
Ambiguous Surface Structure and Phonetic Form in French

Author(s): W. J. M. Levelt, W. Zwanenburg, G. R. E. Ouweneel

Source: *Foundations of Language*, Vol. 6, No. 2 (May, 1970), pp. 260-273

Published by: Springer

Stable URL: <http://www.jstor.org/stable/25000453>

Accessed: 22/01/2010 09:45

Your use of the JSTOR archive indicates your acceptance of JSTOR's Terms and Conditions of Use, available at <http://www.jstor.org/page/info/about/policies/terms.jsp>. JSTOR's Terms and Conditions of Use provides, in part, that unless you have obtained prior permission, you may not download an entire issue of a journal or multiple copies of articles, and you may use content in the JSTOR archive only for your personal, non-commercial use.

Please contact the publisher regarding any further use of this work. Publisher contact information may be obtained at <http://www.jstor.org/action/showPublisher?publisherCode=springer>.

Each copy of any part of a JSTOR transmission must contain the same copyright notice that appears on the screen or printed page of such transmission.

JSTOR is a not-for-profit service that helps scholars, researchers, and students discover, use, and build upon a wide range of content in a trusted digital archive. We use information technology and tools to increase productivity and facilitate new forms of scholarship. For more information about JSTOR, please contact support@jstor.org.



Springer is collaborating with JSTOR to digitize, preserve and extend access to *Foundations of Language*.

AMBIGUOUS SURFACE STRUCTURE AND PHONETIC FORM IN FRENCH*

Abstract. In modern approaches to phonology a lack of clarity exists on the issue of whether phonetic facts are psychological or physical realities. The results from an experiment suggest that phonetic facts can be considered as psychological realities, but with the restriction that they *can* (but not necessarily always *do*) take acoustical shape. More specifically, the syntactic material consisted of ambiguous French sentences of the following sort: *On a tourné ce film intéressant pour les étudiants*. They were spoken (a) in disambiguating contexts, without the (four) readers noticing the ambiguities, and (b) without context, but with the instruction to make a conscious effort to disambiguate. By tape splicing, the contexts were removed from the context-embedded sentences. Twenty-eight native speakers of French listened to the sentences and judged whether one or the other meaning had been intended by the speaker. Subjects performed significantly above chance: 60% correct identifications for context-embedded sentences, 75% for context-free sentences. Pitch-amplitude analyses were made to determine the acoustical differences involved.

I. INTRODUCTION

In spite of much refinement, recent notions on phonology and the character of phonetic representations are often inconsistent with respect to the 'reality-status' of phonetic descriptions. Originally, of course, a phonetic representation was supposed to be a formal statement about linguistically relevant physical facts. The reality denoted by such statements was acoustical or physiological, but essentially it was presumed to be of a *physical* nature.

In recent years also the opinion has been defended that a phonetic representation is a psychological reality rather than a physical one. Psychological in either the sense of *perceptual* or *psychomotor*. The perceptual version says that the listener internally generates an approximation to the speech signal that is presented. This approximation is possible on the basis of his linguistic competence and in fact constitutes the form of his perception. The psychomotor version takes the phonetic representation to be an abstract description of the commands a speaker is relaying to his vocal apparatus. Finally, it is an obvious step then, to identify part of this command system and the internal generation system that is postulated in the perceptual theory. The

* The authors wish to express their gratitude to Mr. L. Docquir of the *Lycée Français de la Haye*, who supplied subjects and room for this experiment, to Mrs. G. Bolhuis-Schwartz, Mrs. A. M. de Both-Diez, Mrs. E. M. R. Sloodman-Bouly and Mrs. T. Hidding-Carolus Barré, who were readers, to Mr. A. van Katwijk of the *Institute for Perception Research*, Eindhoven, who was in charge of the acoustical analysis of the sentences, and to Mrs. A. Salverda-Meletopoulos, to Mr. C. Keers and Mr. J. van der Sman, who assisted throughout the experiment.

motor theory of speech perception, as developed at the Haskins Laboratory, is making precisely this point (for a recent statement of this theory, see Liberman *et al.*, 1967).

In view of this merging of perceptual and psychomotor viewpoints it is sufficient for our present discussion to distinguish simply the *physical* and the *psychological* approaches in phonology.

