

## Simultaneous Consecutive Interpreting: A New Technique Put to the Test

Miriam Hamidi et Franz Pöchhacker

Volume 52, numéro 2, juin 2007

URI : <https://id.erudit.org/iderudit/016070ar>

DOI : <https://doi.org/10.7202/016070ar>

[Aller au sommaire du numéro](#)

Éditeur(s)

Les Presses de l'Université de Montréal

ISSN

0026-0452 (imprimé)

1492-1421 (numérique)

[Découvrir la revue](#)

Citer cet article

Hamidi, M. & Pöchhacker, F. (2007). Simultaneous Consecutive Interpreting: A New Technique Put to the Test. *Meta*, 52(2), 276–289.  
<https://doi.org/10.7202/016070ar>

Résumé de l'article

L'article présente une étude expérimentale à petite échelle ayant pour objectif de tester la viabilité, voire la supériorité de l'interprétation consécutive assistée par technologie numérique, comme nouvelle méthode de travail pour l'interprétation de conférence. Cette technique, utilisée pour la première fois en 1999 par un interprète de la Direction générale de l'interprétation de l'Union européenne, consiste à enregistrer le discours original à l'aide d'un dictaphone numérique pour permettre à l'interprète, muni d'écouteurs, de le réécouter et le reproduire en simultanée. Les productions consécutives et « techno-consécutives » de trois interprètes professionnels et expérimentés (français-allemand) ont été examinées et comparées à partir de l'analyse des transcriptions, de l'autoévaluation des interprètes et de la réaction des auditeurs participant à l'expérience. Les résultats montrent que la « simultanée consécutive » génère de meilleures performances, notamment en ce qui concerne la fluidité du discours, la correspondance source cible et les déviations prosodiques. Bien que l'expérience personnelle des interprètes ainsi que leurs préférences quant au mode de travail semblent avoir influencé considérablement leur performance, les trois sujets ont facilement adopté l'interprétation consécutive assistée par technologie numérique et la considèrent comme une technique viable.

# Simultaneous Consecutive Interpreting: A New Technique Put to the Test

**MIRIAM HAMIDI**

*University of Vienna, Vienna, Austria*  
miriamhamidi@gmx.at

**FRANZ PÖCHHACKER**

*University of Vienna, Vienna, Austria*  
franz.poechhacker@univie.ac.at

## RÉSUMÉ

L'article présente une étude expérimentale à petite échelle ayant pour objectif de tester la viabilité, voire la supériorité de l'interprétation consécutive assistée par technologie numérique, comme nouvelle méthode de travail pour l'interprétation de conférence. Cette technique, utilisée pour la première fois en 1999 par un interprète de la Direction générale de l'interprétation de l'Union européenne, consiste à enregistrer le discours original à l'aide d'un dictaphone numérique pour permettre à l'interprète, muni d'écouteurs, de le réécouter et le reproduire en simultanée. Les productions consécutives et «techno-consécutives» de trois interprètes professionnels et expérimentés (français-allemand) ont été examinées et comparées à partir de l'analyse des transcriptions, de l'autoévaluation des interprètes et de la réaction des auditeurs participant à l'expérience. Les résultats montrent que la «simultanée consécutive» génère de meilleures performances, notamment en ce qui concerne la fluidité du discours, la correspondance source cible et les déviations prosodiques. Bien que l'expérience personnelle des interprètes ainsi que leurs préférences quant au mode de travail semblent avoir influencé considérablement leur performance, les trois sujets ont facilement adopté l'interprétation consécutive assistée par technologie numérique et la considèrent comme une technique viable.

## ABSTRACT

The paper reports a small-scale experimental study to test the viability or even superiority of technology-assisted consecutive interpreting as a new working method for conference interpreters. In this technique, pioneered in 1999 by an EU staff interpreter, a digital voice recorder is used to record the original speech which the interpreter then plays back into earphones and renders in the simultaneous mode. The performances of three experienced professional interpreters (French-German) in the conventional consecutive and the 'simultaneous consecutive' mode were assessed and compared on the basis of transcript analysis, self-assessment and audience response. Our findings suggest that simultaneous consecutive permits enhanced interpreting performances, as reflected in more fluent delivery, closer source–target correspondence, and fewer prosodic deviations. Though the interpreters' personal working experience and preferences appeared to have a significant influence on their performance, all three subjects easily adopted the technology-assisted interpreting mode and considered it a viable technique.

## MOTS-CLÉS/KEYWORDS

consecutive interpreting, simultaneous interpreting, technology-assisted consecutive interpreting, interpreting quality assessment, comparative analysis

## Introduction

For most of its history, interpreting did not require any technical tools. Except for the interpreter's notepad and pen, this was true until well into the twentieth century, when professional conference interpreting began to emerge. It was the introduction of new technologies, such as the use of electro-acoustic transmission systems for simultaneous interpreting, dating back to the mid-1920's (Baigorri-Jalón 1999), that gave the profession a major boost. The advent of remote simultaneous interpreting, which goes back to the late 1970s and has become a focus of attention in recent years (e.g., Mouzourakis 2006), is another impressive example of the impact of technology on the practice of interpreting.

