

CANADIAN THESES ON MICROFICHE

THÈSES CANADIENNES SUR MICROFICHE



National Library of Canada
Collections Development Branch

Canadian Theses on
Microfiche Service

Ottawa, Canada
K1A 0N4

Bibliothèque nationale du Canada
Direction du développement des collections

Service des thèses canadiennes
sur microfiche

NOTICE

The quality of this microfiche is heavily dependent upon the quality of the original thesis submitted for microfilming. Every effort has been made to ensure the highest quality of reproduction possible.

If pages are missing, contact the university which granted the degree.

Some pages may have indistinct print especially if the original pages were typed with a poor typewriter ribbon or if the university sent us an inferior photocopy.

Previously copyrighted materials (journal articles, published tests, etc.) are not filmed.

Reproduction in full or in part of this film is governed by the Canadian Copyright Act, R.S.C. 1970, c. C-30. Please read the authorization forms which accompany this thesis.

THIS DISSERTATION
HAS BEEN MICROFILMED
EXACTLY AS RECEIVED

AVIS

La qualité de cette microfiche dépend grandement de la qualité de la thèse soumise au microfilmage. Nous avons tout fait pour assurer une qualité supérieure de reproduction.

S'il manque des pages, veuillez communiquer avec l'université qui a conféré le grade.

La qualité d'impression de certaines pages peut laisser à désirer, surtout si les pages originales ont été dactylographiées à l'aide d'un ruban usé ou si l'université nous a fait parvenir une photocopie de qualité inférieure.

Les documents qui font déjà l'objet d'un droit d'auteur (articles de revue, examens publiés, etc.) ne sont pas microfilmés.

La reproduction, même partielle, de ce microfilm est soumise à la Loi canadienne sur le droit d'auteur, SRC 1970, c. C-30. Veuillez prendre connaissance des formules d'autorisation qui accompagnent cette thèse.

LA THÈSE A ÉTÉ
MICROFILMÉE TELLE QUE
NOUS L'AVONS REÇUE

Canada

The Nature of the 'Production Grammar' Syllable

by

Ann Stuart Laubstein

A thesis submitted to
the School of Graduate Studies and Research
of the University of Ottawa
in partial fulfilment of the requirements for
the degree of Doctor of Philosophy

Ottawa, Ontario

Canada

March, 1985



Ann Stuart Laubstein, Ottawa, Canada, 1985



UNIVERSITÉ D'OTTAWA
UNIVERSITY OF OTTAWA

To Stuart, Andrew and Karl.

The Nature of the 'Production Grammar' Syllable

Ann Stuart Laubstein

Abstract

This dissertation is concerned with the nature and role of the syllable in on-line speech production. I argue that the syllable is a necessary construct in a production model by illustrating that sub-morphemic segments or segment sequences which interact in perseverations, anticipations or exchanges always belong to the same syllabic position. This constraint on segmental interactions is unaccounted for without recognizing the syllable as a production unit. Moreover, we argue that this constraint is explained if the programmer is assumed to be selecting syllabic constituents at the processing level at which these errors occur. A segmental error is the result of a mis-selection from among elements of the same type - namely the syllabic constituents: onset, peak or coda, or the sub-constituents of the onset and coda. We show that this mis-selection of identical constituents also characterizes word blends.

In recent years, the syllable has been the focus of a great deal of research in the competence domain. The structures proposed have been many and diverse, aimed at accounting for various aspects of competence. We refer to the 'production grammar' in our title wishing to distinguish the units and processes involved in on-line production from the rules and representations of the 'competence grammar' - believing how close the relationship is between the two to be an empirical question. We consider in detail a number of proposed 'competence syllables' in Chapter 1, setting the stage for our investigation of the 'production syllable' in Chapter 2.

This Chapter, the major portion of the dissertation, is aimed at deriving the internal constituent structure and the external boundaries of the production syllable on the basis of those errors which do and do not occur. For instance we argue against the rhyme as a viable constituent of the production unit on the basis of the fact that this sequence of segments is very rarely involved in errors; indeed onset-peak sequence errors are almost as frequent. Sequences of segments are involved in errors only if they are

constituents of some type: phrase, word, morpheme, etc., or as we demonstrate, onset, peak or coda. Not all constituents are involved in errors, i.e. this is not a criterion for constituency, only one kind of evidence - but it would appear that there is no aspect of segmental errors that requires the unit 'rhyme' in its explanation. One common argument in support of the rhyme has been based on the fact that errors more often involve onsets than peaks or codas - the claim being that such errors point to a bifurcation of the root syllable node into onset and rhyme. We believe that this high frequency of onset errors is better understood as implicating word structure, and not syllable structure. There are more peak errors than word-medial onset errors, for instance.

We propose a syllabification algorithm for generating the production syllable we derive; a structure which is a composite of various characteristics proposed by competence theorists but not isomorphic to any single one. We also propose error-generating devices for producing segmental errors and word-blends: devices which we suggest mirror the production process when these errors occur.

In the final portion of this dissertation we consider how our hypothesis, i.e. that syllabic constituents are being processed when segmental errors and word blends occur, could be incorporated into some current models of the production process. Our hypothesis can easily be integrated into Garrett's most recent model and it thereby provides more support for the frequent claim that items in the mental lexicon are syllabified. We conclude that sub-morphemic segmental errors occur in the 'fleshing out' of a skeletal prosodic frame whose bottom level consists of a string of empty syllable templates constantly available for fleshing out.

ACKNOWLEDGEMENTS

I should like to thank a number of people who have helped me over the years: Ian Pringle, who introduced me to linguistics, Stan Jones, who has been my reference library and librarian from the very beginning, and Doug Walker and John Jensen who sustained my initial interest in phonology. I particularly want to express my gratitude to P.G. Patel, without whom this would never have been written. I would also like to thank Jean-Luc Nespoulous for raising pertinent and interesting questions at my thesis defense - questions which will form part of the focus of my future research. Finally I want to thank Deborah Higgins who somehow managed to translate my handwriting.

TABLE OF CONTENTS

	PAGE
ABSTRACT.....	i
ACKNOWLEDGEMENTS.....	iii
INTRODUCTION.....	1
FOOTNOTES TO INTRODUCTION.....	11
CHAPTER 1: THE SYLLABLE: WHY DO WE NEED IT? AND WHAT DOES IT LOOK LIKE?.....	12
1. Introduction.....	12
2. External and Internal Arguments.....	13
(i) External Arguments.....	13
(ii) Internal Arguments.....	20
3. Syllable Templates.....	27
(i) Introduction.....	27
(ii) Fudge's Template (1969).....	28
(iii) Selkirk's Template (1980).....	32
(iv) The Cairns and Feinstein Template (1982)...	44
(v) Hooper's Template (1976).....	48
(vi) Kahn's Template (1976).....	52
(vii) The Clements and Keyser Template (1983)....	57
(viii) The Halle and Vergnaud Template (1980).....	62
(ix) The Fujimura and Lovins Template (1978)....	66
(x) Conclusion.....	68

	PAGE
(xi) Characteristics Which Shall be the Focus of Chapter 2.....	69
(xii) 'victrolas' Syllabified.....	71
APPENDIX.....	73
FOOTNOTES TO CHAPTER 1.....	74
 CHAPTER 2: THE PRODUCTION SYLLABLE: THE ARGUMENT FROM SPEECH ERRORS.....	79
1. Introduction.....	79
2. Errors as Evidence: Problems and Pseudo-Problems..	88
(i) Falsifiability: A Pseudo-Problem.....	88
(ii) The Reliability of the Data: A Problem.....	95
3. Assumptions Underlying Our Investigation.....	100
4. The Data Base of Our Investigation.....	111
5. Giving Empirical Substance to the 'Syllabic- Position-Constraint' Hypothesis.....	126
6. Testing the Hypothesis.....	135
(i) Results.....	135
(ii) Discussion: General.....	139
(iii) Problematical 'V+r' Errors.....	146
(iv) Excursus: Syllabifying Intervocalic Glides..	152
(v) Problematical Onset/Coda Errors.....	158
7. The Rhyme: Where Is It?.....	173
8. The Ambisyllabic Segment Isn't.....	187

	PAGE
9. Onsets: Internal Structure.....	196
(i) Data Selected for Analysis.....	196
(ii) Results.....	199
(iii) Discussion.....	204
10. Codas: Internal Structure.....	219
(i) The Paucity of Data.....	219
(ii) Speculations.....	220
11. Word Blends.....	232
(i) Analysis: Method and Results.....	232
(ii) Discussion.....	234
12. Within-Word Errors.....	251
13. Summing Up Chapter 2: Comparing the Production and Competence Syllables.....	254
(i) Introduction.....	254
(ii) External Boundaries and Root-Dominated Nodes.....	255
(iii) The Structure of Onsets, Codas and Peaks....	264
FOOTNOTES TO CHAPTER 2.....	269
CHAPTER 3: INTEGRATING THE SYLLABLE INTO PRODUCTION MODELS.....	275
1. Introduction.....	275
2. The Shattuck-Hufnagel Account of Segmental Errors: the Role of the Syllable.....	282

	PAGE
3. Incorporating Our Hypothesis Into Garrett's Model.....	295
(i) Introduction.....	295
(ii) The Model: General.....	296
(iii) The Positional Level Slots.....	299
(iv) Slots and Syllabification.....	301
(v) Slot Construction.....	305
(vi) Segmental Slots and Stranding Errors.....	310
(vii) Segmental Slots and Segmental Errors.....	322
(viii) Conclusion.....	335
4. The Fromkin Account of Segmental Errors.....	342
5. Crompton's Account of Segmental Errors.....	347
6. Conclusion.....	354
FOOTNOTES TO CHAPTER 3.....	366
CONCLUSION.....	369
FOOTNOTES TO CONCLUSION.....	379
REFERENCES.....	380

0. INTRODUCTION

Linguists working within the Chomskyan paradigm attempt to explain how it is, that on the basis of relatively poor input data from experience, all people (barring gross pathology) acquire a natural language relatively quickly, and without overt teaching. The assumption is that man comes prepared with an innate structure whose parameters are fixed on the basis of experience, a structure which is 'triggered' by experience, not 'shaped' in any stimulus-response sense. The linguist attempts to specify the innate structure on the basis of which, given the input of experience, the adult grammar is deduced - "...a grammar that may be very intricate and will in general lack grounding in experience in the sense of an inductive basis" (Chomsky, 1982, p. 4). The rules and representations that make up this adult grammar are intended to characterize the linguistic knowledge which, along with knowledge from other domains, underlies the speaker-hearer's capacity to actually use language. That is, the term 'grammar' is used to refer to "...the system of rules represented in the mind of the speaker-hearer...used in the production and interpretation of utterances..." (Chomsky and Halle, 1968, p. 4).

This grammar, this internalized linguistic knowledge, is referred to as the 'competence' of the speaker-hearer.

Competence is distinct from, but a prerequisite for, speech production and perception - capacities which belong to a theory of performance and which will include other systems of knowledge, constraints due to limits on time and memory and so on. It is competence which underlies the speaker's ability to make judgements regarding the well-formedness of an (in principle) infinite set of sentences, judge paraphrase, ambiguity and anaphoric relations among and within sentences, recognize possible versus impossible neologisms, and so forth. From this it is clear that within the Chomskyan paradigm linguistic theory belongs to a more all-embracing theory of mind, and that a grammar is more than an account of the distributional facts with regard to the elements of some language. Within this paradigm a crucial distinction is made between an observationally adequate grammar and one which could be considered descriptively adequate. The former would account for all and only the grammatical sentences of some language, whereas the rules and representations of a descriptively adequate grammar, besides achieving this observational level, would also account for the speaker-hearer's above-mentioned meta-linguistic capacities. The generative grammarian not only aims at producing such a descriptively adequate grammar, but also attempts to explain how such a highly articulated and complex rule system can be uniformly acquired given the limited

evidence available. A linguistic theory will reach this level of explanatory adequacy if it specifies an innate deductive structure rich enough to determine the particular systems attained.

The equation of the linguist's grammar with the mental grammar has been the impetus for a great deal of psycholinguistic research in the last couple of decades - research aimed at getting performance support for the psychological reality of the constructs of the competence grammar. It must be remembered, however, that it is not incumbent upon these constructs that they receive this kind of support; behaviour is not a criterion for possession of knowledge, only one kind of evidence for it. It is assumed that the grammar, that is, the speaker-hearer's competence, will be incorporated as a basic component in a theory of use. But how close the relationship is, between the units and processes of production and perception, and the rules and representations of a competence grammar, is an empirical question. The mapping may be more or less indirect or direct.

For instance, in on-line processing does the brain prefer to select from a store or compute? As a concrete example from the conscious level, consider mathematical computation. To find the sum of 673 and 984 one does not go

to a store and select the correct sum; but the store of basic number facts such as $3 + 4 = 7$ and $7 + 8 = 15$ are used to compute the larger sum. One could use finger counting or some such thing to compute the basic facts but one need not. It is more efficient in this case to combine computation and selection. Similarly, a descriptively adequate grammar will have to account for the fact that it is within the competence of a speaker of English to predict (or compute) the value of the feature [round] given the feature [-back], or the shape of the regular past tense morpheme given the shape of the root-final segment. Native speakers are competent with regard to the vowel inventory and the morphophonemic processes of English, and the grammar must account for this knowledge. But whether real-time processing involves selection from a store of fully specified morphemes or computation of the redundancies is an empirical question. It is plausible that there is an optimal value for both the size of the stores and the amount of computation when it concerns fast, efficient processing. Therefore the types of unit or representation stored and selected need not be, and probably are not, minimally redundant from a competence grammarian's point of view. Thus the relationship between competence constructs and performance constructs could be quite complex. Constructs required to capture linguistic knowledge may be quite far removed from those required for processing efficiency.

Another interesting question regarding the relation of grammars to processors concerns the variety of languages, hence grammars. If grammar differences are, for the most part, a function of the different values assigned to a set of open parameters, one is forced to consider whether the mechanisms for processing different languages reflect these different choices or are independent of them. For instance, how closely related are the mechanisms employed in the production and perception of so-called free-word-order languages like Walpiri, and fixed-order languages like English? Could the processing of a language like English, whose syllables include final consonant clusters, crucially involve or make reference to these final segments, given that numerous languages allow only open syllables? To be more concrete. Visual perception appears to be highly modular. Specific neurons fire in response to highly specific visual stimuli (Marr, 1984). Could certain neurons function only in response to codas despite the fact that not all languages have codas? If language perception at some level, is, like visual perception, common across the species, then, given the different grammars, we have another reason to expect that the mapping between performance constructs and those of competence may be less than direct.

These considerations take nothing away from the fact that the best working hypothesis regarding processing models



is that "...the computational decomposition of language processing does, indeed, reflect the grammatical decomposition of the language faculty" (Garrett, 1982, p. 21). Our observations merely serve to reiterate Chomsky's point that behaviour is not a criterion for possession of knowledge, only one kind of evidence for it. Similarly, competence constructs are not a criterion for the forms of processing. The rules and representations of competence can not be used as evidence against the structures of performance theories, nor can the latter be used against competence. Both may only be used in support of one another. The highly articulated grammars being produced within the generative paradigm provide a reasonable starting point for processing models. These grammars begin to have the requisite elaboration and richness of detail regarding the complex system that must underly the capacity for speech. As such they provide the best working hypotheses for processing theorists.

In this thesis we will examine how recent developments in generative theory regarding the syllable relate to production. We consider in detail the form of the various syllable templates proposed (Chapter 1), we will illustrate their strengths and weaknesses as viable units of production (Chapter 2), and finally we consider how the 'production syllable' can be incorporated into current production models

(Chapter 3). The data we use as evidence in support of a production syllable and the form it takes, are spontaneous, submorphemic segmental errors from normal populations. Our data sources have included errors from the published literature¹ and a couple of hundred errors from our own corpus.

We have chosen spontaneous errors from normals recognizing that as a result we inherit not only the assumptions that are implicit in the use of these data, but also the problems associated with their collection and the inferences derived from them. The use of errors as data for deducing characteristics of the production model assumes that errors manifest the workings of the normal production mechanism. That is, errors and error-free speech are both the output of the normal mechanism. The route from idea to articulation is the same in both cases. Some analyses of pathological speech have suggested that certain errors may be the output of new routes, that is, some other than normal mechanism may be used. (See Butterworth, 1980 and Caplan, 1983). However, the linguistic analysis of normal errors has illustrated the highly regular, non-random nature of such errors, providing no evidence against the assumption that the regular mechanism is responsible for them. Indeed the regularity of error types has provided the framework upon which most models of production have been built. Moreover

the analysis of errors has vindicated the assumption that many of the units and processes of the competence grammar are used in speech production. For instance, errors have been used to support the reality of phonological constructs such as segments and features - units not confirmable at the level of acoustic analysis. The independence of stem and affix, of phrasal and word stress from segments, and of open and closed class vocabulary shows up in errors; again there is nothing in the error-free utterance that can allow the inference of such independence. The modularity of the grammar has also shown up in production. Errors suggest that there are separate, temporally ordered processing levels.

Similarly, we find that errors are highly instructive regarding the syllable. We will show that the syllable is as necessary in the production model as it appears to be in the competence grammar. The great majority of those errors that have been traditionally analysed as segmental errors are better interpreted as syllable constituent errors, if the constraints governing their occurrence are to be explained. It also appears that open class items in the 'production lexicon' are syllabified. (We will use the terms 'production' or 'mental' lexicon to refer to the lexicon used by the production mechanism in on-line processing.)

It is clear that spontaneous error corpuses can provide a rich source of data for the production modeler, but there are well-known problems in the use of such data.² First there are the methodological problems relating to data collection, involving such things as selective attention in the collector and the differential degrees of detectability of different errors. In addition, the problem of interpreting frequency distributions is exacerbated by the minimal amount of surrounding context recorded along with the errors.

Probably one of the most difficult problems with regard to data facing any theorist in any domain, is that of distinguishing those facts which are crucial from those which are best deemed as 'noise'. Without a well-motivated and fairly intricately structured model of the production process this distinction among errors is probably impossible to make. But it appears that we find ourselves in a vicious circle at this point. Production models are to a large extent derived on the basis of errors; therefore these models can not be used to distinguish the 'crucial' errors from those which are random. This is where theories of competence are indispensable. They supply the production modeler with independently motivated and sufficiently elaborate hypotheses regarding the structures and processes underlying production and perception. These provide a framework within which

processing models may find a reasonable starting point for determining which facts to account for, at least in the beginning. They are the best working hypotheses.

The domains that competence and performance theories attempt to account for are different but they intersect. We may expect that in the end the two will benefit each other. Indeed, even now it is not completely a one-way street. Production errors can be used to help to choose between alternate proposals by competence theorists.³ Facts which Newtonian physics could not account for came to be explained within the theory of aerodynamics. But aerodynamic theory could only be developed on the basis of Newtonian insights (cf. Kahn, 1976). We might liken Newtonian and aerodynamic theory to the theories of competence and performance respectively.

In this thesis we attempt to show that production facts lead to the selection of a particular model of the syllable, similar to, but not isomorphic to, the various templates proposed by competence theorists.

NOTES TO INTRODUCTION

¹See in particular Fromkin 1973, 1980, Garrett 1976, 1980, 1982, Shattuck-Hufnagel 1984, Mackay 1970, and Goldstein, M., 1968.

²See Section 1 Chapter 2 for more extensive discussion.

³For instance there are different notions regarding the underlying form of grammatical morphemes. Hooper (1976) suggests that allomorphs are all listed with the appropriate environment specified. So for instance the regular plural morpheme would include the following list: [əz], [z], [s], and their environments. Chomsky would have one underlying form, [z], phonological rules determining the others. It appears that, in general, the processor is liable to errors when it is selecting from among items in a set, whether the set includes syntactic, semantic or phonological types. That is, the production mechanism is liable to mis-select when it has a choice among similar types. It never mis-selects the plural morpheme, however. Consider the following errors:

(i) pots and pans -- pons and pats
 [s] [z] [z] [s]

(ii) bloody students -- bloodent studies.
 [s] [z]

Since there is no phonotactic constraint against [panəz] or [studiys] ('the prawn is' and 'stewed Easter eggs') this immunity to mis-selection is surprising. If all three allomorphs were available one would expect the wrong choice now and then. The production data is more in line with the assumption of a single underlying representation.

CHAPTER 1

THE SYLLABLE: WHY DO WE NEED IT? AND WHAT DOES IT LOOK LIKE?

1. INTRODUCTION

Within the generative paradigm, as exemplified in the Sound Pattern of English (SPE) (Chomsky and Halle, 1968) the phonemic representation of a sentence is a linear string of segments and syntactic boundary markers, the latter specifying such things as word and morpheme boundaries. Segments at this level are bundles of phonetic features, each of which may take one of two values, plus or minus. These representations include all of the information required for the correct application of the phonological rules, whose output will be a phonetic representation. This latter representation is also a string of segments, but boundary markers have been removed and phonetic features may be multi-valued. Thus, within the SPE framework, the claim is made that all generalizations regarding the sound system of languages can be captured using linearly ordered phonetic feature matrices, and morpho-syntactic boundary features. In the past decade these claims have been challenged. Arguments have come from both within and outside the grammar for the re-introduction of the syllable into phonological theory. Various proposals have been made regarding its incorporation

into generative theory, ranging from boundary marker (Hooper, 1972) to a binary branching, multi-tiered structure (Kiparsky, 1979). In the next section we will summarize the arguments advocating the inclusion of the syllable from outside, and within the grammar. In the final section of this chapter we will consider in detail various proposals regarding the exact form the English syllable should take. In Chapter 2 we will see how these structures relate to the 'production syllable' - a structure we will derive on the basis of sub-morphemic segmental errors. Finally, in Chapter 3, we will suggest how this structure is to be incorporated into a model of production, in light of its capacity to explain certain characteristics of segmental errors.

2. External and Internal Arguments

(i) External Arguments

As discussed earlier, a generative grammar attempts not only to generate all and only the grammatical sentences of some language, but also aims at accounting for a speaker's knowledge about his language.¹ One aspect of this meta-linguistic competence, not directly captured within the SPE framework, is the speaker's knowledge with regard to syllables. Although no satisfactory physiological or acoustic correlate of the syllable has been found (see Kahn,

1976, pp. 15-17) speakers know what counts as a syllable in their own language. (Although cf. Bell, 1975) Speakers of English judge 'ala' as bisyllabic and 'apt' as monosyllabic despite the data of chest pulses which may indicate 'ala' is monosyllabic (Pike, 1947). The ability to count syllables is part of the meta-linguistic competence of even very young children and indeed Mehler et al. (1981) argue that "...the syllable is one of the basic unanalysed initial structures the child uses" (p. 295). It also seems to be the case that speakers from different languages hearing the same utterance will assign it different numbers of syllables. Chinese speakers apparently would assign three syllables to 'skates' (Pike, 1947), whereas we, as speakers of English 'KNOW' that it has only one. This kind of 'knowledge' would seem to be on a par with speaker competence regarding what counts as a distinctive opposition in his own language. The physically measurable difference between an aspirated 't' and an unaspirated one is irrelevant at the level of phonemic analysis in English, as is the chest pulse measure of the syllable in characterizing syllable competence.

McNeill and Brown (1966) showed that in 48% of the cases where speakers were in a tip-of-the-tongue (TOT) state they were able to correctly remember the number of syllables in the searched-for word. Similarly, aphasics searching for words can often remember the number of syllables but not the

segments (Jakobson, 1968). Fudge (1969) also points out that young children will often completely botch up the phonemes in a syllable but they won't forget the syllable's presence. In addition, Fay and Cutler (1977) argue that the syllable structures of malapropisms are closely related to the syllable structures of the intended form, as in 'apartment' for 'appointment'. The data from TOT, lexical access and recall studies all converge. It appears that the 'production' or 'mental' lexicon includes information about syllable structure.

Arguments for including the syllable have also come from phonetics. It appears that the length of segments in a syllable is inversely related to the number of segments in the syllable (Selkirk, 1978). For instance, a CVC syllable is about as long as a CVCC syllable despite having fewer segments. Using the syllable rather than the segment as a concatenative unit has also led to a high degree of success in speech synthesis (Fujimura and Lovins, 1982).

It seems too, that the syllable is the most permutable unit in language games (Sherzer, 1976), and syllable-by-syllable concatenation has been used to explain the staccato quality of apraxic and dysarthric speech (Kent and Rosenbek, 1982). Sussman (1984) even speculates that canonical syllable forms are equivalent to neuronal templates, networks

of neurons "...where each consonant and vowel position is associated with a specific cell assembly" (p. 169). He believes that the relatively inviolate nature of phonotactic constraints in aphasic speech and left hemisphere dominance for language may be explicable on this basis. Native language syllable structure constraints can also be used to explain certain second language learner errors. Native Spanish speakers will pronounce English 'skate' as [ɛskɛyt] because there is no syllable initial 'sk' cluster in Spanish. Saudi students will pronounce English 'party' with an initial 'b' since there is no 'p' in Saudi-Arabic, and English speakers will pronounce the initial cluster in 'Nkrumah' as [ɔn.k] (the dot (.) indicates syllable boundary) since there are no syllable-initial [nk] clusters in English.

It seems clear that the syllable is a psychologically real unit of our meta-linguistic competence, that it is implicated in the structure of the mental lexicon, that it can be useful in explaining characteristics of aphasic speech and second language forms, and in accounting for surface phonetic facts, where it seems to be an independent concatenative unit of some sort. As pointed out earlier, however, grammar-external arguments are not a criterion for determining the constructs of the competence grammar. They are only suggestive. The syllable could be a processing construct, functionless at the competence level. The grammar

must account for the speaker's meta-judgements regarding the syllable but it need not do so directly. This aspect of competence may already be accounted for indirectly. For example, it may be captured in the sequences of C's and V's already marked for the feature [syllabic], with the C's in pre- or post-syllabic position. Of course, in this case if a formative consisted of more than one syllabic segment the medial C's could not be distinguished as pre- or post-. But is this information necessary? If not, then syllables are reducible to other 'kinds', independently necessary in themselves. 'Kinds' are the classes of things and events about which pertinent generalizations can be made (cf. Fodor, 1975).

Kohler (1966) argues that, for this reason, the syllable is an "unnecessary concept" (p. 207). He also laments that if a C (cluster) occurs intervocalically, one that is capable of occurring both pre- and post- then a syllabic division is impossible. This would be the case in forms like 'any' and 'basket' where the following divisions are possible: an.y, a.ny, bas.ket, bask.et, and ba.skot. This, the problem of 'fuzzy' boundaries (not Kohler's term), makes the syllable impossible to define according to Kohler (op cit.). Indeed, speakers do disagree on the syllabic division of intervocalic C's. Syllabification decisions will often be a function of perceived morpheme boundaries (Davidson-Neilsen, 1974). For

instance English speakers would be more apt to syllabify a form like 'discomfort' between the 's' and the 'c', i.e. between the morphemes - dis.com.fort, than they would a form like 'discuss' which would more likely be judged monomorphemic and syllabified as di.scuss. This is related to Kohler's third objection to including the syllable in phonological theory. He notes that syllabifying on the basis of terminal allophones leads to clashes between morphological and syllabic segmentation. For instance the dark 'l' in 'coolish' would lead to a division after the 'l' whereas the light 'l' in foolish would define a boundary before the 'l'. It's not clear to me how this argument is relevant. Phoneme boundaries clash with morphological boundaries too, but presumably Kohler wouldn't suggest that phonemes be eliminated.² This problem relates to a more general problem - the problem of where/when syllabification takes place (and the problem of which items are lexically listed (see note 2). Pointing out that syllabification does not appear to be distinctive, some competence theorists argue against the lexicon including syllable structure (cf. Bell and Hooper, 1978, p. 6 for instance). However, if the feature [syllabic] is conceived of in terms of the syllabic node to which a segment is attached (e.g. Kiparsky, 1979, and Clements and Keyser, 1983), or if stress is distinctive and unpredictable, and is captured in the forms of prosodic structure, syllable

and foot, then the lexicon will have to include syllable structure. Selkirk (1980) argues along these lines. Whatever turns out to be the best analysis, Kohler's argument against the syllable on the basis of the potential clash between syllable and morpheme boundary remains irrelevant.

More interesting is the problem of fuzzy boundaries. Kohler points out that an important aspect of the phonological structure of English is the frequency of indeterminate syllable boundaries with regard to intervocalic C's. Syllabification, he argues, would only obscure this. Substantial arguments can be found to diminish the force of this and his other pertinent objection - i.e. that the independently necessary sequences of C's and V's indirectly can account for generalizations attributed to a syllabic unit. Arguments will be summarized that demonstrate that, within the competence theory, syllable boundaries must be delimited, and too, that the syllable may be a hierarchical structure and not just a linear sequence of C's and V's. Secondly, fuzzy boundaries can be insightfully captured with the appropriate formalism and distinguished from those boundaries which are not indeterminate (cf. Kahn, 1976). If fuzzy boundaries are an important aspect of the structure of English as Kohler claims, then this is an important distinction to capture. The syllable conditions numerous phonological processes whose expression is significantly

simplified, and whose phonetic motivation is concomitantly accounted for, with formal recognition of the syllable (Vennemann, 1972). And too, there are good arguments for stating phonotactic constraints in terms of the syllable.

(ii) Internal Arguments

Within the SPE framework phonotactic constraints are expressed at the level of the morpheme. Stanley (1967) recognized the need for the statement of possible syllable structure conditions to supplement the Morpheme Structure Conditions (MSC's). Many Semitic morphemes, for instance, consist of three C's, which makes the morpheme an inappropriate level for the statement of collocational restrictions. MSC's are intended to capture the speaker's ability to recognize that certain phoneme sequences are not possible morphemes in his language - the standard example being 'bnik' in English. Sommerstein (1974) has pointed out that it makes no sense to consider this aspect of competence at any level other than close to surface. Because 'bnik' is an opaque form, in that the relationship between the systematic phonetic and the systematic phonemic levels is not biunique, there is no possibility of exercising this ability except at the surface. Forms such as Spanish '-ndo' and German 'Rad' are rejected out of hand by speakers, because like 'bnik' they are impossible surface forms. Sommerstein

argues that this level of competence can best be captured in terms of surface syllable forms. This argument is well-taken with regard to accounting for the 'bnik' competence, but it is not clear that speakers are only cognizant of what forms are possible surface forms. It would seem reasonable to distinguish such knowledge from the knowledge speakers may have, conscious or unconscious, regarding the shape of possible morphemes, including bound and free. At some level, speakers of English know or 'cognize' (Chomsky's term in Chomsky, 1980, p. 70) that morphemes may be single non-syllabic C's (e.g. plural 'z' and past tense 'd') or even $\emptyset$ (e.g. two sheep), and at some level Arabic speakers 'cognize' that morphemes may consist of sequences of three C's (e.g. 'ktb' for 'write'). Presumably speakers would nevertheless reject such morphemes as possible free surface forms, indicating that these two types of knowledge are different. Because the grammar aims to account for all purely linguistic knowledge possessed by speakers, both these types will have to be accounted for - probably in different ways. Speakers have access to the results of mental processes - i.e. they can make judgements regarding pronounceable surface forms. Surely, however, speakers have no privileged access to the processes themselves - i.e. there is no reason to assume that speakers may be capable of judging the underlying forms of morphemes or the rules which

determine their surface shape. What the best way is, of accounting for both types of knowledge, is not obvious, but there are good arguments for stating sequential constraints in terms of the syllable, not the morpheme. Collocational restrictions seem to have universal characteristics which appear to be a property of syllables, not morphemes. Most convincing arguments in this regard have come from Hooper (1976) and Mohanan (1979) among others. Hooper demonstrates that numerous sequence restrictions can be accounted for in terms of the strength and weakness of different positions in the syllable, interacting with the inherent strength or weakness of different C's. Her theory accounts for many of the universals of clustering listed in Greenberg (1963). For example, in an initial C_1C_2 cluster, C_1 is in the stronger position, so you will tend to find stronger C's in this position. For instance in English we have 'tree' but not 'rtee'. Her notion of strength at the limits corresponds fairly closely to Jespersen's (1909) notion of sonority, and Kiparsky's (1979), and Saussure's (1959) of aperture - the earliest reference we know of, to this quality, being Sievers (1893). Syllable final position is weak so a relatively stronger C will be required in C_2 position. We have 'art' but no 'atr', for instance. The following Greenberg implications support her claims: if Nasal, Liquid (NL) then Consonant, Liquid (CL) (read 'obstruent' for 'consonant'), if

CN then CL, and with respect to syllable final clusters, if LN then LC, and if L $\begin{bmatrix} C \\ +voice \end{bmatrix}$ then L $\begin{bmatrix} C \\ -voice \end{bmatrix}$. As Greenberg points out, the thousands of permutations of C's possible are by no means utilized by the languages of the world. Clearly there are enormous constraints on clustering. Mohanan (1979) argues that the sonority hierarchy built into a universal template, can only capture the universal tendencies of syllable structure - it can not distinguish what he considers to be absolute, universal restrictions from mere tendencies. For instance, obstruents don't function as syllabic nuclei (although cf. Hoard, 1978) and there are restrictions on the occurrence of syllabic consonants and non-syllabic glides (Pike's (1947) contoids and vocoids, respectively). He argues that to capture these restrictions and the tendencies regarding sonority requires the use of markedness conventions on the attachment of segments to syllabic or non-syllabic nodes. Cairns and Feinstein (1982) argue that both a universal template and markedness conventions are required. The latter will define marked and unmarked attachments to particular syllable nodes in a hierarchical structure. Fudge (1969) agrees that collocational restrictions require a hierarchical structure for their statement - MSC's, being finite state in character, are inappropriate. Hooper has also pointed out that MSC's are often duplicated in the phonological rule component

(except that the latter change features and the former do not). It seems indisputable that there are numerous sequential constraints which are universal and that the syllable and not the morpheme is the implicated unit when it comes to the perspicuous statement of these universal generalizations.

Equally persuasive arguments for positing a syllable 'kind' revolve around the dynamic phonological processes of languages, where the syllable often seems implicated in both segmental and suprasegmental domains. Reference to syllable boundaries, initial and final, syllable constituents, and to the whole syllable is often required in the statement of phonological rules. Devoicing of obstruents in German (/Rad/ → [Rat]) and point assimilation of nasals to following consonants in Spanish (/un gato/ → [uŋgato]) take place when these consonants are in syllable final position. Segment loss is common in final position, an example from French Canadian being /tablo/ → [tabl] → [tab]. In general, weakening processes are characteristic of syllable final position. Kahn points out that in order to designate syllable final (or next to final) position, SPE must use the environment (C) $\begin{Bmatrix} C \\ \# \end{Bmatrix}$. He notes that this violates the requirement that environments in phonological rules be natural classes. A natural class will be specified by a conjunction of features, common to the members of the class.

- an unnatural class will require a disjunction (Halle and Clements, 1983). Segment and word boundary, having no features in common, constitute a most unnatural class. The allophonic variants of vowels will often require the use of this environment. For instance in Spanish [e] occurs in open syllables and [ɛ] in closed syllables, and in Danish [amerika] alternates with [emrika] due to the closed syllable which results with optional deletion of the second vowel. Syllable initial position is often required for picking out the environments for strengthening processes. Chomsky and Halle (1968) note that there is a general strengthening of initial C's in Tswana which can't be handled simply within the SPE model. The aspiration of voiceless stops in English occurs in syllable initial position and indeed the strong and weak allophones of 't' in English can only be accounted for with reference to syllable positions (Kahn, 1976).

Donegan and Stampe (1978) claim that they have not found a phonological "...process limited to the domain segment" (p. 28) and argue that the syllable is often the implicated domain.³ For instance the domain of regressive nasalization of sonorants in English appears to be, obligatorily, the syllable and optionally, the foot (their 'measure'). Hence we have [wil.ỹōm] or [wĩĩ.ỹā̃m] but not *[wil.yā̃m]. Similarly, the domain of pharyngealization in colloquial Egyptian Arabic appears to be the syllable. (Broselow,

1976, cited in Selkirk, 1980, p. 577). For handling such processes, the whole syllable must be designated, not simply initial or final syllable boundaries.

We find too, that syllabic constituents may be the focus of phonological rules, particularly in the area of suprasegmental phonology. It is well known that stress rules in many languages are sensitive to syllable quantity. Languages may differ in how they measure heaviness (see Prince, 1983), that is, heaviness may be any VC tautosyllabic sequence, or only VR (R=sonorant) or VV tautosyllabic sequences. What is universal however, is that only tautosyllabic segments following vowels count in the measurement - never preceding tautosyllabic segments. Many current analyses of stress systems account for this property in terms of the syllable constituent, 'rhyme'.

To sum up, there appears to be no doubt that the maximally simple and explanatory statement of a language's phonotactic constraints and phonological processes will require the unit, 'syllable', to refer to. The syllable itself, its boundaries, and its constituents are all implicated when it comes to dealing with the phonological systems of natural languages. Earlier, we gave examples of arguments for the syllable which were outside grammar internal considerations. On the generative assumption that

the grammar lies at the root of our capacity to use language, it is not surprising that the evidence from these two sources converges - both pointing towards the syllable. In the following section we will consider various proposals regarding the precise form this construct should take.

3. Syllable Templates

(i) Introduction

We shall see that all of the templates described in this section reflect both the theoretical framework within which the structures are formulated, and also the particular facts their formulators are interested in accounting for. For instance, the Hooper template comes out of a linear framework, and the Selkirk and Kahn proposals are couched in the language of the recently developed non-linear theories. At the same time, the structures proposed by Fudge (1969) and Cairns and Feinstein (1982) for instance, reflect a focus on collocational restrictions, whereas the Kahn syllable reflects his interest in segmental processes.

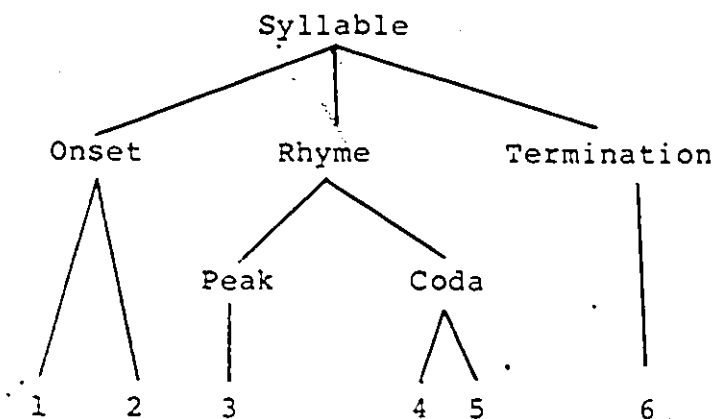
The templates we have chosen to consider in detail are clearly not all of those proposed in the literature. They are a subset which we feel highlight the interesting and different claims regarding when syllabification takes place, and what syllabification algorithm is used. These algorithms

produce different syllable boundaries and different internal structures. We will not reiterate arguments which duplicate one another. For instance, arguments supporting the Rhyme constituent on the basis of collocational facts will be advanced only once, despite their having been advanced by different theorists. Following our description of the different templates we will summarize the differences, pointing to what we will look for in the speech error data - the subject matter of Chapter 2.

(ii) Fudge's (1969) Template

Fudge believes that syllables have two functions: to provide the basis for dealing with prosodic phenomena, and to account for collocational restrictions. He focuses on the latter in 'Syllables' and the structure he proposes reflects this. He argues that collocational restrictions hold within constituents, and hence are not stateable in the finite state terms of the SPE Morpheme Structure Conditions (MSC). Therefore the constituent structure he assigns to his syllable reflects his notions as to where constraints on segment sequences are tightest, i.e. within onset, rhyme and coda, as illustrated in Figure 1.⁴ No constraints are considered to hold between onset and rhyme or rhyme and the termination node.

Figure 1
Fudge's Template



To account for most constraints, he assigns sets of segments, his 'systems', to the various positions. All consonants can occur in positions 1 and 5 unless 2 and 4 are occupied by sonorants.⁵ 'sC' clusters are given a uni-segmental status and are found in positions 1 and 5. He argues for this analysis because (i) otherwise he would require an initial position limited to just one segment 's', (ii) it allows a symmetrical treatment of 'sC' clusters initially and finally and (iii) he feels it solves the problem of whether [sp] for instance, is underlying /sb/ or /sp/. His Termination node is normally a morphological place for such things as plural and past tense, and can only be filled word finally. By

including the Termination node, he believes he can explain why *'glansp' is an impossible form and 'glanst' is not - i.e. more generally, why the length of the rhyme is limited except when the final C's are word final and coronal, and usually are preceded by a morpheme boundary. He requires a special constraint, however, to rule out *'glaymp'. The structure he proposes is not capable of handling this. His peak node consists of one position because his is a phonological, not a phonetic syllable, and he derives most diphthongs from tense vowels. The diphthongs 'oi' and 'au', though not so derived, fit into position 3 and presumably are considered, like 'sC' clusters, unisegmental structurally.

For constraints that can not be predicted on the basis of his systems he uses a set of violable and non-violable constraints which apply to elements within the systems and take the form of implications. For instance, if 2('r') $\rightarrow$ -4('r'). That is, if you have an 'r' in position 2 you can't have one in position 4. Although he has no constraints that involve the onset and peak positions, he has a number which apply across constituents, to positions 2 and 4 and 1 and 5, for instance. It's not clear how he can justify his cross-constituent constraints, given that he derives his constituent structure on the basis of the assumption that constraints are internal to constituents. And, in fact, although he lists none, there do exist some (albeit few)

restrictions on onsets and peaks in English, and other languages. For instance, "voiced fricatives and /Cl/ clusters are excluded before /ʊ/ (and).../Cl/ clusters are excluded before /Vl/..." (Clements and Keyser, 1983, pp. 20-21); hence we find no forms such as [ɔʊt] or [blʌl].

By allowing unfilled nodes, he requires no proliferation of syllable types - one type suffices. Positions 1 and 5, one might term the heads of the onset and coda; they are to be filled if there is a choice between 1 and 2 or between 4 and 5. So, for instance, the 'r' in 'run' would be assigned to position 1 and the 'n' would be put in position 5 although these could occur in positions 2 and 4. He assigns them this head-like status on the basis of constraint facts. For instance, there are not as many restrictions on 'n', 'm', 'l' and 'r' when they are in position 1. 'lull' is a possible form but *'Clull' is not. This head status also captures the symmetry between positions 1 and 5; all consonants occur in each.

In short, Fudge's syllable is derived on the basis of constraints that hold between C's in pre- and post-vocalic clusters and on those which are often found between vowels and post-vocalic clusters. We shall see in the following section that Selkirk's syllable also reflects this interest.

(i) Selkirk's Template

The syllable structure Fudge proposes reflects his interest in one of the two functions he attributes to the syllable (i.e. accounting for collocational restrictions). He ignores the other function, that of providing the basis for prosodic phenomena. If the same structure, namely the syllable, is to perform these two seemingly quite disparate functions we should find that facts from both areas point to the same type of structure. And indeed, Selkirk, who is interested in both these functions (among other functions)⁶ proposes a structure which is, as she points out, similar to Fudge's. First we'll consider her arguments from the area of prosody and subsequently those, not advanced by Fudge, in the area of phonotactics.

Selkirk's theory of prosodic structure is an extension of the non-linear metrical framework introduced by Liberman (1975) and applied in Prince (1975) and Liberman and Prince (1977). Their intention was to explain the unique characteristics of the feature [stress], and the stress rules of the SPE framework, by incorporating a metrical tree and grid into the theory. Relations of relative prominence were to be captured in a binary branching tree, where prominence was a strong/weak (s/w) relation on sister nodes, and these trees were subsequently mapped onto a metrical grid.

Although their system accounted for relative prominence, they maintained the stress feature to account for absolute prominence (i.e. plus stress versus minus stress). Selkirk argues that absolute stress should also be defined in terms of prosodic tree structure (cf. also Halle and Vergnaud, 1980, and McCarthy, 1979), specifically the stress foot - one of her six prosodic categories.⁷

Her claim is then, that all 'stress phenomena are a function of prosodic structure - and in particular that all generalizations and idiosyncrasies of the English stress system are to be accounted for in terms of prosodic units. Relative prominence within the word will depend on how feet are grouped together⁸ and absolute stress will be a matter of internal foot structure.

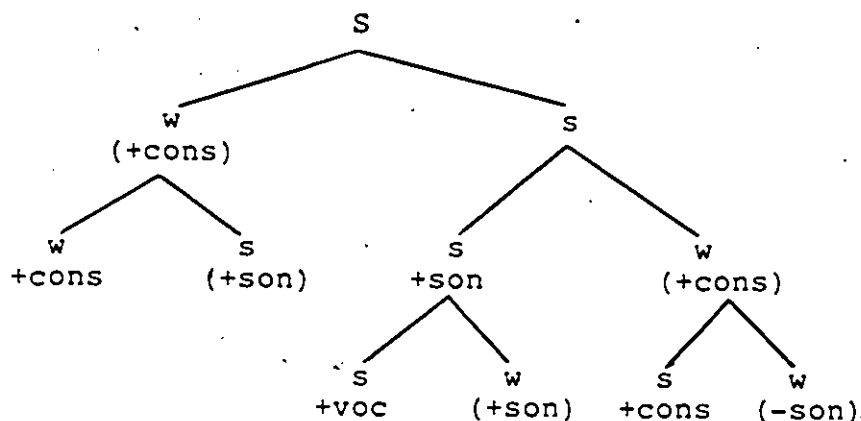
There are three characteristics of the English stress system with regard to absolute stress that Selkirk wants to account for (listed in (1)):

- (1) (i) all tense vowels are stressed,
- (ii) some lax vowels in closed syllable are stressed and some are not,
- and (iii) lax vowels in open syllables are never stressed
 - (a) word medially before a stressed syllable or
 - (b) word finally.

Since all generalizations regarding absolute stress are to be captured in terms of the stress foot, these characteristics of the English system will have to be accounted for in terms of the internal structure of the stressed syllable of the stress foot. If stress is a matter of syllable geometry, then generalizations such as (1) (i) above regarding tense vowels, will have to be captured in terms of syllable structure and not in terms of phonetic features, for instance. (See Prince, 1983, for arguments against this geometric notion of the stress values of diphthongs and long vowels.)

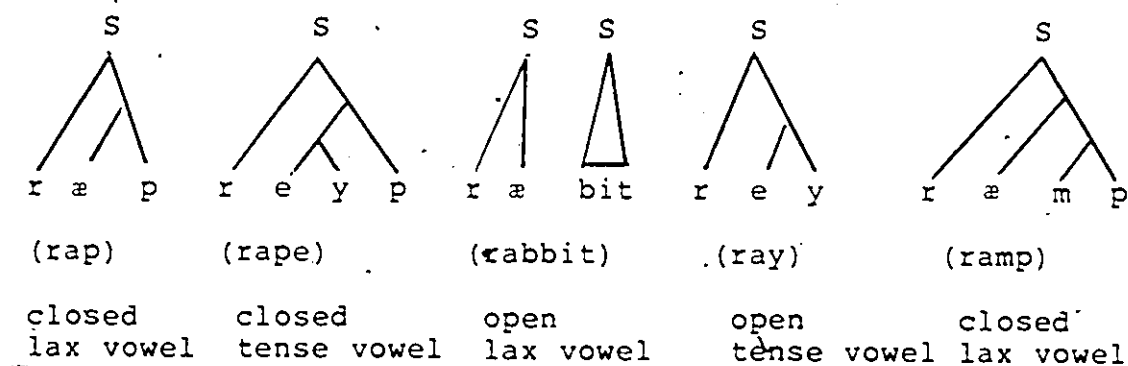
Indeed the syllable template Selkirk proposes could be used to distinguish tense vowels from lax vowels. However, because she argues against labelled nodes, and because she considers that $C\bar{V}$ (tense vowel) and CVC are both branching rhyme structures (1980, p. 583, fn. 11),⁹ and because her template functions as a well-formedness condition on Underlying Representations (UR's), the structure that results from the actual syllabification of a lexical item will often be unable to make this distinction. To see why this is so, consider the difference between her template and the actual syllable structures assigned to lexical items, as illustrated in Figures 2 and 3 respectively.

Figure 2
Selkirk's Template



- N.B. (i) only glides may be sisters of [+voc]
(ii) s and w reflect the sonority hierarchy

Figure 3
Some Selkirk Syllabifications



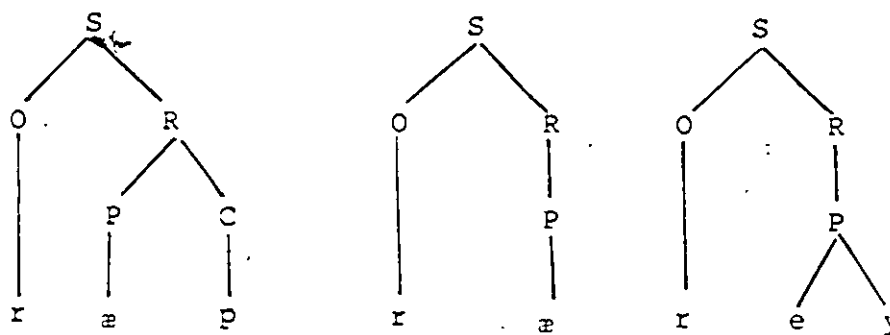
It is clear from Figure 3 that the structure assigned to a lax vowel syllable closed by one C is identical to that

assigned to an open tense vowel syllable ('rap' versus 'ray'). Thus the syllable structure per se can not be used to distinguish tense vowels from lax vowels unless the lax vowel is in an open syllable, or the tense vowel is in a closed syllable ('ra.bbit', 'ramp' and 'rape').

Within her template per se the distinctions between closed lax vowel syllables and open tense vowel syllables are possible, but the syllabification of lexical items must include labelled nodes to make these distinctions, and thereby capture the generalization (i) - (iii) in (1) above. In Figure 4 the nodes are labelled and the distinctions are there.

Figure 4

Selkirk's Nodes Labelled



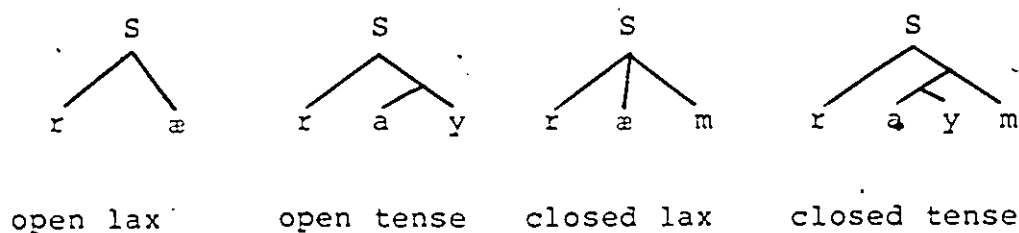
A syllable with a branching peak will always be the left branch of a bi- or tri-syllabic foot, or the single branch of a monosyllabic foot. That is, it will always be stressed. A

syllable whose peak and rhyme don't branch (an open lax syllable) will never be stressed unless it is the left branch of a foot or an initial monosyllabic foot. A syllable whose rhyme branches and whose peak does not (i.e. a lax vowel closed syllable) will be the left or right branch of a foot, depending on whether it is stressed or not. Selkirk argues against the Halle & Vergnaud (H&V) claim that both relative and absolute stress values can be wholly accounted for in terms of s/w relations and branchingness; she maintains that absolute prominence is a matter of foot structure, which is only partially predictable on the basis of branching. In particular, the behaviour of closed (branching) lax vowel syllables is unpredictable. For instance the final syllable in 'gymnast' is a stress foot whereas the final syllable in 'modest' is not - yet these syllables are identical structurally. The HV type analysis would have to claim arbitrarily that the final syllable in 'modest' was unbranching. Selkirk believes that this irregularity in the stressing of such syllables is best accounted for in the repository, available for all unpredictable facts - the lexicon. That is, she proposes that foot structure is lexical. Generalizations (1) (i) - (iii) above will be made in terms of foot structure. As we have suggested, in order to capture these generalizations in this way it would appear that she will have to allow labelled nodes in the lexical

structures she assigns. (Some of the nodes in her template are labelled with features but these can not be used. See below.) In fact, Selkirk explicitly argues against labelled nodes because syllable constituents "...have no independent privileges of distribution" (1980, p. 573). Feet and syllables are building blocks whereas the coda, for instance, is only "...the second position within the rhyme" (ibid.).¹⁰ Labelling the rhyme, peak and coda is sufficient but not necessary. She could use a flatter syllable structure and leave nodes unlabelled, as in Figure 5.

