

CAPITOLO IV

Dall'analisi strategica alla comparazione dei risultati

4.1 Introduzione

Il capitolo tre ha permesso di introdurre il materiale di studio e la metodologia adottati nel corso della ricerca in modo che il presente capitolo potesse essere interamente dedicato all'analisi strategica dei sottotitoli interlinguistici e al confronto dei risultati ottenuti da tale analisi con i risultati dell'analisi strategica dei sottotitoli intralinguistici.

È stato già specificato nell'introduzione e nel capitolo precedente che i sottotitoli interlinguistici analizzati sono il risultato di una sperimentazione sul *respeaking* interlinguistico effettuata grazie al lavoro di tre rispeaker, quattro *editor* e un suggeritore, ma è opportuno ricordare che ogni rispeaker ha sottotitolato in diretta tutto il materiale (a eccezione di dieci minuti per il rispeaker 1 e di cinque minuti per il rispeaker 2 a causa di problemi tecnici) e in semidiretta una trasmissione breve a scelta (dai due ai quattro minuti). Perciò, mentre i sottotitoli intralinguistici a disposizione sono il risultato di 46 minuti totali di rispeakeraggio in francese, i sottotitoli interlinguistici analizzati sono il risultato di 131:52 minuti di rispeakeraggio in italiano, di cui 123:28 in diretta e appena 8:24 in semidiretta.

In particolare, nel paragrafo 4.2, si presenterà l'analisi strategica dei sottotitoli interlinguistici risultanti da ogni singola prova in modo da avere un riscontro il più possibile esaustivo riguardo le diverse soluzioni strategiche e difficoltà adottate e riscontrate dai *respeaker* a fronte di diverse condizioni di lavoro: diversi generi e argomenti; TP precedentemente programmato o a braccio; diversi ritmi dei TP; difficoltà della lingua utilizzata (grammatica, registro, lessico, sintassi, idioletti); preponderanza o meno delle immagini; presenza di uno, due o più oratori; diversi atti comunicativi; sovrapposizione o successione dei turni di parola; velocità e prevedibilità nel cambio turno; diversa durata della trasmissione e quindi del turno di lavoro (un rispeaker che lavora 2 minuti per sottotitolare il meteo ha una *performance* diversa rispetto a quello che lavora 15 minuti di fila). A conferma dell'oggettività dell'analisi effettuata, si precisa che ogni passaggio (TP > TM1 > TM2 > TA) è stato registrato e debitamente preso in considerazione al fine di determinare la reale strategia utilizzata e di imputare gli errori in maniera precisa e corretta al rispeaker, all'*editor* e/o al *software*. Il paragrafo successivo, invece, si occuperà dei miglioramenti (o peggioramenti)

riscontrati in una prova in semidiretta¹⁵⁸. Infine, si passerà alla comparazione dei migliori risultati ottenuti per ogni trasmissione (quindi, per ogni trasmissione, verrà selezionata la migliore tra le tre versioni sottotitolate in diretta) con i rispettivi risultati dell'analisi strategica dei sottotitoli intralinguistici della stessa trasmissione. Questa selezione intende colmare il divario tra l'una e l'altra attività di *respeaking* (quella intralinguistica in francese e quella interlinguistica in italiano), comparabili a livello tecnico in quanto relative alle stesse trasmissioni ed espletate in condizioni di lavoro simili (stessa modalità di ricezione del TP, conoscenza della terminologia tecnica, uso dell'*editing*, ecc.), ma caratterizzate da un diverso grado di difficoltà dovuto al passaggio, per il *respeaking* interlinguistico, da una lingua e da una cultura all'altra.

4.2 Analisi strategica dei sottotitoli interlinguistici: trasmissioni in diretta

Prima di passare alla descrizione dei risultati ottenuti dall'analisi strategica dei sottotitoli interlinguistici realizzati in ogni prova, è opportuno fare un breve raffronto riguardo alle *performance* dei rispeaker. Questo perché, per quanto si tratti di un esperimento assolutamente nuovo, la formazione, le competenze e l'approccio di ognuno hanno influito, positivamente o negativamente, sui risultati della ricerca. Ai fini di interpretare correttamente i dati ottenuti, occorre tenere in considerazione innanzitutto il fatto che i rispeaker non avevano alcuna familiarità con la postazione di sottotitolazione utilizzata e tanto meno esperienza in questo preciso campo, cioè il *respeaking* interlinguistico. Inoltre, bisogna aggiungere che i rispeaker 1 e 3 sono interpreti professionisti con conoscenze sul *respeaking*, mentre il rispeaker 2 è un interprete e un rispeaker intralinguistico professionista.

Il rispeaker 1, a fronte di un'interpretazione fedele e quasi completa (seppure molto letterale), ha mostrato diversi problemi nell'interazione con la macchina che si sono tradotti in un elevato numero di errori imputabili al *software*. Le ragioni sono da ricercarsi proprio nell'eloquio del rispeaker: uso di una voce insicura, stressata e a volte troppo animata; esitazioni ('ehm', 'eeee', 'cheee', ecc.); ritmo di eloquio a volte troppo veloce¹⁵⁹; parole 'mangiate' o pronunciate in modo non troppo chiaro; introduzione di esclamazioni e tratti dell'oralità ('hey', 'oh', 'ah', ecc.). Tutto ciò ha provocato, in generale, un cattivo

¹⁵⁸ Le altre due prove in semidiretta sono contenute nelle tabelle in appendice, ma non sono state analizzate.

¹⁵⁹ Secondo Romero-Fresco (2011:116), la velocità dell'eloquio nel *respeaking* varia tra le 106 e le 190 parole al minuto ed è generalmente inferiore rispetto a quella del TP (da cui l'impossibilità di realizzare sottotitoli totalmente *verbatim*). Solitamente, "respeakers lag 0-20 wpm behind speakers who speak at up to 180 wpm and approximately 40 wpm behind speakers who speak faster."

riconoscimento da parte del *software* e un conseguente maggior carico di lavoro per l'*editor*. Quest'ultimo, non avendo a disposizione il TP e trovandosi davanti a troppi errori di non riconoscimento, in molti casi non è riuscito a dare una logica al testo e ha preferito cancellare interi sottotitoli, in altri ha cercato di rimediare aumentando a dismisura il ritardo nell'invio degli stessi. Inoltre, la ripetizione di parole o addirittura di intere frasi da parte del rispeaker, seppur trascritte bene dal *software*, hanno contribuito a confondere l'*editor*. In tutte queste occasioni, la presenza del suggeritore in appoggio all'*editor* si è rivelata un'ottima scelta. Infine, le macro-unità rese grazie a una traduzione letterale sono veramente poche in rapporto agli errori. Un eccessivo attaccamento al testo nel tentativo di tradurre tutto diminuisce il numero di omissioni, ma aumenta quello degli errori poiché i calchi e la struttura della frase non corrispondono ai modelli di linguaggio, alle grammatiche e ai tassi di frequenza contenuti nel *software*.

Il rispeaker 2, l'unico ad aver già testato la macchina in precedenza e ad aver introdotto anche la punteggiatura nella produzione del TM1, ha dimostrato maggior dimestichezza e affinità col *software*. Infatti, ha cercato di effettuare il *pre-editing* in maniera corretta; ha utilizzato una voce molto chiara, ferma e rilassata; ha tenuto un ritmo costante; non ha mostrato esitazioni; ha monitorato il TM2 (a volte perdendo pochi secondi del TP per suggerire una correzione all'*editor*). Questo si è tradotto in un basso numero di errori dovuti al *software* (con l'eccezione di due trasmissioni) e in un conseguente incremento di errori imputabili al rispeaker, meno attento al TP ed eccessivamente attaccato al testo in alcuni punti. Ciò non significa, però, che i risultati ne abbiano eccessivamente risentito. Al contrario, le prove del rispeaker 2 sono da considerarsi, in generale, le migliori.

Per quanto riguarda il rispeaker 3, questi si è rivelato essere una via di mezzo rispetto agli altri due. Agevolato nella conoscenza delle trasmissioni, ha fatto pochi errori di senso compensati da un alto numero di espansioni (non sempre corrette). Inoltre, il rispeaker 3 è stato l'unico a segnalare il cambio locutore con ottimi risultati. Se da un lato il riconoscimento da parte del *software* è stato perlopiù soddisfacente, dall'altro si è notata un'eccessiva impostazione da interprete a causa dell'introduzione di molti tratti dell'oralità e connettori (*allora, quindi, ecc.*) e di un ritmo di eloquio a volte troppo veloce per stare al passo col TP. Tutto ciò si è tradotto in un TM2 non sempre coeso e con parti mancanti. Infatti, l'eccessiva velocità di eloquio non dà tempo al *software* di elaborare intere stringhe di parole e di distinguere bene i termini (soprattutto monosillabi e bisillabi), accumulando nel frattempo sempre più sottotitoli nel *buffer*. Di conseguenza, l'*editor* è portato a cancellare interi sottotitoli errati aumentando la velocità di *editing* per recuperare. Però, a volte, nella fretta l'*editor* cancella anche i sottotitoli giusti che non riesce a distinguere bene all'interno del testo

trascritto. Questo porta a un aumento del numero di macro-unità non rese per errore del *software* (in caso di non riconoscimento) e/o dell'*editor* (in caso di eliminazione di sottotitoli corretti). Nel momento in cui una parte del TP risulta essere poco chiara, il rispeaker 3 tende a completare le frasi ma in modo insicuro, con numerose ripetizioni ed esitazioni, aggiungendo spesso informazioni sbagliate che si traducono in un aumento delle espansioni, ma anche in un aumento delle macro-unità non rese per errore del rispeaker. Di conseguenza, aumentano anche gli errori da parte del *software* e le difficoltà per l'*editor*.

4.2.1 *Météo*

Il meteo si è rivelata la trasmissione più semplice da sottotitolare per diversi motivi. Innanzitutto, dura appena 2:10 minuti perciò il *respeaker* non rischia di affaticarsi eccessivamente. In secondo luogo, gran parte delle macro-unità contengono informazioni ridondanti rispetto alla componente video e quindi possono essere omesse e/o compresse senza comportare una perdita di informazione. In compenso, la preponderanza delle immagini richiede un ritardo minimo nella trasmissione dei sottotitoli rispetto al TP per evitare che questi siano proiettati sotto le immagini sbagliate creando confusione nel telespettatore. Ciò implica un elevato numero di riduzioni (tra l'82.15% e il 95% delle alterazioni), mentre il ritmo di eloquio sostenuto, l'alta densità lessicale e la bassa complessità grammaticale non forniscono il tempo necessario per applicare la strategia dell'espansione (tra il 5% e il 12.5%) o modificare eccessivamente la struttura della frase, sfociando in una traduzione perlopiù letterale e ricca di calchi¹⁶⁰. Tutto ciò risulta, infine, in un maggior numero di omissioni a livello di micro-unità (tra il 48.3% e il 67.5%) a cui il *rispeaker* ricorre per non sovraccaricare troppo la propria memoria a breve termine o per recuperare il passo col TP. Per quanto riguarda gli errori, non si può trarre una tendenza generale perché questo dato è fortemente dipendente dalla comprensione del TP da parte del *rispeaker*, dal riconoscimento da parte del *software* e dalla politica seguita dall'*editor*. Perciò saranno trattati a parte per ogni prova.

¹⁶⁰ I tecnicismi non consentono di scegliere tra più possibilità di traduzione e la povertà grammaticale non consente grandi cambiamenti, nonostante ciò la traduzione di un testo a ritmi così elevati implica necessariamente una compressione dello stesso, che oscilla, infatti, tra il 32.5% e il 51.7% all'interno delle riduzioni.

Respeaker 1

MACRO-UNITÀ NON RESE							MACRO-UNITÀ RESE				TOT
omissioni			errori			Tot	traduzioni letterali		alterazioni	Tot	46 100%
16 69.6%			7 30.4%			23 50%	3 13%		20 87%	23 50%	
sem.	non sem.	Tot	umani		software	Tot					
5 31.25%	11 68.75%	16 100%	R	E	3 42.85%	7 100%					
			1 14.3%	3 42.85%							

alterazioni													
riduzioni						espansioni			errori			Tot	
19 95%						1 5%			0 0%			20 100%	
omissioni			compressioni			Tot (micro-unità)	S	NS	Tot (micro-unità)	umani		software	Tot (micro-unità)
27 67.5%			13 32.5%			40 100%	1 14.3%	6 85.7%	7 100%	R	E	2 40%	5 100%
sem.	non sem.	Tot	sem.	non sem.	Tot						2 40%	1 20%	
3 11.1%	24 88.9%	27 100%	1 7.7%	12 92.3%	13 100%								

Figura 14 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *Météo* – respeaker 1.

A conferma di quanto appena detto, questa è stata la prova migliore del rispeaker 1. Nonostante la perfetta suddivisione tra macro-unità rese (50%) e non rese (50%), ciò non ha implicato una perdita eccessiva di contenuto del TP. Infatti, solo il 31.25% delle omissioni sono semiotiche, cioè il 21.7% delle macro-unità non rese, contro il 68.75% di quelle non semiotiche, cioè il 47.8% delle macro-unità non rese. Le macro-unità omesse corrispondono soprattutto a passaggi in cui il ritmo dell'eloquio aumenta e la memoria a breve termine del rispeaker rischia di sovraccaricarsi oppure a informazioni inferibili dalle immagini (per esempio, le temperature sul territorio) e quindi ridondanti, che rischiano solamente di distrarre il telespettatore dalla componente video non verbale. Inoltre, il rispeaker tende a omettere quelle informazioni già contenute all'interno di macro-unità precedentemente rese e quindi inutili per la comprensione del testo. Un dato importante è sicuramente rappresentato dagli errori (30.4% delle macro-unità non rese), soprattutto imputabili al *software* (42.85%) e all'*editor* (42.85%). I motivi per questi risultati sono già stati ampiamente spiegati. Vediamo un esempio:

[01]

TP: Pour les températures, cet après-midi. /¹⁶¹ Alors, ça baisse par le nord-ouest avec l'arrivée d'air plus frais / amené par la perturbation.

TM1: le temperature questo pomeriggio ci sarà un abbassamento a causa di una perturbazione

TM2: le temperature questo pomeriggio ci sarà un abbassamento a casa la perturbazione

TA: Le temperature questo pomeriggio / si abbasseranno.

Di fronte a un evidente errore di riconoscimento da parte del *software* (in rosso), l'*editor* ha preferito cancellare il finale della frase omettendo un'informazione non inferibile dalle immagini, che riportano invece le temperature su tutto il territorio francese. In questo stesso esempio possiamo notare una espansione non semantica (in verde) e delle omissioni non semiotiche di micro-unità (in blu). Prendendo in considerazione le espansioni, queste rappresentano solo il 5% delle alterazioni (1 su 20), mentre riscontriamo in tutto 7 casi di espansione a livello di micro-unità, di cui solo una semantica, a conferma di quanto affermato nelle pagine precedenti. Il pochissimo tempo a disposizione, i tecnicismi e l'elevata densità lessicale non permettono una espansione in termini di caratteri e tantomeno l'aggiunta di informazioni. Le immagini parlano da sole e, anzi, in questo caso aiutano il rispeaker nel suo lavoro se perde o non capisce un passaggio del TP. Le sole espansioni possibili sono quindi quelle grammaticali e sintattiche, che rendono più coeso e scorrevole il testo e che sono naturali in interpretazione (per esempio, l'aggiunta del verbo nelle frasi nominali) e delle precisazioni quasi automatiche come la seguente:

[02]

TP: Tout cela va se renforcer. / 90 prévus cet après-midi sur les côtes

TA: Si rinforzeranno / a 90 chilometri all'ora sulle coste

A parte l'evidenza della compressione non semiotica (in arancione) e dell'omissione semiotica (in viola), il rispeaker ha automaticamente aggiunto l'unità di misura inferibile dal contesto (si parla di venti). Ciò accade anche in tutte le altre prove, a conferma dell'automaticità dell'operazione. Nonostante la notevole espansione in termini di caratteri, come in questo caso, le strategie di riduzione hanno sempre la meglio a livello di macro-unità e contano, infatti, per il 95% delle alterazioni, con una preponderanza di omissioni (67.5%) rispetto alle compressioni (32.5%), sempre a conferma di quanto già spiegato nelle pagine

¹⁶¹ Lo *slash* serve a delimitare le macro-unità.

precedenti. Le omissioni non semiotiche costituiscono l'88.9% perché coinvolgono perlopiù elementi tipici della lingua orale (che non si ritrovano nella lingua scritta) e informazioni ridondanti: notizie deducibili dalle immagini o dal contesto (*'par le nord-ouest'* – la presentatrice indica il nord-ovest con le mani – *'pour lundi'* – in alto a destra c'è la scritta *lundi*); intercalari e connettori (*'alors'*, *'et on voit que'*, *'également'*, *'donc'*); rafforzativi (*'très bon appétit'*); ripetizioni (*'beaucoup beaucoup'*). Omettere queste informazioni è fondamentale per la *performance* del rispeaker, per il carico cognitivo del telespettatore (inutilmente appesantito nella lettura) e per il *software*, che non può riconoscere certi tratti dell'oralità e che dovrebbe elaborare un testo non coeso e non lineare generando più errori. Le omissioni semiotiche, invece, contano solo per l'11.1%, ma sono importanti in quanto, cancellando un'informazione non recuperabile dal contesto, solitamente implicano una generalizzazione delle macro-unità di cui fanno parte, come nell'esempio appena citato [02]. Infatti, se il rispeaker non precisa che il vento arriverà a 90 km/h nel pomeriggio, la notizia rimarrà incompleta e, per il telespettatore, il vento potrebbe rafforzarsi, per esempio, durante la notte. I motivi per cui il rispeaker sceglie di omettere una micro-unità anziché un'altra possono essere diversi: insufficiente memoria a breve termine, ritmo di eloquio del TP troppo veloce, troppe informazioni importanti che si susseguono, eccessivo ritardo dei sottotitoli, inizio di un nuovo argomento o di una nuova sottofase per cui si 'taglia' l'ultima notizia, ecc. Per quanto riguarda le compressioni (32.5%), solo il 7.7% di queste sono semiotiche, cioè implicano una minima perdita di informazione pur non omettendo niente, mentre il restante 92.3% sono non semiotiche. Le compressioni si dividono in due micro-strategie: la compressione di singoli lessemi (sinonimia orizzontale e verticale, riformulazione e focalizzazione) e quella di intere locuzioni o frasi (riformulazione, sintesi, dislocazione e lessicalizzazione)¹⁶². Prendiamo un esempio di sinonimia:

[03]

TP: Les perturbations se succèdent.