This paper deals with yet another innovative mode of interpreting, referred to as simultaneous consecutive interpreting (or "SimConsec"; Pöschhacker, *in press*), in which storage of the original message in the interpreter's notes and long-term memory is replaced by a digital recording of the original speech, which the interpreter plays back into earphones and renders in the simultaneous mode. We describe a small-scale experimental study of conventional and technology-assisted consecutive interpreting from French into German carried out at the Center for Translation Studies of the University of Vienna. Performance quality in either mode will be analysed on the basis of transcript-based assessment and audience response.

## Pioneers and previous work

### *Ferrari and SCIC*

The idea for simultaneous consecutive interpreting is claimed by several professional interpreters. The first successful use of "digitally mastered consecutive" was reported by EU staff interpreter Michele Ferrari (2001: 3), who used his palm-size PC to interpret a press conference in Rome by Neil Kinnock, then Vice-president of the European Commission, in March 1999. Shortly after his report of this first trial in a real setting, Ferrari carried out a series of tests within the EU Joint Interpreting and Conference Service (SCIC; now DG Interpretation) (Ferrari 2002).

The aim of these tests was not only to compare traditional and simultaneous consecutive interpreting, but also to examine different devices, such as a handheld PC, a notebook with digital audio-editing software, and a digital voice recorder. The results can be summarized as follows: while simultaneous consecutive is objectively more complete and precise, it sounds "too artificial" for some language combinations (Ferrari 2002: 7). Especially when the source and target languages belong to the same linguistic family, as is the case for Spanish and Italian, for instance, Ferrari concludes that it is essential to find an appropriate playback acceleration rate in order to sound natural.

The second series of tests was also carried out within SCIC and took place in June 2003. This session, parts of which were also video-recorded, led to the general conclusion that simultaneous consecutive was a "viable possibility" and that the use of electronic devices required some practice (Vivas 2003). It is also interesting to note that the report of this second comparison session contains a note by J. Martin, another SCIC interpreter, about his successful use of a digital voice recorder in a real interpreting assignment.

*Lombardi and Camayd-Freixas*

Two court interpreters in the United States, John Lombardi (2003) and Erik Camayd-Freixas (2005), have also reported efforts to use a digital voice recorder for consecutive interpreting in court. While Lombardi tested “digital recorder assisted consecutive (DRAC)” interpreting informally and published an article about the functioning, advantages and possible drawbacks of the method in the spring edition of the NAJIT newsletter (Lombardi 2003), Camayd-Freixas carried out an experiment at Florida International University and even established an equipment label marketed by his language consulting firm (Camayd-Freixas 2005).

The experiment at Florida International University was conducted in April 2003 with 24 advanced interpreting students and beginning professional interpreters with the language combination English/Spanish. The subjects were divided into two groups and presented with a series of unrelated statements of increasing length to be interpreted in either mode and direction, i.e., from English into Spanish and Spanish into English, using either notes or the digital voice recorder. The experiment, in which both groups served as their own controls, was centred on accuracy, measured in terms of “the percentage of words missed in each statement” (Camayd-Freixas 2005: 43). Results showed higher accuracy rates when using the digital voice recorder and a decline in accuracy with conventional note-taking as statement length increased (Camayd-Freixas 2005: 43).

All interpreters who tested the simultaneous consecutive interpreting technique have reported an improvement in quality. Technology-assisted consecutive interpreting is generally praised for its increased accuracy and completeness (Gomes 2002: 8). The resulting rendition is also considered more faithful to the original in terms of intonation and liveliness (Camayd-Freixas 2005: 42). The fact that note-taking is no longer necessary allows the interpreter to devote more attention to listening and comprehension.

Technology-assisted consecutive interpreting is said to be especially useful in court, even though ethical issues, such as the question of confidentiality, have to be considered.

The fact that the interpreter would not need to ask a witness to pause for interpretation and could more accurately convey expressions of emotion is seen as one of the most striking advantages of using an electronic recording and playback device in court interpreting (Lombardi 2003: 8).

**The study**

Based in particular on Ferrari’s (2001, 2003) tests of technology-assisted consecutive interpreting, a small-scale experiment was carried out at the Vienna University Center for Translation Studies in the context of a Master’s thesis (Hamidi 2006). The study focused on the language combination French/German and was designed to address three main questions:

1. Does technology-assisted consecutive interpreting yield better results than the conventional consecutive method?
2. How does the audience respond to the new consecutive technique compared to the traditional one?

3. Is simultaneous consecutive likely to be adopted by professional interpreters as an interpreting method in its own right?