Figure 5

Syllabic Structure Distinctions Without Labelling



In this syllable the onset, peak and coda are all dominated immediately by the root syllable node. The rhyme analysis makes different claims, however. It captures the fact that prosodic systems in general (if they are quantity sensitive systems) are sensitive to the segments which follow the initial C's - that is, the vowel and following material. This suggests an initial bipartite division into onset and

rhyme. Selkirk observes that the branching of syllables into onsets and rhymes is never prosodically relevant; it is "...only the rhyme's branching which appears to make a difference"¹¹ (1980, p. 572, fn. 7).

Selkirk could maintain her rhyme structure, and leave nodes unlabelled, and make the generalizations she aims at (i.e. (1) (i) - (iii) above) if she were to leave empty nodes. This is how Fudge syllabifies. He argues that this prevents a plethora of syllable types. Selkirk opposes this because she believes that every instance of a syllable token should not be asked to characterize the notion 'possible syllable in L' any more than every NP should have to exemplify all possible NP's.

We see then that the structure that Selkirk believes is required to handle prosodic phenomena is indeed very like the structure motivated independently on phontactic grounds by Fudge. As mentioned earlier, Selkirk's syllable has more than a prosodic function within the grammar. It is to account for the 'bnik' competence, i.e. phonotactic constraints. An impossible form of English is one which can not be parsed into well formed syllables and feet. The notion 'possible syllable in L' will be defined by the template in Figure 2 and a set of collocational restrictions "...somewhat in the spirit of Fudge, 1969...but with

differences" (1978, p. 10). There are indeed differences in the two proposals. Although Selkirk's syllable reflects what she refers to as the IC principle of phonotactics (the more closely related structurally the "...more subject to phonotactic constraints the two position slots are" (1978, p. 2) as does Fudge's, it's not clear how she will state the Fudge-type collocational restrictions without labelled nodes. Fudge even requires sub-constituents of the onsets and coda to state his restrictions. He identifies six sub-constituents, i.e. six position slots and states a number of constraints in terms of these sub-constituents. For instance in his restriction '2(4) $\rightarrow$ -4(4)' the parenthesized numbers refer to sub-systems of segments and the 2 and 4 designate the second and first positions in onset and coda respectively. It is non-adjacent constraints like these that lead him to conclude that a finite-state-type account of phonotactics won't do. Selkirk's use of s/w relations will mean that she won't have to distinguish the onset from the coda when she rules out sequences like 'lt-', 'rd-', '-pm' and so on. In the onset, the labelling is w/s and in the coda it is s/w, labels which can handle such facts. Since s/w is a relation on sisters it can not handle the above non-adjacent type of constraint, however. This may, though, be a positive aspect of Selkirk's model and not a failing at all. It's not clear that syllable constituent structure is

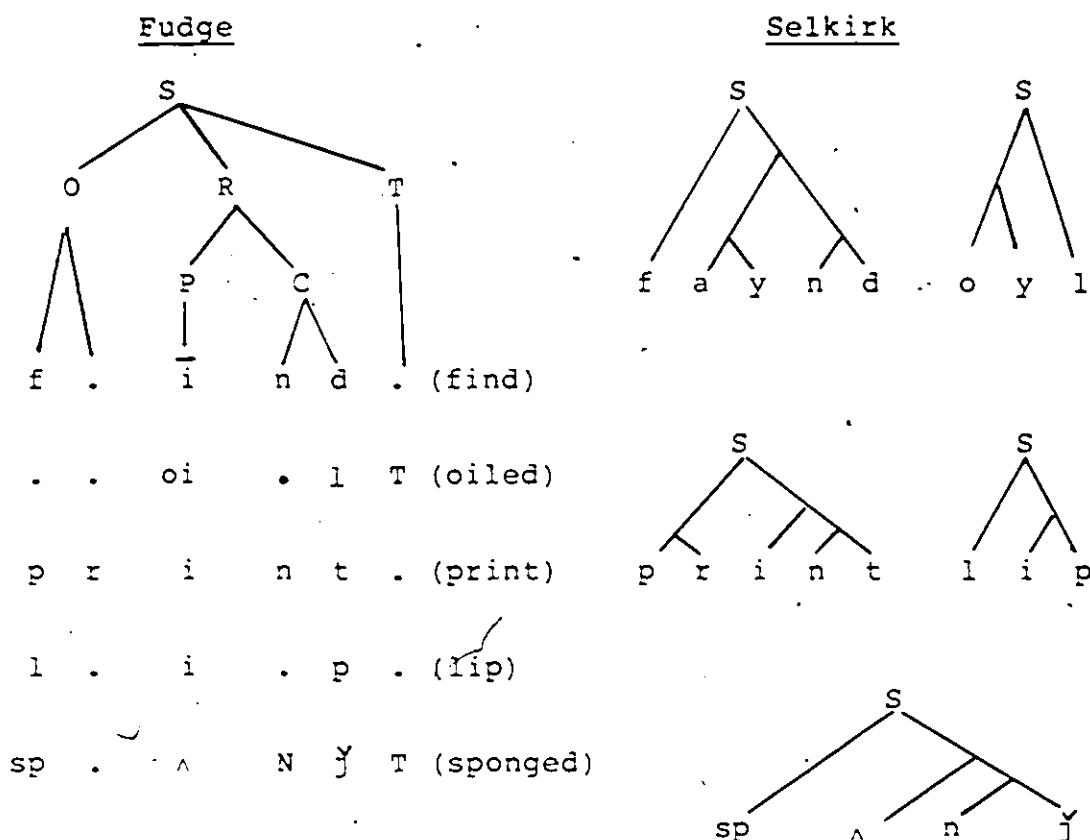
at issue here. Positions 2 and 4 do not form a constituent of any kind and yet there are constraints holding between the two. This goes against the IC principle of phonotactics, which would seem to be a not unreasonable principle. This type of non-adjacent constraint holds within the constituent 'syllable,' not within any subconstituent. Perhaps this type of constraint will find an explanatory account within the 'syllable-is-segment' theory of Fujimura and Lovins (F&L) (1978). They suggest that at the phonetic level, syllables, not segments, are bundles of features. Possibly there are constraints on what features may co-occur in such 'segments'. (See below.) Selkirk's syllable also differs from Fudge's in that it includes a bi-positional peak node¹² and no termination node. Her syllabification domain is the non-branching stem (a syllable boundary occurs between all stems and neutral affixes) and presumably morphemes like past and plural are not included in the stem so she won't often need this node. Forms like 'axe' will have to be treated as exceptional.

Fudge and Selkirk also differ on which branch of the coda is filled when the coda is a single segment. Fudge assigns all such to position 5, the right branch of the coda, whereas this branch is parenthesized in Selkirk's syllable. They agree that the left branch of the onset is chosen - in Selkirk's syllable - the [+cons] branch of the onset. (Since

her nodes are unlabelled the results of syllabifications involving single C's will never indicate which node was chosen.)

This brings up another difference in the two syllables. Selkirk doesn't label constituents but she assigns major class features to all constituents but the rhyme. These are only a part of the template, however, and not included in the prosodic structure of individual lexical items. This feature labelling at least partially takes over the function of Fudge's assignment of systems to positions, however it's not clear how she would syllabify words with syllable initial glides since they are not [+cons] and presumably shouldn't be allowed in the left branch of the onset. It's not clear how glides can occur even in the right branch of the onset, since although they are [+son] they are not [+cons] and the right branch is dominated by a [+cons] node. Normally we would expect this feature to percolate downwards to all nodes it dominates (see Vergnaud, 1977). And presumably this is the reason the rhyme has not been labelled. Because of these problems, it doesn't seem possible to let features take over the function of constituent labels.

The structural differences in the Fudge and Selkirk syllables can perhaps best be seen in Figure 6 where various words are syllabified according to their different notions.

Figure 6

N.B. (i) the dots, (.) indicate empty positions.

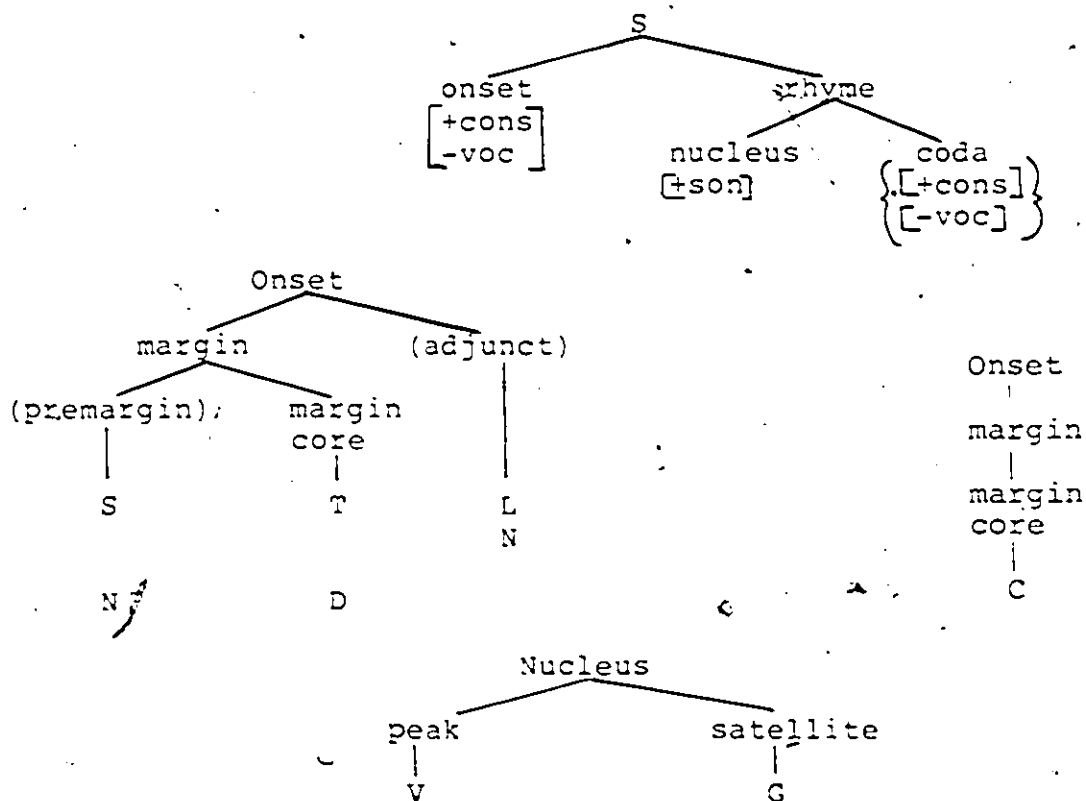
In contrast to Selkirk, the proponents of the next syllable we will look at argue that labelled syllable nodes are crucial; they are the basis for word-formation rules and stress rules, besides forming the foundation for a theory of markedness. And markedness, in turn, determines how forms are syllabified.

(v) The Cairns and Feinstein Template (1982)

Cairns and Feinstein (C&F), 1982, propose a syllable which is similar structurally to Selkirk's (and Fudge's (cf. also Pike and Pike (1947) and Hockett (1955))), the major differences being (i) their insistence that nodes are labelled and (ii) the hierarchical structure assigned to the onset. They do not argue in support of any particular coda structure; however, they suggest that it might be a mirror image of the onset. The structure they propose is illustrated in Figure 7.

Figure 7

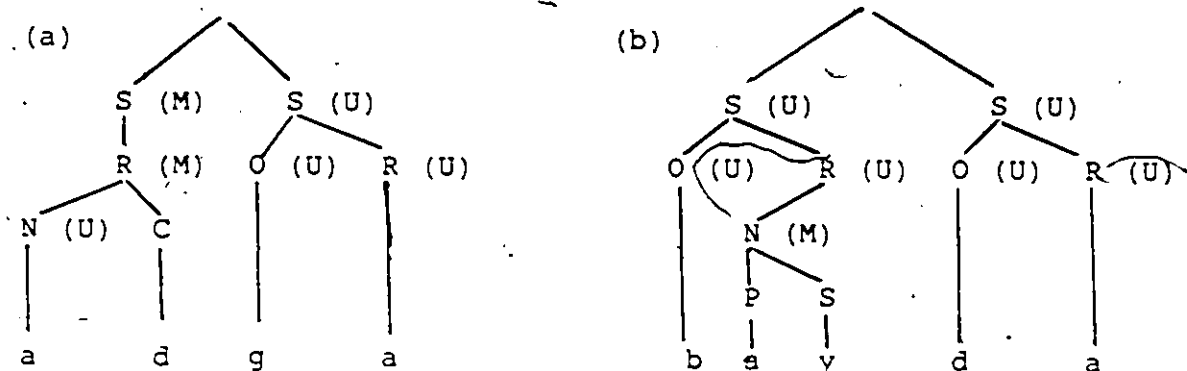
Cairns and Feinstein Template



Their goal is to define the universal unmarked and marked syllable, by specifying the structures always found, those that are usually found and those that are never found. In other words they aim at identifying the parameters that are left open and the values that these parameters may take. The structure they propose reflects this concern. They argue that all languages aim at unmarked syllabification in the lexicon, and markedness determines the syllabification, - not the sonority hierarchy. They propose that syllables and their constituents be generated by Phrase Structure (PS) type rules whose expansions are marked or unmarked. For instance (2) - (4) generate the marked and unmarked syllables, rhymes and nuclei.

- | | | | | |
|-----|---|---|----------------|----------|
| (2) | S | → | Onset Rhyme | Unmarked |
| | S | → | Rhyme | Marked |
| (3) | R | → | Nucleus | Unmarked |
| | R | → | Nucleus Coda | Marked |
| (4) | N | → | Peak | Unmarked |
| | N | → | Peak Satellite | Marked |

A language having the lexical items 'adga' and 'bayda' would syllabify them such that the fewest nodes were marked, as illustrated in Figures 8a-b.

Figure 8

N.B. (i) Node expansions have been omitted that are not relevant here.

In 8(a) the 'd' has not been assigned to the onset of the second syllable because the only unmarked segments in the premargin are fricatives or nasals followed by stops and voiced stops respectively (S, T, N, and D in Figure 7). As 8(a) shows, this language allows codas in 'd' but syllabifying 'bayda' such that 'd' belonged to the first syllable would result in a marked rhyme in the first syllable of 'bayda' and a marked root syllable node in its second syllable. Such a syllabification would not be maximally unmarked so it is an impossible syllabification of 'bayda'.

They argue too that features are marked according to which syllabic position they occupy. For instance (4) claims that the unmarked value of sonorant is minus if it is found in the margin core.

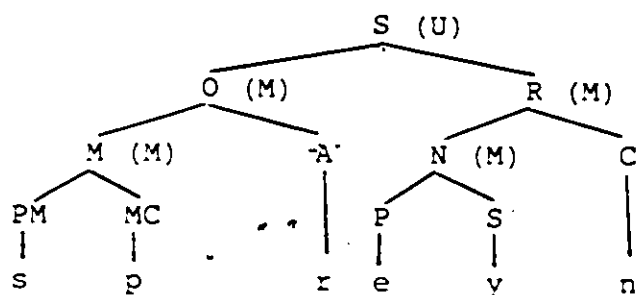
(5) [U son] → [-son] / [M^{Core}]

Labelled nodes, they argue, are not only required for markedness theory but also for stress rules and word-formation rules. For instance in Yidin CVV is heavy but CVC is not. That is, the stress rules must refer to marked nucleus and not to branching rhymes. This is similar to the distinction in English, discussed above regarding lax and tense vowels. They claim that reduplication in Sanskrit and Gothic involves the constituents margin core and margin respectively.

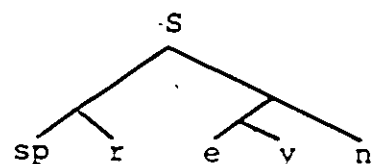
Thus they give four different reasons for having labelled, rather than unlabelled nodes, one of the two differences we noted between theirs and Selkirk's templates. The other difference is the structure assigned to the onset. They argue that Kiparsky, whose syllable is based on the sonority hierarchy (Kiparsky, 1979) must give 'special dispensation' to 'sC' clusters in English, yet such clusters are fairly common. They argue therefore, that the premargin constituent is a universally available option, albeit a marked one, and 's' is an unmarked segment in this constituent. Like Selkirk, they argue that syllabification is lexical; she, because stress in English is unpredictable and they, because morphological rules must make reference to syllable structure. The different lexical structures they

would assign to 'sprain' are illustrated in Figure 9.

Figure 9



Cairns & Feinstein



Selkirk

Fudge, Selkirk and Cairns and Feinstein have all proposed similar hierarchical structures on the basis of arguments from independent domains. Hooper, whose structure we will next consider, aims at accounting for some of the same data and yet she does so within a linear framework.

(v) Hooper's Template (1976)

Like Fudge, Selkirk, and Cairns and Feinstein, Hooper argues that the syllable is essential in accounting for co-occurrence facts, but unlike them she does not recognize any IC principle of phonotactics. Sequential constraints reflect the fact that different segments are inherently weaker or stronger than others - the less sonorous the stronger the segment (cf. Foley, 1970, Kiparsky, 1979, Vennemann, 1972, among others). A strength feature is

assigned to all segments (with some language specific differences) and this feature interacts with the strength value of different syllable positions. Not only is syllable initial position stronger than syllable final position, but there is also a gradual weakening of position strength as we move towards the vowel and a gradual strengthening of position as we move rightwards away from the vowel. Strong → Weak → Strong. This is illustrated in Figure 10.

Figure 10

Hooper's Template

C C C C V C C C
m n p q r s t

where $m > n > p > q$

& $r < s < t$

& $m > t$

The universal facts with regard to distribution are to be explained on the basis of the interaction of position strength and segment strength; she claims stronger allophones will occur initially than finally - the aspirated allophone of 't' syllable initially versus the unaspirated allophone, syllable finally in English exemplifying this interaction. Since she argues that distributional facts are never the result of one segment interacting with another it's not clear how she can explain why 't' can be unreleased or glottal

after the relatively weak sonorants, as in 'dent' and 'belt' and 'Burt' but is unlikely to be, after the stronger obstruents in words such as 'last' and 'act' and 'apt'. These sequences are assigned to the same positions so such differences should not occur. Because she assigns no constituent structure to the syllable, constituency membership can not be used to explain why constraints are much tighter between C_m and C_n for instance, than they are between C_q and V. Her syllable makes the claim that there is no IC principle of phonotactics - only the interaction of segment and position strength.

Besides predicting clustering universals, Hooper wants her syllable to explain the segmental processes that often go under the labels of strengthening and weakening. Strengthening processes such as the aspiration of stops in English, and the $[y] \sim [j]$ ¹³ alternation in some Spanish dialects occur in initial position - the strong syllable position, whereas weakening processes like assimilation and loss occur to segments in syllable final position - the weak position. For instance Spanish /un gato/ $\rightarrow$ [un gato] and Latin /pulsum/ $\rightarrow$ Spanish /poso/. The segmental processes she deals with all implicate syllable boundary positions, and not whole onsets or rhymes or codas for instance, so there is no need for her syllable to include these constituents. Indeed Clements and Keyser (C&K), 1983, argue against the

onset and coda positions for this reason, believing that the notion of syllable boundary is all that is required in stating the structural conditions for these processes. (See below Section (viii)).

Because Hooper believes that syllabification is not distinctive, syllable structure is not lexical; UR's are syllabified at the beginning of a derivation and subsequently throughout. Syllable boundaries are inserted in accordance with what Selkirk refers to as the 'Maximal Onset Principle' (see the following section). These boundaries (\$) are inserted in the string of segments, hence Hooper's representations, like those of SPE are linear - the difference being that hers include syllable boundaries. She would assign boundaries to a word like 'iconoclast' as in (5).¹⁴

(5) \$ i \$ co \$ no \$ clast \$

The initial consonants are all occupying the C_m position; if there were a single consonant post-vocalically it would occupy the C_r position. She argues for this assignment because C_m and C_r are the strongest and weakest syllable positions respectively. This explains why a weak sonorant like [y] can become a fricative [j]; it is in a strong position in Spanish 'yates', for instance.

The next syllable we will consider reflects a similar focus on allophonic variation due to boundary conditions, but the structure proposed is different in that, although relatively flat, it is non-linear.

(vi) Kahn's Template (1976)

As previously discussed, it was in the domain of suprasegmental phonology that the SPE linear framework was found least adequate. In the area of stress this inadequacy had its reflex in metrical theory; in the domain of tone, autosegmental theory. Metrical phonology focusses on relations of constituency and binary prominence between contiguous elements, and these relations are formalized in binary branching, complex tree structures. The Selkirk syllable exemplifies this approach. Autosegmental theory, on the other hand, separates out features of sometimes non-contiguous elements (e.g. tone, nasality) and arrays these linearly on separate independent tiers, leaving no 'segments' in the SPE sense. These separate tiers are joined to one another by association lines and the resulting representations are n-ary branching simplex trees. In the next three syllable templates we will see two that are outgrowths of the autosegmental framework, Kahn's and Clements and Keyser's (1983), and one which amalgamates the notions of the two theories, Halle and Vergnaud's (1980).

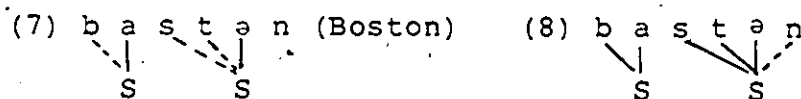
Kahn uses evidence from segmental processes in English to motivate the structure he proposes, and the structure he defines reflects the particular processes that concern him. His main interest is to account for the processes which, within the SPE framework, required the unnatural environment (C) $\begin{Bmatrix} C \\ \# \end{Bmatrix}$. (See Section 2 above.) That is, boundary phenomena are his main concern. His structure comes out of the autosegmental framework proposed in Goldsmith, 1976. He argues for a non-linear representation including two levels, the segmental and the syllable tiers. The syllable itself, however, is not assigned a hierarchical structure. There is no internal constituent structure. The representation of the two tiers is the standard simplex, n-ary branching tree of autosegmental theory, where there are no constraints requiring that single members of one tier be linked to single members on another tier. That is, the relationship between tiers may be many-to-one. It is with this formal device that Kahn captures the 'fuzzy-boundary' characteristics of English discussed in Section (2). He proposes, like Hooper, that syllabification takes place at the beginning of a derivation.¹⁵

Five rules account for his syllabifications.¹⁶ Rule 1 associates segments which are [+syll] with S in autosegmental fashion. An association line is drawn between the two as in (6).

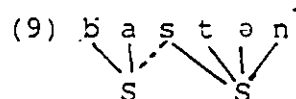
(6) b a s t ə n
 | |
 s s

The results of the application of Rule 2 to lexical items will be the same as the syllabification that results from Selkirk's 'Maximal Onset Principle' (for syllable boundaries). Our purpose here will not be to debate the relative merits of the different proposals regarding what considerations should determine the assignment of word medial C's to initial or final syllable positions. Kahn uses word initial cluster possibilities to determine the membership of medial C's and supports the syllabifications which result on the basis of terminal allophones. Selkirk (1978) maximizes onsets, seemingly constrained by the sonority hierarchy, like Kiparsky (1979), and C&F use markedness theory (cf. also Mohanan, 1979, for this latter approach). For our purposes, what is of interest is the results of the differently motivated algorithms - i.e. which syllables medial C's are assigned to. The results of initial syllabification for Selkirk, C&F, Hooper and Kahn's Rule 2 all appear to be the same.

Kahn's Rule 2 lists permissible clusters and first associates as many consonants as possible with the following S as in (7); then Rule 2 associates the left-over C's with the preceding S as in (8).



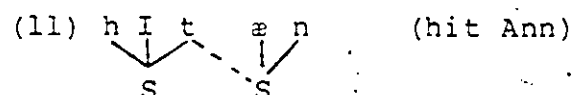
This is the syllabification for slow speech; for normal or connected speech three more rules apply creating ambisyllabic segments. It is with these three rules that his syllabification diverges from Selkirk's. Rule 3 joins syllable initial C's to preceding syllables if the syllable it is originally attached to is unstressed as in (9).



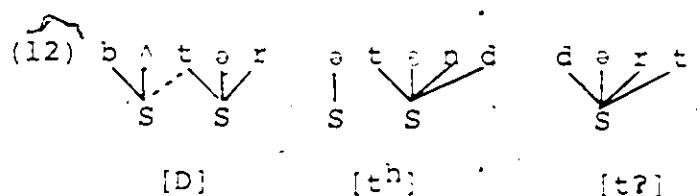
Rule 4 joins syllable final C's to following unstressed syllables,¹⁷ constrained this time by universal clustering constraints rather than the language specific constraints which apply at Rule 2. This means that Rule 4 would not apply to a form like 'bodkin' but it would apply to 'after' and 'hanger' as in (10).



The domain of Rules 1 - 4 is the word. Rule 5, his final rule, joins word final C's to vowel initial words and applies to connected speech as in (11).



With these syllabification rules, Kahn can account for a lot of allophonic variation - variation which he convincingly argues implicates all three types of syllable boundaries the rules produce. For instance the various allophones of 't' reflect the syllabic position of 't'. A 't' is flapped [D] iff it is ambisyllabic, aspirated [t^h] iff it is syllable initial and not syllable final, and glottal [tʔ] iff it is syllable final and not syllable initial. For instance the allophones in 'butter', 'attend' and 'dirt' respectively. See (12).



Kiparsky (1979) and Selkirk (1980) are both against the use of ambisyllabic segments. Kiparsky explains the flap allophone (so-called ambisyllabic phenomena in general) as a function of foot structure, i.e. flaps occur internally in a foot. Selkirk would resyllabify 'butter' assigning the 't' to the initial syllable. It has been suggested (Selkirk, 1980) that a three-way distinction is never necessary - i.e. a distinction among final, initial, and final/initial (ambisyllabic) positions. It would appear that Kahn's coverage of the data is more complete than either Selkirk's or Kiparsky's, however.¹⁸ In addition, as Kohler observes,

'fuzzy' boundaries are characteristic of English. For this reason, we will look at production errors to determine whether there is any evidence which suggests that the production mechanism distinguishes among these three positions. Since production errors are usually a phenomenon of connected speech such a distinction could well show up.

Kahn's syllable is a claim (like Hooper's) that processes will not need to distinguish internal boundary positions, or internal syllabic constituents from one another. The template we will next consider makes the claim that only one internal subsequence is necessary - the nucleus.

(vii) The Clements and Keyser Template (1983)

Clements and Keyser (C&K) adopt Kahn's syllabification rules and syllable structure except that they propose that a phonological representation should consist of two more tiers, a CV tier and a nucleus tier. These tiers, and the syllable tier, are considered to be structural units with no acoustic or articulatory correlates. (Although they do claim that units on the CV tier are primitive units of timing.) These three structural tiers and the segmental tier are all part of a lexical representation. Syllabification is done at the lexical level and continues throughout a derivation.

The CV tier has been used in accounts of segmental, suprasegmental and morphological phenomena. (See, for instance, Halle & Vergnaud, 1980, McCarthy, 1979, and Clements and Keyser, 1983). The CV skeleton is conceived of as a sequence of slots, i.e. structural units, labelled C and V. These labels are not variables which take on segments as values, as do the C's and V's in Hooper's syllable. Metaphorically speaking the CV skeleton is rather like the rings in a loose leaf binder. Each page that is attached to the CV rings represents a different tier, hence the CV tier mediates between the different tiers in a representation. In C&K's theory its function with respect to the syllable is two-fold.

C&K argue that Kahn's notion of the syllable is lacking in two ways: (i) syllabification requires the feature [syllabic] and (ii) the syllabic peak is not distinguished from the margins. They propose that by incorporating the independently motivated CV tier (They refer to McCarthy's (1981) use of the tier in word formation and give their own arguments for it in dealing with suffix allomorphy in Turkish, characterizing moras, etc..) as part of a phonological representation both these failings will disappear. The CV tier can subsume both functions. Peaks will be those segments dominated by V and a margin will be

any segment dominated only by C. Syllables are prelinked to V's (rather than the feature [syll]) in the CV tier and these V's may freely dominate any [-cons] segment.¹⁹ Since the tiers are considered separate planes the nucleus is not to be considered a subconstituent of the syllable despite its being deriveable only on the basis of the syllable. The nucleus is any tautosyllabic V(X) where X is either a single C or a single V. (When we refer to C and V in this and the next section we will be referring to elements of the CV tier. When we want to refer to consonants and vowels we will refer to them as such.) The V of the nucleus may not be preceded by any tautosyllabic V. (See footnote 19.) A lexical representation for 'loud' would look something like Figure 12.

Figure 12

Lexical Representation in Clements and Keyser



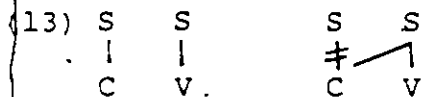
It's not clear to us what the empirical claim is in assigning the nucleus a non-constituent status. Would a nucleus that was a subconstituent of the syllable behave differently in phonological rules than this nucleus? Possibly, considering the nucleus a constituent would imply that any C's flanking

it on either side were also constituents. They explicitly reject this notion as will be clear presently in our discussion of (13).

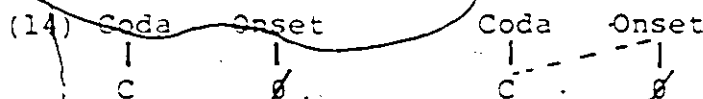
They are against the rhyme and for the nucleus, on the argument that stress systems respond to whether there is or is not a tautosyllabic segment following the peak and not whether there is one, two or three tautosyllabic segments.²¹ The feature make-up of the segment that the X is attached to, may be important, but the claim is that the number of C elements is not. They argue too that since collocational restrictions exist between onsets and peaks, not only between peaks and codas, i.e. within rhymes, phonotactic arguments can not be considered a diagnostic for constituency. Similarly, the fact that there is an upper limit to the length of the rhyme is no basis for assuming a rhyme constituent because this characteristic is not unique to rhymes. There is an upper limit to the length of onsets and codas too.

They argue against the onset and coda claiming that segmental processes and phonotactic constraints involving these constituents can be captured using opening and closing brackets to indicate syllable boundaries. Onsets and codas would be unnecessary complications. For instance they believe that the statement of the common resyllabification

rule across word boundaries, would only be complicated if onsets and codas were included in a representation. They state the rule as (13).²²



How they would express the rule using coda and onset is not clear, but a formulation such as (14) would not appear to be unduly complicated and seems to get the job done.



Clements and Keyser then, are claiming that phonological generalizations that involve the syllable will never have to make reference to the constituents onset, coda or rhyme. On the other hand they claim that the theory must recognize the nucleus - that recognition of this unit will allow generalizations otherwise difficult or impossible to formulate. It appears, however, that if stress rules or rules of vowel reduction for instance, refer to the nucleus, and these rules have to distinguish between long vowels, vowel glide sequences and vowel consonant sequences, the structural unit 'nucleus' by itself will be inadequate because it includes only the structural CV units and no feature information. They are claiming then, that for the

most part rules will only have to distinguish $\begin{array}{c} N \\ V \end{array}$ from $\begin{array}{c} N \\ V \quad C \end{array}$. That is, rules will usually only have to distinguish between branching and non-branching nuclei. This does not seem to be true of English stress (see Section 3. (iii)), nor of stress reduction in either Yupik (Reed et al., 1977, cited in Prince, 1983) or Estonian (Eek, 1975, cited in Prince, 1983).

With regard to production errors then, we will want to find some indication that there is a nucleus - any V(X) sequence (but see footnote 19), and that there is no other relevant unit, other than the syllable itself, including its boundaries.

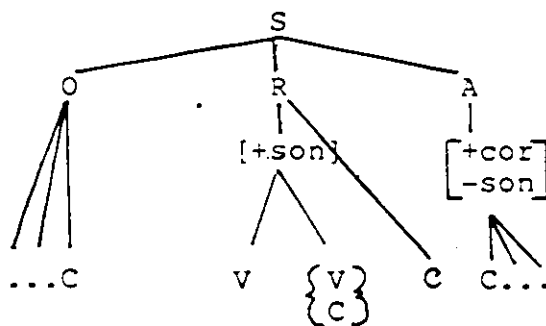
The next non-linear syllable we will look at incorporates notions from both metrical and autosegmental phonology, the hierarchical constituent structure of the former and the CV tier of the latter.

(viii) The Halle Vergnaud Template (1980)

Although Halle & Vergnaud (H&V) refer to their theory of syllable structure as an elaboration of the Liberman and Prince (1977) and Selkirk notions they take the notion of a CV skeleton from autosegmental theory. "The CV skeleton in all languages are divided into subsequences to which the term SYLLABLE has traditionally been attached" (H&V, 1980, p. 93).

Furthermore these subsequences possess internal constituent structure: an obligatory Rime constituent,^o which dominates at least one V in the skeleton, an Onset - an optional subsequence of C's on the left of the rime, and an optional Appendix - a subsequence of C's on the right of the rime, only occurring in word final position. The organization of the syllable is not expressed with Hooper-type boundary markers, but rather by means of trees anchored in the skeleton, but on a separate plane. In other words, their theory includes the separate tiers of autosegmental theory, but one of their tiers is itself non-linear (unless each level of their syllable is considered to be a separate tier). Their template for English is illustrated in Figure 13.

Figure 13
The Halle Vergnaud Template



N.B. (i) C and V are elements of the CV structural tier and not variables whose values are segments.

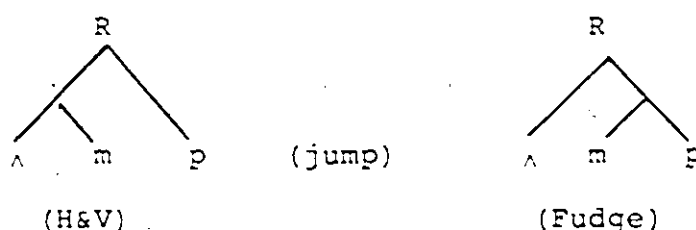
They support the structure they propose on the basis of segmental, suprasegmental processes and phonotactic constraints. For example, the onset is distinguished from the rime on the familiar argument that stress systems are never sensitive to onsets, only to rimes. (But cf. footnote 23.) It also accounts for the fact that stress systems often treat closed syllables as equivalent to long vowels. They also argue that an account of vowel lengthening in Luganda would be almost impossible without the rime constituent. To account for the fact that in English and German, for instance, the second position in the rime is limited to sonorants, (presumably analysing 'sk' as a single segment like Fudge and Selkirk) they give it a left branching structure with [+son] percolating down to both nodes. (Notice that by allowing all sonorants in the second position they will not be able to distinguish tense from lax vowels in terms of structure per se.)

Their arguments in support of the Appendix are convincing. It can be used to account for the differential treatment of stress in certain closed syllables - stress being sensitive only to rimes and not Appendices (in Malayalam, for instance). It can explain the different sequential constraints involving coronals and non-coronals in English. The segments which will be attached to Fudge's Termination node are those that will come under the Appendix

in the H&V syllable, and the two nodes will only occur in word final position. The two are couched in the notions of different theoretical frameworks but we will not attempt to choose between them. For our purposes, any facts from the production data which support the one will be equal support for the other. Although Fudge and H&V both include three rhyme positions the two templates differ regarding the grouping of these rhyme constituents as illustrated in Figure 14.

Figure 14

Rhymes in H&V and Fudge



We will see whether the production data can distinguish between these two claims regarding the constituent structure of the rhyme.

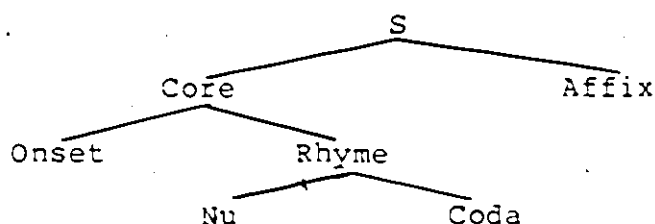
The final syllable we will consider is a reflex of phonetic concerns, and is particularly interesting for us because it has a 'surface level' motivation and is quite different hierarchically from the hierarchical structures we've been looking at.

(ix) The Fujimura & Lovins Template (1978)

Fujimura and Lovins (F&L) (1978) use phonetic arguments in support of a hierarchical syllable whose initial bifurcation is not between onset and rhyme, but between Core and Affix as indicated in Figure 15.

Figure 15

The Fujimura & Lovins Template



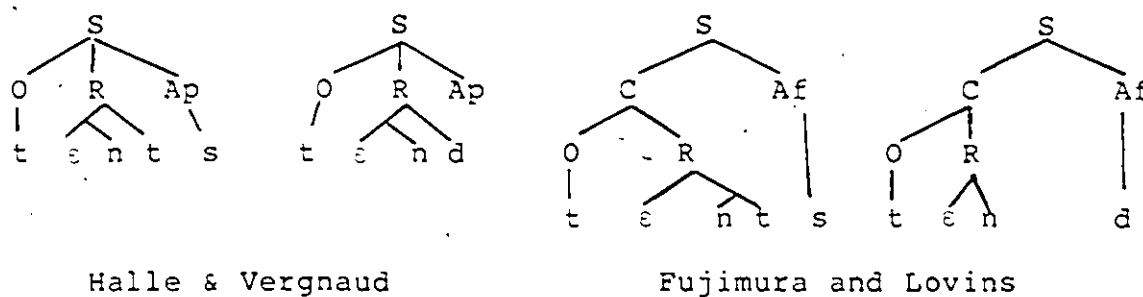
All elements cohesive to the nucleus are core elements and all stressed cores must have one final element. Affixes are stable phonetically and independent of the core; they always agree in voice with the final element in the core. The 'n' in 'tent' is very short because of the core final 't' whereas the 'd' in 'tend' does not affect the core. 't' is part of the core because it does not agree in voicing with 'n'. On the other hand 'd' does agree in the voice feature so it is an affix. The sonority hierarchy is claimed to hold within the core only, which explains why, in English, there are clusters [zd] and [dz] ('rads' and 'razzed') for instance,

which seem to defy the sonority hierarchy. They argue that phonetic realization rules will involve the concatenation of these syllables, and not segments. The core will include three subsets of features for the onset, nucleus and coda, temporally ordered in terms of the sonority hierarchy. Thus syllables will be bundles of features - a phonetic unit which constitutes an integrated event, and the core of these syllables will be the domain of phonotactic constraints.

Although the F&L Affix looks very like the H&V Appendix it is in fact quite different. The syllable structure that both would assign to 'tents' is similar but the structure assigned to 'tend' would be quite different as illustrated in Figure 16.

Figure 16

H&V and F&L Affix and Appendix Compared



N.B. (i) Only relevant nodes have been expanded.

(x) Conclusion

We have considered in detail a number of proposed syllable templates, and the different arguments used to support these various proposals. Our goal here was to highlight the similarities and differences in order to be able to attack the data of production errors within a richly elaborated, highly specific, and well motivated framework. We have seen a number of different structures, all of which seem to be well motivated to a greater or lesser extent.

To claim that syllables, or noun phrases, or any structures, are part of the production mechanism is meaningless without well defined notions as to what these units are. As we mentioned in our section on external evidence, numerous psycholinguists have argued in support of the syllable. But, as we have seen here, the syllable comes in many shapes and sizes. It is a kind of hydra. Our purpose in Chapter 2 will be to determine whether any of its 'heads' surfaces in production, and if so which.

To conclude this chapter we list the characteristics which will be the focus of our attention, and illustrate how each of the theorists we have looked at would syllabify the word 'victrolas'.²⁴

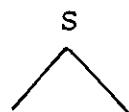
(xi) Characteristics Which Shall be Our Focus in Chapter 2

1. internal constituent structure versus flatter structure in all nodes:

- (a) n-ary, binary or tri-nary branching S node



Kahn



Selkirk

F&L, C&F



Fudge

H&V

- (b) onsets



Selkirk



Fudge



C&F



H&V

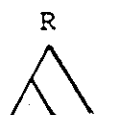
- (c) rhymes



Selkirk



Fudge



H&V

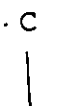
- (d) codas



Fudge



C&F



H&V

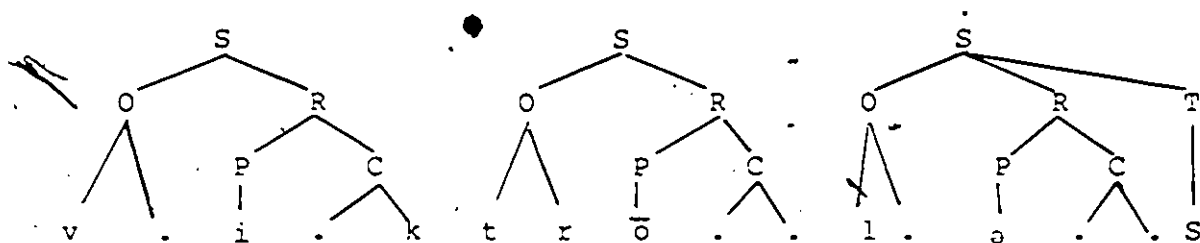
(e) peaks

P
|
Fudge

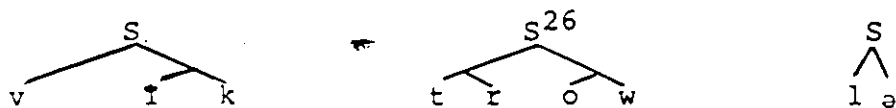
P
/ \
Selkirk
C&F

(f) Appendix (termination) versus Affix

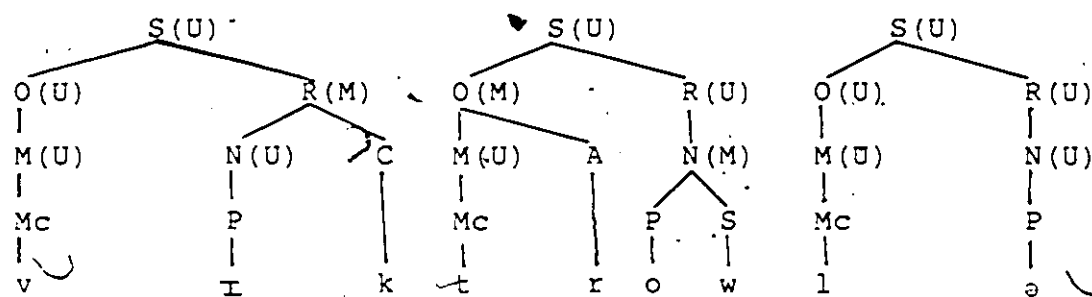
2. labelled versus unlabelled nodes (Selkirk vs. C&F)
3. lexical versus non-lexical syllabification. (Selkirk vs. Hooper)
4. domain of syllabification (word or stem or across word boundary)
5. uni-syllabic versus ambisyllabic C's (Selkirk vs. Kahn)
6. no internal constituent structure (Kahn, Hooper) or almost none (C&K)
7. feature composition of node segments, for instance does peak or rhyme dominate [ay] in an open syllable?

(xii) 'victrolas' Syllabified(a) Fudge²⁵

(b) Selkirk



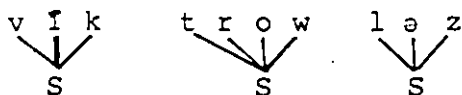
(c) C&F



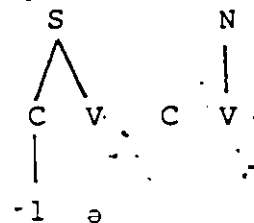
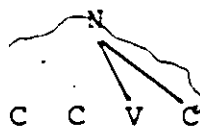
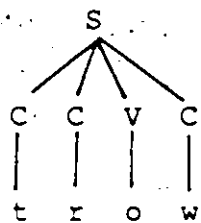
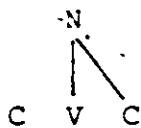
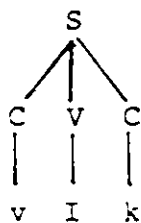
(d) Hooper

§14
~~§~~ v i k ~~§~~ t r o w l ə z ~~§~~

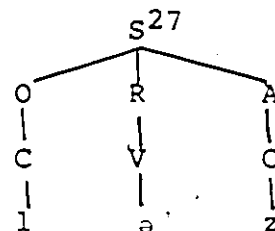
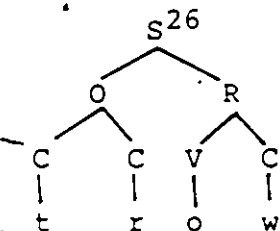
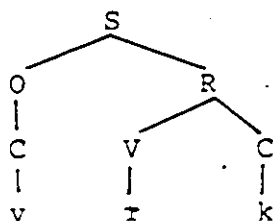
(e) Kahn



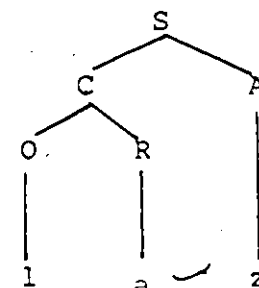
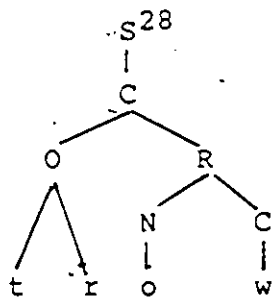
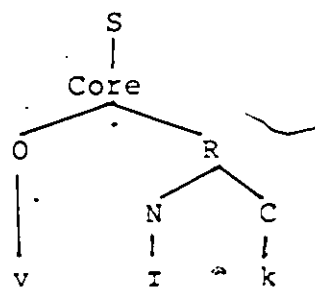
(f) C&K



(g) H&V



(h) F&L



APPENDIX TO CHAPTER 1

Kahn's rules:

Rule 1

$$\begin{array}{c} [+syl] \rightarrow [+syl] \\ | \\ S \end{array}$$
Rule 2

$$\begin{array}{c} a. C_1 \dots C_n \quad V \rightarrow C_1 \dots C_i C_{i+1} \dots C_n V \\ | \qquad \qquad \qquad \searrow \swarrow \nwarrow \nearrow \\ S \qquad \qquad \qquad S \end{array}$$

where $C_{i+1} \dots C_n$ is a member of the set of permissible initial clusters but $C_i C_{i+1} \dots C_n$ is not.

$$\begin{array}{c} b. V \quad C_1 \dots C_n \rightarrow V \quad C_1 \dots C_i C_{i+1} \dots C_n \\ | \qquad \quad | \qquad \quad | \qquad \quad \searrow \swarrow \nwarrow \nearrow \\ S \qquad \quad x \dots x \qquad \quad S \end{array}$$

where $C_1 \dots C_i$ is a member of the set of permissible final clusters but $C_1 \dots C_i C_{i+1}$ is not.

Rule 3

$$\begin{array}{c} \text{In } [-cons] \quad C \quad C_0 \quad \left[\begin{array}{c} V \\ -stress \end{array} \right] \text{ associate } C \text{ and } S_1 \\ | \qquad \quad \searrow \swarrow \nwarrow \nearrow \\ S_1 \qquad \qquad S_2 \end{array}$$
Rule 4

$$\begin{array}{c} \text{In } C \quad C_0 \quad \left[\begin{array}{c} V \\ -stress \end{array} \right] \text{ associate } C \text{ and } S_2 \\ | \qquad \quad \searrow \swarrow \nwarrow \nearrow \\ S_1 \qquad \qquad S_2 \end{array}$$

Rule 5 In $C \quad V$ associate C and S

$$| \\ S$$

NOTES TO CHAPTER 1

¹See Dretske, 1974, for further discussion of this distinction.

²The different allophones of 'l' here, do indicate that 'foolish' and 'coolish' are syllabified differently at the surface level. This can be explained within Garrett's production model. See Chapter 3. If 'foolish' is stored as a major class lexical item and not as the output of a morphological rule it would be syllabified as foo.lish and inserted as such into his phrasal frame. If 'coolish', on the other hand, is the output of a rule it would be inserted as 'cool' and 'ish' would be attached later. This separation of stem and grammatical (speaking loosely) morpheme explains errors like 'The weatherman calls for cloudy skies → ...clouds for cally skies'. In Chapter 2 we will show that there is evidence from speech errors that such root/stem final C's are syllabified as codas and not syllable initial onsets. See Chapter 2, Section 6.

³They distinguish between rules and processes - the former apply to segments. For instance a rule, and not a process, assimilates the 'n' in the prefix 'in' to following C's. We have 'irreplaceable' but 'unrelenting'.

⁴He doesn't talk about restrictions within the T node because he gives it only one position (which may hold up to three segments) and constraints only hold between positions.

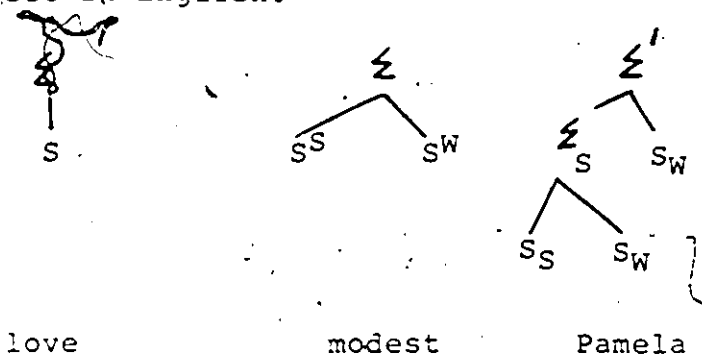
⁵Except that there is no 'h' in position 5 and 'ŋ' is derived from 'nC' so it doesn't occur in position 1. Prevocalic 'y' is derived from $\bar{a}$ and occurs in the peak. Post-vocalic 'w' and 'y' are also in the peak.

⁶She gives examples of segmental rules which require the unit 'syllable' and even suggests that her prosodic units may be the units of production and perception.

⁷These include the syllable, foot, phonological phrase, prosodic word, intonational phrase and utterance. They each have their own principles of prominence, construction and syntactic domain. For instance foot and syllable structure are both assigned to the simple non-branching stem, and relations of prominence are the binary s/w relation. Their principles of construction are unique to themselves (see footnote 8).

⁸Feet may be composed of one, two or three syllables only one of which is stressed. For instance 'love' 'modest' and

'Pamela' all consist of one foot, exhausting the class of possible feet in English.



⁹This means that $C\bar{V}$ in an open syllable is a branching rhyme whereas it is a branching peak in a $C\bar{V}C$ syllable. Intuitively, not a nice result. See discussion in text following.

¹⁰It's not clear how this type of argument is to be understood. Phonetic features seem to have 'no independent privileges of distribution'; they seem to be the defining characteristics of segments or larger phonological units and yet in some sense not 'independent' in the sense above. Yet they must be distinguished from one another.

¹¹The rhyme constituent has long been used to characterize the heavy/light syllable distinction. Kurylowicz, 1948, Pike, 1967, and more recently, Kiparsky, 1979, and Halle & Vergnaud, 1980.

¹²As indicated, a bipositional peak appears to be necessary to handle the prosodic facts and she also gives collocational arguments for distinguishing peaks from rhymes.

¹³[j] designates a voiced palatal fricative. We are not claiming that syllable initial position is always a position of strengthening. Word internal intervocalic consonants are often weakened despite their being syllable initial. E.g. - Spanish lago, [layo], and English water, [waDr].

