



AVERTISSEMENT

Ce document est le fruit d'un long travail approuvé par le jury de soutenance et mis à disposition de l'ensemble de la communauté universitaire élargie.

Il est soumis à la propriété intellectuelle de l'auteur. Ceci implique une obligation de citation et de référencement lors de l'utilisation de ce document.

D'autre part, toute contrefaçon, plagiat, reproduction illicite encourt une poursuite pénale.

Contact : ddoc-memoires-contact@univ-lorraine.fr

LIENS

Code de la Propriété Intellectuelle. articles L 122. 4

Code de la Propriété Intellectuelle. articles L 335.2- L 335.10

http://www.cfcopies.com/V2/leg/leg_droi.php

<http://www.culture.gouv.fr/culture/infos-pratiques/droits/protection.htm>

UNIVERSITE DE NICE-SOPHIA-ANTIPOLIS
FACULTE DE MEDECINE
ECOLE D'ORTHOPHONIE

MEMOIRE PRESENTE POUR L'OBTENTION DU CERTIFICAT DE CAPACITE D'ORTHOPHONISTE

VOIX ET IMITATION

UTILISATION DU LOGICIEL PRAAT DANS LA MISE EN EVIDENCE
DE LA ZONE DE SECURITE PROSODIQUE

ANNE-LAURE COLLIOUD

née le 12 Octobre 1986 à Grenoble

Directeurs : **A.OSTA – Orthophoniste et**
S. CHRISTIAN - Orthophoniste

Co-Directeur : **P.O.VEDRINE – O.R.L-Phoniatre**

NICE – 2010

UNIVERSITE DE NICE-SOPHIA-ANTIPOLIS
FACULTE DE MEDECINE
ECOLE D'ORTHOPHONIE

MEMOIRE PRESENTE POUR L'OBTENTION DU CERTIFICAT DE CAPACITE D'ORTHOPHONISTE

VOIX ET IMITATION

UTILISATION DU LOGICIEL PRAAT DANS LA MISE EN EVIDENCE
DE LA ZONE DE SECURITE PROSODIQUE

ANNE-LAURE COLLIOUD

née le 12 Octobre 1986 à Grenoble

Directeurs : **A.OSTA – Orthophoniste et**
S. CHRISTIAN - Orthophoniste

Co-Directeur : **P.O.VEDRINE – O.R.L-Phoniatre**

NICE – 2010

REMERCIEMENTS

Je tiens à remercier les personnes qui m'ont guidée et soutenue tout au long de ma formation et de l'élaboration de ce mémoire :

- Monsieur Sébastien CHRISTIAN, orthophoniste, pour avoir été à l'origine de ce travail.
- Monsieur le docteur Pierre-Olivier VEDRINE, chirurgien ORL - phoniatre, pour son accueil, sa gentillesse, l'intérêt porté à ce travail et sa grande disponibilité quant à la réalisation du protocole de ce mémoire.
- Madame Arlette OSTA, orthophoniste, pour sa présence dans le jury de ce mémoire et sa relecture attentive.
- Madame Chloé MARSHALL et Madame Danielle JANIAUD, orthophonistes, pour leur présence dans le jury de ce mémoire.
- Madame Hélène LOEVENBRUCK, chercheur dans le domaine parole et cognition au CNRS, pour sa gentillesse, son aide précieuse et sa capacité à trouver les solutions de problèmes épineux.
- Monsieur le docteur Jean ABITBOL, chirurgien ORL – phoniatre, pour son accueil et pour m'avoir aimablement permis d'enregistrer un imitateur.
- Monsieur Michaël GREGORIO, imitateur, pour son aimable participation.
- Madame Claudine TERLE-VITALE, orthophoniste, pour m'avoir conseillée au début de la réalisation de ce travail.
- A toutes les personnes ayant gentiment accepté de participer aux enregistrements vocaux de ce mémoire.
- A mes parents, ma famille, mes amis qui m'ont aidée et soutenue tout au long de la réalisation de ce mémoire et plus largement, au cours de ces quatre années d'études.

TABLE DES MATIERES

INTRODUCTION	1
PARTIE THEORIQUE.....	3
<u>1. IMITATION</u>	4
1.1 QUAND PARLER D'IMITATION ?.....	4
1.1.1 CE QU'ELLE N'EST PAS.....	4
1.1.2 DEFINITION SELON DIFFERENTS POINTS DE VUE.....	5
1.1.2.1 Dimension cognitive.....	5
1.1.2.2 Dimension sociale.....	6
1.2 ETHOLOGIE.....	7
1.2.1 IMITATION CHEZ LES PRIMATES ET NON-PRIMATES.....	7
1.2.2 IMITATION CHEZ LES OISEAUX.....	8
1.2.2.1 Appareil vocal, imitation, acquisition.....	8
1.2.2.2 Cas du perroquet.....	9
1.3 IMPACT SOCIO-CULTUREL SUR L'IMITATION.....	10
1.3.1 THEORIE DE L'APPRENTISSAGE SOCIAL.....	10
1.3.2 THEORIE DE LA COMMUNICATION ACCOMMODANTE.....	11
1.4 IMITATION ET ANTHROPOLOGIE.....	13
1.4.1 IMITATION, UNE COMPETENCE INNEE ?.....	13
1.4.2 DEVELOPPEMENT DU NOURRISSON, IMITATION VOCALE.....	14

1.5 THEORIE DES NEURONES MIROIRS	16
1.5.1 NEURONES MIROIRS CHEZ LES PRIMATES.....	16
1.5.2 NEURONES MIROIRS CHEZ L'HOMME.....	18
1.5.2.1 Situation.....	18
1.5.2.2 Comparaison homme/singe.....	19
1.5.2.3 Actes non présents dans notre « vocabulaire d'actes ».....	19
1.5.2.4 Neurones miroirs, imitation, langage.....	21
 <u>2. CAS DES IMITATEURS</u>	23
2.1 PARTICULARITES ANATOMIQUES	23
2.1.1 CONTORSIONNISTES DU LARYNX ET DES RESONATEURS.....	23
2.1.2 BOUCLE AUDIO-PHONATOIRE	24
2.2 STRATEGIES	25
2.2.1 ASPECT AUDITIF.....	25
2.2.2 ASPECT ACOUSTIQUE.....	26
2.2.2.1 Fréquence fondamentale.....	26
2.2.2.2 Formants.....	27
2.2.2.3 Durée.....	28
2.2.2.4 Articulation.....	28
2.2.3 RECIT D'UN IMITATEUR.....	29
 <u>3. LOGICIEL PRAAT</u>	31
 <u>4. ZONE DE SECURITE PROSODIQUE</u>	35

<u>5. ANATOMIE-PHYSIOLOGIE</u>	36
5.1 PRODUCTION DES SONS	36
5.1.1 POSTURE.....	36
5.1.1.1 Position d'équilibre, verticalité du corps.....	36
5.1.1.2 Enjeu postural.....	37
5.1.2 RESPIRATION.....	38
5.1.2.1 Musculature respiratoire et innervation.....	38
5.1.2.1.1 Muscles dits inspireurs.....	38
5.1.2.1.2 Muscles dits expirateurs.....	39
5.1.2.2 En phonation.....	39
5.1.3 LARYNX.....	41
5.1.3.1 Musculature.....	41
5.1.3.1.1 Muscles extrinsèques.....	42
5.1.3.1.2 Muscles intrinsèques.....	44
5.1.3.2 Commande nerveuse.....	47
5.1.3.2.1 Du cerveau au larynx.....	47
5.1.3.2.2 Contrôle réflexe.....	50
5.1.3.2.3 Innervation des muscles laryngés.....	50
5.1.3.3 Vibration laryngée.....	52
5.1.3.3.1 Théorie de la vibration laryngée.....	52
5.1.3.3.2 Nature du son laryngé produit.....	54
5.1.3.3.3 Différents mécanismes laryngés.....	54
5.1.3.3.4 Attaques vocaliques.....	56
5.1.3.3.5 Contrôle des paramètres de la voix.....	56
5.1.4 RESONATEURS.....	58
5.1.4.1 Cavités de résonance.....	58
5.1.4.2 Eléments mobiles.....	59
5.1.4.3 Commande nerveuse.....	60
5.1.4.4 Leur rôle.....	61
5.1.4.5 Cas des chants diphonique et harmonique.....	63

5.2 PERCEPTION DES SONS	65
5.2.1 RECEPTEUR DE LA VOIX : L'OREILLE.....	65
5.1.2.1 Anatomie de l'oreille.....	65
5.1.2.2 Voies auditives.....	66
5.1.2.3 Développement et audition.....	67
5.1.2.4 Perception des sons.....	68
5.1.2.4.1 Cérébrale.....	68
5.1.2.4.2 Hauteur.....	69
5.2.2 CAS DE L'OREILLE ABSOLUE.....	70
 5.3 LIEN ENTRE PRODUCTION ET PERCEPTION.....	 71
5.3.1 CONDUCTION DU SON PRODUIT.....	71
5.3.2 CONTROLE DU FEEDBACK AUDITIF.....	72
5.3.3 PERCEPTION ET AJUSTEMENT VOCAL.....	73
 PARTIE PRATIQUE.....	 75
 <u>1. MATERIEL ET METHODE.....</u>	 <u>78</u>
1.1 POPULATION.....	78
1.2 PROTOCOLE ET RECUEIL DES ENREGISTREMENTS.....	79
1.2.1 MATERIEL.....	79
1.2.1.1 Logiciel.....	79
1.2.1.2 Eléments périphériques.....	79
1.2.2 EPREUVES.....	80
1.2.3 CONDITIONS D'ENREGISTREMENT.....	86

<u>2. RECUEIL ET ANALYSE DES DONNEES</u>	87
2.1 IMITATEUR	87
2.1.1 TON MELODIQUE	87
2.1.2 TON MONOTONE	90
2.1.3 ANALYSE DES RESULTATS ET RESENTI DE L'IMITATEUR	92
2.2 GROUPE DES TEMOINS	93
2.2.1 DESCRIPTION DES TEMOINS	93
2.2.1.1 Ton mélodique	95
2.2.1.2 Ton monotone	111
2.2.2 RESENTI DES TEMOINS FACE AUX EPREUVES	127
2.3 GROUPE DES PATIENTS DYSPHONIQUES	127
2.3.1 DESCRIPTION DES PATIENTS DYSPHONIQUES	128
2.3.1.1 Ton mélodique	129
2.3.1.2 Ton monotone	136
2.3.2 RESENTI DES PATIENTS FACE AUX EPREUVES	142
2.4 SYNTHESE GLOBALE DES RESULTATS	143
<u>3. DISCUSSION</u>	151
3.1 CRITIQUES METHODOLOGIQUES	151
3.1.1 PARTIE THEORIQUE	151

3.1.2 PARTIE PRATIQUE.....	152
3.2 CRITIQUES DES PRINCIPAUX RESULTATS.....	154
3.3 OUVERTURE AU CHAMP ORTHOPHONIQUE.....	158
CONCLUSION.....	159
BIBLIOGRAPHIE.....	160
ANNEXES.....	169

INTRODUCTION

Dans leur exercice professionnel, les orthophonistes sont souvent confrontés à des patients présentant des troubles de la voix. Un bilan de la fonction vocale permet d'évaluer les altérations vocales et d'orienter la rééducation. Cette rééducation est menée par une voix modèle, la voix de l'orthophoniste qui guide la voix pathologique, celle du patient dysphonique. La voix modèle doit amener la voix pathologique à acquérir ou retrouver un confort vocal et adopter une meilleure technique vocale.

La voix humaine est un signal riche et complexe, au caractère unique et subjectif. Il est difficile voire impossible de fournir un modèle orthophonique respectant exactement les paramètres acoustiques de la voix d'autrui. De plus, le patient peut éprouver des difficultés à transposer sur lui-même les éléments observés ou entendus chez une tierce personne. Pourtant, les exercices vocaux reposant sur une notion d'imitation du modèle orthophonique sont nombreux. Mais nous ne savons si tout être humain est capable de réaliser une imitation.

Ainsi, les capacités d'imitation sont-elles présentes chez tout individu ? Le modèle orthophonique est-il toujours le plus adéquat pour orienter la voix de l'individu vers un meilleur confort vocal et une meilleure technique vocale ? Ne serait-il pas intéressant d'essayer d'objectiver la voix du sujet grâce à l'outil informatique et plus largement d'objectiver la voix que le sujet devrait avoir pour favoriser le meilleur confort vocal ? En ce sens, d'observer dans quelles fréquences sa voix serait la plus confortable pour créer un modèle vocal spécifiquement adapté à la voix du sujet que celui-ci devrait imiter et dans lequel celui-ci pourrait demeurer ou se replier pour retrouver ou acquérir une efficacité vocale améliorer ou prévenir l'apparition de lésions consécutives à une mauvaise utilisation de la voix ?

Sans remettre en question l'efficacité et le bien fondé des méthodes de rééducation actuelles, il s'agit ici de proposer un moyen supplémentaire permettant d'objectiver la voix dans la rééducation orthophonique. Pour ce faire, la démarche suivie s'est faite en quatre étapes.

Tout d'abord, nous essaierons de comprendre ce qu'est l'imitation. Pour cela, nous tenterons de la définir. Nous ferons également appel au domaine éthologique, anthropologique, socio-culturel, neurologique pour mieux cerner ce que recouvre cette notion.

Ensuite, nous nous intéresserons au cas des imitateurs. Nous tenterons de mettre en évidence les éléments participant à la réalisation d'une imitation. Nous nous appuierons également sur le témoignage d'un imitateur.

Puis, nous traiterons du logiciel PRAAT et nous expliquerons ce qu'est la zone de sécurité prosodique.

Egalement, nous ferons un rappel des éléments anatomiques et physiologiques qui concernent la voix et susceptibles d'intervenir dans l'imitation. Nous verrons donc comment la voix est produite mais aussi comment elle est perçue. Nous expliquerons quel lien existe entre la production et la perception de la voix.

Enfin, après ces éléments théoriques, nous verrons au travers d'enregistrements de témoins, de patients et de l'imitateur rencontré, si tout individu est capable d'imitation. Nous objectiverons la voix des sujets, nous définirons la zone fréquentielle la mieux adaptée pour chacun et nous observerons si le sujet réalise la meilleure imitation pour le modèle orthophonique ou pour le modèle de sa propre voix.

PARTIE THEORIQUE

1. IMITATION

L'imitation touche divers domaines. Au fil du temps, de nombreuses significations ont été données au terme d'imitation, étant parfois même opposées selon qu'elles relevaient de l'éthologie, de la psychologie expérimentale, etc. L'imitation a souvent été définie sous un aspect cognitif mais son aspect social ne peut être négligé pour autant. Nous évoquerons ici les éléments les plus pertinents pour notre étude. Dans le cadre de la rééducation orthophonique, l'imitation doit être un moyen pour le patient d'apprendre à moduler sa production vocale pour justement disposer au mieux de ses facultés vocales. L'imitation en tant que moyen pour modifier un comportement est donc essentielle.

1.1. QUAND PARLER D'IMITATION ?

1.1.1 CE QU'ELLE N'EST PAS

Pour parler d'imitation, il convient avant tout de mettre en évidence les mécanismes se distinguant d'elle. Notons tout d'abord que l'acte imitatif diffère de l'acte réflexe défini par K.LORENZ¹ dont l'improbabilité du déclenchement exclut toute forme d'apprentissage.

Egalement, il ne peut y avoir imitation si aucun comportement nouveau n'est appris. Pour A.K.BARD², l'imitation constitue une copie fidèle des moyens et des fins.

Un geste copié pour lui-même, sans conscience du lien entre le geste et la conséquence du geste ne peut s'apparenter à la notion d'imitation. Cette compréhension est essentielle lorsque l'imitation est différée par exemple. Elle fait intervenir des processus cognitifs plus élaborés qu'une simple reproduction automatique.

¹ K. LORENZ, *Companionship in bird life*, 1957 in C.H. SCHILLER, K.S. ASHLEY, *Instinctive behavior: The development of a modern concept*.

² A.K. BARD, C.L. RUSSEL, *Evolutionary foundations of imitation: Social, cognitive and developmental aspects of imitative processes in non-human primates*, 1999, in J. NADEL & G.BUTTERWORTH (Edit.), *Imitation in infancy*

La notion d'intentionnalité doit aussi être présente dans l'imitation dite véritable. C'est ce qui la différencie par ailleurs de l'émulation. Le mimétisme quant à lui ne constitue qu'un niveau basique de l'imitation, un niveau sensori-moteur sans réel traitement cognitif.³

De même, certains auteurs tels que P.M.BAUDONNIERE⁴ distinguent mimétisme comportemental et imitation. L'imitation serait le propre de l'homme par les caractères de sélectivité : « *on n'imité pas n'importe qui, n'importe quoi, n'importe quand, n'importe où* », de conscience de soi et d'autrui qui sont présents dans l'espèce humaine et absents chez l'animal. Le mimétisme se retrouve chez l'animal et l'humain mais ne serait qu'instinctif (par opposition à l'imitation qui nécessite une intentionnalité).

1.1.2 DEFINITIONS SELON DIFFERENTS POINTS DE VUE

1.1.2.1 Dimension cognitive

Comme nous l'avons déjà souligné, bon nombre de définitions existent au sujet de l'imitation. D'une manière générale, l'imitation est considérée comme un phénomène basé sur l'appropriation de savoirs ou de savoir-faire par un individu. F. WINNYKAMEN et L. LAFONT⁵ ont rassemblé différents points de vue :

Selon D.CLEMENTSON-MOHR⁶, lors d'une imitation, le sujet imitant observe les moyens engendrés par le sujet imité pour atteindre un objectif ; il se servira ainsi des éléments repérés pour agir à son tour.

³ E. ZENTAL, *Defining imitation by exclusion*, 1996, in B.G GALEF, JR. & C.M. HEYES (Edit.), *Social learning in animals : The roots of culture*.

⁴ P.M. BAUDONNIERE, *Le mimétisme et l'imitation*, 1997

⁵ F. WINNYKAMEN, L. LAFONT *Place de l'imitation-modélisation parmi les modalités relationnelles d'acquisition : cas des habiletés motrices*, 1990.

⁶ D. CLEMENTSON-MOHR, *Toward a Social-cognitive explanation of imitation development*, in G. BUTTERWORTH and P. UGHET (eds), *Social cognition: studies of the development of understanding*. 1982

Nous avons vu précédemment que le caractère intentionnel dans l'imitation est important. Effectivement, dans l'acte imitatif, BRUNER⁷ souligne qu'il permet l'autoguidage de l'activité de l'individu.

F.WINNYKAMEN⁸ définit l'imitation comme l'usage intentionnel de l'acte observé, celui-ci constituant un renseignement pour atteindre un objectif. Cela laisse entendre qu'il existe bien une compréhension ne serait-ce que partielle des buts du modèle et que l'imitation offre la possibilité d'augmenter cette compréhension.

Tous ces éléments soulignent bien la dimension cognitive intervenant dans l'imitation mais oublient quelque peu l'aspect social.

1.1.2.2 Dimension sociale

Même si nous retrouverons tout au long du mémoire des éléments de l'aspect social de l'imitation, il convient déjà d'évoquer ici toute son importance, en poursuivant par exemple avec le point de vue de F.WINNYKAMEN⁹. Nous lui devons entre autres le concept « d'imitation-modélisation interactive » où elle associe les deux pôles de l'interaction et considère non seulement les attitudes du modèle mais aussi celles développées par le sujet imitant en fonction de ce qu'il aura vu.

Elle dénombre trois points de vue dans la manière de définir l'imitation :

- du point de vue « temporel », l'imitation peut-être immédiate ou décalée. Décalée, c'est-à-dire avec un laps de temps court : l'acte imitatif commence alors que l'action du modèle n'est pas terminée. Elle évoque que pour des auteurs tels que WALLON ou PIAGET, il existe aussi une imitation différée c'est-à-dire où les délais cités sont plus longs.

⁷ J.S. BRUNER, *Le Développement de l'Enfant : Savoir faire, savoir dire*, 1983.

⁸ F. WINNVKAMEN, *Imitation et acquisitions par observation : études récentes et perspectives*, In J. BIDEAUD & RICHELLE (eds), *Psychologie Développementale. Problèmes et Réalités*, 1985.

⁹ F. WINNYKAMEN, L. LAFONT, op. cit. , pp.23-30.

- du point de vue « formel », l'imitation peut être exacte. La production est conforme au modèle. Rien n'est ajouté ou enlevé. Si l'individu produit davantage que le modèle initial, il s'agira d'imitation en expansion. Si au contraire l'individu ne reproduit que quelques éléments, il s'agira d'imitation en réduction. Par ailleurs, une imitation peut être spontanée ou commandée.

- du point de vue « fonctionnel », l'imitation joue un rôle de communication, de contact social et d'acquisition. Leur importance respective et leurs relations varient selon le niveau de développement de l'individu.

1.2 ETHOLOGIE

1.2.1 IMITATION CHEZ LES PRIMATES ET NON-PRIMATES

L'étude des différentes espèces animales offre de nombreux renseignements concernant les facultés d'imitation. Des expériences menées sur des espèces primates et non-primates montrent que des facteurs biologiques sont susceptibles d'affecter les capacités d'imitation. *« Au sein d'une même espèce sachant imiter, la qualité de l'imitation peut être variable. Par exemple, les grands singes et les dauphins se distinguent par leur capacité d'imitation généraliste, en pouvant imiter une large gamme de comportements moteurs et vocaux alors que d'autres espèces imitent un éventail d'actes plus restreint. Les dauphins et les grands singes tendent aussi à imiter en temps réel alors que le perroquet imite seulement après des manifestations répétées et cela au bout d'un certain temps. La motivation, les procédures employées, la tolérance vis-à-vis de l'humain, les préférences sociales, le niveau de développement, d'attention ou encore de perception, les capacités de mémorisation, ... constituent des variables différenciant les individus. La qualité de l'imitation dépend donc de paramètres cognitifs, sociaux et fonctionnels et plus précisément, de l'affect, de la motivation, de la mémoire et de la perception. Egaleme nt, l'imitation vocale peut être utilisée par les dauphins à bec pour acquérir leur sifflement propre, de même que chez les oiseaux, elle interviendrait dans les rencontres sociales. »¹⁰*

¹⁰ L.J. ROGERS, G.T. KAPLAN, *Comparative vertebrate cognition : are primates superior to non-primate ?*, 2004, pp.115-117.

1.2.2 IMITATION CHEZ LES OISEAUX

1.2.2.1 Appareil vocal, imitation, acquisition

Même s'ils ne possèdent pas de cordes vocales, les oiseaux possèdent un organe phonatoire similaire à notre larynx, nommé syrinx. Comme chez l'humain, la production des sons est dépendante de l'appareil respiratoire. L'aptitude à reproduire des sons pour les espèces les plus évoluées suggère la présence d'une syrinx élaborée, d'un développement suffisant des facultés cérébrales pour permettre la compréhension de ces sons et un système de commandes nerveuses permettant de manipuler suffisamment la syrinx pour une reproduction adéquate de ces sons.

L'imitation chez l'oiseau consiste en l'adoption dans son répertoire sonore de vocalisations entendues. Elle constitue donc une des nombreuses possibilités d'apprentissage.

Pour reprendre l'exemple des pinsons cité par M.F.CASTAREDE¹¹, c'est entre la naissance et quatre mois que ces oiseaux entendent et gardent en mémoire les chants qu'ils auront dans leur répertoire. L'acquisition est impossible s'ils subissent un isolement acoustique.

L'oisillon mémorisera d'abord des modèles puis il établira son répertoire en les imitant. L'audition est essentielle dans cette acquisition. Aussi, comme chez l'humain, l'aspect social de l'imitation intervient. En effet, *« c'est le lien social qui guide l'oiseau vers le modèle à imiter. Et si les œufs d'un oiseau chanteur sont couvés par des oiseaux non chanteurs, les oisillons qui naissent ne chantent pas. »*

Dans l'espèce humaine, le bain acoustique reçu au cours du développement est tout aussi important. Il influencera les capacités de perception et de production des sons.

Rappelons cependant que P.M.BAUDONNIERE¹² parle de mimétisme comportemental pour les animaux et l'humain mais réserve le terme d'imitation à l'humain.

¹¹ M.F. CASTAREDE, op. cit. , p.15.

¹² P.M. BAUDONNIERE, op. cit.

1.2.2.2 Cas du perroquet

Il a été souligné depuis de très nombreuses années que certains perroquets imitent plus facilement la voix des enfants que la voix grave d'un adulte. Cela par le fait que la voix des enfants correspond davantage à la portée de leur organe vocal. L'imitation d'une voix grave est plus difficile ; la distinction des paroles est plus confuse. Ainsi, le registre vocal est important dans la qualité de l'imitation.¹³

Le perroquet intrigue. En général, ce sont les similitudes entre l'homme et le singe qui sont soulignées. Pourtant, le perroquet gris, originaire du Gabon serait l'un des animaux se rapprochant le plus du jeune enfant. Il est capable de réaliser une activité de tri, de calculer, de répéter mais il lui est impossible de créer des nouveaux mots. Le perroquet est capable de produire des sons tels que l'homme sait le faire. Nous savons depuis peu que son anatomie particulière est à l'origine d'une telle fonction. En effet, comme mentionné précédemment, tout oiseau possède une syrinx. Le perroquet gris quant à lui, bénéficie non seulement d'une syrinx mais aussi d'un larynx. Les sons sont produits au niveau de la syrinx, située en bas de la trachée et sont modulés dans le larynx, situé en haut de la trachée et en arrière de la langue. Plus précisément, il existe au niveau de la syrinx deux bronches rejoignant la trachée, sous la coupe de muscles et membranes. Cet ensemble comporte une pseudo-corde vocale permettant le son. Grâce à ses deux bronches, l'oiseau peut produire deux sons vibratoires de même fréquence ou non. Comme l'explique J.ABITBOL : « *Cette information sonore crée une sinusoïde qui vient heurter notre oreille et nous donner l'illusion vibratoire de la voix.* »¹⁴

Mais comme l'explique D.K.WARREN, la colonne d'air est capturée entre le larynx et la syrinx, ce qui constitue une caisse de résonance supplémentaire, qui ne se retrouve pas chez l'homme.¹⁵

¹³ G.L. LECLERC-BUFFON, *Œuvres choisies de Buffon, tome 2*, 1845, pp.525-526.

¹⁴ J. ABITBOL, op. cit. , p.466-467

¹⁵ D.K. WARREN, D.K. PATTERSON, I.M. PEPPERBERG, *Mechanisms of American English vowel production in a grey parrot (Psittacus erithacus)*, 1996, pp.41-46.

Bien que son appareil vocal diffère de celui de l'homme, le perroquet gris est capable de parler. L'étude, par la scientifique I.M.PEPPERBERG, du désormais célèbre perroquet *Alex*, aux capacités exceptionnelles, a permis de comprendre davantage ce phénomène. Alex sait parler, c'est-à-dire qu'il est capable d'effectuer un contrôle moteur sur son conduit vocal. L'analyse de sa production au spectrogramme montre qu'il produit des formants sur certains mots ressemblant particulièrement à ceux de l'homme. La forme de sa syrinx aux parois rectilignes et positionnée en avant lui permet d'agir sur la fréquence lors de la production vocale. Il fait aussi varier la longueur de sa trachée. Il ajuste parfaitement la caisse de résonance pour imiter au mieux la voix de l'homme. Il possède un os hyoïde sur lequel s'insèrent les muscles de la langue et de la mâchoire permettant une grande modulation de la caisse de résonance supra-laryngée. Son larynx ne servirait qu'à transformer en phonèmes existants, les vibrations provenant de la syrinx et jouerait en quelque sorte, à l'aide de sa langue, le rôle de nos lèvres. Sa mandibule, la caisse de résonance à l'étage nasal ont également toute leur importance. « *L'instrument vocal du perroquet peut être considéré comme un instrument à vent mais ouvert aux deux bouts.* »¹⁶

1.3 IMPACT SOCIO-CULTUREL SUR L'IMITATION

1.3.1 THEORIE DE L'APPRENTISSAGE SOCIAL

Cette théorie met un accent particulier sur l'importance du rôle des variables sociales dans le développement de la personnalité et du comportement. En basant ses principes sur l'étude de la relation interindividuelle, elle s'écarte un peu du behaviourisme traditionnel dont les expériences portent en général sur des sujets uniques. On la doit au psychologue social A. BANDURA.

¹⁶ J. ABITBOL, op. cit. , p.470.

Selon lui, le sujet doit être motivé, sécurisé par des renforcements sociaux pour permettre l'apprentissage par imitation. L'expression des conduites imitatives peut dépendre de différents paramètres : le comportement du modèle doit correspondre aux capacités perceptives et motrices du sujet ; le taux de présentation peut également jouer. Selon l'auteur, « *des variables motivationnelles peuvent altérer les seuils perceptifs et par conséquent faciliter, gêner ou canaliser les réponses.* »¹⁷

Les comportements imités peuvent être sélectionnés selon des critères intrinsèques particuliers tels que l'intensité. De plus, certaines caractéristiques du modèle acquerront du relief en fonction du conditionnement social précoce du sujet. On pourra retrouver notamment l'attractivité du modèle, son caractère récompensant, prestigieux, compétent, puissant, autant d'éléments pouvant retenir l'attention et entraîner par la suite des comportements d'imitation, davantage que les modèles ne présentant pas ces attributs.

En conclusion, la théorie de l'apprentissage social propose quatre éléments intervenant dans l'apprentissage par observation, à savoir : des processus attentionnels, motivationnels, de rétention, et de reproduction motrice.

1.3.2 THEORIE DE LA COMMUNICATION ACCOMMODANTE

La théorie de la communication accommodante ou *Communication Accommodation Theory* *RCAT* a été élaborée par GILES, COUPLAND et COUPLAND. L'accent est mis sur la notion d'identité sociale, c'est-à-dire une identité découlant de l'appartenance à un groupe, et reconnue comme élément essentiel dans la communication. Il apparaît que l'identification au groupe agirait de façon conséquente sur les perceptions, les attitudes et les comportements. Dans toute situation d'échange transparaîtrait donc l'identité groupale des deux protagonistes. Les auteurs soulignent également que les stratégies de communication du locuteur s'adaptent au caractère positif ou négatif de l'identité sociale de l'interlocuteur et inversement. Si l'identité sociale de ce dernier est positive alors le locuteur utilisera une stratégie « accommodante » pour diminuer la distance entre les deux protagonistes. Il pourra par exemple être évalué plus favorablement par son homologue en réduisant les différences qui

¹⁷ A. BANDURA, *Social Learning Through Imitation*, in M.R. JONES (Edit), *Nebraska Symposium on Motivation*, 1962, p.262.

les séparent. En revanche, si elle est considérée comme négative, il emploiera une stratégie « non accommodante », marquant une distance sociale. Ainsi, l'accommodation constitue la « *tendance générale d'ajuster nos actions de communication envers celles de nos partenaires de conversation* ». ¹⁸

Deux types de comportements se retrouvent : l'un dit « convergent » et l'autre dit « divergent ». La convergence est la « *stratégie par laquelle les individus s'adaptent aux comportements de communication des autres, qui couvre un large panel de caractéristiques linguistiques, prosodiques et non verbales comme la vitesse d'élocution, les phénomènes de pause, la longueur des énoncés, les variantes phonologiques, le sourire, le regard, et ainsi de suite.* » alors que la divergence est « *la façon dont les locuteurs accentuent les différences d'ordre verbal et non-verbal entre eux et les autres locuteurs.* » ¹⁹

Caroline JUILLARD explique « *qu'une des propositions de la théorie de l'accommodation suggère que ceux qui désirent l'approbation des autres s'adapteront davantage que ceux qui n'en ont pas besoin ou ne la désirent pas.* » ²⁰ Convergence et divergence constituent la dimension objective de l'accommodation. On retrouve en outre la dimension subjective qui se rapporte à la croyance de l'émetteur concernant ces phénomènes.

Les modes d'accommodation se retrouvent donc dans le cadre des relations patient-thérapeute. Il faut donc en tenir compte dans la prise en charge orthophonique et plus spécifiquement si celle-ci est basée sur l'imitation.

¹⁸ H. Giles, N. Coupland, & J. Coupland, *Accommodation theory : Communication, context, and consequence*, in H. Giles, N. Coupland, & J. Coupland, editors, *Contexts of Accommodation : Developments in Applied Sociolinguistics*, 1991, p.60.

¹⁹ *ibid.* , p.7.

²⁰ C. JUILLARD dans M.L. MOREAU, *Sociolinguistique : les concepts de base*, 1997, p.13

1.4 IMITATION ET ANTHROPOLOGIE

1.4.1 IMITATION, UNE COMPETENCE INNEE ?

La question de l'innéité de l'imitation fait depuis longtemps débat. Pour certains auteurs, le nourrisson apprend à imiter d'abord seul puis grâce à autrui. D'autres suggèrent que cette capacité est à caractère social, héréditaire et faisant partie intégrante du psychisme humain. Dans son ouvrage, Nélías DIAS²¹ montre certaines divergences entre les auteurs. Elle met notamment en évidence celles existant entre PIAGET et MELTZOFF et MOORE.

J.PIAGET considère que l'imitation joue un rôle non négligeable dans le développement cognitif de l'enfant et décrit six stades dans le développement de l'imitation : tout d'abord l'absence d'imitation, puis l'imitation sporadique, l'imitation vocale et gestuelle entre 6 et 9 mois, l'imitation des mouvements non visibles sur le corps et de modèles nouveaux, et enfin l'imitation représentative, fondamentale dans la formation du processus symbolique survenant vers 18 mois. Il pense que le nourrisson s'auto-imite (phénomène intrapersonnel) pour ensuite imiter les autres (phénomène interpersonnel). Pour lui, « *l'enfant apprend à imiter.* »²²

Des auteurs comme A.N.MELTZOFF et M.KEITH MOORE s'opposent à la vision piagétienne de l'imitation. Selon eux, « *elle ne résulte pas d'un apprentissage mais plutôt un mécanisme, propre à l'espèce, d'apprentissage social et de transmission intergénérationnelle de caractères acquis.* »²³ Ils ont démontré que les nourrissons ont la capacité d'imiter des actes humains dès les premières semaines de vie. L'imitation constitue une capacité innée de l'espèce humaine, c'est-à-dire présente « *à la naissance et préalable à l'expérience de l'apprentissage d'une association particulière entre un stimulus (par exemple l'expression faciale de l'adulte) et une réaction (sa réplique faciale chez l'enfant).* »²⁴

²¹ N. DIAS, *Imitation et anthropologie*, 2005, pp.13-15.

²² J. PIAGET, *La formation du symbole chez l'enfant, Imitation jeu et rêve, image et représentation*, 1970, p.11.

²³ A.N. MELTZOFF, M.K. MOORE, *Imitation et développement humain : les premiers temps de la vie*, 2005, pp.72-73.

²⁴ *ibid.* , pp.72-73.

Aussi, suite aux résultats d'expériences attestant que le nourrisson est capable de s'ajuster et de s'auto-corriger pour imiter au mieux l'attitude de l'adulte observé, ils ajoutent et défendent l'idée d'une intentionnalité dans l'imitation infantile précoce. En imitant, le nouveau-né montre qu'il fait le lien entre les actions qu'il réalise et celles de l'adulte observé, entre sa disposition corporelle et celle de l'adulte.

Par ailleurs, « *l'observation du mécanisme de l'imitation néonatale aboutit à la découverte d'une capacité initiative du nouveau-né, appelée 'provocation'. Les nouveau-nés recommencent le même geste d'imitation, en l'attente d'une réponse de l'examineur.* »²⁵ Il existerait donc une relation de subjectivité, où le nourrisson doit montrer qu'il présente les bases d'une conscience individuelle et intentionnelle et une relation d'intersubjectivité où il doit adapter sa subjectivité à celle des autres.²⁶

Toutefois, les connaissances relatives aux capacités d'imitation précoces restent encore parfois contradictoires selon la position théorique adoptée. Un certain scepticisme peut demeurer selon les éléments postulés. Même si la présence de conduites d'imitation à la naissance ne fait plus doute, l'origine et la cause de ces conduites font encore débat. En effet, tous les auteurs ne s'accordent pas sur les mécanismes sous-jacents de l'imitation néonatale.

1.4.2 DEVELOPPEMENT DU NOURRISSON, IMITATION VOCALE

Concernant la place de l'imitation vocale dans le développement du nourrisson, différents éléments ont été mis en évidence.

Il a été démontré que dès les premiers temps, le bébé est capable de discrimination phonologique. Mais il est aussi capable d'imitation dans ses productions vocales, notamment lors de ses babillages vers le 8^{ème} mois, où ressortent les colorations de sa langue maternelle. En réalité, il porte en lui ces possibilités dès la naissance. Aux alentours de 3-4 jours, il est

²⁵ E.NAGY, P, MOLNAR, *Homo Imitans or Homo Provocans? Human imprinting model of neonatal imitation*, 2003, dans P.Trevarthen, D.Aitken, *Intersubjectivité chez le nourrisson : recherche, théorie et application clinique*, 2003-2004, p.309-428.

²⁶ P.Trevarthen, D.Aitken, *Intersubjectivité chez le nourrisson : recherche, théorie et application clinique*, 2003-2004, p.309-428

déjà capable « *d'imitation vocale sélective* » comme l'expliquent R.LEBIB et P.M.BAUDONNIERE²⁷. Les capacités d'imitation vocale des nourrissons vont se peaufiner et croître avec l'âge. Il leur sera possible vers 3-4 mois d'allier une articulation labiale à sa vocalise spécifique. Ils seront également sensibles à la « *covariation : ils imitent en effet préférentiellement une vocalisation lorsque le stimulus sonore et le stimulus visuel adéquat sont présents* ». Ainsi, ils seront gênés par un discours mimé ou l'écoute d'un son sans voir le mouvement labial. « *Leurs possibilités de traitement multimodal et intermodal vont donc se développer et s'affiner, et ceci est nécessaire pour induire un comportement imitatif à la base, rappelons-le, de l'acquisition du langage et donc de la communication chez l'être humain.* »

D.N.STERN²⁸ évoque également la notion de réciprocité dans les échanges entre un adulte et un bébé. C'est la sensibilité du nourrisson à l'enveloppe émotionnelle dynamique produite par les gestes et les vocalisations de la mère qui rendent cette réciprocité possible. Le bébé adopte les actions de l'adulte en repérant de façon immédiate ses émotions et expressions.

Pour M.PAPOUSEK et H.PAPOUSEK²⁹, les imitations parentales occupent une place essentielle dans le support d'apprentissage que les parents constituent intuitivement par rapport au bébé. Les imitations maternelles sont considérées comme un miroir biologique. Elles possèdent ce double aspect de réponse, quand elles surviennent immédiatement après le comportement-modèle du nourrisson ; et de modèle, quand la mère offre au bébé des stimulations comportementales directement inspirées du répertoire de l'enfant. Dans ce cas, la mère réalise une imitation différée d'une attitude récemment apprise par son enfant et lui soumet un modèle davantage reproductible. Par ailleurs, contrairement aux résultats de certaines études, ils constatent que le nombre de cas d'imitations où c'est le bébé qui initie et la mère qui imite est équivalent au cas inverse c'est-à-dire lorsque c'est la mère qui initie et le bébé qui imite. Il apparaîtrait donc que dès le début, les séances d'imitations possèdent un caractère réciproque tant au niveau de la démonstration que de l'imitation.

²⁷ R.LEBIB et P.M. BAUDONNIERE, *L'imitation vocale du nouveau-né au nourrisson de 3-4 mois: les «premiers pas» vers l'acquisition du langage*, 1989, pp. 37-52.

²⁸ D.N. STERN, *The role of feelings for an interpersonal self*, in U. NEISSER (Ed.), *The perceived self*, 1993, pp.205-215.

²⁹ M. PAPOUSEK, H. PAPOUSEK, *Forms and functions of vocal matching in interactions between mothers and their pre-canonical infants*, 1989, pp.137-158.

Ces conclusions sont également partagées par FONTAINE « *C'est autant la mère que l'enfant qui imite.* »³⁰ Il explique que la mère moule ses mimiques sur le modèle des expressions du bébé et les lui renvoie en miroir en accentuant les caractéristiques. Elle procèdera de même pour la prosodie langagière. L'imitation est donc une communication interpersonnelle.

1.5 THEORIE DES NEURONES MIROIRS

La production et la perception des sons sont des phénomènes étroitement liés. Par ailleurs, dans la communication, pour que l'échange entre le locuteur et l'interlocuteur soit possible, il faut que les deux protagonistes partagent une représentation identique.

Une voie de recherche assez récente propose que des neurones particuliers, les neurones miroirs, rendent possible ce lien entre production et perception, permettant le code commun entre l'émetteur et le récepteur.

Bien que peu diffusée, cette théorie des neurones miroirs semble très prometteuse quant aux nombreuses découvertes qui en découlent et pourront en découler. Ces découvertes, comme l'a écrit V.S RAMACHANDRAN, directeur du *Center for Brain and Cognition* de l'université de Californie : « *vont fournir un cadre unifiant et aider à expliquer une quantité de dispositions mentales qui jusqu'à maintenant restaient mystérieuses et inaccessibles à l'empirisme.* »

1.5.1 NEURONES MIROIRS CHEZ LES PRIMATES

Le système des neurones miroirs a tout d'abord été mis en évidence chez les primates. Nous devons cette découverte à une équipe de neurologues italiens, sous la direction de Giacomo RIZZOLATTI, dirigeant du département de neurosciences de l'université de Parme, dans les années 90.

³⁰ R. Fontaine, *Imitative skills between birth and six months*, 1984, pp.323-333, dans J.A. RONDAL, *Manuel de psychologie de l'enfant*, 1999, p.458

« Il a été montré que des neurones situés dans la zone ventrale F5 du cortex prémoteur qui étaient activés lorsqu'un singe effectuait un mouvement avec un but précis (saisir un objet par exemple), s'activaient également quand le même singe observait ce mouvement réalisé par un congénère ou un humain. En quelque sorte, cette appellation de neurones miroirs vient du fait que les neurones dans le cerveau de l'observateur imitent ceux de la personne observée. Des neurones avec des caractéristiques semblables ont également été trouvés dans le lobule pariétal inférieur des singes. Ces deux aires s'incluent donc dans le système de circuits pariéto-frontaux, participant à l'organisation des actions. C'est par l'intermédiaire des interactions avec le sulcus temporal supérieur que les informations visuelles des actes moteurs atteignent l'aire F5. Ces neurones permettraient donc au singe la compréhension systématique de l'acte moteur réalisé par un semblable, en l'harmonisant à son propre répertoire moteur. Les résultats d'autres expériences vont en ce sens. En effet, il a été montré que les neurones miroirs de la zone F5 s'activent aussi quand le primate ne dispose pas de l'intégralité des constituants d'un acte moteur mais en perçoit suffisamment pour comprendre le but de l'acte. »³¹

« L'étude de l'organisation motrice du lobule pariétal inférieur révèle que les neurones moteurs n'interviennent pas seulement dans la simple compréhension de l'acte moteur. Ils ont la capacité d'encoder différemment le même acte moteur en fonction de l'action dans laquelle il intervient. Ainsi, l'encodage sera différent selon que le singe voit un individu saisissant un objet pour le manger ou saisissant un objet pour le placer. Ils encodent donc non seulement l'acte moteur observé mais aussi son but. »³²

³¹ L. CATTANEO, G. RIZZOLATTI, *The Mirror Neuron System*, 2009, pp.557-560.

³² L. CATTANEO, G. RIZZOLATTI, op. cit. , p.557.

1.5.2 NEURONES MIROIRS CHEZ L'HOMME

1.5.2.1 Situation

La découverte des neurones miroirs chez le singe a immédiatement laissé supposer l'existence d'un système similaire chez l'homme. Les premières preuves, bien qu'indirectes, de l'existence d'un tel système proviennent d'études menées à partir d'électro-encéphalographie sur la réactivité des rythme cérébraux à l'observation de certains mouvements. Diverses expériences ont été réalisées par la suite pour essayer de déterminer s'il existait chez l'homme une zone similaire à celle du système des neurones miroirs du singe.

Ce n'est que récemment, grâce à l'Imagerie par résonance magnétique fonctionnelle (IRMf), qu'une localisation a pu être mise en évidence plus précisément. Comme l'indique G.RIZZOLATTI³³, il apparaît que les aires s'activant lors de l'observation des actions d'autrui et donc s'impliquant dans le système des neurones miroirs se situent au niveau de la partie rostrale du lobe pariétal inférieur, dans la région inférieure du gyrus précentral et dans la région postérieure du gyrus frontal inférieur. Dans certaines conditions expérimentales, il a également été constaté qu'un secteur plus antérieur du gyrus frontal inférieur et du cortex prémoteur dorsal s'activait.

Ainsi, la région activée dans le lobe pariétal inférieur correspondrait à l'aire 40 de Brodmann et plus largement, serait l'homologue de l'aire constituant une partie du système des neurones miroirs chez le singe. Quant à la région postérieure du gyrus frontal inférieur, il semblerait qu'elle corresponde à la zone postérieure de l'aire de Broca qui assure le contrôle des mouvements buccaux lors de l'expression verbale. Des études récentes réalisées par PETRIDES et PANDYA³⁴ ont montré que cette zone pouvait constituer l'homologue humain de l'aire F5 du singe. Ainsi, il y a bien « *présence chez l'homme de mécanismes de résonance « analogue » à celui du singe* ». ³⁵

³³ G. RIZZOLATTI, C. SINIGAGLIA, *Les Neurones Miroirs*, 2008, pp.19-182.

³⁴ M.PETRIDES, D.N. PANDYA, *Comparative architectonic analysis of the human and the macaque frontal cortex*, in F. BOLLER, J. GRAFMAN (sld.), *Handbook of Neuropsychology*, 1997, pp.17-58.

1.5.2.2 Comparaisons homme / singe

Le système des neurones miroirs humain possède des caractéristiques non présentes chez le primate. En effet, chez l'homme, ses fonctions sont plus diverses : il permet de coder des actes moteurs transitifs et intransitifs ; il peut sélectionner à la fois le type d'acte ou la séquence des mouvements qui le composent ; pour finir, il ne nécessite pas une interaction effective avec les objets, il s'active aussi quand l'action est seulement mimée.

Cependant, tout comme le primate, lorsqu'un humain observe des actes réalisés par un autre individu, cela implique d'emblée les aires motrices dévolues à l'organisation et l'exécution de ces mêmes actes. De plus, la compréhension des actes « *chez l'homme ne concerne pas seulement des actes particuliers, mais des chaînes entières d'actes [...].* »³⁶

1.5.2.3 Actes non présents dans notre « vocabulaire d'actes »

Que se passe-t-il lorsqu'un individu assiste à des actions qui n'appartiennent pas au patrimoine moteur de son espèce ou qu'il n'est simplement pas capable de réaliser ? L'étude IRMf de B. CALVO-MERINO *et al.*³⁷ a mis en évidence que la vision d'actes réalisés par une tierce personne provoque une activité cérébrale différente en fonction des compétences motrices spécifiques des sujets présents. Ainsi, l'activation du système moteur des neurones miroirs serait modulée par l'expérience motrice. Cela souligne le rôle essentiel de la connaissance motrice pour pouvoir comprendre la signification des actes effectués par un autre individu.

³⁵ G. RIZZOLATTI, C. SINIGAGLIA, op. cit. , p137.

³⁶ *ibid.* , p.138.

³⁷ B. CALVO-MERINO, D.E. GLASER, J. GREZES, R.E. PASSINGHAM, P. HAGGARD, *Action observation and acquired motor skills : an fMRI study with expert dancers*, 2005, pp.1243-1249.

Etudions d'abord la situation d'imitation où l'individu doit reproduire un acte qui appartient à son patrimoine moteur, après l'avoir vu réalisé par autrui. Les recherches de PRINZ³⁸ mais aussi de BEKKERING *et al.*³⁹ soulignent que lorsque l'observateur perçoit un acte, au plus celui-ci ressemblera à l'acte présent dans son patrimoine moteur, au plus il induira l'exécution. Il faut donc un « schéma représentationnel commun » pour la perception et la réalisation des actions, modulé par la compréhension de l'observateur en ce qui concerne le but et la finalité des mouvements exécutés par l'individu observé. La découverte de l'existence des neurones miroirs permettrait de considérer ce schéma comme un mécanisme de transformation direct des informations visuelles en actes moteurs potentiels. Ainsi, il apparaît que le rôle fondamental de ces neurones dans l'imitation serait de coder l'action visualisée en termes moteurs, rendant possible sa reproduction.

Etudions maintenant le cas où la situation d'imitation nécessite « l'apprentissage d'un nouveau type d'action ». Pour expliquer cela, G.RIZZOLATTI retient tout particulièrement le modèle proposé par l'éthologue R.BYRNE⁴⁰ postulant que l'intégration de deux processus distincts seraient nécessaires à l'apprentissage par imitation : *« le premier permettrait à l'observateur de segmenter l'action qu'il doit imiter en éléments particuliers qui la composent, c'est-à-dire de convertir le flux continu des mouvements observés en chaîne d'actes appartenant à son patrimoine moteur ; le deuxième lui permettrait d'accomplir les actes moteurs ainsi codés dans la séquence la plus appropriée afin que l'action exécutée reflète celle du démonstrateur. »*⁴¹ Ce seraient les neurones miroirs situés dans les lobes frontal et pariétal inférieur qui traduiraient en termes moteurs les actes élémentaires caractérisant l'action observée. Il semblerait par ailleurs que l'aire 46 de Brodmann prendrait en charge une recombinaison des actes moteurs particuliers et la réalisation d'une configuration d'action nouvelle, se rapprochant au maximum de celle réalisée par l'observé.

³⁸ W. PRINZ, *A common-coding approach to perception and action*, in O. NEUMANN, W. PRINZ (sld.) *Relationship between Perception and action : Current Approaches*, 1990, pp.167-203.

³⁹ H. BEKKERING, A. WOHLSCHLÄGER, M. GATTIS, *Imitation of gestures in children in goal-directed*, 2000, pp. 153-164

⁴⁰ R. BYRNE, A.E. RUSSON, *Learning by imitation : a hierarchical approach*, 1998, pp.667-712.

R. BYRNE, *Seeing actions as hierarchically organized structures : great ape manual skills*, 2002, pp.102-140.

R. BYRNE, *Imitation as behavior parsing*, 2003, pp.529-536.

⁴¹ G.RIZZOLATTI, C. SINIGAGLIA, *op. cit.* , p.158

Il convient de souligner cependant que ni la qualité du patrimoine moteur, ni la seule présence du système des neurones miroirs ne peut déterminer exclusivement les capacités d'imitation. De plus, l'existence d'un système de contrôle des neurones miroirs est nécessaire, sans quoi, toute perception d'un acte quelconque se matérialiserait immédiatement par une reproduction. Il ne doit permettre le passage d'une action potentielle à l'exécution d'un acte moteur via le codage des neurones miroirs seulement lorsque cela est opportun pour l'observateur.

1.5.2.4 Neurones miroirs, imitation, langage

A l'origine du langage serait le geste. Tel est le postulat vers lequel tendent diverses disciplines. C'est notamment la présence de représentations motrices différentes (oro-faciales, oro-laryngées et brachio-manuelles) au niveau de l'aire F5 et de l'aire de Broca qui laisse envisager que la communication entre individus est apparue via l'intégration progressive de modalités motrices différentes (gestes faciaux, brachio-manuels et enfin vocaux) et non uniques, avec en parallèle l'émergence de leurs systèmes de neurones miroirs respectifs. Selon M.ARBIB⁴², ce sont les capacités des hominidés à exécuter en premier lieu des comportements imitatifs, puis des mimiques d'actes et enfin, des « *protosignes* » brachio-manuels essentiels pour une communication efficace, qui ont permis une évolution vers le système communicatif que nous connaissons aujourd'hui.

De plus, c'est probablement grâce à l'émergence du système communicatif brachio-manuel que les vocalisations ont pris de l'importance et ont pu être contrôlées. Aussi, le développement de l'intentionnalité dans l'usage des sons avec présence de mimiques et protosignes a dû également contribuer à cette évolution. La précision requise par ces nouvelles pratiques communicatives explique vraisemblablement l'apparition de l'aire de Broca à partir d'une aire semblable à F5, dotée elle aussi d'une représentation des mouvements oro-laryngés, d'une connexion avec le cortex moteur primaire voisin, de

⁴² M.A. ARBIB, *Beyond the mirror system : imitation and evolution of language*, in C. NEHANIV, K. DAUTENHAN (sld.), *Imitation in Animals and Artifacts*, 2002, pp.229-280.

M.A. ARBIB, *From monkey-like action recognition to human language : an evolution framework for neurolinguistics*, 2005, pp. 105-167.

propriétés miroirs. Ainsi, dans l'évolution, les gestes manuels semblent à l'origine tant des mouvements buccaux les plus simples, que des mouvements oro-laryngés plus complexes s'inscrivant dans la syllabation. Le langage serait donc apparu par étapes : *« l'intégration d'un système oro-facial à un système manuel, la formation d'un appareil de protosignes gestuels à matrice principalement mimique, l'émergence d'un protolangage bimodal (gestes et sons) et enfin, l'apparition d'un système principalement vocal. »*⁴³

M.C.CORBALIS⁴⁴ ajoute que *« Le plus difficile à comprendre, comme le souligne BURLING⁴⁵ n'est pas tant le passage du visuel au langage sonore mais l'incorporation des vocalisations dans le système des neurones miroirs. Chez le singe, certains neurones dans l'aire F5 latérale, surnommées "neurones miroir de communication", paraissent répondre à la fois à l'action et à l'observation des gestes de la bouche (protrusion des lèvres, de la langue par exemple). Mais ces actions ne sont en général que visuelles et ne font pas intervenir les vocalisations. »*⁴⁶ Les études de HICKOK et al.⁴⁷ ont montré que *« les aires frontales et temporo-pariétales sont activées lors de la production et de la perception du langage parlé, ce qui laisserait penser selon CORBALIS que la parole vocale est intégrée dans le système des neurones miroirs, principalement dans l'hémisphère gauche. »*

⁴³ G. RIZZOLATTI, op. cit. , p.181.

⁴⁴ M.C. CORBALIS, *Mirror neurons and the evolution of language*, 2009, p.6.

⁴⁵ R. BURLING, *The talking ape*, 2005.

⁴⁶ P. FERRARI, V. GALLESE, G. RIZZOLATTI, L. FOGASSI, *Mirror neurons responding to the information of ingestive and communicative mouth actions in the monkey ventral premotor cortex*, 2003, pp.1703-1704.

⁴⁷ G. HICKOK, B. BUCHSBAUM, B., C. HUMPHRIES, T. MUFTULER, *Auditory-motor interaction revealed by fMRI : Speech, music, and working memory in area Spt.* , 2003, pp.673-682.

G. HICKOK, D. POEPEL, *The cortical organization of speech processing*, 2007, pp.393-402.

2. CAS DES IMITATEURS

La popularité de l'imitation vocale ne cesse de croître dans le divertissement radiotélévisé et le monde du spectacle. Le public connaît un engouement certain pour cette faculté vocale qui fascine tant. Selon la définition, un imitateur est capable d'imiter la voix d'autres personnes. Il convient de distinguer toutefois l'imitation (où le sujet recherche la plus grande fidélité par rapport au timbre de voix original et une justesse maximale) de la caricature qui ne prend pas en compte la voix dans son intégralité mais crée seulement une variante reconnaissable. Nous tenterons donc de mettre en lumière les différents éléments permettant à l'imitateur de réaliser cette prouesse vocale qu'est l'imitation. Nous nous appuierons également sur les informations recueillies suite au questionnaire proposé à l'imitateur Michaël GREGORIO, rencontré dans la cadre de la partie pratique (annexe n°5).

2.1 PARTICULARITES ANATOMIQUES

2.1.1 CONTORSIONNISTES DU LARYNX ET DES RESONATEURS

L'imitateur, comme tout un chacun, possède une disposition laryngée classique avec deux cordes vocales et un système articulaire. Mais ce qui fait sa différence et sa force, c'est l'asymétrie de son larynx. En effet, il peut par exemple ajuster sa corde vocale droite par rapport à la gauche, celle-ci pouvant être aussi recouverte partiellement par la droite, ou encore, il peut allonger la corde vocale droite davantage que la gauche et inversement, ce qui offre une gamme d'harmoniques différents et multiples. Notons que l'imitateur conserve son harmonique vocal propre, sa propre empreinte vocale. C'est le système articulaire des cordes vocales qui lui permet d'accentuer, au gré des imitations, l'asymétrie laryngée. Cela requiert une grande rapidité et une grande précision dans l'exécution. L'épiglotte intervient également pour modifier le timbre. Elle peut se rapprocher en arrière du pharynx et des aryténoïdes selon

les cas. Les ligaments situés entre cette dernière et les cordes vocales constituent une « cathédrale vocale puissante et adaptée à chaque type d'imitation. »⁴⁸

Pour qu'une imitation soit réussie, l'imitateur doit assouplir et développer tant que possible ses caisses de résonance, sa musculature pharyngée et laryngée. Outre les cordes vocales, les résonateurs que sont le pharynx, le voile du palais, la cavité buccale, la langue, les lèvres participent grandement à la réalisation d'une imitation. M.GREGORIO s'appuie particulièrement sur les résonateurs pour parvenir à imiter les voix chantées de son répertoire et même certains instruments de musique. Ainsi, lorsqu'ils ne fonctionnent pas correctement pour cause de facteur pathologique par exemple, il n'y a aucune compensation possible.

2.1.2 BOUCLE AUDIO-PHONATOIRE

L'autre élément essentiel présent chez l'imitateur est sa capacité d'écoute et plus précisément une écoute de sa propre voix. L'imitateur percevrait sa voix sans déformation. Lorsqu'un individu quelconque entend sa voix sur un enregistrement par exemple, il ne se reconnaît pas et peut être surpris par sa production. Au contraire, l'imitateur s'entend comme il parle : le son est restitué tel qu'il est produit, via la boucle audio-phonatoire. « *Il entend la voix de l'autre comme si c'était lui qui parlait.* »⁴⁹

⁴⁸ J. ABITBOL, op. cit. , p.436.

⁴⁹ J. ABITBOL, op. cit. , p.438.

2.2 STRATEGIES

2.2.1 ASPECT AUDITIF

En vue de cerner davantage le potentiel de flexibilité de la voix humaine et de la parole, Elisabeth ZETTERHOLM⁵⁰ s'est intéressée aux imitateurs et à leurs productions vocales. Elle affirme que « *Pour réussir une imitation, l'imitateur doit sélectionner et copier diverses caractéristiques laryngées et supra-laryngées présentes chez son modèle original. Plus largement, il doit aussi mettre l'accent sur certains éléments présents dans le discours de la personne imitée comme l'emploi d'expressions propres au sujet.* »⁵¹ [...] « *Egalement, il existe des variations entre les individus, dues à des caractéristiques biologiques, auxquelles l'imitateur devra s'adapter. Il pourra s'agir de la position des dents, la taille et la forme générale des cordes vocales, les dialectes régionaux ou sociaux, l'état de santé, l'attitude de la personne, l'intention communicative, les émotions, ...* »⁵² « *Il doit donc être conscient des marqueurs généraux témoignant d'une identité groupale (dialecte régional ou social) mais aussi des marqueurs personnels dans le discours tels que la prononciation ou l'articulation.* »⁵³ Cela implique que l'imitateur doit changer ses propres paramètres et notamment articulatoires pour s'adapter au comportement d'autrui tant au niveau de la voix que de la parole. »⁵⁴

L'analyse des productions vocales d'un imitateur réalisée par Elisabeth ZETTERHOLM⁵⁵ montre qu'il est capable de grandes variations de hauteur offrant une grande souplesse d'intonation. Il est aussi capable d'adopter la bonne intonation et la bonne prononciation de dialectes qu'il ne pratique pas. Pour parfaire son imitation, il en exagère parfois le trait. Il exagère également les hésitations et les pauses.

⁵⁰ E. ZETTERHOLM, *Same speaker – different voices. A study of one impersonator and some of his different imitations*, 2006, p.70-75.

⁵¹ *ibid.* , pp.70-71.

⁵² J. PITTAM, *Voice in Social Interaction, An Interdisciplinary Approach*, 1994.

⁵³ K.R. SCHERER, *Personality markers in speech*, in K.R SCHERER and H. GILLES, (eds.), *Social markers in speech*, 1979, pp.147-209.

⁵⁴ E. ZETTERHOLM, *op. cit.* , p.71.

2.2.2 ASPECT ACOUSTIQUE

2.2.2.1 Fréquence fondamentale

Dans ses différentes productions vocales, l'imitateur est capable de faire varier la fréquence fondamentale F0 de façon importante. L'analyse menée par E.ZETTERHOLM⁵⁶ témoigne de ce phénomène. En effet, la fréquence fondamentale de l'imitateur étudié varie entre 97 et 255 Hz alors que sa propre fréquence fondamentale moyenne se situe à 118 Hz. Par ailleurs, elle souligne que l'imitateur est aussi bien capable de diminuer sa fréquence fondamentale que de l'augmenter. Le comportement F0 semble facilement imitable. L'imitateur pourrait adopter deux techniques : soit approcher le niveau global de sa voix vers le FO moyen de la parole imitée, soit suivre de façon la plus proche possible les contours intonatifs tout en gardant le niveau F0 de sa voix naturelle.⁵⁷

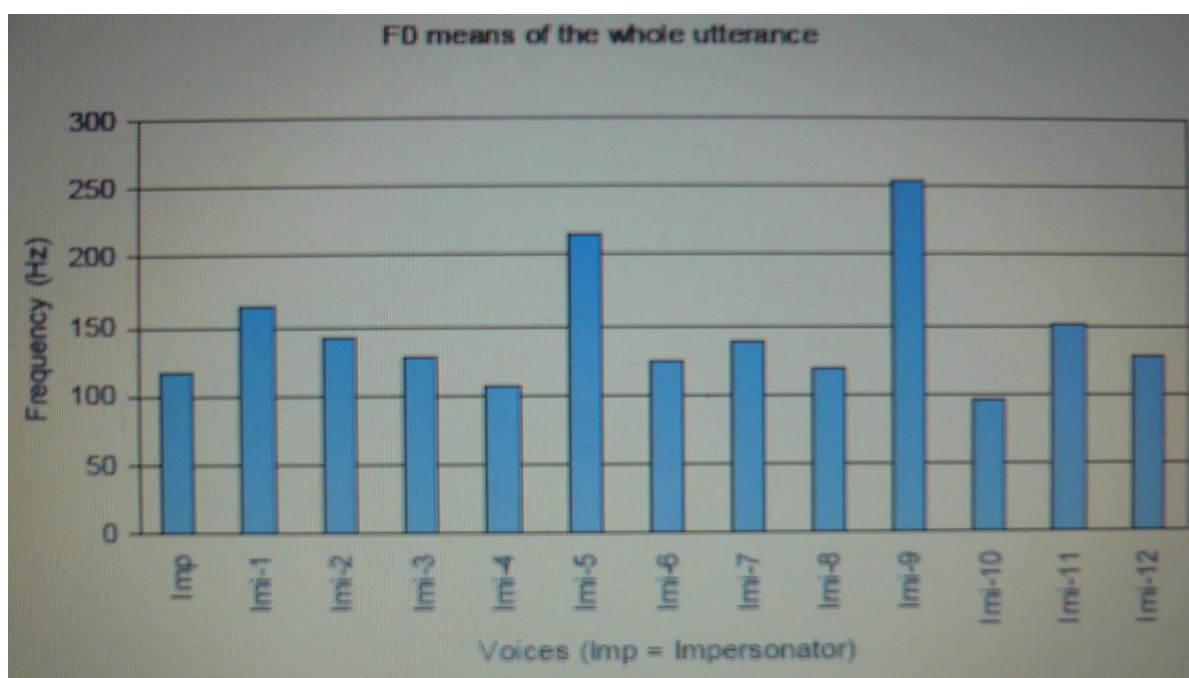


Figure 5 : Fréquence fondamentale de la voix propre de l'imitateur (Imp) et de chaque imitation (Imi-)⁵⁸

⁵⁵ Ibid, p. 75.

⁵⁶ ibid, p.73.

⁵⁷ J. MEJVALDOVA, *Caractéristiques temporelles de la parole imitée*.

⁵⁸ E. ZETTERHOLM, op. cit. , p.73.

2.2.2.2 Formants

Qu'en est-il d'une éventuelle modification de fréquence des formants des voyelles entre la voix de l'imitateur et la voix produite lors d'une imitation ? L'analyse menée par E.ZETTERHOLM⁵⁹ se base notamment sur l'étude des formants F1 à F4. La distribution des formants F1 et F2 montre qu'il y a bien un ajustement des formants des voyelles chez l'imitateur même si cela ne conduit pas à une production identique à la voix cible.

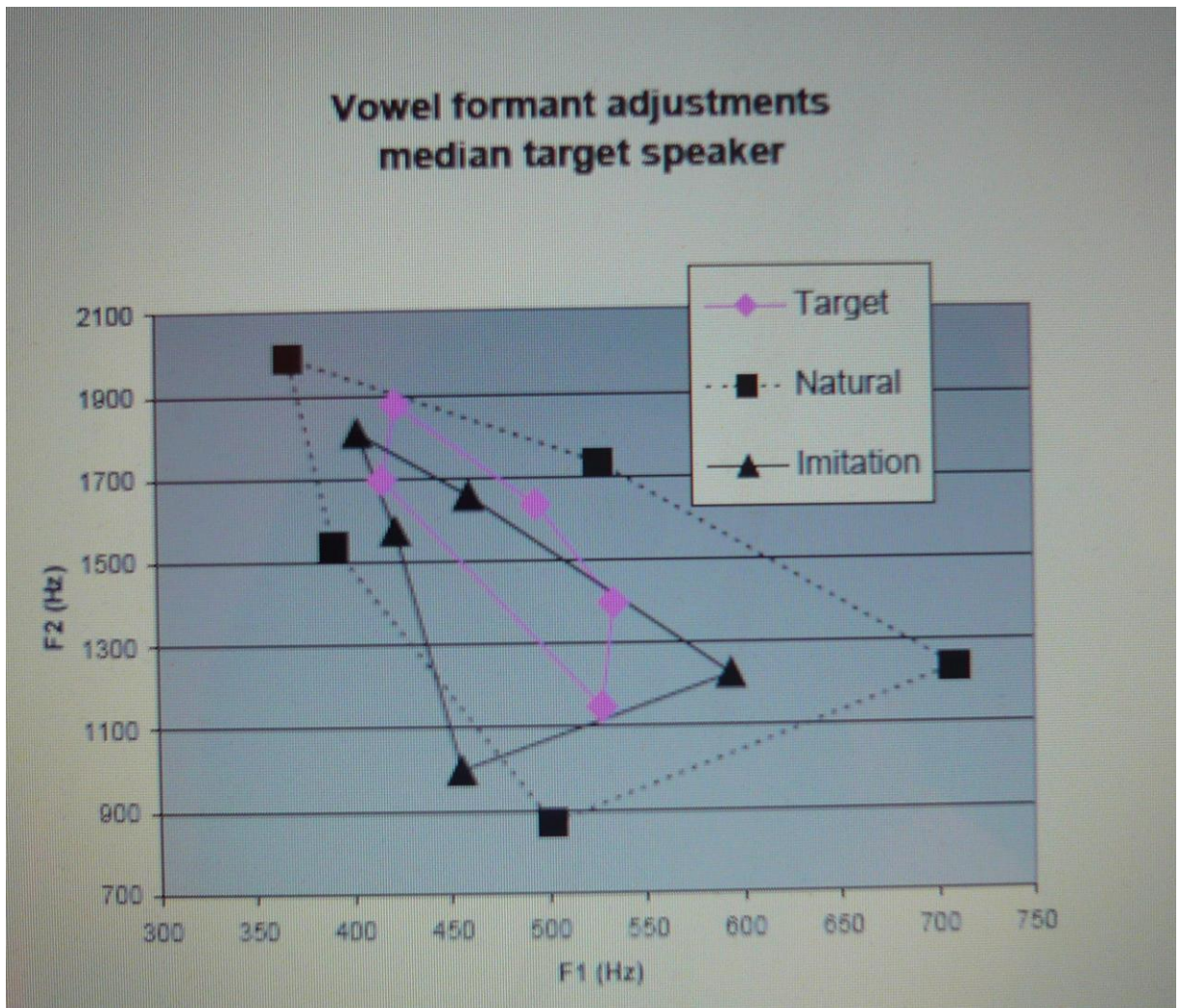


Figure 4 : Distribution des formants des voyelles lors de la production de l'émetteur cible imité (en rose), lors de la production propre de l'imitateur (carré noir), et lors de l'imitation réalisée (triangle noir).⁶⁰

⁵⁹ ibid. , p.74.

⁶⁰ E. ZETTERHOLM, D. ELENIOUS, M. BLOMBERG, A comparison between human perception and a speaker verification system score of a voice imitation, 2004.

2.2.2.3 Durée

Les analyses destinées à mettre en évidence d'éventuelles variations de durée entre la production de la personne imitée, la production personnelle de l'imitateur et la production réalisée par imitation convergent vers le même résultat. Qu'elles soient à l'échelle du mot ou du phonème, il apparaît dans tous les cas que la durée dans l'imitation est plus proche de celle de la voix de l'imitateur que de la voix imitée. A.ERICKSON et P.WRETLING⁶¹ ont étudié le paramètre de la durée : « *La durée articulatoire à un niveau plus local semble être plus rigide.* »

Les propriétés temporelles, la durée globale ainsi que la durée proportionnelle des segments de la parole s'avèreraient être des caractéristiques stables du locuteur. Ces phénomènes seraient donc difficilement imitables. L'imitateur est contraint soit de rester dans ses habitudes discursives, soit d'exagérer davantage. Le contrôle moteur de la durée semble donc être moins aisé que le contrôle moteur de la production des caractéristiques spectrales de la parole.⁶²

2.2.2.4 Articulation

E.ZETTERHOLM a également mis en évidence une stratégie d'adaptation au niveau de l'articulation, cela se produisant lorsque l'imitateur imite le tempo rapide ou lent d'un discours : « *Il y a adaptation de l'articulation selon que l'imitateur imite un discours au tempo rapide ou lent.* »⁶³ En effet, l'étude d'un imitateur dont le taux d'articulation est de 5,8 syllabes par seconde (en temps normal) montre qu'il existe une variabilité de ce taux lors des imitations. En effet, une de ses imitations se situent à un taux d'articulation supérieur (plus de 6 syllabes par seconde) alors que l'ensemble des autres imitations analysées se situe à un taux d'articulation inférieur (entre 3 et 5,5 syllabes par seconde). Un discours prononcé lentement, des hésitations sur certains sons, des voyelles prolongées peuvent expliquer un taux d'articulation moins élevé.

⁶¹ A. ERIKSSON, P. WRETLING, *How flexible is the human voice ? A case study of mimicry*, 1997, pp.1043-1046.

⁶² J. MEJVALDOVA, op. cit.

⁶³ E. ZETTERHOLM, op. cit. ,p.73.

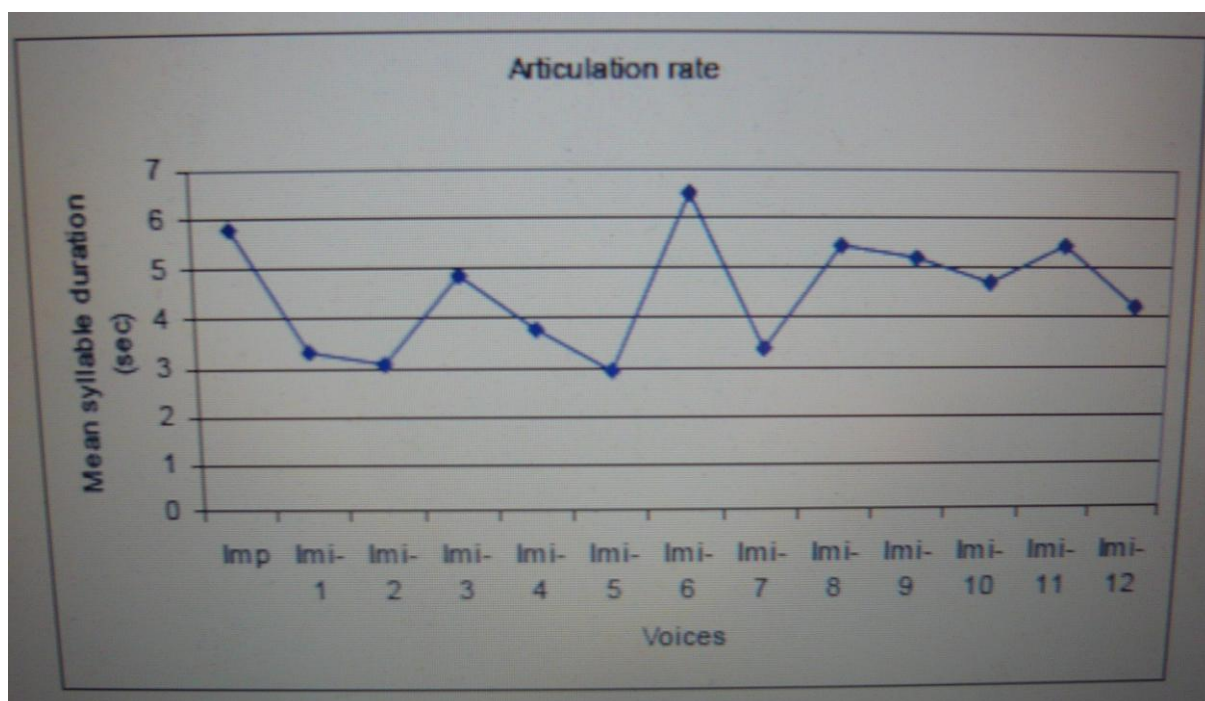


Figure 5 : Taux d'articulation (nombre de syllabes par seconde) chez l'imitateur (Imp) et pour chacune de ses imitations (Imi-).⁶⁴

2.3 RECIT D'UN IMITATEUR :

M.GREGORIO raconte qu'il a réalisé sa première imitation vers l'âge de 15 ans, sans réel élément déclenchant. Comme bon nombre d'adolescents, il était passionné de musique et écoutait divers groupes tels que Radiohead, Muse ... En chantant leurs chansons, il s'est rendu compte qu'il avait tendance à chanter comme eux, « à la manière » de ses idoles. C'est ainsi qu'il a découvert sa « prédisposition » pour l'imitation. Michaël ne parle pas de don. Pour lui, l'imitation est une prédisposition qu'il faut travailler. Ce n'est pas une science, la technique pour tenter de se rapprocher au mieux de l'objectif est quelque part propre à chacun. Son répertoire vocal présente plus de quarante voix pour lesquelles il ne note pas de fil conducteur particulier. Il explique par ailleurs qu'il n'y a pas de limite établie pour réaliser une voix. Il n'y a rien de concret. Sa voix va s'adapter ou parfois non. Ainsi, il arrive que certains éléments empêchent l'imitation d'une voix bien que celle-ci se situe dans sa tessiture.

⁶⁴ E. ZETTERHOLM, op. cit., p.74.

Il souligne également qu'une voix n'est jamais acquise, elle peut s'améliorer ou régresser. L'imitation est subjective, c'est une « suggestion » qui sera appréciée différemment selon que le public connaît très bien la voix imitée ou non. Il n'y a donc pas de règle pour travailler une imitation. Michaël essaie d'écouter la voix ciblée au maximum pour la reproduire au mieux. Il s'enregistre et recueille l'avis de ses proches quant à la qualité de la suggestion. Le travail d'écoute est donc très important. Enfin, Michaël ne suit pas d'entraînement vocal particulier. Il veille seulement à protéger sa voix et à l'échauffer avant tout travail.

Nous exposerons dans la partie pratique son ressenti face au protocole que nous lui avons proposé.

3. LOGICIEL PRAAT

PRAAT est un logiciel d'analyse de la parole développé par Paul Boersma et David Weenink de l'Institut des Sciences phonétiques de l'Université d'Amsterdam en 1996. Il est téléchargeable gratuitement sur le site www.praat.org, est disponible sur les plates-formes Windows, Macintosh et Unix et bénéficie de fréquentes mises à jour. L'intérêt de ce logiciel est d'offrir une approche objective de la voix avec une étude précise des différents paramètres vocaux, ce qui est central dans notre étude.

PRAAT est notamment un logiciel de transcription, d'analyse et de modélisation en phonétique/phonologie. Il propose de nombreux modes de représentation graphique du signal sonore, tels que les enveloppes sonores, les spectres instantanés ou moyennés, ou encore les spectrogrammes à bande large et à bande étroite. Il permet aussi de mesurer un certain nombre de grandeurs utiles aux spécialistes de la voix tels que le fondamental usuel, les fréquences des formants des voyelles, le degré de perturbation de la fréquence ou de l'intensité, et encore bien d'autres, grâce à sa grande souplesse d'utilisation. Les sons à analyser peuvent être directement enregistrés sous Praat à l'aide d'un micro, ou chargés en mémoire à partir d'un fichier au format AIFF, WAV, AU ou NIST.