TA: Le perturbazioni continuano.

Questa soluzione permette di risparmiare tempo nell'enunciazione del TM1 e caratteri nel sottotitolo senza, tuttavia, perdere il significato iniziale. Per quanto riguarda la compressione frastica, si è già fornito un esempio [02] in cui l'intera frase è stata resa con un singolo verbo. Sintetizzare il testo senza eliminare niente è molto utile non tanto per il rispeaker, che deve fare uno sforzo in più per riformulare il TP, quanto per il telespettatore che si ritrova a leggere un testo più breve, lineare e coeso, ma allo stesso tempo completo a

¹⁶² Cfr. Eugeni 2008b:172,175.

livello informativo, col minimo sforzo. Nel *respeaking* interlinguistico questo è ancora più vero perché la traduzione già *per se* implica una riformulazione, a meno che non si decida di tradurre letteralmente il testo (più semplice per il rispeaker) rendendolo più ostico al riconoscimento del *software* e alla lettura. In effetti, le uniche traduzioni letterali riscontrate (3 su 46) riprendono per lo più saluti ed espressioni formulaiche brevi e quindi facilmente riconoscibili (*'Bonjour'*, *'bon appétit'*, *'pour les températures'*), mentre quando il rispeaker ha tentato la traduzione letterale di passaggi più lunghi, questi si sono tramutati sistematicamente in errori. A proposito di errori, quello più frequente, riscontrato anche nelle altre prove, è stato il mal riconoscimento del sintagma *'est'*, confuso con *sta* [05] oppure con *destra* [04]. Mentre nel primo caso l'errore salta subito all'occhio (l'*editor* l'ha corretto nella seconda occorrenza – 05), nel secondo caso, la fortuna vuole che il senso non cambi di molto in quanto l'*est* si trova effettivamente *'a destra'*. Paradossalmente, questo non può considerarsi un errore a tutti gli effetti, tanto che l'*editor* non lo corregge:

[04]

TP: et puis on a des températures très douces, très estivales sur la partie sud-est.

TM1: e temperature estive nella parte sud est

TM2: e temperature estive dalla parte verso destra

TA: e temperature estive dalla parte verso destra

Più problematica la situazione in cui il rispeaker confonde *'est'* e *'ouest'* generando un errore più grave ma non percepibile da chi non conosce la lingua:

[05]

TP: En fait, ça reste un peu frais sur le nord-ouest

TM1: fresco al nord est

TM2: fresco al nord sta

TA: al nord est

Nel correggere il sottotitolo, l'*editor* ha cancellato anche *'fresco'* collegandolo in logica al sottotitolo precedente. A ogni modo gli errori riscontrati sono pochi (5) per via del basso numero di macro-unità rese e di un buon lavoro da parte dell'*editor* (solo 1 errore su 5).

Respeaker 2

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE						TOT
omissioni			errori			Tot	traduzioni letterali		alterazioni		Tot	46 100%
15 78.95%			4 21.05%			19 41.3%	3 11.1%		24 88.9%		27 58.7%	
sem.	non sem.	Tot	umani		software	Tot						
7 46.7%	8 53.3%	15 100%	R	E	2 50%	4 100%						
			1 25%	1 25%								

alterazioni													
riduzioni						espansioni			errori		Tot		
20 83.3%						3 12.5%			1 4.2%		24 100%		
omissioni			compressioni			Tot (micro-unità)	S	NS	Tot (micro-unità)	umani	software	Tot (micro-unità)	
31 57.4%			23 42.6%			54 100%	1 11.1%	8 88.9%	9 100%	R	E	2 25%	8 100%
sem.	non sem.	Tot	sem.	non sem.	Tot					4 50%	2 25%		
3 9.7%	28 90.3%	31 100%	4 17.4%	19 82.6%	23 100%								

Figura 15 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *Météo* – respeaker 2.

La prova del meteo non è la migliore del rispeaker 2. Forse perché è stata l'ultima e quindi sia la memoria del rispeaker, sia quella del *software* erano in sovraccarico. Ciò nonostante, si può notare un leggero miglioramento sia nella totalità delle macro-unità rese (8.7% in più rispetto al rispeaker 1), sia nel numero di macro-unità non rese per errore (9.35% in meno). Controcorrente, invece, la relazione tra le omissioni semiotiche e quelle non semiotiche a livello di macro-unità, rispettivamente il 46.7% e il 53.3%. Questo dato può essere spiegato dalla minore attenzione prestata alla componente video del TP da parte del rispeaker, che è stato più attento al *pre-editing* e ha quindi omesso informazioni non ritenute basilari per la ricezione del TA, ma al contempo non compensate dal contesto. Gli errori, però, non sono stati evitati del tutto, un po' per l'eccessiva velocità dell'*editor* che per recuperare il ritardo ha cancellato una frase intera (*schiarite nelle regioni del Centro e del sud*), un po' per colpa del *software* che addirittura non riconosce il saluto iniziale (*Buongiorno.*) e un po' per colpa del rispeaker stesso che in alcuni punti ostenta una sicurezza eccessiva ('costa' anziché 'Corsica', 'delle perturbazione'). Ancora una volta la traduzione letterale (3 su 46) si limita a saluti ed espressioni formulaiche, le stesse, tra l'altro, della prova precedente. Per quanto riguarda le alterazioni, si conferma l'alto tasso di riduzioni (83.3%),

accompagnato da un leggero aumento delle espansioni (7.5% in più) e un solo errore (4.2% delle alterazioni). Le due espansioni in più si sono registrate nel momento immediatamente precedente la proiezione di nuove immagini:

[06]

TP: Pour les températures, cet après-midi.

TA: Per quanto riguarda le temperature, nel pomeriggio

Il rispeaker, probabilmente, ha capito che le sole immagini sarebbero state sufficienti a veicolare l'informazione successiva e ne ha approfittato per esplicitare la frase. Infatti, il sottotitolo successivo si risolve in una compressione e due omissioni non semiotiche in quanto tutto ciò che non viene detto è inferibile dal contesto:

[07]

TP: Alors, ça baisse par le nord-ouest avec l'arrivée d'air plus frais

TA: ci sarà più fresco

A livello di micro-unità, le espansioni (solo 2 in più rispetto alla prova precedente) riguardano esplicitazioni (a 90 chilometri sulle coste), elementi di coesione (*Poi, ancora*) e aggiunta del verbo nelle frasi nominali per una maggiore scorrevolezza del testo:

[08]

TP: À l'avant, des pluies un petit peu plus éparées,

TA: Poi al nord abbiamo delle nuvole un po' più sparse,

In questo esempio vi sono anche 3 casi di compressione non semiotica. Mentre la compressione di '*petit peu*' in *po'* è più che naturale (*un pochino* sarebbe una forzatura e occuperebbe più spazio in termini di caratteri), la compressione di '*à l'avant*' e '*pluies*' sono evidentemente dovute a una riformulazione a partire dalle immagini (per questo non vi è perdita di informazione) perché forse il rispeaker non ha ben capito l'audio o non lo ricordava con quelle esatte parole. In generale, le omissioni (57.4%) e le compressioni (42.6%) delle micro-unità riflettono la tendenza a eliminare o riformulare una o più parti non rilevanti a livello semiotico all'interno delle macro-unità, generalizzandole e rendendole più facilmente accessibili. A ciò contribuisce anche il *software* che si rivela un vero e proprio collaboratore del rispeaker non riconoscendo elementi poco importanti al fine della comprensione del TA ('nuove perturbazioni' > *perturbazioni*; 'proteggerà le regioni del sud' > *proteggerà il sud*). Anche qui, per le omissioni non semiotiche (90.3%) si registra l'eliminazione di tratti dell'oralità e di informazioni ridondanti, mentre quelle semiotiche sono veramente limitate

(9.7%). Al contrario, il più alto numero di errori rispetto alla prova precedente è dovuto a un più alto numero di macro-unità rese e al non riconoscimento soprattutto di monosillabi ('e' > ne), preposizioni ('degli' > due e gli) e, a volte, della punteggiatura ('punto' > appunto, punto, pronto) che confondono l'editor. Degno di nota un caso più unico che raro in cui, pur riconoscendo un monosillabo anziché un altro (ovest al posto di 'est'), il software trasmette il giusto senso del TP:

[09]

TP: En fait, ça reste un peu frais sur le nord-ouest

TM1: resta fresco nel nord est

TM2: resta fosco nel nord-ovest

TA: resta fresco nel nord-ovest

Il software riconosce *fosco* invece di 'fresco' che, però, è corretto dall'editor.

Respeaker 3

MACRO-UNITÀ NON RESE							MACRO-UNITÀ RESE				TOT
omissioni			errori			Tot	traduzioni letterali		alterazioni	Tot	46 100%
10 66.7%			5 33.3%			15 32.6%	3 9.7%		28 90.3%	31 67.4%	
sem.	non sem.	Tot	umani		software	Tot					
2 20%	8 80%	10 100%	R	E	3 60%	5 100%					
			2 40%	1 20%							

alterazioni												
riduzioni						espansioni			errori		Tot	
23 82.15%						3 10.7%			2 7.15%		28 100%	
omissioni			compressioni			Tot (micro-unità)	S	NS	Tot (micro-unità)	umani	software	Tot (micro-unità)
28 48.3%			30 51.7%			58 100%	4 23.5%	13 76.5%	17 100%	R	E	8 100%
sem.	non sem.	Tot	sem.	non sem.	Tot					2 25%	3 37.5%	
2 7.14%	26 92.86%	28 100%	9 30%	21 70%	30 100%						3 37.5%	

Figura 16 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *Météo* – respeaker 3.

Il respeaker 3 ha il più alto numero di macro-unità rese (31 su 46) con solo il 13.3% delle omissioni semiotiche all'interno delle macro-unità non rese (20% delle omissioni), contro il

53.3% delle omissioni non semiotiche (80% delle omissioni). Purtroppo, ha anche il più alto tasso di errori (33.3%). In realtà, le macro-unità non rese per errore sono solo 5, tuttavia il loro peso è evidente all'interno della *performance*, per esempio:

[10]

TP: et on voit qu'elle arrive avec une formation de dépression ici au large de la Bretagne.

TM1: e vediamo che arrivano con una formazione di depressione qui al largo della Bretagna

TM2: e vediamo che era finita non con una formazione di depressione quella della Spagna

TA: Una formazione di depressione arriva dalla Spagna.

[11]

TP: Partout, d'ailleurs, c'est l'amélioration. / Un petit peu de vent en Méditerranée

TA: Miglioramenti nella zona del Mediterraneo

In quest'ultimo caso, chiaramente il rispeaker, forse in ritardo con l'enunciazione del TM1 e concentrato troppo sulla resa (rompendo l'equilibrio degli '*Efforts*'¹⁶³), ha percepito o ricordato solo le parole '*amélioration*' e '*Méditerranée*' generando un sottotitolo errato che automaticamente si traduce in due macro-unità non rese per errore. Nonostante l'elevato numero di errori in proporzione alle unità non rese (un terzo del totale), questi dati mostrano come il rispeaker 3 abbia cercato di eliminare il minimo indispensabile adottando soprattutto strategie di compressione e fusione frastica. Le traduzioni letterali, invece, sono sempre le stesse 3, mentre il numero di alterazioni aumenta (28), con un conseguente aumento delle riduzioni (23), 3 espansioni e 2 errori. Bisogna dire che le espansioni sono il punto forte del rispeaker 3 che cerca sempre di dare al testo una forma coesa e precisa anche quando non capisce il TP, però, in questo caso, il poco tempo a disposizione implica una preponderanza di riduzioni rispetto alle espansioni, che si ritrovano numerose solo a livello di micro-unità (17):

[12]

TP: On les regagne aujourd'hui avec des températures estivales

TM1: li riacquistiamo oggi con queste temperature quasi estive

TM2: lire riacquisti Monti con queste temperature quasi estive

TA: ma lo riacquistiamo sui monti con queste temperature quasi estive.

¹⁶³ Cfr. sottopar. 2.5.2.

Le prime due espansioni sono non semantiche, cioè aumentano il numero di caratteri senza introdurre informazioni nuove, mentre l'ultima (*quasi estive*) aggiunge un significato che non ritroviamo nel TP (*'estivales'*) e può essere quindi classificata come espansione semantica (in verde chiaro). Un ottimo esempio di riduzione è il seguente:

[13]

TP: En revanche, cette partie sud-ouest, ca donne de la grisaille avec quelques gouttes

TA: Alcune gocce qui,

Il rispeaker ha decisamente ridotto la macro-unità in questione con tre semplici parole, applicando in due casi la strategia dell'omissione non semiotica e in due casi la strategia della compressione non semiotica, contando sul supporto delle immagini e quindi sull'immediata trasmissione del sottotitolo in modo tale che il '*qui*' non corrisponda alle indicazioni di un'altra zona sulla cartina. Il '*grigiore*', deducibile dal contesto, è ritenuto poco importante dal rispeaker, che sceglie di evidenziare invece il fatto che farà qualche goccia. Omissioni (48.3%) e compressioni (51.7%) sono utilizzate pressoché in egual modo, raggiungendo un equilibrio impressionante. All'interno delle omissioni ovviamente prevalgono quelle non semiotiche (92.86%), mentre le omissioni semiotiche sono solo due (7.14%). All'interno delle compressioni l'andamento è lo stesso, ma con una leggera variazione: compressioni semiotiche per il 30% e compressioni non semiotiche per il 70%. Vediamo un esempio di compressione semiotica (in marrone):

[14]

TP: Poches de soleil sur la Méditerranée, également la Corse et aussi sur cette partie du nord-est.

TA: Sole sul Mediterraneo e in Corsica così come in questa parte del nord est.

'degli sprazzi di sole' è stato reso semplicemente con *sole*, il che implica una minima perdita di informazione. Infine, gli errori sono imputabili perlopiù al *software* e all'*editor*. Questo fatto può essere dovuto all'eccessiva quantità di informazioni dettate al *software* seppur a un ritmo contenuto. Infatti, il rispeaker 3, eliminando poche macro-unità, ha accumulato un numero sempre crescente di sottotitoli nel *buffer*, sovraccaricando la memoria del *software* (e quindi favorendo gli errori da parte di questo), aumentando il lavoro da parte dell'*editor* (che non può ricordarsi perfettamente tutto ciò che si sente del TM1) e perciò aumentando anche il ritardo nella trasmissione dei sottotitoli. Un esempio già riportato [12] mostra, infatti, come un errore di riconoscimento (*lire riacquisti Monti* invece di 'li

riacquistiamo oggi') induca l'*editor* a correggere la frase con la forma logica più vicina e coerente col sottotitolo precedente, ma scorretta (*lo riacquistiamo sui monti*).

4.2.2 *Télématin_JT*

Sicuramente più difficile da sottotitolare del meteo per il ruolo marginale delle immagini, l'importanza della componente audio verbale, la varietà di argomenti trattati e la presenza di servizi preregistrati, questa trasmissione, anche se breve, ha posto delle notevoli difficoltà a tutti i rispeaker. Infatti, i risultati sono chiari: le macro-unità non rese (70.6%, 56.5% e 60%) sono di gran lunga superiori alle macro-unità rese (29.4%, 43.5% e 40%) e il tasso di omissioni semiotiche è molto alto rispetto a prima (50%, 53.85% e 45.7%), mentre quello degli errori è simile (36.7%, 18.75% e 31.4%). L'unica spiegazione possibile è che tutti i rispeaker abbiano preferito omettere le macro-unità non capite piuttosto che creare sottotitoli sbagliati e veicolare informazioni erranee, limitandosi giusto a dare la notizia principale senza troppi particolari. In particolare, il telegiornale parla di quattro notizie: l'arresto di Dominique Strauss-Kahn, la scomparsa di tre persone, l'intossicazione di sette bambini e il cambio di direzione presso la società *Aréva*. Mentre le due notizie di portata internazionale, cioè l'arresto di Dominique Strauss-Kahn e l'intossicazione, sono state rese più facilmente (il servizio su Strauss-Kahn è inoltre supportato dalla presenza di didascalie che riportano ciò che viene detto), le altre due notizie (e soprattutto l'ultima) hanno creato grandi difficoltà. La prima perché, non essendoci un servizio preregistrato dopo, non c'è stato neanche modo di capire e spiegare meglio la notizia dopo una prima introduzione; la seconda, trovandosi alla fine del telegiornale, probabilmente non è stata ben capita perché la memoria e l'attenzione dei rispeaker iniziavano a risentire della fatica, oltre alla presenza di diversi nomi propri e alle lacune degli stessi rispeaker su tale argomento. Infatti, conoscere un argomento è fondamentale per il rispeaker in quanto si può destreggiare meglio tra nomi propri, tecnicismi e successione degli avvenimenti. In ogni caso, il fatto che quasi ogni servizio preregistrato fosse anticipato da un'introduzione alla notizia ha perlomeno permesso di veicolare le informazioni essenziali grazie a una notevole riformulazione del testo (96%, 97.3% e 94.1% di alterazioni all'interno delle macro-unità rese). Siccome i rispeaker non avevano elementi sufficienti per dare informazioni in più rispetto al TP e siccome il TP è formato da testi scritti in precedenza e letti oralmente, le espansioni risultano in una percentuale bassa (16.7%, 19.4% e 15.625%), mentre le riduzioni sono ancora una volta molto alte (70.8%, 80.6% e 84.375%). Il discostamento del primo rispeaker rispetto agli altri è dovuto semplicemente alla presenza di errori (12.5%). Inoltre, si nota come tutti i rispeaker abbiano tralasciato le

espressioni formulaiche di introduzione ai servizi preregistrati e le informazioni contenute nei sottopancia degli stessi (per esempio, il nome dei *reporter*). Infine, la presenza dell'interprete della *LSF* non è di grande aiuto, a meno che il telespettatore la conosca.

