

タイトル	Traditional Phonetics and Neo-Macro-Phonetics
著者	WAIN, Peter
引用	北海学園大学人文論集, 7: 31-45
発行日	1996-10-31

Traditional Phonetics and Neo-Macro-Phonetics

Peter WAIN

Summary

Notwithstanding current trends towards electronic visual representation of language, communication by speech continues its primacy. This paper outlines the aspects of phonetic performance of a traditional nature which a university graduate who has specialised in English should be capable of, and the understanding of speech performance that should be at the disposal of the teacher of English who is preparing undergraduates to perform with an adequate functional range.

Key Words: segments tunes timbres

Part I Traditional phonetics

Perhaps in the far distant future English as a language spoken internationally may effectively be a new pidgin (the implication of Sadahiko Ikeura et al. [1]). However, for the foreseeable future — the period at the end of which the graduands of today retire from work, say forty years — the overwhelming need will be for a spoken performance which is recognisably similar to one of the standard Englishes now spoken, or to a variant of these, or to a composite of these, for example “mid-Atlantic”. Let us call these received dialects.

1.1. Segmental phonetics

These received dialects, although they may differ widely in their productions of allophones, nevertheless conform to the same sound system. To be more precise, British Englishes, American Englishes, Australian, New Zealand, Canadian and South African Englishes when spoken as native languages have each within its segmental phonetic system a sub-system of twelve vowels distinguishable by their distinctive features. To these twelve vowels are added eight diphthongs deriving from the twelve vowels.

Setting aside allophones, which we do for all languages by definition of allophone, Japanese has a much reduced system of vowels and diphthongs with comparatively little similarity to that of the English system in all its varieties. While katakana vowels may serve satisfactorily in an interim system for “English conversation” and for stages of learning up to high school (and even in university and beyond for people for whom English is not a specialisation) it is certainly not adequate for the specialist in the pursuit of excellence. There is also a sociolinguistic dimension: while occasional departures from the phonetic system of a target language may not impede communication, and may even be considered charming and fetching, an extensive disregard of the phonetic system is looked upon as funny, and a speaker of such a marked variant is a funny foreigner, not a desirable status sociolinguistically. While we are not training spies we should not, either, be preparing people to speak a low status interlanguage. Sociolinguistic aspects will be considered again with functional phonetics later in this paper.

Sadahiko Ikeura et al. [1] offer two hypotheses: one: that some English vowel sounds are similar to the five Japanese vowels, and two: that those Japanese vowels resembling English vowels are acceptable

for communicating with the international community of English speakers.

Their work is inconclusive for two reasons.

Firstly, although in their conclusion they find 7 English vowels the “same” as a Japanese vowel, 4 English vowels “similar”, only 3 “suspicious”, and none “different”, their findings must be fudged since for three Japanese “same” vowels two quite different English vowels are postulated in each case. Therefore we are expected by a simple extension to accept that three sets of minimal pairs are not distinguishable in English — clearly a nonsense.

Secondly, the 230 informants (all children, Australian, 61, Canadian, 58, English, 51, U.S.A., 60) were offered isolated monosyllabic words in both the tests, and in the second test were given three bites at the cherry, e.g. “earn, turn, bird”; even then only 1% of the children recognised English vowel number 11 from the Japanese “equivalent”, which is therefore classed as “suspicious”. Clearly one cannot base a hypothesis on such data, or decide on the nature of phonetic training.

What such minimal pair type tests do not probe is the effect of series of suspicious and different vowels on the processing system of the hearer. With “earn, turn, bird” there are no other data to process. The three words can be held in short-term memory for a decision to be made, and they never carry the meaning of anything to be transferred to long-term memory.

Speaking and hearing are never like the models used for such simplistic tests. Furthermore Ikeura does not — cannot — include the most common vowel schwa, number twelve in the system, which is always unstressed and thus gives rhythm to the language as well as “pointing” the other eleven vowels; it is about 25% of the vowel sounds in any utterance.

While we are discussing segmental phonetics we should bear in mind that the vowel system is not the only problem. As well as differences in the vowel systems there are differences in the consonant systems. The difference between Japanese bilabial fricative and English labio-dental fricative, and the difference between English lateral liquids and palatal flap are not necessarily among the most serious difficulties communicatively or socio-linguistically. Many English speakers cannot pronounce distinctly /l/ and /r/. The late Right Honourable Sir Alec Douglas Home couldn't, but it didn't stop him from becoming Prime Minister. Nobody found this impediment embarrassing except perhaps when he said that the whole country was looking forward to a general erection.

