

## Des enregistrements aux corpus : transcription et extraction de données d'interprétation en milieu médical

Natacha Niemants

Volume 63, Number 3, December 2018

Traductologie de corpus : 20 ans après

URI: <https://id.erudit.org/iderudit/1060168ar>

DOI: <https://doi.org/10.7202/1060168ar>

[See table of contents](#)

Publisher(s)

Les Presses de l'Université de Montréal

ISSN

0026-0452 (print)

1492-1421 (digital)

[Explore this journal](#)

Cite this article

Niemants, N. (2018). Des enregistrements aux corpus : transcription et extraction de données d'interprétation en milieu médical. *Meta*, 63(3), 665–694.  
<https://doi.org/10.7202/1060168ar>

Article abstract

In the field of dialogue interpreting (DI), both the use and usefulness of corpora is controversial. While on the one hand, DI research is increasingly orienting to the study of naturally-occurring interaction with a consequent increase in recorded and transcribed materials which might well be explored as corpora, on the other, DI interest in “interaction” rather than “text” has made the traditional methods of corpus linguistics less attractive for DI studies. In this paper, we report on our experience with a rather large DI collection of data (the AIM corpus), which we are implementing for corpus use. In particular, we have experimented the *EXMARaLDA* and *ELAN* software on a sub-set of 65 encounters whose transcripts are aligned with audio. So far, our experiments have taken two directions: (a) exploring potentially interesting lexical items for the analysis of structures of interaction; (b) extracting audio-aligned segments for use in e.g. specific training or learning activities. The paper thereby aims to show the use and utility of a DI corpus in the healthcare environment.

# Des enregistrements aux corpus : transcription et extraction de données d'interprétation en milieu médical

**NATACHA NIEMANTS**

*Université de Bologne, Forlì, Italie\**

*natacha.niemants@unibo.it*

## RÉSUMÉ

Dans le domaine de l'interprétation de dialogue (ID), l'utilisation et l'utilité des corpus sont autant de sujets controversés. D'une part, la recherche est de plus en plus orientée vers l'étude de l'interaction authentique, ce qui entraîne une augmentation des données enregistrées et transcrites pouvant être explorées comme corpus ; de l'autre, l'intérêt de l'ID envers l'«interaction» plutôt que le «texte» diminue l'attrait des méthodes traditionnelles de la linguistique de corpus pour l'étude des données d'interprétation. Dans cet article, nous présenterons une collection de données d'ID assez importante (corpus AIM), que nous cherchons à outiller. Nous avons jusqu'à présent testé les logiciels EXMARaLDA et ELAN sur un sous-ensemble de 65 rencontres, dont la transcription est alignée avec l'audio, et nous avons entamé deux pistes de recherche que nous allons illustrer : (a) l'exploration d'éléments lexicaux potentiellement intéressants pour l'analyse des structures de l'interaction ; (b) l'extraction de segments alignés à l'audio que l'on puisse entre autres utiliser dans des activités d'enseignement et d'apprentissage. Nous souhaitons montrer ainsi tant l'utilisation que l'utilité d'un corpus d'ID en milieu médical.

## ABSTRACT

In the field of dialogue interpreting (DI), both the use and usefulness of corpora is controversial. While on the one hand, DI research is increasingly orienting to the study of naturally-occurring interaction with a consequent increase in recorded and transcribed materials which might well be explored as corpora, on the other, DI interest in “interaction” rather than “text” has made the traditional methods of corpus linguistics less attractive for DI studies. In this paper, we report on our experience with a rather large DI collection of data (the AIM corpus), which we are implementing for corpus use. In particular, we have experimented the EXMARaLDA and ELAN software on a sub-set of 65 encounters whose transcripts are aligned with audio. So far, our experiments have taken two directions: (a) exploring potentially interesting lexical items for the analysis of structures of interaction; (b) extracting audio-aligned segments for use in e.g. specific training or learning activities. The paper thereby aims to show the use and utility of a DI corpus in the healthcare environment.

## RESUMEN

En el ámbito de la interpretación bilateral (IB), tanto el uso como la utilidad de los corpus son temas controvertidos. Mientras, por un lado, la investigación en IB se está orientando cada vez más hacia el estudio de la interacción natural, con un consiguiente incremento de materiales grabados y transcritos que podrían explorarse como corpus, por otro lado, el interés de la IB hacia la «interacción» más que hacia el «texto» ha hecho que los métodos tradicionales de la lingüística de corpus no fueran tan atractivos para los estudios de interpretación bilateral. En el presente estudio presentamos nuestra experiencia con una amplia colección de datos de IB (el corpus AIM), que estamos implementando para el uso de corpus. En particular, hemos experimentado los softwares EXMARaLDA

y ELAN en un subconjunto de 65 encuentros, cuya transcripción se ha alineado con el audio. Hasta ahora, nuestros experimentos han seguido dos direcciones principales, que ilustraremos: (a) explorar términos de búsqueda potencialmente interesantes para el análisis de las estructuras de la interacción; (b) extraer segmentos alineados con el audio para su uso en actividades específicas de formación o aprendizaje. El objetivo de este trabajo es el de demostrar el uso y la utilidad de un corpus de IB de ámbito sanitario.

#### MOTS-CLÉS/KEYWORDS/PALABRAS CLAVE

interprétation de dialogue, corpus, transcription, ELAN, EXMARaLDA

dialogue interpreting, corpus, transcription, ELAN, EXMARaLDA

interpretación bilateral, corpus, transcripción, ELAN, EXMARaLDA

Like all speech, interpreting dies on the air. In order to study it,  
we need to resurrect the corpse by recording and transcribing it,  
thereby transforming the corpse into a corpus.

(Cencini et Aston 2005: 23)

### 1. Introduction

Dans les études et la pratique de la traduction écrite, les corpus sont de plus en plus exploités pour mieux comprendre les patrons lexicaux équivalents et observer l'utilisation des termes et des unités phraséologiques dans des types/genres de textes comparables (Zanettin 2012). En outre, il existe des tentatives de modélisation des corpus pour la recherche en interprétation de conférence (Cencini et Aston 2002), où l'emploi des corpus pour la création de glossaires (Williams 2008) pourrait être utile pour la pratique professionnelle.

Pour ce qui est, par contre, de l'interprétation de dialogue (ID), une étiquette proposée par Mason (1999) pour décrire des échanges impliquant au moins trois participants, dont un interprète qui aide les autres à communiquer<sup>1</sup>, l'utilisation et l'utilité des corpus sont autant de sujets controversés. D'une part, les recherches sont de plus en plus orientées vers l'étude des interactions authentiques, ce qui entraîne une augmentation des données enregistrées et transcrites qui pourraient être exploitées comme corpus. En effet, depuis les publications pionnières de Wadensjö (1998) et Mason (1999, 2001), plusieurs chercheurs s'interrogent sur l'apport d'études sur l'interaction à l'analyse de l'ID en langue orale<sup>2</sup>, surtout dans des contextes socio-médicaux (Bolden 2000; Davidson 2000, 2002; Traverso 2003a; Angelelli 2004; Bot 2005; Merlini et Favaron 2005; Pasquandrea 2011; Niemants 2015a; Krystallidou 2016; Zahn et Zeng 2017) et socio-juridiques (Kadric 2000; Hale 2004; Mason 2006; Keselman, Cederborg, *et al.* 2010; Amato et Mack 2015; Ticca et Traverso 2017), mais aussi éducatifs (Davitti 2013) et télévisés (Straniero Sergio 2007)<sup>3</sup>. De l'autre, l'intérêt de l'ID envers l'«interaction» plutôt que le «texte» (Wadensjö 1998; Baraldi et Gavioli 2012), c'est-à-dire pour l'alternance, la co-construction et la temporalité des tours de parole outre que pour les mots dont chaque tour se compose, diminue l'attrait des méthodes traditionnelles de la linguistique de corpus, qui généralement privilégient la dimension textuelle (ce qui est dit) par rapport à la dimension temporelle (quand cela est dit) de la langue orale. C'est justement la temporalité de la parole en interaction, avec et sans interprète, qui conduit à problématiser le passage des enregistrements authentiques aux corpus de transcriptions de l'oral (Niemants 2012;

Straniero Sergio et Falbo 2012; Angermeyer, Meyer, *et al.* 2012). Qu'il s'agisse de « spoken corpus », c'est-à-dire de collections de transcriptions qui ne contiennent pas les données audio ou vidéo, ou de « speech corpus » où les données primaires, de pair avec leurs transcriptions, sont intégrées au corpus (Straniero Sergio et Falbo 2012 : 31-32)<sup>4</sup>, les spécificités d'une ID qui se co-construit dans le temps doivent être d'une manière ou d'une autre préservées.

Dans cet essai, nous partagerons notre expérience avec une collection de données d'ID assez importante (Corpus AIM, environ 550 rencontres, plus de 100 heures d'enregistrement), que nous cherchons maintenant à outiller. Nous commencerons par présenter les données qui ont été collectées au cours des quinze dernières années, en problématisant le processus qui porte à la constitution d'un corpus et en examinant deux logiciels qui ont été testés sur un sous-ensemble de 65 rencontres, dont la transcription est alignée avec l'audio (section 2). Nous illustrerons ensuite deux pistes de recherche que nous avons entamées, à savoir : l'exploration d'éléments lexicaux potentiellement intéressants pour l'analyse des structures de l'interaction avec interprète (section 3) ; l'extraction de segments alignés à l'audio que l'on puisse entre autres utiliser dans des activités d'enseignement et d'apprentissage (section 4). Nous espérons montrer tant l'utilisation que l'utilité d'un corpus d'ID en milieu médical et nous finirons par dresser un bilan de ce qui a été fait et de ce qui reste encore à faire pour tirer le plus grand profit d'une entreprise collective qui dépasse notre travail individuel.

## **2. La constitution d'un corpus d'interprétation de dialogue en milieu médical**

Dans le but de compléter une boîte à outils traditionnellement qualitatifs, les études sur l'interaction se sont récemment interrogées sur la constitution de corpus de données orales pour des analyses quantitatives, voire statistiques (Groupe ICOR 2010; Heritage, Elliot, *et al.* 2010; Stivers 2015). Or, l'observation des interactions authentiques qui a eu lieu au cours des vingt dernières années a non seulement contribué à l'essor de l'ID comme domaine de recherche, mais a également soulevé cette même question, à savoir : comment exploiter des données d'interprétation transcrites par ordinateur et donc susceptibles d'être analysées quantitativement<sup>5</sup> ?

Malgré l'absence de grands corpus d'ID (à cause, entre autres, de la difficulté d'accès aux données, du temps nécessaire pour les transcrire, ainsi que de l'impossibilité d'automatiser certaines requêtes)<sup>6</sup>, plusieurs études menées dans le cadre de projets ou laboratoires individuels ont dégagé des pistes nouvelles de recherche et commencé à combler l'écart entre ce que l'interprète de dialogue devrait théoriquement faire et sa pratique professionnelle (Merlini 2015), ainsi qu'entre la profession et la formation (Davitti et Pasquandrea 2015; Cirillo et Niemants 2017). Toujours en équilibre instable entre un désir de partage global et une exigence d'utilisation locale, ces démarches plus ou moins isolées ont parfois cherché à unir leurs forces en devenant des entreprises collectives, et ce, particulièrement là où les chercheurs commençaient à disposer d'un certain nombre d'enregistrements transcrits.

Cela a premièrement été le cas de Philipp Angermeyer (York University, Canada), Bernd Meyer (Universität Hamburg, Allemagne) et Thomas Schmidt (Archiv für Gesprochenes Deutsch, Allemagne), qui ont rassemblé des données provenant de

trois différents sous-corpus – le DiK en milieu socio-sanitaire (Bührig et Meyer 2004), le liSCC en milieu socio-juridique (Angermeyer 2006) et le SimDiK en milieu éducatif (Bührig, Kliche, *et al.* 2012) – et créé le ComInDat (Community Interpreting Database)<sup>7</sup>, dont les objectifs et les détails sont décrits dans Angermeyer, Meyer, *et al.* (2012). Ce projet vise à couvrir plusieurs contextes et plusieurs langues, à intégrer différents formats et conventions de transcription, à développer un standard commun pour l’annotation des données, à en faciliter l’échange, ainsi qu’à explorer le potentiel de systèmes de transcription intégrant le texte, l’audio et la vidéo. Le ComInDat est ainsi une grande première dans le domaine de l’ID et ses créateurs ont certainement contribué à ouvrir le débat sur les opportunités offertes par l’analyse qualitative et quantitative des corpus d’interprétation, et sur les défis que ces analyses présentent.

2.1. Les enregistrements du centre AIM

Le centre interuniversitaire d’analyse de l’interaction et de la médiation (dorénavant centre AIM) voit la participation de dix universités italiennes (Bologne, Gênes, Macerata, Modène et Reggio d’Émilie, Naples-L’Orientale, Pérouse Statale, Pérouse Stranieri, Rome 3, Sienne, Trieste) et d’une soixantaine de membres qui, tout en ayant des disciplines, des points de vue et des méthodologies différents, partageant un intérêt commun envers l’interaction, avec ou sans interprète.