Nous n'emploierons que certaines de ses fonctionnalités : les analyses phonétiques d'une grande précision, les divers modes de représentation graphique du signal sonore notamment les spectrogrammes, le module de synthèse vocale très perfectionné aux possibilités variées de filtrage et de modification des sons. C'est particulièrement cette possibilité de modification des sons que nous allons utiliser pour définir la zone de sécurité prosodique. Nous développerons ce principe dans la partie réservée au protocole dans la partie pratique.

Lors du lancement du logiciel, deux fenêtres s'affichent à l'écran. A gauche, la fenêtre « Objects » et à droite, la fenêtre « Picture ». La fenêtre Picture sert uniquement à exporter des graphiques vers d'autres programmes (logiciels de dessin ou de traitement de texte). Il n'est pas nécessaire de la laisser ouverte. La fenêtre Objects est en revanche très importante. C'est à partir de cette fenêtre que nous allons pouvoir charger en mémoire les sons que nous voulons analyser, enregistrer des sons avec un micro, ou encore, manipuler les différents objets graphiques (spectres, spectrogrammes...) que nous serons amenés à créer. Le logiciel

permet soit d'enregistrer directement un son à l'aide d'un micro, soit de charger un son déjà enregistré et gardé en mémoire.

Voici l'environnement de travail de PRAAT :

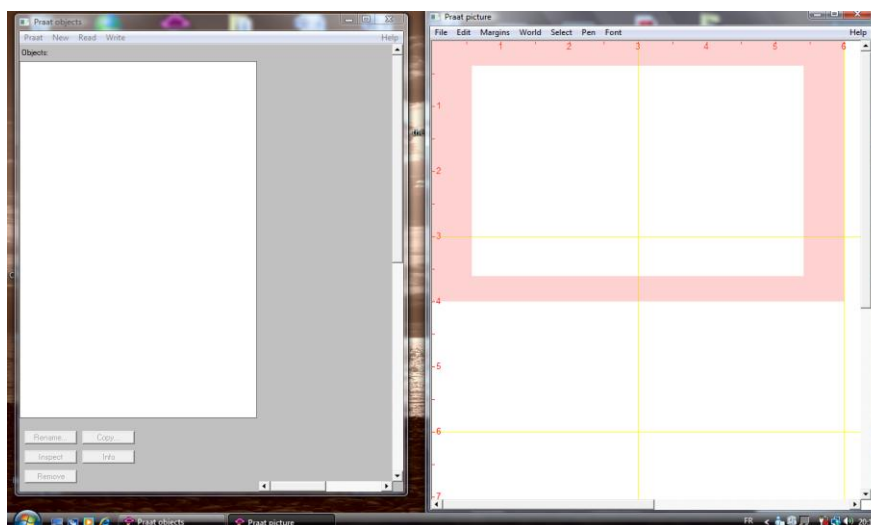


Figure 6 : Environnement de travail de PRAAT.

Nous avons vu que le logiciel offre diverses représentations graphiques du signal sonore. Nous nous intéresserons particulièrement aux spectrogrammes. Ils représentent le son en trois dimensions : en abscisse nous trouvons le temps, en ordonnées nous trouvons la fréquence et enfin, l'intensité est symbolisée par l'aspect plus ou moins foncé du tracé. Pour générer un spectrogramme, PRAAT utilise un algorithme couramment utilisé dans le traitement numérique du signal. Il s'agit de la Transformée Rapide de FOURIER. A intervalles de temps réguliers, le programme isole une petite portion du signal sonore. Cette portion est appelée « fenêtre » par les physiciens. Pour chacune de ces fenêtres, PRAAT calcule le spectre d'amplitude correspondant. La juxtaposition sur un axe temporel de tous les spectres d'amplitude calculés forme le spectrogramme. Deux facteurs principaux déterminent l'apparence du spectrogramme : le type de fenêtre utilisé, et la taille de cette fenêtre. On distingue ainsi deux grands types de spectrogrammes : - les spectrogrammes à bande large où le programme utilise une fenêtre très courte (quelques millisecondes) pour les calculer. Ceux-ci ne permettent pas de distinguer le fondamental et ses harmoniques sur le programme obtenu.

- les spectrogrammes à bandes étroites où le programme utilise une fenêtre plus longue (quelques centièmes de secondes) pour les calculer. Ils permettent de distinguer le fondamental et ses harmoniques.

En pratique, les spectrogrammes à bande large sont surtout utilisés par les phonéticiens, parce qu'ils permettent une bonne visualisation des formants et des transitions entre les phonèmes. L'analyse spectrale de la voix s'appuie plus volontiers sur l'analyse des spectres à bande étroite, puisque ce type de tracé est plus adapté à l'étude de la F_0 et de la composante harmonique relative.

Sur le spectrogramme, selon les cas, différentes courbes en surimpression peuvent être visualisées : la fréquence fondamentale, les formants ou encore les périodes.

Voici une illustration de la représentation en spectrogramme que nous offre PRAAT et sur laquelle nous nous appuierons au cours de l'étude :

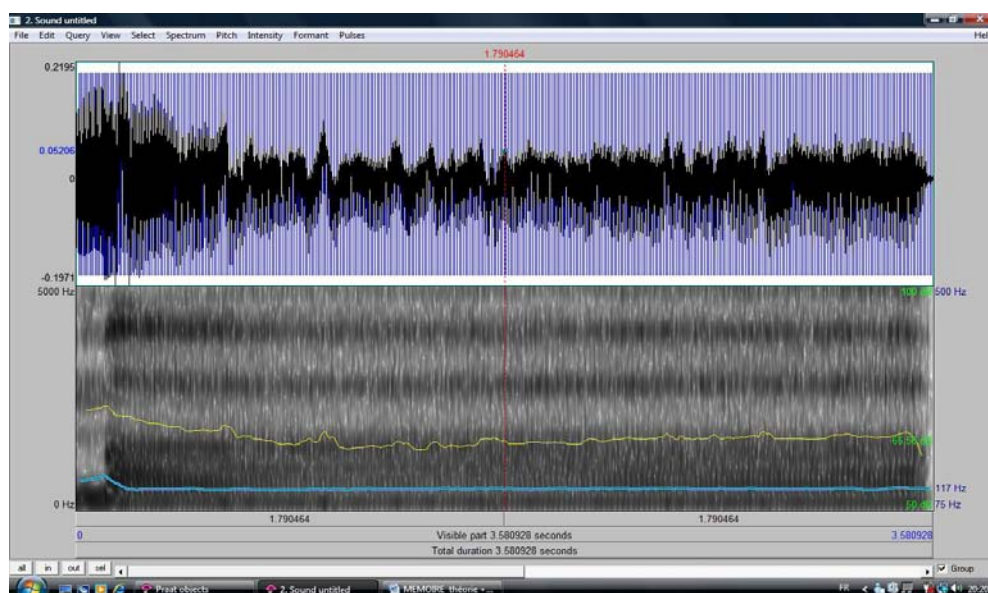


Figure 7 : Fenêtre spectrogramme

La partie supérieure de la fenêtre montre la forme d'onde (en noir), c'est-à-dire l'amplitude du signal en fonction du temps. Cela facilite grandement la lecture et le repérage des moments où il se passe quelque chose. Les lignes bleues verticales (menu Pulses) signalent les occlusions glottiques (en tout cas le repérage théorique de celles-ci).

La partie inférieure de la fenêtre montre différents éléments d'analyse du signal dont, en cyan, la hauteur de la fondamentale (menu Pitch). Cette donnée est centrale dans notre étude. Notons que l'échelle qui correspond à cette courbe est à droite de la figure (ici entre 75 et 500 Hz). Les fréquences moyennes, minima et maxima peuvent être récupérées (menu Pulses, « *voice report* »). Egalement, nous retrouvons en noir et blanc, le spectrogramme (menu Spectrum). Notons que l'échelle correspondant à cette représentation est à gauche de la figure (ici entre 0 et 5000 Hz). Les paramètres de l'ensemble de ces analyses sont modifiables dans les menus correspondants.

Nous nous servons également d'une représentation sous forme de ligne mélodique qui permet de voir le signal et chaque point prosodique de la production d'un sujet. Il est possible de manipuler la prosodie : chaque petit rond vert peut être déplacé soit manuellement, soit par un calcul (comme nous le verrons dans la partie pratique).

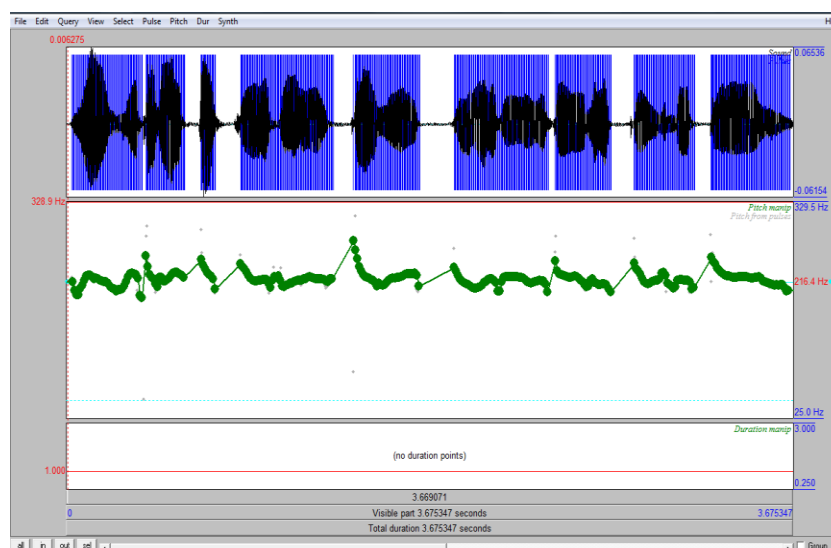


Figure 8 : Fenêtre ligne mélodique

Cette fenêtre (obtenue grâce à la fonction « *To manipulation* ») contient le signal et en dessous la ligne mélodique correspondant au signal. C'est sur cette ligne mélodique que nous travaillerons pour définir la zone de sécurité prosodique du sujet.

Notons qu'à partir de la ligne mélodique, nous pouvons aussi obtenir le spectrogramme correspondant grâce à la fonction « *Publish resynthesis* » dans le menu File).

4. ZONE DE SECURITE PROSODIQUE

Pour définir la zone de sécurité prosodique, nous nous sommes basés sur certaines constatations. Tout d'abord, nous savons que l'analyse de la dynamique vocale, c'est-à-dire l'analyse du comportement de la voix lors des variations de hauteur, permet d'objectiver des éléments tels que des instabilités majeures, des harmoniques mal dessinés ou encore les zones de transition entre les différents registres des cordes vocales (passage entre voix de poitrine et voix de tête notamment). Notons que dans les dysphonies, il y a la plupart du temps altération de ces transitions, qui interviennent précocement (décalage vers les graves), de manière brutale et involontaire. Mais elle permet aussi de visualiser où se situe la zone de bon rendement des voyelles notamment. Cette zone de bon rendement se matérialise par un bon rapport entre l'énergie dépensée par le patient et la qualité du son produit. Elle correspond à la zone où les cordes vocales ont la meilleure vibration possible et un bon accord avec les résonateurs. Pour la définir, il faut observer et sélectionner sur le spectrogramme la zone la plus large possible, compatible avec la hauteur de voix conversationnelle que souhaite le patient, où il n'y a pas de transition, pas d'instabilité majeure, où les harmoniques sont bien dessinés. Cette sélection permet de relever la valeur minimum et la valeur maximum de cette zone de bon rendement pour une voyelle donnée et d'obtenir ainsi l'intervalle de fréquences où le rendement est optimal.

Puis, nous savons qu'il existe le triangle vocalique à l'intérieur duquel toutes les voyelles de la langue française sont produites. Les extrémités de ce triangle étant constituées des voyelles [a], [i] et [u].

La zone de sécurité prosodique se définit donc par un intervalle minimum constituant l'intersection des zones de bon rendement obtenues pour l'émission des trois voyelles du triangle vocalique, les autres voyelles étant forcément contenues dans cet intervalle (le protocole dans la partie pratique et l'annexe n°2 préciseront ces éléments). En d'autres termes, la zone de sécurité prosodique est une zone réduite de la dynamique vocale, au cœur du triangle vocalique, limitant les amplitudes fréquentielles et les variations mélodiques susceptibles d'être néfastes pour la voix du sujet. L'intervalle correspondant à la zone de sécurité prosodique permettra alors d'« écraser » la voix du sujet pour diminuer l'importance des variations prosodiques présentes dans sa production spontanée.

5. ANATOMIE - PHYSIOLOGIE

Nous nous attacherons ici à rappeler les éléments anatomiques et physiologiques intervenant dans la phonation et qui peuvent participer et/ou être modulés de façon consciente ou inconsciente lors d'une imitation.

5.1 PRODUCTION DES SONS

Une bonne émission vocale, qu'elle soit parlée ou chantée, spontanée ou imitée, nécessite l'intégrité fonctionnelle des organes phonateurs. Elle requiert également une synergie de fonctionnement et une bonne coordination entre respiration, vibreur laryngé, corps vocal.

Ainsi, trois équilibres sont nécessaires, en lien avec la notion de verticalité du corps : équilibre diaphragme-abdominaux, équilibre souffle-vibreur, équilibre glotto-résonantiel.

5.1.1 POSTURE

La notion de verticalité du corps est primordiale dans la production vocale. Elle va impacter l'ensemble des organes, quels qu'ils soient. Comme nous le verrons plus loin, la phonation nécessite une liberté thoracique et abdominale ; la qualité de celle-ci dépend inévitablement de l'attitude corporelle. En adoptant une position et des appuis convenables, les sollicitations scapulaires excessives, les tensions cervicales ou faciales pourront alors être évitées.

5.1.1.1 Position d'équilibre, verticalité du corps

La position d'équilibre passe par la recherche d'un axe débutant au sommet du crâne et se terminant entre les deux pieds. La verticalité permet une harmonisation de l'appareil vocal (soufflerie, vibreurs et résonateurs). Ceci est important car le larynx est suspendu à des structures avoisinantes. Pour une statique optimale, différents points doivent être respectés.

G.CORNUT⁶⁵ souligne l'importance d'un appui sur les deux pieds légèrement écartés, ancrés dans le sol de façon stable. Il est particulièrement opportun pour limiter l'effort et l'état de tension au niveau du cou lors d'une projection vocale. En effet, le sujet centre l'effort dans le bas de son corps. Il libère ainsi ses abdominaux.

Aussi, les genoux doivent être en légère flexion, en harmonie avec une bascule du bassin modérée. La cambrure lombaire ne doit pas être exagérée pour ne pas gêner la statique de l'étage corporel supérieur.

La colonne vertébrale suit l'axe vertical, favorisant l'ouverture thoracique. Concernant la colonne cervicale, elle est un peu redressée. Le menton est quelque peu relâché pour permettre une position confortable de la tête par rapport à la colonne vertébrale et libérer le larynx. L'idéal est de s'aider du regard pour s'assurer d'un placement céphalique correct. Il doit se porter légèrement en-dessous de la ligne de l'horizon.

5.1.1.2 Enjeu postural

Il ne doit pas y avoir de confusion entre verticalité et rigidité, sous peine de bloquer l'amplitude respiratoire. En agissant de façon inappropriée sur la dynamique musculaire, les parties rigides osseuses et les parties souples qui les articulent ne peuvent interagir de façon harmonieuse. Il faut donc savoir alterner entre travail et détente. Le visage ne doit pas être projeté en avant exagérément de même que l'on décommande toute tension musculaire et déséquilibre. J.ABITBOL⁶⁶ cite l'exemple, tiré de son expérience, d'une hyper-extension du membre supérieur droit ayant provoqué un déséquilibre des muscles pectoraux, un déséquilibre des muscles laryngés, puis une hyper-contraction de l'articulation crico-aryténoïdienne, qui *ipso facto* était coincée par les muscles et cartilages du cou.

⁶⁵ G. CORNUT, *La Voix*, 1990, p.102.

⁶⁶ J. ABITBOL, *L'Odyssée de la Voix*, 2005, p.296.

De plus, il faut aussi savoir jouer entre tonicité et fluidité. La tonicité cervicale doit être peu prononcée quelle que soit la production vocale comme l'évoque J.SARFATI⁶⁷. Elle met en évidence l'opposition du tonus fort/faible dans la réalisation des consonnes, le tout sous la coupe d'un geste fluide.

5.1.2 RESPIRATION

Pour que les cordes vocales entrent en vibration lorsqu'elles sont en position phonique, l'intervention de la soufflerie pulmonaire est nécessaire. C'est le moteur de leur vibration. La respiration fait appel à un ensemble de phénomènes physiologiques par lesquels l'air pénètre dans les alvéoles pulmonaires à l'inspiration et en sort à l'expiration, cela grâce à l'intervention de différents muscles.

5.1.2.1 Musculature respiratoire et innervation

La musculature respiratoire est constituée de muscles antagonistes devant trouver le bon équilibre. On retrouve :

5.1.2.1.1 Muscles dits inspireurs

- Tout d'abord le **diaphragme**, muscle inspireur principal. A l'inspiration, ses fibres se contractent, se raccourcissent, entraînant sa descente. Une dépression se crée dans la cage thoracique, elle s'écarte, la pression est négative, permettant alors à l'air de pénétrer.

Lors de l'expiration, les fibres se rallongent, il se relâche et remonte. Son mouvement doit se coordonner avec les autres muscles, notamment avec les abdominaux qui participent à l'expiration de l'air.

Son innervation motrice est assurée par le nerf phrénique, branche du plexus cervical profond.

⁶⁷ J. SARFATI, *Soigner la Voix*, 2002, p.72.

- Les muscles **intercostaux externes**⁶⁸ et **moyens** élèvent l'arc costal. Leur innervation provient des branches musculaires collatérales des nerfs intercostaux.

- Les muscles inspireurs accessoires que sont les **scalènes** et **le sterno-cléido-mastoïdien** agissent notamment en cas d'inspiration profonde. Le diaphragme est contracté, la cage thoracique se soulève dans sa partie supérieure au niveau sternum et s'élargit grâce aux muscles inspireurs accessoires.

Le sterno-cléido-mastoïdien est innervé par la racine spinale du nerf accessoire (XI) et les scalènes par un rameau nerveux des cinquième et sixième nerfs cervicaux.

5.1.2.1.2 Muscles dits expirateurs

- Les muscles de la sangle abdominale, à savoir le **transverse**, **le petit** et **le grand oblique**. Ils permettent l'expulsion de l'air. Antagonistes du diaphragme, leur action doit être synchronisée.

Le transverse et le petit oblique dépendent des rameaux des nerfs spinaux T 8 à T 12 et des nerfs ilio-hypogastrique et ilio-inguinal. Le grand oblique quant à lui dépend des rameaux des nerfs spinaux T 7 à T 12 et du nerf ilio-hypogastrique.⁶⁹

- Les muscles **intercostaux internes**⁷⁰ abaissent l'arc costal. Leur innervation provient également des branches musculaires collatérales des nerfs intercostaux.

5.1.2.2 En phonation

« La phonation entraîne l'adoption d'un rythme respiratoire particulier, fondamentalement différent de celui de la respiration calme. Dans la respiration calme, en effet, (...) les deux temps respiratoires sont de durée comparable. Dans la phonation, (...) »

⁶⁸ N.R. BORLEY, R.H. WHITAKER, *Anatomie : Angiologie & Nerfs crâniens et nerfs rachidiens & Organigrammes*, 2003, p.181

⁶⁹ G.J. TORTORA, S.R. GRABOWSKI, *Principes d'anatomie et de physiologie*, 2002, p.345.

⁷⁰ Ibid.

l'inspiration se raccourcit considérablement et prend la signification d'un élan du geste phonatoire (...). C'est le temps expiratoire qui est phonatoire. La voix peut ainsi être conçue comme une expiration sonorisée. Il convient cependant de bien distinguer la voix, du courant d'air qui la produit. »⁷¹

Par ailleurs, la pression pulmonaire doit être réglée continuellement. En effet, lorsque les cordes vocales sont en adduction (position phonatoire), la pression doit être augmentée pour vaincre leur résistance. Elle doit accroître de façon proportionnelle à la tension des cordes vocales. Ainsi, plus la voix est aiguë, plus les cordes sont rigides et requièrent donc une pression pulmonaire d'autant plus importante pour maintenir la même amplitude vibratoire. *« Le maintien d'une pression sous-glottique constante nécessite une adaptation fine des muscles respirateurs. Au début de l'expiration, les muscles inspiratoires restent contractés et freinent le mouvement de fermeture de la cage thoracique. Au fur et à mesure que le volume pulmonaire décroît, cette activité diminue et la pression est maintenue par l'action croissante des muscles expirateurs. »⁷²*

Une respiration buccale et costo-abdominale semble la plus adéquate dans la voix parlée et plus largement dans la voix chantée. En effet, la respiration buccale permet une prise d'air plus rapide et discrète par rapport à la respiration nasale. L'ouverture buccale détend et fait légèrement descendre le larynx. Concernant la respiration costo-abdominale, comme le rappelle M.F CASTAREDE, le geste inspiratoire permet l'union en continu d'un abaissement diaphragmatique accompagné d'une détente abdominale et d'un mouvement d'écartement costal qui entraîne une dilatation de toute la cage thoracique, en particulier au niveau des côtes inférieures. Cette respiration permet un remplissage pulmonaire efficace ; ses bénéfices sont nombreux notamment au niveau phonatoire : elle améliore le fonctionnement du diaphragme et rend également possible la remontée progressive de ce dernier. Elle offre donc un meilleur contrôle de l'arrivée d'air vers les cordes vocales en laissant la musculature du cou et des épaules libre de toute tension.

⁷¹ F. LE HUCHE, A. ALLALI, *La Voix-Tome I, Anatomie et Physiologie des Organes de la Voix et de la Parole*, 2001, pp. 48-49.

⁷² J.A. RONDAL, X. SERON, *Troubles du langage ; Bases théoriques, Diagnostic et Rééducation*, 2000, p.100

« Ce contrôle s'effectue comme suit : les abdominaux tempèrent leur poussée expiratrice en continuant d'opposer une certaine résistance à l'élévation du diaphragme ; cela donne puissance et précision au mouvement global. L'avantage fonctionnel de ce mode respiratoire n'est pas des moindres : il permet au larynx de jouer son rôle de vibreur avec le maximum de souplesse et de liberté ; il assure une arrivée d'air à la fois adaptable, puissante et précise, ce qui réunit les conditions d'un fonctionnement vocal optimum. »⁷³

La sangle abdominale rend possible la tenue du son, la durée du discours, les variations d'intensité, le soutien des aigus. On peut agir sur elle de façon consciente à la différence du diaphragme par exemple.

Ainsi, une bonne coordination pneumo-phonique sous la coupe de la notion de verticalité du corps est essentielle pour la production des sons.

5.1.3 LARYNX

Le larynx, organe de la phonation, est un conduit fibro-musculo-cartilagineux, se situant entre le pharynx, la base de la langue qui le surplombe et la trachée qui le prolonge vers le bas. Il intervient dans les fonctions de phonation, déglutition, respiration. Nous nous centrerons ici sur sa fonction d'organe phonateur. Pour que le larynx puisse jouer pleinement son rôle dans la phonation, et plus largement dans les facultés d'imitation, il faut une bonne coordination entre les différents muscles laryngés, la vibration des cordes vocales et les commandes nerveuses.

5.1.3.1 Musculature

Deux types de muscles affectent la fonction du larynx. D'une part, les muscles extrinsèques qui se fixent sur le larynx et ont un autre point de fixation à l'extérieur de celui-ci. D'autre part, les muscles intrinsèques qui ont leurs deux attachements à l'intérieur du larynx.

⁷³ M.F. CASTAREDE, *La Voix et ses Sortilèges*, 2004, pp.109-110.

5.1.3.1.1 Muscles extrinsèques⁷⁴

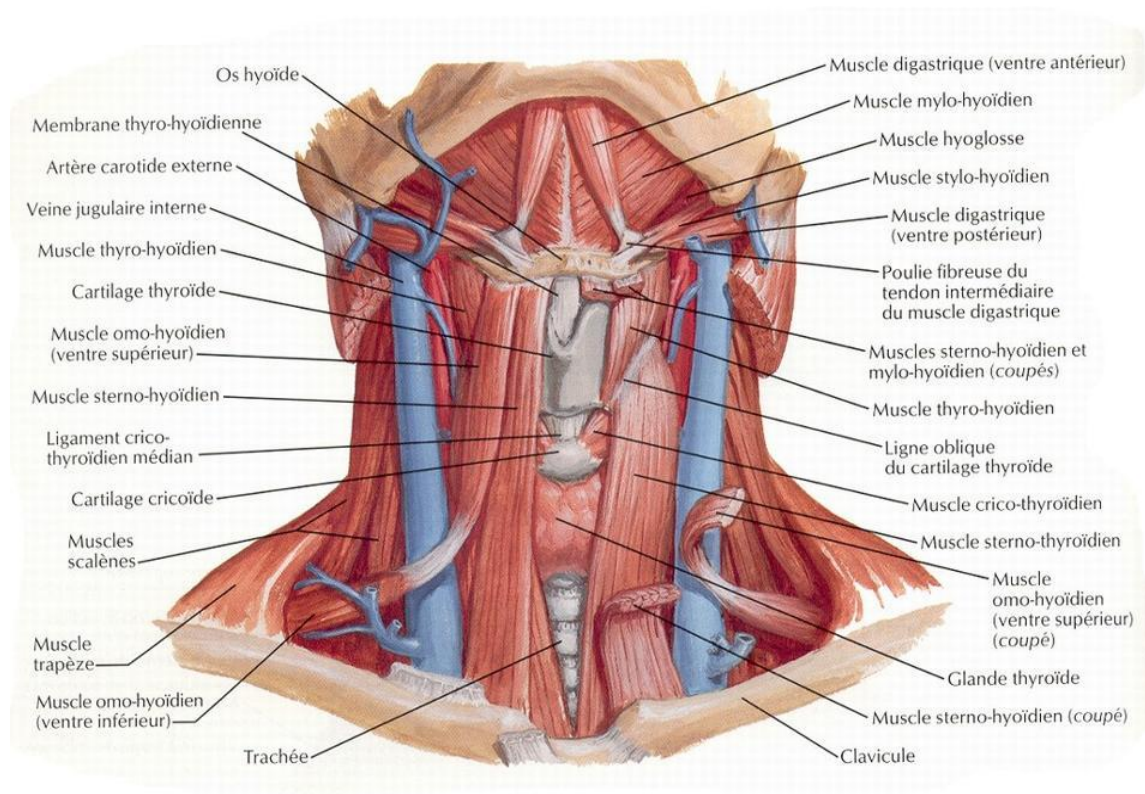


Figure 1 : Muscles sous et sus-hyoïdiens⁷⁵

Ils sont responsables de la suspension et de la mobilité de l'ensemble du larynx. Ils sont tous rattachés à l'os hyoïde et se fixent à la mandibule, à la mastoïde, ou au thorax. Il s'agit d'un ensemble de lanières qui maintiennent la position du larynx dans le cou pour permettre une bonne efficacité de l'action des muscles endolaryngés intrinsèques. On retrouve :

5.1.3.1.1.1 Muscles sous-hyoïdiens

. Le **thyro-hyoïdien** est situé sous la portion supérieure du sterno-hyoïdien. Il prend naissance sur la crête oblique du cartilage thyroïde et sur le bord inférieur de la grande corne de l'os hyoïde.

⁷⁴ D.H. MCFARLAND, F.H. NETTER, *L'Anatomie en Orthophonie*, 2006, pp. 92 – 103.

⁷⁵ Ibid. , p.104.

Sa contraction rapproche l'os hyoïde et le cartilage thyroïde surtout en avant. En fonction de l'action simultanée des autres muscles, il peut abaisser l'os hyoïde ou soulever le thyroïde.

. Le **sterno-cléido-hyoïdien** prend naissance au niveau de la clavicule et de la face postérieure du manubrium sternal. Il s'insère en haut au niveau du bord inférieur du corps de l'os hyoïde.

Sa contraction abaisse l'os hyoïde et le larynx.

. L'**omo-hyoïdien** prolonge en dehors le précédent. Il s'insère en bas sur le bord supérieur de l'omoplate.

Son action est également d'abaisser l'os hyoïde.

. Le **sterno-thyroïdien**, situé sous le sterno-hyoïdien, prend naissance sur le bord postérieur du manubrium sternal et s'insère en haut sur la crête oblique.

Sa contraction abaisse le thyroïde et le larynx.

5.1.3.1.1.2 Muscles sus-hyoïdiens

Ils permettent l'ouverture de la bouche en tirant la mandibule vers le bas, l'élévation de l'os hyoïde, le mouvement du larynx vers l'arrière, le haut ou l'avant.

. Le **mylo-hyoïdien** forme le plancher de la bouche. Il prend son origine le long de la ligne oblique interne sur la surface interne de la mandibule et s'insère sur un raphé médian avec les fibres controlatérales.

Son action est d'élever l'os hyoïde et l'attirer vers l'avant.

. Le **digastrique** est constitué d'un ventre antérieur et d'un ventre postérieur. Le premier prend naissance au niveau du bord inférieur de la mandibule près de la symphyse. Le second va de l'apophyse mastoïde au tendon intermédiaire de l'os hyoïde.

Leur rôle est respectivement d'abaisser la mandibule en bas et en arrière ; de tirer l'os hyoïde postérieurement et vers le haut.

. Le **génio-hyoïdien** se situe au-dessus de l'os hyoïde. Il s'étend des épines mentonnières et de la symphyse mentonnière de la mandibule jusqu'à la surface antérieure du corps de l'os hyoïde.

Il a pour but l'élévation et la protraction de l'os hyoïde.

. Le **stylo-hyoïdien** est parallèle au ventre postérieur du digastrique. Il prend son origine sur l'apophyse styloïde de l'os temporal et s'insère sur le corps de l'os hyoïde.

Il élève et attire vers l'arrière l'os hyoïde.

5.1.3.1.2 Muscles intrinsèques⁷⁶

Ils sont responsables de l'adduction, de l'abduction, et du réglage de la tension des cordes vocales. On retrouve :

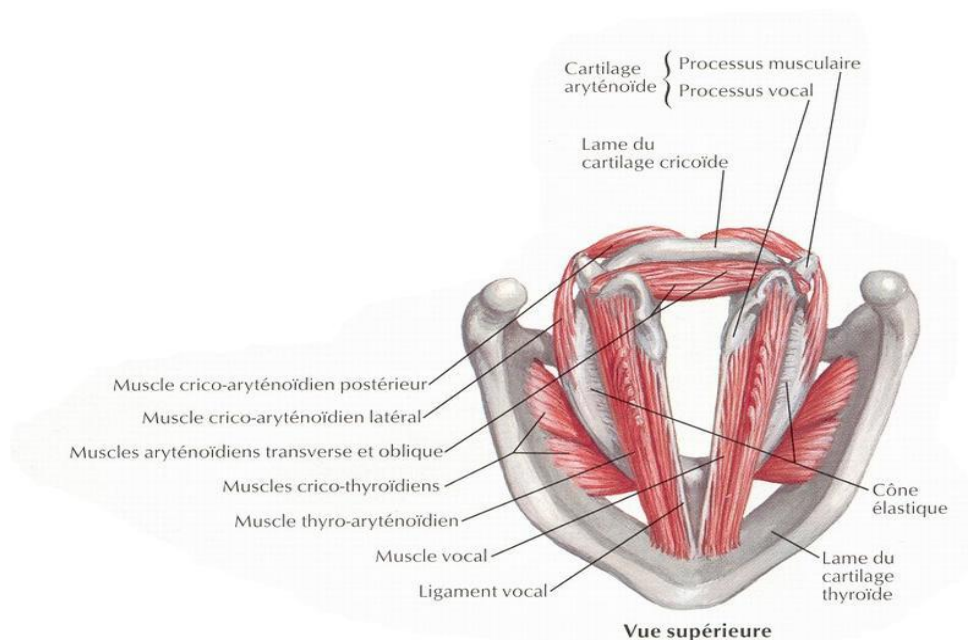


Figure 2 : Muscles intrinsèques du larynx.⁷⁷

⁷⁶ Ibid. pp. 93-101.

⁷⁷ Ibid. , p.95.

5.1.3.1.2.1 Dits tenseurs

Le ***crico-thyroïdien*** prend son origine sur la portion antérieure et latérale de l'anneau cricoïdien. Il comprend deux parties : la pars obliqua qui s'insère sur la moitié postérieure de la face interne de l'aile thyroïdienne et la pars recta qui va directement s'insérer sur le bord inférieur du cartilage thyroïde.

Sa contraction, est responsable d'un abaissement, d'une élongation et d'une mise en tension de la corde vocale. Cette contraction contribue à l'adduction des cordes vocales en position paramédiane ainsi qu'à une augmentation simultanée de la tension longitudinale. Le muscle crico-thyroïdien rend ainsi plus fin le bord libre de la corde vocale. Cela a pour conséquence d'augmenter la tonalité de la voix et rétrécir la glotte.

5.1.3.1.2.2 Dits dilatateurs (abducteurs)

Le ***crico-aryténoïdien postérieur*** est en forme d'éventail. Il prend son origine sur une large zone de la surface postérieure du chaton cricoïdien et s'insère sur la face postérieure de l'apophyse musculaire de l'aryténoïde.

Il est responsable de l'abduction de cordes vocales, de leur élévation et de leur allongement grâce à une rotation du cartilage aryténoïde latéralement et vers l'arrière. Son action augmente la tension de l'ensemble des couches de la corde.

5.1.3.1.2.3 Dits constricteurs (adducteurs)

. Le ***crico-aryténoïdien latéral*** prend son origine au niveau du bord supérieur de la face latérale de l'anneau cricoïdien et s'insère sur la face latérale de l'apophyse musculaire de l'aryténoïde.

Il est responsable d'une rotation en dehors de l'apophyse musculaire et donc d'une rotation en dedans de l'extrémité antérieure de l'apophyse vocale. C'est un muscle adducteur qui abaisse, allonge et affine la corde vocale. Toutes les couches sont ainsi raidies et le bord libre de la corde prend une forme plus triangulaire.

. L'**inter-aryténoïdien** est constitué de fibres transverses et de fibres obliques. Les fibres transverses s'orientent de la marge latérale du cartilage aryténoïde jusqu'à la marge latérale de l'autre aryténoïde. Les fibres obliques vont de l'apex de l'aryténoïde jusqu'à la portion latérale de la base de l'autre aryténoïde.

Il est plus spécifiquement responsable de l'adduction de la partie cartilagineuse des cordes vocales. Il est donc particulièrement important pour la fermeture de la partie postérieure de la glotte.

. Le muscle vocal ou **thyro-aryténoïdien** prend son origine au niveau de la partie basse de la face interne du cartilage thyroïde. Il s'insère de l'autre côté sur la face latérale du cartilage aryténoïde depuis le processus vocal jusqu'au processus musculaire. Il forme la masse majeure des plis vocaux (cordes vocales).

En se contractant, il raccourcit le corps de la corde vocale en l'épaississant. Au total, il contribue à augmenter la raideur de la corde et à augmenter l'adduction, en particulier de la portion membraneuse de la corde. Il semble que la plupart des fibres aient une orientation globalement parallèle au bord libre.

Il est divisé en deux chefs musculaires. Le plus interne ou compartiment médial appelé « vocalis » serait plus riche en fibre musculaires « lentes ». Le compartiment externe est appelé thyro-aryténoïdien latéral et serait plus riche en fibres « rapides » et de ce fait serait plus impliqué dans le phénomène d'adduction tandis que le vocalis serait plus impliqué dans le réglage de la tension des fibres pendant la phonation.

Enfin, le larynx est tapissé d'une tunique muqueuse (face interne). A sa partie moyenne, il présente deux replis superposés, dirigés d'avant en arrière :

. Les bandes ventriculaires, bourrelets latéraux en dehors et au-dessus des cordes vocales, constituées du thyro-aryténoïdien supérieur. Chez l'homme, même si leur rôle phonatoire est assez restreint, elles interviennent notamment dans la résistance du tractus vocal supraglottique de même qu'en tant que facteur de résonance.

. Les cordes vocales sont constituées du thyro-aryténoïdien, d'une sous-muqueuse riche en fibres élastiques (ligament cordal) et d'une muqueuse composée de l'épithélium et du chorion. Entre le chorion et le ligament cordal se trouve l'espace de Reinke. Grâce à lui, la vibration de la corde vocale par ondulation du bord libre et de la muqueuse de la face supérieure peut se produire.

Entre ces deux formations se trouve le ventricule de Morgani.

5.1.3.2 Commande nerveuse

5.1.3.2.1 Du cerveau au larynx

A l'origine de la commande musculaire du larynx se trouve le neurone. Cette cellule nerveuse, située dans la matière grise cérébrale, est composée d'un corps cellulaire et de deux types de prolongement : l'axone qui conduit « l'influx nerveux » de façon centrifuge, et le ou les dendrites. Les axones (ou fibres nerveuses) sont rassemblées en faisceaux, eux-mêmes reliés par du tissu conjonctif et formant notamment les nerfs. La zone de contact fonctionnelle entre un neurone et une autre cellule telle qu'une cellule musculaire est nommée synapse. C'est elle qui permet la transmission d'informations entre le neurone et la cellule musculaire, via l'utilisation de neurotransmetteurs. Ces substances biochimiques libérées par les neurones et agissant sur les cellules musculaires permettent notamment l'activité musculaire. L'un d'eux, le transmetteur cholinergique, est responsable de l'action musculaire au niveau des cordes vocales. C'est lui qui permet la contraction.

Ainsi, le message de contraction musculaire des cordes vocales débute au niveau du corps cellulaire du neurone, emprunte l'axone (de direction centrifuge) qui se termine par une synapse en relation avec le muscle. C'est donc une information motrice. Un autre axone (dit sensitif, de direction centripète) va de l'extérieur vers l'intérieur du cerveau pour l'informer de ce qui est réalisé. C'est donc une information sensorielle. Le cerveau présente donc des fibres efférentes, centrifuges (voies motrices) pour exporter l'information et afférentes, centripètes (voies sensitives) pour l'importer. Au niveau de la synapse et du muscle, la sécrétion du transmetteur cholinergique permettra la contraction musculaire.

Il faut noter que la longueur, la modulation et la durée de contraction des cordes vocales varient en fonction de l'ampleur de l'impulsion nerveuse provenant de la commande cérébrale. En fonction de l'importance des impulsions électriques, il y aura une certaine quantité de substance chimique libérée au niveau des synapses qui permettra la contraction des cordes vocales. Cette *« admirable gradation de la réponse possible par ces microparticules chimiques permet un dosage adapté à la note chantée que nous aimerions émettre lors de la contraction de la corde vocale au huitième de ton près, du vibrato au pianissimo souhaité. »*⁷⁸

. **Voies motrices primaires** (voie pyramidale)

Elles débutent au niveau de l'aire motrice primaire (Penfield) située au pied de la circonvolution de la frontale ascendante et de l'aire motrice secondaire, à la face interne de l'hémisphère cérébral et débordant un peu à sa face externe. Les cellules pyramidales du cortex (neurones d'où débute l'information) développent leurs fibres dans le faisceau pyramidal (nommé faisceau géniculé à l'étage encéphalique). Il prend naissance au niveau du centre ovale, passe par la capsule interne, descend dans le pédoncule cérébral, fait relais dans le noyau ambigu (nerfs IX, X, XI) et dispose de connexions avec les noyaux du XII et du VII. Il est influencé par le système extra-pyramidal et cérébelleux au niveau des noyaux bulbaires.

*« Les centres corticaux interviennent dans le contrôle volontaire de la voix. Leur stimulation provoque l'émission de cris. »*⁷⁹

. **Voies motrices secondaires** (voie extra-pyramidale)

Les noyaux gris de la base entrent en action et empruntent les voies cortico-strio-pallido-rubro-bulbaires. Leur stimulation provoque l'émission en rythme de sons inarticulés et monotones. Elles n'auraient qu'un rôle essentiellement facilitateur.

⁷⁸ J.ABITBOL, op. cit. , p.106.

⁷⁹ J.A. RONDAL, X. SERON, op. cit. , p.98.

Le cervelet participe quant à lui par les voies cortico-ponto-cérébello-rubro-bulbaires.

En coordonnant les différents éléments en action pendant la phonation et en ajustant la tonicité musculaire, ces systèmes interviennent dans ce que le langage a d'automatique.

. Voies sensitives

On retrouve des récepteurs sensitifs dans les structures intervenant lors de la phonation. Les afflux afférents permanents qu'ils émettent sont transmis vers les centres de la phonation via différents nerfs tels que les nerfs trijumeaux, vagues...

« De plus, il existe des relations importantes entre l'oreille et le larynx, corticales mais aussi bulbo-protubérantielles, à la base d'un réflexe cochléo-récurrentiel passant du noyau cochléaire au noyau ambigu. »⁸⁰

Mais soulignons que *« la mise en jeu des diverses parties de l'appareil vocal se fait par l'intermédiaire du système nerveux qui fonctionne comme une boucle cybernétique complexe. »⁸¹*

En effet, comme le souligne M.F CASTAREDE, *« en dehors de l'appareil moteur et de l'appareil sensitif, qui correspondent à des aires bien définies dans le cortex, il y a un autre niveau de l'encéphale, directement impliqué dans la voix : le diencephale. Les états affectifs et les émotions jouent un grand rôle dans la phonation pour moduler tonalité, timbre et intensité de la voix. Le diencephale sert de transmetteur en direction :*

- *du cortex : l'émotion est ressentie et devient une expérience émotionnelle ;*
- *du corps : l'expression corporelle de l'émotion se traduit notamment par des modifications de la respiration et de la voix. »⁸²*

⁸⁰ *ibid.* , p. 99.

⁸¹ G. CORNUT, *La Voix*, 1983, p.38.

⁸² M.F. CASTAREDE, *op. cit.* , p.26.

5.1.3.2.2 Contrôle réflexe

Des ajustements peuvent survenir au niveau préphonatoire et lors de l'émission vocale. L'ajustement préphonatoire ne dépend pas du contrôle audiophonatoire (qui sera traité dans une partie ultérieure). En voix chantée, il est possible de produire un son qui soit d'emblée à la hauteur et à l'intensité voulues. Ce réglage préphonatoire est d'origine corticale. Il dépend notamment des informations fournies par les récepteurs laryngés qui donnent aux centres des informations sur l'état de tension et la position des différents muscles et des différentes articulations. Lors de la phonation, ces informations permettent les ajustements instantanés nécessités par le maintien d'une configuration glottique donnée. D'autres arcs réflexes d'origine abdominale, thoracique, cervicale, linguale,... contribuent probablement aux ajustements permanents agissant par feed-back sur le larynx au cours de la phonation.

5.1.3.2.3 Innervation des muscles laryngés

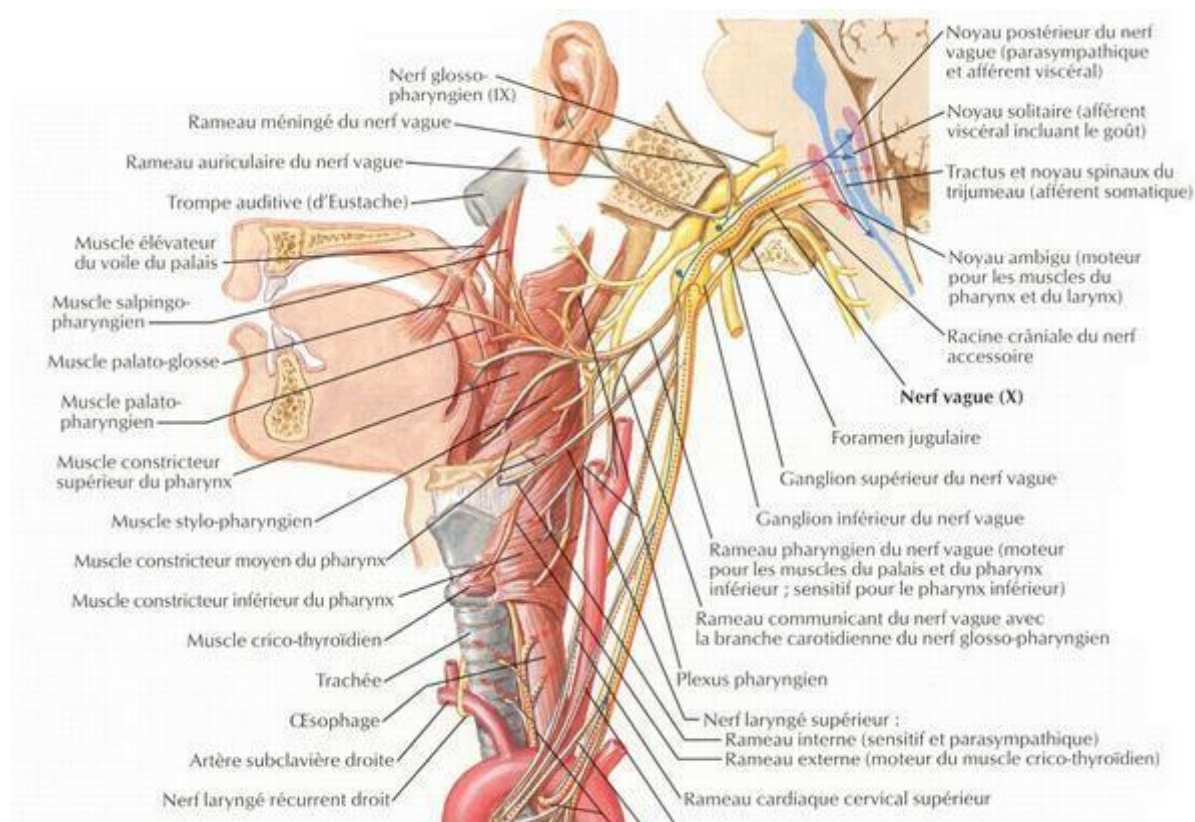


Figure 3 : innervation des différents muscles laryngés, notamment le nerf vague (X).⁸³

⁸³ D.H. MCFARLAND, H. NETTER, op. cit. , p.213.

. Les muscles extrinsèques du larynx :

Les muscles sous-hyoïdiens sont sous la dépendance du nerf hypoglosse constituant la XII^e paire des nerfs crâniens. Ce nerf exclusivement moteur contient sur certains segments de son trajet des fibres d'emprunt du plexus cervical profond qui constituent entre autres des rameaux moteurs pour les muscles sous-hyoïdiens. Au niveau du cou, le nerf traverse la région carotidienne en passant entre l'artère carotide interne et la veine jugulaire interne. Sa courbe accompagne le bord inférieur du ventre postérieur du muscle digastrique. Il pénètre ensuite dans la région sus-hyoïdienne latérale au-dessus de la grande corne de l'os hyoïde. Il passe au-dessus du mylo-hyoïdien et pénètre dans la région sub-linguale, où il donne toutes ses branches motrices. Le muscle génio-hyoïdien au niveau sus-hyoïdien dépend lui aussi du nerf hypoglosse via la branche génio-hyoïdienne.

Les autres muscles sus-hyoïdiens quant à eux dépendent du nerf facial (VII^e paire crânienne) pour le ventre postérieur du digastrique et le stylo-hyoïdien, et du nerf trijumeau (V^e paire crânienne) pour le ventre antérieur du digastrique et le mylo-hyoïdien.

Le nerf facial naît entre la protubérance annulaire et le bulbe rachidien, passe par le conduit auditif interne, emprunte le canal du facial et sort au niveau du foramen stylo-mastoïdien. Des rameaux moteurs sont alors détachés pour les muscles laryngés qu'il innerve.

Le nerf trijumeau apparaît à la face antéro-latérale de la protubérance annulaire du tronc cérébral en deux faisceaux : l'un sensitif, l'autre moteur (plus petit). Dès sa sortie du tronc cérébral, il passe au-dessus du rocher de l'os temporal et se divise en trois branches sensitives principales et une branche motrice. La branche mandibulaire assure l'innervation des muscles laryngés concernés.

. Les muscles intrinsèques :

Leur innervation est assurée par le nerf vague (ou pneumo-gastrique), constituant la X^e paire des nerfs crâniens. Il est le seul nerf commandant à la fois l'ouverture, la fermeture et la sensibilité des cordes vocales. Deux de ses branches interviennent plus particulièrement : le nerf laryngé supérieur et le nerf laryngé inférieur (ou nerf récurrent). Les fibres motrices sont

issues du noyau ambigu situé dans la partie caudale du plancher du IV^e ventricule et empruntent le nerf vague. Les fibres sensibles se terminent dans la partie bulbaire du noyau du tractus solitaire.

Le nerf laryngé supérieur naît de l'extrémité inférieure du ganglion supérieur du nerf vague, sous la base du crâne. Il suit un trajet vertical, passant en arrière puis en dedans de l'artère carotide interne, puis au sein de l'artère carotide externe. Au contact de la membrane thyro-hyoïdienne, il se divise en deux branches terminales. La branche terminale interne pénètre dans le larynx avec le pédicule artériel laryngé supérieur. Il s'agit d'une branche sensible pour la muqueuse du vestibule laryngé. La branche terminale externe longe la ligne oblique de la face antérieure du cartilage thyroïde, innerve le muscle crico-thyroïdien puis traverse le cône élastique pour donner l'innervation sensible du ventricule laryngé à l'étage infra-glottique. Ainsi, le nerf laryngé supérieur est sensitif pour l'ensemble de la muqueuse laryngée et moteur pour le muscle crico-thyroïdien.

Le nerf laryngé inférieur naît, pour le côté gauche, du nerf vague en avant de l'arc de l'aorte, est thoracique pour sa partie basse et cervical pour sa partie haute ; pour le côté droit, du nerf vague droit en avant de l'artère subclavière et est exclusivement cervical. Quel que soit le côté concerné, ce nerf est en rapport étroit avec l'artère thyroïdienne inférieure. A l'extrémité supérieure de la trachée, il passe sous le bord inférieur du muscle constricteur du pharynx puis se distribue au larynx pour innerver tous les muscles intrinsèques à l'exception du muscle crico-thyroïdien.

5.1.3.3 Vibration laryngée

5.1.3.3.1 Théorie de la vibration laryngée

L'air provenant des poumons passe par la trachée, traverse le larynx, les cavités pharyngées et sort par la cavité buccale. Un son est alors produit. Différentes théories ont apporté un éclairage de plus en plus précis sur le fonctionnement des cordes vocales pendant la phonation, notamment la théorie oscillo-impédantielle de P.DEJONCKERE en 1981. La corde vocale vibrante est associée à « *un oscillateur harmonique très amorti entretenu, dont la*

fréquence est déterminée par l'inertie et l'élasticité de la portion vibrante et dont l'apport périodique d'énergie est constitué par l'onde de pression sous-glottique, celle-ci présentant une fréquence identique à celle du fondamental laryngé émis. » Cela constitue le phénomène dit de Bernouilli.⁸⁴

Selon la loi de Bernouilli, *« lorsqu'un courant d'air circule rapidement dans une zone rétrécie, il y crée une pression négative après son passage qui tend à aspirer les berges, ici, la muqueuse des cordes vocales. »*⁸⁵

En pratique, il y a adduction des cordes vocales grâce à la rotation des aryténoïdes constituant un obstacle à l'air expiratoire. L'aspect en fuseau de la glotte permet l'apparition de turbulences (avec à ce niveau une pression atmosphérique approximativement normale) ainsi que l'apparition de l'effet de Bernouilli caractérisé par une rétro-aspiration sur les cordes vocales ; celle-ci est due à une pression négative dans les cavités supra-laryngées après le passage rapide de l'air expiratoire. La première oscillation est quant à elle engendrée par un déséquilibre entre pression sous-glottique et tension de fermeture des cordes vocales. Les cordes vocales s'écartent sous l'effet de la pression d'air, laissent s'échapper l'air, puis grâce à la diminution de la pression, à l'effet de Bernouilli et à leurs propriétés de vibration, les bords libres des cordes vocales se rejoignent. Le phénomène se reproduit alors tant que l'air pulmonaire est présent pour permettre la pression-sous glottique.

P.DEJONCKERE met en évidence qu'un mouvement vibratoire normal s'effectue dans les trois dimensions de l'espace et notamment horizontalement, c'est-à-dire par l'amplitude d'écartement de surface des cordes vocales. Par ailleurs, c'est la pression sous-glottique qui détermine l'intensité sonore et c'est le débit qui entre en jeu pour les sons de fréquence aiguë.

Aussi, une augmentation de la fréquence et donc de la hauteur tonale sera obtenue par :

- un amincissement du bord libre qui dépend de l'activité tonique de la musculature intrinsèque du larynx, notamment des cordes vocales.
- une diminution de l'amplitude vibratoire de la dimension horizontale.
- une latéralisation légère des axes d'oscillation du bord libre des cordes vocales.

⁸⁴ J.A. RONDAL, X. SERON, op. cit. , p.102

⁸⁵ J. HEUILLET-MARTIN, H. GARSON-BAVARD, A. LEGRE, *Une Voix Pour Tous, Tome I*, 1995, p.32.

5.1.3.3.2 Nature du son laryngé produit

Lorsque le larynx entre en vibration, il produit un son complexe.

Ce que l'on nomme « son » est « *une vibration acoustique capable d'éveiller une sensation acoustique*⁸⁶ ». Cette vibration est dite périodique, c'est-à-dire qu'elle présente une certaine régularité. Nous désignons comme période du son l'intervalle de temps au bout duquel l'onde se répète à l'identique. De façon opposée, une production apériodique correspond à un bruit.

On distingue deux types de sons : les sons purs et les sons complexes. La production d'un son pur, tel que peut le faire un diapason par exemple, entraîne une oscillation des molécules d'air qui suivent alors une sinusoïde pure. En revanche, lors de la production d'un son complexe, l'onde est toujours périodique mais elle n'est pas purement sinusoïdale. En réalité, selon la loi de FOURIER, un son complexe équivaut à une somme de sons purs, dont chacun peut être représenté par une sinusoïde. La fréquence de chacun de ces sons purs est un multiple entier de la fréquence du son pur (nommé son fondamental), dont la fréquence est la plus basse (nommée fréquence fondamentale). Ces sons purs dont la fréquence est un multiple entier de la fréquence fondamentale sont les harmoniques.

5.1.3.3.3 Différents mécanismes laryngés

Au niveau du larynx, il existe plusieurs modes vibratoires nommés « mécanismes laryngés ». Effectivement, l'ensemble des notes pouvant être produites par un individu, ne peut être réalisé au travers d'un seul mécanisme. Il en existe quatre :

- Le mécanisme 0 ou pulsé permet de réaliser les fréquences inférieures à 80-100 Hz ; il se manifeste grâce à une vibration laryngée sur toute la longueur de la fente glottique. Il y a contraction des muscles interaryténoïdiens, rapprochement de la partie postérieure des aryténoïdes et détente des cordes vocales.

⁸⁶ P. DEJONCKERE, Précis de Pathologie et de Thérapeutique de la Voix, 1980, p. 17.

- Le mécanisme I ou de poitrine permet la production des fréquences comprises entre 80-100 Hz et 300-400Hz ; il s'obtient par la contraction du muscle vocal et la détente simultanée du muscle crico-thyroïdien. Les cordes vocales sont relâchées, raccourcies et épaissies.

- Le mécanisme II ou de tête ou de fausset rend possible la réalisation des fréquences allant de 300-400 Hz jusqu'à 600-700 Hz ; il s'obtient à l'inverse par une contraction du muscle crico-thyroïdien, avec un raccourcissement de la partie vibrante par augmentation de la force de rapprochement des cartilages aryténoïdes et une détente simultanée du muscle vocal. Les cordes vocales sont tendues, allongées et plus minces, avec une amplitude et une ondulation de la muqueuse limitées.

- Enfin, le mécanisme III ou de sifflet permet de réaliser les fréquences supérieures à 1000 Hz ; il est dû à l'absence de vibration des cordes vocales qui ne s'accrochent plus. La glotte est fermée par contraction des muscles crico-aryténoïdiens latéraux mais présente une ouverture en arrière par relâchement des interaryténoïdiens.

Bien sûr, ces éléments peuvent varier d'un individu à l'autre. De plus, une note peut être à la frontière entre deux registres, elle sera donc produite par l'un ou l'autre. Les mécanismes I et II sont davantage employés pour la voix parlée et chantée : le premier est usité par les hommes en voix parlée et chantée, par les femmes en voix parlée ; le second par les femmes en voix parlée et chantée et par les contre-ténors.⁸⁷

Le terme de mécanisme est préférentiellement utilisé pour décrire les modes de fonctionnements laryngés par rapport à celui, plus flou, de registre qui correspond davantage aux sensations ressenties. Ces mécanismes sont également à distinguer de la tessiture, ensemble des notes dans lequel une voix s'exprime le mieux.

⁸⁷ X.RENARD, Précis d'Audioprothèse : Production, Phonétique Acoustique et Perception de la Parole, 2008, pp. 21-22.

5.1.3.3.4 Attaques vocaliques

L'attaque du son revêt une grande importance. La qualité de la voix sera d'autant plus grande que l'attaque sera précise. On retrouve deux types d'attaques vocaliques⁸⁸ :

- l'attaque dure qui se caractérise par un accolement des cordes vocales alors que les plis vestibulaires s'écartent. La vibration commence glotte fermée avec une pression sous-glottique augmentant progressivement pour forcer le passage. S'il y a exagération, c'est-à-dire présence d'un coup de glotte, cela provoquera un manque de précision tonale, de netteté et la sonorité présentera un caractère explosif. Une attaque aspirée peut aussi se retrouver.

- l'attaque douce où les cordes vocales sont en position d'adduction mais laissant une fente fusiforme entre elles. La vibration commence glotte ouverte. Grâce à l'effet de Bernouilli, l'air s'écoulant à travers permet d'augmenter l'amplitude de vibration jusqu'à provoquer l'accolement. S'il y a exagération, c'est-à-dire attaque soufflée, il y aura un bruit aérien parasite et audible.

5.1.3.3.5 Contrôle des paramètres de la voix

- **Contrôle de la hauteur** : la hauteur tonale est la fréquence du son fondamental laryngé. Elle s'exprime en hertz, ou pour les chanteurs, par la note correspondant à cette fréquence. Elle change en permanence lorsque nous parlons ou chantons.

Son contrôle s'effectue grâce à la pression sous-glottique et à l'action des muscles intrinsèques et extrinsèques du larynx⁸⁹, comme mentionné précédemment.⁹⁰

Une pression inspiratoire plus ou moins intense et la résistance des cordes vocales au passage de l'air peuvent faire varier la pression sous-glottique. Ainsi, lorsque celle-ci n'est pas suffisante et que la production vocale est longue, la voix aura tendance à diminuer s'il n'y a pas une reprise inspiratoire.

⁸⁸ J.A. RONDAL, X. SERON, op. cit. , p.102-103.

⁸⁹ C. PAINTER, Laryngeal Functions in Speech, in Otolaryngology, Head and Neck, p.1752.

⁹⁰ Voir « 1.1.3.1 Musculature », p.8.

Rappelons que l'impact musculaire sur les cordes vocales existe au niveau antéropostérieur, transversal et vertical : l'équilibre de tensions entre les muscles vocaux d'une part et les muscles crico-thyroïdiens et crico-aryténoïdiens d'autre part entraîne une variation de longueur des cordes vocales (sens antéropostérieur). Leur allongement et leur raccourcissement provoquent respectivement une élévation et un abaissement de la hauteur tonale. La force d'accolement des cordes vocales (sens transversal) dépend d'une contraction plus ou moins importante des interaryténoïdiens et crico-aryténoïdiens latéraux. Ces muscles augmentent la force de résistance des cordes vocales au passage de l'air. Cela intervient dans la production des sons aigus. Enfin, les muscles suspenseurs permettent de déplacer le larynx (sens vertical). Son élévation augmente la raideur des cordes vocales (production des sons aigus), son abaissement épaissit les cordes vocales qui vibrent plus lentement (production des sons graves).

- **Contrôle de l'intensité** : l'intensité correspond à la puissance de la voix. Elle s'exprime en décibels. L'intensité de la parole doit être appropriée aux dimensions de l'espace dans lequel nous produisons le son, au nombre d'auditeurs auxquels nous nous adressons, à la situation, ... Si ce n'est pas le cas, il y aura alors risque de forçage vocal.

L'intensité est liée à la hauteur. En effet, elles dépendent toutes les deux de la pression sous-glottique et de la tension des cordes vocales. Quand nous parlons plus fort, intensité et hauteur augmentent. Cependant, la grande variété des modifications possibles de la hauteur tonale permet de contrôler différemment ces deux paramètres. L'intensité est largement modulée grâce aux muscles expiratoires. La respiration costo-abdominale en améliore grandement le contrôle.

- **Contrôle du timbre** : le timbre est une notion complexe qui ne peut être quantifiée complètement. Il dépend de la distribution et de l'intensité relative des différents harmoniques du son complexe qu'est l'émission vocale. La distribution des harmoniques quant à elle est tributaire de la vibration laryngée et de sa transformation dans les résonateurs. On observe deux types de timbre : le timbre vocalique qui inclut globalement les harmoniques de 0 à 2500 Hz, contient les deux premiers formants et permet la différenciation vocalique ; le timbre extra-vocalique qui contient les harmoniques supérieurs à 2500 Hz et caractérise la voix de chacun.

Il dépend de la qualité du son laryngé et de son amplification par les résonateurs. Il est donc lié à l'épaisseur, la tension et le bon accolement des cordes vocales. Plus largement, l'articulation peut aussi impacter sur le timbre de la voix (cas de la voix nasonnée par exemple).

5.1.4 RESONATEURS

A lui seul, le son laryngé n'a pas une grande signification. Il doit sa coloration finale aux résonateurs, cavités agissant sur le son qui les traverse. Il s'agit du pharynx, du nasopharynx, des fosses nasales et de la cavité buccale. Leur volume, leur longueur et leur surface ainsi que les éléments mobiles voisins (mandibule, langue, voile du palais et lèvres) contribuent à moduler le son laryngé. Celui-ci se charge des fréquences propres des résonateurs traversés ; il se complique et s'enrichit alors.

5.1.4.1 Cavités de résonance

- Le pharynx : conduit musculo-membraneux vertical qui s'étend du rhinopharynx jusqu'à la bouche oesophagienne, il est composé de trois étages. Au niveau inférieur, c'est-à-dire autour et en arrière du larynx, nous retrouvons l'hypopharynx, mais il ne joue aucun rôle dans la phonation. Au-dessus se situe l'oropharynx, qui correspond au carrefour aérodigestif ; il s'agit du résonateur primaire pharyngé. Enfin, le rhinopharynx, ou cavum, joue un rôle important dans la respiration et dans la phonation où il agit en tant que résonateur secondaire nasal. Le pharynx comprend deux types de muscles : les muscles élévateurs qui agissent sur sa longueur, et les muscles constricteurs qui agissent sur son diamètre.

- Le nasopharynx et les fosses nasales : cet ensemble se divise en trois étages : le méat supérieur, zone de l'odorat, le méat moyen qui contrôle les débits respiratoires et conditionne l'air inspiré, et le méat inférieur qui oriente le flux d'air. Les fosses nasales interviennent dans la respiration. Une muqueuse ciliée sécrétoire tapisse l'intérieur des narines et permet de filtrer et de réchauffer l'air inspiré.

- La cavité buccale : elle est délimitée par l'orifice buccal en avant, les joues sur les côtés, la voûte palatine en haut, l'isthme du gosier en arrière, l'arc mandibulaire et le plancher buccal en dessous. Son rôle de résonateur est surtout défini par les éléments mobiles qui la composent : la mandibule, les lèvres et la langue.

5.1.4.2 Eléments mobiles

- La mandibule : elle est reliée au crâne par l'articulation temporo-mandibulaire. Ce sont les muscles masticateurs qui régissent ses mouvements d'élévation et d'abaissement, de protrusion et de rétraction. Elle définit la longueur et le diamètre du conduit vocal lors d'une émission vocale. En effet, comme elle est en relation avec les mouvements des lèvres, de la langue et du crâne, son ouverture provoque un agrandissement de la cavité buccale par abaissement du plancher buccal et un abaissement laryngé qui entraîne un élargissement du pharynx. Le timbre vocal sera tendu et l'intensité vocale faible si la mandibule perd de sa mobilité.

- La langue : elle se compose de deux parties : la partie mobile se situe dans la cavité buccale et la partie fixe ou base de langue dans l'oropharynx. La langue mobile s'étale dans la cavité buccale, épousant l'articulé dentaire. Dix-sept muscles la composent. Ils lui permettent une grande mobilité, limitée seulement par le frein lingual fixé au plancher buccal, et des possibilités de rétraction, nécessaires pour l'articulation de la parole. La base de langue commence en arrière du V lingual et se lie dans sa partie inférieure à l'épiglotte. Sa masse musculaire s'insère sur l'os hyoïde en arrière et sur la face interne des angles mandibulaires en avant. Selon sa posture avancée ou non, elle a une importance sur le calibre du résonateur pharyngé primaire qu'est l'oropharynx.

- Le voile du palais : élément musculo-membraneux mobile entre le cavum et l'oropharynx, il prolonge le palais. Pendant la déglutition, sa fonction est d'éviter le reflux nasal. En phonation, il se lève pour diriger l'air vers la cavité buccale et ainsi permettre l'oralisation. Dans le cas d'une nasalisation, il s'abaisse pour entraîner l'air vers les fosses nasales.

- Les lèvres : elles délimitent l'orifice buccal et sont retenues par un frein fibro-muqueux. Lors de la phonation, elles forment des mouvements complexes qui modifient la longueur et le degré d'ouverture de la cavité buccale. En s'écartant des dents, elles offrent une cavité de résonance indispensable à la rondeur du son notamment dans le chant.

5.1.4.3 Commande nerveuse

L'innervation motrice et la majeure partie de l'innervation sensitive du pharynx proviennent du plexus nerveux pharyngien (formé par des éléments du nerf crânien IX, X et du sympathique cervical). Les fibres motrices du plexus sont originaires de la racine crâniale du nerf accessoire (ou nerf spinal) ; elles accompagnent le nerf vague (par l'intermédiaire de sa ou ses branches pharyngiennes) et se distribuent à tous les muscles du pharynx et du voile du palais (sauf pour le stylo-pharyngien, innervé par le nerf crânien IX et pour le tenseur du voile du palais, innervé par le nerf crânien V). Le constricteur supérieur du pharynx bénéficie également de quelques fibres des nerfs laryngé et récurrent, branches du nerf vague. Les fibres sensibles du plexus rejoignent le nerf glosso-pharyngien (IX). Elles recueillent la sensibilité muqueuse des trois parties du pharynx. L'innervation sensitive de la muqueuse des parties antérieure et supérieure du nasopharynx dépend principalement des nerfs maxillaires (nerfs crâniens V 2), exclusivement sensitifs.⁹¹

Le muscle tenseur du voile du palais dépend de la branche mandibulaire (V 3) ; cette branche innerve également la mandibule. Les autres muscles du voile dépendraient du plexus pharyngien dont essentiellement des fibres du nerf vague (X). La muqueuse de la face supérieure du voile dépendrait des nerfs crâniens IX et X. Celle de la face postérieure dépendrait comme la muqueuse postérieure des fosses nasales et de la voûte du palais, des branches du nerf maxillaire (V 2).

Concernant la langue, l'innervation motrice est sous la coupe du nerf hypoglosse (XII) et l'innervation sensitive est sous la dépendance du nerf trijumeau (V).⁹²

⁹¹ K.L. MOORE, A.F. DALLEY, *Anatomie médicale : Aspects fondamentaux et applications cliniques*, 2006, p.1058.

⁹² S.R. DHILLON, C.A. EAST, *Oto-rhino-laryngologie : Et chirurgie cervico-faciale*, 2008, p.74.

L'innervation des muscles des lèvres dépend de la branche buccale, zygomatique ou mandibulaire du nerf facial (VII).

Il faut noter que la vitesse du débit de parole et la précision des points d'articulation requièrent des adaptations à la fois très rapides et très fines, aussi par le fait d'une possible proximité entre ces points d'articulation. Ainsi, l'espace entre différentes structures telles que la langue et le palais définit le niveau de constriction du conduit vocal. Des variations minimales de distance entre langue et palais permettent donc de produire des sons aussi différents qu'une voyelle, une consonne sourde ou encore une consonne occlusive. C'est grâce à l'action des arcs réflexes rapides reposant sur la sensibilité tactile et proprioceptive de la langue, du palais et des lèvres que ces adaptations peuvent être réalisées.

5.1.4.4 Leur rôle

Comme leur nom l'indique, les résonateurs interviennent dans le phénomène de résonance, c'est-à-dire une modification du son résultant de l'enrichissement ou de l'appauvrissement de certains de ses harmoniques. Chaque cavité possède une fréquence ou un son propre, pouvant être un son fondamental ou un de ses harmoniques, telle qu'elle résonne sous l'influence de ce son et le renforce. En général, une cavité grande aura un son propre plus grave qu'une cavité plus petite. Cependant, il convient de préciser que les cavités ne créent pas de nouveaux harmoniques. Aussi, deux cavités de résonance communiquant l'une avec l'autre peuvent modifier leurs propriétés réciproques. Dans ce « couplage », la plus grave s'aggrave et la plus aiguë devient plus aiguë.⁹³

Les résonateurs peuvent être modulés dans leur forme et leur taille ce qui a pour conséquence de modifier leur fréquence propre sans limite. Les possibilités de couleurs sonores sont donc nombreuses. Si la fréquence des résonateurs s'adapte à celle du son laryngé ou à celle de ses harmoniques, il se produit une amplification. Pour une même intensité, il y aura donc un effort moindre de la part du larynx et de la soufflerie. En revanche, s'il existe une grande différence, il se produira un assourdissement et le son laryngé sera plutôt étouffé.⁹⁴

⁹³ J.A. RONDAL, X. SERON, op. cit., p. 104-105.

⁹⁴ J. VERMETTE, R. CLOUTIER, *La Parole en Public : Savoir Etre, Savoir Faire*, 1992, p. 8.

J.A.RONDAL⁹⁵ explique « *qu'ainsi, les variations de la tension musculaire intrinsèque et extrinsèque, de la pression sous-glottique et de l'impédance de l'onde sonore réfléchie au larynx modifient continuellement le mouvement vibratoire glottique.* »

Pour obtenir une production optimale, une coordination parfaite entre résonateurs et vibrateur est nécessaire. « *Les phénomènes résonantiels s'amplifient lorsqu'il existe un rapport harmonique entre le rythme impulsionnel (c'est-à-dire la fréquence du son laryngé) et la fréquence propre des résonateurs.* »⁹⁶

J.VERMETTE⁹⁷ ajoute que la coordination entre résonateurs et son laryngé prend toute son importance dans la production des voyelles notamment. En effet, leur réalisation nécessite un parfait ajustement entre ces éléments.

Les résonateurs ont donc une fonction double : ils permettent la formation des harmoniques mais aussi la formation de notre langage articulé. Les caisses de résonances accentuent certains harmoniques et changent le timbre vocal. Les éléments malléables telles que la langue, la bouche, la cavité nasale ou encore le pharynx modulent, déforment, sculptent le son devenant voix. Le maxillaire inférieur, la langue et les lèvres sont les acteurs principaux dans la modification de la caisse de résonance.⁹⁸

Claire PILLOT⁹⁹ souligne l'importance des résonateurs au travers d'une étude sur les raisons du singing-formant.¹⁰⁰ Une expérience, où l'utilisation de l'Imagerie par Résonance Magnétique (IRM) a permis de recueillir 21 coupes sagittales médianes d'une basse professionnelle produisant les voyelles [a], [i], [o] en voix parlée puis chantée, a montré « *de la parole au chant : un abaissement laryngé, un orifice ary-épiglottique plus étroit, un abaissement de l'os hyoïde, une postériorisation du point de constriction lingual, une ouverture mandibulaire et labiale, une protrusion labiale, une réduction de la courbure cervicale et des différences angulaires entre le palais dur et le voile. On constate en outre une*

⁹⁵ J.A. RONDAL, X. SERON, op. cit. , p.105.

⁹⁶ G. CORNUT, op. cit. , p.35.

⁹⁷ J. VERMETTE, op. cit.

⁹⁸ J. ABITBOL, op. cit. , pp. 201-207.

⁹⁹ C. PILLOT, *Pourquoi le Singing-formant ?*, 1998, pp. 2-7.

¹⁰⁰ Renforcement d'énergie dans le spectre vocal entre 2800 Hz et 4000 Hz.

augmentation de la longueur du conduit vocal de la parole au chant ainsi qu'une baisse du rapport entre les sections laryngées et pharyngées. Certains auteurs contestent cependant cet abaissement laryngé parmi lesquels Wang¹⁰¹ et Sundberg¹⁰². »

C.FOURNIER recommande l'usage des résonateurs afin d'éviter le forçage vocal. Il explique qu' « *adapter en permanence l'ouverture du résonateur permet un gain de puissance qui nécessite beaucoup moins d'effort que le forçage de l'émission.* »¹⁰³

5.1.4.5 Cas des chants diphonique et harmonique

Pour illustrer l'importance des résonateurs dans la production sonore, nous présentons la technique du chant diphonique. Celle-ci consiste en la production de plusieurs notes simultanément. Avec un même organe vocal, le sujet mélange plusieurs types de voix et différents positionnement de langue ou de lèvres. Il y a donc une utilisation et un positionnement particulier des cavités bucco-nasales. Le chant harmonique (ou chant de gorge) permet aussi de produire plusieurs sons à la fois.

Cette façon de procéder se retrouve dans de nombreuses musiques traditionnelles du monde, notamment en Haute-Asie (Mongols, Tibétains, ...) mais aussi en Inde, Italie, etc. Différentes techniques de chant diphonique existent comme le *kargyraa* qui se traduit par la mise en vibration de certains tissus se situant au-dessus des cordes vocales, produisant une note grave, une octave en dessous de la note chantée. La technique de gorge porte le nom de *sygyt*. La voix se caractérise alors par la production simultanée de deux sons : l'un dit « son fondamental » ou bourdon qui est maintenu à la même hauteur le temps de l'expiration pendant que l'autre plus aigu dit « son harmonique » varie au gré du chanteur. Le bourdon existe car son intensité dépasse celle des harmoniques, il n'est pas filtré par les résonateurs et bénéficie d'un rayonnement buccal et nasal. Si la cavité nasale venait à être fermée, l'intensité du bourdon diminuerait cela par l'absence de source de rayonnement, par une diminution du débit d'air et de l'intensité sonore. Le chant diphonique exige des formants aigus que seul le couplage entre plusieurs cavités permet d'obtenir. Ainsi, la production de deux sortes de chant

¹⁰¹ S. WANG, *Singing voice : bright timbre, singer's formants and larynx positions*, 1983, p.313-322.

¹⁰² J. SUNDBERG, *The Science of the Singing Voice*, 1987, p.92.

¹⁰³ C.FOURNIER, *La Voix, un Art et un Métier*, 1989, pp.119-120

diphonique a été expérimentée : l'une avec la bouche comme unique cavité (sans intervention de la langue qui reste en position basse) ; l'autre, avec utilisation de la pointe de langue en contact avec la voûte palatine, divisant la bouche en deux cavités. Dans le premier cas, on entend convenablement le bourdon mais les formants aigus se dispersent sur plusieurs harmoniques, la deuxième voix est difficile à percevoir et l'effet diphonique se perd. Dans le second cas, le formant aigu et intense réapparaît.

Un réseau de résonateurs sélectifs filtrant uniquement les harmoniques souhaités par le chanteur est donc nécessaire pour ce type de production vocale. Pour une résonance unique très aiguë, il faut un couplage serré entre les deux cavités. En relâchant ce couplage, l'intensité du formant décroît, l'énergie sonore est répandue dans le spectre. Lorsqu'il n'y a qu'une seule cavité, le chant diphonique devient alors beaucoup plus diffus comme souligné précédemment.

Il existe trois façons différentes pour produire deux sons simultanément :

- avec une seule cavité buccale : soit la langue est en position de repos, soit la base de la langue est légèrement remontée (elle ne doit pas entrer en contact avec la partie molle du palais). Il y a uniquement un mouvement de la bouche et des lèvres. Par cette modulation de la cavité buccale en prononçant les deux voyelles [ü] (qui n'existe pas en français) et [i] liées sans interruption (comme si l'on disait "oui" en français), une légère mélodie des harmoniques est perceptible, elle n'excède pas l'harmonique 8 en général.