Respeaker 1

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE			TOT
omissioni			errori			Tot	traduzioni letterali	alterazioni	Tot
38			22			60	1	24	25
63.3%			36.7%			70.6%	4%	96%	29.4%
sem.	non sem.	Tot	umani		software	Tot			
19	19	38	R	E	16	22			
50%	50%	100%	6	0	72.7%	100%			
			27.3%	0%					

alterazioni									
riduzioni					espansioni		errori		Tot
17					4		3		24
70.8%					16.7%		12.5%		100%
omissioni		compressioni		Tot (micro-unità)	S	NS	Tot (micro-unità)	umani	software
18		20		38	4	6	10	R	E
47.4%		52.6%		100%	40%	60%	100%		
sem.	non sem.	Tot	sem.	non sem.	Tot				
7	11	18	6	14	20	4	2	10	16
38.9%	61.1%	100%	30%	70%	100%	25%	12.5%	62.5%	100%

Figura 17 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *Télématin_JT* – respeaker 1.

Innanzitutto, l'*output* di questa prova non è utilizzabile a causa dell'alto numero di errori riscontrato (25 a livello di macro-unità e 16 a livello di micro-unità) e di omissioni semiotiche (19 a livello di macro-unità e 7 a livello di micro-unità) per niente compensate dalle espansioni (4 a livello di macro-unità e 4 espansioni semantiche a livello di micro-unità). Tutto ciò significa che le macro-unità senza nemmeno un errore sono 12 su 85, di cui 2 sono il saluto iniziale (l'unica traduzione letterale) e quello finale, 7 sono riduzioni (5 compressioni non semiotiche e una semiotica) e 2 espansioni non semantiche (*che cosa è successo* invece di *cos'è successo*; *più di 30 morti* invece di *39 morti*). Gli errori sono tra i più vari della categoria: *resaconto* ('*compte rendu*'), 'immunità demografica' ('*immunité diplomatique*'), il vecchio *Presidente* ('*ancien*'), la gente ('*le policier*'), che hanno / *che sono scomparse* / *facendo il canone in* ('*emportées par une vague / alors qu'elles faisaient du canyoning en Savoie.*'). Le 'colpe' sono da attribuirsi al respeaker che ha mostrato troppe esitazioni ed è

stato poco attento nella resa del TM1 (poteva dire ‘il poliziotto’ o ‘la polizia’ invece di ‘l’agente’, mentre ‘immunità demografica’ è un *lapsus*); al *software* che ha generato un numero ingestibile di errori (*canyonig* è diventato *canone in* anche se la parola era stata aggiunta nel vocabolario); infine l’*editor* non è riuscito a districarsi in mezzo a tanti errori, inviando sottotitoli del tipo *che hanno che sono scomparse facendo (...)* o modificando solo in parte gli errori di riconoscimento del *software* (*resa conto* è diventato *resaconto*). Per il poco che rimane, si possono notare numerose compressioni (52.6%) al fine di ridurre all’osso l’alta densità lessicale tipica del telegiornale:

[15]

TP: Sept enfants sont toujours hospitalisés ce matin à Lille. / Tous sont soignés pour une intoxication

TA: Sette bambini sono stati curati per intossicazione

Al di là dell’evidente calco semantico (*‘soignés’ > curati*), si può notare la fusione delle due frasi con una perdita di informazione: i bambini sono ancora ricoverati, mentre la macro-unità resa rimane generica. La perdita di informazione ‘à Lille’ è invece recuperata dal sottopancia durante il servizio. Le omissioni in effetti non sono molte meno (47.4%), ma si caratterizzano per una percentuale inferiore di omissioni non semiotiche (61.1%) dovuta al fatto che i testi preparati in precedenza e il poco tempo a disposizione non lasciano spazio ai tratti dell’oralità.

Respeaker 2

MACRO-UNITÀ NON RESE				MACRO-UNITÀ RESE			TOT
omissioni			errori			Tot	85 100%
39 81.25%			9 11.25%			48 56.5%	
sem.			umani			traduzioni letterali	37 43.5%
non sem.			software			alterazioni	
Tot			Tot			Tot	36 97.3%
21 53.85%	18 46.15%	39 100%	R	E	3 33.3%	1 2.7%	
			6 66.7%	0 0%	9 100%		

alterazioni										
riduzioni						espansioni			errori	
29 80.6%						7 19.4%			0 0%	
Tot (micro-unità)			Tot (micro-unità)			Tot (micro-unità)			Tot (micro-unità)	
omissioni			compressioni			S NS			umani software	
39 60%			26 40%			5 15 25% 75%			1 2 50% 100%	
sem.	non sem.	Tot	sem.	non sem.	Tot	R E			1 0 50% 0%	
13 33.3%	26 66.6%	39 100%	10 38.5%	16 61.5%	26 100%					

Figura 18 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *Télématin_JT* – *respeaker 2*.

Nonostante l'alto numero di omissioni semiotiche a livello di macro-unità (53.85% contro il 46.15% delle omissioni non semiotiche), il rispeaker 2 ha trasmesso tutte le notizie principali senza entrare troppo nei particolari, ma restituendo un TA accettabile per il pubblico italiano. Questa è, infatti, la migliore tra le tre prove. L'alto numero di errori (9) a livello di macro-unità è dovuto a una poca concentrazione sul TP per uno sforzo maggiore nella produzione del TM1 e nel controllo del TM2 (il 66.7% degli errori sono del rispeaker). Gli errori, infatti, non portano a sottotitoli insensati o incompleti, bensì a sottotitoli lineari e coerenti col contesto che però non dicono quello che si dice nel TP:

[16]

TP: La catastrophe de Fukushima l'avait même remise sur le devant de la scène,

TA: La catastrophe di Fukushima insegna.

È probabile che il rispeaker abbia sentito solo le parole iniziali e non volendo perdere tempo ha completato la frase restando sul generico per dedicarsi direttamente alla macro-unità successiva. Infatti, è stato l'unico a rendere perfettamente quest'ultima:

[17]

TP: mais l'État lui a préféré Luc Oursel, l'homme du compromis.

TA: Lo Stato ha preferito a lei Luc Oursel.

Bisogna ricordare che si tratta di una traduzione e il ritmo di eloquio è sostenuto (180 wpm) per cui il lavoro non è semplice e il fatto che i sottotitoli siano comunque comprensibili e rendano correttamente l'informazione è da considerarsi già un buon risultato. Le macro-unità rese sono tutte alterazioni, tranne una che non viene modificata sicuramente per mancanza di tempo e per un migliore riconoscimento da parte del *software*. Un alto numero di

riduzioni (80.6%) conferma la politica di selezione sempre per mancanza di tempo e al fine di semplificare il più possibile il TA, mentre le espansioni sono poche a causa degli stessi motivi (19.4%). Degno di nota è l'esempio seguente:

[18]

TP: mais cette infection n'aurait pas de liens avec l'épidémie / qui a déjà fait 39 morts.

TA: Questo però non avrebbe niente a che vedere con l'epidemia di Escherichia coli.

Il rispeaker evidentemente aveva una reminiscenza del 'batterio Escherichia coli' presente nel glossario oppure delle notizie sentite al telegiornale italiano mesi prima. Ha specificato quindi il nome dell'epidemia senza che il TP fornisse informazioni specifiche a riguardo (il batterio E-coli viene citato dopo). Si tratta perciò di una espansione semantica a tutti gli effetti. Interessante notare anche che, per precisare il tipo di epidemia, si è omessa un'informazione importante sulla gravità dell'epidemia stessa. Ciò non risulta essere anormale dal momento in cui i numeri e i nomi propri (come anche in interpretazione) sono gli elementi più difficili da tenere in mente e perciò è più semplice ripeterli subito o ometterli, come in questo caso, o generalizzare (il rispeaker 1 ha reso la frase in questo modo: *che ha già fatto più di 30 morti*). A proposito di nomi propri, il rispeaker 2, citati i nomi una prima volta, successivamente li omette o li comprime utilizzando il pronome soggetto oppure usa altri appellativi facilmente riconoscibili ('*Henri Proglio, patron d'EDF*' > *il capo di EDF*). C'è una sola eccezione che si traduce in una esplicitazione tramite l'uso dell'iperonimia:

[19]

TP: DSK: Que se passe-t-il?

TA: Dominique Strauss-Kahn. Che succede?

Il motivo è semplice: il *software* e l'*editor* non avrebbero riconosciuto *DSK* e probabilmente il '*tag*' non sarebbe comparso nel sottotitolo finale. Una scelta del tutto arbitraria del rispeaker è invece quella di espandere la sigla FMI (*fondo monetario internazionale*), riconosciuta perfettamente nelle altre prove e comprensibile da parte del telespettatore medio. Il rispeaker forse ha pensato che il *software* non avrebbe riconosciuto la sigla e ha preferito non rischiare. Le compressioni contano per il 40%, le omissioni per il 60% e la loro suddivisione interna rispecchia una politica di selezione volta a tagliare e/o eliminare le informazioni meno importanti. Gli errori, infine, forniscono sempre gli esempi più curiosi: 'in aereo' diventa *il reo*; 'l'uomo protesta' *ha luogo una protesta*; 'Aréva' *aveva*. Ora, per

quanto questi siano tutti errori di riconoscimento, la loro perfetta collocazione all'interno del sottotitolo non ci consente di considerarli come tali perché il senso trasmesso è quello giusto: il reo si è fatto fermare dalla Polizia (si può intendere Dominique Strauss-Kahn e può essere una compressione); ma in Commissariato ha luogo una protesta (può essere considerata una espansione, ma non c'è nulla di sbagliato); perché aveva uno dei più grandi gruppi del mondo (sottotitolo lineare e sensato con soggetto sottinteso, cioè il capo di Aréva). Una spiegazione plausibile a questo fenomeno è che il *software*, se non riconosce un termine, attraverso un'analisi semantico-sintattica e di frequenza delle occorrenze, sceglie la parola più adatta tenendo conto di una o più parole successive e precedenti al termine in questione. Perciò, un errore non sfocia mai in una 'non parola' (a eccezione dei monosillabi o di un'eccessiva velocità nell'eloquio, in qual caso il *software* elabora stringhe di testo e non singole parole), ma in omofoni o termini simili che si adattano meglio al 'mini-contesto' creato da 2 o 3 parole della frase.

Respeaker 3

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE				TOT	
omissioni			errori			Tot	traduzioni letterali		alterazioni	Tot	85 100%
35 68.6%			16 31.4%			51 60%	2 5.9%		32 94.1%	34 40%	
sem.	non sem.	Tot	umani		software	Tot					
16 45.7%	19 54.3%	35 100%	R	E	12 75%	16 100%					
			2 12.5%	1 6.3%							

alterazioni										Tot	
riduzioni					espansioni			errori			
27 84.375%					5 15.625%			0 0%		32 100%	
omissioni			compressioni		Tot (micro-unità)	S	NS	Tot (micro-unità)	umani	software	Tot (micro-unità)
29 46%			34 54%		63 100%	3 13.6%	19 86.4%	22 100%	R	E	
sem.	non sem.	Tot	sem.	non sem.	Tot				0 0%	2 50%	4 100%
10 34.5%	19 65.5%	29 100%	12 35.3%	22 64.7%	34 100%						

Figura 19 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *Télématin_JT* – respeaker 3.

La prova del rispeaker 3 è da considerarsi una via di mezzo tra le altre due. Accettabile nella prima parte, mostra un calo eccessivo nella seconda, a un punto tale da renderla

inutilizzabile. L'ultima notizia non viene resa in alcun modo: è qui che si concentrano molte omissioni e molti errori a livello di macro-unità. Il rispeaker non capisce il fulcro della notizia, cerca di tradurre letteralmente il testo, ma esita e ripete le stesse cose. Di conseguenza, il *software* accumula errori e l'*editor*, già in ritardo, non riesce a tenere il passo.

[20]

TP: Depuis des mois, Anne Lauvergeon faisait campagne pour un troisième mandat.

TM1: e delle campagne che abbiamo alle quali abbiamo assistito nel corso degli ultimi mesi

TM2: e di delle campagne che abbiamo alle spalle per l'assistito nel corso degli ultimi mesi

TA: e di delle campagne che abbiamo alle spalle per l'assistito nel corso degli ultimi mesi

L'errore più grave è che l'*editor* ha inviato il sottotitolo, ma trovandosi praticamente alla fine della trasmissione è probabile che abbia deciso di accelerare non apportando più modifiche e salvando solo quelle parti di testo più 'logiche'. Nella confusione generale che un testo così ricco di informazioni può creare, il rispeaker non riesce a distinguere per bene quali siano le macro-unità non rilevanti semioticamente, per cui, anche in questa prova, le omissioni semiotiche sono numerose sia a livello di macro-unità (45.7% contro il 20% del meteo), sia a livello di micro-unità (34.5% contro il 7.14% del meteo). A parte due traduzioni letterali, le restanti macro-unità sono alterate con una preponderanza (non a sorpresa) di riduzioni (84.375%). Numerose le compressioni (54%) che implicano perlopiù delle generalizzazioni delle macro-unità:

[21]

TP: Le parquet de New York vient de publier le compte rendu détaillé / de tout ce qui s'est passé le 14 mai dernier.

TA: Da New York hanno pubblicato un resoconto / di tutto quanto è accaduto.

Le espansioni semantiche non sono tante (13.6%) perché non c'è tempo per aggiungere informazioni e, per di più, è già tanto se le informazioni veicolate dal TP vengono trasmesse. Predominano invece le espansioni non semantiche (86.4%) in quanto le poche cose assimilate dal rispeaker sono ripetute e spiegate in maniera precisa e lineare. Per quanto riguarda gli errori, l'unico caso degno di nota è uno sbaglio nell'*editing*:

[22]

TP: L'ancien directeur du FMI a préféré garder le silence d'après le conseil de son avocat. / Accusé de viol, / il a plaidé non coupable le 6 juin dernier. / En France, les recherches ont repris ce matin pour tenter de retrouver les trois personnes

TA: e su consiglio del suo avvocato ha preferito restare in silenzio. / Dopo l'accusa di violenza / questa mattina sono continuate le ricerche per cercare le tre persone

Il rispeaker non ha capito o non ricorda il finale della frase e decide di non proseguire per dedicarsi alla notizia immediatamente successiva. *Dopo l'accusa di violenza*, tuttavia, sarebbe potuta essere la degna conclusione della frase precedente se l'*editor* non avesse inserito un punto. Così facendo, è diventato un sottotitolo a sé in cui il telespettatore nota perfettamente la mancanza di un pezzo, soprattutto perché dopo appare un sottotitolo (senza punteggiatura e senza maiuscola) con la nuova notizia. Infine, è da notare il fatto che i rispeaker iniziano con una traduzione più attaccata al TP, continuano con delle strategie di riformulazione (sia compressione, sia espansione) e finiscono con delle omissioni. Questo succede per tutti e tre i rispeaker. La ragione può risiedere nel fatto che all'inizio la notizia è nuova e quindi il rispeaker segue da vicino il testo; una volta capito di cosa si parla, riformula il testo con parole proprie e lo rende in un buon italiano; alla fine, tralascia le ultime frasi perché nel frattempo il giornalista è passato a una nuova notizia e, per non dimenticarsi l'inizio, ricomincia la traduzione più letterale e così di seguito. Perciò le omissioni semiotiche (sia di macro- che di micro-unità) si trovano quasi sempre in posizione finale.

4.2.3 *Comment ça va bien ?_LA DEMO !*

Questo programma è decisamente diverso rispetto ai precedenti. Siamo partiti dal meteo, in cui una persona legge un testo scritto complementare alle immagini di fondo e semioticamente molto ridondante. Siamo passati poi al telegiornale, in cui c'è più di un locutore, ma i turni di parola sono prevedibili e non si sovrappongono. Anche in questo caso il testo era scritto per essere letto. In entrambi i casi la densità lessicale e informativa era tale da non permettere grandi modifiche e il registro utilizzato era formale. *Comment ça va bien ?* invece è un programma di intrattenimento, in cui c'è più di un locutore, i turni spesso si sovrappongono e sono imprevedibili. Per di più, le immagini sono molto importanti nella

sottofase in cui viene spiegato come si costruisce l'oggetto del giorno¹⁶⁴ perciò i sottotitoli devono avere un ritardo minimo nella trasmissione per essere di supporto alle stesse immagini. Il ritmo di eloquio, infine, è abbastanza elevato (198 wpm), le battute sono brevi e lessicalmente povere. Tutto ciò ha delle conseguenze importanti sui dati della ricerca. Da una parte cresce a dismisura il numero delle macro-unità non rese (83.9%, 72.1% e 74.7%) e in particolare delle omissioni non semiotiche (87.5%, 86.5% e 85.6%) perché la maggior parte delle macro-unità è composta da numerosi tratti del parlato ('*Ah bon?*', '*D'accord*'), ripetizioni di quanto dice il cronista ('– (...) *on va fabriquer aujourd'hui une étagère pour enfants. – Une étagère pour enfants, d'accord.*') ed esclamazioni ('*C'est mignon!*'), tutti elementi semioticamente irrilevanti; dall'altra cresce il numero delle espansioni (25.4%, 22.8%, 32.32%) perché il registro informale dà il tempo al rispeaker di aggiungere più elementi ed essere più chiaro nella spiegazione dei tecnicismi. Le compressioni, invece, subiscono un netto calo (38.1%, 37.7% e 33.1%) a causa della brevità delle battute. Infatti, se prima il rispeaker aveva bisogno di ricorrere soprattutto alla fusione per rendere un concetto all'interno di uno o più sottotitoli, adesso non si ha più un discorso continuo da poter sintetizzare. A non permettere la compressione frastica o lessicale è anche l'utilizzo di tecnicismi per la costruzione dell'oggetto del giorno. La bassa densità lessicale e la preponderanza semiotica delle immagini rispetto alla componente audio verbale è visibile dalle percentuali riguardanti le omissioni semiotiche sia a livello di macro-unità (12.5%, 13.5%, 14.4%), sia a livello di micro-unità (3.85%, 6.25% e 1.8%). Invece, la forte presenza di tratti dell'oralità tipici di un linguaggio colloquiale fa aumentare a dismisura le omissioni non semiotiche a livello di micro-unità (96.15%, 93.75% e 98.2%). Se da una parte le immagini richiedono la realizzazione di sottotitoli in un tempo veramente limitato, dall'altra queste aiutano il rispeaker nel suo lavoro, suggerendo a volte la soluzione più corretta se non si sente l'audio oppure permettendo l'eliminazione di una o più frasi senza grande perdita di informazione.

¹⁶⁴ Cfr. sottopar. 3.5.3.