Entre 2004 et 2018, les membres du Département d’études linguistiques et culturelles de l’Université de Modène et Reggio d’Émilie ont collecté, à eux seuls, plus de 120 heures d’interprétations en milieu médical, auxquelles s’ajoutent plus de 9 heures de conversations directes entre le personnel soignant et des patients étrangers. Le tableau 1 ci-dessous détaille les langues parlées dans 591 enregistrements audio (la vidéo n’ayant jusqu’à présent pas été acceptée dans ce contexte sensible), ainsi que la durée des échanges et le nombre de médiatrices – car il ne s’agit que de femmes – impliquées.

TABLEAU 1  
Toute la collection 2004-2018

| Paires de langues | Nombre de rencontres | Durée échange (minutes) | Nombre de médiatrices |
|-------------------|----------------------|-------------------------|-----------------------|
| Italien-Anglais   | 309                  | 3127’                   | 7                     |
| Italien-Arabe     | 169                  | 2498’                   | 5                     |
| Italien-Chinois   | 81                   | 1195’                   | 2                     |
| Italien-Français  | 32                   | 359’                    | 5                     |
| Total             | 591                  | 7179’                   | 19                    |

Tout en nous plaçant volontairement en deçà du vif débat idéologique et terminologique qui oppose l’*interprétation* à la *médiation* (pour un approfondissement, voir Mack et Russo 2005; Baraldi et Gavioli 2012; Falbo 2013; Vecchiato 2015), nous nous devons de souligner que l’appellation du centre AIM est en phase avec le choix des institutions italiennes de recourir à des *médiateurs interculturels* bilingues qualifiés par leur expérience de travail au sein des institutions, plutôt qu’à des interprètes diplômés d’une université et/ou certifiés par un organisme accrédité. En général, ces

médiateurs médient non pas les différends, mais les différences dans la conversation (Guillaume-Hofnung 2015 : 68) et puisque ces différences sont avant tout d'ordre linguistique, leur activité de médiation interculturelle se mélange, parfois se limite même, à une activité de traduction. Il en découle que la distinction entre *interprète* et *médiateur* n'est pas toujours facile à cerner, ce qui, d'une part, contribue au flou d'une profession qui peine à s'affirmer, mais de l'autre, attise l'intérêt d'un public beaucoup plus vaste que celui de l'Italie ou de l'Espagne – où les médiateurs sont effectivement employés – pour le corpus AIM. Par ailleurs, si on y regarde de près, il arrive que les deux mandats coexistent chez la même personne et donc la question est peut-être moins de savoir si on est médiateur ou interprète que d'avoir conscience de quand on passe de l'un à l'autre (Niemants 2017). En partant de l'idée que le rôle d'interprète de dialogue n'existe pas en dehors des activités qui sont construites dans l'interaction, il peut être éclairant d'observer ce que font ces intermédiaires linguistiques dans l'un des milieux où ils – et le plus souvent elles – exercent habituellement leur profession, à savoir les institutions socio-sanitaires répondant aux demandes d'assistance et de soins des migrant(e)s.

## 2.2. Des enregistrements au(x) corpus

Comme le rappellent Cencini et Aston (2005) dans l'exergue, la parole vive meurt alors qu'elle est prononcée et l'interprétation ne fait certes pas exception. Pour étudier l'oralité, il faut donc l'enregistrer et la transcrire, passant ainsi des données authentiques aux corpus (Cook 1995 ; Boulton et Tyne 2014). Ce passage retient, depuis des années, toute notre attention, comme cela se dégage des quelques publications qui ont paru (Niemants 2012, 2015b ; Niemants et Pallotti 2017 : 103-106), ainsi que des présentations qui en sont restées à la forme orale (Gavioli, Baraldi, *et al.* 2015<sup>8</sup>, 2016<sup>9</sup>). Sans trop rentrer dans le détail d'un processus quelque peu complexe, nous présenterons brièvement les trois questions principales qu'il soulève.

La première question est celle de la sélection, car tout comme un cartographe, le transcrip-teur doit distinguer ce qui est à retenir de ce qui est à exclure de la transcription (Cook 1995 : 45). Le processus de sélection à partir d'un enregistrement audio touche au moins à cinq éléments : les participants, la structure de la conversation, les traits linguistiques et paralinguistiques, les traits prosodiques et les silences qui (ne) peuvent (pas) être attribués. Pour chacun d'entre eux, le transcrip-teur se doit de décider si écrire, ou pas, ce qu'il/elle entend. Même si « toute fidélité à l'oral n'est qu'illusion » (Galazzi 2002 : 142), il convient de s'atteler à cette « tâche paradoxale » (Falbo 2005), car non seulement elle permet de produire des observables qui peuvent ensuite être partagés (Ayaß 2015), mais elle aide à familiariser avec les données, en stimulant des réflexions d'ordre méthodologique et théorique qui sont à la base de leur interprétation (Gavioli et Mansfield 1990). Le fait de participer au terrain et de travailler sur les données primaires (la transcription de l'audio ou de la vidéo étant toujours une donnée secondaire) contribuerait d'ailleurs, selon Traverso (2003b : 18), à un attachement au corpus « qui évite l'ennui de l'analyste » et qui s'apparente au « caractère émotionnel de la relation que l'ethnographe au groupe qu'il étudie ».

La deuxième question est celle de la représentation, car une fois décidé ce que l'on veut préserver de l'enregistrement, il reste encore à établir comment le faire. La réponse à cette question a tout d'abord affaire à la disposition des données sur l'écran,

car on peut transcrire verticalement comme dans un scénario, horizontalement comme dans une partition de musique, ou encore en attribuant une colonne à chaque interlocuteur (Edwards et Lampert 1993: 11). Chaque disposition présente des avantages et des inconvénients, sur lesquels nous ne pouvons pas nous attarder ici (voir éventuellement Parisse et Morgenstern 2010). Bornons-nous à constater qu'alors que le format vertical a probablement été, jusqu'à présent, le plus utilisé dans la transcription orthographique (c'est-à-dire non phonétique) de données audios, et ce, spécialement en analyse conversationnelle (Sacks, Schegloff, *et al.* 1974; Traverso 1999), le système à partition paraît aujourd'hui le mieux équipé pour transcrire les données vidéo, car plusieurs lignes peuvent être utilisées pour représenter les multiples dimensions de la communication, en rendant compte de leur simultanéité (Niemants et Pallotti 2017).

La troisième et dernière question dont il faut tenir compte pendant la phase de préparation du corpus est celle de l'exploitation et de la quantification des transcriptions par ordinateur.

Ce saut qualitatif est aujourd'hui indispensable dans la perspective de la linguistique de corpus : il est inutile de stocker de grandes masses de données si celles-ci ne sont pas exploitables par des moteurs de recherche. (Groupe ICOR 2010: 26)

Pour outiller un corpus et le rendre exploitable, il faut tout d'abord que la transcription soit non seulement écrite à l'ordinateur mais également lisible par la machine grâce à des balises (Thompson 2005). À la suite de ce balisage, le transcrip-teur peut rajouter des informations supplémentaires, que l'on appelle *annotations*, et à la fois les balises et les annotations peuvent être ensuite comptées par l'analyste, ce qui permet de saisir la fréquence de certains traits linguistiques et interactionnels (Groupe ICOR 2007, 2010).

Gavioli, Baraldi, *et al.* (2015, 2016) ont affirmé, à cet égard, que l'annotation des corpus d'ID devrait tout d'abord viser à extraire des données intéressantes pour l'analyse, par exemple des types de rencontres ou des séquences interactionnelles, et seulement dans un second temps à compter ce qui a été extrait, comme les éléments lexicaux ou les séquences triadiques et dyadiques. Il ne faut d'ailleurs pas oublier, comme le conclut Schegloff (1993: 114), que la quantification ne remplace pas l'analyse et que certaines particules discursives ne sont pas compréhensibles en relation à leur nombre par minutes de conversation, mais seulement en relation à leur contexte interactionnel (1993: 106). Pour ne citer qu'un exemple, le fait de compter la fréquence de *|okay|*<sup>10</sup> dans un corpus, comme ce sera le cas ici, ne nous assure pas qu'on est en train de compter le même objet: non différemment des autres « petits mots » de l'interaction qui ont été étudiés par le Groupe ICOR (2007, 2008, 2010), ce marqueur exerce en fait de multiples fonctions dans l'interaction et, tout comme sa transcription, sa quantification nécessite également un « intelligent agent » (O'Connell et Kowal 1999: 104) qui interprète les résultats, en faisant tout d'abord la distinction entre le *|okay|* de continuation (signifiant à l'interlocuteur qu'il peut continuer) et celui de transition (signifiant que son tour est traité comme complet), pour descendre ensuite dans le détail de leurs usages<sup>11</sup>. En effet, pour des phénomènes aussi complexes que les particules à l'oral, qui comme le rappelle le Groupe ICOR (2010: 30) font intervenir des « faisceaux de propriétés » (entre autres leur nombre, leur place dans le tour ou la longueur de ce dernier, ainsi que leur relation avec le contexte institutionnel de

leur emploi), les « moteurs de recherche constituent une aide indispensable, mais ne sauraient remplacer le travail du chercheur dans l'évaluation des résultats et dans l'analyse fine des extraits obtenus en réponse » (Groupe ICOR 2010: 30).

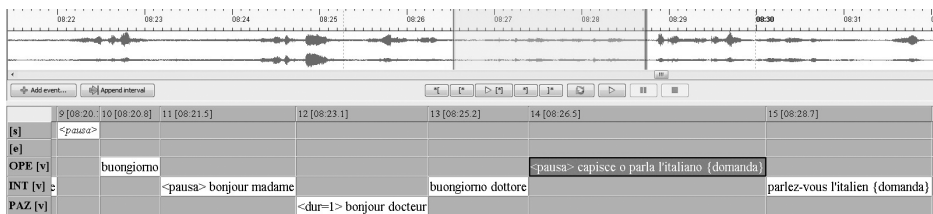
Les considérations de Gavioli, Baraldi, *et al.* (2015, 2016) se dégagent de l'utilisation du corpus AIM qui, tout en étant très éloigné du niveau d'annotation ajouté au ComInDat, témoigne bien des efforts nécessaires pour transformer une collection de transcriptions de l'oral dans un corpus où les données audio ou vidéo sont intégrées et donc analysables. Avant de passer à l'analyse, nous présenterons les deux logiciels qui ont justement permis ce saut qualitatif.

### 2.2.1. Le sous-corpus EXMARaLDA

L'idée de l'utilisation d'EXMARaLDA<sup>12</sup> a germé en 2011 au cours de nos recherches doctorales, lorsque nous avons comparé plusieurs outils de transcription et conclu qu'il représentait une bonne solution de compromis, dans la mesure où il fonctionne comme dans une partition de musique, où à chaque interlocuteur correspond une ligne et où d'autres lignes peuvent être ajoutées pour plus d'informations et/ou cachées lorsque celles-ci n'intéressent pas (Niemants 2013).

FIGURE 1

#### La partition dans EXMARaLDA



EXMARaLDA permet un accès immédiat de la transcription à l'audio/vidéo si les liens ont été créés et il dispose également d'un outil d'interrogation intégré, EXAKT, qui permet à tout instant de calculer une concordance et de vérifier la fréquence d'un fait linguistique. Le dernier avantage, et non des moindres, est la possibilité d'exporter les transcriptions en plusieurs formats et de créer à la fois une collection de transcriptions de données audio ou vidéo (« spoken corpus ») et un corpus multimodal où les transcriptions et les fichiers audio/vidéo font partie du corpus lui-même (« speech corpus »).

Ce que nous appelons sous-corpus EX (pour EXMARaLDA) contient, à présent, 19 enregistrements italien-français collectés en 2010 dans des centres de santé en Italie et en Belgique (pour un total de 284 minutes), dont 5 où les médiatrices sont seules avec le patient – enregistrements qui ont été transcrits par une étudiante dans le cadre de son mémoire (Pederzoli 2011) –, alors que le reste a été transcrit par nos soins, en misant sur la collaboration avec Bernd Meyer et Thomas Schmidt pour que l'exportation en XML réponde au désir initial d'interroger le corpus avec XAIRA<sup>13</sup>. Malgré toutes les implémentations techniques, pour lesquelles nous tenons encore une fois à remercier Thomas Schmidt, nous avons fini par abandonner l'idée du corpus en XML et par générer et interroger le corpus en utilisant EXAKT. Mais cette décision n'enlève rien à la valeur d'un logiciel conçu et constamment développé par



dit, il faut au minimum décider les noms des interlocuteurs (acteurs), la typologie de ce qui est contenu dans un acteur (type linguistique), le lien qui existe entre un sous-acteur et son acteur principal (stéréotype), ainsi que l'ensemble de codes partagés par l'équipe de transpositeurs (vocabulaire contrôlé), si l'on veut être à même de rechercher les tours de parole d'un ou plusieurs interlocuteurs, de faire des concordances, ou d'observer des phénomènes qui se produisent simultanément, pour ne faire que quelques exemples de requêtes possibles.

