de filtrage adapté, ou encore filtrage optimal. Elle consiste à comparer à un seuil prédéterminé l'expression :

$$S = (h|h_T) = 2 \int_0^\infty \frac{h(f)h_T^*(f) + h^*(f)h_T(f)}{S_n(f)} \cdot df \quad (5.9)$$

où $h(f)$ et $h_T(f)$ sont les transformées de Fourier des signaux expérimental et théorique. S est une sorte de produit scalaire entre signaux, qui mesure leur corrélation normalisée par la densité de bruit.

En l'absence de signal dans le bruit, S a donc une distribution gaussienne centrée en 0 et de largeur :

$$(\Delta S)^2 = \langle (n|h_T)^2 \rangle = (h_T|h_T)$$

tandis qu'en présence de signal, S n'est plus centrée en 0 mais en :

$$\langle S \rangle = \langle (s|h_T) + (n|h_T) \rangle = (s|h_T)$$

Le rapport signal sur bruit (SNR) de l'observable S vaut donc :

$$SNR \equiv \frac{\langle S \rangle}{\sqrt{(\Delta S)^2}} = \frac{(s|h_T)}{\sqrt{(h_T|h_T)}} \quad (5.10)$$

Dans le cas - espéré - où le signal présent s est identique au signal recherché h_T , le SNR est donné par la formule de Wiener :

$$SNR = \sqrt{(h_T|h_T)} = \sqrt{4 \int_0^\infty \frac{|h_T(f)|^2}{S_n(f)} \cdot df} \quad (5.11)$$

On l'appelle alors SNR optimal, car la forme (5.9) du filtrage adapté peut s'obtenir en recherchant, parmi l'ensemble des filtres linéaires quelconques, celui pour lequel le SNR est maximum.

En réalité, les bornes de l'intégrale (5.11) sont finies: la limite basse f_l est imposée par la durée finie du signal de coalescence utilisé par l'analyse, tandis que la limite haute f_u est imposée par l'échantillonnage des données. Le SNR s'écrit alors:

$$SNR = \sqrt{4 \int_{f_l}^{f_u} \frac{|h_T(f)|^2}{S_n(f)} \cdot df} \quad (5.12)$$

Les limites f_l et f_u doivent donc être choisies afin d'éviter une dégradation conséquente du SNR optimal.

Pour appliquer le filtrage adapté à la recherche de coalescences binaires dans Virgo, on doit aussi tenir compte des nombreuses inconnues dont dépend le signal attendu h_T :

- La position et l'orientation de la source vis-à-vis du détecteur,
- Le temps d'arrivée du signal,
- La polarisation (phase initial ϕ_0),
- Les masses des deux astres.

Pour ne pas passer à côté du signal quelle que soient les valeurs prises par ces différents paramètres, on doit donc en principe répéter le calcul de S pour "toutes" leurs valeurs possibles. En pratique quelques petites modifications de l'équation (5.9) suffisent à simplifier considérablement le balayage des trois premiers paramètres.

Position et orientation de la source

Il n'est pas nécessaire de répéter le filtrage pour différentes positions et orientations. En effet le signal h_T émis par une source quelconque s'écrit:

$$h_T(t) = \Theta \left(\frac{1 \text{Mpc}}{D} \right) T(t)$$

où $T(t)$ est le signal émis par une source d'orientation optimale ($\Theta = 1$) située à 1 Mpc du détecteur. Le résultat du filtrage adapté pour une source quelconque s'écrit donc:

$$(h|h_T) = \Theta \left(\frac{1 \text{Mpc}}{D} \right) (h|T)$$

C'est pourquoi dans la pratique il est suffisant de ne calculer que $S = (h|T)$; le rapport signal sur bruit pour un signal quelconque présent dans les données est alors donné par (eq. 5.10):

$$SNR = \Theta \left(\frac{1 \text{Mpc}}{D} \right) \sqrt{4 \int_0^\infty \frac{|T(f)|^2}{S_n(f)} \cdot df} \quad (5.13)$$

d'où l'on déduit la distance maximale à laquelle Virgo peut voir une coalescence d'orientation optimale au dessus d'un seuil donné SNR_{min} :

$$\frac{D_{max}}{1Mpc} = \frac{1}{SNR_{min}} \sqrt{4 \int_0^\infty \frac{|T(f)|^2}{S_n(f)} \cdot df}. \quad (5.14)$$

Appliqué à la sensibilité nominale de Virgo, cette expression prédit, pour un SNR_{min} de 10, une distance maximale de ~ 25 Mpc pour un système binaire d'étoiles à neutron ou de l'ordre de $100Mpc$ pour un systèmes de trous noirs.

Temps d'arrivée du signal

Le temps d'arrivée d'un éventuel signal détectable est également inconnu. Le filtrage adapté effectué avec un signal de référence T_{t_0} d'instant initial t_0 donné ne donne pas d'informations sur la présence éventuel d'un signal commençant à $t_0 + \Delta t$. Pour ne pas manquer ce signal on doit reprendre le calcul de S pour chaque instant t d'une prise de données, pouvant correspondre à l'instant initial t_0 d'un signal de coalescence éventuellement présent dans les données. Le résultat du filtrage adapté devient alors un signal à la fréquence d'échantillonnage de $h(t)$:

$$S(t) = (h|T_{t_0=t})$$

qui présentera des pics aux instants t correspondant au démarrage de signaux de coalescence dans les données.

Cet opération répétitive serait techniquement irréalisable sur des longues prises de données, sans tirer parti de la forme simple que prend le signal de référence retardé $T_{t_0=t}$ dans l'espace des fréquences:

$$T_{t_0=t}(f) = T_{t_0=0}(f)e^{i2\pi ft}$$

qui permet d'écrire:

$$S(t) = 2 \int_0^\infty \left(\frac{h(f)T_0^*(f)}{S_n(f)} e^{-i2\pi ft} + c.c. \right) df \quad (5.15)$$

Autrement dit, il suffit de calculer la transformé inverse - au lieu d'une simple intégrale - de $\frac{h(f)T_0^*(f)}{S_n(f)}$ pour obtenir $S(t)$.

Dans la pratique les deux FFTs et la FFT inverse ne peuvent être calculées sur des durée infinies ou même sur toute la durée d'une prise de données. Les limitations sur la durée des FFTs sont aussi bien imposées par l'aspect technique du calcul, que par le souhait de traiter les données en ligne et sans trop de retard, donc par paquets de durée restreinte.

Ainsi appliqué à une "fenêtre" de données, le calcul de $S(t)$ n'est évidemment valable que lorsque l'intégralité du *template* est contenu dans cette fenêtre. Ceci implique que:

- La taille T de la fenêtre d'analyse (durée des FFTs) doit être plus grande que la durée τ du *template*

- Le résultats n'est pas valide pour les τ dernières secondes de la fenêtre d'analyse, pour lesquels le *template* est sorti² de la fenêtre. Les segments successifs doivent donc au moins se chevaucher de la durée du *template*.

La durée T semble alors être un compromis entre:

- une durée suffisamment petite pour ne pas avoir à stocker en mémoire et à calculer les FFTs sur des longueurs de *template* trop importantes .
- Une durée suffisamment grande comparée à celle du *template* , pour limiter la perte de temps liée au chevauchement des segments.

En pratique T est surtout limité par les tailles de cache des CPUs actuels. Le choix canonique est alors d'adopter une fenêtre double de la durée du *template*.

Phase initiale ϕ_0

La sortie du filtre ainsi construit dépend encore de la phase initiale ϕ_0 choisie pour le signal de référence. Comme pour une démodulation classique, l'inconnue sur la phase du signal à détecter est levée en filtrant deux fois les données:

1. avec un signal de référence T_0 de phase initial 0 (signal en phase),
2. avec un signal de référence $T_{\frac{\pi}{2}}$ de phase initial $\frac{\pi}{2}$ (signal en quadrature).

On obtient ainsi deux signaux filtrés:

$$\begin{aligned} S_p(t) &= (h|T_{t,0}) \\ S_q(t) &= (h|T_{t,\frac{\pi}{2}}) \end{aligned}$$

qui contiennent à eux deux toute l'information nécessaire pour reconstruire la sortie d'un filtre de phase quelconque. En effet tout signal de référence T_{ϕ_0} se projette sur la base des signaux en phase et en quadrature suivant:

$$T(t, \phi_0) = \cos(\phi_0)T_{t,0} - \sin(\phi_0)T_{t,\frac{\pi}{2}}$$

Par linéarité, la sortie d'un filtre quelconque s'écrit donc:

$$S_{\phi_0}(t) = \cos(\phi_0)S_0(t) - \sin(\phi_0)S_{\frac{\pi}{2}}(t)$$

Parmi tous ces filtres, celui dont la phase ϕ_0 donne le plus grand signal de sortie S est celui qui correspond à la valeur de phase la plus probable d'un éventuel signal détecté. Son signal de sortie vaut:

$$\rho(t) = \max_{\phi_0} S_{\phi_0}(t) = \sqrt{S_p^2(t) + S_q^2(t)} \quad (5.16)$$

Il est donc suffisant de rechercher des pics dans $\rho(t)$ pour être sûr de ne pas manquer un signal, quelle que soit sa phase initiale.

2. En fait les FFTs supposent par définition le signal T -périodique. La partie du *template* qui est sortie de la fenêtre est donc réapparue au début de celle-ci. En tous les cas ceci revient à chercher un signal qui n'a rien à voir avec le signal de binaire.

Après l'opération de somme quadratique décrite à l'instant, la sortie du filtrage adapté en présence de bruit seul n'est plus gaussienne, mais obéit à la distribution de Rayleigh (χ^2 à deux degrés de liberté):

$$P_{\chi^2}(\rho) = \rho \exp\left(-\frac{\rho^2}{2\sigma^2}\right) \quad (5.17)$$

dont la largeur σ n'est autre que la largeur ΔS_0 de la sortie d'un simple filtre adapté sans les deux quadratures (eq.5.10).

En présence de signal on a un χ^2 non-central, dont la distribution est très proche, pour un SNR dépassant 3 ou 4, de la gaussienne obtenue dans le cas du filtre à une seule quadrature. C'est pourquoi on conserve par convention la définition du SNR donné par l'équation 5.10.

5.2.2 Grilles de filtres

En résumé, le signal $\rho(t)$ est le résultat d'une procédure améliorée de filtrage adapté, qui reste optimale quelque soit la position et l'orientation de la source, le temps d'arrivée du signal ou encore sa phase initiale. Reste alors uniquement le problème des masses des deux astres dont dépend le signal de référence.

Un balayage continu de ces deux paramètres est impossible : on est donc contraint à répéter la recherche pour un nombre fini de couple de masses, en appliquant aux données une grille de filtres adaptés, couvrant le mieux possible l'espace de masses que l'on souhaite explorer ($\sim 1M_\odot - 30M_\odot$).

Les coalescences binaires détectables par Virgo n'ont aucune raison de correspondre précisément aux couples de masses de la grille. Avant de la construire, il est donc indispensable de quantifier la perte de SNR induite par l'utilisation d'un filtre "inadapté", dont les masses sont différentes de celles du signal détecté.

Supposons donc qu'un signal s , émis par une coalescence de masses (m_1, m_2) , soit présent dans les données :

$$s = h_T = \Theta \left(\frac{1Mpc}{D} \right) T$$

et que celles-ci soient filtrées avec un signal de référence T' associé aux masses (m'_1, m'_2) . A temps d'arrivée et phase initiale identiques, le SNR est alors donné par (eq. 5.10):

$$SNR = \Theta \left(\frac{1Mpc}{D} \right) \frac{(T|T')}{\sqrt{(T'|T')}}$$

soit une fraction du SNR optimal (eq. 5.11):

$$\frac{SNR}{SNR_{opt}} = \frac{(T|T')}{\sqrt{(T|T)}\sqrt{(T'|T')}}$$

La fraction de SNR optimal effectivement récupérable dans la pratique, après balayage par le filtre $\rho(t)$ des différents temps d'arrivée t_0 et phases initiales ϕ_0 possibles, s'écrit:

$$M = \max_{\phi_0, t_0} \frac{(T|T'_{\phi_0, t_0})}{\sqrt{(T|T)}\sqrt{(T'|T')}} \quad (5.18)$$

Cette expression est symétrique par échange de T et T' : elle caractérise la similitude, du point de vue du filtrage adapté, de ces deux signaux, et prend sa valeur maximum de 1 lorsqu'ils sont identiques. On l'appelle *recouvrement* - ou *match* - entre deux couples de masses.

Le signal étant en $\frac{1}{D}$, La distance maximale de détection (5.14) diminue proportionnellement au SNR; en considérant en première approximation que les coalescences binaires sont uniformément réparties dans l'univers, la fraction d'événements détectés malgré la sous-optimalité du filtre est donc:

$$\frac{N}{N_{opt}} = M^3 \quad (5.19)$$

Le recouvrement permet donc de définir un critère de construction de la grille de *templates*. Pour limiter le nombre d'événements passant à travers les mailles du filet, on doit s'assurer que n'importe quel couple de masses faisant partie de l'espace à couvrir, affiche un recouvrement avec le *template* de la grille le plus proche au moins supérieure à un certain *recouvrement minimal* MM , fixé au préalable afin que MM^3 soit comparable au taux maximum d'événements que l'on accepte de sacrifier.

Une fois fixé ce critère on peut définir la zone de couverture d'un *template*: elle est délimitée par un contour formé des points dont le recouvrement avec ce *template* est égal au recouvrement minimal: $M = MM$. Le choix des couples de masses de la grille s'apparente à un jeu de pavage dont le but est de couvrir l'espace des masses avec les zones de couvertures des *template* qu'on choisit. Le nombre de *template* dépend donc des tailles de ces zones de couvertures et de l'efficacité du pavage.

On préfère décrire les *templates* par leurs coordonnées $\vec{\tau} = (\tau_0, \tau_{1.5})$, équivalentes aux masses (eq. 5.4), en fonction desquelles les variations du recouvrement $M(\vec{\tau}, \vec{\tau}')$ et donc les formes des zones de couvertures des *templates* sont plus régulières. Dans la limite des forts recouvrements - en pratique pour $M \gtrsim 95\%$ d'après [36] - on peut développer la fonction M autour de son maximum de 1 obtenu pour deux *templates* identiques:

$$M(\vec{\tau}, \vec{\tau} + \Delta\vec{\tau}) \simeq 1 - g_{ij} \Delta\tau^i \Delta\tau^j \quad (5.20)$$

avec

$$g_{ij}(\vec{\tau}) = -\frac{1}{2} \frac{\partial^2 M(\vec{\tau}, \vec{\tau}')}{\partial \tau^i \partial \tau'^j} \bigg|_{\tau'^j = \tau^k} \quad (5.21)$$

La zone de couverture ($M \leq MM$) d'un *template* a alors la forme d'une ellipse³, que l'on peut aussi voir comme un cercle dans un espace $(\tau_0, \tau_{1.5})$ courbe dont la métrique est la matrice g_{ij} , la distance entre deux points étant mesurée par $\sqrt{1 - M}$.

Pour estimer le nombre de *templates* à l'issue du pavage, on utilise la métrique pour mesurer la surface totale de l'espace à couvrir par les ellipses :

$$S_{tot} = \int d^2\tau \sqrt{\det \|g_{ij}\|}$$

En tenant compte du chevauchement des ellipses, la surface couverte ΔS par un seul *template* de la grille dépend de l'efficacité du pavage. Dans le cas simple d'une grille rectiligne, la surface efficace est celle d'un carré centré sur le *template* :

$$\Delta S = \left(2\sqrt{\frac{1 - MM}{2}} \right)^2 = 2(1 - MM)$$

On peut alors estimer le nombre d'ellipse ou de *templates* nécessaire à la couverture de l'espace souhaité comme le rapport des deux surfaces :

$$N = \frac{S_{tot}}{\Delta S}$$

Dans le cas de Virgo [37], on dispose ainsi d'une estimation du nombre de *templates* nécessaire allant de 10^4 à 10^6 , suivant que l'on cherche ou non à étendre la recherche à des objets plus légers que les étoiles à neutrons. A ces nombres on associe une estimation très approximative - et très dépendantes des choix finals d'implémentation - des moyens de calculs nécessaires, allant de 10 à 300 GFlops et de 2×10^9 à 5×10^{11} mots pour le stockage en mémoire des *templates*.

Les moyens de calculs seront donc l'un des principaux facteurs limitant la recherche de coalescences binaires. C'est pourquoi de nombreuses méthodes alternatives, moins coûteuses en moyen de calculs, ont pu être envisagées. Le prix à payer est généralement une plus grande perte de SNR, ces méthodes sont alors qualifiées de sous-optimales. La recherche multi-bande, qui va être développée à présent, est en ce sens spéciale, puisque elle permet de réduire les moyens de calculs nécessaires tout en conservant l'optimalité du filtrage adapté.

5.3 Recherche multi-bandes

5.3.1 Principe

L'objectif de la recherche multi-bandes est de permettre d'effectuer une recherche de coalescences binaires par filtrage adapté avec grille de *templates*, sans

3. dont les paramètres - grand et petit axes, orientation - s'obtiennent en diagonalisant g_{ij} .

nécessairement disposer des moyens de calculs cités ci-dessus. Comme son nom l'indique, le principe de cette méthode est de séparer l'analyse en plusieurs bandes de fréquence, chaque bande étant analysée indépendamment par la méthode du filtrage adapté avec grille de *templates* décrite dans la partie précédente. Les résultats obtenus dans chaque bande (signaux filtrés par les différents *templates*) doivent alors être recombinaés pour obtenir des résultats comparables à ceux qu'aurait donnés une analyse sur la bande totale de détection. Nous allons voir l'intérêt d'une telle méthode.

Cas d'un *template* unique

Considérons d'abord le cas d'une recherche avec un *template* unique. Le temps de calcul et la mémoire nécessaire sont directement liés au nombre de points de FFTs à calculer, donc à la durée des *templates* et leur fréquence d'échantillonnage. Or la durée des *templates* de coalescences binaires est essentiellement imposée par la partie basse fréquence du signal, alors que leur fréquence d'échantillonnage est au contraire imposée par la plus haute fréquence de l'analyse. La plus longue partie du *template* est donc inutilement sur-échantillonnée.

En séparant l'analyse en bandes de fréquence, on résout cette contradiction puisque:

- Le *template* pour la bande haute fréquence est très court.
- Le *template* pour la bande basse fréquence (et éventuellement pour la bande moyenne fréquence) peut être sous-échantillonné.

Ainsi, pour une analyse couvrant la bande de détection nominale de Virgo, les nombres de points des FFTs à calculer pour chaque bande après une séparation en trois bandes sont de l'ordre des tailles de cache des CPUs, et donc proches des valeurs optimales pour ce genre de calculs.

Les résultats de ces deux ou trois analyses "optimisées" sont ensuite recombinaés exactement pour obtenir le signal filtré tel que l'aurait donné une analyse à une seule bande, d'après l'expression:

$$\int_{f_l}^{f_u} \frac{h(f)T^*(f)}{S_n(f)} df = \sum_{i=1}^N \int_{f_i}^{f_{i+1}} \frac{h(f)T^*(f)}{S_n(f)} df \quad (5.22)$$

avec N le nombre de bandes et f_i les limites inférieures des bandes, $f_1 = f_l$ et $f_{N+1} = f_u$. Le SNR doit donc être le même que dans la recherche à une bande.

Grilles de *templates*

Dans le cas général où l'on utilise une grille de *template*, chaque *template* bénéficie des avantages déjà évoqués de la recherche multi-bande. De plus l'ensemble de la grille bénéficie d'un avantage lié au nombre total de *templates* utilisé.

Ce nombre est en effet déterminé par la taille des ellipses qui représente les zone couvertes par les *templates*. Le recouvrement entre deux *templates* est plus élevé si la durée des *templates*, c'est à dire la bande de fréquence de l'analyse, est petite; en première approximation la taille des ellipses est inversement proportionnelle à la durée des *templates*. Les ellipses sont donc plus grandes pour les *templates* des bandes que pour la bande totale, et le nombre de *templates* dans les grilles associées à chaque bande est réduit par rapport à la bande totale. On s'attend donc à des nombres de *templates* plus petits que dans une analyse à une seule bande, surtout pour les *templates* de la bande haute fréquence, qui sont plus court.

On peut donc réduire encore les moyens nécessaires en construisant indépendamment les grilles de *templates* dans chaque bande de fréquence. La recombinaison des résultats fera alors intervenir des *templates* de chaque bande dont les masses ne coïncident pas nécessairement. Mais puisque par construction des grilles il est garanti de retrouver au moins une fraction donnée MM du SNR de chaque bande, on doit s'attendre à retrouver au moins cette même fraction MM sur la bande totale.

La recherche multi-bande devrait donc offrir la même optimalité que la recherche à une seule bande. La validation de cet hypothèse sera l'objet du prochain chapitre. Mais avant, passons à une description plus détaillée de l'implémentation de cette méthode.

5.3.2 Implémentation

La recherche multi-bande a été implémentée, sur le principe énoncé ci-dessus, sous la forme d'un code en C appelé MBTA (Multi-Band Template Analysis). Les principaux éléments qui le constituent sont :

- Les banques de *templates* correspondant à chacune des deux ou trois bandes de fréquence. A ces *templates*, que l'on qualifie de *réels*, sont associées les méthodes de filtrage adapté qui s'appliqueront aux données.
- La banque de *templates* correspondant à la bande de fréquence totale, qui serait utilisés pour le filtrage dans une analyse mono-bande. A ces *templates* sont associées les méthodes de recombinaison des résultats issus des *templates* réels. On les qualifie de "virtuels" car ils ne servent pas au filtrage; en terme de mémoire ils ne représente finalement que quelques paramètres de recombinaison, et non les signaux tout entiers.
- Le suivi de la sensibilité du détecteur, à travers la densité spectrale de puissance du bruit (PSD) utilisée pour le filtrage adapté.

templates réels

Chaque segment de données est d'abord filtré par les banques de *templates* réels.

Les banques de *templates* réels ne sont construites qu'une seule fois, à l'initialisation de l'analyse. Le choix des couples de masses comme le calcul des signaux en phase et quadrature sont effectués par des fonctions de la librairie *inspiral* développé dans Virgo. Pour éviter certains effets de bord, les *templates* sont calculés dans le domaine temporel sur une bande de fréquence plus large de 5% que leur bande d'analyse, convertis dans le domaine fréquentiel par FFT, puis ramenés à leur bande d'analyse en tronquant des extrémités.

Les banques de *templates* réels associés à chaque bande ont des tailles de fenêtres différentes; les segments de données leur sont donc distribués de façon asynchrone. Mais au sein d'une même banque tous les *templates* ont la même taille de fenêtre - par défaut le double de la longueur du *template* le plus long, arrondi à une puissance entière de deux pour garantir un calcul rapide des FFTs.

Chaque banque procède sur chaque nouveau segment reçu de la façon suivante:

1. Pour limiter les effets de bord, le segment est préalablement multiplié par une fenêtre de lissage qui vaut 1 partout sauf sur les premiers et derniers 5% de la fenêtre ou une transition douce - sinusoïdale - vers 0 est assurée. Les $2 \times 5\%$ de la fenêtre affectés seront jetés à l'issue du filtrage.
2. La FFT des données - pour calculer $h(f)$ dans l'équation (5.15) - n'est calculée qu'une seule fois pour toute la banque.
3. Le sous-échantillonnage des données - hormis pour la bande haute fréquence - est alors effectué dans le domaine fréquentiel en tronquant la partie supérieure de la FFT. Le facteur de sous-échantillonnage est entier; on le choisit le plus grand possible, pourvu que la limite supérieure de la bande - définie par l'utilisateur - reste inférieure à la nouvelle fréquence de Nyquist.
4. Le filtrage en phase et quadrature est ensuite effectué pour chaque *template* réel suivant (5.15), à l'aide d'une PSD contenant le même nombre de points que le segment. On calcule donc une FFT inverse pour chaque quadrature de chaque *template* réel.
5. Une mesure glissante de la largeur σ du signal filtré est également mise à jour. Pour servir à la normalisation du SNR suivant la définition standard (eq. 5.11), σ doit correspondre à la largeur d'une seule quadrature; on utilise en pratique une moyenne quadratique des largeurs des deux quadratures:

$$\sigma = \sqrt{\frac{1}{2} [\sigma_p^2 + \sigma_q^2]} \quad (5.23)$$

Les nouvelles données filtrées $S_p^{(i)}(t)$ et $S_q^{(i)}(t)$ ainsi obtenues sont ensuite recombinaées par les *templates* virtuels, au fur et à mesure qu'elles sont simultanément disponibles dans toutes les bandes i .

***templates* virtuels**

La banque de *templates* virtuels garantit la recombinaison optimale des signaux filtrés par les *templates* réels. Ils permettent en effet de répondre à deux difficultés cachées par l'aspect formellement simple de l'équation (5.22) sur laquelle repose la recombinaison.

D'une part, il faut savoir au préalable quels doublets ou triplets $\{T^{(i)}\}$ de *templates* réels issus de chaque bande i doivent être recombinaisonnés pour couvrir au mieux la recherche de signaux sur la bande totale. Les *templates* virtuels permettent de former ces associations, en recherchant pour chaque *template* virtuel de la banque le *template* réel de chaque bande qui correspond le mieux, c'est à dire qui donne le meilleur recouvrement.

D'autre part, au sein d'un de ces doublets ou triplets, chaque *template* réel détecte une partie différente du signal potentiellement présent dans les données. La détection du signal par les bandes successives i se produit donc avec un retard croissant τ_i et un déphasage ϕ_i , correspondant respectivement au temps mis par le signal pour atteindre chaque bande et à la phase accumulée au moment où il l'atteint. Pour additionner les signaux filtrés $S_p^{(i)}(t)$ et $S_q^{(i)}(t)$ de façon cohérente, il doivent être préalablement recalés et remis en phase:

$$\begin{aligned} S_p(t) &= \sum_{i=1}^{nBand} \cos(-\phi_i) S_p^{(i)}(t + \tau_i) - \sin(-\phi_i) S_q^{(i)}(t + \tau_i) \\ S_q(t) &= \sum_{i=1}^{nBand} \sin(-\phi_i) S_p^{(i)}(t + \tau_i) + \cos(-\phi_i) S_q^{(i)}(t + \tau_i) \\ \rho(t) &= \sqrt{S_p(t)^2 + S_q(t)^2} \end{aligned} \quad (5.24)$$

où $\rho(t)$ désigne le signal recombinaisonné associé au *template* virtuel. On comprend alors la deuxième fonction des *templates* virtuels: ils permettent, lors de l'initialisation, de déterminer les paramètres τ_i et ϕ_i indispensables à la recombinaison; il suffit pour cela de leur faire jouer le rôle du signal à détecter, et de mesurer les retards et déphasages dans chaque bande à la détection d'un tel signal.