In his recent writings, Chomsky has strongly endorsed the latter approach. Let us take one of several statements he has made on this issue: "A person who knows the language should 'hear' the predicted phonetic shapes. In particular, the careful and sophisticated impressionistic phonetician who knows the language should be able to bring this perceptual reality to the level of awareness, and there is ample evidence that phoneticians are capable of doing this. We take for granted then, that phonetic representations describe a perceptual reality. Our problem is to provide an explanation for this fact. *Notice, however, that there is nothing to suggest that these phonetic representations also describe a physical or acoustic reality in any detail*" (N. Chomsky and M. Halle, 1968, p. 25; italics ours).

This position is clear, and attractive in the eyes of the perception theorist who is familiar with concepts such as constancies, i.e. with the relative salience of psychological over physical factors in perception. But a consequent application of this way of thinking in phonology raises some difficulties. For example, on page 297 of the same source we find: "*In fact, the phonetic features are physical scales* and may thus assume numerous coefficients, as determined by the rules of the phonological component" (italics ours). The contradiction with the above citation may be a result of the double authorship of this book. But even if this is true, it reflects a problem of phonology that can only be solved by experimentation. The situation is not essentially different from others in perception theory. The general empirical problem is precisely how stimulation and perceptual coding are related. To what degree are perceptual codes unaffected by changes in physical stimulation? And on the other hand: what physical change triggers a perceptual change of code? The study of such questions can shed light on the active role perception plays in imposing structure on the physical world. This is also true for phonology. If the phonetic representation is taken to be the formulation of a perceptual reality, how then is this related to the physical stimulus? Is the relation such a close one that phonetic features can indeed be conceived of as physical scales? Only careful experimentation can decide such issues.

II. PURPOSE OF THIS STUDY

If a string of words can have two different phrase structures, both of which

can be constructed as grammatical sentences, we say that the string (or 'sentence') is surface structure ambiguous, i.e. there are two possible surface structures for the same string of words. If indeed the surface structure is the input in the phonological component, one would expect these different surface structures to result in different phonetic representations. Moreover, one would expect the hearer to be able to distinguish these two sentences on the basis of their different phonetic shapes. But it should be noted that these are no logical necessities but assumptions about empirical facts. To sum up, three assumptions are being made:

(1) Different phrase structures correspond to different phonetic forms.

(2) These phonetic forms entail not only perceptual, but also physical differences, which are realized by the speaker.

(3) The hearer is able to detect and interpret these differences correctly.

The present study is concerned with a further empirical investigation of these assumptions.¹

The first assumption: different surface structures correspond to different phonetic representations. Dow (1966) takes this to be Chomsky's view, as expounded in his MIT-lectures. To our knowledge there is no such statement in his published writings.

There is both positive and negative evidence for the second assumption. On the one hand there are physical differences in the signals, corresponding to different word groupings for Bolinger and Gerstman's (1957) example: (*light house*) (*keeper*), vs. (*light*) (*house keeper*), as shown by Bolinger's and by Lieberman's (1967) analyses of these cases. The difference is mainly in disjuncture, i.e. in the length of the interval between the vowels (Lieberman, p. 153). On the other hand the findings by Garrett *et al.* (1966) provide negative evidence. They induced a perceptual change in constituent boundaries without any change in the physical stimulus. Using the click-procedure, they show that the perceptual phrasing for the sentences $\left\{ \begin{smallmatrix} \text{his} \\ \text{in her} \end{smallmatrix} \right\} \text{hope of marrying}$ *Anna was surely impractical* corresponds to the grammatical constituents. Their method of tape splicing obviated the possibility of physical differences in the two sentences from the word 'hope' on. Hence, perceptual and physical distinctions need not be concomitant. Finally, if we do find that subjects are able to identify the intended syntactic structure by listening to the isolated sentence, then, of course, not only the third assumption, but also the second and first are correct with respect to the ambiguity in question. If independent subjects can decide on the speaker's intention, then there must be a physical difference between the two cases (assumption 2). And if there exists a

¹ The experiment to be reported has also been motivated independently by certain problems in French phonology. We will report elsewhere on these issues.

phonetic difference that is syntactically motivated, then the two surface structures are indeed mapped onto two different phonetic forms (assumption 1). Problems of interpretation arise only in the event that subjects are *not* able to identify the meaning intended by the speaker. This is what happened to Anne Dow, who on the basis of such negative findings challenged assumption 1. However, she did not reject the possibility that her negative findings could have been due to lack of motivation or of perceptual sensitivity on the part of the listeners. The present experiment is designed to further investigate this problem.