Respeaker 1

MACRO-UNITÀ NON RESE							MACRO-UNITÀ RESE				TOT
omissioni			errori			Tot	traduzioni letterali		alterazioni	Tot	391 100%
280 85.4%			48 14.6%			328 83.9%	0 0%		63 100%	63 16.1%	
sem.	non sem.	Tot	umani		software	Tot					
35 12.5%	245 87.5%	280 100%	R	E	34 70.8%	48 100%					
			7 14.6%	7 14.6%							

alterazioni													
riduzioni						espansioni			errori			Tot	
42 66.7%						16 25.4%			5 7.9%			63 100%	
omissioni			compressioni			Tot (micro-unità)	S	NS	Tot (micro-unità)	umani		software	Tot (micro-unità)
52 61.9%			32 38.1%			84 100%	7 25.9%	20 74.1%	27 100%	R	E	4 26.7%	15 100%
sem.	non sem.	Tot	sem.	non sem.	Tot					5 33.3%	6 40%		
2 3.85%	50 96.15%	52 100%	2 6.25%	30 93.75%	32 100%								

Figura 20 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *Comment ça va bien ?_LA DEMO !* – respeaker 1.

Anche se il numero di macro-unità non rese è veramente elevato (328 su 391), così come il numero di errori (53 a livello di macro-unità e 15 a livello di micro-unità), questa *performance* può essere utilizzata. I concetti espressi sono corretti, per cui la maggior parte degli errori a livello di macro-unità non si sono tramutati in sottotitoli, ma sono stati cancellati dall'*editor* oppure non sono stati riconosciuti dal *software*. Tutti i tentativi di traduzione letterale, invece, sono sfociati in errori di riconoscimento, tanto che non vi è nemmeno un esempio di traduzione letterale corretta:

[23]

TP: Exactement, c'est une petite feuille de plastique un peu rigolote

TA: con una foglia di plastica

Per questo motivo, tutte le macro-unità rese sono alterazioni. Ovviamente, le riduzioni sono le più numerose (42 su 63) con una prevalenza delle omissioni, come è già stato spiegato, sulle compressioni e delle riduzioni non semiotiche sulle semiotiche (2 su 52 per le omissioni e 2 su 32 per le compressioni). In più, il numero di espansioni (7 semantiche e 20

non semantiche) compensa di gran lunga le poche omissioni e compressioni semiotiche. Le macro-unità rese si limitano a dire l'essenziale e contengono perlopiù le battute del presentatore e del cronista. Non è riportata, invece, la maggior parte dei commenti puntigliosi, scherzosi e sarcastici degli altri locutori. Questo perché le voci si sovrappongono, vengono da fuori quadro e non sono da ritenersi importanti rispetto a quello che è l'oggetto della trasmissione. Non ci sono esempi particolarmente eclatanti, il rispeaker ha tutto il tempo necessario per riformulare adeguatamente le frasi adattandole ai sottotitoli. Si può riportare, però, un esempio interessante di fusione che finora non è stata mai incontrata:

[24]

TP: - Vous, vous les collez pas? /

- Non, non, non.

TA: - Non sono incollati nel legno.¹⁶⁵

In un sottotitolo sono riassunte due battute (domanda e risposta) di due persone diverse, veicolando l'informazione in modo corretto e addirittura espandendola con una esplicitazione (*nel legno*). Se in un film ciò sarebbe suonato strano, nella sottotitolazione in tempo reale ciò è permesso in quanto permette di risparmiare tempo e spazio col minimo sforzo e perché, non essendoci perfetta sincronia con le immagini, non compromette la semiotica del TP. Nello stesso esempio troviamo anche due omissioni non semiotiche, a conferma della tendenza a eliminare i tratti dell'oralità (due ripetizioni in questo caso). Un altro caso interessante riguarda una espansione semantica:

[25]

TP: Là, vous avez suivi, jusque-là? / Formidable!

TA: Se non avete seguito / ci sarà un riassunto sul sito.

In questo passaggio il presentatore aumenta improvvisamente il ritmo di eloquio pronunciando, nel seguito, una serie di misure (*'C'est 10x20x15x...'*) prima di essere interrotto dal cronista. Evidentemente, il rispeaker ha iniziato a riformulare la frase e poi ha finito nel modo più idoneo ignorando il sarcasmo del presentatore, che mal si adattava al completamento della frase. La cosa curiosa è che il rispeaker ha anticipato qualcosa che effettivamente si dirà a fine trasmissione. Le ragioni possono essere due: o è stato fortunato oppure si era già informato sulla trasmissione e sapeva che su Internet ci sono i riassunti delle puntate del programma e le istruzioni per costruire gli oggetti in questione. Quanto agli errori a livello di micro-unità, sono due i più significativi. Il primo:

¹⁶⁵ I trattini sono stati aggiunti a fini esplicativi.

[26]

TP: Mais quelque part, avec le système d'attache et la vague, (...)

TM1: però col sistema di attaccamento delle onde

TM2: però sistema di attaccamento delle donneTA: Con lunil sistema di attaccamento delle onde (...)

Nel tentativo di inserire l'articolo 'il', l'*editor* ha sbagliato, premendo troppi tasti in contemporanea: la 'u' e la 'l' sono vicine alla 'i' nella tastiera, un po' meno logico è l'inserimento di una 'n', a meno che l'*editor* non fosse indeciso sull'uso dell'articolo determinativo o indeterminativo e li ha messi entrambi dimenticando di cancellarne uno. D'altronde, nella fretta, l'*editor* lascia passare altri errori (sul sito www.Fra www.france2.fr). Il secondo:

[27]

TP: François, fais gaffe, tes doigts.TA: Francesco attento!

Non si può considerare come macro-unità non resa e alcuni non lo classificherebbero nemmeno come un errore (il telespettatore non ci farebbe caso, anzi, forse sarebbe contento di non ritrovarsi davanti un nome francese). Il rispeaker, concentrato nella traduzione, ha tradotto in automatico anche il nome proprio. Questo è un caso isolato, prima di tutto perché i nomi propri non si traducono, in secondo luogo perché solitamente i nomi sono omessi o sostituiti da pronomi soggetto o complemento.

Respeaker 2

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE			TOT
omissioni			errori			Tot	traduzioni letterali	alterazioni	Tot
274 97.2%			8 2.8%			282 72.1%	8 7.3%	101 92.7%	109 27.9%
sem.	non sem.	Tot	umani		software	Tot			
37 13.5%	237 86.5%	274 100%	R	E	1	8			
			5 62.5%	2 25%	1 12.5%	8 100%			

alterazioni										
riduzioni						espansioni			errori	
75 74.2%						23 22.8%			3 3%	
Tot (micro-unità)						Tot (micro-unità)			Tot (micro-unità)	
omissioni			compressioni			S	NS		umani	software
96 62.3%			58 37.7%			14	29	43	R	E
100%			100%			32.56%	67.44%	100%	5	1
sem.	non sem.	Tot	sem.	non sem.	Tot				55.6%	11.1%
6	90	96	17	41	58				3	9
6.25%	93.75%	100%	29.3%	70.7%	100%				33.3%	100%

Figura 21 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *Comment ça va bien ?_LA DEMO !* – respeaker 2.

Anche questa volta, la prova del rispeaker 2 si è rivelata la migliore, nonostante alcune parti non propriamente fedeli al TP e alcuni errori di senso. Anche la prova del rispeaker 3 è buona, se non fosse per alcune macro-unità rese e fedeli al testo, ma poco comprensibili e di difficile lettura (non a caso gli errori a livello di macro-unità sono più del triplo rispetto al rispeaker 2). Per cui, la sicurezza dimostrata dal rispeaker 2, anche se a volte porta a soluzioni sbagliate, è utile al fine di creare un TA coeso e fruibile da parte del telespettatore. I dati mostrati nelle tabelle confermano quanto si è detto nella presentazione del programma, perciò si hanno una preponderanza delle omissioni non semiotiche (237 su 274) su quelle semiotiche (37 su 274), un alto numero di riduzioni (75 su 101 alterazioni), ma anche un numero più elevato di espansioni rispetto alle prove precedenti (23 su 101 alterazioni) e solo 8 errori su 282 macro-unità non rese e 3 errori su 101 alterazioni. Inoltre, salta subito all'occhio il numero di traduzioni letterali corrette (8 rispetto allo 0 degli altri due rispeaker) contenenti frasi brevi e semplici. A differenza del rispeaker 1, il rispeaker 2 cerca anche di trasporre le battute e i commenti sarcastici, a volte con successo, a volte con una soluzione infelice, come nel caso seguente:

[28]

TP: - Aux Pâques ou à la Noël. /

- Ou à la Trinité.

TA: - A Pasqua o all'Epifania.

I personaggi fanno un gioco di parole a partire dal termine 'opaque' (in riferimento al colore del polipropilene usato per costruire la mensola, oggetto del programma), omofono di 'Aux Pâques'. In italiano, lo stesso gioco di parole non è possibile e non c'è il tempo sufficiente per pensare a un'alternativa valida, perciò sarebbe meglio ignorare la battuta

altrimenti si rischia di creare un sottotitolo insensato come in questo caso. Il telespettatore che non conosce il francese si chiederà, infatti, cosa hanno a che fare la Pasqua e l'Epifania con l'oggetto che si sta costruendo. Di grande effetto, invece, la battuta sul negozio svedese. Il presentatore ironizza sul fatto che l'oggetto che costruisce il cronista è venduto anche da una famosa azienda svedese e alla fine della rubrica dice:

[29]

TP: - Alors, la, la, ce qui est intéressant c'est que ça revient, ça revient encore moins cher. /

- Oui /

- En plus, au prix où est l'essence, / on n'est plus d'aller, obligé d'aller à la boutique suédoise. / Et, et donc c'est encore moins cher.

TA: La cosa interessante / è che non c'è più bisogno di andare da IKEA, / quindi costa ancora meno.

A parte l'ottimo esempio di compressione frastica non semiotica, il famoso negozio svedese è diventato l'IKEA. Il telespettatore italiano capisce al volo la battuta col minor sforzo possibile. Sfortunatamente, poche macro-unità dopo sbaglia il prezzo del materiale che da 12 euro passa a 2 euro. In questo stesso esempio notiamo una omissione semiotica e una omissione non semiotica a livello di macro-unità e quattro omissioni non semiotiche che riguardano un tratto dell'oralità ('*alors*') e tre ripetizioni ('*la, la*', '*ça revient, ça revient encore moins cher*', '*et, et*') a livello di micro-unità. Le omissioni non semiotiche a livello di micro-unità sono tantissime (90 su 96 omissioni totali) a causa dell'elevato numero di tratti tipici della lingua orale: ripetizioni ('*Il faut être un peu bricoleur.* / - *Oui, il faut être un tout petit peu bricoleur, quand même.* / - *Il faut être un peu bricoleur.*'), connettori ('*Donc*', '*et puis*'), esitazioni ('*pol...polypropilène*', '*vous, vous les collez pas*', '*ça, ça, c'est*'), esclamazioni ('*Bien sûr!*', '*Bravo!*'), false partenze ('*Est-ce que...alors, montrez-moi.*'), intercalari ('*quand même*', '*bon*'), frasi interrotte ('*Et pour bien couper droit...*'). Le omissioni semiotiche, invece, sono solo 6 su 96 perché effettivamente la maggior parte di tutto quello che è stato omesso è mostrato dalle immagini. Solo in un caso bisognava essere veramente precisi nel rendere l'informazione in maniera letterale, cioè quando il cronista dà le misure per fare degli intagli nel legno. Nessuno dei rispeaker è stato abbastanza preciso:

[30]

TP: Là, j'ai fait des, des, des coupes. / Et là j'ai mis 10 centimètres, 20 centimètres, 20 centimètres, 20 centimètres, 20 centimètres et 10 centimètres.

TA – *respeaker* 2: Ho fatto dei tagli , / ho messo dieci centimetri, 20 centimetri, 20 centimetri, 20 centimetri, dieci centimetri ,

TA – *respeaker* 3: Ho fatto dei tagli / di circa dieci centimetri, 20 centimetri, 30 centimetri. / Ho fatto delle tacche fino ad arrivare a un metro e 20 / di profondità.

Il *rispeaker* 1 non ha reso le macro-unità, il *rispeaker* 2 si è interrotto a metà, il *rispeaker* 3 ha sbagliato in un primo momento e poi ha corretto il tiro. Peccato che l'informazione seguente sia errata: i centimetri di profondità erano 15, mentre il 'metro e 20' era di lunghezza. Nonostante il cronista mostri il pezzo di legno con gli intagli, dalle immagini le misure non sono percepibili, perciò è un'informazione in meno o comunque sbagliata trasmessa al telespettatore. Anche le compressioni rispecchiano l'andamento generale, nonostante le compressioni semiotiche siano più numerose rispetto agli altri risultati (17 su 58). Le espansioni semantiche hanno un'alta percentuale (32.56%) e l'esempio seguente ci può far capire il perché:

[31]

TP: Je montre?

TA: Adesso vi insegno come si fa. Non è complicato.

Il *rispeaker* 2 ha tutto il tempo necessario per formulare un sottotitolo completo e ricco di informazioni. Infine, gli errori a livello di micro-unità sono solo 9, di cui 5 imputabili al *rispeaker*. La '*scie sauteuse*' diventa un *taglierino* invece di 'sega a nastro'; '*votre*', *vostro* invece che 'suo'. Sono perlopiù errori di distrazione o *lapsus* che influiscono a livello di micro-unità. A creare errori di senso contribuisce anche l'*editor* che, per esempio, corregge un sottotitolo inserendo *tiene* invece di 'cade', cioè l'esatto contrario, veicolando l'informazione sbagliata. Tra gli errori del *software*, degni di nota sono il *ventaglio* al posto dell''intaglio' e la *mensa* al posto della 'mensola'.

Respeaker 3

MACRO-UNITÀ NON RESE							MACRO-UNITÀ RESE					TOT
omissioni			errori			Tot	traduzioni letterali		alterazioni		Tot	391 100%
264 90.4%			28 9.6%			292 74.7%	0 0%		99 100%		99 25.3%	
sem.	non sem.	Tot	umani		software	Tot						
38 14.4%	226 85.6%	264 100%	R	E	13 46.4%	28 100%						
			7 25%	8 28.6%								

alterazioni													
riduzioni						espansioni			errori		Tot		
65 65.66%						32 32.32%			2 2.02%		99 100%		
omissioni			compressioni			Tot (micro-unità)	S	NS	Tot (micro-unità)	umani	software	Tot (micro-unità)	
113 66.9%			56 33.1%			169 100%	7 8.75%	73 91.25%	80 100%	R	E	4 36.4%	11 100%
sem.	non sem.	Tot	sem.	non sem.	Tot					3 27.2%	4 36.4%		
2 1.8%	111 98.2%	113 100%	9 16.07%	47 83.93%	56 100%								

Figura 22 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *Comment ça va bien ?_LA DEMO !* – respeaker 3.

Come già accennato, anche questa prova è valida. Il problema è che il rispeaker 3, ancora una volta, cerca di non tralasciare niente, ripetizioni comprese (– *Si tratta di una mensola per bambini – Una mensola per bambini con legno e plastica. – Quindi un miscuglio tra plastica e legno.*). Queste ripetizioni, a parte aumentare il numero di macro-unità espanse (32.32%), comportano un maggior numero di sottotitoli da elaborare, verificare ed eventualmente correggere, ma anche da leggere. Ciò porta a un incremento degli errori generati dal *software* e dall'*editor* e a un maggior impegno da parte del telespettatore. Su 292 macro-unità non rese, 28 sono errori, di cui il 46.4% imputabile al mal riconoscimento del *software*, il 25% imputabile al rispeaker (quindi errori di senso) e il 28.6% imputabile all'*editor*. Non ci sono casi eclatanti. Giusto per citare un esempio:

[32]

TP: Et aussi chez certains marchands de bricolage.

TA: e gli strumenti che servono per il bricolage.

Ancora una volta, un caso in cui il rispeaker percepisce l'ultima parola e collegandosi alle frasi precedenti produce un sottotitolo di senso compiuto e coerente, ma diverso dal TP. Gli errori imputabili all'*editor*, invece, si trovano perlopiù alla fine della trasmissione perché, considerato il ritardo dei sottotitoli, accelera il ritmo del proprio lavoro cancellando sottotitoli corretti. Nel tentativo di non tralasciare niente, il rispeaker 3 riporta brillantemente alcune macro-unità completamente ignorate dagli altri due rispeaker: allora bisogna essere un po' bricoleur per fare questa cosa.; Attenzione a non tagliare fino in fondo il legno altrimenti si farebbero molte piccole mensole / e sarebbe un'altra cosa. Come si può ben notare sono tutti esempi di espansione non semantica, una strategia di cui il rispeaker 3 abusa (73 su 80). Un esempio di espansione semantica, invece, può essere il seguente:

[33]

TP: On apprécie votre technique. / (...) / Vous commencez en douceur et puis vous finissez en...

TA: abbiamo apprezzato molto la tecnica / di iniziare lentamente e poi dare il colpo di grazia.

Il rispeaker esplicita qualcosa che nel TP resta in sospeso, completando la battuta. Per quanto riguarda le riduzioni (65.66%), queste riflettono l'andamento generale, cioè più omissioni (66.9%) e meno compressioni (33.1%), più riduzioni non semiotiche (98.2% di omissioni e 83.93% di compressioni) e meno riduzioni semiotiche (1.8% di omissioni e 16.07% di compressioni). Un esempio per tutti:

[34]

TP: En fait, ça, ça, vous allez voir, ça tient par, par coïncidence, tout simplement.

TA: Vedrai che tiene.

Il rispeaker elimina i tratti dell'oralità e gli elementi inutili, comprime l'informazione più importante, ma ne omette un'altra altrettanto importante, cioè il modo in cui tiene l'oggetto. Tra gli errori a livello di micro-unità degni di nota, abbiamo *una plastica al polipropilene (editor)*, *Questa è quello che ho creato (respeaker)*, per una volta il software non è colpevole di un errore di concordanza grammaticale) e *oppure come questa in parte* (il software non riconosce 'opaca'). Infine, occorre ricordare che il rispeaker 3 ha introdotto anche il cambio locutore tramite l'inserimento di trattini (-) a inizio sottotitolo con ottimi risultati. C'è da dire, perciò, che, seppur minimo, uno sforzo in più rispetto agli altri due rispeaker c'è stato.