Plus techniquement, cette banque est au départ initialisée comme le serait l'unique banque de *templates* d'une recherche à une seule bande: la grille de couples de masses et les *templates* en phase et quadrature sont déterminés sur la bande de fréquence totale. La transformation de ces *templates* en *templates* virtuels se fait ensuite:

1. On associe au *template* virtuel un *template* réel de chaque bande, en cherchant, parmi les *templates* réels de chaque banque, celui qui donne le meilleur recouvrement. Le recouvrement n'est pas calculé pour tous les *templates* réels; l'outil de construction de grille permet une pré-sélection rapide des plus proches voisins du *template* virtuel en terme de recouvrement.

2. Le recouvrement résultant d'une maximisation sur le temps et la phase (eq.5.18), on note les décalages optimaux - en temps et en phase - correspondant au pic de détection du *template* virtuel par le *template* réel, qui ne sont autres que les paramètres de recombinaison τ_i et ϕ_i .
3. Les signaux en phase et quadrature sont alors détruits. On ne conserve donc essentiellement que les masses, les identités des *templates* réels $\{T^{(i)}\}$ associés, et les deux paramètres de recombinaison.

Lors du traitement des données, la recombinaison se fait, pour chaque *template* virtuel, par segments de tailles variables au cours du temps, correspondant aux nouvelles données filtrées $S_p^{(i)}(t)$ et $S_q^{(i)}(t)$ simultanément disponibles dans toutes les bandes, en tenant compte des retards à appliquer. Elle a lieu chaque fois que possible, donc la durée de ce segment ne dépasse jamais la plus courte des tailles de segments des banques de *templates* réels. Pour chaque segment - et chaque *template* virtuel - on retrouve:

1. Une pré-sélection évitant la recombinaison de segments qui n'ont aucune chance de contenir un dépassement du seuil requis SNR_{min} par le signal recombinaison $\rho(t)$. On vérifie pour cela qu'au moins l'une des amplitude $\rho_i(t)$ des signaux filtrés des N bandes contient un dépassement du seuil "par bande":

$$SNR_{min}^{(i)} = \frac{SNR_{min}}{\sqrt{N}}$$

En effet, même si un pic dépassant SNR_{min} se répartit de façon parfaitement équitable sur les N bandes, il ne peut descendre en dessous de $SNR_{min}^{(i)}$ dans toutes les bandes. Ensuite, à moins que le segment ne soit éliminé, on passe aux étapes suivantes.

2. L'interpolation des signaux filtrés de toutes les bandes sauf la dernière, pour les ramener à la fréquence d'échantillonnage maximum et permettre le calcul de la recombinaison échantillon par échantillon. Les signaux en phase et quadrature $S_p^{(i)}(t)$ et $S_q^{(i)}(t)$ pouvant osciller rapidement notamment en présence de signal, l'interpolation se fait sur leur puissance $\rho_i^2(t)$ et phase

$$\Phi_i(t) = \arctan \frac{S_q^{(i)}(t)}{S_p^{(i)}(t)}$$

qui varient plus doucement; l'interpolation est quadratique sur l'amplitude et linéaire sur la phase.

3. La recombinaison, effectuée échantillon par échantillon suivant l'équation (5.24).
4. La recherche de candidats sur le signal. La normalisation en SNR de $\rho(t)$ suivant la définition standard (eq. 5.11) implique de recombinaison également

les largeurs glissantes σ_i de chaque bande:

$$\sigma = \sqrt{\sum_{i=1}^N \sigma_i^2} \quad (5.25)$$

Pour chaque maximum dépassant le seuil SNR_{min} , un candidat est enregistré en format *frame*, contenant notamment le temps du pic, et les paramètres de masses du *template* virtuel. Une zone morte, arbitrairement fixé à ± 7.5 ms, est exclue autour du pic. La recherche de maximum continue tant qu'on trouve des candidats au dessus du seuil.

Suivi de la sensibilité du détecteur

La sensibilité du détecteur intervient à deux niveaux dans la recherche multi-bande, comme elle le ferait dans une recherche à une seule bande :

- Lors de la construction des grilles de *templates*, puisque le recouvrement et donc la métrique qui définit les ellipses couvertes par chaque *template* dépendent de la PSD
- Lors du filtrage - par les *templates* réels - puisque la PSD sert de normalisation de la corrélation entre le *template* et les données (eq. 5.9).

Il est donc important de permettre à la recherche de s'adapter en cours de route pour suivre d'éventuelles variations de la sensibilité du détecteur.

Des effets sur la grille de *templates* ne sont pas à craindre dans un premier temps; en particulier celle-ci n'est pas affectée par des variations d'un facteur global de la sensibilité, mais uniquement par des changements de forme, dont elle ne dépend qu'au second ordre. De plus reconstruire les grilles de *templates* en cours de route risquerait de s'avérer extrêmement lourd, puisque l'initialisation des *templates* réels et virtuels serait alors à refaire. On choisit donc de conserver les mêmes grilles de *templates* sur toute la durée de l'analyse, quitte à choisir dès le départ une grille suffisamment fine pour amortir toutes les variations possibles.

Le filtrage adapté ne peut par contre s'épargner une prise en compte de ces variations, car l'utilisation d'une PSD fautive dans l'équation (5.15) serait synonyme d'écart aux conditions du filtrage optimal et donc de perte de SNR. Un premier niveau de suivi, déjà évoqué ci-dessus, est donné par l'utilisation de σ (5.23) glissant - dont les constantes de temps sont typiquement choisies de quelques minutes - pour la normalisation des signaux filtrés et recombinaés en SNR. On s'adapte ainsi aux variations d'un facteur global de la sensibilité, évitant par exemple une soudaine augmentation du nombre de candidats lorsque la sensibilité diminue globalement.

Les variations de forme de la sensibilité - comme l'apparition ou la disparition d'une source de bruit donné sur une certaine bande de fréquence - font également l'objet d'un suivi au cours du traitement des données. Le code de la recherche multi-bande contient un module PSD par banque de *templates* réels, donc par

bande de fréquence. Pour éviter le calcul de FFTs supplémentaires, chaque module utilise la FFT des données déjà calculée par sa banque pour le filtrage, qui a naturellement le bon nombre de points pour le calcul de l'intégrale 5.15. A chaque itération, les FFTs successives sont moyennées en puissance avec un facteur d'oubli Λ :

$$PSD \leftarrow \Lambda \times PSD + (1 - \Lambda)|FFT|^2$$

correspondant au premier ordre à une moyenne glissante de n FFTs, avec $n = (1 - \Lambda)^{-1}$.

Des tests préliminaires réalisés avant l'implémentation de MBTA avec un code très simple tournant sur du bruit blanc ont montré que n devait dépasser plusieurs dizaines pour que la dégradation du SNR par introduction du bruit de mesure de la PSD dans le filtre ne dépasse pas $\sim 1\%$. Dans le cas des bandes basses fréquences qui contiennent l'essentiel de la durée du signal, les segments peuvent atteindre quelques minutes. Le suivi souhaité - généralement de l'ordre de 15 à 30 minutes pour pouvoir suivre le détecteur tout en restant stable à l'échelle de la durée d'un événement - commencerait donc à se montrer problématique en terme de SNR. Puisqu'il n'y a pas de raison particulière pour que ces bandes aient besoin d'une meilleure résolution sur la PSD que la bande haute fréquence, on ramène alors leurs résolutions à celle de la bande haute fréquence en lissant les PSD en fréquence:

$$PSD(f) = \frac{1}{k} \sum_{f' f - \frac{k}{2} * df}^{f' f - \frac{k}{2} * df} PSD(f') \quad (5.26)$$

où $df = \frac{1}{T}$ est la résolution initiale de la PSD, et k le facteur de sous-échantillonnage de la bande concernée. Les mêmes études préliminaires ont montré qu'un tel lissage sur k points est équivalent voir meilleur que la moyenne sur n FFT.

A la fin de ces deux opérations - moyenne en temps et en fréquence - toutes les PSDs peuvent s'adapter avec la même constante de temps que celle de la bande haute fréquence qui est donc commune à toute l'analyse, et qui peut être fixée à 15 ou 30 mn sans risquer de dégrader le SNR. Elles ont également toutes la même résolution effective, même si leurs nombres de points sont différents, c'est-à-dire la résolution de la bande haute fréquence.

Une trace de ce suivi est enregistrée dans les données: la distance maximale de détection vue par le détecteur, calculée pour chaque *template* virtuel suivant l'équation (5.14), en utilisant en pratique les *templates* réels associés et les PSDs de chaque bandes pour intégrer le SNR théorique.

5.3.3 Regroupement de candidats

A l'issue de la recherche, la présence d'un signal au sein du bruit se manifeste par la présence parmi les candidats sélectionnés d'un candidats affichant un SNR élevé.

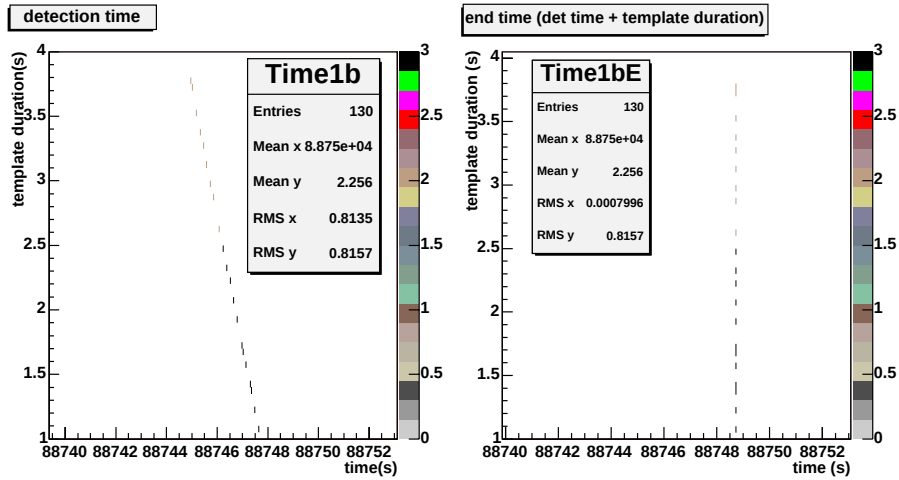


FIG. 5.1 – *Distribution des candidats engendrés par différents templates, pour la détection d'un même signal immergé dans du bruit simulé, suivant la durée du template (axe Y) et le temps du candidat (axe X). A droite, avec le temps du pic de détection; à gauche avec le temps de fin.*

En pratique un signal de binaires dans les données entraîne la création de bien plus qu'un candidat:

- D'une part parce que chaque *template* (virtuel) peut générer plusieurs candidats pour un même événement, suivant l'auto-corrélation du signal filtré (ou plutôt recombinaison dans le cas de la recherche multi-bande),
- D'autre part parce que plusieurs *templates* (virtuels) de la grille sont sensibles à un même signal, le SNR décroissant avec l'écart aux masses du signal.

Techniquement, cette multiplicité des candidats complique l'analyse des résultats et le discernement des vraies détections parmi tous les candidats. De plus on souhaite pouvoir regrouper tous les candidats associés à un même événement car cet ensemble donne des informations pouvant servir à une reconstruction des paramètres de la source.

C'est pourquoi un algorithme de regroupement des candidats pouvant correspondre à un même événement a été développé en aval de la recherche multi-bande. Le critère choisi pour ce regroupement est le temps d'arrivée des candidats. Cependant ce temps diffère fortement lorsque les candidats associés aux mêmes événements sont issus de *templates* de durée différente. Pour corriger cet effet on considère non pas le temps d'arrivée correspondant au début de l'événement (pic dans le signal filtré), mais le temps de fin (temps d'arrivée + durée du *template*). Ce temps de fin est en effet très proche du temps auquel le signal entre dans la bande de fréquence pour laquelle la densité de SNR est maximale, et qui compte

donc le plus pour la détection. Une illustration de ce rôle particulier du temps de fin est visible sur la figure 5.1, qui montre la détection, avec une petite grille de *templates*, d'un signal immergé dans du bruit simulé.

Pour qu'il n'y ait pas d'ambiguïté sur la façon de regrouper les candidats, on impose que les "groupes" soient toujours centrés - en temps de fin - autour du candidat de plus fort SNR, et inclue tous les autres candidats situés à moins d'une certaine demi-largeur de ce candidat central. En cas d'ambiguïté, un sous-candidat pouvant appartenir à deux groupes sera toujours placé dans celui dont le candidat central a le plus fort SNR.

5.3.4 Perspectives

La version actuelle du code MBTA, dont l'implémentation vient d'être décrite, a essentiellement pour but de démontrer l'équivalence de cette méthode et de la méthode habituelle à une bande en terme de détection d'événements et de réjection du bruit. Elle n'est pas réellement optimisée en terme de moyens de calculs, c'est pourquoi cet aspect n'est pas plus développé ici. On notera néanmoins que, sur l'analyse typique utilisant la sensibilité nominale de Virgo qui sera présentée au chapitre 8, une réduction de l'occupation de mémoire d'un facteur 10 par rapport à l'analyse à une seule bande est observé.

Outre cette optimisation de MBTA qui reste à faire, on doit également envisager d'en étendre l'implémentation au cas des *templates* qui sont générés directement dans le domaine fréquentiel et ne nécessite donc pas de FFTs à l'initialisation, et ce dans un double but:

- Accélérer l'initialisation dans le cas des modèles de *template* qui peuvent être générés aux choix dans l'un ou l'autre domaine grâce à l'approximation de phase stationnaire.
- Ouvrir la recherche à des nouveaux modèles de *templates* qui ne peuvent être générés que dans le domaine fréquentiel.

Chapitre 6

Validation de la recherche multi-bandes

Après la description de la recherche multi-bandes, ce nouveau chapitre décrit sa validation et l'étude de ses principales caractéristiques, sur des données simulées.

L'analyse de données simulées contenant des événements permet de tester les deux principaux aspects de l'analyse : le risque de fausses alarmes en l'absence de signal, et l'efficacité de détection du signal lorsqu'il est bien présent. Elle permet également de mesurer la précision avec laquelle sont reconstruits les paramètres du signal: distance effective, temps d'arrivée et masses de la source.

La recherche multi-bandes a ainsi été testée - et a évolué en réponse aux problèmes rencontrés - tout au long des 6 MDCs Virgo. Plutôt que de s'attarder sur cette évolution, ce chapitre en montre le résultat : la validation, sur données simulées, de la recherche multi-bandes dans sa version la plus aboutie à ce jour .

Pour plus de clarté, ces résultats sont d'abord présentés dans le cas du *template* exact, cas fictif où les masses de la source du signal à détecter seraient connues à l'avance. On attend alors une statistique de détection optimale, telle qu'elle est prédite par la théorie de filtrage adapté.

Dans la deuxième partie, on étudie le cas réaliste où la recherche multi-bandes utilise des grilles de *templates*. La perte d'optimalité est alors inévitable mais doit être inférieure ou comparable à la perte tolérée, pour laquelle la grille est définie.

6.1 Recherche avec le *template* exact

6.1.1 Fausses alarmes

On s'intéresse dans un premier temps au risque de fausses alarmes, autrement dit au comportement de la recherche en présence de bruit seul. Le bruit est simulé, à l'aide de SIESTA, par une configuration rapide qui permet de reproduire les

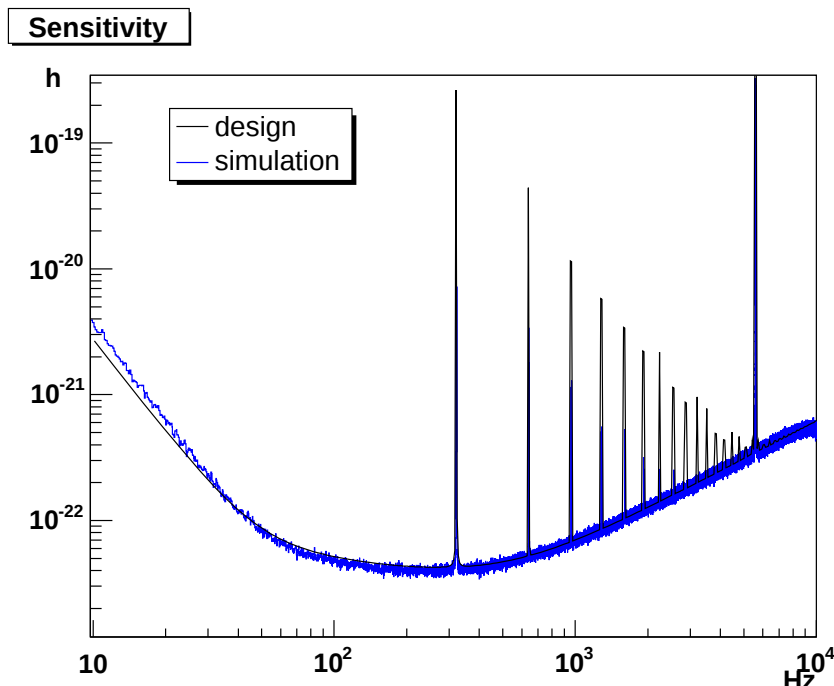


FIG. 6.1 – Spectre du bruit simulé (en bleu) comparé à la sensibilité théorique de Virgo (en noir).

principales composante spectrales de la sensibilité théorique de Virgo¹ : bruits thermiques divers et bruit de photons (figure 6.1). Il est ensuite sous-échantillonné à 4 kHz après application d'un filtre anti-repliement à 2kHz similaire à celui utilisé dans le code de reconstruction.

Ces données sont ensuite analysées, par une recherche en trois bandes avec un *template* unique. Le *template* choisi correspond au cas standard d'une coalescence de deux étoiles à neutrons de masses $1.4M_{\odot}$. Il est calculé à l'ordre PN2 (cf. 5.1.3), ordre le plus haut qui alors disponible dans le code *inspiral* de Virgo [34]. La bande d'analyse - et donc le *template* - va de 24 Hz à 2kHz. Le SNR perdu entre 0 Hz et 24 Hz est estimé inférieur à 1% du SNR total d'après l'équation (5.12), tandis qu'au dessus de 2 kHz il est complètement négligeable.

Pour vérifier que la recherche multi-bandes se comporte en l'absence de signal comme le ferait une recherche mono-bande, on veut comparer la distribution du signal de sortie $\rho(t)$ - obtenu par recombinaison du *template* virtuel - à celle

1. lors des premiers MDCs, la simulation du bruit reproduisait le spectre du dernier *run* du CITF, dans un souci de réalisme vis à vis de la sensibilité réelle du détecteur en cours de mise en route ou même celle attendue pour les premières recherches astrophysiques. La recherche des problèmes techniques a progressivement montré que ceux-ci sont d'autant plus nombreux que la bande passante analysée est large. C'est pourquoi par sécurité on utilise à présent un spectre Virgo, dont la bande passante est largement supérieure à celle des spectres plus réalistes.

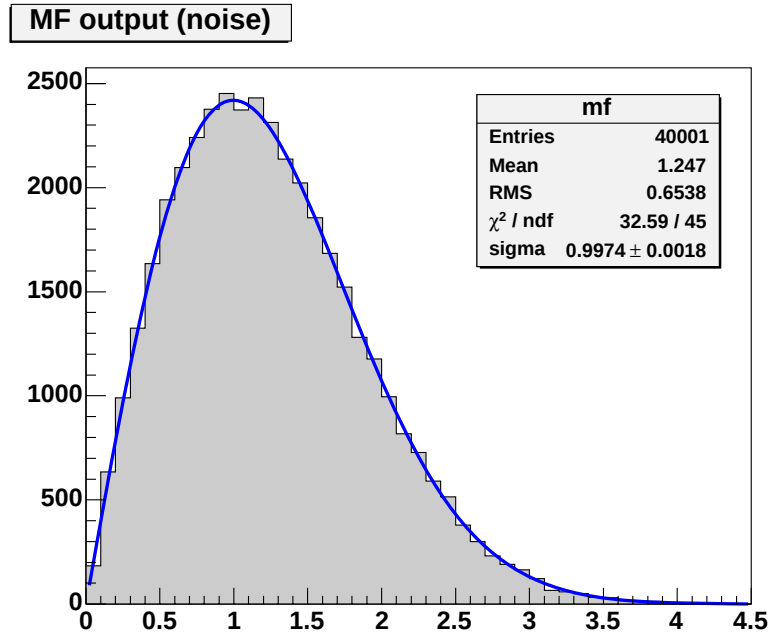


FIG. 6.2 – Distribution du signal en sortie de l'analyse trois bandes en présence de bruit seul, et ajustement par une Rayleigh de largeur $\sigma = 0.9974 \pm 0.0018$.

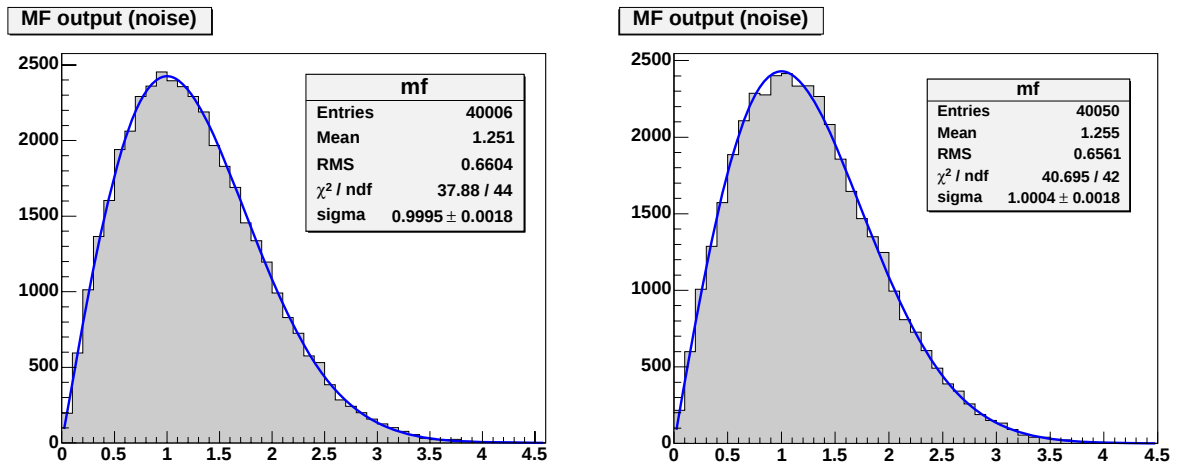


FIG. 6.3 – idem pour une analyse deux bandes (à gauche; $\sigma = 0.9995 \pm 0.0018$) et une bande (à droite; $\sigma = 1.0004 \pm 0.0018$)

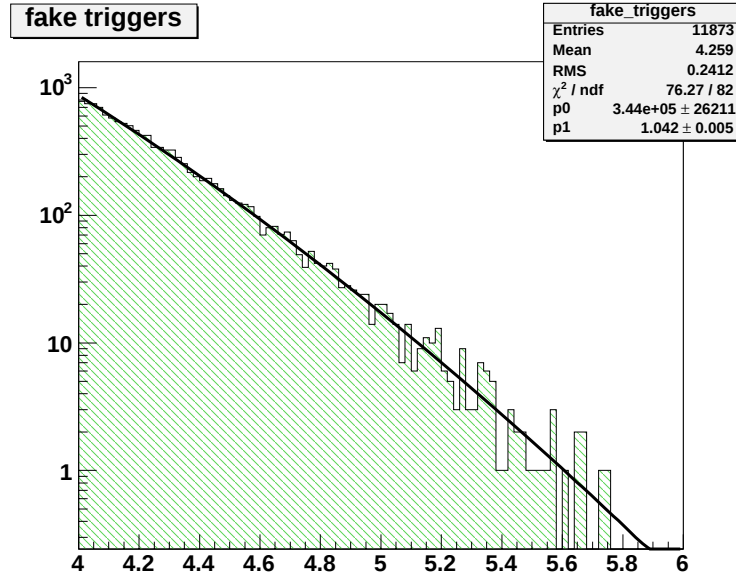


FIG. 6.4 – Distribution de SNR des candidats de l'analyse à 3 bandes sur le bruit seul. Ajustement par une Rayleigh

attendue (eq. 5.17). Pour l'évaluer, on commence par sous-échantillonner $\rho(t)$ à moins d'une valeur par demi-seconde. On s'assure ainsi de disposer de valeurs successives décorréées, la durée d'auto-corrélation du signal $\rho(t)$ étant de l'ordre de ~ 5 ms.

On peut alors ajuster l'histogramme des valeurs du signal sous-échantillonné par la distribution théorique. Le facteur d'échelle de la distribution étant fixé par le nombre d'entrées, le seul paramètre libre de l'ajustement est la largeur σ de la distribution. La distribution et l'ajustement pour environ 9 h de données sont représentés sur la figure 6.2.

$\rho(t)$ est préalablement normalisé en SNR (§5.3.2), la valeur attendue pour σ est donc 1; on obtient effectivement par l'ajustement une valeur compatible avec 1. Des résultats similaires (figure 6.3) sont obtenus pour des analyses à deux bandes ou à une seule bande.

Le filtrage du bruit seul par l'analyse multi-bande est donc aussi performant que prévu par la théorie du filtrage adapté. Caractériser cette performance par un taux de fausses alarmes par unité de temps est délicat, dans la mesure où sa définition dépend des conventions de comptage des fausses alarmes, notamment du choix de la durée devant séparer deux fausses alarmes pour les considérer comme distinctes. On notera d'ailleurs que si la distribution du signal $\rho(t)$ correspond bien à une Rayleigh de $\sigma=1$, la distribution de SNR des candidats 6.4 sélectionnés à partir de ce signal ne s'en approche qu'approximativement ($\sigma = 1.04$), déformée par le processus de sélection et de regroupement des candidats.