- avec deux cavités buccales : la voix de gorge est utilisée pour chanter. Une fois le son [l] prononcé, il s'agit de maintenir en position la pointe de la langue au milieu de la voûte palatine puis de prononcer la voyelle [ü] en conservant toujours la pointe de la langue collée fermement à la jonction entre palais dur et palais mou. Les muscles du cou et la musculature abdominale doivent intervenir pendant le chant comme s'il fallait soulever un objet lourd. Le timbre est quant à lui nasalisé et amplifié grâce aux fosses nasales. Les deux voyelles [i] et [ü] (ou [o] et [a]) liées sont émises l'une après l'autre en plusieurs fois. Cela permet au chanteur de bénéficier du bourdon et des harmoniques sur une pente ascendante ou descendante. Afin de modifier la mélodie des harmoniques, le chanteur doit intervenir sur la position de ses lèvres et de sa langue. Une concentration musculaire importante accroît la clarté harmonique.

- la dernière méthode nécessite de placer la base de langue de façon à ce qu'elle soit « mordue » par les molaires supérieures, tout en produisant le son de gorge sur les deux

voyelles [i] et [ü] liées et reproduites à maintes reprises pour entraîner une suite d'harmoniques descendants et ascendants. Cependant, cette méthode n'offre pas un contrôle de la mélodie formantique et ne constitue qu'une démonstration expérimentale sur les possibilités de timbre harmonique.

Tout cela rend compte de l'action et du contrôle que peut avoir le chanteur et plus largement tout individu sur ses résonateurs.

5.1 PERCEPTION DES SONS

5.2.1 RECEPTEUR DE LA VOIX : L'OREILLE

5.2.1.1 Anatomie de l'oreille

L'oreille humaine est constituée de trois parties¹⁰⁴ :

- l'oreille externe comprend le pavillon et le conduit auditif externe fermé à l'intérieur par le tympan. Le pavillon capte les ondes et les oriente vers le conduit auditif externe. Ce dernier les amplifie un peu et les dirige vers le tympan. Sous leur effet, la membrane du tympan entre en vibration.

- l'oreille moyenne comprend la caisse du tympan, traversée par la chaîne des osselets (le marteau, l'enclume et l'étrier). Les ondes sonores provenant du tympan arrivent aux osselets, sont amplifiées à nouveau, puis progressent vers la fenêtre ovale, point de jonction entre l'étrier et la cochlée, organe de l'audition. L'oreille moyenne convertit donc la vibration de l'air en énergie mécanique.

- l'oreille interne est à la fois l'organe de l'équilibre grâce au vestibule mais aussi de l'audition grâce à la cochlée. Cet organe creux, rempli de liquide, en forme de spirale hélicoïdale est divisé en trois rampes : vestibulaire, cochléaire et tympanique. Entre rampe

¹⁰⁴ B. KOLB, I.Q. WHISHAW, Cerveau et Comportement, 2002, pp.326-336.

tympanique et cochléaire se trouve la membrane basilaire. A sa surface, il y a l'organe de Corti qui comprend les cellules ciliées, récepteurs sensibles aux vibrations sonores. Il en existe de deux sortes : les cellules ciliées dites « internes » sont les récepteurs auditifs et transmettent l'influx nerveux aux fibres du nerf auditif ; les cellules ciliées dites « externes » sont au contact de la membrane tectoriale et permettent d'augmenter la puissance du pouvoir de résolution de la cochlée. Ainsi, les vibrations sonores provenant de l'étrier provoquent un mouvement du liquide cochléaire puis un déplacement des membranes qui inclinent les cellules ciliées dans un sens ou dans l'autre, transformant l'onde sonore en influx nerveux.

Cependant, il faut souligner que la sensibilité aux fréquences sonores n'est pas la même en tout point de la cochlée. Il existe une localisation des fréquences dans la cochlée. La résonance de la membrane basilaire se fait à haute fréquence près de sa base (sons aigus) et à basse fréquence près de son extrémité (sons graves). Cette caractéristique due à son épaisseur non homogène offre une détection fine des différences de fréquence.

5.2.1.2 Voies auditives

Il existe trois circuits¹⁰⁵ distincts au niveau des voies auditives. Tout d'abord, la voie sensorielle à proprement parler qui est formée de trois neurones :

- le premier va de l'oreille interne jusqu'au tronc cérébral en empruntant le trajet du nerf auditif (VIII) pour se terminer dans les noyaux cochléaires dorsal et ventral. Il est articulé avec l'organe de Corti dans la cochlée membraneuse. C'est dans le ganglion spiral de Corti que se situe son corps cellulaire.

- le deuxième a son corps cellulaire présent dans les noyaux cochléaires du tronc cérébral. De ces noyaux émergent des fibres nerveuses formant le corps trapézoïde au niveau de la partie centrale de la protubérance. Puis cet ensemble fibreux monte dans le tronc cérébral et se termine dans le corps genouillé interne.

- le troisième, thalamo-cortical, a son corps cellulaire dans le corps genouillé interne. Il se termine au niveau du cortex de la première circonvolution temporale.

¹⁰⁵ A. BOUCHET, J. CUIILLERET, *Anatomie topographique, descriptive et fonctionnelle*, 1991, p. 302.

C'est au niveau de la première circonvolution temporale, sur l'aire 41 que se situe le centre de l'audition. La réception des sons aigus repose sur partie la plus interne de cette circonvolution alors que la perception des sons graves dépend de la partie plus externe.

Ensuite, il existe un circuit réflexe, qui intervient dans le cadre d'une réaction réflexe aux bruits. Il est branché en dérivation sur le circuit des voies sensorielles précédemment évoquées. C'est à partir du corps trapézoïde que certaines fibres vont se projeter entre autres sur les noyaux moteurs du tronc cérébral.

Enfin les voies commissurales qui permettent aux voies auditives gauches et droites d'entrer en relation. Ces voies empruntent la commissure de GUDDEN et le corps calleux.

5.2.1.3 Développement et audition

Des expériences de stimulations acoustiques, basées sur les réactions cardiaques du fœtus ont montré que pendant la vie intra-utérine, le fœtus peut déjà discriminer des stimulations sonores variant dans leur dimension de hauteur, durée, ... Il est capable de percevoir « *la permutation d'ordre d'une paire de syllabes, le passage d'un locuteur masculin à un locuteur féminin, ou le changement de hauteur d'une note de musique. [...] Cela fait intervenir des capacités d'attention auditive sélective et de mémoire à court-terme, et peut être considéré comme marquant un apprentissage élémentaire. L'expérience auditive fœtale peut aussi participer à l'élaboration des capacités perceptives (acuité, discrimination, préférences) et préparer le cerveau à traiter des indices acoustiques qui seront pertinents pendant la vie post-natale.* »¹⁰⁶

Puis, le nouveau-né est capable d'identifier les voix. D'autres études basées sur la succion du bébé lorsqu'il entend la voix de sa mère ou celle d'une autre femme ont montré que le taux de succion est supérieur quand l'enfant entend la voix maternelle.¹⁰⁷

¹⁰⁶ E. SALIBA, S. HAMAMAH, F. Gold, *Médecine et biologie du développement: du gène au nouveau-né*, 2001, p.206.

¹⁰⁷ A.J. DE CASPER, W.P. FIFER, *Of human bonding: newborns prefer their mothers' voices*, 1980, p.1174

L'enfant préfère la voix de sa mère à toute autre voix. Au fil des mois, il saura devenir attentif à son interlocuteur, écouter et reconnaître le timbre vocal, comprendre son nom, répéter des sons et mots entendus, ... Ainsi, la prosodie, le rythme vocal de sa langue maternelle se fixent dans le cortex cérébral grâce à l'oreille. Celle-ci s'affine peu à peu et d'autant plus avec la pratique d'un instrument. « *A sept ans, il sait imiter, moduler et jouer avec sa voix.* »¹⁰⁸

5.2.1.4 Perception des sons

5.2.1.4.1 Cérébrale

Entre les deux hémisphères cérébraux, la perception des sons et des harmoniques varie. Il apparaît que du côté gauche, l'information sonore recueille des bandes de fréquences larges, en faible quantité mais pouvant se modifier promptement. Alors que l'hémisphère droit intervient davantage dans la réception des harmoniques, grâce à un filtrage plus performant. L'expression de l'information est peu rapide. « *Ainsi, l'oreille de la parole et de l'analyse du solfège est à gauche. L'oreille des harmoniques et de la mélodie est à droite.* »¹⁰⁹

Le cortex temporal gauche est donc associé à l'analyse des sons du langage alors que le cortex temporal droit est préférentiellement associé à l'analyse des sons musicaux. Cette asymétrie fonctionnelle est corrélée à une asymétrie anatomique, présente dans le cortex auditif. Chez le droitier, le planum temporal est plus marqué à gauche alors que le gyrus de Heschl (aire auditive primaire) est plus conséquent à droite. Cette organisation se retrouve chez la plupart des gauchers.¹¹⁰

¹⁰⁸ J.ABITBOL, op. cit. , p.114.

¹⁰⁹ Ibid, p.121.

¹¹⁰ B. KOLB, I.Q. WHISHAW, op. cit. , p.332.

5.2.1.4.2 Hauteur

Dans notre étude, la hauteur est le paramètre le plus important sur lequel le sujet devra jouer. C'est pourquoi nous nous centrerons ici sur la description de la perception de la hauteur essentiellement.

La hauteur est principalement déterminée par la fréquence des vibrations composant le son, à savoir le nombre de vibrations par seconde, exprimé en Hz. Ce sont les cellules de la cochlée qui renseignent sur la hauteur des sons comme nous l'avons vu (les basses fréquences activant les cellules proches de l'apex et les fréquences aiguës activant celles à proximité de la base de la cochlée). Cette représentation est dite « tonotopique ». Chaque axone du nerf cochléaire est connecté à une seule cellule ciliée. Cela offre une grande précision quant à la zone de la membrane basilaire stimulée. Même si chacun des axones transmet une information seulement sur une petite partie du spectre auditif, les cellules peuvent répondre à une gamme de fréquences plus large à la condition que l'intensité du son augmente. Les fibres répondant aux fréquences élevées peuvent aussi répondre aux basses fréquences. Il y a donc plus de fibres pouvant répondre aux basses fréquences. « *Le mode d'activité de chaque fibre du nerf cochléaire semble dépendre de la partie de la cochlée d'où elle provient.* »¹¹¹ Cette organisation spatiale se retrouve dans les noyaux cochléaires du tronc cérébral, jusqu'à l'aire du cortex auditif primaire.

Notons cependant, que pour les très basses fréquences, ce système tonotopique n'est pas autant respecté. Dans ce cas, toutes les cellules situées à l'extrémité apicale de la membrane basilaire répondent aux mouvements de celle-ci et de façon proportionnelle au son incident.

¹¹¹ *ibid.* , pp.332-333.

5.2.2 CAS DE L'OREILLE ABSOLUE

Un individu possède l'oreille absolue lorsqu'il a la capacité d'identifier une note musicale sans disposer de référence préalable. Cela est possible grâce à l'activité des cellules ciliées offrant une capacité de discrimination auditive extrêmement fine des fréquences mais aussi grâce à une mémoire sonore particulièrement développée. Le sujet est capable de reconnaître des notes très proches et de courte durée, de nommer la hauteur et le timbre des tonalités entendues. Il existe une oreille absolue « passive » avec laquelle la personne peut identifier une note entendue sans référence autre mais est incapable de chanter avec justesse une note demandée. La personne bénéficiant d'une oreille absolue « active » est capable non seulement de nommer la note entendue mais aussi de chanter avec justesse une note donnée.

Un autre fait s'observe : les chanteurs lyriques notamment, bénéficiant de l'oreille absolue, identifieraient les hauteurs sonores grâce à leur mémoire proprioceptive et non auditive. *« C'est en positionnant mentalement leurs organes phonatoires qu'ils retrouvent la hauteur exacte de la note qu'ils vont chanter. »*¹¹²

La question de l'oreille absolue innée ou acquise se pose et fait débat. Cette capacité auditive et mnésique semble être favorisée par l'apprentissage de la musique durant l'enfance. Mais il apparaît également qu'elle se heurte à un âge critique : 6 ans. Les études réalisées montrent que l'entraînement et le contact musical joue incontestablement dans le développement des capacités de reconnaissance. Cependant, les diverses différences entre les individus laissent supposer qu'un don inné intervient lui aussi.¹¹³

¹¹² N.SCOTTO DI CARLO, *Oreille absolue et mémoire proprioceptive*, 2003, pp.7-15.

¹¹³ F.A. Russo et al., *Learning the "special note": Evidence for a critical period for absolute pitch acquisition*, 2003.

5.3 LIEN ENTRE PRODUCTION ET PERCEPTION

5.3.1 CONDUCTION DU SON PRODUIT

L.I.SHUSTER et J.D.DURRANT¹¹⁴ ont essayé de mieux comprendre comment le locuteur perçoit son message vocal. Pour cela, ils se basent sur différentes études.

Lorsque des personnes écoutent leur voix sur un enregistrement, la plupart disent percevoir une différence entre le discours enregistré et leur parole telle qu'elles l'entendent en direct. Selon J.TONNDORF¹¹⁵ : « *La voie utilisée pour transmettre le son n'est pas la même selon le mode d'écoute.* » Le son produit par l'enregistrement est entendu grâce à la conduction aérienne seule tandis que lors de l'écoute de son propre discours, le son est transmis à l'oreille par conduction aérienne et osseuse.

G.VON BESEKY¹¹⁶ « *dit que la communication, tant chez l'homme que chez de nombreux animaux, a été conçue de façon à maximiser la capacité de transport de la voix tout en réduisant l'intensité avec laquelle l'organisme l'entend lui-même. Il soutient que l'oreille moyenne est construite de telle sorte que le son généré par le larynx n'est pas transféré par les mécanismes auditifs de la même manière qu'un son provenant de l'extérieur. Par conséquent, la voix atteint son but de communication envers l'auditeur sans pour autant être perçue trop forte par le locuteur. Il a également constaté que, lors du passage bouche-conduit auditif externe, il existe une perte d'intensité pour les fréquences élevées, pendant l'audition du discours autoproduit.* »

¹¹⁴ L.I. SHUSTER, J.D. DURRANT, *Toward a better understanding of the perception of self-produced speech*, 2003.

¹¹⁵ J.TONNDORF, *Bone conduction*, 1972, pp.197-237.

¹¹⁶ G.VON BESEKY, *The structure of the middle ear and the hearing's of one's own voice by bone conduction*, 1949.

J.TONNDORF¹¹⁷ a étudié « l'efficacité de la conduction osseuse et aérienne. Il apparaît qu'à basses fréquences, les deux voies sont d'une efficacité semblable. Cependant, à des fréquences plus élevées, la conduction osseuse est moins performante. C'est pourquoi selon lui l'enregistrement de sa propre voix sonne différemment. Il manque l'accent de basse fréquence à laquelle nous sommes habitués au cours de notre production en temps réel. Le son serait sans doute plus naturel si l'accent de basse fréquence pouvait être fourni sur l'enregistrement avec un filtre électronique.

5.3.2 CONTROLE DU FEEDBACK AUDITIF

Le feedback auditif joue un rôle important dans la parole. Lorsque nous parlons, la hauteur, la pression acoustique sont volontairement ou non contrôlées grâce au feedback auditif.

A.TOYOMURA et al.¹¹⁸ ont essayé de mettre en évidence les mécanismes neuronaux impliqués dans le feedback auditif pour le contrôle de la hauteur notamment. Pour cela, l'imagerie par résonance magnétique fonctionnelle a été utilisée. Douze sujets droitiers ont dû vocaliser un [a] pendant 5 secondes tout en écoutant leur voix grâce à des écouteurs. Dans un premier temps, le feedback auditif perçu par les sujets était artificiel¹¹⁹, c'est-à-dire modifié volontairement. Si la hauteur du feedback entendu était déplacée vers le haut, ils devaient adapter leur émission en conséquence. Dans un second temps, le feedback perçu n'a pas été modifié¹²⁰, les sujets ont dû vocaliser en continu le [a] de façon normale. Ils ont donc comparé l'activité cérébrale au cours de la vocalisation émise avec et sans modification de la hauteur du feedback perçu.

¹¹⁷ J. TONNDORF, op. cit.

¹¹⁸ A.TOYOMURA, S. KOYAMA, T. MIYAMAOTO, A. TERAOKA, T. OMORI, H.MUROHASHI AND S. KURIKI, *Neural Correlates Of Auditory Feedback Control In Human*, 2007, pp.499-503.

¹¹⁹ Transformed Auditory Feedback (TAF)

¹²⁰ Non-Transformed Auditory Feedback (non-TAF)

« Les différences observées ont été principalement ciblées dans l'hémisphère droit. Bien que ne pouvant justifier dans quelle mesure ces différences sont attribuables aux variations des conditions acoustiques, ils émettent l'hypothèse que l'activation observée au niveau du sillon intrapariétal droit serait due à ces variations. Il apparaîtrait aussi que le gyrus supramarginal droit serait davantage impliqué dans le feedback pour le contrôle de la hauteur plutôt que dans la perception de la voix elle-même. D'autres éléments s'avèrent également impliqués dans le feedback : il s'agirait de la zone préfrontale droite, de l'aire temporale supérieure droite et de l'insula antérieure droite. Le cortex prémoteur de l'hémisphère gauche semble lui aussi actif malgré une dominance certaine de l'hémisphère droit pour cette fonction. »¹²¹

5.3.3 PERCEPTION ET AJUSTEMENT VOCAL

R. BEHROOZMAND et al.¹²² ont essayé de cerner quel pouvait être l'impact du feedback auditif sur l'ajustement vocal : *« Le contrôle de la fréquence fondamentale (F0) joue un rôle important dans la production de la parole et permet une transmission efficace des signes linguistiques et non-linguistiques dans la communication humaine. Deux mécanismes ont été mis en évidence : tout d'abord, un contrôle antérieur où le système moteur ajuste les paramètres biomécaniques des muscles du larynx à partir des commandes motrices apprises auparavant, puis les informations sensorielles (kinesthésiques et auditives) sont utilisées pour minimiser le décalage entre l'émission vocale prévue et réelle.¹²³ »*

Suite à leur étude, les auteurs expliquent que le cortex auditif est plus sensible aux variations du retour auditif pendant la vocalisation que lorsqu'il y a écoute passive de la voix via un enregistrement. *« La modulation interne des neurones auditifs fournirait un mécanisme*

¹²¹ A.TOYOMURA, S. KOYAMA, T. MIYAMAOTO, A. TERAOKA, T. OMORI, H.MUROHASHI AND S. KURIKI, op. cit. , pp.501-502.

¹²² R. BEHROOZMAND et al., Vocalization-induced enhancement of the auditory cortex responsiveness during voice F0 feedback perturbation, 2009, p.1303.

¹²³ R. BEHROOZMAND et al., op. cit. , p.1303.

pour augmenter la sensibilité des neurones afin de mieux détecter et aider le système moteur vocal à corriger les déviations apparaissant dans le feed-back de la voix générée. »¹²⁴

Le feed-back permettrait donc au système de production vocale d'être informé de la réalisation des vocalisations en vue d'une adaptation, si d'éventuelles erreurs survenaient pendant l'émission vocale.

¹²⁴ *ibid.* , p.1312.

PARTIE PRATIQUE

IMITATION ET ZONE DE SECURITE

PROSODIQUE

Dans la rééducation orthophonique, et plus particulièrement dans la rééducation de la voix, il est fréquent que le thérapeute demande au patient d'imiter ce qu'il réalise. L'imitation constitue un support non négligeable pour amener un sujet à modifier ou moduler certains comportements. Or, sommes-nous tous capables de réaliser une modulation de notre voix pour imiter un modèle donné ? Dans un premier temps, nous allons donc tenter de vérifier l'hypothèse selon laquelle les capacités de modulation vocale en vue d'une imitation sont présentes chez tout un chacun, que le sujet soit sain ou dysphonique.

Partant de ce principe, l'orthophoniste constitue donc un modèle que le sujet doit imiter au mieux. Cependant, pour fournir un modèle prosodique spécialement adapté au sujet, le modèle orthophonique est-il toujours adapté ? Il paraît difficile voire impossible pour l'orthophoniste de produire le modèle correspondant exactement à la voix et à la façon de parler d'une personne puisque chaque voix est unique et chaque façon de s'exprimer l'est aussi. Par ailleurs, il existe pour chaque personne une zone que nous appelons « Zone de Sécurité Prosodique » (ZSP) dans laquelle la voix a le meilleur rendement et qui permet d'éviter au maximum le malmenage vocal. Nous savons que de grandes variations de fréquences sont néfastes pour la voix. Pour limiter ces variations, cette zone est définie au cœur du triangle vocalique propre à chaque personne (toutes les voyelles étant contenues à l'intérieur de ce triangle). La voix du sujet garde sa mélodie mais reste dans un intervalle fréquentiel réduit et imposé. La diminution de l'amplitude fréquentielle grâce à l'usage de cette zone diminue logiquement les facteurs de malmenage vocal. Or, le modèle orthophonique peut-il orienter systématiquement tout sujet vers sa zone sécurisée ?

En définissant la zone de sécurité prosodique pour chacun des sujets rencontrés nous allons donc, dans un second temps, tenter de vérifier l'hypothèse selon laquelle un sujet sain comme un sujet dysphonique réalise un meilleur contrôle de sa prosodie lorsqu'il « s'auto-imité »,

c'est-à-dire lorsqu'il imite sa propre voix parlée que lorsqu'il imite l'orthophoniste produisant le modèle et que de surcroît, cette auto-imitation lui permet de rester plus proche de sa zone de sécurité prosodique. Dans le cas où il imite sa propre voix, les paramètres de celle-ci sont alors modifiés (grâce au logiciel PRAAT) de façon à respecter la zone de sécurité prosodique. Dans le cas d'un sujet dysphonique, pour définir cette zone, le sujet doit quand même bénéficier d'une voix et non d'un souffle.

Nous comparerons alors, grâce aux enregistrements effectués dans différents cas de figure (voix mélodique et voix monotone), les écarts entre la ligne mélodique du modèle et celle de l'imitation réalisée par le sujet.

Nous verrons ainsi si son imitation se rapproche le plus du modèle orthophonique ou du modèle de sa propre voix et dans quel cas il reste le plus proche de sa zone de sécurité prosodique.

Nous avons également souhaité appliquer ce protocole à un professionnel de l'imitation afin de voir à quel point une personne aux capacités maximales d'imitation parvient à se rapprocher des modèles proposés. Bien sûr, le sujet n'a pas été choisi selon les critères de sélection de la population.

1. MATERIEL ET METHODE

Nous proposons, dans ce chapitre, de décrire toutes les conditions nécessaires à l'élaboration de notre expérimentation avec notamment le type de population retenu, les épreuves proposées, le matériel utilisé.

1.1 POPULATION

La population est constituée de sujets « témoins » sélectionnés au hasard et de patients dysphoniques ayant été suivis par le docteur Pierre-Olivier VEDRINE, phoniatre à l'hôpital de CANNES. Pour cela, nous nous sommes donc limités à la zone géographique de la région PACA.

Nous avons choisi d'étudier des personnes adultes, de sexe féminin et masculin. Concernant les patients dysphoniques, nous avons retenu des patients ayant souffert ou souffrant d'un épisode dysphonique mais pour lesquels la voix n'est pas exagérément soufflée. En effet, pour trouver la zone de sécurité prosodique, il faut qu'une variation mélodique soit possible et que l'exercice proposé n'ajoute pas des tensions vocales supplémentaires. Nous avons donc exclu de notre étude, les personnes présentant une paralysie récurrentielle, une dysphonie spasmodique, une dysarthrie, un trouble ou retard cognitif majeur, une surdité, une laryngectomie partielle ou totale.

Les personnes devaient être coopérantes au moment de l'examen et maîtriser suffisamment la langue française pour comprendre les consignes et réaliser les épreuves.

Au total, 19 personnes ont subi les enregistrements audio du protocole dont 9 hommes et 10 femmes, âgés de 18 à 64 ans. Nous retrouvons 14 témoins (7 hommes / 7 femmes) et 4 patients dysphoniques (1 homme / 3 femmes). L'étude a été menée sur 3 mois.

Nous avons par ailleurs étudié un imitateur professionnel, Michaël GREGORIO, suivi par le docteur Jean ABITBOL à PARIS.

1.2 PROTOCOLE ET RECUEIL DES ENREGISTREMENTS

1.2.1 MATERIEL

1.2.1.1 Le logiciel

Pour enregistrer et effectuer l'analyse acoustique objective des productions de chaque individu, nous avons utilisé le logiciel informatique Praat créé par Paul Boersma et David Weenink de l'Institute of Phonetic Sciences de l'Université d'Amsterdam en 1996.

L'ordinateur que nous avons utilisé pour réaliser les enregistrements de l'imitateur et des patients dysphoniques est un ordinateur portable classique de type Asus. L'ordinateur utilisé pour l'enregistrement des témoins est quant à lui de type Packard Bell.

1.2.1.2 Les éléments périphériques

Nous utilisons un micro afin d'enregistrer les productions des sujets et de transformer les variations de la pression de l'air, due au passage de l'onde sonore, en un signal électrique. Pour ce faire, il doit être de bonne qualité, générant le moins de bruit possible, captant convenablement les sons produits et sensible à la gamme de fréquence nécessaire (de 50 à 15000 Hz). Pour l'imitateur et les patients dysphoniques, nous avons utilisé l'enregistreur numérique Zoom H4N. Le H4N enregistre jusqu'à 4 canaux simultanément sur des cartes SD et SDHC Flash media jusqu'à 32GB en 24 bit 96KHz. La paire de micros X/Y intégrés au H4n offre une capture spatiale de l'enregistrement de grande qualité. Les enregistrements sont naturels, profonds et précis. Pour la population témoin, nous avons utilisé le micro SHURE 849. C'est un microphone électrostatique, à électret de haute qualité, à réponse en fréquence régulière et linéaire et à configuration cardioïde. Il offre une reproduction sonore de haute qualité. Sa réponse en fréquence (c'est-à-dire la plus ou moins grande sensibilité du micro à la fréquence des vibrations ; idéalement, elle doit être la même pour toutes les fréquences audibles par l'oreille humaine, comprises entre 20 Hz et 20 kHz) va de 40 à 16 000 Hz. Sa courbe de directivité (distribution spatiale de la sensibilité) est donc de type cardioïde. Sa sensibilité (tension électrique que le micro délivre à partir d'une pression donnée. Plus la sensibilité est élevée, plus la tension produite est importante, à pression égale, et plus le micro

peut donc détecter des variations fines de pression) à 1 000 Hz est de – 51dBV/Pa (2,8 mV). Notre micro peut donc être considéré comme sensible.

Pour une écoute de bonne qualité, un ampli associé à des haut-parleurs branchés sur l'ordinateur est nécessaire. Cela nous permet, à partir de Praat, de faire écouter au patient sa production lors des épreuves concernées.

1.2.2 EPREUVES

Les épreuves présentées doivent être identiques pour toutes les personnes et être proposées dans le même ordre afin d'éviter au maximum la présence de biais difficiles à corriger et à évaluer. Le protocole est donc standard et permet de prendre en compte les mêmes paramètres pour comparer les différents échantillons de voix. Les résultats extraits au cours des différentes épreuves pour calculer les modifications prosodiques à apporter sont notés au fur et à mesure de la passation sur une feuille prévue à cet effet (annexe n°1).

Le protocole comporte 4 épreuves à partir de PRAAT (nous le résumons et donnons un exemple dans l'annexe n°2) :

1 - Enregistrement comportant les trois voyelles les plus extrêmes du Français [a], [i], [u] : nous choisissons ce triangle vocalique car les données recueillies incluront forcément le reste des voyelles utilisées dans la production orale.

« Prendre une inspiration de façon à produire une glissade montante sur un [ma], de durée suffisante, d'intensité et de hauteur confortables, sans passer en voix chantée. »

Pour éviter toute confusion, nous donnons l'exemple d'une glissade. Nous proposons le [ma] pour éviter d'éventuels coups de glotte. Nous enregistrons la production du sujet grâce à la manipulation suivante : « New » (dans le menu principal)

« Record mono sound » (pour tous les enregistrements, nous définissons le « Sampling frequency » à 44000 Hz)

« Record »

« Stop » puis « Rename » pour le nommer tel qu'on le souhaite.

« Save to list ». Nous obtenons un objet PRAAT « *sound [ma]* » ajouté à la liste des objets dans la fenêtre principale.

Puis, nous éditons l'objet avec la fonction « Edit » et sélectionnons une portion riche en harmoniques du [a] produit. Cette zone se définit de façon subjective. Il s'agit d'observer visuellement où les harmoniques commencent et où ils terminent (détecter la zone de transition, c'est-à-dire la zone de désorganisation des harmoniques). La zone adéquate doit être très noire (cf zone rose sur l'image qui suit). Par ailleurs, il y a également un travail auditif, une analyse de l'oreille pour éviter de sélectionner une zone de souffle ou de forçage. La zone sélectionnée est d'une durée d'une seconde environ pour tous les cas de figure et pour chacune des personnes.

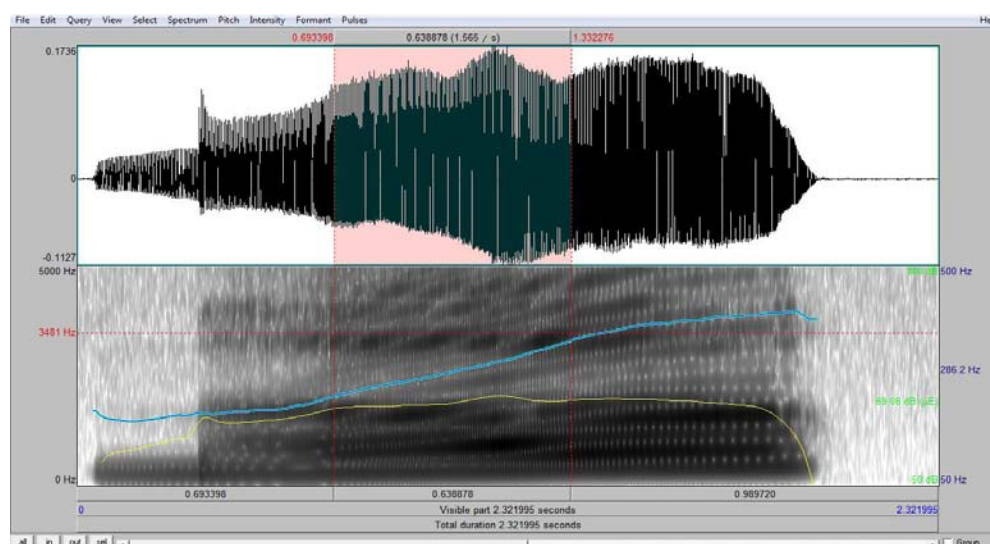


Figure 9 : Zone sélectionnée riche en harmoniques

Nous observons et relevons quelles sont les fréquences minimale et maximale du [a] avec la fonction « voice report » (située dans la rubrique « Pulses » de la fenêtre *sound [ma]*). Nous procédons de la même façon pour les voyelles [i] et [u]. Nous obtenons donc un objet « *sound [mi]* » et un objet « *sound [mu]* ».

Puis, nous retenons la fréquence la plus haute parmi les trois fréquences minimales obtenues et la fréquence la plus basse parmi les trois fréquences maximales obtenues. Nous obtenons ainsi une Fréquence maximale voulue (F max voulue) et une Fréquence minimale voulue (F min voulue). Cela nous indique l'intervalle de fréquences dans lequel la voix du sujet doit s'inscrire pour rester dans la zone de sécurité prosodique. Il s'agit donc de

l'intervalle [F max voulue ; F min voulue]. L'écart entre ces deux fréquences (F max voulue – F min voulue) constitue la valeur nommée « δ voulu » et nous servira plus tard pour trouver l'une des constantes permettant de maintenir la voix de l'individu dans cette zone de sécurité prosodique définie.

2 – Production spontanée en deux temps :

« Lire la phrase inscrite sur la feuille présente devant vous sur un ton monotone. »

La feuille donnée figure dans la partie Annexes (annexe n°3). Nous ne lisons pas la phrase afin d'éviter tout mimétisme et de ne pas influencer le sujet quant à sa production. La phrase est la suivante : « Je voudrais partir en Italie pour admirer la Tour de Pise. » Elle comporte intentionnellement les trois voyelles les plus extrêmes du Français. Pour aider le sujet à comprendre ce qu'est le ton monotone, nous lui indiquons qu'il est similaire au ton attribué au robot. Nous pouvons si nécessaire, donner un exemple en utilisant une autre phrase qui est « Le soleil brille aujourd'hui. »

Nous enregistrons la production de la personne de la même façon que les glissades. Nous obtenons un objet « *sound phrase monotone* » Cette épreuve nous permet de voir où se situe spontanément le sujet si on lui impose un ton monotone.

« Lire la même phrase sur un ton mélodique. »

Si nécessaire, nous lui donnons l'exemple avec la phrase « Le soleil brille aujourd'hui. »

Nous enregistrons la production du sujet et obtenons un objet « *sound phrase mélodique* ». Cette épreuve nous permet de voir où se situe spontanément la personne si on lui impose un ton mélodique et plus largement, elle nous permet de relever dans quel champ de fréquences se situe le sujet. Ainsi, nous sélectionnons la phrase produite, puis nous appliquons la fonction « voice report » et obtenons une Fréquence maximale observée (F max observée) et une Fréquence minimale observée (F min observée). L'écart entre ces deux valeurs (F max observée – F min observée) constitue la valeur nommée « δ observé » et nous servira plus tard pour trouver l'une des constantes permettant de maintenir la voix de l'individu dans cette zone de sécurité prosodique définie.

3 - Imitation du modèle orthophonique en deux temps :

« J'utilise la même phrase que précédemment. Imitiez le modèle monotone que je vous donne. Vous pouvez imiter dès que j'ai terminé la phrase. »

Nous enregistrons le modèle orthophonique puis la production du sujet à la suite. Nous pourrions ainsi comparer leurs différences. Nous obtenons un objet « *sound imitation ortho monotone*. »

« J'utilise la même phrase que précédemment. Imitiez le modèle mélodique que je vous donne. Vous pouvez imiter dès que j'ai terminé la phrase. »

Nous procédons de la même façon, en enregistrant d'abord le modèle orthophonique puis la production du sujet. Nous obtenons un objet « *sound imitation ortho mélodique*. »

4 - Imitation de sa voix transformée respectant la zone de sécurité prosodique en deux temps :

« Vous allez entendre votre voix enregistrée lors de l'épreuve sur consigne où vous lisez sur un ton mélodique. Votre voix a subi des modifications de façon à rester dans votre zone de sécurité prosodique. Imitiez votre voix. »

Pour appliquer la zone de sécurité prosodique c'est-à-dire le champ de fréquences [F max voulue ; F min voulue] à la voix mélodique spontanée du sujet, il nous faut trouver deux constantes : α et β .

$$\alpha = \delta \text{ voulu} / \delta \text{ observé.}$$

$$\beta = F \text{ max voulue} - (F \text{ max observé} \times \alpha)$$

Une fois ces deux constantes établies, nous reprenons la phrase mélodique spontanée de la personne, c'est-à-dire le « *sound phrase mélodique* » et lui appliquons la fonction « to manipulation » (située dans la fenêtre du menu principal). Nous obtenons la ligne mélodique de sa voix en vert. Nous appliquons les constantes. Pour cela, nous sélectionnons la ligne mélodique. Nous imposons la valeur α dans la fonction « multiply pitch frequencies » (située

dans la rubrique « Pitch »). Puis, nous imposons la valeur β dans la fonction « shift pitch frequencies » (située elle aussi dans la rubrique « Pitch »). Nous disposons alors d'un nouvel objet « *manipulation phrase mélodique sécurisée* ». C'est cet objet qui constitue la voix dite « sécurisée ».

Bien que les données soient objectives, soulignons que le jugement de l'orthophoniste intervient également. Il peut ajuster la voix pour que celle-ci corresponde toujours au sujet. L'enjeu est de trouver un compromis pour que la voix reste dans la zone de sécurité prosodique et que la personne la reconnaisse comme sienne.

Nous demandons donc à l'individu d'imiter cette « nouvelle voix » et l'enregistrons après une écoute (une seconde écoute peut être tolérée si nécessaire). Nous obtenons l'objet « *sound imitation phrase mélodique sécurisée* ».

« Vous allez entendre votre voix enregistrée lors de l'épreuve sur consigne où vous lisez sur un ton monotone. Votre voix a subi des modifications de façon à rester dans la zone de sécurité prosodique. Imitiez votre voix. »

Nous recherchons la ligne monotone optimale pour le sujet. Plus celui-ci se rapprochera de cette ligne monotone, plus il approchera la zone de sécurité prosodique maximale. Bien sûr, cette ligne est dénuée de mélodie, elle ne peut donc pas être utilisée comme « nouvelle voix ». Elle permet cependant de donner au patient la ligne mélodique autour de laquelle il faut faire varier la voix.

Celle-ci s'obtient de cette façon : nous reprenons la phrase mélodique spontanée du sujet, c'est-à-dire le « *sound phrase mélodique* » et lui appliquons la fonction « to manipulation ». Nous obtenons à nouveau la ligne mélodique de sa voix en vert. Puis, nous définissons la fréquence moyenne (F moyenne) où la ligne monotone doit se situer pour le sujet.

$$F \text{ moyenne} = F \text{ min voulue} + (\delta \text{ voulu} / 2)$$

Nous appliquons cette F moyenne. Pour cela, nous sélectionnons la ligne mélodique. Nous utilisons la fonction « Add pitch point at » (située dans la rubrique « Pitch ») pour laquelle nous imposons un « Time » débutant avant le début de la ligne mélodique et une « Frequence » correspondant à la valeur F moyenne. Puis, nous utilisons la fonction

« Remove pitch point » (située elle aussi dans la rubrique « Pitch »). Nous obtenons un nouvel objet « *manipulation phrase monotone sécurisée* ».

Nous demandons donc à la personne d'imiter cette « nouvelle voix » et l'enregistrons après une écoute (une seconde écoute peut être tolérée si nécessaire). Nous obtenons l'objet « *sound imitation phrase monotone sécurisée* ».

A partir des résultats obtenus dans chacune des épreuves, nous comparerons alors si l'imitation réalisée par le sujet se rapproche plus du modèle orthophonique ou du modèle basé sur sa propre voix, tant sur le ton monotone que mélodique.

Nous avons choisi de proposer également un questionnaire (annexe n°4) afin de mieux connaître le sujet et cerner son parcours, son ressenti face à sa voix, face à la nouvelle voix proposée, face à l'imitation. Les réponses obtenues sont associées au paragraphe de la partie pratique correspondant. Le questionnaire comporte deux parties pour la population témoin. La dernière partie est ajoutée pour la population dysphonique :

- Généralités : c'est une prise de contact avec la personne. Nous lui demandons son âge, sexe, profession, s'il utilise fréquemment sa voix au cours de la journée.
- Impressions concernant l'exercice proposé : savoir s'il lui est difficile de placer sa voix en fonction de celle de l'orthophoniste, ce qu'il pense de la nouvelle voix proposée.
- Cause(s) de la dysphonie.

Nous avons interrogé le sujet à la suite des épreuves d'enregistrement. Le mode de récupération du questionnaire s'est fait en main propre, le jour de la passation des épreuves.

Toutes les épreuves d'enregistrement ont également été appliquées au professionnel de l'imitation. Par ailleurs, un questionnaire (annexe n°5) lui aura été proposé afin de connaître son parcours, son ressenti face à cette épreuve et éventuellement en retirer des pistes pour aider le patient à imiter au mieux. Nous avons inclus les réponses correspondant au protocole dans le paragraphe de la partie pratique consacré à l'imitateur. Nous avons précédemment exposé le reste des réponses dans la partie théorique traitant de l'imitateur en général. Le questionnaire comporte quinze questions. Il est divisé en trois parties :

- Généralités : pour mieux connaître le parcours de l'imitateur et les raisons l'ayant conduit à l'imitation.
- Pratique générale : pour cerner comment l'imitateur gère le travail de la voix.
- Impressions sur l'exercice proposé : pour avoir l'avis d'un professionnel de l'imitation sur les épreuves présentées.

Ne disposant que de peu de temps le jour de cette rencontre, nous avons préféré opter pour une réponse aux questions par téléphone. Elles sont retranscrites dans le questionnaire en annexes.

1.2.3 CONDITIONS D'ENREGISTREMENT

Pour chaque personne, nous réalisons un enregistrement audio pour l'ensemble des épreuves. Les deux premières épreuves ne nécessitent que l'enregistrement de la production du sujet contrairement aux deux suivantes qui requièrent de disposer à la fois de l'enregistrement du modèle donné et de la production de l'individu.

Bien entendu, il est recommandé d'effectuer les enregistrements de voix dans une pièce aussi calme que possible. Ceci est particulièrement important pour l'analyse objective, car contrairement à l'oreille humaine, le logiciel utilisé ne saura pas distinguer la voix du sujet des éventuels bruits extérieurs captés par le micro. Dans notre étude, nous tâchons donc de réduire au maximum les sources sonores parasites.

Même si notre étude ne se centre pas sur l'intensité, il est important d'essayer de placer le micro toujours à la même distance. Avant le premier enregistrement, nous leur expliquons le déroulement de l'ensemble des épreuves en lisant les consignes qui leur seront redonnées par la suite au fur et mesure des épreuves.

Concernant la sélection de l'imitateur et les conditions matérielles d'enregistrement, elles auront été définies par le docteur ABITBOL qui a eu l'amabilité de nous accueillir dans son cabinet. Nous avons essayé de nous rapprocher au mieux des conditions types retenues dans ce mémoire.

2.RECUEIL ET ANALYSE DES DONNEES

Dans cette partie, nous exposerons tous les éléments obtenus à l'issue des questionnaires proposés aux sujets étudiés et du protocole expérimental mis en place. Nous présenterons dans un premier temps les résultats de l'imitateur puis de la population témoin et enfin de la population dysphonique. Nous réaliserons dans un tableau, une synthèse de l'ensemble de ces résultats à la fin de ce chapitre. Il est important de souligner que dans le recueil des résultats via les spectrogrammes, nous ne prendrons pas en compte les fréquences susceptibles d'être des artefacts causés par les conditions d'enregistrement, une émission ponctuelle de consonne parasitant la réalisation vocalique...

2.1 IMITATEUR

Nous évoquons ici les résultats recueillis suite aux enregistrements de M.GREGORIO. Nous traitons d'abord des épreuves sur ton mélodique puis celles sur ton monotone. L'étude d'une telle voix offre des données où la performance d'imitation est sensée être la meilleure. Ainsi, nous choisissons de mettre les résultats, les spectrogrammes et les lignes mélodiques obtenues pour l'ensemble des épreuves de M.GREGORIO afin d'illustrer de façon complète l'ensemble des éléments que nous exploitons dans notre travail (annexe n°6).

Pour le reste de la population étudiée, nous ne mettrons que les éléments les plus pertinents, à savoir ceux permettant de calculer les constantes pour la zone de sécurité prosodique et les lignes mélodiques obtenues pour pouvoir comparer modèle orthophonique / imitation du sujet et modèle de la voix transformée / imitation du sujet.

2.1.1 TON MELODIQUE

Dans un premier temps, nous illustrons à l'aide d'un histogramme, les résultats de la fréquence fondamentale moyenne obtenus pour M.GREGORIO et chacun des modèles en fonction des différentes épreuves. Cela nous permet notamment de comparer la fréquence moyenne entre M.GREGORIO et le modèle orthophonique puis entre M.GREGORIO et le modèle de sa voix transformée (là où la voix se situe dans la zone de sécurité prosodique).

Dans un second temps, nous illustrons à l'aide d'un histogramme, les intervalles de fréquences dans lesquels se situent M.GREGORIO et le modèle selon chaque épreuve. Nous indiquons également la zone de sécurité prosodique afin de comparer où se situe M.GREGORIO par rapport à cette zone.

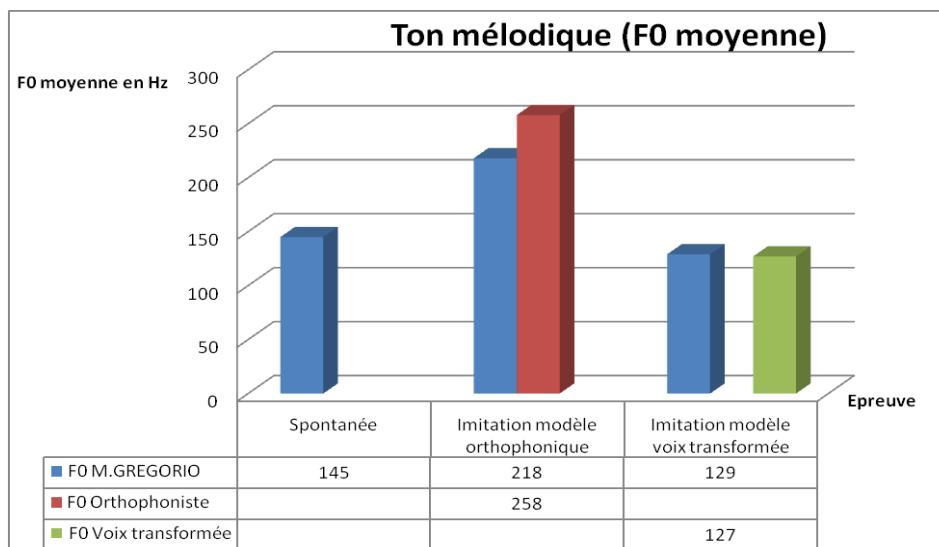


Figure 10 : Histogramme des fréquences fondamentales moyennes pour chaque modèle et pour les productions de M.GREGORIO dans les différentes épreuves en voix mélodique.

Nous constatons tout d'abord que, dans l'épreuve sans modèle où nous demandons au sujet de prononcer la phrase sur un ton mélodique, l'imitateur parle spontanément à une fréquence moyenne de 145 Hz.

Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 258 Hz, nous voyons que la fréquence moyenne de l'imitateur est de 218 Hz. Nous pouvons observer qu'il tente d'imiter la hauteur du modèle orthophonique. Lors de l'enregistrement, M.GREGORIO a demandé la précision suivante : « Faut-il imiter exactement la voix ? » Question à laquelle nous avons répondu : « Vous devez procéder comme il vous semble opportun. » C'est son habitude professionnelle qui l'a sans doute conduit à adopter une voix de tête pour essayer de reproduire au mieux le modèle orthophonique.

L'épreuve d'imitation la mieux réussie est celle de l'imitation de la voix transformée : le modèle est à une fréquence moyenne de 127 Hz et M.GREGORIO réalise une imitation à 129 Hz.

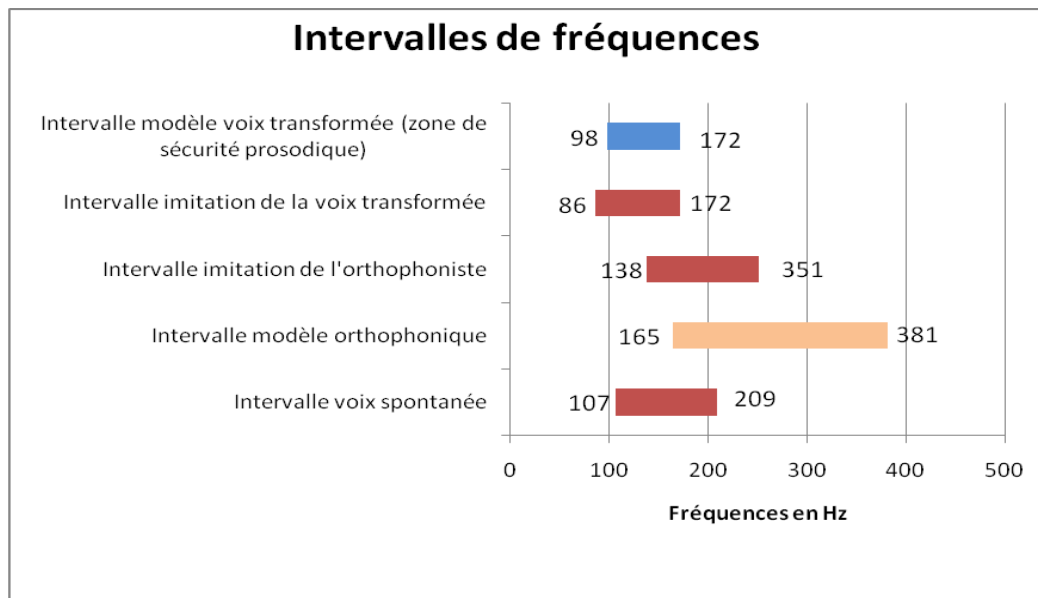


Figure 11 : Histogramme des intervalles de fréquences de M.GREGORIO selon l'épreuve.

Par ailleurs, il est également intéressant d'étudier sous l'angle de l'étendue fréquentielle, les productions réalisées par le sujet ou le modèle, grâce à l'histogramme ci-dessus. Cela nous permet particulièrement de voir dans quelle épreuve l'imitateur parvient à réaliser la meilleure imitation tout en se rapprochant au mieux de la zone de sécurité prosodique. De ce point de vue, nous observons que **l'imitation la mieux réussie est là encore celle de la voix transformée. M.GREGORIO se situe entre 86 et 172 Hz dans son imitation, pour un modèle oscillant entre 98 et 172 Hz ce qui lui permet de rester très proche de sa zone de sécurité prosodique.** Nous voyons également qu'il adapte sa voix puisqu'en épreuve spontanée, il présente un intervalle plus étendu [107 Hz - 209 Hz].

Concernant le modèle orthophonique, nous constatons qu'il est moins bien imité, quand bien même l'imitateur passe en voix de tête. Par ailleurs, cela n'entraîne pas un usage sécurisé de la voix puisque celle-ci dépasse largement la zone de sécurité prosodique. Les variations mélodiques sont donc plutôt importantes dans sa production. Ce sont de tels écarts prosodiques qui sont susceptibles de fatiguer la voix en cas de souffrance vocale.

Ainsi, il apparaît que l'imitateur est capable d'imiter sa propre voix mélodique tout en restant dans sa zone de sécurité prosodique. Malgré ses talents d'imitation, l'épreuve la mieux réussie est donc l'auto-imitation qui favorise les adaptations fréquentielles souhaitées.

2.1.2 TON MONOTONE

Dans un premier temps, nous illustrons à l'aide d'un histogramme, les résultats de la fréquence fondamentale moyenne obtenus pour M.GREGORIO et chacun des modèles en fonction des différentes épreuves. Cela nous permet notamment de comparer la fréquence moyenne entre M.GREGORIO et le modèle orthophonique puis entre M.GREGORIO et le modèle de sa voix transformée (là où la voix se situe dans la zone de sécurité prosodique).

Dans un second temps, nous illustrons à l'aide d'un histogramme, les intervalles de fréquences dans lesquels se situent M.GREGORIO et les modèles selon chaque épreuve. L'intérêt de l'épreuve sur ton monotone est aussi de voir comment l'individu peut osciller autour du point central (milieu) de la zone de sécurité prosodique. Nous indiquons donc le point central de l'intervalle de la zone de sécurité prosodique, point autour duquel la voix de l'imitateur doit tenter de se rapprocher (et qui n'est autre que la fréquence moyenne calculée pour la ligne monotone de la voix transformée - ici, 135 Hz).

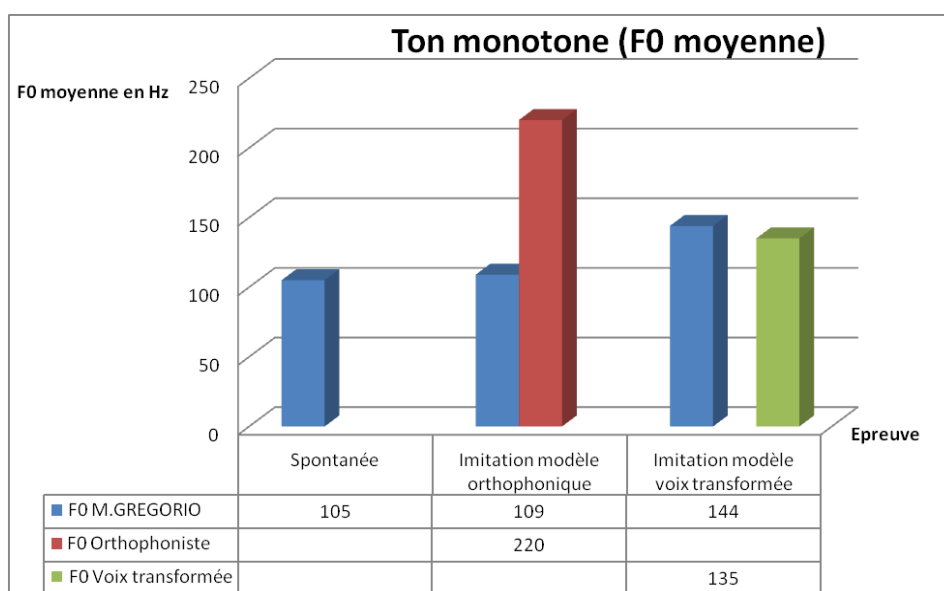


Figure 12 : Histogramme des fréquences fondamentales moyennes pour chaque modèle et pour les productions de M.GREGORIO dans les différentes épreuves en voix monotone.

Soulignons que la fréquence moyenne s'obtient grâce à la fonction « *voice report* ».

Nous constatons tout d'abord que, dans l'épreuve sans modèle où nous demandons au sujet de prononcer la phrase sur un ton monotone, l'imitateur parle spontanément à une fréquence moyenne de 105 Hz.

Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 220 Hz, nous voyons que la fréquence moyenne de l'imitateur est de 109 Hz. Il reste donc plutôt proche de sa fréquence moyenne spontanée et ne se rapproche pas du modèle orthophonique.

Dans le cas où il imite sa propre voix respectant la zone de sécurité prosodique imposée, nous pouvons constater que les fréquences moyennes sont très proches : **le modèle est à 135 Hz et M.GREGORIO réalise une imitation à 144 Hz**. Il adapte sa voix puisque dans l'épreuve spontanée, il se situe à une fréquence moyenne de 105 Hz.

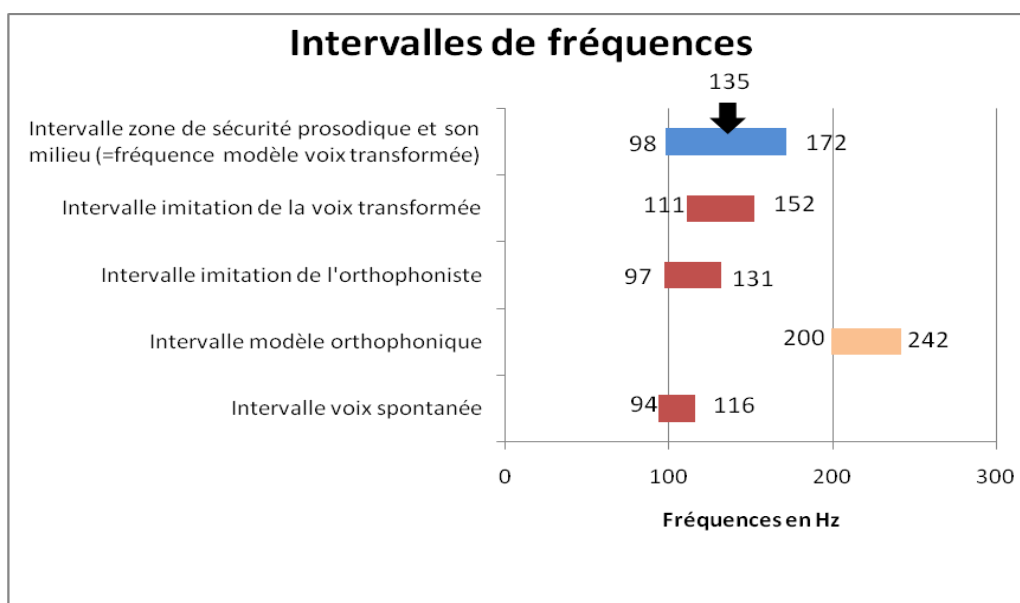


Figure 13 : Histogramme des intervalles de fréquences de M.GREGORIO selon l'épreuve.

En étudiant l'histogramme ci-dessus, nous pouvons observer plusieurs choses. Tout d'abord, il est vrai que dans chacune des épreuves, M.GREGORIO se situe globalement dans sa zone de sécurité prosodique (entre 98 et 172 Hz). Cependant, ce qui nous intéresse, c'est de voir dans quelle épreuve l'imitateur parvient à réaliser la meilleure imitation tout en se rapprochant au mieux de la zone de sécurité prosodique (et plus particulièrement ici du point central de la zone de sécurité prosodique, à savoir 135 Hz). Il apparaît que **l'imitation la mieux réussie, sous l'angle des étendues fréquentielles, est celle de l'imitation de la voix transformée. En effet, dans cette épreuve, la voix de l'imitateur s'adapte et oscille entre 111 et 152 Hz, intervalle se rapprochant du point central de 135 Hz**. Contrairement aux deux autres épreuves où nous voyons des intervalles beaucoup plus bas en fréquences.

De plus, nous voyons que l'imitation du modèle orthophonique est peu probante. L'imitateur reste plutôt proche de son intervalle en voix spontanée. Il n'y a donc pas de réelle adaptation au modèle orthophonique.

Ainsi, malgré ses talents d'imitation, nous constatons à nouveau que l'imitateur réalise de meilleures performances lorsqu'il s'agit d'imiter sa propre voix, et que cette imitation est la plus propice pour respecter au mieux la zone de sécurité prosodique. Les adaptations fréquentielles souhaitées sont favorisées par l'auto-imitation.

2.1.3 ANALYSE DES RESULTATS ET RESSENTI DE L'IMITATEUR

Les résultats des épreuves monotone et mélodique vont dans le même sens. Nous retrouvons une adaptation au modèle orthophonique et au modèle de la voix transformée, tant sur le ton mélodique que monotone. Il était évident que M. GREGORIO serait capable d'imitation mais nous ne savions où se situeraient ses meilleures performances. Il en ressort que l'imitateur est plus performant dans l'épreuve d'auto-imitation. Epreuve pour laquelle la zone de sécurité prosodique est également la mieux respectée. Cela montre que même la personne la plus qualifiée pour réaliser une imitation est davantage réceptive à sa propre voix plutôt qu'à un modèle externe. Sa boucle audio-phonatoire est plus performante lorsqu'il s'agit d'imiter sa prosodie personnelle. Même si l'imitateur est connu pour disposer d'un contrôle moteur très performant sur sa musculature laryngée et ses résonateurs, il n'en demeure pas moins que c'est lors de l'auto-imitation que ce contrôle est le meilleur.

Concernant ce protocole, M.GREGORIO n'a pas ressenti de difficultés particulières. Pour réaliser au mieux cet exercice, il conseille aux personnes et patients éventuels d'écouter la voix, d'entendre sa mélodie, de la visualiser et de bien la timbrer pour essayer d'imiter de la meilleure façon qu'il soit la voix transformée respectant la zone de sécurité prosodique. L'enjeu étant que celle-ci soit la plus naturelle possible.

2.2 GROUPE DES TEMOINS

Nous exposerons tout d'abord dans un tableau les éléments recueillis grâce au questionnaire, afin de présenter la population témoin.

Puis nous évoquerons les résultats recueillis suite aux enregistrements des différents sujets témoins pris dans la population. Nous traitons d'abord des épreuves sur ton mélodique puis de celles sur ton monotone.

Nous procédons de la même façon que pour M.GREGORIO. Les histogrammes obtenus pour chacun des témoins sont disposés les uns à la suite des autres. Nous ferons une brève analyse des résultats au cas par cas pour faire ressortir les éléments pertinents. Cela nous servira à réaliser une analyse globale des résultats obtenus pour le ton mélodique puis pour le ton monotone.

Nous mettrons dans la partie des annexes pour chaque témoin les résultats obtenus pour les différents calculs et les lignes mélodiques respectives permettant de comparer modèle orthophonique / imitation du sujet et modèle de la voix transformée / imitation du sujet.

2.2.1 DESCRIPTION DES TEMOINS

Nous avons recueilli 14 questionnaires. Dans l'analyse, nous nous appuyons sur les résultats bruts des questionnaires que nous mettons en relation avec d'éventuelles informations que les patients auront pu ajouter. Cela nous permet d'avoir une idée précise de la population étudiée. Nous indiquons dans le tableau récapitulatif qui suit : l'âge, le sexe, la profession, l'usage de la voix au quotidien, la pratique de la voix chantée, d'éventuels antécédents vocaux et le numéro des annexes où trouver les résultats et les lignes mélodiques correspondant à chaque sujet. Nous traiterons du ressenti des témoins face à l'exercice proposé lors de la synthèse des résultats obtenus.

Témoïn	Sexe	Age	Profession	Usage de la voix	Pratique voix chantée	Antécédents vocaux	Résultats et lignes mélodiques
N°1	Masc	57	Cadre	Normal	Non	RAS	Annexe n°7
N°2	Masc	31	Artisan	Normal	Non	RAS	Annexe n°8
N°3	Fém	65	Retraitée	Normal	Non	RAS	Annexe n°9
N°4	Fém	45	Educatrice de jeunes enfants	Fréquent	Non	RAS	Annexe n°10
N°5	Fém	55	Sans emploi	Normal	Non	RAS	Annexe n°11
N°6	Masc	18	Lycéen	Normal	Non	RAS	Annexe n°12
N°7	Masc	23	Vendeur	Fréquent	Non	RAS	Annexe n°13
N°8	Masc	63	Retraité	Normal	Non	RAS	Annexe n°14
N°9	Masc	22	Etudiant	Normal	Non	RAS	Annexe n°15
N°10	Fém	23	Etudiante	Normal	Non	RAS	Annexe n°16
N°11	Fém	24	Etudiante	Normal	Non	Dysphonie suite à nodules	Annexe n°17
N°12	Fém	51	Agent administratif	Fréquent	Non	RAS	Annexe n°18
N°13	Fém	24	Etudiante	Normal	Non	RAS	Annexe n°19
N°14	Masc	31	Ingénieur	Normal	Non	RAS	Annexe n°20

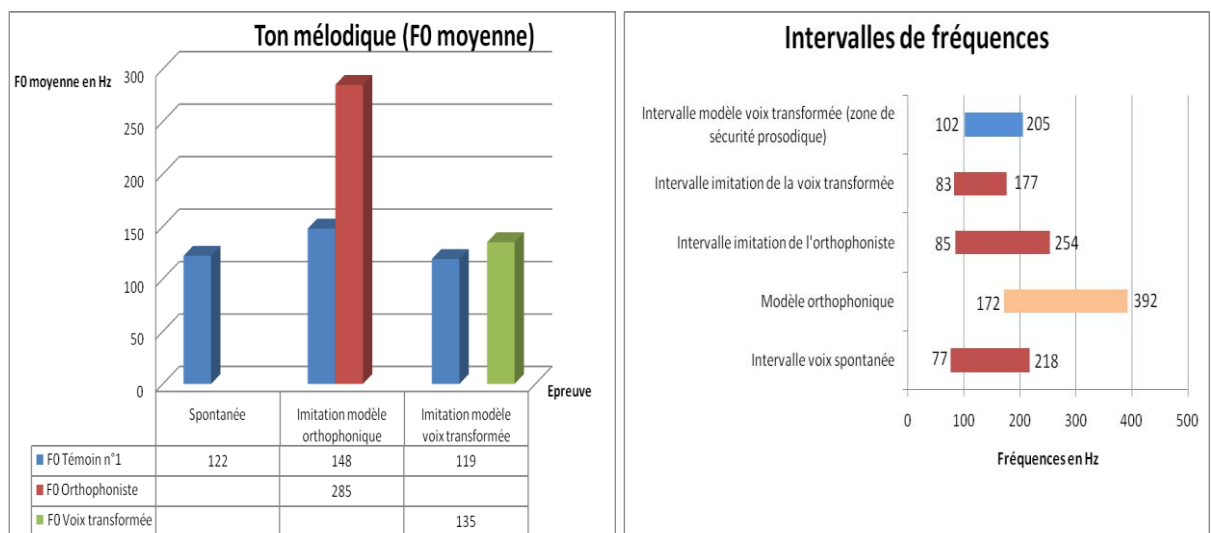
2.2.1.1 Ton mélodique

Dans un premier temps, nous illustrons à l'aide d'un histogramme, les résultats de la fréquence fondamentale moyenne obtenus pour le témoin et chacun des modèles en fonction des différentes épreuves. Cela nous permet notamment de comparer la fréquence moyenne entre le témoin et le modèle orthophonique puis entre le témoin et le modèle de sa voix transformée (là où la voix se situe dans la zone de sécurité prosodique).

Dans un second temps, nous illustrons à l'aide d'un histogramme placé à côté, les intervalles de fréquences dans lesquels se situent le témoin et le modèle selon chaque épreuve. Nous indiquons également la zone de sécurité prosodique afin de comparer où se situe le sujet par rapport à cette zone.

Nous ferons une synthèse des résultats à la fin de cette partie consacrée au ton mélodique.

- Témoin n°1 :

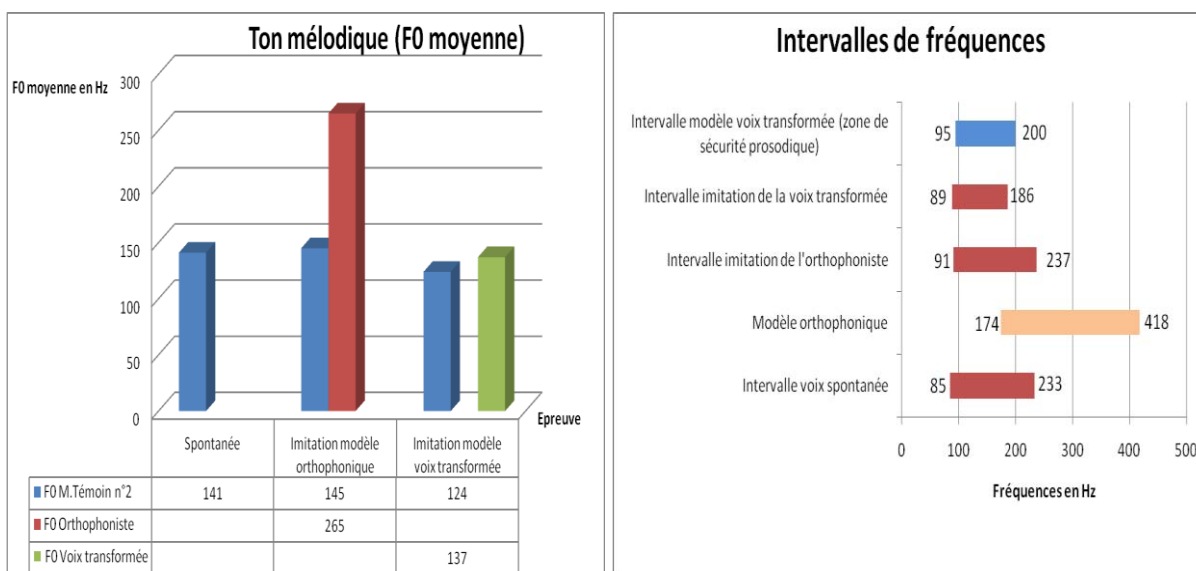


Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 122 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 285 Hz, nous voyons que la fréquence moyenne du sujet est de 148 Hz. Il y a une adaptation au modèle orthophonique. Mais **l'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une**

production de 119 Hz pour un modèle à 135 Hz. Cependant, nous pouvons observer que les fréquences moyennes sont très proches entre l'émission spontanée et le modèle de la voix transformée. Ainsi, ces premiers résultats ne nous permettent pas d'affirmer qu'il y a une réelle adaptation vocale au modèle puisque les fréquences moyennes sont très voisines.

L'observation du second histogramme précise ces éléments. **L'adaptation vocale dans l'imitation de la voix transformée est confirmée.** En effet, nous constatons que l'étendue fréquentielle se réduit entre l'émission spontanée et l'imitation de la voix transformée (au niveau des fréquences maximales). Le sujet passe de 218 Hz à 177 Hz. Cela lui permet également de rester davantage dans sa zone de sécurité prosodique. Soulignons qu'il y a bien tentative d'imitation du modèle orthophonique car l'étendue fréquentielle varie. Mais cette imitation le contraint à sortir de la zone sécurisée puisque sa fréquence maximale est à 254 Hz alors que la zone sécurisée ne dépasse pas 205 Hz.

- Témoin n°2 :

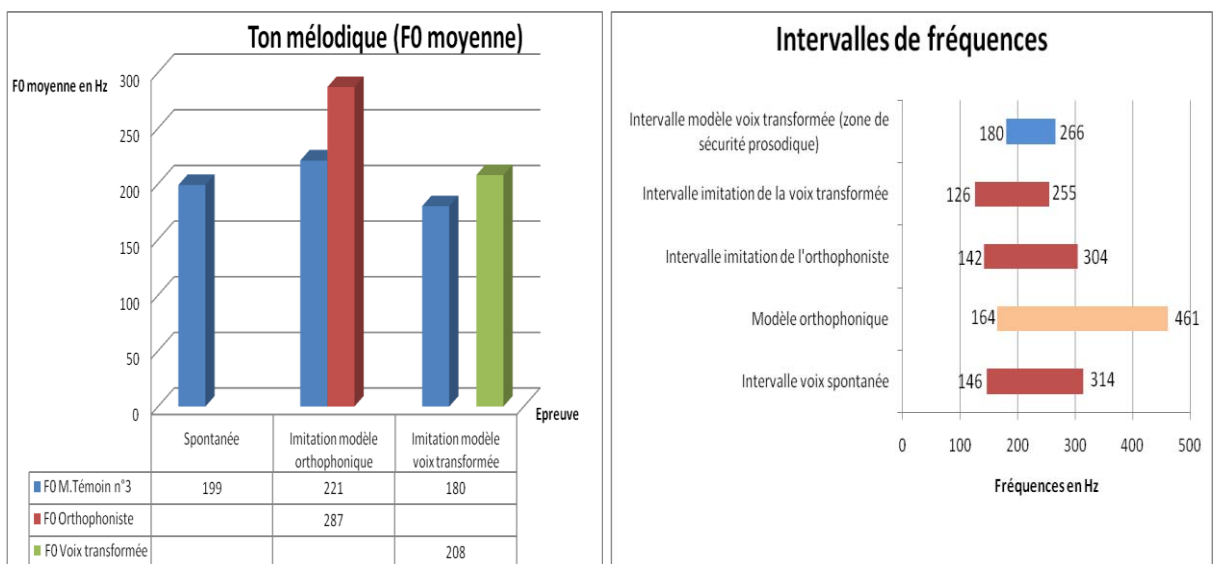


Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 141 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 265 Hz, nous voyons que la fréquence moyenne du sujet est de 145 Hz. Il n'y a pas d'adaptation au modèle orthophonique. **L'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une**

production de 124 Hz pour un modèle à 137 Hz. Même si la fréquence fondamentale moyenne du modèle de la voix transformée est assez proche de la fréquence spontanée du sujet, à l'écoute du modèle proposé, nous entendons que l'étendue fréquentielle est globalement plus réduite que la production spontanée. Pour l'imiter, il fallait donc une légère diminution de l'amplitude tonale. Chose que le sujet a réalisé puisque son imitation a une fréquence moyenne de 124 Hz alors que sa production spontanée est à 141 Hz.

L'observation du second histogramme montre que **l'adaptation vocale dans l'imitation de la voix transformée est confirmée.** En effet, nous constatons que l'étendue fréquentielle se réduit entre l'émission spontanée et l'imitation de la voix transformée (au niveau des fréquences maximales). Le sujet passe de 233 Hz à 186 Hz. Cela lui permet également de rester davantage dans sa zone de sécurité prosodique. Pour l'imitation du modèle orthophonique, il y a une légère modulation de la voix même si le sujet reste proche de son étendue fréquentielle spontanée. Il est donc de ce fait plus éloigné de sa zone de sécurité prosodique.

- Témoin n°3 :

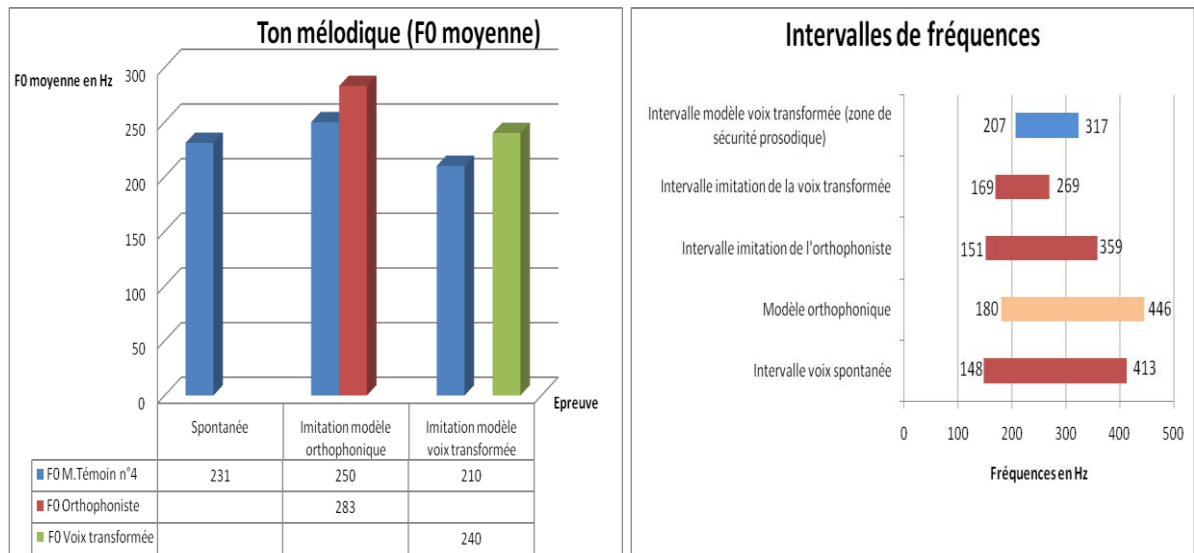


Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 199 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 287 Hz, nous voyons que la fréquence moyenne du sujet est de 221 Hz. Il y a donc une adaptation au

modèle orthophonique. Mais **l'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 180 Hz pour un modèle à 208 Hz.** Cependant, nous pouvons observer que ce sujet aggrave sa voix pour tenter d'imiter le modèle de la voix transformée qui n'est pourtant pas très éloigné de la fréquence moyenne de sa voix spontanée.

L'observation du second histogramme montre que **la meilleure imitation est bien celle de la voix transformée.** En effet, nous constatons que l'étendue fréquentielle se réduit entre l'émission spontanée et l'imitation de la voix transformée (au niveau des fréquences maximales). Le sujet passe de 314 Hz à 255 Hz. Toutefois, nous constatons que l'étendue fréquentielle de l'imitation se déplace vers des fréquences inférieures. Le sujet passe de 146 Hz en voix spontanée à 126 Hz lors de l'imitation. Même si cette imitation reste la meilleure, le sujet s'éloigne de sa zone de sécurité prosodique dans les fréquences minimales, en tentant de s'approcher du modèle. La personne a sans doute bien perçu que l'écart prosodique de sa voix se réduisait dans le modèle de la voix transformée. Et en tentant de s'adapter au mieux, sa voix s'est ainsi décalée vers des fréquences inférieures. Soulignons également qu'il y a bien tentative d'imitation du modèle orthophonique car l'étendue fréquentielle varie. Mais cette imitation la contraint à sortir considérablement de la zone sécurisée.

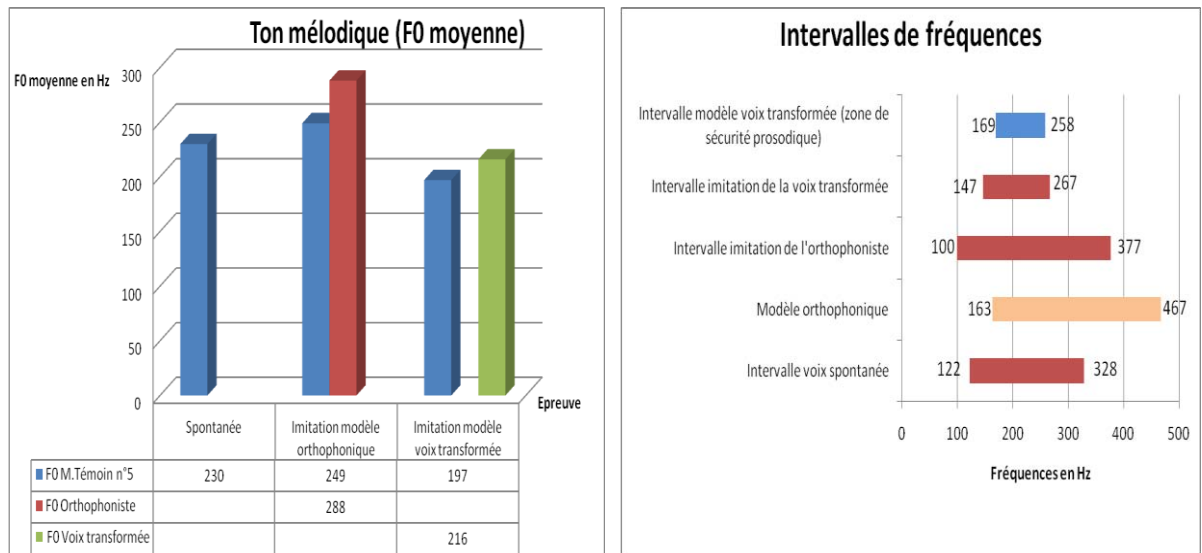
- Témoin n°4 :



Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 231 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 283 Hz, nous voyons que la fréquence moyenne du sujet est de 250 Hz. Il y a une adaptation au modèle orthophonique. Mais **l'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 210 Hz pour un modèle à 240 Hz**. Cependant, nous pouvons observer que ce sujet aggrave sa voix pour tenter d'imiter le modèle de la voix transformée qui n'est pourtant pas très éloigné de la fréquence moyenne de sa voix spontanée.

