

## Information Processing among Conference Interpreters A Test of the Depth-of-Processing Hypothesis

Sylvie Lambert

Volume 33, numéro 3, septembre 1988

URI : <https://id.erudit.org/iderudit/003380ar>

DOI : <https://doi.org/10.7202/003380ar>

[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

Lambert, S. (1988). Information Processing among Conference Interpreters: A Test of the Depth-of-Processing Hypothesis. *Meta*, 33(3), 377–387.  
<https://doi.org/10.7202/003380ar>

# INFORMATION PROCESSING AMONG CONFERENCE INTERPRETERS : A TEST OF THE DEPTH-OF-PROCESSING HYPOTHESIS

SYLVIE LAMBERT

*University of Ottawa, Ottawa, Ontario, Canada*

In the thinking of Craik and Lockhart (1972), the existing dichotomy in cognitive psychology between short- and long-term memory stores is more a function of different forms of coding processing than it is a substantive difference in separate stores. For them, trace durability is a function of the way in which the material is encoded. It is, in other words, the depth of the analyses required to encode the input which determines retention, and greater degrees of semantic or cognitive analyses are supposedly performed at deeper levels in the hierarchy. Their depth-of-processing hypothesis is therefore presented as a hierarchical series of processing stages through which incoming information passes.

By applying the Craik and Lockhart model to tasks which are highly familiar to conference interpreters, namely, listening, shadowing, simultaneous interpretation, and consecutive interpretation (see definitions below), and by measuring the retentive ability of interpreters following each task, an attempt was made, with the present study, at determining which specific tasks require deeper or shallower processing for the interpreter.

The hypothesis that some tasks may require deeper levels of processing than others is justified by the features of each task as provided below.

**Shadowing** is a paced, auditory tracking task which involves the immediate vocalization of auditorily presented stimuli, in other words, repeating word-for-word, and in the same language, a message presented to a subject through headphones. This technique has often been used as means of studying selective attention, but is also part of the training for interpretation candidates who must learn to listen and speak simultaneously before attempting any code-switching as in simultaneous interpretation.

Norman (1976) distinguishes two types of shadowing, namely *phonemic shadowing*, where the subject repeats each sound as soon as s/he hears it, without waiting for the completion of a proposition or even an entire word, so that the shadower stays "right on top" of the speaker. The second form of shadowing, *phrase shadowing*, involves repetition of speech at longer latencies — more precisely from 250 milliseconds and up — with the shadower waiting for a chunk of verbal information or a propositional phrase before shadowing.

A study by Chistovitch, Aliakrinskii and Abilian (1960) suggested that subjects who shadowed at longer latencies showed superior recall of the shadowed material. It was hypothesized that subjects, when asked to shadow with understanding, used this lag to analyze the content of the material, as opposed to those who shadowed without. Therefore, in the present experiment, subjects were asked to shadow phonemically in

order to minimize the lag as much as possible and thus decrease their likelihood of analyzing the content of the presented material.

**Simultaneous interpretation** at first glance may seem like an extension of phrase shadowing in that the interpreter begins speaking once s/he has heard a chunk or propositional phrase in the stimulus passage, except that, in this case, the stimulus material is in one language, usually the interpreter's passive or B language, and s/he is expected to "translate" or interpret it simultaneously into his or her dominant language (also known as "A" language).

The mere addition of the translation variable gives rise to one of the most demanding human information processing tasks in that the interpreter has to juggle the following activities :

- a) s/he receives part of sentence ("chunk") through the headphones ;
- b) s/he begins translating and conveying chunk 1 ;
- c) at the same time as s/he is vocalizing chunk 1, chunk 2 is being processed auditorily and stored until chunk 1 has been dealt with. According to Gerver (1964), the interpreter must be able to "hold" chunk 2 in some sort of echoic or phonemic store until chunk 1 has been transmitted. Furthermore, while emitting the translation of chunk 1, the interpreter is continuously monitoring his/her output to ensure its correctness.

In 1978, Barbara Moser presented a working model of simultaneous interpretation based on Massaro's 1975 model. Moser refers to the Massaro model as an attempt to describe the temporal flow of auditory information, beginning with the acoustic signal that arrives at the ear of the listener and ending with some form of mental representation of the message.