¹⁴Post 1976, Hooper allows ambisyllabic segments, so the 'n' would be η to show that it was ambisyllabic.

¹⁵The stress feature is required in the syllabification rules for normal speech (i.e. rules 3 - 4) so it might appear that Kahn believes stress to be lexical in English since syllabification begins at the beginning of a derivation.

"...rules of syllable structure assignment (are A.L.) ordered prior to the phonological rules proper" (p. 9) and the rules apply in a "block" (p. 131). It is not absolutely clear what Kahn's exact position is, however, because "...the stress assignment rule must be ordered between syllable-structure assignment rules 11 and 111" (p. 87).

¹⁶See Appendix to Chapter 1 for the formally specified rules.

¹⁷This rule is intuitively unsatisfying, in that it moves C's away from stress towards non-stressed syllables, and doesn't allow C's in unstressed syllables to move towards stressed syllables. Kahn supports this, claiming that in the two forms 'after' and 'Haftonium' there is a clear boundary only between the first two syllables in 'Haftonium'. Firstly, our intuitions don't agree with his here, and secondly it's not clear that the presumed difference should be considered a function of the 'stressedness' (i.e. the plus stress value) of the syllable 'ton'. This goes against the often noted fact that stress attracts. This rule would assign the 'n' in 'entry' to the following syllable, but not the 'n' in 'interior'. In my judgement the opposite assignment seems right. This could perhaps be decided by observations regarding progressive nasalization in the two forms.

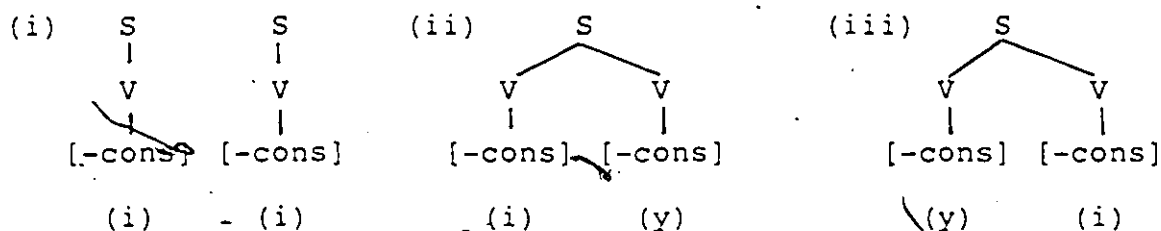
¹⁸The 't' in 'hit Ann' can be [t?] or [D], but never [t^h]. We have a 't' here which can be re-syllabified with the following word - the only resyllabification rule which Kiparsky accepts. If the resulting structure is



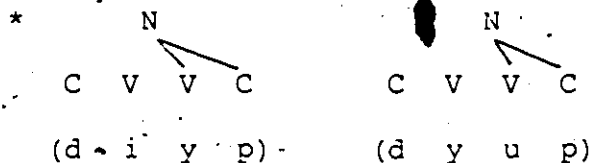
we would expect the [t^h] allophone to be possible. The flap variant here can not be a matter of foot structure because we have two adjacent feet. It could be argued that 'hit' is defooted here but this won't go through because the same facts are true of 'hate Ian' and diphthongs presumably can't ever lose their foot status. Selkirk would claim that the flap variant is due to the 't' being in final position in the syllable. But this won't explain why the flap variant is unavailable to the first 't' in 'Ottenbrite'. There are two variants possible here; in very slow citation-type speech the [t^h] allophone, and in normal speech the [t?] variant. In other words, it appears that this 't' can be syllable initial, or syllable final but not both initial and final as can the 't' in 'water' or 'hate Ian'. This raises interesting questions about the underlying status of the 'n' in such words which we will not attempt to solve here. But see

Chapter 2 Section 11 for some discussion of syllabic liquids and nasals. Notice, for instance that the [D] allophone is possible in 'bought in (Canada)'. Kahn's rules would produce the wrong results if the second syllable in 'Ottenbrite' were assumed to include the sequence [ən]. Rule 2 would attach the 't' to the second syllable and Rule 3 would link it to the first syllable, predicting the [D] variant, incorrectly. Nevertheless, the relevant point here is that Kahn appears to be right in claiming that a three-way distinction is necessary. To account for some allophonic variants we must know whether a segment is syllable initial, syllable final, or both.

¹⁹It's not clear to me how this is to be understood. In their theory, glides are not distinguished at the segmental level from their syllabic counterparts. If V's may be freely linked to [-cons] how would a representation distinguish (i) - (iii)?



That is, how does the syllable tier know that the VV in (i) belong to two syllables, whereas in (ii) and (iii) the VV belongs to one syllable? Similarly, how does one distinguish the [yi] from the [iy]? There is no information in the representation to distinguish them. It would appear that the [-syll] value of [y] must be accounted for by pre-linking it to a C. This also leads to problems with their definition of the nucleus. It appears that the V₂ C in 'deep' - C V V₂ C would be correctly ruled out as the nucleus, but the V₂ C in 'dupe' - C V V₂ C - would also be ruled out, incorrectly in this case.



Perhaps we have misunderstood their arguments.

²⁰It's not clear why the nucleus is not simply rather than



²¹The actual number of segments would appear to be important in word final position, at least to account for the fact that quantity sensitive systems will sometimes ignore VC word finally but will not ignore VCC. This may not be a problem for them since they recognize extrasyllabic C's.

²²See footnote 18 for arguments against this breaking of the association line between S and C.

²³C&K, 1983, point to Roca's (1982) analysis of stress in which stress assignment must refer to the onset. (See footnote 16, p. 52).

²⁴We are making no claims regarding the underlying phonemic shape but assume that all would assign three syllables to the word.

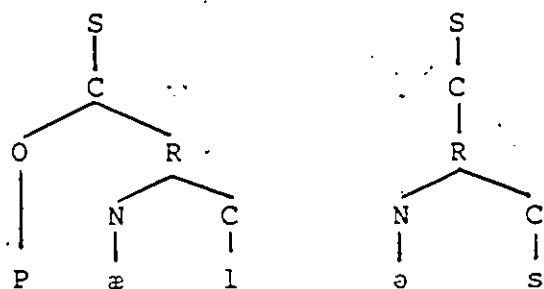
²⁵It's not clear whether Fudge would generate an empty internal T node or not.

²⁶H&V (and Selkirk) refer to structures like this as branching rhymes, so presumably this is the minimal structure that they would assign, despite their [+son] position on the rhyme's left branch.

²⁷It's not clear to us when H&V would attach grammatical morphemes but presumably not in lexical structure. We include the [z] because it points out the function of the Appendix.

²⁸F&L insist that a stressed syllable in English must "be closed by a non-nuclear final element..." (p. 114). If the vowel were not a diphthong in the second syllable the "l" would be assigned to the coda. For example 'palace' would be as illustrated in (i).

(i)



CHAPTER 2

THE PRODUCTION SYLLABLE: THE ARGUMENT FROM SPEECH ERRORS

1. INTRODUCTION

Our intention, in this Chapter, is to determine whether or not the syllable is required as a construct in a model of production, and if so, to determine the type of structure this production syllable has; and, how closely it relates to the templates proposed by competence theorists as described in the preceding chapter. The basis for our conclusions will be spontaneous errors in normal populations.

In Section 2 we discuss various problems that are associated with the use of such data as the basis for inference. We point out that some of the problems are not unique to speech error theorists; indeed they reflect the more general problem of the underdetermination of theory by evidence. (See Section 2.1.) We argue that other problems relating specifically to speech error corpora and their collection, are not intractable. (See Section 2.2.) We conclude that speech errors provide a sound basis for reaching conclusions with respect to the units and processes of the production mechanisms.

In Section 3 we outline the assumptions which underlie our investigation. In particular, we assume that speech errors reflect the processing mechanism responsible for error-free speech and secondly that the best working hypothesis is Garrett's assumption that processing will "...reflect the grammatical decomposition of the language faculty" (1982, p. 27). From this second assumption we show how it follows that we must therefore assume that elements involved in errors must be analysable as constituents of some type, and if two elements interact in an error they must be constituents of the same type. We point out, however, that it does not follow that all processing types must be subject to errors of selection. That is, we make no assumptions regarding a criterial form of evidence for processing types.

The errors which form the data of our investigation are described in Section 4. Since syllables themselves are not subject to mis-selection, arguments in support of the syllable's viability as a processing type will have to be based on its explanatory power with respect to other data. Various theorists have observed that segmental exchanges seem to respect syllabic position, although none have substantiated this observation with a precise definition of syllabic position. If this claim should turn out to be correct when given some empirical content, we would have a good argument for maintaining that the syllable is a

processing unit - it would provide the basis for explaining this constraint. On our assumptions, since submorphemic segmental errors all appear to involve the same processing type, namely, segments, this constraint, if viable, should apply to all types of segmental errors; not only exchanges. The only errors which can provide a basis for testing this claim are those in which two string positions may be compared, however. Thus we are limited to errors which have been labelled in the literature as segmental exchanges, anticipations, perseverations, movements and additions and word blends. We explain on what basis these are the labels which have been used in the literature - i.e. what has motivated making a distinction among these error types. We indicate how all these types are relevant for our investigation and end the section with examples of the types which form our data base ((18-22)).

In Section 5 we give empirical content to the claim mentioned above regarding a 'syllabic-position-constraint' in segmental errors. We specify a syllabification algorithm (Figure 1) for determining which syllabic positions segments come from. We distinguish three positions, onset, peak and coda using the word as the domain of syllabification. With this algorithm we may identify precisely over 25 possible syllabic position interactions including those which involve identical positions, i.e. onset/onset, coda/coda and peak/peak. In this section we also point to the ambiguous

status of segments in word-initial and word-final positions. In order to test the hypothesis regarding a syllabic position constraint, errors which involve at least one word medial consonant are required. Otherwise conclusions regarding syllabic structure could be an artefact of lexical structure. Having specified only three positions we can not distinguish errors which involve single segments from consonant clusters in onset or coda position, from those which involve entire clusters. Thus our first test of the hypothesis will use both the errors (1) and (2) as support for the hypothesis, although in (1) only part of an onset is involved and in (2) the whole onset.

(1) strive for perfection.

sprive . . .

(2) central choir

[kwəntrel] sire

(Phonetic spellings will be given only when the conventional orthography would be unclear.)

Section 5, therefore, provides us with a means of empirically testing the 'syllabic-position-constraint' hypothesis. This we proceed to do in Section 6. We find that the hypothesis is supported: Eighty-eight percent (88%) of segmental errors involve the interaction of segments from onset positions, peak positions and coda positions. Three percent (3%) are analysable as errors involving identical sequences of positions (peak-coda, and onset-peak). (This

constraint on syllabic position was also evident in the errors involving at least one word-medial consonant.) We specify an algorithm for generating these errors (see Figure 3). This algorithm, Possible Segmental Error (PSE), reflects our hypothesis that this constraint operates on segmental errors because at the processing level at which these errors occur, the processor is selecting and mis-selecting syllabic constituents, not unstructured segments. We give examples of PSE at Work (Figure 4) but note that it is only a first approximation. It can only account for errors which involve entire identical constituents, that is, 90% of the errors.

Four percent of the errors that are problematical for our hypothesis are errors which involve 'V+r' sequences. We demonstrate that the production mechanism treats 'r' as a glide and these errors are really peak/peak errors.

After our discussion of the peak constituent we break briefly for an excursus. The syllabification of intervocalic glides will always assign the glide as onset to the second vowel syllable. Although we had no errors whose analysis was jeopardized by this syllabification it is not clear that all intervocalic glides should be syllabified as onsets (see Figure 7).

The final set of problematical errors were those (4%, n=25) that involved the interaction of segments from onsets with those from codas. We find that at least ten of these could be analysed as coda/coda interactions if the 'root' (or stem) is taken as the domain of syllabification. We illustrate that this domain holds up when we re-analyse the rest of our errors using this domain rather than the word (see Figure 11). Seven of the onset/coda interactions are within-word errors. We suggest that a number of within-word errors occur at a different locus in the production process. The pattern they exhibit is different enough from between-word errors to warrant this conclusion. (See Section 12.) By factoring out the within-word errors we find that 95% of errors involve interactions between the single constituents onset, peak and coda.

As the reader will have noticed, the rhyme constituent, so strongly supported by competence theorists, has not played a part in our discussion. Is there a production rhyme constituent? Section 7 we devote to answering this question. We argue that there is not - only 11 errors involve a peak-coda sequence and only 7 of these include the entire coda. That is there are only 7 apparent rhyme errors - although with the root as domain there are 10. There are almost as many onset-peak errors - (n=8). We can account quite simply for both types of error as complex errors

involving morpheme exchanges and onset exchanges, both common error types.

In Section 8 we consider the viability of ambisyllabic segments as defined by Kahn's normal Speech Rules III-V. See Appendix to Chapter 1. We find we can predict precisely when an ambisyllabic segment will function like a coda and when it will function like an onset. In other words at the level of processing where between-word segmental errors occur there is no evidence of ambisyllabicity.

Section 9 is devoted to the errors which unambiguously involved only parts of onsets. We were interested in determining whether such errors supported a hierarchical sub-constituent structure within the onset. We also wanted to determine whether we could generate these part errors. We found that the type of error which entire onsets participated in was duplicated by the 'part' errors. They were involved in movements, exchanges, anticipations, and perseverations and the errors always involved the interaction of similar types, sonorants, obstruents or the sibilants 's' and [ʃ]. We found no evidence of internal hierarchical structure but the similar pattern exhibited by the part errors and syllabic constituent errors suggested that they be given an analagous account. We proposed that the onset includes three linearly ordered constituents (like the syllable) and specified a rule

for generating these errors - identical to the rule for generating the syllabic constituent errors; only the constituent labels were changed (see Figure 15). This analysis, besides falling in line with the assumption that mis-selection errors will always involve constituents of the same type, allows us to account for other characteristics of segmental errors - characteristics which a segmental analysis (cf. Fromkin, for instance) can not account for.

Section 10 we devote to codas. Unfortunately we can only speculate about the internal structure of these constituents. There were only 12 errors which unambiguously involved parts of tautomorphemic coda clusters. There were, however, 11 errors which involved clusters which included a morpheme boundary and these were informative. Such clusters never moved as units, supporting our claim that roots should be the domain of syllabification. Indeed the segments preceding the morpheme boundary were the ones to move. This independence of coronal morphemes from the preceding consonant suggested that the Halle and Vergnaud Appendix might be a viable syllabic constituent of the production template. We argue against this conclusion however; the main argument against it being that it does not behave like the other three constituents. When it dominates a morpheme it is not involved in exchanges, anticipations or perseverations and coronal segments which do not follow a morpheme boundary

are not independent of their root, unlike the coronals which are concomitantly morphemes. On the basis of the twelve part errors involving tautomorphemic clusters we speculated that the coda, like the onset might consist of three linearly ordered positions, and coda constituent errors might involve an account analogous to the account of onset sub-constituent errors.

In Section 11 we consider whether word blends can support or disconfirm the constituent structure we have assigned to syllables on the basis of segmental errors. We find that all blends which involve the simple juxtaposition of two words can be explained as the mis-selection of identical syllabic constituents, those very constituents which feature in segmental errors.

Before summing up in our final section we illustrate (in Section 12) how a number of within-word errors do not exhibit the ~~characteristics~~ of between-word errors. If within-word segmental errors are syllabic constituent errors like between-word errors their domain of syllabification must be the word and not the root. And even with this modification a number of errors (17%, n=7) can not be construed as errors involving the mis-selection of identical syllabic constituents.

2. Errors as Evidence: Problems and Pseudo-Problems

(i) Falsifiability: A Pseudo-Problem

In this section we will consider why the claim that the methods and data of speech error analysis are fundamentally unscientific is unfounded. Besides the well-known problems relating to the collection of speech errors (which will be reviewed in the next section), it has been claimed (see Kupin, 1982 and references cited therein) that there is a more fundamental problem in the use of speech errors as evidence - one which concerns the falsifiability of hypotheses regarding the production process derived on the basis of such errors. It has been suggested that these hypotheses are not falsifiable, that some evidence will always be interpretable as support for any type of unit proposed and no evidence will be able to falsify the type proposed. In other words, the hypotheses are mere speculation, having moved outside the bounds of science and into the realm of fantasy.

What the reasons are, for this accusation, are not clear to us, but what seems to be at issue here is what Chomsky has discussed under the labels of "the indeterminacy thesis" and the related "bifurcation thesis" (Rules and Representations, p. 15 ff., Reflections on Language, p. 179 ff.). By the former Chomsky means "...the thesis that whatever evidence we

may accumulate in support of some hypothesis, there always exist alternative hypotheses inconsistent with ours but compatible with the evidence" (Rules and Representations, p. 258). As Chomsky points out this is true, but "...the conclusion that there is, therefore, "no fact of matter (Quine's conclusion, (A:L.)) is unwarranted, particularly if one is unwilling to draw the same conclusion for physics on the same grounds" (ibid.). "no argument at all has been presented" (Rules and Representations, p. 16) for the bifurcation thesis - the thesis that "theories of...psychology are faced with a problem of indeterminacy that is qualitatively different in some way from the underdetermination of theory by evidence in the natural sciences" (ibid.). This qualitative distinction is based on the notion that the quality of the evidence used by linguists is not of the same standard as that used by physicists; and this difference in quality defines a boundary between the natural sciences and linguistics. Chomsky argues against this with an analogy from physics. Suppose astronomers have no direct access to the thermo-nuclear reactions of the sun's interior, (comparable to the brain). Analysis of the sun's internal behaviour will have to be made on the basis of the sun's periphery - analogous to speech behaviour, including speaker intuitions. The only evidence available in both cases is indirect, and it is not conclusive in either case,

but "no empirical evidence can be conclusive" (Rules and Representations, p. 190); we are dealing with science, not mathematics. Evidence can be more or less direct, and "needless to say, the evidence which supports the linguist's constructions is incomparably less satisfying than that available to the physicist. But in essence the problems are the same..." (op cit., pp. 191-192). Theories in both domains are underdetermined by evidence.

Problems of indeterminacy are particularly apparent regarding speech production hypotheses because most of the arguments will of necessity be theory internal, given the current 'state of the art'. For instance Fromkin would analyse the error in (3) as the perseveration of the feature [-cor]¹ whereas it would be analysed by Nootboom as the substitution of 'b' for 'd'.

(3) flipping his lid	lib
[-cor] [+cor]	[-cor]

Indeed, the evidence in support of one or the other analysis is "incomparably less satisfying" than one would like, given that "...it seems impossible to find independent segments in the time domain that are units of phoneme size..." (Fujimura and Lovins, 1982, p. 187) let alone units of feature size. And even if we could acoustically or articulatorily identify such units in the speech chain, they would only be indirect

evidence with regard to the neurological processes involved in the error. The most direct evaluation of the two hypotheses would involve the examination of these processes - such evidence is obviously not available at the moment. But such a contingency is not unheard of in the hard sciences either. A number of years ago physicists predicted the existence of certain elementary building blocks of protons and neutrons, at the time having no way of directly testing the prediction. Only recently, with the development of powerful particle accelerators could the predictions be tested. Indirect evidence for the quark's existence has been supplemented by the more direct evidence provided by the accelerators.² Unfortunately production theorists are still a long way from having a device like the particle accelerator to test their predictions directly. They will have to be satisfied with less direct evidence for the time being. The conflicting analyses will have to be tested, evaluated from a theory-internal perspective,³ or by considering which analysis has independent support in theories from other domains such as language acquisition, perception, and so forth. The Fromkin and Nooteboom hypotheses are, in this case, conflicting hypotheses based on the same data but such conflicting analyses are not confined to the 'soft' sciences. Light has been analysed both as particle and as wave: contradictory analyses which are both supported by the

evidence. The two analyses reflect different theoretical perspectives; the wave analysis has not been disconfirmed but the particle analysis has more explanatory power within modern-day particle physics. The 'feature' versus 'segment' analysis will have to be evaluated on similar grounds.

Theory-internal arguments regarding the speech processor are particularly problematical because they are often highly dependent on one or another theories of competence, and will be, to some extent, as good or as bad as the competence theory which is their source. As discussed in the Introduction (p. 1 ff.), the highly articulated hypotheses of competence theorists provide a rich source for production theorists. No modeler of the production process begins without at least some more or less abstract notions from grammar. Reference is always made to such century-old, familiar notions as word, noun, segment, etc. But some theorists explicitly embrace constructs more specifically related to one particular theory. Fay (1980) for instance, suggests that certain errors involving particles may be the result of a failure to apply the deletion part of a movement transformation. An example of such an error is (4) where Fay suggests that the particle is copied in sentence-final position but post-verbal deletion has been neglected.

(4) ~~Do~~ I have to put on my seat belt on?

This analysis assumes (i) the viability of transformations in general, (ii) the existence of a particular transformation (Particle Movement) and (iii) that transformations involve copying and deletion - not just movement. If any one of these assumptions turns out to be wrong, his analysis founders. This hypothesis regarding the production process - that it involves transformations - will be supported or disconfirmed to a large extent on the basis of evidence relevant for competence theories, not evidence from production at all. At the same time, however, his analysis provides support for the transformational component of the competence theory. Obviously the inter-relationships and interdependencies are very complex but there would not appear to be anything fundamentally at variance with normal scientific analysis here. To return to Chomsky's analogy and extend it, predictions regarding the interior thermo-nuclear reactions are made on the basis of the sun's periphery, but there may be more than one theory accounting for behaviour at the periphery. The hypotheses regarding the interior will reflect one of these and will founder if the espoused theory fails.

On the other hand, processing theories include assumptions, methods of analysis, and data which are unique to them, and therefore their explanatory power is not only a reflection of the competence theory they embrace. Speech

error theorists working out of the same competence framework may arrive at quite distinct conclusions on the basis of the same data. Garrett, for instance, like many others, is working out of the T.G. framework and yet his analysis of segmental errors, is, so far as I know, unique to his own theory of the production process. He classifies segmental exchanges and morpheme exchanges (what he calls 'stranding errors') as errors of one type and classifies segmental anticipations and perseverations as errors of a different type. Klatt (1981) and Fromkin (1971), working within the same theoretical framework, don't differentiate among these types. And indeed, we know of no other theorist who does so.⁴ In this case, choosing between Garrett and Klatt or Fromkin will, for the most part, involve a direct examination of the empirical coverage of the production models proposed, the competence theory which provides the framework for the models will be less crucially involved than was the case in Fay's transformational analysis of particle errors. If two theories of production are found to have equal empirical scope then choosing between them will involve recourse to "...concepts of simplicity, insight and explanatory power that are not at all understood..." (Rules and Representations, p. 22). Such methods of evaluation are clearly not unique to the 'soft' sciences.

In this section we have attempted to indicate that the problems regarding falsifiability alleged to characterize theories of production based on speech error data are 'pseudo problems'. At least they are problems faced by any scientist in any domain and they are 'pseudo-problems' if they are held to be specific and unique to production theorists. The methods of production theorists are well within the normal procedures of scientific analysis. Questions regarding indeterminacy, questions which concern the very foundations of science should be raised "...where the hope of gaining illumination is the highest" (Rules and Representations, p. 22) i.e. in the hard sciences, not psychology - "nothing much comes of the discussion..." (ibid.) if such questions cannot even be answered in the case of physics.

We will now turn to problems which are not 'pseudo-problems' in this sense - problems which relate directly to the data of speech error theorists - the spontaneous errors of normal populations.

(ii) The Reliability of the Data: A Problem

In the Introduction (p. 1 ff.) we alluded to some of the often discussed problems associated with the use of speech errors as evidence for speech production models. (See the discussion in Cutler, 1982, in particular, and also Fromkin,

1971, among others). Firstly, it is not clear to what extent error collections reflect characteristics of the perception mechanism rather than the production mechanism. This is due both to the possible selective attention of the different collectors, and to the differential degree of detectability of various error types. The first of these two problems is probably relatively unimportant. As Garrett (1980) points out, the principle distributional regularities of the many different error collections are robust. That is, different collectors appear to 'select' the same data to attend to, which means in effect that the selective attention of individual collectors is not a problem. There would appear to be no effect for selective attention in error distributions. The second problem, that of differential degrees of detectability, is indeed a problem - a crucial and very complex one. For instance, with regard to segmental errors, experimental work indicates that there are at least three variables which are the source of this differential, and there may be interaction effects among these three:

- (i) word position (initial, medial and final)
- (ii) stress (plus or minus)
- and (iii) type of error (anticipation, perseveration and exchange)

Studies on word recognition (Taft and Forster, 1976, Marlsen-Wilson and Welsh, 1978), word recall (Horowitz et al., 1968, 1969) and 'tip of the tongue' states (Brown and McNeill, 1966) indicate that word initial segments carry the largest functional load. And also, hearing errors are most common in word medial positions (Browman, 1980). This makes conclusions based on the relatively high frequency of errors in word-initial positions difficult, because they may be an artefact of perception, not production. Similarly, errors in unstressed syllables are more likely to go unnoticed (Cutler, 1981). Added to this is the fact that if an error involved the vowels of two unstressed syllables they would be impossible to detect since they would involve phonetically identical segments. There is no reason to assume a priori that such errors do not occur since there are consonant errors in unstressed syllables as exemplified in (5).

(5) a Tanadian from Toronto Canadian

Therefore, reaching conclusions based on the preponderance of stressed syllable errors is not altogether safe. Interesting in this regard is the fact that Kupin, 1982, for instance, finds that more errors occur in unstressed than stressed syllables in his experiments with tongue-twisters. Finally, it seems that perseverations are less likely to be noticed than anticipations (Tent and Clark, 1980, and Cohen, 1980).

We find that error distributions show an effect for these three variables. Relatively high proportions of errors in word initial position, and stressed syllables, characterize all speech error collections (that we know of) that have been studied in any detail. These collections also include more anticipatory errors than perseverations. Do these frequencies reflect the perception mechanisms of collectors or the productive mechanisms of speakers? Probably both, but clearly inferences based on these frequency distributions can only be tentative.

There is another problem regarding error-rate based conclusions. Nothing can be deduced given a number of errors involving two particular units for instance, unless the number correct is also given for each of them (see van den Broecke and Goldstein, 1980). Relevant here are both the distributional patterns of the language, and the overall context in which the error occurred. As an example, suppose the initial consonants 'p' and 'y' are seldom involved in an exchange error whereas initial 'p' and 't' interchange relatively frequently. If there are 50 times as many occurrences of 't' as there are of 'y' we would expect there to be fifty times as many errors with 't' as there are with 'y'. Thus a language's distributional patterns must be taken into account, or better still, the context of the error. Unfortunately it is likely that spontaneous error corpuses do

not include 'enough' context to measure the proportion of right to wrong.⁵ Whatever 'enough' context might turn out to be.

Despite these problems the error distributions which one would predict on the basis of differential detectability would not in general reflect the distributions we find in error corpuses. Since certain error types that are relatively difficult to detect turn up more frequently in corpuses than errors more easily detected "...the most frequently reported sound substitutions are those between sounds which are most like one another" (Cutler, p. 572, 1982) and yet these are exactly the errors that go undetected in shadowing tasks. For instance, simple consonant errors which differ in only one feature (Marlsen-Wilson and Welsh, 1978), and consonant errors, as opposed to vowel errors (Cohen, 1980). In fact in a study by Tent and Clark (1980) anticipatory phonetic level errors in general were detected less than 28% of the time, whereas word and syllable level anticipations were almost always detected - 98% of the time. Despite this high level of perceptibility speech error corpuses are singularly lacking in syllable errors (unless they are also analysable as morphemes). Apparently then, speech error corpuses can, and do often reflect production phenomena and conclusions about production processes may be reached on the basis of them - keeping in mind the various possible confounding

factors, and tempering our conclusions accordingly. The more general problem of indeterminacy discussed earlier merely reflects the fact that theories in general are underdetermined by evidence, and thus poses no unique problem for production theories.

In the next section we will outline our aims and working hypotheses and follow this up with a description of the speech error data which will provide the basis for our inferences regarding the production syllable.

3. Assumptions Underlying Our Investigation

As mentioned above, in this study we are interested first in determining whether speech error data confirm the viability of the syllable as a unit of production. Secondly, if we find evidence for a production syllable, we would like to determine how closely it resembles the various constructs proposed by competence theorists. This means we shall want to define its boundaries, fuzzy or clear (see Chapter 1, Section 2.2) and determine when syllabification takes place - or in other words, determine the domain of the syllabification algorithm. We shall also want to see whether this algorithm should define any internal structure and if so, determine how closely it corresponds to the proposed competence units.

(Our final goal will be to determine how our conclusions regarding these aspects of the syllable may be integrated into current models of the entire production process. This we will leave to our final chapter.)

In this section we should like to outline the assumptions which will underly this investigation.

The clearest characteristic of speech errors is their non-random nature. It is this characteristic that provides the basis for the assumption that errors occur during the normal production process and can lead to insights regarding this process. This non-random quality is particularly obvious when sequences of segments occur out of their intended position. Such segment sequences are always analysable as constituents of some given sort - speech errors, like languages appear to be structure dependent. For instance in (6) - (9) we have errors which involve, respectively, phrases, words, free morphemes and bound morphemes. (Dots (...) in the right hand column indicate that the material not reprinted from the left hand column was correctly uttered. Dashes (--) indicate a pause on the part of the speaker.)

Phrase: (6) (a) I have to smoke ...smoke my coffee
 [a cigarette] with with a cigarette
 NP
 [my coffee]
 NP NP

(b) If you'll [stick ...meet him you'll
 VP stick around
 around] you'll
 [meet him]
 VP

(c) He [removed the he facilitated what
 VP he was doing to
 barricade] to remove the barricade
 VP
 [facilitate what
 VP
 he was doing]

Word: (7) (a) I want to see the ...kid King
 Burgher King kid
 (b) You got to try to you try to got...
 make a play⁶
 (c) Back later with back more with later
 more
 (d) We have to bring ...some cat for the
 some milk for the cat milk

Free Morphemes:

(8) (a) A hard childhood a child hard hood
 (b) It's really thirst ...quench thirsting
 quenching
 (c) Let's sit some place ...some side out
 outside-like place - like
 (d) A bucketful of purrful of buckets
 purrs

Bound Morphemes:

(9) (a) ...had forgotten had forgot abouten
 about that that

- | | |
|---|-------------------|
| (b) ... <u>pointed</u> out | point outed |
| (c) <u>adding</u> ten | add tening |
| (d) You might as well
try to do something
with a service
<u>return</u> | ...reservice turn |

Note that all the errors in (6) - (8) involve two units out of position; more specifically they involve the exchange of these two units, and these are constituents of the same type. That is, it seems that we do not find errors which exchange units of different types. The errors in (10) - (11) appear to be impossible errors.

- | | |
|--|--|
| (10) (a) *We have to bring
some milk for the
cat | ...some [the cat] for
NP
[milk]
N |
| (b) | *some [the] for [milk]
art N
cat |
| (11) (a) *It's really [thirst]
quench[ing] | really ing quench
N thirst |
| (b) *a bucketful of
<u>purrs</u> | a bucket purr of fulls |

These facts suggest that it may be a reasonable working hypothesis to assume that when two elements are involved in an error they will always be elements of the same type. The errors in (10) - (11) involve exchanges of different types,

nouns with noun phrases and affixes with roots and it is this difference in type which makes the errors impossible. Garrett works on this assumption, which he refers to as the two principles of error analysis: (1976, p. 237) elements which interact in errors must be describable in terms of one type, and the constraints on errors of a given type, must also be characterized in terms of one type. By this latter principle he is referring to the errors' privilege of occurrence in terms of the limits on type and/or length of string which may occur between the error elements involved. Ideally, one would like to be able to characterize both the error elements and its environment in terms of the same linguistic type. So, for instance, if heads of syntactic categories are involved in an error, one would like to be able to specify their environment in terms of syntactic notions like clause or phrase, and not in terms of prosodic notions like foot or syllable. These principles of Garrett's are not ad hoc but follow naturally from his basic theoretical assumption. Namely that "...the computational description of language processing does, indeed, reflect the grammatical decomposition of the language faculty" (1982, p. 27). The grammar is modular hence the 'production grammar' will be modular. The different modules of the grammar Garrett equates with different processing levels of the production mechanism. Errors may occur at one or another of

these independent levels. Since different processing types will be handled at different processing levels and since the levels are independent of each other, errors must involve elements of a single processing type. In other words Garrett's principles of error analysis are not separate working hypotheses but are corollaries of his basic assumptions regarding modularity, which in turn reflect his assumption regarding the relationship between the competence grammar and processing. Of course it is not obvious what the processor analyses as elements of the same type, nor how closely competence types correspond to processing types. Nor is it necessarily the case that only elements of the same type may be involved in exchanges (i.e. the assumption that the production grammar is modular may be wrong. For instance (12) - (14) appear to go against this assumption.

- | | |
|----------------------------|---------------------|
| (12) Did his [team] [win]? | ...win team? |
| N N V | |
| (13) The weather report | ...clouds for cally |
| [calls] | skies |
| V | |
| for [cloud]y skies. | |
| N | |
| (14) He's wearing too | ...too short tights |
| [tight]-[shorts] | |
| Adj N | |

Here we see the interchange of nouns, verbs and adjectives. Notice however that they are all open class roots⁷ (and at

the same $\bar{X}$ level). We will make the assumption here (like Garrett) that the 'production grammar' is modular. Hence if an error involves two elements, these two elements must be analysable as tokens of one type, because the error occurs when this particular type is being processed. This is distinct from the claim that any type identified as an element in the production grammar may be involved as a unit in an exchange error. It is also distinct from the claim that any type of processing unit is subject to errors. It is tempting here to assume the biconditional: i.e. two elements may be involved in an error iff they are elements of the same type.

We may not assume this however; it appears to be a necessary condition on errors but not a sufficient condition. Apparently we may have elements of the same type which are not involved in exchange errors. For instance Garrett argues that bound morphemes do not exchange with one another, although they are subject to movement errors as in (9) above. We also may have processing types which are not involved as units in errors at all. For instance, as we will do our best to illustrate in this thesis, the syllable.

The syllable behaves quite differently from constituents like the word, the phrase, and the morpheme, and as a result we have two diametrically opposed views with regard to its

status as a processing type. Shattuck (1975) (as do others) observes that errors only very rarely involve syllables and hence Klatt (1981) conjectures that "the mental lexicon used for speaking represents words in terms of sequences of phonemes (rather than in terms of syllables or distinctive features)" (p. 11). On the other hand, Fromkin (1971) and MacKay (1970, Boomer and Laver 1968, and Nooteboom, 1969) all believe that speech errors support the viability of the syllable as a processing unit. One reason for this is that "There are of course many errors which involve the substitution, omission, replacement, addition etc. of one or more whole syllables..." (Fromkin, 1971, p. 229).

Notice that this seems to contradict the Shattuck findings referred to above. This is because what Fromkin classifies as syllable errors are all analysable as morpheme replacements or deletions, blends or haplogogies. For instance (15):

- | | |
|--|-------------------------------|
| (15) (a) butterfly and
catterpillar | butterpillar and
catterfly |
| (b) clarinet/viola | clarinola |
| (c) shrimp and egg
souffle | shrig souffle |

Because only that subset of syllables is involved in errors that are concomitantly analysable as morphemes, Klatt

concludes that "Syllables...play at most a limited role in representing morphemes" (ibid.). But this implies that 'moveability' (including under this, deletion and replacement) is a criterion for 'unit-hood'. That is, if some hypothesized type of construct or unit is not involved in errors of deletion, addition or replacement, it is not supported, on this view. Surely this is one kind of evidence for some hypothesized type but not the only kind. As Prince (1983) so succinctly puts it "It is a rule of thumb of practical ontology that a thing exists to the extent that other things interact with it, make use of it" (p. 31)). That is, if generalizations regarding speech errors can only be made on the basis of the syllable, and if the syllable can, at the same time, allow the formulation of generalizations about other aspects of the production process, generalizations not otherwise formulable, then the syllable is supported. Indeed part of Fromkin's argument for the syllable is based on MacKay's claim that the "syllabic position of reversed consonants was almost invariably identical" (MacKay, 1969) and she observes that "the only examples in (her) data which do not support this finding are..." (1971, p. 328).

(16) (a) whisper

whipser

(b) ask

aks

It's not at all clear to me how she syllabifies (she doesn't define her algorithm) because in her within-word errors alone (Appendix) she has four other reversals which seem to contradict this. Namely, (17):

(17) fish	shiff
medal	melad [mɛlɔd]
reverse	reserve
shilling	shingle

And in her large corpus of between-word errors it is certainly not always obvious that the identical syllabic position is involved. For instance (18):

(18) western Mexico	wexsern Mestico
---------------------	-----------------

Under most syllabifications the 'st' of 'western' would be syllable initial and the 'x' of 'Mexico' would include a syllable final 'k' and syllable initial 's'. It is also not clear under Fromkin's analysis what particular syllabic position is involved in (19) (a)-(b),

(19) (a) a heater switch	a sweeter hitch
(b) a pedal steel guitar	stedal peel

since Fromkin seems to identify syllable position with segment position. For instance she claims that the syllabic positions of the 's' and 'k' are distinct in 'ask', which

suggests that she would consider the above clusters to be filling two syllabic positions. It is not clear how the two positions in (19) (a)-(b) can be considered identical to one position under this type of uni-level analysis. Another indication that she identifies syllabic position and sequential position is her hypothesis that such misorderings (as in (19)) occur when "segments are 'sent' into the short-term memory buffer, (and, AL) segment 1 of syllable 1 may be substituted for segment 1 of syllable 4" (1971, p. 240). This formulation is clearly not at all satisfactory. It would incorrectly predict the following errors for example.⁸

- (20) (a) *[hi][sto][ry] *aylstory hdeology
 1 2 3
 [i]deology
 4
- (b) *democratic remocdatic

We believe however that the basic insight is correct, that syllabic structure is important with regard to consonant reversals. Indeed it is the very purpose of our thesis to give substance to this insight, to formalize it and extend it. Syllabic structure, we will maintain, is crucially involved in the explanation of all types of between-word submorphemic segmental errors.

In our attempt to support this claim, we will, as we have just discussed, assume as working hypotheses:

(i) that speech errors reflect the normal production process,

and (ii) Garrett's hypothesis, that the processing mechanism will reflect the grammatical decomposition of the language faculty.

From these assumptions it follows that 'production grammars' are structure dependent; any unit involved in an error will be analysable as a constituent of some sort, and if two elements are involved in a single error they will be constituents of the same type. But it does not follow from these assumptions that all constituents may themselves be subject to mis-selection errors. There is no assumption made regarding a criterial form of evidence.

A constituent type will be supported to the extent that it has explanatory power.

4. The Data Base of Our Investigation

As mentioned above, unless it is a morpheme, the syllable as a unit is only very rarely involved in errors; this has led to conflicting views as to its viability as a unit in the production process. To realize our aims

regarding the syllable's status and structure we will have to use segmental errors which are unambiguously non-morphemic. But since syllables per se are apparently not subject to mis-selection we would appear to be left with no data on which to base arguments regarding the viability of the syllable as a processing unit. This is not the case however. A type is viable in so far as it has explanatory capacity. Various theorists have noticed that segmental exchanges are constrained in that they respect syllabic position.⁹ This informal observation, if true, could form the basis for arguments in support of the syllable as a processing unit. The constraint could only be explained if the syllable were recognized. On our assumptions, it is natural to assume that the so-called segmental errors (or 'phonemic' errors, Nooteboom, 1969) all occur at the same level of processing because, superficially at least, they all appear to involve the same processing type, - namely, the segment. We want, if possible, to give explicit formal content to the informal insight that consonant exchanges¹⁰ respect syllabic position, but this characteristic should not be unique to consonant exchanges if all segmental errors occur at the same level of processing. What is characteristic of one type of segmental error should be characteristic of the others; if not, this will have to be explained. Hence we will want to study not

only the reversals but also the other types of segmental errors.

But we are limited to studying those errors where one position may be compared to another position. Exchanges are ideal sources for testing the claim that syllabic position is maintained because the string position of both elements is clear. Another clear case of two string positions is the class of errors which Fromkin calls 'movement' errors (see below). Perseverations and anticipations also provide us with string positions that may be compared, and although they are not clear-cut, i.e. we can never be sure that the error is not merely a substitution, it appears that some characteristics exhibited by reversals are also common to perseverations and anticipations;¹¹ (see below) hence it is not unreasonable to include them as a source of evidence. We shall also include some of those errors which Fromkin calls 'additions' and some which she refers to as 'substitutions'.

In the literature, speech errors have been given numerous labels; the criteria according to which errors have been classified have been diverse and sometimes baffling. In general terms the defined classes will reflect the theoretical perspective of the classifier (see discussion regarding Garrett's classification of exchanges in Section 2.1) but more particularly and ignoring simple terminological

differences,¹² we may distinguish four main criteria that have played the biggest role in the classifications. First, the focus may be the type of unit involved: word, noun, morpheme, segment, vowel, consonant, feature, etc. Secondly, the focus may be the resulting output string. Some unit may be added to, omitted, replaced or moved in the output. Thirdly the classification may reflect the classifier's hypothesis regarding the source of an error. Finally, and probably most importantly, the classes may reflect the production modeler's hypothesis as to what the processor is doing when a particular error occurs. The classes in Fromkin's Appendix reflect all of these to a greater or lesser extent. For instance she classifies the following errors as additions, deletions, substitutions and movements, respectively, focussing on the resulting output strings (the second criteria above).

(21) additions:

(a) Try to understand	...understrand
(b) significantly	significlantly
(c) the optimal number	...moptimal...
(d) think through	thrink...
(e) all sorts of errors	...sports...
(f) enjoying it	enjoyding

(22) deletions:

- | | |
|-----------------------|-----------|
| (a) stretch your legs | retch... |
| (b) posed nude | poed nude |
| (c) below the glottis | ...gottis |

(23) substitutions:

- | | |
|--|-------------------|
| (a) all these magnificent
sights | ...bagnificent... |
| (b) What does the course
consist of | ...gorse... |
| (c) phonetic data | pholetic... |

(24) movements:

- | | |
|-------------------------------------|---------------------------------|
| (a) price to pay | pry to pace |
| (b) Ed and Bessy | Bed and Essy |
| (c) the feature back
plus labial | ...black puss... |
| (d) a copy of the
bibliography | a copley of the
bibliography |

In (21) and (22) the output string contains too many or too few segments, in (23) there is an incorrect segment replacing an intended segment and in (24) the output string includes a segment out of position.

Consider now the errors in (25) - (27) which she labels as anticipations, perseverations and reversals, respectively:

(25) anticipations:

- | | |
|-------------------------|---------------|
| (a) write it or type it | ripe it... |
| (b) paddle tennis | taddle tennis |
| (c) with a brush | wish... |

(26) perseverations:

- | | |
|--|----------|
| (a) I can't cook worth
a damn | ...cam |
| (b) thick slab | ...slack |
| (c) have you given your
paper title | ...titer |

(27) reversals;

- | | |
|---------------------|-----------------|
| (a) hash or grass | hass or grash |
| (b) left hemisphere | heft lemisphere |
| (c) top shelf | tof shelp |

The errors in (25) - (26) have been distinguished from those in (23) because although there is an incorrect segment replacing another, as was the case in (23), in (25) - (26) there is a presumed source for this incorrect choice. In (23) there was no segment in the preceding or following context which could be considered the source of the error. (27) and (24) also appear to be similar. In (27) there are two segments out of position compared to the one segment in (24). But they have to be considered distinct error types by Fromkin. Because Fromkin believes that the processor is substituting one segment for the other in (27) and since the notion of a null segment is meaningless, the procesor can not be exchanging one segment for another in (24). This distinction between 'movements' and 'reversals' is therefore a reflection of Fromkin's hypothesis as to the type of unit being processed (the first classification criteria mentioned

above), interacting with her hypothesis concerning the functioning of the processing mechanism when the error occurs (the fourth classification criteria). She is clearly correct in assuming that reversals are not simply two independent movements which by chance move each of the segments into the other segment's position. Reversals are far too common for this interpretation. Thus she is forced to distinguish the two types of error. She is faced with a similar problem when it comes to the errors she classified as additions. The errors in (21) (a) - (d), look suspiciously like the perseverations and anticipations in (25) - (26) whereas (21) (e) - (f) are like her consonant substitutions in (23). In (21) (a) - (d) we find segments from preceding or following context appearing in the string whereas in (21) (e) - (f) we see no contextual source for the errors. Fromkin does not classify (21) (a) - (d) as perseverations and anticipations because she defines these error types as the replacement of some segment (or feature) by some segment (or feature) where the segment is from the preceding or following context. Since there is no segment to be replaced in the positions indicated by a dash (--) in 'underst--and', 'signific--antly', '--optimal', and 'th--ink', these cannot be classed as anticipations or perseverations. Similarly there is no segment in the indicated positions in 's--orts' and 'enjoy--ing' and since she defines substitutions as the

substitution of one segment (or feature) for another, these must be classed as additions.

To this point we have given examples of all the error types which involve consonants. Errors which are of interest to us include two positions which may be compared, namely reversals, perseverations, anticipations, movements, and those errors called 'additions' by Fromkin, which appear to be perseverations or anticipations. Errors involving vowels are also of interest to us, although the fact that consonants and vowels never interchange cannot be taken as evidence for the syllable. Fromkin, for instance, believes that this fact is due to a major delineation between the two types based on the features [consonantal] and [vocalic] (Fromkin, 1971). Errors which involve vowels are crucial for determining the structure of the peak, or nucleus constituent, however. We find vowels involved in exchanges, perseverations, anticipations as exemplified in (29) - (31) although not movements.¹³ For instance we have no errors like (28) (a) analogous to the consonant movement in (28) (b).

(28) (a) *ideas come and go deas icome and go

(b) Anders Jarryd Janders Arryd
 [ændərs] [yærid]. [yændərs] [ærid]

Exchange:

- | | |
|---------------------------|-------------------|
| (29) (a) feet moving | fuwt meeving |
| (b) David, feed the pooch | ...food the peach |
| (c) ad hoc | odd hack |

Anticipations:

- | | |
|--|--------------|
| (30) (a) you picked out such
a nice book
[bvk] | [pʊkt]... |
| (b) several cars reported | ...[korz]... |
| (c) available for
exploitation | available... |

Perseverations:

- | | |
|--|---------------------|
| (31) (a) the data is derived
[deɪ] | ...deraved
[rey] |
| (b) beef noodle soup | ...needle... |
| (c) environment before
[vay]
voice | ...vice
[vays] |

The errors in (29) - (31) are all from Fromkin class S (1971, p. 251) and for some reason this class of vowel errors is labelled as consisting of reversals, additions, deletions and substitutions, although the analogous errors involving consonants would be classed as anticipations and perseverations. Although she considers those in (29) to be reversals, she does not label those in (30) - (31) as perseverations and anticipations. These she appears to consider substitutions, and those in (32) appear to involve deletion and addition, besides substitution (under her analysis).

- | | |
|--|------------------------|
| (32) (a) foreground items | ...grind... |
| (b) the committee will
have to meet more
often | ...commeetee... |
| (c) competing rules | computing...
[pyuw] |

Again, because she believes the segment (or feature) is being processed the addition and deletion of glides in (32) forces her not to analyse the errors in (32) as anticipations and perseverations. However it is not at all clear why (30) - (31) are 'substitutions' since there is a clear contextual source for the substitution. We shall refer to the error types in (29) - (32) as reversals, perseverations and anticipations.

We have not yet referred to multisegmental errors such as (33).

- | | |
|---------------------------------------|------------------|
| (33) (a) squeeze fresh oranges | freeze squesh... |
| (b) I can knit but I can't
crochet | ...krit... |
| (c) steady state vowels | ...stowels |
| (d) in the next ten minutes | nen text... |
| (e) pussy cat | cassy put |
| (f) heap of junk | hunk of jeep |

As we shall see below, there are a number of errors involving sequences of C's which could be classified as exchanges, anticipations and perseverations as in (33) (a) - (d), but we don't often find sequences of C's and V's which are perseverated or anticipated. These sequences are invariably exchanges like (33) (e) - (f) unless the two sequences involved have identical vowels such as the errors in (34), in which case we do find such sequences perseverated or anticipated or exchanged.

- | | | |
|----------|--|-----------------------------|
| (34) (a) | I see <u>lots</u> of initial
<u>minutes</u> | ... <u>im</u> itial... |
| (b) | last <u>gasp</u> | <u>lasp</u> gasp |
| (c) | ...lost to Hana
<u>Mandlikova</u> | ... <u>Ma</u> na Mandlikova |
| (d) | I'm <u>washing</u> my watch | ...my <u>wash</u> |
| (e) | talk to you in a
<u>little</u> bit | ... <u>blit</u> |
| (f) | <u>white</u> <u>lion</u> | <u>light</u> <u>wion</u> |
| (g) | <u>quite</u> <u>right</u> | <u>krite</u> <u>wight</u> |

These errors highlight the final constraint on which errors can serve as the basis for our analysis. A number of segmental errors are ambiguous either with respect to the actual elements involved in the error, errors (34) (a) - (g) above, or with regard to a perseverative or anticipatory source as in (35) (a) - (c). (The possible sources are underlined.)

- (35) (a) hard to get good ...hood...
 help
- (b) I don't see how anyone ...plean...
 could eat a plain
- (c) Well, Mum was with me ...mith...

Ambiguous errors such as (35) (a) - (c) we have not included because we want to defend our claim that anticipations, perseverations and exchanges occur at the same level of processing and thus we require clear cases of the three types. We clearly cannot include the ambiguity type exemplified in (36)

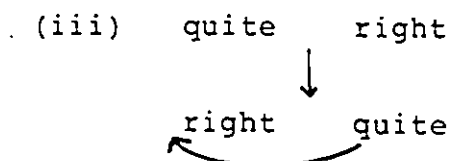
- (36) the optimal number ...moptimal

because the syllabic position of the source can't be determined.

The ambiguous errors in (34) have been included however, for reasons which we will immediately explain.

All types of segmental errors have certain characteristics in common. For instance the feature makeup of segments involved in errors is often more similar than would be predicted by chance (cf. van den Broecke and Goldstein, 1980, Nooteboom, 1969, among others) and so the replacement of 'n' by 'm' in 34 (a) could be the anticipation of [-cor] rather than anticipation of the entire segment.

Similarly, it has been noted by a number of observers that "The environments of moved elements...are similar" (Garrett, 1980, p. 184). This applies to vowels but even more strongly to the environments of consonants as illustrated in (34). Is the error in (34) (g), (i) the reversal of 'r' and 'u' (quite wight), or (ii) the exchange of the whole strings 'uite' and 'right' or (iii) is there a morpheme exchange followed by a movement?



Errors such as these can only be interpreted in the light of the patterns we find in the apparently¹⁴ unambiguous errors. (We have here the familiar problem of the undertermination of theories by evidence (discussed above).) We have consistently assigned the 'minimal' interpretation to the types of errors in (34), i.e. if the error can be analysed as the exchange, reversal or anticipation of one segment rather than two segments it has been analysed as such.¹⁵

Another error type which can provide us with information regarding syllabic structure is the type commonly referred to as 'word blends'. Blends involve words similar in meaning which coalesce to form a nonexistent word consisting of part of one word and part of the other. For instance (37) (a) - (c):

- | | |
|-----------------------|----------|
| (37) (a) sunny/hot | sot |
| (b) terrible/horrible | torrible |
| (c) makes/means | manes |

Almost all blends involve the simple juxtaposition of the two competing words and hence the syllabic position of the breaking point in the first word and the continuation point in the second word may be compared. These errors too will be analysed following our analysis of segmental errors.