Perhaps the biggest impediment to native speakers' perception of Japanese English speech is the marked separateness on the lenis-fortis scale. To put it simply, the native English hearer who is accustomed to hearing a fortis language has difficulty perceiving many consonants in a lenis language. While speakers of a lenis dialect are not impeded in their perception by "excessive" noise, speakers of a fortis dialect are impeded by hearing insufficient noise. The importance of fortis English for communication with other non-native speakers of English must also be stressed. A fortis Japanese English will be more easily comprehensible to German and Russian speakers of English than a lenis Japanese variant would be, although French speakers of English will have less difficulty with lenis Japanese English.

1.ii. Suprasegmental and syntactic phonetics

Compared with segmental phonetics syntactic phonetics is easily accessible and varies little between Englishes, and where there are variations in word stress, e.g. *témporarily/temporárilý*, these variations

do not obstruct meaning and are not strongly marked sociolinguistically.

The word “intonation” has different meanings according to which English teacher is using it so that to avoid confusion with timbre suprasyntactics with which I will deal later I will use the term “syntactic phonetics”. Similarly the word “tone” has different meanings and I will use the term “tune”. The tune system of English is basically one of rising tunes and falling tunes. Rising tunes are used for asking yes/no questions, falling tunes for statements and WH questions. Speech usually links tunes together for syntactic and semantic purposes. Rising is followed by falling for oppositional focus, which includes series rising with final falling for lists. Add to this sentence prominence stress. e.g. Q. “Did you say he is in a bad temper?” A. “No, she is in a bad temper.” This last, however, can be dealt with theoretically by given/new prominence tunes

This is the whole range of syntactic phonetic resources necessary for communication of information including specifying subject/predicate, theme/rheme, and given/new.

Control of these phonetic features of English would seem to be essential for adequate communication of information.

Part II Neo-Macro-Phonetics

It may have surprised some readers that I have dismissed syntactic phonetics as such a simple and accessible performance. It is clear that speaking English, or any language, cannot be as simple as this.

The apparent simplicity or the apparent complexity of syntactic phonetics arises from the position of the teacher/researcher with respect to theoretical linguistics.

A model of language as something used for giving, soliciting and receiving data will be quite adequately served by the limited range of tunes that I have outlined above.

Communication, however, is very far from being limited to information giving and getting. Information is only one of the functions of language, and while it is not very often absent from an utterance it is also not very often the only function of an utterance, or even the most important function. Halliday's functional system includes a representational function only after his exposition of the instrumental, regulatory, interactional, personal, heuristic and imaginative functions.

In the adult language these seven functions are usually performed in various combinations, and it is the components other than the representational (content/information) function that commonly determine the phonetic performance.

This performance is as yet little researched and not widely understood, but work has been going on in this field for many years. It may seem a bewildering task to extend and elaborate the representational/syntactic system to produce models which are applicable to all the combinations. Separation into syntactic and suprasyntactic studies, however, affords a much less unwieldy tool. The late Dr. Masao Onishi [2] called his studies "speechology". V.Stefanescu-Drăganesti [3] calls his studies "neo-macro-phonetics". Olga Mindrul [4] calls her work "timbre suprasyntactics." Olga Akhmanova and Peter Wain [5] ventured into what we called "philological phonetics". A variety of terms exists, but all these and other study variants have in common that they are concerned with functional speech, that is speech which communicates other human functions as well as data.

For this paper, as for the basis from which I teach speech to university students, I find Mindrul's timbre suprasyntactics to be the

most immediately practical.

Most of the work on timbre suprasyntactics has been done in what was the Soviet Union, where Marxist-Leninist dialectical materialism prevailed for so long. It is not surprising, then, that for Mindrul's functional phonetics the functions of language are reduced to two. One function is called the "intellective" function and parallels what I have called the information function which Halliday [6] calls the representational function. The second function is the function of impact, which subsumes all the other functions.

That the political, economic and social systems of Marxism have been thrown out does not mean that useful tools such as this dialectical opposition should also be thrown out with the bath water.

II.i. Timbre I

The term "timbre" is commonly used to refer to the speech sound of an individual. A person's timbre is composed of his normal pitch range, pitch movement, loudness, voice quality, speed of speech, and pausation. It is this timbre that permits us to recognise a person when we hear just a small fragment of speech on the telephone. Loudness, pitch, pitch movement, speed and pausation are easily accessible and demonstrable concepts. Voice quality refers to the way in which vocal resonance is amplified in distribution among the chest, the mouth, the pharynx and the nasal cavity; it also refers to the place of speech on the lenis-fortis scale. This timbre is static and will be discussed later under the heading of speech portrayal. It is the functional basis of articulation for information or the intellective function.