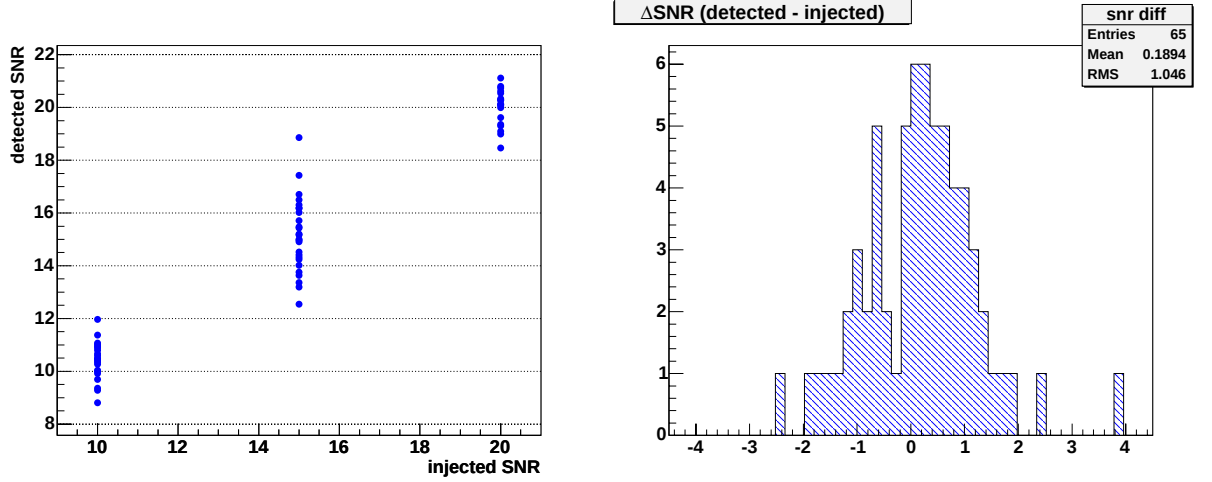


FIG. 6.5 – Comparaison entre SNR détectés et SNR injectés pour la recherche avec le template exact en trois bandes. À gauche: SNR détecté en fonction du SNR injecté. À droite: distribution de la différence SNR détecté – SNR injecté.

Mais on peut dès à présent conclure que ce taux sera identique à celui du filtrage adapté, quelle que soit la convention adopté.

6.1.2 Détection des événements

Le deuxième aspect à valider concerne la détection des signaux lorsque ceux-ci sont effectivement présents. Pour ce faire, des signaux de coalescences d'étoiles à neutrons sont ensuite injectés dans le bruit, répondant aux caractéristiques suivantes :

- On les suppose optimalement orientées par rapport au détecteur ($\Theta = 1$);
- Leurs distances peuvent prendre trois valeurs précises possibles, correspondant à des SNR optimaux (eq. 5.11) de 10, 15 ou 20.
- Leurs phases initiales ϕ_0 sont variables et balayent approximativement 2π .
- Puisqu'on teste ici la recherche avec un *template* exact, les signaux correspondent tous à l'unique couple de masses ($1.4M_{\odot} - 1.4M_{\odot}$).
- Les signaux sont calculés à l'ordre PN2, comme le *template* de l'analyse, entre 24 Hz et la fréquence de dernière orbite stable (~ 3 kHz). Ils sont sous-échantillonnés de la même façon que le bruit à 4 kHz.
- Les temps d'injections sont tirés au sort en visant une fréquence moyenne d'une injection toutes les cinq minutes; un veto empêche l'injection si deux signaux se chevauchent.

Au total 65 événements sont injectés sur 9 h de données.

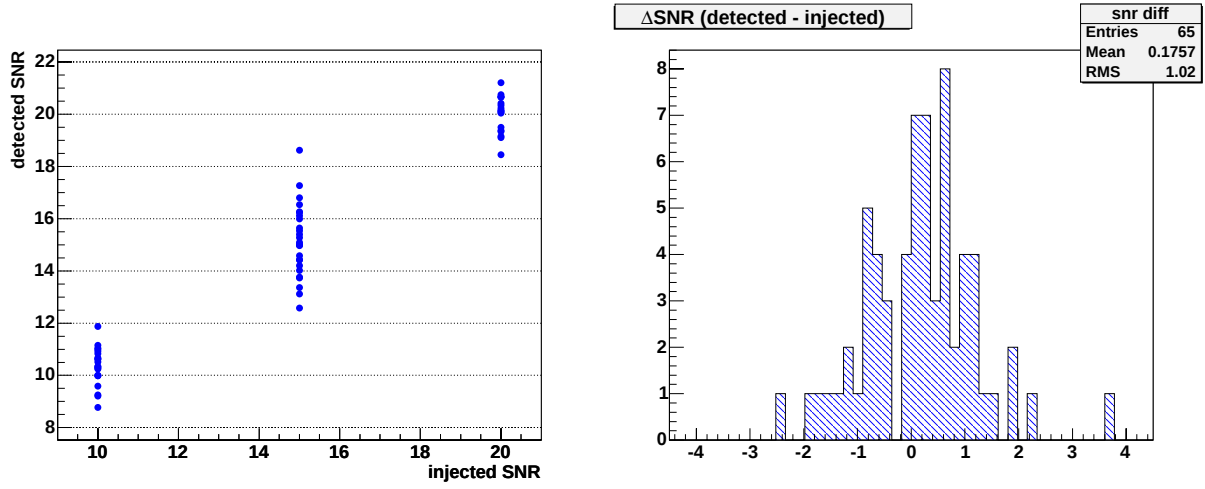


FIG. 6.6 – Comparaison entre SNR détectés et SNR injectés pour la recherche avec le template exact en deux bandes. À gauche: SNR détecté en fonction du SNR injecté. À droite: distribution de la différence SNR détecté – SNR injecté.

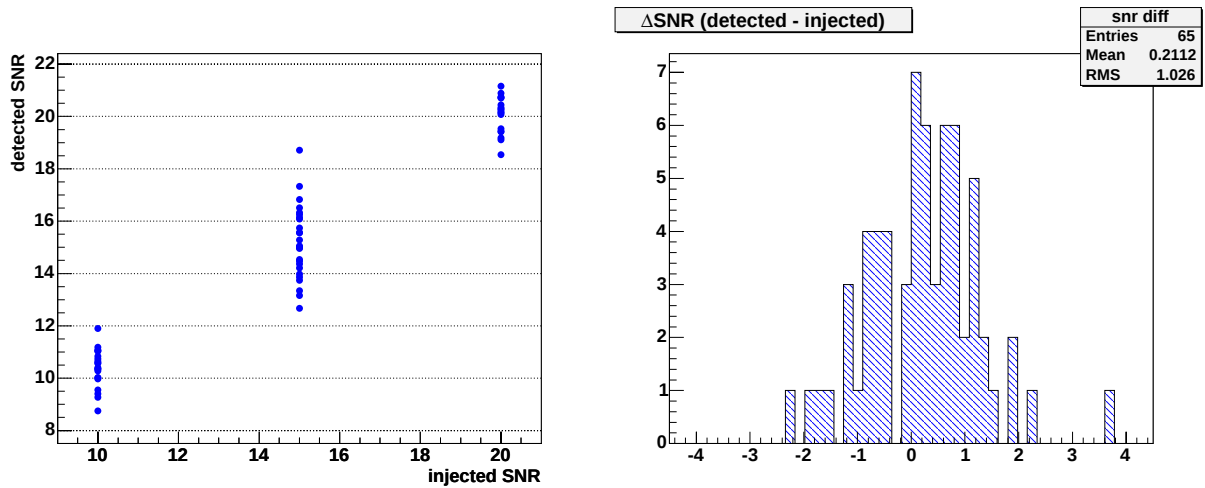


FIG. 6.7 – Comparaison entre SNR détectés et SNR injectés pour la recherche avec le template exact en une bande (filtrage adapté). À gauche: SNR détecté en fonction du SNR injecté. À droite: distribution de la différence SNR détecté – SNR injecté.

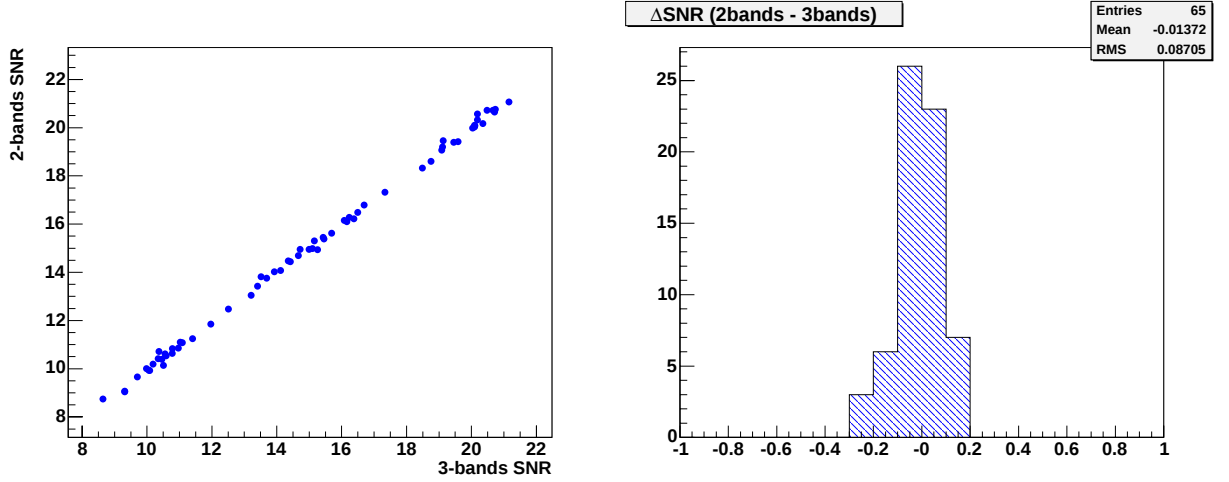


FIG. 6.8 – Comparaison entre les analyses trois bandes et deux bandes pour la recherche avec le template exact. À gauche: SNR "2 bandes" en fonction du SNR "3 bandes". À droite: distribution de la différence SNR "2 bandes" – "3 bandes".

Les données sont alors analysées avec la même configuration que pour le bruit seul. Pour s'assurer que tous les signaux sont correctement détectés, on recherche autour de chaque injection un candidat généré par l'analyse dont le temps de détection coïncide avec le temps d'injection à 1 ms près; si plusieurs candidats correspondent, on garde celui affichant le meilleur SNR. On compare ensuite le SNR détecté au SNR injecté (10, 15 ou 20).

La figure 6.5 présente le résultat de cette comparaison. La dispersion de l'écart entre SNR détectés et SNR injectés correspondants est proche de 1 comme attendu, et l'écart moyen vers le haut observé dans cet exemple ($m \sim 0.2$) est compatible avec les fluctuations du bruit ($\Delta m = \frac{1}{\sqrt{65}} \simeq 0.12$).

Des figures similaires sont obtenues avec l'analyse à deux bandes (figure 6.6) ou une bande (figure 6.7). Pour affiner la comparaison, on évalue événement par événement la dispersion entre les analyses trois bandes et deux bandes (figure 6.8), ou trois bandes et une bande (figure 6.9).

Les trois analyses sont fortement corrélées comme le montre la petite dispersion du SNR (~ 0.1). Cette dispersion résiduelle montre néanmoins que les trois analyses ne voient pas exactement le même bruit. Ces différences sont interprétées comme un effet secondaire du calcul des FFTs par segments successifs de données, qui introduit une dépendance du filtrage dans la taille et la position des segments. La figure 6.10 illustre cet effet. Elle montre la dispersion entre l'analyse à une seule bande, et une analyse identique "décalée" (qui démarre une demi-longueur de segment plus tard, décalant ainsi chaque segment d'une demi-longueur jusqu'à la fin). Cette dispersion est comparable aux dispersions résiduelles observés entre

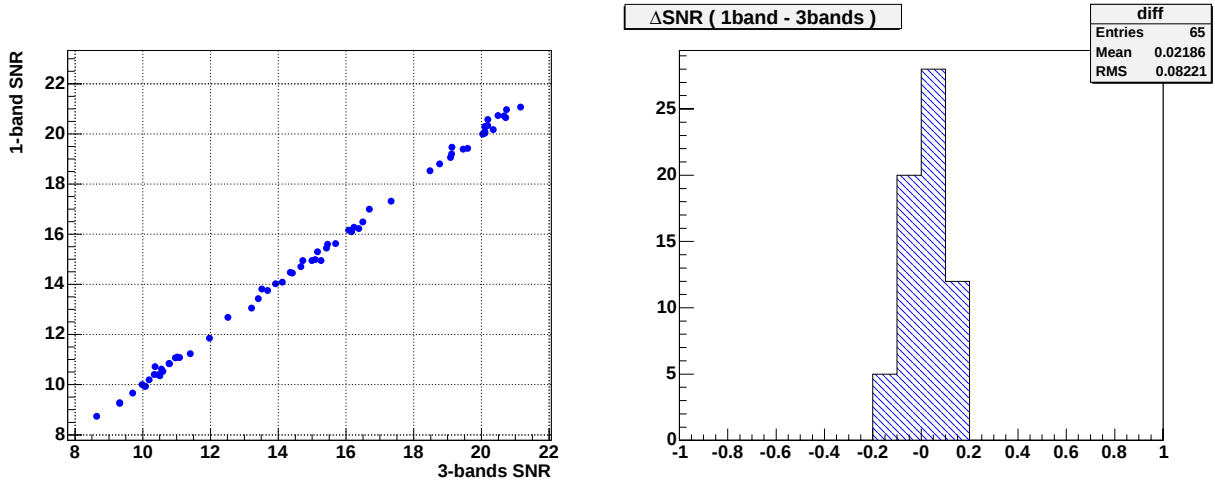


FIG. 6.9 – Comparaison entre les analyses trois bandes et une bande (filtrage adapté) pour la recherche avec le template exact. À gauche: SNR "1 bande" en fonction du SNR "3 bandes". À droite: distribution de la différence SNR "1 bande" – "3 bandes".

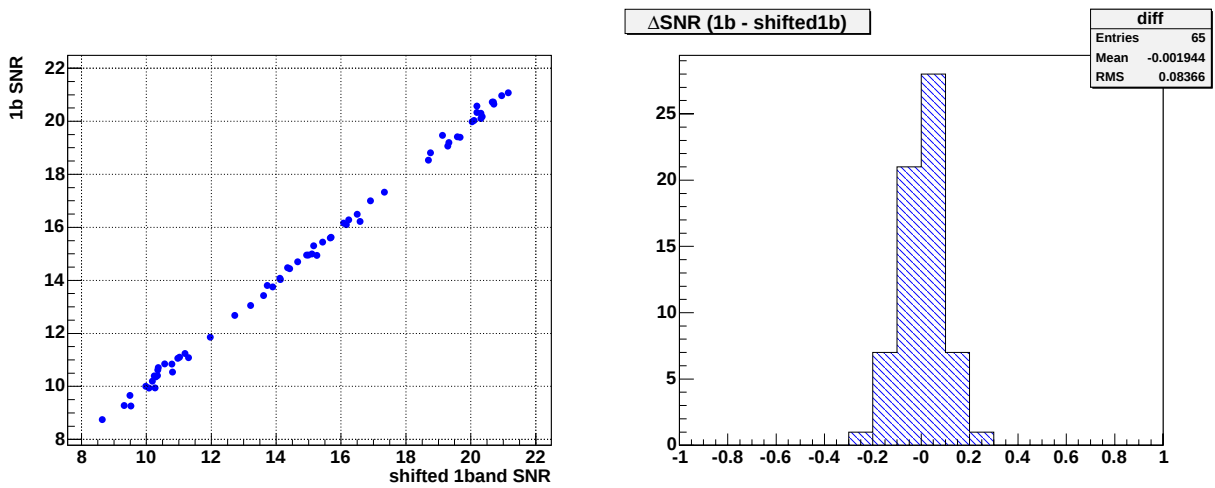


FIG. 6.10 – Comparaison entre l'analyse à une seule bande, et une analyse identique dont les segments d'analyse sont décalés d'une demi-longueur (recherche avec le template exact). À gauche: SNR normal en fonction du SNR de l'analyse décalée. À droite: distribution de la différence SNR "normal" – "décalé".

les trois analyses; celles-ci ne sont donc pas propres à la recherche multi-bandes.

L'écart de SNR moyen (~ 0.04) entre l'analyse deux-bandes et le filtrage adapté est un peu élevé pour être justifié ainsi; on peut donc soupçonner des imperfections résiduelles dans le mécanisme de recombinaison, par exemple dans l'interpolation des signaux ou la prise en compte de leurs retards. Toutefois un tel écart n'est pas considéré comme un problème; il ne représente au pire - pour un SNR de 8 - qu'une perte de SNR de 0,5%.

6.1.3 Résolution en temps et en distance effective

Dans le cas du *template* exact où les masses sont supposées connues à l'avance, les seuls paramètres qui peuvent être extraits de la recherche sont le temps d'arrivée de l'événement et sa distance effective.

La distance effective est directement liée au SNR détecté; on s'attend donc à ce qu'elle soit correctement reconstruite, puisque c'est déjà le cas du SNR. La figure (6.12) montre la comparaison entre distance reconstruite et distance simulée, et confirme cette hypothèse.

La figure 6.11 présente, pour les trois analyses, les écarts entre temps d'injection et temps de détection. On compare en fait le temps de fin d'événement (temps du pic + durée du *template*) donné par la recherche multi-bandes, au temps d'injection défini par SIESTA comme le temps où le signal est maximal; ces deux temps ne diffèrent que de deux ou trois échantillons. Les dispersions respectives des trois analyses sont comparables entre elles, et proches de l'échantillonnage $dt = \frac{1}{4k\text{Hz}} = 0,25 \text{ ms}$ qui constitue à priori une limite inférieure de la résolution en temps. Elles sont relativement faibles comparées au $\sim 27 \text{ ms}$ qui séparent Virgo de l'un ou l'autre des deux sites LIGO, donc peuvent être considérées comme satisfaisantes dans l'optique d'une recherche en coïncidence avec LIGO.

6.2 Recherche avec une grille de *templates*

On aborde à présent le cas réel pour lequel la recherche multi-bandes est conçue: une recherche utilisant une grille de *templates* pour détecter des signaux dont les masses sont inconnues au préalable.

6.2.1 Fausses alarmes

La validation de la recherche multi-bandes avec grille de *templates* au regard des fausses alarmes est identique à celle faite dans le cas du *template* exact: on vérifie simplement que la distribution du signal de sortie $\rho(t)$ obtenu par recombinaison d'un *template* virtuel est toujours conforme à celle attendue dans le cas du filtrage adapté, autrement dit que la recombinaison du *template* virtuel à partir

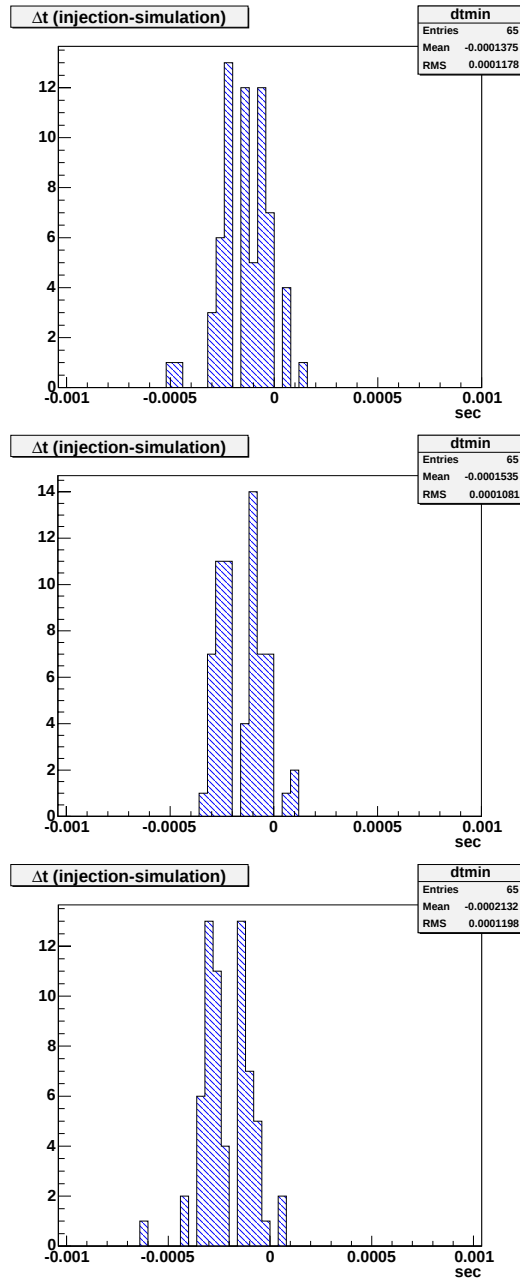


FIG. 6.11 – Différence entre le temps de détection et le temps d'injection, pour trois bandes (en haut), deux bandes, (milieu), une bande (en bas).

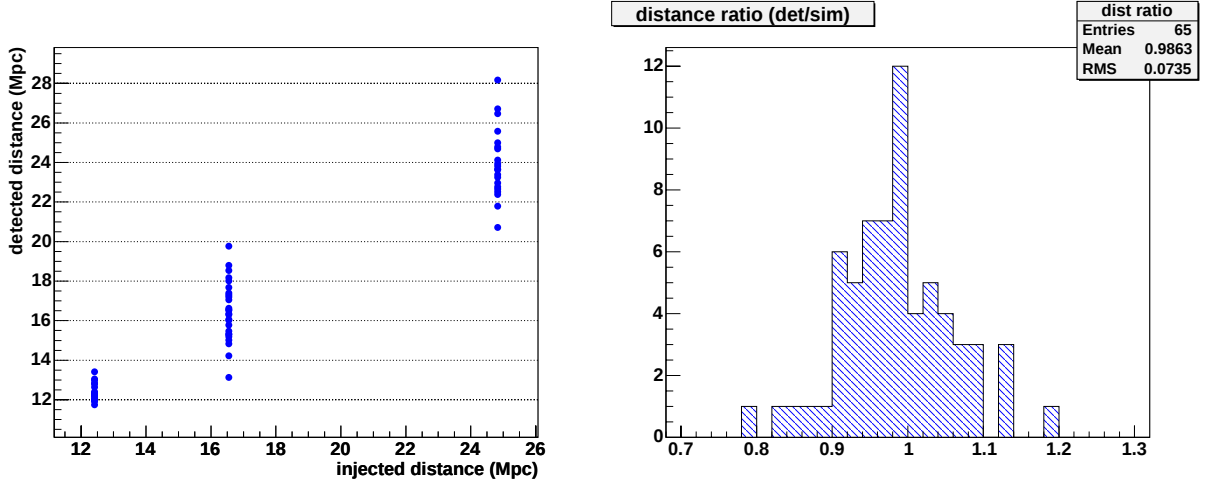


FIG. 6.12 – Comparaison entre distance effectives détectées et injectées pour la recherche avec le template exact en trois bandes. À gauche: distance détectée en fonction de la distance injectée. À droite: distribution du rapport distance détectée / distance injectée.

de templates réels correspondant à des masses distinctes n'introduit pas d'excès de bruit dans la distribution.

Le bruit simulé utilisé est le même que dans le cas précédent. La figure 6.13 montre la distribution obtenue par analyse de bruit seul pour un template virtuel de masses $(10.107M_{\odot}, 10.108M_{\odot})$, associé dans une recherche à deux bandes aux templates réels $(9.809M_{\odot}, 9.809M_{\odot})$ et $(10.148M_{\odot}, 10.149M_{\odot})$, compatible comme dans le cas du template exact avec la forme théorique 5.17.

6.2.2 Détection des événements

Comme précédemment, on souhaite à présent analyser des données simulées contenant des injections d'événements. Pour tester la recherche avec une grille, il n'est plus nécessaire d'injecter un seul couple de masses; on préfère au contraire varier les masses pour sonder plusieurs points de la grille. De plus on souhaite effectuer cette validation à la fois pour les binaires d'étoiles à neutrons ou de trous noirs. Deux jeux indépendants d'injections sont donc générés:

- 77 injections similaires à celles utilisées pour le cas du template exact, mais dont les couples de masses varient entre 1.3 et 1.6 $M_{\odot}$
- 90 injections similaires à celles utilisées pour le cas du template exact, mais générées par la méthode EOB (cf. 5.1.3) à l'ordre 3, et dont les couples de masses varient entre 4 et 8 $M_{\odot}$.

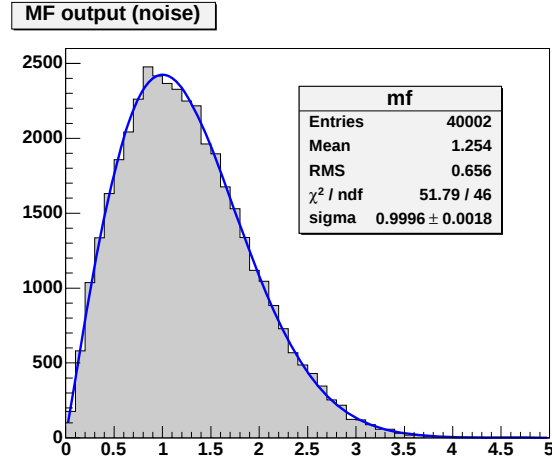


FIG. 6.13 – Distribution du signal en sortie de l'analyse deux bandes avec grille, en présence de bruit seul, et ajustement par une Rayleigh de largeur $\sigma = 1.00092 \pm 0,00018$.

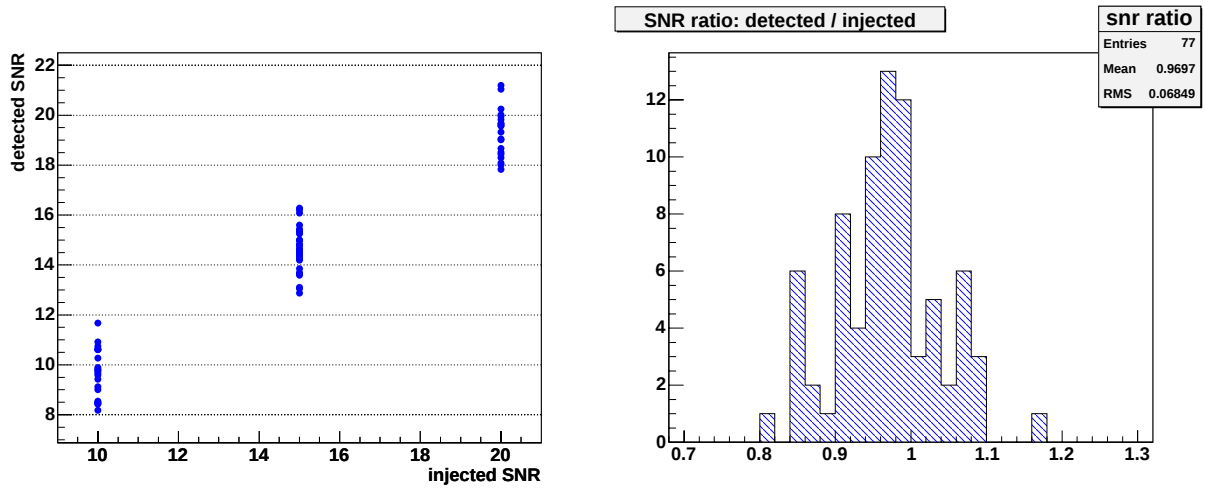


FIG. 6.14 – Comparaison entre SNR détectés et SNR injectés pour la recherche avec grille (étoiles à neutrons) en deux bandes. À gauche: SNR détecté en fonction du SNR injecté. À droite: distribution du rapport SNR détecté sur SNR injecté.