To sum up, our data will consist of blends and those submorphemic errors analysed by Fromkin as segmental reversals, perseverations, anticipations and movements, and some of those which she would analyse as substitutions, additions and deletions. The error is included if we can unambiguously relate some segment or sequence of segments to two string positions (always choosing the minimal analysis, as explained). We found five hundred and fifty-nine segmental errors which met our criteria, and 113 word blends. The major sources of our data have been the Fromkin Appendix and our own corpus (n=36 blends, n=171 segmental errors). Errors from other sources have been included with Garrett, 1980, 1982, Shattuck-Hufnagel, 1979, 1984, Goldstein, 1968, and Eay and Cutler, 1977, providing the greatest number. In (38) - (42) we list examples of all the error types which

will provide the basis for our study. (The two relevant string positions are indicated by arrows.)

Reversals:

- | | |
|-----------------------------|----------------------|
| (38) (a) Lloyd Moseby | Lloyz Modeby |
| (b) It hurts Tate | ...turts hate |
| (c) he's always goofing off | ...goffing oof [uwf] |

Perseverations:

- | | |
|--|------------------|
| (39) (a) practice teaching | ...preaching |
| (b) quadruple | qwadwuple |
| (c) group one, group two and group three | ...group through |
| (d) innate urges | ...nurges |

Anticipations:

- | | |
|----------------------------------|----------------|
| (40) (a) grapefruit flavour | ...fluit... |
| (b) Did he play for the Redwings | ...wedwings |
| (c) people who keep the same... | ...cape... |
| (d) Cathy Squires | Cwathy Squires |

Movements:

- | | |
|----------------------|--------------|
| (41) (a) re-energize | regenerize |
| (b) parking spot | sparking pot |
| (c) battle axe | attle baxe |
| (d) egg yolk | olk yegg |

Blends:

(42) (a)	tines/prongs	prines
	↑ ↑	
(b)	Orrie/Eric	Orric
	↑ ↑	
(c)	space/place	splace
	↑ ↑	

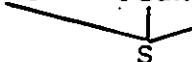
5. Giving Empirical Substance to the 'Syllabic-Position-Constraint' Hypothesis

Since syllables themselves are not involved as units in errors we must argue for its recognition as a unit of production on the basis of its explanatory power. We have suggested above that it must be recognized as a processing unit if segmental errors only involve segments from identical syllabic positions. This claim regarding a 'syllabic-position-constraint', i.e. Boomer and Laver's claim that "...segmental slips obey a structural law with regard to syllable-place; that is, initial segments in the origin syllable replace initial segments in the target syllable, nuclear replace nuclear, and final replace final" (cited in Fromkin, 1971, p. 228) must be made precise in order to test it. We must have a precise algorithm for determining what segments are initial and final, which will entail having a precise domain to which the algorithm applies. For our first analysis of the data we will use Kahn's slow speech rules, (Rules I and II),¹⁶ to determine syllable boundaries. These rules represent one version of the maximization of onsets and the results of their application in terms of the boundaries

defined would appear in most cases to be indistinguishable from the boundaries which would be defined by Selkirk, C&F, and Hooper for instance, and we had no data which made the differences crucial in the analysis.¹⁷ Kahn's rules apply at the beginning of a derivation to words and we shall, accordingly, take the word as the domain of application. Kahn's rules do not assign any internal structure to the syllable however, and in order to test Boomer and Laver's claim we shall need to specify precisely what segments belong to initial, nuclear and final positions. We shall modify Kahn's rules as indicated in Figure 1 below, calling the position that precedes the [+syll] segment, the onset position, and the position which follows the [+syll] segment the coda position. At the same time we shall call the position that the [+syll] segment occupies the Peak position,¹⁸ and shall include in this position any glide which has not been assigned to the onset position by Rule 2 (a) in Figure 1.¹⁹ In other words Kahn takes as given only the S node, whereas we take as given the S node and three positions as specified by Rule 1 (a).

Figure 1Kahn's Slow Speech Rules Modified

Rule 1': (a) $S \rightarrow \text{Onset} \quad \text{Peak} \quad \text{Coda}$



(b) $[+syll] \rightarrow [+syll]$



Rule 2': (a) $C_1 \dots C_n [+syll] \rightarrow C_1 \dots C_i C_i + 1 \dots C_n [+syll]$



where $C_i + 1 \dots C_n$ is a member of the set of permissible initial clusters, and

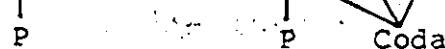
$C_i C_i + 1 \dots C_n$ is not, and C is a variable ranging over any $[-syll]$ segment.

(b) $[+syll] \begin{bmatrix} -cons \\ +son \end{bmatrix} \rightarrow [+syll] \begin{bmatrix} -cons \\ +son \end{bmatrix}$



where $\begin{bmatrix} -cons \\ +son \end{bmatrix}$ is unattached.

(c) $X C_1 \dots C_n \rightarrow X C_1 \dots C_i C_i + 1 \dots C_n$



where $C_1 \dots C_i$ is a member of the set of

permissible final clusters and $C_1 \dots C_i + 1$

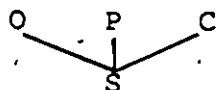
is not, and $C_1 \dots C_i$ are unattached.

Examples of the structures assigned by the syllabification rules are seen in Figure 2.

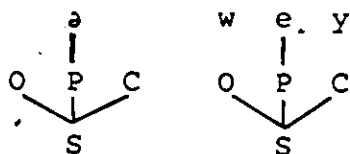
Figure 2

Structures Assigned by Rules 1' and 2'

by 1' (a):

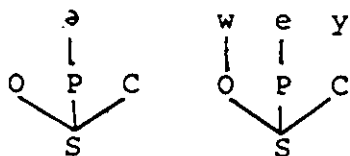


by 1' (b):

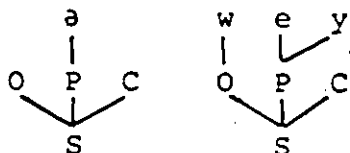


(i) away

by 2' (a):

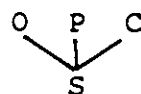


by 2' (b):



by 2' (c):

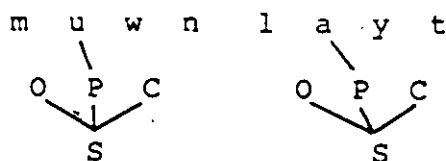
NA



(ii) moonlight

by 1' (a):

by 1' (b):



by 2' (a):

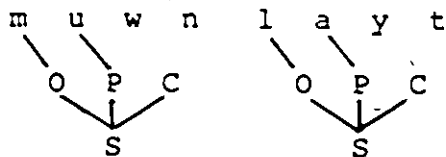
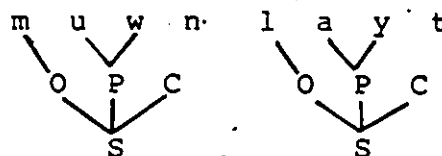
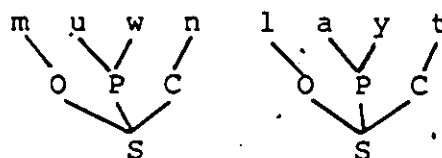


Figure 2 (continued).

by 2' (b):



by 2' (c)



On the basis of this syllabification algorithm we can identify five sub-syllabic positions which segments or sequences of segments involved in errors could belong to:

onset (O)

peak (P)

coda (C)

onset-peak (O-P)

peak-coda (P-C)

In terms of these five positions there are $5 \times 5 = 25$ substitution possibilities (ignoring the additional possibilities including non-contiguous sub-syllabic positions). That is, source onset might substitute for target onset, target peak, target coda, target onset-peak or target peak-coda. Similarly for each of the other four.

We will designate the two positions involved in an error using a slash (/) between the two positions. So for instance the error in (43)

(43) phone call

cone fall

would be designated as an onset/onset error whereas the error in (44)

(44) Dexie Blaff

Blaxie Deff

would be designated as an O-P/O-P error. A hyphen (-), then, indicates a sequence of syllable positions, whereas a slash (/) separates the positions interacting in errors.

In our first analysis any error which involved part of an onset (or coda) or an entire onset (or coda) was considered an onset (coda) position error whether the source or target was implicated. For instance in (45) we have only part of an onset position involved in the errors

(45) (a) strive for perfection sprive...

(b) the gutters are ...glutters are
flooded fuddled

whereas in (46) we have two entire onsets.

(46) (a) central choir

kwentral sire

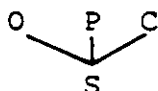
(b) highway driving

dryway hiving

All of these will be considered onset/onset errors in this first analysis.²⁰ Also, errors analysed as movements or additions by Fromkin such as (47) - (48) respectively,

- | | |
|---|---------------|
| (47) (a) Anders Jarryd | Janders Arryd |
| (b) ...toast will pass | toas...past |
| (48) (a) Talk to you in a
little bit | ...blit |
| (b) If you've got 'ing' | ...'ging' |

where there is only one segment involved (i.e. there is no substitution of one segment by another) will also be analysable in terms of syllabic positions, because we begin with a three position structure



where the onset and coda positions may be empty. Thus the error in (47) (a) will be identified as onset/onset and (47) (b) as coda/coda. That is, they are analysable as onset and coda position errors respectively.

A crucial part of our analysis will be the analysis of those errors which involve consonants which are not exclusively from word-initial or word-final positions. For instance the errors in (49) - (50)

- | | |
|--------------------------|------------|
| (49) It's a real mystery | ...meal... |
|--------------------------|------------|

(50) It's the red book

...bood .
[bud]

appear to be onset/onset and coda/coda errors. However they could be simply the interaction of word-initial and word-final positions respectively. Since syllabic constituent positions can not be distinguished unambiguously from word-initial and word-final positions no conclusions regarding the truth of the 'syllabic-position-constraint' (i.e. Boomer and Laver's claim) can be based exclusively on errors which involve segments from these positions. Similarly no conclusions can be reached on the basis of the fact that nucleuses ('prescriptivists' read 'nuclei') never interact with margins - this could be explained as due to the major class difference between vowels and consonants. Fromkin believes, for instance, that consonants and vowels are distinguished by the processor on the basis of the features [consonantal] and [vocalic]. The same set of segments appear in onset and coda positions (except for [h] and [ŋ]) hence if the programmer distinguishes consonants in onsets from those in coda positions it cannot be a function of features. (Low level allophonic rules apply after segmental errors occur (see Chapter 3) and so they can't be distinguished allophonically either.) Therefore, if the programmer distinguishes onset consonants from coda consonants it cannot be attributed to feature differences, but such errors could be attributed to

word position. Without analysing the errors which do not involve word-initial or word-final position, any conclusions regarding syllable structure could be an artefact of lexical structure and not syllable structure per se. Therefore we will also analyse separately, those errors in which at least one consonant is from a word-medial position. The errors (51) - (54) exemplify this type of error.

- | | | |
|------|--|-----------------|
| (51) | Ivan Lendl overpowered
John McEnroe in three sets | ...McEnthrow... |
| (52) | Vague category; they
can't even define it | ...devine |
| (53) | ...likes to come Monday
Wednesday and Friday | ...Monsday... |
| (54) | I'm missing mash | ...mishing... |

To sum up, in our first analysis we will syllabify words involved in segmental errors according to the algorithm in Figure 1 above, with the word as the domain. We will compare the syllabic positions of the segments involved in errors. We will first analyse all the errors and then separate out the errors which involve at least one consonant from a word medial position. By specifying precisely the domain and algorithm for syllabification we have made the claim regarding a syllabic-position-constraint on segmental errors an empirically testable hypothesis. If it holds up we should find that of the over 25 possible syllable position

interactions only a very few are found - only those which involve identical positions, as defined by our algorithm in Figure 1, above.

6. Testing the Hypothesis

(i) Results

Segmental errors involve segments from identical syllabic positions in 91% of the cases. Of the more than 25 possible syllabic position interactions only three occur with any frequency as indicated in Table 1, where errors involving interactions between onsets, peaks or codas make up 88% of the total. Onsets are the most frequently implicated positions, accounting for 63% of the errors.²¹

Table 1
Syllabic Position of Segments Interacting in Errors

Syllabic Positions		Frequency (Number)	Frequency (Percent)
Most	onset/onset	352	63%
Common	coda/coda	59	11%
Errors	peak/peak	81	14%
Subtotal		492	88%
Other Errors	'v+r'/peak	14	3%
	'v+r'/'v+r'	5	1%
	peak-coda/peak-coda (rhyme) (rhyme)	11	2%
	onset-peak/onset-peak	8	1%
	onset/coda	25	4%
	other	4	1%
Subtotal		67	12%
TOTAL		559	100%

N.B. Percents are rounded off.

The eight percent (8%) of errors which do not involve identical syllabic positions include onset/coda errors (4%) as in (55)

(55) I'm washing my watch ...watching...

and errors involving sequences of vowel plus 'r' (V+r) interacting with peaks, both monophthongs and diphthongs as in (56) - (57)

(56) the shot didn't hurt much ...shirt...hot...

(57) support the movement sepuwt
 [səp^huwt]

There are also other 'V+r' errors which involved another 'V+r' sequence as in (58):

(58) (a) cards sorted cords sorted
 (b) took part in the first pert in the first

These, (58) (a) - (b), could have been classed with the other peak errors, since I have always used the minimal analysis (as discussed above), and these appear to be errors of vowels with other vowels. But I have separated them out for reasons which will become clear below in our discussion of the peak constituent.

It is clear from Table 1 that segmental errors very rarely implicate a sequence of syllabic constituents as defined by our algorithm in Figure 2. We have only eleven errors which include a sequence of peak and coda as in (59)

(59) heap of junk

hunk of jeep

and only eight which involve the sequence onset-peak as in (60).

(60) pussy cat

cassy put

As mentioned above the crucial errors for coming to conclusions regarding syllabic constituents are those which involve medial consonants. These errors exhibit the same pattern as those which include word-initial and final positions (see Table 2) but the proportion of onset/coda exchanges is higher as indicated in Table 2. It is 16% rather than the 5% onset/coda errors when word-initial and word-final positions are included.²² The proportion of onset/onset errors is high 72% but almost 10% lower: 72% as compared to 81%. The proportion of coda/coda errors is lower: 9% compared to 14%.

Table 2
Syllabic Position of Errors Involving at Least One
Word-Medial Consonant

Syllabic Position	Frequency	Frequency (%)
onset/onset	84	72%
coda/coda	11	9%
onset/coda	18	16%
other	3	3%
TOTAL	116	100%

(ii) Discussion

From the results tabulated in Table 1 (and confirmed in Table 2), it would appear that since 88% of segmental errors involve segments from one and the same syllabic constituent, syllabic constituent structure constrains segmental errors. That is segments do not substitute for one another randomly and the pattern of substitutions cannot be explained simply as a function of the feature makeup of segments. If feature similarity were the relevant variable there would be no explanation for the fact that usually onsets substitute for

onsets, and codas for codas. One way of accounting for our results is to assume that the processor is manipulating (or, equivalently, selecting, or processing) syllabic constituents when segmental errors occur. And, just as the processor can only mis-select among identical syntactic types (noun for noun, verb for verb, etc.) it can only mis-select among identical syllabic constituent types.

That is, as a first approximation we could characterize segment substitution errors thus:

Figure 3

Possible Segment Substitution Errors (PSE)

Anticipation:

$$\begin{array}{l} \text{If: } x \underset{\vdots}{O_1} y \underset{\vdots}{O_2} z \quad \text{If: } x \underset{\vdots}{P_1} y \underset{\vdots}{P_2} z \quad \text{If: } x \underset{\vdots}{C_1} y \underset{\vdots}{C_2} z \\ \text{Then: } (O_2) \quad (P_2) \quad (C_2) \end{array}$$

Perseveration:

$$\begin{array}{lll} \text{If: } x \text{ } O_1 \text{ } y \text{ } O_2 \text{ } z & \text{If: } x \text{ } P_1 \text{ } y \text{ } P_2 \text{ } z & \text{If: } x \text{ } C_1 \text{ } y \text{ } C_2 \text{ } z \\ \text{Then: } & & \\ & (O_1) & \sim (P_1) \quad (C_1) \end{array}$$

Reversal (and Movement):

$$\begin{array}{lll} \text{If: } x \text{ } O_1 \text{ } y \text{ } O_2 \text{ } z & \text{If: } x \text{ } P_1 \text{ } y \text{ } P_2 \text{ } z & \text{If: } x \text{ } C_1 \text{ } y \text{ } C_2 \text{ } z \\ \text{Then: } (O_2) (O_1) & (P_2) (P_1) & (C_2) (C_1) \end{array}$$

These could be conflated as:

PSE

If: $X \begin{pmatrix} O_1 \\ P_1 \\ C_1 \end{pmatrix} \quad Y \begin{pmatrix} O_2 \\ P_2 \\ C_2 \end{pmatrix} \quad Z$

Then: $X \begin{pmatrix} O_2 \\ P_2 \\ C_2 \end{pmatrix} \quad Y \begin{pmatrix} O_1 \\ P_1 \\ C_1 \end{pmatrix} \quad Z$

Examples of this at work are given in Figure 4 where the errors of (61) (below) are exemplified. (Only implicated nodes are linked.)

(61) (a) anticipation:

(i) what are you going to do in the Monroe booth ...Menroo...

(ii) I see lots of initial minutes ...imital...

(iii) last gasp lasp...

(b) perseverations:

(i) Have you seen Tango, Andrew ...Andro

(ii) Because he felt he wasn't getting enough ice time ...nice...

(iii) sip of my coke ...cope

(c) reversals (movements):

(i) the roof slopes ...rofe sloops

(ii) white lions light wions

(iii) Lloyd Moseby Lloyz Modeby

(iv) battle axe attle baxe

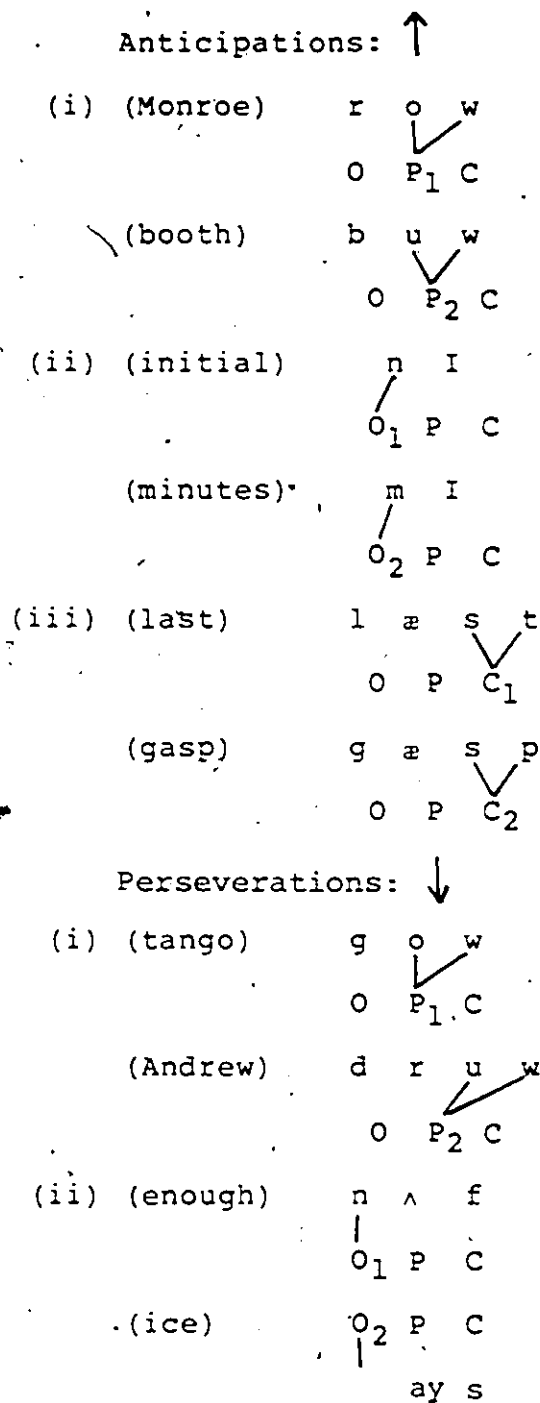
Figure 4Operation of PSE on Errors of (41)

Figure 4 (continued)

(iii) (sip) s I P
 O P C₁

(coke) k o w k
 O P C₂

Reversals (and movements):



(i) (roof) r u w f
 O P₁ C

(slopes) s l o w p s
 O P₂ C

(ii) (white) w a y t
 O₁ P C

(lions) l a
 O₂ P C

(iii) (Lloyd) l o y d
 O P C₁

(Moseby) m o w z
 O P C₂

(iv) (battle) b æ
 O₁ P C

(Axe) æ ks
 O₂ P C

Possible Segment Substitution Errors (PSE) is only a first approximation of the mechanism whereby segmental errors occur, although as such it accounts for 90% of the errors.²³ We have not yet distinguished errors which involve entire onsets and codas from those which only involve parts of these constituents. We will return to this below when we consider these constituents in detail (Sections 9 and 10). We also are assuming for the moment, that the few errors (3% of the errors) which involve the sequences of constituents, o-p and p-c involve the occurrence of two errors, onset then peak, and peak then coda. That is, for the moment these will have to be viewed as a conjunction of errors.

More importantly we have not accounted for almost 8% of the errors, the -V+r errors and the coda/onset errors. PSE cannot generate these errors. There seem to be certain clear ways of getting around this and explaining these errors. First, it is possible that segmental errors occur at more than one processing level and when the processor mis-selects among non-identical syllabic constituents it is not processing these constituents but some other processing type (e.g. segments or features). Secondly, our algorithm for syllabifying could be (slightly) inaccurate - it might define incorrect external boundaries, or it might pick out the wrong types internally, or both. Thirdly, the algorithm might be accurate but the domain of its application could be wrong,

i.e. it could involve some domain other than the word. Finally, there may be a combination of all or some of these. Let us first consider how 'V+r' errors may be accounted for. We will argue that 'V+r' belongs to the Peak constituent.

(iii) Problematical 'V+r' Errors

The clearest characteristic of the Peak is that it is an inviolable unit. There are no errors which split tautosyllabic vowels and glides. That is errors such as (62) would seem to be impossible.

- (62) (a) *The [boy] [baws] [bay] [bows]
 (b) *The moon's might mawn's muyt

We find that the same thing is true of tautosyllabic 'V+r' sequences. We don't find errors like (63).

- (63) *The horse bayed [hers] [boyd]

Secondly, as we have seen, errors involve identical constituents (or sequences of identical constituents). This is true of 91% of the errors and as we shall see below this number is conservative: a number of onset/coda errors are actually coda/coda errors. This suggests that 'V+r' is a constituent of some sort. Since the only constituent which 'V+r' interacts with is the peak constituent (other than other 'V+r' sequences) it is reasonable to conclude that 'V+r' is a Peak.

Kahn (1976) claims that 'r' is a glide in English. To support his claim, Kahn cites distributional and articulatory facts. The major class features [consonantal] and [sonorant] are defined in terms of relative obstruction in the oral part of the vocal tract ([cons]), and relative pressure build-up in the vocal tract ([son]). Kahn observes there is neither obstruction nor pressure build-up in the English 'r'. In other words it is [-cons], in fact a glide. Distributional facts characteristic of 'y' and 'w' are also true of 'r':

- (i) they occur pre-vocalically - yet, wet, red
- (ii) they are plus syllabic under stress - beet, boot, burr
- (iii) they are the second element in a diphthong - toy, tow, tore
- (iv) they only occur after tense vowels - say, so, care
- (v) no two glides may be tautosyllabically contiguous - *sowry this restricts 'r' too, hence we have 'tire' in two syllables not one - [tay.r]
- (vi) following [-cons] segments 't's are unreleased; after [+cons] they are released - hat, height, heart versus apt, mist, mint, belt

It is interesting too, that in SPE a weak cluster is characterized as one including a lax vowel plus C, or a lax vowel plus $C\left\{\begin{smallmatrix} r \\ w \end{smallmatrix}\right\}$. In Ijo, apparently 'r' patterns with 'w' and 'y' (Hooper, 1976) and in Sanskrit 'r' was weaker than 'l' and "not substantially stronger than glides" (p. 205, *ibid.*). Clements and Keyser (1981) note that a vowel plus 'r' "presents a steady state r-coloured vowel segment occupying the syllable peak..." (p. 16).²⁴ The data from speech errors leaves no doubt - 'r' is a glide. Just as onsets interchange with onsets and codas with codas, V+r switches with peaks, not with peaks and codas. For instance, consider the following V+r error, and the two possible syllable positions for 'r' - that is, either as part of peak (glide) or as part of coda ([+cons]), as illustrated in Figure 5.

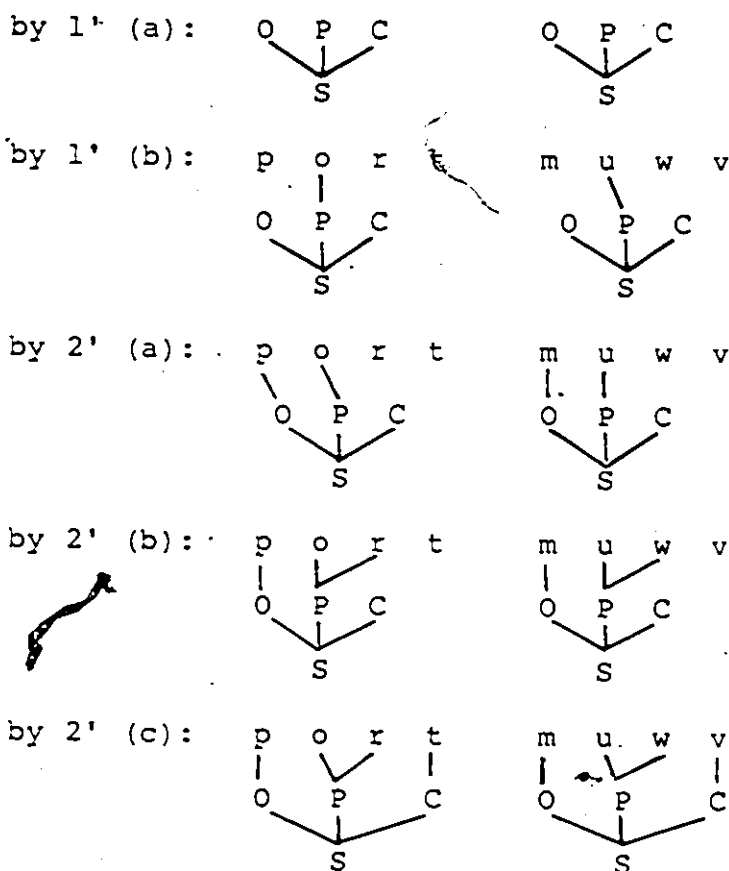
(64) Kansas City Cheifs
remained the first

Kansas City [čerfs]...

would be syllabified by our algorithm (Figure 1 above) as indicated in Figure 6.

Figure 6

Syllabification of 'port' and 'move'



By assigning 'V+r' errors to the peak/peak errors we can now claim that 92% of segmental errors involve the interaction of segments from the three constituents, onset, peak and coda, as defined by the algorithm in Figure 1 above. The revised Table, including 'V+r' errors as peak errors is:

Table 3
Syllabic Position of Segments Involved in Errors

Syllabic Position	Frequency (Number)	Frequency (Percent)
onset/onset	352	63%
coda/coda	59	11%
peak/peak	100	18%
onset/coda	25	4%
other	23	4%
TOTAL	559	100%

N.B. 'V+r' analysed as Peak.

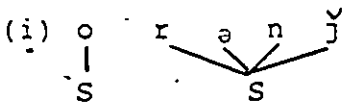
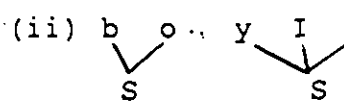
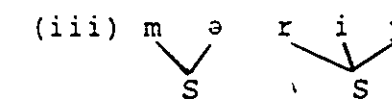
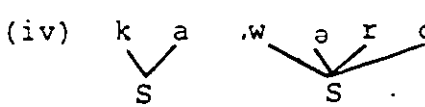
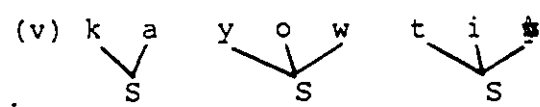
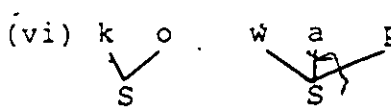
This brings us to a problem concerning the syllabification of glides - a problem that applies to all three, 'y', 'w' and 'r'. How are they syllabified when they are intervocalic? Any algorithm which maximizes onsets, i.e. Kahn's, ours (Clements and Keyser's, Selkirk's, etc.) would assign intervocalic glides as onsets to the second vowel syllable. It's not clear that this is always the correct assignment. And unfortunately the production data - our errors - do not provide us with any solution.

(iv) Excursus: Syllabifying Intervocalic Glides

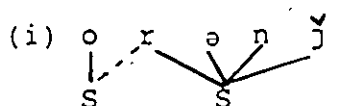
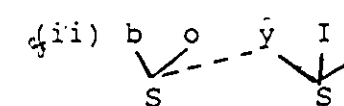
In the syllabification of the words 'orange', 'boyish', 'Murray', 'coward', 'coyote', 'co-op' the intervocalic glides are assigned to the syllable they precede by Kahn's Rule II (and our 2' (a)) since he maximizes onsets and VG sequences are considered to be a sequence of vowel and consonant. Subsequently the glides in 'orange', 'Murray', 'boyish' and 'coward' would be associated with the preceding vowels by Kahn's Rule III. They would be assigned this ambisyllabic status because they are originally associated with an unstressed syllable; the intervocalic glides in 'coyote' and 'co-op' would not be attached to the first syllable because they are associated with a stressed syllable. The structures Kahn's rules would create are illustrated in Figure 7.

Figure 7

Syllabifying Intervocalic Glides by Kahn's Rules I-IIIRules I and II

- (i)  (orange)
- (ii)  (boyish)
- (iii)  (Murray)
- (iv)  (coward)
- (v)  (coyote)
- (vi)  (co-op)

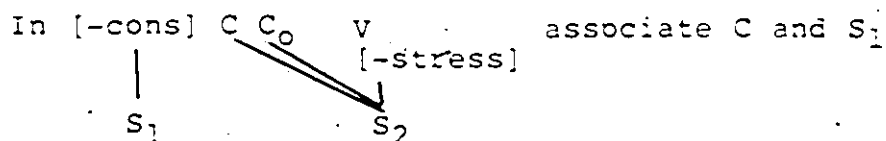
Rule III

- (i)  (ii) 
- (iii)  (iv) 

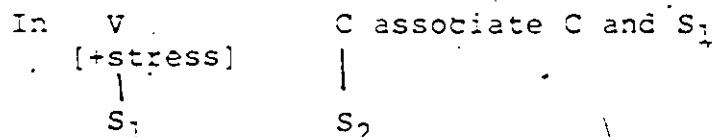
Intuitively, these syllabifications are wrong, particularly in very slow, citation-type speech. (There appears to be a great deal of dialect-related variability in the intuitive syllabification of these forms, however.) The glides, in each of these words, in slow speech, seem to belong to the initial syllable as indicated: [or.ənʃ], [boy.Iʃ], [mər.iy], [kay.ow.tiy], [kaw.ərd], and [kow.ap].

The ambisyllabicity of these segments (and as we shall see in Section 8 below, ambisyllabicity appears to be a phenomenon of forms posterior to the occurrence of segmental errors) in normal speech would appear to be a function of the common rule which joins syllable final consonants to vowel-initial syllables across word boundary.. Kahn's Rule III gets the desired result for only four of these. To change his Rule III such that 'coyote' and 'co-op' could also be resyllabified is possible. We could make the association a function of the '[+stress]' value of the initial syllable rather than a function of the '[-stress]' value of the syllable to which it is joined by Rule II. That is we could have Rule III' rather than Rule III.

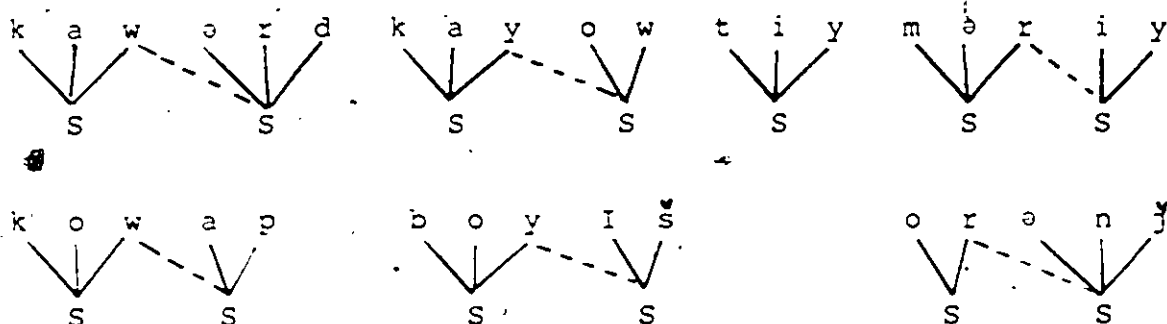
Kahn's Rule III



Kahn's Rule III'



We could get the desired effect this way but it would be wrong. We want the reverse order of application. In super-slow speech the glide belongs to the preceding vowel and in normal speech it associates with the following vowel - the order of application we want is illustrated in Figure 8.

Figure 8Reverse (Intuitive) Order of Glide Association

Kahn's sequence, Rule II and then Rule III, is correct for the intervocalic glides in words such as 'away', 'yo-yo' and 'arrest' where, due to the stress in the final syllable, no change is effected by Rule II.

It has often been suggested that stressed syllables attract (cf. Prince, 1983, Selkirk, 1980, and F&L, 1982). Let us assume that in fact this is what is happening in the words in Figure 7. We could incorporate a sub-rule into Kahn's Rule II which applies before consonants are attached as onsets. It would join glides to stressed preceding vowels as specified in Figure 9. (In Kahn's system this would be impossible since Rule II precedes stress assignment.)

Figure 9Incorporating [Stress] into Kahn's Rule II

Rule II': In $\begin{bmatrix} +\text{syll} \\ +\text{stress} \end{bmatrix}$ $[-\text{cons}]$ associate $[-\text{cons}]$ and S

This rule is ad hoc, however. There is no reason to assume that stressed syllables only attract glides when they precede them, and indeed in words such as 'beattitutde', 'coyote' and 'co-op' the syllables preceding and following the glides are stressed, and the glides appear to belong to both syllables. We could conclude that the rule is a mirror image rule and in such cases Rule II'' would build two lines creating ambisyllabic glides. But this is not correct. Firstly as was discussed above (see Figure 8) in super-slow speech these glides are not ambisyllabic; they belong to the preceding syllable.

And secondly, although it may be the case that stressed syllables attract, it would appear that they do not have this power in general, in super-slow citation-type speech. For instance the first syllable of 'attitude' is stressed but the citation form is 'a.tti.tude' and not 'att.i.tude'. In normal speech the first syllable attracts and this explains the flap allophone of the initial 't'. In citation form this 't' is aspirated however and attached to the second syllable. There is no 'fuzzy boundary' here. Why, in super-slow speech, should glides be attracted to stressed syllables (in words such as 'coyote', etc.) and yet not other consonants?

A plausible answer to this would seem to be that the glides belong to these syllables to begin with, and their appearance here is not a function of the stress value of the syllable. A word like '-forensic' supports this conclusion. In my speech, the super-slow form of this may be either [fær.ən.zIk] or [fæ.rən.zIk] and (perhaps) preferably the former, or so it would seem. It would appear to be generally the case that monomorphemic VCV (and frequently polymorphemic VCV) sequences are syllabified V.CV. That is, the C is attached as onset, not as coda. Why then is the coda attachment preferred in the above forms with intervocalic glides even when the second syllable is stressed. One clear possibility is that these VG sequences are complex segments like the German [tʰ] in 'Zeit' and the English [tʃ] in 'church'. We will not attempt a formal account of this here. In fact as it turned out the problem regarding the syllabification of intervocalic glides did not affect our analysis because we only found two errors (66) (a) - (b) which involved an intervocalic glide and the assignment assigned by Kahn's Rule II (hence ours) was the correct one in this case; in (66) (a) - (b) onset 's' replaces the onsets 'w' and 'r' respectively:

(66) (a) been away been abay

(b) carrot and cabbage cabbot and carriage

The lack of such errors may be significant but it may simply reflect the relative infrequency of word-medial errors in general.

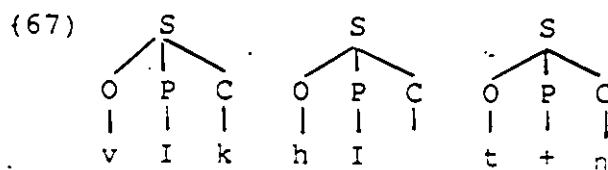
We have no way of knowing whether the production mechanism would support our intuitions regarding the syllable associations of some intervocalic glides. We will have to be satisfied with what we do find from the errors we have. It seems clear that syllable peaks in the production grammar may consist of a vowel and an optional glide, 'y', 'w' or 'r'. These glides are never split away from their vowels. In other words the peak constituent of the production syllable consists of one single position. This matches the peak proposed by Fudge (see Chapter 1, Section 3.) although he would not assign 'r' to the peak. We will see that this inviolable unity is unique to the peak constituent, quite unlike the other syllabic constituents.

(v) Problematical Onset/Coda Errors

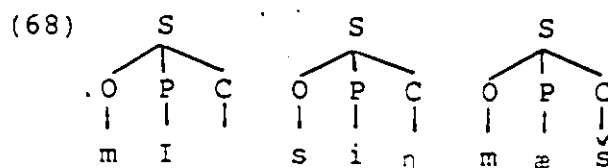
We mentioned in Section (iv) above four possible ways of accounting for the 'exceptional errors': i.e. the 'V+r' errors and the onset/coda errors; errors which apparently did not implicate identical syllabic constituents. The 'V+r' errors turned out not to be 'exceptional'. Our algorithm had neglected to pick out 'r' as a glide and hence the 'r' was

originally assigned to the wrong constituent - the coda rather than the peak. Explaining a number of the onset/coda errors will not involve the algorithm per se but rather its domain of application. We borrowed Kahn's assumption - that syllabification was the first phonological rule in a derivation. The word is the domain of Kahn's Rules I-IV and Rule V applies across word boundaries. It appears that this is incorrect, that the root, or stem should have been chosen as the domain for Rules I and II.

Of our total number of errors involving only onsets and codas (n=436) 352 (81%) involve onset/onset interactions, 59 (14%) involve coda/coda interactions and 25 (6%) appear to implicate constituents of different types, i.e. onsets and codas. These errors do not adhere to our assumption that 'like constituents interact in errors'. Consider now the following ten errors from this 'exceptional' set of 25.



Have you seen Vik hitting? ...vit hitting



I'm missing Mash

...mishing...

- (69)
- | | | |
|------------|------------|------------|
| S
/ \ \ | S
/ \ \ | S
/ \ \ |
| O P C | O P C | O P C |
| | | |
| w a | s i ŋ | w a ʃ |
- I'm washing my watch ...watching...

- (70)
- | | | | | |
|------------|------------|------------|------------|------------|
| S
/ \ \ | S
/ \ \ | S
/ \ \ | S
/ \ \ | S
/ \ \ |
| O P C | O P C | O P C | O P C | O P C |
| | | | | |
| d ə | l i y | s ə n | d ə | l i y t |
- If you have deletion rules you can only delete ...[dəliyʃ] deleesh

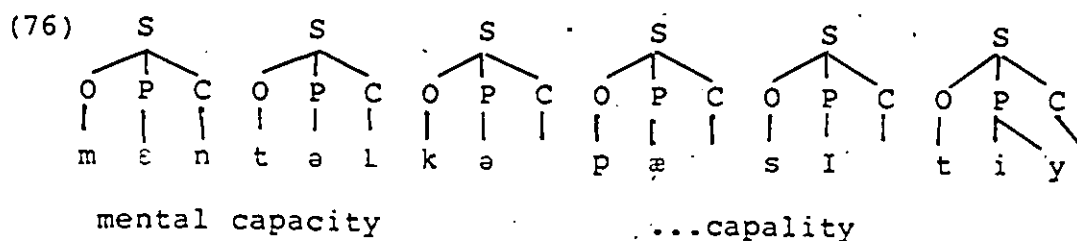
- (71) is ambiguous (see below) involving underlying V's ...[viyvz]

- (72)
- | | | |
|------------|------------|------------|
| S
/ \ \ | S
/ \ \ | S
/ \ \ |
| O P C | O P C | O P C |
| | | |
| h e lp | l æ | f I ŋ |
- I can't help laughing helf lapping

- (73)
- | | | |
|------------|------------|------------|
| S
/ \ \ | S
/ \ \ | S
/ \ \ |
| O P C | O P C | O P C |
| | | |
| r o w | l i ŋ | p i n |
- rolling pin rolling pill

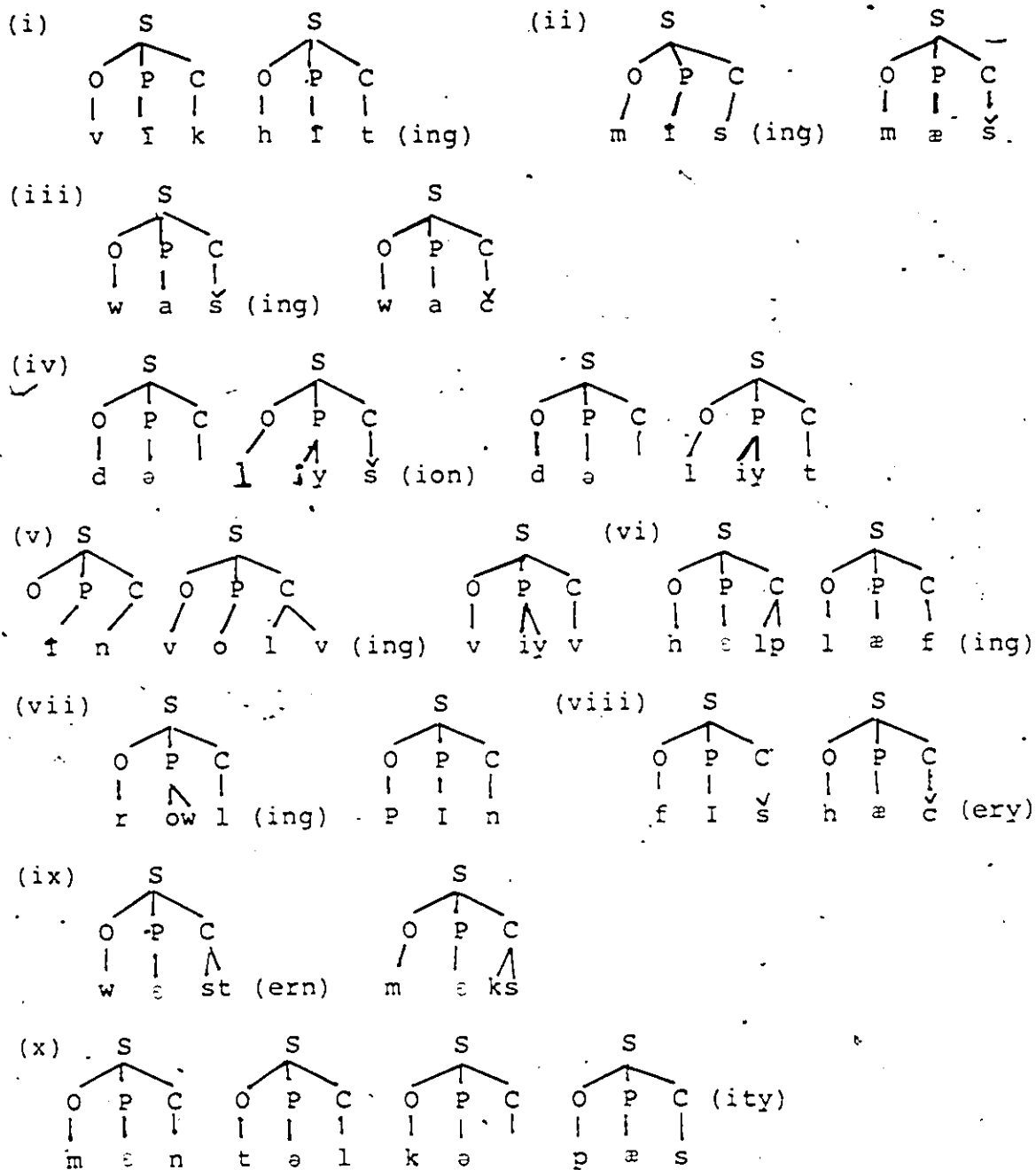
- (74)
- | | | | |
|------------|------------|------------|------------|
| S
/ \ \ | S
/ \ \ | S
/ \ \ | S
/ \ \ |
| O P C | O P C | O P C | O P C |
| | | | |
| f ɪ ʃ | h æ | ʃ ə | r i y |
- fish hatchery [fiʃ] hatchery

- (75)
- | | | | | |
|------------|------------|------------|------------|------------|
| S
/ \ \ | S
/ \ \ | S
/ \ \ | S
/ \ \ | S
/ \ \ |
| O P C | O P C | O P C | O P C | O P C |
| | | | | |
| w e | st æ r n | m e k | s i | k o w |
- western Mexico wexern Mestico



It is clear from the diagrams that in each of these we have codas and onsets interacting. In (75) we have the peculiar interaction of segments from two syllables, 'k' and 's', exchanging with the onset 'st' of one syllable. What is also interesting about these errors, is the fact that none of them involves only monomorphemic words. And more importantly, the problem constituent is, in each case, the final segment(s) of the root - which, when attached to the inflectional and derivational suffix(es) becomes an onset.²⁵ If the domain of syllabification were the root, and not the entire word, all of these errors would involve two identical constituents - codas, as illustrated in Figure 10.

Figure 10

Root Syllabification of c/o Errors

Error number (71) could have its source in any one of the three V's so let us ignore it, but the others are all clear. They involve the anticipation, perseveration or exchange of coda constituents. In addition, we don't find any root final consonants which interact with other onsets. That is we do not have errors like (77):

(77) I'm missing Barney Miller I'm mibbing...

As will be discussed in detail in Chapter 3, the level at which sub-morphemic segmental errors occur appears to be one in which roots are being processed, more particularly, Garrett argues, when roots are being assigned to their phrasal positions (cf. Garrett, 1980, 1982). It is well known that roots and affixes are treated differently by the programmer and these errors confirm this. That is, the errors were apparent counter examples to the claim that when these errors occurred the programmer could only mis-select among identical syllable constituents; but they turn out to be support for the claim. Moreover, they are additional support for Garrett's hypothesis that such errors occur when roots are being manipulated, not full words. Two more of the apparent c/o errors may involve what Garrett (1975) refers to as 'pseudo' morphs. They are phonologically like bound morphemes but are not in fact morphemes. For instance in the error

(78) I carved a pumpkin

I pumped a carven

the pseudo morph 'n' is left behind as is the past tense morpheme when the 'roots' are exchanged.²⁶ Errors (79) and (80) may involve such pseudo morphs.

(79) cup of coffee

cuff... (y?)

(80) Jack Rabbit

racket (it?)

That is, if the processor analyses 'coffee' and 'rabbit' as containing 'roots', 'coff' and 'rabb' these two errors could involve the codas 'p', 'f' and 'b', 'k' respectively. There are other possible analyses of both of these errors, of course.²⁷ Of the remaining thirteen onset/coda errors 4 involve the word-final syllabic sonorants.

(81) strange reasons

strain regions

(82) rank order

rand orker

(83) ocean perch

ocean perch

(84) I'll call it spasm

cause it spasm

They will be considered in detail in our sub-section on blends below. We shall see that these errors can also be analysed as coda/coda errors.

Of the last nine coda/onset errors seven are word-internal errors. These errors will be discussed below in our sub-section on such errors. Our reasons for

separating out these errors from the between-word errors will also be explained. We are left with two inexplicable onset/coda anticipations: one of which may be an onset substitution whose source is not the coda in the following word,²⁸

(85) particle shift

farticle shift

and another (86)

(86) Don't make a fuss
over nothing

...fun...

which may involve the intrusion of the complex verb 'make fun of', among other possibilities.

How does this change in our domain of the syllabification rules affect the analysis of the rest of our data? That is, will errors that involved identical constituents using the word as domain turn out to be exceptions using the root as domain? In fact the new domain holds up. First of all, so many of the onset errors involve word-initial onsets and so many of the coda errors involve root-final codas such as (87)

(87) steak or chops

steak or chawks

that we found very few errors that required re-analysis. The two analyses, word domain, root domain, for three of these errors (88) - (90) are illustrated in Figure 11.

(88) prefixes...According
to Vicky

...vicksy

(89) ...inning with a
tender elbow

inding...

(90) dressing taster

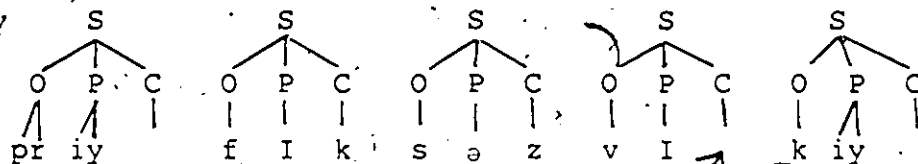
dressting taser

Figure 11

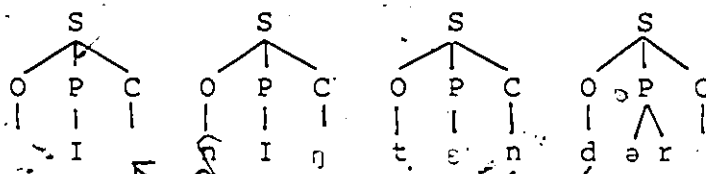
Analysis of (88) - (90) Using Root and Word Domains

(a) Word Domain

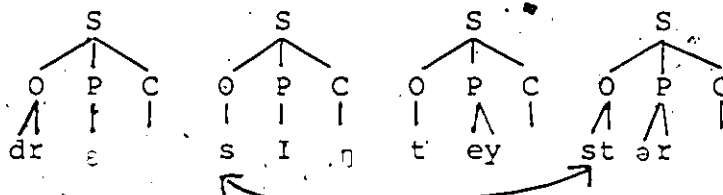
(i) (prefixes)
Vicky



(ii) (inning)
tender

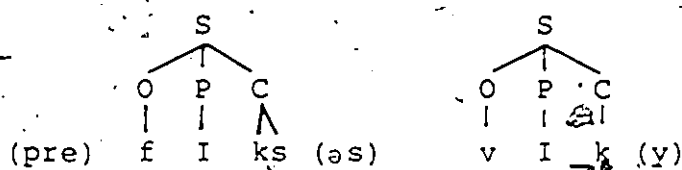


(iii) (dressing)
taster

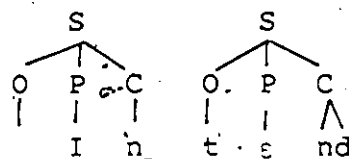


(b) Root Domain

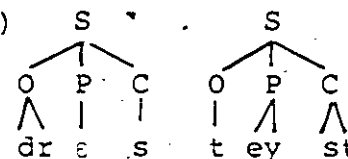
(i)



(ii)



(iii)



As can be seen from Figure 11 the errors (88) and (89) required a conjunction-of-errors analysis using the word as domain. For instance in (89) the coda 'n' from 'tender' had to move into the empty coda position in the initial syllable of 'i.nning' and the 'd' from the onset position in tender had to replace the onset 'n' in the second syllable of inning. Using the root as domain we have a single movement of a coda (or coda part). The root analysis is thus more appealing, being simpler.