TABEAU 2  
La structure des données

| Acteur                                                                           | Type linguistique | Stéréotype           | Vocabulaire contrôlé |
|----------------------------------------------------------------------------------|-------------------|----------------------|----------------------|
| OBSf (Obstetrician)<br>GYNf (Gynechologist)<br>MEDf (Mediator)<br>PATf (Patient) | (default)         |                      |                      |
| OBSf-W<br>MEDf-W<br>PATf-W                                                       | Words             | Time subdivision     |                      |
| OBSf-T<br>MEDf-T<br>PATf-T                                                       | Translation       | Symbolic association |                      |
| OBSf-L<br>MEDf-L<br>PATf-L                                                       | Language spoken   | Symbolic subdivision | Language spoken      |
| Event                                                                            | Audible events    |                      | Audible events       |

Dans ce tableau, la structure des données est en anglais, langue partagée par tous les partenaires internationaux du projet FAR pour lequel elle a été initialement conçue. Pour ce qui est des acteurs, les trois lettres suivies de -f ou -m correspondent au nom et au sexe de la ligne dans la notation (GYNf, une gynécologue), alors que le nom entre parenthèses correspond au type de participant (gynécologue, sans distinction entre hommes et femmes). Cette duplicité permet ainsi de rechercher (dans) ce qui est dit par les gynécologues sans distinction de sexe, ou bien de limiter la recherche aux hommes et/ou aux femmes afin de faire des études de genre. En ce qui concerne le type linguistique, dans notre cas cela peut être la transcription de ce qui est dit (default), les mots individuels (Words), la traduction (Translation), la langue parlée (Language spoken), ou des événements audibles qui ne sont pas attribuables à un participant (Audible events). Dans la colonne « Stéréotype », qui décrit le lien qui existe entre un sous-acteur (*enfant*) et son acteur principal (*parent*), la Time subdivision est une division dont les éléments sont situés temporellement, ce qui signifie qu'il est éventuellement possible d'associer chaque mot (Word) d'une annotation au son qui lui correspond. La Symbolic association est par contre une division en éléments sans indications temporelles : la traduction (Translation) en anglais de ce qui est dit en italien ou en français, que nous avons faite pour les partenaires internationaux du projet, est donc uniquement liée à l'original qui lui correspond par une relation 1 à 1, et elle n'est pas directement associée à la bande son. Enfin, la Symbolic subdivision, que nous avons choisie pour coder la ou les langues parlées dans l'annotation, est l'association d'une valeur en conservant la même temporalité que l'acteur principal et en permettant de faire des subdivisions ultérieures (cela nous permet par

exemple de diviser les annotations à présent codées comme mixed en distinguant entre ce qui est dit en français et en italien). Enfin, le vocabulaire partagé par l'équipe de transcripteurs permet de contrôler les valeurs des éléments d'une catégorie de codage et donc d'être sûrs que tout le monde codera la langue parlée et les événements audibles de la même façon (par exemple, Italian lorsque l'acteur parle en Italien, ou typing lorsque quelqu'un tape à l'ordinateur).

Nous nous devons de souligner que le corpus analysé dans cet article est toujours en augmentation, car outre les 40 interactions du FAR, qui ont été directement transcrites avec *ELAN*, nous avons également: (a) les 19 interactions du sous-corpus EX, qui ont été exportées dans le format .eaf pour les rendre interrogeables par la recherche structurée d'*ELAN*, qui comme on le verra est un outil beaucoup plus puissant que le simple concordancier *EXAKT*; (b) les interactions de deux autres sous-corpus AIM, initialement transcrites en *MS Word* par deux étudiantes et récemment retranscrites dans *ELAN* pour aligner son et parole, à savoir le sous-corpus Luppi (qui contient 10 interactions italien-anglais pour un total de 143 minutes) et le sous-corpus Fatima (qui contient 21 interactions italien-arabe pour un total de 105 minutes); (c) une trentaine d'interactions que d'autres étudiantes de Modène sont actuellement en train de transcrire dans le cadre de leur mémoire. Il en découle que l'extraction des données d'interprétation en milieu médical qui suit est faite à partir d'un corpus de 65 interactions issues de différents sous-corpus (en sont exclues les 20 interactions sans médiateur du FAR, les 5 interactions transcrites par Pederzoli, ainsi que les interactions qui à l'époque de notre communication au colloque SOFT étaient encore en cours de transcription), pour un total de 1124 minutes (presque 19 heures) d'enregistrements qui ont été transcrits par différents moyens et logiciels, mais qui sont ici tous analysés avec les outils de requête d'*ELAN*. Nous l'avons choisi car il permet de faire aisément des recherches avec des limites temporelles, en misant sur le lien qui se crée entre les annotations et la bande-son, et donc d'extraire des clusters de formes linguistiques (telles que |okay|) et de phénomènes interactionnels (tels que les silences, les chevauchements, ou l'alternance des tours) qui apparaissent de manière concomitante, un peu comme cela est possible à partir de la banque de données CLAPI du Groupe ICOR (2007, 2008, 2010), qui bien avant le Centre AIM s'est posé la question de comment passer d'un ensemble de corpus archivés à un ensemble de corpus outillés afin d'analyser les particules discursives en interaction.

### 3. L'utilisation du ou des corpus: l'analyse de |okay| dans les tours de parole des médiatrices

Lorsque nous avons commencé à utiliser *ELAN*, nous pensions que la façon la meilleure d'extraire des données intéressantes d'un point de vue interactionnel était de prévoir un acteur supplémentaire et d'y insérer des annotations correspondant aux phénomènes étudiés. La distinction entre séquences dyadiques (impliquant le médiateur et l'un des deux autres participants) et triadiques (impliquant tous les participants) étant pertinente dans les études sur l'ID, on a donc songé à ajouter une annotation pour chaque séquence et à les extraire (et éventuellement compter) par la suite.

Nous avons toutefois réalisé que cette opération, outre qu'elle demandait beaucoup d'énergie et de temps, nous amenait *de facto* à analyser les données: une fois

extraites, elles n'auraient pas été de nouvelles données à analyser, mais une sorte de mémoire de classifications précédentes pouvant être confirmées ou mieux définies. Nous avons donc changé de perspective et au lieu de nous demander ce que nous voulions faire avec *ELAN*, nous nous sommes demandée ce qu'*ELAN* pouvait bien nous offrir en termes de visualisation et d'annotation des interactions en milieu socio-médical, où la temporalité et la co-construction des tours de parole peuvent avoir des retombées sur la relation de confiance entre prestataires et bénéficiaires des soins<sup>16</sup>.

Pour ce qui est de la visualisation, nous avons tiré profit des multiples exportations offertes par ce logiciel et cherché à observer les mêmes données à la fois horizontalement (c'est-à-dire en partition de musique) et verticalement (c'est-à-dire dans un format tour-par-tour), ce qui nous a donné des idées pour la segmentation et l'extraction de certaines séquences. En ce qui concerne l'annotation, nous avons décidé de nous limiter au minimum et d'utiliser la structure présentée dans le tableau 2 pour extraire des tours et des séquences potentiellement intéressants.

Considérant l'utilité de vérifier les analyses effectuées manuellement sur des corpus restreints, en étudiant les mêmes particules discursives dans d'autres situations et/ou corpus<sup>17</sup>, nous allons revenir sur les résultats de l'étude de Gavioli (2012) sur le rôle des tours régulateurs dans les interactions du corpus AIM tel qu'il était à l'époque, en montrant le potentiel d'*ELAN* dans la recherche des |okay| des médiatrices des 65 interactions avec médiateur maintenant alignées à l'audio. Nous le ferons par stades, en allant de l'extraction des |okay| comme annotation, ce qui dans le vocabulaire d'*ELAN* équivaut à un segment isolé de transcription, pour en venir aux fonctions de |okay| dans l'annotation, c'est-à-dire lorsque |okay| est l'un des éléments du segment transcrit, et finir par l'observation de |okay| dans deux types de séquences interactionnelles qui s'avèrent importantes pour l'ID, à savoir les séquences triadiques et dyadiques impliquant le médiateur. Mais avant de ce faire, quelques repères théoriques s'imposent afin d'insérer les considérations qui vont suivre dans le cadre des études sur les tours régulateurs dans la conversation et dans l'interprétation, ainsi que sur le rôle d'un marqueur compris, voire produit, par des locuteurs de langues maternelles différentes.

Les études linguistiques de l'oral ont établi depuis longtemps l'importance et les multiples fonctions de « tokens » tels que |mm hm|, |okay|, *ben*, *voilà*, *donc*, *alors*, etc. dans la construction de l'interaction, avec une contribution importante de la part de l'analyse conversationnelle (Gardner 2001 : 4). Il s'agit généralement de productions brèves, émises par le ou les récepteurs en réponse au tour du locuteur, qui mettent en évidence la nature co-construite de la conversation, ainsi que le travail de coordination que les participants mettent en place pour signifier qu'ils suivent et comprennent, ou pour ajuster leur discours. Parmi leurs fonctions, on distingue par exemple celles d'accusé de réception et de continueurs, qui comme le rappelle Traverso (2016 : 41) se distinguent de deux autres productions fréquentes, à savoir les marques d'accord et les évaluations (voir également la classification plus précise de Gardner 2001 : 2). Or, les recherches sur les fonctions pragmatiques et interactionnelles de ces « petits mots » ont montré que leur rôle est loin d'être marginal dans la conversation ordinaire (Maynard 1986 : 1080) et qu'ils contribuent activement à la négociation des actions pertinentes à accomplir dans des interactions institutionnelles telles que celles en milieu sanitaire (Czyzewski 1995). Moins clair est leur rôle

dans l'ID, où tout en n'étant pas les principaux destinataires de ce que disent les interlocuteurs primaires, les interprètes bilingues sont, dans la plupart des cas, les premiers à recevoir leurs énoncés. Leur activité d'auditeurs – dont le but n'est certes pas le même que celui des interlocuteurs primaires puisqu'ils ou elles écoutent pour traduire (Englund Dimitrova 1997: 160) – peut être observée dans l'utilisation de tours régulateurs tels que *[okay]*, qui contribuent de deux manières à la traduction. Comme Gavioli (2012) l'a bien montré dans son étude, d'une part ils servent à marquer le passage entre l'activité d'écoute et l'activité de traduction, c'est-à-dire qu'en prononçant un *[okay]*, l'interprète (a) manifeste sa propre compréhension de ce qu'un premier locuteur a dit et (b) projette le début de la traduction pour l'autre, contribuant ainsi à définir l'unité de traduction et à coordonner l'interaction. D'autre part, les tours régulateurs fonctionnent comme des invitations à continuer l'activité en cours afin d'atteindre un objectif interactionnel, par exemple celui de faire en sorte que l'énoncé d'un participant soit traduisible. En réfléchissant sur le caractère international d'un marqueur comme *[okay]*, Gavioli (2012: 224) suppose que certains tours régulateurs peuvent servir non seulement à promouvoir le récit du premier locuteur et à construire, ainsi, la « traduisibilité » de ce qu'il est en train de dire, mais également à signaler aux autres participants qu'ils peuvent commencer à parler, en fonctionnant comme une sorte de « feu vert » que nous allons retrouver dans les extraits ci-dessous.

### 3.1 *[Okay]* comme annotation

L'outil de recherche structurée d'ELAN offre trois possibilités de recherche, à savoir : « Substring Search », qui permet de chercher une chaîne de caractères dans tout le corpus ; « Single Layer Search », qui permet de chercher une chaîne de caractères sur un acteur spécifique ; « Multiple Layer Search », qui permet de faire des requêtes plus complexes que nous détaillerons dans §3.3. Pour extraire des exemples de *[okay]* isolés de la médiatrice, il faut donc opter pour la « Single Layer Search » et sélectionner l'acteur MEDf, puis cocher la case « exact match », c'est-à-dire tous les cas où une annotation de la médiatrice ne contient que le mot *[okay]*.

FIGURE 3

« Single Layer Search » de *[okay]*



En double-cliquant sur l’une des concordances, on peut observer l’occurrence dans la fenêtre d’*ELAN* que nous avons montrée auparavant (figure 2), et exporter le [okay] dans le format « Traditional Transcript Text », de façon à bien voir la séquence de tours dans laquelle il s’insère.