At smaller gatherings, where there is no booth equipment available, and/or where the material is confidential and/or highly technical, the other service available is *consecutive interpretation*. During consecutive interpretation, the interpreter hears a speech delivered in French, for example. The interpreter takes notes concurrently, normally into his or her target language (English, for example). When the speaker chooses to pause, or when the entire statement has been pronounced, the interpreter is called on to render a consecutive interpretation of the speech delivered by the delegate.

Under normal working conditions, interpreters use their notes to recall and reconstruct the original statement. But once the delivery is over, the notes can be discarded and the message forgotten since the interpreter is not asked to reproduce that message a second time. In this experiment however, following the consecutive delivery, (*i.e.* restitution of original message using the notes as cues), subjects were asked to turn their notes over to the experimenter and then recall as much of the speech as they could remember, something that would normally never happen in professional settings.

All forms of oral translation activity have common basic tasks such as listening, translating, and rendering. However, *consecutive interpretation* is special because of added factors such as the temporal sequence of the translation ; the fact that the consecutive delivery can be considered as some type of interpolated activity prior to recall ; the fact that subjects are exposed to one complete rehearsal of the text during the consecutive delivery ; and finally, that the fact the notes *per se* serve as visual cues. Thus, consecutive interpretation draws on cognitive faculties of memory and attention which are not typical of other forms of translation.