Using the root as the domain of syllabification and separating out the 45 word-internal errors from our set of 559, 95% of our segmental errors may be accounted for as the mis-selection from among the single syllable constituents onset, peak and coda as illustrated in Table 4.²⁹

Table 4Syllabic Position of Between-Word Segmental Errors:Root as Domain

Syllabic Position	Frequency (Number)	Frequency (Percent)
onset/onset_	312	61%
coda/coda	78	15%
peak/peak	100	19%
other	24	5%
TOTAL	514	100%

You will recall that we felt that in order to determine whether the processor was processing segments or syllabic constituents the errors which were crucial were those which involved at least one consonant which was word-medial. The fact that errors never included interactions between vowels and consonants did not necessarily implicate syllabic nucleus versus margins; and the high proportion of errors involving word-initial and word-final positions were not unambiguously interpretable as errors involving syllabic onsets and codas. They could be telling us something about lexical structure rather than syllabic structure. Because of this we analysed all those errors which included a word-medial consonant

separately from those which did not. See Table 2 above. We had 116 such errors 95 (81%) of which involved onset/onset or coda/coda interactions but a relatively high percentage 16% (n=18) were onset/coda errors. Of these 18, two were within-word errors:

- | | |
|--------------|---------|
| (91) reserve | reverse |
| (92) medal | meled |

The rest were the onset/coda interactions in (67) - (84) which we have just discussed. There is good reason to believe that 14 of these are in fact coda/coda errors (although our arguments regarding (81) - (84) are yet to come in our section on blends below). The three errors classed as 'other' errors were word internal

- | | |
|---------------|---------|
| (93) shilling | shingle |
| (94) whisper | whipser |
| (95) fixed | sixed |

Ignoring the 5 word internal errors for now we see that there are 111 errors that involve at least one word medial onset or coda. Of these 84 (76%) are onset/onset errors, 25 (23%) are coda/coda errors and 2 (2%), the errors in (79) - (80) are possibly coda/coda errors but may not be. These facts are found in Table 5.

Table 5Syllabic Position of Medial Consonant Errors:Between Word Errors

Syllabic Position	Frequency (Number)	Frequency (Percent)
onset/onset	84	76%
coda/coda	25	23%
other	2	2%
TOTAL	111	101%

N.B. (i) Errors involve at least one word medial consonant.

(ii) 101% due to rounding off.

The evidence seems⁴ clear. Segments that interact in between-word errors always belong to the same syllabic position. This fact is totally inexplicable if the programmer is processing unstructured segments at the time that these errors occur. On the other hand, if the processor is selecting syllabic constituents, this constraint on between-word errors is explained. Moreover, errors which involve sequences of segments always involve segments belonging to one constituent, onset, peak or coda, except for the few errors (n=19) involving sequences of onset-peaks and

peak-codas. (These errors will be discussed below in our section on rhymes.) The assumption that errors will involve constituents of the same type is supported; in this case the types involved are onset, peak and coda as defined by our algorithm (see Figure 1) and whose domain is the root. PSE (Figure 3 above) mirrors what the processor is doing when segmental errors occur, although only approximately. Peak errors always involve whole peaks so PSE accurately depicts these errors but onset and coda errors are not (so 'holistic'. PSE will generate errors like

- | | | | |
|------|-----|-------------------------|------------------|
| (96) | (a) | squeeze fresh oranges | freeze squesh... |
| | (b) | central choir | kwentral sire |
| (97) | (a) | in the next ten minutes | nen text... |
| | (b) | check the set | chet...seck |

but what of errors like (98) - (99)?

- | | | | |
|------|-----|----------------|----------------|
| (98) | (a) | flame broils | frame bloils |
| | (b) | deep structure | steep dructure |
| (99) | (a) | pinch hit | pitch hint |
| | (b) | sink a ship | simp a shick |

Onsets and codas appear to include substructure - they are not indivisible. The errors in (98) - (99) involve onset and coda positions respectively, but only parts of these positions. PSE can only generate errors which involve the

entire constituent, not those which involve sub-constituent parts. We would like to be able to account for these errors as well.

We will attempt to do so below but first we would like to consider the errors which involve what we have regarded as sequences of two constituents onset-peak, and peak-coda, as exemplified in (100) - (101):

(100) pussy cat	cassy put
(101) heap of junk	hunk of jeep

Any error which requires a 'conjunction' analysis is less than satisfactory. That is, we don't expect sequences of segments to be involved in errors unless they are constituents of some sort. We have not assigned constituent status to onset-peaks, or to peak-codas. Therefore these errors are startling. Why should the processor randomly select these particular sequences in these particular words and just happen by chance to exchange each one with the other. That is we have analysed the errors as a conjunction of onset/onset plus peak/peak (100), (or peak/peak plus coda/coda (101)) where the processor just randomly selects these particular constituents and they happen to be side by side.

Peak-coda sequences are often analysed as the constituent 'rhyme' and we know of one study (Kozhevnikov and Chistovich, 1965) which concludes that onset-peaks are constituents. In the following section we will consider whether our syllable should be enriched, i.e. whether it should include more internal structure in order to account for these onset-peak and peak-coda errors.

7. The Rhyme: Where Is It?

As noted above we found eight (nine, if we include the complex error in (vii)) onset-peak/onset-peak errors and eleven peak-coda/peak-coda errors. These errors are listed in Tables 6 - 7 respectively.

Table 6

Onset-Peak/Onset-Peak Errors

(i) edge of his wit	widge of his et
(ii) stress and pitch	piss and stretch
(iii) lost and found	forust and lond
(iv) stick in the mud	smuck in the tid
(v) pussy cat	cassy put
(vi) Dixie Blaff	Blaxie Deff
(vii) loose leaf note book	[nows niyf luwf] book
(viii) Woodie Allen	Eddie Woolen
(ix) Thorndike Guthrie and Estes	...Ethrie and Gusters

Table 7.Peak-Coda/Peak-Coda Errors

(i) Salt Lake City	Sake Lalt City
(ii) he will lead the way	...lay the weed
(iii) washing his feet in the sink	...fink in the seat
(iv) I don't care which	...kich where
(v) paint cans	pan kaints
(vi) spice racks	spack rices
(vii) monkey's uncle	monkle's unkey
(viii) basketball court	basket bore callt
(ix) books for Vietnam	bomb for Viet Nooks
(x) the juice is still on the table Is that enough?	...enuice?
(xi) Yankee doodle dandy	Yankle...

The first characteristic of the errors in Tables 6 and 7 which we would like to call to your attention is the fact that with the exception of only a few ((vi) and (ix) in Table 6, and (viii), (ix) and (x) in Table 7) they are all simple exchanges. There is a very common type of exchange error in the literature - what Garrett refers to as 'stranding errors'. Stranding errors involve the exchange of morphemes within phrases, irrespective of their grammatical class, membership (see Chapter 3). These morpheme exchanges are

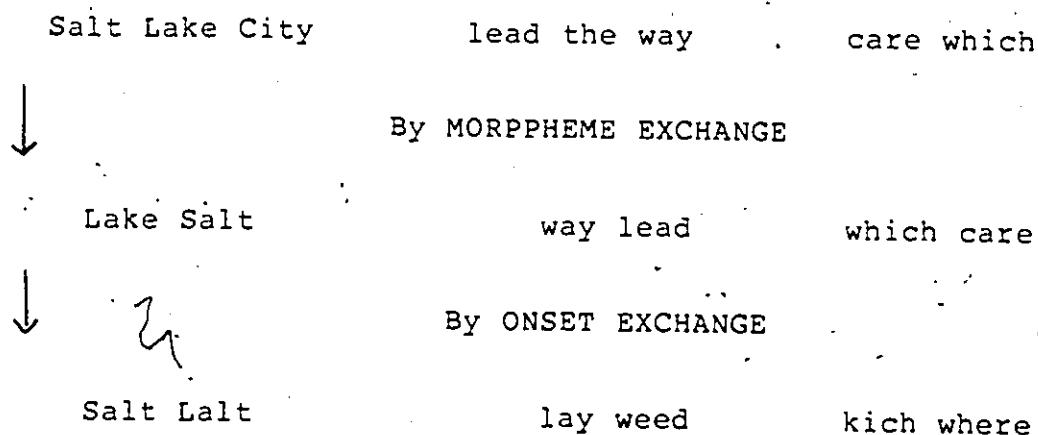
referred to as 'stranding errors' because they leave grammatical morphemes stranded when there is one. For instance (102) (a) - (c).

- | | |
|--------------------------|------------------|
| (102) (a) the guy's name | ...name's guy |
| (b) I cooked a roast | I roasted a cook |
| (c) his team rested | his rest teamed |

Consider first the errors (i) - (iv) in Table 7. All of these could be analysed as a morpheme exchange followed by the most common type of error, the onset/onset error. This sequence is illustrated in Figure 12, for three of the errors. (The reader may verify the others himself.)

Figure 12

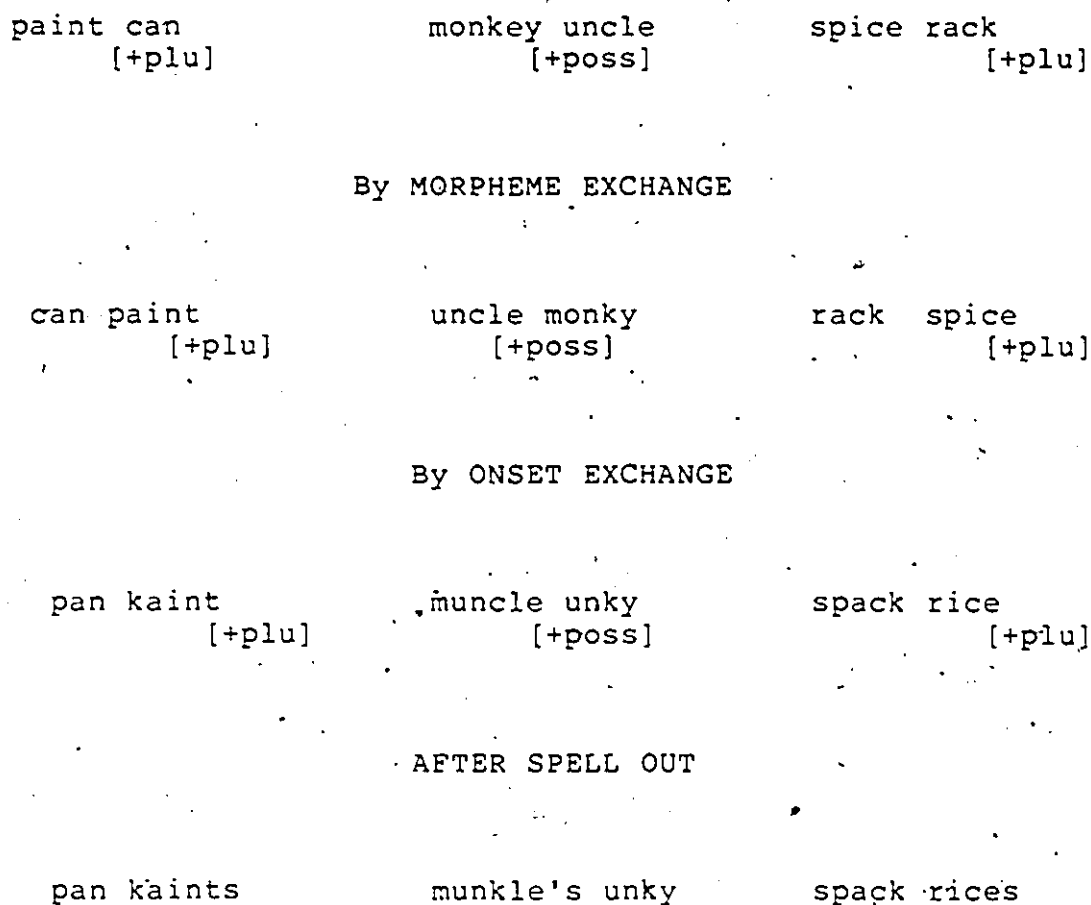
'Peak-Coda' Errors



Consider now the errors (v) - (viii) in Table 7. These appear to be peak-coda errors but part of the coda is not involved. Thus if the final element of these codas is part

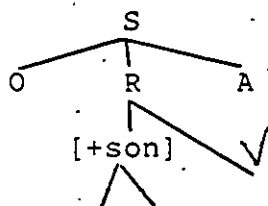
of the rhyme these errors do not involve a rhyme exchange - so they would be evidence against a rhyme constituent, not for it. We notice however that in (v) to (vii) the stranded coda parts are in fact morphemes, plural and possessive. We appear to have 'stranding errors' that involve parts of rhymes rather than morphemes. Clearly a simpler analysis is that these too, are morpheme exchanges (Garrett's stranding errors) with an onset exchange superimposed on them, followed by spell-out.³⁰ (In Chapter 3 we will show how this sequence of events can be predicted by Garrett's model with minimal changes to the model.) The sequence is illustrated in Figure 13.

Figure 13
'Peak-Coda' Stranding Errors



The error in (viii) does not strand a morpheme however.
 Is this evidence for Halle and Vergnaud's [+son] constituent
 illustrated in Figure 14?

Figure 14



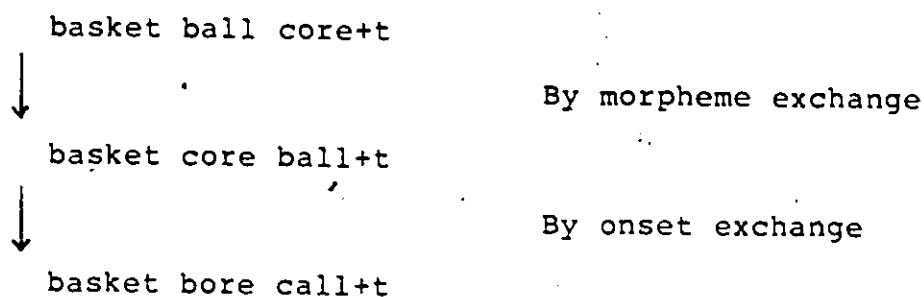
In Halle and Vergnaud's syllable glides and the liquids and nasals go in the right branch of this constituent. That is, the 'or' of 'court' and the 'all' of 'ball' would both be dominated by this node. We argued above that 'V+r' belonged to the peak with other vowel-glide complexes. Do the other sonorants perhaps belong to this constituent as suggested by the Halle and Vergnaud syllable? The error in (viii) suggests that the processor might treat 'V+r' and 'V+l' as equivalent at some level. It would appear not, however. There are a number of indications that 'r' is treated differently from 'l' and 'n' by the production mechanism.

Firstly, post-vocalic 'r' does not substitute for consonants.³¹ Post-vocalic 'n' and 'l' do. For instance in (103) 'n' exchanges with 't' and 'l' with 'c'.

- | | | |
|-----------|-----------------|-----------------|
| (103) (a) | pots and pans | pons and pats |
| (b) | mental capacity | mental capality |

Secondly, post-vocalic nasals and liquids may be involved in errors without the preceding vowel, as in (103) above; and vowels preceding nasals or 'l' may move, leaving 'n' and 'l' intact. For instance (104):

- | | | |
|-----------|---------------------|---------------------|
| (104) (a) | Wang's bibliography | Wing's babliography |
| (b) | living with Jean | leaving with Jean |

Figure 15A Pseudq-Morph Account of (viii)

The error in (ix) is similar to the morpheme exchanges in (i) - (viii) except that there is no morpheme stranded. The plural morpheme moves with 'book'. There are errors in which the morphemes are not stranded. For instance (107):

(107) Well you can cut trees ...rain in the trees
 in the rain

Garrett locates these word errors at another level in his model (about which more later in Chapter 3). The error in (ix) then is analogous to (i) - (viii): books and Nam exchange and then the onsets 'b' and 'n'.

The final two peak-coda errors in Table 7, (x) and (xi), can not be explained as morpheme exchanges followed by onset exchanges - (x) has the characteristics of blends and hapologies. Segmental anticipations, perseverations and exchanges are almost always within-phrase errors (see Chapter 3). This 'perseveration' uncharacteristically crosses a

sentence boundary; it is more likely a blend or hapology of 'enough' and 'juice'. The error in (x) may involve a pseudo morph 'le'. There is a stranding error which suggests this possibility, namely (108):

(108) The single biggest biggle singest...
 problem

There have been numerous good arguments in support of a rhyme constituent in the literature (see Chapter 1). If the production mechanism recognizes such a constituent it does so at a level different from the level at which it recognizes onsets, peaks and codas. It seems clear that these latter three constituents constrain which segments may interact in errors and which submorphemic sequences of segments may be involved as a unit in errors. The rhyme does not play this role. It does not interact with other rhymes; the few errors which could be analysed as rhyme errors have a better explanation because they can be analysed as complex errors involving constituents which regularly interact in errors. As mentioned above, moveability is not a criterion for viability of a processing unit. The syllable is not moved around and yet it must be recognized as a unit in a processing model because onsets, peaks and codas are defined in terms of this unit. The Rhyme could be like the syllable in this respect. However we do not know of any constituent whose definition is dependent on the rhyme. So far as we can

determine, from segmental speech errors, there is no evidence which supports a rhyme constituent.

We come now to those few errors in Table 6 (n=8) which seem to involve an onset-coda constituent. Kozhevnikov and Chistovich have argued that articulatory syllables break after vowels. All consonants following a vowel belong to the onset of the following syllable which ends at the vowel. Eight out of 559 errors would hardly be strong support for such a unit even if these errors had no other obvious explanation. However they can be accounted for without recognizing them as a constituent of some sort. Consider first the errors (i) - (iii). All of these can be analysed as morpheme exchanges with a subsequent exchange of codas as illustrated in Figure 16.

Figure 16

Onset/Peak Errors

edge of his wit	stress and pitch	lost and found
↓		
By MORPHEME EXCHANGE		
wit of his edge	pitch and stress	found and lost
↓		
By CODA EXCHANGE		
widge of his et	piss and stretch	foust and loud

Now consider the error in (iv). It also could be a morpheme exchange of 'mud' and 'stick' followed by both a coda exchange of 'd' and 'k' and an onset movement of 's'. Such movements of 's' are found; for instance (109):

(109) paint in the studio spaint in the tudio

The errors (v) and (vi) appear to involve roots and the morpheme 'y' is stranded as illustrated in Figure 17.

Figure 17

Onset/Peak Stranding Errors

puss +y cat

Dex +y

Blaff



By MORPPHEME EXCHANGE

cat +y puss

Blaff +y

Dex



By CODA EXCHANGE

cas +y put

Blax +y

Deff

(cassy put)

(Blaxie Deff)

This stranding of the morpheme 'y' is attested in the error (110):

(110) is the kitty purring?

...purry kitting

The error in (vii) is very complex³² but it is nevertheless made up of all commonly occurring errors as indicated in Figure 18.

Figure 18.

loose leaf note -- [nows] [niyf] [luwf]

- (a) by morpheme exchange: note leaf loose
- (b) by coda anticipation: nows leaf loose
- (c) by onset perseveration: nows niyf loose
- (d) by coda perseveration: nows niyf luwf

It appears then that no interaction of the onset-peaks '[luw]' and '[now]' is required in the analysis of this error. Their analysis is complex but plausible. The error, after all is complex.

The final two onset-peak errors (viii) - (ix), we will not attempt to analyse although they appear to involve morphemes and pseudo-morphs.

We have given a plausible account of at least sixteen of the 19 errors which seemed to involve onset-peak, or peak-coda constituents. We have analysed them as compound errors in all cases exemplifying common error types in the literature. In all cases we find morpheme exchanges compounded by syllabic constituent errors. In most cases the

errors were exchanges of onsets or codas although we also found a couple of perseverations and anticipations.

In our original analysis whose results were tabulated in Table 1 we found a number of errors which were not analysable as the replacement of identical syllable constituents types for one another. A closer analysis revealed that 'V+r' errors were actually peak/peak substitutions and a number of the coda/onset interactions were in fact coda/coda errors. In the one case we argued 'r' was a glide and in the latter we argued that the syllabification domain was the root. Now we have argued that the onset-peak and peak-coda errors are also errors in which identical syllabic constituents replace each other - but that they occur after a morpheme exchange has occurred. Some of the analyses of coda/onset errors, and two of the onset-peak analyses were not absolutely convincing. Nevertheless we may conclude with some conviction that submorphemic segmental errors are in fact syllabic constituent errors and that the constituents are those defined by the syllabification rules in Figure 1 using the root as domain. This account applies only to between-word segmental errors; we have not yet considered word-internal errors. But it accounts for about 98% of the between-word errors. The results including a redistribution of the peak-coda, and onset-peak errors are listed in Table 8.

Table 8Syllabic Position of Between-Word Errors

Syllabic Position	Frequency (Number)	Frequency (Percent)
onset/onset	321	63%
coda/coda	84	16%
peak/peak	100	19%
other	8 ³³	2%
TOTAL	513	100%

PSE will not generate all of these errors. It can only generate those which include entire onsets, peaks or codas. As we shall see below there are 63 errors which unambiguously involve parts of onsets and 23 which involve parts of codas. That is of the 513 between-word errors PSE will only generate 427 (=83%) of them. We will consider below how the 'part' errors in onsets and codas might be accounted for but first we would like to consider whether Kahn's hypothesis that there are ambisyllabic segments receives any support from the speech error data.

That is, is there any evidence that the programmer distinguishes ambisyllabic segments from unisyllabic

segments? Does the processor distinguish a segment which is both onset and coda from others?

8. The Ambisyllabic Segment Isn't

Since our syllabification rules, derived from Kahn's slow speech rules, have identified 98% of segmental errors as the mis-selection of segments from identical constituents we don't require Kahn's normal speech rules. That is, his slow speech boundaries are the right boundaries for 98% of the errors. What of the 2% - the 'other' errors? In fact his normal speech rules cannot help us in explaining our 'other' errors, which are listed below in Table 9.

Table 9³⁴'Other' Errors

(i) Have you given your paper a title?	...titer
(ii) mirror image	mirage immer [aj]
(iii) Yankee Doodle Dandy	Yankle...
(iv) It was Gretsky that saved it	it was gravy--it was Gretsky...
(v) I want lessons with Greg	I want legs--lessons with Greg
(vi) That happened tonight in Hamilton	...hammled [hæməld] --happened...
(vii) particle shift	farticle shift
(viii) Don't make a fuss over nothing	...fun...
(ix) Woodie Allen	Eddie Woolen
(x) Thorndike Guthrie and Estes	...Ethrie and Gusties
(xi) People would come in with their prepared talk on some topic	...talkit on some topic

The errors (i) - (iii) are a problem for us because they involve the replacement of peaks by codas. They will be discussed below in our section on Blends. Kahn's Rule V could be used to explain (viii) but the others remain inexplicable because Kahn's normal speech rules could only help in accounting for errors which involved an exchange of onsets and codas.

Recall that Kahn's Rules III - V create the ambisyllabic segments which Kahn believes characterize normal fluent speech (see Appendix to Chapter 1). Rule III attaches syllable-initial consonants to preceding syllables if the syllable they belong to is unstressed. Rule IV attaches syllable final segments to the following syllables if this following syllable is unstressed and if the resulting cluster would not violate universal phonotactic constraints. Rules III - IV have the word as their domain, Rule V applies across word boundaries attaching syllable final segments to vowel initial words.

We had a number of errors after our first analysis of the data (which used the word as domain) which seemed to be evidence against our claim that the processor could only mis-select segments from identical syllabic constituents. In particular there were a number of errors (n=18) which apparently involved interactions of coda and onset constituents. We argued that the root should be the syllabification domain, not the word, and under this change fourteen of the eighteen errors became explicable as coda/coda interactions. It appears, however, that this is not the only solution to these apparent exceptions to our claim. Indeed if we maintain with Kahn that the word is the correct domain and that the syllabification rules should

include Kahn's normal speech rules, we can account for seventeen of the eighteen errors. Since errors are characteristic of normal speech this would seem to be a better account particularly as it can account more simply for more of the errors.³⁵

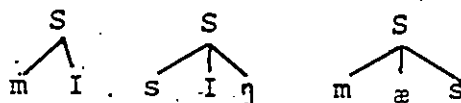
To see how his rules would allow the interpretation of these errors as exchanges of segments from identical constituents we will syllabify a couple of these exceptional errors using his normal speech rules. The problematical errors discussed above, we abbreviate here in Table 10, focussing on the 'problem' segments. Examples of syllabification by slow speech rules I and II and normal speech rules III - V are given in Figure 19.

Table 10Errors Abbreviated

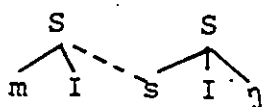
(i) Vik <u>h</u> itting	Vit
(ii) miss <u>i</u> ng Mash	mishing
(iii) washing wat <u>c</u> h	wash
(iv) deletion delet <u>e</u>	[dɛliyʃ]
(v) involving V' <u>s</u>	[viyvz]
(vi) ment <u>a</u> l capacity	capality
(vii) fish <u>h</u> atchery	[fɪʃ]
(viii) help <u>l</u> aughing	helf lapping
(ix) west <u>e</u> rn Mex <u>i</u> co	wexern Mestico
(x) rolling pin <u>e</u>	pill
(xi) strange <u>r</u> easons	strain regions
(xii) rank <u>o</u> rd <u>e</u> r	rand orker
(xiii) o <u>c</u> ean perch	[oʃən]
(xiv) call <u>i</u> it spasm	cause
(xv) cup <u>o</u> f coffee	cuff
(xvi) Jack rabb <u>i</u> t	...racket
(xvii) <u>p</u> article shift	...farticle
(xviii) f <u>u</u> ss over nothing	fun

Figure 19Kahn's Normal and Slow Speech Rules

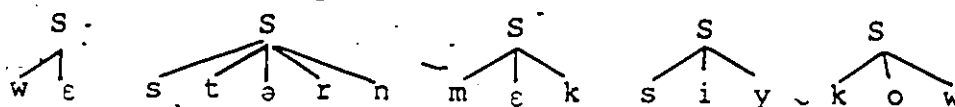
Rule I - II (missing mash):



Rule III:



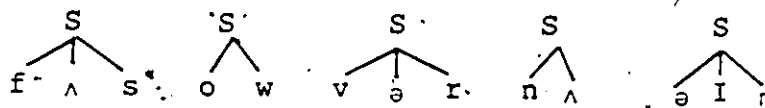
Rules I - II (western Mexico):



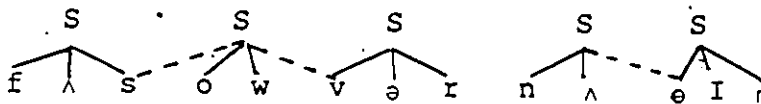
Rules III - IV:



Rules I - II (fuss over nothing):



Rules III - V:



The reader can verify that all of the errors in Table 11 except (xvii) can be analysed as involving identical syllabic constituents since all of the errors involve ambisyllabic segments as defined by Rules III - V. That is we could argue

that segmental errors involved the replacement of segments from identical syllabic constituents for one another, but the use of all of Kahn's rules is required. Why don't we argue for this analysis? There are a number of reasons.

First, the only generally accepted resyllabification rule is Rule V (it is the only one accepted by Kiparsky, for instance (Kiparsky, 1979)) and yet we do not find errors like (111):

(111) steak or chops steak or cops

where Rule V would give the final 'k' in 'steak' an onset role before 'or'. The lack of such errors is surprising given the frequency of onset/onset errors in general. However this rule applies across word boundary and we could explain this lack by limiting our domain to the word. Rule IV creates onsets. We only have three errors which could be analysed as onset/onset errors after the application of this rule. For instance in (112)

(112) Monday Wednesday Monday...

the '[z]' in Wednesday could be made ambisyllabic by Rule IV and this error could be considered the insertion of part of an onset into part of an onset rather than the analysis with codas. (See also the second error in Figure 19). That is, we do not find any errors involving consonants that require

such an analysis and only three out of over 400 errors which could be given such an analysis. Thus the ambisyllabic onsets created by Rules IV and V receive no support. What of the ambisyllabic codas created by Rule III. They apparently are viable, being able to account for fifteen of the eighteen onset/coda errors above. Should we therefore syllabify by Rules I - III using the word as domain? It appears not.

To demonstrate the need for ambisyllabic segments, you must demonstrate that their behaviour is unique to them. For instance Kahn distinguishes among syllable-initial, syllable-final and ambisyllabic segments because the three have distinct patterns phonologically - unique to each of them. We don't find this with regard to ambisyllabic segments in errors - sometimes the ambisyllabic segments behave like codas (15 of the errors in Table 10) and sometimes they behave like onsets. We found five between-word errors which we analysed as onset/onset errors which involved ambisyllabic onsets as defined by Kahn's Rule III. For instance (113) - (115).

- | | |
|--|------------------------|
| (113) <u>m</u> ushroom <u>c</u> asserole | casherole |
| (114) Peterborough leading
<u>O</u> shawa | Peshe--Peterborough... |
| (115) Angela Davis | Andela Javis |

In (113) - (114) both source and target are ambisyllabic. But in (115) only the 'g' in 'Angela' is ambisyllabic by Rule III (if the vowel is nasalized and the 'n' deleted). That is, we have an ambisyllabic segment ('g') interacting with a unisyllabic segment, 'd'.

That is we have an ambisyllabic segment acting like a unisyllabic onset (115), or like unisyllabic codas (errors (xv) and (xvi) among others in Table 11), or like other ambisyllabic segments, (113) - (114). Ambisyllabic segments, one might conclude, are therefore truly ambiguous. But they are not. An ambisyllabic segment behaves like a coda if it is root final - otherwise it is indistinguishable from an onset. The behaviour of ambisyllabic segments is not unique to them. One can predict precisely when an ambisyllabic segment will function like a coda and when it will not. Therefore we must conclude that if the processor distinguishes ambisyllabic from unisyllabic segments, it must do so at a point downstream from the level at which segmental errors occur. Our hypothesis that roots are syllabified by the rules in Figure 1, which do not create ambisyllabic segments, is all that is required. We need not add rules on the model of Kahn's Rules III - V.

As mentioned, our claim that segmental errors involve segments from identical syllabic constituents accounts for 98% of the errors but we can only generate 83% of them by

PSE, because PSE will not generate those errors which involve parts of onsets or codas. We will now consider these two constituents in detail and attempt to remedy this.

9. Onsets: Internal Structure

(i) Data Selected For Analysis

Onset errors which involve entire onsets, whether they be multi- or uni-segmental, can tell us nothing about the internal structure of onsets. They are only evidence for the unit itself. It is clear that errors which involve sequences of segments are interpretable as errors which involve constituents of some sort, that is, the processor does not randomly move sequences of segments. Segment sequences which move are analysable as words, phrases, morphemes or, as we have just seen onsets, codas and peaks. An error such as (116). (a)

(116) (a) central choir [kwentrəl] sire

would be truly amazing if it involved the movement of '[k]', of '[w]', and of '[s]', which randomly landed in exactly the correct onset positions for each of them. It becomes less remarkable if it involves the simple exchange of two identical constituents - in this case onsets. If onsets include internal hierarchical structure like that proposed by

C&F (see Chapter 1, Section 3) for instance, then we might expect to find the sub-constituents so defined interacting in errors. We will therefore look at all errors which involve parts of onset clusters; but only those which are unambiguously 'part' errors. Certain cluster errors are ambiguous between a 'whole' and a 'part' interpretation. For instance, the errors in (117).

- | | |
|---|------------------------|
| (117) (a) fewer inflections
'l' or 'f' and 'fl'? | flewer infections |
| (b) Fred Flinstone
'l' and 'r' or 'fr'
and 'fl'? | Fled Frintstone |
| (c) along came a spider
and sat down beside
her
's' and 'sp' or 'p'? | ...sider...bespide her |

The errors in (118) - (120), on the other hand, are unambiguously 'part' errors; either the source of the error is part of a cluster (118) or only part of a target cluster is affected (119), or both, (120).

- | | |
|---|--------------------------------------|
| (118) (a) the <u>third</u> and
surviving brother | the sird--the bird
and--the third |
| (b) <u>streak</u> of bad luck | ...brad... |
| (119) (a) the flame <u>trees</u> of
<u>Thika</u> | ...threes of Thika |
| (b) grapefruit <u>juice</u> | grapejruit... |
| (120) (a) <u>flame</u> <u>broils</u> | frame bloils |
| (b) alfalfa <u>sprouts</u> | alspalfa frouts |

There are other errors which we have designated as 'part' errors - the remaining onset errors which PSE cannot account for. These are the errors which appear to juxtapose entire onsets. That is, we don't have the mis-selection of one onset constituent for another, but rather the juxtaposition of two onsets. For instance in (121) the entire onset 'sh' of 'short' is perseverated but it does not replace an onset (segmentally filled or empty) but rather attaches to it.

(121) a short lady

a short shlady

This type of error appears to be turning a whole onset into a part of an onset. These errors, then, we have also designated as 'part' errors and will be added to the data analysed.

Designating 'part' errors this way we found 149 errors that involved at least one cluster, 63 of which were unambiguously 'part' errors. The others were 'whole' errors or ambiguously 'whole' errors as indicated in Table 11.

Table 11
Errors Involving Onset Clusters

Whole	Part	Ambiguous	Total
41	63	45	149
(28%)	(42%)	(30%)	(100%)

- N.B.** (i) 'whole' examples - squeeze freeze squesh
 fresh
- (ii) 'part' examples - flame broil frame bloil
- (iii) 'ambiguous' example - trees prees are tritty
 are pretty

Thus we had 63 errors on which to base our analysis and which could not be generated by PSE. We noted the type of segment and string position of the segments involved.

(ii) Results

Of the 63 'part' onset errors, thirty-six (36) involve the movement of the non-nasal sonorants 'l', 'r', and 'w', ten (10) of them involve 's' or 'sh' movement and 13 involve single obstruents or nasals. Two errors involved a SC sequence of a three-member cluster. (There were ten errors which involved three member clusters; three, involved the entire cluster, five involved single segments and two involved two of the three segments - the SC in both cases.)

Finally, two errors involved movement and deletion. This information is found in Table 12.

Table 12
Part Errors

Segment Type	Frequency (Number)	Frequency (Percent)
glides 'l', 'r', 'w'	36	57%
single obstruents or nasals	13	21%
's' or 'z'	10	16%
other	4	6%
TOTAL	63	100%

An example of each type of onset error, whole and part, is found in Figure 20 where the arrows indicate the string position of the segment(s) involved.

Figure 20Onset Errors

Anticipations:

- | | |
|---|------------|
| (a) [↑] prince [↑] charming | chints... |
| (b) [↑] strive for [↑] perfection | sprive... |
| (c) [↑] playing [↑] Friday | praying... |
| (d) [↑] pot [↑] shot | shpot shot |

Perseverations:

- | | |
|--|-----------------------|
| (a) steady [↑] state [↑] vowels | stowels |
| (b) [↑] grew up in [↑] dirty water | ...girty |
| (c) [↑] brisket and [↑] potatoes | ...protatoes |
| (d) [↑] mushroom [↑] casseroles | casheroles [šərowlɪz] |

Exchanges:

- | | |
|---|---------------|
| (a) [↑] squeeze [↑] fresh oranges | freeze squesh |
| (b) [↑] pink [↑] flamingo | fink plamingo |
| (c) [↑] quite [↑] right | krite wight |
| (d) [↑] stick [↑] shift | shtick sift |

Movements:

- | | |
|---|------------------|
| (a) [↑] dinner at [↑] eight | inner at date |
| (b) [↑] an ice [↑] cream cone | a kays ream cone |
| (c) [↑] skip [↑] class | sklip cass |
| (d) [↑] paint the [↑] studio | spaint the tudio |

As is clear from the errors in Figure 20, the fact that entire onsets may be anticipated, perseverated, exchanged or moved (the (a) examples), is not unique to them, but is

duplicate~~d~~ by the onset parts in (b) - (d). In the (b), (c) and (d) examples we see the anticipation, perseveration, exchange and movement of obstruents, non-nasal sonorants and sibilants respectively. The pattern of errors that we found in errors involving parts of the syllable we see here, in errors involving parts of the onset.

In syllable constituent errors we did not find codas appearing in onset positions or vice versa. Similarly with onset parts. We do not find sonorants ending up in front of onset obstruents, nor do we find sibilants ending up behind obstruents or sonorants. That is we don't have errors like (122) - (124) (analogous to the actual anticipation in Figure 20).

- | | |
|-------------------------------|------------|
| (122) *strive for perfection. | psrive... |
| (123) *playing Friday | rpaying... |
| (124) *pot shot | pshot... |

Although each of the segments indicated above moves in the (b) - (d) anticipations in Figure 20, they only move into phonotactically possible positions. Cluster collocational restrictions are not violated. These positional constraints are highly reminiscent of the constraints on onset, peak and coda errors. Another important characteristic of onset errors related to co-occurrence restrictions is the fact that we do not find voiced obstruents moving into positions after

'p'. So, for instance, although we have the 'p' replacing the 't' in (125)

(125) strive for perfection sprive...

we do not find errors like (126)

(126) *strive for goodness sgrive for goodness

Indeed, we do not even find (or so it appears from the data at our disposal) errors like (127) (a) - (c), analogous to (126) above and (128).

(127) (a) *strive for goodness skrive...

(b) *bought shot shpot shot

(c) *buy the studio spy the tudio

(128) (a) pot shot ſpot shot

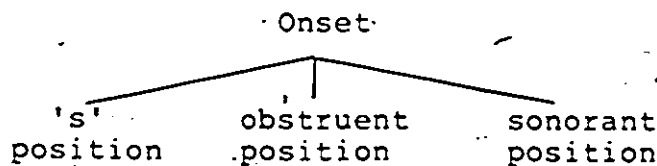
(b) paint the studio spaint the tudio

The lack of these errors is surprising, since they would simply involve voicing assimilation subsequent to the error. Apparently this voice assimilation is a low level allophonic rule and these rules occur downstream of segmental errors (see Chapter 3).

In the following section we will consider what the various characteristics of onset errors suggest regarding the constituent structure of onsets.

(iii) Discussion

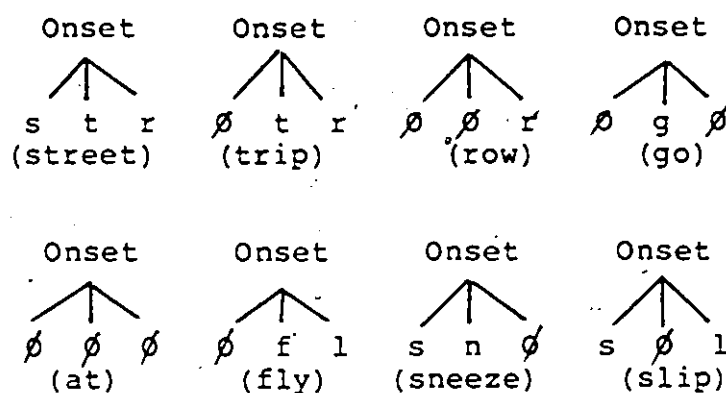
The clear parallel between syllable constituent errors and onset 'part' errors immediately suggests an analogous account. That is, if we assume the onset, like the syllable, is a sequence of three constituents, all of the 'part' errors may be analysed as the replacement of one constituent by the identical constituent where any of the constituents may be phonetically empty, or filled by some segment. That is, we suggest that the onset is a linear concatenation of what we will refer to as the 's' position', the 'obstruent position' and the 'sonorant position' as illustrated in Figure 21. The 's' position may or may not be filled by [s] or [š], the obstruent position may or may not be filled by any nasal or obstruent other than [s] or [š], and the sonorant position may or may not be filled by [l], [r], [w] or [y].

Figure 21The Onset

Examples of some English onsets are found in Figure 22.

Figure 22

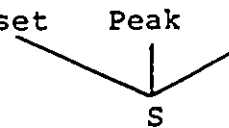
Some Onset Tokens



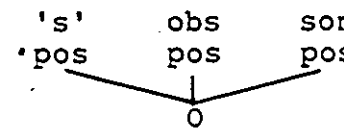
Errors involving parts of onsets may then be accounted for by Possible Segment ERROR - Onsets (PSE-O) illustrated in Figure 24 if we amend our syllabification algorithm as indicated in Figure 23.

Figure 23Incorporating Onset Structure Into Our Syllabification Rules(Figure 1)

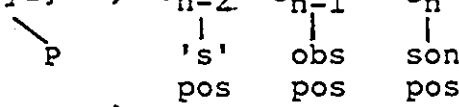
Rule 1': (a) $S \rightarrow \text{Onset} \quad \text{Peak} \quad \text{Coda}$



(b) $\text{onset} \rightarrow \begin{matrix} 's' & \text{obs} & \text{son} \\ \text{'pos} & \text{pos} & \text{pos} \end{matrix}$



Rule 2': (i) $C_1 \dots C_n \text{ [+syl]} \rightarrow C_{n-2} \quad C_{n-1} \quad C_n$



Where: C_n is $\begin{bmatrix} +\text{son} \\ -\text{nas} \end{bmatrix}$

C_{n-1} is $\begin{bmatrix} [-\text{son}] \\ [+nas] \end{bmatrix}$

C_{n-2} is $\begin{bmatrix} +\text{cor} \\ +\text{str} \end{bmatrix}$

Figure 24Possible Segmental Errors Involving Parts of Onset PSE-0

If: $X \left(\begin{matrix} 's' o_i \\ \text{obs}_{o_i} \\ \text{son}_{o_i} \end{matrix} \right) Y \left(\begin{matrix} 's' o_j \\ \text{obs}_{o_j} \\ \text{son}_{o_j} \end{matrix} \right) Z$

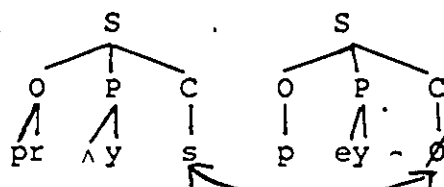
Then: $X \left(\begin{matrix} 's' o_j \\ \text{obs}_{o_j} \\ \text{son}_{o_j} \end{matrix} \right) Y \left(\begin{matrix} 's' o_i \\ \text{obs}_{o_i} \\ \text{son}_{o_i} \end{matrix} \right)$

PSE-O accounts for onset constituent errors whereas PSE (which we will rename PSE-S(yllable)) accounts for syllable constituent errors: Empty syllabic constituents may be replaced by filled constituents when they are constituents of the same type, as indicated in Figure 25.

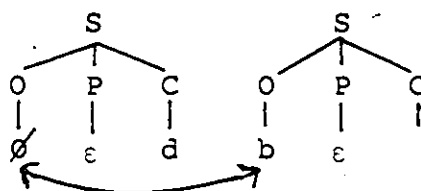
Figure 25

Empty Syllabic Constituent Replacements

price to pay → pry to pace



Ed and Bessy → Bed and Essy

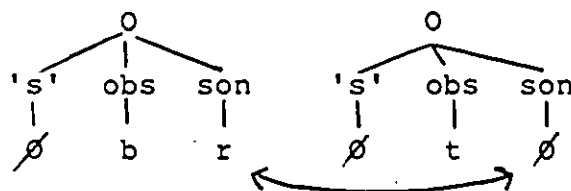


Similarly, empty onset sub-constituents may be replaced by filled constituents when they are constituents of the same type, as illustrated in Figure 26.

Figure 26Empty Onset Constituent Replacements

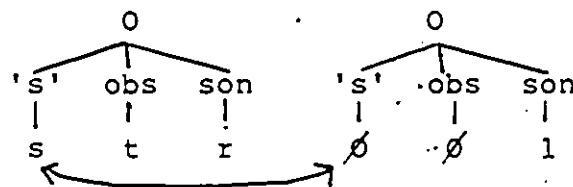
sonorant position error:

brush your teeth -- bush your treeth



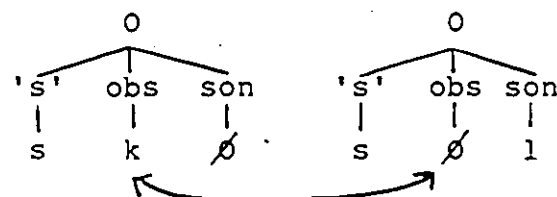
's' position error:

strong and long -- trrong and slong



obstruent position error:

ski slopes

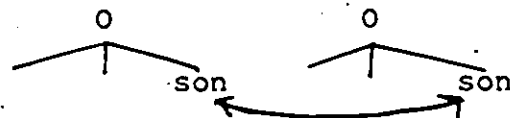
sni sklopes³⁶

And of course filled constituents may replace filled constituents as illustrated schematically in Figure 27.

Figure 27

Filled Onset Constituent Replacements

sonorant position errors:

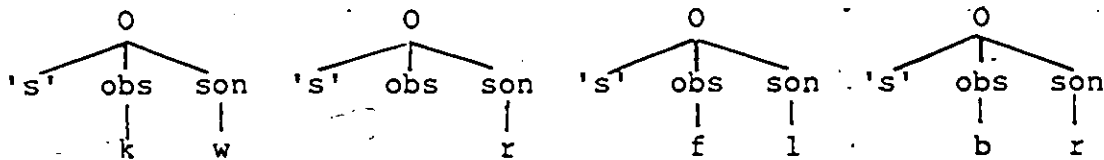


quite right

krite wight

flame broils

frame bloils



obstruent position errors:

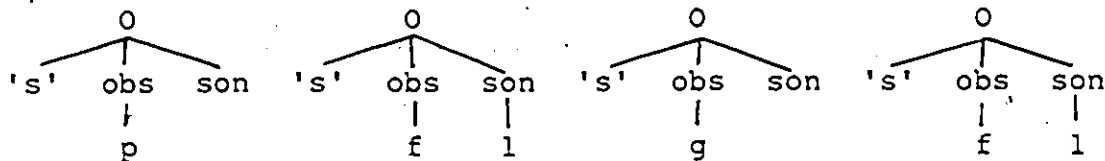


pink flamingo

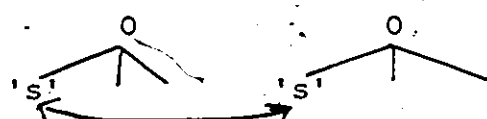
fink plamingo

golden fleece

folden gleece

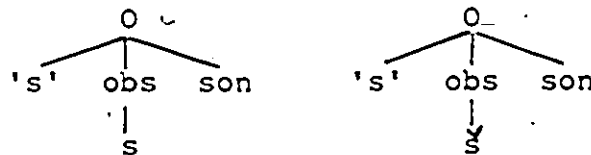


's' position errors: 37



stick shift

shtick sift



The error types in Figures 26 and 27 can all be considered exchange errors, where sometimes the exchange involved a phonetically empty constituent (Figure 26). Onset subconstituents are also perseverated and anticipated as illustrated in Figure 28.

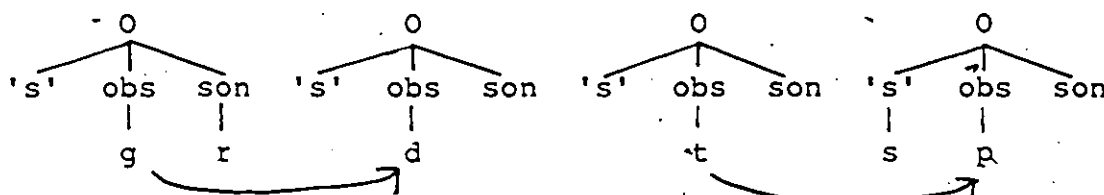
Figure 28

Perseveration and Anticipations of Onset Constituents

Perseverations:

I grew up in dirty water ...girty

how much time you want
to spend ...stend



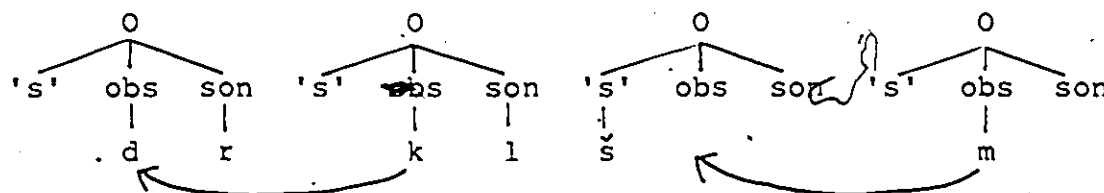
Anticipations:

draperies cleaned

craperies...

shut his mouth

shmut...



This symmetrical account of syllabic constituent and onset constituent errors is appealing for a number of reasons. Firstly, all these segmental errors may be considered to be the mis-selection of elements of identical types, characteristic of errors in general. Secondly, it has

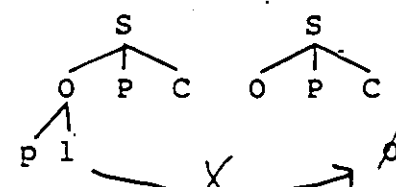
often been observed that segments interacting in errors share features more often than would be predicted by chance (cf. van den Broecke and Goldstein, 1980, Garrett, 1980, Nootboom, 1969, among others)). If the processor were only susceptible to substituting similar feature complexes one could infer that the implicated variable was indeed the feature. However the 'pairs' of feature complexes indicated in (129) appear to be radically different under any feature analysis that we are aware of and yet they interact in errors, as indicated.

- | | |
|--|------------------|
| (129) (a) j/w - John's work | Juan's jerk |
| (b) st/fl - over inflated
state | ..instated flate |
| (c) p/h - guinea pig hair | ...hig pair |
| (d) gr/m - do you have
May's grading? | ...grey's mading |
| (e) l/j - lawfully joined | jawfully loined |
| (f) skw/fr - squeeze
fresh oranges | freeze squesh... |

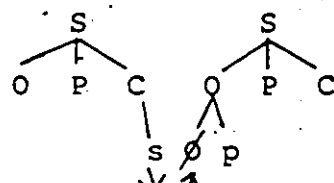
Similarly it's not clear how the processor could mistake no features for feature complexes. Errors involving positions filled by no segment are not at all uncommon (see Figure 26) and clearly cannot be considered the random movement of segments. The segments which move in such cases always move into the syllabic position which they come from. We do not

find errors like (130) (a) - (c) even though they would not violate phonotactic constraints.

(130) (a) *place to pay pace pail



(b) *price to pay pry to spay



(c) *streak of bad luck ...bard

If we assume that constituents are being processed, the 'impossible' status of these errors is explained. They would entail an interaction of constituents of different types, something that apparently doesn't happen. The errors in (130) have their analogues in actually attested errors, as indicated in (131).

(131) (a) place to pay pace to play

(b) price to pay pry to pace

(c) streak of bad luck ...brad...

In these we find movement into empty positions as we did in (130) but in this case the segments move into positions identical to those they have vacated.³⁸ Not only does the

constituent analysis. explain why very different feature complexes may be involved in one error, it also explains why errors often involve segments that are similar under a feature analysis. Any sub-constituent onset error will involve segments of similar types, sibilants, obstruents or sonorants. That is, the 'feature similarity' effect, or at least part of it, may be an artefact of the constituency of onsets, and, more generally, syllables. If feature makeup were crucial we would expect (132) an exchange of d/t, but not (133) an exchange of dr/h. In fact the reverse is what we find.

- | | |
|----------------------------|---------------|
| (132) d/t *driving Pat | triving Pad |
| (133) dr/h highway driving | dryway hiving |

The PSE-0 account of onset part errors also accounts for why there are no errors like (134) (a), which would violate phonotactic constraints, but (134) (b) is attested.

- | | |
|-------------------------------|--------------------------|
| (134) (a)*piece of pie please | piece of lpie pease |
| (b) piece pie please | plie pease ³⁸ |

'l' is a member of the set which can fill the sonorant constituent therefore it cannot turn up in the 's' position.