L’exemple 1 (tiré du sous-corpus directement transcrit avec *ELAN*) montre la première des deux façons dont les [okay] des médiatrices peuvent, selon Gavioli (2012), contribuer à la traduction, en marquant le passage entre l’écoute et la restitution de ce qui a été écouté. Par son [okay], MEDf réagit en effet au tour de la patiente (PATf) et projette le début de la traduction pour la sage-femme (OBSf), qui par la répétition d’[okay] dans son tour montre qu’elle est prête à l’écouter (son attention étant généralement partagée entre ce qui se déroule dans le cabinet et ce qu’elle en écrit à l’ordinateur).

| 1)   | EXTRAIT                                | GLOSE                                     |
|------|----------------------------------------|-------------------------------------------|
| PATf | <u>I don't like too much pepper</u>    | ---                                       |
| MEDf | [okay]                                 | ---                                       |
| OBSf | okay                                   | okay                                      |
| MEDf | eh mangia però non: (.) [[non troppo]] | eh elle mange, mais pas: (.) [[pas trop]] |
| OBSf | [[non troppo]]                         | [[pas trop]]                              |

(EL: CSFS 4)<sup>18</sup>

Lorsque le [okay] réagit par contre à un tour de parole d’un prestataire (dans l’exemple 2, tiré du corpus initialement transcrit avec *EXMARaLDA*, PREf est une femme au guichet), il est souvent suivi d’un autre tour plein du même prestataire, qui ajoute des détails supplémentaires.

|    |      |                                                                             |
|----|------|-----------------------------------------------------------------------------|
| 2) | PREF | sinon je contacte directement le service voir ce qu’il peut faire pour nous |
|    | MEDf | okay                                                                        |
|    | PREF | c’est un peu du cas par cas ça dépend aussi ce qu’a monsieur                |

(EX: 100531\_000)

Comme cela a été mis en évidence par d’autres recherches sur le corpus AIM, le [okay] de la médiatrice peut s’insérer dans un long tour de parole du prestataire, souvent introduit par « alors on lui explique/dit que » (Gavioli 2015), où il ou elle donne un certain nombre d’informations, en déléguant implicitement à la médiatrice le soin de les traduire de la façon qu’elle considère la plus adaptée. Le tour du prestataire étant composé de deux ou plusieurs unités, il est généralement ponctué de continueurs au travers desquels la médiatrice affiche sa compréhension d’une partie de ce long tour, et *okay* s’avère l’un des plus fréquents.

Le troisième et dernier exemple de [okay] comme tour isolé nous amène vers le deuxième niveau d’analyse, dans la mesure où il présente un premier **okay** à valeur de tour qui accuse réception d’une partie de l’explication de la sage-femme et un autre **okay** où la médiatrice fait écho à la sage-femme en reprenant, en disant<sup>19</sup> en même temps qu’elle, l’expression **con calma** [avec calme].

| 3)   | EXTRAIT                                                                                              | GLOSE                                                                                                      |
|------|------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------|
| OBSf | <b>le visite con ecografia non le fanno le ostetriche ma le fanno i ginecologi</b>                   | les visites avec échographie ne sont pas faites par les sages-femmes, mais par les gynécologues            |
| MEDf | <b>okay</b>                                                                                          | okay                                                                                                       |
| OBSf | <b>okay? perciò quello che io posso fare è cercare un posto per fare una visita [ginecologica]</b>   | okay? donc ce que je peux faire c'est chercher une place pour faire une visite [gynécologique]             |
| MEDf | <b>[ginecologica]</b>                                                                                | [gynécologique]                                                                                            |
| OBSf | <b>con ecografia (.) non è necessario farla (0,6) nel giro di un mese o due mesi [con ca:lma la]</b> | avec échographie (.) il n'est pas nécessaire de la faire (0,6) d'ici un mois ou deux mois [avec ca:lme la] |
| MEDf | <b>[okay con calma]</b>                                                                              | [okay avec calme]                                                                                          |

(EL: 111217\_001)

La reprise, de pair avec |okay|, de mots qui viennent d'être prononcés ou qui peuvent l'être est assez récurrente dans notre corpus et c'est un premier exemple de |okay| dans l'annotation. Nous en avons repéré d'autres en utilisant la recherche « Simple Layer », à une différence près.

3.2. |Okay| dans l'annotation

Alors que pour extraire les |okay| isolés nous avons sélectionné « exact match » du menu déroulant d'ELAN, pour chercher les |okay| qui figurent dans une annotation de MEDf avec d'autres éléments transcrits il faut sélectionner « substring match »<sup>20</sup>.

On repère tout d'abord des cas qui se trouvent à mi-chemin entre les deux premiers niveaux d'analyse, dans la mesure où |okay| n'est pas tout à fait seul, mais précédé ou suivi d'un élément qui n'en change pas la fonction d'accusé de réception.

| 4)   | EXTRAIT                                                                | GLOSE                                                                      |
|------|------------------------------------------------------------------------|----------------------------------------------------------------------------|
| PATm | il y a toujours quelque chose dans le sang pour que ehm ehm            | ---                                                                        |
| MEDf | ah okay                                                                | ---                                                                        |
| PATm | mais c'était en baisse (.) sans sans médicaments                       | ---                                                                        |
| MEDf | ah okay                                                                | ---                                                                        |
| PATm | sans médicaments ça ça baisse                                          | ---                                                                        |
| MEDf | <b>dice che ehm la diagnosi l'aveva fatta il suo medico curante...</b> | il dit que ehm le diagnostic avait été fait par son médecin généraliste... |

(EX: 02)

MEDf, qui dans ce cas-ci est une interprète diplômée, dit en effet **ah okay**, tout comme d'autres médiatrices disent ailleurs |hm okay|, *oui okay*, **okay sì**, en s'adressant tant aux bénéficiaires qu'aux prestataires des services de santé. Au-delà de ce cas limite, dans notre corpus, |okay| est le plus souvent en début de tour et nous en montrerons deux cas emblématiques.

Dans l'exemple 5, la sage-femme est en train d'expliquer les rendez-vous successifs, dont une prise de sang qui doit être faite le 11, à dix heures, à jeun. Le |okay| que la médiatrice prononce, après avoir vérifié sa compréhension, inaugure sa restitution

de ce que la sage-femme a dit, et a donc, encore une fois, une valeur de transition à la traduction.

| 5)    | EXTRAIT                                                          | GLOSE                               |
|-------|------------------------------------------------------------------|-------------------------------------|
| OBSf  | <b>u: no</b> [(??3syll) ore] <b>dieci</b>                        | un [(??3syll) heures] dix           |
| MEDf  | <b>[che fa la cartella giusto?]</b>                              | [qu'elle fait le dossier c'est ça?] |
| (1,0) |                                                                  |                                     |
| MEDf  | okay  <u>this one (.) is on</u><br><u>[the: e- eleven yes]</u>   | ---                                 |
| PATf  | <u>[on the eleven without eating]</u>                            | ---                                 |
| MEDf  | <u>yes (.) to do pregnancy file</u><br><u>[with the midwife]</u> | ---                                 |
| PATf  | <u>[file hm]</u> <u>okay</u>                                     | ---                                 |
| MEDf  | <u>okay?</u>                                                     | ---                                 |

(EL: B005 14)

Après avoir co-construit la traduction avec la patiente, la médiatrice prononce un okay avec intonation montante qui lui permet de s'assurer de la compréhension de son interlocutrice. Il s'agit là d'une autre fonction assez fréquente de |okay| comme tour individuel dans le corpus AIM, que nous apprécions davantage en regardant sa place dans la séquence dans laquelle il est produit.

Dans l'exemple 6, le |okay| suit par contre une longue séquence dyadique où la médiatrice et le patient évaluent combien ce dernier gagne lorsqu'il travaille comme maçon, car cela est important pour vérifier s'il est en mesure de payer ses frais de santé.

| 6)    | EXTRAIT                                                                                                                             | GLOSE                                             |
|-------|-------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------|
| MEDf  | <b>e quando</b>                                                                                                                     | et quand                                          |
| PATm  | <b>ehm non</b>                                                                                                                      | ehm je ne                                         |
| MEDf  | <b>lavorate quanto guadagnate più o meno</b>                                                                                        | vous travaillez combien vous gagnez plus ou moins |
| (.)   |                                                                                                                                     |                                                   |
| PATm  | <b>ah</b>                                                                                                                           | ah                                                |
| MEDf  | <b>al giorno</b>                                                                                                                    | par jour                                          |
| PATm  | <b>al giorno (.) ehm trenta (.) trenta ehm</b>                                                                                      | par jour (.) ehm trente (.) trente ehm            |
| MEDf  | <b>trenta euro?</b>                                                                                                                 | trente euros?                                     |
| PATm  | <b>euro</b>                                                                                                                         | euros                                             |
| (1,0) |                                                                                                                                     |                                                   |
| MEDf  | <b>°(okay)° ma il</b>                                                                                                               | °(okay)°, mais le                                 |
| PATm  | <b>trenta cinque</b>                                                                                                                | trente-cinq                                       |
| (0,5) |                                                                                                                                     |                                                   |
| MEDf  | okay  (.) donc quand il fait le maçon ça dépend hein ça dépend de ce qu'il (y) a (.) ben il gagne trente trente-cinq euros par jour | ---                                               |

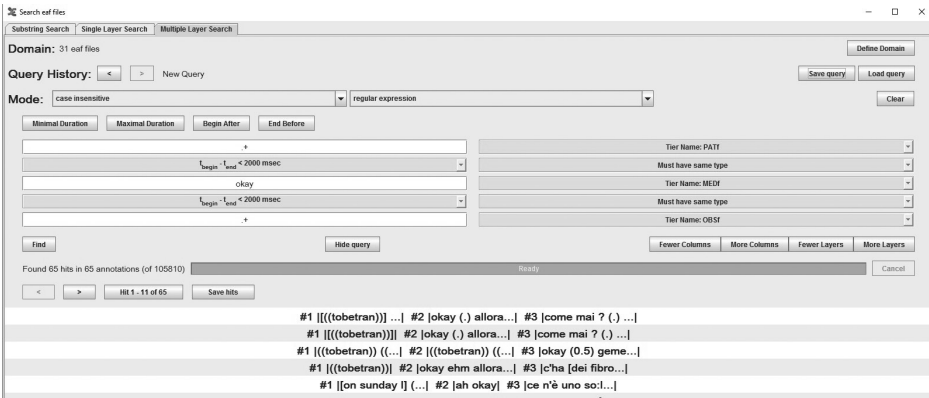
(EX: 100330\_000)

|Okay| se trouve à nouveau en début de tour et a une valeur de transition à la traduction qu’il inaugure. Tout comme d’autres régulateurs à caractère international, il peut avoir la fonction de feu rouge ou vert signalant aux interlocuteurs de s’arrêter ou de continuer à parler (Gavioli 2012: 224). Il joue donc un rôle majeur dans la création d’un espace conversationnel pour la traduction d’informations qui ne sont pas nécessairement contenues dans un seul tour original, mais parfois co-construites dans toute la séquence qui précède, comme cela est justement le cas de l’exemple ci-dessus.

3.3. |Okay| dans la séquence

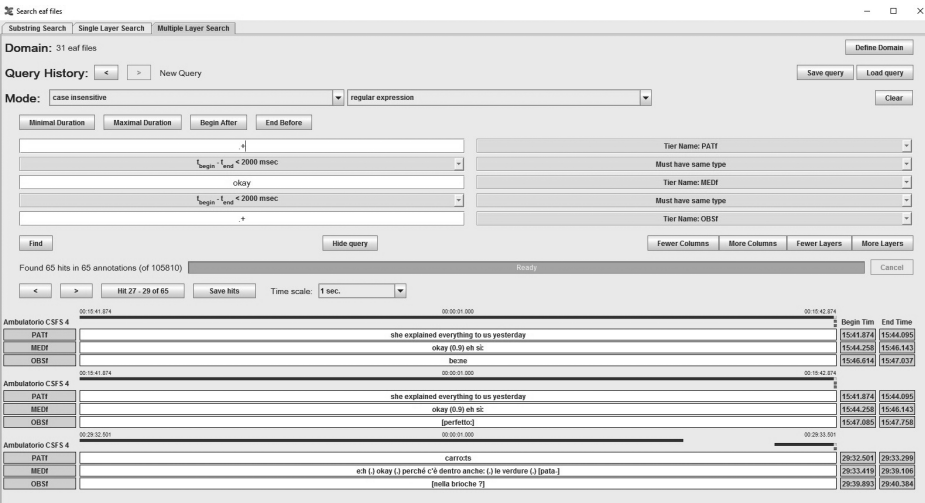
Nous en venons au dernier niveau d’analyse et donc à la recherche multiple structurée, qui permet d’extraire |okay| dans une certaine séquence d’annotations. La requête affichée dans la figure 4 est réalisée en utilisant les expressions régulières et vise à extraire une séquence triadique classique du type PATf-MEDf-OBSf, où la médiatrice traduit en italien pour la sage-femme tout de suite après le tour de parole de la patiente. Le cas échéant, l’écart temporel entre les tours est inférieur à deux secondes<sup>21</sup> et la médiatrice doit prononcer un |okay|, alors que les deux autres participantes peuvent dire n’importe quoi.

FIGURE 4  
« Multiple Layer Search » de |okay| : triadique 1



Nous nous devons de rappeler que les résultats peuvent être affichés de plusieurs manières, dont celle qui suit, permettant de voir les annotations en entier et d’apprécier la progression temporelle de la séquence, à droite.