4.2.4 *C à vous_L'invité & Le dîner*

C à vous è un programma di intrattenimento, proprio come il precedente, di cui è stata selezionata una parte relativa a un'intervista informale. Perciò, anche se sono presenti altri personaggi che intervengono con qualche battuta, il dialogo principale si svolge tra la presentatrice e l'intervistato. Di conseguenza, il discorso è meno confuso per il rispeaker che è così facilitato nella selezione delle battute da sottotitolare. Di contro, il linguaggio utilizzato in questo programma è ancora più colloquiale e, soprattutto nella sottofase iniziale, si ricorre a una serie di battute circostanziali molto veloci. La componente video è preponderante in certi tratti e complementare a ciò che si dice in altri. Non è quindi sempre necessario accorciare il ritardo nella trasmissione dei sottotitoli, fatta eccezione per alcuni punti precisi. Detto questo, pare ovvio prevedere che i risultati siano simili al programma precedente. Infine, bisogna ricordare che questo programma è stato diviso in due parti: la prima parte (*L'invité*) è stata sottotitolata da tutti i rispeaker; la seconda parte (*Le dîner*) è stata sottotitolata interamente dal rispeaker 3 e parzialmente dal rispeaker 2. Considerando la difficoltà della sottofase iniziale (e l'inutilità nel sottotitolarla visto che dal video si capisce benissimo che si tratta di due persone che si salutano), i tecnicismi, i numerosi tratti dell'oralità e la velocità di eloquio (212 wpm, la più alta tra le diverse trasmissioni), gli unici dati veramente importanti da sottolineare sono, rispetto alla prova precedente, il maggior tasso di macro-unità rese¹⁶⁶ (20.74%, 34.8% e 31.1% per *L'invité*; 32.3% e 29.4% per *Le dîner*) e la diminuzione delle espansioni (9.1% e 26.83% per *L'invité*, rispettivamente per il rispeaker 2 e il rispeaker 3, mentre il rispeaker 1 è in controtendenza con il 33.3%; in *Le dîner*, 26.4% per il rispeaker 2 e in controtendenza il rispeaker 3 con il 33.7%¹⁶⁷). Per il resto, le percentuali variano di qualche punto in più o in meno, ma la tendenza è la stessa di *LA DEMO* !. Nell'insieme, nessuna di queste prove è presentabile come prodotto finale. Se nessun rispeaker (in nessuna delle due parti e quindi anche con turni di lavoro più o meno lunghi) è riuscito a realizzare buoni risultati, significa che il TP era forse troppo difficile per essere inserito all'interno di questo stadio della sperimentazione. Con una formazione appropriata, il giusto addestramento del *software* e dell'esercizio, i risultati sarebbero stati sicuramente accettabili (soprattutto per le prove del rispeaker 2 e del rispeaker 3).

¹⁶⁶ Sicuramente perché *L'invité* dura solo quattro minuti per cui il rispeaker non si stanca eccessivamente.

¹⁶⁷ Ciò non stupisce più di tanto visto la tendenza del rispeaker 3 a utilizzare la strategia dell'espansione. Mentre nella prima parte doveva prendere dimestichezza col programma TV, nella seconda parte possedeva già le conoscenze necessarie per riformulare il testo.

Respeaker 1

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE			TOT
omissioni			errori			Tot	traduzioni letterali	alterazioni	Tot
82 76.6%			25 23.4%			107 79.26%	1 3.6%	27 96.4%	28 20.74%
sem.	non sem.	Tot	umani	software	Tot				
14 17.1%	68 82.9%	82 100%	R	E	15 60%	25 100%			
			5 20%	5 20%					

alterazioni									
riduzioni					espansioni			errori	
17 63%					9 33.3%			1 3.7%	
omissioni		compressioni			Tot (micro-unità)	S	NS	Tot (micro-unità)	
30 69.8%		13 30.2%			43 100%	5 31.25%	11 68.75%	16 100%	
sem.	non sem.	Tot	sem.	non sem.	Tot	umani	software	Tot	
1 3.3%	29 96.7%	30 100%	1 7.7%	12 92.3%	13 100%	R	E	5 55.6%	9 100%
						1 11.1%	3 33.3%		

Figura 23 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *C à vous_L'invité* – respeaker 1.

Per quanto alcune parti siano rese bene, nell'insieme, la sottotitolazione non è utilizzabile. Infatti, solo 28 macro-unità su 135 sono rese e di queste la maggior parte è minimalista (su 27 macro-unità alterate si contano 43 riduzioni a livello di micro-unità, 9 errori e solo 16 espansioni). Da sottolineare l'eccessiva enfasi utilizzata dal rispeaker durante la formulazione del TM1 ('ooh', 'aah', 'che belli!'), la quale ha contribuito a incrementare il numero di errori imputabili al *software* (60% a livello di macro-unità e 55.6% a livello di micro-unità). Le parti informative non deducibili dalle immagini sono state sbagliate o contraddette durante la sottotitolazione. A sua volta, l'*editor*, non capendo i turni di parola, in alcuni punti ha inserito male la punteggiatura e il *software* non ha di certo contribuito alla riuscita della prova. Le percentuali delle omissioni e delle compressioni semiotiche sono molto basse (17.1% di omissioni semiotiche a livello di macro-unità, 3.3% di omissioni semiotiche e 7.7% di compressioni semiotiche a livello di micro-unità) perché la mancanza di una sottotitolazione adeguata è stata compensata dalle immagini. Le omissioni (69.8%) sono più del doppio rispetto alle compressioni (30.2%) perché è più facile ignorare le battute che comprimerle, anche se non mancano degli esempi di compressione [35]. La percentuale più alta di espansioni rispetto alle prove precedenti (dal 7.9% al 28.3% in più) è dovuta al fatto

che il rispeaker ha cercato di recuperare le informazioni perse, ma ha anche esplicitato informazioni che erano già chiare:

[35]

TP: Bref, c'est une grande star dans son domaine.

TA: è una grande star nel suo ambito, il pattinaggio in linea.

Gli unici altri due esempi degni di nota, riguardano errori a livello di macro-unità. Il primo è un 'concorso in errore', cioè sia il rispeaker, sia il *software*, sia l'*editor* (con 'l'aiuto' del suggeritore) hanno contribuito a non rendere questa macro-unità:

[36]

TP: Plus de 100 compétitions professionnelles. 75 victoires. (...)

TM1: Ha fatto più di cento gare /eeee / ha fatto cento gare professionali, 75 vittorie

TM2: ha fatto più di cento ad Korea e ha fatto sette dovere professionale 5 vittorie

TA: Ha fatto più di cento salti. E ha fatto sette gare professionali. 5 vittorie

Il rispeaker ha esitato troppo, ha ripetuto la frase e ha pronunciato male '75'; il *software* ha riconosciuto male l'intera frase; l'*editor*, di fronte a tanti errori e immemore del TM1, ha fatto quello che ha potuto procedendo per logica, ma ottenendo risultati errati. Il secondo, è un esempio che fa sorridere: *ha fatto un salto dal primo piano / e per farlo è stata in Cina*. Questo sottotitolo scorre sotto le immagini del salto dalla torre Eiffel. Semplicemente, il *software* ha riconosciuto *in Cina* al posto di 'allucinante' e l'*editor* ha fuso questa frase con la precedente.

Respeaker 2

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE			TOT
omissioni			errori			Tot	traduzioni letterali	alterazioni	Tot
79			9			88	3	44	47
89.8%			10.2%			65.2%	6.4%	93.6%	34.8%
sem.	non sem.	Tot	umani		software	Tot			
13	66	79	R	E	3	9			
16.46%	83.54%	100%	6	0	33.3%	100%			
			66.7%	0%					

alterazioni										
riduzioni						espansioni			errori	
39 88.6%						4 9.1%			1 2.3%	
Tot			Tot			Tot			Tot	
44 100%			44 100%			44 100%			44 100%	
omissioni			compressioni			S	NS	Tot (micro-unità)	umani	software
62 70.5%			26 29.5%			5 25%	15 75%	20 100%	R	E
sem.	non sem.	Tot	sem.	non sem.	Tot				2 40%	1 20%
5 8.1%	57 91.9%	62 100%	8 30.8%	18 69.2%	26 100%				2 40%	
									5 100%	

Figura 24 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *C à vous_L'invité* – respeaker 2.

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE			TOT	
omissioni			errori			Tot	traduzioni letterali	alterazioni	Tot	167 100%
96 84.96%			17 13.04%			113 67.7%	1 1.85%	53 98.15%	54 32.3%	
sem.	non sem.	Tot	umani		software	Tot				
25 26.04%	71 73.96%	96 100%	R	E	6 35.3%	17 100%				
			10 58.8%	1 5.9%						

alterazioni										
riduzioni						espansioni			errori	
39 73.6%						14 26.4%			0 0%	
Tot			Tot			Tot			Tot	
53 100%			53 100%			53 100%			53 100%	
omissioni			compressioni			S	NS	Tot (micro-unità)	umani	software
80 72.7%			30 27.3%			10 19.6%	41 80.4%	51 100%	R	E
sem.	non sem.	Tot	sem.	non sem.	Tot				0 0%	0 0%
11 13.75%	69 86.25%	80 100%	6 20%	24 80%	30 100%				1 100%	
									1 100%	

Figura 25 – Risultati dell'analisi strategica dei sottotitoli interlinguistici¹⁶⁸ – *C à vous_Le dîner* – respeaker 2.

Entrambe le parti della trasmissione potrebbero andar bene se confrontate ad altre prove decisamente incomprensibili. Però, per mantenere un livello di accettabilità medio-alto, è necessario scartare entrambe le sottotitolazioni come prodotto finale e mantenerle, tuttavia, in sede di analisi in quanto contengono spunti interessanti. Partendo dalle macro-unità non rese

¹⁶⁸ Questa è l'unica prova analizzata senza il supporto della registrazione video. Inoltre, è stato effettuato il *respeaking* solo dei primi 5:28 minuti del filmato.

(65.2% e 67.7%), queste sono praticamente il doppio rispetto a quelle rese (34.8% e 32.3%). Ciò accadeva anche nel programma precedente, addirittura con uno scarto maggiore, perciò questo risultato non ci meraviglia. Il problema di questi risultati risiede nella percentuale di errori all'interno delle macro-unità non rese (10.2% e 15.04%). A prima vista, gli errori in questione hanno percentuali basse, ma se andiamo a vedere il video con i relativi sottotitoli, ci accorgeremmo che più di un sottotitolo è sbagliato. Le macro-unità non rese per errore, a differenza della prova precedente, si tramutano per la maggior parte in sottotitoli visibili, ma non fruibili. Alcuni contengono errori grammaticali e ortografici (meno gravi): In Ci sono tutti i miei allenamenti; altri sono sensati ma non sono fedeli a ciò che dice il TP (mediamente gravi):

[37]

TP: *parce que c'est des sports extrêmes, / mais c'est des sports que vous essayez de faire rentrer aussi dans les vraies mentalités sportives, c'est à dire.*

TA: e questo è uno sport che lei cerca di diffondere più possibile.

Altri sono insensati o comunque non sono collocati all'interno di un discorso logico con un inizio, uno svolgimento e una fine (più gravi):

[38]

TP: *Mais nous, moi, je vous vouvoie, quand même. / Vous pouvez me tutoyer, Taïg.*

TA: Lei può dare del lei, / può dare del tu.

Per quanto riguarda le macro-unità rese, nella prima sottofase (quella dei saluti iniziali), i sottotitoli sono pochi, brevi ed essenziali. Questa soluzione può considerarsi ottimale dal momento in cui le immagini sono sufficienti a capire cosa sta succedendo e quale può essere il dialogo ('ciao, come stai?', 'tutto bene?', 'bene, grazie'). Appare pertanto inutile sottotitolare ogni singola battuta. Tra l'altro, il cambio turno è molto frequente, le voci si sovrappongono e il rispeaker dovrebbe fare uno sforzo notevole proprio all'inizio del suo turno per capire e sottotitolare tutti questi passaggi, stancandosi più facilmente e rischiando di compromettere il resto della *performance*. Infine, dal momento che i sottotitoli non sono sincronizzati con le battute e le immagini in rapida successione, è ulteriormente inutile tenere il passo con il TP. Le parti che dovrebbero essere rese meglio sono, invece, quelle più esplicative e lunghe. Da notare anche che le poche traduzioni letterali corrette riguardano, ancora una volta, battute brevi e concise (*Come stai?, che ha il record del mondo, Venga*). All'interno delle alterazioni (44 su 135 e 53 su 167), predominano ancora le riduzioni (88.6%

e 73.6%), con una decisiva preponderanza delle omissioni (70.5% e 72.7%) sulle compressioni (29.5% e 27.3%). Da ciò si capisce che i sottotitoli sono minimalisti, soprattutto se si tiene in considerazione la percentuale delle espansioni in *L'invité* (9.1%). La produzione di sottotitoli essenziali se da una parte comporta una maggiore perdita di notizie (8.1% e 13.75% di omissioni semiotiche; 30.8% e 20% di compressioni semiotiche), dall'altra può rivelarsi molto efficace, come in questo caso:

[39]

TP: - Mais alors, comment on s'entraîne, Taïg, concrètement? /

- Alors, plein de façons différentes.

TA: - Come ci si allena? /

- In modi diversi.

Le espansioni aumentano del 17.3% in *Le dîner*. Il rispeaker, conoscendo già l'argomento, acquista una maggiore sicurezza e utilizza più frequentemente la strategia dell'espansione, recuperando anche informazioni che aveva ignorato precedentemente:

[40]

TP: Moi, j'ai vraiment décroché tout l'aspect...

TA: Io non mi occupo più della parte organizzativa / quindi non so quanto costi questa impresa.

In questo esempio, il rispeaker ha completato la macro-unità lasciata in sospeso in francese e ha inoltre ripreso una macro-unità, che non aveva tradotto in precedenza, senza la quale il sottotitolo non avrebbe avuto senso. Infine, per quanto riguarda gli errori a livello di micro-unità, i più importanti sono: *la porta diversa ai* invece della 'Porta di Versailles'; *ho ricevuto anche il Presidente della Repubblica* invece di 'mi ha ricevuto anche il Presidente della Repubblica'.

Respeaker 3

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE			TOT
omissioni			errori			Tot	traduzioni letterali	alterazioni	Tot
81 87.1%			12 12.9%			93 68.9%	1 2.4%	41 97.6%	42 31.1%
sem.	non sem.	Tot	umani	software	Tot				
11 13.6%	70 86.4%	81 100%	R	E	1 12 100%				
			9 75%	2 16.7%					

alterazioni													
riduzioni						espansioni			errori		Tot		
26 63.41%						11 26.83%			4 9.76%		41 100%		
omissioni			compressioni			Tot (micro-unità)	S	NS	Tot (micro-unità)	umani		software	Tot (micro-unità)
46 63.9%			26 36.1%			72 100%	3 8.8%	31 91.2%	34 100%	R	E	7 53.8%	13 100%
sem.	non sem.	Tot	sem.	non sem.	Tot					1 7.7%	5 38.5%		
3 6.5%	43 93.5%	46 100%	9 34.6%	17 65.4%	26 100%								

19.1% per *Le dîner*). Il rispeaker 3, però, mostra una tendenza alla riformulazione piuttosto che all'eliminazione perché le percentuali relative alle compressioni e alle espansioni sono maggiori rispetto a quelle del rispeaker 2 (il 6.6% in più delle compressioni e il 17.73% in più delle espansioni per *L'invité*; il 5.6% in più delle compressioni e il 7.3% in più delle espansioni per *Le dîner*). Vediamo un esempio di compressione semiotica:

[41]

TP: Plus de 100 compétitions professionnelles. 75 victoires. Vous êtes triple champion du monde de roller sur rampe, vainqueur de X Game, de Gravity Games et du tournoi ASA.

TA: campione del mondo ha vinto moltissime competizioni,

Pur cui, con una pesante perdita semiotica, tutto ciò è stato reso con una semplice frase riassuntiva e priva di errori di senso o di traduzione. Un esempio di espansione non semantica è invece il seguente:

[42]

TP: et ça passe.

TA: si arriva a risolvere tutto senza problemi.

Nonostante i dialoghi siano molto veloci, il rispeaker (che, ricordiamo, introduce anche i trattini al cambio oratore, ma non detta la punteggiatura) trova tempo anche per questo tipo di esplicitazioni. In *L'invité*, la sottofase iniziale (quella dei saluti) è resa molto bene, con sottotitoli brevi ed essenziali. Il resto della prima parte e *Le dîner* presentano, invece, diversi errori ben evidenti all'interno dei sottotitoli. Citiamo tre esempi. In riferimento all'esempio [38], la soluzione del rispeaker 3 è: *possiamo darci del tu?* Un pubblico italiano non troverebbe strana questa domanda, nonostante riporti un'altra cosa rispetto al testo francese. Il problema è che per tutto il resto del programma il rispeaker deve ricordarsi di dare del tu (cosa che non succede, anzi, a un certo punto dà del 'voi'), mentre il testo francese continuerà a utilizzare il 'vous'. Il secondo esempio: 12 metri e 50 di caduta libera sono diventati *150 metri*, a conferma che i numeri sono sempre difficili da capire e tradurre. Il terzo esempio è una frase indecifrabile: *e ho detto gli strumenti ne permetterà di atterrare con quel che arrivare senza correre alcun rischio*. Una parola di difficile riconoscimento per il software è 'autorizzazioni', riconosciuta più volte come *prestazioni*. Ciò accade più volte con il rispeaker 2 e una volta con il rispeaker 3. Da notare, infine, gli ultimi sottotitoli giusti nel TM2, ma cancellati dall'*editor* per il ritardo accumulato nel verificare, correggere e inviare i sottotitoli.