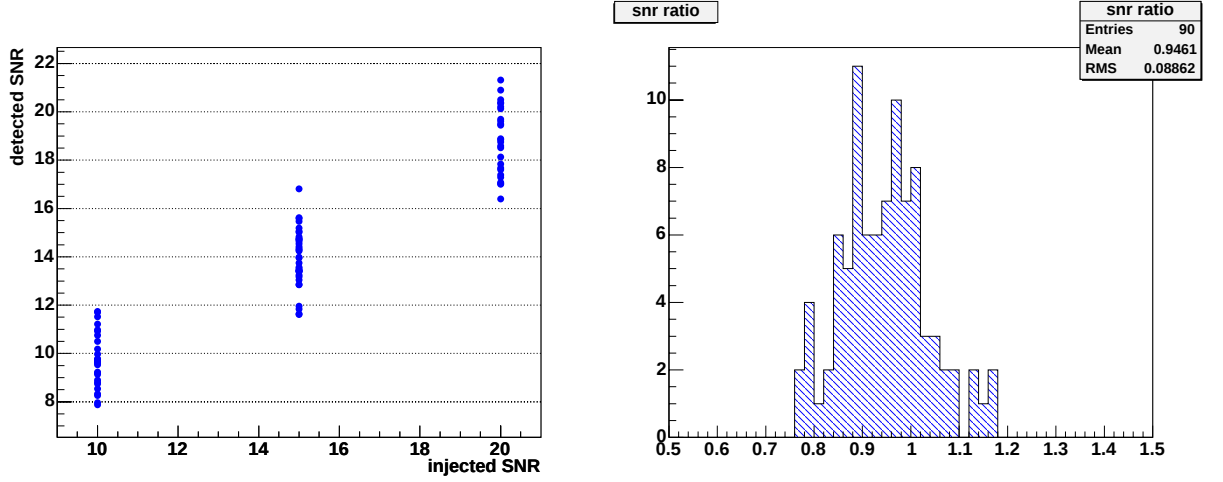


FIG. 6.15 – *Comparaison entre SNR détectés et SNR injectés pour la recherche avec grille (trou noirs) en deux bandes. À gauche: SNR détecté en fonction du SNR injecté. À droite: distribution du rapport SNR détecté sur SNR injecté.*

Pour les binaires d'étoiles à neutrons, une analyse à deux-bandes est appliquée aux données - bruit + signal - couvrant un intervalle de masses allant de $1.2M_{\odot}$ à $2M_{\odot}$, avec un recouvrement minimum $MM = 95\%$. Les résultats - comparaisons entre SNR injectés et détectés - sont présentés sur la figure 6.14; on s'intéresse dans ce cas au rapport entre SNR détecté sur injecté, car il doit être comparé au recouvrement minimal. La valeur moyenne obtenue de $\sim 0.97\%$ est, comme attendu, supérieure au recouvrement minimal.

Pour les trous noirs, l'analyse à deux-bandes couvre un intervalle de masses allant de $4M_{\odot}$ à $11M_{\odot}$, avec un recouvrement minimum $MM = 95\%$. Les résultats sont présentés sur la figure 6.15. Le rapport de SNR moyen ($\sim 0.95\%$) est juste à la limite du recouvrement minimal, ce qui n'est pas surprenant : la grille étant définie à partir de la métrique de l'espace des paramètres - calculée pour des signaux PN2 et non EOB3 - peut présenter un défaut de couverture. À long terme ces effets devront être quantifiés plus généralement, pour savoir s'ils imposent d'implémenter un calcul de métrique spécifique pour EOB, ou si ils peuvent être pris en compte par l'introduction d'une marge de sécurité dans la définition du recouvrement minimale MM de la grille; des problèmes similaires se poseront d'ailleurs pour tout les futurs générateurs de signaux que l'on souhaitera intégrer dans l'analyse.

6.2.3 Estimation des masses de la source

Les masses du *template* (virtuel) de la grille ayant donné le plus fort SNR lors de la détection d'un signal de coalescence binaire constituent une première estimation des masses de cette dernière. On compare donc événement par événement ces

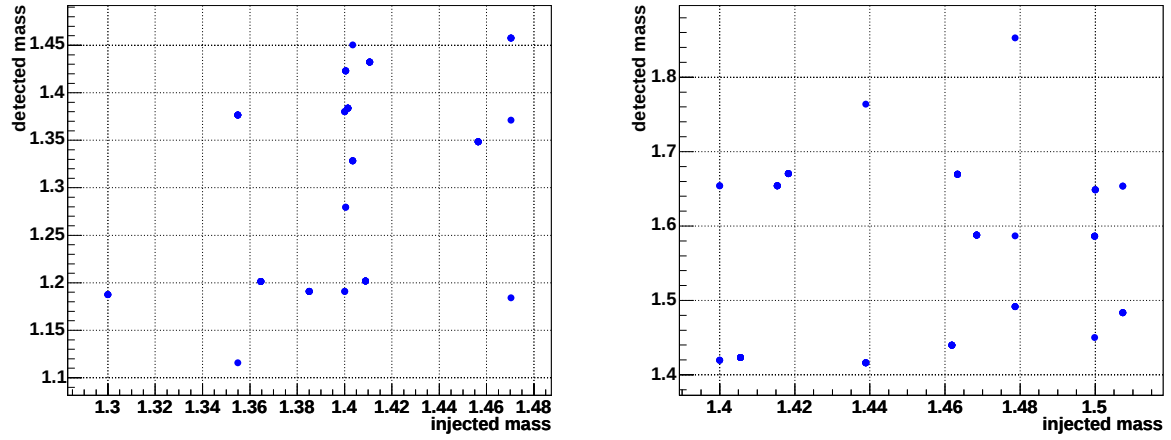


FIG. 6.16 – *Masses du template maximum en fonction des masses injectées correspondante, pour la recherche avec grille (étoiles à neutrons) en deux bandes. À gauche: plus petite masse du couple. À droite: plus grande masse du couple.*

masses à celles des événements injectés. La figure 6.16 montre cette comparaison dans le cas de la grille d'étoiles à neutrons; la grille étant plus serrée à basses masses on doit obtenir une meilleure précision dans ce cas. Cette précision est en fait très mauvaise, la corrélation entre les masses injectées et détectées semblant presque inexistante à l'échelle de la grille utilisée.

Pour interpréter ce résultat, qui n'est pas propre à la recherche multi-bande, on doit se rappeler que le signal de coalescence binaire à l'ordre Newtonien, qui constitue la base du signal, ne dépend pas des deux masses mais seulement de la durée de coalescence τ_0 , qui lui même n'est que fonction de la combinaison des deux masses K . La dépendance dans les deux masses est introduites par les corrections post-Newtoniennes qui sont importantes pour la détection, mais pas suffisamment pour permettre de les résoudre.

La précision doit par contre devenir considérablement meilleure si l'on considère le paramètre K , que l'on peut ramener à une masse commune - symétrique entre les deux astres - appelé *chirp mass*:

$$\mathcal{M} = K^{3/5} \quad (6.1)$$

La figure 6.17 montre la très bonne résolution ($\sim 10^{-4}$) obtenue pour la *chirp mass*. Dans le cas des coalescences de trous noirs, la résolution reste meilleure que 1%, bien qu'elle diminue puisque que la grille est plus éparse.

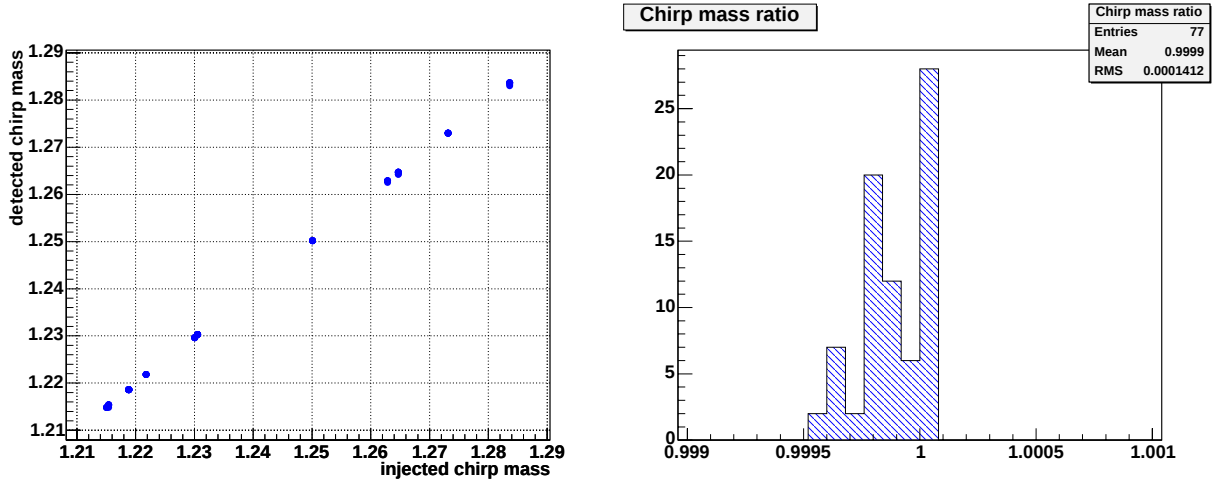


FIG. 6.17 – Comparaison entre chirp mass détectées et injectées pour la recherche avec grille (étoiles à neutrons) en deux bandes. À gauche: chirp mass détectée en fonction de celle injectée. À droite: distribution du rapport chirp mass détectée sur injectée.

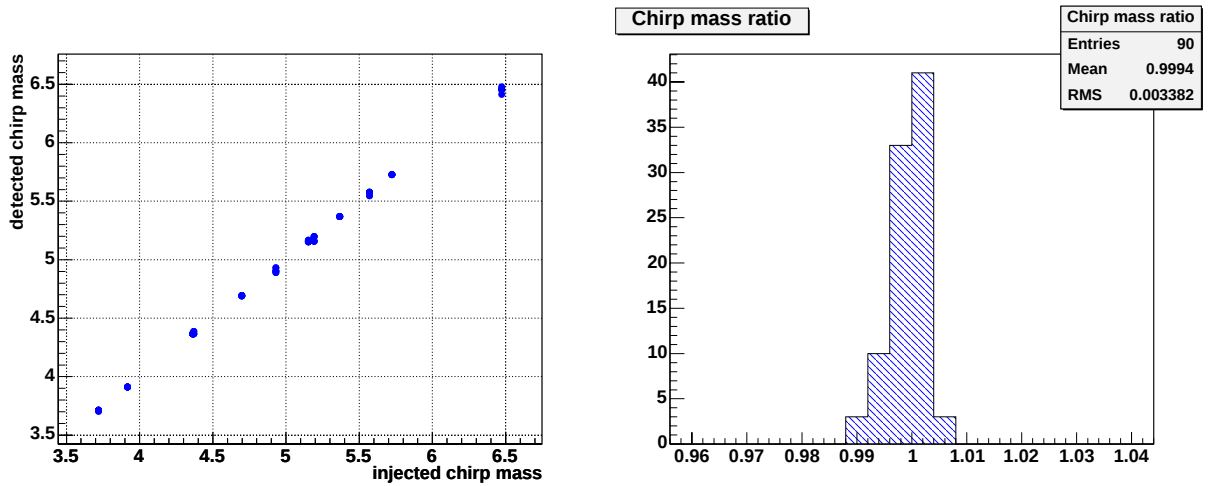


FIG. 6.18 – Comparaison entre chirp mass détectées et injectées pour la recherche avec grille (trous noirs) en deux bandes. À gauche: chirp mass détectée en fonction de celle injectée. À droite: distribution du rapport chirp mass détectée sur injectée.

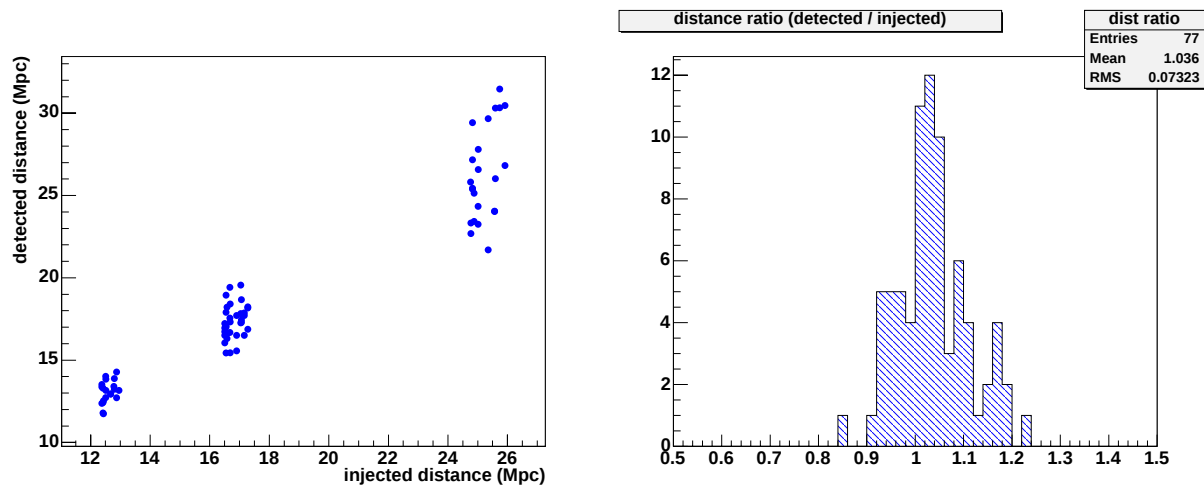


FIG. 6.19 – Comparaison entre distance effectives détectées et injectées pour la recherche avec grille (étoiles à neutrons) en deux bandes. À gauche: distance détectée en fonction de la distance injectée. À droite: distribution du rapport distance détectée / distance injectée.

6.2.4 Résolution en temps et en distance effective

L’usage de la grille a des effets sur la résolution en temps comme en distance effective.

La distance effective étant inversement proportionnelle au SNR détecté, les pertes de SNR liées au recouvrement minimal de la grille entraînent une hausse de la distance estimée (les signaux plus faiblement détectés sont interprétés comme venant de plus loin). Cet effet est visible sur les figures 6.19 et 6.20 qui montrent respectivement une hausse moyenne de la distance effective reconstruite de 3,6% pour les binaires d’étoiles à neutrons et 4,5% pour les binaires de trous noirs. On note, par comparaison, au cas à un seul *template*, que les distance injectées sont plus dispersées, puisqu’en faisant varier les masses des injections on fait varier la distance correspondant à un SNR donné.

La comparaison des temps d’arrivée se fait sur le même principe que dans le cas du *template* exact, à un détail près: dans le cas des trous noirs on ne peut pas comparer le temps de fin d’événement donné par la recherche multi-bande au temps d’amplitude maximal du signal donné par SIESTA, car ces deux définitions sont trop différentes pour des signaux calculés avec la méthode EOB - alors calculés au-delà de la fréquence de dernière orbite stable. On corrige donc le temps d’injection, après comptage dans les fichiers des événements injectés du nombre d’échantillons entre le temps d’amplitude maximal et le temps de fin.

L’utilisation d’une grille *templates* ne coïncidant pas exactement avec le signal injecté introduit des écarts supplémentaires entre les temps d’arrivée injectés et

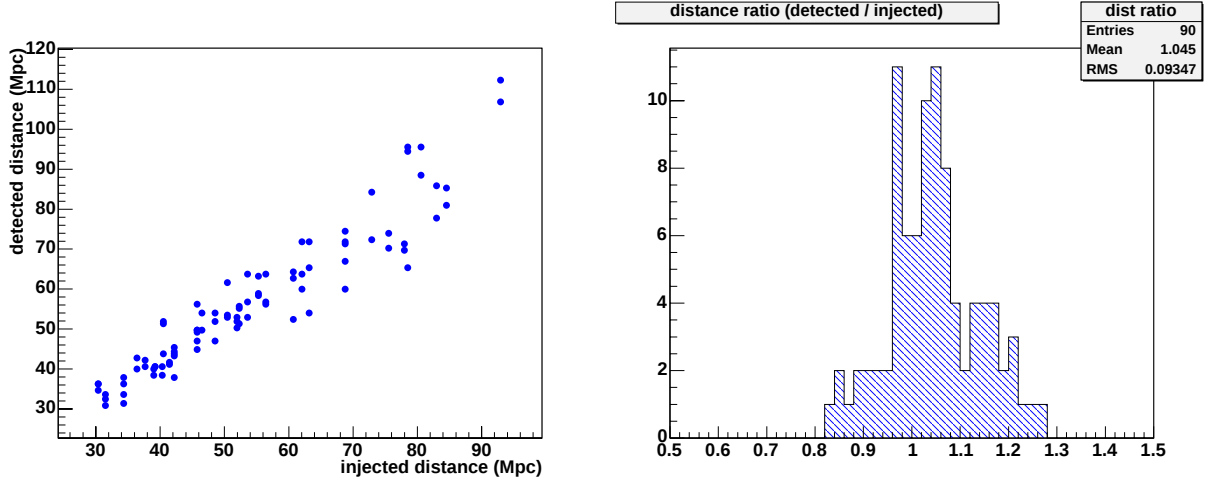


FIG. 6.20 – Comparaison entre distance effectives détectées et injectées pour la recherche avec grille (trous noirs) en deux bandes. À gauche: distance détectée en fonction de la distance injectée. À droite: distribution du rapport distance détectée / distance injectée.

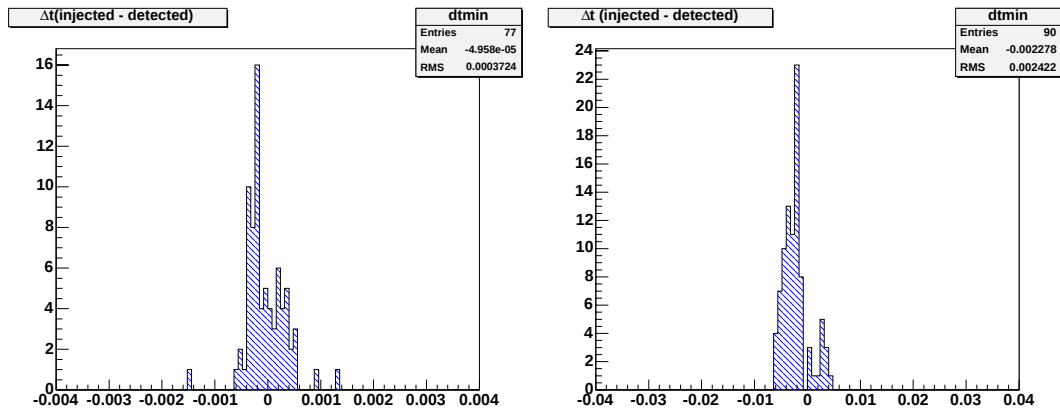


FIG. 6.21 – Différence entre le temps de détection et le temps d'injection, pour la recherche avec grille en deux bandes. À gauche: étoiles à neutrons, à droite: trous noirs.

déTECTÉS, qui ne sont qu'imparfaitement compensés par le fait de s'intéresser au temps de fin de l'événement plutôt que le temps où il commence. La figure 6.21 le confirme: la résolution en temps empire d'un facteur ~ 3 pour les étoiles à neutrons et ~ 20 pour les trous noirs, dont la grille est plus clairsemée entraînant de plus grands décalages en temps entre le signal et le *template*. Cette dernière résolution reste toutefois un ordre de grandeur en dessous de la séparation LIGO-Virgo.

6.3 Conclusion

Ces résultats montrent le bilan positif des MDCs et plus généralement de la mise en oeuvre de la recherche multi-bandes en simulation.

Ils démontrent l'optimalité de la recherche multi-bandes, aussi bien en terme de fausses alarmes en présence de bruit seul que pour la détection des événements. Aucune dégradation du SNR optimal n'est observé pour la recherche avec le *template* exact dont les masses coïncident avec celles du signal injecté; les pertes de SNR pour la recherche avec une grille de *templates* ne dépassent pas celles tolérées par le recouvrement minimal de la grille.

Les résultats sur l'estimation des paramètres des événements détectés sont également satisfaisants. Les précisions obtenues sur le temps d'arrivée comme sur la distance sont - inévitablement - limitées par les effets de grilles (discrétisation) et dépendent donc de la finesse de celle-ci. L'estimation des masses de la source répond aux prévisions théoriques: alors que les deux masses sont individuellement impossible à déterminer, une grande précision est mesurée pour la *chirp mass*, combinaison des deux masses qui caractérise le signal à l'ordre Newtonien.

A l'issue de cette démarche de validation, la recherche multi-bande peut donc servir à l'analyse de données issues du détecteur.

Chapitre 7

Application aux données actuelles de l'interféromètre

La chaîne d'analyse développée au cours des MDCs peut à présent être appliquée aux données réelles de l'interféromètre.

A ce jour, la sensibilité de ce dernier n'est pas suffisante pour légitimement espérer une véritable détection d'ondes gravitationnelles. Le but d'une recherche de signaux de coalescences binaires, dans les données prises en cours de progression vers la sensibilité nominale, est donc plutôt la détection de "faux événements", c'est à dire d'artefacts instrumentaux qui pourrait faire croire à une détection de coalescence. Indispensable à la caractérisation du détecteur, cette recherche dresse un nouveau tableau du détecteur, complémentaire de celui offert par la mesure de sensibilité: elle rend compte de l'aspect gaussien et/ou stationnaire des données, non visible sur la courbe.

Pour dresser correctement ce tableau, on doit s'assurer dans un premier temps du bon fonctionnement de la chaîne d'analyse utilisée. La recherche multi-bande ayant fait ses preuves en simulation, il s'agira de vérifier que sa capacité de détection n'est pas dégradée par son intégration à la suite de la reconstruction ou par le passage à des données réels, et non issues de modèle de bruit simple. Cette vérification, décrite dans la première partie du chapitre, passe par l'injection de signaux d'ondes gravitationnelles dans l'interféromètre et leur recherche dans les données reconstruites.

L'étude de l'aspect gaussien ou non - et plus généralement propre à l'analyse - du bruit vient ensuite, avec la recherche de "faux événements" engendrés par les bruits instrumentaux.

Enfin la stabilité du détecteur sera abordée, à travers ses variations de sensibilité au cours du temps et leurs implications pour la recherche de binaires.

7.1 Recherche d'injections

7.1.1 Les injections

Des expériences d'injections de signaux d'ondes gravitationnelles dans l'interféromètre ont donc été réalisées, au cours des *run* C4 et C5. Durant C4 les injections ont été effectuées sur une période limitée - la première nuit du run - tandis qu'elles étaient réparties sur toute la première partie du *run* utilisant l'interféromètre recombinaison, hormis une période "silencieuse".

Ces injections sont effectuées à l'aide des actionneurs électromagnétiques du couple masse de référence - miroir, comme les injections de bruit de la calibration. Elles "traversent" ensuite la chaîne de détection au sens le plus large: interféromètre, reconstruction, recherche multi-bande.

Techniquement, les injections ne peuvent pas être mélangées à un signal de correction; elle serait alors éliminées par la reconstruction, lorsque celle-ci cherche à supprimer les effets des contrôles. On choisit donc un miroir d'entrée de bras (NI) sur lequel n'est envoyé aucun signal de correction, et qui n'est donc pas pris en compte par la reconstruction. De plus on évite ainsi de gêner, au niveau du programme d'injection CaRT, l'injection des lignes de calibration sur les miroirs de bout de bras.

Comme pour la calibration, les signaux envoyés à l'électronique qui pilote les bobines sont des signaux en Volts et doivent tenir compte de la réponse des actionneurs, c'est à dire la réponse du pendule, le pôle des bobines, et le gain absolue en m/Volts, dont l'étude¹ a été présenté au chapitre 3.

Les signaux injectées sont des coalescences binaires d'étoiles à neutron de masses ($1.4M_{\odot}, 1.4M_{\odot}$) à l'ordre PN2. L'expérience étant commune au groupe d'analyse de données de Virgo, on trouve donc aussi des événements impulsifs de type supernovae.

Dans le cas de C4, les amplitudes des signaux sont choisies à priori pour des valeurs de SNR de 7, 10 et 15. Le calcul de SNR d'après la formule (5.11) utilise une courbe de sensibilité issue d'un étalonnage du détecteur antérieur au *run*, moins bonne que la vraie. De plus on a vu que la reconstruction peut améliorer cette sensibilité par soustractions d'une partie des bruits de contrôles. Ces valeurs de SNR sont donc des limites inférieures, qui peuvent être corrigés à l'analyse des résultats.

Dans le cas de C5, l'utilisation d'un actionneur instable - l'électronique de pilotage des bobines ayant été modifiée et non testée quelques jours avant le run - a rendu pratiquement impossible d'établir une liste de référence incluant le SNR des injections.

1. En fait, la donnée du gain en m/Volts est le seul élément des chapitres précédents qui n'est pas sondé par les expériences d'injections, puisque elle est utilisé à la fois par les injections et la reconstruction.