### *Design and procedure*

Answering the research questions required an experimental comparison of the two consecutive interpreting modes. Professional conference interpreters were to interpret two comparable speeches (of just under ten minutes' duration in two takes of roughly equal length), rendering the first one in the conventional consecutive mode and the second with the aid of a digital voice recorder.

To ensure the same conditions for all participating subjects, the two speeches were videotaped and the recordings subsequently shown as source material in three separate experimental sessions, which took place in late June 2005. At the end of each session, a debriefing interview was conducted with the interpreter on the basis of seven questions relating to his/her professional experience and preferred working mode as well as his/her own assessment of the two performances and the new technique.

In order to study audience response, we asked two dozen students in the B.A. programme in International Communication at the Center for Translation Studies to serve as listeners in the experimental sessions. As a result of time constraints, however, the majority of them, busy with end-of-term exams, were unable to participate. Therefore, the video-recorded interpreting performances were presented to a small and heterogeneous experimental audience, with each of three groups responding to one interpreter's performances on the basis of a questionnaire.

### *Subjects and material*

Out of the relatively small group of interpreters who met our criteria for participation (university-degree in interpreting; German as A language, French B; extensive professional experience in both interpreting modes), only three were able and willing to participate in the experiment in June 2005. The subjects, two women and one man, had at least ten years of professional experience. Two of them are members of AIIC and work mainly in the simultaneous mode, following a decline in the number of consecutive assignments in the course of their careers. The third subject does nearly as much consecutive as simultaneous interpreting. Two interpreters expressed a preference for the consecutive mode, indicating that it permits more contact with the audience and can give greater professional satisfaction. One interpreter had previously experimented with using a digital tape recorder in a real consecutive setting.

The source speeches, which were scripted for the purpose of the experiment on the basis of authentic material, were comparable in terms of content and level of language use. Both were delivered by a French native-speaker at a rate of some 130 words per minute (articulation rate: 138 and 144 wpm, respectively). The speeches dealt with French-American relations (Speech 1) and France's role in the European Union (Speech 2) and contained a comparable number of proper names (8/10), dates (6/9) and figures (6/7). On account of the differences in the number of such potential problem triggers as well as the higher articulation rate, Speech 2 may have been slightly more difficult, thus raising the challenge for working in the simultaneous consecutive mode.

The audience, who had hardly any knowledge of interpreting, consisted of a sample of nine persons aged 20 to 56 years and included two translation students, two lawyers, two psychoanalysts, a student of Arabic, a student of mathematics and a computer scientist. The audience was divided into three groups, with two female and one male listener for each experimental session.

The eight-page questionnaire given to the listeners in the course of the experimental sessions first elicited their general impression of the interpreter's performances. They were asked to indicate whether they had been able to follow the German interpretation of the original French speech 'very well,' 'well,' 'with some difficulty,' 'badly' or 'very badly' and were given space to add comments. The questionnaire also included three comprehension questions (one free recall and two multiple-choice questions) and seven-point Likert scales for rating the following six quality criteria: 'fluency of delivery,' 'clarity and cohesion,' 'quality of expression (language),' 'intonation and emphasis,' 'contact with audience' and 'confidence and professionalism.' (The seven points on the scale, labeled only with anchors for the lowest (-) and the highest assessment (+), were interpreted as numerical ratings – from '1' to '7' – in the subsequent analysis.) Having evaluated the two performances, the listeners were asked to state which of the interpretations they liked better, and why.

The electronic device used in our experiment had to be user-friendly and reasonably priced. We opted for an Olympus VN-480PC digital voice recorder with a built-in flash memory (64 MB) and a recording time capacity of up to 177 minutes in the high-quality mode. During the experiment the subjects had to operate three buttons, REC, STOP and PLAY, aside from controlling the volume; they were given a short warm-up period to test the equipment. Replay speed was set to "normal." To offset the loss in sound quality resulting from the video playback of the source speech, we used a free-standing condenser microphone instead of the clip-on microphone supplied with the voice-recorder.

### *Pilot study*

The experimental setup and materials were tested in a pilot study involving a recent interpreting graduate and a live student audience, four of whom (second semester; no French) filled in the questionnaire. No problems with the procedure were found.

### *Output analysis*

In addition to the post-test interviews conducted with the interpreters and the questionnaire-based evaluation by the audience, we used output analysis to permit triangulation of data and implement a multi-method approach to interpreting quality assessment (cf. Pöchhacker 2001). For this purpose, the source speeches and video-recorded interpretations were transcribed (orthographically) and analysed with regard to such parameters as quality of expression, fluency and prosody. All pauses with a duration of more than one second were marked in the transcript and measured with the help of audio analysis software (PRAAT). In the case of filled pauses, the hesitation sounds were also noted. Pitch movement was first analysed on an auditive basis and subsequently checked with PRAAT. (For a detailed description and illustrations, see Ahrens 2005).