III. THE EXPERIMENT

Syntactic Material. The experimental sentences were surface structure ambiguous French sentences. All had the same syntactic ambiguity. A typical example is *On a tourné ce film intéressant pour les étudiants*. There are two possible constituent structures in this case: *pour les étudiants* may belong to *intéressant*; the movie, then, is interesting for the students. But *pour les étudiants* may as well modify the verb *tourné*, in the sense that the movie is shown to the students. In the latter case *film* and *intéressant* belong together. In short, these sentences are distinguished by the prepositional phrase (i.e. *pour les étudiants*), which may modify the verb (*tourné*) or the adjective (*intéressant*). We will therefore indicate these alternatives by *verbal* and *adjectival* form of the sentence, or by the symbols V- and A-form. Forty-eight such sentences were used in this experiment. They are given in Appendix A, below. These 48 sentences were selected from a preliminary pool of 96. Our reasoning was that the sensitivity of the experiment could be considerably increased by excluding sentences which, in spite of their ambiguity, are actually biased towards one or the other interpretation. In a pre-experiment the original 96 sentences were presented to 38 Dutch students of French. This was done in written form in order to exclude prosodic information, and to concentrate on the semantic bias. The students were instructed to indicate the first interpretation that came to mind upon reading a sentence. Results indicated that certain sentences nearly always led to the same interpretation, whereas others were satisfactorily ambiguous. We took the 48 most ambiguous sentences as test sentences for the main experiment. The maximal bias in this set was 2:1, i.e. 25 students marking one interpretation for the sentence and 13 the other one. Moreover the actual biases were equally divided over the V- and A-forms.

Each of the 48 sentences was embedded in two short anecdotal contexts. The function of these contexts was to disambiguate the sentence. One story induced the V-form of the sentence, the other one the A-form. As an example

we give the two contexts for *on a tourné ce film intéressant pour les étudiants*.

Verbal form: Enfin un film intéressant, après tout le fatras que nous avons eu ces derniers mois. Je me vois déjà courir au cinéma. Mais qu'est-ce qui se passe? Il n'est pas destiné au grand public; *on va tourner ce film intéressant pour les étudiants*. S'est à s'arracher les cheveux.

Adjectival form: Ce film a spécialement été fait pour les étudiants, et n'intéresse vraiment personne d'autre. Or, devinez ce qu'on va faire pour le centenaire de notre association? Devant un public d'épiciers *on va tourner ce film intéressant pour les étudiants*. Bien malin qui y comprend quelque chose.

Speakers Design. Four adult women, native speakers of French, acted as speakers. Each of them completed the following speakers programme: in the first place they read the stories. They were not informed about the aim of the experiment or the ambiguity of the embedded sentences. Each speaker was given 12 of the 48 sentences each in A- and V-context. In this and also in the next phase the 12 sentences were at first read in one version (A or V, randomly assigned) and then all were read in the alternative version (V or A). At the end of this phase the speaker was asked about what he thought the purpose of the experiment might be. Only one of them had noticed that there were certain ambiguous sentences. We did an independent analysis on the data that were obtained from this speaker's sentences. No noticeable difference from the other three speakers could be found.

In the second phase we informed the speakers about the ambiguity of the test sentences and then asked them to pronounce 12 new sentences (from the pool of 48) twice without context, once with the adjectival intention and once with the verbal intention. It was stressed that they should make it as easy as possible for an eventual listener to detect their intention.

All speaking was done in a sound proof room and high quality recordings were made (Sennheiser microphone, Revox recorder).

In this way we obtained four versions of each sentence: an adjectival (A) and a verbal (V) version spoken in context (C^+) and also both versions intentionally produced in a context-free (C^-) condition. We will consequently denote these versions by C^+A , C^+V , C^-A , and C^-V , respectively. In a latin square design each version of each sentence was read by one speaker.

Listeners Design. All C^+ versions were cut out of their context tapes. Together with the C^- versions we had 192 different tape-segments in total (48 from each speaker). These 192 segments were distributed over four test-tapes. Each test-tape had the following properties:

- (i) It contained one and only one version of each of the 48 test sentences.
- (ii) It contained 12 sentences from each speaker. The 12 sentences from a speaker were kept together in order to give maximal opportunity for the listener to get accustomed to an individual speaker. The sequence of speakers on the tape was determined according to a 4×4 latin square design.
- (iii) Within each 'speaker segment' $3 C^+A$, $3 C^+V$, $3 C^-A$, and $3 C^-V$ versions occurred in random order.

Subjects. 28 pupils of the Lycée Français de La Haye, all native speakers of French, participated as subjects. Their age range was from 12 to 17; there were 21 girls and 7 boys, randomly assigned to 4 groups of 7.