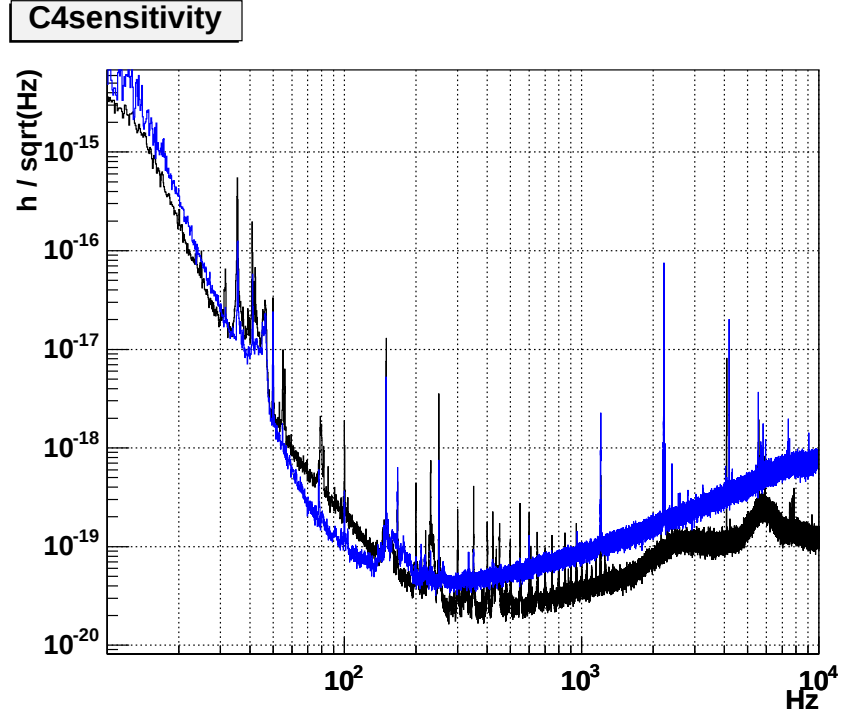


FIG. 7.1 – Sensibilité en mode recombiné pour les runs C4 (en noir) et C5 (en bleu).

7.1.2 Résultats

Au vue des injections à détecter, les configurations d'analyses utilisées doivent évidemment couvrir le couple $(1.4M_{\odot}, 1.4M_{\odot})$. Les sensibilités des deux *run* étant comparables (figure 7.1), les configurations d'analyse sont choisi identiques:

- Une grille restreinte autour du couple visée, allant de 1.35 à $1.45M_{\odot}$.
- Une bande totale d'analyse allant de 60 Hz à 2 kHz -largement suffisante pour les sensibilités en question.
- Le recouvrement minimal habituel $MM = 95\%$.

Comme en simulation, il faut ensuite identifier les candidats correspondant à des injections; pour cela on doit d'abord savoir si les critères d'associations injections-candidats sur le temps de fin sont encore valables. La figure 7.2 montre la distribution des écarts entre injections et candidats les plus proches pour l'analyse de C4; des écarts d'une ou deux secondes sont systématiquement présents. La première seconde d'écarts est provoqué par le code de reconstruction, dans sa version utilisée à l'époque. La deuxième - aléatoire - est attribué à des problèmes connus de synchronisation du programme d'injections CaRT. La fenêtre d'associations peut donc être élargie à plusieurs secondes, ce qui amène à conclure que les 35 candidats de la figure 7.2 correspondent aux 35 injections. On peut

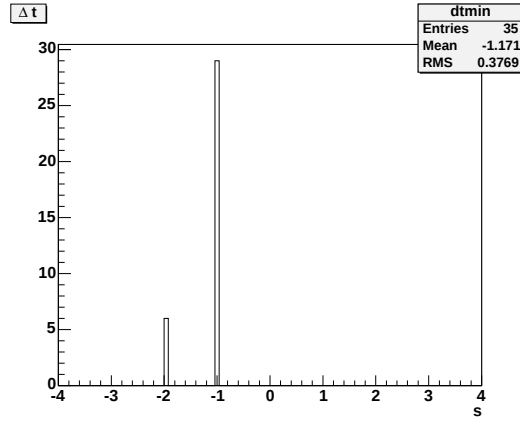


FIG. 7.2 – *Distribution des écarts entre injections et candidats les plus proches pour l'analyse de C_4*

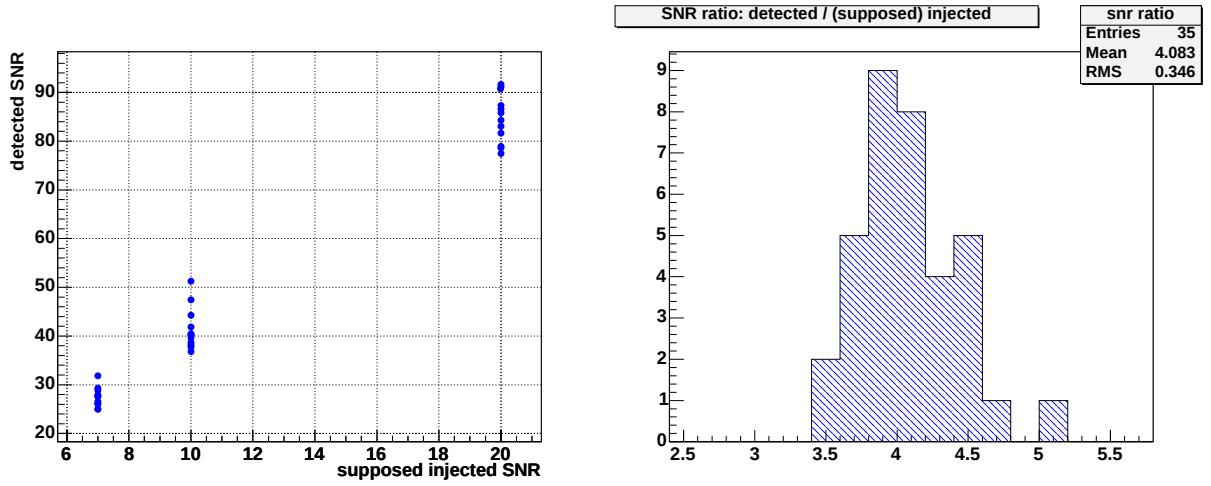


FIG. 7.3 – *Comparaison du SNR candidats-injections dans C_4 .*

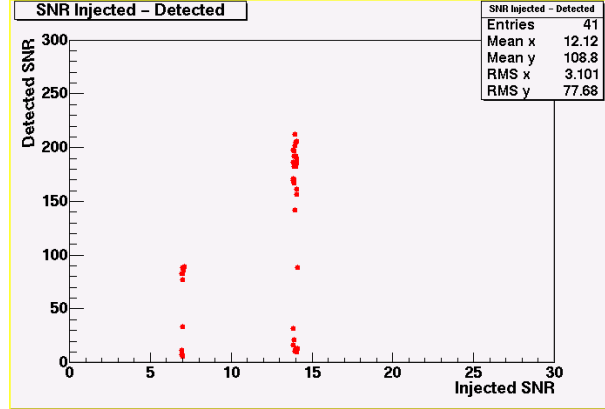


FIG. 7.4 – Comparaison du SNR candidats-injections dans C5.

alors comparer les valeurs de SNR (figure 7.3) à celles injectés. Comme prévu, les SNR détectés excèdent largement les SNR injectés, puisque ceci ne sont pas normalisés par la bonne courbe de sensibilité. Le calcul du facteur correctif [38] nécessaire pour tenir compte de cette erreur donne une valeur proche de 4, donc proche du rapport de SNR moyen observé de 4,08. La figure montre par ailleurs une très grande dispersion des SNR détectés, de l'ordre de $\sim 10\sigma$. Cette dispersion très différente de celle mesurée en simulation peut être due aux fluctuations de sensibilité comme à l'erreur de la reconstruction, qui évolue entre 10% et 20%. Dans la limite de cette erreur, on peut considérer que la chaîne d'analyse dans son ensemble fonctionne correctement sur les données réelles de l'interféromètre. Dans le cas de C5, la comparaison des SNR (figure 7.4) n'est pas vraiment interprétable pour les raisons exposées précédemment; on parvient néanmoins pour chaque groupe d'injections de SNR donné à distinguer deux groupes de SNR détectés, correspondant à deux modes de fonctionnement intermittents - haut ou bas gain - de l'actionneur durant C5.

7.2 Bruit du détecteur: Faux événements

7.2.1 Runs complets

La chaîne d'analyse ayant fait - en partie - ses preuves sur les injections, elle peut servir à la caractérisation du détecteur en l'absence de telles injections, via la recherche de candidats à fort SNR - simulant une détection de coalescence - engendrés par les bruit instrumentaux non-gaussiens.

On utilise d'abord les analyses décrites précédemment, qui incluent les injections. On doit donc éliminer les queues de distributions associées à ces événements, ce pourquoi on exclue - largement - tous les candidats situés à moins de dix secondes d'une injection.

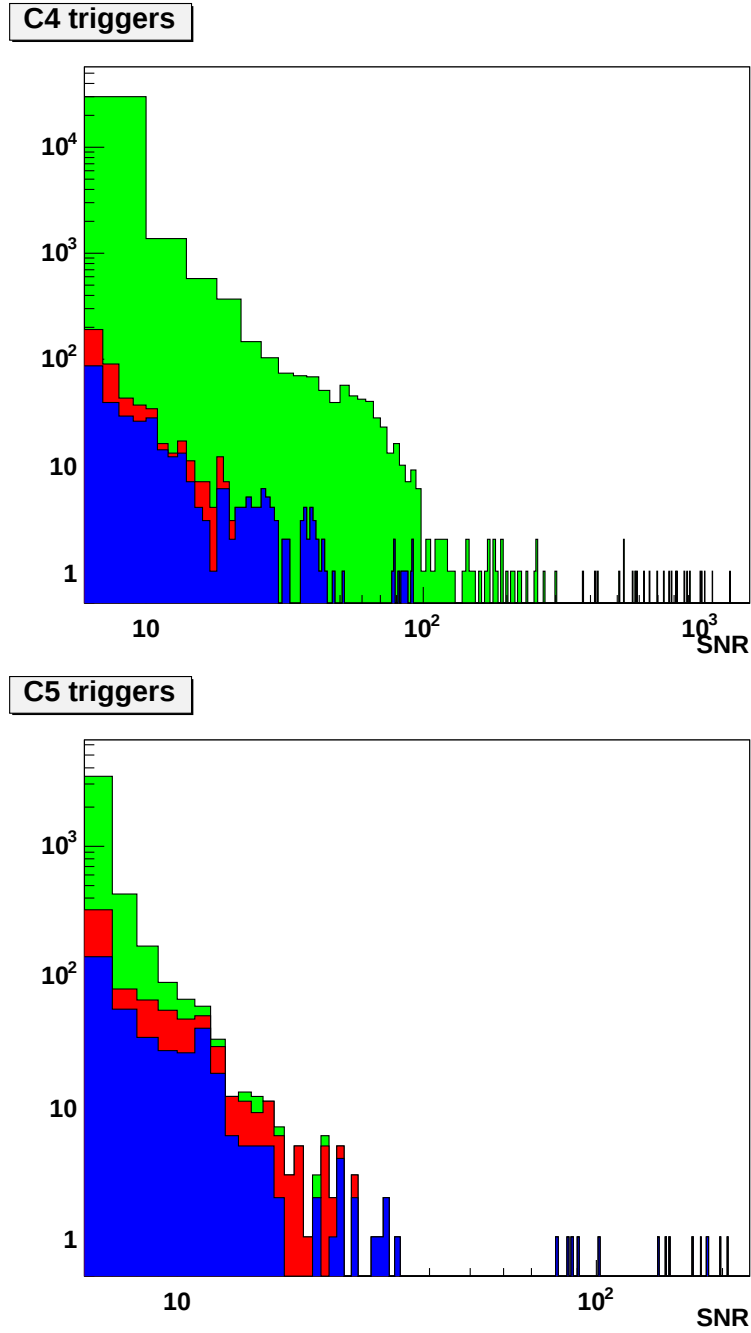


FIG. 7.5 – Distributions de SNR des candidats sur les analyses complètes. La part des candidats situés à moins de 10 s d’une injection de binaire (ou de burst) apparaît en bleu (ou en rouge,) la partie verte restant correspondant donc aux faux événements. En haut, C4, en bas C5.

Les distributions de SNR des candidats obtenus sur toute la durée des deux *runs* sont présentées figure 7.5. Les deux distributions sont très différentes. Après exclusion de la part dues aux injections d'événements impulsifs (*bursts*) et de coalescences binaires, les excès d'événements sont très important dans C4, avec des SNR dépassant 1000. Ils sont au contraire relativement rares dans le cas de C5 et ne dépassent pas quelques dizaines. Cette différence s'explique par l'utilisation dans le cas de C5 de véto en amont de la recherche de signaux, interdisant notamment l'analyse des données:

- Lorsque le label de qualité de la reconstruction (cf. 4.4.2) n'est pas au maximum, ce qui exclue en particulier les débuts de segments de *lock* ou la reconstruction n'est pas stabilisée, et les périodes où les lignes de calibration ne ressortent pas suffisamment du bruit.
- Lors de la manipulation d'éléments optiques de la détection à l'aide des "picomoteurs".
- Lorsque en salle de contrôle le bouton "science mode" qui indique l'absence de toutes interventions humaines n'est pas enclenché.

Cette comparaison montre donc avant tout l'importance de ces vetos qui améliore considérablement la sensibilité - au sens large - du détecteur.

7.2.2 Période silencieuse

Pour caractériser plus précisément les données plus propres de C5, on s'intéresse à une période "silencieuse" constituée de 5 heures de prises de données en continu lors de la première nuit. On profite ainsi d'un faible risque de perturbations humaines, et de plus cette période est exempt d'injections, ce qui supprime tout doute sur la possibilité d'une exclusion insuffisante de ces dernières.

Pour cette période, une nouvelle analyse a été effectuée sur une espace des paramètres moins restreint. On s'attend en effet à ce que les *templates* plus court associés aux masses supérieures à 2 ou 3 $M_{\odot}$ réponde plus facilement aux bruits instrumentaux. L'analyse restreinte autour de $(1.4M_{\odot}, 1.4M_{\odot})$ risque donc de donner une image trop favorable du bruit du détecteur. La nouvelle configuration, allant de 1 à 5 $M_{\odot}$, est plus conforme à ce que pourrait être une première recherche astrophysique avec Virgo. La figure 7.6 montre la distribution de SNR des candidats sur cette période et leur répartition dans le temps. Peu de candidats sont détecté au dessus de 8, valeur standard de coupure pour une recherche de signaux; ces candidats légèrement excessifs sont réparti dans le temps de façon homogène, sans qu'aucune zone d'accumulation ne ressorte nettement.

Cette période sélectionné pour ses bonnes conditions montre donc un détecteur très propre, et renforce la validité de la sensibilité mesurée.

Elle permet également de confirmer partiellement - et avec peu de statistiques - l'idée que les *templates* plus court engendre une plus grande part des événements en excès (figure 7.7). Ils sera donc intéressant de voir si la fréquence de

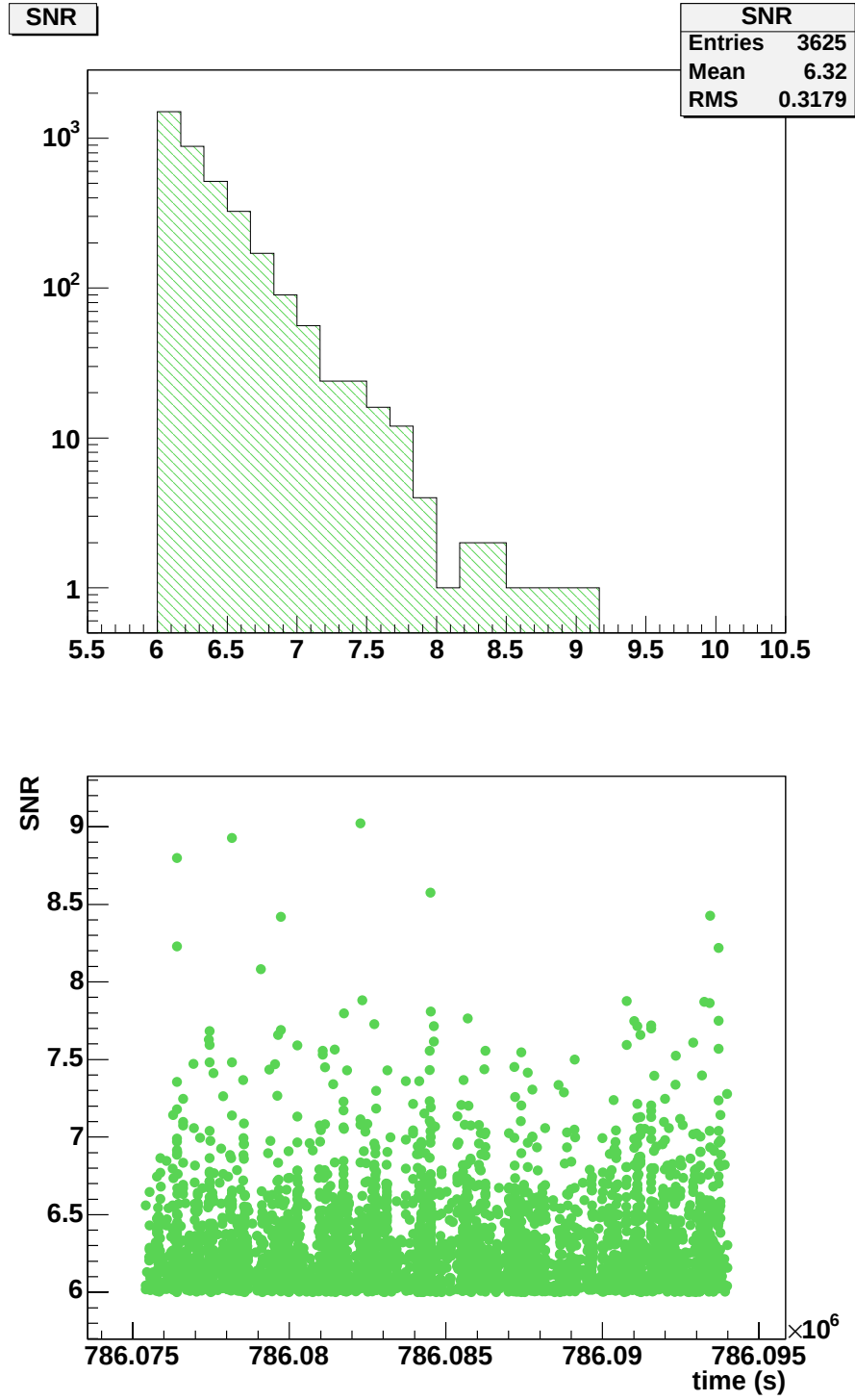


FIG. 7.6 – SNR des candidats sur la période silencieuse analysée entre 1 et 5 M_{odot} (distribution de SNR en haut et SNR en fonction du temps en bas).

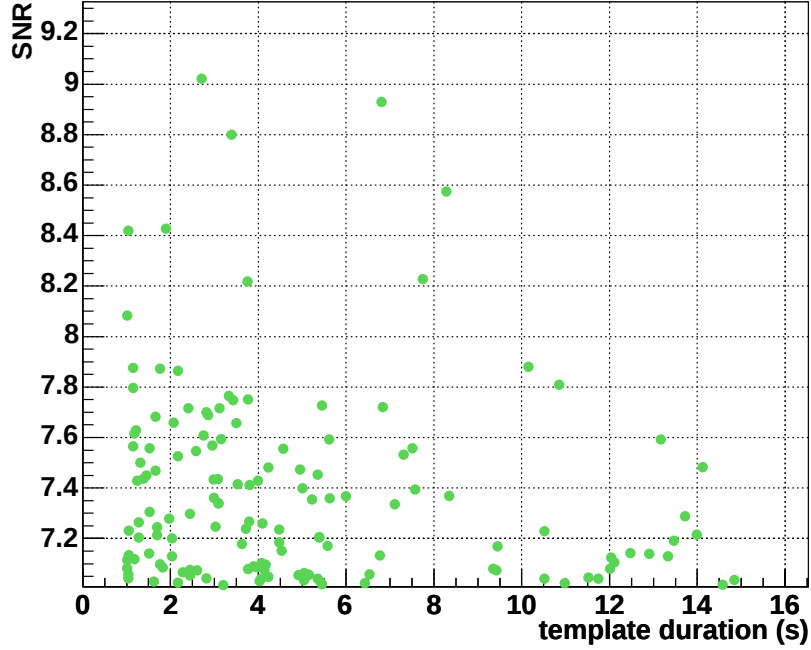


FIG. 7.7 – *Candidats en excès (>7). SNR en fonction de la durée du template*

ces candidats diminue à l'avenir, lorsque la bande passante de Virgo et donc la longueur des *templates* augmenteront. Cet effet peut cependant s'avérer préoccupant, si l'on considère les coalescences de trous noirs comme les candidats les plus prometteurs, car leur recherche risque en effet d'être plus sensible aux bruits instrumentaux non-gaussien. Ainsi, une recherche sur la même période silencieuse pour une coalescence de trous noirs, à $(10M_{\odot}, 10M_{\odot})$ avec un *template* EOB, produit des candidats sensiblement différents (figure 7.8), dont le plus fort atteint un SNR de 10.2, donc supérieures aux valeurs précédentes.

7.3 Stabilité du détecteur

7.3.1 Distance maximale de détection

Après les événements instrumentaux non-gaussien, l'autre différence du bruit de Virgo par rapport à celui de la simulation est sa non-stationnarité, c'est à dire les variations au cours du temps de la sensibilité du détecteur.

Pour une recherche de coalescences binaires, la sensibilité à un instant donné² est essentiellement représenté par la distance maximale de détection (eq. 5.14),

2. plus justement pour une période donnée, assez longue au regard du signal recherché

seule grandeur caractéristique du détecteur dont dépend le taux d'événements attendu.

Le suivi de cette distance est effectué pour chaque *template* dans la recherche multi-bande, à une échelle de temps qui est celle de la mise à jour des PSDs. Dans toutes les analyses décrites jusqu'ici celle-ci est de 30 min, choisie pour être sur que les PSDs ne soit pas bruitée. .

Les figures 7.9 et 7.10 montre l'évolution au cours des *run* C4 et C5 de la distance maximale de détection, en moyenne sur tous les *templates* . Celle-ci est dans l'ensemble environ deux fois moins élevé dans C5 que dans C4, cette baisse de sensibilité³ était d'ailleurs déjà visible sur les courbes issue de l'étalonnage en début de *run* (figure 7.1).

Les variations sur toute la durée du *run* sont également importantes, jusque à rendre insignifiantes certaines périodes où la sensibilité chute en dessous de la moitié de sa valeur initiale, soit - en première approximation - près d'un ordre de grandeur en nombre d'événements détectables (0.5^3). Elles ne correspondent pas de façon évidente au cycle jour/nuit comme c'est souvent le cas dans les détecteurs plus proches de leur sensibilité nominale. Des variations peuvent être observées d'un segment de lock à un autre, mais aussi sur la durée d'un seul segment. Les long segments stables sont plus rares dans C5, en particulier à cause des lignes de calibration souvent trop faibles, qui déclenchent le veto sur l'analyse.

En se penchant plus particulièrement sur la première nuit de C5 (période silencieuse), on constate (figure 7.11) que la distance maximale de détection reste stable, à part de faibles oscillations qui mettent en évidence la constante de temps du suivi (30 mn). On peut donc conclure comme pour les faux événements que le détecteur se comporte actuellement de façon satisfaisante sur une période choisie, mais pas sur l'ensemble d'un *run*.

7.3.2 Effets sur les grilles de *templates*

Un autre manière d'évaluer la stabilité du détecteur sur la durée d'un *run* est de mesurer les impacts de l'évolution de la sensibilité sur le placement des *templates* d'une grille dans l'espace $(\tau_0, \tau_{1.5})$. On a deux raisons pour le faire:

- On peut suivre ainsi un type particulier de variations du spectre, la grille n'étant sensible qu'aux variations de forme et non aux variations d'échelle absolue.
- Contrairement à ce qui est fait dans la recherche multi-bande, dans les chaînes d'analyse d'autres expériences comme LIGO et TAMA, la grille est régulièrement reconstruites pour tenir compte des changements de sensibilité; les variations observés du nombre de *templates* atteigne alors parfois

3. principalement due à la baisse de puissance dans l'interféromètre entre les deux *runs* après l'installation d'un atténuateurs entre l'injection et le reste de l'interféromètre en réponse à des problèmes de lumière rétro-diffusée par le miroir de recyclage.