4.2.5 *Journal_Parole directe*

Da un'intervista informale si passa a un'intervista formale e, più precisamente, politica. In questa trasmissione le immagini non svolgono un ruolo importante¹⁶⁹, perciò la componente audio verbale è preponderante. A conferma di ciò, notiamo subito un sostanziale aumento delle omissioni semiotiche a livello di macro-unità (41%, 37.4% e 39.9% delle omissioni) e a livello di micro-unità (11.8%, 15% e 19% delle omissioni)¹⁷⁰, a fronte di una percentuale inferiore di macro-unità non rese (74.6%, 48.6% e 58.5%)¹⁷¹ rispetto agli altri programmi. Per quanto riguarda la lingua, sono presenti molti tecnicismi tipici del politichese, le pause sono brevi e il testo è lessicalmente denso. I turni, a differenza dei due programmi precedenti, sono abbastanza prevedibili e si limitano a pochi locutori. La velocità di eloquio (198 wpm) non lascia grande spazio alle espansioni, anche se queste non mancano (8.5%, 22.1% e 30.93% delle alterazioni). Aumentano nuovamente le compressioni (37.5%, 46.6% e 39.1% delle riduzioni) proprio come nel meteo e nel telegiornale, a testimonianza di un TP più lungo e discorsivo e non frammentato in tante brevi battute. Anche il numero delle traduzioni letterali corrette aumentano rispetto ai due programmi precedenti (7, 19 e 8 su 354 macro-unità), segno della presenza di espressioni formulaiche, facili da tradurre parola per parola. In ultima analisi, non mancano i tratti dell'oralità e la ripetizione degli stessi concetti, testimoniati dalla sempre alta percentuale delle omissioni non semiotiche (59%, 62.6% e 60.1% a livello di macro unità; 88.2%, 85% e 81% a livello di micro-unità). Per chi conosce la politica francese e l'attualità, questi 15 minuti di trasmissione non sono eccessivamente difficili da sottotitolare, ma la lunghezza del turno di lavoro, la varietà degli argomenti trattati e la terminologia del politichese pesano a livello psico-cognitivo. Perciò il rispeaker non riesce a tenere lo stesso livello di concentrazione, facendo molto bene alcune parti e calando vistosamente in altre.

¹⁶⁹ Cfr. sottopar. 3.5.3.

¹⁷⁰ Le percentuali sono seconde solo a quelle del JT. Non a caso, *Parole directe*, come si è detto, è una rubrica del telegiornale. Quindi è normale che ci siano delle similitudini tra i due generi.

¹⁷¹ Le percentuali sono seconde solo a quelle del *Météo*, in cui le immagini, la semplicità dell'argomento e la linearità del TP aiutavano senza dubbio il rispeaker nella resa delle macro-unità.

Respeaker 1

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE			TOT
omissioni			errori			Tot	traduzioni letterali	alterazioni	Tot
156 59.1%			108 40.9%			264 74.6%	7 7.8%	83 92.2%	90 25.4%
sem.	non sem.	Tot	umani	software	Tot				
64 41%	92 59%	156 100%	R	E	68 63%	108 100%			
			21 19.4%	19 17.6%					

alterazioni									
riduzioni					espansioni			errori	
67 80.7%					7 8.5%			9 10.8%	
omissioni		compressioni		Tot (micro-unità)	S	NS	Tot (micro-unità)	umani	software
85 62.5%		51 37.5%		136 100%	9 32.14%	19 67.86%	28 100%	R	E
sem.	non sem.	Tot	sem.	non sem.	Tot			3 8.3%	6 16.7%
10 11.8%	75 88.2%	85 100%	8 15.7%	43 84.3%	51 100%			27 75%	36 100%

Figura 28 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *Journal_Parole directe* – respeaker 1.

Citando un dato per tutti, si può capire l'assoluta inaccettabilità della prova: su 354 macro-unità, 90 sono rese, di cui 54 sono completamente corrette, mentre 36 contengono almeno un errore. Ciò che maggiormente rende impossibile l'analisi a livello di macro-unità in questa prova è il fatto che le poche macro-unità rese sono sparse nel testo. Ciò significa che non ci sono più di 4-5 macro-unità giuste di seguito e molto spesso sono addirittura isolate. Comparando la registrazione audio ai sottotitoli, si è scoperto che in realtà il respeaker 1, a parte in alcune occasioni, non ha fatto errori di senso rilevanti e ha seguito abbastanza bene il testo tralasciando poche frasi. Non c'è una logica selettiva per cui il respeaker ha deciso di omettere delle informazioni e renderne delle altre. Non regge la scusante secondo la quale in certi argomenti il respeaker è più ferrato rispetto ad altri. Il problema quindi risiede nel mal riconoscimento del TM1 da parte del *software* (che in gran parte, comunque, dipende dalla *performance* del respeaker) e, infatti, ciò è visibile nei risultati finali dell'analisi: 68 macro-unità non rese per errori di riconoscimento, 19 errori di *editing* (per quanto bravo, l'*editor* non può stare dietro a tanti errori) e 21 errori di *respeaking* per un totale di 108. Non meno gravi gli errori a livello di micro-unità: rispettivamente 27, 6 e 3 errori per un totale di 36 (su 83 macro-unità alterate). In alcuni punti l'*editor* è talmente veloce che oltre ai sottotitoli sbagliati

ne cancella anche alcuni giusti. Questo succede quando il rispeaker aumenta il ritmo di eloquio. Di fronte a tali risultati, rimane poco da analizzare. Si confermano la predominanza delle riduzioni (80.7%) e, al suo interno, delle omissioni (62.5%) rispetto alle compressioni (37.5%), risultato dell'applicazione del principio 'si fa prima a eliminare che a riformulare'. Infatti, anche le espansioni (8.5%) sono poche rispetto alle prove degli altri due rispeaker (22.1% e 30.93%). Le riduzioni semiotiche sono dell'11.8% per le omissioni e del 15.7% per le compressioni. Si eliminano quindi molti elementi ridondanti, ma, in proporzione alle macro-unità rese, si eliminano e comprimono anche tante informazioni irrecuperabili importanti. Passiamo ora ad alcuni esempi:

[43]

TP: Cette situation est intolérable / et l'est d'autant plus que le Président de la République avait fait du problème de la sécurité un objectif dans ses discours,

TA: Questa situazione è intollerante / anche perché il Presidente aveva fatto del problema sicurezza il suo obiettivo

Ritroviamo qui tre casi di compressione non semiotica, uno di compressione semiotica e uno di omissione non semiotica. Nel complesso, le macro-unità sono ben rese a parte per l'errore iniziale nella prima delle due macro-unità.

[44]

TP: deux milliards d'euros en moins dans les caisses de l'État.

TA: e ci saranno due miliardi di euro in meno nelle casse dello Stato

L'aggiunta del verbo rende più scorrevole la frase alla lettura. A parte l'espansione non semantica, notiamo che il resto della frase è tradotta letteralmente senza che ciò generi errori di senso o di riconoscimento. Questo succede anche in tante altre macro-unità e in particolare nelle 4 rese solo grazie alla traduzione letterale. Ciò significa che con un opportuno controllo del tono della voce, dello stress e del ritmo di eloquio, il *software* riconosce bene gli enunciati. Tutto quello di cui si ha bisogno è quindi un po' di esercizio e di familiarità col *software*, dopodiché il lavoro avrà dei risultati decisamente migliori, come nel caso del *Météo* e de *LA DEMO* ! per il rispeaker 1.

Respeaker 2

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE			TOT
omissioni			errori			Tot	traduzioni letterali	alterazioni	Tot
147 85.5%			25 14.5%			172 48.6%	19 10.4%	163 89.6%	182 51.4%
sem.	non sem.	Tot	umani	software	Tot				
55 37.4%	92 62.6%	147 100%	R	E	17 68%	25 100%			
			8 12%	0 0%					

alterazioni									
riduzioni					espansioni			errori	
124 76.1%					36 22.1%			3 1.8%	
omissioni		compressioni		Tot (micro-unità)	S	NS	Tot (micro-unità)	umani	software
127 53.4%		111 46.6%		238 100%	16 28.6%	40 71.4%	56 100%	R	E
sem.	non sem.	Tot	sem.	non sem.	Tot	11 35.5%		7 22.6%	13 41.9%
19 15%	108 85%	127 100%	27 24.3%	84 75.7%	111 100%				31 100%

Figura 29 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *Journal_Parole directe* – respeaker 2.

Insieme al meteo, questa è l'unica prova in cui le macro-unità rese (182) sono di più rispetto a quelle non rese (172). La differenza è che il meteo dura 2:10 minuti e *Parole directe* ne dura 15. Se consideriamo poi che il rispeaker ha voluto continuare il rispeakeraggio fino alla fine della trasmissione¹⁷², arriviamo a 23 minuti totali. Ciò che è ancora più importante dire è che il rispeaker ha trattato quasi tutti gli argomenti (anche se alcuni con un po' di imprecisione) e solo in due punti ha mostrato delle difficoltà. Un primo punto corrisponde alla sottofase in cui l'intervistatrice fa delle brevi domande pro/contro all'intervistata, perciò le battute sono più veloci e trattano argomenti diversi nel giro di poche macro-unità. Mentre la prima domanda e la conseguente risposta sono state rese nel TM1 ma non nel TA a causa di un cattivo riconoscimento da parte del *software*, la domanda e la risposta successive non sono state riportate nemmeno dal rispeaker, che ha ripreso il filo dalla domanda seguente. Tutto ciò significa che il rispeakeraggio è stato più semplice nelle parti discorsive ed esplicative (nonostante il politichese non sia un linguaggio facile né da capire, né da tradurre, né da

¹⁷² Si è preferito mettere da parte questi dati perché i risultati sarebbero stati incomparabili a causa della lunghezza del testo. Quindi per poter fare un paragone scientificamente valido era importante paragonare 15 minuti in tutti e tre i casi.

riprodurre), mentre nello scambio veloce di battute si è incontrata qualche difficoltà, forse per un eccessivo controllo del TM2. Il secondo punto in cui il rispeaker 2 ha incontrato delle difficoltà contiene troppi dati e informazioni importanti, con cifre e nomi propri, che solo chi conosce bene la situazione politica francese può ben capire e restituire in sintesi all'interno dei sottotitoli. Il rispeaker ha inizialmente tentato di spiegare ciò che aveva sentito, ma non riuscendo a completare il discorso ha preferito lasciare il sottotitolo in sospeso tramite l'utilizzo dei puntini di sospensione. Nel complesso, la *performance* è andata molto bene. È vero che molte delle macro-unità rese sono generiche a causa di omissioni e compressioni semiotiche (rispettivamente 15% e 24.3%), ma ciò non vuol dire che non esprimano il concetto generale veicolato nel TP, anzi sono più dirette e accessibili da parte del telespettatore:

[45]

TP: Vous savez, 40% des Français disent qu'ils n'ont pas été augmentés depuis trois ans. / C'est, quand même, absolument énorme quand le logement augmente, quand l'énergie, l'essence, tout, le, le gazole augmentent, (...).

TA: molti francesi denunciano che sono tre anni che il loro stipendio non aumenta. / Invece tutto il costo della vita aumenta.

[46]

TP: Et quand on voit ce qui se passe actuellement, / c'est vrai que la situation est intolérable. / Des, des enfants qui s'envoient des coups de couteau. / La jeune, cette jeune petite fille, qui a été violée récemment. / Ce qu'on a vu à Sevran avec des tirs de balles sur, dans, dans, dans des écoles. / À Grigny, encore, récemment, où des policiers se sont battus avec des jeunes. / Cette situation est intolérable

TA: Quando vediamo quello che succede attualmente, / come ragazzine violentate, / sparatorie / e poliziotti picchiati / pensiamo che questa situazione è intollerabile.

Nonostante i due errori nel secondo esempio [46], uno di senso e l'altro di grammatica, e nonostante la forte perdita semiotica in entrambi, il concetto generale di ogni discorso è stato reso. Ricordandoci che si tratta di una sottotitolazione in tempo reale, possiamo affermare che queste soluzioni costituiscono un esempio perfetto di organizzazione del discorso all'interno dei sottotitoli per questo genere di programmi. Ottimo quindi l'adattamento dei sottotitoli a una maggiore fruibilità, come dimostrato anche dal caso seguente:

[47]

TP: Mais, moi, je réponds d'abord qu'il faut créer des emplois / et aujourd'hui l'emploi ça reste la priorité majeure pour les Français. / Et le gouvernement devrait s'en occuper matin, midi et soir.

TA: Innanzitutto bisogna creare occupazione. / I posti di lavoro sono la priorità per i francesi. / Il Governo se ne deve occupare 24 ore su 24.

È anche vero che alcune macro-unità sono molto precise e recuperano alcune delle informazioni perse in precedenza (22.1% di macro-unità espanse e 28.6% di espansioni semantiche). Inoltre, le frasi-chiave tipiche della politica non solo vengono riportate, ma anche esplicitate:

[48]

TP: C'est pour ça qu'il faut des choix tout à fait à l'inverse.

TA: Noi faremo delle scelte contrarie a questo Governo.

Dai quattro esempi precedenti notiamo che, comunque, non mancano le omissioni non semiotiche di macro-unità (62.6%) e le riduzioni non semiotiche (85% di omissioni e 75.7% di compressioni). L'alta percentuale delle omissioni semiotiche di macro-unità, inoltre, è dovuta alla densità dei concetti espressi e delle informazioni fornite dall'intervistata, che ovviamente non possono essere compensate dall'inquadratura della stessa. Un'ultima nota prima di passare all'analisi degli errori (questa volta unica per le macro-unità e per le micro-unità perché nel complesso non contano più di tanto) riguarda le traduzioni letterali. Il rispeaker tende a seguire più da vicino il testo nelle parti in cui il ritmo dell'eloquio cala (168 wmp), cioè nei servizi preregistrati. Con un minimo adattamento alla lingua italiana (e quindi col minimo sforzo in quanto il rispeaker è madrelingua) e il giusto ritmo di produzione del TM1, intere macro-unità sono perfettamente riconosciute dal *software*. Gli errori contano per il 14.5% a livello di macro-unità non rese e per l'1.8% a livello di macro-unità rese, mentre a livello di micro-unità sono solo 31 all'interno di 163 macro-unità. Si nota poi una netta preponderanza di errori da parte del *software* a livello di macro-unità (68%) e di micro-unità (41.9%), quindi di cattivo o non riconoscimento. Ma a differenza del rispeaker 1, ciò non è dovuto a una voce stressata e/o insicura. C'è da precisare che gli errori sono 17 a livello di macro-unità e quindi presumibilmente 17 parole (o poco più) riconosciute male, da aggiungere ai 13 errori a livello di micro-unità (compresi quindi gli errori ortografici e grammaticali), contro un totale di circa 1.340 parole nel TA. Arrotondando per eccesso, la percentuale di errori rispetto al testo contenuto nei sottotitoli è del 3%, dato che corrisponde

perfettamente al margine di accuratezza garantito oggi dalla maggior parte delle compagnie che producono *software* di riconoscimento del parlato (97-98%¹⁷³). I dati degli errori nella tabella sono irrilevanti all'interno dell'intera sottotitolazione. Di certo la prova non è perfetta e i telespettatori avrebbero molto da ridire, ma non è da considerarsi fallimentare. Più preoccupanti gli errori del rispeaker (32% a livello di macro-unità e 35.5% a livello di micro-unità), che si conferma essere a volte troppo distratto generando errori evitabili: *Buongiorno* invece di 'buonasera' (quando poi si corregge al saluto successivo, l'*editor* riscrive *Buongiorno*); *molto francesi*; *subiscono* invece di 'subiscano'; *dei anziani*. A volte fa anche errori di senso, ma niente di eclatante, così come l'*editor* (7 errori su 31 totali a livello di micro-unità).

Respeaker 3

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE				TOT
omissioni			errori			Tot	traduzioni letterali	alterazioni	Tot	354 100%
143 69.1%			64 30.9%			207 58.5%	8 5.44%	139 94.56%	147 41.5%	
sem.	non sem.	Tot	umani		software	Tot				
57 39.9%	86 60.1%	143 100%	R	E	22 34.4%	64 100%				
			38 59.4%	4 6.2%						

alterazioni											Tot		
riduzioni						espansioni			errori				
90 64.75%						43 30.93%			6 4.32%		139 100%		
omissioni			compressioni			Tot (micro- unità)	S	NS	Tot (micro- unità)	umani	software	Tot (micro- unità)	
137 60.9%			88 39.1%			225 100%	28 19.2%	118 80.8%	146 100%	R	E	32 65.3%	49 100%
sem.	non sem.	Tot	sem.	non sem.	Tot					9 18.4%	8 16.3%		
26 19%	111 81%	137 100%	25 28.4%	63 71.6%	88 100%								

Figura 30 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *Journal_Parole directe* – *respeaker 3*.

In tutto il TA, ci sono solo due parti rese in maniera accettabile: dalla macro-unità 133 alla 166 (comprese) e dalla 316 fino alla fine (354)¹⁷⁴. Le rimanenti macro-unità rese sono sparse, per cui nel complesso la prova non è utilizzabile come *output*. Ciò significa che il

¹⁷³ Cfr. Aliprandi & Verruso 2006.

¹⁷⁴ Cfr. appendice – Tabella 6.

rispeaker ha incontrato minori difficoltà in certe tematiche e maggiori difficoltà nelle altre. A conferma di ciò, vediamo che i dati sugli errori del rispeaker a livello di macro-unità (59.4%), non per la prima volta, superano quelli della macchina nonostante il riconoscimento da parte del *software* non sia stato ottimale (34.4% a livello di macro-unità e 65.3% a livello di micro-unità). Bisogna riconoscere che il rispeaker 3, ancora una volta, cerca di non tralasciare niente e di spiegare bene i concetti (30.93% macro-unità espanse e 146 micro-unità espanse contro 225 – compresi i tratti della lingua orale – ridotte), ma ciò si tramuta in una catena di sottotitoli sbagliati e a volte indecifrabili: *chiede del Presidente di Cometa per quartiere di Marsiglia*; oppure in sottotitoli corretti e coerenti col testo, ma che aggiungono qualcosa che nel TP non c'è (19.2% di espansioni semantiche):

[49]

TP: Deux milliards d'euros, c'est à peu près le coût de la prise en charge de la dépendance. / On pourrait en faire 200.000 emplois jeunes. / On pourrait conserver 200.000 logements sociaux.

TA: ee questa cifra di 3 miliardi di euro / potrebbe permettere di aumentare gli impieghi per i giovani trovare nuovi lavori / questo potrebbe essere qualcosa che potrebbe aiutare per risolvere il debito

A parte gli errori più o meno gravi, si può notare come due macro-unità importanti siano state omesse¹⁷⁵ e, invece, si è aggiunta una macro-unità corretta nel contesto, ma non espressa nel TP¹⁷⁶. Si nota anche la reiterazione di un concetto già espresso (*aumentare gli impieghi per i giovani e trovare nuovi lavori*), cosa che avviene spesso quando il rispeaker ha ben chiaro il senso del discorso, ma non riesce a completare la traduzione perché cosciente di aver perso qualcosa:

[50]

TP: Il faut des efforts, / mais ils les veulent justement répartis.