## Results

### *Interpreters*

Except for the sound quality, which was considered quite poor in two of the three experimental sessions, all three subjects were relatively comfortable with the simultaneous consecutive interpreting technique. The post-test interviews also showed that Int-1 and Int-3 felt more comfortable with the second speech and interpreting mode, whereas Int-2 felt better with Speech 1 and the conventional consecutive mode. Asked which of their performances had been superior, Int-1 and Int-3 indicated for Speech 2, interpreted in the simultaneous consecutive mode.

According to the interpreters, the main advantages of the simultaneous consecutive technique are that it is "less strenuous" and that it permits the interpreter to listen to the original twice, which is especially relevant when speakers do not express themselves in their mother tongue and have a strong accent. The fact that one has to "translate everything," even in long and redundant speeches, was seen as the most important drawback. One interpreter criticised a certain loss of the "human element" in simultaneous consecutive. Nevertheless, all three subjects agreed that they could imagine using digital equipment for real consecutive interpreting assignments.

### *Audience response*

#### *Int-1*

For all criteria under study the experimental audience gave a more favourable assessment to Int-1's second interpreting performance (Speech 2). All three listeners indicated that they had been able to follow Speech 2 'very well,' an impression that was corroborated by the results of the comprehension test. Speech 1, however, elicited a response of 'very well' only once, and one listener was able to follow it only 'very badly.'

Whereas the best rating achieved by Int-1 for Speech 1 was 6 (once for 'contact with audience'), the rendition of Speech 2 received the highest rating (7 on the seven-point scale) at least once for each criterion. The lowest rating for Speech 1 was 2 (for 'fluency'), while the worst rating for Speech 2 was 4 (once for each criterion).

Asked about their preferences for one or the other of Int-1's performances, the audience liked the simultaneous consecutive version better, as it was regarded as more fluent and understandable, and much livelier. It is also interesting to note that the two psychoanalysts in the audience for Int-1 explained their preference for the second performance by observing that it reflected less tension and appeared to preserve the "affect" of the original.

#### *Int-2*

With a rating of at least 4, both performances by Int-2 received a consistently high average assessment. While Speech 1 yielded the top rating five times (twice for 'fluency,' once for 'intonation,' and three times for 'confidence and professionalism'), the second interpreting performance was rated as highly only four times ('fluency,' 'clarity,' 'expression,' and 'confidence and professionalism').

Although the two renditions by Int-2 reached relatively similar results, the average assessment by the audience showed a slight preference for the conventional consecutive

version. Two out of three raters liked the second performance better, and one listener found both speeches equally good. The answers given to the comprehension questions showed that the audience was able to follow both performances well.

The preference for the interpretation of Speech 1 was mainly attributed to the natural and fluent presentation and to better target-language expression. One listener found it irritating to hear the original of Speech 2 through the interpreter's earphones. Int-2 was indeed the only interpreter to use just one earphone during playback for simultaneous consecutive.

### *Int-3*

While a single listener was able to follow the first rendition 'very well' (the two others indicated that they were only able to do so 'with some difficulty'), the second performance received a better assessment ('very well,' 'well,' 'with some difficulty'). This overall impression was confirmed by the ratings on the seven-point scales: the first performance achieved the top rating five times (incl. three times for 'quality of expression'), whereas the rendition of Speech 2 was rated as highly eight times and in six instances received the second-highest rating on the seven-point scale.

In both cases, each of the three raters assigned the lowest rating to 'contact with audience,' which is understandable as Int-3 did not look up even once. Given the high(er) rate of correct answers to the comprehension questions for Speech 2, it can be assumed that the audience could follow the second rendition more easily. Two of the three listeners indeed preferred the second performance because of the more fluent and lively delivery. According to one listener, the interpreter obviously felt far more comfortable during the rendition of Speech 2.

### *Transcript- and video-based analysis*

#### *Fluency*

The complex criterion of fluency (cf. Pradas Macías 2006) was analysed with reference to speaking speed (in words per minute) and pauses. The 'articulation rate' was calculated by dividing the total number of words by total speech time minus pause time (i.e., the total duration of pauses lasting more than one second). Results are shown in Table 1.

The comparison between the German interpretations and the French originals yielded an unexpected result: while no relevant differences in duration were found in the simultaneous consecutive mode, two renditions of Speech 1 were considerably longer than the source speech – by 37% (13'45 vs. 9'45) for Int-1 and by 40% (13'25 vs. 9'35) for Int-3. Int-2's simultaneous consecutive interpretation, totalling 1151 words, was slightly shorter than the conventional consecutive performance (9'30 vs. 9'35).