II.ii. Timbre II

This term, "timbre II" is used to refer to all the departures from

the static norm of timbre I that go to make up the resources of functional phonetic performance. Departures are in:

a) loudness

By making our voices more or less loud than normal we communicate that the listener should perceive the function of impact, or one or more of the Hallidayan functions subsumed. (We thunder, we murmur.)

b) pitch

Similarly a departure from the norm above or below the pitch range of timbre I communicates the function of impact. (We squeak, we rumble.)

c) pitch movement

When the degree of pitch movement used in timbre I is extended beyond the norm for the intellective function we again communicate impact. This resource is further exploited by reduplication and inversion of the tunes of syntactic phonetics. (We swoop and dive.)

Q 1 Have you seen the foreign papers?

A 1 Yes, they're on the table.

Q 2 Have you seen the foreign papers?

A 2 Yes! Isn't it terrible!

In this case Q 2 not only reduplicates the intellective tune but also inverts it. Similarly Shylock does not ask "Hath not a Jew eyes?" but "Hath not a Jew eyes?" inverting the tune for a yes/no question.

d) voice quality

For impact voice quality will also change. As an example we can hear the voice change from its normal timbre I to a pharyngeal voice which sounds husky. Similarly if speech changes on the lenis/fortis scale some function other than information is implied.

e) speed of speech

When the speaker accelerates or decelerates the speed of speech this produces a departure from information giving to impact. Even when

very marked deceleration in speech is necessary for the communication of information, possibly because of the complicatedness of the information or because of an unfavourable acoustic environment, this implies the inclusion of another complex function.

f) pausation

The use of pausation is similar to changing speed of speech. I include it separately since it is quite possible to speak very fast between long pauses to produce another combination of functions.

II.ii. Minus timbre

This term is used to refer to the creation of impact by the deliberate omission of timbre II where it would be expected. (It is almost always expected.) It is perhaps a misleading term since it might suggest that the absence of timbre II returns us to the default position of timbre I. It is almost always present in ironical speech; it is what is often referred to as deadpan humour; just as the deadpan comic shows no expression on his face, so his speech betrays no function other than information. W.C.Fields's "There is no such thing as a tough child. Boiled for a couple of hours...." is the classic example of such speech. When Calvin (below) says, "I believe angels are everywhere." Hobbes answers,



~~“You do?”~~ Hobbes is not checking whether he has heard correctly, he is using timbre II to express the personal and interpersonal functions.

When Holly Martins [7] says, ~~“You do.”~~ he is not confirming that Mr. Crabbin represents the CRS of GHQ, he is using minus timbre to convey that he is not interested in Crabbin, the CRS or GHQ.

Minus timbre is a useful concept and much work remains to be done in finding the ways in which native speakers recognise irony from speech apparently independently of the content.

Part III Speech portrayal

It is unusual to find in published audio English language courses any adequate treatment of functional phonetic performance. Similarly, English by radio deals with timbre suprasyntactics casually as an unspecified divergence from syntactic phonetics.

Even exercises in syntactic phonetics are sometimes perverted from their purpose by modelling tunes in the first model sentence and in subsequent drills innocently offering given and new tunes or list tunes. This is because to the native speaker who is modelling the tunes the core of the sample becomes the given and the lexical variations in the subsequent drills become the new. This is because the native speakers are tempted to make their performance more exciting. Let's face it: drills, if they really are to be drills, are dull since by their very nature they are repetitive.

The excitedness of the native speakers modelling syntactic phonetics sometimes wanders into timbre II producing quite inappropriate effects. I offer as an example the following:

~~“I saw your mother at the swimming pool the other day.”~~

With the pitch movement I have indicated here the syntactic pattern necessary to convey the information becomes suggestive of some secret knowledge, conspiratorial, even salacious. It sounds as if the least important thing to be communicated was the information when it was, in fact, the most important.

So it seems that we have two learning situations, the first in which syntactic phonetics is taught but not timbre II, the second in which syntactic phonetics is taught with an inappropriate random accretion of timbre II. What would benefit the advanced learner of English is training specifically in timbre II, so that he or she can know which features of timbre II to produce for which function and which features to avoid.

I suggest that study of speech in fiction is a superlatively rich source of information on both syntactic phonetics and timbre phonetics. Most speech portrayal offers both timbre I and timbre II. Timbre I is the static or characterising aspect of speech portrayal, timbre II is the functional/dramatic aspect of speech portrayal.

III.i. Timbre I — static speech portrayal

Static speech portrayal is that indication of speech which the writer uses to establish character. Initially in fiction this is rather a stereotyping than the creation of original characters. The powerful have deep pitched voices, the weak have high pitched voices, squeaky even. Relaxed and confident people have low, well modulated voices, the arrogant speak with a drawl, the authoritative clip their speech, and the aristocrat speaks with clip and drawl.