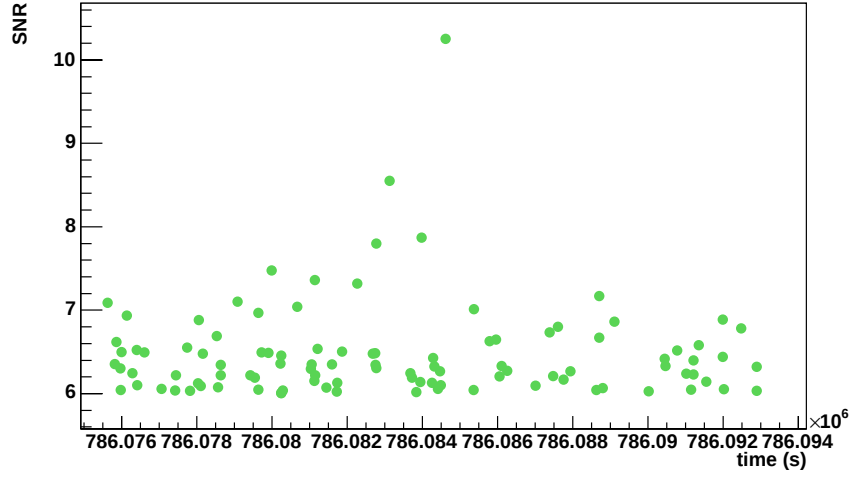


FIG. 7.8 – *SNR des candidats issue d'une analyse sur la même période silencieuse pour une coalescence de trous noirs, à $(10M_{\odot}, 10M_{\odot})$ avec un template EOB. SNR en fonction du temps.*

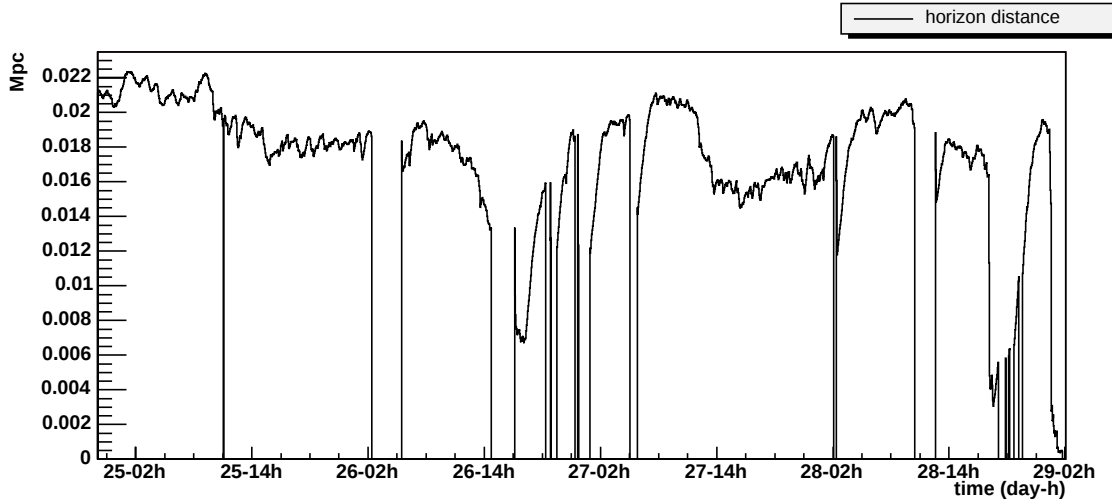


FIG. 7.9 – *Évolution au cours du run C4 de la distance maximale de détection, en moyenne sur tous les templates de l'analyse.*

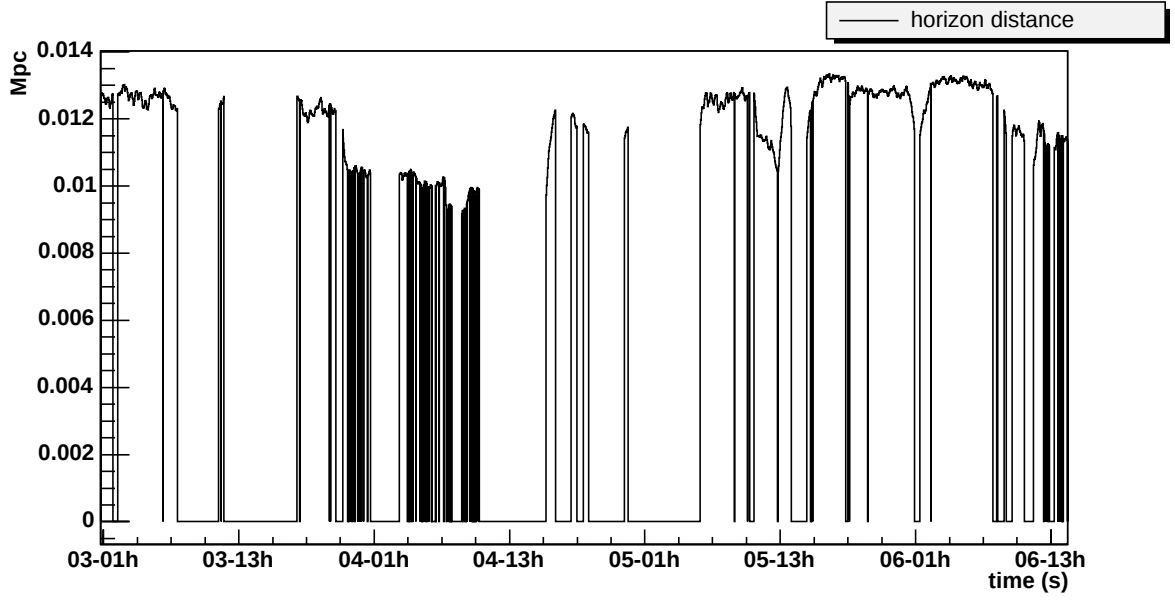


FIG. 7.10 – Évolution au cours du run C5 de la distance maximale de détection, en moyenne sur tous les templates de l'analyse.

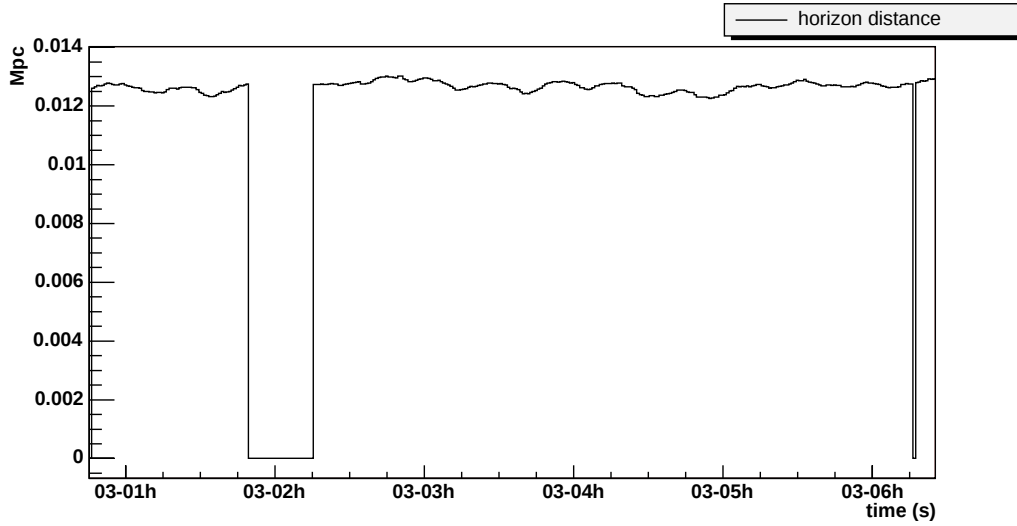


FIG. 7.11 – Évolution au cours de la première nuit du run C5 de la distance maximale de détection, en moyenne sur tous les templates de l'analyse.

un facteur 2. Dans notre cas où les grilles ne sont construites qu'une seule fois au moment de l'initialisation problématique, un tel impact serait problématique, les grilles initiales utilisées devenant rapidement très différentes de celle qu'il *faudrait* utiliser.

Pour savoir ce qu'il en est, on doit mesurer les effets des variations de sensibilité au cours d'un *run* sur le calcul de la grille. On peut pour cela suivre l'évolution de l'ellipse couverte par un *template* donné. Ce suivi a été effectué à posteriori sur les données des deux *runs*, en considérant l'ellipse couverte par un *template* de masses ($1.4M_{\odot}, 1.4M_{\odot}$) dans l'espace $(\tau_0, \tau_{1.5})$; les ellipses successives sont calculés à l'aide d'une PSD glissante et de la librairie LAL, également utilisée par le code de calcul des grilles. Les figures 7.12 et 7.13 montrent cette évolution à travers les trois paramètres: grand axe, petit axe, et angle donnant l'orientation de l'ellipse. On constate ainsi que:

- les ellipses et donc les grilles sont moins sensibles aux variations de sensibilité du détecteur que sa distance maximale de détection, puisqu'elles sont globalement plus stables.
- La forte corrélation entre petit axe et grand axe indiquent que les ellipses ne se déforment pas à mesure qu'elles changent de tailles.
- Les variations les plus importantes sont observées pour C4, et sont associées à des débuts ou fin de segments de *lock* et correspondent à des périodes où la sensibilité est de toute façon mauvaise⁴. En dehors de ces périodes les variations de tailles des ellipses sont de l'ordre de 10% à 20%.
- Toujours dans C4, des variations résiduelles plus importantes et soudaines (pics) de l'orientation de l'ellipse ne sont pas comprises.
- L'ellipse est plus stable pour C5, les différents paramètres se maintenant systématiquement à $\pm 5\%$, même pour les plus grandes variations de sensibilité.

Dans l'ensemble les variations observées sont plus faibles que ce qui était craint au vu des premiers résultats des expériences LIGO et TAMA. De plus il n'y a pas de déformations majeures des ellipses, on doit donc pouvoir trouver une "plus petite ellipse" - de toutes les sensibilités possibles - qu'il suffirait d'utiliser au départ pour être sûr de ne jamais avoir de trous dans la grille. Les pics de l'angle d'orientation vus dans C4 restent par contre à surveiller s'ils réapparaissent au cours de prochains *runs*.

Pour s'assurer que ces résultats ne dépendent que faiblement du *template* considéré, le suivi pour C5 a été effectué à nouveau dans le cas d'un *template* de masses ($10M_{\odot}, 10M_{\odot}$). On vérifie ainsi (figure 7.14) que les ordres de grandeurs des variations de l'ellipse ne changent pas même pour ce *template* très éloigné du précédent.

4. les périodes du même type sont d'ailleurs exclues dans C5 par la mise en place des vetos

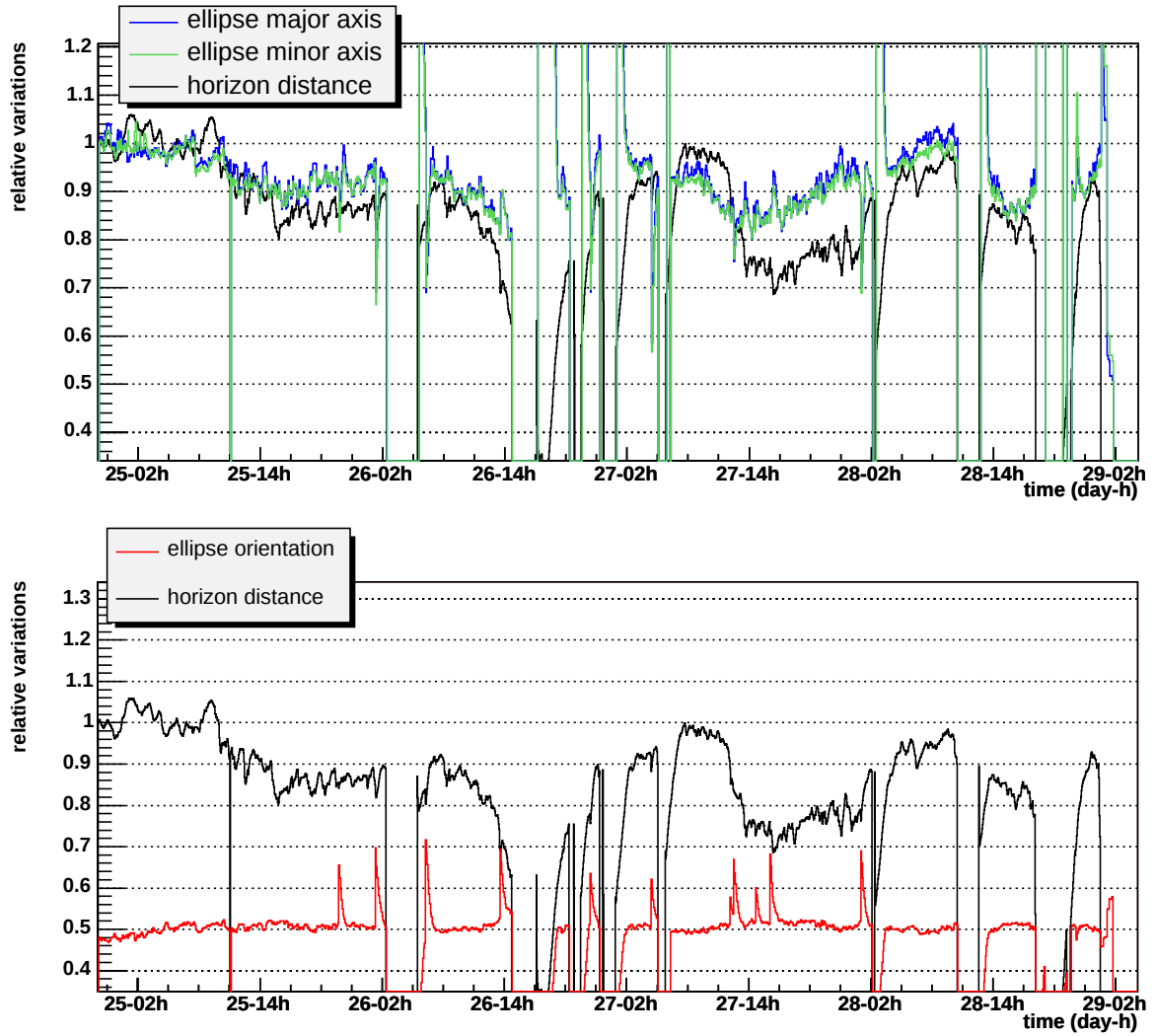


FIG. 7.12 – Évolution au cours du run $C4$ de l'ellipse couverte par un template à $(1.4M_{\odot}, 1.4M_{\odot})$: Variations relatives du grand axe (en haut, en bleu) et du petit axe (en haut, en vert), angle donnant l'orientation de l'ellipse (en bas, en rouge), ainsi que de la distance horizon (en haut et en bas, en noir)

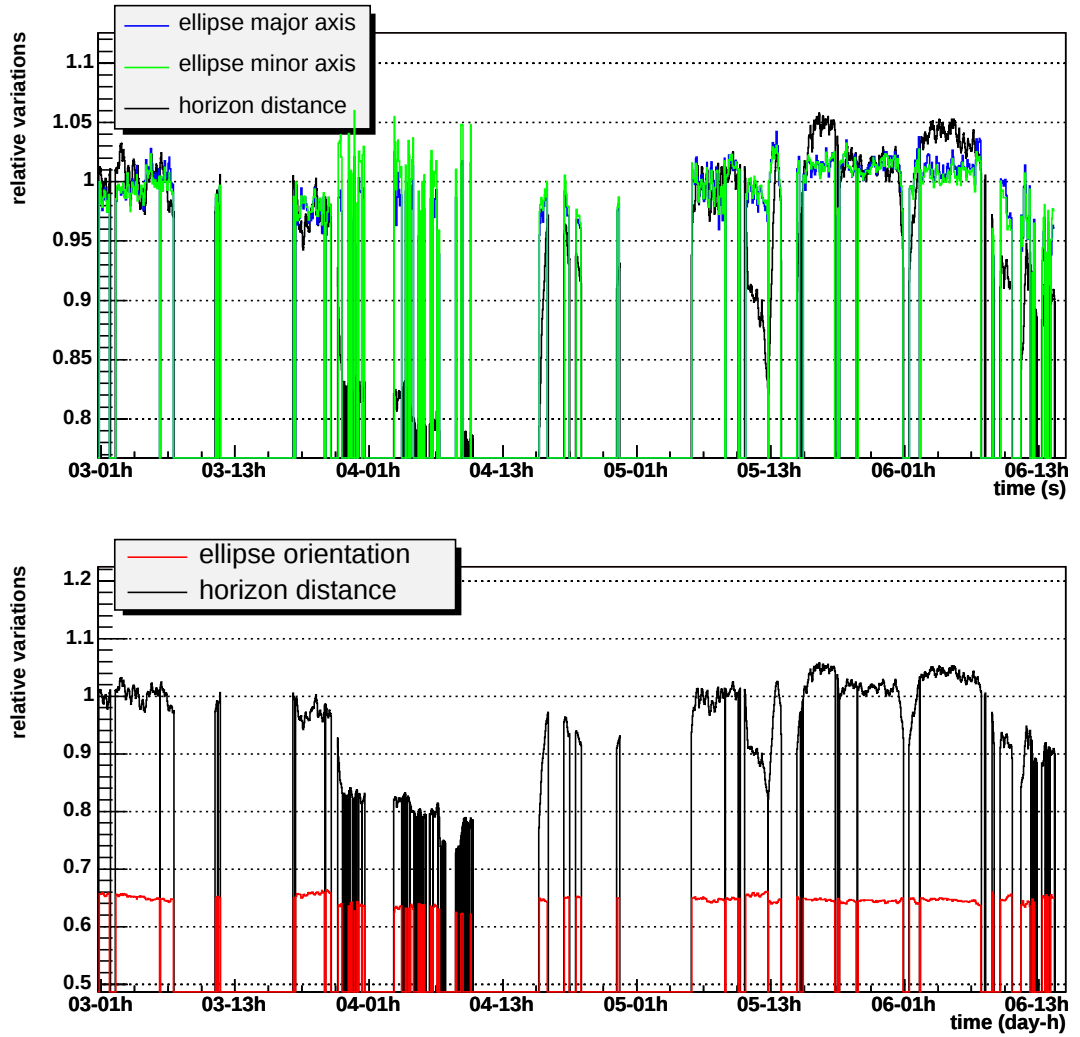


FIG. 7.13 – Évolution au cours du run C5 de l'ellipse couverte par un template à $(1.4M_{\odot}, 1.4M_{\odot})$: Variations relatives du grand axe (en haut, en bleu) et du petit axe (en haut, en vert), angle donnant l'orientation de l'ellipse (en bas, en rouge), ainsi que de la distance horizon (en haut et en bas, en noir)

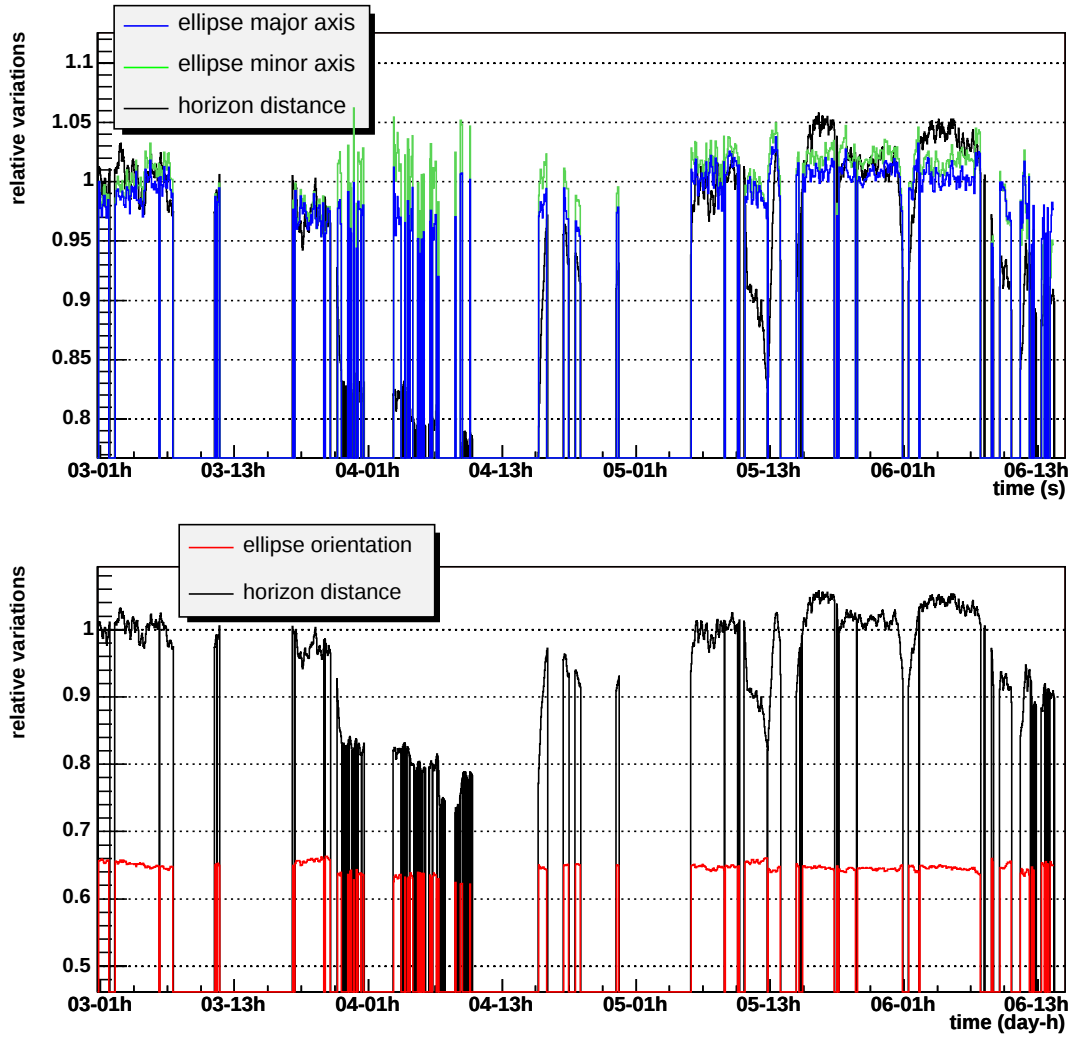


FIG. 7.14 – Évolution au cours du run C5 de l'ellipse couverte par un template à $(10M_{\odot}, 10M_{\odot})$: Variations relatives du grand axe (en haut, en bleu) et du petit axe (en haut, en vert), angle donnant l'orientation de l'ellipse (en bas, en rouge), ainsi que de la distance horizon (en haut et en bas, en noir)

Une étude plus poussée est envisageable, consistant à mesurer régulièrement par MC la perte de SNR en chaque point de l'espace des masses avec la grille initiale, pour différentes courbes de sensibilités correspondant à différents moments d'un *run*. Cependant les résultats de cette première analyse sont des premiers indices indiquant que le choix fait dans MBTA de garder la grille initiale pour toute l'analyse ne devrait pas poser de problèmes.

7.4 Perspectives: Vetos

Cette première analyse a permis de s'assurer du bon fonctionnement de la recherche multi-bande sur les données actuelles de l'interféromètre, avec la détection des signaux injectés.

Elle a également permis d'offrir une première caractérisation du bruit en montrant un détecteur qui peut-être relativement propre et stable, mais sur une période sélectionnée et après prises en compte de certains vetos liés à l'état du détecteur, qui permettent d'exclure une grande partie des "faux événements" engendrés par l'instrument.

C'est donc vers cette principale perspective que doit se tourner la suite de ce travail: la définition et la mise en oeuvre d'un nombre suffisant de vetos pour nettoyer la distribution de SNR de ces faux événements et permettre ainsi une détection optimale.

Une première catégorie de vetos, dans la lignée de ceux déjà utilisés, consiste à détecter un possible événement instrumental à l'aide des canaux auxiliaires de l'interféromètre qui informent sur l'état des différents sous-systèmes. Ce type de veto implique:

- l'identification préalable des cause des faux événements, pour que le veto soit efficace.
- de pouvoir garantir que le canal observé ne réagira pas également en cas de vrai détection, pour qu'il soit sure.

et nécessite donc une étude approfondie.

Plus éloignée de l'instrument et plus proche du filtrage adapté, l'autre catégorie de vetos vise à tester la cohérence de l'événement détecté vis à vis du signal attendu.

Le plus classique [39] consiste en un test de χ^2 qui compare la répartition du SNR de détection par bandes de fréquences à celle attendue; on élimine ainsi les bruits localisés en fréquence. Ce test est déjà naturellement implémenté dans la recherche multi-bande, mais n'a pas encore fait l'objet d'un étude approfondie. On s'attend de plus à devoir l'étendre à un plus grand nombre de bandes (de l'ordre d'une dizaine) pour qu'il devienne le plus sensible possible. Cependant un premier aperçu de la répartition du SNR sur les deux bandes utilisées pour C5 (figure 7.15) est déjà prometteur: avec une coupure entre les deux bandes choisie pour un

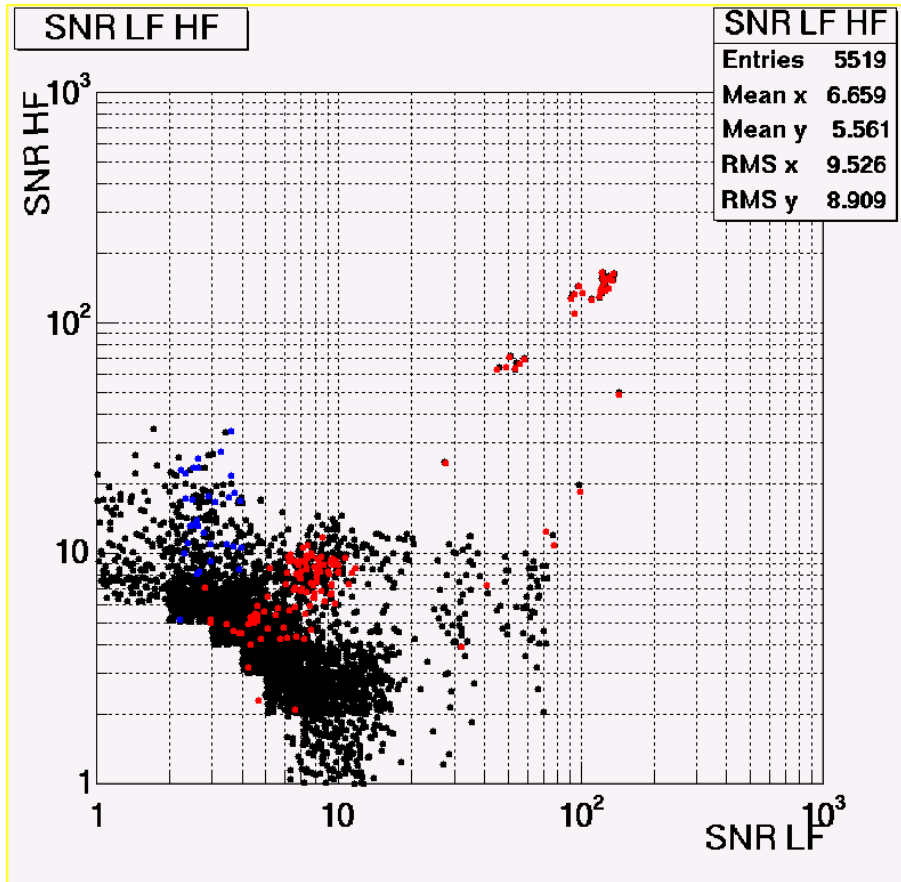


FIG. 7.15 – *Candidates de C5: SNR sur les bandes basses fréquences (axe X) et hautes fréquences (axe Y). Les candidats en rouge sont associés à des injections de coalescences, en bleu à des bursts.*

SNR optimal identique dans chaque bande, on constate que les candidats associés à des injections de coalescences binaires occupent essentiellement la ligne d'équi-répartition du SNR, tandis que ceux déclenchés par les événements impulsifs sont massivement localisés sur la bande supérieure, et qu'une partie du bruit en excès est au contraire concentrée dans la bande inférieure.

Une alternative peut-être plus simple dont la mise en oeuvre a débuté récemment consiste à tester la répartition du SNR de détection sur les différents *template* de la grille; les bruits localisés en fréquence devraient plus "s'étaler" sur la grille, puisqu'il voient moins les différences entre *templates*.