TABLE 1

**Length and speed**

|                         | Int-1 | Int-2 | Int-3 |
|-------------------------|-------|-------|-------|
| <b>Speech 1</b>         |       |       |       |
| Duration (min'sec)      | 13'45 | 9'35  | 13'25 |
| Number of words         | 1362  | 1151  | 1390  |
| Speech rate (wpm)       | 99    | 120.1 | 103.6 |
| Articulation rate (wpm) | 111   | 128.6 | 119   |
| <b>Speech 2</b>         |       |       |       |
| Duration (min'sec)      | 9'45  | 9'30  | 9'35  |
| Number of words         | 975   | 993   | 1110  |
| Speech rate (wpm)       | 100   | 104.1 | 115.8 |
| Articulation rate (wpm) | 113   | 127.2 | 121.9 |

In the traditional consecutive mode, Int-2 used 15% fewer words than Int-1 (1151 vs. 1362) and 17% less than Int-3's 1390 words. Int-3 was also the wordiest in the simultaneous consecutive mode (1110 words), while Int-1 and Int-2 used 975 and 993 words respectively. The speech rates for the simultaneous consecutive performances (i.e., word count divided by duration in minutes) range from 100 (Int-1) to 116 (Int-3) words per minute. Adjusted for pauses, these figures correspond to articulations rates that are quite consistent for each subject in the two experimental conditions: 111/113 wpm for Int-1, 129/127 wpm for Int-2 and 119/122 for Int-3.

The speech and articulation rates suggest that Int-2's rendition of Speech 1 was very fluent and free of longer pauses (both figures are similarly high and higher than those for the second performance). The difference between Int-2's speech and articulation rates for the second rendition (104.5 vs. 127 wpm) reveals that this interpreter was not particularly at ease during the interpretation of Speech 2. In the case of Int-3 there is a substantial difference between these rates in the consecutive mode, which indicates a significant number of pauses.

The overall pattern of pauses in the interpreters' output (Table 2) shows that Int-1 and Int-3 interrupted their delivery more often during their traditional consecutive renditions (54 and 39 times, resp.), whereas Int-2 paused more often in the simultaneous consecutive mode (59). In contrast to Int-2, Int-3 also made more long pauses (> 1.5 sec) during the first rendition. Int-3's second performance contains only seven long pauses, compared to 32 in the interpretation of Speech 1. While Int-3 also has a higher pause time in the rendition of Speech 1 (13% vs. 5%), this ratio is lower for the conventional consecutive rendition by Int-2 (6,6% vs. 17,8%). Int-1 and Int-3 made their longest pauses (4.11 and 14.89 seconds, resp.) in their first consecutive interpretation. By the same token, the higher number of voiced hesitations for each of the three interpreters was recorded for Speech 1.

TABLE 2

## Pauses

|                                           | Int-1    |          | Int-2    |          | Int-3    |          |
|-------------------------------------------|----------|----------|----------|----------|----------|----------|
|                                           | Speech 1 | Speech 2 | Speech 1 | Speech 2 | Speech 1 | Speech 2 |
| Number of pauses                          | 54       | 43       | 27       | 59       | 38       | 18       |
| Percentage of total duration              | 10.8     | 11.6     | 6.6      | 17.8     | 13       | 5        |
| Number of long pauses (>1.5 sec)          | 28       | 28       | 10       | 35       | 32       | 7        |
| Longest pause (sec)                       | 4.11     | 2.71     | 2.26     | 3.05     | 14.89    | 3.46     |
| No. of filled pauses (voiced hesitations) | 179      | 150      | 22       | 9        | 43       | 22       |

*Source–target correspondence*

Source–target correspondence was assessed in terms of substantial departures of the interpretation from the original. Based on the three classic types of deviations identified by Barik (1975/2002), the transcripts were examined for two levels of omission, addition, and substitution: those that do not substantially alter the meaning, and those that do ('meaning-relevant' changes).

Example of a meaning-relevant omission:

- Original: ... *ce sont deux Français qui ont initié le projet de construction européenne*,  
(it was two Frenchmen who began the project of building Europe,  
*Robert Schuman et Jean Monnet....*  
Robert Schuman and Jean Monnet)  
Int-1: *Zwei Namen ...äh... sind da zu nennen: Schuman und Monnet...*  
(Two names ...uh... must be mentioned: Schuman and Monnet)

The following example, in contrast, was not considered a meaning-relevant omission:

- Original: *La crise de Suez en 1956, les épisodes de décolonisation dans les années 60*  
(The Suez crisis in 1956, the decolonization events in the 60s  
*ou encore la guerre du Vietnam...*  
or later on the Vietnam war)  
Int-1: *Ich verweise auf die Suez-Krise, auf die Entkolonialisierung in den sechziger*  
(I refer you to the Suez crisis, to decolonization in the 60s,  
*Jahren, auf den Vietnamkrieg.*  
to the Vietnam war.)