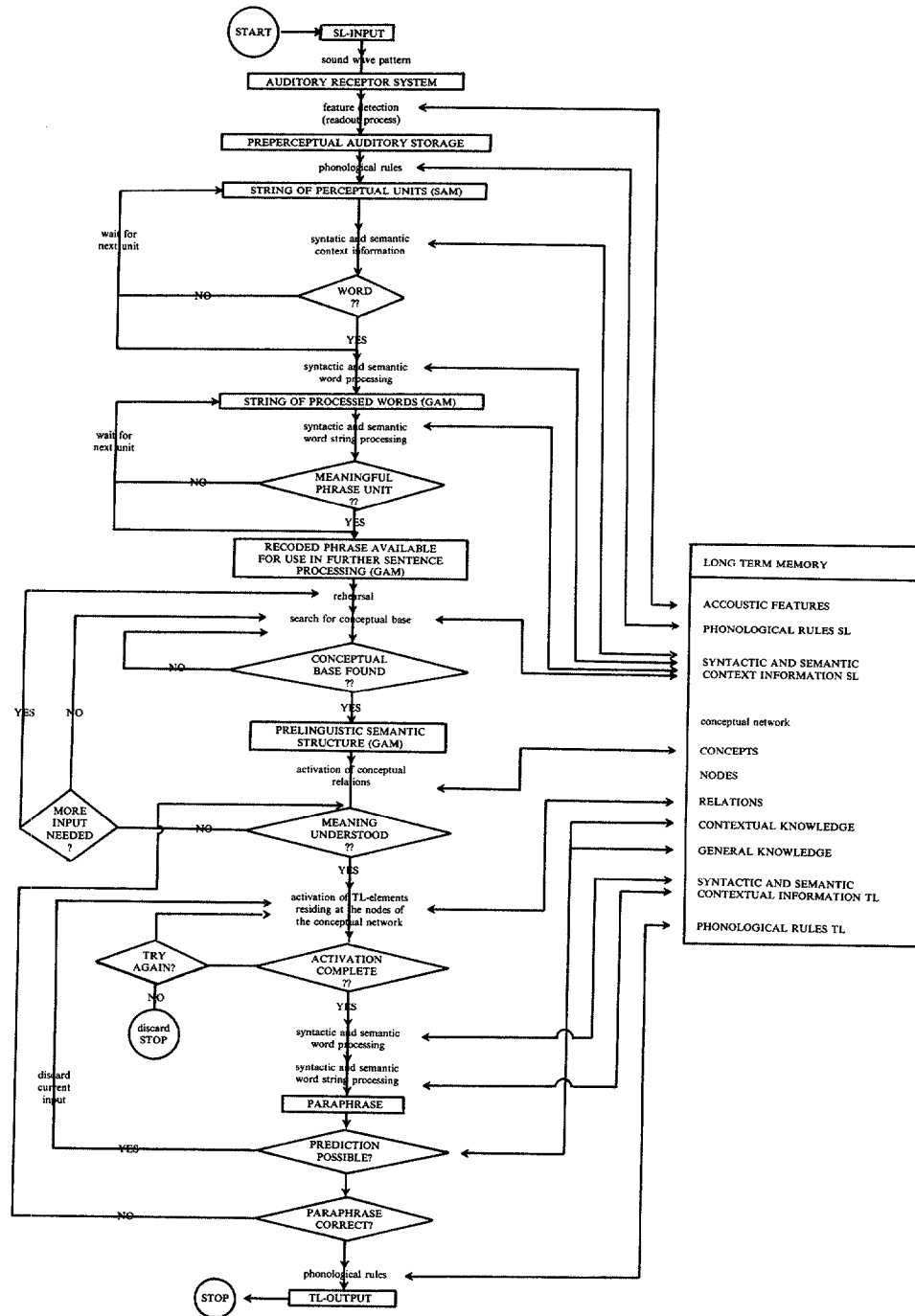


Fig. 1: A Processing Model of Simultaneous Interpretation (Barbara Moser, University of Innsbruck)

Finally, **listening** is a task implicit in shadowing, simultaneous and consecutive interpretation, assuming that the interpreter listens in order to shadow and interpret simultaneously or consecutively. Therefore, listening was used in the present experiment as a control condition, or as the common denominator for all the other conditions in the experiment. By having subjects simply listen to a text without performing any other concurrent task, recall following listening should thus serve as a baseline for establishing which tasks required more or less deep processing for the interpreter.

#### HYPOTHESES UNDERLYING THE STUDY

In the present experiment, an attempt was made to retain the actual real-life features of each task. For instance, for the consecutive interpretation condition, subjects were left free to rehearse the input material during the consecutive delivery *per se*. Since the rehearsal of the delivery itself is interpolated between the end of the stimulus input and onset of recall, it was hypothesized that recall following consecutive interpretation would yield the highest retention scores. In other words, consecutive interpretation would represent the deepest form of processing for the interpreter because it demands the greatest amount of sustained attention and effort. In addition, consecutive interpretation involves an additional effort demanded by note-taking, an activity more likely to strengthen the learning process than detract from it, and where visual cues are provided by the notes *per se*.

It was also hypothesized that recall following simultaneous interpretation would be greater than that following shadowing because it apparently calls for deeper processing than shadowing. This hypothesis was based on several studies carried out on the shadowing technique, three of which will be mentioned here.

One of the first to use the shadowing technique was Colin Cherry in 1953. Cherry found that when subjects shadowed a message presented in one ear, they were apparently unaware of and oblivious to the message presented in the other unattended ear. Moreover he found that subjects remembered very little of the shadowed message. According to Cherry, therefore, although a message can be shadowed accurately, little seems to be retained of its informational content. This finding has been confirmed by Waugh and Norman (1965) who concluded that, "...when subjects shadow verbal material, they have little retention of the material just shadowed even though this material must have undergone considerable processing by the nervous system in order for it to be heard and repeated verbally" (p. 18).

Since the shadowing technique, in contrast to listening, involves the overt vocalization of a message, Carey (1971) wondered whether shadowing would yield higher retention scores than listening. To answer this question, he tested a group of 36 listeners and 36 shadowers to determine whether shadowing a story improves retention of the details of that story over the level achieved by listeners. Most of Carey's testing procedure was replicated in the present experiment and will therefore be presented below in some detail.

Six passages of prose, each approximately 250 words long, were recorded at speeds of 1 word per second, 2 words per second, and 3 words per second. For each of the six stories, 3 recognition tests were administered following the listening or shadowing task : 1) a word or lexical recognition test ; 2) a semantic content recognition test ; and 3) a syntax recognition test. In the *lexical* recognition test, 3 nouns, 3 verbs, and 3 adjectives were selected from the text. In addition, synonyms for 3 other nouns, 3 other verbs, and 3 other adjectives were introduced. These 18 words were arranged in a vertical column in random order and subjects were asked to indicate whether or not the items had appeared in the original passage. In the *semantic* recognition test, 14 multiple-choice questions on the informational content of the story were presented, with 4

possible choices per question. The order of the questions did not correspond to the order in which the relevant information occurred in the story. Finally, the *syntax* recognition test consisted of a phrase or sentence taken *verbatim* from the original text, in addition to 2 paraphrases of the same phrase or sentence which still preserved the meaning of the original. Subjects were asked to listen or shadow texts presented at various speeds and then complete the lexical, semantic, and syntax recognition tests.