One common account of the non-violation of phonotactic constraints in errors is the 'Filter' account. For instance, Fromkin suggests that "...errors which violate phonotactic

constraints are 'corrected' after the substitution occurs" (1973, p. 22). This type of explanation has its problems. Besides being a 'mysterious black box' type of explanation - i.e. no real explanation at all - but rather a device which can be used to account for anything that can't be explained, it's not clear how such a device would work in on-line production. For instance, if the production mechanism moves the 'r' in 'str' in 'streak of bad luck' and randomly inserts it in the string, constrained only in that it must insert it in an onset, it could insert it at the beginning of 'of' or 'luck' or 'bad', producing 'rof' or 'rluck' or 'rbad' or at the end of the onset cluster producing 'rof' or 'lruck' or 'brad'. On the filter hypothesis 'rof' and 'brad' would be the only two to get through.³⁹

If the filter won't let the others pass, what happens? Presumably the whole production process has to be reversed in order to find a possible - the correct - output. This would surely be fairly time consuming and would presumably lead to fairly non-fluent hesitant speech output rather frequently. In this case, for instance a random placement of 'r' in the available onsets gives six possible errors - only one of which is (probably) possible. We should in this case get 5 reversals of the production process for one allowable error. Intuitively this is unsatisfactory. Surely the production mechanism is not so inefficient.

Ideally we would like to be able to specify the details of the production model such that errors and error free speech are both a reflection of its internal workings - such that the mechanism itself constrains the types of errors we find and don't find.

Our account attributes the non-occurrence of these errors to the actual mechanism - no filter is required, hence no backing up. Our account fails to account for why errors such as (135) don't occur, however.

(135) sporting a guru beard skorting a puru

We can offer no explanation for this lack. Since we find errors in which voiced and voiceless consonants interact, i.e. the obstruent position interactions, as in

(136) (a) call the girl gall the curl
 (b) draperies cleaned craperies...

It would appear that errors like (135) should occur with automatic voicing assimilation. It is possible that an error such as (136) (a) is a feature exchange [voice] (as Fromkin believes) or an entire onset exchange. But (136) (b) can't be explained this way.

There is another possible solution to this. Fudge and Selkirk both assign single segment status to 'sC' onsets.

This would predict the non-occurrence of errors like (135). This unitary status won't hold up however. Such clusters are often separated (unlike vowel-glide sequences in peaks) as exemplified in (137):

(137) (a) the sticky point	spicky point
(b) stick shift	shtick sift
(c) long and strong	trong and slong
(d) paint the studio	spaint the tudio

Fudge's notion that the peak is unitary is supported, but SC clusters appear to be treated as clusters by the production mechanism. We have a number of errors which are ambiguously entire SC cluster errors but these are analysable as onset/onset errors and do not require Fudge's unitary analysis.

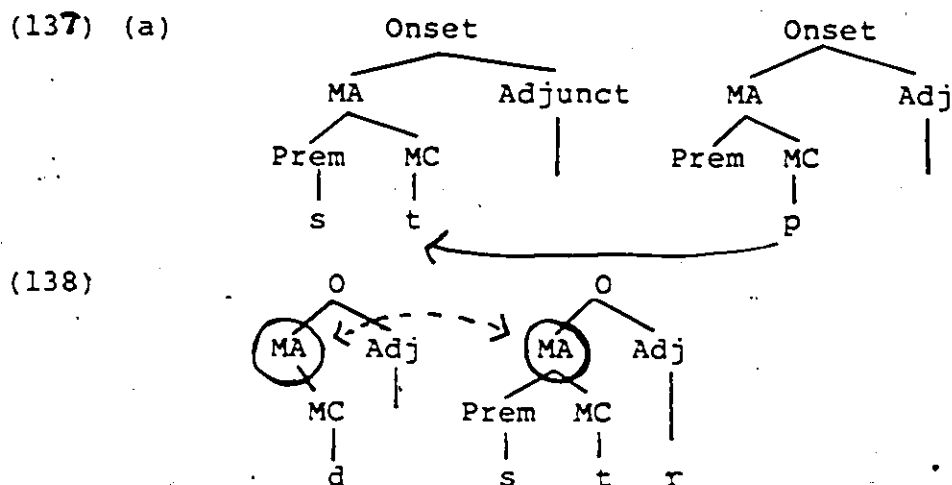
There are two SC errors (138) - (139) which can be used to support Fudge's notion, or the Margin Constituent hypothesized by Cairns and Feinstein.

(138) deep structure	steep dructure
(139) alfalfa sprouts	alspalfa frouts

The Cairns and Feinstein analysis allows the possibility of the errors in (137) and (138) as indicated in Figure 29.

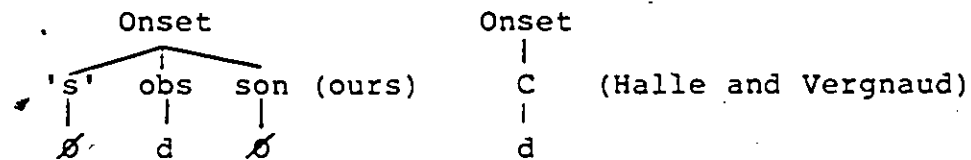
Figure 29

An Account of (137) (a) and (138) Using Cairns and
Feinstein's Onsets



Two errors is hardly strong support for the Cairns and Feinstein Margin constituent, but since we only have 10 errors which involve three-segment clusters we cannot come to any conclusion. We will therefore assume that these two errors are the conjunction of an 's' position and obstruent position error.

Halle and Vergnaud assign the onset a flat structure like our own. However they don't suggest that onset nodes be labelled. Indeed they would not build branches to dominate no material. Our analysis of the onset in 'deep' would involve two empty positions, theirs would include only a single branch, as indicated in Figure 30.

Figure. 30Onset Analysis of 'Deep'

Our onset reflects Fudge's notion in this regard although with differences. Without it we cannot explain segmental errors as the mis-selection of identical syllabic constituents by the processor. Although there are remaining problems to be worked out, we believe we are on the right track, that our empirical coverage is extensive.

However, there is one more problem with our onset constituent analysis, a problem of a conceptual, rather than an empirical nature. Codas and onsets do not interact as in (132) above and since the outputs would be acceptable if they did, it is reasonable to assume that the processor handles these constituents separately. The positions within onsets cannot be reversed however. Hence the programmer need not process these constituents separately. Our analysis in some sense, assigns more than the required modularity to the processor. Unfortunately we see no way of maintaining our empirical coverage and eliminating this problem.

10. Codas: Internal Structure

(i) The Paucity of Data

Errors in which there is an interaction of entire codas are uninformative with regard to the internal consistency of the coda constituent, whether the errors involve single segments or clusters. Unfortunately errors involving coda clusters are rare, particularly if one uses the uninflected root as the domain of syllabification. We have a total of 36 errors in which at least one cluster is involved, eleven of which include a morpheme boundary.⁴⁰ Of the 25 tautomorphemic cluster errors, four are 'whole' errors, and nine are ambiguously 'whole' errors.⁴¹ For instance (140) - (141) respectively.

(140) in the next ten minutes ...nen·text...

(141) (a) ...Monday Wednesday ...Monday...
 and Friday

(b) toast will pass toas will past

It would be, at best, premature to come to any conclusions regarding the internal constituency of codas on the basis of such a small corpus - twelve tautomorphemic part cluster errors, and eleven clusters which include a morpheme boundary. Nevertheless we find the errors suggestive, particularly those which include a morpheme boundary, and we will speculate briefly.

(ii) Speculations

It is noteworthy that segmental errors don't span morpheme boundaries.⁴² That is, contiguous segments from separate morphemes do not move as units. This suggests at first glance that the Halle and Vergnaud Appendix and the Fudge Termination nodes are viable syllabic constituents. For instance eight of our errors involve the segments preceding the third person plural, or possessive morphemes. These morphemes would all be assigned to these two nodes. The errors are listed in Figure 31 with the syllable structure they would be assigned by Halle and Vergnaud (only relevant nodes are expanded and we ignore the CV tier).

Figure 31Cluster Errors with Morpheme Boundary: The H&V Appendix

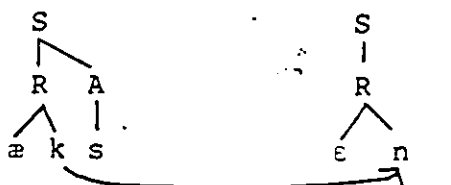
(i) steak or chops

chalks



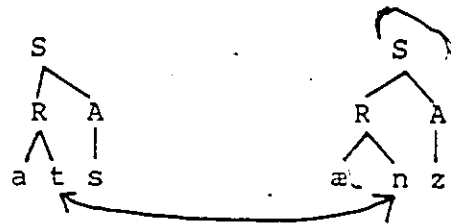
(ii) Jack's pen

Jack's peck



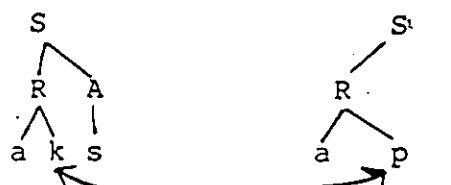
(iii) pots and pans

pons and pats



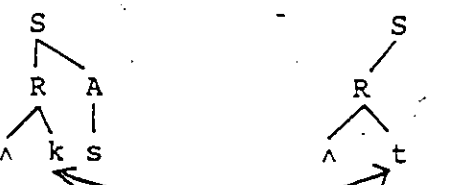
(iv) talks non-stop

tops non-stalk



(v) duck's butt

dutt's buck



(vi) rock ropes

rop rokes

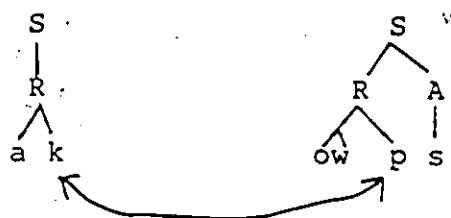
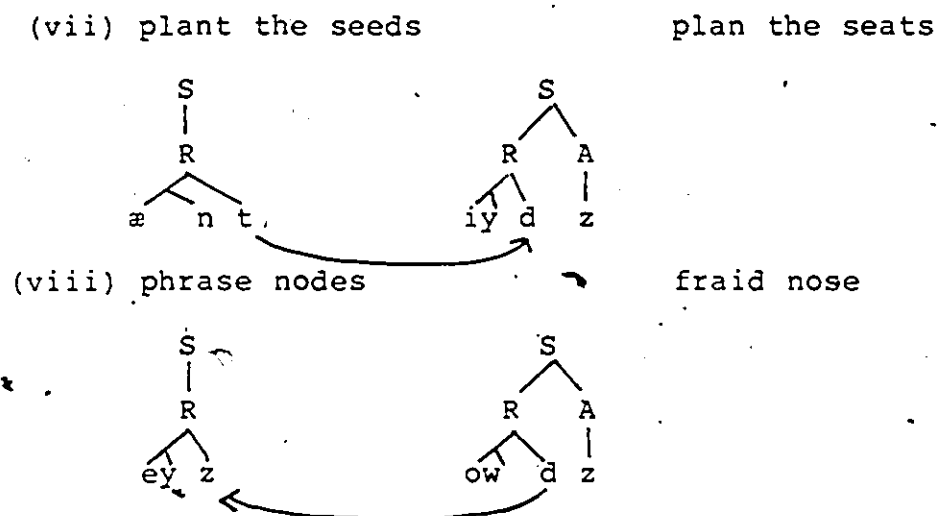
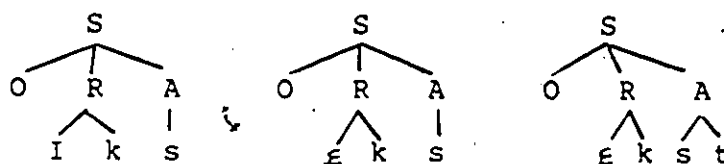


Figure 31 (continued)



The fact that these morphemes behave independently of the roots appears to be strong support for the H&V Appendix, and equivalently the Fudge Termination node. However, although the Appendix and Termination nodes are limited to word final position they are not limited to separate morphemes. That is any word-final coronal segments which are cluster final go in this node. For the most part these will be morphemes but not always. For instance 'next', 'Mex', 'fix' would be assigned the structures in Figure 32 where it is clear that these nodes may dominate coronal segments that are not affixes. (At least 'fix' and 'Mex' are monomorphemic.)

Figure 32

The error discussed above and repeated here as (142)

(142) Western Mexico

Wexern Mestico

appears to be an exchange of the codas [st] and [ks]. By using the H&V Appendix for the 's' the 's' acquires the same independent status the morphemes have in the errors in Figure 31. Unlike these, this 's' does not behave independently, however. (The same applies to the errors above involving 'fix' and 'next'.) Thus it would appear that the Appendix should be limited to coronals which are morphemes. As will be discussed in Chapter 3, Garrett points out that inflectional morphemes like past tense and plural are not involved in exchange errors. They may be attached in the wrong place as in (143) (a) - (d) but not exchanged.

- | | |
|---------------------------|----------------------|
| (143) (a) I had forgotten | ...forgot abouten... |
| about that | |
| (b) pointed out | point outed |
| (c) gets out | get outs |
| (d) oxen's yoke | oxen yokes |

These errors he refers to as shifts. Another characteristic of such morphemes is that they get stranded, left behind in Garrett's 'stranding errors'. For instance (144) (a) - (b).

- (144) (a) both kids are sick sicks are kid
 (b) I turned in a change ...changed...turn...
 of address

In addition, such morphemes do not turn up in word-internal coda positions. For instance we have word-final [d] exchanging with word-internal [z] in (145).

- (145) Lloyd Moseby Lloyz Modeby

This would not happen if the 'd' were the past tense morpheme apparently.

It appears then that the Appendix (or Termination Node) is supported if it is limited to segments following a morpheme boundary. Such segments have a unique pattern of behaviour. But in this case should it be considered a syllabic constituent? It doesn't behave like one. It does not appear to be involved in exchanges, perseverations or anticipations like other syllabic constituents. And as we shall see more fully in Chapter 3, Garrett's model of production accounts rather nicely for this differential behaviour of grammatical morphemes.

The other three errors we have which include a morpheme boundary argue against Garrett's analysis. However, unlike 6 of the 8 errors above they are not exchanges and hence are at least not clear-cut counter-examples. The errors are listed in (146)

- | | |
|------------------------|----------------------|
| (146) (a) french fried | friend fried |
| (b) nineteen millionth | nineteenth millionth |
| (c) third and fourth | thirth and fourth |

None of these segments would be assigned to the H&V Appendix so they cannot be construed as support for the Appendix. However these morphemes would be assigned to the Fudge Termination node since he allows empty positions. But the 'c' in 'French' and 't' in 'current' would not be assigned to this node, thus the errors could not be considered straightforward mis-selections of the Termination node, under Fudge's analysis.

We leave these as an enigma.

Let us briefly consider now the twelve part errors in tautomorphemic clusters we found. They are listed below as (147) - (158).

- | | |
|------------------------|--------------|
| (147) sink a ship | simp a shick |
| (148) kitchen sink | kinchen... |
| - (149) spent in Welsh | spelt... |

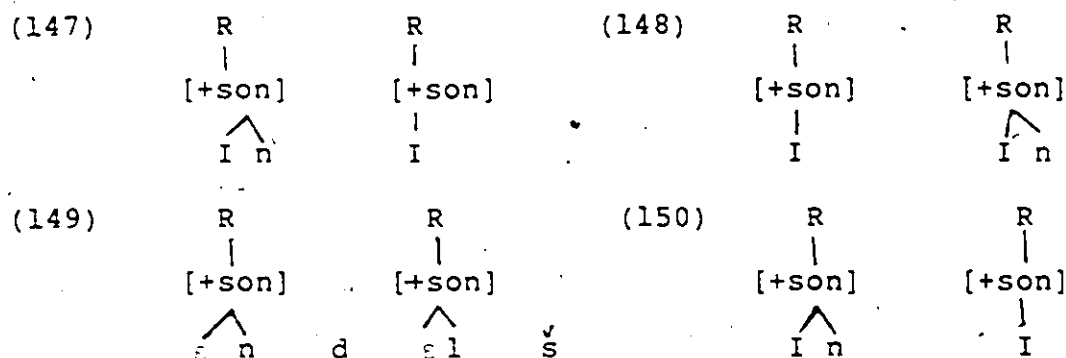
(150) pinch hit	pitch hint
(151) pest in every class	peess...clatt
(152) stink up the sea	stin up the seak [siyk]
(153) bank will pay	ban will payk [peyk]
(154) top shelf	tof shelp
(155) involving, underlying V's	...[viyvs]
(156) help laughing	helf lapping
(157) strange reasons	strain regions
(158) rank order	rand orker

The errors (155) - (156) are referred to above and the analysis of (157) - (158) as coda errors will be discussed below in our section on blends.

The errors in (147) - (150) could be analysed as errors of Halle and Vergnaud's [+son] constituent which dominates vowels and sonorant consonants within their rime, as indicated in Figure 33.

Figure 33

The H&V [+son] Constituent and Errors (147) - (150)



This is not likely a tenable analysis however, since such errors do not occur when the vowels are not identical. That is, we have no errors like (159).

(159) *pinch hat

pač hint

It appears clear that errors are more frequent when similar segments occur in context (see discussion above). The errors in (147) - (150) seem to reflect this characteristic of errors rather than provide support for the H&V [+son] node.

The coda errors obey collocational restrictions as do onset errors. We do not find errors like (160) (cf. (154)).

(160) *top shelf

tof shepl

Obstruents move into appropriate positions for obstruents: (151) - (156), and sonorants do the same: (148) - (150). We find obstruents appearing in positions where there are no segments: (152) - (153), (155) and ~~sonorants~~ doing the same: (148) - (149).

That is, members of the coda may be moved; but they may also be exchanged: (154), (156) and (158). As we argued above, exchanges and movements appear to be the same phenomenon, the difference being that in movements the exchange involves a phonetically empty constituent. Members of the coda may also be anticipated, (150), (146) (a) - (c).

Interestingly, although we find entire codas perseverated as in (161)

(161) thick slab

red book

thick slack

red [bud]

we find no errors that unambiguously are 'part' coda perseverations. We have ambiguous perseverations however as in (162) and so we assume this is an accidental gap given that we have so few part errors.

(162) a conjunction reduction

...redunction

All these characteristics are familiar. They suggest that the coda might be handled in a fashion analogous to our treatment of syllabic constituents and onset constituents: that is, we might find the coda to be a structure like the one depicted in Figure 34 and that the errors involve the interaction of identical constituents as in Figure 35.

Figure 34

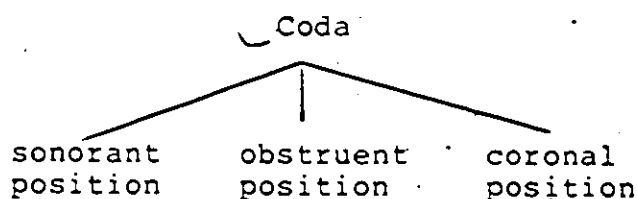


Figure 35PSE-Codas

$$\text{If: } X \begin{pmatrix} \text{son}_{c_i} \\ \text{obs}_{c_i} \\ \text{cor}_{c_i} \end{pmatrix} Y \begin{pmatrix} \text{son}_{c_j} \\ \text{obs}_{c_j} \\ \text{cor}_{c_j} \end{pmatrix} Z$$

$$\text{Then: } X \begin{pmatrix} \text{son}_{c_j} \\ \text{obs}_{c_j} \\ \text{cor}_{c_j} \end{pmatrix} Y \begin{pmatrix} \text{son}_{c_i} \\ \text{obs}_{c_i} \\ \text{cor}_{c_i} \end{pmatrix} Z$$

We propose that the coda dominates [+cons] segments and the laterals and nasals belong in the sonorant position,⁴³ the obstruents in the obstruent position and coronals in the final position. We must modify our syllabification rules as indicated in Figure 36.

Figure 36Incorporating Coda Structure Into Our Syllabification Rules(Figure 1)

Rule 1': (c) coda $\rightarrow$ $\begin{array}{ccc} \text{son} & \text{obs} & \text{cor} \\ \text{pos} & \text{pos} & \text{pos} \end{array}$

C

Rule 2': (c) $\begin{array}{c} X \\ | \\ p \end{array} C_1 \dots C_n \rightarrow \begin{array}{ccc} C_1 & C_2 & C_3 \\ | & | & | \\ \text{son} & \text{obs} & \text{cor} \\ \text{pos} & \text{pos} & \text{pos} \end{array}$

Where: C_1 is [+son]

C_2 is [-son]⁴³

C_3 is $\begin{bmatrix} +\text{cor} \\ -\text{son} \end{bmatrix}$

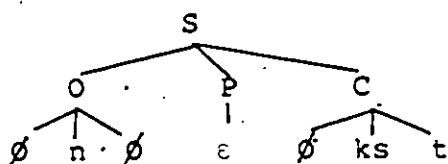
and C_1 - C_3 are unattached

We have no examples of monomorphemic [ks] (orthographic 'x') breaking apart, unlike 'k+s' in the errors (ii) and (iv) in Figure 20. Thus 'x' may be assigned to the middle coda constituent giving the syllabification for 'next' illustrated in Figure 37. Similarly we have no examples of 'sk' or 'sp' splitting up so they too may be assigned to the obstruent constituent.⁴⁴ 'st' clusters, on the other hand, do split apart as in (151) repeated here:

(151) pest in every class pes...clat

hence it would appear that these two segments are dominated by separate nodes, the obstruent node and the coronal node.

Figure 37



PSE-S, PSE- $\bar{O}$, and PSE-C will generate 98% of the submorphemic segmental anticipations, perseverations, and exchanges (including movements as a variant of exchanges). The symmetrical accounts of the syllable constituents, and the sub-constituents of onsets and codas is appealing as is the symmetrical account of the errors involving these constituents. We feel our arguments are strong for the constituent structure we have assigned to the syllable itself and to the onset, but the scarcity of part errors in codas has prevented us from coming to any conclusions regarding the internal structure of codas. Our account of how errors are generated which involve the syllabic constituents onset, peak and coda, i.e. PSE-S, we also find convincing. We are not as satisfied with our account of the errors involving subconstituents of onsets and codas, i.e. PSE-O and PSE-C.

Nevertheless we believe the account is superior to accounts which consider such errors to be mis-selections of segments per se. Such accounts can not account for the fact

that phonotactic constraints are adhered to without the use of a Filter. Neither can they account for why segments which interact in errors are often similar, but at the same time may often be quite dissimilar. Neither can they explain why segments turn up in positions where there was no segment to start with. That is, without the notion of constituents which may be empty or filled, so-called 'movements' or 'additions' must be considered to be a completely different kind of error from exchanges, perseverations, or anticipations. PSE-O and PSE-C are not thoroughly justified but as we have seen they are not without support either.

The other type of speech error which points to two string positions is the type referred to as word blends. In the following section we will consider to what extent the errors support or disconfirm the syllabic structure we have derived on the basis of segmental errors.

11. Word Blends

(i) Analysis: Method and Results

We had 113 blends at our disposal, including 65 from Fromkin's Appendix, Class U. We syllabified each of the two words involved in the blend noting the syllabic position at the point of union. Of our 113 blends 107 involved the simple juxtaposition of two words. The intersection point

was often ambiguous because a number of words involved in blends intersect at a common segment(s). Indeed this was true of 41% of our corpus (n=46). For instance (163):

- | | | |
|-----------|--------------------------------------|--------------|
| (163) (a) | <u>s</u> pank/ <u>p</u> addle | spaddle |
| (b) | de <u>a</u> ler/sa <u>e</u> lesman | dealsman |
| (c) | momentary/
instant <u>a</u> neous | momentaneous |

We considered the break to be where some segment in the first word, word₁, unambiguously belonged only to word₁.

This decision was not arbitrary. It reflected our recognition of the fact that there is a strong tendency for blends to occur at a point where words share segments. If we were to choose the intersection point as occurring after the shared segments we would not capture this characteristic of blends. Thus in (163) (a) we considered word₁ to end at 's' and word₂ to begin at 'p'. In (163) (b) we considered word₁ to end at the peak [iy], and word₂ to begin at the 'l', and in (163) (c) we considered word₁ to end at 'm' as the vowel following 'm', schwa, [ə], was common to both words. We would therefore consider the point of intersection in (a) to be within the onset, and the point in (b) to be the coda and the point in (c) to be the peak. (See below for why (b) is considered a coda intersection.)

The results of our analysis are found in Table 15.

Table 15Syllabic Position of Intersection Point in Blends

Frequency	Onsets	Peaks	Codas	Other	Total
Number	35	52	18	8	113
Percent	31%	46%	16%	7%	100%

(ii) Discussion

Table 15 makes it absolutely clear that blends are not random juxtapositions of words. In all but 8 errors the point of intersection was the identical syllabic constituent in both words. And five of the eight exceptions were blends that were not simple juxtapositions. These are listed in (164):

- | | |
|--------------------------------|------------|
| (164) (a) coins/change | canes |
| (b) transposed/
transcribed | transpised |
| (c) movie/music | [myuvik] |
| (d) meant/means | [menst] |
| (e) combination/group | combyution |

That is, a very simple algorithm will produce 93% of our blends. This rule, which we will refer to as Possible Blend (PB) is found in Figure 38.

Figure 38

Possible Blend

$$\text{If: } X \begin{pmatrix} O_i \\ P_i \\ C_i \end{pmatrix} \quad Y \begin{pmatrix} O_j \\ P_j \\ C_j \end{pmatrix} \quad Z$$

$$\text{Then: } X \begin{pmatrix} O_j \\ P_j \\ C_j \end{pmatrix} \quad Z$$

Examples of this algorithm at work are found in Figure 39.

Figure 39
Generating Blends

I F	(i)	X p æ r	o _i s	y n piy	o _j p æ l	Z l	person/people
		X p æ r			o _j p æ l	Z l	perple
I F	(ii)	X s	p _i ^	y niy h	p _j a	Z t	sunny/hot
		X s			p _j a	Z t	sot
I F	(iii)	X swI	c _i ct	y cey	c _j njd	Z ø	switched/changed
		X swI			c _j njd		swinged [swInjd]

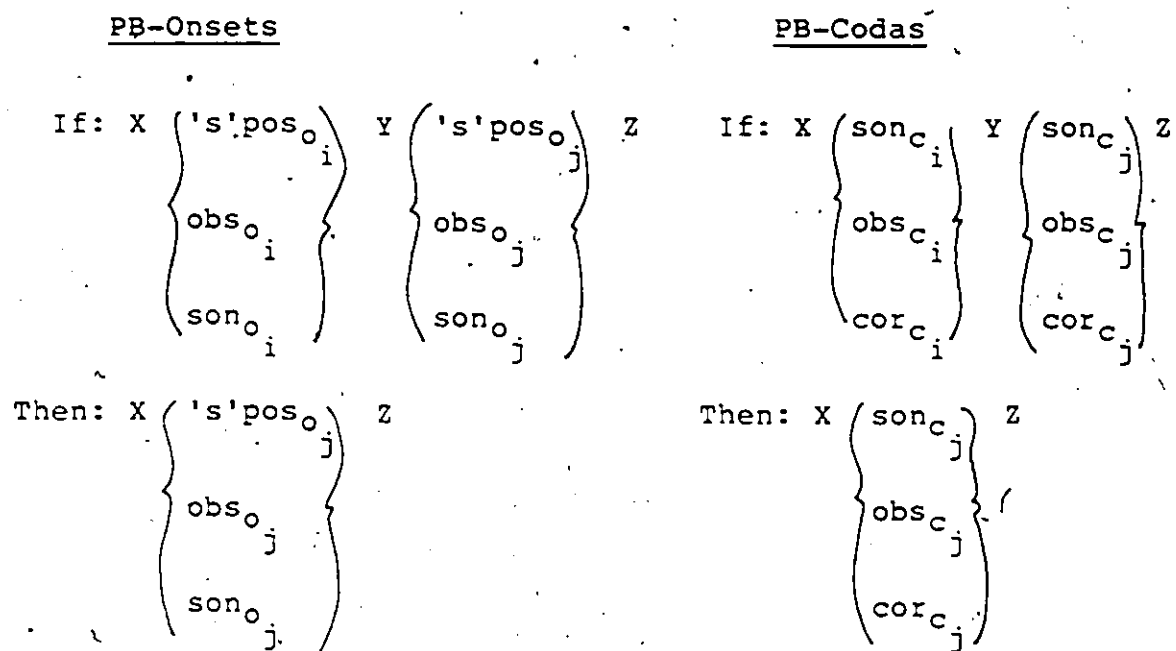
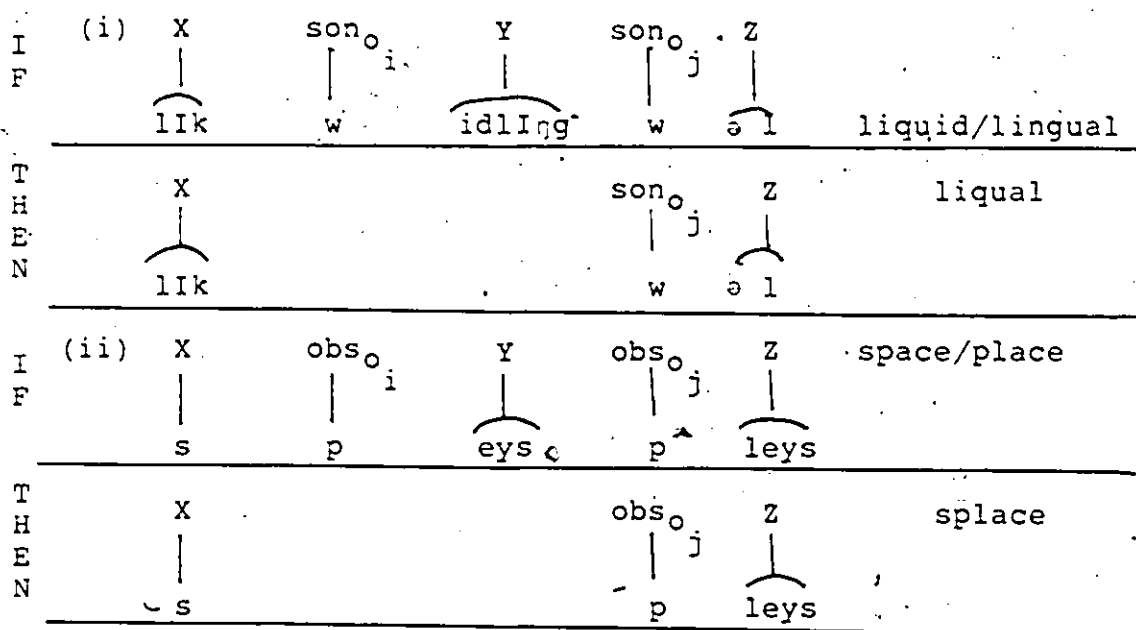
The similarity between the algorithms for generating blends and segmental exchanges, perseverations and anticipations is noteworthy. In all cases we find the mis-selection of identical constituents. Usually entire onsets and codas are involved and always entire peaks. PB can only generate blends of entire onsets, peaks and codas. We do, however, find some onsets that are split up or adjacent in these word-word juxtapositions, but the output is not an impossible cluster. That is we find errors like (165) (a) - (b) but not (166) (a) - (b).

- | | | | |
|-------|-----|------------------|-----------|
| (165) | (a) | move/run | mr̥n |
| | (b) | farce/slapstick | flapstick |
| (166) | (a) | *move/run | rmove |
| | (b) | *farce/slapstick | flapstick |

We found only two errors which involved coda parts, the error in (164) (f), which is not a simple juxtaposition, and one other.

Just as we needed to recognize subconstituents of onsets and codas to generate some segmental errors, we appear to need them for blends. The fact that phonotactic constraints are not violated is explained if segments are assigned to particular positions in onsets and codas, and errors are the result of mis-selecting from among these identical positions.

Thus we appear to require PB-Onsets and PB-Codas (analogous to PSE-O and PSE-C) for these few errors which involve parts of clusters. PB-O and PB-C are illustrated in Figure 40 and Figure 41 presents PB-O at work. (See Figure 46 for PB-C).

Figure 41Figure 42Generating Within-Onset Blends

We should like to reiterate that in our syllabifications we are assuming that all nodes are generated and they may be filled or empty in particular words. That is, we are claiming that mis-selections are mis-selections of syllabic constituents (or sub constituents as in PB-0) which may be filled or empty. This was what we claimed for segmental errors. Examples of empty constituents replacing or being replaced by filled constituents are found in (167) and (168) respectively.

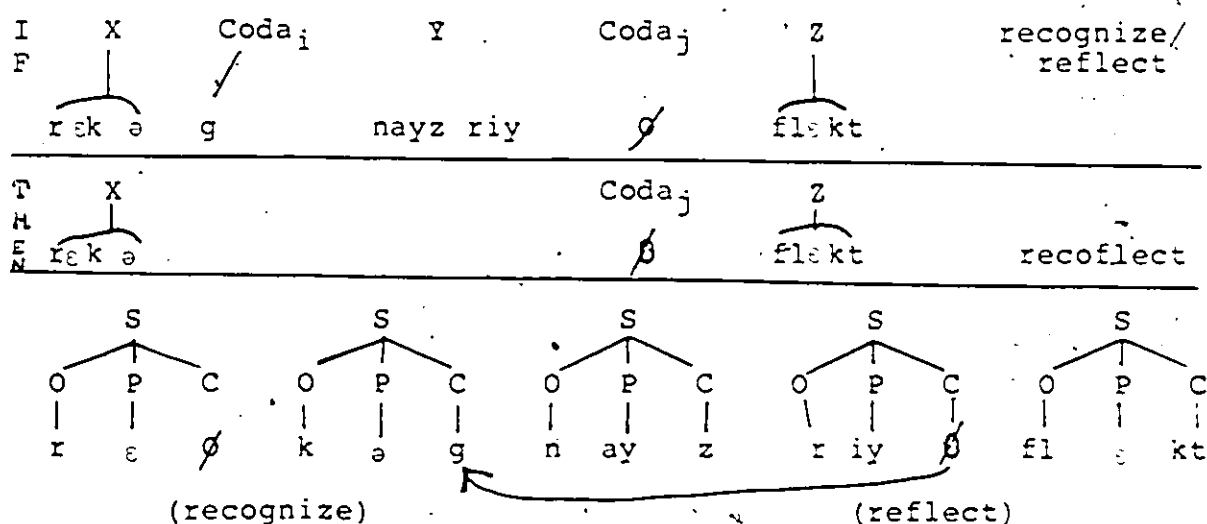
(167) recognize/reflect recollect

```
(168) move/run      mrun
```

How these are generated is illustrated in Figures 43 and 44 respectively.

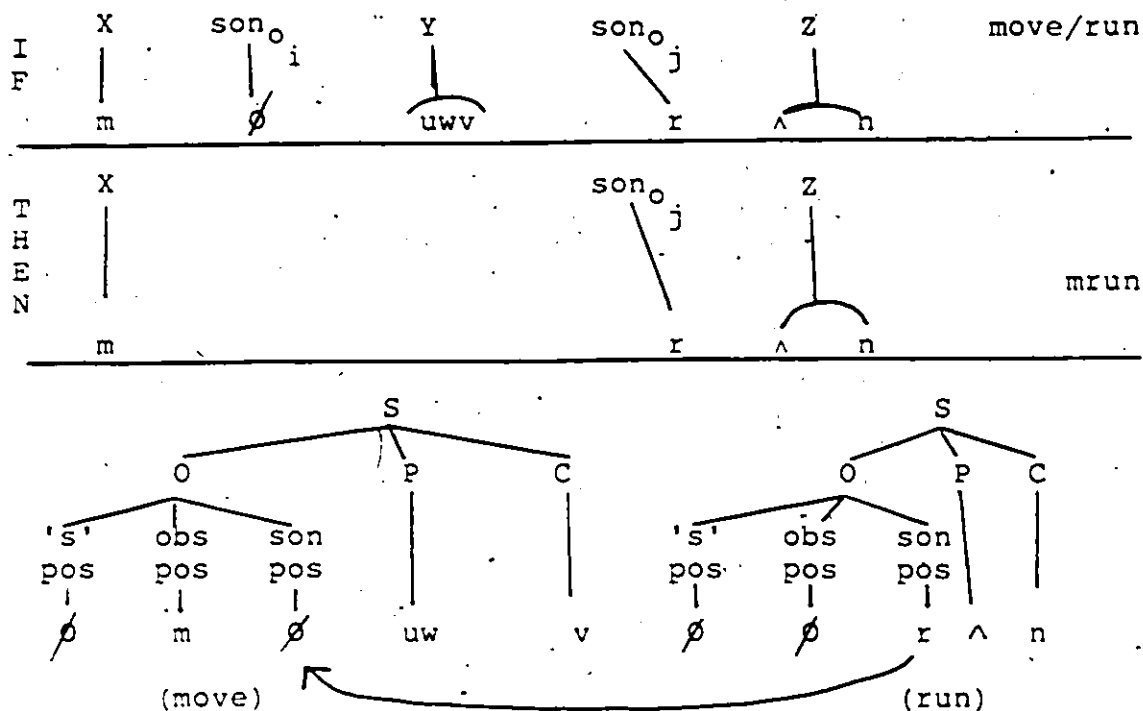
Figure 43

Blend Involving Empty Coda Constituents



N.B. Only relevant nodes are expanded.⁴⁶

Figure 44

Blend Involving Empty Onset Constituents

N.B. Only relevant nodes expanded.

Allowing these empty replacements seems to be giving ourselves too much power but it seems justified. Segmental exchanges seem to clearly require this power and we will see in Chapter 3 how these empty constituents fit in rather nicely with Garrett's recent notions (1982) regarding the production process (borrowed from Shattuck-Hufnagel).

We have not specified the domain of our syllabification algorithm. We have good reason to believe that when segmental errors occur roots are syllabified but, with

respect to blends, there are no strong arguments in favour of either the root or the word, so far as we can determine. The error in (169) suggests that the root is the domain but the error in (170) suggests that it may be the word.

- | | |
|--------------------------|-----------|
| (169) dealer/salesman | dealsman |
| (170) adjoining/adjacent | adjoicent |

To determine the relevant domain would require a richly developed theory of morphology because prefixes and suffixes are often involved in blends. This is not true of segmental errors. Is 'adjac' a root? or 'adja'? or neither? We don't attempt to answer this here. We have assumed the word as the domain. There is no reason to assume a priori that blends and segmental errors occur at the same level of processing and that the same domain is involved.

If we assume the word is the domain how do we explain the error in (169)? The 'l' in 'dealer' appears to be an onset and the 'l' in 'salesman' a coda. (We could assume the replacement of a null coda by a filled one but this would ignore the fact that blends often take place at points where words share the same segment.) We have two other errors which pose a problem for PB, (171) - (172).

- | | | |
|--------------------|---|--------|
| (171) pink/purple |  | pinkle |
| (172) bike/bicycle | | bikle |

In both cases the 'k' is coda but replaces segments which appear to be onsets.

Recall that we had four segmental errors that appeared to involve exchanges of onsets with codas which all had word final syllabic sonorants. The errors are repeated here as (173) - (176).

(173) strange reasons	strain regions
(174) call it spasm	...cause it...
(175) ocean perch	očan...
(176) rank order	rand orker

We also had two errors which we had to class as 'other' errors because on our analysis they involved an interaction of peaks with peak-codas. Namely the errors (177) - (178).

(177) ...paper a title	...titer
(178) mirror image	[mirəʃ] immer.

These were the only two errors we had in which 'r' was treated like a true consonant rather than a glide.

We would like to suggest that these syllabic sonorants are not integrated into the syllabic structure of words when errors occur but rather are unattached elements. Within current models of phonological theory there is a formal mechanism for recognizing end-elements (words, syllables,

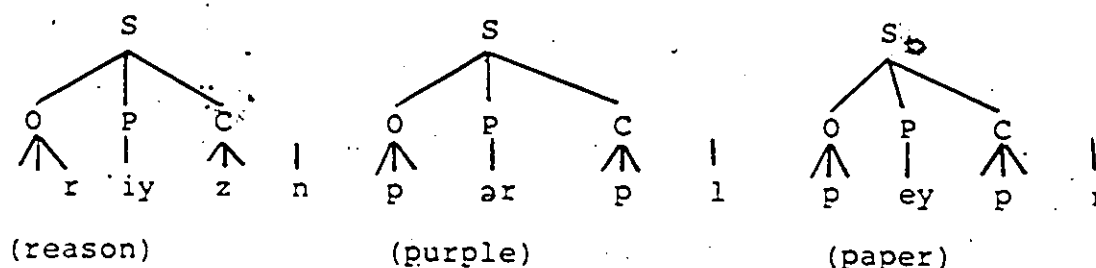
segments) that are not integrated into a particular structure. Such elements in the case of segments are termed 'extrasyllabic', and they are unrooted or dangling segments. (See McCarthy (1979), C&K (1983), Prince (1983), among others.) We would like to suggest the use of this formal device in dealing with syllabic sonorants; we propose that at the point at which segmental errors occur they are what we shall call 'extra-wordic', meaning that they have not been rooted in the word. Production errors suggest that they are un-rooted peaks, hence unrooted syllables, hence 'extra-wordic'.

It has been argued that the syllabicity of sonorant consonants is predictable in terms of their environment, i.e. in terms of whether they precede or follow consonants or vowels.⁴⁷ Mohanan (1979), for instance, claims that sonorants may be syllabic only if they follow a [+cons] segment. This would not seem to be correct however. The liquids may be syllabic following glides: optionally in the case of 'l' and obligatorily in the case of 'r'. Thus we have 'aisle' ([ayl] or [ayl]) and 'bowl' ([bowl] or [bowl]) with one or two syllables, but 'fire' [fayr] not *[fayr], and 'four' pronounced as one syllable if there is no glide, i.e. [for], but obligatorily with two if there is a glide [fowr], as in 'hour' [awr]. There are no minimal pairs in English 'VGL' versus 'VGL' but there are such minimal pairs with

nasals. We have 'Ryan' versus 'Rhine'. It is beyond the scope of this paper to account for the distribution of the sonorants; unfortunately it is not as clear cut as Mohanan claims. For our purposes it is enough to recognize that sonorants following a [+cons] segment are obligatorily peaks if they are word final or they precede a consonant. They are optionally onsets or peaks preceding neutral vowel-initial affixes, and onsets otherwise. Thus we have 'reasonable', 'wondering' and 'bicycling' with optional syllabicity of the sonorants, but 'rhythmic' in non-syllabic 'm'.

The speech errors above suggest that the processor treats post-consonantal sonorants as equivalent in some sense. We suggested above that the coda be limited to [+cons] segments, the first position being open to the sonorant, nasals and lateral 'l', the second being open to obstruents and the third to coronal obstruents. This would mean that there would be no syllabic position for post-consonantal sonorants in roots. That is, in words such as 'reason', 'purple' and 'paper' there would be no position for the sonorants. We would have the structures in Figure 45 with the sonorants extra-syllabic.

Figure 45

Extrasyllabic Sonorants.

Such extra syllabic sonorants become either O or P depending on what follows them. We will not attempt here to specify precisely when, or under what conditions, they become onsets or peaks. However, two things appear clear, firstly they are peaks word finally after consonants, and secondly, at the point at which segmental errors and blends occur they seem to be handled independently of the consonants that precede them. We have no errors which involve a syllabic sonorant and the segment that precedes it in the word. The sonorants themselves are exchanged or moved, or the preceding consonants are exchanged or moved. We don't find the consonants preceding syllabic sonorants moving with them in errors. In other words these sonorants do not belong to the coda which precedes them, only to be made [+syllabic] after the errors occur. In general we do not find errors involving a sequence of two syllabic constituents. What is noteworthy is that the consonants which precede such syllable peaks are not treated as onsets, but rather as codas. That is the

normal syllabification process which maximizes onsets aiming at preferred CV syllable structure has not applied to even single consonants which precede these sonorants when blends or segmental errors occur. We have no errors in which consonants which precede syllabic sonorants interact with onsets. In all cases they interact with codas. (See Figures 46 and 47 below.) We suggest that the generally accepted resyllabification rule which attaches word final consonants to vowel initial words is the rule which is responsible for assigning obstruents in codas their onset role before syllabic sonorants. This rule applies downstream of the locus of between-word segmental errors and blends.

The sonorants cannot be assigned to the coda because they never function as a part of it. We could leave them without a peak node dominating them as in (179), i.e. as true dangling extra syllabic segments.



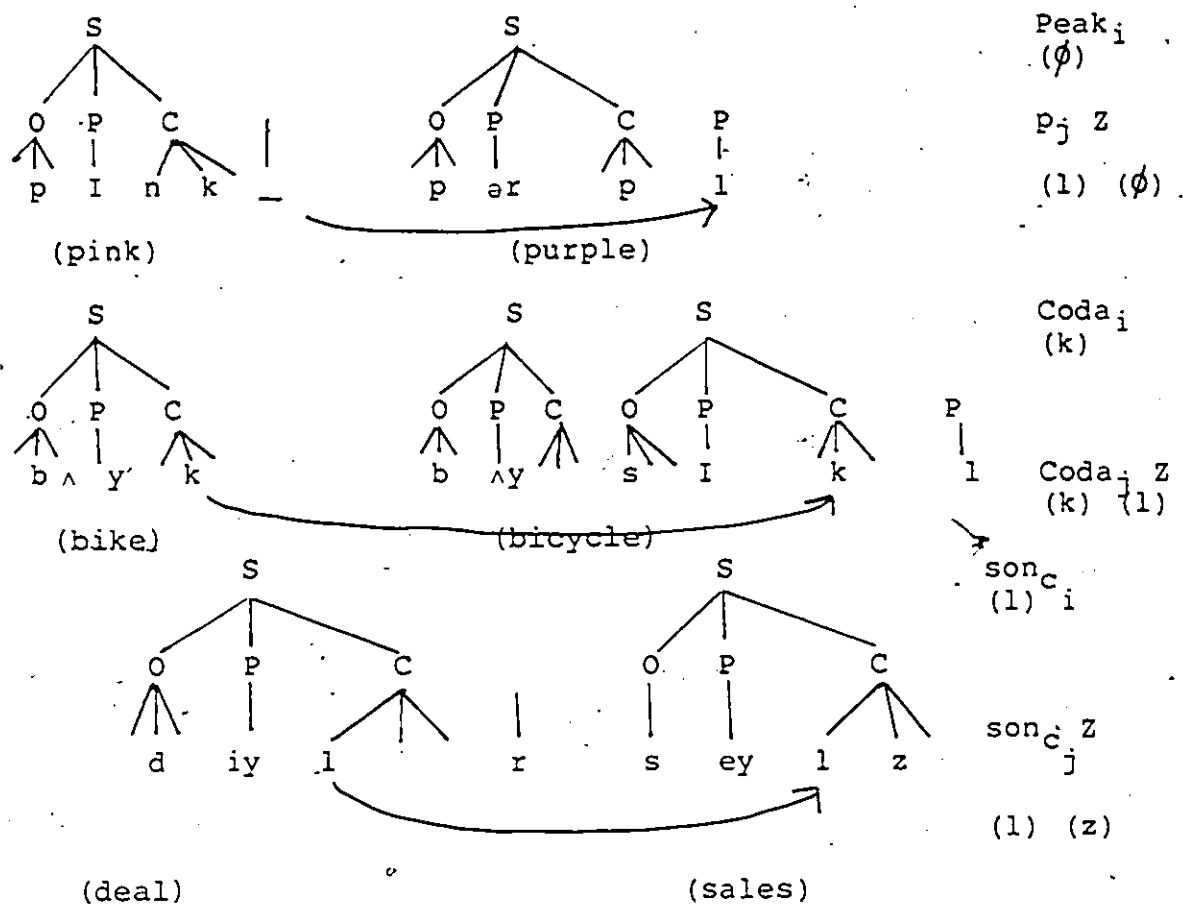
But then we would require extra formalism to account for errors in which they are implicated and are treated like peaks. For instance (180) (a) - (c).

- (180) (a) famous alternations ...fermous
(b) integer interger
(c) Filmore and Gruber Filmer...

It appears that the processor considers such segments to be peaks. (Although perhaps we need some enrichment which gives them a special status because although we find the errors in (180) we don't find analogous errors involving syllabic nasals or the lateral liquid. We would expect such errors if they were peaks on a par with other peaks. We will have to leave this unaccounted for.)

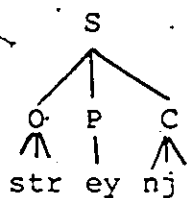
Consider now how the problematical blends can be accounted for within the mechanisms we have proposed. We illustrate this in Figure 46.

Figure 46
Generating Blends

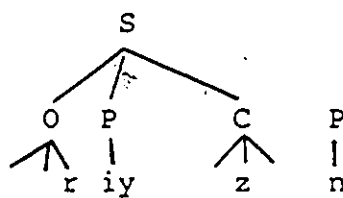


Thus segmental errors (174) - (177) also have a reasonable analysis as illustrated in Figure 47.

Figure 47

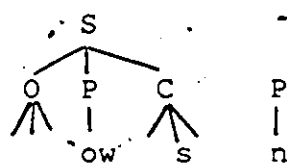
Segmental Errors (174) - (177)

(strange)

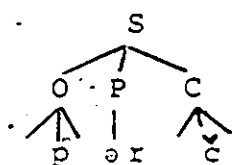


(reason)

obs_{C_i} (j)
↓
obs_{C_j} (z)

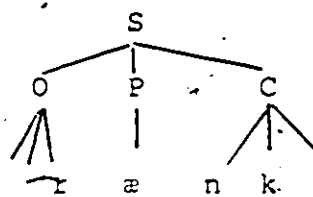


(ocean)

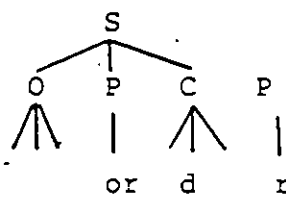


(perch)

coda_i (š)
↓
coda_j (č)

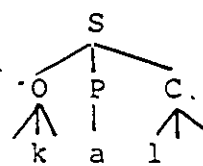


(rank)

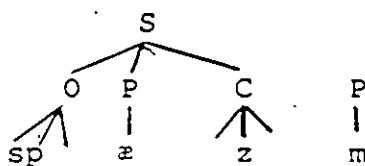


(order)

obs_{C_i} (k) (d)
↕
obs_{C_j} (d) (k)

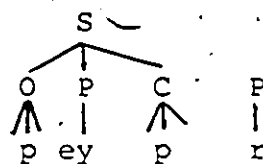


(call)

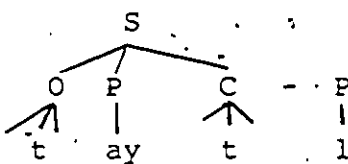


(spasm)

coda_i (l)
↓
coda_j (z)



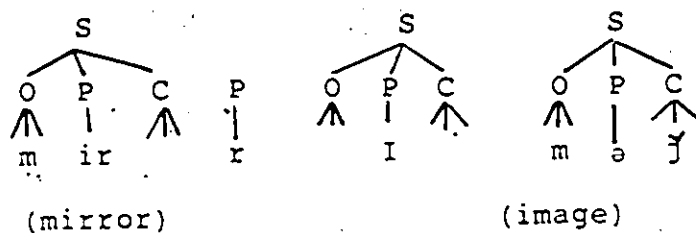
(paper)



(title)

peak_i (r)
peak_j (l)

Figure 47 (continued)



As can be seen from Figure 47, we remain incapable of accounting for the error in (179). We shall have to leave it as a problem.

We have seen that blends and sub-morphemic between-word errors are similarly constrained. We find segments interacting only if they belong to the same syllabic constituent. We have proposed that this constraint can be accounted for on the assumption that they occur when the processor is processing syllabic constituents. Errors involve the mis-selection of identical constituents. We have proposed that our algorithms PSE-S, PSE-O, PSS-C, PB-S, PB-O, PB-C generate the errors and hence mirror the production process at the level at which blends and segmental errors occur. We can account for about 98% of the segmental errors and 93% of the blends, and all of those blends that involve simple juxtapositions. We believe that haplogogies, which involve the coalescence of syntagmatically related words, can be given a similar account. The few at our disposal fit the pattern. We have not, however considered within-word errors.