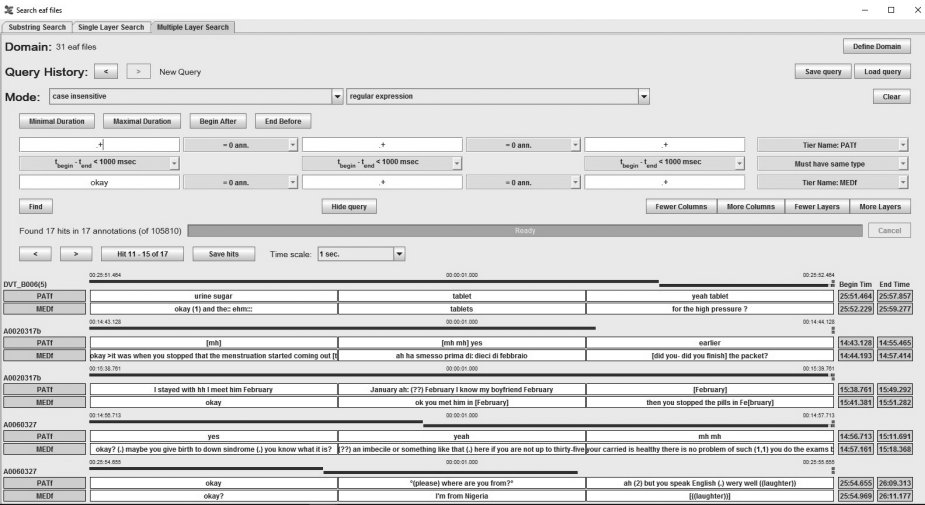
FIGURE 5  
« Multiple Layer Search » de |okay| : triadique 2



Force a été de constater que la séquence triadique PATf-MEDf-OBSf était moins fréquente qu'attendu. Comme l'avait déjà bien montré Davidson (2000) et comme nous l'avons vu dans les trois derniers exemples ci-dessus, il peut en effet s'avérer nécessaire de vérifier la compréhension avant de passer à la traduction, ce qui est fait dans des séquences dyadiques plus ou moins étendues entre la médiatrice et l'un des deux autres participants.

Les séquences dyadiques peuvent à leur tour faire l'objet d'une requête complexe, telle que celle de la figure 6, qui vise à extraire des séquences dyadiques à (au moins) six tours du type PATf-MEDf-PATf-MEDf-PATf-MEDf, où le premier tour de la médiatrice doit nécessairement contenir un |okay|.

FIGURE 6  
« Multiple Layer Search » de |okay| : dyadique



Par la même requête nous avons également repéré, en changeant d'acteur, des dyadiques entre les professionnels de santé et les médiatrices, mais nous nous bornerons à la présentation d'un seul extrait représentatif de dyadique PATf-MEDf, ainsi que d'un seul extrait de triadique PATf-MEDf-OBSf affichant des [okay] qui, si on y regarde de près, ont des fonctions différentes et sont donc un bon condensé de ce que nous avons analysé jusqu'ici.

Dans l'exemple 7, [okay] exerce la fonction d'accusé de réception servant à montrer que la médiatrice a bien compris que le père de la patiente avait du sucre dans ses urines et que la conversation peut donc se poursuivre sur ce terrain partagé. Il n'en est pas ainsi pour les comprimés (*tablets*), un terme qui échappe à la médiatrice (qui fait une pause de 1 seconde suivie d'hésitations) et que la patiente lui suggère par son hétéro-réparation auto-initiée, témoignant ainsi de la co-construction de la parole à traduire et d'un phénomène interactionnel (la réparation) qui se place entre le tour et l'échange et qui n'est donc analysable qu'en prenant en compte la séquence et la temporalité des contributions individuelles (Traverso 2016: 101-119).

- 7)
- |       |                                   |
|-------|-----------------------------------|
| PATf  | <u>urine sugar</u>                |
| MEDf  | <u>okay (1,0) and the:: ehm::</u> |
| PATf  | <u>tablet</u>                     |
| MEDf  | <u>tablets</u>                    |
| PATf  | <u>yeah tablet</u>                |
| (0,6) |                                   |
| MEDf  | <u>for the high pressure?</u>     |

(EL: B006(5))

Dans l'exemple 8, le premier [okay] apparaît dans l'expression anglaise *they are okay*, par laquelle la médiatrice traduit l'italien **stanno bene** et se réfère au bien-être du patient. Il s'agit là d'un cas où [okay] joue un tout autre rôle et que nous avons donc manuellement exclu des concordances au moment où on en analysait les fonctions (sur la vérification manuelle des attestations recensées, voir Colón de Carvajal 2010). Le deuxième [okay], qui, en ne lisant que la transcription, pourrait être pris pour une répétition du précédent, nous semble plutôt, après une écoute attentive de l'audio, avoir la double fonction d'accusé de réception de ce qu'a dit la patiente en anglais et de transition à la traduction pour la sage-femme en italien, **okay (.) sì**.

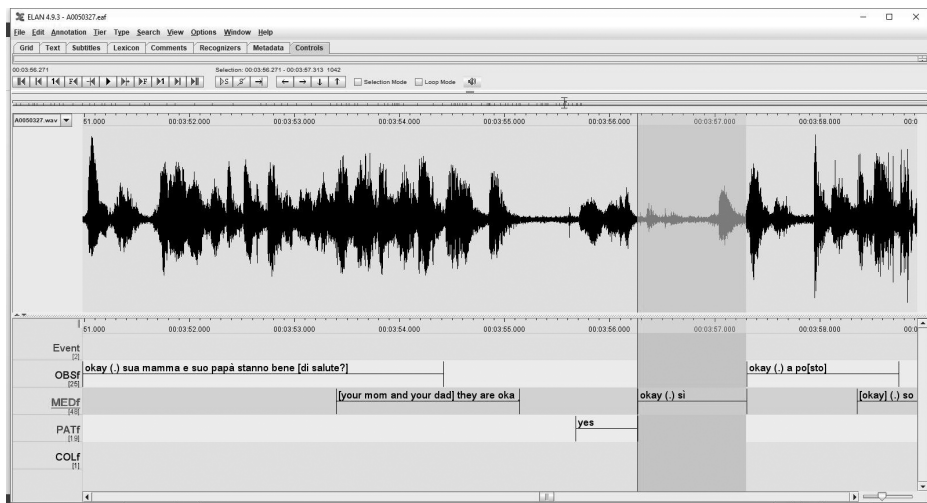
- |    |                                                                    |                                                     |
|----|--------------------------------------------------------------------|-----------------------------------------------------|
| 8) | EXTRAIT                                                            | GLOSE                                               |
|    | OBSf <b>okay (.) sua mamma e suo papà stanno bene [di salute?]</b> | okay (.) sa maman et son papa vont bien [en santé?] |
|    | MEDf [your mom and your dad] they are okay                         | ---                                                 |
|    | (0,5)                                                              |                                                     |
|    | PATf yes                                                           | ---                                                 |
|    | MEDf <b>okay (.) sì</b>                                            | okay (.) oui                                        |
|    | OBSf <b>okay (.) a po [sto]</b>                                    | okay (.) bi[en]                                     |
|    | MEDf [okay] (.) so see you on Wednesday then () okay?              | ---                                                 |

(EL: A0050327)

En effet, ce deuxième [okay] a un contour intonatif assez différent par rapport au premier et, comme le témoigne la forme d'onde ci-dessous, il est prononcé à un volume beaucoup plus bas que ce qui précède et qui suit.

FIGURE 7

### La visualisation de la bande-son



Le troisième et le quatrième [okay] dans cet extrait sont deux exemples assez clairs d'occurrences en début et en fin (du même) tour : alors que le troisième montre l'intersection entre l'activité d'écoute/réception et celle de traduction, le quatrième a une fonction de vérification de la compréhension par son intonation montante. Nous nous devons de remarquer, enfin, que même la sage-femme utilise deux fois ce marqueur en début de tour, en confirmant ainsi tant un caractère international qui peut contribuer à construire un terrain commun entre deux interlocuteurs qui ne partagent pas la même langue, que sa fréquence en position initiale.

## 4. L'utilité du ou des corpus : l'extraction de segments alignés à l'audio

Ayant décrit la création et l'utilisation d'un corpus qui intègre données primaires (enregistrements) et secondaires (transcription), il reste encore à montrer son utilité dans le domaine de l'ID. Avant d'en venir aux conclusions, nous allons donc tirer le bilan de nos transcriptions et de nos analyses, en nous penchant sur l'un des éléments à notre avis les plus utiles tant d'*EXMARALDA* que d'*ELAN*, à savoir la possibilité d'extraire des segments alignés à l'audio pour des fins de recherche ou de formation.

Il ne fait désormais pas de doute que les moyens modernes de transcription représentent une aide considérable au moment de la création d'un corpus, puisqu'ils facilitent la tâche du transcripteur, permettent de créer un lien qui sera utilisé instantanément par la suite et permettent le partage des données primaires et secondaires à la fois (Parisse et Morgenstern 2010). Les deux logiciels que nous avons testés offrent en outre l'avantage de visualiser les données de plusieurs façons – ce qui peut stimuler la réflexion sur les phénomènes à étudier – et ils permettent une multiplicité d'annotations, en partant du principe que la transcription elle-même est

une forme de codage. Ce que nous avons donc cherché à montrer dans les paragraphes qui précèdent, outre que la compatibilité des données qui viennent d'*EXMARaLDA* avec *ELAN*, est comment utiliser la transcription et des codes tels que le nom des participants (MEDf, PATf, OBSf) ou des éléments lexicaux ([okay]) pour extraire des séquences potentiellement intéressantes.

En résumant un peu, nous avons réussi à extraire des séquences avec :

- un *item* comme annotation ;
- un *item* avec une position dans l'annotation (par exemple initiale) ;
- un *item* dans une séquence d'annotations impliquant deux ou trois participants.

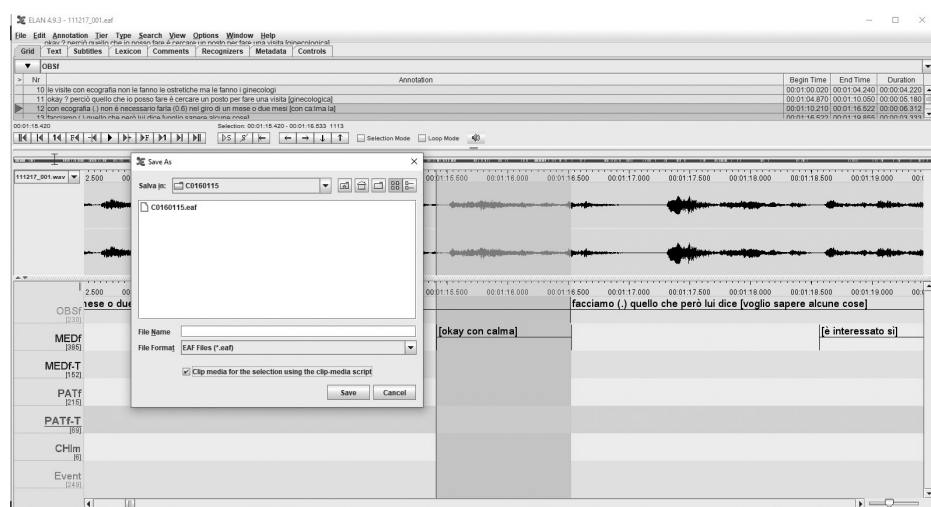
La première utilité du corpus est donc celle d'automatiser la recherche d'un *item*, dans notre cas [okay], dans un ensemble de transcriptions qui étaient auparavant divisées en centaines de documents *MS Word* non alignés, et de simplifier ainsi le repérage et l'analyse des fonctions exercées dans des annotations et des séquences différentes, comme la fonction de vérification de la compréhension lorsque [okay] est dans l'annotation isolée et prononcé avec intonation montante (exemple 5), la fonction de transition à la traduction lorsque [okay] est en position initiale dans l'annotation (exemple 6), ou la fonction d'accusé de réception servant à montrer que la conversation peut se poursuivre sur un terrain partagé lorsque [okay] se trouve dans des séquences dyadiques entre les médiatrices et l'un des deux autres interlocuteurs (exemple 7). Il importe de rappeler que les résultats d'une requête nécessiteront toujours une validation humaine et que malgré la possibilité de faire intervenir le critère de la position et donc de repérer différentes situations syntaxiques, seul l'œil de l'analyste est en mesure de valider la distinction entre they are okay et okay see you on Wednesday du dernier exemple, en décidant de l'exclusion du premier de la liste des exemples pertinents de [okay] dans l'organisation des tours. Bien que les avancées techniques soient quelque part toujours valables, les outils fournissent, comme le dit le groupe ICOR (2007: 7), « des “pistes” pour mettre en évidence des associations qu'il convient de vérifier de manière qualitative pour éviter les erreurs d'interprétation, ils ne sont pas à ce jour capables d'identifier de manière automatique des emplois et sont soumis de ce fait à validation ».