L'observation du second histogramme confirme qu'il y a bien **adaptation vocale dans l'imitation la voix transformée**. En effet, nous constatons que l'intervalle fréquentiel se réduit entre la production spontanée et cette imitation. Il passe de [148 Hz – 413 Hz] à [169 Hz – 269 Hz]. Cette imitation est la mieux réussie et elle reste la plus proche de la zone de sécurité prosodique. Soulignons également qu'il y a bien tentative d'imitation du modèle orthophonique car l'étendue fréquentielle varie. Mais cette imitation contraint le sujet à sortir considérablement de la zone sécurisée.

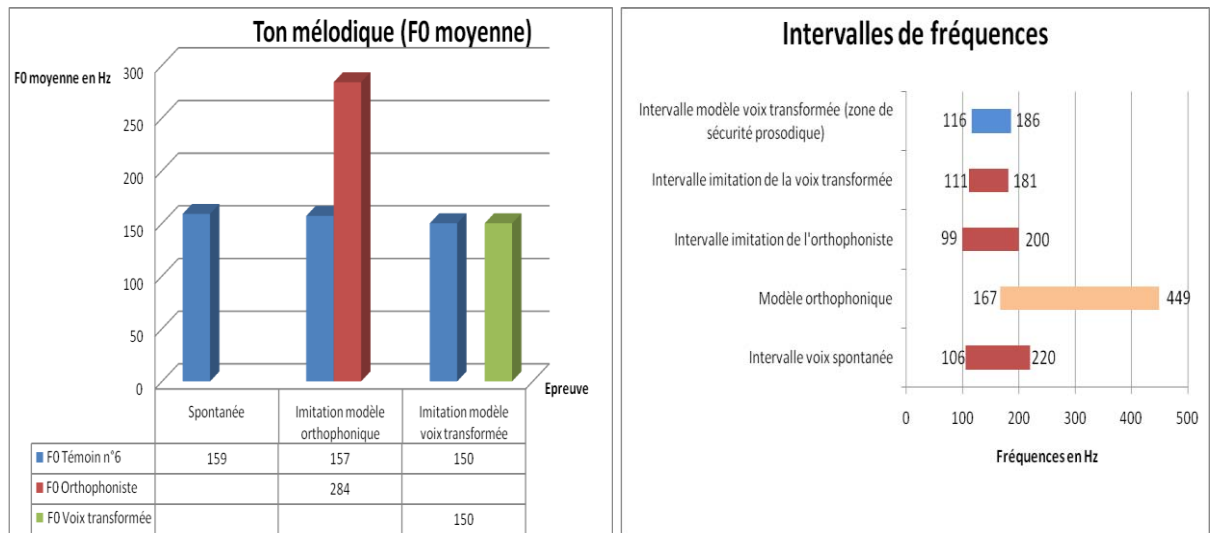
- Témoin n°5 :



Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 230 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 288 Hz, nous voyons que la fréquence moyenne du sujet est de 249 Hz. Il y a une adaptation au modèle orthophonique. Mais **l'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 197 Hz pour un modèle à 216 Hz.** Cependant, nous pouvons observer que ce sujet aggrave sa voix pour tenter d'imiter le modèle de la voix transformée qui n'est pourtant pas très éloigné de la fréquence moyenne de sa voix spontanée.

L'observation du second histogramme précise ces éléments. **L'adaptation vocale dans l'imitation de la voix transformée est confirmée.** En effet, nous constatons que l'étendue fréquentielle se réduit entre l'émission spontanée et l'imitation de la voix transformée notamment au niveau des fréquences maximales. Le sujet passe de 328 Hz à 267 Hz. Cela lui permet également d'osciller davantage dans sa zone de sécurité prosodique. Soulignons qu'il y a bien tentative d'imitation du modèle orthophonique car l'étendue fréquentielle varie. Mais cette imitation le contraint à sortir considérablement de la zone. Dans les fréquences supérieures, le sujet est à 377 Hz alors que la zone de sécurité prosodique ne dépasse pas 258 Hz.

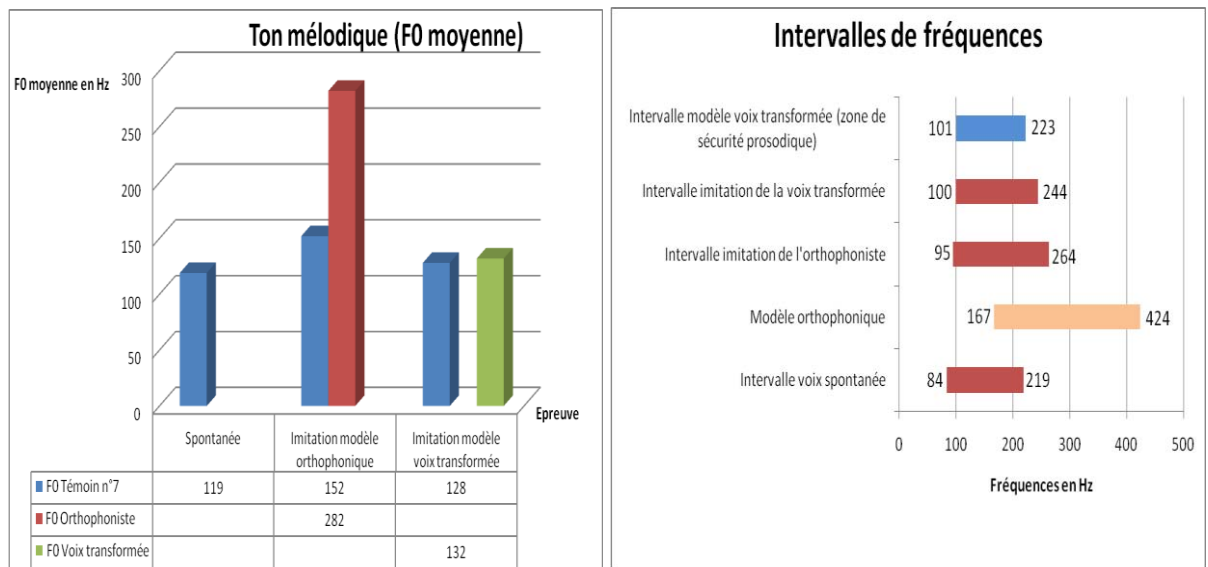
- Témoin n°6 :



Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 159 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 284 Hz, nous voyons que la fréquence moyenne du sujet est de 157 Hz. Il n'y a pas d'adaptation au modèle orthophonique. **L'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 150 Hz pour un modèle à 150 Hz.** A l'écoute, le sujet réalise une très bonne imitation d'un point de vue global. Cependant, nous pouvons observer que les fréquences moyennes sont très proches entre l'émission spontanée et le modèle de la voix transformée. Ainsi, ces premiers résultats ne nous permettent pas d'affirmer qu'il y a une réelle adaptation vocale au modèle puisque les fréquences moyennes sont très voisines.

L'observation du second histogramme nous permet de voir que **l'adaptation vocale dans l'imitation de la voix transformée est confirmée.** En effet, nous constatons que l'étendue fréquentielle se réduit entre l'émission spontanée et l'imitation de la voix transformée (au niveau des fréquences maximales notamment). Le sujet passe de 220 Hz à 186 Hz. Cela lui permet également de rester complètement dans sa zone de sécurité prosodique, à quelques hertz près. L'observation des différentes fréquences montre qu'il y a bien absence d'imitation du modèle orthophonique.

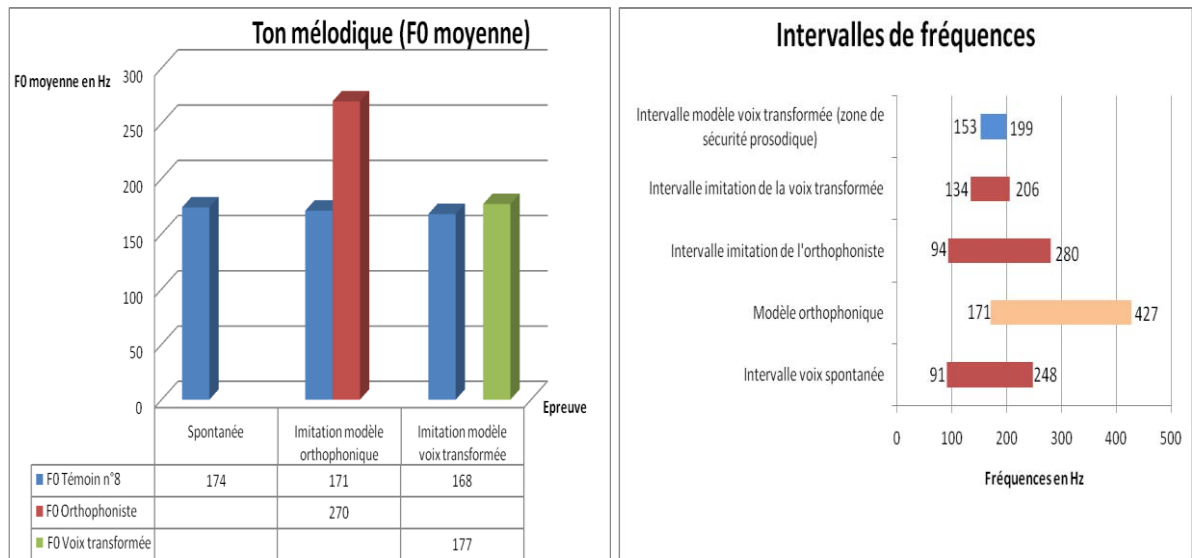
- Témoin n°7 :



Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 119 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 282 Hz, nous voyons que la fréquence moyenne du sujet est de 152 Hz. Il y a une adaptation au modèle orthophonique. Mais **l'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 128 Hz pour un modèle à 132 Hz.**

L'observation du second histogramme montre que **l'adaptation vocale dans l'imitation de la voix transformée est confirmée, particulièrement pour les fréquences minimales.** En effet, nous constatons que l'étendue fréquentielle se réduit entre l'émission spontanée et l'imitation de la voix transformée. Le sujet passe de 84 Hz à 100 Hz. Cela lui permet de se rapprocher davantage de sa zone de sécurité prosodique au niveau minimal. Il s'en éloigne quelque peu au niveau maximal en passant de 219 Hz en spontané à 244 Hz dans l'imitation de la voix transformée mais celle-ci reste la meilleure production. En effet, il y a bien tentative d'imitation du modèle orthophonique car l'étendue fréquentielle varie. Mais cette imitation est peu opportune, elle le contraint d'autant plus à sortir de la zone sécurisée puisque sa fréquence maximale est à 264 Hz alors que la zone sécurisée ne dépasse pas 223 Hz.

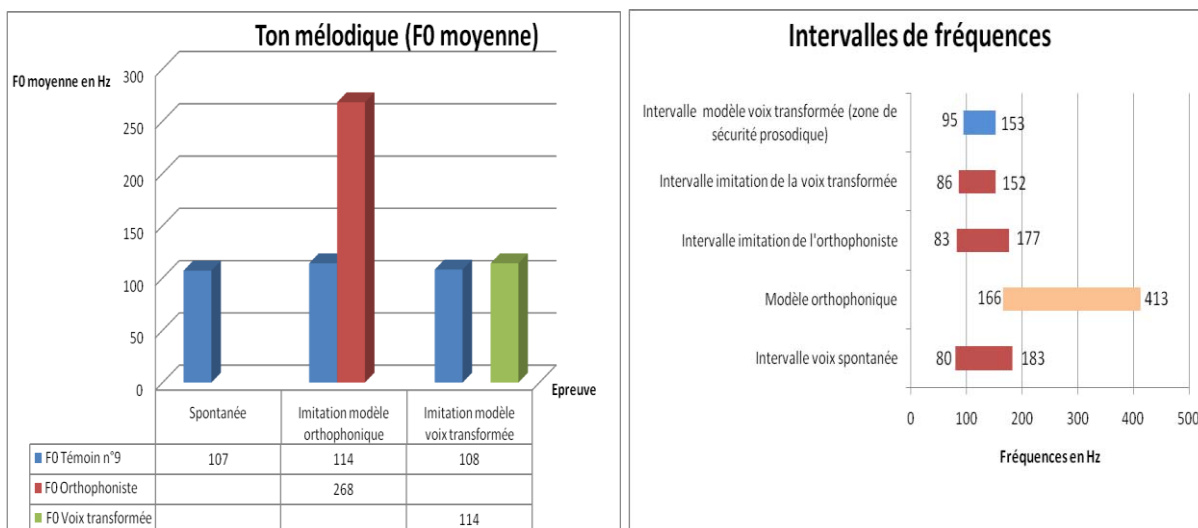
- Témoin n°8 :



Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 174 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 270 Hz, nous voyons que la fréquence moyenne du sujet est de 171 Hz. Il n'y a pas d'adaptation au modèle orthophonique. **L'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 168 Hz pour un modèle à 177 Hz.** Cependant, nous pouvons observer que les fréquences moyennes sont très proches entre l'émission spontanée et le modèle de la voix transformée. Ainsi, ces premiers résultats ne nous permettent pas d'affirmer qu'il y a une réelle adaptation vocale au modèle puisque les fréquences moyennes sont très voisines.

L'observation du second histogramme précise ces éléments. **L'adaptation vocale dans l'imitation de la voix transformée est confirmée.** En effet, nous constatons que l'étendue fréquentielle se réduit considérablement entre l'émission spontanée et l'imitation de la voix transformée. Le sujet passe d'un intervalle de [91 Hz - 248 Hz] en production spontanée à un intervalle de [134 Hz - 206 Hz]. Cela lui permet également de se rapprocher davantage de sa zone de sécurité prosodique. Soulignons qu'il y a bien tentative d'imitation du modèle orthophonique car l'étendue fréquentielle varie. Mais cette imitation le contraint à sortir de la zone sécurisée puisque l'intervalle fréquentiel de l'imitation est de [94 Hz – 280 Hz].

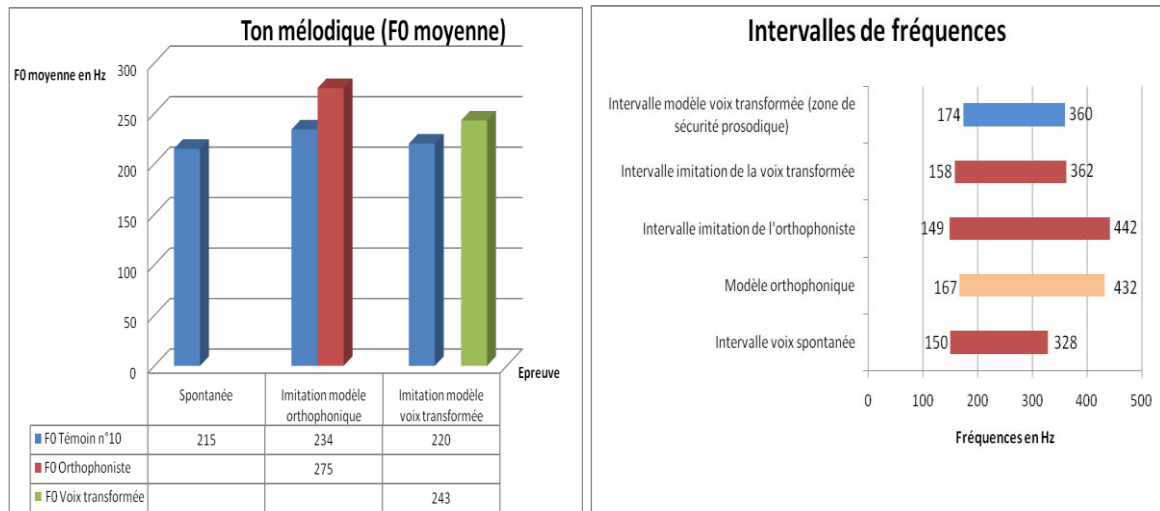
- Témoin n°9 :



Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 107 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 268 Hz, nous voyons que la fréquence moyenne du sujet est de 114 Hz. **L'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 108 Hz pour un modèle à 114 Hz.** Cependant, nous pouvons observer que les fréquences moyennes sont très proches entre l'émission spontanée et le modèle de la voix transformée. Ainsi, ces premiers résultats ne nous permettent pas d'affirmer qu'il y a une réelle adaptation vocale au modèle puisque les fréquences moyennes sont très voisines.

L'observation du second histogramme précise ces éléments. **L'adaptation vocale dans l'imitation de la voix transformée est confirmée.** En effet, nous constatons que l'étendue fréquentielle se réduit considérablement entre l'émission spontanée et l'imitation de la voix transformée. Le sujet passe d'un intervalle de [80 Hz - 183 Hz] en production spontanée à un intervalle de [86 Hz - 152 Hz]. Cela lui permet également de se rapprocher davantage de sa zone de sécurité prosodique. Soulignons qu'il n'y a pas d'imitation du modèle orthophonique puisque l'étendue fréquentielle est diminuée par rapport à la production spontanée, alors que le modèle présente des fréquences beaucoup plus hautes. Cette imitation le contraint à sortir de la zone sécurisée puisque l'intervalle fréquentiel de l'imitation est de [83 Hz - 177 Hz].

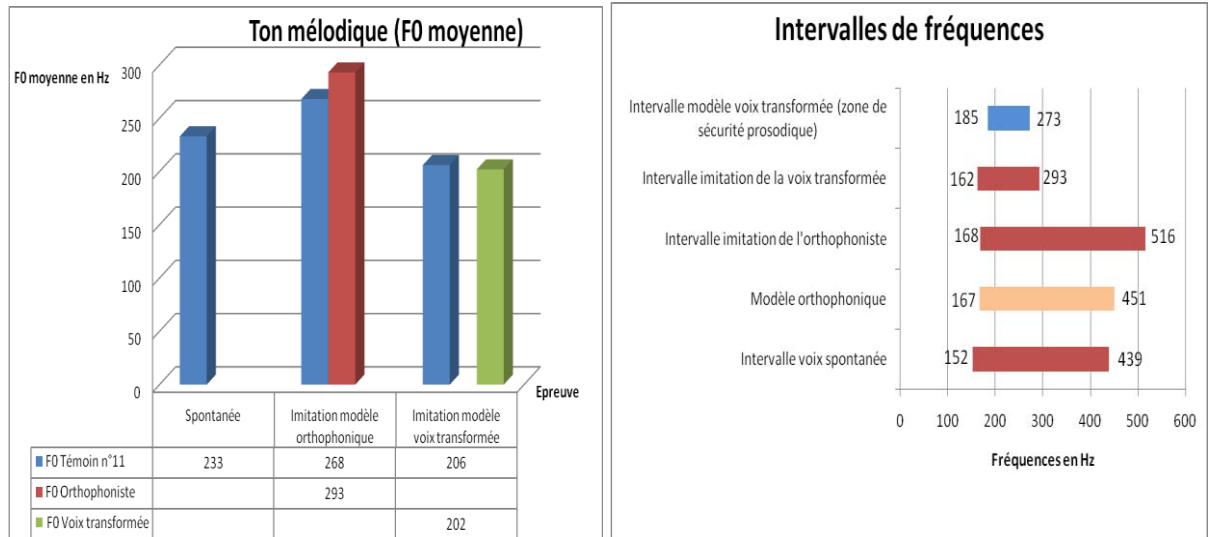
- Témoin°10 :



Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 215 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 275 Hz, nous voyons que la fréquence moyenne du sujet est de 234 Hz. Il y a une adaptation au modèle orthophonique. **L'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 220 Hz pour un modèle à 243 Hz.** Cependant, nous pouvons observer que les fréquences moyennes sont très proches entre l'émission spontanée et le modèle de la voix transformée. Ainsi, ces premiers résultats ne nous permettent pas d'affirmer qu'il y a une réelle adaptation vocale au modèle puisque les fréquences moyennes sont très voisines.

L'observation du second histogramme précise ces éléments. **L'adaptation vocale dans l'imitation de la voix transformée est confirmée.** En effet, nous constatons que l'étendue fréquentielle se déplace vers des fréquences supérieures entre l'émission spontanée et l'imitation de la voix transformée. Le sujet passe d'un intervalle de [150 Hz - 328 Hz] en production spontanée à un intervalle de [158 Hz - 362 Hz]. Il réalise donc une imitation proche du modèle. Cela lui permet également d'osciller davantage de sa zone de sécurité prosodique. Soulignons qu'il y a bien imitation du modèle orthophonique puisque l'étendue fréquentielle est augmentée par rapport à la production spontanée. L'étendue fréquentielle est très proche de celle du modèle orthophonique. Mais cette imitation contraint le sujet à sortir considérablement de la zone sécurisée puisque l'intervalle fréquentiel de l'imitation est de [149 Hz - 442 Hz].

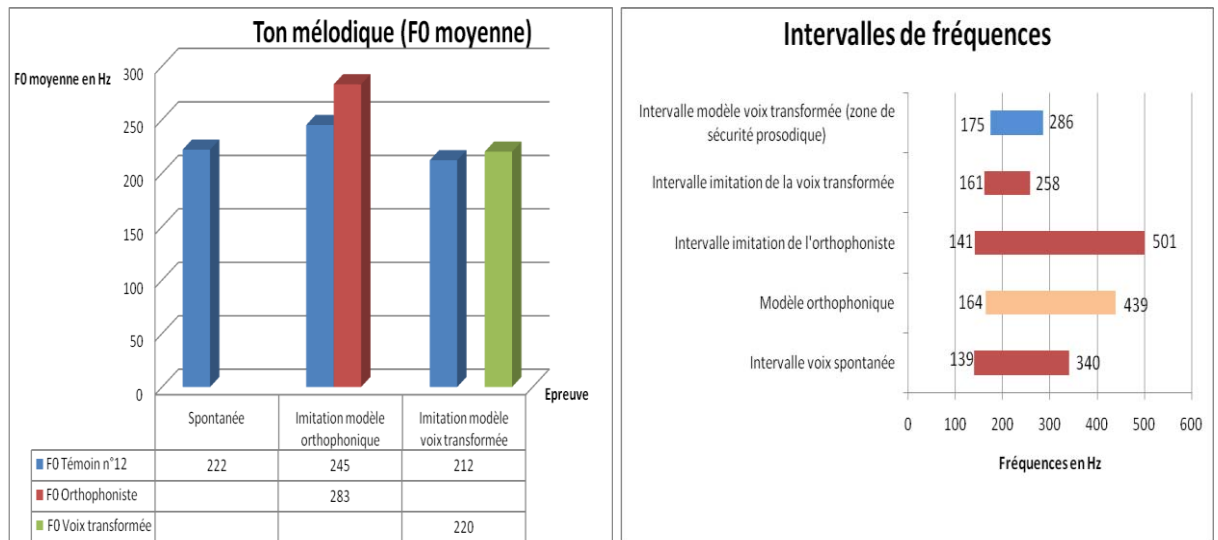
- Témoïn n°11 :



Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 233 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 293 Hz, nous voyons que la fréquence moyenne du sujet est de 268 Hz. Il y a une adaptation au modèle orthophonique. **L'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 202 Hz pour un modèle à 206 Hz.**

L'observation du second histogramme atteste qu'il y a bien **adaptation vocale dans l'imitation de la voix transformée**. En effet, nous constatons que l'étendue fréquentielle se réduit particulièrement entre l'émission spontanée et l'imitation de la voix transformée. Le sujet passe d'un intervalle de [152 Hz - 439 Hz] en production spontanée à un intervalle de [162 Hz - 293 Hz]. Il réalise donc une modulation importante de sa voix pour une imitation proche du modèle. Cela lui permet également de rester beaucoup plus proche de sa zone de sécurité prosodique. Soulignons qu'il y a bien imitation du modèle orthophonique puisque l'étendue fréquentielle est augmentée par rapport à la production spontanée. Mais cette imitation contraint le sujet à sortir considérablement de la zone sécurisée puisque l'intervalle fréquentiel de l'imitation est de [168 Hz – 516 Hz].

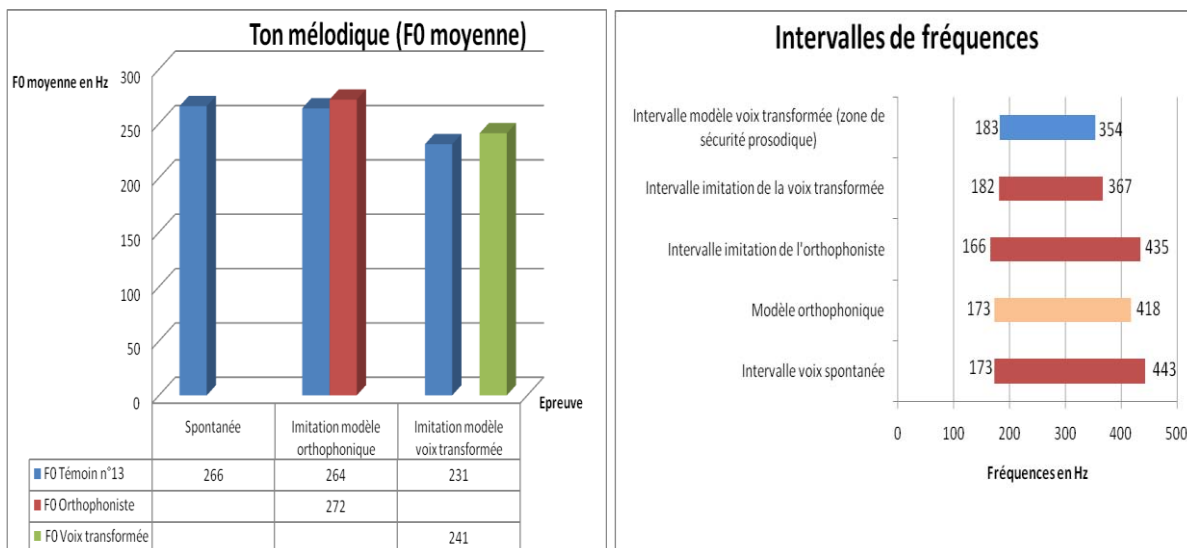
- Témoïn n°12 :



Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 222 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 283 Hz, nous voyons que la fréquence moyenne du sujet est de 245 Hz. Il y a une adaptation au modèle orthophonique. **L'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 212 Hz pour un modèle à 220 Hz.**

L'observation du second histogramme atteste qu'il y a bien **adaptation vocale dans l'imitation de la voix transformée**. En effet, nous constatons que l'étendue fréquentielle se réduit particulièrement entre l'émission spontanée et l'imitation de la voix transformée. Le sujet passe d'un intervalle de [139 Hz - 340 Hz] en production spontanée à un intervalle de [161 Hz - 258 Hz]. Il réalise donc une modulation importante de sa voix pour une imitation proche du modèle. Cela lui permet également de rester beaucoup plus proche de sa zone de sécurité prosodique. Soulignons qu'il y a bien imitation du modèle orthophonique puisque l'étendue fréquentielle est augmentée par rapport à la production spontanée. Mais cette imitation contraint le sujet, comme le sujet n°11, à sortir considérablement de la zone sécurisée puisque l'intervalle fréquentiel de l'imitation est de [141 Hz – 501 Hz].

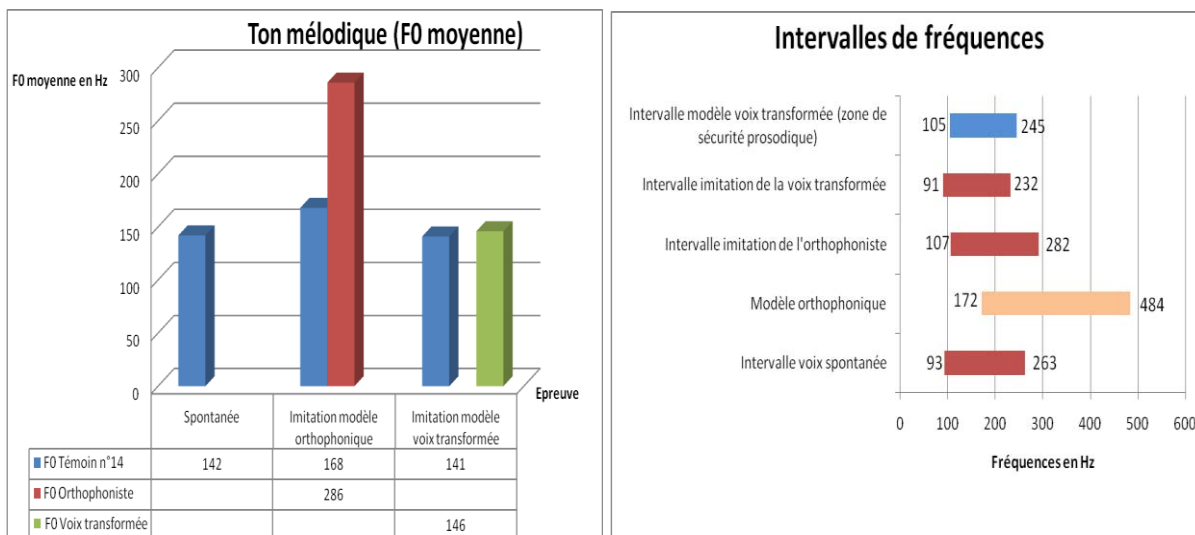
- Témoin n°13 :



Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 266 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 272 Hz, nous voyons que la fréquence moyenne du sujet est de 264 Hz. Il n'y a pas de réelle adaptation au modèle orthophonique. **L'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 231 Hz pour un modèle à 241 Hz.**

L'observation du second histogramme atteste qu'il y a bien **adaptation vocale dans l'imitation de la voix transformée**. En effet, nous constatons que le sujet parvient à réduire son étendue fréquentielle entre l'émission spontanée et l'imitation de la voix transformée. Le sujet passe d'un intervalle de [173 Hz - 443 Hz] en production spontanée à un intervalle de [182 Hz - 367 Hz]. Il réalise donc une modulation de sa voix pour se rapprocher du modèle. Cela lui permet également de rester plus proche de sa zone de sécurité prosodique. Soulignons que la modulation de la voix pour imiter le modèle orthophonique est quand même présente même si très légère, on note une légère diminution au niveau maximal. Cette imitation contraint le sujet à sortir de la zone sécurisée puisque l'intervalle fréquentiel de l'imitation est de [166 Hz – 435 Hz].

- Témoin n°14 :



Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 142 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 286 Hz, nous voyons que la fréquence moyenne du sujet est de 168 Hz. Il y a une adaptation au modèle orthophonique. **L'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 141 Hz pour un modèle à 146 Hz.** Mais la proximité des fréquences moyennes entre la production spontanée et l'imitation de la voix transformée ne permet pas d'affirmer qu'il y a une réelle adaptation de la voix.

L'observation du second histogramme atteste qu'il y a bien **adaptation vocale dans l'imitation de la voix transformée**. En effet, nous constatons que le sujet parvient à réduire son étendue fréquentielle entre l'émission spontanée et l'imitation de la voix transformée. Le sujet passe d'un intervalle de [93 Hz - 263 Hz] en production spontanée à un intervalle de [191 Hz - 232 Hz]. Il réalise donc une modulation de sa voix pour se rapprocher du modèle. Cela lui permet également d'osciller davantage dans sa zone de sécurité prosodique. Soulignons qu'il y a modulation de la voix pour imiter le modèle orthophonique mais elle est peu probante. De plus, cette imitation contraint le sujet à sortir de la zone sécurisée puisque l'intervalle fréquentiel de l'imitation est de [107 Hz – 282 Hz].

- **Synthèse des résultats du ton mélodique :**

Nous synthétisons, dans le tableau suivant, le nombre de témoins ayant adapté leur production au modèle orthophonique d'une part, puis au modèle de la voix transformée d'autre part. Nous nous basons surtout sur le deuxième histogramme pour chacun des sujets. En effet, il permet de voir qu'il y a bien modulation de la voix contrairement au premier histogramme qui reste plus global. Le reste du tableau met en évidence le nombre de sujets se rapprochant le plus de la zone de sécurité prosodique lors de l'imitation du modèle orthophonique et le nombre de sujets se rapprochant le plus de la zone de sécurité prosodique lors de l'imitation du modèle de la voix transformée.

Adaptation de la production spontanée au modèle orthophonique	Adaptation de la production spontanée au modèle de la voix transformée	Cas où l'imitation du modèle orthophonique est la meilleure et se rapproche le plus de la zone de sécurité prosodique	Cas où l'imitation du modèle de la voix transformée est la meilleure et se rapproche le plus de la zone de sécurité prosodique
10 / 14 témoins soit environ 71 % des témoins	14 / 14 témoins soit 100 % des témoins	0 / 14 témoins soit 0 % des témoins	14 / 14 témoins soit 100 % des témoins

Quatre témoins n'ont donc pas réussi à réaliser une réelle modulation de leur voix pour imiter le modèle orthophonique. Il s'agit des témoins n°3, 4, 6, 9. Nous retrouvons deux hommes et deux femmes. Nous avons observé que ces personnes parlent spontanément avec une voix plutôt monocorde ce qui pourrait expliquer la faible modulation de leur voix.

L'ensemble des témoins a su s'adapter au modèle de la voix transformée. L'imitation de la voix transformée est la meilleure pour tous les témoins en terme de proximité par rapport à la zone de sécurité prosodique.

Nous savons qu'un contrôle moteur et des ajustements musculaires au niveau laryngé sont possibles. En effet, nous avons vu qu'il existe un impact musculaire sur les cordes vocales au niveau antéropostérieur, transversal et vertical. Les résonateurs peuvent être également

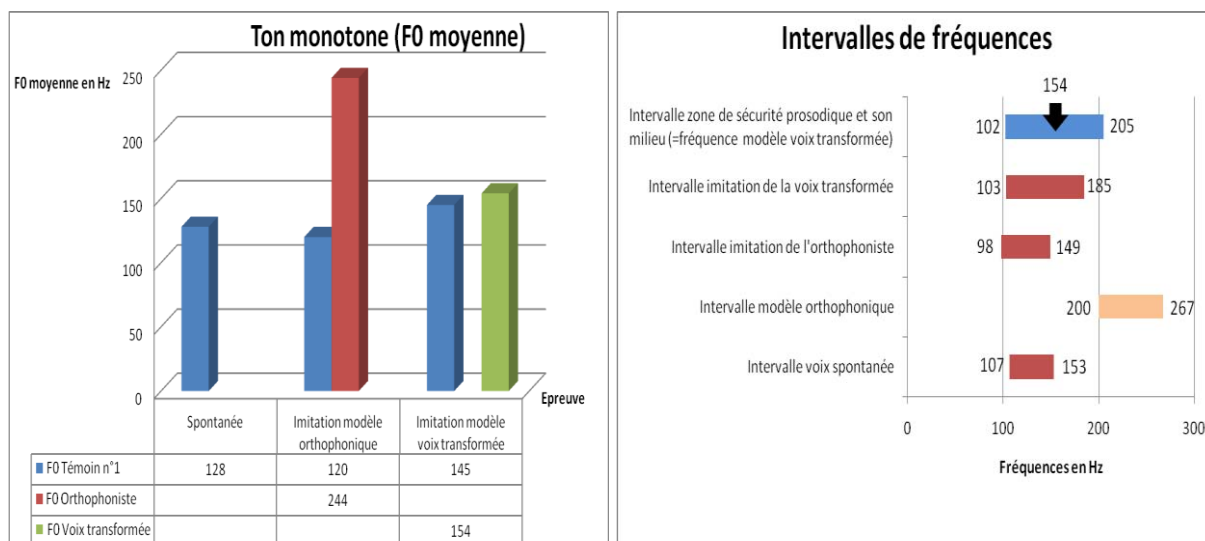
modulés dans leur forme et leur taille. Les résultats observés ici montrent que la plupart des témoins sont capables d'agir sur les structures intervenant dans la phonation pour tenter d'imiter le modèle proposé. Cela laisse penser que le sujet réalise grâce à la boucle audio-phonatoire des ajustements tout au long de sa production pour respecter au mieux le modèle donné. Cependant, seul le modèle de la voix transformée induit systématiquement des ajustements adaptés pour que la production du sujet ne dépasse pas la zone de la dynamique vocale où la voix est en sécurité. Ainsi, les ajustements et le feed-back auditif semblent être plus performants lors de l'imitation de la voix transformée puisque les performances d'adaptation au modèle sont les meilleures.

2.2.1.2 Ton monotone :

Comme pour le ton mélodique, dans un premier temps, nous illustrons à l'aide d'un histogramme, les résultats de la fréquence fondamentale moyenne obtenus pour le témoin et chacun des modèles en fonction des différentes épreuves. Cela nous permet notamment de comparer la fréquence moyenne entre le témoin et le modèle orthophonique puis entre le témoin et le modèle de sa voix transformée (là où la voix se situe dans la zone de sécurité prosodique).

Dans un second temps, nous illustrons à l'aide d'un histogramme placé à côté, les intervalles de fréquences dans lesquels se situe le témoin selon chaque épreuve. L'intérêt de l'épreuve sur ton monotone est aussi de voir comment l'individu peut osciller autour du point central (milieu) de la zone de sécurité prosodique. Nous indiquons donc le point central de l'intervalle de la zone de sécurité prosodique, point autour duquel la voix du sujet doit tenter de se rapprocher (et qui n'est autre que la fréquence moyenne calculée pour la ligne monotone de la voix transformée).

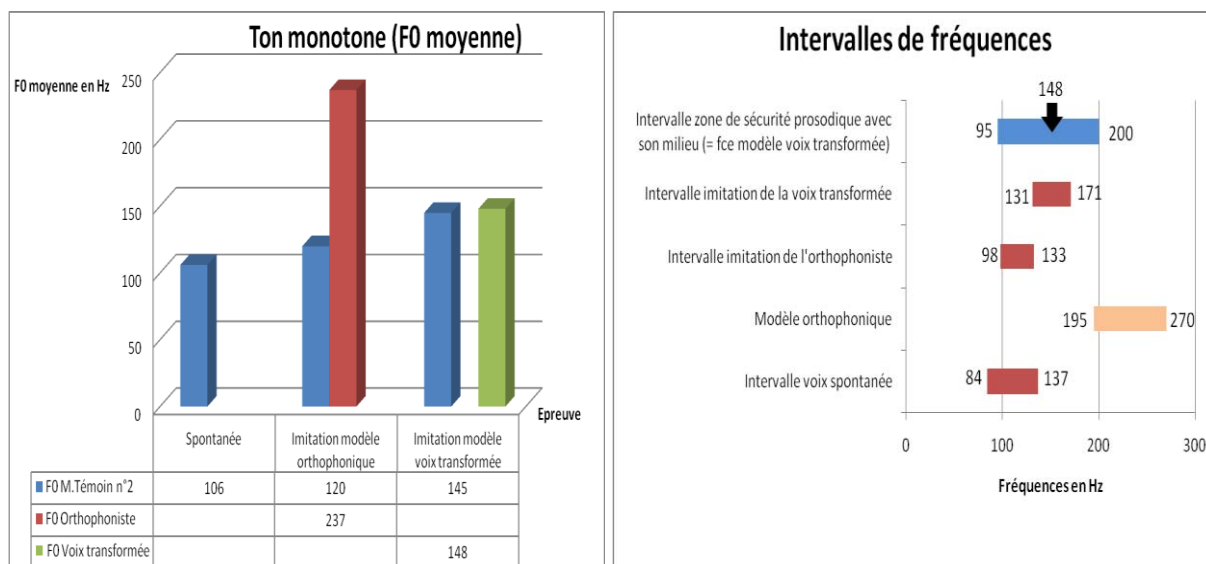
- Témoin n°1 :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 128 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 244 Hz, nous voyons que la fréquence moyenne du sujet est de 120 Hz. Il reste donc plutôt proche de sa fréquence moyenne spontanée et ne se rapproche pas du modèle orthophonique. **L'imitation la mieux réussie en terme de fréquence moyenne est celle de la voix transformée : le modèle est à 154 Hz et le sujet réalise une imitation à 145 Hz.** Mais compte tenu de la proximité des fréquences moyennes observées entre la production et l'imitation de la voix transformée, nous ne pouvons affirmer qu'il y a bien adaptation.

Le second histogramme nous montre effectivement qu'il y a bien **adaptation de la voix du sujet pour imiter le modèle de la voix transformée puisque son étendue fréquentielle augmente au niveau maximum : il passe de 153 Hz en voix spontanée à 185 Hz lors de l'imitation de ce modèle.** L'étendue fréquentielle lors de cette imitation s'équilibre autour du point central à 154 Hz de la zone de sécurité prosodique ce qui témoigne de l'adaptation positive. Notons que dans le reste des épreuves le sujet reste dans sa zone de sécurité prosodique quoique dans la portion plutôt inférieure de cette zone. D'autre part, il n'y a pas de réelle modulation de la voix lors de l'imitation du modèle orthophonique.

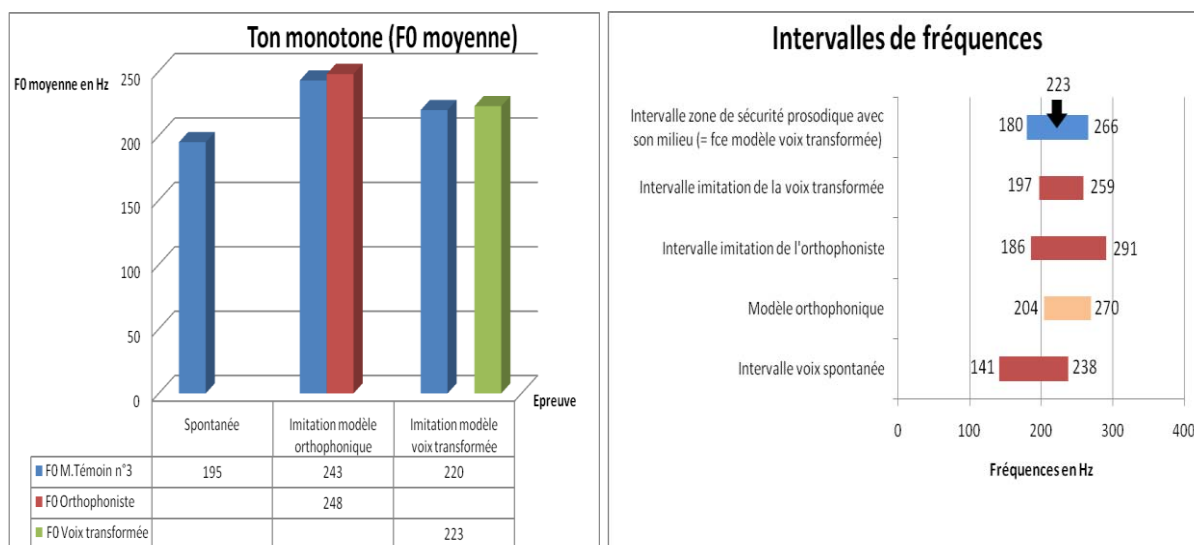
- Témoin n°2 :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 106 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 237 Hz, nous voyons que la fréquence moyenne du sujet est de 120 Hz. D'un point de vue global, il y a donc une légère adaptation au modèle orthophonique. **Mais l'imitation la mieux réussie en terme de fréquence moyenne est celle de la voix transformée : le modèle est à 148 Hz et le sujet réalise une imitation à 145 Hz.** L'imitation réalisée est très proche du modèle et le sujet s'adapte parfaitement à l'élévation tonale imposée par le modèle.

Le second histogramme nous montre effectivement qu'il y a bien **une réelle imitation du modèle de la voix transformée puisque le sujet passe d'un intervalle de [84 Hz – 137 Hz] en voix spontanée à un intervalle de [131 Hz – 171 Hz] lors de l'imitation de ce modèle.** L'étendue fréquentielle lors de cette imitation s'équilibre autour du point central à 148 Hz de la zone de sécurité prosodique ce qui témoigne de l'adaptation positive. Notons que dans le reste des épreuves le sujet reste globalement dans sa zone de sécurité prosodique quoique dans la portion très inférieure de cette zone. D'autre part, il n'y a pas de réelle modulation de la voix lors de l'imitation du modèle orthophonique.

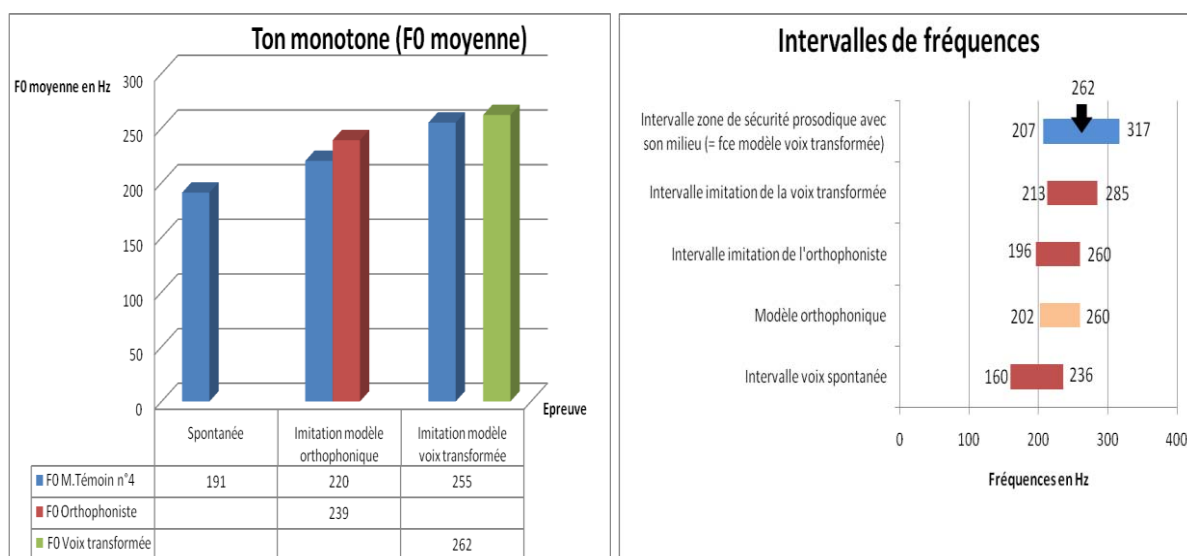
- Témoïn n°3 :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 195 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 248 Hz, nous voyons que la fréquence moyenne du sujet est de 243 Hz. Il y a donc une réelle adaptation au modèle orthophonique. Mais **l'imitation la mieux réussie en terme de fréquence moyenne est celle de la voix transformée : le modèle est à 223 Hz et le sujet réalise une imitation à 220 Hz.**

Le second histogramme atteste que **l'imitation du modèle de la voix transformée est la meilleure**. D'une part, il y a une adaptation de la voix par rapport à la production spontanée (le sujet passe d'un intervalle [141 Hz - 238 Hz] à un intervalle de [197 Hz – 259 Hz]), d'autre part, cette adaptation reste pleinement dans la zone de sécurité prosodique et s'équilibre autour du point central à 223 Hz, ce qui témoigne de l'adaptation positive. Contrairement au reste des épreuves où le sujet sort de sa zone de sécurité. Ainsi, même s'il y a une imitation assez proche du modèle orthophonique, celle-ci reste moins probante et plus éloignée de la zone de sécurité prosodique que l'imitation de la voix transformée.

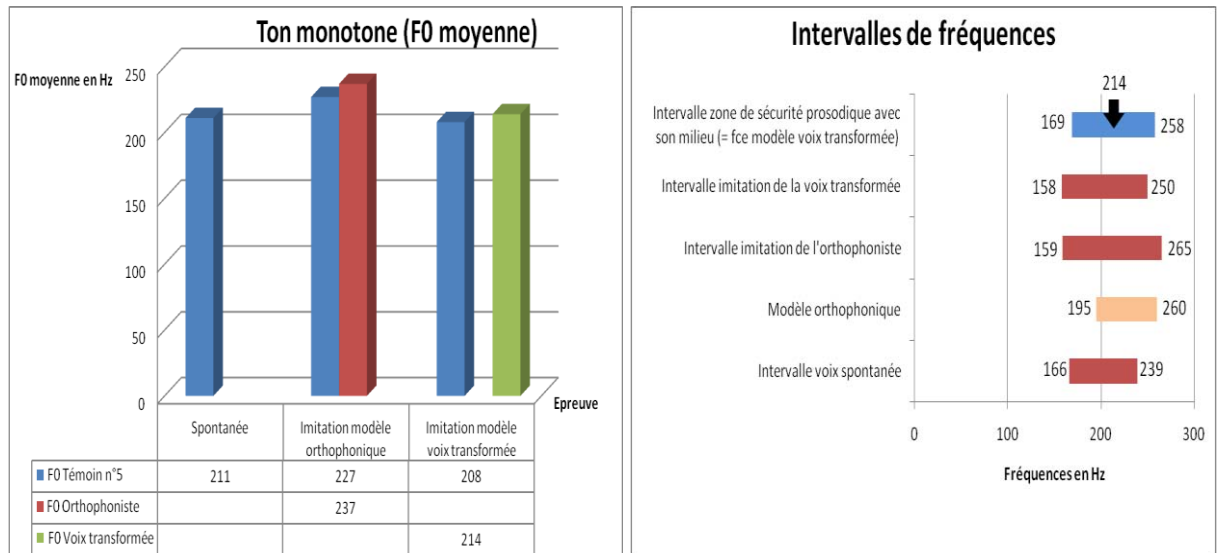
- Témoïn n°4 :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 191 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 239 Hz, nous voyons que la fréquence moyenne du sujet est de 220 Hz. Il y a une adaptation au modèle orthophonique. Mais **l'imitation la mieux réussie en terme de fréquence moyenne est celle de la voix transformée : le modèle est à 262 Hz et le sujet réalise une imitation à 255 Hz.**

Le second histogramme confirme qu'il y a bien **adaptation de la voix du sujet pour imiter le modèle de la voix transformée**. En effet, l'étendue fréquentielle passe d'un intervalle de [160 Hz – 236 Hz] en voix spontanée à un intervalle de [213- Hz - 285 Hz] lors de l'imitation de la voix transformée. Elle s'équilibre assez bien autour du point central à 262 Hz de la zone de sécurité prosodique ce qui témoigne de l'adaptation positive. Notons que l'imitation du modèle orthophonique est réussie et qu'elle s'inscrit également dans la zone de sécurité prosodique mais de façon moins équilibrée que l'imitation de la voix transformée.

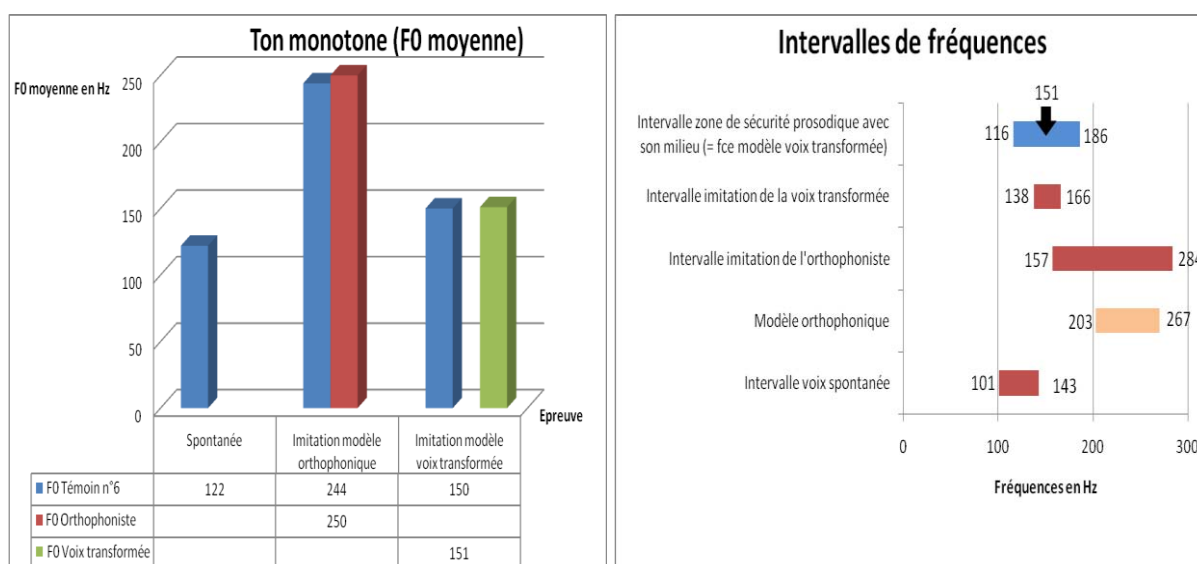
- Témoïn n°5 :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 211 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 237 Hz, nous voyons que la fréquence moyenne du sujet est de 227 Hz. Il y a une adaptation au modèle orthophonique. **L'imitation la mieux réussie en terme de fréquence moyenne est celle de la voix transformée : le modèle est à 214 Hz et le sujet réalise une imitation à 208 Hz.** Mais compte tenu de la très grande proximité des fréquences moyennes observées entre la production et l'imitation de la voix transformée, et pour l'imitation du modèle orthophonique, nous ne pouvons pas affirmer qu'il y a bien adaptation.

Le second histogramme nous montre qu'il y a peu d'adaptations. Le sujet parle spontanément dans une étendue fréquentielle presque totalement dans la zone de sécurité prosodique. Ainsi, les adaptations nécessaires sont très faibles. Ce que nous pouvons constater, c'est que l'imitation du modèle orthophonique est la moins réussie et s'éloigne de la zone de sécurité prosodique.

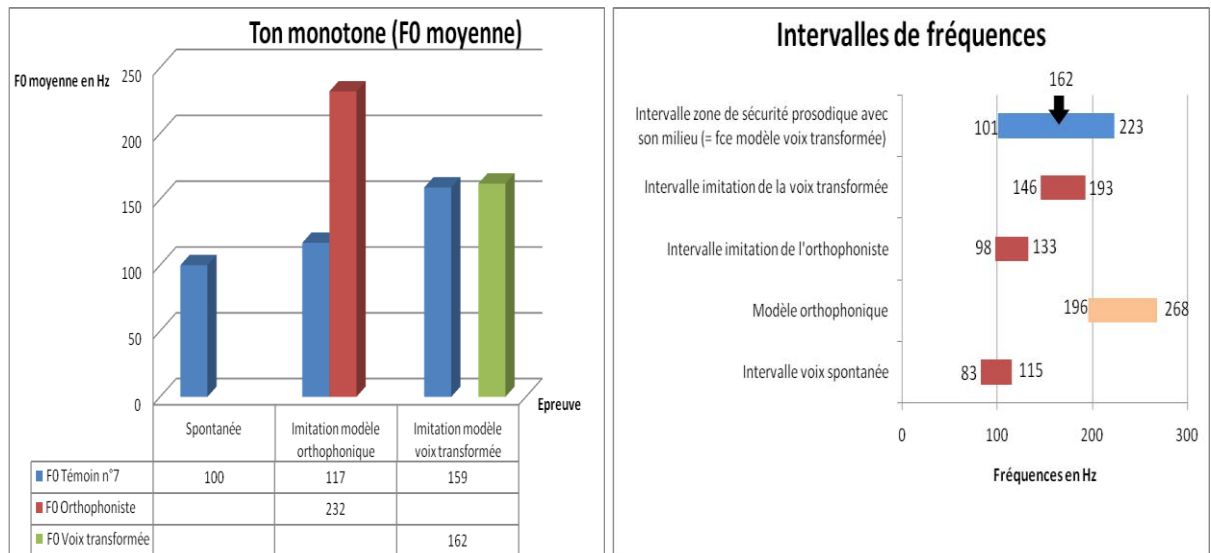
- Témoin n°6 :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 122 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 250 Hz, nous voyons que la fréquence moyenne du sujet est de 244 Hz. Il y a donc forte adaptation au modèle orthophonique. **Mais l'imitation la mieux réussie en terme de fréquence moyenne est celle de la voix transformée : le modèle est à 151 Hz et le sujet réalise une imitation à 150 Hz.**

Le second histogramme nous montre effectivement qu'il y a bien **adaptation de la voix du sujet pour imiter le modèle de la voix transformée**. L'intervalle fréquentiel passe de [101 Hz – 143 Hz] en production spontanée à un intervalle de [138 Hz- 166 Hz] dans l'imitation de la voix transformée. L'étendue fréquentielle lors de cette imitation s'équilibre de façon très précise autour du point central à 151 Hz de la zone de sécurité prosodique ce qui témoigne de l'adaptation très positive. Notons que pour le modèle orthophonique, le sujet tente effectivement de l'imiter mais sa production sort énormément de la zone sécurisée.

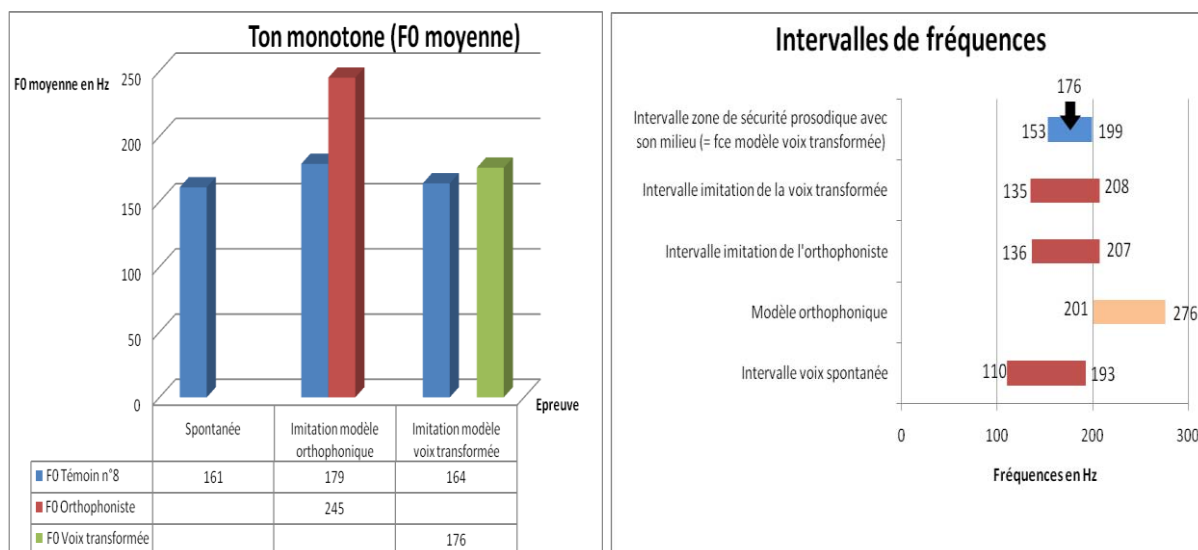
- Témoin n°7 :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 128 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 244 Hz, nous voyons que la fréquence moyenne du sujet est de 120 Hz. Il reste donc plutôt proche de sa fréquence moyenne spontanée et ne se rapproche pas du modèle orthophonique. **L'imitation la mieux réussie en terme de fréquence moyenne est celle de la voix transformée : le modèle est à 154 Hz et le sujet réalise une imitation à 145 Hz.** Mais compte tenu de la proximité des fréquences moyennes observées entre la production et l'imitation de la voix transformée, nous ne pouvons affirmer qu'il y a bien adaptation.

Le second histogramme nous montre effectivement qu'il y a bien **adaptation de la voix du sujet pour imiter le modèle de la voix transformée puisque son étendue fréquentielle augmente au niveau maximum : il passe de 153 Hz en voix spontanée à 185 Hz lors de l'imitation.** L'étendue fréquentielle lors de cette imitation s'équilibre autour du point central à 154 Hz de la zone de sécurité prosodique ce qui témoigne de l'adaptation positive. Notons que dans le reste des épreuves le sujet reste dans sa zone de sécurité prosodique quoique dans la portion plutôt inférieure de cette zone. D'autre part, il n'y a pas de réelle modulation de la voix lors de l'imitation du modèle orthophonique.

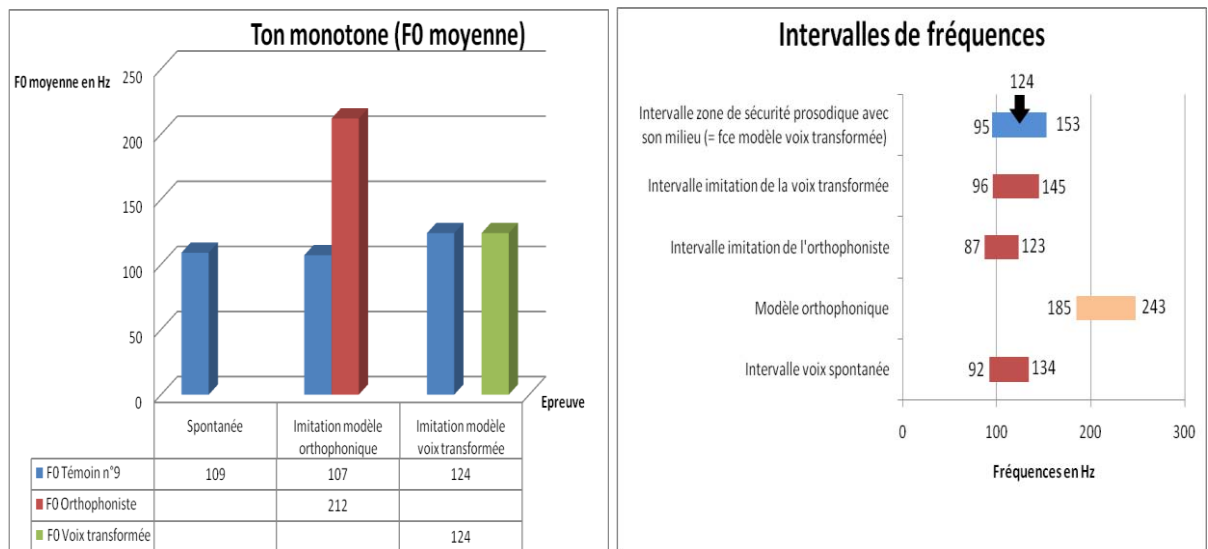
- Témoin N°8 :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 161 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 245 Hz, nous voyons que la fréquence moyenne du sujet est de 179 Hz. Il y a une tentative d'adaptation au modèle orthophonique. **L'imitation la mieux réussie en terme de fréquence moyenne est celle de la voix transformée : le modèle est à 176 Hz et le sujet réalise une imitation à 176 Hz.** Mais compte tenu de la proximité des fréquences moyennes observées entre la production et l'imitation de la voix transformée, nous ne pouvons affirmer qu'il y a bien adaptation.

Le second histogramme nous montre effectivement qu'il y a bien **adaptation de la voix du sujet pour imiter le modèle de la voix transformée puisque son étendue fréquentielle passe de [110 Hz – 193 Hz] en voix spontanée à [135 Hz – 208 Hz] lors de l'imitation.** L'étendue fréquentielle lors de cette imitation s'équilibre davantage autour du point central à 176 Hz de la zone de sécurité prosodique que lors de la production spontanée, ce qui témoigne de l'adaptation positive. D'autre part, il y a modulation de la voix lors de l'imitation du modèle orthophonique. L'étendue fréquentielle est similaire à celle de l'imitation de la voix transformée. Le sujet est resté dans le même champ de fréquences, il a tenté de moduler sa voix mais il n'a pas réellement réussi à l'adapter en fonction de l'épreuve.

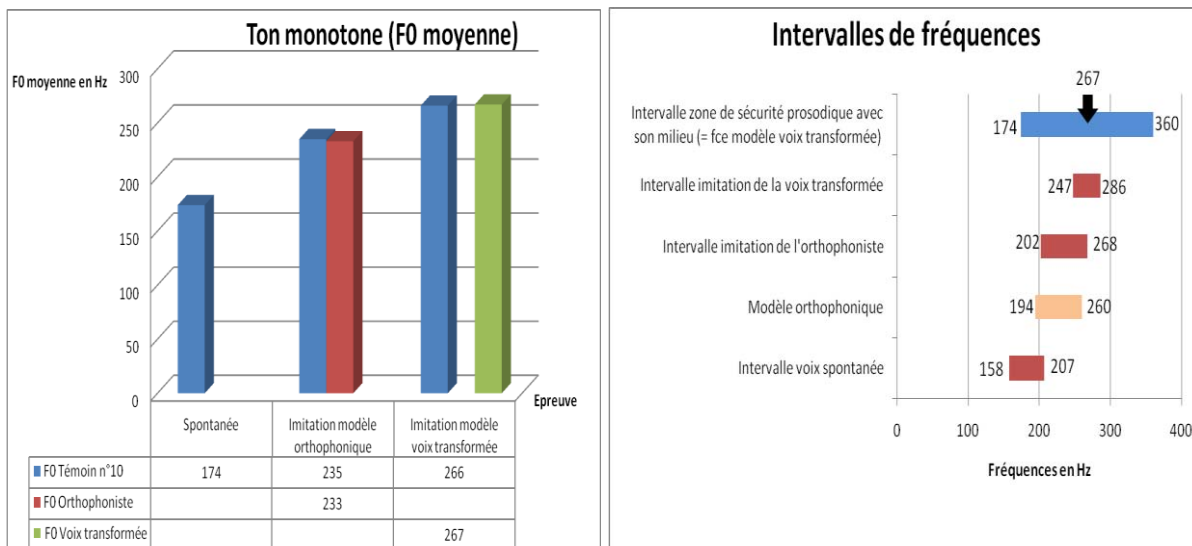
- Témoin n°9 :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 109 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 212 Hz, nous voyons que la fréquence moyenne du sujet est de 107 Hz. Il n'y a pas d'adaptation au modèle orthophonique. **L'imitation la mieux réussie en terme de fréquence moyenne est celle de la voix transformée : le modèle est à 124 Hz et le sujet réalise une imitation à 124 Hz.** Il y a bien une adaptation de la voix.

Le second histogramme montre **que le sujet tente bien d'imiter le modèle de la voix transformée puisque son étendue fréquentielle passe de [92 Hz – 134 Hz] en voix spontanée à [96 Hz – 145 Hz] lors de l'imitation.** Même si en terme de fréquences, les modulations paraissent faibles, l'écoute des productions témoigne des variations. L'étendue fréquentielle lors de cette imitation s'équilibre davantage autour du point central à 124 Hz de la zone de sécurité prosodique que lors de la production spontanée, ce qui témoigne de l'adaptation positive. Concernant le modèle orthophonique, l'imitation est peu probante puisque la voix du sujet baisse en fréquences [87 Hz – 123 Hz] alors que le modèle présente des fréquences plus élevées que sa production spontanée. Du reste, le sujet sort davantage du champ de fréquences de la zone de sécurité prosodique.

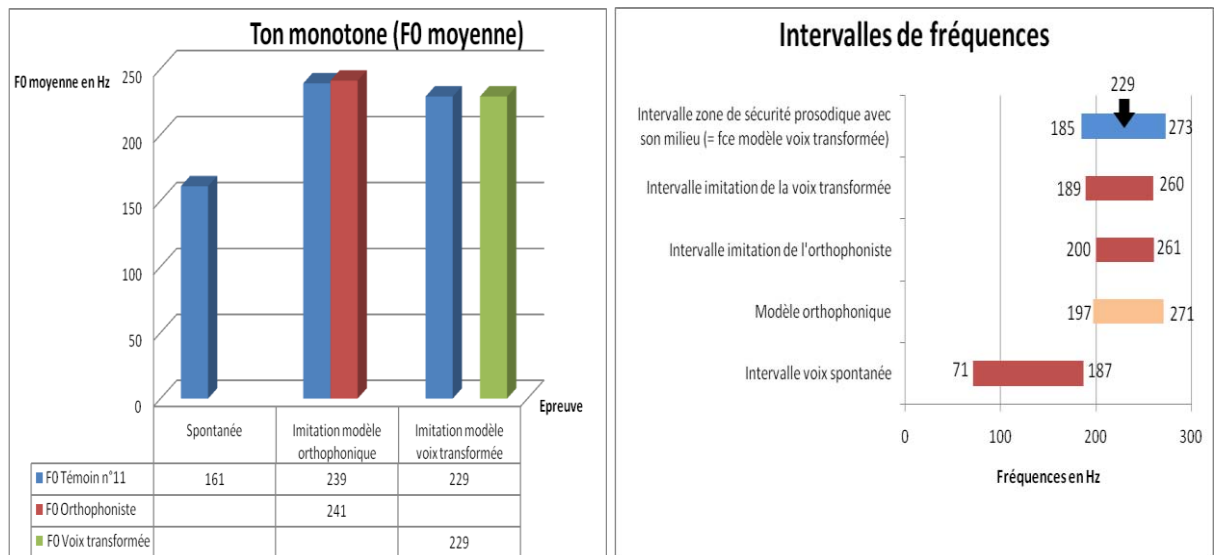
- Témoign n°10 :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 174 Hz. A l'écoute, le sujet prend une voix assez grave par rapport à sa tessiture. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 233 Hz, nous voyons que la fréquence moyenne du sujet est de 235 Hz. Il y a une importante adaptation au modèle orthophonique. **L'imitation la mieux réussie en terme de fréquence moyenne reste celle de la voix transformée : le modèle est à 266 Hz et le sujet réalise une imitation à 267 Hz.** Il y a une grande adaptation de la voix.

Le second histogramme montre **que le sujet imite bien le modèle de la voix transformée puisque son étendue fréquentielle passe de [158 Hz – 207 Hz] en voix spontanée à [247 Hz – 286 Hz] lors de l'imitation.** La modulation de la voix est extrêmement importante ici. L'étendue fréquentielle lors de cette imitation s'équilibre parfaitement autour du point central à 267 Hz de la zone de sécurité prosodique que lors de la production spontanée, ce qui témoigne de l'adaptation plus que positive. Concernant le modèle orthophonique, l'imitation est également très probante puisque l'étendue fréquentielle du sujet passe à [202 Hz – 268 Hz]. Mais, bien que le sujet reste dans le champ de fréquences de la zone de sécurité prosodique, il s'éloigne du point central recherché.

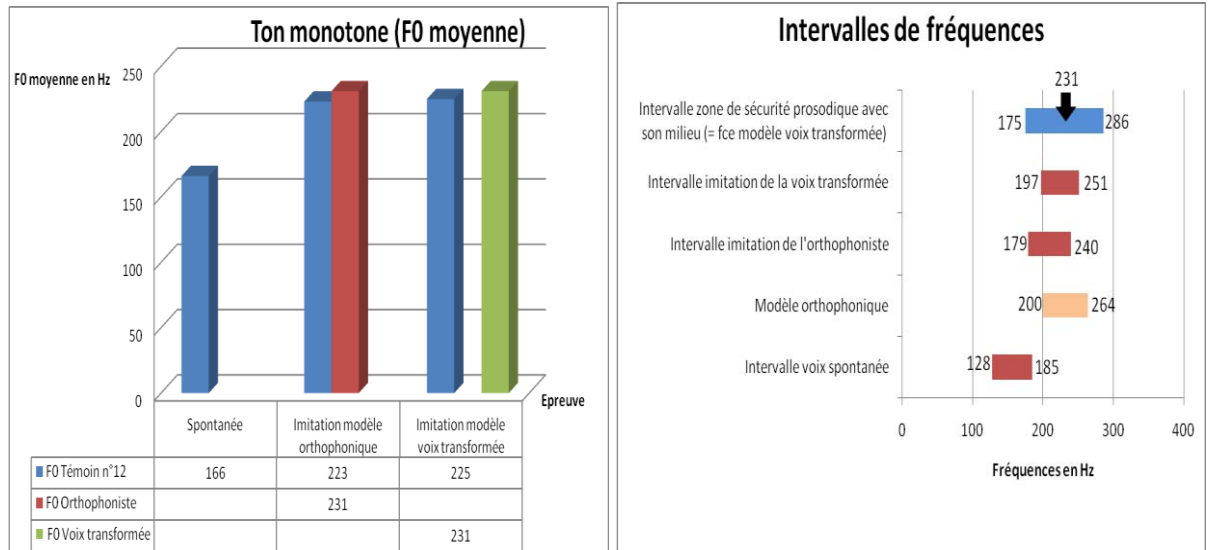
- Témoïn n°11 :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 161 Hz. A l'écoute, le sujet prend une voix assez grave par rapport à sa tessiture. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 241 Hz, nous voyons que la fréquence moyenne du sujet est de 239 Hz. Il y a une importante adaptation au modèle orthophonique. **L'imitation la mieux réussie en terme de fréquence moyenne reste celle de la voix transformée : le modèle est à 229 Hz et le sujet réalise une imitation à 229 Hz.** Il y a une grande adaptation de la voix.

Le second histogramme montre que le sujet imite bien le modèle de la voix transformée puisque son étendue fréquentielle passe de [71 Hz – 187 Hz] en voix spontanée à [189 Hz – 260 Hz] lors de l'imitation. La modulation de la voix est extrêmement importante ici. L'étendue fréquentielle lors de cette imitation s'équilibre particulièrement autour du point central à 229 Hz de la zone de sécurité prosodique que lors de la production spontanée, ce qui témoigne de l'adaptation plus que positive. **Mais l'imitation la mieux réussie à ce niveau est celle du modèle orthophonique, l'imitation est très probante puisque l'étendue fréquentielle du sujet passe à [200 Hz – 261 Hz] et s'équilibre davantage autour du point central.** Le sujet reste bien dans le champ de fréquences de la zone de sécurité prosodique.

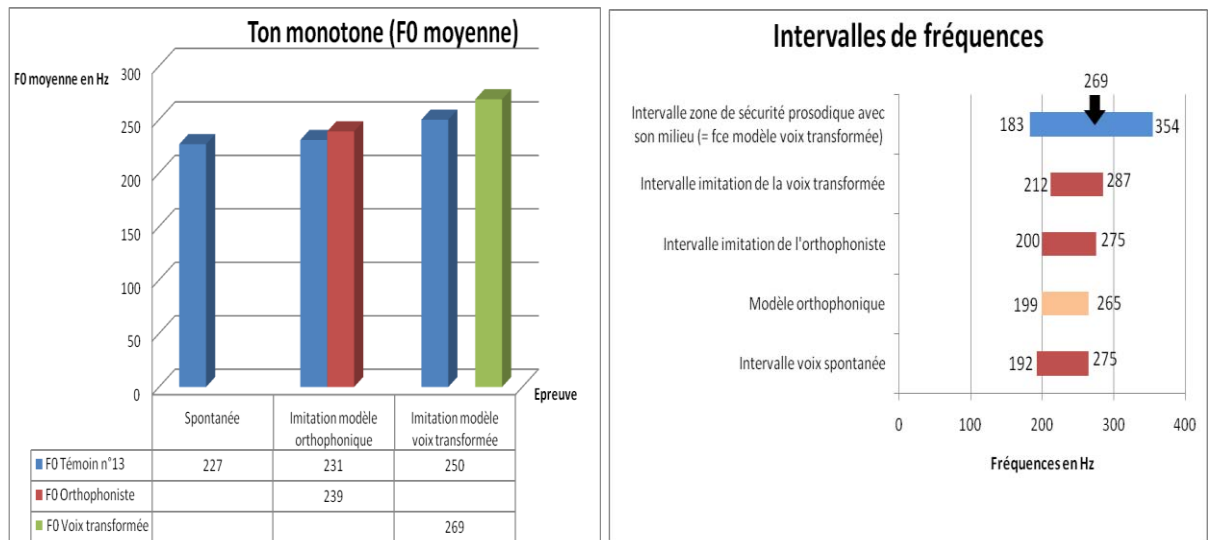
- Témoin n°12 :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 166 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 231 Hz, nous voyons que la fréquence moyenne du sujet est de 224 Hz. Il y a une importante adaptation au modèle orthophonique. **L'imitation la mieux réussie en terme de fréquence moyenne reste celle de la voix transformée : le modèle est à 231 Hz et le sujet réalise une imitation à 225 Hz.** Il y a une grande adaptation de la voix.