Carey's hypothesis, labelled as the "shadowing facilitation hypothesis", states that when shadowing is successful and when the shadowing response monitored by the subject is identical to the input, then shadowing improves retention. Since shadowing requires the strategies of listening as a prerequisite but demands additional processing from the subject, then retention scores after shadowing should be higher than after listening.

Carey's results indicated that increasing the rate of presentation (*i.e.* from 1 to 2, to 3 words per second) generally reduced retention scores for both listeners and shadowers, although the rate effect was more pronounced for shadowers. At the slowest rate (1 word per second), shadowers tend to obtain higher retention scores than listeners on both word and content tests. At 2 words per second, shadowers' scores on word and content tests are still higher than listeners' retention scores, but not significantly so. Finally, at 3 words per second, differences disappear for word recognition but go in the opposite direction for content.

Carey concluded that when shadowing a slow message, subjects have enough time to perceive the words and structure as well as plan and execute the shadowing response; but at faster rates, the subject must both perceive and speak in less time. These results corroborate those of the study conducted by Chistovitch, Aliakrinskii and Abilian (1960) mentioned earlier.

However, Gerver (1974) showed that comprehension was greater following the listening task than after the shadowing task. Gerver's study, unlike Carey's, included translation as a variable. He asked nine trainee interpreters to a) listen to, b) shadow, and c) simultaneously interpret into English, three passages of French prose. Gerver found that significantly higher comprehension scores were obtained following the listening condition (58%) than after simultaneous interpretation (51%) and after shadowing (43%). Gerver argues that although shadowing may be considered as a type of rehearsal or repetition, the possibility of task-sharing and interference should be considered as well. He adds that simultaneity of listening and speaking does not impair comprehension and recall since retention scores following simultaneous interpretation were higher than after shadowing. This outcome indicated that the translation factor ought to be examined more closely.

When explaining the difference between shadowing and simultaneous interpretation, Gerver argues that shadowing does not require the complex analysis and transformation of input necessary in translation and that transferring the input from the auditory to the vocal mode is a comparatively simpler type of processing. Simultaneous interpretation requires analysis of both input and output at a number of levels, so that the need to monitor both input and output can be viewed as a more intensive task than shadowing. Gerver concludes that an interpreter's ability to "comprehend" seems to depend on the level of processing.

What surprised Gerver, however, was that such an intensive and active form of processing as simultaneous interpretation could lead to poorer comprehension scores than the comparatively passive activity of listening. On the other hand, when listening, the subject is able to devote all of his attention, in other words his full channel capacity,

to processing, not having to share his attention among multiple tasks, which could explain the higher comprehension scores for listening.

Mackintosh (1985), in response to results obtained in the author's doctoral thesis (Lambert, 1983), found that simultaneous interpretation imposes the heaviest processing load on interpreters, more so than consecutive interpretation, be it consecutive interpretation with a language switch (from one language into the other) or consecutive without code-switching (from one language into the same). It should be pointed out that, unlike the above-mentioned studies where subjects' performance was based on their ability to recall information following each task, Mackintosh drew her conclusions by examining the number of departures from the standard level of language as an indication of the task demands placed on the interpreter.

#### OVERVIEW OF STUDY

The present study compares four types of information processing, a straightforward listening task, to be used as a control or contrast type, and three experimental conditions, namely 1) shadowing, 2) simultaneous interpretation and 3) consecutive interpretation. By examining and comparing the amount and quality of retention following each processing type, it is hoped that we may better understand what is meant by depth of processing, how deeply each type of message input is processed, and which type requires the greatest or the least amount of effort and attention on the part of the interpreter. It was predicted that both consecutive and simultaneous interpretation would require greater processing than either listening or shadowing. The subjects used, the experimental design, the scoring procedure are presented below.