They appear to be different from between-word errors as we shall see next.

Within-Word Errors

We have separated out the within-word errors from the between-word errors because a number of them do not exhibit the constraints on segmental interactions exhibited by between-word errors and blends. We have forty-one within-word errors, including the twenty in Fromkin's Class M. Of these, 7 involve the interaction of codas and onsets and if the root is taken as the domain of syllabification at least 10 more involve onset/coda interactions. This uncharacteristic distribution of errors is given in Table 16.

Table 16

Within Word Errors: Syllabic Position

(a) Root Domain

Syllabic Position	Frequency (Number)	Frequency (Percent)
onset/onset	17	41%
onset/coda	24	59%
TOTAL	41	100%

Table 16 (continued)

(b) Word Domain

Syllabic Position	Frequency (Number)	Frequency (Percent)
onset/onset	34	83%
onset/coda	7	17%
TOTAL	41	100%

Clearly these distributions are very different from those we found in between-word errors. In Table 17 we list the seven errors which involve onset/coda interactions with the word taken as the domain of syllabification. In Table 18 we give five examples of errors which are onset/onset interactions if the word is the domain, but become onset/coda interactions if the root is the domain.

Table 17Within-Word Onset/Coda Errors

(i) medal	meled [mɛləd]
(ii) reverse	reserve
(iii) fixed	sixed
(iv) fish	shiff
(v) steak	skate
(vi) olive	ovil [avɛl]
(vii) puck	cup

Table 18Within-Word Errors: Onset/Onset or Onset/Coda?

(i) felony	fenoly (fel <u>on</u>)
(ii) conversation	conservation (con <u>verse</u>)
(iii) nativity	navitivity (nat <u>ive</u>) [n vɪdɪv]
(iv) Harrisian	Hasserian (Harr <u>is</u>)
(v) modeling	moleding (mod <u>el</u>) [mɛlədɪŋ]

It appears that the domain of syllabification for within-word errors should be the word and not the root. And even if it is the root we still find that 17% of errors are clear violations of the syllabic position constraints

characteristic of between-word errors. There are also a couple of within-onset } and within-coda reversals which we don't find in between-word errors. Namely (181) (a) - (b).

(181) (a) whisper

whipser

(b) ask

axé

Some of these errors could be malapropisms (for instance (181) (b) and (ii) in Table 18) but most of them can not be accounted for in this way. It appears that a number of within-word errors must occur at a different processing level from the one at which between-word errors occur. Indeed, we will see in Chapter 3 that Garrett's model of the production process can accomodate the conflicting characteristics we find in these two error types.

Summing Up Chapter 2: Comparing the Production and Competence Syllables

(i) Introduction

Working on the assumption that the production process will "reflect the grammatical decomposition of the language faculty" (Garrett, 1982, p. 27), we expected that the syllable would prove to be a necessary construct in a model of production and that speech errors would confirm this. At the conclusion of Chapter 1 we pointed to seven characteristics of the syllables proposed by competence

theorists which would be the focus of our attention when we considered the production grammar construct. These included the domain of syllabification, the external boundaries of the syllable, and its internal constituency, including whether or not such internal constituents were labelled. We concluded that the domain of syllabification at the processing level at which between-word segmental errors occur was the root (or stem) and not the word, that Kahn's slow speech rules accounted for external syllable boundaries, and that the syllable consisted of three constituents, peak, onset, and coda - the latter two being themselves composed of three constituents. We also concluded that constituent and subconstituent nodes were labelled in order to account for the fact that the processor handled these constituents independently, mis-selecting among identical constituents. In this, our final section of Chapter 2, we will summarize our reasons for reaching these conclusions and illustrate in what ways the production template differs from or coincides with those competence templates described in Chapter 1.

(ii) External Boundaries and Root-Dominated Nodes

Our first goal in this Chapter was to determine whether speech errors supported the syllable as a unit of production. The fact that the syllable itself, unlike morphemes, phrases, words and segments, was rarely involved as a unit in

production errors had suggested (to Klatt, 1981, for instance) that it need not be recognized as a unit of the production process (at the processing level where segmental errors occur). On the other hand, theorists like Mackay (1969), and Fromkin (1971) had noticed that segmental errors seemed to involve segments from identical syllabic positions, which suggested that in fact the syllable was a necessary construct in a model of production. In order to determine whether this insight stood up it was necessary to specify precisely what syllabic positions were involved and what segments belonged to what positions. In other words this insight was completely unsubstantiated without a precise algorithm defining syllable boundaries and the internal constituency of syllables. Moreover the algorithm would have to specify a domain for its application.

We found that, indeed, segmental errors did only involve interactions between segments from identical syllabic constituents as defined by our rules in Figure 1. These rules were a modification of Kahn's slow speech Rules I and II where we incorporated an extra tier which identified pre-vocalic consonants as belonging to onsets, post-vocalic consonants to codas, and vowels as belonging to peaks. Initially this extra tier was proposed for expository purposes, thereby allowing consonantal segments interacting in errors to be identified as belonging to the onset position

or the coda position. However it turned out that what began as an expository convenience for identifying the position of our segments turned into a tier - a tier identifying production syllable constituents. We found, not only that segments which interacted in errors come from identical syllabic positions, but also that if a sequence of segments was involved in an error, this sequence was a constituent, either onset, peak or coda. Above the segmental level sequences of segments involved in errors invariably are analysable as constituents of some type: word, morpheme, phrase etc.; we discovered this was true too of submorphemic sequences. Such sequences could be identified as onset, peak or coda. Submorphemic errors often implicate sequences of segments, but not when these cross the constituent boundaries defined by our rules in Figure 1. Fromkin and Mackay's basic insight was correct in so far as it went - identical syllabic positions are implicated in errors, but more specifically the constituents onsets, peak and coda are implicated.

Fromkin, for instance, had suggested that errors always involved identical syllabic positions (without any precise account of what these positions were), but this analysis could not differentiate those sequences of segments which would likely be involved in errors from those that would not. For instance, under her analysis, errors involving sequences

of pre-vocalic consonants were as likely, or unlikely, as those involving sequences of vowels and consonants (VC) or consonants and vowels (CV) as long as the sequences came from the same syllabic position. These positions being only intuitively identified, made this claim impossible to test rigorously but clearly it did not account for the facts with regard to segment sequence errors. We found numerous errors which involved entire onsets for instance but very few which involved onset-peaks or peak-codas. In addition to these facts regarding sequences of segments, we also found that errors which involved empty onset and empty coda positions could be handled simply if these constituents were recognized. Because segmental errors always implicated identical syllabic positions, and because sequences of segments were only involved if they were syllabic constituents, we concluded that segmental errors were in fact syllabic constituent errors. We proposed a device, PSE-S, which we believe mirrors the behaviour of the processor when between-word segmental errors occur. It mis-selects among units of the same constituent type, namely, onset, peak and coda. The lack of internal structure assigned to the syllables of Kahn, Hooper and Clements and Keyser would not appear to be adequate with regard to the production unit.

We found no evidence that the processor selects (or mis-selects) rhymes. That is, if the rhyme is a constituent

of the production grammar syllable, it has a very different status from the onset, peak and coda constituents. The fact that rhymes are not moved as units is not evidence against the rhyme; syllables do not move as units and they are clearly supported because onsets, peaks and codas do move and they are defined in terms of the unit syllable. We do not require the rhyme in order to define any constituent, however. Thus we must conclude that there is no evidence from submorphemic speech errors that the rhyme is a constituent of the production syllable.

Arguments for the rhyme from competence theorists have had two principal areas of support: the Rhyme has been used to account for collocational restrictions between vowels and post-vocalic tautosyllabic consonants, and to account for prosodic phenomena relating to quantity sensitive stress systems. As we shall see in Chapter 3, there is good evidence that both syllabic structure and stress are lexical as far as the production lexicon is concerned, hence the rhyme would not appear to be required in a model of on-line production.

With regard to the H&V Appendix and Fudge's Termination nodes there appeared initially to be strong support. Grammatical morphemes which would be assigned to these nodes such as 'past' and 'plural' never were involved in errors

with preceding tautosyllabic consonants. That is, such segments appeared to be truly independent of the other coda consonants. They were always left behind. At first glance the conclusion would seem to be that the syllable does indeed include this constituent. We argued however that this constituent is a word constituent, not a syllabic constituent. Unlike the other syllabic constituents the members of the Appendix are not involved in selection errors unless they are tautomorphemic with the preceding segments. That is, the independence of the Appendix node was true only when it dominated separate morphemes. Hence we found errors like (182) (a) - (c)

- (182) (a) the toastu will pass toas will past
 (b) pest in every class pess...clatt
 (c) Monday Wednesday and Monday...
 Friday

but coronals which were morphemes got left behind when other segments moved. Indeed, as we shall see in Chapter 3, Garrett argues that grammatical morphemes are not involved in exchange errors at all.⁴⁸ This explains the independence of morphemic coronals dominated by the Appendix. Other syllabic constituents, onsets, codas and peaks, are clearly involved in exchange errors. It would appear that if the Appenndix (or Termination) node is a syllabic constituent it behaves

very - differently from the other constituents. This difference has to be explained. We believe that the difference is due to its not being a syllabic constituent. These arguments apply equally to the Fujimura and Lovins Affix constituent. Segmental errors provide no support for these syllabic constituents.

The external boundaries of onsets and codas were those defined maximizing onsets on the basis of the root (or stem) rather than the word because between-word errors which involved onsets interacting with codas almost invariably involved root final onsets. That is if the root were taken as the domain of syllabification these errors were analysable as coda/coda errors and were not evidence against our claim that identical syllabic constituents were always implicated in between-word errors. At the same time we did not find root-final consonants interacting with other onsets: ergo these root final segments belong to codas at the level at which between-word errors occur.

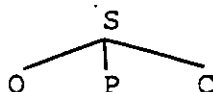
There was no support for 'fuzzy' external boundaries. Segments defined as ambisyllabic by Kahn's normal speech rules (Rules III - V) behaved like onsets or codas but not ambiguously. We could predict when such segments would behave like onsets and when they would behave like codas;

specifically as determined by onset maximization with the root as domain.

By assigning this much internal structure to the syllable, i.e. the structure indicated in Figure 48

Figure 48

Root Dominated Nodes of the Production Syllable



we could generate 83% of segmental errors which involved consonants and all of the vowel (glide) errors. The generation of these errors required that these nodes be labelled and allowed to dominate no phonetic material, in the case of onsets and codas. That is we took over Fudge's notion of empty positions. (See Chapter 1, Section 3.1.) Selkirk argues against the use of such a device on the grounds that not every token syllable should reflect the type 'syllable'. She argues that this device is analogous to asking every token NP to represent the type NP. We believe, in opposition to Selkirk, that every token NP does exemplify the type. Let us assume that every NP is generated by the X-bar rule in Figure 49.

Figure 49The NP 'Type' $\bar{X} \rightarrow \text{Spec}_{\bar{X}} \bar{X}$ $\bar{X} \rightarrow X \text{ Complement}_X$

(or, equivalently,)

 $\bar{N} \rightarrow \text{Spec}_{\bar{N}} \bar{N}$ $\bar{N} \rightarrow N \text{ Complement}_N$

Surely every NP token reflects this type; in some cases the phonetic material dominated will be null and in other cases it will not. We propose that this is analogous to Fudge's notion (and ours) of empty onset and coda positions. Moreover, assuming empty positions allows us to explain why onsets and codas which are filled can end up in positions which are empty as in (183) (a) - (c).

- (183) (a) a history ideology an istory hideology
 (b) an ice cream cone a kice ream cone
 (c) price to pay pry to pace

The remaining errors which involved consonantal segments could not be handled by PSE-S because they involved only parts of onsets and codas, not the entire constituent. On the basis of these 'part' errors we suggested that onsets and codas, unlike peaks were made up of sub-constituents.

(iii) The Structure of Onset, Coda and Peak

We designated as the 'peak' what is often referred to as the nucleus. We did this because when the peak includes a glide, the vowel and glide never separate. It is structurally unitary, as in Fudge's peak, which also includes the glides 'w' and 'y', and which he refers to as 'PEAK'. We argued that the speech error data supported Kahn's claim that the English 'r' is a glide: 'r', like the glides 'w' and 'y', was never separated from the preceding tautosyllabic vowel when it was involved in errors. As a syllabic constituent the peak dominates one position which may be filled by simple vowels or vowel-glide sequences. This indivisibility is not characteristic of onsets or codas however.

The pattern exhibited by entire onsets was duplicated by errors which involved only parts of onsets. We found sonorants, obstruents and nasals involved in exchanges, perseverations and anticipations. They would turn up in positions which had been filled or empty in the target word as was true of entire onset errors. For instance in (184) we find parts of onsets doing what entire onsets do in (185).

(184) (a) quite, right

(b) skip class

(185) (a) squeaky floor

(b) dinner at eight

grite white

sklip cass

fleaky squoor-

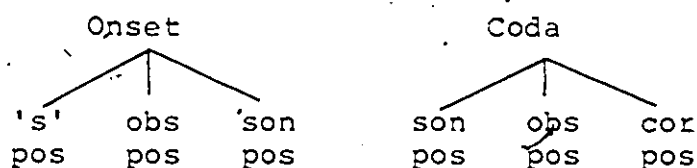
inner at date

Just as onsets did not move into coda positions we did not find sonorants ending up in obstruent positions in errors involving onset parts. Co-occurrence restrictions were not violated. The parallel behaviour of entire onsets and onset parts suggested a parallel analysis. We suggested that the errors be analysed as errors of the three sub-constituents, the 's' position, the obstruent position and the sonorant position. These mis-selections of identical sub-constituents we generated by PSE-O, identical to PSE-S except that the constituents mis-selected were onset constituents rather than syllabic constituents. This could account for the fact that phonotactic constraints were not violated without the use of a 'mysterious black box' filter. It could also explain why segments involved in errors could often be closely related from a feature point of view and at other times, completely unrelated. Errors involving onset sub-constituents would be closely related in features but those involving entire onsets could be quite different. This is the pattern we find. Our empirical coverage with this analysis is good but problems remain, particularly in that our analysis predicts certain errors which apparently do not occur.

Unfortunately our analysis of codas was severely hampered by our lack of unambiguous data involving parts of clusters. The only strong evidence was that pointing to the

independence of word-final grammatical morphemes. This evidence led us to our conclusions regarding the Termination and Appendix nodes mentioned above. We had only 12 part errors in tautomorphic coda clusters. Clearly with such a small sample we could only speculate. Since they exhibited characteristics which duplicated those of syllabic constituent and onset sub-constituent errors we speculated that they might have a structure similar to onsets: that found in Figure 50 where the two are compared.

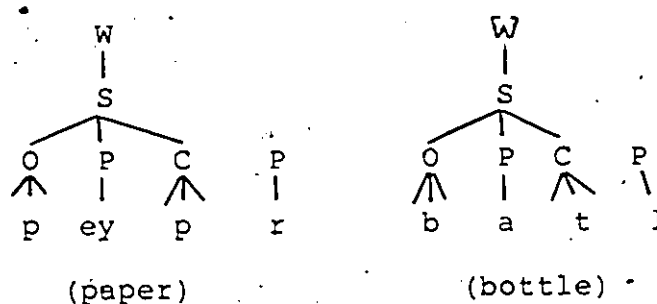
Figure 50
Onset and Coda



We found no evidence which required that either of these constituents be assigned any hierarchical structure like that suggested by Cairns and Feinstein. However our analysis of blends suggested that the English syllable might include an extrasyllabic position for sonorants which are preceded by consonants and followed by consonants, or are in word-final position. We suggested that when segmental errors and blends occur such sonorants have not been integrated into the structure of the word. They are unrooted peaks. We argued for this because consonants preceding such segments do not

behave like onsets. In all errors (blends and segmental errors) which include such consonants we find them behaving as codas not onsets. We argued that to assign the sonorants as members of the coda, on the basis of the coda-like behaviour of the consonants preceding them was not possible because no errors involved such pre-sonorant consonants and these sonorants acting as a unit. We suggested, therefore, that such sonorants were unrooted peaks. The consonants preceding them, the coda consonants, are resyllabified as onsets by the commonly accepted resyllabification rule which assigns consonants to open syllables across word boundary. This rule has not applied when segmental errors occur.⁴⁹ Thus the structure of words like 'paper' or 'bottle' would be that indicated in Figure 51 at the point at which segmental errors and blends occur.

Figure 51



The syllabic structure we have derived on the basis of between-word segmental errors and blends allows us to account for almost all between-word errors and blends as the

mis-selection of identical syllabic constituents by the processor. There were a number of within-word errors which could not be accounted for by our error generators. We concluded that these within-word errors must occur at a different processing level from between-word errors. We shall see in the following Chapter that this is not an unreasonable assumption.

In the following Chapter we will consider whether our observations regarding segmental errors can be integrated into some proposed models of the production process. We will suggest that our observations may be incorporated into Garrett's model not only maintaining its empirical coverage but also extending it.

NOTES TO CHAPTER 2

On the assumption that Fromkin would use this particular feature to distinguish 'b' from 'd'.

²Although apparently, (as of May 1984), scientists have still not be able to create the final undiscovered 'top' quark. (Grunwald, 1984.)

³The feature analysis predicts incorrectly that an error such as (i) should be possible where the feature [+cont] is anticipated producing a non-existent velar fricative (in English, phonemically).

(i) the girl fought the [ɣərɪ]...

Fromkin would have a pre-output Filter correct this. It's not clear why, considering the possible fricative variant of 'Agatha', i.e. [æɣə]. There are bigger problems with the use of a Filter, which will be discussed in the text.

⁴Garrett's classification scheme will be discussed in detail in Chapter 3.

⁵I attempted to remedy this in my own corpus but was unable to. Unfortunately it's only very rarely that one can be sure of what utterance, word-for-word, followed or preceded the error.

⁶The slash (/) above a word indicates the primary stress and/or the intonation peak in the utterance. This will be indicated where available and relevant. For instance in this error neither 'try' nor 'got' was stressed.

⁷This is Garrett's insight (see Chapter 3 for details).

⁸Fromkin might intend this, expecting her filter to correct the violation in 20 (b). In her theory 20 (a) would be impossible because she believes the processor distinguishes vowels from consonants. (See Fromkin, 1971, p. 224.)

⁹E.g. Fromkin, 1973, Boomer and Laver, 1969, MacKay, 1969, among others.

¹⁰We will refer to 'reversals' as 'exchanges' or 'metatheses' interchangeably. See footnote 12.

¹¹Garrett points out their differences. His arguments will be considered in detail in Chapter 3.

¹²For instance one man's 'metathesis' may be another man's 'reversal', 'exchange', 'transposition' or 'misordering'.

¹³This seems to us to be a significant fact regarding the processor's functioning which we can not explain.

¹⁴Of course no error can be given only one analysis - but certain analyses will get more support than others. For instance

(i) the two-yard line

the two lard yine

could be a word substitution 'lard' for 'yard' and a segment substitution of 'y' for 'l'. This is possible but not likely. Evaluation of particular analyses may require evaluation of entire theories - not a simple matter, clearly.

¹⁵However, later, we separate out consonant cluster errors, distinguishing among the whole, part, and ambiguous errors. See Section 9, this Chapter.

¹⁶See, Appendix to Chapter 1 and Chapter 1, Section 3. (vi).

¹⁷Whether one maximizes onsets on the basis of the sonority hierarchy (Selkirk), or strength feature (Hooper), or on the basis of permissible word initial clusters (Kahn), the resulting boundaries will for the most part be the same. However, the sonority hierarchy would allow 'thl-' [ɛl] yet this cluster is not allowed, word initially. It does appear to occur internally in 'a.thle.tic' so the sonority hierarchy appears to be correct here. However, we have chosen Kahn's method because we didn't need to make decisions regarding such forms as 'Atlee' which may be [.t^h.liy] or [t?.liy]. We therefore do not claim that the use of word-initial clusters predicts internal syllabification accurately but it was a precise method, and as it turned out it didn't lead to questionable analysis. There were no errors like (i) to pose problems. For instance if there were an error like

(1) the athlete brought

the ablete brought,

using permissible word initial clusters would require analysing this as a replacement of syllable final [ɛ] by syllable initial [b]. Using the sonority hierarchy we would have initial replacing initial. The first analysis disconfirms Boomer and Laver's claim and the second analysis confirms it. Fortunately we had no errors which were ambiguous in this way.

¹⁸A post hoc decision to refer to this position as 'peak' rather than 'nucleus' because the latter suggests a bipositional unit which is wrong.

¹⁹The initial assignment of post-vocalic glides 'y' and 'w' to the medial 'peak' position reflects the common assignment of these segments to the nucleus (cf. Selkirk, C&F, for instance), but for the moment is arbitrary. It will be justified by our results. The assignment of intervocalic glides poses a real problem which we discuss in Section 6.4 below. For the moment, glides which are not assigned to onsets will be assigned to peaks, away versus Bay.Shore.

²⁰We look at onset errors in detail in Section 9 below.

²¹Part of this high frequency is undoubtedly due to the relatively high perceptibility of errors in word-initial position. See discussion this Chapter, Section 2.2.

²²The following is a table of the syllabic position of consonant errors including only those which involve consonants exclusively.

Syllable Position	Frequency (Number)	Frequency (Percent)
onset/onset	352	81%
coda/coda	59	14%
onset/coda	25	5%
TOTAL	463	100%

²³289 onset/onset, plus 36 coda/coda plus 81 peak/peak plus 11 peak-coda/peak-coda plus 8 onset-peak/onset-peak = 505

$$\text{Percent} = \frac{505}{559} \times 100 = 90\%$$

²⁴In fact C&K are referring to 'schwa+r' ([ər]) as in 'bird', and distinguish this from the 'V+r' sequence in 'beard'. However, these sequences appear to be equivalent to the processor at some level since we have the error

I can sure hear you
[ʃər] [hiyr]

I can sheer her you

Similarly 'or' and 'er' are treated as equivalent in

Fillmore and Gruber

Fillmer...

²⁵To specify precisely what counts as a root, stem, derivational or inflectional affix for the production mechanism is beyond the scope of the current study. We clearly have here a rich area for research. We will have to be content with our intuitive notions regarding these things, unfortunately. It is certainly not obvious that 'Mex' is a root, for instance.

²⁶The error points to how differently the notion 'root' may have to be explicated when it comes to processing theories. See footnote 25.

²⁷For instance, it is not impossible that the error in (79) is a haplology, a not uncommon type of error. In normal speech the 'f' in 'of' is unvoiced before the voiceless 'c' in 'coffee' [kəpəf] → [kəf], i.e. 'cuff'.

²⁸There are numerous substitutions whose sources are not in the output string. For instance we have an error

(i) ski slopes

sni skIopes

where there was apparently an intrusion of a competing plan 'snow' and the movement of the onset 'sk' or the onset-part 'k'. That is, we have a blend of 'snow/ski' and a segmental error combined in one error.

²⁹Still ignoring the fact that onset and coda errors often involve only parts of these constituents.

³⁰Garrett demonstrates that spell-out must follow stranding errors. Hence the form [əz] in 'rices' rather than [s] from 'racks'.

³¹We have found two errors in which this appears not to be the case: (See also note 24.)

(i) mirror image

[mirəj] immer

and (ii) have you given your
paper a title?

...titer

³²and certainly has other possible analyses.

³³This includes the two onset-peak errors (viii) and (ix) Table 6 and the one peak-coda error (ix), Table 7. We deleted the error (x) believing it to be a blend as mentioned.

³⁴We include three 'other' errors here (i-iii) which were not included in Table 7 because we believe two of them to be coda/coda errors. How these are analysed will be illustrated in Section 11 below.

³⁵It does not require a particular analysis of codas which we shall suggest below in our section on blends.

³⁶The error involves a blend of snow/ski → sni and the movement of the unuttered 'k' from 'ski'.

³⁷All other 's' position errors which we have involve movement into empty positions. For instance

paint the studio

spaint the tudio

³⁸(131) (a) and (134) (b) are ambiguous - could be entire onset p/pl.

³⁹'rof' is not a likely error here. - in Garrett's model this is explained. See Chapter 3. Note too that the argument against production filters doesn't apply to competence filters since time-related efficiency concerns are irrelevant.

⁴⁰We have not included any of Fromkin's errors in 'ng' because their analysis is neither straightforward nor universally accepted. For instance in

(i) sing for the man

sig for the [mæn]

where does the [ŋ] come from, if not 'sing'? And why has the underlying 'g' disappeared in

(ii) clinging vine

[klɪŋɪŋ] vine?

⁴¹'whole' and 'part' errors being defined as they were in onsets. See the preceding Section.

⁴²We have an exception to this

(i) french fried

freind...

⁴³We have no errors involving liquid nasal sequences and such nasals we suggest might be assigned to an extra syllabic position. See our section on blends below.

⁴⁴How to formally incorporate such double segments is beyond the scope of this thesis but perhaps like vowel-glides they could be dominated by a single element of a CV tier - i.e. treated as complex segments. See our excursus on intervocalic glides.

⁴⁵We have no examples of 's' position blends.

⁴⁶Since this is an error of the whole coda constituent we have not indicated subconstituents of onsets or codas in this figure.

⁴⁷Kiparsky (1979) would account for their syllabicity in terms of the relative sonority of surrounding segments.

⁴⁸It's not clear that grammatical morphemes are not ever involved in exchanges, and are only moved as Garrett argues.

The single biggest problem
problem

the singest biggle

could be an exchange of such morphemes. There are also a very few examples of grammatical morphemes being anticipated. For instance:

(i) French fried

friend fried

(ii) the third and fourth

thirth...

⁴⁹We found no word-final consonants behaving like onsets. Cf. also Nooteboom, 1969, for the same result.

CHAPTER 3

INTEGRATING THE SYLLABLE INTO PRODUCTION MODELS

1. INTRODUCTION

We have shown that the segments which interact in errors belong to identical syllabic constituents, as defined by our rules (in Chapter 2, Figure 1), with the root as domain. We have suggested that this 'syllabic-position-constraint' can be explained if it is assumed that segmental errors occur when the processor is processing syllabic constituents. In other words, segmental errors are not segmental errors, but rather, syllabic constituent errors. On this assumption segmental errors may be considered to be the result of mis-selection of identical types by the processor. This supports Garrett's notion that the processor in general can only mis-select among elements of the same type, given that they occur when this type is being processed. We are claiming that this mis-selection occurs when the syllabic constituent is computationally relevant. It will be recalled that certain errors involved the manipulation of the entire constituents, onset, peak and coda (generated by PSE-S); others involved sub-constituents of the onset and coda (generated by PSE-O, and PSE-C). Since the sub-constituents are different types from the higher nodes, onset, peak and

coda, there must be two levels of processing at which segmental errors can occur given our assumptions (based on Garrett's - see Chapter 2, Section 3.).

In order for our hypothesis to be more than simply an observationally accurate description of the facts we must determine whether it can be integrated into some model of the production process. In addition we must show that roots are, at some level, computationally relevant and moreover that they are computationally relevant when segmental errors occur. We must also show that syllabic structure is computationally relevant, that is, there must be good reasons for assuming that the production mechanism requires syllable structure information in its processing. Clearly it would be less than explanatory to assume that syllable structure is assigned simply so that the processor can mis-select syllabic constituents.

This leads us to the more general question of levels of adequacy for production models. In the Introduction to this thesis we outlined what is meant for competence theorists to achieve observational, descriptive and explanatory levels of adequacy. Can one specify analagous levels for a production theorist?

A theory of production attempts to specify how it is that a speaker translates an idea into an utterance. In the

long run, the hope is, that the neurological processes which achieve this translation will be able to be specified. Presumably such processes are common to the species so we might consider a theory of production to have achieved a level of explanatory adequacy if it were to specify precisely those aspects of production that were universal, the parameters left open, and the possible values these parameters could take.

A production model of a particular language could be considered to reach a level of observational adequacy if it were to generate all and only the possible utterances in that language. Minimally, therefore, it is incumbent upon the speech modeler to account not only for error-free utterances but also for those that are not error free, and only those. It appears that both the explanatory and the observational levels of adequacy from competence theory have their counterparts in production. But what of descriptive adequacy?¹ It is not at all clear to us what might be suitable criteria for ascribing such a level to models of production. Let us assume, however, that if a model succeeds in pinpointing the actual locus of all possible errors within the normal production mechanism it has gone part of the way toward reaching descriptive adequacy. That is, if the model can specify where and how an error occurs, (not why - which was Freud's concern; hence 'Freudian slip') within a

framework which is required for error-free processing it will be at least partially descriptively adequate. This entails that any mechanism used to account for errors must be a part of normal error-free processing. It can not be ad hoc; it can not be a device which only accounts for errors. With respect to our comments above regarding the necessity for illustrating the computational relevance of roots and syllable structure, we are requiring of our hypothesis that it be descriptively adequate in this sense. We will want independent evidence that the processor handles roots, and we will want evidence that processing requires syllable structure. There is good evidence for both.

Firstly it appears clear that very general low level allophonic processes require the notion 'syllable' for their maximally simple statement. (See Chapter 1 and references therein.) It is also clear that low level allophonic rules apply whether or not errors do, or do not occur. However these general processes are at work after the occurrence of errors. For instance in (1) - (3) the allophonic variants are those appropriate to their error induced environment:

(1) paint the studio

[p^h] [t]

spaint tudio

[p] [t^h]

(2) the party of Pats

[D] [t?]

patty of parts

[D] [t?]

(3) guinea pig hair

...hig paɪr

I [ɛ]

[ɪ] [ɛ]

Since syllable-based allophonic rules apply it seems clear that syllable structure is required, but since they apply after errors occur it is not necessarily the case that syllable structure is available when segmental errors occur. There are good arguments that the production lexicon contains syllable structure however, (see our section on Garrett's model below) so since segmental errors must occur after lexical retrieval, syllable structure is available to the processor when these errors occur. Thus we are on the way towards explanation. Syllabic constituent structure is available when errors occur and it is required for allophonic processes downstream of the error. We must still explain precisely when the processor is selecting and mis-selecting these syllabic constituents and if possible, why. We will see below that within Garrett's model we can account for the 'when', but not the why.

With regard to roots. There is clearly some level of processing when roots are separate from derivational affixes - not only inflectional affixes. Consider the following errors in Table 1 from Fromkin Class R.

Table 1Derivational Morpheme ErrorsFromkin Class R

groupment	(grouping)
intervenient	(intervening)
motionly	(motionless)
oftenly	(often)
acoustal	(acoustic)
explanatings	(explanation)
perceptic	(perceptual)
ambigual	(ambiguous)
horizontical	(horizontal)
peculiaracy	(peculiarity)

It seems clear that none of these "words" is a lexical item. But they all seem to be the product of morphological rules of English. They indicate clearly that at some level the processor distinguishes roots from derivational affixes. (We saw this separation with respect to inflectional affixes in Chapter 2, Section 10.) Thus we have independent evidence that roots are computationally relevant at some level of processing. As was true of the syllable we must show that they are computationally relevant when segmental errors occur. We find that within Garrett's model they may be. In other words, in order for our hypothesis to be supported we must determine whether any current model of the production

process can accomodate it in a non ad hoc manner. It will be our purpose in this chapter to consider whether this accomodation is possible.

A number of speech production models have been derived on the basis of speech errors, three of which we will look at in this chapter. Each one of them pinpoints the locus of segmental errors so it will be our purpose to determine whether our observations regarding the role of the syllable in these errors can be incorporated at these loci.

In Chapter 2 we referred to various aspects of Fromkin's and Garrett's model, arguing for and against certain features in the light of our findings. We will reiterate here for expository purposes, wanting to present a clear picture of how these models account for segmental errors.

In Section 2 we will describe the most recent (1984) version of Shattuck-Hufnagel's model. We consider her model first because Garrett has recently (Garrett, 1982) suggested that her notions regarding segmental slots be included in his model. Our hypothesis cannot be incorporated into her model because she believes that lexical items are not syllabified.

In Section 3 we will discuss Garrett's model in depth. His empirical coverage is extensive and we shall see that our

notions can be integrated into his model and, we believe, extend its empirical coverage.

In Section 4 we will look briefly at Fromkin's account of segmental errors comparing its scope to Garrett's model. We shall see that it does not face certain problems Garrett and Shattuck-Hufnagel face but neither has it got the depth of empirical coverage enjoyed by Garrett's model.

We find that there is one intractable problem for all these models: why is it that segments which must begin in the correct serial position in the lexicon end up out of position in speech? The fact that words turn up in incorrect positions is less mysterious. Sentences have to be built and during construction things can go wrong. But words do not have to be built. Why do they get 'unbuilt'?

In Section 5 we will consider one attempt to deal with this problem, Crompton's (1981). We shall see that he doesn't altogether succeed.

2. The Shattuck-Hufnagel (1984) Account of Segmental Errors and the Role of the Syllable in This Account

The Shattuck-Hufnagel model of the production process was designed specifically with normal speech errors in mind. The aim was to build a model in which one could pinpoint the locus of the various types of speech errors. Of interest to

us here is her account of sub-morphemic segmental errors and the role that the syllable has to play in this account.

Shattuck-Hufnagel points out that a most puzzling characteristic of segmental errors is the fact that although segmental position must be included in lexical representations, (apt $\neq$ tap $\neq$ pat), for some reason segments sometimes appear in incorrect positions in utterances. For instance 'felony' comes out as 'fenoly' and 'feet moving' as 'fuwt meeving'. Shattuck-Hufnagel concludes, on the basis of this, that at some point in the production process feature matrices defining particular segments are separated from their positional slot in the lexical representation. The segments are later reintegrated into their correct position. It is at this point, the point of re-integration, that segmental errors occur.

As she herself realizes, this separation and re-incorporation of segment and positional slot was completely ad hoc in the original model. She believes, however, that recent developments in non-linear phonology provide independent theoretical support for this separation.

The part of her model which corresponds roughly to the phonological component of the grammar, she calls the Phrasal Prosody Component (henceforth the PP component). This component computes all prosodic characteristics and deals

with all allophonic variation in utterances. It integrates all linguistic and extra-linguistic information into its computations. The final out-put would appear to be a motor programme.

Two components in the PP component are of interest to us here. The first contains two representations: (i) an unordered set of segments from the pre-activated lexical items, and (ii) a sequence of ordered slots. Each of the slots carries what she refers to as a 'description' of the segment that occupied it before they were separated. This 'description' must include, minimally, a list of distinctive features, or else the slot and filler must be co-indexed. The second sub-component, the scanner-copyer, will be responsible for re-associating segments and slots.

The slots are part of the phrasal frame (cf. Garrett's (1980), (1982) model) where grammatical morphemes such as 'plural' have not yet been spelled out. Supported by the slots in this phrasal frame is the syllabic structure of the phrase. It is not at all clear how syllabic structure has been assigned but syllabification, and re-syllabification where necessary, is done in the PP component. Shattuck-Hufnagel explicitly rejects the notion that lexical representations include syllable structure but at the same

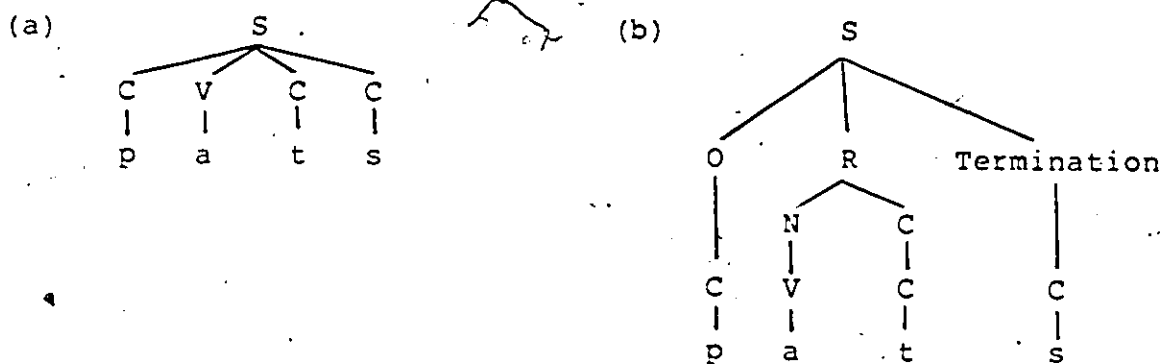
time insists that segments are reinserted into slots which support syllabic structure.

Perhaps slots are assigned syllabic structure on the basis of the 'descriptions' of the segments which the slots carry with them. This rules out the possibility of a simple co-indexing algorithm applying when the two are separated because feature specifications are required for syllable structure assignment.

The scanner-copier re-assigns segments to their slots, by reading the 'descriptions', scanning the segments for a correct match, and finally choosing the appropriate one and re-inserting it. In the old model the scanner-copier's function was simply to re-assign serial order to segments. In the extended model this is only a by-product of its main function which is to assign segments to their position in the syllabic structure associated with the slots. ~~Shattuck~~-Hufnagel rightly believes that a great deal of the processing that goes on in her PP component, such as stress assignment, allophonic variation, etc., is dependent upon syllabic structure. By incorporating it into her slot representation her model is strengthened in being able to handle such things. She believes that at the same time she motivates her separation of segment and slot.

She compares her slots to the CV skeleton in autosegmental models (cf. Goldsmith (1976), McCarthy (1979), Clements and Keyser (1983) and Chapter 1, Section 3.1 and 3.8) and equates her slots with units of timing and cites Halle and Vergnaud (1980) in this regard. The representation which would result after re-assignment by the scanner-copier does indeed resemble a representation in Clements and Keyser's theory. (See Chapter 1, Section 3 above for detailed discussion of the latter.) Hers is a hierarchical syllable however, based on the template in Fudge 1969. The two representations which would be appropriate for 'pats' are compared in Figure 1 (a) (Clements and Keyser) and (b) (Shattuck-Hufnagel).

Figure 1



A Clements & Keyser syllable

A Shattuck-Hufnagel syllable

Although there is independent motivation from theoretical phonology for this separation of levels it is not at all

clear how the formalism of autosegmental phonology can be incorporated into the Shattuck-Hufnagel model.

In autosegmental theory the units on each tier are sequentially ordered within their tier, and these units are associated with the elements on other tiers by highly constrained rules of association. The association rules will vary in particular details depending on the type of tier involved but nevertheless they involve precise algorithms. Similarly, lines of association are broken only with good reason. Assuming that Shattuck-Hufnagel starts with a representation in the lexicon of segments and slots as in Figure 2 (where the C's and V's represent her slots), by the time the representation reaches the first sub-component of the PP component it might look like Figure 3 (a) - (b)

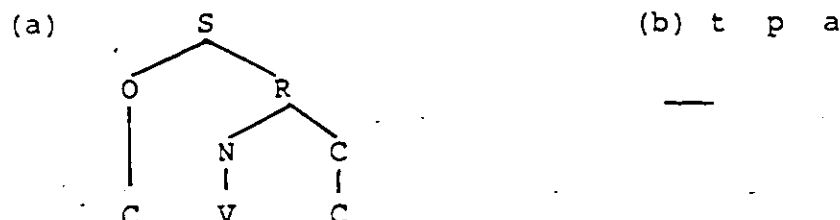
Figure 2

Lexical Representation of Shattuck-Hufnagel

C	V	C
p	a	t

Figure 3

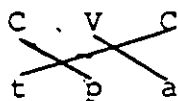
Lexical Item Scrambled



The lines of association between the segmental and the CV levels must be broken or else there will be no explanation for segments appearing in the wrong position in speech errors. For the same reason the segments must be scrambled or they could simply be mapped one-for-one and there would be no errors. The original problem with the model remains - why is the relationship between the two levels completely severed, and why do sequentially ordered segments suddenly become scrambled? The answer remains ad hoc; to account for speech errors. Supposing that some independent reason can be advanced for this there remains the problem of re-association of the two levels within an autosegmental framework. There is no mechanism within this framework which could duplicate the function of the scanner-copier. The scanner simply searches for the correct segment among an unordered set. The result could be a representation like Figure 4.

Figure 4

The Result of the Scanner Copier



This is not a possible representation within autosegmental theory. The analogy between Shattuck-Hufnagel's representations and autosegmental theory is superficial at best. In the former the segmental tier is unordered, an option

unavailable in the theory, and the principles governing the insertion and deletion of association lines are not honoured (cf. Goldsmith, 1976).

We have seen how Shattuck-Hufnagel incorporates the syllable into her model of production. It provides the structural frame into which segments are inserted. What kind of evidence does she give for its inclusion in the actual process of production (in addition to the arguments from theoretical phonology regarding its function in determining stress and allophonic variants)? She bases her argument on 80 polysegmental sub-morphemic errors from the MIT corpus. Of the 42 errors involving onset clusters, 36 involved the entire cluster as in 'cop the star' for 'stop the car', and not just a part of the cluster, as in 'sprit blain' for 'split brain'.² She concludes on this basis that there is a syllable unit 'onset', and thus, by implication, a unit 'syllable'. Similarly, she finds 10 rhyme exchanges³ and argues that the rhyme is therefore a viable unit in the production process. (See Chapter 2, Section 7 for arguments against this constituent.) She doesn't discuss the coda or the affix and is puzzled by the 5 CV (i.e. onset-peak) errors since usually this sequence is not considered a syllabic constituent.⁴

It's not obvious that her evidence for the syllable is very strong. The onset errors are all word initial and this would support a sub-unit of the word, not of the syllable. (See below for Shattuck-Hufnagel's hypothesis regarding this.) Also, she has only 10 errors which involve the rhyme, not an enormous number considering the size of the MIT corpus. (The 80 errors looked at constitute the entire set of polysegmental sub-morphemic exchanges in the corpus - not the whole set of segmental errors.) And too, she ignores the fact that there are only five more rhyme errors than there are onset-peak errors.

In fact, her model predicts that polysegmental errors which actually involve syllabic constituents will be a matter of chance. The scanner-copier is selecting segments, not syllabic constituents, when mis-selections occur. For example, to account for polysegmental errors which involve the rhyme she suggests that VC sequences destined for this unit might sometimes be transferred together even though lexical representations are not syllabified. Since segments, and hence VC sequences have not been syllabified the scanner-copier is not searching for syllabic constituents in this model. This model predicts that there will be no more errors involving constituents than there will be errors involving non-constituents. The relatively high frequency of onset errors compared to CV and VC (rhyme) errors is in fact

evidence against, not for, the model as a whole, because it predicts that errors which involve syllabic constituents as opposed to non-constituents will be a matter of chance. This high frequency of onset errors is support for the onset as a constituent of some kind -- a word or syllable sub-unit and the model must be revised in order that the frequency facts may be predicted. Shattuck-Hufnagel suggests that word-initial onsets may have a special status in lexical representations. They appear to play a special role in lexical access (Marlsen-Wilson and Welch, 1978; Cole et al., 1978) and 80% of exchange errors in Fromkin's UCLA corpus involve word-initial onsets.

This observation of Shattuck-Hufnagel's regarding word-initial onsets is a means of distinguishing the processing level at which onsets, peaks, and codas are selected from the one at which sub-constituents of the onset and coda are selected. The fact that 36/42 of Shattuck-Hufnagel's onset errors involve the entire constituent is interesting. We found many more unambiguous part errors (see Figure 2). This discrepancy in our distributions we believe can be explained partly on the basis of the fact that all of her errors involve word-initial onsets. Indeed we found only one error out of 149 onset cluster errors which involved a medial cluster, and it involved a single consonant being replaced by a sequence:

(4) alfalfa sprouts

alspalfa frouts

The cluster that moves (which is not an entire cluster) is from word-initial position. This suggests to us that lexical look-up may be implicated when entire onsets and codas are involved in errors. We will show how this fits into Garrett's model in Section 3 below.

Since the Shattuck-Hufnagel model defines the locus of errors to be the point at which the scanner-copier selects unsyllabified segments for their re-incorporation into their positional slot, the model predicts that errors involving whole syllables should be strictly a matter of chance and hence relatively rare. The evidence supports the model in this regard - very few errors occur which could be classified as syllable errors. Those that one does find are almost invariably morphemes or 'pseudomorphs' (cf. Garrett, 1975) as Shattuck-Hufnagel points out. Interestingly, this is also true of the German errors in the Meringer corpus (Meringer and Mayer, 1895). For example, 'Lackner and Goldstein' becomes 'Goldner and Lackstein' (MIT corpus) and the German 'Luder ins Loch' becomes 'Locher ins Lud' (Meringer corpus). Both of these involve morphs which are concomitantly syllables - 'Loch', 'Gold', etc. Syllables in this model are immovable structures in the phrasal frame - they are not manipulated or selected by the scanner-copier. Thus they can

not be MIS-selected; no more than sub-syllabic constituents can. The syllable has a passive role in this model; it is the structure used in the determination of the prosodic and allophonic shapes of utterances. It is the domain for the application of various phonological rules. Any segmental or polysegmental errors which appear to involve the syllable or its constituents are, within this model, a mere artefact - a result that is inevitable due to the fact that syllables are made up of segments. And in this model, it is segments or segment sequences that are liable to mis-selection.

There are two main problems with the model:

(i) It is quite clear from the data that we have analysed, that segmental errors are highly constrained by their syllabic position; onsets are involved with onsets, peaks with peaks and codas with codas. Errors always involve elements of the same syllable constituent. A model in which syllable constituents are actively processed by the processor can predict this - a model in which unstructured segments are processed, can not.

If Shattuck-Hufnagel were to allow syllabification before her segments and slots were separated, or include syllable structure as part of a lexical representation, this problem would be overcome. However, if she were to adopt either of these proposals as a solution to the distributional

facts noted in (i) she would concomitantly lose her motivation for separating segments from their slots. Slots support prosodic structure in her model, not segments. It is with this characteristic that she justifies their separation. It is in the assignment of segments to their prosodic structure, built on the slots, that segmental errors are hypothesized to occur. But, as we have seen, this hypothesis is not supported. Syllabic constituency membership is a constraint on which elements may interact in segmental errors. Hence syllabification must be accomplished before the occurrence of errors, not during or after.

To sum up, we can not integrate our hypothesis into Shattuck-Hufnagel's model because it would entail that she lose her motivation for separating segments from their serial position in lexical representations. The role of the syllable in her model is to provide the motivation for this separation. Hence although the syllable is crucial in her account of segmental errors, it can not account for the syllabic-position-constraint.

It appears clear that syllable structure must be assigned to the units involved in production as Shattuck-Hufnagel believes. However, there is no motivation for separating segments from their serial slots when such structure is assigned. We are left with the problem which

Shattuck-Hufnagel was trying to solve in this extension of her model:

(ii) There is no motivation for scrambling segments which must originally be sequentially ordered - other than to allow one to account for misordered segments in speech errors. This is obviously ad hoc. And as we saw above, her belief that the principles and representations of autosegmental theory could be of help in this regard, was unfounded.

3. Incorporating Our Hypothesis into Garrett's Model⁵

(i) Introduction

The Shattuck-Hufnagel and Garrett models of speech production are both derived on the basis of errors in normal populations. The structures and processes of their models are specifically designed so as to be able to account for each common error type. Both models specify the locus of segmental exchanges and our purpose here is to determine whether our observations regarding the syllabic constraints on such errors can be integrated into the models at the point where these errors are hypothesized to occur. We saw that in the Shattuck-Hufnagel model syllabification took place after segmental exchanges took place⁶ and hence the model could not account for the syllabic constituency constraints in such

errors. In the Garrett model, lexical representations appear to be syllabified and exchange errors involve these LRs, so it may be possible to account for our observations in this model.

In the original Garrett model there were no segmental slots but recently (Garrett, 1982) he has suggested integrating the Shattuck-Hufnagel slots into his model, although he doesn't specify how. We will first outline the model's general features and then consider in detail how this integration might take place. We will then look at how this change affects the model's empirical coverage - particularly regarding the errors whose locus is hypothesized to be the location of this integration. We will see that with minor changes our hypothesis can be incorporated and the empirical coverage improved, but the major problem associated with the Shattuck-Hufnagel model remains: Why do segments lose the serial position that they must originally occupy in the lexicon?

(ii) The Model: General

As a working hypothesis, Garrett provisionally assumes that "...the computational decomposition of language processing does, indeed, reflect the grammatical decomposition of the language faculty" (1982, p. 21). Therefore, although his model's observational base is the data of performance, in

particular spontaneous speech errors in normal populations,⁷ his analysis of the data is couched in a Chomskyan framework. His model, like TG grammar is modular. It includes three independent and temporally ordered levels: the message level, the sentence level (the first specifically linguistic level), and the articulatory level. Each of these levels may have its own independent sub-components, that is, sub-representations and processes which apply solely to these representations. Indeed the sentence level, the level his model is principally concerned with, includes at least two levels,⁸ the functional and the positional levels. The respective representations and processes of each of these appear to more or less correspond to those of the syntactic and phonological components of the grammar.

Garrett postulates these two levels, or computational types, on the basis of two things: (i) the surface domain which error elements traverse, and (ii) the grammatical category membership of the elements involved in errors. These two reflect what he refers to (1976, p. 237) as the two principles of error analysis discussed above (Chapter 2, Section 3.).

In his analysis of speech errors Garrett found that he could classify exchange errors into two broad categories: (i) those whose error elements respected grammatical category,

and occurred between phrases (or clauses) and (ii) those whose elements more often involved distinct categories, and occurred within one phrase. He concluded therefore that there were two levels, the functional, where grammatical category was computationally relevant and which was multi-phrasal (a two-clause limit being "...the best working hypothesis" (1982, p. 65)) and the positional, where grammatical category was not implicated and which consisted of one phrase. He found segmental exchanges to have the characteristics associated with this, the positional level. (He doesn't specify the locus of perseveration and anticipation errors, believing they are errors of a different type. (See below for discussion regarding this division.))

The functional level is developed under message level control. Major category lexical items (N, V, A, and P) are assigned to their phrases and given their grammatical role.⁹
¹⁰ A misassignment at this level will lead to an error like
 (5)

(5) Well you can cut rain ...trees in the rain
 in trees

where two nouns have been assigned to the incorrect phrase. In the construction of the functional level representation the processor is manipulating syntactic categories, hence we expect errors at this locus to be insensitive to sound or

meaning. This prediction regarding modularity is borne out. Words which exchange interphrasally belong to the same syntactic category but show no meaning or sound relationship.

Segmental errors, on the other hand, do exhibit effects for sound. "In significantly more cases than is to be expected in a random distribution the two elements involved in a substitution error are phonetically similar to one another" (Nooteboom, 1973, p. 149). (See also Shattuck and Klatt, 1978, Goldstein, 1980). Therefore, they are hypothesized to occur when the processor is working on phonological representations. In Garrett's model the module concerned with these is the positional level hence it is here, in the construction of this level, that we should be able to account for our observations regarding the syllable.

(ii) The Positional Level Slots

The positional level is built under functional level control. A single phrase planning frame is built on the basis of functional level information. In the original model this frame included phrasal stress and serially ordered slots for major class lexical items. Inflectional morphemes and other minor class items such as determiners were also part of the phrasal frame. These latter were attached in their appropriate positions after the major class items had been

inserted into their phrasal slots. This order of events explains why grammatical morphemes have the form appropriate to their error induced environment. For instance in (2)

(2) It pays to wait waits to pay

the present singular morpheme is pronounced [s] and not [z], as it would be if it were attached before the error occurred.¹¹

The positional level has been changed slightly in the most recent (1982) version of the model. It is at this level that Garrett suggests incorporating Shattuck-Hufnagel's ideas regarding serially ordered segmental slots (see Section 2. above). "...a detailed segmental skeleton for each lexical site would be a part of the planning frame and lexical retrieval would provide the detail of that structure" (Garrett, 1982, p. 51). That is, the slots would be built on the basis of the segmental information available at the functional level and later fleshed out on the basis of the second lexical retrieval. (Lexical retrieval in the model will be discussed below.) Therefore the phrasal frame now includes not only slots for major category items but these slots themselves are made up of slots. These slots support syllabic structure and, more generally, higher prosodic structure.