Let us be clear that we are talking here about fiction; in real life some powerful people have high squeaky voices and some quite ineffectual people have deep well modulated voices.

Static portrayal is mostly used as a basis from which dynamic portrayal can be developed and training in this aspect of speech has little value for the learner of a language, who is generally stuck with his natural timbre I.

It is of course immensely important for actors who can create a character by modifying the natural timbre I; Laurence Olivier was a master of this craft and acquired a new timbre I for different characters. It is also possible to meet people whose timbre I changes for speech in a foreign language where a different timbre I from the native language is the norm; some native English speakers produce a French norm of timbre I when speaking French, and there are resonant Russian speakers who comfortably slip into the comparative huskiness of English when they change from one language to the other.

However, for all but the most gifted learner the timbre II of a foreign language is a quite adequate new performance.

III.ii. Dynamic speech portrayal.

Dynamic speech portrayal in fiction is far from being beyond the comprehension of the advanced specialist learner, and the following examples offer excellent challenges.

- 1 he answered in a low voice
- 2 he repeated firmly
- 3 said her husband off-handedly
- 4 said the woman sharply
- 5 said his mother with satisfaction
- 6 exclaimed the boy
- 7 said Rhoda with reluctance
- 8 said she huskily
- 9 she asked tentatively

- 10 said she diffidently
11 continued the conjurer impassively (minus timbre)
12 retorted Carlier with indignation

There is no need to search for such examples; they jump off the page. Sometimes the portrayal reaches to remarkable extravagances with a sense of comedy created by the speech portrayal such as:

- 13 Walter Mitty's voice was like thin ice breaking.

Retrieving from fiction and creating timbre II is for the student a later very advanced development of the skill of reading aloud and with inner voice.

Dynamic speech is quite easily accessible in models. We have professional timbre II users to help us here. Models are available in plenty in the wealth of film videos that can be used for training in speech.

It is important, however, to choose films carefully for their timbre content. Some of the "best" films from the point of view of their literariness, their cultural, moral and educative value and their filmic excellence can prove a millstone. I refer to films which are largely polemic like "My Fair Lady" or diaries like "Dances with Wolves" or declamations in series. Whatever one may think of the film "Casablanca" (and I admire it greatly) there is no gainsaying that it is remarkable for its minus timbre qualities: Q. "What is your nationality?" A. "I'm a drunkard." or "I remember every detail. The Germans wore grey; you wore blue." These are all far from suitable training for functional speech skills. But there are many films that are excellent models for speech and it is for the individual teacher to find his or her own from which to devise appropriate class work.

Conclusion

Language teaching has been successful for centuries without basing teaching on a theoretical linguistic model whereas the recent decline in foreign language skills, although it cannot be ascribed to advances in theoretical linguistics, has paralleled its rise.

The spoken language is a real and material thing and can be described in real terms as I have done above. It is this real language that teachers should study, understand and hand on to their students, and offer to them as a model for learning and a subject of research.

References

- [1] in *Study of Sounds*, Vol XX (1984) ed. Onishi (Phonetic Society of Japan)
- [2] in *Bulletin of The Phonetic Society of Japan*, Nos.138 & 139
- [3] in *Study of Sounds*, Vol XX (1984) ed. Onishi (Phonetic Society of Japan)
- [4] in *Philological Phonetics*, (1986) ed. O.S.Mindrul (Moscow State University)
- [5] in *Study of Sounds*, Vol XX (1984) ed. Onishi (Phonetic Society of Japan)
- [6] Halliday, M.A.K. (1973), *Explorations in the Functions of Language* (Edward Arnold)
- [6] Halliday, M.A.K. (1979), *Language as a Social Semiotic* (Edward Arnold)
- [7] *The Third Man* (film, 1949), dir. by Carol Reed the text by Graham Greene.

Bibliography

- Baker, Ann (1981) *Ship or Sheep?* (CUP)
- Gimson, A.C. (1980), *The Phonetics of English* (Edward Arnold)
- Jones, Daniel (1975), *An Outline of English Phonetics* (CUP)
- Katz, J.J. & Fodor, J.A. (1963), *The Structure of Semantic Theory* (Language no.39)
- Onishi, Masao (1981) *A Grand Dictionary of Phonetics* (Phonetic Society of Japan)
- Pike, K.L. (1967), *Language in Relation to a Unified Theory of the Structure of Human Behaviour*, 2nd revised edition, (The Hague, Mouton)
- Trim, John (1975) (Illustrated by Peter Kneebone) *English Pronunciation Illustrated* (CUP)