Procedure. The experiment was run in groups. The 7 subjects were seated in a classroom. Each of them had a set of written instructions and a test booklet. The instructions had also been tape recorded, and were played over a high-quality loudspeaker. The instructions started with an extensive explanation of the ambiguous nature of the test sentences that would follow. Several examples were discussed. It was then suggested that 'language-sensitive' people could often hear the intended meaning if such sentences were spoken in isolation and the subjects were asked to try this themselves for the sentences which were to follow. They were then instructed how to record their judgments in the test booklets. On each page a test sentence was printed with dotted lines under the verb and the adjective. Subjects were instructed to first study the test sentence in order to realize its two possible meanings, then listen to the tape recorded version of it and finally mark their judgments, i.e. whether the prepositional phrase had been intended to modify the verb or the adjective. The marking could be done by connecting either the dots under the verb or those under the adjective, in accordance with the decision.

A trial sentence was given, the subjects marked their judgments in the test booklets and had the opportunity to ask questions. Finally, a summary of the instruction was given, and the experiment began. However, 4 additional sentences were inserted before the first test sentence in order to insure acquaintance with the task; the subjects were not aware these were only training sentences. The experimenter took care that all subjects had made their judgments and read the next sentence before its recorded version was presented. After the first 24 test sentences a 5-minute pause was given.

Scoring. Subjects' judgments were scored as follows: if a sentence was judged to be intended in his V-version, it got a score 1, *irrespective of the correctness of this judgment*. The A-response always got a score 0. This way

of scoring had certain advantage in the analysis of variance to be performed.

IV. RESULTS

Table I gives a summary of the data.

TABLE I

Experimental results. Cell-values represent the number of times sentences of the particular condition are judged to be of the verbal version (max.: 3)

Subjects	+ Context version								- Context version							
	Verbal version				Adjectival version				Verbal version				Adjectival version			
	Readers				Readers				Readers				Readers			
	1	2	3	4	1	2	3	4	1	2	3	4	1	2	3	4
Group 1: 1	1	2	2	1	1	0	2	1	3	3	3	3	1	0	1	1
2	2	2	0	0	2	1	3	0	3	3	2	3	1	1	1	1
3	2	3	1	1	2	2	1	0	3	3	3	3	1	0	1	1
4	2	0	0	1	1	1	1	1	3	3	3	1	2	2	1	2
5	3	3	2	2	3	2	3	1	0	2	1	1	1	3	3	2
6	2	2	2	2	1	1	2	1	3	3	3	3	0	0	1	1
7	1	0	1	1	1	2	2	0	1	3	3	2	2	0	1	2
Group 2: 8	1	1	2	1	0	1	2	2	2	1	2	0	0	1	1	1
9	1	2	3	1	0	2	0	2	2	2	2	3	2	0	0	1
10	1	2	2	3	2	1	3	1	1	0	2	2	2	1	1	2
11	2	3	2	1	0	3	1	2	3	1	2	2	1	1	2	1
12	1	2	3	3	1	0	1	1	3	3	3	3	0	1	1	3
13	2	0	3	1	1	1	0	1	2	3	3	3	1	2	1	1
14	2	2	2	1	1	1	0	1	3	3	3	3	0	1	0	0
Group 3: 15	3	2	2	2	0	1	0	1	3	2	3	2	1	0	0	1
16	3	2	2	3	1	2	0	1	3	3	3	3	0	0	1	0
17	2	2	3	2	1	2	1	0	3	3	3	3	0	0	1	1
18	2	2	2	1	0	0	2	0	3	3	3	2	0	0	1	1
19	2	1	0	1	2	2	1	2	0	0	0	1	1	1	2	2
20	3	3	3	2	0	1	1	0	3	3	3	3	0	0	0	1
21	2	2	1	1	0	2	0	0	3	0	2	3	0	0	0	2
Group 4: 22	2	2	2	2	3	1	0	2	3	2	3	3	1	0	0	1
23	1	3	2	1	1	0	2	1	3	2	2	3	1	1	1	0
24	3	1	2	1	2	1	1	2	2	2	1	2	0	1	1	0
25	1	3	2	1	1	0	0	2	3	3	3	3	1	0	1	0
26	2	2	0	0	1	1	1	2	0	2	3	1	0	1	2	2
27	1	1	2	2	1	1	1	0	3	2	3	3	1	1	1	1
28	1	3	2	2	3	2	0	3	3	1	2	3	2	0	0	1
Σ	194				127				264				98			

The scores range from 0 to 3. This is because from each reader there were three sentences of a particular version (e.g. the C⁺A-version) on the test tape for a group. If all three sentences were judged to be of the V-version by a particular subject, this subject got a score 3 for that speaker/version condition. High numbers indicate many 'verb'-judgments, low numbers many 'adjective'-judgments. An analysis of variance was performed on these data. The significance levels which follow are derived from this analysis, unless otherwise stated. The main results are:

(1) The V-versions of the sentences are significantly ($p < 0.01$) more often identified as V than the A-versions. On the whole, therefore, listeners seem to be able to identify the intention of the reader. An idea of the size of this effect can be obtained from the percentage of correct judgments. For the whole experiment this is 67% (chance level is 50%).