L'autre utilité, et pas des moindres, est celle d'extraire non seulement la transcription, mais également les données primaires auxquelles elle se relie. Pour banal que cela puisse paraître, le lien avec l'audio ou la vidéo joue un rôle majeur, non seulement en phase d'analyse, mais surtout de divulgation des résultats. Cela permet, d'une part, d'apprécier le contour intonatif d'un *item* ([okay] ayant la valeur de vérification de la compréhension a une intonation montante que nous avons signalée par un ? pour en faciliter le repérage automatique, mais qui peut toujours être écouté), ainsi que de saisir sa différente prononciation en fonction du contexte interactionnel et de la langue de l'interlocuteur (dans nos deux sous-corpus alignés à l'audio, [okay] ayant la valeur de transition à la traduction est souvent prononcé /əʊ'keɪ/ lorsqu'il s'adresse à un anglophone, /okk'ei/ lorsqu'il s'adresse à italophone, et /oke/ quand le destinataire est francophone). D'autre part, le fait de pouvoir rapidement extraire une séquence de tours et l'audio ou la vidéo qui lui correspondent accélère la préparation de diapositives ou d'autres supports à utiliser dans le contexte de colloques scientifiques ou de cours de formation. Il est en effet assez fréquent de devoir remonter à un exemple pertinent dont on ne se souvient que partiellement, ou bien de vouloir

associer l'audio à la transcription insérée dans un .ppt. Avec *ELAN*, il suffit de taper un ou deux mots dans la « Substring Search » pour que l'exemple apparaisse, et d'exporter ensuite la séquence sélectionnée (en couleur plus foncée dans la figure 8) comme .eaf, en précisant que nous souhaitons également le média. Nous obtiendrons ainsi deux fichiers, dont l'un est la transcription en partition reliée à l'audio de la séquence concernée (qui peut toujours être exportée en d'autres formats), alors que l'autre est un fichier audio dont le début et la fin coïncident avec la séquence sélectionnée. Aussi peu scientifique que cet argument puisse paraître aux yeux des analystes de corpus écrits, cette utilité n'est pas des moindres pour les chercheurs en interaction et en ID, qui ont jusqu'à présent souvent archivé séparément leurs données primaires et secondaires.

FIGURE 8

L'extraction de segments alignés à l'audio



Le corpus AIM présente, enfin, l'intérêt non négligeable de reposer sur des données authentiques, et donc de produire des observables qui peuvent être utilisés dans la formation supérieure et professionnelle des interprètes/médiateurs, ainsi que de ceux qui utilisent leurs services.

La « Conversation Analytic Role-play Method » (Stokoe 2011), pour ne citer que l'une des méthodes didactiques qui prévoient l'utilisation de transcriptions authentiques alignées à l'audio/vidéo, a été récemment testée dans plusieurs cours d'ID au niveau universitaire (Niemants et Stokoe 2017, Dal Fovo 2018) et peut sans doute bénéficier de corpus déjà alignés tels nos sous-corpus *EX* et *EL*. La présentation simultanée des données primaires et secondaires peut contribuer à « construire un sentiment d'« y être » » (Pallotti 2002 : 3, reprenant Geertz), faisant accéder des étudiants qui n'ont généralement encore jamais exercé la profession aux spécificités de l'interprétation en différents milieux, et en leur permettant d'observer tant la dimension textuelle que temporelle de la langue orale en interaction, ainsi que de produire et de réfléchir à des manières alternatives de réagir aux tours de parole d'interlocuteurs authentiques enregistrés en milieu naturel.

Lorsqu'on obtient le consentement à l'utilisation des données primaires, rien n'empêche de présenter l'audio des interactions entre prestataires et bénéficiaires des services de santé au moment de la restitution de la recherche. En effet, des cours de formation adressés au personnel soignant sont généralement organisés par plusieurs membres du centre AIM à la suite de la collecte des données et ils représentent une étape importante dans la recherche, dans la mesure où les prestataires qui ont consenti à l'enregistrement peuvent tirer profit des résultats des analyses, en prenant conscience de ce qu'ils ou elles disent et surtout de la façon dont ils ou elles le font. Malgré l'anonymisation des données, ces prestataires n'ont généralement pas de difficultés à se reconnaître dans les transcriptions écrites qui sont distribuées au cours. En cas de recours à l'audio ou à la vidéo, il faudrait donc envisager des altérations du signal sonore et visuel, que ni *ELAN* ni *EXMARaLDA* ne sont jusqu'à présent en mesure de produire directement. Il faut toutefois rappeler qu'*ELAN* permet l'importation de fichiers anonymisés par d'autres moyens (comme *Audacity*)<sup>22</sup>, ainsi que l'activation de ces derniers au lieu des fichiers sources. Le problème de l'anonymisation peut donc être assez aisément contourné et le sous-corpus aligné EL préserve, selon nous, toute son utilité.

L'utilité didactique du corpus AIM dans son ensemble ne saurait certes se résumer à la seule disponibilité de données authentiques, d'autant plus que « reality does not travel with the text » (Widdowson 1998 : 711) et que l'utilisation que l'on en fait, qui dépend des besoins et des objectifs de la formation, pourrait très bien être non authentique. Faute d'espace, nous ne pouvons toutefois pas nous pencher davantage sur cette dimension, sur laquelle nous avons plus longuement réfléchi ailleurs (Niemants 2015 ; Niemants et Cirillo 2016 ; Niemants et Stokoe 2017), donc nous nous bornons à reprendre ici quelques éléments qui se sont dégagés de nos analyses et qui pourraient selon nous trouver leur place dans la formation des interprètes de dialogue.

Nos exemples ont montré que les informations à traduire ne sont pas nécessairement contenues dans un seul tour original, mais souvent co-construites à l'intérieur de séquences dyadiques et triadiques plus ou moins étendues, où l'interprète doit coordonner l'échange et créer un espace conversationnel pour sa traduction. Bien éloignée du monologisme et de l'invisibilité qui ont dominé les études sur l'interprétation avant le tournant déterminé par Wadensjö (1998), la réalité de la profession est donc faite de dialogues auxquels l'interprète participe visiblement, comme récepteur avant même que comme traducteur de ce qui est dit. Il en découle qu'il faudrait sensibiliser les étudiants à la vaste gamme de particules discursives qui témoignent de leur activité d'auditeurs et qui ont souvent tendance à être pénalisées dans la formation et l'évaluation des futurs interprètes (car on vise une production « propre »), alors qu'elles permettent de stimuler l'interlocuteur à en dire davantage, de signaler/vérifier la compréhension, ainsi que de passer à la traduction. Pour banal que cela puisse paraître, l'interprète ne peut traduire que si l'interlocuteur s'exprime, s'il ou elle le comprend et s'il ou elle parvient à créer un espace pour sa restitution, donc l'importance de particules faisant intervenir ces faisceaux d'usages ne saurait être sous-estimée, et les étudiants devraient être amenés à les observer, outre qu'à les produire dans le discours.

## 5. Discussion conclusive

Dans cet article, une place importante a été accordée à la présentation du parcours de création du corpus AIM, parallèlement à celle des premiers résultats auxquels il a abouti. Ce parcours conserve une part irréductiblement unique, mais son explicitation nous paraissait importante pour témoigner d'une entreprise collective qui n'avait pas encore fait l'objet d'une tractation exhaustive et dont l'intérêt technique et théorique dépasse largement les frontières des universités et des disciplines au sein desquelles il a été conçu.

Étant donné la difficulté de constitution d'un corpus d'ID, il est évident que chaque chercheur dans ce domaine souhaite un logiciel à même de répondre à toutes ses exigences, à la fois de transcription et d'extraction des données. Nous étions nous-même partie de ce souhait lorsque nous avons évalué les avantages et les inconvénients de plusieurs logiciels, en finissant par choisir celui qui représentait, à l'époque, la meilleure solution de compromis. Aujourd'hui, nous sommes d'autant plus convaincue de ce choix technique, car s'il est vrai qu'il n'existe aucun logiciel en mesure de tout faire, il est vrai aussi qu'il en existe plusieurs capables de faire une partie du travail, puis de dialoguer entre eux. Cela est justement le cas des deux logiciels que nous avons testés sur le corpus AIM, puisque les transcriptions faites en 2011 avec *EXMARaLDA* ont pu être automatiquement exportées dans le format .eaf afin de bénéficier de la recherche structurée d'*ELAN*, et que rien n'empêche d'importer les transcriptions faites avec *ELAN* dans *EXMARaLDA* et donc dans *EXAKT*, qui, par rapport à l'outil de requête d'*ELAN*, offre moins de possibilités de recherche mais est beaucoup plus proche des concordanciers traditionnellement utilisés dans la linguistique de corpus.

Même si l'interopérabilité technique entre *EXMARaLDA* et *ELAN* ne fait donc, à présent, plus de doute, il faut tenir compte des besoins spécifiques pour lesquels ils ont été conçus, ainsi que du fait que la compatibilité des formats n'équivaut pas à une égalité dans la structure des données, qui dépend de l'utilisation que l'on fait de ces deux logiciels, outre les considérations théoriques qui amènent à les utiliser ainsi. Autrement dit, le fait qu'on puisse exporter les interactions transcrites avec *EXMARaLDA* dans le format .eaf et les interroger au travers d'*ELAN*, comme nous l'avons fait ici, ne signifie pas pour autant que ces données auront la structure décrite dans le tableau 2. En effet, les transcriptions datant de 2011 présentaient des acteurs complètement différents (le médiateur était IM au lieu de MEDf, le patient était Patient au lieu de PATf ou PATm) et les liens entre parents et enfants n'étaient pas nécessairement les mêmes (si MEDf présentait un sous-acteur de type traductif, IM contenait par contre un sous-acteur de type fonctionnel où nous avons codé des phénomènes qui nous intéressaient à l'époque, comme la présence de tours optionnels à l'intérieur de la séquence triadique classique). Il en découle que les recherches que nous avons affichées dans les figures 3-6 ne peuvent pas être effectuées telles quelles dans le sous-corpus EX, où les acteurs MEDf et PATf ne sont pas présents, et que pour parvenir à extraire les [okay] des médiatrices, il faut soit sélectionner IM et Patient dans les menus déroulants de la «Single Layer Search» et de la «Multiple Layer Search», soit modifier le nom des acteurs avant de lancer la recherche. Sans compter que l'*item* recherché doit nécessairement être écrit de la même façon, ce qui n'était malheureusement pas notre cas (c'était «ok» dans le sous-corpus EX et [okay]

dans le sous-corpus EL) et cela nous a donc demandé un effort supplémentaire en termes d'expressions régulières permettant d'afficher les deux versions à la fois.

En conclusion, pour s'assurer d'une plus grande interopérabilité entre les logiciels et les transcriptions, il s'avère nécessaire de définir un ensemble de normes que chaque logiciel et chaque transcripateur vont utiliser dans leur propre fonctionnement. Ces normes ont non seulement affaire à la structure des données (tableau 2) ou à la façon dont ces dernières sont transcrites (voir le cas de [okay]), mais elles touchent également à la segmentation de l'oralité en unités de base, qui peuvent être grammaticales (comme les mots), acoustiques (comme les phonèmes) ou interactionnelles (comme les tours de parole), ainsi qu'au traitement des chevauchements entre deux ou plusieurs unités. Ces normes sont autant de défis qui montrent « l'imbrication fondamentale des questions théoriques et analytiques avec les questions techniques et technologiques, qui est un trait définitoire des entreprises de corpus » (Groupe ICOR 2010: 32). Or, il n'existe malheureusement pas de normes universellement applicables pour transcrire et extraire des données d'interaction orale, et encore moins pour des corpus d'ID qui n'en sont qu'à leurs premiers balbutiements. Mais il est souhaitable de s'accorder (au minimum) sur une base commune aux individus (transcripteurs et chercheurs) et aux logiciels (de transcription et d'extraction) qui participent à la construction d'un grand corpus d'ID, faute de quoi ce corpus ne sera jamais réellement exploitable. Puisque tant les individus que les logiciels sont fortement influencés par leurs objectifs, le dialogue entre les uns et les autres représente un atout majeur pour la recherche assistée par ordinateur, ce qui ne fait que confirmer la valeur ajoutée d'un logiciel comme *EXMARaLDA*, qui de la collaboration entre chercheurs et techniciens fait sa raison d'être.

## REMERCIEMENTS

Nous tenons à remercier Guy Aston et Laura Gavioli, car même si leurs noms ne figurent pas dans la liste des auteurs et s'ils ne partagent pas la responsabilité de ce que nous avons écrit, cet article s'inspire d'une présentation faite avec Laura dans le cadre d'un atelier coordonné par Guy: il est donc le fruit d'une contamination entre disciplines et méthodologies qui nous dépassent largement.