While the three interpreters made relatively few meaning-relevant departures from the original, the number of substantial additions was generally higher in conventional consecutive than in the simultaneous consecutive mode. For each of the interpreters, fewer additions were found in Speech 2.

While Int-1 made more additions and substitutions in Speech 1 (56 and 49), the number of omissions is quite even for both speeches, with more meaning-relevant omissions in Speech 2 than in Speech 1 (8 vs. 1).

Similar results were found for Int-3, who made more additions (58), substitutions (56) and omissions (24) in the conventional consecutive version. The number of meaning-relevant additions in the rendition of Speech 1 (12) is three times higher than for the simultaneous consecutive mode (4).

The incidence of Int-2's meaning-relevant omissions (8 in Speech 1 vs. 12 in Speech 2), total substitutions (71 vs. 50) and total additions (39 vs. 27) is comparable to those of Int-1, though the difference between the two renditions is less pronounced.

*Quality of expression*

The quality of the interpreters' target-language production was analysed with reference to grammatical, syntactical and lexical errors, false starts, repetitions and slips of the tongue (Table 3).

TABLE 3

**Quality of expression**

|                     | Int-1    |          | Int-2    |          | Int-3    |          |
|---------------------|----------|----------|----------|----------|----------|----------|
|                     | Speech 1 | Speech 2 | Speech 1 | Speech 2 | Speech 1 | Speech 2 |
| Grammatical errors  | 8        | 8        | 9        | 10       | 9        | 5        |
| Syntactical errors  | 3        | 5        | 8        | 10       | 8        | 5        |
| Lexical errors      | 5        | 2        | 8        | 11       | 3        | 0        |
| False starts        | 17       | 13       | 10       | 13       | 20       | 17       |
| Repetitions         | 11       | 2        | 7        | 8        | 10       | 10       |
| Slips of the tongue | 1        | 3        | 4        | 5        | 3        | 6        |
| Reformulations      | 5        | 0        | 3        | 3        | 9        | 2        |
| Total               | 50       | 33       | 49       | 60       | 62       | 45       |

Overall, there is a predominance of expressive flaws in the conventional consecutive mode for Int-1 and Int-3, while Int-2 had a higher number in the simultaneous consecutive mode. As can be seen from Table 3, the most frequent phenomenon affecting quality of expression were false starts. All interpreters made at least 10 false starts in each of their performances. While Int-1 and Int-3 made fewer false starts in their second renditions, Int-2 had fewer in the traditional consecutive mode.

With regard to repetitions, another frequent phenomenon affecting the interpreters' output quality, Int-3's interpreting performances yielded an entirely even result: their number is the same in either mode. Int-2, however, had one repetition more in Speech 2, while Int-1 made only two repetitions in the simultaneous consecutive mode (vs. 11 in conventional consecutive).

Table 3 shows that Int-1 and Int-2 made almost the same number of grammatical mistakes in both speeches, whereas Int-3's performance in this respect is noticeably better in Speech 2. A similar pattern was found for Int-3 regarding syntactical errors: unlike the two other subjects, who made more errors in their second renditions, Int-3 made fewer syntactical errors in the rendition of Speech 2.

While not a single lexical error is in evidence in Int-3's performance in the simultaneous consecutive mode, Int-2's second rendition features a higher number of errors at the lexical level. Int-1's results are similar to those of Int-3, who made more lexical errors in Speech 1.

The number of speech errors (slips of the tongue) is generally low, and their distribution is similar for the three interpreters, that is, fewer slips were committed in the traditional consecutive mode. The highest number of slips was recorded for Int-3, who made 6 such speech errors in the rendition of Speech 2.

With the exception of Int-2, repairs were more frequent in the conventional consecutive mode, with Int-1 and Int-3 repairing roughly five times as often as in the simultaneous consecutive mode.

### *Prosody*

Prosodic phenomena were analysed with reference to the features used in Shlesinger's (1994) study of "interpretational intonation," in particular tentative final pauses, pauses within constituents, incompatible stress, segment lengthening and acceleration (Table 4).

TABLE 4  
**Prosody**

|                                       | <b>Int-1</b> |          | <b>Int-2</b> |          | <b>Int-3</b> |          |
|---------------------------------------|--------------|----------|--------------|----------|--------------|----------|
|                                       | Speech 1     | Speech 2 | Speech 1     | Speech 2 | Speech 1     | Speech 2 |
| Tentative final pauses                | 19           | 20       | 7            | 6        | 13           | 10       |
| Pauses within constituents            | 28           | 22       | 10           | 28       | 18           | 12       |
| Incompatible stress                   | 21           | 16       | 6            | 5        | 6            | 4        |
| Elevated pitch at end of meaning unit | 33           | 30       | 8            | 7        | 7            | 7        |
| Segment lengthening                   | 20           | 4        | 0            | 0        | 0            | 3        |
| Acceleration                          | 2            | 2        | 0            | 0        | 0            | 2        |

The overall pattern of prosodic deviations is rather variable, but a number of categories show fewer occurrences in the simultaneous consecutive mode. In the interpretations by Int-1, where such phenomena were most numerous, there were fewer instances of lengthened segments (20 vs. 4) and incompatibility of stress with semantic contrast (21 vs. 16) in Speech 2. The number of accelerations was the same in both speeches, as was the number of tentative final pauses (19 vs. 20).