(2) For the intentionally spoken sentences, i.e. the C⁻ versions, there is 75% correct identification. This is significantly more ($p < 0.02$) than for the sentences spoken in context; for these C⁺ versions there is only 60% correct identification. For both conditions, however, correct identification is significantly above chance ($p < 0.01$, $p < 0.025$ respectively, Scheffé post hoc comparisons). So, for both the context-free and the context-embedded versions listeners perform above chance level, but they are significantly better in the context-free condition.

Further results show that the interpretation of the two main effects can be straightforward because possibly interfering effects are minimal:

(3) Subjects show a balance in their V- and A-judgments. There are 48% V-answers in the C⁺ condition and 54% in the C⁻ condition.

(4) There is no significant difference due to speakers, i.e. the sentences read by different speakers have not led to different results. This is also true if we look into the C⁺ and C⁻ conditions separately.

(5) Though the four experimental groups do not differ significantly in percentage of V-judgments, they are significantly different in percentage of correct identifications ($p < 0.01$). The results provide no satisfactory explanation for these differences. We checked several possibilities, such as age level, grades, experimenter, etc. Though the groups are too small for determining any systematic effects, we found slight indications for an influence of age and intelligence (grades) on performance. However, apart from the overall level of performance, the groups do not differ significantly.

V. DISCUSSION

If a reader consciously attempts to disambiguate a surface structure ambiguous sentence, he will frequently succeed. In 75% of the cases listeners correct-

ly identify the intention of the speaker. This single result gives positive evidence for all three assumptions made in the introduction:

(1) Different phrase structures should be reflected in different phonetic shapes.

(2) This difference is not only perceptual, but also physical. All other cues for a perceptual, non-physical difference (i.e. by context) have been excluded in this experiment. There must have been a characteristic difference in sound pattern between the A- and V-forms of the sentences.

(3) The listeners have been able to detect and correctly interpret this difference.

For the ambiguity under concern we can, therefore, not challenge assumption 1, as Anne Dow did. The French language apparently does allow for a phonetic difference corresponding to the difference in surface structure.

It is in this light that we should interpret the findings for the sentences read in context. Here we found 60% correct identifications. Although this result is above chance, it is significantly less than for the context-free sentences.

The interpretation of this finding can be quite specific because of what we already know from the context-free sentences: contrary to Anne Dow we already have evidence for the correctness of the first assumption. Moreover, we do know that the listeners are sensitive to a certain level of acoustical difference between the two sentence forms (assumption 3). Otherwise they would not have been able to perform at the 75% level in the context-free case. By exclusion, then, the less pronounced effect in the C⁺ condition must be due to the fact that the phonetic shapes of the A- and V-forms are not sufficiently determined acoustically in the speech of the readers. Stated in other terms: for context-embedded speech the assumption that a phonetic difference is actually expressed acoustically (assumption 2) cannot be maintained as a general truth. This could very well be the explanation for Anne Dow's negative findings.

Two points remain to be discussed. The first concerns the acoustic cues that differentiate between the two surface structures, the second concerns the relevance of the present findings for a theory of speech perception and specifically the reality status of phonetic structure.

A. Acoustic Analysis

For an acoustic analysis we selected a few sentence pairs (A- and V-forms) for which a maximum of correct identifications had been obtained in the experiment. These sentences were subjected to a fundamental pitch and amplitude analysis.² Though we considered the results for five pairs of

² Performed at the *Institute for Perception Research, Eindhoven*.

sentences, Figures 1-4 give data for the two optimal pairs only. For the context-embedded condition the sentence that was most often correctly identified (in fact in 86% of the cases) was *il veut vendre cet objet volé à son ami*. For the V-form the data are shown in Figure 1. The A-form analysis is given in Figure 2.