Le second histogramme montre **que le sujet imite bien le modèle de la voix transformée puisque son étendue fréquentielle passe de [128 Hz – 185 Hz] en voix spontanée à [197 Hz – 251 Hz] lors de l'imitation.** La modulation de la voix est extrêmement importante ici. L'étendue fréquentielle lors de cette imitation s'équilibre particulièrement autour du point central à 231 Hz de la zone de sécurité prosodique que lors de la production spontanée, ce qui témoigne de l'adaptation plus que positive. Concernant le modèle orthophonique, il y a bien modulation de la voix puisque l'étendue fréquentielle du sujet passe à [179 Hz – 240 Hz] mais l'imitation est plus éloignée du modèle donné. Par ailleurs, le sujet reste dans le champ de fréquences de la zone de sécurité prosodique, mais son étendue est moins équilibrée autour du point central recherché.

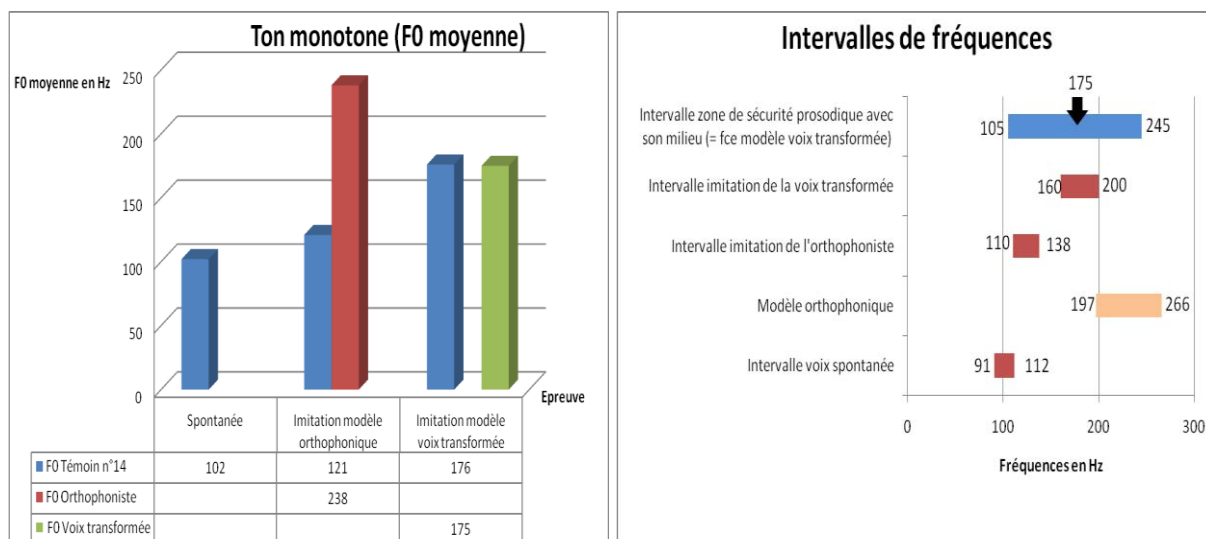
- Témoin n°13 :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 227 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 239 Hz, nous voyons que la fréquence moyenne du sujet est de 231 Hz. Il y a une légère adaptation au modèle orthophonique. **L'imitation la mieux réussie est celle-ci en terme de fréquence moyenne.** Pour l'imitation de la voix transformée : le modèle est à 269 Hz et le sujet réalise une imitation à 250 Hz. On note une légère adaptation de la voix.

Cependant, le second histogramme montre que **l'imitation la mieux réussie du point de vue de l'étendue fréquentielle est l'imitation de la voix transformée.** L'étendue fréquentielle passe de [192 Hz – 275 Hz] en voix spontanée à [212 Hz – 287 Hz] lors de l'imitation. Elle se rapproche du point central de la zone de sécurité prosodique à 269 Hz. Nous voyons cependant que le sujet présente une étendue fréquentielle proche de celle du modèle orthophonique, et il se situe déjà dans sa zone de sécurité prosodique pour parler en production spontanée. Cela explique les faibles modulations nécessaires et présentes. L'imitation du modèle de la voix production entraîne quand même une modulation de la voix.

- Témoign n°14 :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 102 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 238 Hz, nous voyons que la fréquence moyenne du sujet est de 176 Hz. Il y a une importante adaptation au modèle orthophonique. **L'imitation la mieux réussie en terme de fréquence moyenne est celle de l'imitation de la voix transformée** : le modèle est à 175 Hz et le sujet réalise une imitation à 176 Hz. On note une grande adaptation de la voix.

Le second histogramme montre que **l'imitation la mieux réussie du point de vue de l'étendue fréquentielle est bien l'imitation de la voix transformée**. L'étendue fréquentielle passe de [91 Hz – 112 Hz] en voix spontanée à [160 Hz – 200 Hz] lors de l'imitation. La modulation de la voix est extrêmement importante ici. L'étendue fréquentielle lors de cette imitation s'équilibre particulièrement autour du point central à 175 Hz de la zone de sécurité prosodique que lors de la production spontanée, ce qui témoigne de l'adaptation plus que positive. Concernant le modèle orthophonique, il y a bien une modulation de la voix puisque l'étendue fréquentielle du sujet passe à [110 Hz – 138 Hz] mais l'imitation est éloignée du modèle donné. Par ailleurs, le sujet reste dans le champ de fréquences de la zone de sécurité prosodique, mais son étendue est moins équilibrée autour du point central recherché.

- **Synthèse des résultats du ton monotone :**

Nous synthétisons, dans le tableau suivant, le nombre de témoins ayant adapté leur production au modèle orthophonique d'une part, puis au modèle de la voix transformée d'autre part. Nous nous basons surtout sur le deuxième histogramme pour chacun des sujets. En effet, il permet de voir qu'il y a bien modulation de la voix contrairement au premier histogramme qui reste plus global. Le reste du tableau met en évidence le nombre de sujets se rapprochant le plus de la zone de sécurité prosodique lors de l'imitation du modèle orthophonique et le nombre de sujets se rapprochant le plus de la zone de sécurité prosodique lors de l'imitation du modèle de la voix transformée.

Adaptation de la production spontanée au modèle orthophonique	Adaptation de la production spontanée au modèle de la voix transformée	Cas où l'imitation du modèle orthophonique est la meilleure et se rapproche le plus de la zone de sécurité prosodique	Cas où l'imitation du modèle de la voix transformée est la meilleure et se rapproche le plus de la zone de sécurité prosodique
12 / 14 témoins soit environ 85 % des témoins	14 / 14 témoins soit 100 % des témoins	1 / 14 témoins soit 7 % des témoins	13 / 14 témoins soit 93 % des témoins

Deux témoins n'ont donc pas réussi à réaliser une réelle modulation de leur voix pour imiter le modèle orthophonique. Il s'agit des témoins n°1, 9. Nous retrouvons deux hommes dont l'un n'est pas parvenu à moduler sa voix pour le ton mélodique.

L'ensemble des témoins a su s'adapter au modèle de la voix transformée. Les ajustements des muscles du larynx et des résonateurs sont donc très présents lors de cet exercice. L'imitation lors de cette épreuve est également la meilleure pour treize d'entre eux en terme de proximité par rapport à la zone de sécurité prosodique. L'imitation de la voix transformée favorise donc chez les témoins un usage plus sécurisé de la voix. Seul le témoin n°11 réalise une meilleure imitation du modèle orthophonique en terme de proximité par rapport à la zone de sécurité prosodique. Cependant, l'imitation du modèle de la voix transformée ne s'éloigne que de onze hertz supplémentaires au niveau de la borne inférieure. Cela ne constitue donc pas un réel écart.

2.2.2 RESENTI DES TEMOINS FACE AUX EPREUVES

Les personnes enregistrées, n'ont globalement pas ressenti de difficultés particulières dans la réalisation des épreuves proposées. Sur l'ensemble des témoins, tous ont la sensation d'avoir modulé leur voix. Tous soulignent que le modèle de la voix transformée est plus facile à reproduire que le modèle de l'orthophoniste et perçoivent à l'oreille les variations mélodiques entre leur voix et le modèle de la voix transformée.

12 des témoins affirment que le modèle de la voix transformée restant dans leur zone de sécurité prosodique est probant. Les 2 autres témoins trouvent que leur voix transformée est trop aiguë mais parviennent quand même à réaliser une imitation proche de leur nouvelle voix.

2.3 GROUPE DES PATIENTS DYSPHONIQUES

Nous exposerons tout d'abord dans un tableau les éléments recueillis grâce au questionnaire, afin de présenter les patients.

Puis nous évoquerons les résultats recueillis suite aux enregistrements des différents sujets dysphoniques pris dans la patientèle du Docteur P-O VEDRINE à CANNES. Nous traitons d'abord des épreuves sur ton mélodique puis de celles sur ton monotone.

Nous procédons de la même façon que pour M.GREGORIO. Les histogrammes obtenus pour chacun des témoins sont disposés les uns à la suite des autres. Nous ferons une brève analyse des résultats au cas par cas pour faire ressortir les éléments pertinents. Cela nous servira à réaliser une analyse globale des résultats obtenus pour le ton mélodique puis pour le ton monotone.

Nous mettrons dans la partie des annexes pour chaque témoin les résultats obtenus pour les différents calculs et les lignes mélodiques respectives permettant de comparer modèle orthophonique / imitation du sujet et modèle de la voix transformée / imitation du sujet.

2.3.1 DESCRIPTION DES PATIENTS DYSPHONIQUES

Nous avons recueilli 4 questionnaires. Dans l'analyse, nous nous appuyons sur les résultats bruts des questionnaires que nous mettons en relation avec d'éventuelles informations que les patients auront pu ajouter. Cela nous permet d'avoir une idée précise de la population étudiée. Nous indiquons dans le tableau récapitulatif qui suit : l'âge, le sexe, la profession, l'usage de la voix au quotidien, la pratique de la voix chantée, la cause de la dysphonie et le numéro des annexes où trouver les résultats et les lignes mélodiques correspondant à chaque sujet. Nous traiterons du ressenti des patients face à l'exercice proposé lors de la synthèse des résultats obtenus.

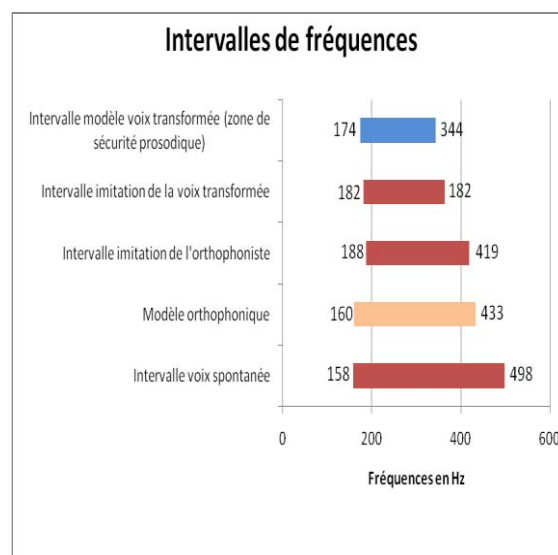
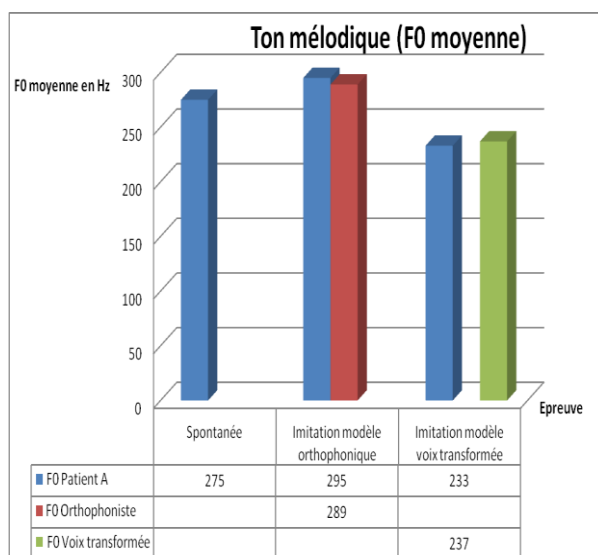
Patient	Sexe	Age	Profession	Usage de la voix	Pratique voix chantée	Cause de la dysphonie	Résultats et lignes mélodiques
A	Fém	56	Employée de bureau	Fréquent	Oui	Kyste sur l'apophyse vocale à droite	Annexe n°21
B	Fém	41	Assistante d'éducation	Fréquent	Oui	Suspicion d'un sulcus sur la corde vocale droite	Annexe n°22
C	Masc	25	Assistant d'éducation	Fréquent	Oui	Angiome sur la corde vocale gauche	Annexe n°23
D	Fém	50	Aide-soignante	Normal	Non	Dysphonie dysfonctionnelle	Annexe n°24

2.3.1.1 Ton mélodique :

Dans un premier temps, nous illustrons à l'aide d'un histogramme, les résultats de la fréquence fondamentale moyenne obtenus pour le patient et chacun des modèles en fonction des différentes épreuves. Cela nous permet notamment de comparer la fréquence moyenne entre le patient et le modèle orthophonique puis entre le patient et le modèle de sa voix transformée (là où la voix se situe dans la zone de sécurité prosodique).

Dans un second temps, nous illustrons à l'aide d'un histogramme placé à côté, les intervalles de fréquences dans lesquels se situent le patient et le modèle selon chaque épreuve. Nous indiquons également la zone de sécurité prosodique afin de comparer où se situe le sujet par rapport à cette zone.

- Patient A :

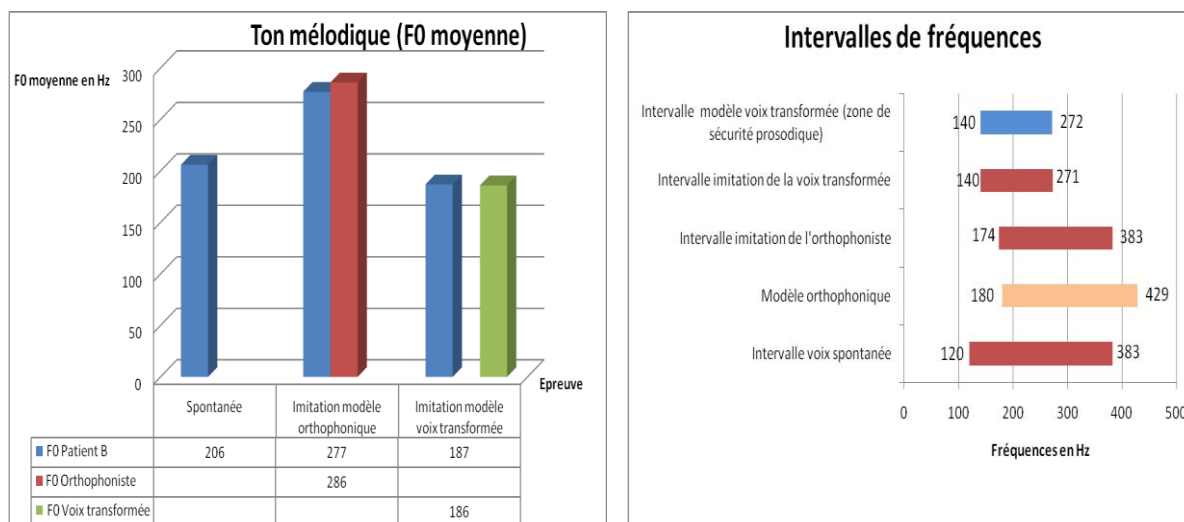


Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 275 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 289 Hz, nous voyons que la fréquence moyenne du sujet est de 295 Hz. Il y a une adaptation au modèle orthophonique avec un léger dépassement. Mais **l'épreuve d'imitation la mieux réussie en**

terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 233 Hz pour un modèle à 237 Hz. Il y a modulation de la voix. De plus, à l'écoute des productions, nous entendons que le sujet réalise de grandes variations mélodiques en émission spontanée et que celles-ci diminuent considérablement lors de l'imitation de la voix transformée. Le sujet reconnaît que dans la vie quotidienne, en fonction des émotions, il a bien tendance à réaliser de grandes variations mélodiques qui selon lui pourraient bien être néfastes pour sa voix.

L'observation du second histogramme permet d'observer qu'il y a bien une **adaptation vocale dans l'imitation de la voix transformée**. En effet, nous constatons que l'étendue fréquentielle se réduit considérablement entre l'émission spontanée et l'imitation de la voix transformée. Le sujet passe d'un intervalle de [158 Hz – 498 Hz] en émission spontanée à un intervalle de [182 Hz – 364 Hz] lors de l'imitation de la voix transformée. Cela lui permet également de rester beaucoup plus proche de sa zone de sécurité prosodique. Son amplitude fréquentielle est quasiment réduite de moitié. Nous voyons bien que l'intervalle fréquentiel en production spontanée est trop important, qui plus est en présence d'une pathologie vocale. La production spontanée du sujet a tendance à aller très haut dans les fréquences aigues. En effet, en écoutant le sujet, ses modulations sont proches de modulations théâtrales. Avec l'imitation de la voix transformée, la voix du sujet est ramenée dans un domaine sécurisé. A l'écoute, les variations mélodiques sont moindres. Le sujet a pris conscience du placement vocal adapté pour sa propre voix. Il a agi en conséquence, en réduisant l'intervalle fréquentiel. Soulignons qu'il y a bien tentative d'imitation du modèle orthophonique car l'étendue fréquentielle varie. Celle-ci est plutôt proche du modèle, mais cette imitation le contraint à sortir de la zone sécurisée puisque sa fréquence maximale est à 419 Hz alors que la zone sécurisée ne dépasse pas 344 Hz.

- Patient B :

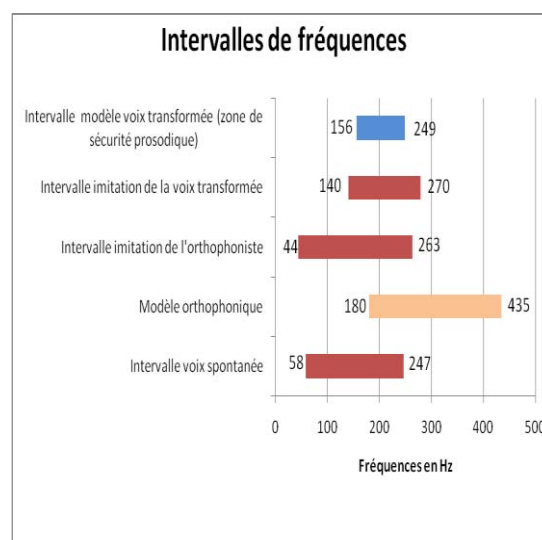
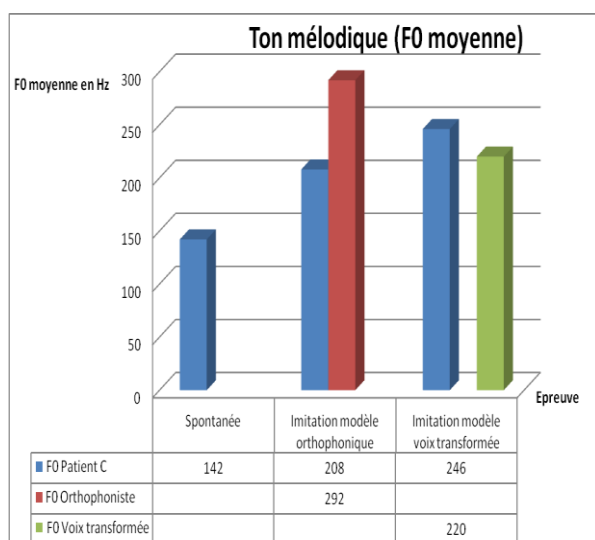


Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 206 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 286 Hz, nous voyons que la fréquence moyenne du sujet est de 277 Hz. Il y a une adaptation au modèle orthophonique. Mais **l'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 187 Hz pour un modèle à 186 Hz.** Il y a modulation de la voix.

L'observation du second histogramme permet d'observer qu'il y a bien une **adaptation vocale dans l'imitation de la voix transformée**. En effet, nous constatons que l'étendue fréquentielle se réduit considérablement entre l'émission spontanée et l'imitation de la voix transformée. Le sujet passe d'un intervalle de [120 Hz – 383 Hz] en émission spontanée à un intervalle de [140 Hz – 271 Hz] lors de l'imitation de la voix transformée. Là encore, l'amplitude fréquentielle est réduite de moitié. Cela souligne que le sujet utilise d'ordinaire sa voix bien au-delà des fréquences considérées comme sécurisées. En effet, à l'écoute des productions, nous entendons que la prosodie du sujet a tendance à descendre sur les fins de phrases. Le sujet reconnaît que dans la vie quotidienne, il parle dans une tonalité assez grave et que sa voix est sujette aux aggravations pendant la conversation. L'écoute d'un modèle objectif le guidant vers la voix idéale en terme de sécurité prosodique lui permet d'ajuster sa production. A l'écoute, l'imitation du modèle limite bien l'alternance pathologique de

fréquences graves et aigues extrêmes dans la production. Le sujet reste beaucoup plus proche de sa zone de sécurité prosodique. Soulignons qu'il y a bien tentative d'imitation du modèle orthophonique car l'étendue fréquentielle varie dans les fréquences basses. Mais cette imitation le contraint à sortir de la zone sécurisée puisque sa fréquence maximale est à 383 Hz alors que la zone sécurisée ne dépasse pas 272 Hz. Ainsi, seule l'imitation de la voix transformée permet d'aiguiller au mieux le sujet vers les fréquences adaptées.

- Patient C :

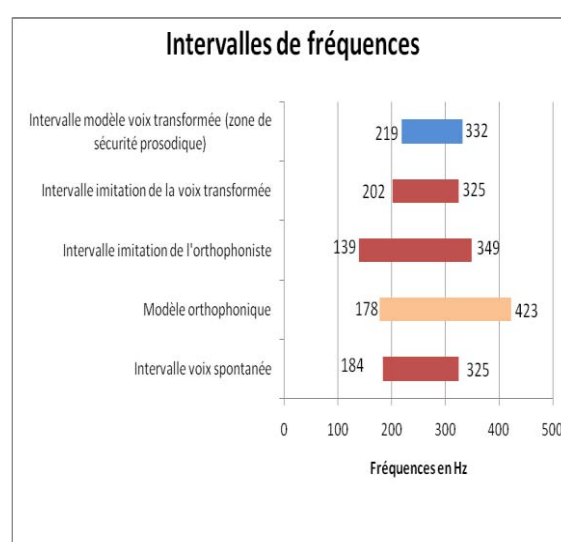
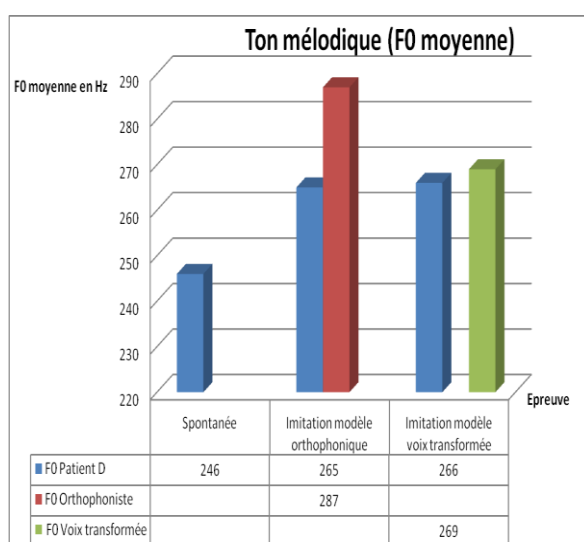


Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 142 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 292 Hz, nous voyons que la fréquence moyenne du sujet est de 208 Hz. Il y a une adaptation au modèle orthophonique. Mais **l'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 220 Hz pour un modèle à 246 Hz.** Il y a modulation de la voix.

L'observation du second histogramme permet d'observer qu'il y a bien une **adaptation vocale dans l'imitation de la voix transformée.** En effet, nous constatons que l'étendue fréquentielle se réduit considérablement entre l'émission spontanée et l'imitation de la voix transformée. Le sujet passe d'un intervalle de [58 Hz – 247 Hz] en émission spontanée à un

intervalle de [140 Hz – 249 Hz] lors de l'imitation de la voix transformée. L'écoute du modèle permet au sujet de déplacer sa production vers des fréquences plus élevées et de corriger le fait de parler dans des fréquences trop basses. Il reste de ce fait beaucoup plus proche de sa zone de sécurité prosodique, l'amplitude fréquentielle est donc diminuée. Soulignons qu'il y a bien tentative d'imitation du modèle orthophonique car l'étendue fréquentielle varie. Mais cette imitation le contraint à sortir considérablement de la zone sécurisée notamment au niveau des fréquences basses : sa fréquence minimale est à 44 Hz alors que la zone sécurisée ne dépasse pas 156 Hz.

- Patient D :



Pour le premier histogramme traitant de la fréquence fondamentale moyenne, nous constatons que le sujet parle spontanément à une fréquence moyenne de 246 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 287 Hz, nous voyons que la fréquence moyenne du sujet est de 265 Hz. Il y a une adaptation au modèle orthophonique avec un léger dépassement. Mais **l'épreuve d'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle de la voix transformée puisque le sujet réalise une production de 266 Hz pour un modèle à 269 Hz.** Il y a modulation de la voix. De plus, à l'écoute des productions, nous entendons que le sujet parle plutôt dans des fréquences assez basses. L'écoute du modèle de la voix transformée lui permet d'augmenter son niveau de fréquence vocale.

L'observation du second histogramme permet d'observer qu'il y a bien une **adaptation vocale dans l'imitation de la voix transformée**. En effet, nous constatons que l'étendue fréquentielle se réduit considérablement entre l'émission spontanée et l'imitation de la voix transformée. Le sujet passe d'un intervalle de [184 Hz – 325 Hz] en émission spontanée à un intervalle de [202 Hz – 325 Hz] lors de l'imitation de la voix transformée. Même si l'amplitude fréquentielle en production spontanée est proche de la zone de sécurité prosodique, nous voyons que le sujet parvient quand même à réaliser une modulation pour se rapprocher davantage de cette zone. Cela souligne d'autant plus l'impact de l'écoute de ce modèle puisque même de petits ajustements sont possibles. Soulignons qu'il y a bien tentative d'imitation du modèle orthophonique car l'étendue fréquentielle varie. Mais cette imitation le contraint à sortir considérablement de la zone sécurisée, à l'image du sujet précédent. En effet, dans les fréquences minimales notamment, sa fréquence est à 139 Hz alors que la zone sécurisée ne dépasse pas 219 Hz.

- **Synthèse des résultats des patients sur le ton mélodique :**

Nous synthétisons, dans le tableau suivant, le nombre de patients ayant adapté leur production au modèle orthophonique d'une part, puis au modèle de la voix transformée d'autre part. Nous nous basons surtout sur le deuxième histogramme pour chacun des patients. En effet, il permet de voir qu'il y a bien modulation de la voix contrairement au premier histogramme qui reste plus global. Le reste du tableau met en évidence le nombre de patients se rapprochant le plus de la zone de sécurité prosodique lors de l'imitation du modèle orthophonique et le nombre de patients se rapprochant le plus de la zone de sécurité prosodique lors de l'imitation du modèle de la voix transformée.

Adaptation de la production spontanée au modèle orthophonique	Adaptation de la production spontanée au modèle de la voix transformée	Cas où l'imitation du modèle orthophonique est la meilleure et se rapproche le plus de la zone de sécurité prosodique	Cas où l'imitation du modèle de la voix transformée est la meilleure et se rapproche le plus de la zone de sécurité prosodique
4 / 4 patients soit 100 % des patients	4 / 4 patients soit 100 % des patients	0 / 4 patient soit 0 % des patients	4 / 4 patients soit 100 % des patients

L'ensemble des patients a su moduler sa voix lors de l'imitation du modèle orthophonique mais aussi lors de l'imitation de la voix transformée. Tous les sujets ont donc été capables d'adapter leur prosodie au modèle donné et d'effectuer les ajustements nécessaires sur la musculature laryngée et les résonateurs. Le contrôle moteur sur ces structures est donc bien présent. Chez tous les patients, le modèle orthophonique favorise d'importantes modulations vocales, contrairement au modèle de la voix transformée qui permet aux patients de se rapprocher de leur zone de sécurité prosodique. Pour mieux se représenter l'importance pour une personne dysphonique de tendre vers cette zone de sécurité prosodique, prenons l'exemple, imagé bien sûr, d'une personne souffrant d'un problème musculaire à la jambe. Si le kinésithérapeute a l'habitude de faire de grands pas pour se déplacer, en montrant l'exemple à son patient, il l'induit à faire de grands pas. Or cela ne convient pas pour ce patient souffrant de la jambe. En faisant des pas plus petits, adaptés à lui, il est sûr de limiter une action néfaste. La zone de sécurité prosodique suit un peu ce principe. L'orthophoniste a sa propre façon de s'exprimer, difficilement adaptable à autrui. En créant le modèle vocal correspondant exactement au patient pour limiter les mauvaises variations vocales et placer convenablement sa voix, nous l'orientons là où la voix est sécurisée et limitons fatigue et malmenage vocal.

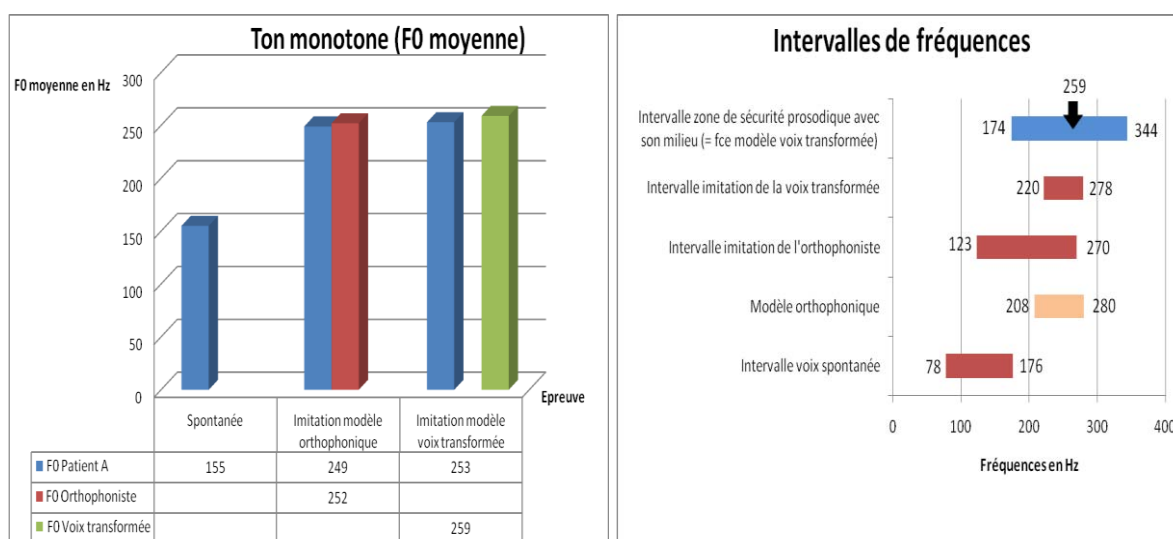
Pour tous les patients, l'imitation de la voix transformée est donc la plus propice pour favoriser une production dans la zone de dynamique vocale sécurisée.

2.3.1.2 Ton monotone :

Comme pour le ton mélodique, dans un premier temps, nous illustrons à l'aide d'un histogramme, les résultats de la fréquence fondamentale moyenne obtenus pour le patient et chacun des modèles en fonction des différentes épreuves. Cela nous permet notamment de comparer la fréquence moyenne entre le patient et le modèle orthophonique puis entre le patient et le modèle de sa voix transformée (là où la voix se situe dans la zone de sécurité prosodique).

Dans un second temps, nous illustrons à l'aide d'un histogramme placé à côté, les intervalles de fréquences dans lesquels se situe le patient selon chaque épreuve. L'intérêt de l'épreuve sur ton monotone est aussi de voir comment l'individu peut osciller autour du point central (milieu) de la zone de sécurité prosodique. Nous indiquons donc le point central de l'intervalle de la zone de sécurité prosodique, point autour duquel la voix du sujet doit tenter de se rapprocher (et qui n'est autre que la fréquence moyenne calculée pour la ligne monotone de la voix transformée).

- Patient A :

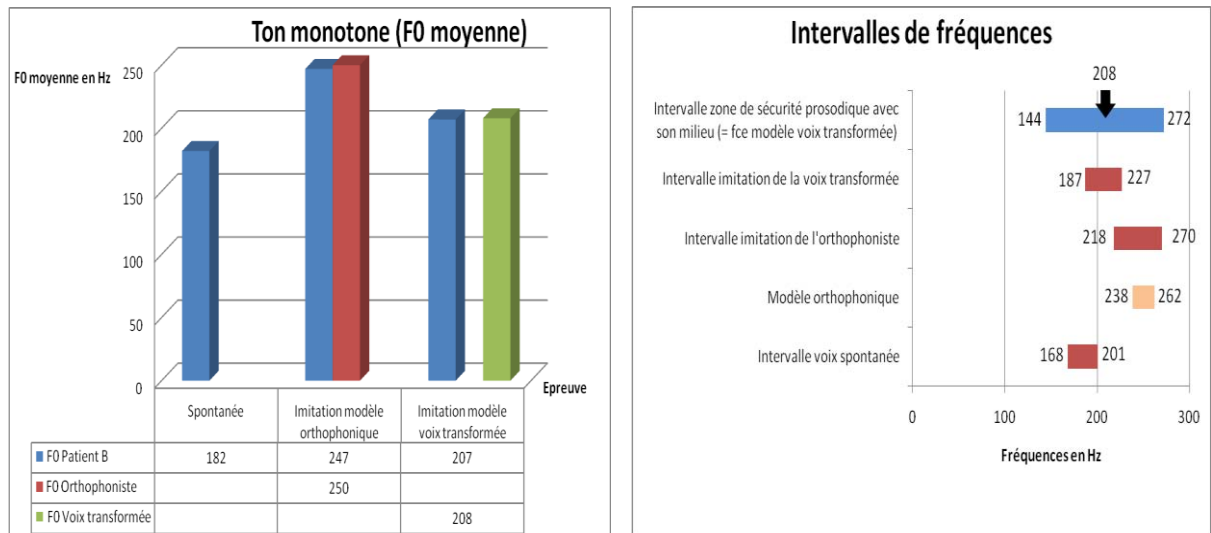


Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 155 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 252 Hz, nous voyons que la fréquence moyenne du sujet est de 249 Hz. Il y a donc une imitation probante du modèle orthophonique puisque le sujet module

considérablement sa voix pour tenter d'imiter ce modèle. Sa production est très proche de ce modèle en terme de fréquence fondamentale moyenne. Le sujet se rapproche d'ailleurs plus du modèle orthophonique que du modèle de la voix transformée, à quelques hertz près. La modulation de la voix est présente dans les deux épreuves d'imitation mais **l'imitation la mieux réussie en terme de fréquence moyenne est celle du modèle orthophonique.**

Cependant, le second histogramme nous indique des éléments différents. **Nous pouvons observer que l'imitation la mieux réussie en terme d'étendue fréquentielle n'est pas celle du modèle orthophonique mais celle du modèle de la voix transformée.** En effet, l'imitation du modèle orthophonique contraint le sujet à sortir davantage de sa zone de sécurité prosodique dans les fréquences inférieures et son étendue fréquentielle est beaucoup plus importante que celle de l'imitation du modèle de la voix transformée comme le montrent les intervalles : [123 Hz – 270 Hz] pour l'imitation du modèle orthophonique et [220 Hz – 278 Hz] pour l'imitation du modèle de la voix transformée. L'imitation de sa voix transformée est donc la mieux réalisée. De plus ce dernier modèle permet au sujet de rester complètement dans sa zone de sécurité prosodique en se rapprochant bien du point central à 259 Hz.

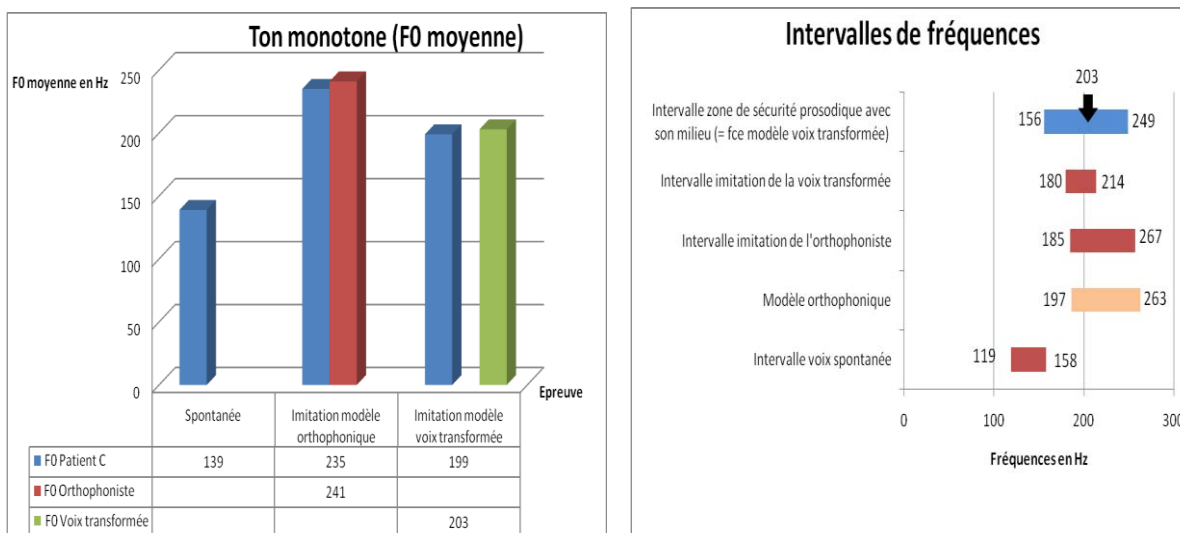
- Patient B :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 182 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 250 Hz, nous voyons que la fréquence moyenne du sujet est de 247 Hz. Il y a donc une imitation probante du modèle orthophonique puisque le sujet module considérablement sa voix pour tenter d'imiter ce modèle. Sa production est très proche de ce modèle en terme de fréquence fondamentale moyenne. **Mais l'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle du modèle de la voix transformée. Le sujet a une fréquence fondamentale à 207 Hz pour un modèle à 207 Hz.**

Le second histogramme confirme ces éléments. Nous pouvons observer que l'imitation la mieux réussie en terme d'étendue fréquentielle est bien celle du modèle de la voix transformée. En effet, l'étendue fréquentielle passe de [168 Hz – 201 Hz] en voix spontanée à [187 Hz – 227 Hz] lors de l'imitation. La modulation de la voix est importante ici. Par rapport à la production spontanée, l'étendue fréquentielle lors de cette imitation s'équilibre davantage autour du point central à 208 Hz de la zone de sécurité prosodique, ce qui témoigne de l'adaptation plus que positive. Concernant le modèle orthophonique, il y a bien une modulation de la voix puisque l'étendue fréquentielle du sujet passe à [218 Hz – 270 Hz] mais le sujet sort considérablement du champ de fréquences de la zone de sécurité prosodique.

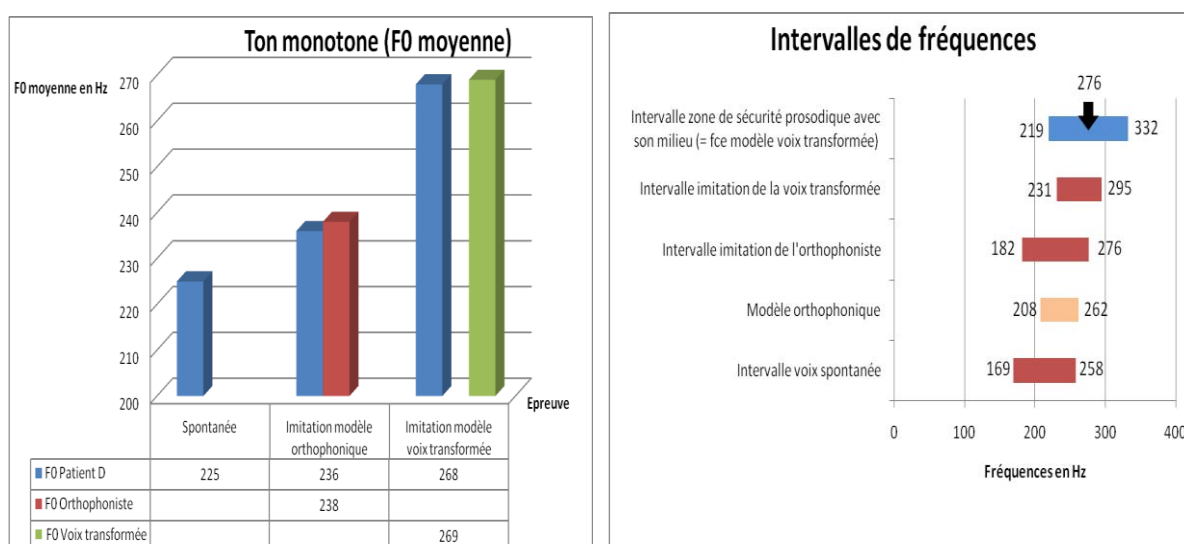
- Patient C :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 139 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 241 Hz, nous voyons que la fréquence moyenne du sujet est de 235 Hz. Il y a donc une imitation probante du modèle orthophonique puisque le sujet module considérablement sa voix pour tenter d'imiter ce modèle. Sa production est très proche de ce modèle en terme de fréquence fondamentale moyenne. **Mais l'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle du modèle de la voix transformée. Le sujet a une fréquence fondamentale à 199 Hz pour un modèle à 203 Hz.**

Le second histogramme confirme ces éléments. **Nous pouvons observer que l'imitation la mieux réussie en terme d'étendue fréquentielle est bien celle du modèle de la voix transformée.** En effet, l'étendue fréquentielle passe de [119 Hz – 158 Hz] en voix spontanée à [180 Hz – 214 Hz] lors de l'imitation. La modulation de la voix est importante ici. Par rapport à la production spontanée, l'étendue fréquentielle lors de cette imitation s'équilibre davantage autour du point central à 203 Hz de la zone de sécurité prosodique, ce qui témoigne de l'adaptation plus que positive. Concernant le modèle orthophonique, il y a bien une modulation de la voix puisque l'étendue fréquentielle du sujet passe à [185 Hz – 267 Hz] mais le sujet sort considérablement du champ de fréquences de la zone de sécurité prosodique.

- Patient D :



Dans le premier histogramme, nous constatons que le sujet parle spontanément à une fréquence moyenne de 225 Hz. Puis, lorsqu'il imite le modèle orthophonique dont la fréquence moyenne se situe à 238 Hz, nous voyons que la fréquence moyenne du sujet est de 236 Hz. Il y a donc une imitation probante du modèle orthophonique puisque le sujet module considérablement sa voix pour tenter d'imiter ce modèle. Sa production est très proche de ce modèle en terme de fréquence fondamentale moyenne. Même si les performances sont présentes pour les deux imitations, **l'imitation la mieux réussie en terme de fréquence fondamentale moyenne est celle du modèle de la voix transformée. Le sujet a une fréquence fondamentale à 268 Hz pour un modèle à 269 Hz.**

Le second histogramme confirme ces éléments différents. **Nous pouvons observer que l'imitation la mieux réussie en terme d'étendue fréquentielle est bien celle du modèle de la voix transformée.** En effet, l'étendue fréquentielle passe de [231 Hz – 295 Hz] en voix spontanée à [169 Hz – 258 Hz] lors de l'imitation. La modulation de la voix est importante ici. Comme dans les cas précédents, l'étendue fréquentielle s'équilibre davantage autour du point central à 276 Hz de la zone de sécurité prosodique lors de cette imitation que lors de la production spontanée, ce qui témoigne de l'adaptation positive. Concernant le modèle orthophonique, il y a bien une modulation de la voix puisque l'étendue fréquentielle du sujet passe à [182 Hz – 276 Hz] mais le sujet sort du champ de fréquences de la zone de sécurité prosodique.

- **Synthèse des résultats des patients pour le ton monotone :**

Nous synthétisons, dans le tableau suivant, le nombre de patients ayant adapté leur production au modèle orthophonique d'une part, puis au modèle de la voix transformée d'autre part. Nous nous basons surtout sur le deuxième histogramme pour chacun des patients. En effet, il permet de voir qu'il y a bien modulation de la voix contrairement au premier histogramme qui reste plus global. Le reste du tableau met en évidence le nombre de patients se rapprochant le plus de la zone de sécurité prosodique lors de l'imitation du modèle orthophonique et le nombre de patients se rapprochant le plus de la zone de sécurité prosodique lors de l'imitation du modèle de la voix transformée.

Adaptation de la production spontanée au modèle orthophonique	Adaptation de la production spontanée au modèle de la voix transformée	Cas où l'imitation du modèle orthophonique est la meilleure et se rapproche le plus de la zone de sécurité prosodique	Cas où l'imitation du modèle de la voix transformée est la meilleure et se rapproche le plus de la zone de sécurité prosodique
4 / 4 patients soit 100 % des patients	4 / 4 patients soit 100 % des patients	0 / 4 patients soit 0 % des patients	4 / 4 patients soit 100 % des patients

L'ensemble des patients a su moduler sa voix lors de l'imitation du modèle orthophonique mais aussi lors de l'imitation de la voix transformée. Tous les sujets ont donc été capables d'adapter leur prosodie au modèle donné et d'effectuer les ajustements nécessaires pour la musculature laryngée et les résonateurs. Chez tous les patients, à l'image du ton mélodique, le modèle orthophonique favorise d'importantes modulations vocales, contrairement au modèle de la voix transformée qui permet aux patients de se rapprocher de leur zone de sécurité prosodique. L'écoute du modèle de la voix transformée sur ton monotone entraîne systématiquement une imitation où la voix oscille autour du point sensé être le plus sécurisé. En effet, il se situe au milieu de la zone de sécurité prosodique. En essayant de se rapprocher de ce point, le sujet ne peut que se rapprocher de cette zone sécurisée. Seule l'écoute et l'imitation de ce modèle permet au sujet de cerner exactement le point central où sa voix doit se situer pour éviter le malmenage.

Pour tous les patients, l'imitation de la voix transformée est donc la plus propice pour respecter la zone de sécurité prosodique.

2.3.2 RESENTI DES PATIENTS FACE AUX EPREUVES

Les personnes enregistrées n'ont globalement pas ressenti de difficultés particulières dans la réalisation des épreuves proposées. Sur l'ensemble des patients, tous perçoivent à l'oreille les variations mélodiques entre leur voix et le modèle de la voix transformée. Tous les patients affirment que le modèle de la voix transformée restant dans leur zone de sécurité prosodique est probant. Les patients expliquent qu'ils ont conscience des variations mélodiques néfastes présentes dans leur production spontanée mais ont des difficultés à les objectiver. Ainsi, les patients soulignent l'intérêt de l'écoute réalisée dans ce travail qui leur a permis de mieux cerner où se situent les variations pathologiques de leur voix et d'avoir un réel modèle de leur voix dite «confortable ». Plus largement, la proposition d'un travail d'écoute de la voix transformée pour l'appliquer à leur quotidien leur semblerait très positive.

2.4 SYNTHÈSE GLOBALE DES RESULTATS

1- Tout d'abord, nous synthétisons, dans le tableau suivant, l'ensemble des résultats précédemment trouvés concernant le nombre de sujets ayant adapté leur production au modèle orthophonique d'une part, puis au modèle de la voix transformée d'autre part, tant sur le ton mélodique que monotone. Le reste du tableau met également en évidence le nombre de sujets se rapprochant le plus de la zone de sécurité prosodique lors de l'imitation du modèle orthophonique et le nombre de sujets se rapprochant le plus de la zone de sécurité prosodique lors de l'imitation du modèle de la voix transformée.

Adaptation de la production spontanée au modèle orthophonique	Adaptation de la production spontanée au modèle de la voix transformée	Cas où l'imitation du modèle orthophonique est la meilleure et se rapproche le plus de la zone de sécurité prosodique	Cas où l'imitation du modèle de la voix transformée est la meilleure et se rapproche le plus de la zone de sécurité prosodique
80 % de la population	100 % de la population	3 % de la population	97 % de la population

- **Cas de l'adaptation de la production spontanée au modèle orthophonique :**

Les résultats obtenus pour l'adaptation de la production au modèle orthophonique montrent que la modulation de la voix est observée pour plus des trois-quarts de la population, tant sur le ton mélodique que monotone. L'imitateur et la population de patients arrivent tous à moduler leur voix dans ce cas. Seuls quelques témoins ne sont pas parvenus à réaliser une réelle modulation de leur voix pour imiter le modèle orthophonique. Ces résultats montrent toutefois que la modulation de la voix pour imiter l'orthophoniste est possible pour la majorité de la population étudiée. Les sujets sont sensibles à une hauteur tonale différente de la leur. L'écoute du modèle entraîne bien des ajustements laryngés. Il y a bien un contrôle cérébral conscient dans l'ajustement de la musculature laryngée pour essayer de reproduire le modèle.

Notons que l'adaptation au modèle orthophonique est mieux réussie sur le ton monotone (seuls 2 témoins ne parviennent pas à réellement moduler leur voix) que sur le ton mélodique (où 4 témoins ne parviennent pas à moduler leur voix).

- **Cas de l'adaptation de la production spontanée au modèle de la voix transformée :**

Les résultats obtenus pour l'adaptation de la production au modèle de la voix transformée montrent que la modulation de la voix est présente pour l'ensemble de la population étudiée, tant sur le ton mélodique que monotone. Il y a donc ici une réussite totale. Ces résultats indiquent que l'imitation de la voix transformée est possible pour l'ensemble des personnes étudiées. Ils sont meilleurs que ceux relevés dans le cas de l'adaptation de la production au modèle orthophonique. Les personnes dysphoniques au même titre que les témoins et l'imitateur sont davantage sensibles à leur propre voix et parviennent mieux à ajuster leur production en fonction. La boucle audio-phonatoire et les ajustements des paramètres biomécaniques de la musculature laryngée réalisés par le système moteur sont plus performants lorsque le sujet perçoit sa propre voix.

- **Cas où l'imitation du modèle orthophonique est la meilleure et se rapproche le plus de la zone de sécurité prosodique :**

Nous retrouvons un seul sujet pour lequel l'imitation du modèle orthophonique est la meilleure sur l'ensemble de la population étudiée, cela dans le cas du ton monotone. Il s'agit d'un des témoins. Ce résultat peu s'expliquer par de bonnes capacités d'imitation. Ce témoin a su moduler sa voix dans l'imitation du modèle orthophonique et de la voix transformée mais l'imitation du modèle orthophonique l'emporte de peu. Pour la majorité, il apparaît que le modèle orthophonique étant éloigné de la fréquence fondamentale moyenne des sujets, il n'est pas approprié pour que le sujet reste dans la zone de sécurité prosodique lui correspondant.

- **Cas où l'imitation de la voix transformée est la meilleure et se rapproche le plus de la zone de sécurité prosodique :**

Tous les sujets à l'exception du sujet précédemment cité réalisent la meilleure imitation pour le modèle de la voix transformée. Ce modèle correspondant exactement à la voix du sujet, il lui est plus facile de l'imiter que le modèle orthophonique. De plus, l'étendue fréquentielle du modèle de la voix transformée est calculée pour ne pas dépasser la zone de sécurité prosodique. Le sujet a donc le repère adéquat pour savoir où se trouve cette zone. Il constitue un support beaucoup plus objectif et adapté que le modèle orthophonique et demeure le seul à offrir l'illustration exacte de la zone où le rendement vocal est optimal. Grâce à ce modèle, le sujet perçoit exactement quelle est la zone de sa dynamique vocale où il doit demeurer pour ne pas malmenager sa voix. Ainsi, lors de l'imitation, disposant du modèle exact et avec l'intervention de la boucle audio-phonatoire, les ajustements conscients ou réflexes nécessaires peuvent se produire, qu'ils soient au niveau de la musculature laryngée ou des résonateurs.

2- Puis, nous avons décidé de réaliser deux tableaux plus détaillés pour faire ressortir notamment les oscillations de la voix pour le ton monotone et plus spécifiquement, pour évaluer l'importance des écarts observés par rapport à la zone de sécurité prosodique pour chacun des sujets sur le ton mélodique.

Nous retrouvons : - la zone de sécurité prosodique minimale et maximale,
 - le milieu de la zone de sécurité prosodique,
 - la fréquence moyenne en voix spontanée monotone,
 - le taux de variations (nommé « σ ») dans la voix, relevé pour le ton monotone lors de la production spontanée, lors de l'imitation du modèle orthophonique, et lors de l'imitation du modèle de la voix transformée. Nous obtenons ce taux grâce à la fonction « *voice report* » puis « *standard deviation* » du logiciel PRAAT. Pour illustrer cela, c'est comme s'il existait sur le spectrogramme un trait symbolisant la valeur moyenne de la fréquence fondamentale de la production réalisée. Le logiciel calcule alors en moyenne les écarts de la voix par rapport à cette ligne. Il évalue en quelque sorte le taux d'oscillations de la voix.

Puis : - la fréquence moyenne de la voix spontanée mélodique,
- le taux de variations (nommé « δ ») par rapport à la zone de sécurité prosodique, observé pour la production spontanée, l'imitation du modèle orthophonique, l'imitation du modèle de la voix transformée. Cet écart est calculé par rapport aux étendues fréquentielles. La référence est la zone de sécurité prosodique. Nous observons l'endroit où la production spontanée, l'imitation du modèle orthophonique ou l'imitation du modèle de la voix transformée dépasse la zone sécurisée. S'il y a dépassement au niveau des deux bornes de l'étendue fréquentielle, nous additionnons les deux chiffres trouvés. S'il n'y a dépassement qu'au niveau d'une seule borne, nous utilisons ce chiffre seul. Dans les deux cas, nous divisons le chiffre trouvé par le milieu de la zone de sécurité prosodique. Cela nous permet de voir les variations de la voix par rapport à la zone de sécurité prosodique. Nous avons choisi de mettre en couleur trois niveaux de seuil :

- en gris lorsque le δ est inférieur ou égal à 0,2
- en orange lorsque le δ est strictement supérieur à 0,2 et strictement inférieur à 0,4
- en rose lorsque le δ est supérieur ou égal à 0,4.

Nous avons défini ces bornes en observant les résultats obtenus dans les histogrammes. L'imitation de la voix transformée étant la meilleure, nous avons pris comme référence la valeur obtenue la plus grande qui symbolise l'éloignement maximal de la zone de sécurité prosodique pour cet exercice. Il s'agit de 0,242 que nous arrondissons à 0,2. Au dessous de cette valeur le sujet est proche de sa zone de sécurité prosodique. De ce fait, nous avons décidé d'appliquer un intervalle de progression de 0,2 ; c'est pour cette raison que le deuxième seuil se situe à 0,4 et signale un éloignement plus marqué de la zone sécurisée. Enfin, toutes les personnes obtenant plus de 0,4 s'éloignent réellement de la zone de sécurité prosodique.

Dans l'analyse des résultats, nous détaillerons surtout le ton mélodique qui se rapproche le plus de la production qu'un sujet utilise dans la vie quotidienne.

- **Ton monotone :**

Sujet	Catégorie	ZSP min	ZSP max	ZSP mid	fréquence moyenne voix spontanée monotone	σ monotone spontané	σ monotone sur imitation modèle ortho	σ monotone sur imitation modèle voix transformée
M.G.	Imitateur	98	172	135	105	3,772	5,206	4,878
n°1	Témoin	102	205	154	128	6,682	7,498	11,861
n°2	Témoin	95	200	147,5	106	9,959	6,434	5,619
n°3	Témoin	180	266	223	195	18,568	12,4	8,796
n°4	Témoin	207	317	262	191	12,023	9,037	13,634
n°5	Témoin	169	258	213,5	211	10,632	15,725	13,574
n°6	Témoin	116	186	151	122	6,879	18,583	3,219
n°7	Témoin	101	223	162	100	5,633	4,915	5,697
n°8	Témoin	153	199	176	161	11,552	12,172	12,465
n°9	Témoin	95	153	124	109	5,906	4,639	5,118
n°10	Témoin	174	360	267	174	6,381	8,309	6,914
n°11	Témoin	185	273	229	161	10,898	11,656	10,579
n°12	Témoin	175	286	230,5	166	10,08	10,777	8,769
n°13	Témoin	183	354	268,5	227	13,245	11,331	11,846
n°14	Témoin	105	245	175	102	2,753	4,433	6,28
A	Patient	174	344	259	155	26,109	21,486	6,779
B	Patient	144	172	208	182	6,44	7,994	6,421
C	Patient	156	249	203	139	8,482	10,112	5,609
D	Patient	219	332	275,5	225	16,062	13,509	13,427
	MOYENNE	149	252,3	205,25	155,73684	10,108211	10,327158	8,4992105
	ECART-TYPE	40,87	66,06	99,35	42,37	5,60	4,71	3,42

Le tableau montre que la voix des personnes est sujette à de moins grandes oscillations lors de l'imitation de la voix transformée. En effet, l'écart-type relevé en moyenne est de 3,42 alors que pour l'imitation du modèle de l'orthophoniste il se situe à 4,71 et pour la production spontanée il a une valeur de 5,60 (en bleu clair dans le tableau). La moyenne est quant à elle d'environ 8,5 pour l'imitation de la voix transformée alors qu'elle est de 10,33 pour l'imitation du modèle orthophonique et 10,11 pour la production spontanée (en bleu foncé dans le tableau). Lorsque l'on observe au cas par cas, les chiffres indiquent parfois une imitation du modèle orthophonique où le taux d'oscillations est moins important que celui relevé pour l'imitation du modèle de la voix transformée. Cependant, les chiffres obtenus grâce au logiciel sont sensibles à la zone sélectionnée. Le logiciel délivre un résultat différent en fonction de la zone sélectionnée. Ainsi, les résultats changent si la sélection de la production varie à quelques millimètres près. De plus, lorsque des artefacts sont présents, il

est aussi difficile d'évaluer précisément les variations. Il vaut donc mieux avoir un regard général sur les oscillations. La moyenne et l'écart-type obtenus pour chaque exercice permettent d'avoir cette approche plus globale des oscillations et de se rendre compte sur la totalité de la population quel exercice limite le mieux les oscillations. Il n'était donc pas opportun dans ce cas d'établir des seuils comme il en est le cas pour le ton mélodique.

Ces résultats vont dans le sens du précédent recueil de données (cf partie 2.4 paragraphe 1 de la partie pratique). Les oscillations de la voix sont moins importantes dans l'imitation de la voix transformée. Cette imitation est donc la plus probante au niveau de la sécurité vocale.

- **Ton mélodique :**

Sujet	Catégorie	ZSP min	ZSP max	ZSP mid	fréquence moyenne voix spontanée mélodique	δ mélodique spontané	δ mélodique sur imitation modèle ortho	δ mélodique sur imitation modèle voix transformée
M.G.	Imitateur	98	172	135	145	0,274	1,325	0,089
n°1	Témoin	102	205	154	122	0,246	0,428	0,123
n°2	Témoin	95	200	147,5	141	0,291	0,278	0,04
n°3	Témoin	180	266	223	199	0,367	0,34	0,242
n°4	Témoin	207	317	262	231	0,591	0,374	0,145
n°5	Témoin	169	258	213,5	230	0,547	0,88	0,145
n°6	Témoin	116	186	151	159	0,291	0,204	0,033
n°7	Témoin	101	223	162	119	0,104	0,29	0,135
n°8	Témoin	153	199	176	174	0,63	0,795	0,148
n°9	Témoin	95	153	124	107	0,361	0,289	0,072
n°10	Témoin	174	360	267	215	0,089	0,4	0,066
n°11	Témoin	185	273	229	233	0,868	1,135	0,187
n°12	Témoin	175	286	230,5	222	0,388	1,077	0,06
n°13	Témoin	183	354	268,5	266	0,367	0,364	0,052
n°14	Témoin	105	245	175	142	0,17	0,211	0,08
A	Patient	174	344	259	275	0,776	0,397	0,077
B	Patient	144	172	208	206	0,634	0,534	0
C	Patient	156	249	203	142	0,483	0,62	0,173
D	Patient	219	332	275,5	246	0,127	0,352	0,062
	MOYENNE	149	252,3	205,25	188,10526	0,4002105	0,5417368	0,1015263
	ECART- TYPE	40,87	66,06	99,35	52,70	0,23	0,34	0,06

L'observation du tableau montre différents éléments. Tout d'abord, nous pouvons observer qu'en production spontanée, seuls 4 sujets sur les 19 soit 20% de la population sont réellement proches de leur zone de sécurité prosodique (en gris dans le tableau), il s'agit de 3 témoins et d'un patient. Il ressort que 7 des 19 sujets s'éloignent réellement de la zone de sécurité prosodique soit environ 40 % de la population (en rose dans le tableau). Parmi eux, nous retrouvons 3 patients et 4 témoins. Le reste des personnes enregistrées, soit environ 40 % des sujets, s'éloignent bien de leur zone sécurisée mais d'une façon moindre puisqu'ils sont dans l'intervalle intermédiaire (en orange dans le tableau). L'imitateur fait partie de ces personnes.

Lors de l'imitation du modèle orthophonique, 9 des 19 sujets soit environ 50 % de la population s'éloignent réellement de la zone de sécurité prosodique (en rose dans le tableau). Il s'agit de 6 témoins, 2 patients et de l'imitateur. C'est l'imitateur qui se trouve le plus éloigné de sa zone sécurisée. Son habitude professionnelle peut expliquer ce résultat. En effet, lors de cet exercice, M.GREGORIO a vraiment cherché à imiter la voix de l'orthophoniste. Cela explique ce taux maximal. Le reste de la population, soit les 50 % restants, s'éloignent bien de la zone sécurisée mais d'une façon moindre puisqu'ils sont dans l'intervalle intermédiaire (en orange dans le tableau). Il s'agit de certains témoins et de 2 patients dysphoniques.

Il est intéressant de noter que dans deux cas seulement, le modèle orthophonique permet de se rapprocher davantage de la zone de sécurité prosodique par rapport à la production spontanée. Il s'agit du témoin n°4 et du patient A. Dans 7 cas, les sujets changent de seuils car l'imitation de ce modèle les éloigne davantage de leur zone sécurisée. Il s'agit de l'imitateur, des témoins n°1, 7, 10, 12, 14 et du patient D. Par ailleurs, relevons que 2 témoins et 2 patients présentent un éloignement réel de la zone de sécurité prosodique en voix spontanée et dans l'imitation du modèle orthophonique. Il s'agit des témoins n°5 et 11 et des patients B et C. Soulignons que le témoin n°11 a déjà présenté des nodules sur les cordes vocales, comme nous l'avons vu dans le tableau présentant la population témoin (cf. partie 2.2.1 de la partie pratique). Le modèle orthophonique ne permet donc qu'un ajustement intuitif, il ne constitue en aucun cas un modèle réellement adéquat pour guider la sujet vers un usage sécurisé de sa voix.

Enfin, évidemment, les sujets sont tous plus proches de leur zone de sécurité prosodique lorsqu'ils imitent le modèle de la voix transformée (en gris dans le tableau). En effet, seul le sujet pris comme référence du seuil dépasse de la valeur cible à 0,2 (en orange dans le tableau).

Ces résultats confortent le précédent recueil de données (cf partie 2.4 paragraphe 1 de la partie pratique). Les sujets, en particulier les patients dysphoniques, sont plus proches de leur zone de sécurité prosodique lors de l'imitation du modèle de la voix transformée. L'écart-type (en vert clair dans le tableau) est de 0,06 et la moyenne (en vert foncé dans le tableau) est de 0,10 environ. L'usage de la voix est beaucoup plus sécurisé lors de cette épreuve. La production spontanée révèle que certains des sujets s'éloignent déjà de leur zone de sécurité prosodique dans la vie quotidienne. L'écart-type est de 0,23 et la moyenne de 0,4 environ. L'imitation du modèle orthophonique oriente les sujets loin de leur zone de sécurité prosodique dans tous les cas, sauf pour le témoin n°4 et le patient A. Effectivement, l'écart-type est de 0,34 et la moyenne de 0,55 environ.

Dans l'ensemble, les sujets dysphoniques font un mauvais usage de leur voix. Le modèle orthophonique ne les aide pas à trouver une fréquence vocale adaptée à leur prosodie. Seul le modèle de la voix transformée apporte les meilleurs résultats. En rassemblant les données comme tel, il apparaît à nouveau qu'objectiver la zone de sécurité prosodique est le seul moyen pour que le sujet perçoive les modulations vocales à supprimer dans sa production et qu'il réalise les ajustements requis. Les muscles intrinsèques et extrinsèques du larynx qui agissent sur la hauteur tonale notamment sont plus sensibles à l'auto-imitation. Cette dernière rend également la boucle audio-phonatoire plus performante puisque les productions sont proches du modèle donné. Il est donc possible, via l'imitation, d'induire les ajustements musculaires adaptés, entre autres, pour se rapprocher d'une zone fréquentielle sans danger. Cette zone, où la production des phonèmes ne doit pas dépasser les fréquences imposées, offre aux cordes vocales une vibration sécurisée. Ainsi, l'imitation est porteuse, en terme de sécurité prosodique, à condition de disposer d'un modèle guidant vers les bonnes fréquences.

3. DISCUSSION :

Dans cette partie, nous nous appuyons sur les résultats observés pour dégager les critiques méthodologiques et les problèmes rencontrés lors de notre travail. Puis nous discuterons ces résultats pour enfin les réintégrer dans le champ de l'orthophonie.

3.1 CRITIQUES METHODOLOGIQUES

3.1.1 PARTIE THEORIQUE

Ce mémoire traitant de la voix et de l'imitation, nous avons choisi de rassembler les différentes notions et éclairages connus sur ces thèmes.

Concernant le chapitre consacré à l'imitation, nous avons souhaité explorer un large champ de connaissances pour tenter de mieux cerner ce que recouvre ce terme. Pour cela, nous avons retenu les domaines les plus pertinents pour notre étude et intégré les dernières découvertes apportant une ouverture certaine mais non exclusive.

Il nous a semblé important de consacrer un chapitre sur le cas particulier des imitateurs pour tenter d'illustrer les processus mis en œuvre dans l'imitation vocale. En effet, il nous paraissait important d'essayer de comprendre comment procèdent ces sujets aux facultés d'imitation reconnues et pourtant si mystérieuses. Le but étant, plus largement, de savoir ce qui serait susceptible de se passer à moindre échelle chez une personne non professionnelle.

Le chapitre sur le logiciel PRAAT a pour objectif de montrer les possibilités qu'offre le logiciel pour objectiver la voix et également, d'illustrer l'environnement de travail qui nous a permis de définir la zone de sécurité prosodique.

Le chapitre sur la zone de sécurité prosodique permet de mieux cerner en quoi elle consiste.

Nous terminons par le chapitre relatif à la voix, les éléments apportés sont connus de tous les professionnels s’y intéressant mais il nous semblait important de les rappeler afin de permettre la meilleure compréhension possible pour un public non averti. L’évocation des éléments anatomiques nous semblait importante pour mettre en évidence les structures et processus susceptibles d’intervenir dans l’imitation.

3.1.2 PARTIE PRATIQUE

Dans cette étude, nous avons choisi de nous centrer exclusivement sur la voix parlée mais nous n’omettons pas que la voix chantée a également son importance, de même que l’émotion véhiculée par la voix.

Concernant le nombre de personnes enregistrées, nous avons rencontré des contraintes d’organisation qui n’ont pas permis d’augmenter ce nombre autant que souhaité. Cependant, nous avons eu la volonté et la possibilité de respecter un panel représentatif : cas d’un imitateur, population de témoins, population de patients dysphoniques. Il nous semblait important de pouvoir analyser et confronter les résultats de personnes aux capacités d’imitation et aux possibilités physiologiques différentes. Par ailleurs, les différentes tranches d’âge recherchées pour notre étude sont présentes.

Il nous a été difficile de respecter systématiquement des conditions d’enregistrement similaires, la localisation géographique étant différente pour la plupart des personnes sollicitées. De plus, lors des enregistrements des témoins, il est arrivé à trois reprises que deux personnes se trouvent dans la même pièce pour réaliser la passation sans pouvoir travailler individuellement. Ceci a sans doute aiguillé la seconde personne dans la réalisation des épreuves pour anticiper le comportement vocal requis.

Dans cette partie, nous avons réalisé l’ensemble des passations, nous avons donc pu contrôler le bon déroulement des enregistrements. Ceci nous a permis de détecter et écarter d’éventuels artefacts susceptibles de fausser les différentes données nécessaires aux calculs.

L'évaluation des personnes sur le ton mélodique était indispensable puisque c'est le ton se rapprochant du ton employé dans la vie courante. Il met en évidence les variations mélodiques de la voix et permet de voir quels peuvent être les extrêmes fréquentiels dans la prosodie de la personne. Il nous a semblé aussi intéressant d'inclure le ton monotone pour observer quelles sont les adaptations lorsque les variations mélodiques sont réduites. Plus largement, il nous a semblé important de pouvoir montrer au sujet, grâce à l'imitation de la voix transformée, le point central autour duquel sa voix devrait osciller pour être pleinement dans la zone de sécurité prosodique. Dans l'analyse globale des résultats, nous avons choisi de nous centrer davantage sur le ton mélodique qui reste le plus employé dans la vie courante. Le ton monotone ayant finalement plus un rôle d'illustration et de repère pour une imprégnation maximale, dans le cas de l'imitation de la voix transformée.

Il a été problématique également de toujours fournir un modèle orthophonique à la même fréquence fondamentale. Nous aurions pu nous enregistrer une fois et faire écouter le même enregistrement à tous les sujets mais il nous semblait important de respecter des conditions similaires à la rééducation orthophonique où le thérapeute est face au sujet et où sa production est spontanée et unique. De plus, la valeur hertzienne du modèle orthophonique n'est pas centrale dans notre étude. Ce qui importe, c'est l'adaptation du sujet au modèle donné, quelle que soit la fréquence du modèle.

Lorsque nous avons déterminé la meilleure imitation au niveau de la fréquence fondamentale moyenne entre l'imitation du modèle orthophonique et l'imitation du modèle de la voix transformée (premier histogramme des tons mélodique et monotone pour chacun des sujets), il y a parfois eu des cas où les productions étaient très proches en terme de hertz. Lorsque cela s'est produit, comme nous l'avons précisé dans le recueil et l'analyse des résultats, c'est le second histogramme qui nous a permis d'avoir une idée plus marquée.

Le protocole a été élaboré à partir de données fournies par un professionnel de l'orthophonie, il n'émane pas de notre propre expérience. Le protocole est basé sur l'utilisation du logiciel PRAAT, ce qui implique d'avoir la maîtrise des fonctionnalités utilisées.

3.2 CRITIQUES DES PRINCIPAUX RESULTATS

Nous n'avons pu rattacher les résultats présentés dans la partie pratique à des connaissances préexistantes de la littérature car il n'existe pas de publications récentes s'apparentant à notre démarche expérimentale (détermination d'une zone de sécurité prosodique entre autres). Nous tenterons donc d'expliquer les principaux résultats obtenus à l'aide de nos propres connaissances.

- **Adaptation de la production au modèle orthophonique :**

Nous supposons fortement que l'adaptation au modèle orthophonique serait possible pour l'imitateur. Cependant, nous ne savons pas ce qu'il en serait pour le reste des sujets étudiés. En effet, la voix du modèle orthophonique étant assez différente de celle de la majorité des sujets en terme de hauteur tonale notamment, nous ne savons pas qu'elle serait l'adaptation vocale des personnes vis-à-vis de cet exercice. La première personne enregistrée dans la réalisation de ce protocole a été l'imitateur qui, sur le ton mélodique, a réellement cherché à imiter la voix de l'orthophoniste, du fait de sa profession sans doute. Nous savons pertinemment que les témoins et patients masculins en particulier auraient du mal à réellement imiter le modèle orthophonique tel qu'a pu le faire l'imitateur. Il nous importait de voir comment ce modèle pouvait influencer leur production. Les résultats observés pour cet exercice laissent bien apparaître que le comportement d'imitation est possible chez tout un chacun y compris les patients dysphoniques, puisque les trois-quarts de la population étudiée parviennent à moduler leur production vocale. Cependant, il est bien sûr inadéquat pour une personne et particulièrement un patient de chercher à adopter ou tout du moins se rapprocher de la hauteur tonale de l'orthophoniste (surtout s'il s'agit d'un homme). Cela entraînera forcément un mauvais usage de la voix.

Dans cet exercice, la meilleure réalisation de l'imitation sur ton monotone nous semble logique puisque la variation de la voix est moins importante que sur le ton mélodique. Le sujet doit percevoir la hauteur imposée par le modèle orthophonique et s'y maintenir de façon stable. Sur le ton mélodique, les variations vocales étant plus importantes, les repères sont peu stables.

- **Adaptation de la production au modèle de la voix transformée :**

Les résultats obtenus pour cet exercice sont sans appel. Tous les sujets adaptent leur production au modèle de la voix transformée, y compris les patients dysphoniques. Ainsi, l'imitation est certes possible dans les deux cas (modèle orthophonique et modèle voix transformée) mais seul le modèle de la voix transformée induit systématiquement une modulation de la voix chez tout sujet. Les résultats montrent que la perception auditive et les ajustements dans la production vocale sont davantage réceptifs à la propre voix du sujet qu'à l'imitation d'un modèle externe. Nous pouvons apporter un élément de réponse expliquant ces résultats : lors de l'imitation du modèle orthophonique, le sujet doit se concentrer non seulement sur la hauteur tonale mais aussi sur l'articulation, la mélodie, le rythme du modèle. Or, dans l'imitation du modèle de la voix transformée il ne se concentre que sur la hauteur tonale imposée puisque le modèle comporte déjà sa propre articulation, sa propre mélodie et son propre rythme. Cela favorise la réussite.

Par ailleurs, même s'il est toujours surprenant de prime abord d'entendre sa voix émanant d'un enregistrement (du fait de l'absence de conduction osseuse entre autres), la voix est toujours mieux perçue par le sujet que la voix d'une tierce personne, puisque sienne.

- **Cas où l'imitation du modèle orthophonique est la meilleure et se rapproche le plus de la zone de sécurité prosodique :**

Nous retrouvons un seul sujet (un témoin) pour lequel l'imitation du modèle orthophonique est la meilleure (cela sur le ton monotone) sur l'ensemble de la population étudiée. Il nous semblait très probable que l'imitation du modèle orthophonique ne serait pas la meilleure. En effet, il est difficile pour l'orthophoniste de produire spontanément un modèle correspondant exactement au sujet. De plus, même en disposant des données de la zone de sécurité prosodique nécessaire, comment l'orthophoniste pourrait produire l'exact modèle de la voix attendue ? Et quand bien même le thérapeute s'imprégnerait du modèle respectant la zone de sécurité prosodique, il lui serait difficile de s'adapter à la fréquence vocale voulue pour le patient.

- **Cas où l'imitation du modèle de la voix transformée est la meilleure et se rapproche le plus de la zone de sécurité prosodique :**

Tous les sujets à l'exception du sujet précédemment cité réalisent donc la meilleure performance pour l'imitation de la voix transformée. Ainsi, cette réussite s'observe bien chez les sujets dysphoniques. C'est lors de cette imitation que les sujets se rapprochent le plus du modèle présenté qui leur permet *ipso facto* d'être au plus proche de leur zone de sécurité prosodique. Ce modèle objective l'étendue dans laquelle leur production doit se trouver pour rester confortable. L'intérêt de ce modèle est de donner l'exacte illustration de ce que le sujet doit produire contrairement au modèle orthophonique qui n'est pas adapté pour orienter vers la zone de sécurité prosodique.

Dans ce mémoire, pour rester le plus objectif possible et conserver la même façon de procéder pour tous les sujets, nous avons décidé de ne pas apporter d'importantes modifications manuelles aux lignes mélodiques obtenues pour le modèle de la voix transformée, même si celles-ci pouvaient être encore améliorées. Nous voulions toujours garder l'étendue fréquentielle calculée bien que celle-ci aurait pu être encore diminuée pour le confort vocal. Il est également arrivé que quelques sujets trouvent leur voix transformée un peu trop aigüe. Mais l'objectif ici était de voir si le sujet était capable d'imiter sa propre voix. Dans l'absolu, il faudrait bien sûr travailler la ligne mélodique pour l'adapter au mieux à la perception de la personne sans pour autant s'éloigner de la zone de sécurité prosodique.

- **Ecart-types et moyennes :**

Le tableau récapitulatif mettant notamment en évidence les écart-types et les moyennes obtenus pour les différents exercices sur ton mélodique et monotone nous permet d'avoir une idée plus précise des oscillations de la voix et de l'éloignement des productions par rapport à la zone de sécurité prosodique. Du fait des premiers résultats obtenus, nous savions que les résultats seraient certainement meilleurs pour l'imitation de la voix transformée. C'est effectivement le cas. Cependant il nous semblait opportun de pouvoir objectiver davantage le taux des oscillations vocales et l'éloignement de la zone de sécurité prosodique.