Int-2 and Int-3 showed fewer salient prosodic features, their number being indeed negligible for Int-2, whose renditions contained neither lengthening of segments nor accelerations. For Int-3, the number of tentative final pauses (13 vs. 10) and instances of stress incompatibility (6 vs. 4) is higher in Speech 1; the opposite is true for accelerations (0 vs. 2) and segment lengthening (0 vs. 3). Pauses within constituents ("non-functional pauses") were predominant in the conventional consecutive performances of Int-1 (28) and Int-3 (18), whereas Int-2 paused more often within constituents in Speech 2 (10 vs. 28).

### *Eye contact with the audience*

The video recordings of the interpreters' performances enabled us to analyse the subjects' eye contact with the audience. While Int-3 did not look up even once, neither in the conventional consecutive nor in the simultaneous mode, Int-1 and Int-2 made more or less regular eye contact with their listeners.

In this respect, Int-1 was the most communicative by far, with 160 instances of eye contact in Speech 1 and 63 in Speech 2 (Table 5). Int-2 had 51 instances of eye contact during Speech 1 and only 9 during Speech 2. Most of these glances at the

audience were of very short duration, but there were also quite a few 'long eye contacts,' lasting some three or four seconds. Most of these – as well as the longest – were recorded for Int-1. Int-2 had 12 instances of long eye contact during conventional consecutive and only 3 in the simultaneous consecutive mode.

TABLE 5

**Eye contact**

|                                 | <b>Int-1</b> |          | <b>Int-2</b> |          |
|---------------------------------|--------------|----------|--------------|----------|
|                                 | Speech 1     | Speech 2 | Speech 1     | Speech 2 |
| Instances of eye contact        | 160          | 163      | 51           | 9        |
| No. of long eye contacts        | 20           | 22       | 12           | 3        |
| Percentage of long eye contacts | 12.5         | 35       | 24           | 33       |
| No. of eye contacts/minute      | 11.6         | 6.4      | 5.3          | 0.9      |
| Longest eye contact (sec)       | 15           | 4        | 3            | 2        |

*Confidence and professionalism*

The video recordings were a rich source for the analysis of the interpreters' body language (used as an indicator of confidence and professionalism) and showed that body movements were generally more vivid and apparently unconscious in the simultaneous consecutive mode.

The most conspicuous kinesic behaviour was recorded for Int-1, who made 17 different types of body movement during each rendition. These movements, which also involved nearby objects, were rather frequent (about 40 for each speech) and reflected some regularity. While Int-1's "object-adaptors" (Poyatos 1987/2002) in the conventional consecutive mode were directed to utensils like pencil, glasses and notepad, they mainly involved the voice recorder and its accessories during Speech 2. The most salient feature was the incidence of apparently unconscious movements during simultaneous consecutive (e.g., foot-tapping), which conveyed a sense of agitation that was noted also in the audience assessment.

Int-2, by comparison, seemed calmer and showed far less kinesic behaviour. During the renditions of the two speeches, Int-2 only used "language-markers" (Poyatos 1987/2002), i.e., finger or hand movements to support the meaning of the concurrently pronounced words. Int-2's body language was nonetheless more intense during simultaneous consecutive. While interpreting Speech 2, Int-2 also touched the notebook and other objects on the table.

Int-3's monotonous intonation and the absence of eye contact in both interpreting situations was matched by relatively passive body language. The few salient movements made by this interpreter seemed nevertheless more vivid in the rendition of Speech 2 (pulling at voice recorder and microphone cord, pushing back pullover sleeves, rocking the chair).

**Conclusions**

The study reported here sought to test the viability or even superiority of digital voice recorder-assisted consecutive interpreting as a new working method for conference interpreters. In our attempt to compare experienced professionals' performance in

conventional consecutive interpreting with notes and the new simultaneous consecutive mode, we came up against the vexing issue of quality assessment in interpreting (cf. Pöchhacker 2001). Nevertheless, based on the triangulation of data collected from multiple sources, we managed to obtain a fairly robust pattern of results. Even though the trials were conducted on a relatively modest scale and under some unfortunate methodological constraints, our findings for the performances of three experienced professional conference interpreters working from French into their A language, German, should be of interest to practitioners and researchers alike.