#### METHOD AND PROCEDURE

##### Subjects

The subjects were 16 interpreters, 8 of whom were professional conference interpreters of the A.I.I.C. (Association internationale d'interprètes de conférence), all having interpreted for more than five years. The remaining eight were trainee-interpreters, all enrolled in a six-month intensive formation course at the Polytechnic of Central London, in the hope of obtaining the Diploma in Conference Interpretation Techniques. None of the 8 trainees had interpreted for less than 3 months or for more than 6 months. Among the 8 professional interpreters, 4 were female and 4 male; among the 8 trainees, 4 were female, and 4 male. The mother tongue of all 16 subjects was English and all claimed French as their B language, in other words, their strongest passive language, one which they readily interpret out of, but not necessarily into.

##### Experimental Design

A 4 x 4 Graeco-Latin Square design was used in an intentional learning paradigm (subjects were told in advance what the experiment consisted of and what was expected of them), where each subject was asked to listen to, shadow, interpret simultaneously and interpret consecutively, four French prose passages of equal length. Subjects were asked to recall what they had processed immediately following the completion of each task. Texts and tasks were balanced across subjects so that text and task effect were kept to a minimum. In other words, had all subjects begun with shadowing for example, their scores following that particular condition might have been higher simply because they were fresher and subject to less inter-text interference than after the fourth condition. All responses were recorded for subsequent analysis. Following recall, three recognition tests were administered in the following fixed order: 1) lexical, 2) semantic, and 3) syntactic recognition tests (see Carey 1971). For recall, the input was in subjects' pas-

sive language with recall in mother tongue (L2 into L1) and for recognition, both input and recognition test were in subjects' passive language (L2 into L2).

### Scoring Procedure

Following the methodology proposed by Kintsch and van Dijk (1978), the passages were broken down into a structured list of propositions and then matched against the subject's oral recall. Two independent scorers evaluated each subject's recall to ensure reliability.

### RESULTS

With regard to recall scores, a two-way analysis of variance showed that the type of task required of subjects clearly influenced their subsequent recall beyond the 0,001 level ( $F = 9,04^{***}$ ; d.f. = 3,36). Post-hoc tests of multiple comparisons indicated that recall scores were significantly higher following the listening condition ( $X = 45,19\%$ ) than after shadowing ( $X = 30,19\%$ ) (Scheffé  $F = 6,48^{**}$ ;  $p \leq 0,05$ ; d.f. = 2,36). Similar post-hoc tests also revealed that recall scores following consecutive interpretation were significantly higher than those following shadowing ( $X = 45,00\%$  versus  $30,19\%$ ,  $F = 6,31^{**}$ ;  $p \leq 0,05$ ; d.f. = 2,36). No significant differences were found between listening, consecutive interpretation and simultaneous interpretation, judging by recall scores.

Figure 2 displays the mean percentages (ranging from 50% to 80%) of recognition scores which combine scores for lexical, semantic, and syntactic recognition tests for each condition. The highest results were obtained under the listening condition ( $X = 72,35$ ), followed by consecutive interpretation ( $X = 70,27$ ), followed in turn by simultaneous interpretation ( $X = 67,06$ ), followed by shadowing ( $X = 65,02$ ).