Enfin un troisième veto, développé au sein de la collaboration, compare la forme du pic de détection - dans le signal en sortie du filtrage adapté - à celle attendue pour un vrai signal. Il doit donc pouvoir être appliqué au signal recombinaison des *template* virtuels de la recherche multi-bande.

Il reste donc un travail important avant une recherche astrophysique; il béné-

ficiera dans un avenir proche des plus longues durées - quelques semaines - qui caractériseront les nouvelles prises de données. L'impact positif de l'utilisation des premiers vetos entre C4 et C5 ainsi que la propreté de la distribution de SNR sur la période silencieuse de C5 laisse néanmoins espérer une progression rapide.

Chapitre 8

Comparaison avec la chaîne d'analyse de LIGO

La recherche multi-bande a également été utilisée dans le cadre de la collaboration avec l'expérience LIGO, pour préparer la recherche de coalescences binaires lors des futures prises de données en coïncidence des deux interféromètres LIGO et de Virgo. La détection de binaires avec un réseau LIGO-Virgo présente un double intérêt:

- 3 détecteurs mesurant chacun un temps d'arrivée et une distance effective permettent de reconstruire la position de la source dans le ciel.
- La détection en coïncidence sur plusieurs détecteurs permet de rejeter des artefacts instrumentaux qui pourrait simuler un faux signal de binaire sur un seul d'entre eux.

Selon toute vraisemblance, une telle analyse commencerait par l'échange des candidats - au-dessus d'un certain seuil en SNR - de chaque détecteur, tels qu'ils sont produits par les chaînes d'analyse respective des deux expériences. Le premier projet du groupe de travail LIGO-Virgo pour les coalescences binaires a donc consisté à comparer ces deux chaînes d'analyses. Cette comparaison s'est faite au cours d'un premier MDC commun, sur des données simulés par les deux collaborations, incluant bruit et injections de signaux; elle portait sur la capacité des chaînes à retrouver les événements injectés, mais aussi sur l'estimation de leurs différents paramètres: temps d'arrivée, distance effective, masses de la sources.

Ce premier MDC LIGO-Virgo s'est déroulé durant l'automne 2004; la simulation des données et leur traitement par les chaînes d'analyses des deux collaborations est décrite dans la première partie de ce chapitre.

Le travail de comparaison des candidats produits a été effectué dans la foulée pour être présenté en décembre 2004 (GWDAW9)[40]. Les résultats - réactualisés pour prendre en compte les modifications les plus récentes - sont présentés dans la deuxième partie.

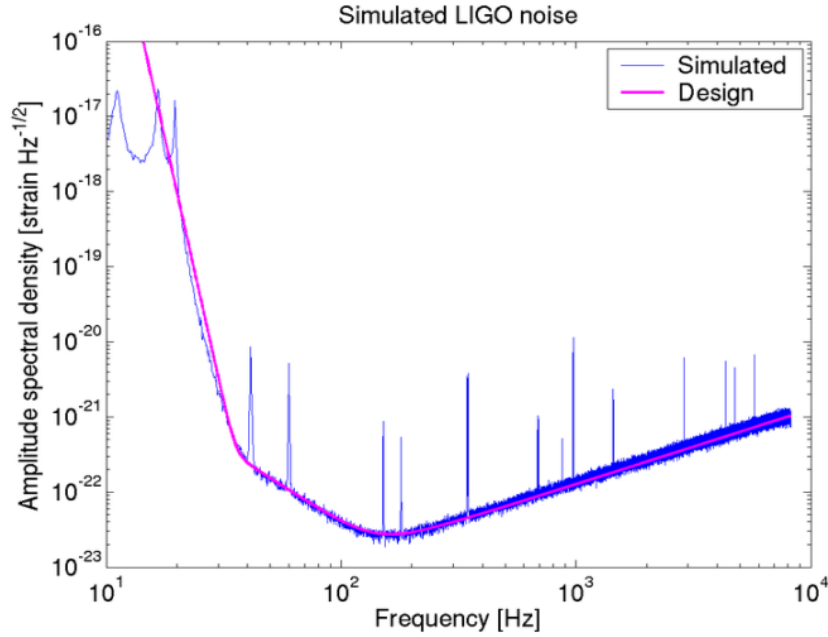


FIG. 8.1 – *Spectre du bruit simulé de LIGO (en violet) comparé à la sensibilité théorique (en rose).*

8.1 Le premier MDC LIGO-Virgo

8.1.1 Les deux jeux de données simulés

Pour cette première comparaison, des données ont été simulées par les deux collaborations, chacune simulant son propre détecteur à sa sensibilité nominale. On obtient ainsi deux jeux indépendants de 3 h de données chacun.

Du côté de Virgo, le bruit est simulé par SIESTA avec une configuration identique à celle présentée au chapitre 6. Des signaux de coalescences binaires sont également injectés aléatoirement à raison d'un événement toutes les cinq minutes en moyenne. Il s'agit de formes d'onde calculées à l'ordre PN2, pour l'unique couple de masses ($1.4M_{\odot}, 1.4M_{\odot}$), toujours à partir de 24 Hz et jusqu'à la dernière orbite stable. Ils sont tous normalisés pour un SNR de 10, soit une distance - pour des coalescences optimalement orienté - de 24.8 Mpc entre la source et le détecteur.

Le bruit produit par LIGO est lui aussi obtenu par une simulation rapide visant à reproduire la sensibilité nominale (figure 8.1). Les événements sont au nombre de 26 (environ un toutes les 400 s), identiques à ceux du jeu Virgo, à ceci près que deux couples de masses sont utilisés: ($1.4M_{\odot}, 1.4M_{\odot}$) et ($1M_{\odot}, 2M_{\odot}$). Ils sont normalisés en distance avec plusieurs valeurs possible: 20, 25, 30, ou 35 Mpc, qui correspondent à des valeurs de SNR entre 7 et 13. Ce MDC sert également de test technique sur la possibilité d'échanger des données, et la maniabilité par

chacun des données issues de l'autre collaboration. C'est pourquoi les deux jeux de données étaient stockés dans le format *frame* commun aux deux expériences, et respectivement échantillonnés aux fréquences des expériences: 16384 Hz pour LIGO et 20000 Hz pour Virgo. Seule concession, la longueur des *frames* était 60 s - longueur utilisée par LIGO - pour les 2 jeux de données, en raison de difficultés techniques pour LIGO à stocker des *frames* de 1 s.

8.1.2 Les deux chaînes d'analyse

Ces deux jeux de données doivent ensuite être analysés afin de comparer les deux chaînes d'analyses. Elles reposent toutes deux sur le principe de filtrage adapté et utilisent des *templates* à l'ordre PN2 - comme les signaux injectés.

La chaîne d'analyse de la LSC est décrite dans [41]. Il s'agit en résumé d'une chaîne de filtrage adapté standard, avec grille de *templates*. Contrairement à la recherche multi-bande elle n'est pas conçue pour analyser les données en ligne. L'ensemble des données à analyser sont séparées en blocs de 15 segments de données se chevauchant partiellement, dont la durée est au minimum quatre fois la durée du plus long *template* de la grille. Chaque bloc est analysé indépendamment - éventuellement en parallèle:

1. Estimation du spectre sur le bloc entier. Avant utilisation, le spectre est partiellement tronqué dans le domaine temporel, dans le but de réduire les effets relatifs aux tailles et positions de fenêtres - effets observés dans le cas de MBTA mais jugés négligeables, cf §6.1.2.
2. Définition de la grille dans l'espace des paramètres; elle est calculée dans l'approximation d'une métrique uniforme, dont la valeur est calculée au point du *template* le plus long (dont l'ellipse de couverture est la plus petite).
3. Création de la banque de *templates* correspondante. Le calcul se fait directement dans le domaine fréquentiel, via l'approximation de phase stationnaire.
4. Filtrage adapté des données par les différents *templates*, au moyen habituel de FFTs segments par segments.
5. Enregistrement pour chaque *template* des candidats correspondant au dépassement d'un seuil donné en SNR. Un regroupement est effectué *template* par *template*, qui revient à considérer deux candidats séparés par moins que la durée du *template* comme un candidat unique. Par contre aucun regroupement n'est effectué entre candidats issues de différents *templates* de la grille.

Pour la collaboration Virgo, deux chaînes de filtrage adapté sont en développement: la chaîne parallélisée *Merlino* [42], et la recherche multi-bandes, largement détaillée dans les chapitres précédents. Pour cette comparaison c'est essentiellement cette dernière qui fut utilisée¹.

1. Quelques résultats préliminaires obtenus avec *Merlino* vont globalement dans le même sens que ceux qui seront présentés ici avec la recherche multi-bande

8.1.3 La production des candidats

Ce MDC a donné lieu à quatre analyses distinctes, chaque collaboration analysant chacun des deux jeux de données.

La recherche multi-bande est effectuée en deux bandes; au-delà de cette spécificité, les principaux paramètres communs des analyses sont fixés au préalable, pour permettre une comparaison directe des candidats:

- Les banques de *templates* sont calculées à l'ordre PN2 (comme les injections)
- L'espace de masses couvert va de 1 à 3 $M_{\odot}$
- Le recouvrement minimal pour définir la grille est de 95%
- L'analyse commence à 40 Hz pour les données LIGO et 30 Hz pour les données Virgo, soit un *template* le plus long de ~ 45 s dans le premier cas et ~ 96 s dans le deuxième.
- Les candidats sont sélectionnés pour un SNR supérieur à 6, coupure permettant - en simulation avec du bruit gaussien - d'éliminer la totalité ou presque de la distribution de bruit (Figure 6.4).

Malgré ces paramètres communs, de nombreuses différences - outre le caractère particulier de la recherche multi-bande - subsistent entre les analyses.

D'une part, les grilles sont générées par des méthodes différentes; l'approximation d'une métrique constante utilisée pour la définition de la grille de LIGO engendre une grille largement sur-couverte, l'espacement entre *templates* étant imposé sur tout l'espace des paramètres par la plus petite ellipse de couverture. Sur les données LIGO, la grille de la LSC contient ~ 3500 *templates*, à comparer aux ~ 1900 *templates* virtuels de la recherche multi-bande. De même sur les données Virgo, où la chaîne de la LSC utilise ~ 9500 *templates* alors que la recherche multi-bande n'en a que ~ 6100 .

D'autre part, de par la nature des algorithmes, la répartition de chaque analyse en plusieurs processus se fait de façons différentes. Dans le cas de la LSC on découpe les données en blocs proportionnels à la durée du *template*. Il y a donc plus de processus (6) pour les données LIGO que pour les données Virgo (3). Pour la recherche multi-bande on découpe - par valeurs de la *chirp mass* - des portions d'espaces des paramètres dans l'espace total à couvrir. Chaque portion doit contenir un nombre de *templates* n'excédant pas la mémoire disponible. Le nombre de processus augmente donc avec le nombre total de *templates*, et donc avec la bande de fréquence d'analyse: 7 pour les données Virgo, un seul pour les données LIGO.

8.2 Comparaison des résultats

8.2.1 Injections détectées

Pour comparer les listes de candidats, on doit tout d'abord reconnaître ceux qui correspondent à des injections. Dans chaque liste, on considère comme "vrai" tout candidat dont le temps de fin correspond au temps de fin d'une injection à $\pm 20ms$ près; cette durée est choisi pour englober largement les résolutions temporels respectives des deux chaînes d'analyse.

On vérifie ainsi que les deux chaînes détectent les mêmes injections:

- Toutes les injections (11/11) dans le cas des données Virgo
- 25 injections sur 26 dans le cas des données LIGO; l'événement manquant, qui est le même pour les deux chaînes, est injecté à une distance de 35 Mpc, soit un SNR de 6.4, proche du seuil de déclenchement de 6 adopté pour toutes les analyses.

Les listes de candidats produites par la recherche multi-bande et la chaîne d'analyse de LIGO sont néanmoins assez différentes, comme le montrent les distributions de SNR de chaque analyse sur les données LIGO (figure 8.2) et Virgo (figure 8.3). Ces différences ne sont pas dues au fausses alarmes - qui présentent des distributions en SNR similaires d'une analyse à l'autre - mais aux "vrai" candidats, à cause des différences dans les algorithmes de regroupement des candidats des deux chaînes. Ils sont beaucoup plus nombreux dans les listes de la chaîne de LIGO pour laquelle les candidats issue de *templates* distincts mais associé à une même injection ne sont pas regroupés. Pour les listes de la recherche multi-bande ils sont aux nombres de un candidats pour une injection sur les données LIGO, un peu plus sur les données Virgo pour lesquels le réglage initial de la largeur de regroupement lors de la production était trop faible.

Pour pouvoir comparer un à un les paramètres des "vrais" candidats, on choisit, pour chaque injection, le candidat correspondant de chaque analyse offrant le meilleur SNR.

8.2.2 Rapport signal sur bruit

Pour comparer avec plus de précision - dans la limite de la statistique disponible - l'efficacité de détection des deux chaînes, on commence par comparer les valeurs de SNR avec lesquelles les différentes injections sont détectés.

Les figures 8.4 et 8.5 montrent les résultats de ces comparaisons pour les données Virgo et LIGO. La comparaison montre une bonne corrélation et un accord général des deux chaînes.

Les écarts moyens observés sont attribués aux différences de grilles, qui peuvent fortement favoriser une analyse ou l'autre au vu de la faible diversité des couples de masses utilisés pour les injections. Le cas des données Virgo constitue un cas

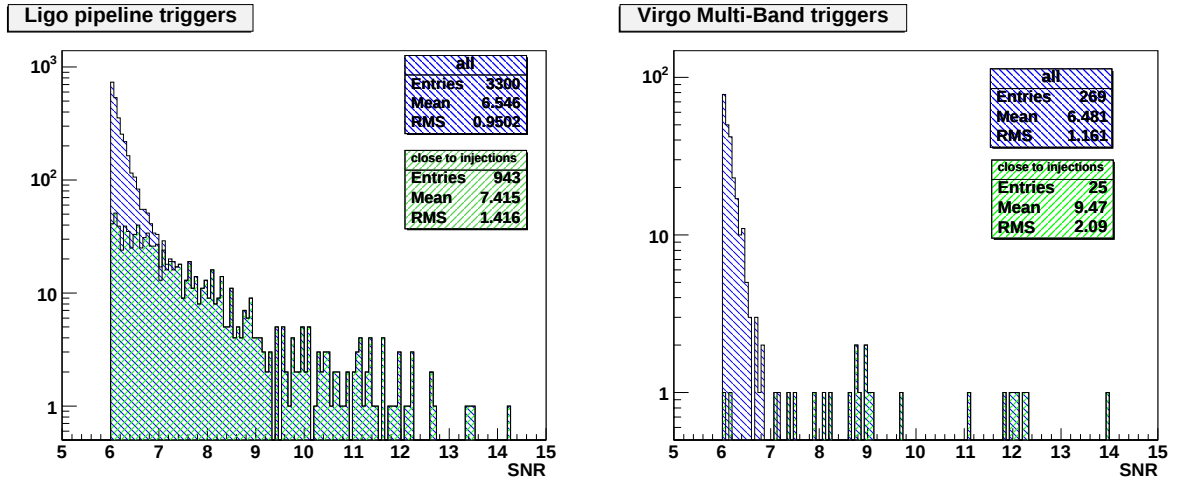


FIG. 8.2 – Distributions en SNR des candidats produits lors de l’analyse des données LIGO par la recherche multi-bande (à gauche) et la chaîne d’analyse de LIGO (à droite). Candidats situés à moins de 20 ms d’une injection en vert, superposé à la distribution totale en bleu.

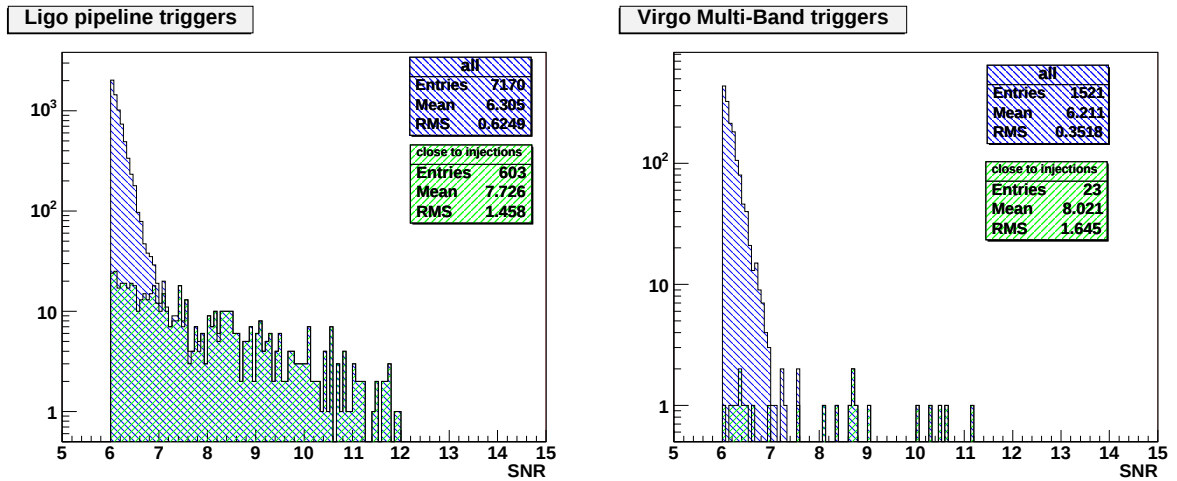


FIG. 8.3 – Distributions en SNR des candidats produits lors de l’analyse des données Virgo par la recherche multi-bande (à gauche) et la chaîne d’analyse de LIGO (à droite). Candidats situés à moins de 20 ms d’une injection en vert, superposé à la distribution totale en bleu.

extrême: alors que la grille de la chaîne de LIGO possède un *template* très près de $(1.4M_{\odot}, 1.4M_{\odot})$, le *template* le plus proche de la grille de *templates* virtuels de la recherche multi-bande donne un recouvrement avec ce signal proche du recouvrement minimal de 95% pour lequel les grilles sont définis.

8.2.3 Distance

La comparaison des distances effectives mesurées par les deux chaînes (figure 8.7 et 8.6) donne des résultats similaires: les valeurs obtenues sont corrélées pour les deux jeu de donnés, et présentent des écarts systématiques compris dans les 5% autorisés par le recouvrement minimal des grilles. Cependant ces écarts ne sont pas cohérent, dans le cas des données LIGO, avec les écarts de SNR: la distance estimée devrait être plus faible pour la recherche multi-bande, qui détecte plus fort les événements. Ce cas reste donc à comprendre.

8.2.4 Temps d'arrivée

Le temps d'arrivée joue un rôle particulièrement important dans une analyse en réseaux, puisque c'est sur la base de ce paramètre qu'une détection en coïncidence peut être identifiée. L'accord des deux chaînes sur ce paramètre constitue donc un pré-requis à une analyse en coïncidence des deux détecteurs.

On a pu voir au chapitre 5 que la dispersion en temps est - au moins dans le cas de la recherche multi-bande - dominée par les effets de grilles. On ne s'attend donc pas à une forte corrélation entre les deux chaînes d'analyse. La comparaison ne montre dans les deux cas (figure 8.8) et 8.9 qu'une très légère corrélation, qui disparaît complètement pour certains événements détectés par des *templates* particuliers.

Sur l'ensemble, les deux chaînes sont en accord à 1 ms ou 2 ms près, ce qui est comparable à leurs résolutions intrinsèques.

8.2.5 Masses de la source

La dernière comparaison effectué porte sur l'estimation des masses de la source. Il s'agit de comparer pour chaque événement détecté les masses des *templates* ayant donné les plus fort SNR. De tels paramètres peuvent être utilisés comme critère de coïncidence entre deux événements lors d'une analyse en réseau.

Les résultats obtenus rappelle ceux du chapitre 5: aucune des deux chaînes n'obtient une bonne précision sur la détermination des deux masses indépendamment l'une de l'autre (figures 8.10 et 8.11) - et aucune corrélation significative n'est visible entre les deux chaînes.

Par contre une bonne précision est systématiquement obtenue pour la *chirp mass*, avec des écarts ne dépassant pas 4.10^{-3} pour les données LIGO (figure 8.13)- qui sont analysés avec des grilles moins fines et des *templates* plus court - et 1.10^{-3}

8.2. Comparaison des résultats

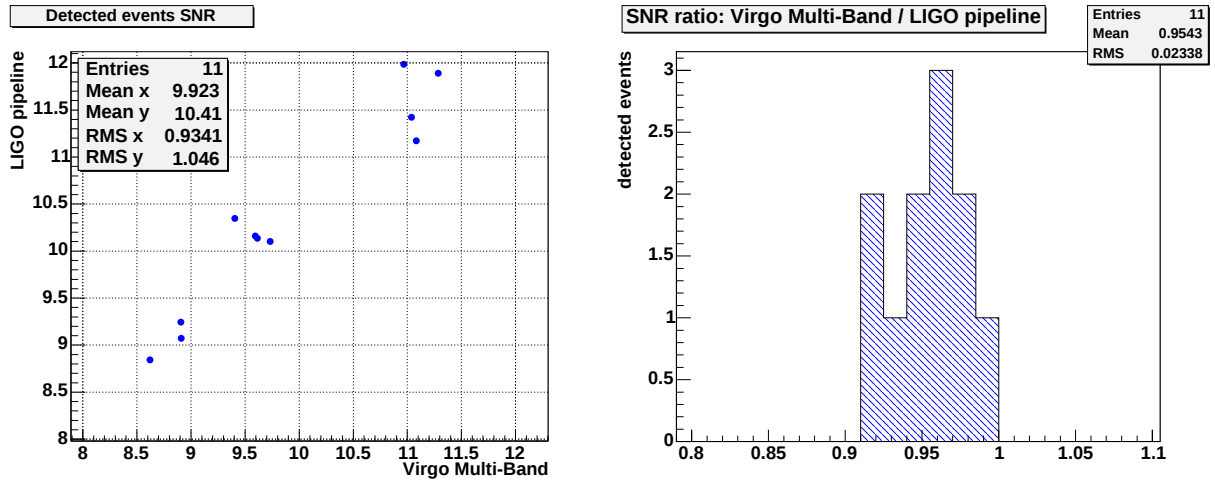


FIG. 8.4 – Comparaison des valeurs de SNR obtenues pour chaque injection par les deux chaînes d'analyses sur les données Virgo. À gauche: SNR de la chaîne de LIGO en fonction du SNR de la recherche multi-bande. À droite: distribution du rapport SNR "multi-bandes" / SNR "chaîne LIGO" .

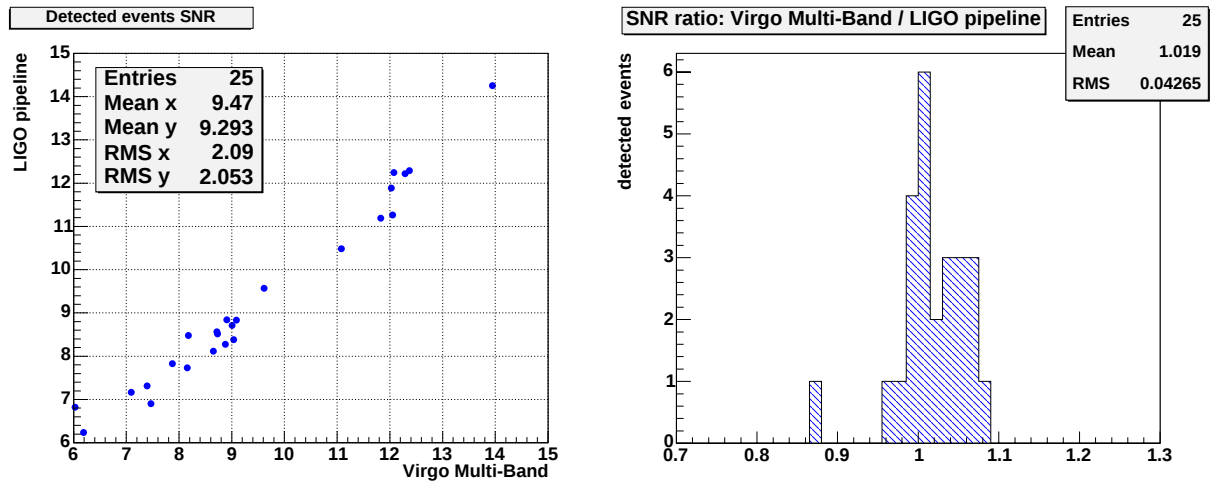


FIG. 8.5 – Comparaison des valeurs de SNR obtenues pour chaque injection par les deux chaînes d'analyses sur les données LIGO. À gauche: SNR de la chaîne de LIGO en fonction du SNR de la recherche multi-bande. À droite: distribution du rapport SNR "multi-bandes" / SNR "chaîne LIGO" .

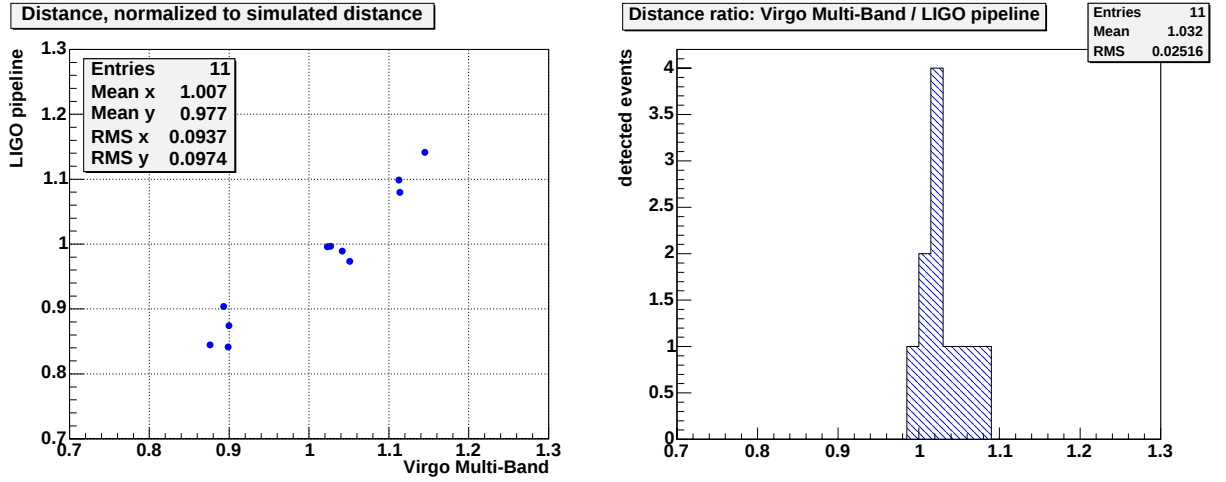


FIG. 8.6 – Comparaison des distances effectives obtenues pour chaque injection par les deux chaînes d'analyses sur les données Virgo. À gauche: distance de la chaîne de LIGO en fonction de la distance de la recherche multi-bande, toutes deux normalisées à la distance de l'injection. À droite: distribution du rapport distance "multi-bandes" / distance "chaîne LIGO".