Ainsi, pour le ton mélodique, il est intéressant de voir que dans leur production spontanée, certains sujets sont déjà très éloignés de leur zone sécurisée. Cela souligne un mauvais comportement vocal dans la vie quotidienne pouvant entraîner un malmenage vocal et favoriser l'apparition de pathologies vocales. Qui plus est, cela se retrouve pour 3 patients dysphoniques et le témoin ayant déjà souffert d'un épisode dysphonique. Il aurait été propice de pouvoir enregistrer d'autres témoins ayant souffert de dysphonie pour voir si ce fait était systématique. Quoi qu'il en soit, ces données montrent que les personnes ayant souffert ou souffrant de dysphonie dans notre population ont spontanément tendance à mal placer leur voix par rapport à la zone de sécurité prosodique. De plus, le modèle orthophonique ne les aide absolument pas à se rapprocher de cette zone. Seul le modèle de la voix transformée permet de s'en rapprocher considérablement.

Globalement, les données attestent bien que l'imitation du modèle orthophonique augmente l'importance des oscillations vocales par rapport à l'imitation de la voix transformée tout en favorisant l'éloignement de la zone de sécurité prosodique.

A partir de l'ensemble de ces résultats, nous pouvons affirmer que tout individu, même dysphonique, est capable de moduler sa voix en vue de réaliser une imitation. En ce sens, tout sujet est capable de réaliser un travail d'écoute, de réaliser un contrôle moteur sur ses organes phonateurs. Chez tout individu, la boucle audio-phonatoire peut entrer en action pour favoriser l'ajustement de la musculature laryngée et des résonateurs. Il apparaît également que le sujet sain comme dysphonique réalise un meilleur contrôle de sa prosodie lorsqu'il s'auto-imité. De ce fait, sa production est plus proche de la zone de sécurité prosodique définie. La perception de sa propre voix permet d'induire de meilleurs ajustements laryngés notamment. Dans la zone sécurisée, tous les phonèmes du français ont un rendement optimal. Elle constitue une partie de la dynamique vocale où il est certain que la voix du sujet est le moins en souffrance. Les capacités d'imitation étant présentes chez tout un chacun, elles constituent un bon tremplin pour prendre conscience et adopter l'intervalle fréquentiel favorable pour la voix du sujet. L'imitation de la voix respectant la zone sécurisée est donc le meilleur moyen pour émettre dans l'intervalle fréquentiel non pathologique. Cela est important pour la voix du patient dysphonique qui devra se replier au sein de cet intervalle pour être préservée.

3.3 OUVERTURE AU CHAMP ORTHOPHONIQUE

Le but premier de ce mémoire était de voir si les capacités d'imitation étaient présentes chez tout un chacun y compris les personnes dysphoniques. L'objectif était également de montrer que tout sujet, même dysphonique, réalise un meilleur contrôle de sa prosodie et se retrouve plus proche de sa zone de sécurité prosodique lorsqu'il s'auto-imité.

Il apparaît, à l'issue de cette étude, que le patient dysphonique est effectivement capable de moduler sa voix en vue de réaliser une imitation. Par ailleurs, l'écoute de sa propre voix respectant la zone de sécurité prosodique définie, favorise un meilleur contrôle de sa prosodie et permet de diminuer les variations prosodiques extrêmes et néfastes pour sa voix.

L'intérêt de ce travail pour la rééducation orthophonique de la voix réside en différents points. Tout d'abord, il s'agit d'offrir au praticien un moyen de se familiariser avec un logiciel permettant de bénéficier d'une analyse objective de la voix. Puis, à partir de l'outil informatique, de disposer d'un moyen supplémentaire d'objectivation de la voix tout en déterminant la zone de sécurité prosodique adaptée à chaque patient pour lui permettre d'avoir le modèle exact d'une voix confortable (en complément et non en remplacement des techniques actuelles de rééducation orthophonique dans le domaine de la voix). Enfin, de pouvoir mettre à disposition des patients un support audio (sur clé USB par exemple) comme référence et modèle dans leur vie quotidienne pour favoriser l'imprégnation de la voix sécurisée. L'objectif étant à long terme d'entraîner le patient à utiliser sa zone de sécurité prosodique dans la conversation spontanée (grâce à une écoute répétée de l'enregistrement audio fourni) et de s'y tenir dès qu'il ressent une fatigue vocale.

CONCLUSION

A travers ce mémoire, nous nous sommes attachés à mieux connaître ce que recouvre la notion d'imitation. Nous avons pour cela tenté de la définir et de l'aborder sous différents aspects en nous orientant vers des domaines divers et variés. Il s'est avéré que tout individu présente des capacités d'imitation, lesquelles peuvent être utilisées à bon escient pour travailler la voix. Puis, partant du principe qu'il existe une zone de sécurité prosodique pour chaque personne où la voix a le meilleur rendement, où chaque phonème produit reste dans des fréquences confortables pour la voix du sujet, où toute production demeure dans la zone de la dynamique vocale la plus sécurisée, nous avons cherché à savoir si l'imitation de cette zone entraînerait une production plus adaptée que l'imitation du modèle orthophonique traditionnel en terme d'illustration de confort vocal. Nous avons observé que les sujets sont davantage réceptifs à l'auto-imitation. En effet, seule l'auto-imitation permet à la fois de réaliser des productions très proches du modèle grâce à l'intervention de la boucle audio-phonatoire favorisant des ajustements musculaires pertinents et à la fois d'objectiver les fréquences à respecter pour ne pas malmener la voix. La définition de cette zone de sécurité prosodique constitue donc un support non négligeable pour une prise de conscience des modulations vocales adaptées. Ce moyen constitue une aide supplémentaire à disposition du thérapeute pour donner le repère vocal dont le patient a besoin pour objectiver sa voix et l'utiliser dans les meilleures fréquences. L'objectif premier étant d'amener le patient à demeurer ou se replier dans la zone vocale sécurisée lui correspondant, pour plus largement, l'aider à retrouver ou acquérir une efficacité vocale, améliorer ou prévenir l'apparition de lésions consécutives à une mauvaise utilisation de la voix sur le long terme.

Cependant, nous avons choisi dans cette étude d'observer uniquement l'auto-imitation à un temps « t ». Il serait donc intéressant dans une autre étude, d'évaluer après un certain temps d'entraînement comment le sujet s'est approprié cette voix idéale. Et éventuellement, d'affiner sa zone de sécurité prosodique dans la suite de la rééducation orthophonique, en fonction de l'adaptation positive.

BIBLIOGRAPHIE

Ouvrages :

- ABITBOL J. « L'odyssée de la voix », Paris, éditions Robert Laffont, 2005, 515 p.
- BAUDONNIERE P.M. « Le mimétisme et l'imitation », Paris, éditions Flammarion, 1997, 128 p.
- BONNEVILLE L. et GROSJEAN S. « Repenser la communication dans les organisations », Paris, éditions L'Harmattan, 2007, 300 p.
- BOUCHE A. et CUILLERET J. « Anatomie topographique, descriptive et fonctionnelle », 2^{ème} édition SIMEP, 1998, 598 p.
- LECLERC BUFFON G.L. « Oeuvres choisies de Buffon », tome 2, Paris, Librairie de Firmin Didot Frères.
- CASTAREDE M.F. « La voix et ses sortilèges », Paris, éditions Les belles lettres, 2004, 280 p.
- CORNUT G. « La voix », Paris, éditions Presses Universitaires de France, 3^{ème} édition, collection « Que sais-je ? », septembre 1990, 128 p.
- DEJONCKERE Ph., ESTIENNE F. et BARBAIX M.T. « Précis de pathologie et de thérapeutique de la voix », Paris, éditions P.Delarge, 1980.
- DHILLON R.S. et EAST C.A. « Oto-rhino-laryngologie et chirurgie cervico-faciale », Paris, éditions Elsevier, 2008, 164 p.
- GOSLING J.A., HARRIS P.F., WHITMORE I. et WILLAN P.F.T. « Anatomie humaine : atlas en couleurs », Bruxelles, éditions de Boeck, 2003, 396 p.
- HEUILLET-MARTIN G., GARSON-BAVARD H. et LEGRE A. « Une voix pour tous » tome I., Paris, éditions Solal, 1995, 215 p.
- FOURNIER C. « La voix un art et un métier », CCL Editions, 1989, collection Jardins d'Isère, 258 p.
- KOLB B. et WHISHAW I.Q. « Cerveau et comportement », Bruxelles, éditions De Boeck, 2002, 672 p.
- LE HUCHE F. et ALLALI A. « La voix tome 1. Anatomie et physiologie des organes de la voix et de la parole », Paris, éditions Masson, 2001, 199 p.
- MCFARLAND D.H. et NETTER F.H. « L'anatomie en orthophonie », Paris, éditions Masson, 2006, 226 p.
- MOORE K.L. et DALLEY F. « Anatomie médicale, aspects fondamentaux et applications cliniques », Bruxelles, édition De Boeck, 2006, 1216 p.

- MOREAU M.L. « Sociolinguistique : les concepts de base », Wavre Belgique, éditions Mardaga (2), 1997, 312 p.
- PIAGET J. « La formation du symbole chez l'enfant, imitation jeu et rêve, image et représentation », Neuchâtel, éditions Delachaux et Niestle, 1970.
- RENARD X. « Précis d'audioprothèse : production, phonétique acoustique et perception de la parole », Paris, éditions Masson, 2008, 411 p.
- RIZZOLATTI G. et SINIGAGLIA C. « Les neurones miroirs », Paris, éditions Odile Jacob, 2008, 236 p.
- RONDAL J.A. et SERON X. « Troubles du langage ; Bases théoriques, diagnostic et rééducation », Wavre Belgique, éditions Mardaga, 2000, 840 p.
- RONDAL J.A. « Manuel de psychologie de l'enfant, Wavre Belgique, éditions Mardaga, 1999, 644 p.
- SALIBA E., HAMAMAH S. et GOLD F. « Médecine et biologie du développement : du gène au nourrisson » Paris, édition Masson, collection Gynécologie obstétrique, 2001, 431 p.
- SARFATI J. « Soigner la voix », éditions Solal, collection Voix parole langage, 2002, 127 p.
- SUNDBERG J. « The Science of the Singing Voice », Northern Illinois, University Press, 1987, 216 p.
- THOMAS R.M et MICHEL C. « Théories du développement de l'enfant : études comparatives », Bruxelles, éditions De Boeck, 1994, 576 p.
- TORTORA G.J. et GRABOWSKI S.R. « Principes d'anatomie et de physiologie », Bruxelles, éditions De Boeck, 2002, 1256 p.
- VERMETTE J. et CLOUTIER R. « La parole en public : savoir être, savoir faire », Québec Canada, éditions Les presses de l'université Laval, 1992, 208 p.
- WANG S. « Singing voice : bright timbre, singer's formants and larynx positions », Stockh. Mus. Acoust. Congress, 1983.
- WHITAKER R.H. et BORLEY N.R. « L'anatomie », Bruxelles, éditions De Boeck, 2003, 236 p.

Articles / Revues :

- ADRIEN J.L. « Autisme du jeune enfant: développement psychologique et régulation de l'activité », collection Autisme, expansion scientifique française, 1996, 193 p.
- ARBIB M.A. « Beyond the mirror system : imitation and evolution of language », in C. NEHANIV, K. DAUTENHAN (sld.), Imitation in Animals and Artifacts, 2002, Boston (MA), MIT Press, pp.229-280.
- ARBIB M.A. « From monkey-like action recognition to human language : an evolution framework for neurolinguistics », 2005, Behavioral and Brain Sciences, 28, pp. 105-167.
- BANDURA A. « Social learning through imitation », 1962, in M.R. JONES (Edit), Nebraska Symposium on Motivation, Lincol: University of Nebraska Press.
- BARD A.K. et RUSSEL C.L. « Evolutionary foundations of imitation: Social, cognitive and developmental aspects of imitative processes in non-human primates », 1999, in J. NADEL & G.BUTTERWORTH (Edit.), Imitation in infancy, Cambridge University Press, pp.89-123.
- BEHROOZMAND R. et al. « Vocalization-induced enhancement of the auditory cortex responsiveness during voice F0 feedback perturbation », International Federation of Clinical Neurophysiology, Elsevier 2009, pp.1303-1312.
- BEKKERING H., WOHLSCHLÄGER A. et GATTIS M. « Imitation of gestures in children in goal-directed », 2000, Quarterly Journal of Experimental Psychology, 53A, pp. 153-164.
- BURLING R. « The talking ape », 2005, New York, Oxford University Press.
- BYRNE R. « Seeing actions as hierarchically organized structures : great ape manual skills », 2002, in PRINZ W., MELTZOFF A.N. (sld.), The imitative Mind : Development, Evolution and Brain Bases, Cambridge, Cambridge University Press pp.122-140.
- BYRNE R. « Imitation as behavior parsing », 2003, in Philosophical Transactions of the Royal Society of London, Sreie B, 358, pp.529-536.
- BYRNE R. et RUSSON A.E. « Learning by imitation : a hierarchical approach », 1998, Behavioral Brain Sciences, 21, pp.667-712.
- CALVO-MERINO B., GLASER D.E., GREZES J., PASSINGHAM J. et HAGGARD P. « Action observation and acquired motor skills : an fMRI study with expert dancers », 2005, Cerebral Cortex, 15, 8, pp.1243-1249.

- CATTANEO L. et RIZZOLATTI G. « The mirror neuron system », Arch neurol, vol 66 (NO.5), mai 09, pp.557-560.
- CLEMENTSON-MOHR D. « Toward a social-cognitive explanation of imitation development » in G. BUTTERWORTH & P. UGHT (eds), Social cognition: studies of the development of understanding. Brighton, Harvester Press, 1982.
- CORBALLIS M.C. « Mirror neurons and the evolution of language », Brain and Language, 2009, doi:10.1016/j.bandl.2009.02.002.
- DE CASPER A.J. et FIFER W.P. « Of human bonding : newborns prefer their mothers' voice », Science, Vol 208, Issue 4448, 1174-1176, 1980.
- DIAS N. « Imitation et anthropologie », éditions MSH, terrain 44, 2005, 170 p.
- ERIKSSON A. et WRETLING P. « How flexible is the human voice? – A case study of mimicry », 1997, in Proc. EUROSPEECH '97, Vol. 2, 1043–1046
- FERRARI P., GALLESE V., RIZZOLATTI G. et FOGASSI L. « Mirror neurons responding to the information of ingestive and communicative mouth actions in the monkey ventral premotor cortex », 2003, European Journal of Neuroscience, 17, pp.1703-1704.
- FONTAINE R. « Imitative skills between birth and six months », 1984, Infant Behavior and Development, n° 7, pp. 323-333.
- FRANK A.R. et al. « Learning the "special note": Evidence for a critical period for absolute pitch acquisition », Music Perception, vol. XXI, n° 1, Automne 2003.
- GILES H., COUPLAND N. et COUPLAND J. « Accommodation theory: Communication, context, and consequence » in H. Giles, N. Coupland, & J. Coupland, editors, Contexts of Accommodation: Developments in Applied Sociolinguistics, pages 1–68. Cambridge University Press, Cambridge, UK, 1991.
- GUENTHER F.H., GHOSH S.S. et TOURVILLE J.A. « Neural modeling and imaging of the cortical interactions underlying syllable production », Brain Lang 2006; 96:280–301.
- HICKOK G., BUCHSBAUM B., HUMPHRIES C., et MUFTULER T. « Auditory-motor interaction revealed by fMRI: Speech, music, and working memory in area Spt.», 2003, Journal of Cognitive Neuroscience, 15, 673–682.
- HICKOK G. et POEPEL D. « The cortical organization of speech processing » 2007, Nature reviews Neuroscience, 8, pp.393-402.
- LEBIB R. et BAUDONNIERE P.M. « L'imitation vocale du nouveau-né au nourrisson de 3-4 mois : les «premiers pas» vers l'acquisition du langage », Devenir n°2, 1999, vol. 11, n°2, pp. 37-52.

- LORENZ K. « Companionship in bird life », 1957, in C.H. SCHILLER & K.S. ASHLEY, *Instinctive behavior : The development of a modern concept*, International Universities Press, New-York, pp.83-128.
- MEJVALDOVA J. « Caractéristiques temporelles de la parole imitée », Institut de Phonétique, Université Charles nám. J. Palacha 2, Prague, République Tchèque.
- MELTZOFF A.N. et MOORE M.K. « Imitation et développement humain : les premiers temps de la vie », *Terrain* n° 44, 2005.
- NAGY E. et MOLNÁR P. « Homo Imitans or Homo Provocans? Human imprinting model of neonatal imitation », *Science Direct, Infant Behavior and Development*, vol 27, février 2004, pp.54-63.
- PAPOUSEK M. et PAPOUSEK H. « Forms and functions of vocal matching in interactions between mothers and their pre-canonical infants », 1989, *First Language*, 9, pp.137-158.
- PILLOT C., Institut de Phonétique de Paris et hôpital Laënnec, « Pourquoi le singing-formant ? Données scientifiques et hypothèse musicologiques relatives à son apparition », *Médecine des Arts*, 1998.
- PAINTER C. « Laryngeal functions in speech », in CUMMINGS, C., FREDRICKSON, J., MARKER, L., KRAUSE, C., SCHULLER, D., Dirs, *Otolaryngology, Head and Neck Surgery*. St. Louis : Mosby, 1986.
- PETRIDES M. et PANDYA D.N. « Comparative architectonic analysis of the human and the macaque frontal cortex », in F. BOLLER, J. GRAFMAN (sld.), *Handbook of Neuropsychology*, 1997, pp.17-58.
- PITTAM J. « Voice in Social Interaction ; An Interdisciplinary Approach », 1994, Thousand Oaks, SAGE Publications.
- PRINZ W. « A common-coding approach to perception and action », in O. NEUMANN, W. PRINZ (sld.) *Relationship between Perception and action : Current Approaches*, Berlin, Springer, 1990, pp.167-203.
- ROGERS L.S. et KAPLAN G.T. « Comparative vertebrate cognition : are primates superior to non-primates ? », 2004, Kluwer academie, Plenum Publishers, 386 p. pp.115 à 118.
- SCOTTO DI CARLO N. « Oreille absolue et mémoire proprioceptive », *Medecine des Arts*, 2003, n°44, pp.7-15.
- SCHERER K.R. « Personality markers in speech », 1979, in K.R.Scherer and H. Giles (eds.) *Social markers in speech*. Cambridge, Cambridge University Press: 147-209.

- SHUSTER L.I. et DURRANT J.D. « Toward a better understanding of the perception of self-produced speech », 2003, Journal of communication disorders 36, Elsevier, 1-11.
- SRUNER J.S. « Le développement de l'enfant. Savoir-faire, savoir-dire » (traduction: M. Deleau). Paris, PUF, 1983.
- STERN D.N. « The role of feelings for an interpersonal self », 1993, in U. NEISSER (Ed.), The perceived self, New-York, Cambridge University Press, 205-215.
- TREVARTHEN P. et AITKEN D. « Intersubjectivité chez le nourrisson : recherche, théorie et application clinique », *Devenir* 2003/4, Volume 34, pp. 309-428.
- TONNDORF J. « Bone conduction », 1972, in J.V.Tobias (ED.) Foundations of modern auditory theory (vol 2, pp 197-237) New York : academic press.
- TOYOMURA A., KOYAMA S., MIYAMAOTO T., TERAOKA A., OMORI T., MUROHASHI H. et KURIKI S. « Neural correlates of auditory feedback control in human », 2007, Neuroscience 146, pp.499–503.
- VON BESEKY G. « The structure of the middle ear and the hearing's of one's own voice by bone conduction », Journal of the acoustical society of America, 21, 1949, pp.217-232.
- WARREN D.K., PATTERSON D.K., PEPPERBERG I.M., « Mechanisms of American English vowel production in a grey parrot (psittacus erithacus) », The Auk, Washington, janvier 1996, Vol 113, Iss 1, pg 41, 4 pgs.
- WINNYKAMEN F. et FAFONT L. « Place de l'imitation-modélisation parmi les modalités relationnelles d'acquisition : cas des habiletés motrices », Revue française de pédagogie, n°92 juillet-août-septembre 1990, pp.23-30.
- WINNYKAMEN F. « Imitation et acquisitions par observation : études récentes et perspectives » in J. BIDEAUD & RICHELLE (eds), Psychologie Développementale. Problèmes et Réalités, Bruxelles, Mardaga, 1985.
- WRETLING P. et ERIKSON A. « Is articulatory timing speaker specified? – Evidence from imitated voices », 1998, in Proc Phonetic, Department of Phonetics, Umea University, Sweden.
- ZENTALL E. « Defining imitation by exclusion », 1996, in B.G GALEF, JR. & C.M. HEYES (Edit.), Social learning in animals: The roots of culture, New York: Academic Press, pp.221-243.
- ZETTERHOLM E. « Same speaker – different voices. A study of one impersonator and some of his different imitations », 2006, Proceedings SST2006, Auckland, New Zealand, Dec 6-8 2006.

- ZETTERHOLM E., ELENIOUS D. et BLOMBERG M. « A comparison between human perception and a speaker verification system score of a voice imitation », 2004, Proceedings SST2004, Sydney, Australia, Dec 8-10 2004: 393-397.

Mémoires pour l'obtention du certificat de capacité d'orthophonie :

- COUDIERE C. « De l'utilité des logiciels pour la voix en rééducation de dysphonies dysfonctionnelles », Toulouse, 2003.
- GUITTOT L. et PERON L. « La rééducation vocale de la personne transsexuelle : intérêts et limites de l'orthophonie », Lille, 2003.
- JACQUEMIN F. « Les imitateurs, comment font-ils ? Essai de mise en évidence des procédés utilisés lors de l'imitation de quatre voix parlées », Nancy, 1989.
- ROUBLOT P. « Analyse comparative subjective et objective de la voix avant et après bloc interscalénique du plexus brachial. Implications pratiques dans la prise en charge post-opératoire des patients anesthésiés », Nancy, 2003.

Thèse :

- DEVOUCHE E. « La situation d'imitation à 8 et 12 mois. Aspects sociaux et cognitifs », Thèse de doctorat, sous la direction de Mme Marie-Germaine PECHEUX, (janvier 2000), université paris V – René DESCARTES, UFR Institut de Psychologie.

Matériel multimédia :

- BOERSMA P. et WEENINK D., Logiciel PRAAT, Institut de sciences phonétiques de l'université d'Amsterdam, 1996.

Sites consultés

- <http://www.medix.free.fr/cours/physiologie-phonation.php> consulté en juillet 2009.
- <http://orlinfo.free.fr/anat/larynx/extr.htm> consulté en juillet 2009.

- <http://www.automatesintelligents.com/labo/2005/mar/neuronesmiroir.html> consulté le 15 août 2009.
- <http://www.medix.free.fr/cours/physiologie-phonation.php> consulté en octobre 2009.
- http://www.chambon.ac-versailles.fr/science/faune/phy_a/resp/voixoiseau.htm consulté en décembre 2009.
- <http://fr.wikipedia.org>, consulté en janvier 2010.
- http://aune.lpl.univ-aix.fr/~ilc/16janv09_Fiasson.pdf, consulté en janvier 2010.

ANNEXES

RESULTATS :

[ma] = } min : max :

[mi] = } min : max : Ecart voulu : [/]

[mu] = } min : max :

F max voulue :

F max observée :

F min voulue :

F min observée :

δ voulu (F max – F min) :

δ observé :

$\alpha = \delta \text{ voulu} / \delta \text{ observé} =$

$\beta = F \text{ max voulue} - (F \text{ max observée} \times \alpha) =$

Moyenne pour la ligne monotone = $F \text{ min voulue} + (\delta \text{ voulu} / 2) =$

PROTOCOLE AVEC L'EXEMPLE DE M.GREGORIO :

1. Enregistrement des glissades [ma], [mi], [mu]

- **Consigne** : « Prenez une inspiration de façon à produire une glissade montante sur un [ma] de durée suffisante, d'intensité et de hauteur confortables, sans passer en voix chantée. »

- **Procédure** :
- New
 - Record mono sound
 - Sampling frequency 44 000 Hz
 - Record =} enregistrer la glissade
 - Stop
 - Rename « sound ma »
 - Save to list

Faire la même chose pour [mi] et [mu].

On a donc trois objets : « sound ma », « sound mi », « sound mou ».

2. Enregistrement d'une phrase sur consigne

- **Consigne** : « Lire la phrase inscrite sur la feuille présente devant vous sur un ton monotone. » (La phrase est : « Je voudrais partir en Italie pour admirer la Tour de Pise »)

Enregistrer comme pour les glissades, on obtient un 4^{ème} objet : « sound phrase monotone. »

- **Consigne** : « Lire la même phrase sur un ton mélodique »

Enregistrer comme précédemment, on obtient un 5^{ème} objet : « sound phrase mélodique. »

C'est de cet objet dont on se servira plus tard pour transformer la voix du sujet et définir sa zone de sécurité prosodique.

3. Enregistrement avec modèle orthophonique

- **Consigne** : « J'utilise la même phrase que précédemment. Imitiez le modèle monotone que je vous donne. Vous pouvez imitez dès que j'ai terminé la phrase. »

Enregistrer la production de l'orthophoniste suivie immédiatement de la production du patient, on obtient un 6^{ème} objet : « **sound imitation ortho monotone.** »

- **Consigne** : « J'utilise la même phrase que précédemment. Imitiez le modèle mélodique que je vous donne. Vous pouvez imiter dès que j'ai terminé la phrase. »

Enregistrer la production de l'orthophoniste suivie immédiatement de la production du patient, on obtient un 7^{ème} objet : « **sound imitation ortho mélodique.** »

4. Calculs manipulation courbe intonative du patient

- **Procédure pour chaque glissade** :

- Edit la glissade souhaitée
- Show pitch (pour voir F0 en bleu)
- Si la courbe présente de trop grosses ruptures par rapport à ce qui est entendu, faire :
 - Menu pitch
 - Pitch settings
 - Modifier la valeur max et/ou min pour que la voix soit bien comprise dedans (en fonction des valeurs affichées à droite en Hz).
- Sélectionner le [ma]
- Menu Pulses
- Voice Report

On note sur la feuille prévue à cet effet (annexe n°1) le min et le max pour chaque voyelle. On peut alors déduire l'écart prosodique sécurisé.

173

- **Calcul des constantes pour écraser la voix** :

- $\alpha = \delta \text{ voulu} / \delta \text{ observé}$

Dans l'exemple : il s'agit de $\alpha = 74 / 102 = 0,72$

- $\beta = F_{\text{max}} \text{ voulue} - [F_{\text{max}} \text{ observée} \times \alpha]$

Dans l'exemple : il s'agit de $\beta = 172 - [209 \times 0.72] = 22$

- **Application à la phrase mélodique** :

- Sélectionner l'ensemble de la phrase
- Menu Pitch
- Multiply pitch frequencies = rentrer le nombre trouvé pour α (ici 0.72)
- Ok
- Menu Pitch
- Shift pitch frequencies = rentrer le nombre trouvé pour β (ici 22)
- Ok

La voix se transforme et reste ainsi dans le domaine prosodique souhaité.

On obtient donc un 8^{ème} objet : « **manipulation voix transformée mélodique** »

5. Enregistrement imitation voix transformée

- **Consigne** : « Vous allez entendre votre voix enregistrée lors de l'épreuve sur consigne où vous lisez sur un ton mélodique. Votre voix a subi des modifications de façon à rester dans la zone de sécurité prosodique. Imitiez votre voix dès mon signal. »

Enregistrer la production du patient. On obtient un 9^{ème} objet : « **sound imitation patient voix transformée mélodique**. »

6. Modification de la voix pour la rendre monotone

- Utilisation de la voix mélodique enregistrée sur consigne (étape 2), c'est toujours la même.
- To manipulation
 - Partir de l'écart de confort établi et trouver le milieu de cet écart (*dans l'exemple*, [98 Hz – 172 Hz] soit 74 Hz d'écart)
 - Prendre le milieu de cet écart (*ici*, $74 / 2 = 37$)
 - Ajouter à la fréquence min (*ici*, $98 + 37 = 135$ Hz)

Il faut donc que la ligne monotone se situe à **135** Hz.

- Surligner la phrase
- Menu pitch
- Pitch
 - Add pitch point at : - Time = mettre un peu avant le début de la phrase
 - Frequency = fréquence définie par l'écart (*ici*, 135 Hz)
- ok
- Menu pitch
- Remove pitch point
- ok

On obtient un 10^{ème} objet : « **manipulation voix transformée monotone.** »

7. Enregistrement de l'imitation de la voix transformée monotone

- **Consigne** : « Vous allez entendre votre voix enregistrée lors de l'épreuve sur consigne où vous lisez sur un ton mélodique. Votre voix a subi des modifications de façon à rester dans un ton monotone, dans la zone de sécurité prosodique. Imitiez votre voix dès mon signal. »

Enregistrer la production du patient. On obtient un 11^{ème} objet : « **sound imitation patient voix transformée monotone.** »

TABLEAU RECAPITULATIF DES RESULTATS POUR L'EXEMPLE :

[ma] = } min : 95 Hz max : 172 Hz

[mi] = } min : 98 Hz max : 175 Hz

Ecart voulu : [98 Hz / 172 Hz]

[mu] = } min : 96 Hz max : 219 Hz

Fmax voulue = 172 Hz

Fmax observée = 209 Hz

Fmin voulue = 98 Hz

Fmin observée = 107 Hz

δ voulu = 172 – 98 = 74 Hz

δ observé = 209 – 107 = 102 Hz

- α = δ voulu / δ observé = 74 / 102 = 0,72

- β = Fmax voulue – [Fmax observée x α] = 172 – [209 x 0.72]
= 22

Ligne monotone obtenue avec : Fmin voulue + [δ voulu / 2] = 98 + 74 / 2
= 98 + 37
= 135 Hz

**« JE VOUDRAIS PARTIR EN ITALIE
POUR ADMIRER LA TOUR DE PISE. »**

QUESTIONNAIRE PATIENTS :

I – Généralités :

- 1- Initiales :
- 2- Age :
- 3- Sexe :
- 4- Profession :
- 5- Usage de la voix au quotidien :
- 6- Pratique de la voix chantée :
- 7- Antécédents vocaux :

II- Impressions concernant l'exercice proposé :

- 8- Trouvez-vous difficile d'imiter l'orthophoniste, de placer votre voix en fonction ?
- 9- Que pensez-vous de la nouvelle voix proposée ?

III- Parcours orthophonique :

- 10- Cause(s) de la dysphonie ayant amené à consulter :

QUESTIONNAIRE IMITATEUR :

I- GENERALITES :

- 1- A quel âge avez-vous réalisé votre première imitation ?
- 2- Pourquoi ? De façon spontanée, ludique, contrainte ... ?
- 3- Quelle a été la première voix imitée et pourquoi ?
- 4- De combien de voix disposez-vous ? Y a-t-il un fil conducteur entre ces voix ?
- 5- Qu'est-ce qui rend une voix impossible ou difficile à réaliser ?
- 6- Quand peut-on dire qu'une voix est acquise (si elle peut l'être dans l'absolu) ?

II- PRATIQUE GENERALE :

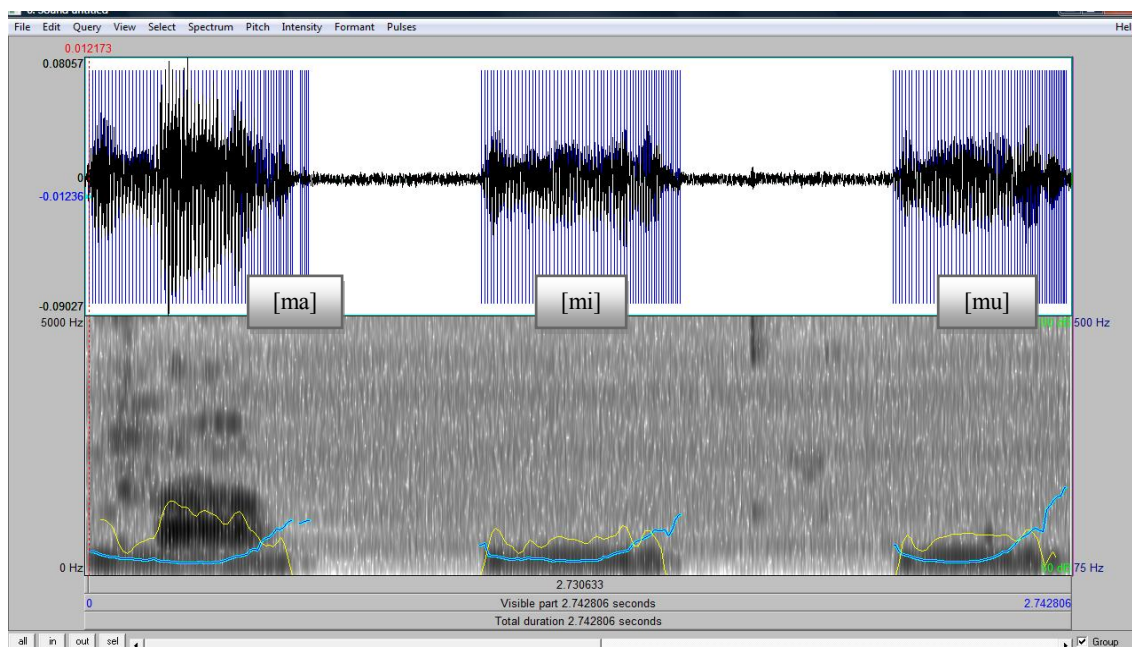
- 7- Suivez-vous un entraînement vocal ?
- 8- Comment travaillez-vous une voix ? A quelle fréquence ?
- 9- La procédure utilisée pour réaliser/s'approprier une voix est-elle toujours la même ?
- 10- Comment vous adaptez-vous quand vous subissez un facteur pathologique vocal susceptible de diminuer les performances ?
- 11- Quelles sont les possibilités de progression dans l'imitation en général ? Liées au facteur temps, à l'entraînement ... ?
- 12- A votre avis, même si vous ne pratiquez pas l'imitation de la voix parlée, celle-ci vous semble-t-elle plus difficile que l'imitation de la voix chantée ? Pourquoi ?

III- IMPRESSIONS SUR L'EXERCICE :

- 12- Qu'est-ce qui vous a semblé difficile à appréhender ?
- 13- Comment avez-vous agi en conséquence ? Quelle adaptation ?
- 14- Quel(s) conseil(s) donneriez-vous à un patient potentiel pour s'approprier sa nouvelle voix, c'est-à-dire imiter la voix "sécurisée" qu'on lui impose ?

SPECTROGRAMMES ET LIGNES MELODIQUES TYPES UTILISEES DANS NOTRE ETUDE (illustration à partir des productions de M.Grégorio)

1- Spectrogramme des glissades [ma], [mi], [mu] :



[ma] : min = 95 Hz

max = 172 Hz

[mi] : min = 98 Hz

max = 175 Hz

[mu] : min = 96 Hz

max = 219 Hz

F max voulue : 172 Hz

F min voulue : 98 Hz

δ voulu : 74 Hz

Fmax observée : 209 Hz

F min observée : 107 Hz

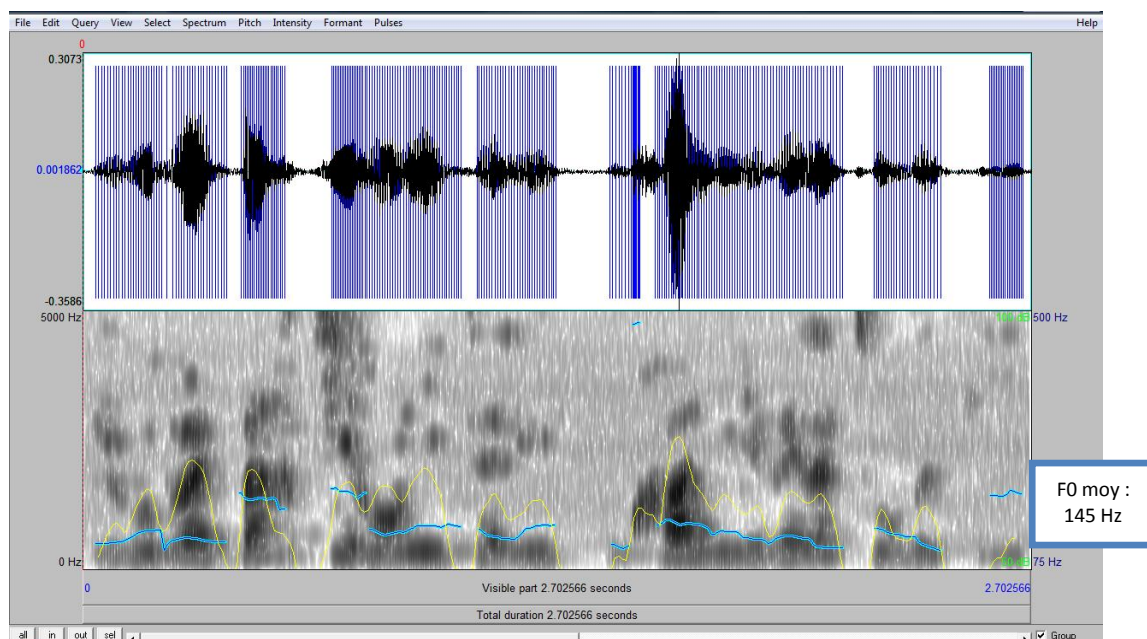
δ observé : 102 Hz

$$\alpha = 74 / 102 = 0.72 \text{ Hz}$$

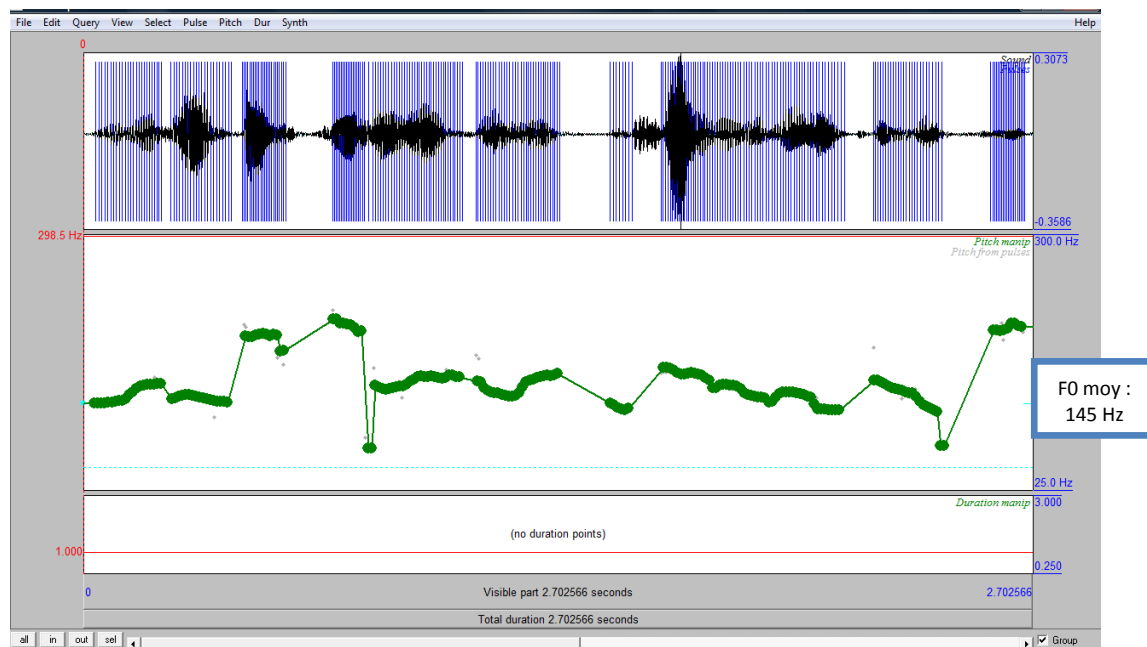
$$\beta = 172 - (209 \times 0.72) = 22$$

Moyenne pour la ligne monotone = $98 + 74/2 = 135 \text{ Hz}$

1- Spectrogramme et ligne mélodique de M.GREGORIO pour l'épreuve spontanée mélodique :



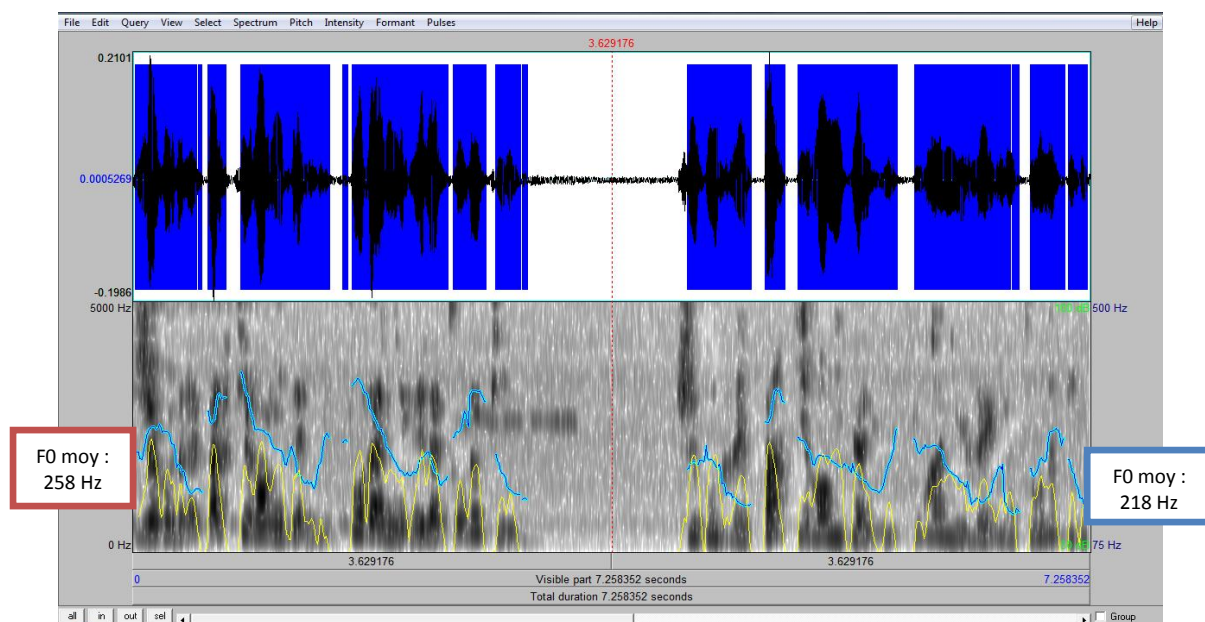
Spectrogramme de la production mélodique spontanée de M.GREGORIO



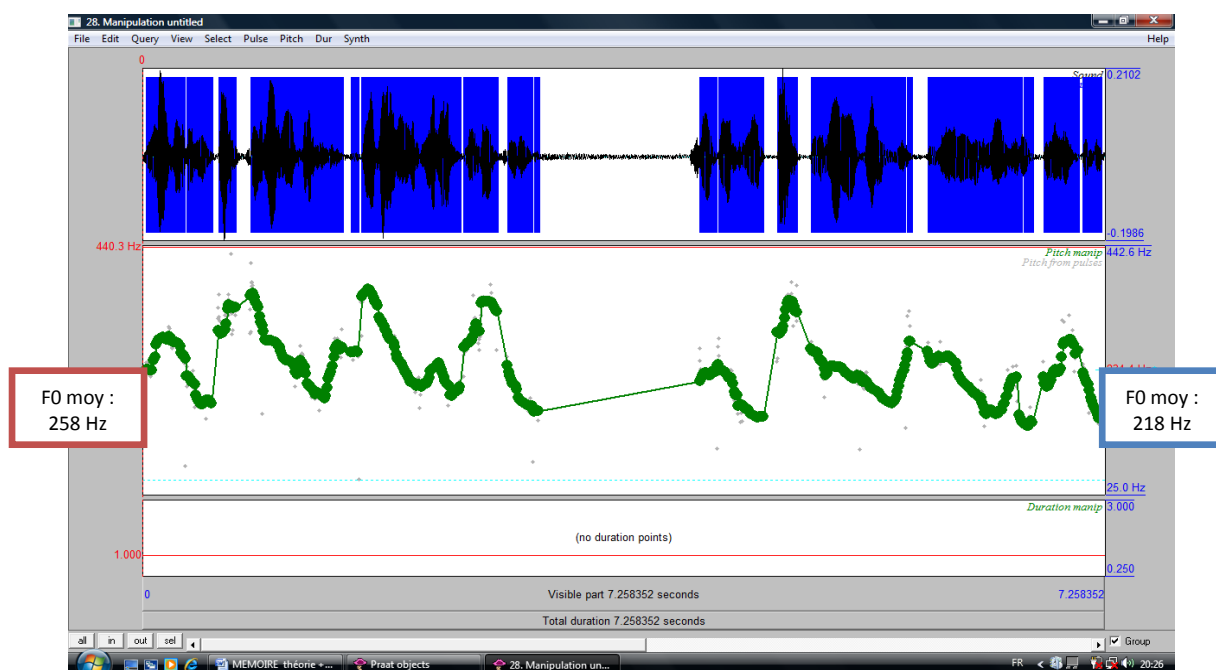
Ligne mélodique de la production spontanée de M.GREGORIO

2- Spectrogramme puis ligne mélodique pour l'épreuve d'imitation du modèle orthophonique mélodique :

Nous avons réalisé les enregistrements à la suite, avec en premier le modèle de l'orthophoniste (à gauche) suivi de l'imitation de M.GREGORIO (à droite).



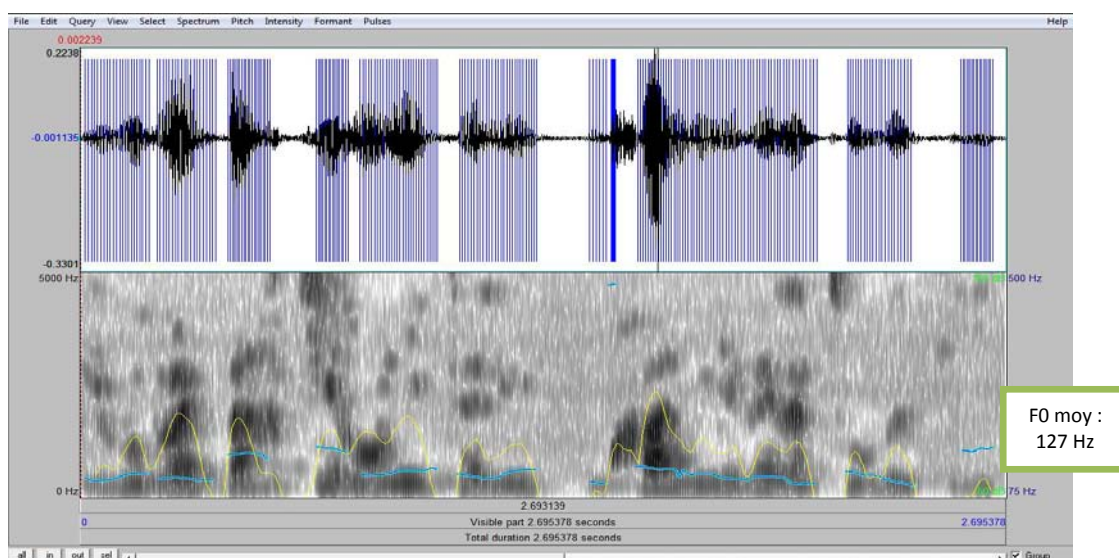
Spectrogramme du modèle orthophonique mélodique suivi de l'imitation de M.GREGORIO



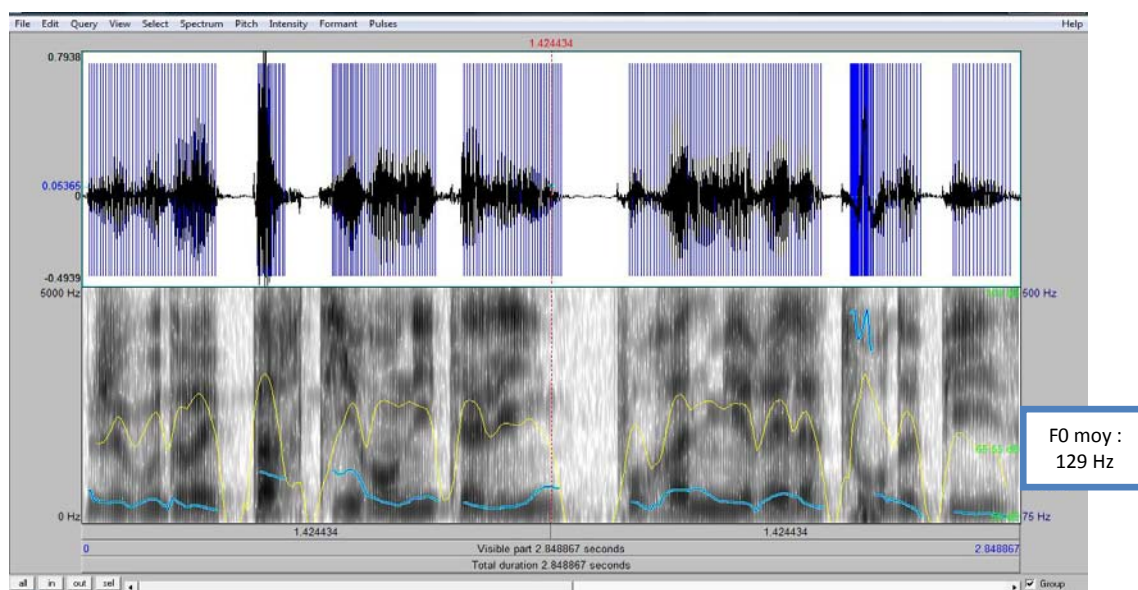
Ligne mélodique du modèle orthophonique suivi de l'imitation de M.GREGORIO

3- Spectrogramme et ligne mélodique pour l'épreuve d'imitation de la voix transformée mélodique :

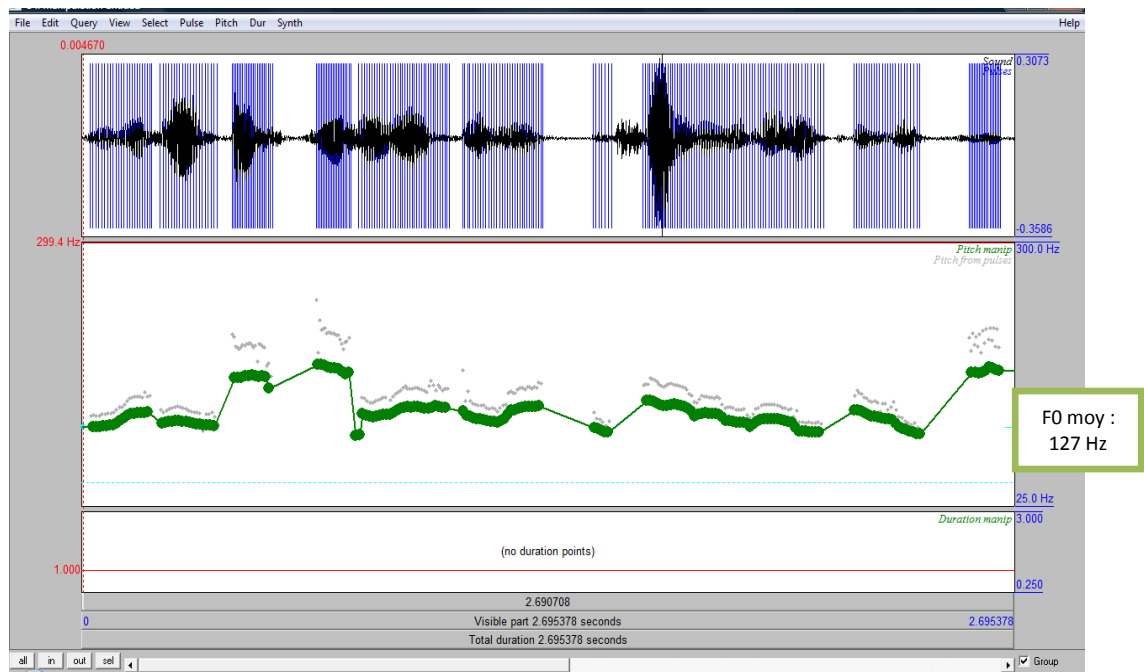
Comme nous obtenons la voix transformée à partir de la production mélodique de M.GREGORIO, nous ne pouvons disposer les enregistrements sur un seul spectrogramme comme précédemment. Nous exposons donc dans un premier temps le spectrogramme obtenu après calcul de la voix mélodique imposée suivi du spectrogramme représentant l'imitation de M.GREGORIO. Dans un second temps, nous faisons de même avec les lignes mélodiques respectives.



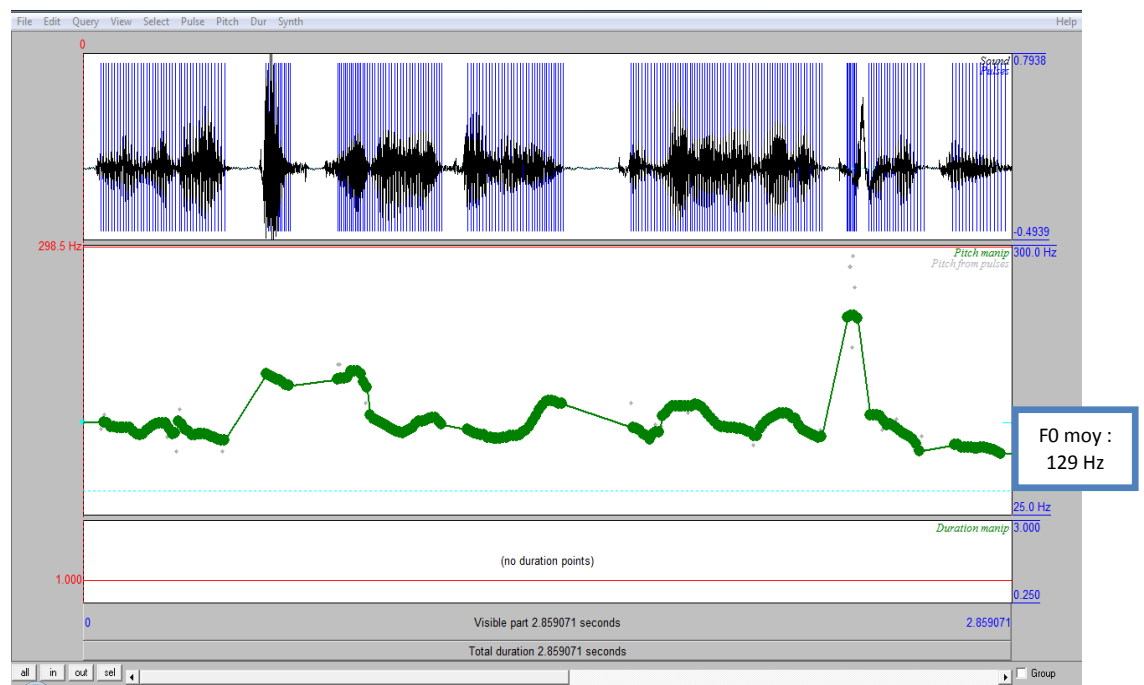
Spectrogramme de la voix transformée mélodique



Spectrogramme de l'imitation de M.GREGORIO



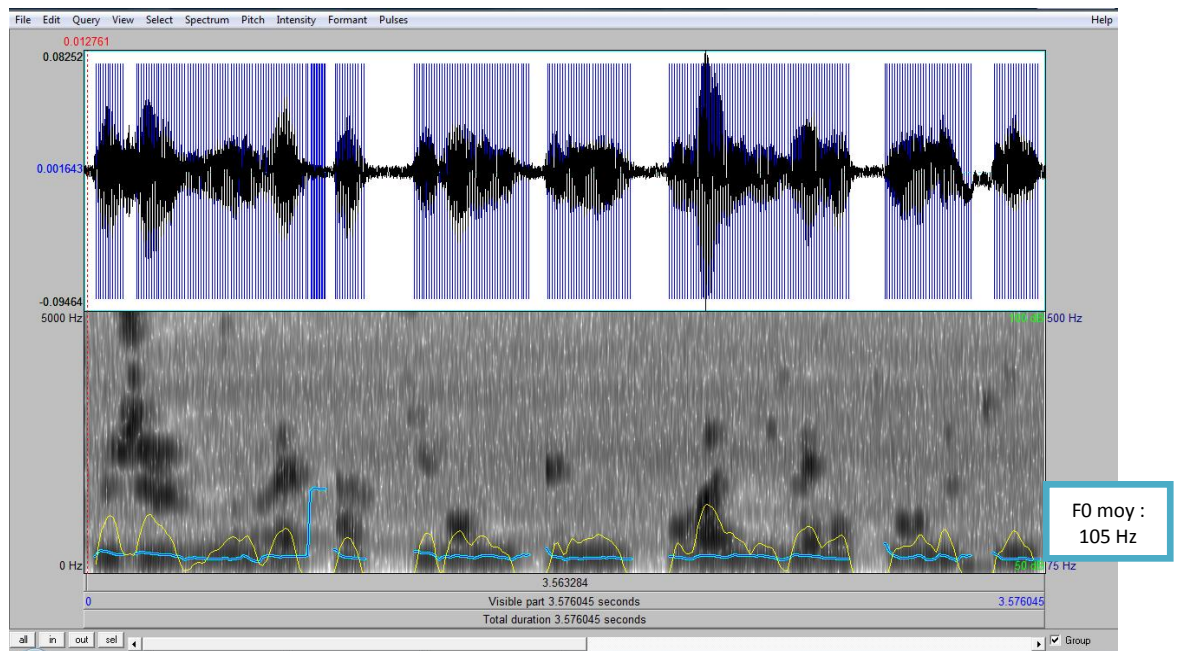
Ligne mélodique de la voix transformée mélodique



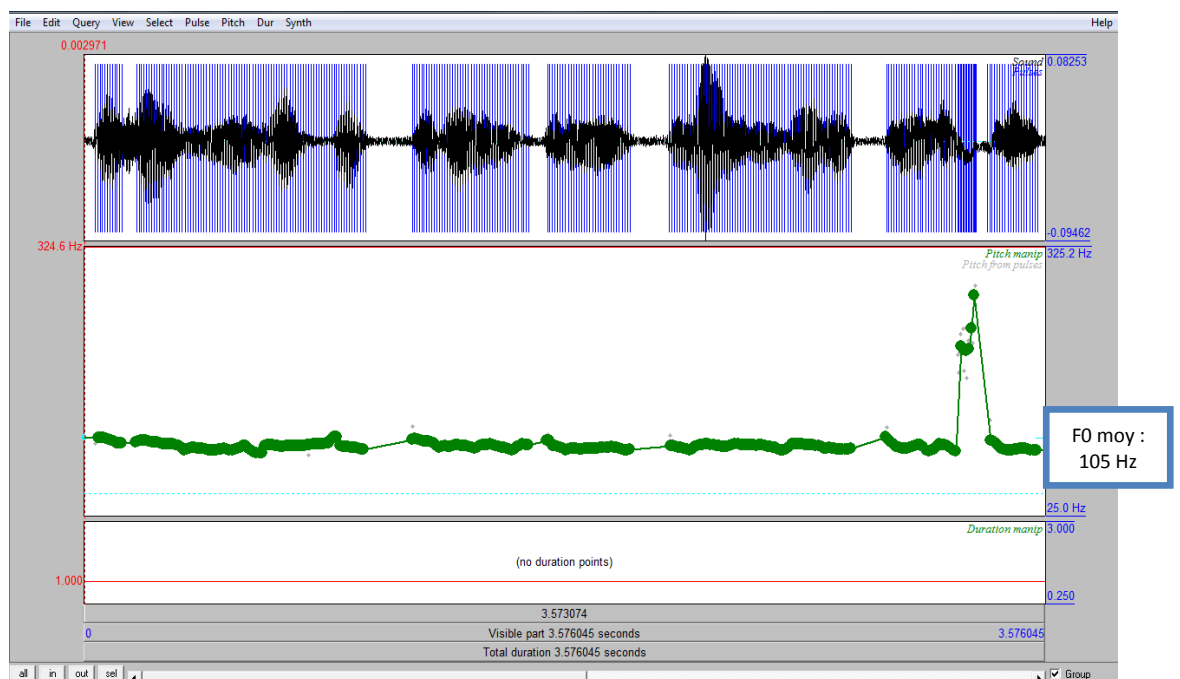
Ligne mélodique de l'imitation de M.GREGORIO

4- Spectrogramme et ligne mélodique de M.GREGORIO pour l'épreuve spontanée monotone :

Nous n'utilisons pas la ligne mélodique de cette épreuve dans le protocole donc nous ne l'exposons pas.



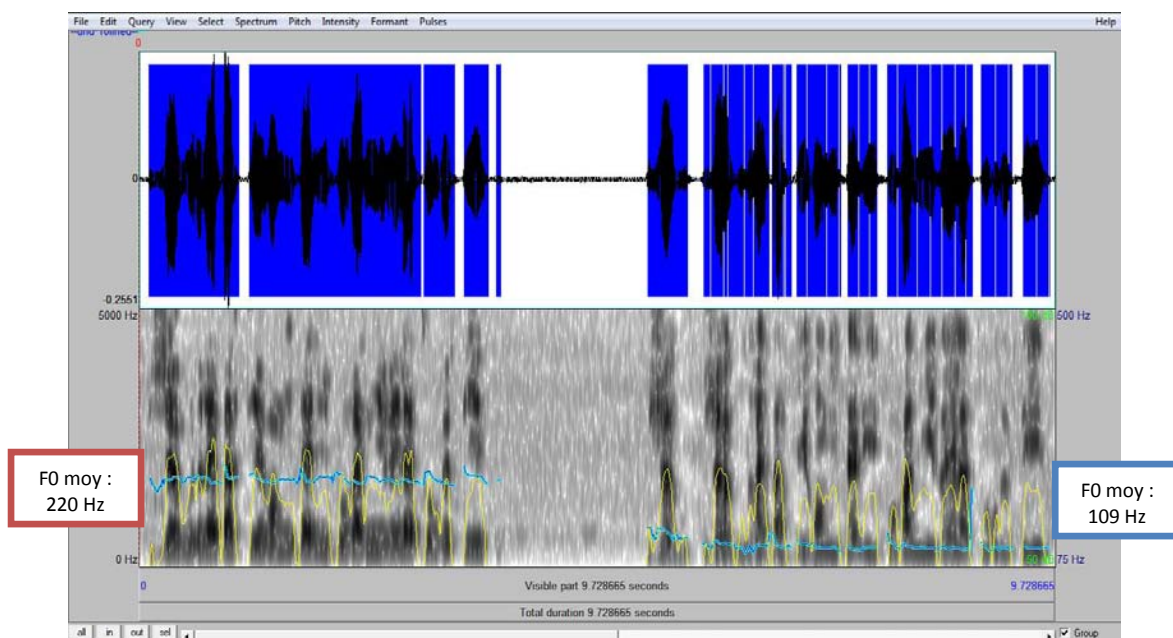
Spectrogramme de la production monotone spontanée de M.GREGORIO



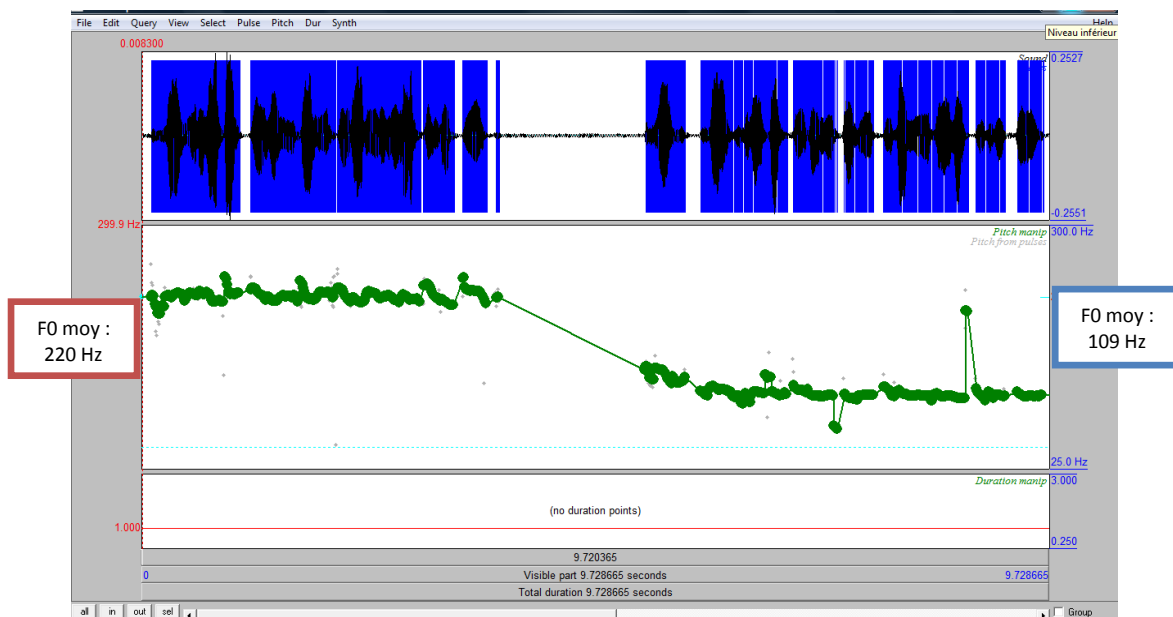
Ligne mélodique de la production spontanée de M.GREGORIO

5- Spectrogramme puis ligne mélodique pour l'épreuve d'imitation du modèle orthophonique monotone :

Nous avons réalisé les enregistrements à la suite, avec en premier le modèle de l'orthophoniste (à gauche) suivi de l'imitation de M.GREGORIO (à droite).



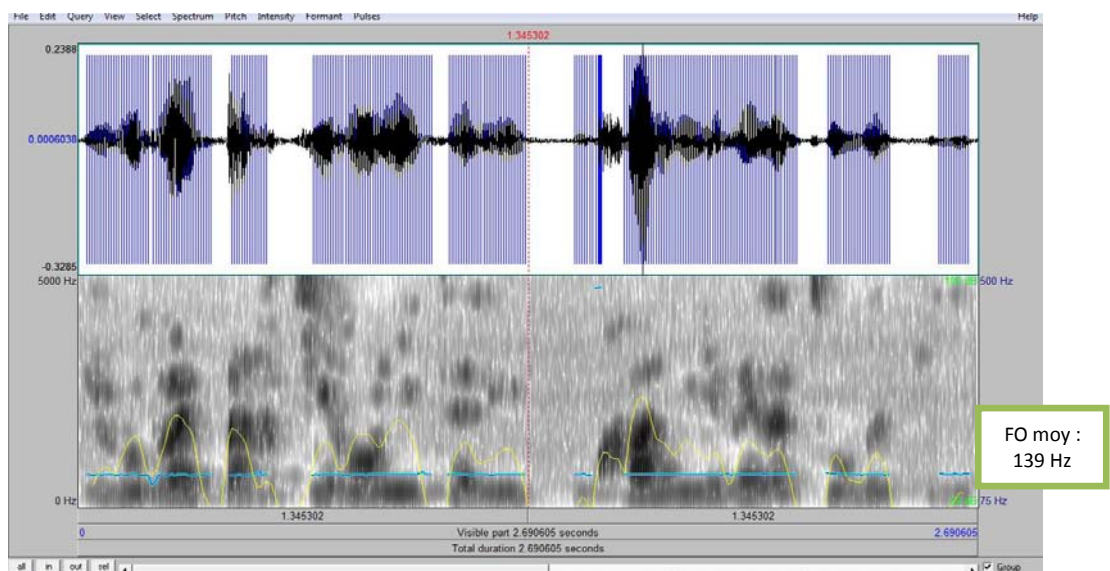
Spectrogramme du modèle orthophonique monotone suivi de l'imitation de M.GREGORIO



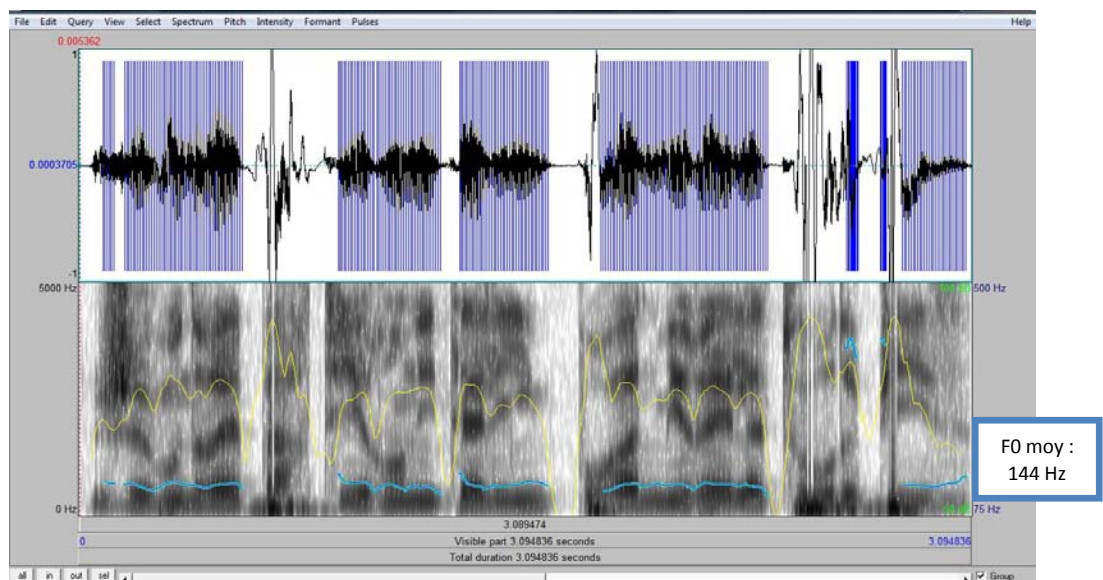
Ligne mélodique du modèle orthophonique monotone suivi de l'imitation de M.GREGORIO

6- Spectrogramme et ligne mélodique pour l'épreuve d'imitation de la voix transformée monotone :

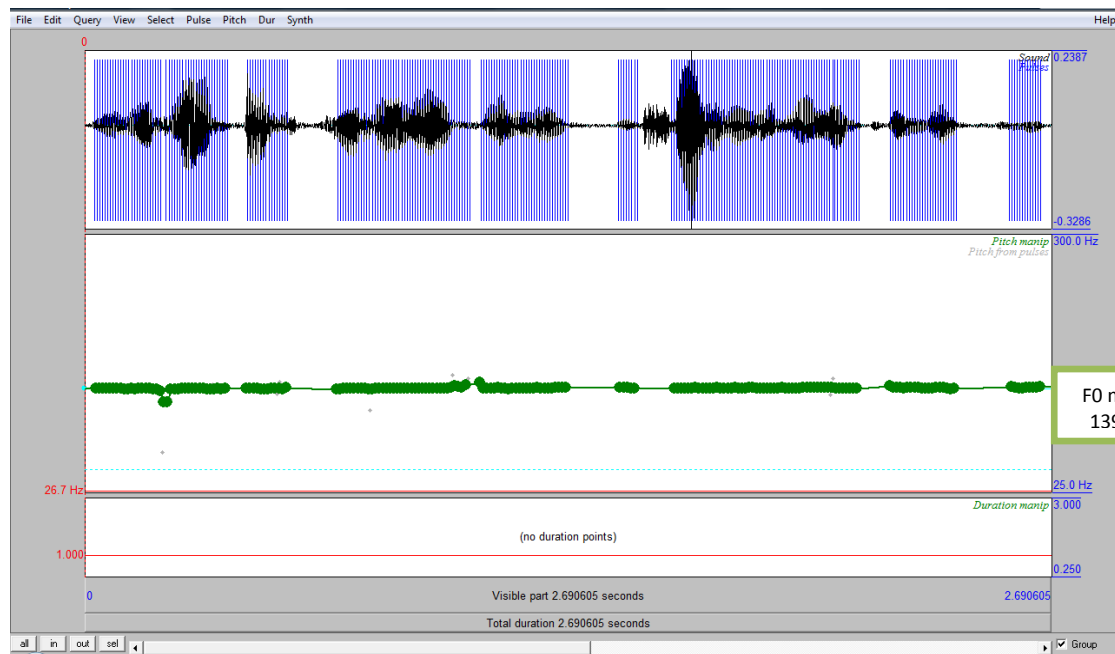
Comme nous obtenons la voix transformée à partir de la production mélodique de M.GREGORIO, nous ne pouvons disposer les enregistrements sur un seul spectrogramme comme précédemment. Nous exposons donc dans un premier temps le spectrogramme obtenu après calcul de la voix monotone imposée suivi du spectrogramme représentant l'imitation de M.GREGORIO. Dans un second temps, nous faisons de même avec les lignes mélodiques respectives.