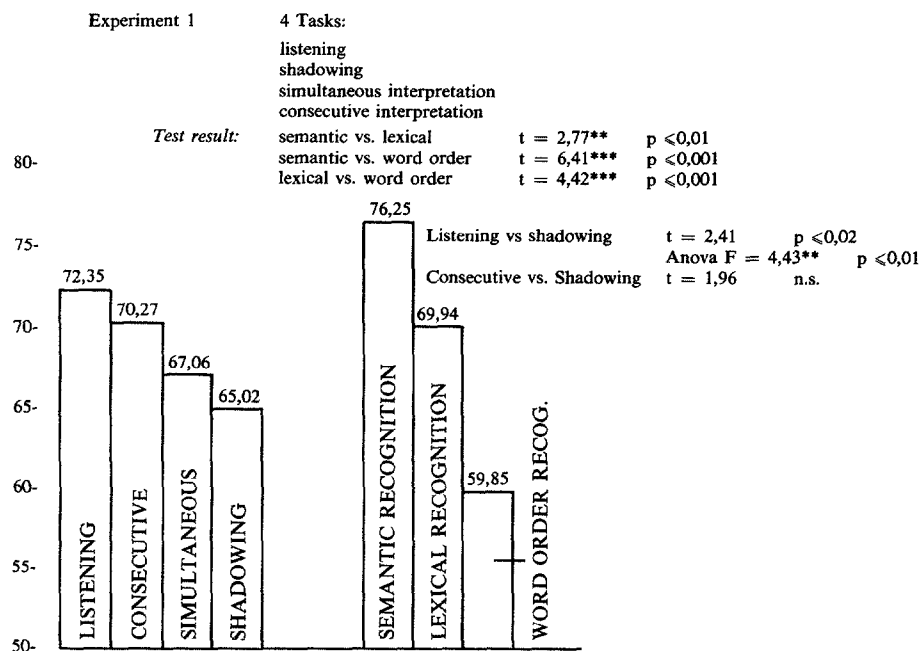


Fig. 2: Main Percentage of Recognition Scores



Following a significant treatment effect in the Graeco-Latin Square design, two tailed t-tests performed on listening vs. shadowing scores revealed that listening yielded significantly higher recognition scores than shadowing ( $X = 72,35\%$  versus  $65,02\%$ ;  $t = 2,42$ ;  $p \leq 0,02$ ); the difference between consecutive interpretation and shadowing approached significance ( $X = 70,27\%$  versus  $65,02\%$ ;  $t = 1,96$ ). No other mean differences were statistically significant.

When we break the overall recognition score down into its three components (see Figure 3), it is clear that the semantic recognition component is the most discriminating. Furthermore, listening yields the highest semantic recognition scores ( $X = 87,50\%$ ); consecutive interpretation yields the second highest set of results ( $X = 80,63\%$ ); simultaneous interpretation yields the third highest set of results ( $X = 71,86\%$ ); the lowest set of results followed the shadowing condition ( $X = 65,00\%$ ).

Two-tailed t-tests showed that there was no significant difference between listening and consecutive interpretation in terms of how these tasks affected subsequent semantic recognition tests, although results approached significance ( $t = 1,83$ ). But when listening and simultaneous interpretation were compared, listening yielded significantly higher semantic recognition scores than simultaneous interpretation ( $t = 3,44^{***}$ ;  $p \leq 0,001$ ), and higher scores than shadowing ( $t = 6,07^{***}$ ;  $p \leq 0,001$ ). Consecutive interpretation also produced higher semantic recognition scores than did shadowing ( $t = 3,99^{***}$ ;  $p \leq 0,001$ ). There were no significant differences found between listening and consecutive interpretation, nor between consecutive interpretation and simultaneous interpretation.