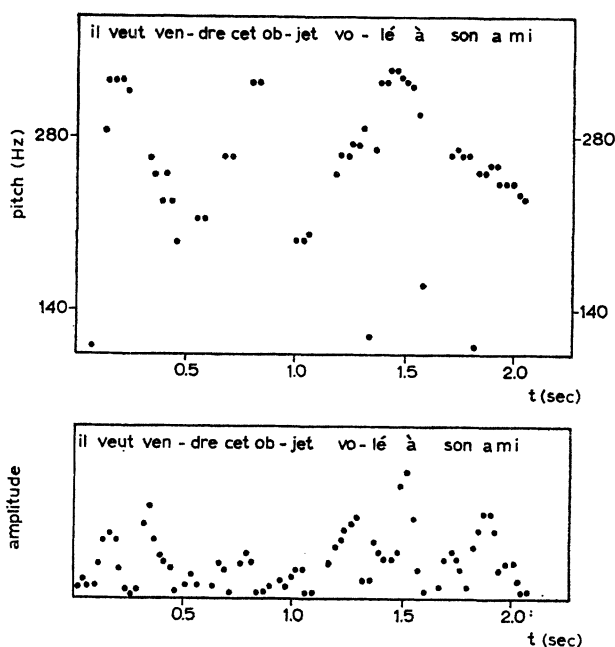


Fig. 1. Pitch and amplitude analysis of *il veut vendre (cet objet volé) (à son ami)* (verbal version).

The most striking acoustic differences between these sentences are:

(1) The disjuncture (in Lieberman's sense) between *objet* and *volé* is longer for the A-form (Figure 2) than for the V-form (Figure 1). This is consonant with the respective constituent structures.

(2) For the V-form the intonation increases from *-jet* to *vo-*. For the A-form it is just the reverse. The falling intonation in the latter case may set apart the phrase *volé à son ami*.

These are the two differences that *mutatis mutandis* are most characteristic for all the cases analysed. They also hold for the best context-free sentence pair: *il faut préparer cet âme impénétrable à la grâce* (100% correct). This pair is shown in Figures 3 and 4. For the adjectival form (Figure 4) the disjuncture, i.e. the vowel-vowel interval, between *âme* and *impénétrable* is

relatively large. It also shows the proportionately high pitch for *âme*.

There are other acoustical differences that occur in some cases. In the V-forms the speaker often makes a rather long disjuncture between adjective and preposition, corresponding to the phrase structure.

None of these physical differences, however, seem to be absolutely necessary for correct identification. Our general impression from the analyses is

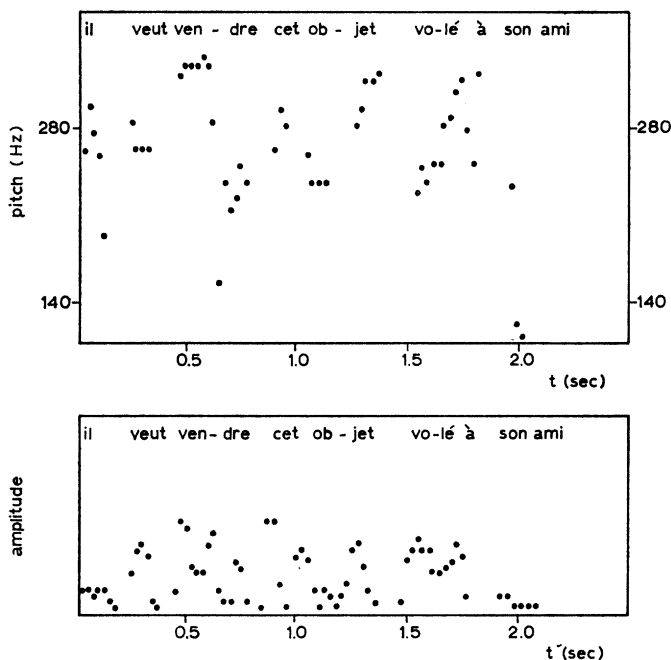


Fig. 2. Pitch and amplitude analysis of *il veut vendre (cet objet) (volé à son ami)* (adjectival version).

that these cues are mutually substitutable. Independent manipulation of these variables will of course be required to substantiate this impression. Moreover, one should keep in mind that these sentences were read, not spontaneously spoken. The acoustical pattern may be different in the latter case.

One interesting detail – which is not immediately relevant for the present discussion – concerns the relations between intonation and amplitude pattern. Though intonation and amplitude generally covary, there is a notable exception for *volé*, which ends at a rising intonation but a falling amplitude. Less clearly, a similar pattern occurs for *impénétrable*.