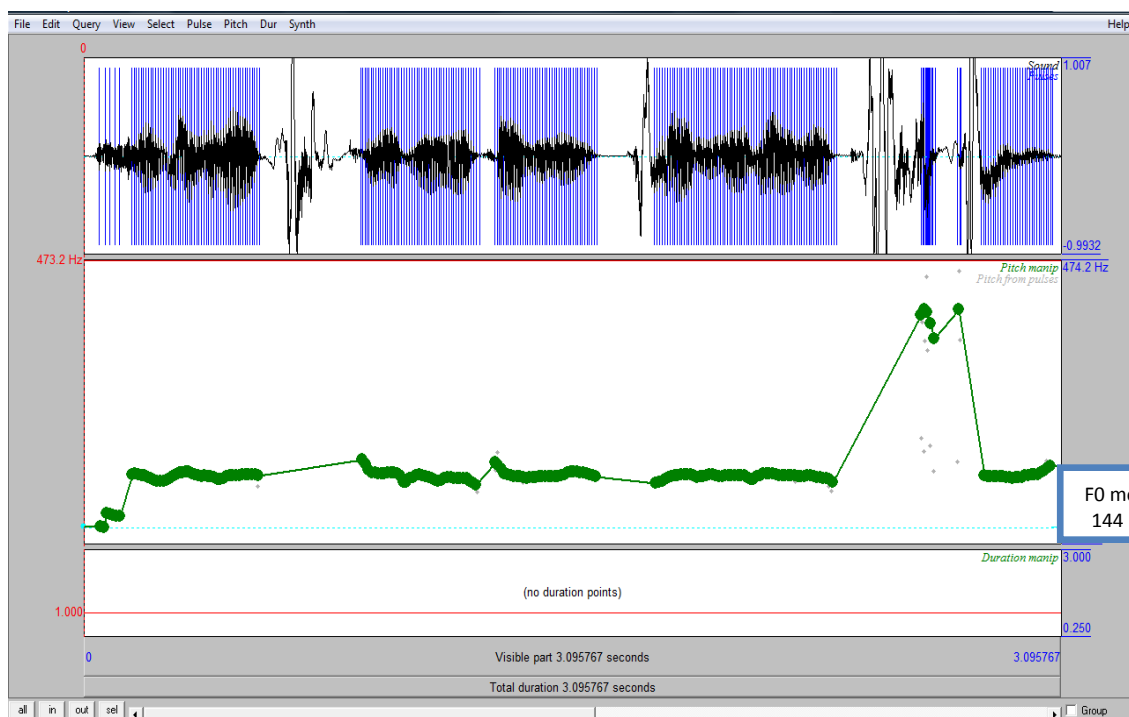
Spectrogramme de la voix transformée monotone



Spectrogramme de l'imitation de M.GREGORIO



Ligne mélodique de la voix transformée monotone



Ligne mélodique de l'imitation de M.GREGORIO

RESULTATS ET LIGNES MELODIQUES TEMOIN N°1

[ma] =} min : 102 Hz max : 205 Hz
 [mi] =} min : 91 Hz max : 304 Hz
 [mu] =} min : 95 Hz max : 235 Hz

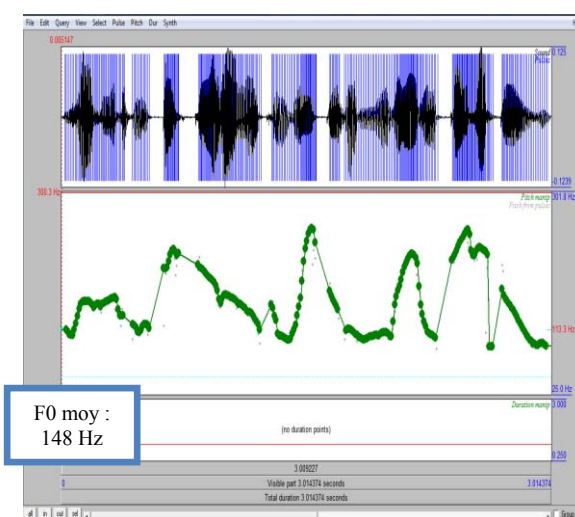
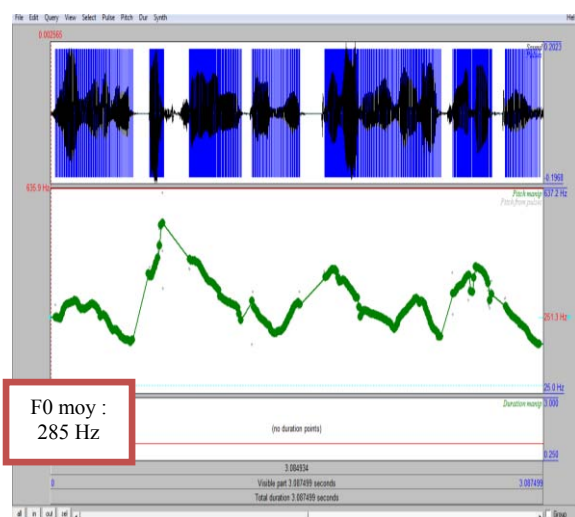
Ecart voulu : [102 Hz / 205 Hz]

Fmax voulue = 205 Hz
 Fmin voulue = 102 Hz
 δ voulu = 205 – 102 = 103 Hz

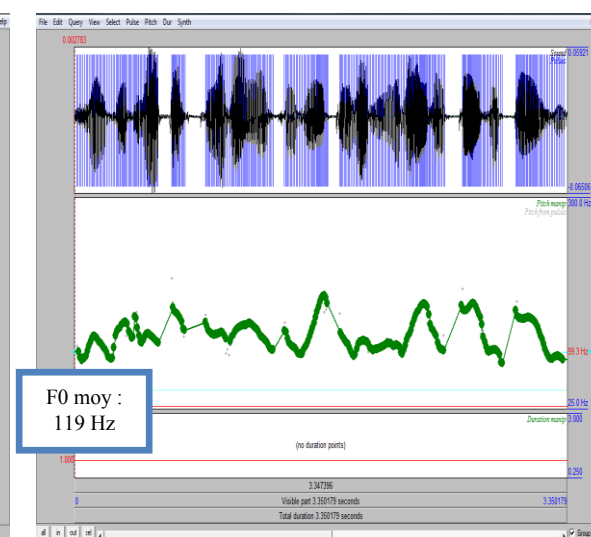
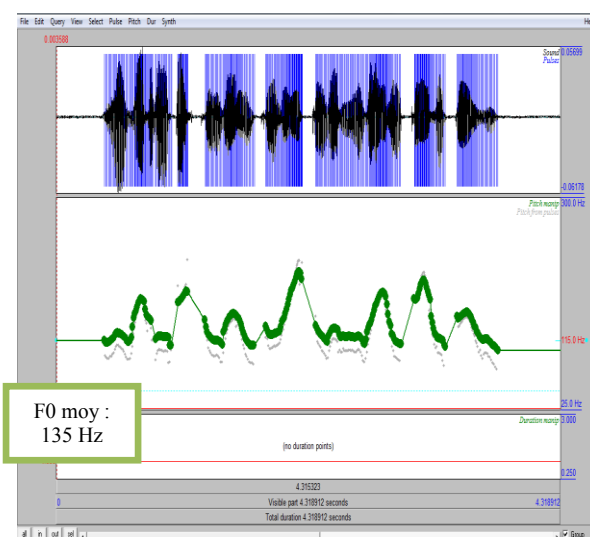
Fmax observée = 218 Hz
 Fmin observée = 77 Hz
 δ observé = 218 – 77 = 141 Hz

- α = 103 / 141 = 0,73
 - β = 205 – [218 x 0,73] = 45,76

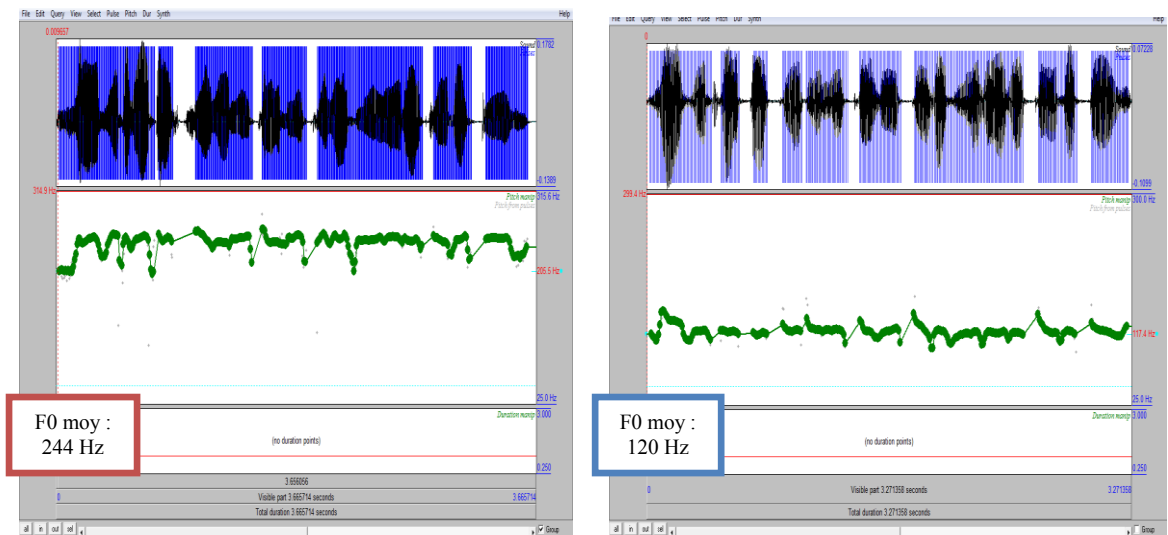
Ligne monotone = 102 + 103 / 2 = 153,5 Hz (que l'on arrondit à 154 Hz).



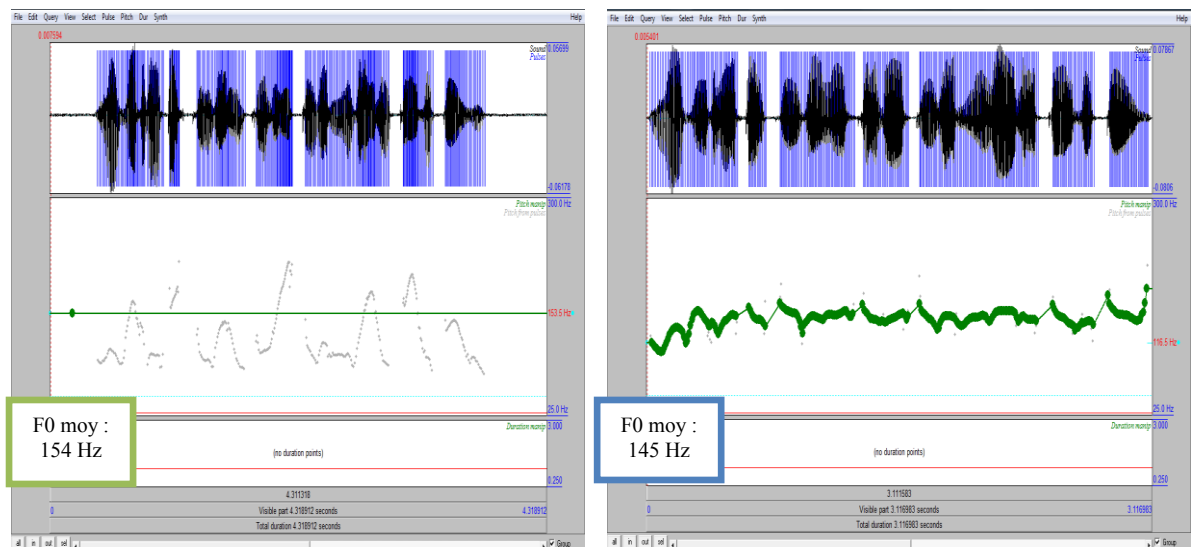
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du témoin n°1



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du témoin n°1



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du témoin n°1



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du témoin n°1

RESULTATS ET LIGNES MELODIQUES TEMOIN N°2

[ma] =} min : 88 Hz max : 208 Hz
 [mi] =} min : **95** Hz max : **200** Hz
 [mu] =} min : 88 Hz max : 223 Hz

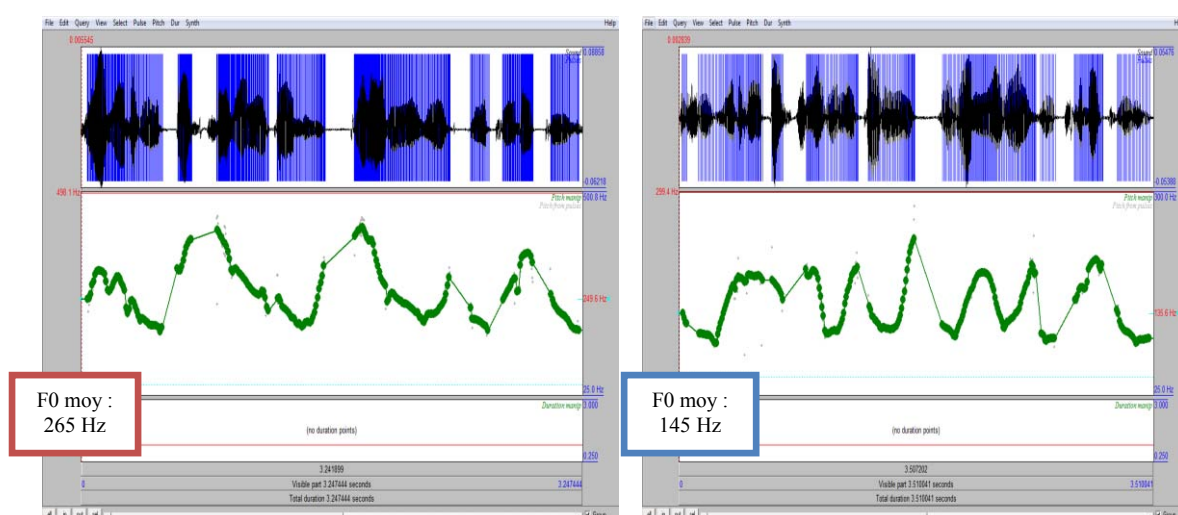
Ecart voulu : [200 Hz / 95 Hz]

Fmax voulue = 200 Hz
 Fmin voulue = 95 Hz
 δ voulu = 200 – 95 = **105** Hz

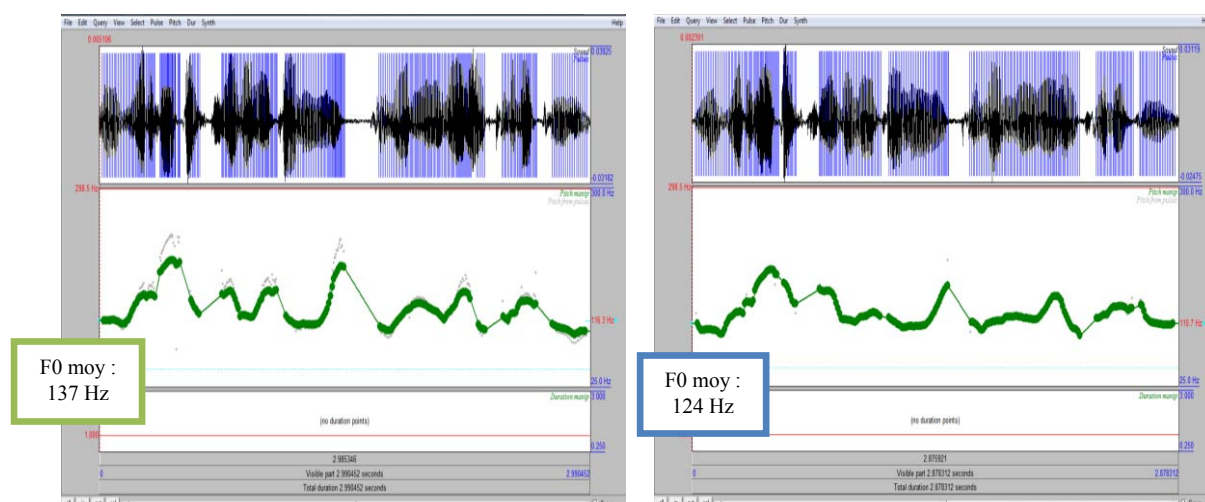
Fmax observée = 233 Hz
 Fmin observée = 84 Hz
 δ observé = 233 – 84 = **149** Hz

- $\alpha = 105 / 149 = 0,7$
 - $\beta = 200 - [233 \times 0.7] = 36,9$

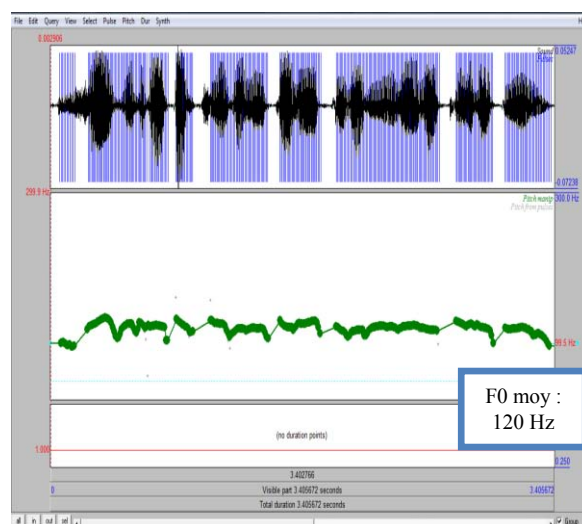
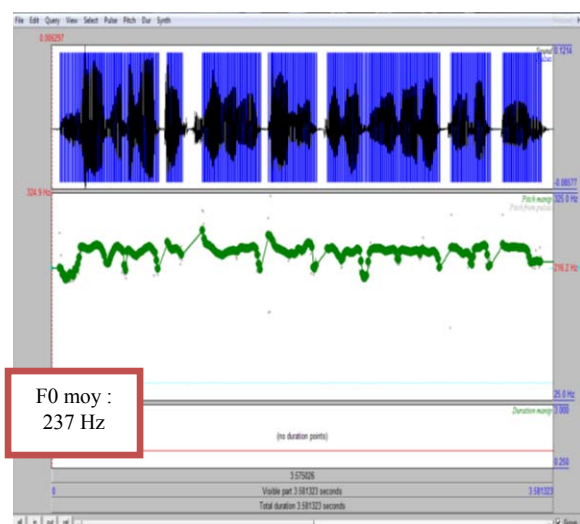
Ligne monotone = $95 + 105 / 2 = 147,5$ Hz (que l'on arrondit à 148 Hz).



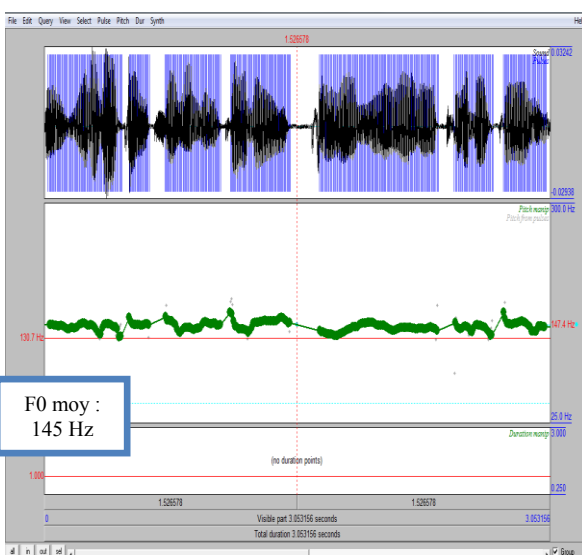
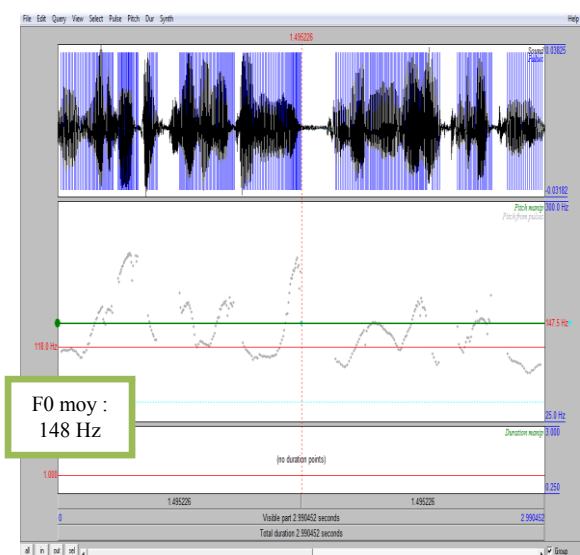
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du témoin n°2



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du témoin n°2



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du témoin n°2



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du témoin n°2

RESULTATS ET LIGNES MELODIQUES TEMOIN N°3

[ma] =} min : 158 Hz max : 266 Hz
 [mi] =} min : 155 Hz max : 322 Hz
 [mu] =} min : 180 Hz max : 291 Hz

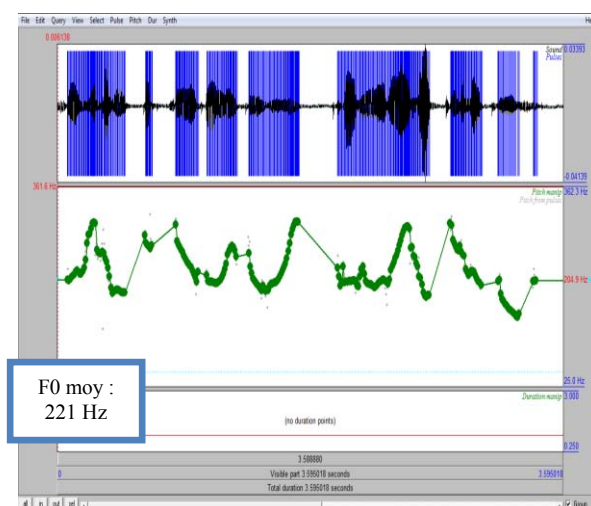
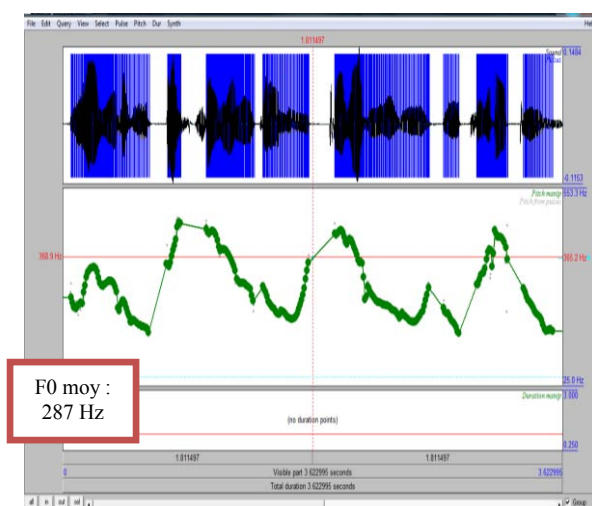
Ecart voulu : [180 Hz / 266 Hz]

Fmax voulue = 266 Hz
 Fmin voulue = 180 Hz
 δ voulu = 266 – 180 = 86 Hz

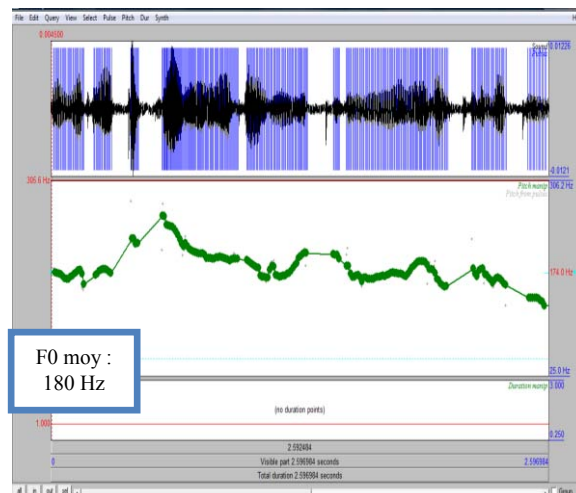
Fmax observée = 313 Hz
 Fmin observée = 146 Hz
 δ observé = 313 – 146 = 167 Hz

- $\alpha = 86 / 167 = 0,51$
 - $\beta = 266 - [313 \times 0,51] = 106$

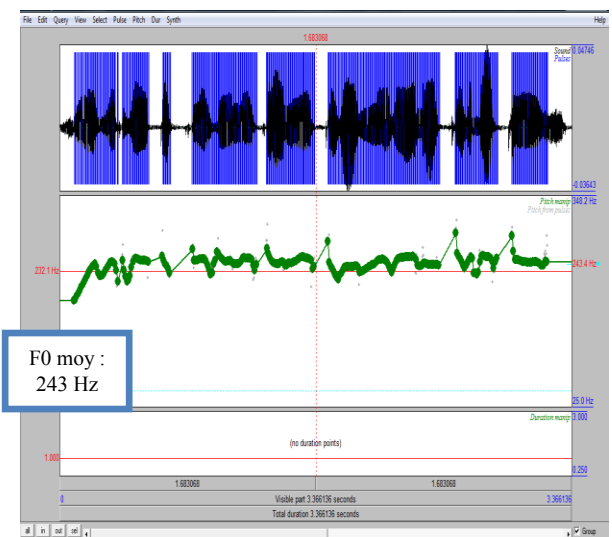
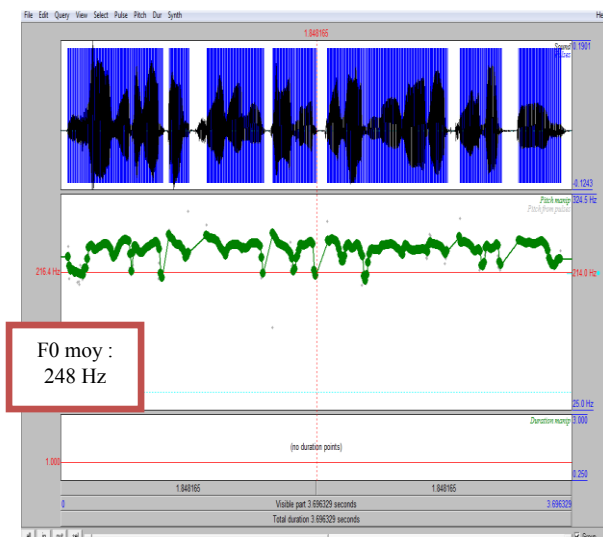
Ligne monotone = 180 + 86/2 = 223 Hz.



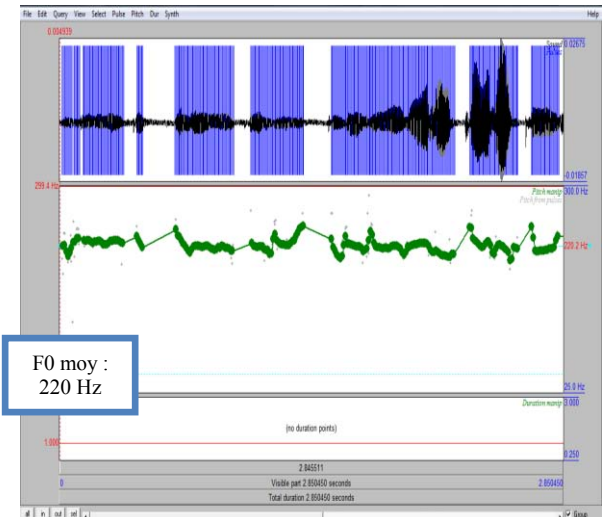
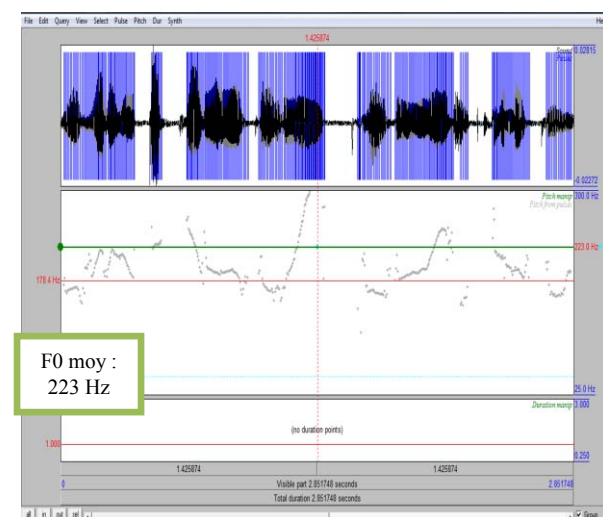
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du témoin n°3



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du témoin n°3



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du témoin n°3



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du témoin n°3

RESULTATS ET LIGNES MELODIQUES TEMOIN N°4

[ma] =} min : 207 Hz max : 338 Hz
 [mi] =} min : 189 Hz max : 317 Hz
 [mu] =} min : 200 Hz max : 370 Hz

Ecart voulu : [207 Hz / 317 Hz]

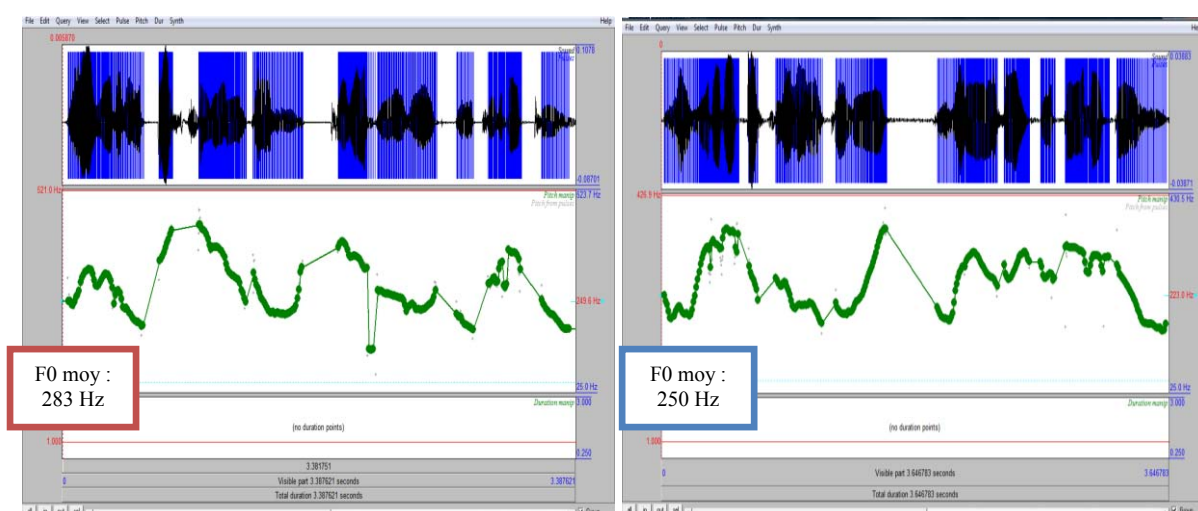
Fmax voulue = 317 Hz
 Fmin voulue = 207 Hz
 δ voulu = 317 – 207 = 110 Hz

Fmax observée = 413 Hz
 Fmin observée = 148 Hz
 δ observé = 413 – 148 = 265 Hz

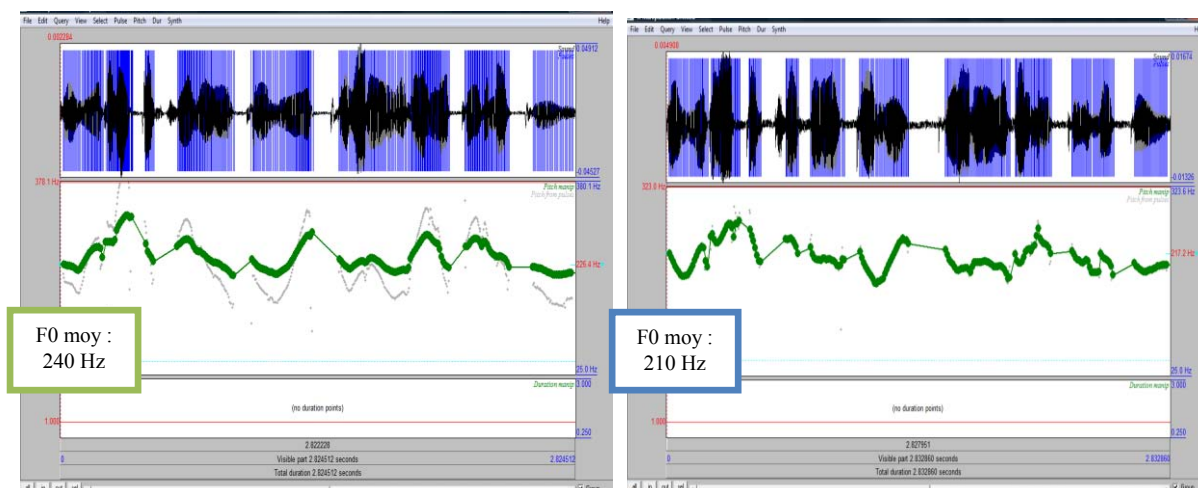
- $\alpha = 110 / 265 = 0,42$

- $\beta = 317 - [413 \times 0.42] = 143,54$

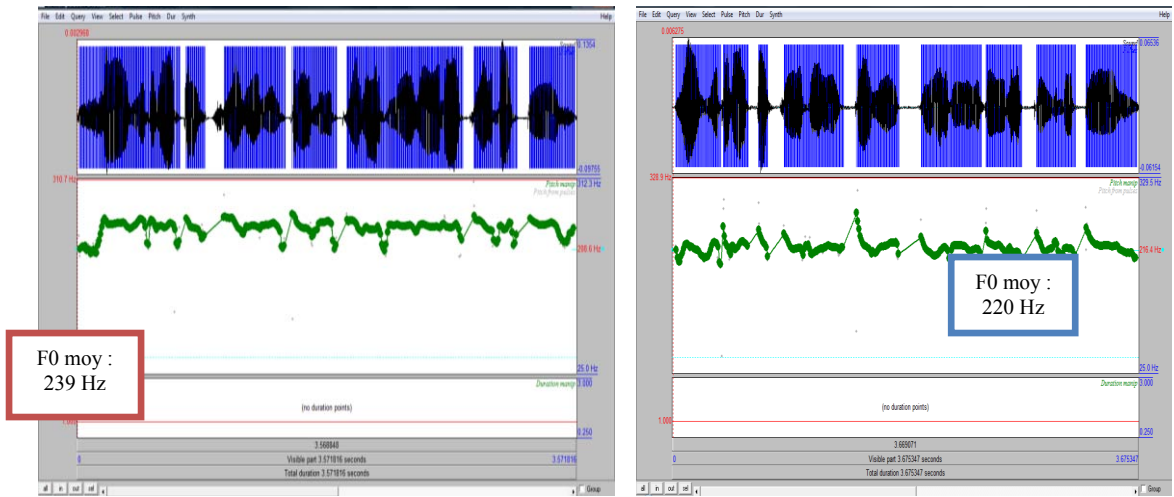
Ligne monotone = 207 + 110/2 = 262 Hz.



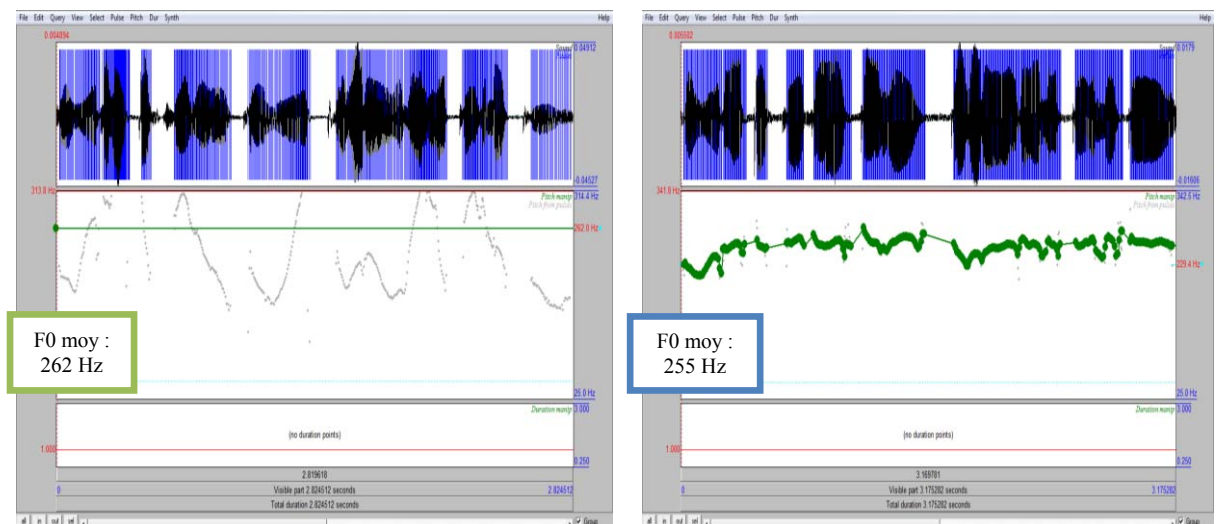
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du témoin n°4



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du témoin n°4



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du témoin n°4



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du témoin n°4

RESULTATS ET LIGNES MELODIQUES TEMOIN N°5

[ma] =} min : 160 Hz max : 288 Hz
 [mi] =} min : 169 Hz max : 270 Hz
 [mu] =} min : 152 Hz max : 258 Hz

Ecart voulu : [169 Hz / 258 Hz]

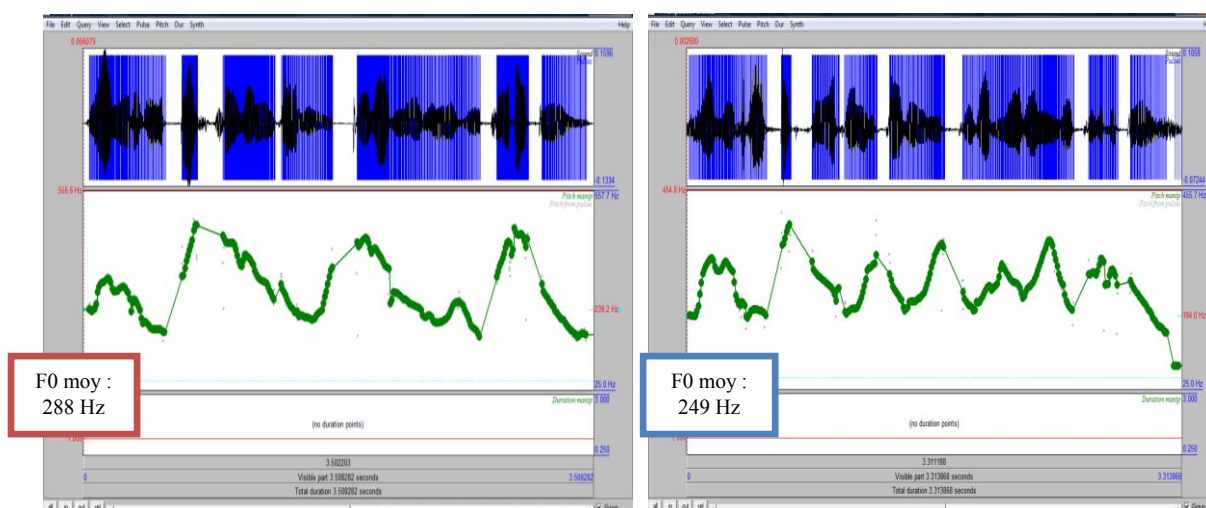
Fmax voulue = 258 Hz
 Fmin voulue = 169 Hz
 δ voulu = 258 – 169 = 89 Hz

Fmax observée = 328 Hz
 Fmin observée = 122 Hz
 δ observé = 328 – 122 = 206 Hz

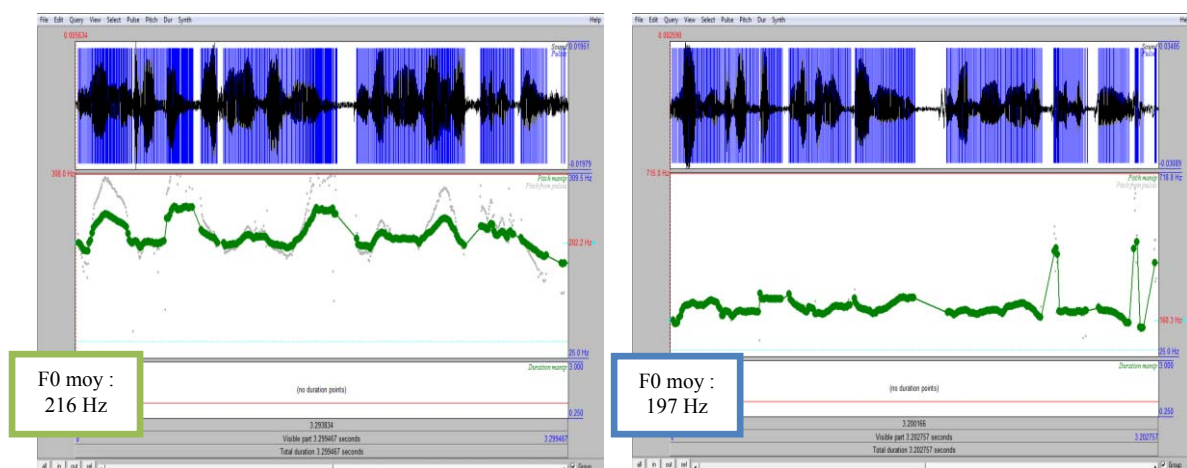
- $\alpha = 89 / 206 = 0,43$

- $\beta = 258 - [328 \times 0.43] = 116,96$

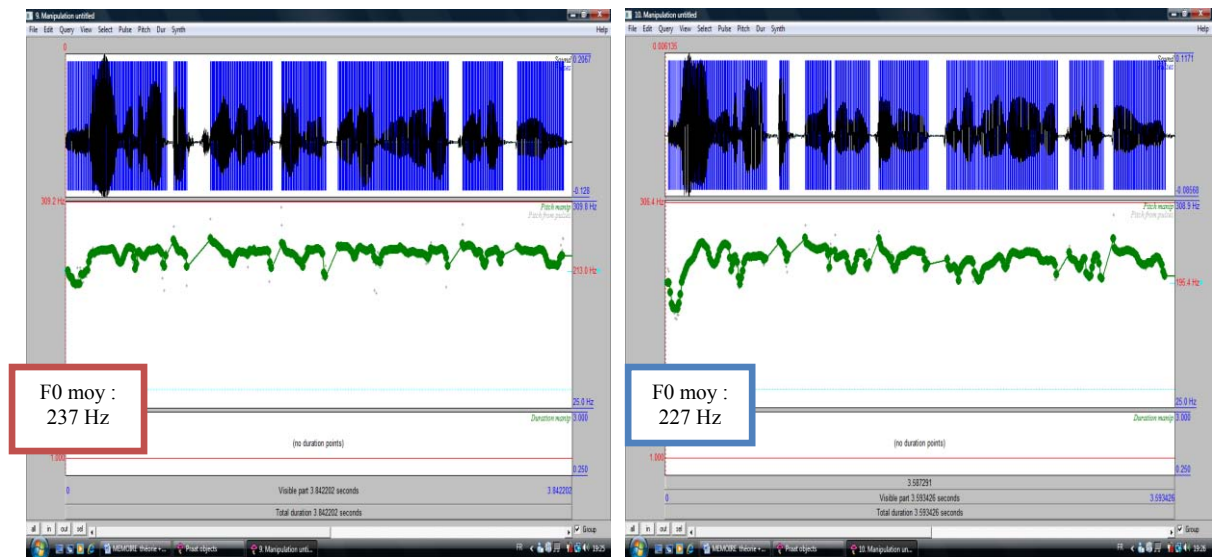
Ligne monotone = $169 + 89/2 = 213,5$ Hz (que l'on arrondit à 214 Hz).



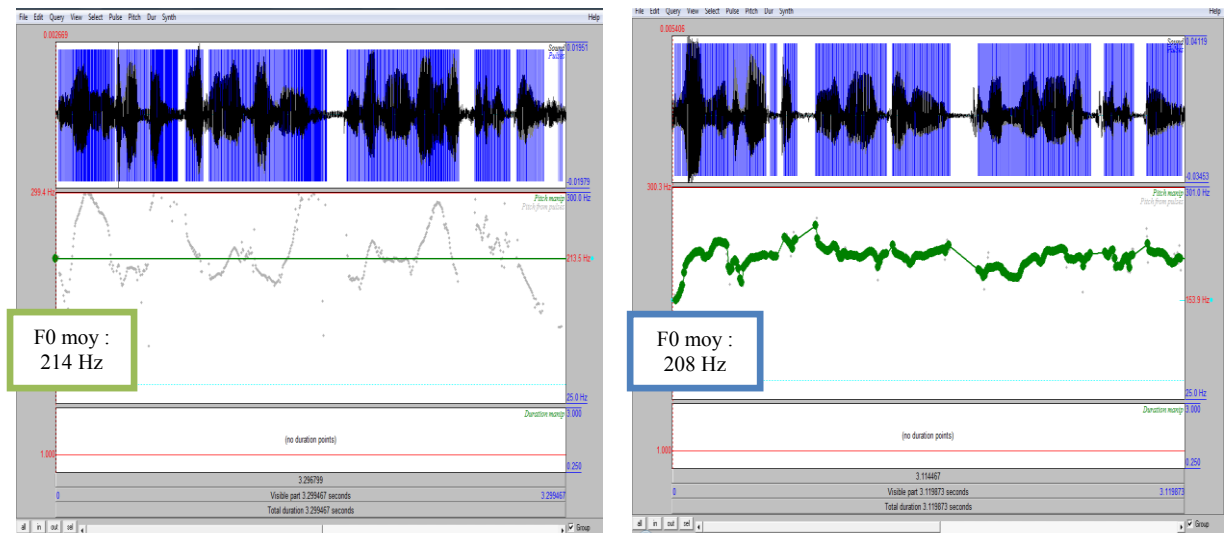
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du témoin n°5



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du témoin n°5



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du témoin n°5



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du témoin n°5

RESULTATS ET LIGNES MELODIQUES TEMOIN N°6

[ma] =} min : 100 Hz max : 205 Hz
 [mi] =} min : 110 Hz max : 186 Hz
 [mu] =} min : 116 Hz max : 213 Hz

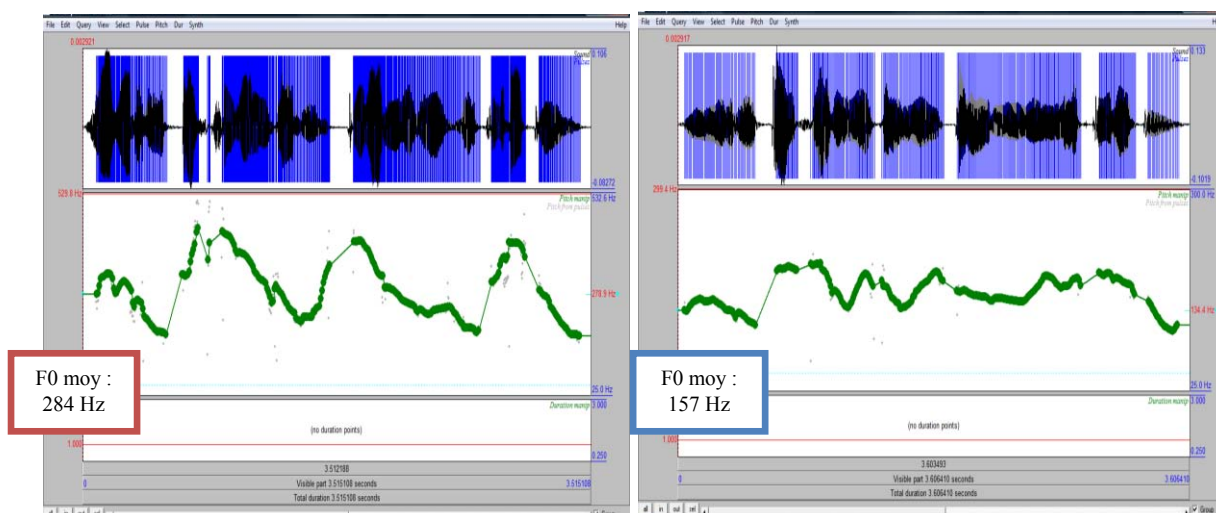
Ecart voulu : [116 Hz / 186 Hz]

Fmax voulue = 186 Hz
 Fmin voulue = 116 Hz
 δ voulu = 186 – 116 = 70 Hz

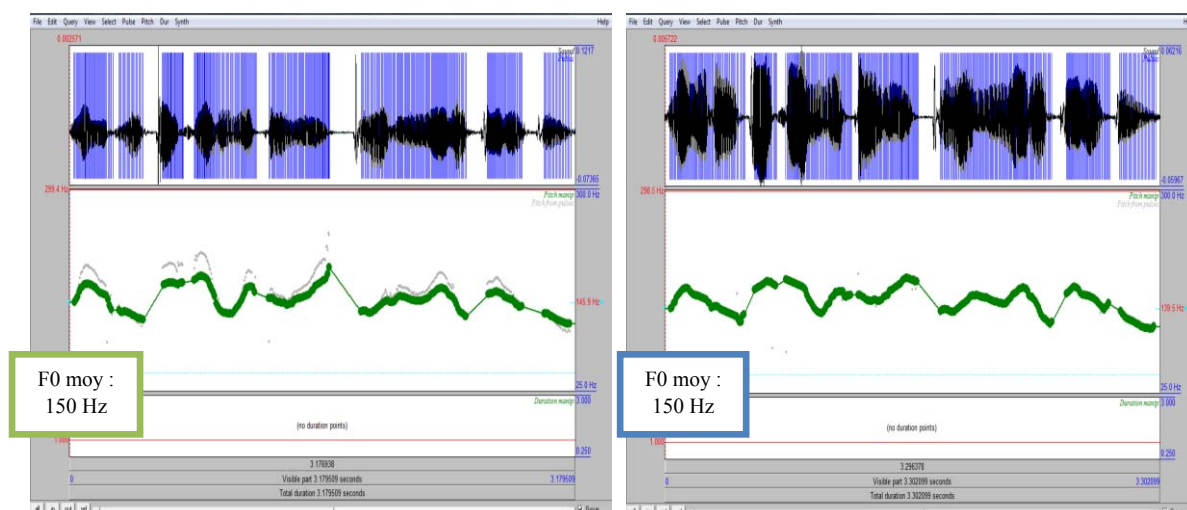
Fmax observée = 220 Hz
 Fmin observée = 106 Hz
 δ observé = 220 – 106 = 114 Hz

- $\alpha = 70 / 114 = 0,61$
 - $\beta = 186 - [220 \times 0,61] = 51,8$

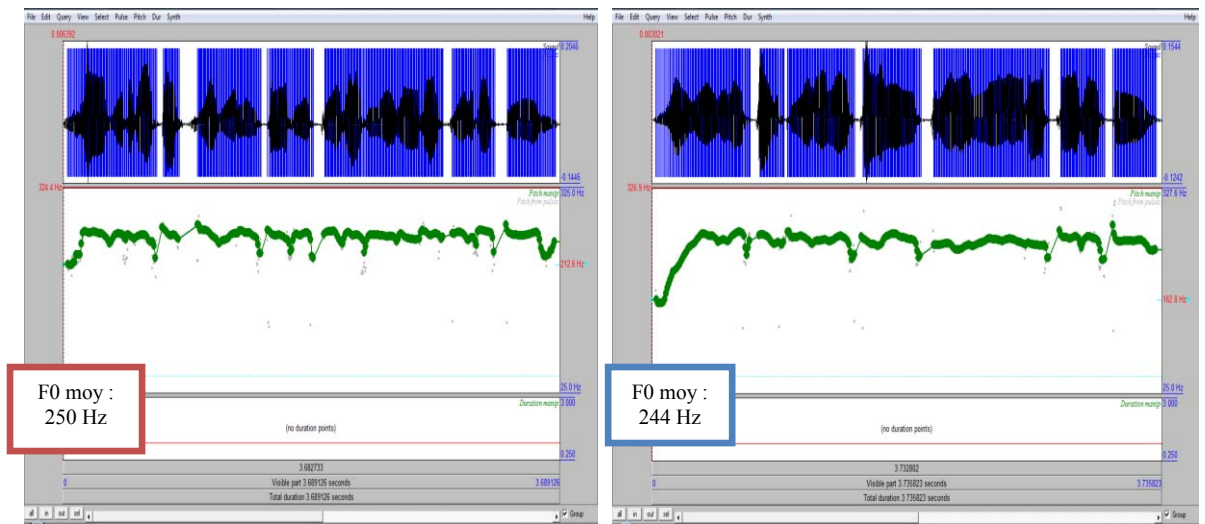
Ligne monotone = 116 + 70/2 = 151 Hz.



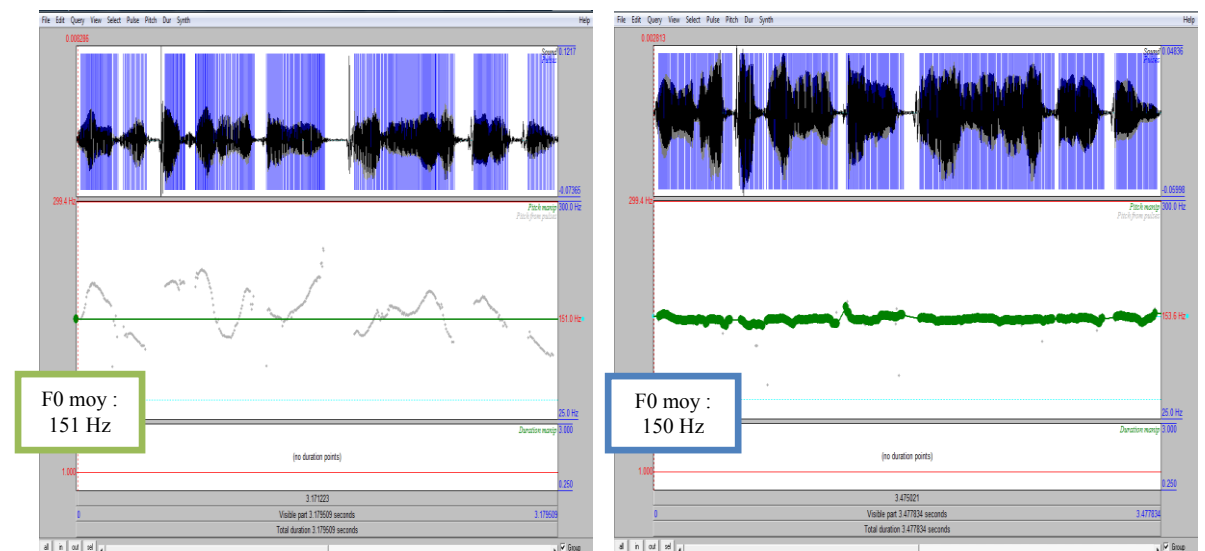
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du témoin n°6



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du témoin n°6



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du témoin n°6



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du témoin n°6

RESULTATS ET LIGNES MELODIQUES TEMOIN N°7

[ma] =} min : 92 Hz max : **223** Hz
 [mi] =} min : **101** Hz max : 237 Hz
 [mu] =} min : 101 Hz max : 306 Hz

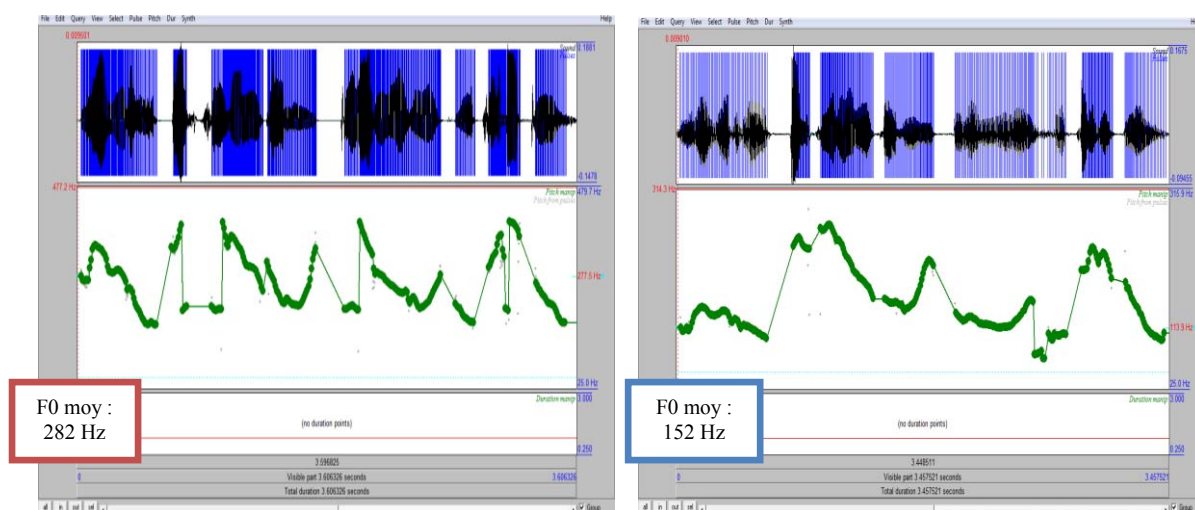
Ecart voulu : [101 Hz / 223 Hz]

Fmax voulue = 223 Hz
 Fmin voulue = 101 Hz
 δ voulu = 223 – 101 = **122** Hz

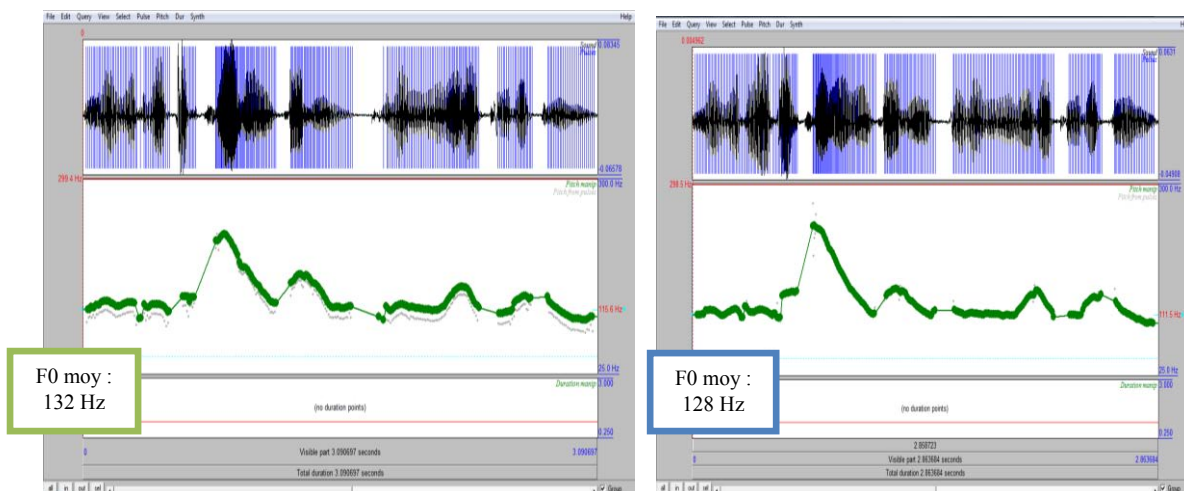
Fmax observée = 219 Hz
 Fmin observée = 84 Hz
 δ observé = 219 – 84 = **135** Hz

- $\alpha = 122 / 135 = 0,9$
 - $\beta = 223 - [219 \times 0.9] = 25$

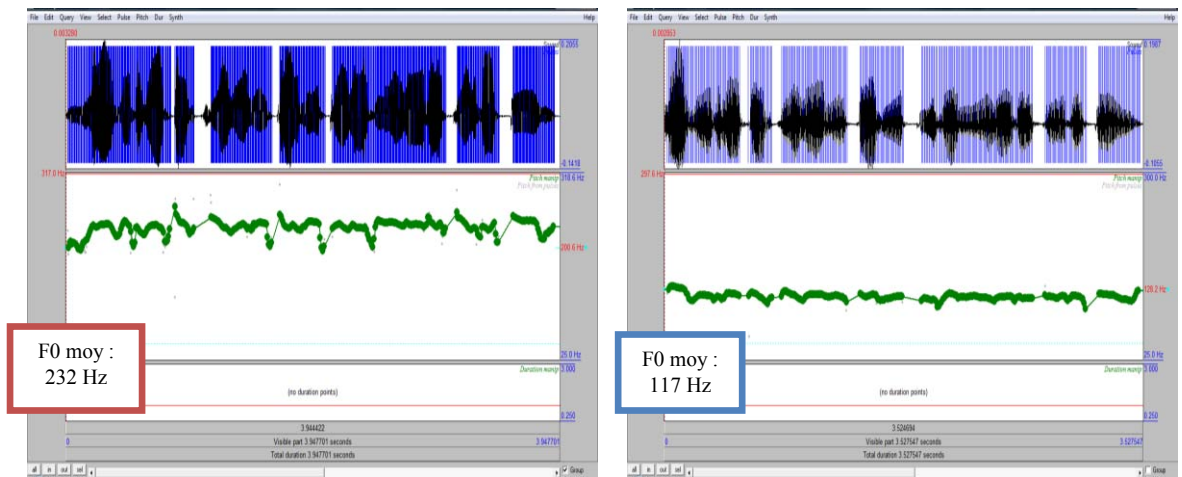
Ligne monotone = 101 + 122/2 = **162** Hz.



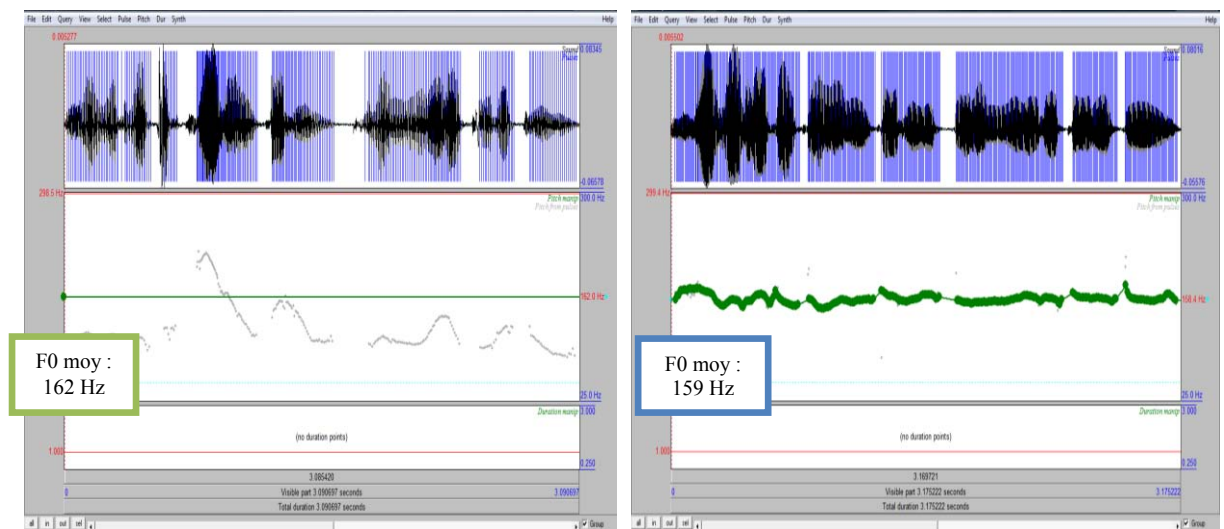
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du témoin n°7



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du témoin n°7



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du témoin n°7



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du témoin n°7

RESULTATS ET LIGNES MELODIQUES TEMOIN N°8

[ma] =} min : 153 Hz max : 199 Hz
 [mi] =} min : 127 Hz max : 208 Hz
 [mu] =} min : 121 Hz max : 210 Hz

Ecart voulu : [153 Hz / 199 Hz]

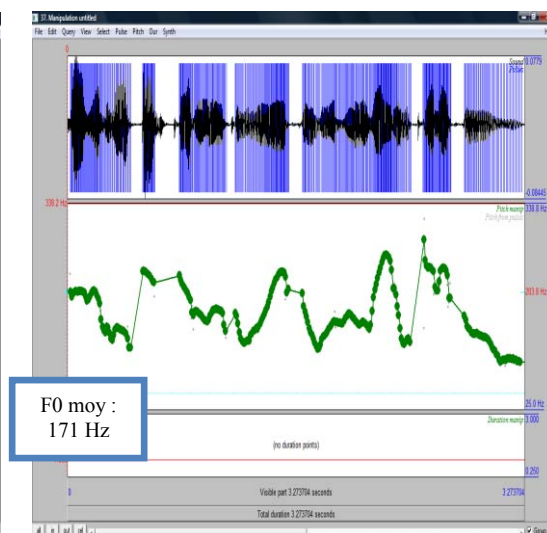
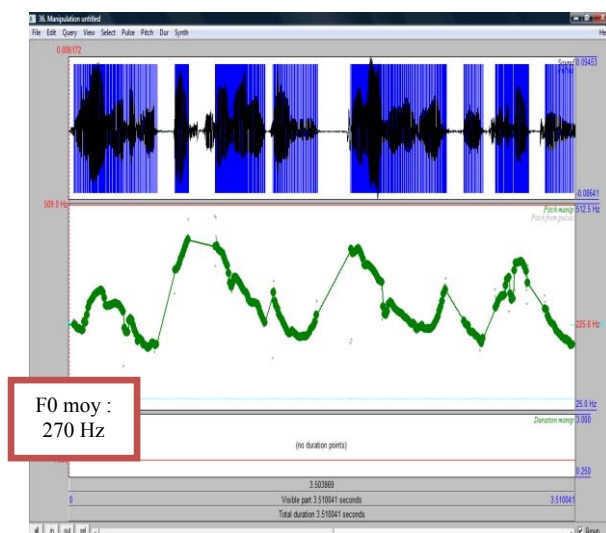
Fmax voulue = 199 Hz
 Fmin voulue = 153 Hz
 δ voulu = 199 – 153 = 46 Hz

Fmax observée = 248 Hz
 Fmin observée = 91 Hz
 δ observé = 248 – 91 = 157 Hz

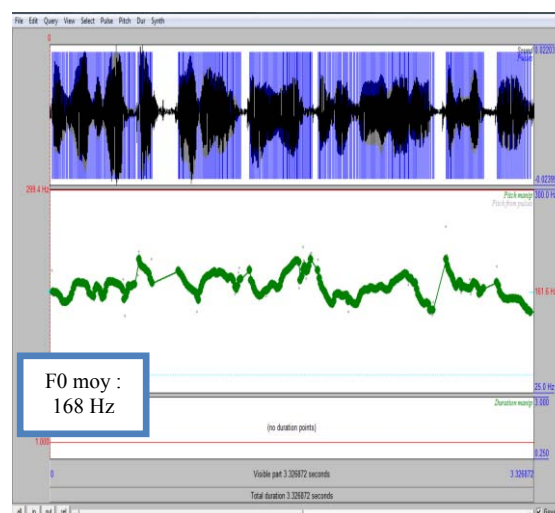
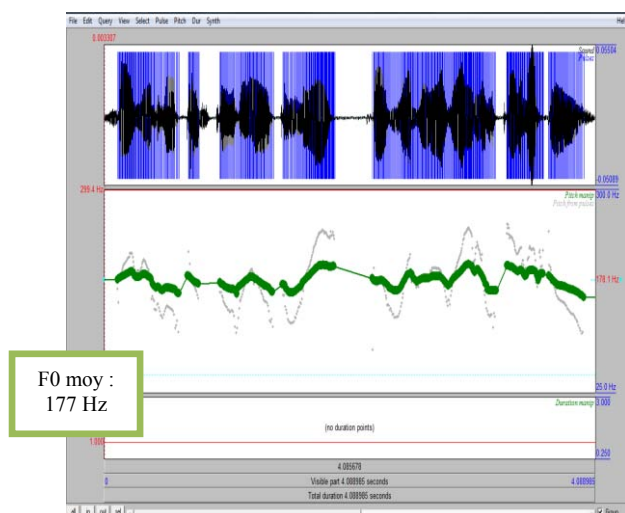
- $\alpha = 46 / 157 = 0,3$

- $\beta = 199 - [248 \times 0.3] = 124,6$

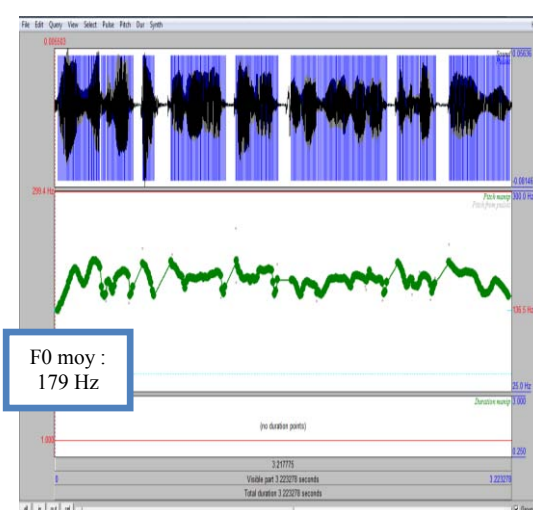
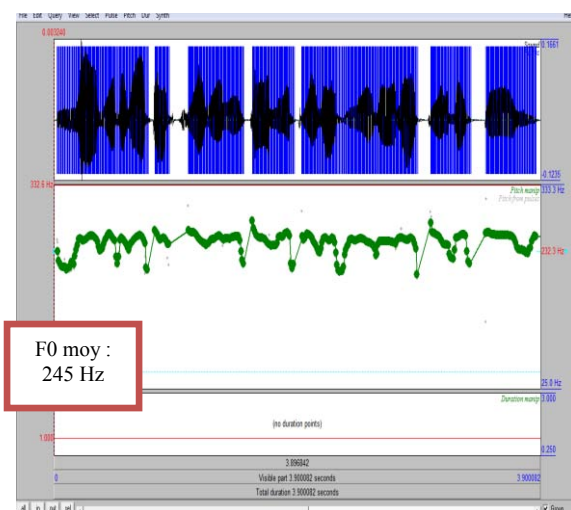
Ligne monotone = 153 + 46/2 = 176 Hz.



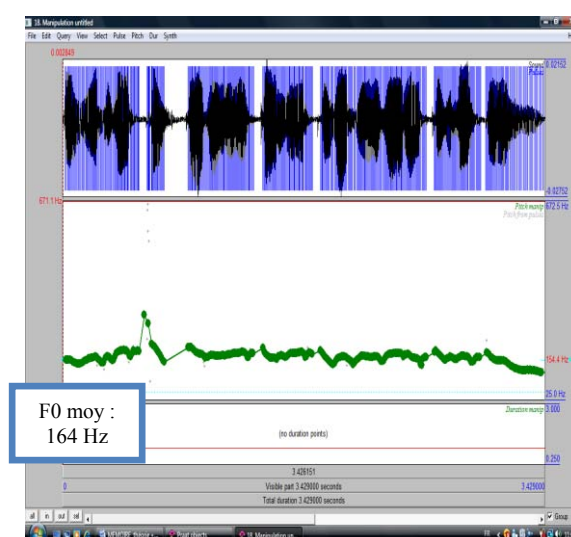
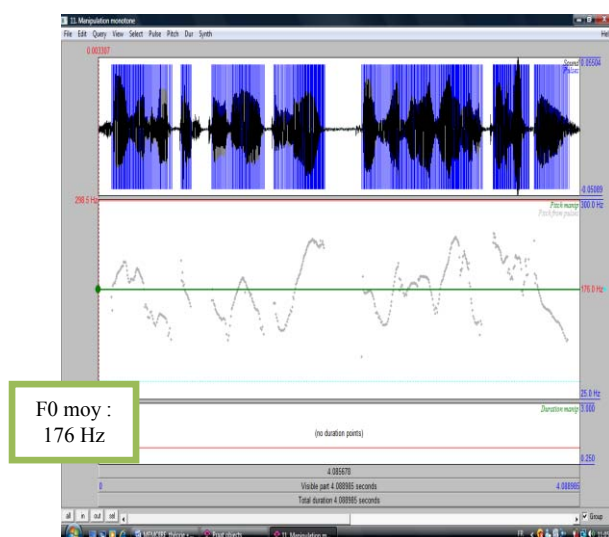
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du témoin n°8



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du témoin n°8



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du témoin n°8



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du témoin n°8

RESULTATS ET LIGNES MELODIQUES TEMOIN N°9

[ma] =} min : 85 Hz max : **153** Hz
 [mi] =} min : 89 Hz max : 160 Hz
 [mu] =} min : **95** Hz max : 158 Hz

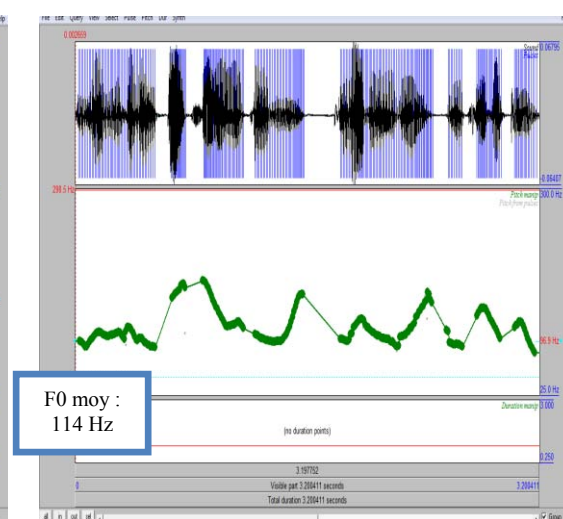
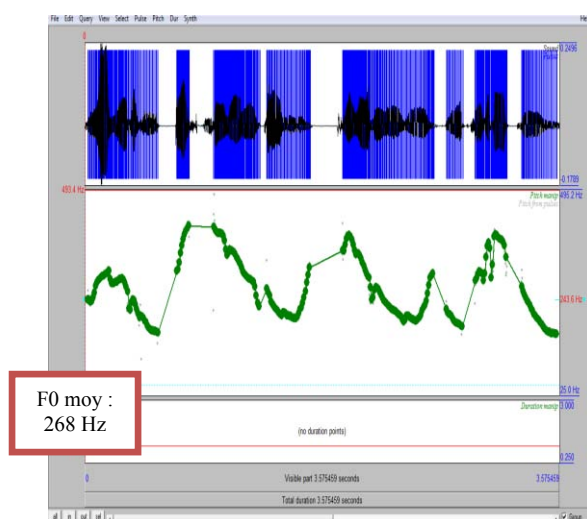
Ecart voulu : [95 Hz / 153 Hz]

Fmax voulue = 153 Hz
 Fmin voulue = 95 Hz
 δ voulu = $153 - 95 = 58$ Hz

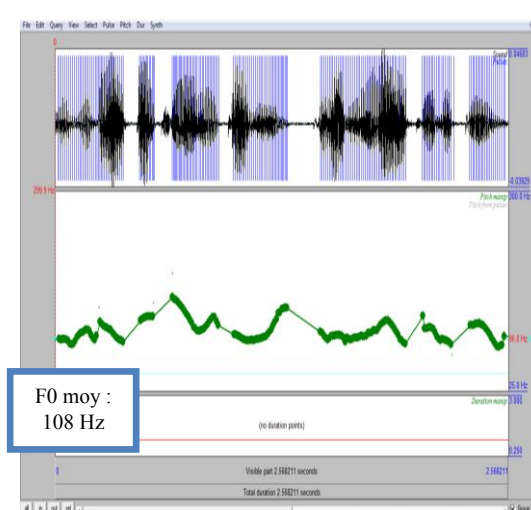
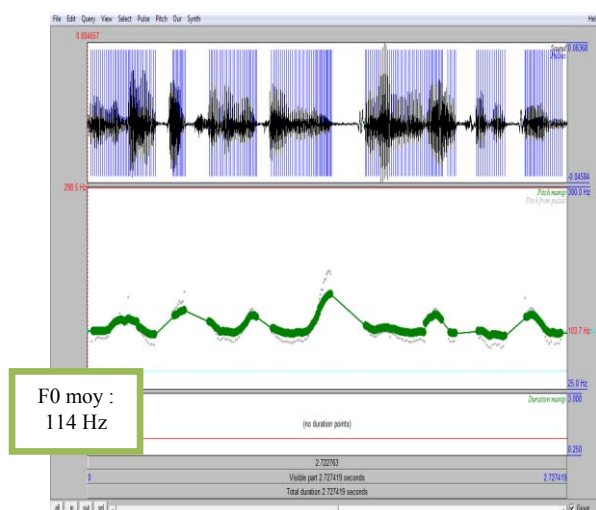
Fmax observée = 183 Hz
 Fmin observée = 80 Hz
 δ observé = $183 - 80 = 103$ Hz

- $\alpha = 58 / 103 = 0,56$
 - $\beta = 153 - [183 \times 0.56] = 51$

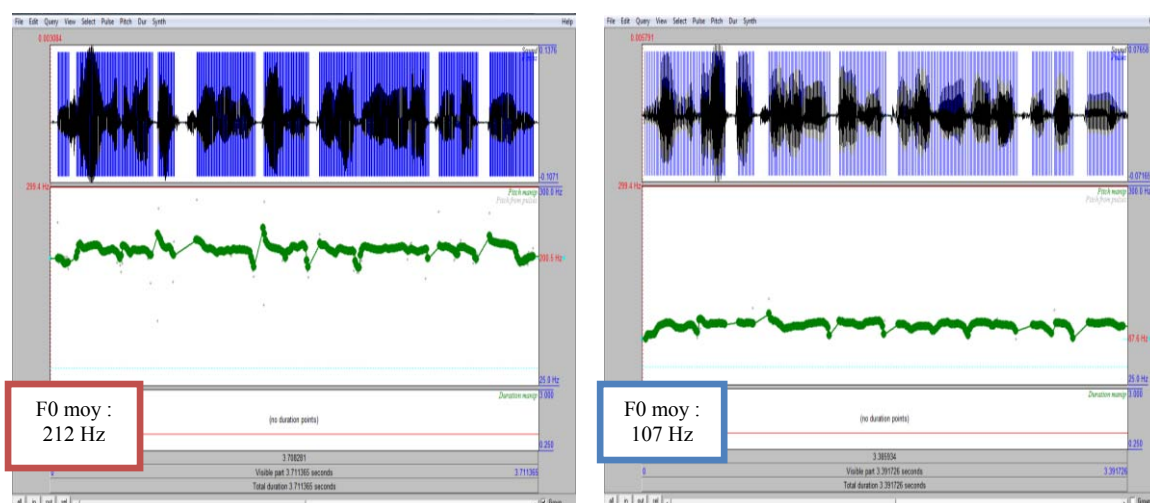
Ligne monotone = $95 + 58/2 = 124$ Hz.