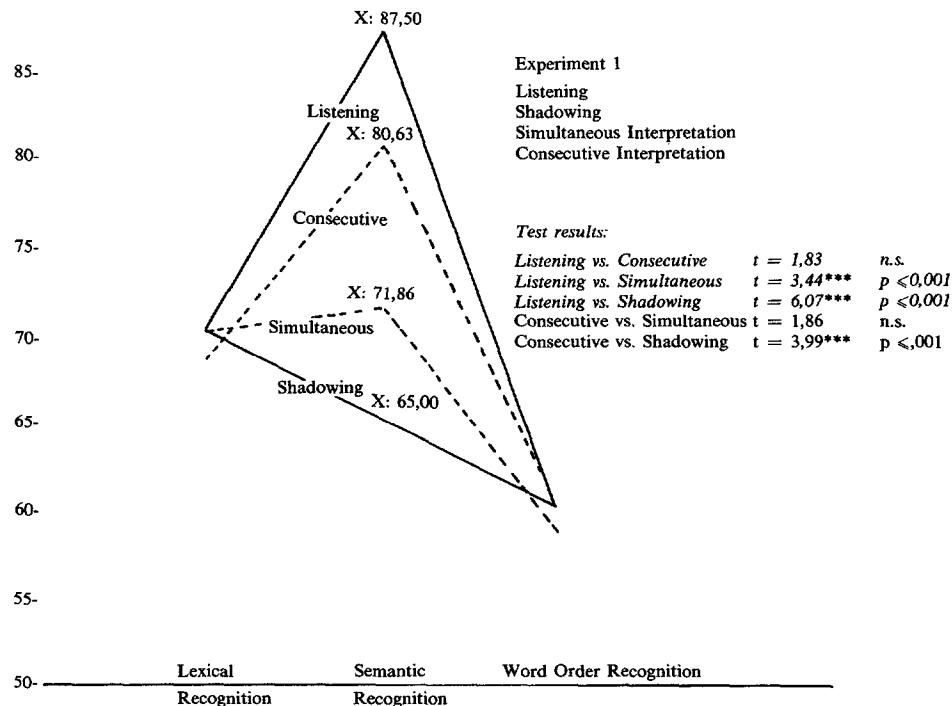


Fig. 3: Recognition Scores Including Treatment Effect

## DISCUSSION

If we compare the results obtained in Gerver's experiment to those obtained in the present experiment, certain similarities emerge, as indicated by the semantic recognition scores used in both studies : the listening condition yielded the highest mean percentage, followed by simultaneous interpretation, which, in turn, was followed by shadowing.

|                              | Gerver | Lambert |
|------------------------------|--------|---------|
| Listening                    | 58%    | 87,50%  |
| Simultaneous interpretation* | 51%    | 75,63%  |
| Shadowing                    | 43%    | 68,13%  |

\* There was no consecutive interpretation condition in the Gerver study.

Thus, Gerver found significantly higher comprehension scores for listening in contrast to simultaneous interpretation, and higher scores for simultaneous interpretation than for shadowing. Although Gerver's measure of comprehension and Lambert's semantic recognition tests are not identical, both were designed to evaluate an interpreter's ability to understand a message he or she has just interpreted.

There is then a clear agreement in the Gerver and Lambert studies with regard to listening which yields the highest recognition scores in both cases. This is interesting because *listening* can be considered as a relatively passive form of processing when compared to either shadowing, simultaneous interpretation or consecutive interpretation in that there is no ongoing overt and tangible rehearsal. This unexpected finding may be interpreted thus : listeners are able to devote their whole attention to the processing task, in other words, their full channel capacity to processing of input, not having to share attention between multiple tasks as in the cases of shadowing, simultaneous or consecutive interpreting. A more far-fetched hypothesis proposes that in order to accommodate both speed and fluency when interpreting simultaneously or consecutively, interpreters may not "listen" in the same way as those subjects told to "listen" in silence.

If we set aside the listening condition for a moment, we are left with three other forms of processing which bear directly on tasks required of an interpreter, namely, shadowing, simultaneous interpretation, and consecutive interpretation.

Of the three other forms of processing, *consecutive interpretation* yields the highest recognition scores, as originally hypothesized. Consecutive interpretation was the only condition that allowed the subject to take notes as the stimulus material was being presented. This active and visual form of rehearsal may serve to reinforce the learning activity. Furthermore, prior to recall, subjects in this condition had all given a consecutive delivery with the use of their notes, thus exposing them to a complete rehearsal of the text, from beginning to end. It was only after their consecutive deliveries that subjects were asked to hand over their notes and recall what they had just finished interpreting. Thus a strong argument can be made that the consecutive processing aided learning and memory through the overt rehearsal of the passage combined with the use of notes. On the other hand, the consecutive delivery can also be viewed as a type of interpolated activity, delaying the onset of recall, differentiating it from the other three conditions where subjects began recalling as soon as the stimulus material ended. Thus there is more chance for trace decay in consecutive processing. Since recall scores fol-

lowing consecutive delivery were higher than those following simultaneous and shadowing, it would appear that consecutive delivery is a form of rehearsal which is beneficial for, rather than detrimental to, subsequent recall. Consecutive interpretation, therefore, apparently represents a deeper form of processing due to such factors as additional rehearsal time, longer exposure to the information, the visual cues provided by the graphic notes and the aural feedback when rendering the consecutive delivery.

We are now left to explain the effects of simultaneous interpretation and shadowing. The translation factor present in simultaneous interpretation and absent in shadowing may be the distinguishing feature between the two tasks. Translation takes time (Treisman 1965). According to the Craik and Lockhart model, it is not only the depth of analysis which determines retention, but the fact that greater depth of processing usually implies more processing of the stimulus, thus requiring more time to carry out the subsequent operations for a deeper level of analysis. If the total processing time appears to covary with degree of retention, this would explain the better degree of retention following simultaneous interpretation. In terms of milliseconds, the processing time involved during simultaneous interpretation is longer than that involved during phonemic shadowing.