B. Reality of Phonetic Structure

We have shown that the French language supplies the phonological means to disambiguate sentences of the type studied in this paper. A phonology of French should therefore assign different phonetic structures to the A- and V-forms of these sentences.

This difference is clearly perceptual, but our results indicate that although a physical component may be present, it need not necessarily occur in

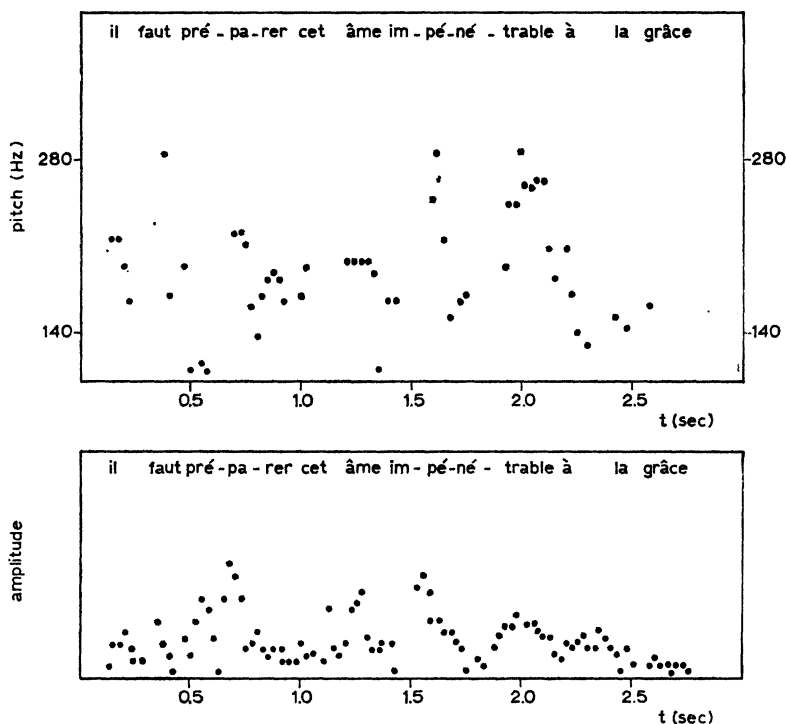


Fig. 3. Pitch and amplitude analysis of *il faut préparer (cet âme impénétrable) (à la grâce)* (verbal version).

normal speech. In fact, the acoustical realization occurs mainly in the case where it is not possible for the listener's perceptions to be guided by contextual information. In other words, in our case the speaker makes minimal use of the prosodic features of his language as long as he is assured that a listener will correctly interpret his speech. The speaker is apparently assuming that for the same perception to occur semantic information can replace prosodic information. This intuition is in full agreement with Garrett *et al.*'s findings that have been discussed above.

Turning back now to our introductory discussion on the physical vs. psychological reality status of phonetic descriptions, we would suggest the following convention:

For a linguistic fact to be phonetic, it should be *virtually* acoustical, in the sense that it *can* take physical shape. This is the only acoustical limitation on an otherwise psychological approach in phonology. This means that phonetic aspects of normal speech will often be only psychological, i.e. that they are inferred from context or meaning as substitutes for a possible but not actually realized physical form.

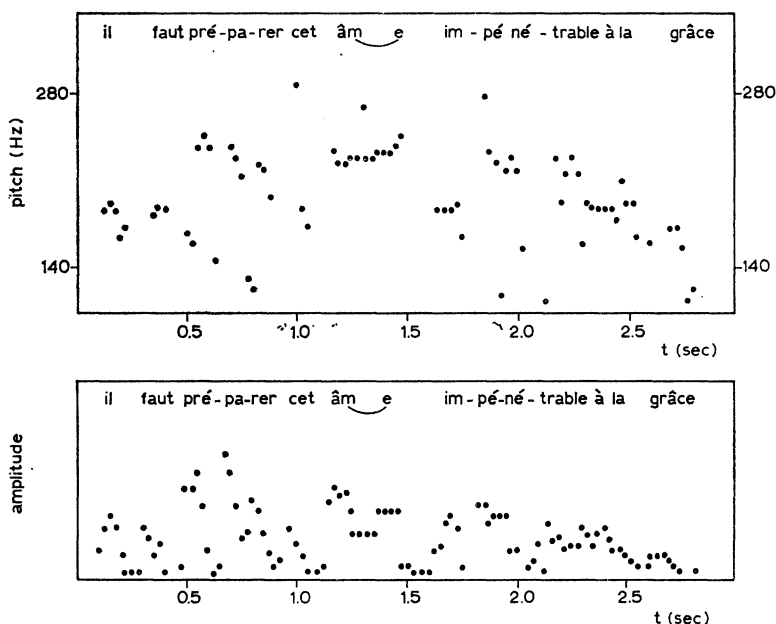


Fig. 4. Pitch and amplitude analysis of *il faut préparer (cet âme) (impénétrable à la grâce)* (adjectival version).