TA: è il problema oggi che servono degli sforzi, / servono degli sforzi da parte di tutti

È ottima, anche in questa prova, la soluzione delle riduzioni (64.75%, di cui 60.9% di omissioni e 39.1% di compressioni). In riferimento all'esempio [47], la soluzione del

¹⁷⁵ Si parla del pensiero di una possibile candidata alle elezioni presidenziali quindi alcuni concetti non possono essere semplicemente eliminati o compressi senza perdita di informazione.

¹⁷⁶ Come nel caso precedente, non si possono mettere parole in bocca a chi non le ha pronunciate, anche se il senso è quello giusto, in un contesto così delicato. Il rispeaker dovrebbe essere il più fedele possibile al TP, benché non sia possibile tradurlo parola per parola.

rispeaker 3 è: - M. Aubry: per me il problema è creare degli impieghi / il Governo dovrebbe occuparsene a tempo pieno. Le generalizzazioni non sono meno importanti:

[51]

TP: Par exemple, un enfant qui met le feu à une poubelle,

TA: per esempio un bambino che fa qualcosa che non va bene

Ottima soluzione nel caso il rispeaker non abbia sentito bene o non voglia incorrere in difficoltà future. Inutile citare gli errori riscontrati (70 a livello di macro-unità e 49 a livello di micro-unità). È sufficiente dire che il rispeaker ogni tanto dà del 'voi' e che il *software* fa errori del tipo *fare e folk vere* invece di 'far evolvere', che non sono stati corretti dall'*editor* per il troppo carico di lavoro.

4.3 Analisi strategica dei sottotitoli interlinguistici: *Météo* in semidiretta e comparazione dei risultati

La prova in semidiretta migliore tra tutte è stata quella del meteo effettuata dal rispeaker 3. A nostro avviso, la comparazione dei risultati dell'analisi strategica di questa prova con la prova dello stesso rispeaker in diretta, è sufficiente per capire l'importanza che ha la possibilità di poter visualizzare il TP in anticipo. Tra l'altro, si precisa che il rispeaker non ha avuto la *chance* di concentrarsi totalmente sul TP per una visualizzazione completa della trasmissione prima della sottotitolazione, ma ha innanzitutto effettuato la sottotitolazione in diretta (quindi senza aver preso visione del TP in precedenza) e poi ha effettuato nuovamente la sottotitolazione in semidiretta. I risultati dell'analisi strategica effettuata sulla sottotitolazione in semidiretta sono stati i seguenti:

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE			TOT
omissioni			errori			Tot	traduzioni letterali	alterazioni	Tot
7 77.8%			2 22.2%			9 19.6%	3 8.1%	34 91.9%	37 80.4%
sem.	non sem.	Tot	umani		software	Tot			
2 28.6%	5 71.4%	7 100%	R	E	0 0%	2 100%			
			1 50%	1 50%					

alterazioni										
riduzioni						espansioni			errori	
24 70.6%						9 26.5%			1 2.9%	
Tot			Tot			Tot			Tot	
34 100%			34 100%			34 100%			34 100%	
omissioni			compressioni			S	NS	Tot (micro-unità)	umani	software
30 53.6%			26 46.4%			1 3.6%	27 96.4%	28 100%	R	E
sem.	non sem.	Tot	sem.	non sem.	Tot				2 66.7%	1 33.3%
1 3.3%	29 96.7%	30 100%	4 15.4%	22 84.6%	26 100%				0 0%	3 100%

Figura 31 – Risultati dell'analisi strategica dei sottotitoli interlinguistici – *Météo* in semidiretta – *respeaker* 3.

La percentuale delle macro-unità rese è salita del 13%, mentre quella degli errori è scesa del 15.35% a livello di macro-unità e da 8 a 3 nelle micro-unità. Il numero di errori cala perché il *rispeaker* e l'*editor* correggono il tiro degli sbagli fatti nella versione precedente, perché sanno già quali parole verranno meglio riconosciute dal *software* e quali, invece, sono più ostiche e come (pre-)editarle. Infatti, per la prima volta il *software* ha lo 0% di errori. L'unico errore di senso del *rispeaker* è stato quello già visto nel *rispeakeraggio* in diretta [11]. Il resto degli errori sono veramente insignificanti. Tra tutti, si cita solamente un errore di battitura dell'*editor*: *avanzerannoù*. Se già la prova in diretta era valida, questa è più che valida. Rasenta quasi la perfezione se non ci fossero ancora dei piccoli accorgimenti a cui bisognerebbe prestare più attenzione. Essendo una prova molto breve, la conoscenza del TP ha portato a un maggior numero di espansioni (da 3 a 9 macro-unità e da 17 a 28 micro-unità) senza rischiare di generare errori. Il *rispeaker*, infatti, realizza sottotitoli più completi e precisi senza preoccuparsi di non capire il passaggio successivo del TP. Di conseguenza, può impegnarsi più nella realizzazione del TM1 e diminuire gli sforzi dell'ascolto/comprendimento e della memoria a breve termine, giocando più su quella a 'medio' termine, cioè ciò che ha già assimilato nella visualizzazione precedente. Solo una espansione su 28 è semantica, ma ciò non deve stupire perché se si parla di previsioni meteorologiche e si segue da vicino il testo, è effettivamente difficile aggiungere ulteriori informazioni. All'interno delle riduzioni (sempre elevate, 70.6%), omissioni e compressioni hanno uno scarto del 7.2% con la prevalenza delle prime sulle seconde. Le riduzioni semiotiche sono poco meno della metà e quelle non semiotiche aumentano leggermente (3 omissioni e una compressione in più). Il motivo di tutto questo è una migliore selezione delle informazioni da veicolare o meno, ma ricordiamo che il meteo già *per se* non 'consente' perdita di informazione grazie alla preponderanza delle immagini sulla componente audio verbale. Per citare alcuni esempi di quanto detto finora, è

utile prendere spunto dagli stessi esempi citati nell'analisi strategica della sottotitolazione in diretta, al fine di vedere quali diverse strategie sono state usate¹⁷⁷.

[10]

TP: et on voit qu'elle arrive avec une formation de dépression ici au large de la Bretagne.

D: Una formazione di depressione arriva dalla Spagna.

SD: Arrivano con la formazione di depressioni qui al largo della Bretagna.

[12]

TP: On les regagne aujourd'hui avec des températures estivales

D: ma lo riacquistiamo sui monti con queste temperature quasi estive.

SD: e li riacquistiamo oggi con temperature pressoché estive

[13]

TP: En revanche, cette partie sud-ouest, ça donne de la grisaille avec quelques gouttes

D: Alcune gocce qui,

SD: OMISSIONE NON SEMIOTICA (informazione deducibile dalle immagini)

[14]

TP: Poches de soleil sur la Méditerranée, également la Corse et aussi sur cette partie du nord-est.

D: Sole sul Mediterraneo e in Corsica così come in questa parte del nord est.

SD: zone soleggiate sia nel Mediterraneo che in Corsica così come in questa zona del nord est.

4.4 Analisi strategica dei sottotitoli intralinguistici: comparazione con i risultati dell'analisi strategica dei sottotitoli interlinguistici

Come è stato anticipato nell'introduzione, nel presente paragrafo si procederà alla comparazione dei migliori risultati ottenuti dall'analisi strategica dei sottotitoli interlinguistici prodotti in diretta con i risultati dell'analisi strategica dei sottotitoli intralinguistici degli stessi programmi. Per ogni programma rispeakerato, si è scelta la migliore sottotitolazione tra le tre

¹⁷⁷ Nei seguenti esempi saranno utilizzate le abbreviazioni 'D' (diretta) e 'SD' (semidiretta).

prove in diretta: il meteo del rispeaker 1; il telegiornale, *La Demo !*, *L'invité* e *Parole directe* del rispeaker 2. Non saranno forniti esempi esplicativi della sottotitolazione intralinguistica in quanto non è interesse specifico della presente tesi sapere come sono stati realizzati i sottotitoli francesi. L'analisi strategica di questi ultimi è stata esclusivamente effettuata nell'ottica di avere un termine di paragone concreto e oggettivo¹⁷⁸ per poter stabilire definitivamente la riuscita o meno degli esperimenti interlinguistici. Bisogna precisare tre cose:

- le imperfezioni quali la mancanza degli accenti sulle maiuscole e della *cédille* sulla 'C' maiuscola non sono state considerate come errori in quanto dovute ad aspetti tecnici del sistema di trasmissione dei sottotitoli;
- la presenza dei trattini e/o dei nomi propri per il cambio interlocutore e delle informazioni para- e/o extralinguistiche non sono state prese in considerazione ai fini di questa analisi, pur costituendo uno sforzo in più per il rispeaker o per l'editor;
- il ritardo nella trasmissione dei sottotitoli rispetto al TP non è stato considerato in nessuna delle due sottotitolazioni.

4.4.1 Météo

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE				Tot
omissioni			errori			Tot	ripetizioni	alterazioni	Tot	46 100%
15 100%			0 0%			15 32.6%	5 16.1%	26 83.9%	31 67.4%	
sem.	non sem.	Tot	umani	software	Tot					
6 40%	9 60%	15 100%	0 0%	0 0%	0 0%					

alterazioni											
riduzioni					espansioni			errori		Tot	
24 92.3%					1 3.85%			1 3.85%		26 100%	
omissioni			compressioni		Tot (micro-unità)	S	NS	Tot (micro-unità)	umani	software	Tot (micro-unità)
30 77%			9 23%		39 100%	2 28.6%	5 71.4%	7 100%	0 0%	1 100%	1 100%
sem.	non sem.	Tot	sem.	non sem.	Tot						
3 10%	27 90%	30 100%	2 22.2%	7 77.8%	9 100%						

Figura 32 – Risultati dell'analisi strategica dei sottotitoli intralinguistici – *Météo*.

¹⁷⁸ Si ricorda che i sottotitoli intralinguistici sono quelli diffusi in diretta dalla televisione francese.

La sottotitolazione intralinguistica del meteo contiene il 17.4% di macro-unità rese in più rispetto alla prova del rispeaker 1. Tuttavia, se la nostra attenzione si sposta sul numero di omissioni, si nota una sola omissione in più da parte del rispeaker 1 rispetto all'intralinguistico e, se si guarda la suddivisione interna delle omissioni, ci si accorge che il rispeaker 1 ha omesso 2 macro-unità non semiotiche in più e una macro-unità semiotica in meno rispetto all'intralinguistico. Ciò significa, prima di tutto, che nell'interlinguistico è stata fornita un'informazione in più non deducibile dalle immagini e dal contesto¹⁷⁹. In secondo luogo, togliendo le 2 omissioni non semiotiche in più rispetto all'intralinguistico e aggiungendo l'omissione semiotica in meno (da uno scarto di 8 macro-unità non rese) risulta che la non resa delle restanti 7 macro-unità è completamente dovuta agli errori. Per cui, se si trovasse il modo di evitare questi errori, i risultati delle due prove non sarebbero molto diversi. Infatti, il resto dei dati non differisce più di tanto. In entrambe le sottotitolazioni vi è una sola espansione a livello di macro-unità e 7 espansioni a livello di micro-unità (l'interlinguistico ha una espansione semantica in meno e, di conseguenza, una espansione non semantica in più). Addirittura, nell'intralinguistico troviamo un errore a livello di macro-unità che non c'è nell'interlinguistico, mentre a livello di micro-unità ci sono 4 errori in meno. Anche all'interno delle riduzioni i risultati cambiano di poco: 39 nell'intralinguistico e 40 nell'interlinguistico. A cambiare è la situazione interna: il rispeaker 1 preferisce comprimere (9.5% di compressioni in più) piuttosto che eliminare (9.5% di omissioni in meno). Le riduzioni semiotiche sono quasi le stesse (3 omissioni e una compressione nell'interlinguistico; 3 omissioni e 2 compressioni nell'intralinguistico), mentre naturalmente cambiano i dati di quelle non semiotiche (rispettivamente 24 e 12 nell'interlinguistico contro 27 e 7 nell'intralinguistico). Infine, nell'intralinguistico vi sono 5 ripetizioni a fronte di 3 traduzioni letterali nell'interlinguistico. Se si considera che il rispeaker francese potrebbe aver avuto in precedenza il testo trascritto¹⁸⁰, possiamo concludere questa comparazione confermando la piena riuscita della prova interlinguistica¹⁸¹ del rispeaker 1, che addirittura ha preferito riformulare le informazioni piuttosto che eliminarle, scegliendo la strada più difficile al contrario del *respeaker* francese.

¹⁷⁹ Considerato che ci sono anche 2 macro-unità non rilevanti semioticamente in meno, il TA è più leggibile, perciò il TA interlinguistico è migliore di quello intralinguistico perché ha un'informazione semioticamente rilevante in più e 2 irrilevanti in meno. Quindi ha meno testo (il che lo rende più fruibile), ma dà più informazioni.

¹⁸⁰ A un certo punto delle parole compaiono sul video prima che le stesse vengano pronunciate. È un caso isolato e le parole potevano essere anticipate guardando le immagini, ma c'è la concreta possibilità che il testo fosse a disposizione del rispeaker prima della diretta. In ogni caso, si esclude la realizzazione dei sottotitoli in precedenza perché oltre a non essere sincronizzati, questi sono anche incompleti.

¹⁸¹ Si tiene a sottolineare che la prova del meteo del rispeaker 1 è stata la prima prova in assoluto dell'intera sperimentazione, con tutti i problemi tecnici e le difficoltà che una prima può comportare, senza una preparazione precedente del rispeaker.

4.4.2 *Télématin_JT*

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE			TOT
omissioni			errori			Tot	ripetizioni	alterazioni	Tot
39 97.5%			1 2.5%			40 47.1%	14 31.1%	31 68.9%	45 52.9%
									85 100%
sem.	non sem.	Tot	umani	software	Tot				
14 35.9%	25 64.1%	39 100%	0 0%	1 100%	1 100%				

alterazioni								
riduzioni			espansioni			errori		
29 93.55%			2 6.45%			0 0%		
omissioni	compressioni		Tot (micro-unità)	S	NS	Tot (micro-unità)	umani	software
21 48.8%	22 51.2%		43 100%	0 0%	4 100%	4 100%	0 0%	0 0%
sem.	non sem.	Tot	sem.	non sem.	Tot			
9 42.86%	12 57.14%	21 100%	4 18.2%	18 81.8%	22 100%			

Figura 33 – Risultati dell'analisi strategica dei sottotitoli intralinguistici – *Télématin_JT*.

La sottotitolazione intralinguistica del telegiornale contiene il 9.4% di macro-unità rese in più rispetto alla prova del rispeaker 2. Se ci spostiamo sul numero di omissioni, ritroviamo però lo stesso identico dato, mentre cambia il numero di errori (9 nell'interlinguistico e uno imputabile al *software* nell'intralinguistico). Se si guarda la suddivisione interna delle omissioni, ci si accorge che il rispeaker 2 ha fatto delle scelte diverse rispetto al collega francese: nell'intralinguistico le omissioni non semiotiche sono più delle semiotiche (come è normale che sia), mentre nell'interlinguistico avviene il contrario per i motivi che abbiamo già spiegato al sottoparagrafo 4.2.2. Tuttavia, questi dati sono da riconsiderare. Nel TP, a un certo punto, appare la trascrizione dell'audio verbale. Il rispeaker francese non ha bisogno di sottotitolare queste parti (13 macro-unità), mentre il rispeaker italiano deve farlo perché sono comunque scritte in francese e deve farlo sulla base di un ritmo di eloquio più elevato del normale. Queste 13 macro-unità non si trovano una di seguito all'altra, ma sono inframmezzate dai commenti del *reporter*. A ogni modo, oltre ad avere lavoro in meno, il rispeaker francese ha anche il tempo di recuperare l'eventuale ritardo accumulato nella sottotitolazione e di riprendere tranquillamente a produrre il TM1 dai commenti del *reporter*. Tornando ai nostri dati, se dalle 25 macro-unità non semioticamente rilevanti non rese

nell'intralinguistico togliamo queste 13 macro-unità, allora la situazione cambia: 14 omissioni semiotiche e 12 non semiotiche. La stessa operazione è da effettuarsi nell'interlinguistico. Di queste 13 macro-unità, 4 sono alterate, 4 sono omissioni semiotiche e 5 sono omissioni non semiotiche. 'Correggendo' i dati, il risultato è il seguente: 17 omissioni semiotiche e 13 omissioni non semiotiche nell'interlinguistico. Notiamo quindi che il rispeaker 2 ha uno scarto di sole 3 macro-unità semioticamente rilevanti non rese in più rispetto al collega e ha addirittura una macro-unità non semioticamente rilevante non resa in più, il che agevola la lettura dei sottotitoli. Purtroppo, ciò che fa realmente la differenza tra macro-unità rese e non rese è il numero di errori. Una differenza importante si riscontra anche tra le ripetizioni (14) e le traduzioni letterali (1). Dal momento che il rispeaker francese, da quello che si può evincere dalla sottotitolazione, cerca di realizzare sottotitoli *verbatim* (cosa impossibile in un passaggio da una lingua a un'altra) e il rispeaker 2 tende a riformulare il testo, questi dati non possono essere paragonati oggettivamente. Per quanto riguarda le alterazioni, se già il rispeaker 2 riformulando il testo non aveva il tempo di fare troppe espansioni (19.4% a livello di macro-unità e 20 a livello di micro-unità), il rispeaker francese che crea sottotitoli il più possibile *verbatim* ha ancora meno tempo (6.45% a livello di macro-unità e 4 a livello di micro-unità, tutte non semantiche¹⁸²). La diversa politica di sottotitolazione spiega anche perché il rispeaker 2 effettua più compressioni (4 in più) ed elimina più elementi (18 in più) rispetto al rispeaker francese. In controtendenza, invece, le compressioni non semiotiche (2 in meno). Ciò può essere spiegato solo dalla difficoltà del passaggio da una lingua a un'altra, che spinge il rispeaker 2 a ridurre il testo ma con una maggiore perdita di informazione (38.5% di compressioni semiotiche) rispetto al rispeaker francese (18.2% di compressioni semiotiche). Considerando anche in questo caso che probabilmente il rispeaker francese ha avuto in precedenza il testo trascritto, tenendo conto della difficoltà nel passaggio da una lingua all'altra e considerando che la maggior parte degli errori realizzati dal rispeaker 2 possono essere evitati, possiamo concludere questa comparazione confermando, anche in questo caso, la piena riuscita della prova.