In two out of three cases, the transcript- and video-based analysis showed that digital voice recorder-assisted consecutive permits enhanced interpreting performances, reflected in more fluent delivery, closer source–target correspondence, and fewer prosodic deviations. The findings from our comprehensive output-based assessment were corroborated by the favourable response of the experimental audience to performances in the technology-assisted mode. Despite the small size and heterogeneity of the groups, audience ratings for overall impression and various aspects of presentation were surprisingly consistent. These judgments were furthermore confirmed by the self-assessment of the three interpreters participating in this study. Although their personal working experience and preferences appeared to have a significant influence on their performance, all three subjects easily adopted the technology-assisted interpreting mode and considered it a viable technique.

While these findings are in line with those of previous trials, there is clearly a need for more research to test the scope of application of this new interpreting mode. Various constraints and variables, such as settings, delivery styles, input speed, language combinations and equipment features should be built into the design of future studies. Furthermore, the implications of this innovative interpreting mode for training and professional practice, in conference as well as community-based settings, need to be assessed.

## REFERENCES

- AHRENS, B. (2005): "Prosodic phenomena in simultaneous interpreting: A corpus-based analysis," *Interpreting* 7-1, p. 51-76.
- BAIGORRI-JALÓN, J. (1999): "Conference Interpreting: From Modern Times to Space Technology," *Interpreting* 4-1, p. 29-40.
- BARIK, H. C. (1975/2002): "Simultaneous Interpretation: Qualitative and Linguistic Data," in PÖCHHACKER, F. and M. SHLESINGER (eds.), *The Interpreting Studies Reader*, London/New York, Routledge, p. 79-91.
- CAMAYD-FREIXAS, E. (2005): "A Revolution in Consecutive Interpretation: Digital Voice-Recorder-Assisted CI," *The ATA Chronicle* 34, p. 40-46.
- FERRARI, M. (2001): "Consecutive simultaneous?," *SCIC News* 26, p. 2-4, <[http://scic.cec.eu.int/scicnews/2001/011121/default\\_26.htm](http://scic.cec.eu.int/scicnews/2001/011121/default_26.htm)> (accessed 28 September 2006).
- FERRARI, M. (2002): "Traditional vs. 'simultaneous' consecutive," *SCIC News* 29, p. 6-7, <[http://scic.cec.eu.int/scicnews/2002/020130/default\\_29.htm](http://scic.cec.eu.int/scicnews/2002/020130/default_29.htm)> (accessed 28 September 2006).
- GOMES, M. (2002): "Digitally mastered consecutive. An interview with Michele Ferrari," *Lingua franca* 5-6, p. 6-10, <[http://www.europarl.europa.eu/interp/online/lf99\\_one/v05\\_no6/page2.html](http://www.europarl.europa.eu/interp/online/lf99_one/v05_no6/page2.html)> (accessed 28 September 2006).
- HAMIDI, M. (2006): *Simultanes Konsekutivdolmetschen. Ein experimenteller Vergleich im Sprachenpaar Französisch-Deutsch*, MA thesis, University of Vienna.

- LOMBARDI, J. (2003): "DRAC Interpreting: Coming Soon To A Courthouse Near You?," *Proteus* 12-2, p. 7-9, <[http://www.najit.org/proteus/PDFVersions/Proteus\\_Spr03%20web.pdf](http://www.najit.org/proteus/PDFVersions/Proteus_Spr03%20web.pdf)> (accessed 28 September 2006).
- MOUZOURAKIS, P. (2006): "Remote Interpreting: A technical perspective on recent experiments," *Interpreting* 8-1, p. 45-66.
- POYATOS, F. (1987/2002): "Nonverbal Communication in Simultaneous and Consecutive Interpretation: A theoretical model and new perspectives," in PÖCHHACKER, F. and M. SHLESINGER (eds.), *The Interpreting Studies Reader*, London/New York, Routledge, p. 235-246.
- PÖCHHACKER, F. (2001): "Quality Assessment in Conference and Community Interpreting," *Meta* 46-2, p. 410-425.
- PÖCHHACKER, F. (in press): "'Going simul?' Technology-assisted consecutive interpreting," in BAO, C. et al. (eds.) *Proceedings of the MIIS Anniversary Conference*, 9-11 September 2005.
- PRADAS MACÍAS, M. (2006): "Probing Quality Criteria in Simultaneous Interpreting: The role of silent pauses in fluency," *Interpreting* 8-1, p. 25-43.
- SHLESINGER, M. (1994): "Intonation in the Production and Perception of Simultaneous Interpretation," in LAMBERT, S. and B. MOSER-MERCER (eds.), *Bridging the Gap. Empirical Research in Simultaneous Interpretation*, Amsterdam/Philadelphia, John Benjamins, p. 225-236.
- VIVAS, J. (2003): "Simultaneous consecutive: Report on the comparison session of June 11. 2003," SCIC B4/JV D2003, Brussels, European Commission, Joint Interpreting and Conference Service.