APPENDIX

1. On ne peut pas bannir cet homme absent de sa résidence.
2. On veut empêcher cette évolution accélérée par l'intervention du gouvernement.
3. On ne peut pas refuser cette subvention déjà accordée aux communes.
4. Cela peut rajeunir un homme âgé de trente ans.
5. Il faut promettre des mesures agréables aux élèves.
6. Il veut fatiguer ces troupes aguerries par de nombreux exercices.
7. On ne doit pas encourager ce jeune homme âpre à exiger son dû.
8. Il est inutile d'exhorter cette personne attentive à ne mécontenter personne.
9. Il va consoler l'enfant attristé par ces paroles.
10. Il ne faut pas combler ce vieillard avare de louanges.
11. Il ne faut pas combler ce vieillard avide d'honneurs.

12. Il veut accuser son ami complice du vol.
13. Il faut féliciter ce vieillard content de sa décision.
14. On ne saurait réduire ces esprits contraires à la raison.
15. Il ne veut pas parler à cette femme curieuse des secrets d'autrui.
16. Il doit confirmer cette nouvelle désagréable aux intéressés.
17. On va combler ce vieillard digne d'honneurs.
18. Elle sait amuser cet esprit facilement distrait par des choses imprévues.
19. Il faut mettre fin à cette discussion échauffée par une simple parole.
20. Il doit promettre la somme empruntée à son ami.
21. Il va assoupir l'auditoire ennuyé par une musique trop lente.
22. Elle fait douter son mari envieux de l'honneur de son voisin.
23. Il veut suggérer cette idée peu familière au grand public.
24. Il veut gagner cette jeune fille gâtée par des propos flatteurs.
25. Il ne faut pas encourager cet esprit habile à tromper les autres.
26. Il sait décider ce client hésitant à acheter.
27. Il veut avertir son ami ignorant des richesses de ses parents.
28. Il faut préparer cette âme impénétrable à la grâce.
29. Il doit demander cette information importante pour ses amis.
30. Il va abandonner cet ami importun à lui-même.
31. Il est inutile d'exhorter cette armée impuissante à se retirer.
32. Il faut éloigner ce fils indigne d'un tel père.
33. Il faut absoudre cette personne innocente du crime.
34. Il veut accuser son ami innocent du crime.
35. Il est inutile de prévenir cet esprit inquiet de ces événements.
36. On va tourner ce film intéressant pour les étudiants.
37. Il ne sait pas gagner cette personne intimidée par la familiarité.
38. Il faut prévenir cet ami jaloux de sa réputation.
39. On doit empêcher ce mal menaçant de ruiner le pays.
40. Il est inutile de comparer ces deux intérêts parallèles l'un à l'autre.
41. Il va préparer une situation pénible à ses amis.
42. Il faut comparer les langues postérieures au latin.
43. On ne doit pas admettre ces étudiants peu préparés à l'examen.
44. Elle sait réduire un coeur rebelle à l'amour.
45. On peut distinguer ces deux mots synonymes l'un de l'autre.
46. Elle veut acheter un cadeau utile à son fils.
47. On doit transmettre des connaissances utiles à la génération suivante.
48. Il veut vendre cet objet volé à son ami.

Groningen University, The Netherlands

BIBLIOGRAPHY

- Bolinger, D. L. and Gerstman, L. J.: 1957, 'Disjuncture as a Cue to Constructs', *Word* 13, 246-255.
- Chomsky, N. and Halle, M.: 1968, *The Sound Pattern of English*, Harper and Row, New York.
- Dow, A. R.: 1966, *The Relation of Prosodic Features to Syntax*, Master's thesis Harvard University.
- Garrett, M., Bever, T., and Fodor, J. A.: 1966, 'The Active Use of Grammar in Speech Perception', *Perception and Psychophysics* 1, 30-32.
- Lieberman, A. M., Cooper, F. S., Shankweiler, P. P., and Studdert-Kennedy, M.: 1967, 'Perception of the Speech Code', *Psychological Review* 74, 431-461.
- Lieberman, P.: 1967, *Intonation, Perception and Language*, M.I.T. Press, Cambridge, Mass.