By examining the results obtained on the recall measures and grouping them into two categories, an interesting hypothesis emerges. The higher retention scores were obtained following listening and consecutive interpretation and the lower scores followed shadowing and simultaneous interpretation. Both listening and consecutive interpretation allow the subject to listen to the incoming message without emitting any concurrent vocal sound. In other words, when the subject is listening, no other interfering vocal activity is required. In the case of consecutive interpretation, although the subject is taking notes as the information is being processed, no vocalization is required. On the contrary, both shadowing and simultaneous interpretation require simultaneous vocalization on the part of the subject, possibly interfering with his ability to process material to any great depth. This concurrent vocal activity may in fact be the source of conflict which prevents the interpreter from processing the material to any greater extent.

With this hypothesis in mind, let us reconsider the results of the Mackintosh study (1985), where simultaneous interpretation was thought to impose a heavier processing load than consecutive interpretation, based on the number of departures from standard English. It would appear that the greater number of departures under the simultaneous interpretation condition may have been due to the simultaneity of listening, translation and speaking, in other words, conflicting activities which prevent the interpreter from processing material as deeply as under consecutive interpretation conditions.

In conclusion, by weighing the retention scores obtained by interpreters following four tasks, it would appear that deeper processing of incoming material occurs during listening and consecutive interpretation, followed by simultaneous interpretation and lastly, by shadowing. These findings, when followed by further research, may have important implications for future training of conference interpreters. The findings may also shed some light on procedures for increasing reading proficiency, for improving studying techniques, language teaching techniques, the use of audio-visual means of presentation. More generally, it is hoped the results may help us understand better which types of cognitive processing strategies enhance attention in learners.

#### REFERENCES

- CAREY, P. (1971) : "Verbal Retention After Shadowing and After Listening", *Perception and Psychophysics*, 9, 1B, pp. 79-83.  
 CHERRY, C. (1953) : "Some Experiments on the Recognition of Speech With One and Two Ears", *Journal of the Acoustic Society of America*, 25, pp. 975-979.

- CHISTOVICH, L.A., V.V. ALIAKRINSKII and V.A. ABILIAN (1960) : "Time Delays in Speech Repetition", *Questions of Psychology*, 1, pp. 64-70.
- CRAIK, F.I.M. and R.S. LOCKHART (1972) : "Levels of Processing : A Framework for Memory Research", *Journal of Verbal Learning and Verbal Behavior*, 11, pp. 671-684.
- GERVER, D. (1974) : "Simultaneous Listening and Speaking and Retention of Prose", *Quarterly Journal of Experimental Psychology*, 26, pp. 337-341.
- KINTSCH, W. and T.A. van DIJK (1975) : "Comment on rappelle et on résumé des histoires", *Langages*, 9, pp. 98-116.
- LAMBERT, S.M. (1983) : *Recall and Recognition among Conference Interpreters*, unpublished doctoral dissertation, University of Stirling, Stirling, Scotland.
- MACKINTOSH, J. (1985) : "The Kintsch and van Dijk Model of Discourse Comprehension and Production Applied to the Interpretation Process", *META*, 30:1, pp. 37-43.
- MASSARO, D. W. (1975) : "Language and Information Processing", in D. W. Massaro (ed.), *Understanding Language*, New York, Academic Press.
- MOSER, B. (1978) : "Simultaneous Interpretation : A Hypothetical Model and Its Practical Application", in D. Gerver and H.W. Sinaiko (eds.), *Language Interpretation and Communication*, New York, London, Plenum Press.
- NORMAN, D.A. (1976) : *Memory and Attention*, New York, John Wiley and Sons, Inc.
- TREISMAN, A. (1965) : "Verbal Responses and Contextual Constraints in Language", *Journal of Verbal Learning and Verbal Behavior*, 4, pp. 118-128.
- WAUGH, N.C. and D.A. NORMAN (1965) : "Primary Memory", *Psychology Review*, 72, pp. 89-104.