## NOTES

- \* Membre du Centre interuniversitaire d'analyse de l'interaction et de la médiation ([www.aim.unimore.it](http://www.aim.unimore.it)).
- 1. Contrairement aux étiquettes qui rendent compte du contexte dans lequel l'interaction se déroule, par exemple *interprétation communautaire* (Délizée 2015) ou *interprétation de service public* (Pointurier 2016), l'ID met l'accent sur le fonctionnement effectif de l'échange et sur sa nature dialogale (Falbo 2009), pouvant couvrir des milieux aussi variés que les hôpitaux, les tribunaux, les bureaux d'immigration, ainsi que d'autres institutions publiques et privées.
- 2. Nous nous devons de préciser qu'un nombre de chercheurs tout aussi nombreux s'est penché sur l'ID en langue des signes, de l'ouvrage pionnier de Metzger (1999) jusqu'aux études plus récentes de Turner et Merrison (2016), en passant par plusieurs volumes incontournables publiés par la Gallaudet University Press.
- 3. Pour des recueils, voir: Valero-Garcés et Martin (2008); Baraldi et Gavioli (2012); Davitti et Pasquandrea (2014); Dal Fovo et Niemants (2015); Russo, Bendazzoli, *et al.* (2017); Bendazzoli, Russo, *et al.* (2018).
- 4. Sur la distinction entre données primaires et secondaires et sur les avantages de leur alignement pour l'étude de l'oral en interaction, voir aussi la publication du Groupe ICOR (2010: 21-24).

5. Pour une description des méthodes et des outils de recherche quantitative pouvant être utilisés dans les études sur la traduction et l'interprétation, voir Mellinger et Hanson (2017).
6. Contrairement à ce qui se produit dans le cas des corpus parallèles de discours monologiques et de leurs traductions (écrites ou orales), où il est aujourd'hui assez aisé de repérer les unités en langue source et leurs traductions correspondantes en langue cible, dans les corpus d'ID se pose le problème de l'unité de traduction, car la parole à traduire est négociée localement par les participants à la rencontre et la parole traduite ne suit pas nécessairement, ni de près, son correspondant dans la langue originale. Il en découle qu'il est impossible d'automatiser la recherche des tours originaux et de leurs restitutions par l'interprète, et que seul l'œil humain est en mesure de faire la distinction, d'ailleurs non pas sans difficultés, entre *renditions* et *non-renditions*. Voir Wadensjö (1998) pour cette distinction théorique et Castagnoli et Niemants (2018) pour une application pratique de cette distinction à un corpus d'interprétation téléphonique.
7. ANGERMEYER, Philip S., MEYER, Bernd et SCHMIDT, Thomas (2012- ) : *ComInDat. Community Interpreting Database*. Hambourg : Hamburger Zentrum für Sprachkorpora (Universität Hamburg). Consulté le 4 décembre 2018, <<http://www.yorku.ca/comindat/comindat.htm>>.
8. GAVIOLI, Laura, BARALDI, Claudio et NIEMANTS, Natacha (2015) : The AIM corpus of dialogue interpreting in healthcare : some thoughts on coding, counting, searching. *First Forlì International Workshop – Corpus-based Interpreting Studies: The State of the Art*. Forlì, 7-8 mai 2015.
9. GAVIOLI, Laura, BARALDI, Claudio et NIEMANTS, Natacha (2016) : From archives to corpora : extracting data for analysis, from a collection of interpreter-mediated interactions. *CL8: 8<sup>th</sup> International Critical Link Conference*. Édimbourg, 29 juin-1 juillet 2016.
10. Les deux barres verticales dans [okay] indiquent qu'il s'agit d'un homographe interlangue. Ce symbole est également employé lorsque la langue d'expression est incertaine (cf. [okay] de transition). Si la langue est connue ou peut être inférée, la typographie propre à chaque langue sera appliquée.
11. Sur la fonction de transition de [okay], voir par exemple Beach (1993).
12. HAMBURGER ZENTRUM FÜR SPRACHKORPORA (9 mai 2017) : *EXMARALDA*. Hambourg : Universität Hamburg. Consulté le 30 novembre 2018, <<http://exmaralda.org>>.
13. BURNARD, Lou et DODD, Tony (6 août 2010) : *XAIRA*. Version 1.26. Consulté le 3 octobre 2018, <<http://xaira.sourceforge.net/>>.
14. THE LANGUAGE ARCHIVE (4 avril 2018) : *ELAN*. Version 5.2. Nimègue : Max Planck Institute for Psycholinguistics. Consulté le 22 septembre 2018, <<https://tla.mpi.nl/tools/tla-tools/elan/>>.
15. P.I. Prof. Claudio Baraldi, Università di Modena e Reggio Emilia, financé par le programme compétitif FAR 2014 et clôturé par un séminaire international qui s'est tenu à Modène le 13 décembre 2016.
16. Il s'agit de deux éléments qui font à présent l'objet d'un autre projet financé par le programme compétitif FAR (FAR 2017, P.I. Prof. Claudio Baraldi, Università di Modena e Reggio Emilia), dont le but est celui d'analyser l'interaction médecin-patient pendant la consultation andrologique, en observant les mécanismes de participation et de communication centrés sur le patient.
17. Cette utilité est bien montrée par les analyses de *voilà* que Bruxelles et Traverso (2006) ont fait en se focalisant sur une réunion entre architectes et que le groupe ICOR a ensuite mené sur l'ensemble de CLAPI).
18. Toutes les données sont traduites par nos soins.
19. Même si cela sort de la portée de ce travail, nous signalons qu'une fonctionnalité intéressante d'*ELAN* est celle qui permet d'analyser des phénomènes simultanés qui se produisent sur deux ou plusieurs acteurs et sous-acteurs, et donc de chercher des annotations et/ou des mots qui se superposent totalement ou qui se chevauchent partiellement (à droite ou à gauche).
20. Nous nous devons de rappeler que ce menu déroulant offre également une troisième option, à savoir les expressions régulières, mais nous les réservons à des cas plus complexes, comme ceux où l'on cherche plusieurs mots en succession dans une même annotation, ou bien ceux où les occurrences sont particulièrement nombreuses et où on a donc besoin de filtrer les résultats. À titre d'exemple, si on veut restreindre la recherche aux [okay] en début de tour, il suffit de chercher \bokay\b.+ ; si, ou contraire, on cherche les [okay] en fin de tour, on peut utiliser l'expression .+ \bokay\b ; avec l'expression .+ \bokay\b.+ , on obtient, enfin, tous les [okay] qui sont précédés et suivis de quelque chose, donc au milieu du tour.
21. Le fait de spécifier la durée de l'écart n'est que l'une des options d'un menu déroulant qui contient une multiplicité de liens possibles entre les acteurs recherchés. On ne saurait qu'encourager les

- usagers à tester ces options pour vérifier celle(s) qui permet(tent) d'extraire le plus grand nombre d'occurrences (dans notre cas, c'est l'écart de deux secondes qui a produit les meilleurs résultats).
22. L'ÉQUIPE AUDACITY (6 décembre 2017): *Audacity*. (Créé par Dominic MAZZONI et Roger DANNENBERG) Version 2.2.1. Consulté le 3 janvier 2018, <<https://www.audacityteam.org/>>.

## RÉFÉRENCES

- AMATO, Amalia et MACK, Gabriele (2015): *Comunicare tramite interprete nelle indagini di polizia. Implicazioni didattiche di un'analisi linguistica* [Communiquer par l'entremise d'interprète lors d'enquêtes policières. Implications pédagogiques d'une analyse linguistique]. Bologne: Bononia University Press.
- ANGELELLI, Claudia (2004): *Medical Interpreting and Cross-cultural Communication*. Cambridge: Cambridge University Press.
- ANGERMEYER, Philip, MEYER, Bernd et SCHMIDT, Thomas (2012): Sharing community interpreting corpora. In: Thomas SCHMIDT et Kai WÖRNER, dir. *Multilingual Corpora and Multilingual Corpus Analysis*. Amsterdam/Philadelphie: John Benjamins, 275-294.
- ASTON, Guy et CENCINI, Marco (2002): Resurrecting the corp(us|se): Towards an encoding standard for interpreting data. In: Giuliana GARZONE et Maurizio VIEZZI, dir. *Interpreting in the 21st century: Challenges and Opportunities*. Amsterdam/Philadelphie: John Benjamins, 47-62.
- AYASS, Ruth (2015): Doing data: the status of transcripts in conversation analysis. *Discourse Studies*. 17(5):505-528.
- BARALDI, Claudio, et GAVIOLI, Laura, dir. (2012): *Coordinating Participation in Dialogue Interpreting*. Amsterdam/Philadelphie: John Benjamins.
- BEACH, Wayne (1993): Transitional regularities for 'casual' "okay" usages. *Journal of Pragmatics*. 19:325-352.
- BENDAZZOLI, Claudio, RUSSO, Mariachiara et DEFRANCQ, Bart (2018): Introduction. In: Claudio BENDAZZOLI, Mariachiara RUSSO et Bart DEFRANCQ, dir. *New Findings in Corpus-based Interpreting Studies*. inTRAlinea. Consulté le 14 décembre 2018, <<http://www.intralineat.org/specials/cbis>>.
- BOLDEN, Galina (2000): Toward understanding practices of medical interpreting: interpreters' involvement in history taking. *Discourse Studies*. 2(4):387-419.
- BOT, Hanneke (2005): *Dialogue Interpreting in Mental Health*. Amsterdam/New York: Rodopi.
- BOULTON, Alex et TYNE, Henry (2014): *Des documents authentiques aux corpus*. Paris: Didier.
- BRUXELLES, Sylvie et TRAVERSO, Véronique (2006): Usages de la particule voilà dans une réunion de travail: analyse multimodale. In: Martina DRESCHER et Barbara JOB, dir. *Les marqueurs discursifs dans les langues romanes: approches théoriques et méthodologiques*. Berne: Peter Lang, 71-92.
- BUHRIG, Kristin et MEYER, Bernd (2004): Ad hoc-interpreting and the achievement of communicative purposes in doctor-patient communication. In: Juliane HOUSE et Jochen REHBEIN, dir. *Multilingual Communication*. Amsterdam/Philadelphie: John Benjamins, 43-62.
- BUHRIG, Kristin, Kliche, Ortrun, MEYER, Bernd, et al. (2012): The corpus 'Interpreting in hospitals': Possible applications for research and communication training. In: Thomas SCHMIDT et Kai WÖRNER, dir. *Multilingual Corpora and Multilingual Corpus Analysis*. Amsterdam/Philadelphie: John Benjamins, 305-315.
- CASTAGNOLI, Sara et NIEMANTS, Natacha (2018): Corpora worth creating: A pilot study on telephone interpreting. In: Claudio BENDAZZOLI, Mariachiara RUSSO et Bart DEFRANCQ, dir. *New Findings in Corpus-based Interpreting Studies*. inTRAlinea. Consulté le 14 décembre 2018, <<http://www.intralineat.org/specials/article/2315>>.
- CIRILLO, Letizia et NIEMANTS, Natacha, dir. (2017): *Teaching Dialogue Interpreting. Research-based Proposals for Higher Education*. Amsterdam/Philadelphie: John Benjamins.
- COLÓN DE CARVAJAL, Isabel (2010): Du corpus enregistré au corpus analysé: questions méthodologiques sur l'utilisation d'outils de requêtes informatisés. *Cahiers de praxématique*. 54-55:313-326.