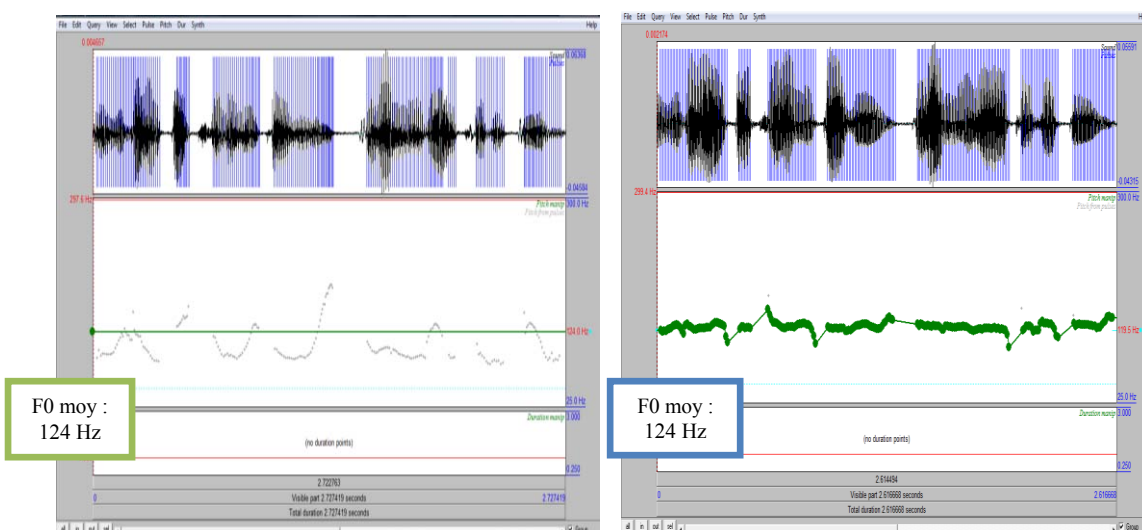
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du témoin n°9



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du témoin n°9



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du témoin n°9



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du témoin n°9

RESULTATS ET LIGNES MELODIQUES TEMOIN N°10

[ma] =} min : 166 Hz max : 383 Hz
 [mi] =} min : 174 Hz max : 373 Hz
 [mu] =} min : 174 Hz max : 360 Hz

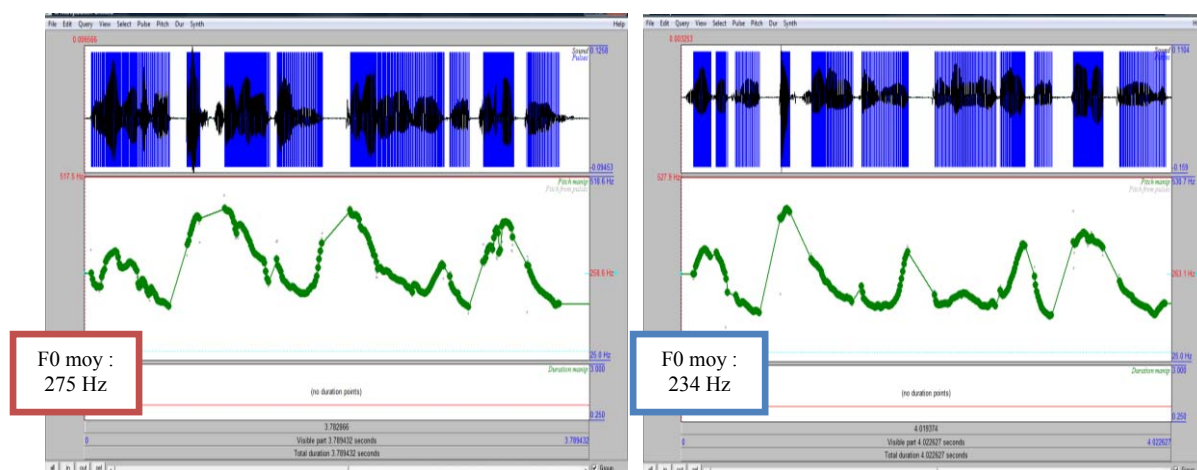
Ecart voulu : [174 Hz / 360 Hz]

Fmax voulue = 360 Hz
 Fmin voulue = 174 Hz
 δ voulu = 360 – 174 = 186 Hz

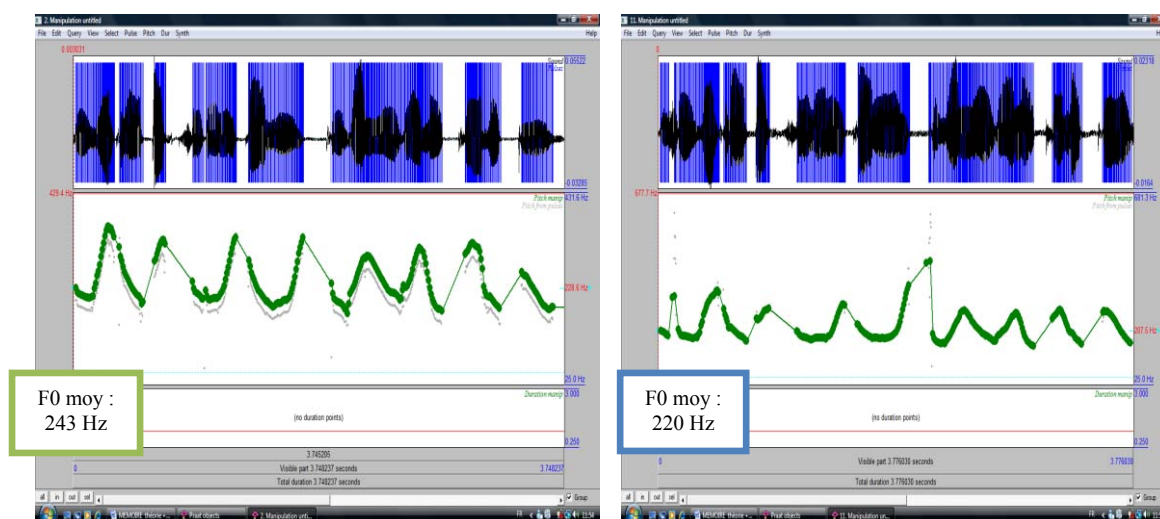
Fmax observée = 328 Hz
 Fmin observée = 150 Hz
 δ observé = 328 – 150 = 178 Hz

- $\alpha = 186 / 178 = 1,04$
 - $\beta = 360 - [328 \times 1,04] = 19$

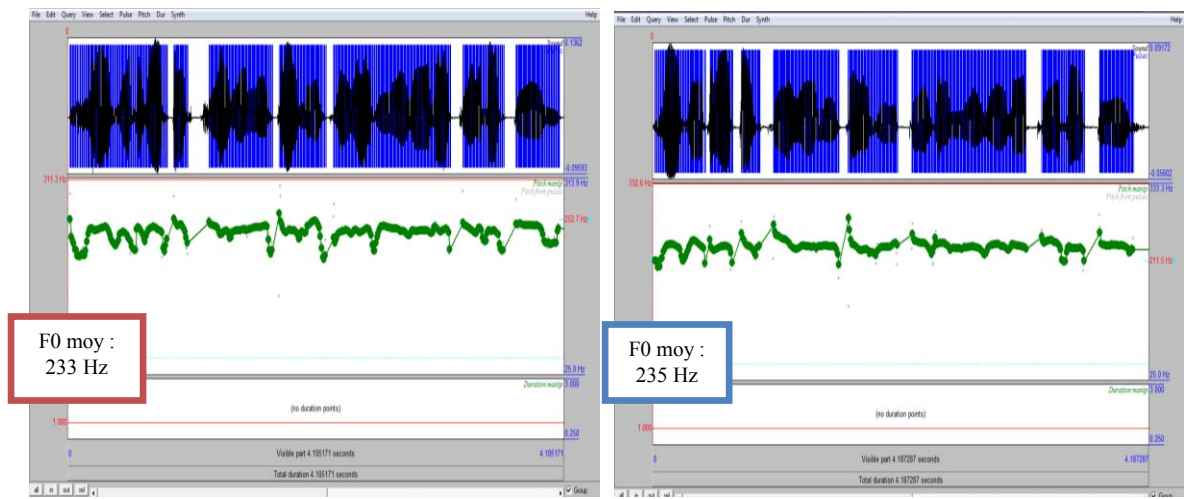
Ligne monotone = 174 + 186/2 = 267 Hz.



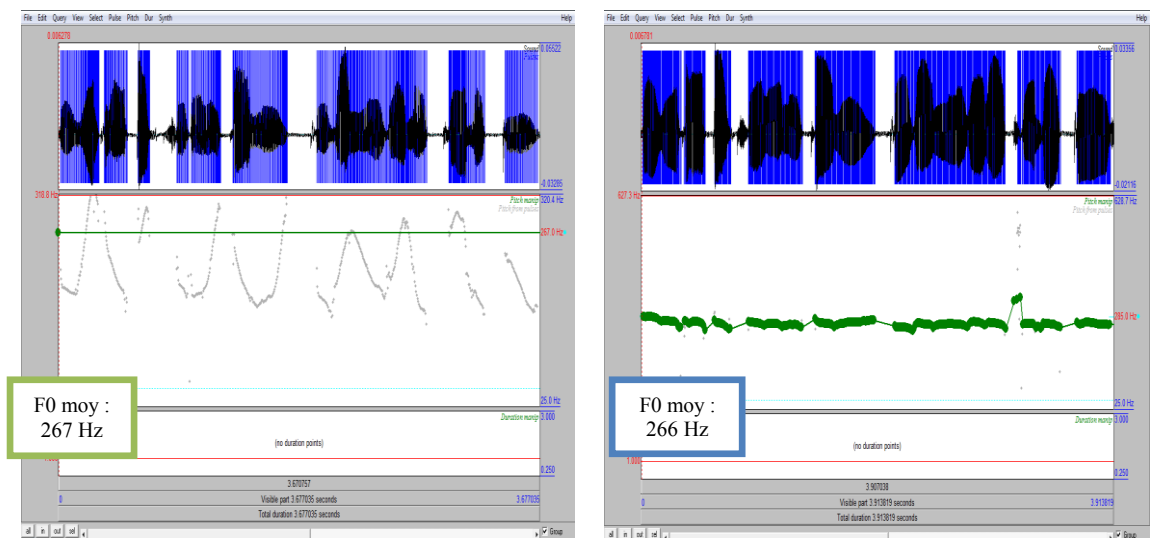
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du témoin n°10



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du témoin n°10



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du témoin n°10



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du témoin n°10

RESULTATS ET LIGNES MELODIQUES TEMOIN N°11

[ma] =} min : 173 Hz max : 293 Hz
 [mi] =} min : 184 Hz max : 383 Hz
 [mu] =} min : 185 Hz max : 273 Hz

Ecart voulu : [185 Hz / 273 Hz]

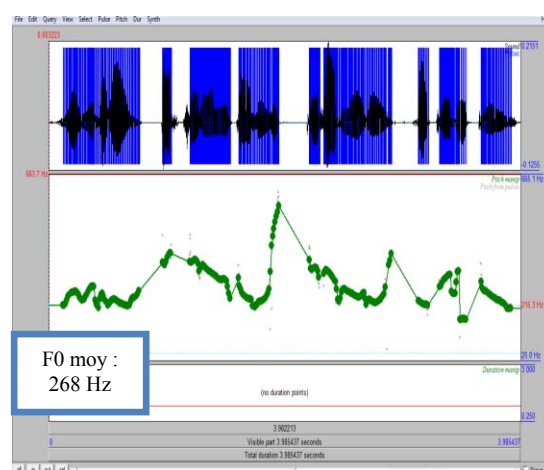
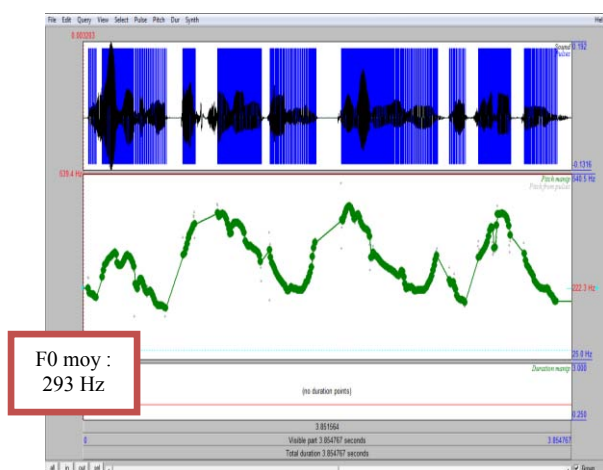
Fmax voulue = 273 Hz
 Fmin voulue = 185 Hz
 δ voulu = 273 – 185 = 98 Hz

Fmax observée = 439 Hz
 Fmin observée = 152 Hz
 δ observé = 439 – 152 = 287 Hz

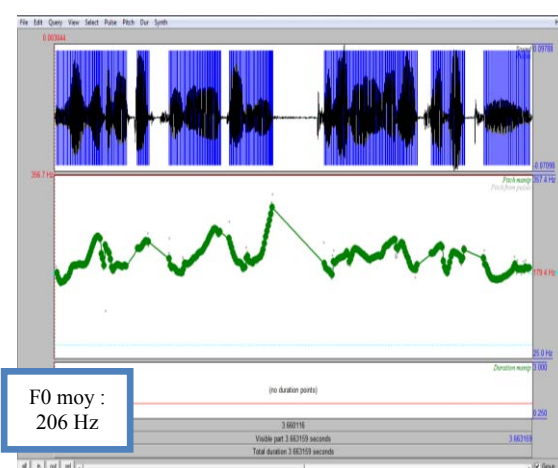
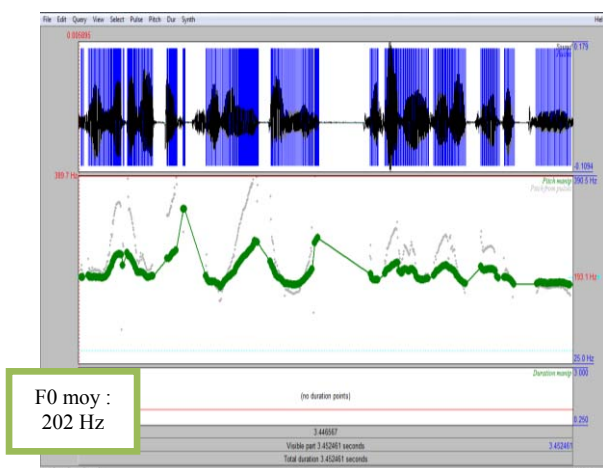
- α = 98 / 287 = 0,34

- β = 273 – [439 x 0,34] = 123

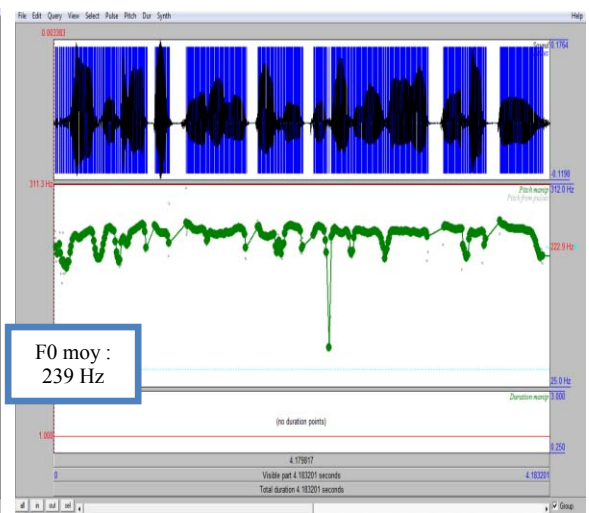
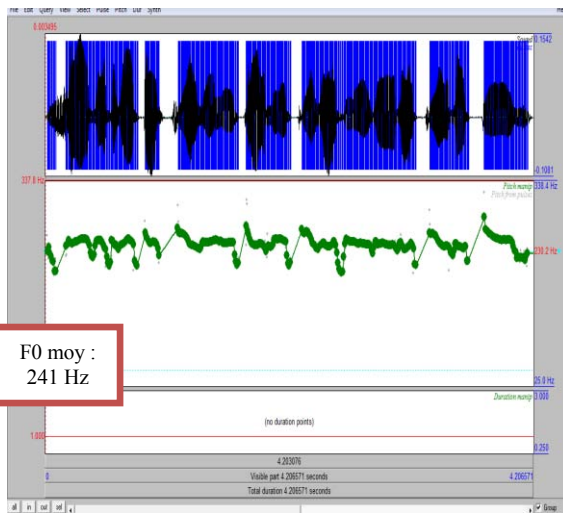
Ligne monotone = 185 + 98/2 = 229 Hz.



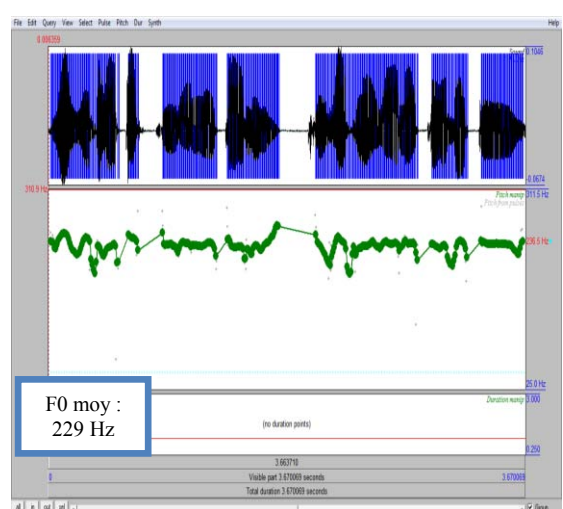
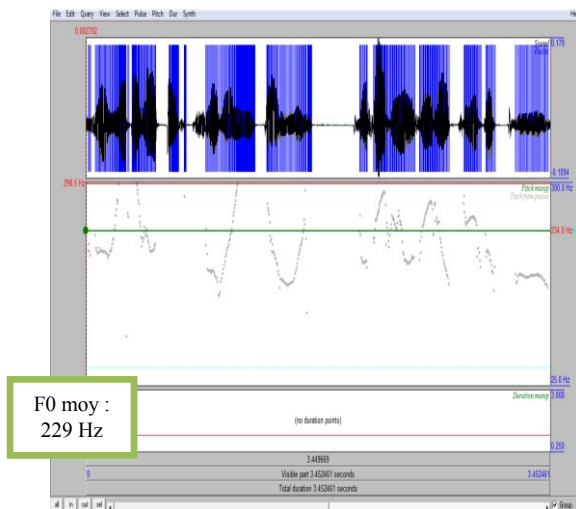
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du témoin n°11



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du témoin n°11



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du témoin n°11



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du témoin n°11

RESULTATS ET LIGNES MELODIQUES TEMOIN N°12

[ma] =} min : 172 Hz max : 357 Hz
 [mi] =} min : 175 Hz max : 286 Hz
 [mu] =} min : 162 Hz max : 311 Hz

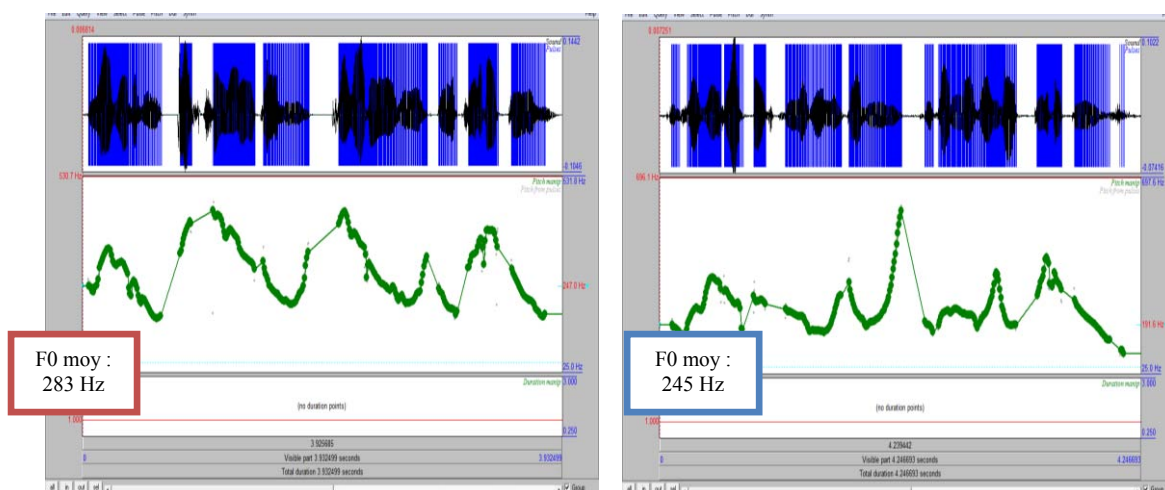
Ecart voulu : [175 Hz / 286 Hz]

Fmax voulue = 286 Hz
 Fmin voulue = 175 Hz
 δ voulu = 286 – 175 = 111 Hz

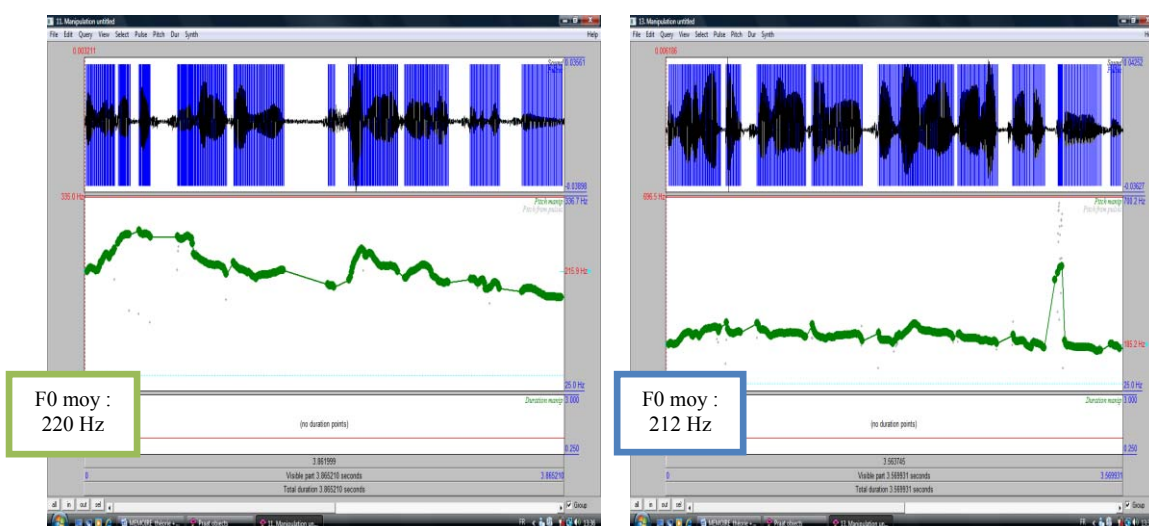
Fmax observée = 340 Hz
 Fmin observée = 139 Hz
 δ observé = 340 – 139 = 201 Hz

- $\alpha = 111 / 201 = 0,55$
 - $\beta = 286 - [340 \times 0,55] = 98$

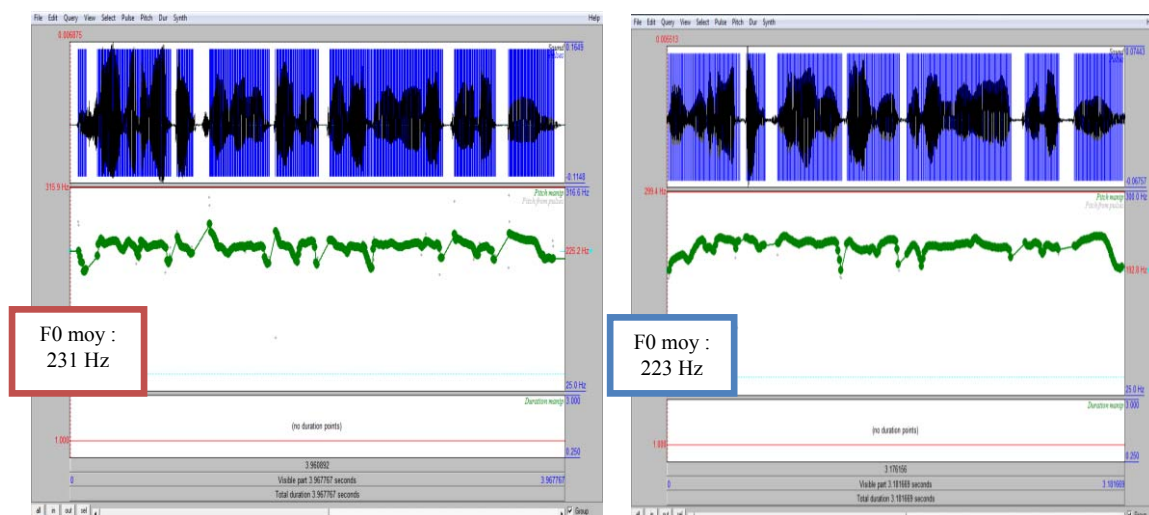
Ligne monotone = 175 + 111/2 = 230,5 Hz (que l'on arrondit à 231 Hz).



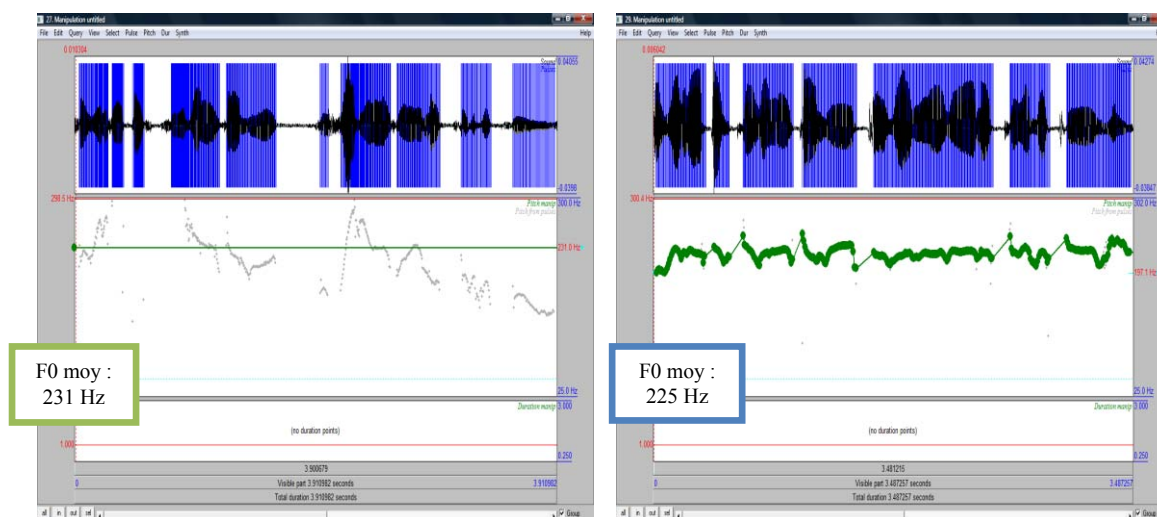
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du témoin n°12



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du témoin n°12



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du témoin n°12



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du témoin n°12

RESULTATS ET LIGNES MELODIQUES TEMOIN N°13

[ma] => min : 182 Hz max : 362 Hz
 [mi] => min : 183 Hz max : 354 Hz
 [mu] => min : 177 Hz max : 436 Hz

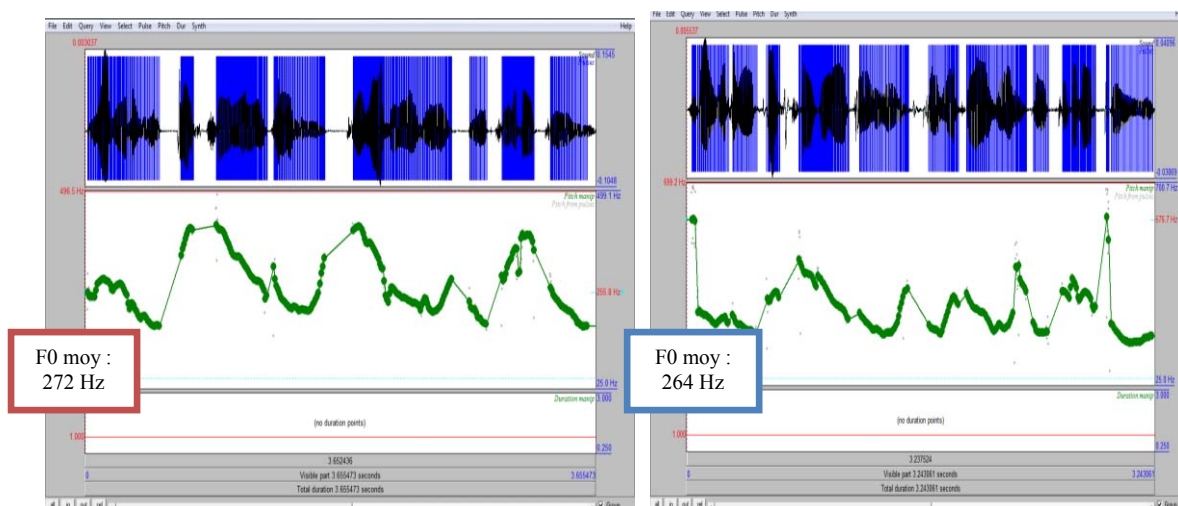
Ecart voulu : [183 Hz / 354 Hz]

Fmax voulue = 354 Hz
 Fmin voulue = 183 Hz
 δ voulu = 354 – 183 = 171 Hz

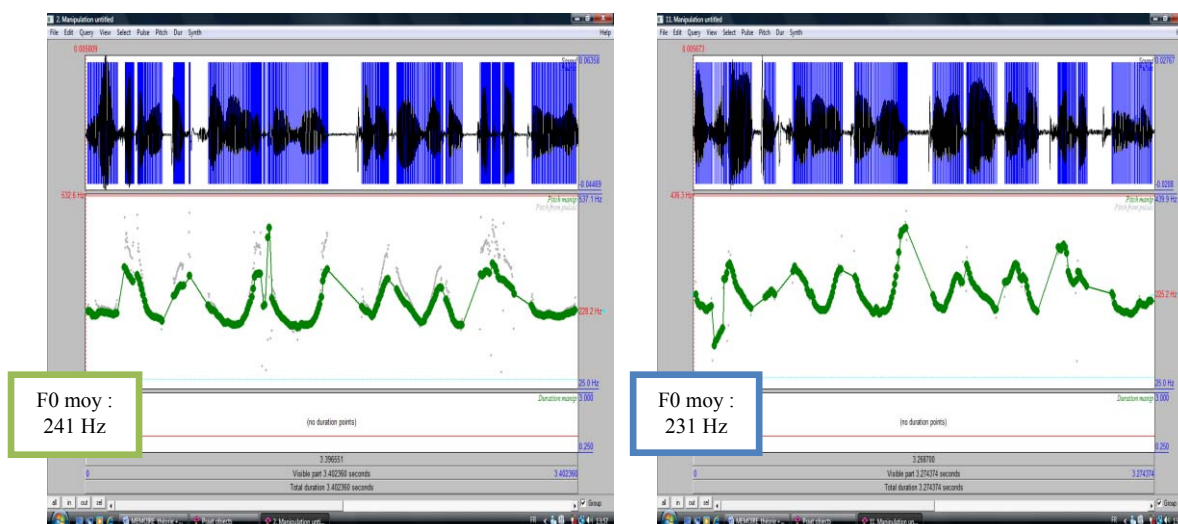
Fmax observée = 443 Hz
 Fmin observée = 172 Hz
 δ observé = 443 – 172 = 271 Hz

- $\alpha = 171 / 271 = 0,63$
 - $\beta = 354 - [443 \times 0,63] = 75$

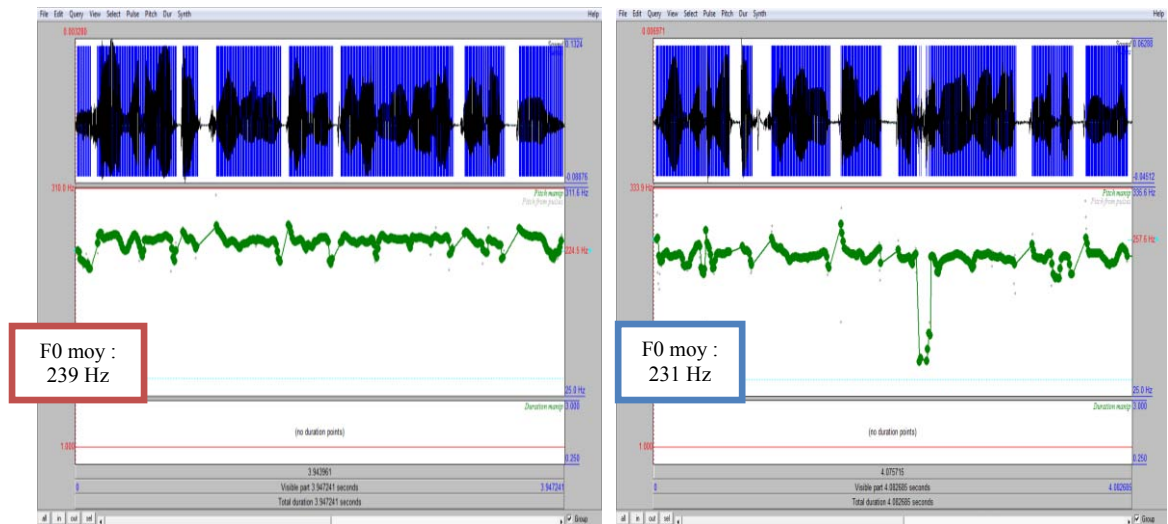
Ligne monotone = 183 + 171/2 = 268,5 Hz (que l'on arrondit à 269 Hz).



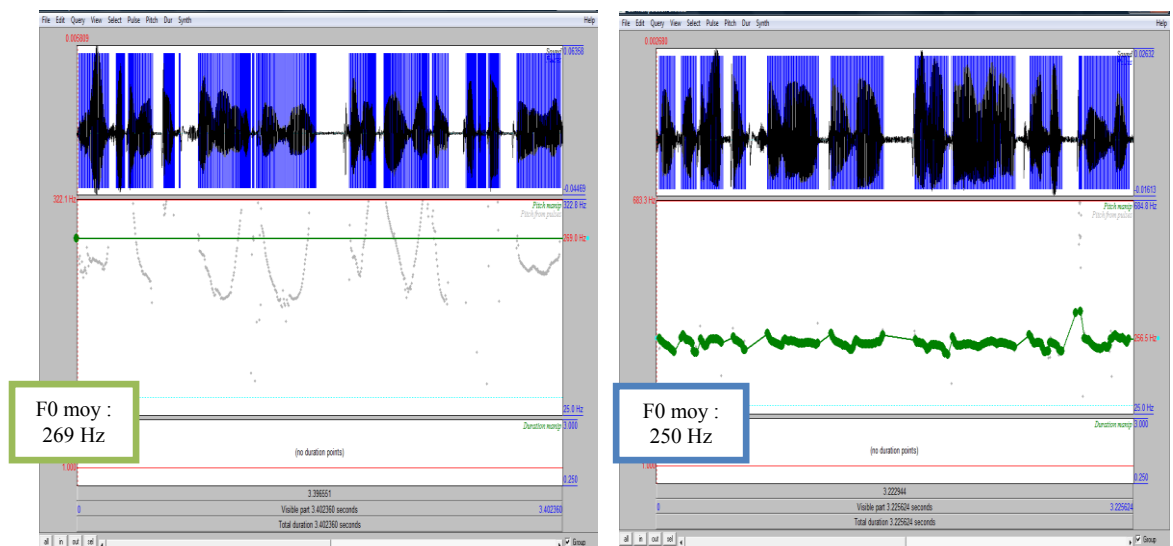
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du témoin n°13



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du témoin n°13



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du témoin n°13



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du témoin n°13

RESULTATS ET LIGNES MELODIQUES TEMOIN N°14

[ma] =} min : 101 Hz max : 245 Hz
 [mi] =} min : 104 Hz max : 282 Hz
 [mu] =} min : 105 Hz max : 308 Hz

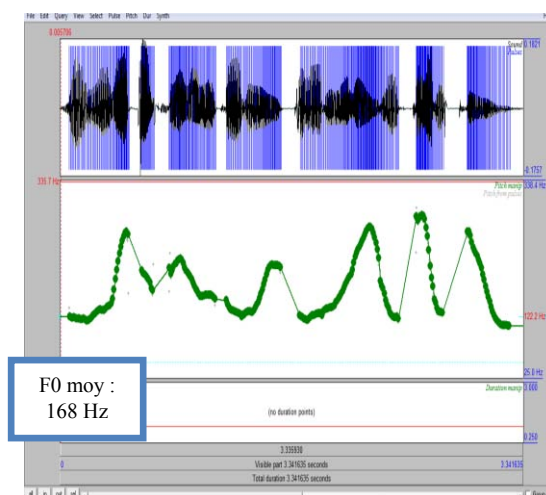
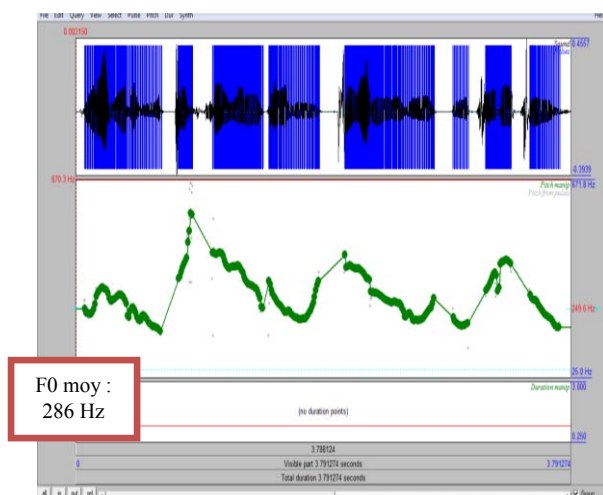
Ecart voulu : [105 Hz / 245 Hz]

Fmax voulue = 245 Hz
 Fmin voulue = 105 Hz
 δ voulu = 245 – 105 = 140 Hz

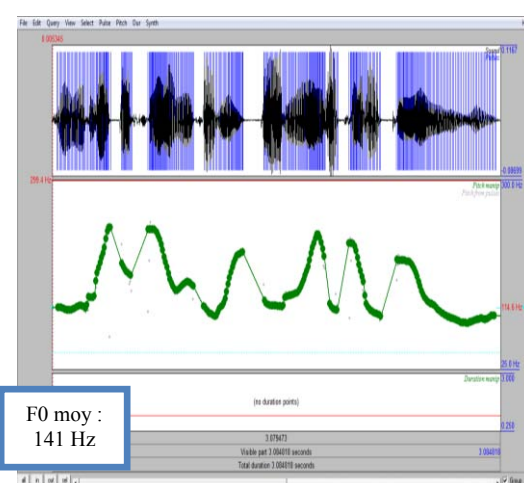
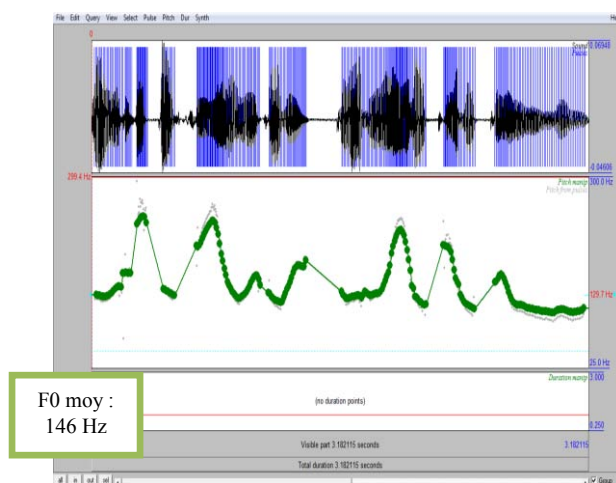
Fmax observée = 263 Hz
 Fmin observée = 93 Hz
 δ observé = 263 – 93 = 170 Hz

- $\alpha = 140 / 170 = 0,82$
 - $\beta = 245 - [263 \times 0,82] = 28$

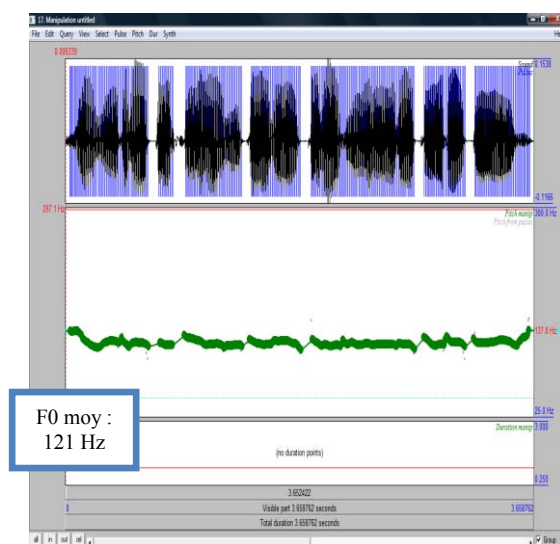
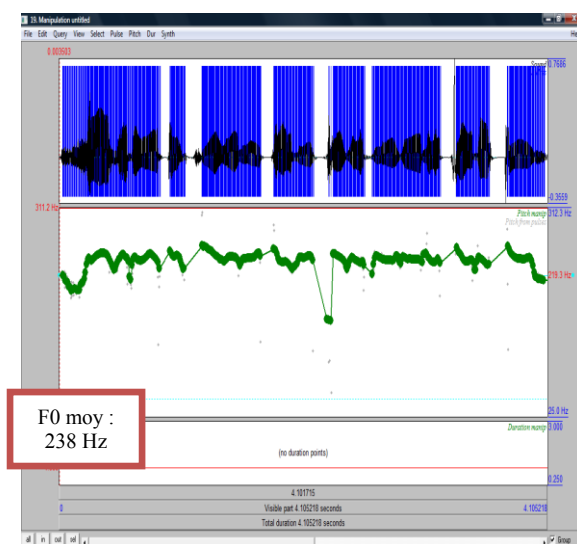
Ligne monotone = 105 + 140/2 = 175 Hz.



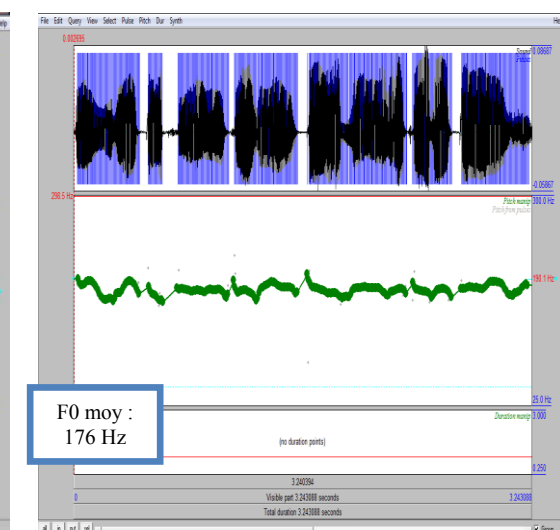
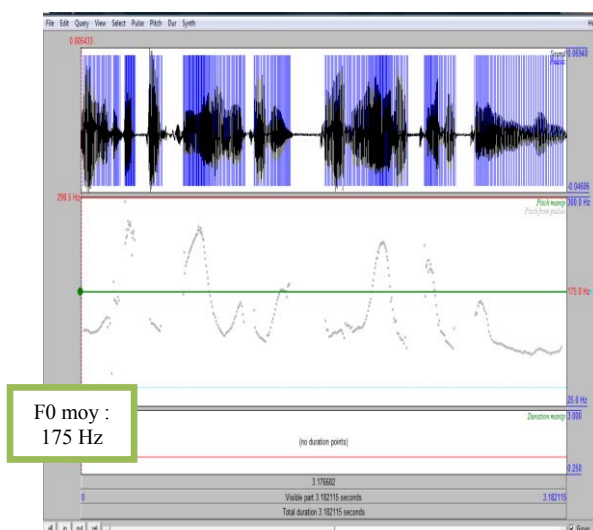
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du témoin n°14



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du témoin n°14



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du témoin n°14



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du témoin n°14

RESULTATS ET LIGNES MELODIQUES PATIENT A

[ma] =} min : 171 Hz max : 344 Hz
 [mi] =} min : 145 Hz max : 368 Hz
 [mu] =} min : 174 Hz max : 395 Hz

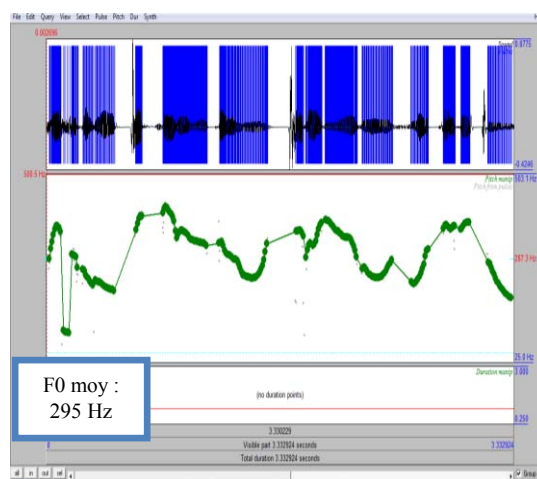
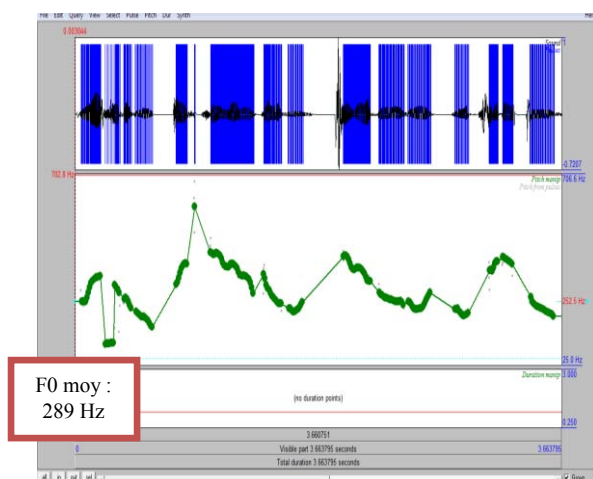
Ecart voulu : [174 Hz / 344 Hz]

Fmax voulue = 344 Hz
 Fmin voulue = 174 Hz
 δ voulu = 344 – 174 = 170 Hz

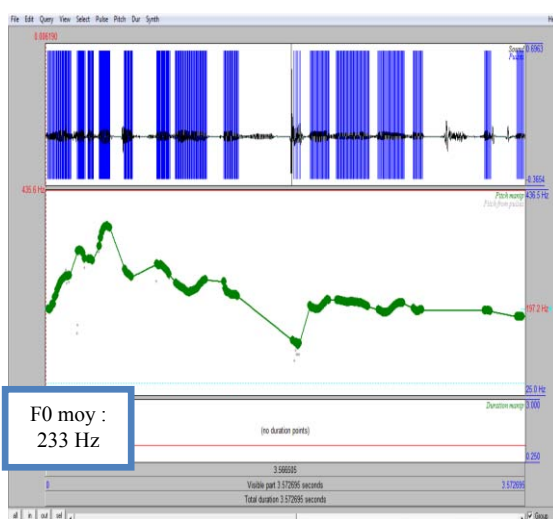
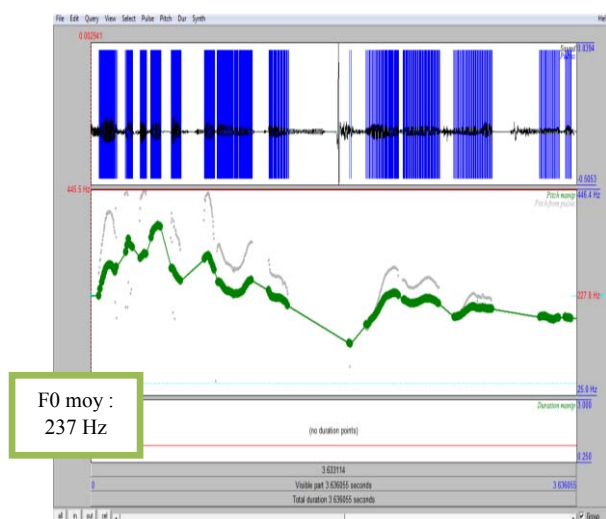
Fmax observée = 498 Hz
 Fmin observée = 158 Hz
 δ observé = 498 – 158 = 330 Hz

- $\alpha = 170 / 330 = 0,51$
 - $\beta = 344 - [498 \times 0,51] = 88$

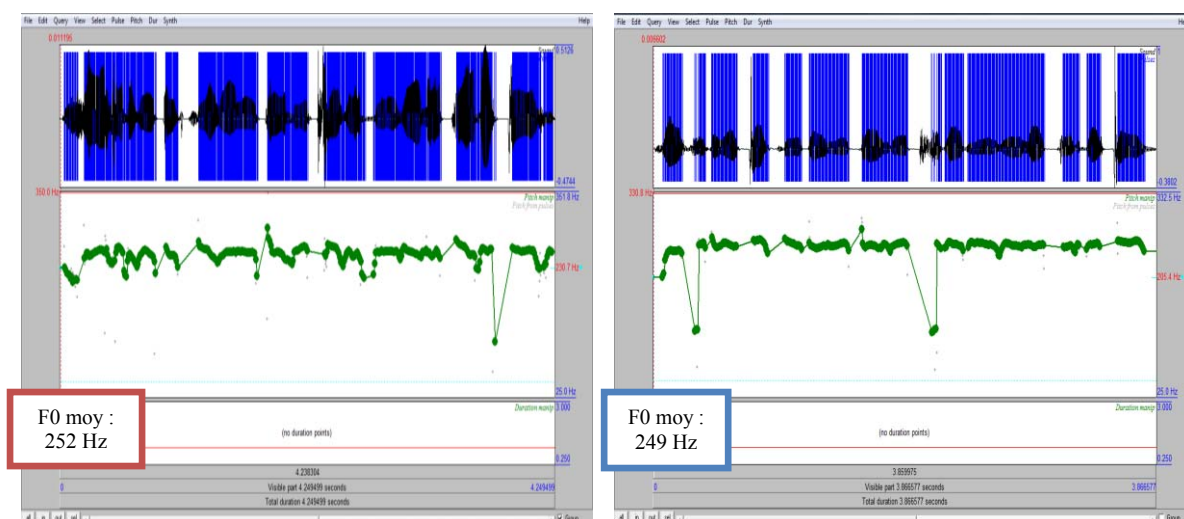
Ligne monotone = 174 + 170/2 = 259 Hz.



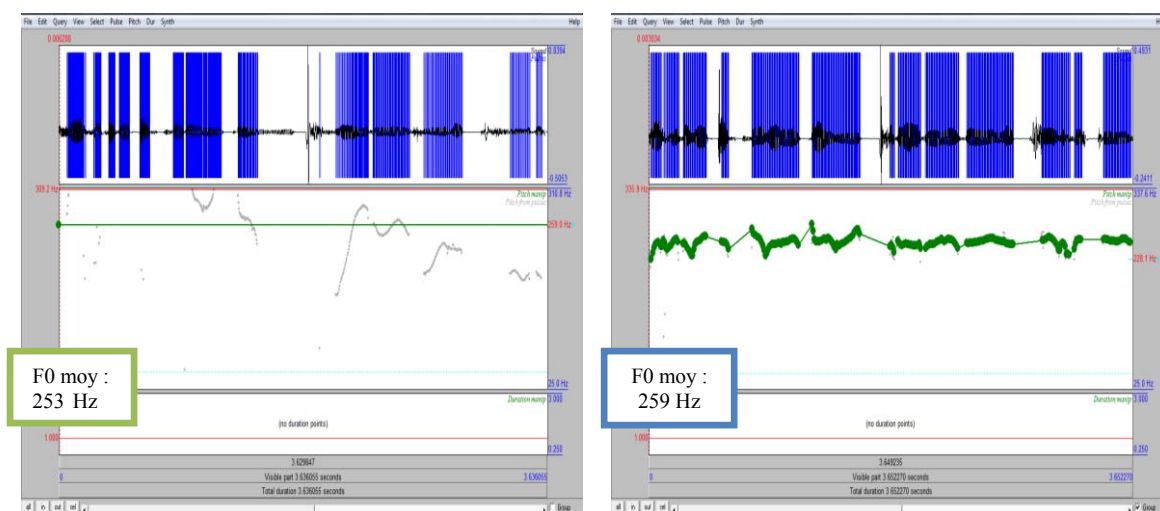
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du patient A



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du patient A



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du patient A



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du patient A

RESULTATS ET LIGNES MELODIQUES PATIENT B

[ma] =} min : 130 Hz max : 274 Hz
 [mi] =} min : 137 Hz max : 273 Hz
 [mu] =} min : 144 Hz max : 272 Hz

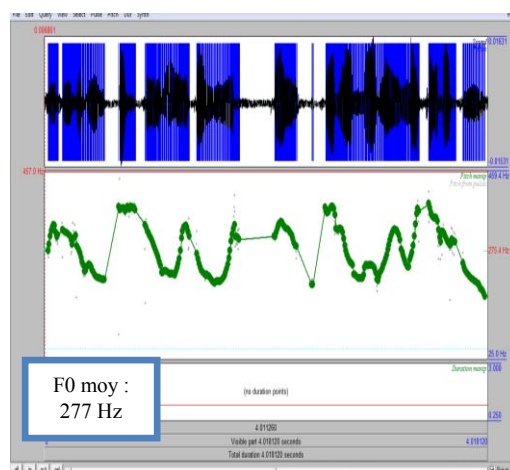
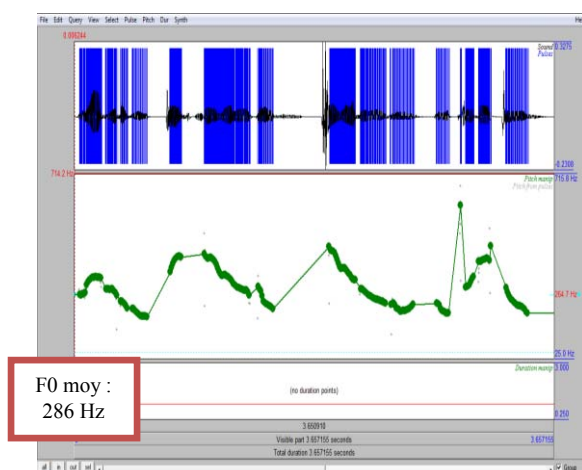
Ecart voulu : [144 Hz / 272 Hz]

Fmax voulue = 272 Hz
 Fmin voulue = 144 Hz
 δ voulu = 272 - 144 = 128 Hz

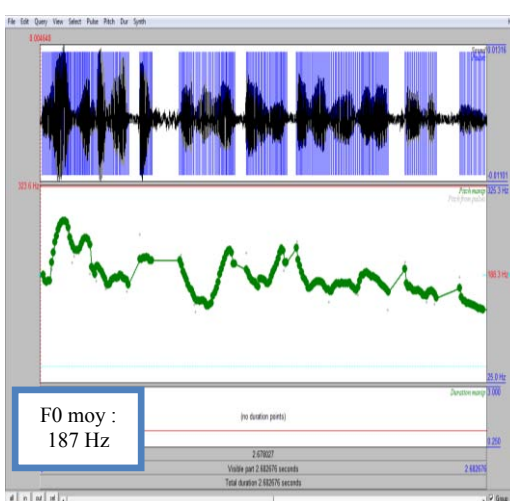
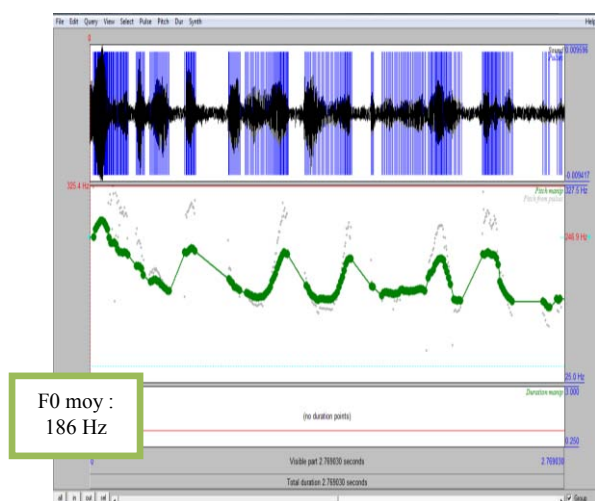
Fmax observée = 382 Hz
 Fmin observée = 119 Hz
 δ observé = 382 - 119 = 263 Hz

- $\alpha = 128 / 263 = 0,49$
 - $\beta = 272 - [382 \times 0,49] = 84,82$

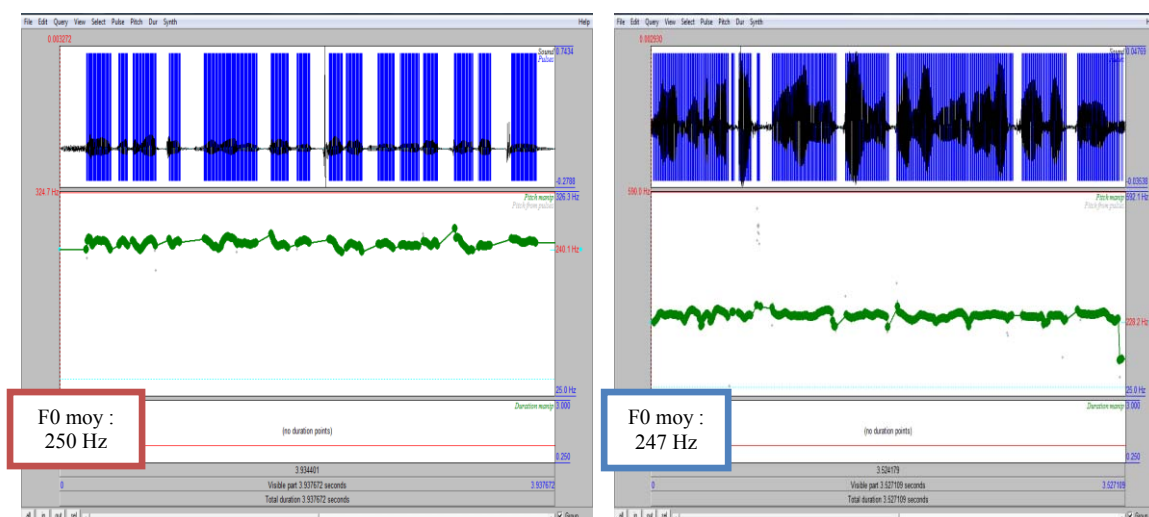
Ligne monotone = 144 + 128/2 = 208 Hz.



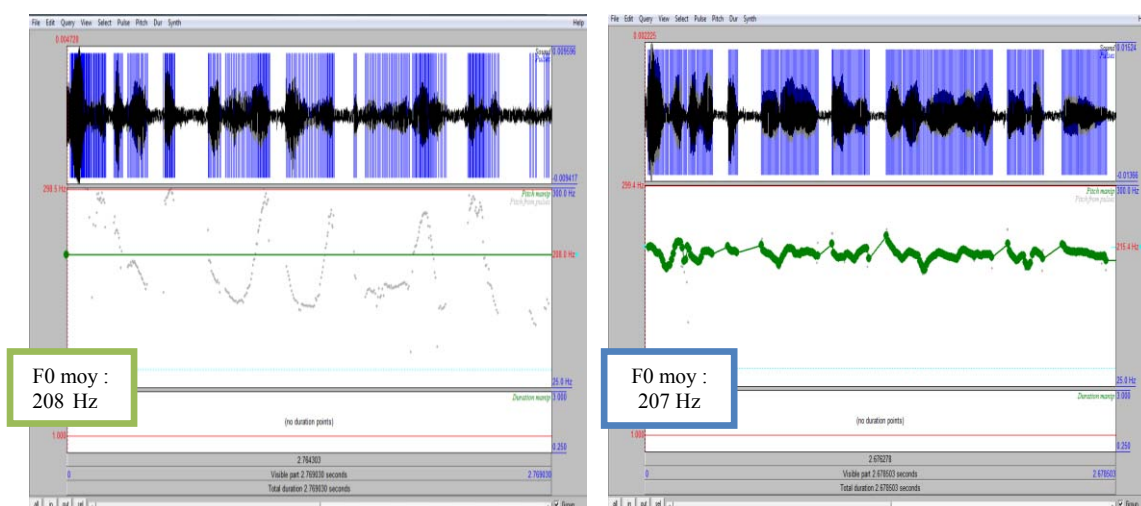
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du patient B



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du patient B



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du patient B



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du patient B

RESULTATS ET LIGNES MELODIQUES PATIENT C

[ma] =} min : 154 Hz max : 253 Hz
 [mi] =} min : 142 Hz max : 249 Hz
 [mu] =} min : 156 Hz max : 261 Hz

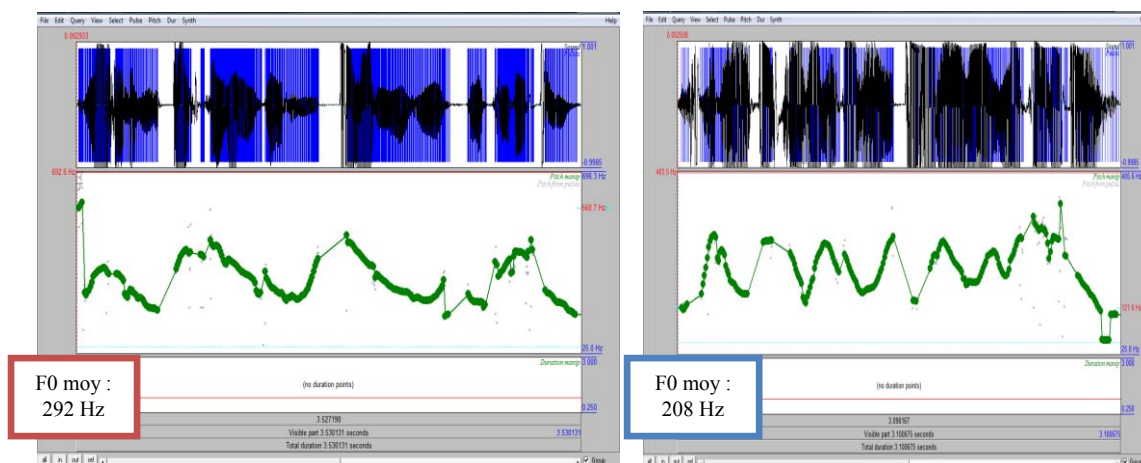
Ecart voulu : [156 Hz / 261 Hz]

Fmax voulue = 249 Hz
 Fmin voulue = 156 Hz
 δ voulu = 249 – 156 = 93 Hz

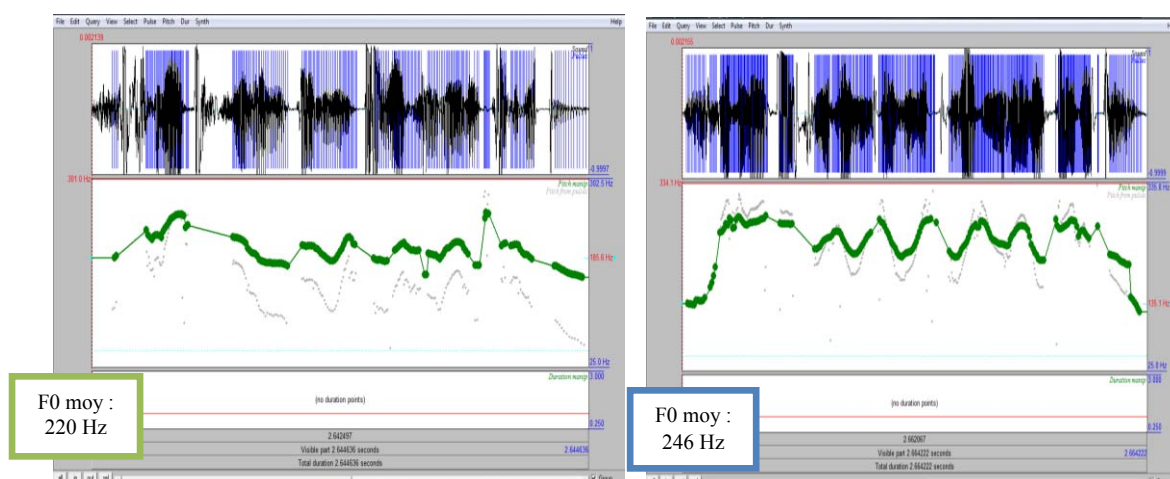
Fmax observée = 247 Hz
 Fmin observée = 58 Hz
 δ observé = 247 – 58 = 189 Hz

- α = 93 / 189 = 0,49
 - β = 249 – [247 x 0,49] = 127,97

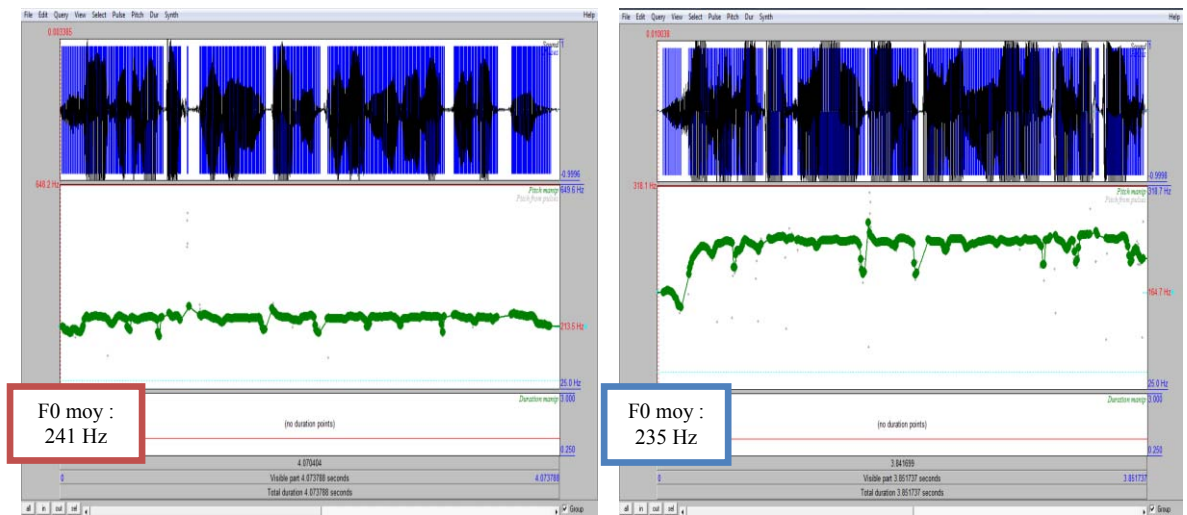
Ligne monotone = 156 + 93/2 = 203 Hz.



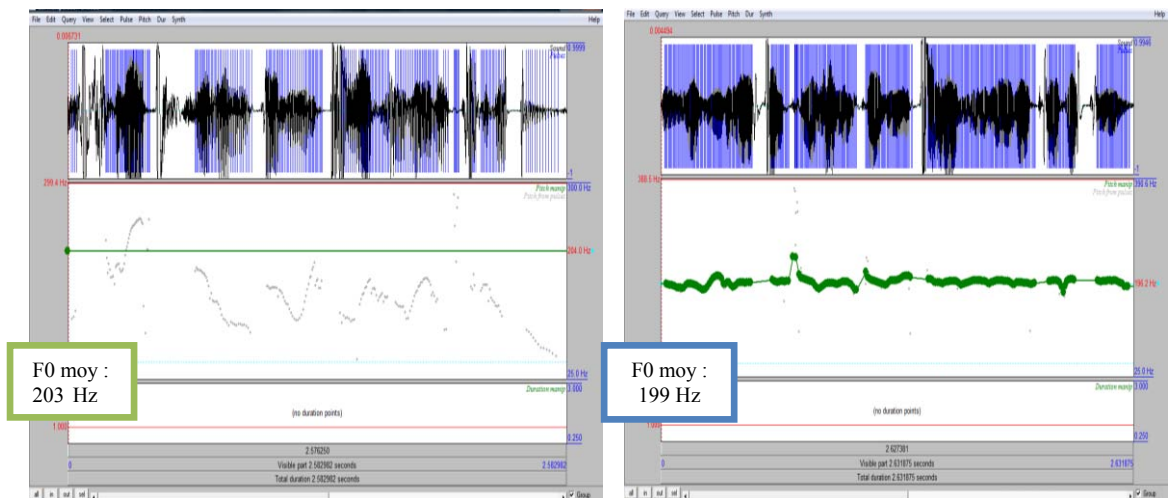
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du patient C



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du patient C



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du patient C



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du patient C

RESULTATS ET LIGNES MELODIQUES PATIENT D

[ma] =} min : 219 Hz max : 332 Hz
 [mi] =} min : 216 Hz max : 338 Hz
 [mu] =} min : 218 Hz max : 361 Hz

Ecart voulu : [219 Hz / 332 Hz]

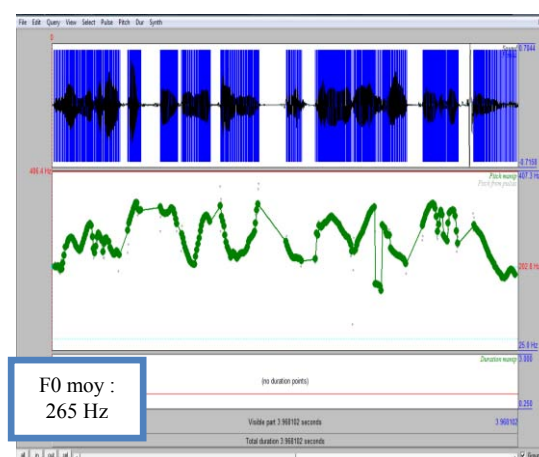
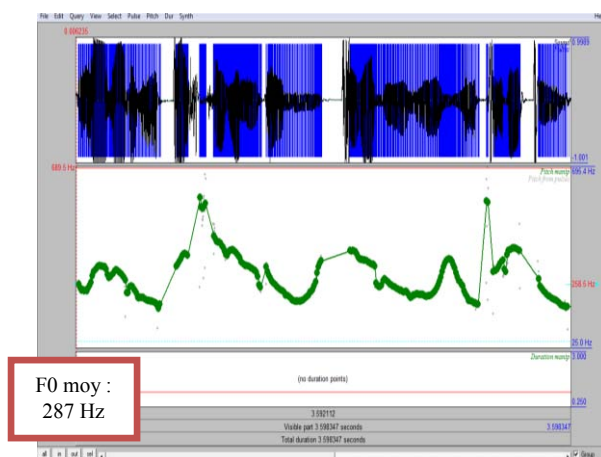
Fmax voulue = 332 Hz
 Fmin voulue = 219 Hz
 δ voulu = 332 – 219 = 113 Hz

Fmax observée = 325 Hz
 Fmin observée = 184 Hz
 δ observé = 325 – 184 = 144 Hz

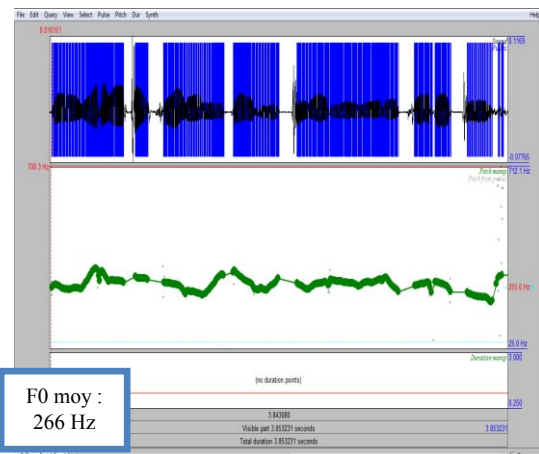
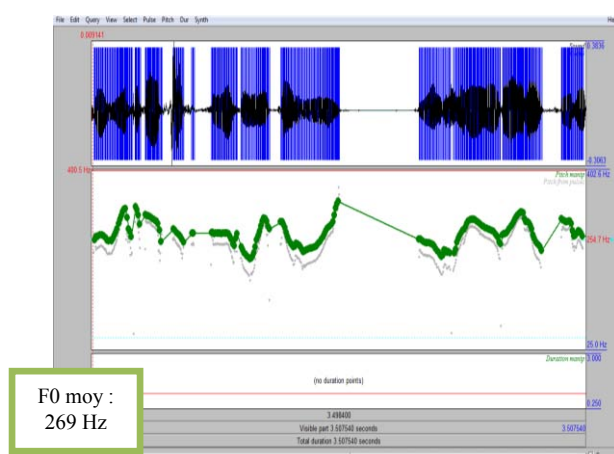
- $\alpha = 113 / 144 = 0,8$

- $\beta = 332 - [325 \times 0,8] = 72$

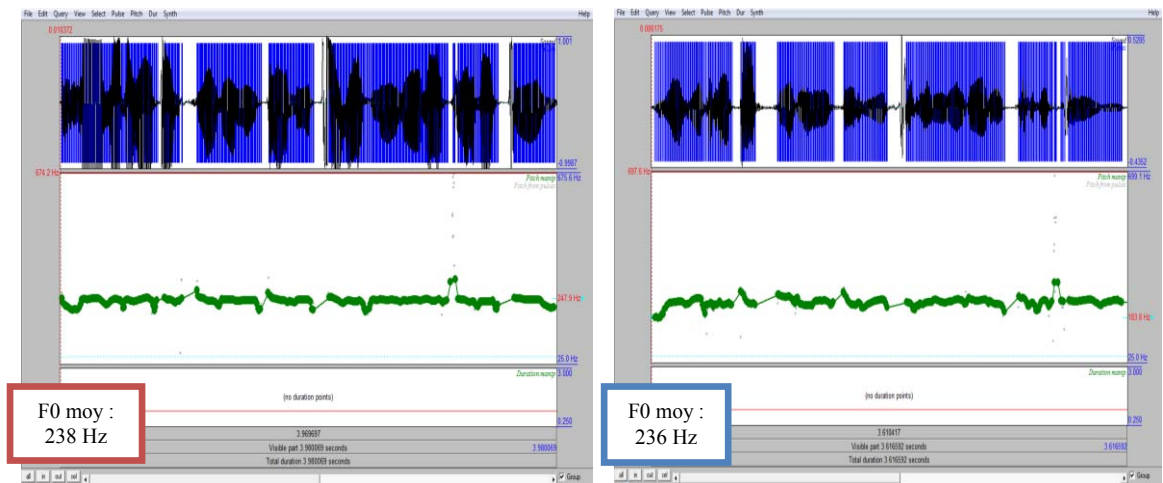
Ligne monotone = 219 + 113/2 = 276 Hz.



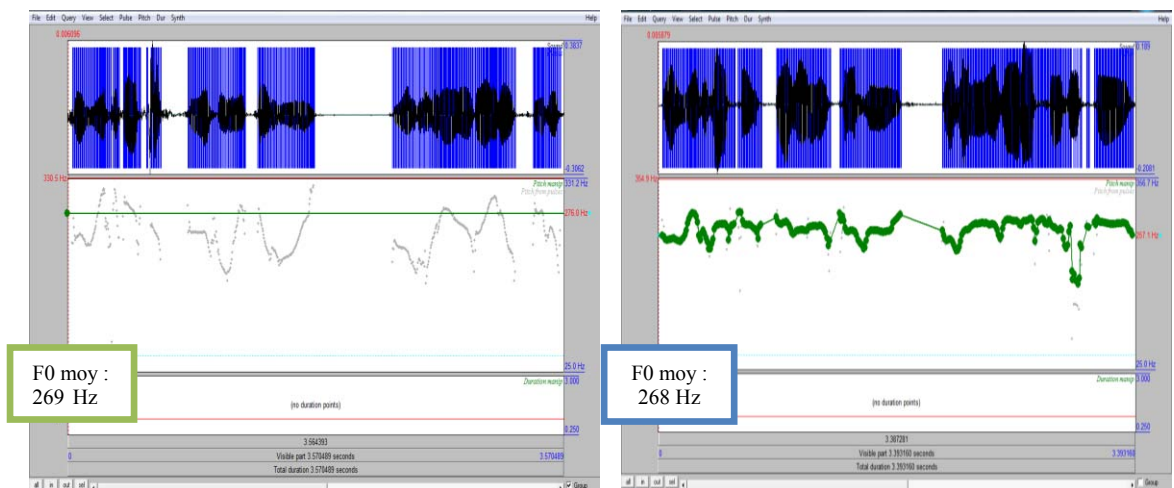
Ligne mélodique du modèle orthophonique mélodique suivie de l'imitation du patient D



Ligne mélodique du modèle voix transformée mélodique suivie de l'imitation du patient D



Ligne mélodique du modèle orthophonique monotone suivie de l'imitation du patient D



Ligne mélodique du modèle voix transformée monotone suivie de l'imitation du patient D

Il s'agit dans ce travail de démontrer que les capacités d'imitation existent chez un sujet sain comme dysphonique. Et plus largement, de définir un modèle spécialement adapté à la voix du sujet respectant une « zone de sécurité prosodique » (zone de la dynamique vocale où la voix a le meilleur rendement, située à l'intérieur du triangle vocalique) afin de montrer que l'imitation de ce modèle est plus adaptée que l'imitation de la voix du thérapeute lorsqu'il s'agit d'amener le sujet dans la zone fréquentielle limitant au maximum le malmenage vocal. Pour ce faire, nous avons réalisé des enregistrements vocaux d'un imitateur, de patients dysphoniques et de témoins grâce au logiciel PRAAT, avec usage d'un micro et de haut-parleurs. Le protocole s'est déroulé en quatre étapes d'enregistrement : glissades sur [ma], [mi], [mu] pour définir la zone de la dynamique vocale dite sécurisée, production spontanée sur ton mélodique et monotone, imitation du modèle orthophonique sur ton mélodique et monotone, imitation du modèle de la voix du sujet respectant la zone de sécurité prosodique sur ton mélodique et monotone. Il ressort que la plupart des sujets sont capables d'imitation et que 97% d'entre eux réalisent une meilleure imitation lorsqu'il s'agit de leur propre voix, leur production se rapprochant davantage de la zone sécurisée. Cela témoigne que les sujets ont des capacités d'imitation et que la présentation d'un modèle spécialement adapté au sujet pour illustrer la zone limitant au maximum le risque de malmenage vocal est plus adéquat et mieux imité que le modèle orthophonique. Ce moyen constitue une aide supplémentaire à disposition du thérapeute pour donner le repère vocal dont le patient a besoin pour objectiver sa voix et l'utiliser dans les meilleures fréquences. L'objectif premier étant d'offrir au patient la possibilité de demeurer ou se replier dans la zone vocale sécurisée lui correspondant pour retrouver ou acquérir une efficacité vocale, améliorer ou prévenir l'apparition de lésions consécutives à une mauvaise utilisation de la voix.

Mots clés : dysphonie ; voix ; évaluation ; expérimentation ; adulte ; imitation.