¹⁸² Le strategie di espansione dei sottotitolatori francesi si riducono quasi esclusivamente all'esplicitazione delle forme contratte ('*c'est*' > *ce est*; '*ça*' > *cela*) e all'aggiunta della particella '*ne*' nella negazione (quasi mai utilizzata nella lingua orale). Quindi sono essenzialmente delle espansioni grammaticali dettate da una politica di chiarezza e correttezza del testo scritto nei sottotitoli.

4.4.3 *Comment ça va bien ?_LA DEMO !*

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE			TOT
omissioni			errori			Tot	ripetizioni	alterazioni	Tot
217 99.1%			2 0.9%			219 56%	44 25.6%	128 74.4%	172 44%
sem.	non sem.	Tot	umani	software	Tot				
17 7.8%	200 92.2%	217 100%	2 100%	0 0%	2 100%				

alterazioni								
riduzioni			espansioni			errori		
121 94.5%			5 3.9%			2 1.6%		
omissioni			compressioni			Tot (micro-unità)	S	NS
172 83.1%			35 16.9%			207 100%	4 21.05%	15 78.95%
sem.	non sem.	Tot	sem.	non sem.	Tot			
5 2.9%	167 97.1%	172 100%	3 8.6%	32 91.4%	35 100%			

Figura 34 – Risultati dell'analisi strategica dei sottotitoli intralinguistici – *Comment ça va bien ?_LA DEMO !*.

La sottotitolazione intralinguistica di *LA DEMO !* contiene il 16.1% di macro-unità rese in più rispetto alla prova del rispeaker 2, che si traducono in 63 macro-unità in più. All'interno delle macro-unità non rese, la percentuale delle omissioni (97.2% nell'interlinguistico e 99.1% nell'intralinguistico) e, di conseguenza, degli errori (0.9% e 2.8%), non sono molto diverse. Cambia, invece, la suddivisione interna delle omissioni visto che il rispeaker francese omette il 5.3%¹⁸³ delle macro-unità semiotiche in meno del rispeaker 2, veicolando più informazioni rilevanti. Per la prima volta, lo scarto del numero di errori non è importante, ma la differenza delle macro-unità rese (63) inizia a essere un dato pesante. Se però si vanno a vedere quali macro-unità in meno sono state restituite dal rispeaker 2 rispetto al collega francese, ci accorgiamo che 42 su 63 sono omissioni non semiotiche, 17 sono omissioni semiotiche e 4 sono errori (di cui 3 si sono tradotti in sottotitoli infedeli al TP e uno non genera un sottotitolo, ma omette una macro-unità non semioticamente rilevante). Il risultato di tutto ciò è che il rispeaker 2, in realtà, ha omesso 17 macro-unità in più rispetto al

¹⁸³ Il calcolo è stato effettuato in proporzione alle macro-unità non rese.

collega francese invece di 63, dato che le restanti 46 sono formate da 43¹⁸⁴ omissioni non semiotiche inutili nella sottotitolazione e da 3 sottotitoli 'infedeli' al TP. Spostando l'attenzione sulle ripetizioni (25.6%) e sulle traduzioni letterali (7.3%), notiamo subito la differenza, a conferma di una resa tendente al *verbatim* (per quanto sia possibile con un linguaggio informale, turni di parola che si sovrappongono, numerosi tratti dell'oralità e battute brevi in rapida successione) da parte del rispeaker francese. Di conseguenza, lo scarto delle macro-unità alterate diminuisce (27 in più nell'intralinguistico, a fronte di 63¹⁸⁵ macro-unità rese in più). La tendenza alla riformulazione nell'interlinguistico si nota dalla percentuale maggiore sia delle espansioni (il 22.8% delle macro-unità e 43 micro-unità contro il 3.9% delle macro-unità e 19 micro-unità), sia delle compressioni (37.7% nell'interlinguistico e 16.9% nell'intralinguistico). Anche nella suddivisione interna le differenze sono notevoli: il rispeaker francese effettua solo 4 espansioni semantiche e 3 compressioni semiotiche all'interno di 128 alterazioni; il rispeaker 2, per i motivi di cui al sottoparagrafo 4.2.3, effettua 14 espansioni semantiche e 17 compressioni semiotiche all'interno di 101 alterazioni. In proporzione, la differenza è enorme, ma, come si è già spiegato, il passaggio da una lingua a un'altra richiede necessariamente una riformulazione e questa non è semplice dal momento in cui le battute da sottotitolare sono brevi e veloci. Irrilevante, invece, il numero di alterazioni per errore (3 su 101 nell'interlinguistico e 2 su 128 nell'intralinguistico). A livello di micro-unità, gli errori imputabili al *software* sono in entrambi i casi 3, mentre nell'interlinguistico vi sono 6 errori umani contro gli 0 dell'intralinguistico. Tra le omissioni, nonostante a prima vista lo scarto delle semiotiche possa apparire minimale (una sola omissione semiotica in più per il rispeaker 2), bisogna fare i conti col numero totale di omissioni. Se nell'intralinguistico su 172 omissioni le riduzioni semiotiche sono 5, nell'interlinguistico su 96 omissioni le riduzioni semiotiche sono 6. Quindi, in proporzione, vi è il 3.35% di omissioni semiotiche in più nell'interlinguistico. In ogni caso, le percentuali delle omissioni semiotiche sono basse perché la maggior parte delle informazioni sono ricavabili dalle immagini, mentre le omissioni non semiotiche sono elevate (90 nell'interlinguistico e addirittura 167 nell'intralinguistico) perché il linguaggio colloquiale contiene molti tratti dell'oralità. Possiamo concludere questa comparazione ammettendo che sicuramente la sottotitolazione intralinguistica è migliore dell'interlinguistica. Tuttavia, dall'analisi appena effettuata, la prova risulta accettabile considerando le difficoltà imposte dalla lingua e la durata superiore della prova (10:18 minuti) rispetto alle precedenti (2:10 e 4:30 minuti). Infatti, benché il rispeaker francese lavori solitamente 15 minuti di fila

¹⁸⁴ Alle 42 omissioni non semiotiche di cui sopra, si aggiunge l'errore che non ha portato alla creazione di un sottotitolo omettendo un'altra macro-unità non semiotica.

¹⁸⁵ Indipendentemente dalla distinzione tra semiotiche e non semiotiche.

alternandosi coi colleghi, fa uno sforzo psico-cognitivo in meno (quello traduttivo) che è importante tenere in considerazione quando si fa *respeaking*. Il fatto che il rispeaker 2 sia arrivato fino alla fine senza mostrare un calo nella *performance* (così come succede in *Parole directe* per 15 minuti di programma) è un risultato importante per la nostra ricerca.

4.4.4 C à vous_L'invité

Dal momento che la migliore tra le prove de *Le dîner* (quella del rispeaker 2) è incompleta e l'unica altra prova interlinguistica a disposizione (quella del rispeaker 3) non è comparabile con l'intralinguistico, si effettuerà soltanto la comparazione relativa alla prima parte del programma, cioè *L'invité*. Si riporta comunque anche la tabella contenente i dati dell'analisi strategica della sottotitolazione intralinguistica di *Le dîner*.

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE			TOT
omissioni			errori			Tot	ripetizioni	alterazioni	Tot
178 98.9%			2 1.1%			180 57%	21 15.44%	115 84.56%	136 43%
sem.	non sem.	Tot	umani	software	Tot				
33 18.5%	145 81.5%	178 100%	2 100%	0 0%	2 100%				

alterazioni									
riduzioni					espansioni			errori	
90 78.3%					25 21.7%			0 0%	
Tot (micro-unità)					S	NS	Tot (micro-unità)	umani	software
115 74.2%					3 7%	40 93%	43 100%	0 0%	0 0%
sem.	non sem.	Tot	sem.	non sem.	Tot				
11 9.6%	104 90.4%	115 100%	9 22.5%	31 77.5%	40 100%				

Figura 35 – Risultati dell'analisi strategica dei sottotitoli intralinguistici – *C à vous_Le dîner*.

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE						TOT
omissioni			errori			Tot	ripetizioni		alterazioni		Tot	135 100%
95 100%			0 0%			95 70.4%	5 12.5%		35 87.5%		40 29.6%	
sem.	non sem.	Tot	umani	software	Tot							
18 18.95%	77 81.05%	95 100%	0 0%	0 0%	0 0%							

alterazioni										
riduzioni					espansioni			errori		Tot
29 82.86%					6 17.14%			0 0%		35 100%
omissioni		compressioni		Tot (micro-unità)	S	NS	Tot (micro-unità)	umani	software	Tot (micro-unità)
43 74.1%		15 25.9%		58 100%	5 35.7%	9 64.3%	14 100%	0 0%	0 0%	0 0%
sem.	non sem.	Tot	sem.	non sem.	Tot					
6 13.95%	37 86.05%	43 100%	8 53.3%	7 46.7%	15 100%					

Figura 36 – Risultati dell'analisi strategica dei sottotitoli intralinguistici – *C à vous_L l'invité*.

Per la prima volta, la sottotitolazione intralinguistica contiene il 5.2% di macro-unità rese in meno rispetto alla prova del rispeaker 2, cioè 7 macro-unità. Dal momento che, di conseguenza, le omissioni sono di più, non sembra strano che ci siano 11 omissioni non semiotiche in più, il che fa pendere la bilancia a favore della sottotitolazione intralinguistica perché c'è meno testo (inutile) da leggere per il telespettatore. Il fatto è che ci sono anche 5 omissioni semiotiche in più. Tuttavia, bisogna spostare l'attenzione sugli errori nell'interlinguistico (9) per stabilire se queste considerazioni sono effettivamente valide. Dei 9 errori del rispeaker 2, 3 sono gravi e risultano in sottotitoli incomprensibili, uno non è grave perché il rispeaker 2 corregge il tiro nei sottotitoli successivi, 5 sono omissioni non semiotiche. Per il rispeaker 2, le macro-unità non rese per errore rimangono 4, mentre le omissioni non semiotiche salgono a 71. La differenza delle omissioni non semiotiche si riduce quindi da 11 a 6. Se da una parte abbiamo 6 omissioni non semiotiche di scarto e dall'altra abbiamo 5 omissioni semiotiche di scarto, la bilancia pende ancora a favore dell'intralinguistico per una macro-unità non semiotica in più omessa. Alla luce di ciò, le due sottotitolazioni sono da considerarsi quasi alla pari in fatto di macro-unità rese e non rese, il che è un ottimo risultato considerato il fatto che nel sottoparagrafo 4.2.4 la prova interlinguistica era stata bocciata per la non presentabilità del TA a un pubblico italiano a causa degli errori. Procedendo, le ripetizioni questa volta sono poche (5 a fronte di 35

alterazioni nell'intralinguistico), segno di una reale difficoltà intrinseca al programma che neanche un rispeaker madrelingua francese riesce a seguire per bene. Anche all'interno delle macro-unità alterate ritroviamo un caso anomalo rispetto alle comparazioni precedenti: le espansioni nell'intralinguistico (17.14%) superano quelle dell'interlinguistico (9.1%). Però, se andiamo a vedere le micro-unità, la situazione torna 'normale' e ci accorgiamo che in realtà le micro-unità espanse sono 14 per l'intralinguistico e 20 per l'interlinguistico. Questa 'incongruenza' è spiegata dalle omissioni (62) e dalle compressioni (26), il cui peso va a controbilanciare la 'leggerezza' delle macro-unità espanse nell'interlinguistico. Evidentemente, il rispeaker 2 ha trovato più semplice riformulare il testo eliminando e comprimendo micro-unità in modo tale da potersi concentrare nell'ascolto/comprendimento del TP, mentre il rispeaker francese, avvantaggiato nella lingua, ha prestato più attenzione alla formulazione del TM1 (non c'è nemmeno un errore infatti) cercando di eliminare (43) e comprimere (15) il meno possibile. Il fatto è che i risultati interni vanno decisamente a favore dell'interlinguistico. A quanto pare l'intralinguistico, a fronte di meno omissioni (19 in meno) e compressioni (11 in meno), conta le stesse compressioni semiotiche (8) e addirittura una omissione semiotica in più. A far riequilibrare l'ago della bilancia sono ancora una volta gli errori dell'interlinguistico: una macro-unità alterata per errore e 5 errori a livello di micro-unità. Per concludere, la comparazione appena effettuata ci porta a riconsiderare la nostra 'bocciatura' iniziale promuovendo anche questa prova interlinguistica.

4.4.5 Journal_Parole directe

MACRO-UNITÀ NON RESE						MACRO-UNITÀ RESE			TOT
omissioni			errori			Tot	ripetizioni	alterazioni	Tot
101			3			104	55	195	250
97.1%			2.9%			29.4%	22%	78%	70.6%
sem.	non sem.	Tot	umani	software	Tot				
38	63	101	3	0	3				
37.6%	62.4%	100%	100%	0%	100%				

alterazioni									
riduzioni			espansioni			errori			Tot
170 87.2%			22 11.3%			3 1.5%			195 100%
omissioni	compressioni		Tot (micro-unità)	S	NS	Tot (micro-unità)	umani	software	Tot (micro-unità)
171 64.045%	96 35.955%		267 100%	4 9.5%	38 90.5%	42 100%	1 25%	3 75%	4 100%
sem.	non sem.	Tot	sem.	non sem.	Tot				
20 11.7%	151 88.3%	171 100%	17 17.7%	79 82.3%	96 100%				

Figura 37 – Risultati dell'analisi strategica dei sottotitoli intralinguistici – *Journal_Parole directe*.

La sottotitolazione intralinguistica di *Parole directe* contiene ben il 19.2% di macro-unità rese in più rispetto alla prova del rispeaker 2, cioè 68 macro-unità di differenza. Trattandosi di una trasmissione in cui la componente audio verbale predomina, è inutile effettuare le analisi contenute nei sottoparagrafi precedenti per accertare che questo scarto sia prevalentemente non semiotico, perché si può già prevedere che non lo è. Tuttavia, se la nostra attenzione si sposta sulla suddivisione interna delle omissioni, si nota che le percentuali sono quasi le stesse e ciò ci rincuora perché i rispeaker hanno effettuato (ovviamente in proporzione) quasi le stesse scelte di omissione: 37.6% di omissioni semiotiche per l'intralinguistico e 37.4% per l'interlinguistico con percentuali ovviamente complementari riguardo le rispettive omissioni non semiotiche. Sono importanti, invece, le differenze percentuali delle macro-unità non rese per errore (11.6% in più per l'interlinguistico) e delle macro-unità rese per ripetizione e traduzione letterale (11.6% in più per l'intralinguistico). Curiosamente i dati combaciano alla perfezione, ma questa volta può solo trattarsi di una coincidenza visto che le macro-unità errate in una sottotitolazione non corrispondono alle macro-unità ripetute nell'altra. Già da questi pochi dati pare ovvio che la prova intralinguistica è nettamente migliore a quella interlinguistica e quindi si dovrebbero fare molti sforzi in più per ottenere risultati simili a quelli della sottotitolazione francese. Bisogna comunque spezzare più lance a favore della sottotitolazione interlinguistica in generale e del rispeaker 2 in particolare:

- in primo luogo, come già spiegato nel sottoparagrafo 4.2.5, l'incidenza degli errori sull'intero TA è veramente bassa;
- in secondo luogo, il rispeaker ha lavorato per 15 minuti di fila così come (probabilmente) il rispeaker francese, ma si deve aggiungere lo sforzo psico-cognitivo della traduzione e la mancanza di un addestramento adeguato;

c) in terzo luogo, in una trasmissione importante come questa, rubrica mensile del telegiornale di *TF1* in cui, soprattutto, è coinvolta una possibile candidata alle elezioni presidenziali, ci deve essere necessariamente una scaletta con tutte le possibili domande o perlomeno gli argomenti che saranno trattati. Se si considera che è interesse dell'emittente televisiva e del programma fornire un buon servizio di sottotitolazione durante una trasmissione con un'*audience* presumibilmente alta, non si esclude di certo che i rispeaker possano aver ricevuto prima una scaletta del programma con tutti gli argomenti.

Detto questo, passiamo a un'analisi più approfondita delle alterazioni (195 nell'intralinguistico e 163 nell'interlinguistico). I dati relativi alle espansioni confermano ciò che si è detto riguardo gli altri programmi: la riformulazione, proprietà intrinseca al passaggio da una lingua all'altra, si espleta tramite le espansioni (10.8% in più nell'interlinguistico a livello di macro-unità) e le compressioni (10.645% in più nell'interlinguistico). Sempre per l'interlinguistico, all'interno delle riformulazioni si notano anche una notevole perdita semiotica (24.3% di compressioni semiotiche) e una notevole aggiunta semantica (28.6% di espansioni semantiche), che potrebbero risultare più che complementari se non ci fosse anche una notevole omissione semiotica (15%). Più lievi la perdita semiotica (17.7% di compressioni semiotiche e 11.7% di omissioni semiotiche) e l'aggiunte semantica (9.5% di espansioni semantiche) nell'intralinguistico. Questi risultati confermano la predilezione per una sottotitolazione *verbatim* nell'intralinguistico. Per quanto riguarda gli errori, la differenza non si fa sentire a livello di macro-unità alterate (1.5% nell'intralinguistico e 1.8% nell'interlinguistico), ma è di peso a livello di micro-unità (4 nell'intralinguistico e 31 nell'interlinguistico). In definitiva, si ribadisce la netta superiorità della sottotitolazione intralinguistica e quindi la non applicabilità della prova interlinguistica come *output*. Le ragioni sono perlopiù da ricercarsi nella preponderanza della componente audio verbale, nell'assenza di immagini di supporto alla comprensione del TP, nella pluralità di argomenti trattati e negli errori da evitare. Tuttavia, si precisa che la prova resta accettabile e perfino incoraggiante a livello sperimentale, per le ragioni di cui sopra e al sottoparagrafo 4.2.5.