- COOK, Guy (1995): Theoretical issues: transcribing the untranscribable. In: Geoffrey LEECH, Greg MYERS et Jenny THOMAS, dir. *Spoken English on Computer*. Londres/New York: Routledge, 35-53.
- DAL FOVO, Eugenia et NIEMANTS, Natacha, dir. (2015): Studying Dialogue Interpreting: An Introduction. In: Eugenia DAL FOVO et Natacha NIEMANTS, dir. *Dialogue Interpreting. The Interpreters' Newsletter*. 20:1-9.
- DAL FOVO, Eugenia (2018): The use of dialogue interpreting corpora in healthcare interpreters' training: taking stock. *The Interpreters' Newsletter*. 23:83-113.
- DAVIDSON, Brad (2000): The interpreter as institutional gatekeeper: the social-linguistic role of interpreters in Spanish-English medical discourse. *Journal of Sociolinguistics*. 4(3):378-405.
- DAVIDSON, Brad (2002): A model for the construction of conversational common ground in interpreted discourse. *Journal of Pragmatics*. 34(9):1273-1300.
- DAVITTI, Elena (2013): Dialogue interpreting as intercultural mediation. Interpreters' use of upgrading moves in parent-teacher meetings. *Interpreting*. 15(2):168-199.
- DAVITTI, Elena et PASQUANDREA, Sergio (2014): Guest Editorial. In: Elena DAVITTI et Sergio PASQUANDREA, dir. *Dialogue Interpreting in Practice: Bridging the gap between empirical research and interpreter education. The Interpreter and Translator Trainer*. 8(3):329-335.
- EDWARDS, Jane et LAMPERT, Martin, dir. (1993): *Talking Data: Transcription and Coding in Discourse Research*. Hillsdale: Lawrence Erlbaum.
- ENGLUND DIMITROVA, Birgitta (1997): Degree of interpreter responsibility in the interaction process in community interpreting. In: Silvana CARR, Roda ROBERTS, Aileen DUFOR, et al., dir. *The Critical Link: Interpreters in the Community*. Amsterdam/Philadelphie: John Benjamins, 147-164.
- FALBO, Caterina (2005): La transcription: une tâche paradoxale. *The Interpreters' Newsletter*. 13:25-38.
- FALBO, Caterina (2009): Un grand corpus d'interprétation: à la recherche d'une stratégie de classification. In: Paola PAISSA et Marta BIAGINI, dir. *Cahiers de recherche de l'École Doctorale en linguistique française*. Milan: Lampi di stampa, 105-120.
- FALBO, Caterina (2013): 'Interprete' et 'mediatore' linguistico-culturale': deux figures professionnelles opposées? In: Giovanni AGRESTI et Cristina SCHIAVONE, dir. *Plurilinguisme et monde du travail: professions, opérateurs et acteurs de la diversité linguistique*. Rome: Aracne, 253-270.
- GALAZZI, Enrica (2002): *Le son à l'école*. Brescia: La Scuola.
- GARDNER, Rod (2001): *When Listeners Talk: Response Tokens and Listeners Stance*. Amsterdam/Philadelphie: John Benjamins.
- GAVIOLI, Laura et MANSFIELD, Gillian, dir. (1990): *The PIXI corpora*. Bologne: Cooperativa Libreria Universitaria Editrice.
- GAVIOLI, Laura (2012): Minimal responses in interpreter-mediated medical talk. In: Claudio BARALDI et Laura GAVIOLI, dir. *Coordinating Participation in Dialogue Interpreting*. Amsterdam/Philadelphie: John Benjamins, 201-228.
- GAVIOLI, Laura (2015): On the distribution of responsibilities in creating critical issues in interpreter-mediating medical consultations: the case of 'le spieghi(amo)'. *Journal of Pragmatics*. 76:169-180.
- GROUPE ICOR (2007): L'étude des particules à l'oral dans différents contextes à partir de la banque de données de Corpus de Langue Parlée en Interaction CLAPI. In: Geoffrey WILLIAMS, dir. *Actes des 5<sup>èmes</sup> Journées de la Linguistique de Corpus*. (5<sup>èmes</sup> Journées de Linguistique de Corpus, Lorient, 13-15 septembre 2007). *Texte et corpus*. 3:233-244.
- GROUPE ICOR (2008): "Oh::, oh là là, oh ben...", les usages du marqueur oh en français parlé en interaction. In: Jacques DURANT, Benoit HABERT et Bernard LACKS, dir. *Congrès Mondial de Linguistique Française – CMLF'08*. (CMLF 2008: 1<sup>er</sup> Congrès mondial de linguistique française, Paris, 9-12 juillet 2008). Paris: Institut de linguistique française, 685-701.
- GROUPE ICOR (2010): Grands corpus et linguistique outillée pour l'étude du français en interaction (plateforme CLAPI et corpus CIEL). *Pratiques*. 147/148:17-34.

- GUILLAUME-HOFNUNG, Michelle (2015): *La médiation*. Paris: PUF.
- HALE, Sandra (2004): *The Discourse of Court Interpreting*. Amsterdam/Philadelphie: John Benjamins.
- HERITAGE, John, ELLIOTT, Marc, STIVERS, Tanya, *et al.* (2010): Reducing inappropriate antibiotics prescribing: the role of online commentary on physical examination findings. *Patient Education and Counselling*. 81(1):119-125.
- KADRIC, Mira (2000): Interpreting in the Austrian courtroom. In: Roda ROBERTS, Silvana CARR, Diana ABRAHAM, *et al.*, dir. *The Critical Link 2: Interpreters in the community*. Amsterdam/Philadelphie: John Benjamins, 153-164.
- KESELMAN, Olga, CEDERBORG, Ann-Christin, LINELL, Per, *et al.* (2010): 'That is not necessary for you to know': Negotiation of participation status of unaccompanied children in interpreter-mediated asylum hearings. *Interpreting*. 12(1):83-104.
- KRYSTALLIDOU, Demi (2016): Investigating the interpreter's role(s). The A.R.T framework. *Interpreting*. 18(2):172-197.
- MERLINI, Raffaella (2015): Mediation. In: Franz PÖCHHACKER, dir. *The Routledge Encyclopedia of Interpreting Studies*. Londres/New York: Routledge.
- MACK, Gabi et RUSSO, Mariachiara, dir. (2005): *Interpretazione di trattativa: la mediazione linguistico-culturale nel contesto formativo e professionale* [Interprétation de liaison: médiation linguistique et culturelle dans le contexte de la formation et de la profession]. Milan: Hoepli.
- MAYNARD, Senko (1986): On back-channel behavior in Japanese and English casual conversation. *Linguistics*. 24:1079-1108.
- MASON, Ian, dir. (1999): Introduction. In: Ian MASON, dir. *Dialogue Interpreting. The Translator*. 5(2):147-160.
- MASON, Ian, dir. (2001): *Triadic Exchanges: Studies in Dialogue Interpreting*. Manchester/Northampton: St. Jerome.
- MASON, Ian (2006): On mutual accessibility of contextual assumptions in dialogue interpreting. *Journal of Pragmatics*. 38(3):359-373.
- MELLINGER, Christopher et HANSON, Thomas (2017): *Quantitative Research Methods in Translation and Interpreting Studies*. Londres/New York: Routledge.
- MERLINI, Raffaella et FAVARON, Roberta (2005): Examining the voice of interpreting in speech pathology. *Interpreting*. 7(2):263-301.
- METZGER, Melanie (1999): *Sign Language Interpreting: Deconstructing the Myth of Neutrality*. Washington (D.C.): Gallaudet University Press.
- NIEMANTS, Natacha (2012): The transcription of interpreting data. *Interpreting*. 14(2):165-191.
- NIEMANTS, Natacha (2013): L'utilisation de corpus d'entretiens cliniques (français/italien) dans la didactique de l'interprétation en milieu médical. In: Dorothee HELLER, Michele SALA et Cécile DESOUTTER, dir. *Corpora in Specialized Communication/Korpora in der Fachkommunikation/Les corpus dans la communication spécialisée*. Bergame: CELSB, 209-235.
- NIEMANTS, Natacha (2015a): *L'interprétation de dialogue en milieu médical: du jeu de rôle à l'exercice d'une responsabilité*. Rome: Aracne.
- NIEMANTS, Natacha (2015b): Transcription. In: Franz PÖCHHACKER, dir. *The Routledge Encyclopedia of Interpreting Studies*. Londres/New York: Routledge, 421-423.
- NIEMANTS, Natacha et CIRILLO, Letizia (2016): Il role-play nella didattica dell'interpretazione dialogica: Focus sull'apprendente [Le jeu de rôle dans l'enseignement de l'interprétation dialogique: l'apprenant sous la loupe]. In: Roberta GRASSI et Cecilia ANDORNO, dir. *Le dinamiche dell'interazione. Prospettive di analisi e contesti applicativi* [Dynamiques de l'interaction. Perspectives d'analyse et contextes d'application]. (XVI Congresso Internazionale di Studi dell'Associazione Italiana di Linguistica Applicata, Modène, 18-20 février 2016). Studi AltLA. 5. Milan: Officinaventuno, 301-317.
- NIEMANTS, Natacha et PALLOTTI, Gabriele (2017): The human viewpoint and the system's viewpoint. In: Paul SEEDHOUSE, dir. *Task-based Language Learning in a Real-world Digital Environment*. Londres: Bloomsbury, 99-134.

- NIEMANTS, Natacha (2017) : La médiation interculturelle : un dialogue au service de la santé. In : Philippe GRECIANO, dir. *La médiation dans un monde sans frontières*. Paris : Mare & Martin, 47-72.
- NIEMANTS, Natacha et STOKOE, Elisabeth (2017) : Using the Conversation Analytic Role-Play Method in healthcare interpreter education. In : Letizia CIRILLO et Natacha NIEMANTS, dir. *Teaching Dialogue Interpreting. Research-based Proposals for Higher Education*. Amsterdam/Philadelphie : John Benjamins.
- O'CONNEL, Daniel et KOWAL, Sabine (1999) : Transcription and the use of standardization. *Journal of Psycholinguistic Research*. 28(2):103-120.
- PALLOTTI, Gabriele (2002) : La classe dans une perspective écologique de l'acquisition. *Acquisition et interaction en langue étrangère*. 16:165-197.
- PARISSE, Christophe et MORGENSTERN, Aliyah (2010) : Transcrire et analyser les corpus d'interactions adulte-enfant. In : Edy VENEZIANO, Anne SALAZAR ORVIG et Josie BERNICOT, dir. *Acquisition du langage et interaction*. Paris : L'Harmattan, 201-222.
- PASQUANDREA, Sergio (2011) : Managing multiple actions through multimodality. Doctors' involvement in interpreter-mediated interactions. *Language in Society*. 40(4):455-481.
- PEDERZOLI, Lucia (2011) : *La mediazione interculturale: un'analisi di interazioni presso istituzioni sanitarie in Belgio* [La médiation interculturelle : une analyse des interactions dans les établissements de santé en Belgique]. Mémoire de maîtrise, non publié. Modène : Università di Modena e Reggio Emilia.
- RUSO, Mariachiara, BENDAZZOLI, Claudio et DEFRANCQ, Bart, dir. (2017) : *Making Way in Corpus-based Interpreting Studies*. Singapore : Springer.
- SACKS, Harvey, SCHEGLOFF, Emmanuel et JEFFERSON, Gail (1974) : A simplest systematics for the organization of turn-taking for conversation. *Language*. 50(4):696-735.
- SCHEGLOFF, Emmanuel (1993) : Reflections on quantification in the study of conversation. *Research on Language and Social Interaction*. 26(1):99-128.
- STIVERS, Tanya (2015) : Coding Social Interaction : a heretical approach in conversation analysis? *Research on Language and Social Interaction*. 48(1):1-19.
- STOKOE, Elisabeth (2011) : Simulated interaction and communication skills training. In : Charles ANTAKI, dir. *Applied Conversation Analysis*. Londres/New York : Palgrave Macmillan, 119-139.
- STRANIERO SERGIO, Francesco (2007) : *La mediazione linguistica nella conversazione spettacolo* [Médiation linguistique dans la conversation spectacle]. Trieste : Edizioni Università di Trieste.
- STRANIERO SERGIO, Francesco et FALBO, Caterina (2012) : Studying interpreting through corpora. An introduction. In : Francesco STRANIERO SERGIO et Caterina FALBO, dir. *Breaking Ground in Corpus-based Interpreting Studies*. Berne : Peter Lang, 9-52.
- THOMPSON, Paul (2005) : Spoken language corpora. In : Martin WYNNE, dir. *Developing Linguistic Corpora : A Guide to Good Practice*. Oxford : Oxbow Books, 59-70.
- TICCA, Anna Claudia et TRAVERSO, Véronique (2017) : Participation in bilingual interactions : translating, interpreting and mediating documents in a French social centre. *Journal of Pragmatics*. 107:129-146.
- TRAVERSO, Véronique (1999) : *L'analyse des conversations*. Paris : Armand Colin.
- TRAVERSO, Véronique (2003a) : Rencontres interculturelles à l'hôpital : la consultation médicale avec interprète. *Tranel*. 36:81-100.
- TRAVERSO, Véronique (2003b) : *Analyse des interactions : questions sur la pratique*. Synthèse d'habilitation à diriger des recherches, non publiée. Lyon : Université Lumière-Lyon-II.
- TRAVERSO, Véronique (2016) : *Décrire le français parlé en interaction*. Paris : Ophrys.
- TURNER, Graham et MERRISON, Andrew (2016) : Doing 'understanding' in dialogue interpreting. *Interpreting*. 18(2):137-171.
- VALERO-GARCÉS, Carmen et MARTIN, Anne, dir. (2008) : *Crossing Borders in Community Interpreting*. Amsterdam/Philadelphie : John Benjamins.

- VECCHIATO, Sara, GEROLIMICH, Sonia et KIMNINOS, Nickolas, dir. (2015): *Plurilingualism in Healthcare. An Insight from Italy*. Londres: Bica.
- WADENSJÖ, Cecilia (1998): *Interpreting as Interaction*. Londres: Longman.
- WILLIAMS, Geoffrey (2008): Traduction et corpus, corpus et recherche. *Cahiers de l'APLIUT*. 27(1):69-79.
- ZAHN, Cheng et ZENG, Lishan (2017): Chinese medical interpreters' visibility through text ownership. *Interpreting*. 19(1):97-117.
- ZANETTIN, Federico (2012): *Translation-Driven Corpora: Corpus Resources for Descriptive and Applied Translation Studies*. Londres/New York: Routledge.

## ANNEXES

### Annexe 1

Conventions de transcription utilisées dans les extraits:

- ? intonation montante, pas nécessairement une question;
- (.) silence inférieur à 0,5 seconde (les chiffres entre parenthèses indiquent la durée des silences plus longs);
- (??) mots ou lettres incompréhensibles, si possible avec indication du nombre de syllabes (??3syll);
- : prolongation du son qui précède;
- [xxx] les crochets entourent le discours simultané entre deux ou plusieurs participants;
- ab- interruption abrupte ou mot incomplet;
- °abc° discours prononcé à un volume peu élevé.