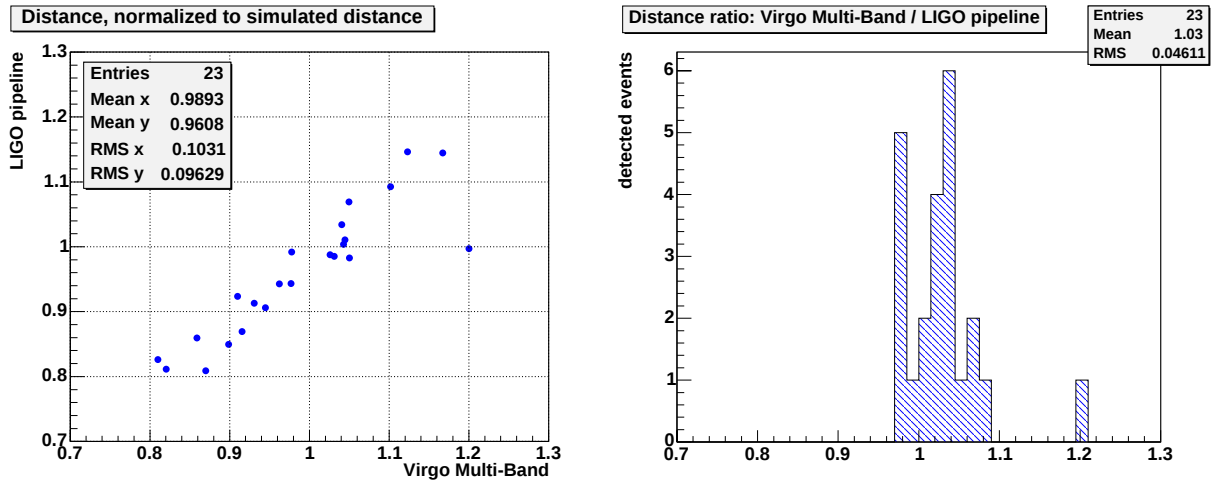


FIG. 8.7 – Comparaison des distances effectives obtenues pour chaque injection par les deux chaînes d'analyses sur les données LIGO. À gauche: distance de la chaîne de LIGO en fonction de la distance de la recherche multi-bande, toutes deux normalisées à la distance de l'injection. À droite: distribution du rapport distance "multi-bandes" / distance "chaîne LIGO".

8.2. Comparaison des résultats

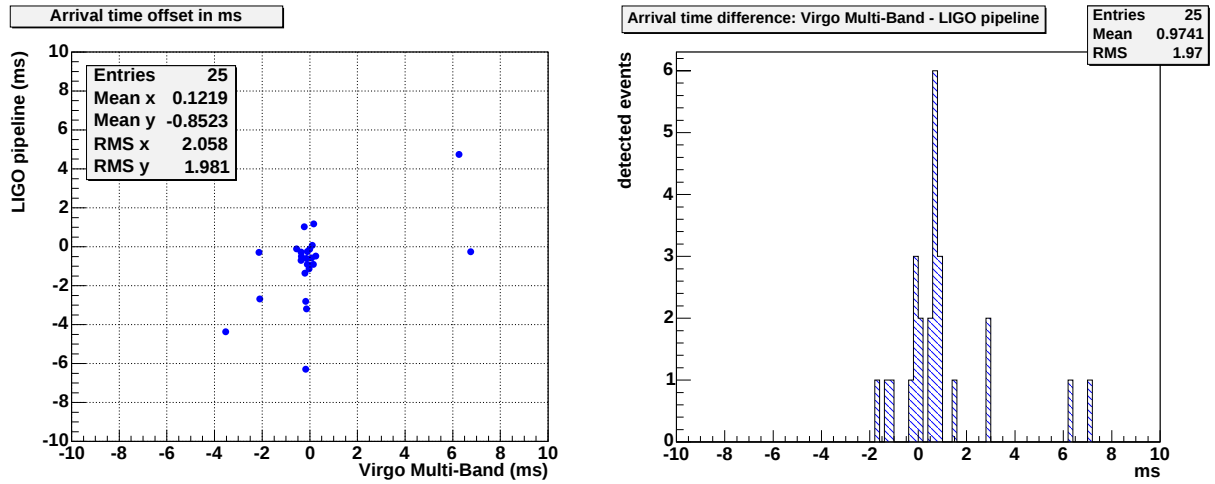


FIG. 8.8 – Comparaison des temps de fin obtenus pour chaque injection par les deux chaînes d'analyses sur les données LIGO. À gauche: temps donné par la chaîne de LIGO en fonction de celui donné par la recherche multi-bande (temps relatifs au temps de fin d'injection). À droite: distribution de la différence en temps "multi-bandes" – "chaîne LIGO" .

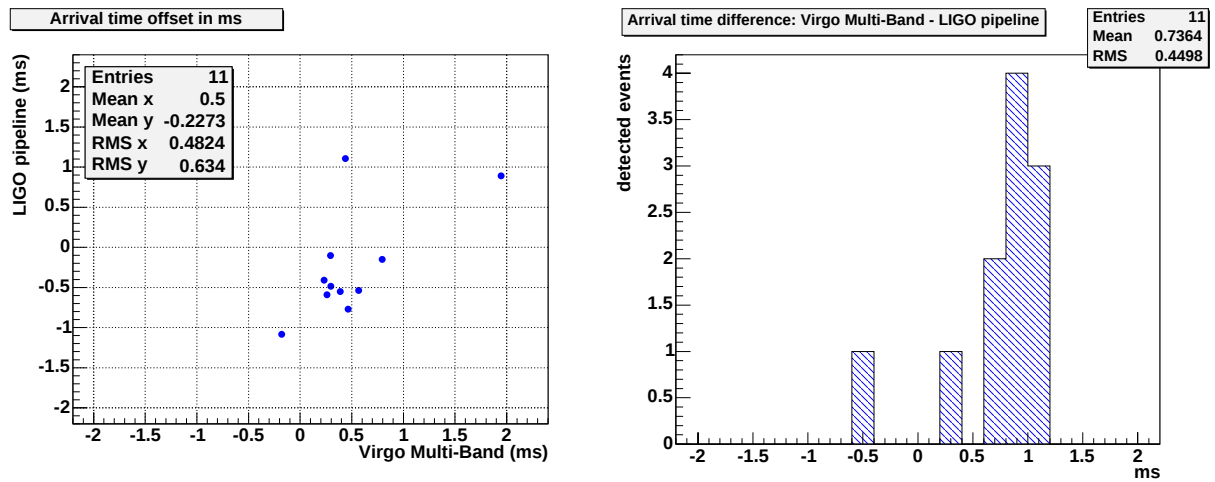


FIG. 8.9 – Comparaison des temps de fin obtenus pour chaque injection par les deux chaînes d'analyses sur les données Virgo. À gauche: temps donné par la chaîne de LIGO en fonction de celui donné par la recherche multi-bande (temps relatifs au temps de fin d'injection). À droite: distribution de la différence en temps "multi-bandes" – "chaîne LIGO" .

pour les données Virgo (figure 8.12). On note que ces faibles erreurs ne sont pas corrélés entre les deux chaînes, cette absence de corrélation étant attribué aux différences des grilles.

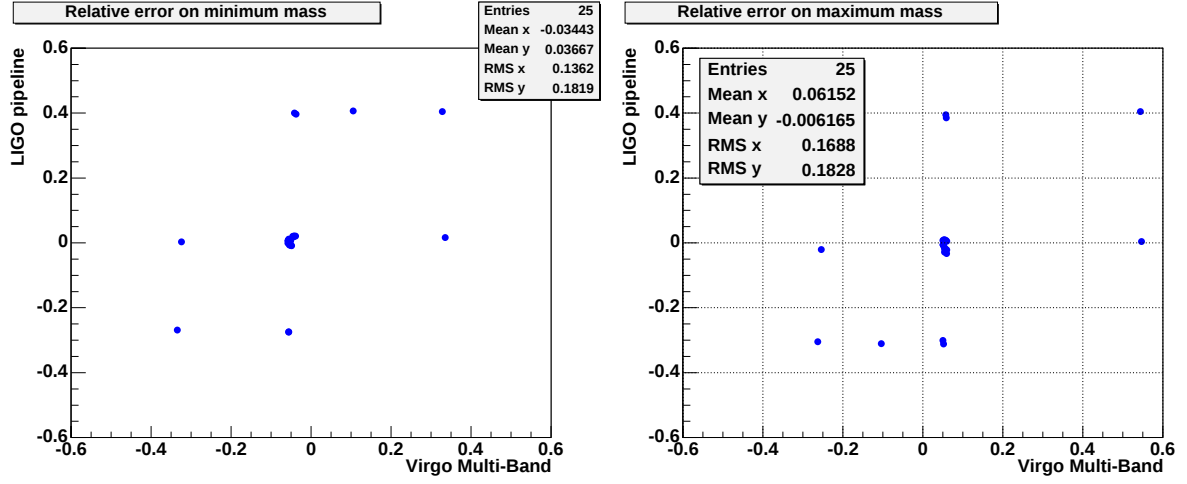


FIG. 8.10 – Comparaison des masses du template maximum pour les données LIGO: erreur (relative à la masse injectée) de la chaîne de LIGO en fonction de celle de la recherche multi-bande. À gauche: plus petite masse du couple. À droite: plus grande masse du couple.

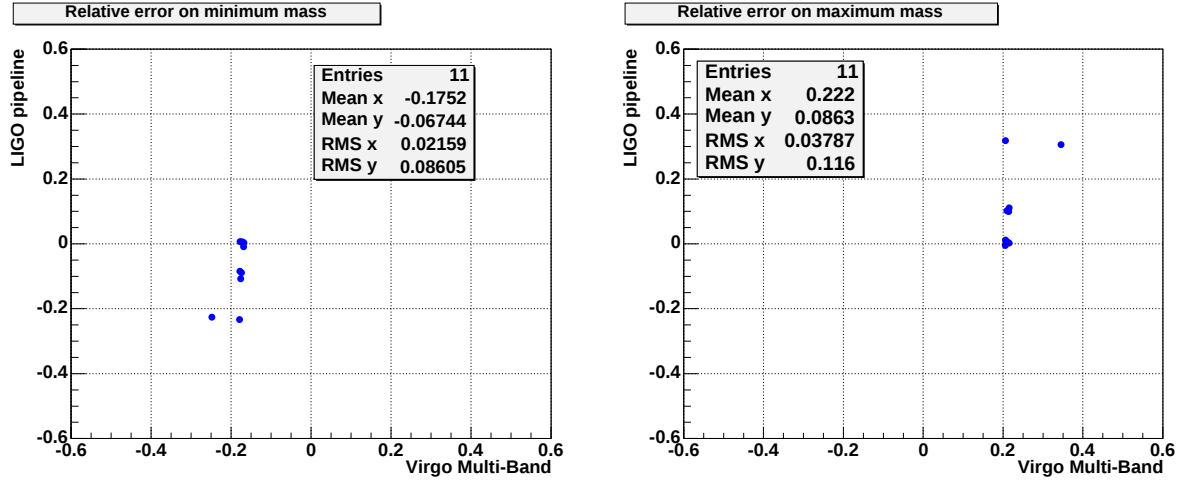


FIG. 8.11 – Comparaison des masses du template maximum pour les données Virgo: chaîne de LIGO en fonction de la recherche multi-bande. À gauche: plus petite masse du couple. À droite: plus grande masse du couple.

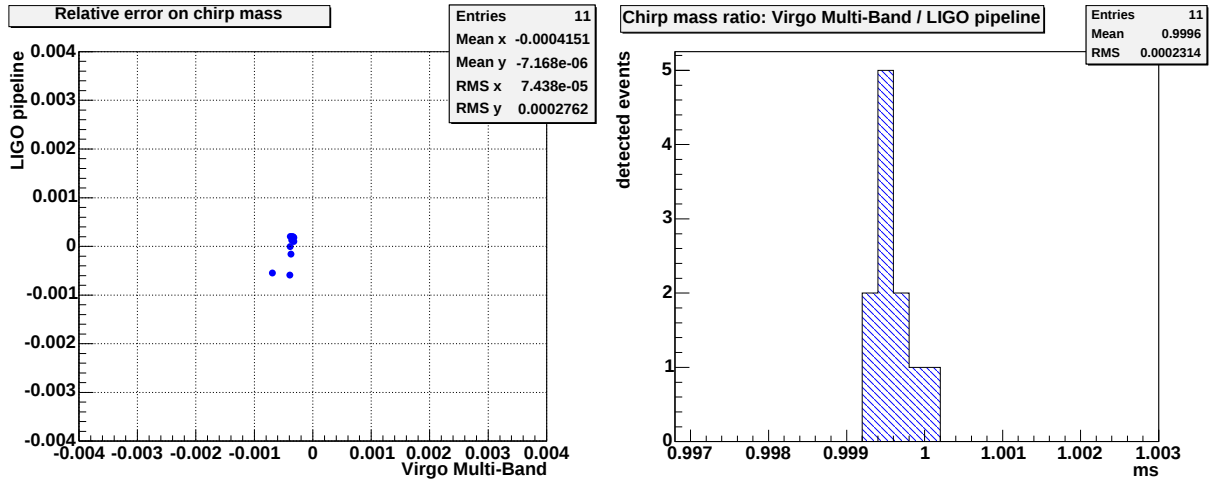


FIG. 8.12 – Comparaison entre chirp mass du template maximum sur les données Virgo. À gauche: erreur (relative à la chirp mass injectée) de la chaîne de LIGO en fonction de celle de la recherche multi-bande. À droite: distribution du rapport de chirp mass: "multi-bandes" / "chaîne LIGO" .

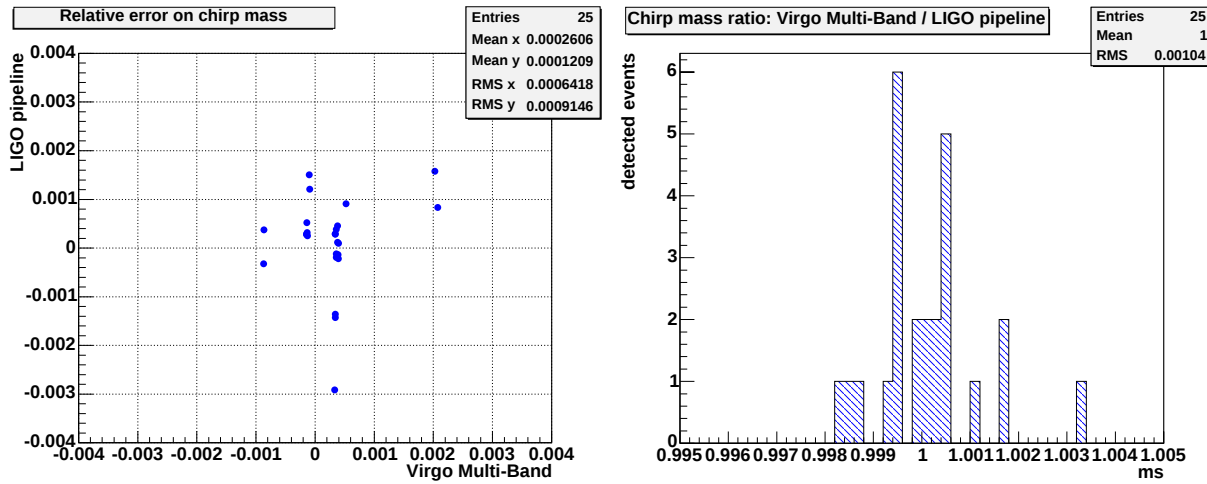


FIG. 8.13 – Comparaison entre chirp mass du template maximum sur les données LIGO. À gauche: erreur (relative à la chirp mass injectée) de la chaîne de LIGO en fonction de celle de la recherche multi-bande. À droite: distribution du rapport de chirp mass: "multi-bandes" / "chaîne LIGO" .

8.3 Conclusion

A l'issue de cette analyse comparative, les objectifs du premier MDC LIGO-Virgo sont atteints.

Le premier est l'objectif technique qui consistait à réussir un premier échange de données, et l'analyse par chaque collaboration de données issues de l'autre collaboration. Grâce au format commun de stockage des données, seul des difficultés techniques mineures ont été rencontrés.

Le travail de comparaison présenté dans ce chapitre met en avant le deuxième volet de cette réussite, puisque les résultats montrent la concordance de la recherche multi-bande et de la chaîne de recherche de coalescences binaires de LIGO. La compréhension de la plupart des écarts résiduels est à l'origine d'une meilleure compréhension par chaque collaboration de l'autre chaîne d'analyse, ainsi que la confiance mutuelle indispensable à l'analyse en coïncidence. Celle-ci bénéficiera entre autres de la connaissance quantitative désormais acquise des différences entre chaînes d'analyse.

Conclusion

Le travail de thèse présenté dans ce mémoire s'inscrit dans le cadre de la mise en route de Virgo, à la fois par une contribution direct au travers de l'étalonnage et de la reconstruction des données, et par la mise au point d'une méthode d'analyse pour la recherche de coalescences binaires, qui permet d'affiner la caractérisation du détecteur au cours de son évolution et de préparer les futures recherches astrophysiques.

L'étalonnage de Virgo, au cours de sa mise en route, a permis la mise au point d'une procédure de mesure de la sensibilité applicable à toutes les configurations de l'interféromètre, qui permettra de suivre la progression de la sensibilité. La nécessaire calibration des actionneurs, principale source d'incertitudes, doit néanmoins faire l'objet d'un suivi et devra gagner en précision à l'approche de la sensibilité nominale.

Un processus de reconstruction adaptative à plusieurs signaux a été mis en oeuvre et utilisé sur les données de l'interféromètre recombinaison; on s'attend à ce qu'il fonctionne de façon analogue sur l'interféromètre recyclé. Combiné à l'algorithme de soustraction de lignes, il a permis d'explorer la possibilité d'améliorer la sensibilité par soustraction de bruits identifiés. On envisage à présent d'optimiser cette soustraction, en augmentant le nombre de bruits pris en compte et en améliorant leurs mesures.

Après la mise au point de la recherche de coalescences binaires, par la méthode alternative "multi-bande" visant à dépasser les limitations imposées par les moyens de calculs, celle-ci affiche en simulation les performances de détection optimales de la méthode standard. Si les gains en terme de moyens de calculs se sont déjà fait sentir, il reste à présent à optimiser le code pour en tirer pleinement parti.

Appliquée aux données actuelles, cette recherche montre une distribution du bruit proche de celle attendue, sur des périodes choisies mais pas sur l'ensemble des prises de données. Elle montre ainsi la nécessité de rejeter les mauvaises périodes et les faux événements engendrés par l'instruments grâce à des vetos divers.

Enfin, la première comparaison de la recherche multi-bande avec la chaîne d'analyse de LIGO a permis de contre-valider les deux chaînes et d'établir les bases indispensables à la suite de la collaboration, qui continue au sein du groupe d'analyse jointe LIGO-Virgo.

Bibliographie

- [1] S. M. Carroll, gr-qc/9712019
Lecture notes on General Relativity
- [2] K. Thorne, cours de physique classique disponible sur:
<http://www.pma.caltech.edu/Courses/ph136/yr2004/>
Applications of classical physics
- [3] P. Fré, VIRGO-SIGRAV SCHOOL ON GRAVITATIONAL WAVES (2005)
Introduction to General Relativity
- [4] V. Ferrari, VIRGO-SIGRAV SCHOOL ON GRAVITATIONAL WAVES (2005)
The quadrupole formalism
- [5] R. A. Hulse, J.H. Taylor, *Astroph. Journal* 195 L51 (1975).
Discovery of a pulsar in a binary system
- [6] J.M Weisberg, J.H Taylor, ASP Conf. Ser. 302: Radio Pulsars 93 (2003).
The relativistic binary pulsar B1913+16
- [7] M. Fierz, W. Pauli, *Proc. R. Soc. A* **173** 211 (1939).
- [8] VIRGO collaboration, (1989)
Proposal for the construction of a large interferometric detector of Gravitational Waves
- [9] A. Abramovici et al, *Science*, 256 325 (1992)
- [10] K. Danzmann et al, *Lectures Notes in Physics*, 410 184 (1992)

- [11] K. Tsubono, Gravitational wave experiments, p.112, Ed. E. Coccia and G. Pizella and F. Ronga, World Scientific, Singapore
- [12] K. Danzmann, Class. Quant. Grav. 13 A247 (1996)
LISA: laser interferometer space antenna for gravitational wave measurements
- [13] P. Jaranowski, A. Królak, B. F. Schutz, Phys. Rev. D 58, 063001 (1998).
Data analysis of gravitational-wave signals from spinning neutron stars: The signal and its detection
- [14] Nicolas Arnaud, thèse de doctorat (2002)
Recherche de signaux impulsionnels: application aux coïncidences entre détecteurs
- [15] E.S Phinney, Astroph. Journal 380 L17 (1991).
The rate of neutron star binary mergers in the universe: minimal predictions for gravity wave detectors
- [16] M. Burgay, et al, Nature 426 531 (2003).
An increased estimate of the merger rate of double neutron stars from observations of a highly relativistic system
- [17] C.Kim, V. Kalogera, D.R. Lorimer, Astroph. Journal 584 985 (2003).
The probability distribution of binary pulsar coalescence rates. I. Double neutron star systems in the galactic field
- [18] V. Kalogera, C. Kim, D. R. Lorimer, M. Burgay, N. D'Amico, A. Possenti, R. N. Manchester, A. G. Lyne, B. C. Joshi, M. A. McLaughlin, M. Kramer, J. M. Sarkissian, F. Camilo, Astroph. Journal 601 L179 (2004).
The Cosmic Coalescence Rates for Double Neutron Star Binaries
- [19] V. Kalogera, R. Narayan, D. N. Spergel, J. H. Taylor,
The coalescence rate of double neutron star systems
- [20] L. P. Grishchuck, V. M. Lipunov, K. A. Postnov, M. E. Prokhorov, B. S. Sathyaprakash, arXiv:astro-ph/0008481 (2001) *Gravitational Wave Astronomy: in Anticipation of First Sources to be Detected*
- [21] K. Belczynski, V. Kalogera, T. Bulik, Astroph. Journal 572 407 (2002).
A Comprehensive Study of Binary Compact Objects as Gravitational Wave Sources: Evolutionary Channels, Rates, and Physical Properties

- [22] L. Derome, thèse de doctorat (1999)
Le système de détection de l'expérience VIRGO dédiée à la recherche d'ondes gravitationnelles
- [23] Flaminio, Heitmann, Physics Letters A, **214** 112-122 (1996).
Longitudinal control of an interferometer for the detection of gravitational waves
- [24] C. Mehmél, R. Flaminio, B. Caron, VIR-NOT-LAP-1380-130 (1995).
Linear dynamical model to the Virgo interferometer
- [25] M. Punturo, VIR-NOT-PER-1390-51
The VIRGO sensitivity curve
<http://www.virgo.infn.it/senscurve/>
- [26] F. Acernese *et al.* [Virgo Collaboration], Astropart. Phys. **21** (2004) 1.
The commissioning of the central interferometer of the Virgo gravitational wave detector
- [27] B. Caron *et al.*, Astropart. Phys. **10** (1999) 369.
SIESTA, a time domain, general purpose simulation program for the VIRGO experiment
- [28] A. Bozzi *et al.*, Review of Scientific Instruments 73 5 2143-2149 (2002).
Last stage control and mechanical transfer function of the Virgo suspensions
- [29] O. Veziant, thèse de doctorat (2003).
Calibration de l'expérience VIRGO: de l'étalonnage du détecteur à la recherche de signaux de coalescences binaires avec l'interféromètre central
- [30] M. Hewitson, H. Grote, G. Heinzel, K. A. Strain, U. Weiland, Review of Scientific Instruments 74 (9), 4184-4190 (2003)
Calibration of the power-recycled gravitational wave detector, GEO 600
- [31] M Hewitson *et al.*, Class. Quantum Grav. 21 S1711-S1722, (2004)
Calibration of the dual-recycled GEO 600 detector for the S3 science run
- [32] Xavier Siemens *et al.*, Class. Quantum Grav. 21 S1723-S1735 (2004)
Making $h(t)$ for LIGO

- [33] R. Flaminio, International Journal of Modern Physics D, 9-3 (2000).
Monitoring and adaptative removal of the power supply harmonics applied to the Virgo read-out noise.
- [34] L. Bosi, D. Buskulic, G. Cella, T.Cokelaer, G.M. Guidi, A. Viceré,
VIR-MAN-CAS-7500-16 (2003).
The inspiral library user's manual
- [35] A. Vicere, Virgo-SIGRAV school (2000).
Data Analysis for Coalescing binaries
- [36] B. Owen, B.S.Sathyaprakash, Phys. Rev. D 60, 022002 (1996).
Matched filtering of gravitationnal waves from inspiraling compact binaries: Computational cost and template placement
- [37] P. Canitrot, L. Milano, A. Viceré, VIR-NOT-PIS-1390 (2000).
Computational costs for coalescing binaries detection in VIRGO using matched filters
- [38] A. Viceré, Réunion de collaboration Virgo.
- [39] Bruce Allen, Phys.Rev. D 71 062001 (2005).
A chi-squared time-frequency discriminator for gravitational wave detection
- [40] L. Blackburn *et al.* [Joint LIGO / Virgo working group],
Soumis à Classical and Quantum Gravity, arXiv:gr-qc/0504050 (2005).
A first comparison between LIGO and Virgo inspiral search pipelines
- [41] LIGO Scientific Collaboration, B.Abbott et al.,
Phys.Rev.D 69 122001 (2004).
Analysis of LIGO data for gravitational waves from binary neutron star
- [42] P.Amico, L.Bosi, C.Cattuto, L.Gammaitoni, F.Marchesoni, M.Punturo,
F.Travasso, H.Vocca, Comp. Phys. Comm. 153, 179 (2003).
A parallel beowulf-based system for the detection of GW in interferometric detectors