This is only a small change in the model but, as might be expected from a theoretically interesting model, this apparently small alteration is neither simple nor straightforward. In a model with a fairly detailed structure, a minor change at one point may require some rather extensive changes at other points if the model is to maintain its empirical coverage. And of course each of these changes may require their own concomitant changes. Whether this snowball effect will be the result of the integration of the segmental slots we cannot foresee but certain minimal changes are definitely required if the Shattuck-Hufnagel notions are to be incorporated.

(iv) Slots and Syllabification

Recall first Shattuck-Hufnagel's motivation for separating segments from their slots. She does this in order to account for errors in which segments are found out of position - a position which they clearly must occupy in the lexicon. For example (3).

(3) left hemisphere heft lemisphere

Recognizing that this is ad hoc she searches for independent reasons for this separation. She believes she finds it in current three-dimensional models of phonology where prosodic structure is on a tier separate from the actual feature

specifications of segments. She argues that since allophonic rules often require syllabic structure in their statement the segments of lexical representations must be assigned a syllabic position. It is during this syllable structure assignment that segmental errors occur. (She gives no explanation for why segments are not assigned a syllable structure directly, rather than first being separated from their linear position and then being reassigned to this position once these positions (i.e. slots) have been assigned a syllabic role.) In Garrett's model however, lexical representations (LRs) are already syllabified, so this aspect of Shattuck-Hufnagel he will not be able to adopt, or at least he will lose some empirical coverage if he does.

Although Garrett doesn't explicitly state that LR's are syllabified, it can be inferred on the basis of other claims. For instance, he argues that the parameters used in the second lexical retrieval¹² are the "...parameters of form which relate target and intrusion in malapropisms" (1980, p. 207). These parameters coincide with those used for word recognition (Marlsen-Wilson, 1980, Browman, 1980) and TOT states (Brown and McNeill, 1966), and Garrett suggests that this might be one point at which production and perception processes intersect. In Garrett, 1980, these parameters include, among others, syllable structure and stress. In Garrett 1982 we find 'number of syllables' referred to rather

than 'syllable structure'. If the feature [syllabic] is maintained then number of syllables could be available without actual syllable structure. Stress, however, remains one of the parameters, and Garrett believes that speech error data confirms the current formal theoretical separation of segments and prosodic structure. The speech error data points to the "...need for a separation of the prosodic features of words and phrases from their particular lexical instantiation" (1980, p. 1977). For instance, in (6) we have a within word peak exchange but the correct syllable remains stressed.

(6) marsúpiál

musárpial

"...there is indeed reason to seek an identification..." (1982, p. 51) of planning frames with the formal representations exemplified in Halle and Vergnaud (1980) (see Chapter 1, Section 8.). Since stress is built into these representations and not segment-bound, and since stress is one of the parameters of look-up, we can conclude that syllable structure is part of a lexical representation in Garrett's model.

If we remove syllable structure from the lexicon in order to match the Shattuck-Hufnagel model in this regard, we lose Garrett's account of malapropisms. To see why this is so we must first outline lexical look-up in his model.

Initially the lexicon is accessed via meaning. This is done before the functional level is spelled out and it is here that meaning-based lexical substitutions occur, as in (7):

(7) He rode his bicycle to school tomorrow. (today)

The item retrieved includes both syntactic and phonological information but only the former is used in the construction of the functional level. Subsequent to the development of this level, lexical items are accessed for the second time. This time the retrieval is sound-based and not meaning-based, and in particular, is driven by the parameters mentioned above. Items are retrieved on the basis of the phonological information acquired in the first retrieval. It is at this point that sound-based lexical substitutions, i.e. malapropisms, occur. For example (8) - (9).

(8) No, I'm amphibian. (ambidextrous)

(9) Because I've got an apartment now. (appointment)

(The errors associated with the two types of lexical retrieval exemplify Garrett's principles of error analysis mentioned above: sound-based retrieval results in sound-related errors and meaning-based retrieval results in meaning-related errors.)

At this point in the model, no phonological computations have taken place and yet as Garrett observes, stress and syllable structure are often implicated in malapropisms. If syllable structure is removed from his LRs his model will no longer be able to account for this. Indeed by having syllabified LRs, his model does not face the problem that Shattuck-Hufnagel does in accounting for our observations regarding syllabic constituency constraints on segmental exchanges.

Let us assume then, that although Garrett will incorporate the segmental slot separation, he will not attempt to justify the separation on the same grounds as Shattuck-Hufnagel. That is we will assume that he would not want to adopt the unsyllabified LR's of her model.

(v) Slot Construction

Now the question becomes, what forms provide the requisite information for slot construction in Garrett's model? Recall that Shattuck-Hufnagel adopts Garrett's notion of a single phrasal frame at the positional level. The slots become a part of this frame. The segments of lexical items lose their serial position when their position and feature composition are encoded in the slot description attached to each slot. That is, they lose their serial order when the

slots are built. This loss is important in the model because it is this, that accounts for the possibility of segments ending up in the wrong positions in utterances.

In Garrett's model, information at the functional level is used to construct the phrasal frame. Functional level information will direct the placement of phrasal stress, and will include the subcategorization facts which will direct the building of the phrasal slots. Presumably the lexical choices at the functional level will now also direct the construction of the segmental slot structure. The requisite information is available because they direct the phonologically based lexical look-up. We will refer to 'building' or 'directing the construction' of the segmental slots because we have only inferred that Garrett includes prosodic structure in the LR. Presumably the inference is correct, and if so then 'selecting slot structure' would be the more perspicuous account of the facts. Indeed Cutler's (1980) observations regarding incorrect stress in speech errors is more support for including prosodic structure in the LR. She notes that if the wrong syllable in a word is given main stress it is always a syllable that has main stress in some derivationally related word. This means that 'psychologist' is a possible stress error but 'psychologist' is not, the former coming from 'psychological'. Further

support for including stress in the LR comes from an error like (10).

(10) ~~R~~ockefeller

Feller~~o~~cker

The regular rules of English¹³ would assign main stress to the penult here and not to the initial syllable, and yet the error maintains initial stress. Assume then that segmental slot structure is built in accordance with the information provided in the functional level lexical choices. When the slot structure is built (or selected) the segments lose their serial position, they become scrambled. But this means that the second lexical retrieval will not be able to take place. Therefore lexical retrieval must take place before the slots are built. Suppose it does; the LR is used to direct lexical retrieval and this same LR is used to direct slot construction afterward. The LR is now scrambled.

There are two things wrong with this scenario. In the first place, the lexical item that gets inserted into the slots in the phrasal frame is the item retrieved in the sound-based look-up. But the segments in this item are not scrambled; it is the segments in the LR that directed the retrieval and the slot construction, that are now scrambled. It is this scrambling that is the source of segmental errors in the Shattuck-Hufnagel model and since Garrett is incorporating this part of the model, by implication, this

scrambling will also be the source of his segmental exchanges. The processor (Shattuck-Hufnagel's scanner-copier) chooses the wrong segment when it is fleshing out the segmental slot skeleton. If the segments of the lexical item that gets inserted are not scrambled the function of the scanner-copier is lost - as is the account of segmental errors in this model. Therefore the result of the second look-up must be scrambled; therefore it must be the form that directs slot construction (or on the basis of which slot structure is selected).

A look at both stranding errors and malapropisms leads to this same conclusion. Recall that in Garrett's model errors like (8) (repeated here)

(8) No, I'm amphibian. (ambidextrous)

are hypothesized to occur as a result of a mis-selection by the processor during the second lexical access. That is, the form retrieved in the meaning-based look-up (in this case 'ambidextrous') directs the sound-based look-up, which can result in retrieval of the wrong form, but it will be a phonologically related form (in this case 'amphibian'). Suppose that the slot representation is built on the basis of the intended 'ambidextrous'. How could the scanner-copier possibly map the segments of the incorrectly retrieved 'amphibian' onto the slots for 'ambidextrous'? The two are

clearly not isomorphic. Either the scanner-copier has a function - that of scanning and selecting segments according to the slot description, or it has, what amounts to no function - that of totally unconstrainedly, mapping any segment into any slot. If the scanner worked this way, we would expect correct forms to surface only rarely, i.e. to be a matter of chance. This is clearly not the case. Again we see that (if we are to retain a device like the scanner-copier) the segmental slot structure will have to be built under the direction of the form retrieved in the second lexical look-up. Alternatively, we can look for another account of malapropisms.

We are suggesting then, that the phrasal frame of the positional level be constructed on the basis of two representations; the functional level representation will direct construction of the syntactic category slots and determine the location of phrasal stress, and the forms retrieved as a result of the second lexical access will direct segmental slot construction of major class slots. This is a reasonable solution to slot construction in Garrett's model because his second lexical look-up is aimed at getting detailed phonological information. Since segmental slots are a reflection of just such detailed information it is plausible that they be constructed on the form retrieved in the phonological look-up.

(vi) Segmental Slots and Stranding Errors

This division of labour also receives some support when it comes to accounting for stranding errors like (119) in the new version of the model.

(11) Fancy getting your model renosed. (nose remodeled)

Stranding errors occur after the second lexical look-up. The stems so retrieved are inserted in the wrong phrasal slot in the phrasal frame, resulting in stranded grammatical morphemes. In the earlier version of the model the processor was selecting major class phonological representations as units from the sound-structured lexicon for insertion in the frame. Selection of the wrong morpheme led to a stranding error. In the new segmental slot model however, the processor (Shattuck-Hufnagel's scanner-copier) is mapping segments into slots. If the segmental slot structure is built for the intended 'nose' how can the bisyllabic 'model' be mapped onto it? Again, without completely trivializing the function of the scanner this mapping could not take place. Therefore the segmental structure selected must be appropriate for 'model' and not 'nose'. That is, the selection error must occur before the slot structure is built. This means that stranding errors in this new model can not occur when morphemes are being inserted in the phrasal frame into major

class slots, because morphemes are no longer inserted, segments (or syllabic constituents) are inserted one by one into segmental (or syllabic constituent) slots. These slots are superimposed on the phrasal slots, but morphemes as such are no longer inserted.

In fact at this particular locus in the new model it is well-nigh impossible to have errors involving whole morphemes, even if the scanner-copier were fairly unconstrained. Consider what it would involve.

We know that when the skeleton is being filled with segments by the scanner it has the segments from more than one major class morpheme in view. In particular it has in sight the segments of each of the major class forms that belong to the single phrase being constructed. Otherwise we would not find errors which involve more than one morpheme such as (12).

(12) moonlight sonata

[man]light [sənuwtə]

At this point, the point at which the skeleton is being fleshed out, the scanner is searching among an unordered set of segments - those segments which belong to all the major class morphemes which are a part of the phrase it is constructing. The scanner inserts each segment as it finds it, in accordance with information attached to each slot. It

would be truly remarkable if the scanner were to MIS-select, one by one, all and only the segments belonging to one morpheme, insert them into the wrong slots, and then do the same thing with all and only the segments belonging to the other morpheme involved in the stranding error. For instance when the slots are being filled with the segments involved in the error in (11) (repeated here)

(11) Fancy getting your model renosed. (nose remodeled)

the scanner has in view, at least the unordered set of segments shown here.

m, n, o, a, d, l, z, (n o z, m a d l) (nose, model)

To produce the error the scanner would have to mis-select all seven segments and somehow insert them in slots appropriate for the two morphemes. This is obviously so improbable as to be impossible.¹⁴ Garrett's notion that these errors were exchanges of morphemes and not individual segments was clearly the correct one. But since his model now incorporates the segmental slot notions he will have to relocate the locus of these errors.

This re-location is fairly straightforward it seems. The mis-selection will have to occur when the processor is manipulating major class morphemes as units, without respecting the morpheme's grammatical category. (See below

for why this latter condition must hold.) There appear to be three points in the model where these conditions are met: (i) when the processor is directing the second lexical look-up¹⁵ (ii) when the processor retrieves the items looked up and (iii) when the processor directs segmental slot construction. We argued above that slot construction (or selection) must be on the basis of the form retrieved in the second lexical look-up and hence (ii) and (iii) appear to be indistinguishable. Segmental slot construction on each of the phrasal slots does not involve any reference to grammatical category, meaning, or phonological features associated with the phrasal slots. Therefore it must be done on some sort of order-of-mention basis. It appears that order of retrieval must determine order of slot construction. It is not clear how to determine whether the order of the morphemes retrieved is due to the functional level directing the retrieval in the wrong order or whether they are retrieved in the wrong order for some other reason. What we seem to have then, is morphemes selected in the wrong order at some point before segmental slot construction begins. This may be at (i), (ii), or (iii).

Whatever the precise location is, this change in the model makes one clear prediction. It predicts that we may have segmental errors superimposed on stranding errors because segmental errors will occur after the occurrence of a

stranding error. We are claiming that stranding errors occur when whole morphemes are being manipulated by the processor, prior to their manipulation as segments. This is required because morphemes are now inserted into the phrasal slots segment by segment and as seen above, errors involving whole morphemes become so improbable that they are virtually impossible. Segmental errors will still occur at this locus in the model, because segments (or syllabic constituents) are the type being processed by the production mechanism at this locus. (See below for detailed discussion of segmental errors.) In other words, segmental errors occur after stranding errors, and there is no reason why stranding errors should be immune to them. And indeed it appears that they are not. Consider the following errors, all of which we believe are stranding errors with segmental errors of some sort superimposed on them. (See also the discussion in Chapter 2, Section 7.).

(12) Airforce Cambridge Research Lab	ambridge careforce--uh Cambridge Airforce-- Airforce Cambridge
(13) room and board	bood and rorm
(14) turn the heat on	hurt the team on
(15) spice racks	spack rices
(16) books for Viet Nam	Bomb for Viet Nooks

In (12) we have a beautifully clear case because the speaker corrects both the errors separately, correcting the most recent (according to the model) error first. First the speaker corrects the segmental error involving the exchange of onsets in Airforce and Cambridge. 'ambridge careforce' is corrected to 'Cambridge Airforce'. This exchange of a $\emptyset$ (null) onset for a segmentally filled onset is not uncommon. For instance (17):

(17) a history ideology an istory hideology

Next the speaker corrects the morpheme exchange, 'Cambridge Airforce' becomes 'Airforce Cambridge'. In the old version of Garrett's model, both of these error types occurred when major category items were being inserted into the positional level frame. Since whole morphemes as units were being inserted there was no explanation as to why purely segmental exchanges should occur. With the introduction of the Shattuck-Hufnagel notion we can pinpoint the locus of both types of error, distinguishing the one from the other. If we were to assume that there was only one error involved in (12), and that it occurred upon insertion into the phrasal slots, as would have been the account in the old model, we would have to assume that there was a mis-selection of each of the morphemes, but for some reason one of the segments in one of the morphemes got left behind. If we were to assume

that the error involved the insertion of the morphemes segment by segment into the segmental slots we would face the problem regarding probability discussed above. That is, it would involve numerous mis-selections of segments rather than one mis-selection of a morpheme, all segments being mis-selected except for the one initial [k] in Cambridge. The probability of such an error is zero. Our account is plausible because we distinguish the types. Notice that we see no grammatical morphemes left stranded here because the forms involved include no overt grammatical markers but this exchange of morphemes within the phrase is nevertheless a stranding error, an error in which morphemes exchange irrespective of their grammatical category. If they belong to the same category it is a matter of chance - it will depend on whether there are two items belonging to the same category within the single phrase.

Consider now (13). This could be analysed as two segmental exchanges, one onset exchange ('b' and 'r') and one coda exchange ('m' and 'd'). Coda exchanges are rare compared to peak and onset exchanges (see Chapter 2, Section 6.). What is the likelihood of finding both a coda and an onset exchange intersection in the same two morphemes? Not very likely, we should imagine. On the other hand the error could involve the relatively common morpheme exchange, followed by the relatively common peak exchange. Since both

of these errors occur at different points the intersection of the two in the same words is not remarkable. What we have then, is

room and board

by morpheme exchange

board and room

by peak exchange

bood and rorm

This peak exchange involving V+r is not uncommon. (See Section 6.3 for 'r' analysed as a glide). For instance (18) - (19).

(18) soup is served

serp is souved

(19) fight very hard

fart very hide

In (14) we also seem to have a morpheme exchange followed by a peak exchange.

turn the heat on

by morpheme exchange

heat the turn on

by peak exchange

hurt the team on

Similarly in (13), although this time we have a morpheme exchange which leaves an overt morpheme stranded, followed by an onset exchange.

spice rack	
[+plu]	by morpheme exchange
rack spice	
[+plu]	by onset exchange
spack rice	
[+plu]	by spell-out
spack rices	

This error could also be analyzed as one involving an exchange of two rhymes 'ack' and 'ice', [æk] and [ays], but as discussed in Chapter 2 above, if the rhyme is on an equal footing with the onset, peak and coda, why are the latter three so much more often involved in errors. The rhyme errors are so rare that they are virtually non-existent. The few rhyme errors there are, can have another explanation as is the case in the error in (15). The error in (16) is interesting. In Garrett's model it could be analysed as a word exchange at the functional level, followed by an onset exchange when slots are being filled at the positional level. This appears to be a word exchange rather than a morpheme exchange (i.e. a stranding error) at the positional level, because the plural morpheme is not left stranded. It is comparable to the error in (20). Compare the two.

(16) book for Viet Nam
 [+plu]

Nam for Viet book
 [+plu]
by word exchange
(at Functional level)

(20) Well you can cut
tree in the rain
[+plu]

rain in the tree
[+plu]
by word exchange
(at Functional level)

If (20) were a morpheme exchange at the positional level the error would strand the plural morpheme giving the non-occurrent (21).

(21) *Well you can cut
rain in the tree
[+plu]

(by spellout)

Well you can cut rains
in the tree.

Therefore, the error is a word exchange at the functional level. This can also be used to account for (16). The word exchange here, is then followed at the positional level, by an onset exchange - the most frequent type of error (see Chapter 2).

(16) book for Viet Nam
[+plu]

by word exchange

Nam for Viet book
[+plu]

by onset exchange

bomb [bam] for Viet
Nook
[+plu]

by spellout

bomb for Viet Nooks

This error, like (15) above, could also be analysed as an exchange of rhymes. In this case the error would involve the whole rhyme 'ooks' and the whole rhyme 'am', or the rhyme

'look' and the rhyme 'am' with a subsequent shift of the plural morpheme. That is we could have either (i) or (ii).

- | | |
|------------------------|-------------------|
| (i) books for Viet Nam | by rhyme exchange |
| bam for Viet Nooks | |

In Garrett's model, this is not a possible account of the error because the plural morpheme has not been spelled out at the point when segments are inserted into phrasal frame slots. His model is correct in this regard. Errors which involve coda clusters which are not monomorphemic invariably break up the clusters. (See Chapter 2, Section 10.) For instance (22) and (23).

- | | |
|---------------------|------------------------|
| (22) steak or chops | steak or choks [ɔ̃aks] |
| (23) Jack's pen | Jack's peck |

In Garrett's model grammatical morphemes may not be exchanged with anything; they may only be shifted, i.e. attached to the wrong phrasal slot. Therefore the error account in (i) is not possible; what is required is the account in (ii).

- | | |
|------------------------|-------------------|
| (ii) book for Viet Nam | by rhyme exchange |
| [+plu] | |
| bam for Viet Nook | by shift |
| [+plu] | |
| bam for Viet Nook | |
| [+plu] | |

As mentioned, this analysis gives a status to rhymes in the production model which, on the basis of their rare involvement in segmental errors, they should not be accorded. If Garrett does not accept that syllabic constituents are being manipulated, rather than segments which are unstructured he would have to analyse this as the occurrence of two segmental exchanges 'oo' for 'a' and 'k' for 'm' followed by a shift of the plural morpheme. This would be the analysis in the Shattuck-Hufnagel model. As demonstrated, this analysis just won't stand up because it misses accounting for the fact that segmental errors involve elements of the same syllabic constituent. The analysis in (16) seems the most reasonable as it involves two fairly frequent error types, a word exchange and an onset exchange.

By positing a modular model of speech production Garrett can predict that different error types may involve the same morpheme at different levels. At the outset of this section, we pointed out that by incorporating the Shattuck-Hufnagel segmental slots, his model was no longer capable of accounting for stranding errors as the mis-assignment of morphemes to major category slots in the phrasal frame. This is because the morphemes are now assigned segment by segment to the phrasal frame (or syllabic constituent by syllabic constituent). We find, however, that morphemes as units must be selected before slot construction begins and morpheme

exchanges can readily be accounted for here. At the same time segmental exchanges, which were inexplicable in the old model are readily explained in the segmental-slot-adapted model. The locus of these errors remains the same as it was in the original model - the point at which phrasal skeletons are assigned morphemes. These errors now receive a satisfactory account because the processor in the new model is processing segments (or syllabic constituents) and hence these are susceptible to error.

(vii) Segmental Slots and Segmental Errors

In the old version of the model it was not clear why uni-segmental exchanges would have the same locus as multi-segmental (i.e. morphemic) exchanges. How could a segmental exchange like (24)

(24) top shelf

tof [taf] shelp

and a morphemic exchange like (25)

(25) You'll have to face it
squarely.

square it facely

be considered errors of the same type? At the same time anticipation and perseveration errors like (24) and (25) respectively, were classified as errors of a different type.

(26) Rafferty is a jackass

Jafferty is a jackass

(27) steady state vowels

steady state stowels

We will go over Garrett's reasons for this typology and suggest that in his new Shattuck-Hufnagel adapted model he can not only distinguish morphemic from segmental exchanges, as just shown, but he can also classify exchanges, anticipations and perseverations as errors of one type. At the same time his insights regarding the similarities of segmental and morphemic exchanges can be maintained.

"...sound exchanges and stranding errors are hypothesized to occur in the processes which associate segmental specifications of major category items with sites in the planning frame" (1982, p. 57). In line with his principles of error analysis mentioned above, Garrett identifies these two error types as one type for three reasons:

- (28) (i) both involve major class items,
 (ii) neither honours grammatical category
 and (iii) both involve elements which belong to the same phrase.

It was this constellation of characteristics which led Garrett to include the positional level in his model. It consists of one phrase, into which major class morphemes are inserted - morphemes which are phonologically specified, not

syntactically. Hence an error at this level is insensitive to grammatical category. Both types of errors do, however, show an effect for sound. According to Garrett: "...sound exchanges and stranding errors...implicate various aspects of the segmental, morphological and stress descriptions of their source words" (1982, p. 57). 'Morphological descriptions' are implicated as in (28) (i) - (ii) above. Presumably the reference to 'stress descriptions' is alluding to the fact that in segmental and stranding errors intonation and stress patterns are not deviant. For instance in (29) the correct syllable is stressed,

(29) marsúpiál musárpial

and in (30) the main stress remains in final position, i.e. on the phrasal head, despite its being occupied by the incorrect morpheme.

(30) tight shórts short tíghts

Consider however how segmental descriptions are relevant in the two types of error. There is abundant evidence that segmental exchanges are more common among similar elements like 'p' and 't' than among phonetically different segments like 'd' and 'y'. (See van den Broecke and Goldstein, 1980, Mackay, 1970, among others). There is not one error in confusion matrices based on over 2,000 errors involving 'd'

and 'y' and yet the same matrices yield about 70 involving 'p' and 't' (van den Broecke and Goldstein, 1980). (But cf. discussion in Chapter 2, Section 2.2.) But it's not evident in what way segmental descriptions are shared by morphemes involved in a stranding error. Certainly the errors in our corpus and in Fromkin show no such similarity. In (31) we see several examples of morphemes that have been exchanged.

(31) George	lab
bunny	steak
maniac	weekend
pump	carve

Under no feature analysis could these be considered similar. This means that the processor is sensitive to feature matrices or something correlated with them when segmental exchanges occur; but it is not, when stranding errors occur. This suggests that different mechanisms are responsible for the two types of error.

Both stranding errors and segmental exchanges involve major class items within one phrase.¹⁶ On this basis Garrett distinguishes segmental exchanges on the one hand, from anticipations and perseverations like (32) - (33) on the other.

(32) a Canadian from Toronto	a Tanadian from Toronto
------------------------------	----------------------------

(33) so that I can start the
tape back up

...start the stape...

He also distinguished the former from the latter because in perseverations and anticipations there is "...occasional involvement of closed class category items..." (42 of the 336 in his corpus) whereas sound exchanges "...do not often involve those classes" (1980, p. 195). This class distinction, he concludes "...seems of much lesser importance to the system which gives rise to errors of anticipation and perseveration" (ibid.). Certainly most anticipations and perseverations, like exchanges, involve major class items; in fact 88% of his corpus. Considering the fact that only in an exchange can we be sure of both the elements involved in ~~the~~ error, i.e. "...the confidence we can place in the identification of the two error loci is less (for anticipations and perseverations A.L.) than for exchanges or shifts" (Garrett, 1982, p. 197), it seems rather strong to state, as Garrett does, that in anticipations and perseverations the class category constraint is of 'much lesser importance'. In fact, we find, unlike Garrett, that exchanges, anticipations and perseverations are all equally rare in closed class items.

We analysed our segmental exchanges, perseverations, and anticipations with respect to the open and closed class

(36) It pays to wait waits to pay

We considered an error to be 'adjacent open class' if both the segments (or morphemes) involved were from open class items which were not separated by any other open class items. They could be separated by closed class items and still be considered 'adjacent open class'. An error was considered to be a closed class error if at least one segment (or morpheme) involved, came from the closed class as defined above. A non-adjacent open class error would be (37).

(37) If George Jackson were [ɛkskeɪp]
trying to escape

An adjacent error is exemplified in (38).

(38) stick around and try to trick around and sigh
see to tea

The results of our analysis are found in Tables 3, 4 and 5.

In Table 3 exchanges are compared to perseverations and anticipations. About 90% of anticipations and perseverations involve open class items, compared to the 96% involvement of the open class in exchanges.

Table 3
Segmental Errors in Open and Closed Classes

(a) Exchanges

Open Adjacent	Open	Closed	Total
242 (95%)	3 (1%)	11 (4%)	256 (100%)

(b) Perseverations

Open Adjacent	Open	Closed	Total
74 (78%)	11 (11.5%)	10 (10.5%)	95 (100%)

(c) Anticipations

Open Adjacent	Open	Closed	Total
128 (84%)	11 (7%)	14 (9%)	256 (100%)

This 6% difference would not seem to warrant Garrett's contention that open class membership is of "much lesser importance" in perseverations and anticipations.

In Tables 4 and 5 can be seen the distribution of segmental errors compared to morpheme exchanges (Garrett's

stranding errors) on the basis of these criteria. In both cases 'open adjacent' items are involved most frequently, (87% and 86% respectively) and the closed class very rarely (6% and 7% respectively).

Table 4

Morpheme Exchanges According to Word Class

Open Adjacent	Open Non-Adjacent	Closed	Total
55 (86%)	5 (8%)	4 (6%)	64 (100%)

(i) Errors from Fromkin (1973) class S and from our corpus.

Table 5

Segmental Errors According to Word Class

Open Adjacent	Open Non-Adjacent	Closed	Total
457 (87%)	30 (6%)	38 (7%)	525 (100%)

(i) Errors are classified as open iff all segments are from open classes. They are classified as closed if at least one segment is from a closed class.

(ii) Open-Adjacent errors involve two open class words, where if there is any unit between the two it belongs to the closed class.

That is, Garrett's analysis (i.e. using a 'within phrase' criterion) brings out a distributional pattern which is

common to stranding errors and exchanges. Our analysis brings out this similarity, but it also brings out the similar distribution pattern of anticipations and perseverations.

It is also noteworthy that in all three types of segmental error the interacting elements belong to the same syllabic constituent. That is, they are similarly constrained. Their relative frequency of occurrence in each of the syllabic constituents is also interesting. Onsets are the most common locus for all three types, and codas the least common. See Table 6.

Table 6

Syllabic Position of Different Error Types

Error Type	Onset/Onset	Peak/Peak	Coda/Coda	Totals
Perseveration	(72) 75%	(14) 15%	(10) 10%	(96) 100%
Anticipation	(103) 66%	(34) 22%	(18) 12%	(155) 100%
Exchange	(177) 69%	(49) 19%	(31) 12%	(257) 100%
Totals	352	97	59	508 100%

(i) 19 'V+r' errors included as peak/peak errors (see text, Chapter 2, Section 6.3)

Nooteboom (1980) points out that a number of anticipations may in fact be corrected exchanges. An example of this in our corpus was (39)

(39) I'm a burn victim

I'm a vurn--

where the speaker was aware that he was about to say 'bictim' but didn't. This is another indication that anticipations and exchanges may have the same source in the production mechanism. Perseverations may be different. They appear to involve a relatively high number of onsets and a low number of peaks compared to the other two. They may have a different source, or more than one source for instance. Perhaps they have one source which is the source of exchanges and anticipations and one source which is unique to themselves.

It seems not unreasonable to hypothesize a similar source for all three types - bearing in mind that perseverations may have more than one source. All three types share a number of characteristics listed in Table 7.

Table 7

Characteristics Shared by Exchanges, Perseverations and
Anticipations

- (i) The interacting segments are often similar in terms of phonetic features.
- (ii) They all involve segments which belong to the same syllabic constituent, i.e. they are all constrained by their syllabic position.
- (iii) They all occur most frequently in onsets.
- (iv) They all usually involve elements from open class items and most often, adjacent open class items.
- and (v) They do not honour grammatical category.

What we notice here is that the final two characteristics are those which led Garrett to hypothesize that exchanges and stranding errors were errors of the same type (although he used 'within one phrase'). We are suggesting that anticipation, perseveration and exchange errors all implicate the same mechanism and that this mechanism is not the source of stranding errors since the first three characteristics are unique to the segmental error types. They do not apply to stranding errors.

How then do we explain Garrett's insight-the final two characteristics? We explain it in the same way that Garrett did. They all occur in the construction of the phrasal frame

of the positional level. The difference in the segmental and the morphemic errors is due to their different locus in the construction of this frame. Stranding errors occur when the processor accesses or retrieves morphemes from the lexicon in the phonologically based lexical look-up. The processor misorders the morphemes before they are used to construct (or select) the slots in the phrasal slots available for these morphemes. In Garrett 1980, before the adoption of the Shattuck-Hufnagel ideas, the locus of the stranding error was the actual insertion of the morpheme into the major class slots. Morphemes were selected in the wrong order for insertion. In the new version of the model this can't be right because segments are hypothesized to be inserted at this point - not whole morphemes. But the change is minimal, that we suggest. We showed that the morphemes selected from the second look-up must direct segmental slot construction. The processor as before, still selects morphemes, and selects them in the wrong order. The difference is that now, in the new model, the processor is selecting morphemes not for direct insertion into phrasal slots, but rather in order to direct construction of the segmental slots in the phrasal slots. This particular selection was not required in the earlier version of the model because there were no segmental slots. Garrett's observations regarding the within phrase and open class category constraints which stranding errors

obey, and their disregard for grammatical category, are still accounted for. Lexical access and retrieval in the old model were preparatory to, and limited by, the phrasal frame under construction. That is, the morphemes that were selected were those that were to be inserted in a single phrasal frame. This is still the case. But now the morphemes are not inserted, they direct slot construction.

The segments of these same morphemes are then inserted into the slots in the phrasal frame. When the processor is selecting these segments (i.e. syllabic constituents) it may mis-select, and what will result will be a perseveration, an anticipation or an exchange of similar syllabic constituents. The fact that these errors are also characterized by the qualities associated with stranding errors is also explained. They are due to the fact that the syllabic constituents being manipulated are those that come from the open class morphemes which have been selected for construction of the positional level phrasal frame.

(viii) Conclusion

In this section we have considered how the segmental slots of the Shattuck-Hufnagel model could be integrated into the Garrett model without a loss in empirical coverage. We have discussed at what point in the model segmental slot

construction must take place and pointed out what changes were concomitantly required to maintain empirical coverage. We have seen that segmental slots must be built, or selected, on the basis of the second lexical look-up if the model is to maintain its capacity to account for malapropisms and stranding errors. We have also demonstrated that the locus of stranding errors in the model must be minimally shifted and distinguished from the locus for segmental errors. By doing this we find that Garrett's insights regarding the morphologically based similarities in the two types of error can be maintained, and at the same time we can explain why certain characteristics of segmental errors are unique to them. Garrett claims that morphological, stress and segmental descriptions are all implicated in both stranding errors and sound exchanges. We find that segmental similarity is not a characteristic of morphemes that exchange, unlike substitutions of morphemes (malapropisms). With regard to the fact that stress is never deviant in these two error types, this applies with equal force to anticipations and perseverations. It also applies to word exchange errors, errors which occur at the functional level. Therefore this observation is not useful in classifying these error types. That is, non-deviant stress can not be used as a diagnostic for these errors because it is not unique to these types of errors.

If Garrett's LRs are in fact syllabified as appears to be the case, one of the two major problems we found in the Shattuck-Hufnagel model is solved. If the processor is assigning syllabic constituents to the slots in the phrasal frame and not unstructured segments, Garrett's model is capable of accounting for the syllabic constituency constraints on errors of anticipation, perseveration and exchange. The other major problem with the Shattuck-Hufnagel model remains however. Why are segments separated from their prosodic structure only to be re-integrated? This separation will account for the facts, i.e. it will be able to describe the facts, but it receives no independent motivation and hence is not explanatory. The mystery remains.

Having seen in broad terms how our hypothesis fits into this model, to end this section, let us consider the integration in more detailed terms. As mentioned, studies of malapropisms, stress errors, TOT states, word recall and recognition all point to lexical items in the production lexicon including syllable structure and stress. (Indeed, from within the competence domain, Selkirk argues that foot structure (hence syllable structure) is lexical. (See Chapter 1, Section 3.3.)) Garrett incorporates these findings into his model hence syllable structure is available to the processor when segmental errors occur.

Secondly Garrett's insights regarding the class of items involved in segmental exchanges and stranding errors we find apply also to perseverations and anticipations. Moreover his hypothesis that such errors occur when major class stems are inserted in the phrasal frame coincides nicely with our observation that roots are the domain of the syllabification rules. Segmental errors do not involve affixes - and root final consonants behave like codas - even before vowel initial affixes.

What of our distinction regarding syllabic constituents, onset, peak, and coda and the sub-constituents of codas and onsets? Is there any point within the model at which this distinction might hold? On our assumptions these errors must occur at different levels of processing. Can we distinguish these levels within Garrett's model? We believe that we can. You will recall Shattuck-Hufnagel's findings regarding onset clusters. Eighty-six percent (86%, $n=36$) of them involved entire clusters and they were all word-initial. Our data was similar. We had no movement of entire word-medial onset clusters. Shattuck-Hufnagel suggested that the high proportion of onset errors in word-initial position is indicative of these onsets having a special lexical status. Lexical access studies have shown that word-initial and word-final positions are more important than word-medial

positions. Consider these observations in relation to Garrett's model, modified to include the segmental slots proposed by Shattuck-Hufnagel. We have shown that morpheme exchanges ('stranding errors') must take place just prior to insertion of lexical items in the segmental slots, either due to morphemes being accessed or retrieved in the wrong order. Suppose that the items so retrieved are accessed or retrieved with word-initial and final segments figuring prominently. It could be that the processor retrieves these morphemes as sequences of syllabic constituents, onsets, peaks and codas. Within Garrett's model the processor could then select each of these constituents for subsequent mapping onto the segmental slots, sub-constituent by sub-constituent. In the selection of the syllabic constituents supplied after lexical retrieval, the processor could mis-select, leading to an error involving entire onset constituents. That is, constituents dominated by the root syllable node, and in particular those in word-initial or final-position are available for selection or mis-selection by the processor prior to their insertion in the segmental slots of the phrasal frame. This can result in errors like (40) (a) - (f).

- | | |
|--------------------------------------|---------------|
| (40) (a) <u>central</u> <u>choir</u> | kwentral sire |
| (b) <u>highway</u> <u>driving</u> | dryway hiving |

- | | |
|--|--------------------|
| (c) <u>s</u> queeze <u>f</u> resh
oranges | freeze squesh |
| (d) west <u>e</u> rn) Mex <u>i</u> co) | Wex(ern) Mest(ico) |
| (e) clutch <u>f</u> irst <u>d</u> own | ...[fərč]... |
| (f) next ten minutes | nen text... |

Within Garrett's model, the processor then must read the sub-constituents in order to build the segmental slots of the phrasal frame. If it mis-reads or mis-selects the sub-constituents it will result in building a slot for a segment in the right sub-constituent position but the wrong root-dominated position, and errors like (41) (a) - (f) will result.

- | | |
|--|-------------------|
| (41) (a) spl <u>i</u> t br <u>a</u> in | sprit blain |
| (b) pounds and fr <u>a</u> ncs | prounds and fancs |
| (c) Cathy Squ <u>i</u> res | Cwathy... |
| (d) top sh <u>e</u> lf | tof shelp |
| (e) sink a sh <u>i</u> p | simp a shick |
| (f) the toast <u>a</u> will pass | ...toas...past |

In Garrett's model, then, there is easily a way of distinguishing the two loci. There seems to be a need for recognizing initial word segments (and segment sequences) (and perhaps final segments) as special in some sense - in particular they appear to be special when lexical look-up is taking place. Thus word-initial and final (cluster) errors

can be attributed to this level. Similarly, since Garrett wants to incorporate segmental slots into his model we have a natural place to account for onset and coda sub-constituent errors since sub-constituents are filled by segments as are the segmental slots. Thus it appears that our hypothesis that segmental errors are syllabic constituent or sub-constituent errors, these constituents being defined on roots, can be incorporated into Garrett's model. But we have only achieved an observational level of adequacy. We have not explained why the processor is manipulating syllabic constituents when it could be processing entire lexical items. Shattuck-Hufnagel attempted to explain this as we saw in the preceding section but failed. By incorporating segmental slots into his model Garrett faces the same problem. Lexical items are retrieved after the phonologically based look-up; they direct the construction of slots losing their serial position as they do so - then these same items are mapped directly onto these slots and the result is segmental errors. There remains no motivation for this aspect of the model with regard to normal processing; it is completely and most unsatisfyingly ad hoc.

We will next consider a model which does not make use of these segmental slots - hence does not fail in this way.

4. The Fromkin Account of Segmental Errors

According to Garrett, his model "...differs only in detail from the outline in Fromkin (1973)" (Garrett, 1982, p. 70). Indeed, segmental errors are attributed to very similar loci in the two models. In Fromkin's model, meaning and sound based word substitutions, and segmental errors (anticipations, perseverations and exchanges) all occur in her fourth stage. The input to this stage is a linear sequence of morphemic slots with their syntactic and semantic features specified. This slot sequence has already been assigned phrasal stress and an intonation contour. At this point the semantically structured lexicon is accessed on the basis of the features attached to each slot. This sub-lexicon specifies the addresses of the phonologically specified morphemes. That is, each constellation of semantic and syntactic features directs the processor to a particular item in "...the semantic class sub-section..." (p. 239) - this item then points to a particular segmental string found in a phonologically organized "...overall vocabulary..." (p. 239). In a semantically related substitution- such as 'racquet' for 'bat' the wrong address has been found due to reading the wrong but related set of semantic features.¹⁷ In a sound related substitution such as 'present' for 'pressure' the correct address is obtained from the semantic lexicon but the selection in the phonological lexicon is incorrect. A

phonological neighbour is selected by mistake. Segmental and multi-segmental errors occur after the items have been selected from the overall vocabulary. The items must then be placed in a buffer, in the word slots waiting for them. It is in this transfer of items to the frame in the buffer, that segmental errors occur. This is similar to the Garrett model before the incorporation of the Shattuck-Hufnagel segmental slots. There are differences, however. Although grammatical morphemes such as plural have not been spelled out the closed class vocabulary is not in general distinguished from the open classes. That is, both classes are inserted here. Although Fromkin doesn't specify either an algorithm for the syllabification of lexical items, or a specific structure internal to the syllable, as discussed in Chapter 2, she is aware of the syllabic constituency constraints on segmental errors. In fact she claims that "the only examples in her data.." (p. 228) which are not thus constrained are two within-word metatheses. (In fact we found a few more than this in her data, perhaps due to differences in our syllabification algorithms.) Because of this she believes that each segment in a vocabulary item is specified both for its feature composition and for its syllabic position. As she says, she is not concerned "...with an explanation of why and how this occurs..." (p. 240) but as the segments are transferred to the buffer "...segment 1 of syllable 1 may be

substituted for segment 1 of syllable 4" (p. 240) and so forth.

It is also at this point that errors involving sequences of segments occur, whether these sequences involve syllable sequences, single syllables, or parts of syllables like CV and VC sequences. (Presumably Garrett's stranding errors would be included here.) Subsequent to the insertion of segments, morphophonemic rules apply, spelling out items like 'plural'. Finally the low level allophonic rules apply. We see that with regard to the localization of segmental errors, Fromkin's and Garrett's models are indeed similar, hers lacking certain details like the distinction between open and closed class vocabulary, and the Shattuck-Hufnagel segmental slots. This lack means her empirical coverage is not as great. She cannot account for the characteristics which stranding errors and segmental errors have in common (see Table 7, Section 3, above). Neither can her model explain why errors which involve certain segment sequences are relatively frequent whereas others are very rare. Segment sequences are commonly involved in errors only if these sequences are constituents of some sort - either morphemes, words (or 'pseudo-morphs' (Garrett, 1975)) etc. or syllabic constituents. In this model, segments from a certain syllabic position are being transferred one by one to the buffer and an error involving a whole morpheme exchange such

as (42) is almost impossible. (See above Section 3 and footnote 14 regarding the improbability of all these segmental errors intersecting.)

(42) Both kids are sick. Both sick's are kid.

Fromkin refers to "segment 1 in syllable 1" but realizes that the implicated unit involves syllable position (see p. 228). As we have seen, syllabic constituents are implicated, as evidenced by the fact that onset clusters often exchange with single segment onsets as in (43) - (44) or with other onset clusters as in (45) - (46).

(43) pedal steel guitar	stedal peel guitar
(44) heater switch	sweeter hitch
(45) squeaky floor	fleaky squoor
(46) when you're old your spine shrinks	...shrine spinks

'Segment 1 in syllable 1 would not exchange with segment 1 in syllable 4' if the segments were from different syllabic constituents. An error involving a peak and a onset exchange is impossible and one involving an onset and a coda is very rare, as we have seen. Fromkin's model could account for syllabic constituency constraints since LRs are syllabified and the errors occur when LRs are inserted into the phrasal slots of the frame in the buffer, and her model doesn't face the problem faced by the Shattuck-Hufnagel and Garrett

models. Namely, why a segment loses its serial position only to be re-assigned it. However, Fromkin's model, the Shattuck-Hufnagel and Garrett models share one problem. In all three, morphemes, although they are retrieved from the lexicon as units, are inserted into the phrasal frame in pieces. In the Shattuck-Hufnagel and Garrett models, the morphemes must be inserted bit by bit, because when the slots are built the segments lose their serial position; that is, the piece-by-piece insertion is motivated. What is not motivated, however, is the initial separation of segment and slot, concomitant with slot construction. The scrambling is unmotivated. In the Fromkin model there is no segmental slot construction and hence no scrambling but as a result there is no motivation for inserting the morphemes into the phrasal frame piece by piece. In all three models this piecemeal insertion is the means whereby speech errors which involve fragments of morphemes may be accounted for. Fragments which must begin in their correct positions in the lexicon somehow end up in the wrong position. An analysis of the initial, correct positions and the final, wrong positions of these fragments indicates that the beginning and end positions are always the same if one analyses them as positions in syllables. That is, a segment which begins in an onset position, for instance, ends up in an onset position. Fragments are in position, despite being out of position.

This suggests that the fragments are syllabic constituents and not merely unstructured segments. The locus of the errors is therefore at some point when syllabic constituents are being processed. For an explanatory account of the errors one would like to motivate the processor's manipulation of syllabic constituents at this point. In the three models, the locus of the errors is assumed to be when morphemes are inserted into phrasal slots. Why would the processor select syllabic constituents or segments at this point rather than whole morphemes. There appears to be no reason for this except for the ad hoc one of describing the speech error data.

Crompton's account of segmental errors does not fail in this regard; segmental errors involve syllabic constituents because they occur when the processor is involved in a syllable look-up, and the syllable is located on the basis of its onset, nucleus and coda.

5. Crompton's Account of Segmental Errors

Crompton (1981), like Fromkin, recognizes the syllabic constraints on segmental errors¹⁸ and argues that these are a reflection of the fact that the articulatory program includes a set of routines, and these routines correspond to syllables. The addresses of these routines, in the library

of routines, are defined by the feature specifications of each of the syllabic constituents, onset, nucleus and coda. (His syllable template in fact includes a rhyme constituent but he doesn't specify how this constituent fits into his model at all - it seems to be his compromise to the theoretical models of competence grammars.) This explains why lexical items are being processed as sequences of syllabic constituents, and why errors involve these constituents. The processor must read individual syllabic constituents rather than whole morphemes because these constituents are used to locate syllable routines.

The retrieval of the routines is characterized by three stages: addressing, activation and incorporation. Segmental exchanges, perseverations and anticipations occur at stage one, the addressing stage. The processor reads the feature complexes in the onset, nucleus and coda and activates sub-areas in the library corresponding to each of these. For example, in locating the routine for the syllable 'bat' all syllables with onset in 'b' would be activated, all syllables with nucleus 'a' would be activated, and all those with coda 't' would be activated. Each of the sub-areas so activated would include as one member the syllable 'bat' which would achieve a required threshold level on being activated a number of times. On reaching threshold it would be available to the articulatory programmer for incorporation. The

incorporator will sequentially order all syllables so activated, apply the rules which will produce the allophonic variants of segments in the environments of other syllables, and incorporate the correct prosodic pattern.

Errors occur in the following manner. The processor reads lexical items in parallel and the onset condition (i.e. the feature complexes) of one syllable may 'contaminate' (p. 683) the onset condition for another syllable. Similarly for the nucleus and coda conditions. When this happens we have a perseveration, exchange or anticipation. In all three cases the misreading of one of the conditions will lead to the activation of the wrong syllable. For this reason Crompton says that segmental errors are in fact syllable errors. More accurately, the result of misreading a syllabic constituent is the exchange, perseveration or anticipation of the wrong syllable.

Although Crompton has motivated the processing of syllabic constituents there are various problems with his model - problems of both an empirical and a conceptual nature.

The articulatory program is the program which gets sent to the articulators and once the syllable routines have been activated ~~also~~ all phonological computation has been completed. In particular grammatical morphemes such as

plural have been spelled out. But there is abundant evidence that spell-out must follow the occurrence of segmental errors and yet in this model spell-out precedes the addressing of routines, which is the locus of segmental errors. There is also good evidence that segmental errors are largely confined to open class items (see above, Section 3.) but this model predicts that these errors will be randomly distributed across all morphemes. The job of the articulatory programmer is to take a text which is specified in 'classical phonemic' terms and convert it into a string of syllable routines, thus the text includes a classical phonemic specification of all morphemes. Crompton explicitly claims that grammatical morphemes are spelled out when the text is constructed "once the shifting and stranding errors have taken place.." (p. 706). He appears to be unaware that grammatical morphemes accomodate not only to environments due to stranding and shift errors such as (47) - (48)

(47) It pays to wait
[z]

It waits to pay
[s]

(48) Ralph's and my
[s]

Ralph and my's
[z]

respectively, but also to the environments provided by segmental errors such as (49) - (50).

(49) a history ideology

an istory hideology

of the syllable as a static organizing frame is correct and that syllables are not processed in the way that Crompton envisages.

Conceptually too, Crompton's model is less than satisfactory, in fact a little overwhelming. The retrieval of a single syllable routine involves an incredible amount of activation - most of which is activation of incorrect syllables - before the required threshold is reached. Besides the onset, nucleus and coda conditions there is also a reduced or unreduced syllable condition, attached to every syllable. This isn't part of the nucleus condition because Crompton believes that reduced and unreduced syllables are stored relatively far apart in separate libraries. Therefore to reach threshold for a syllable like 'bod' all unreduced syllables in the language are activated (including 'bod') then all syllables with onset 'b' (including 'bod'), then all with 'o' nucleus (including 'bod') and finally all syllables with 'd' coda (including 'bod'). That is, to reach threshold, 'bod' is activated four times and all unreduced syllables of English at least once and some of these more than once. This aspect of the model is less than parsimonious although something along these lines may provide an explanation for the fact that segmental errors more often occur when a particular segment recurs within a single phrase. Recall that Crompton's lexical items are addressed

in parallel, so if three items are addressed at once we will have all syllables in the language activated three times, and if two of these items include a nucleus in 'u' we will have a double activation of all syllables in 'u' and so on.

Crompton recognizes the extent of the empirical coverage of the Garrett and Shattuck-Hufnagel models and attempts to adapt their models to his own. It seems to me that any adaptation is impossible. He suggests that Garrett's positional level representation should be identified with his articulatory program. Segmental errors²⁰ are hypothesized to occur when segments from major class items are inserted into the segmental slots of the phrasal slots. Crompton feels that the phonologically based lexical look-up that precedes this insertion should be equated with his addressing stage. But this is impossible, given Garrett's clearly insightful distinction between closed and open classes at this point in the model. The articulatory programmer can not omit closed classes and yet at this point in Garrett's model only open classes are being processed. Crompton also argues that Shattuck-Hufnagel's scan-copier "...IS (his emphasis) the mechanism by which addresses are generated" (p. 703). He argues that "there is no need to assume any independent computation of slots: all the necessary information is available in the text" (ibid.). But the slots are a crucial part of Shattuck-Hufnagel's model. She explicitly rejects

the notion that syllables are discrete units open to manipulation. Syllables form the static organizing frame for utterances - they are a part of the skeletal prosodic structure. They are not independent units subject to disorderings. The Shattuck-Hufnagel model makes the correct predictions in this regard whereas the Crompton model does not. She rejects what would seem to be the central core of the Crompton model and any attempt to integrate the two seems doomed. The Crompton model is different in critical respects from both the Garrett and Shattuck-Hufnagel models and in attempting to adapt their model to his, Crompton wipes out the major insights of both. The empirical coverage enjoyed by both models is lost in the adaptation and Crompton is misguided in believing that "A consequence of this (adaptation) is that whatever success can be claimed for Garrett's model in accounting for speech errors can now also be claimed for (his)" (p. 699).

6. Conclusion

It would appear that the Garrett account of segmental errors²¹ has greater empirical scope than the Fromkin, Shattuck-Hufnagel or Crompton accounts. Each one has their own strong points relative to one another and their own weaknesses. The main problem with the Garrett and Shattuck-Hufnagel accounts is their use of completely

unmotivated scrambling; Fromkin doesn't scramble but she doesn't motivate the transfer of morphemes segment by segment to the buffer. All three of these fail with respect to the device which they employ to account for the fact that segments which must begin in the correct serial position in the lexicon, for some reason can lose that position. Crompton's account is appealing in that he motivates the device he uses. He gives good arguments in support of his notion that the programmer has access to a set of routines and that these correspond to syllables. This will account for the fact that what appear to be syllable-based allophonic rules are never violated in utterances. His hypothesis is also supported in that the best speech synthesis results have come through the use of syllables as the concatenative unit. (See Fujimura and Lovins, 1982).

As we have seen, there are numerous problems with his account, however - the main ones being his two predictions: (i) that spell-out occurs before the occurrence of segmental errors and (ii) that syllables themselves will be subject to misordering. For Shattuck-Hufnagel, syllables are an organizing frame and so they are not subject to movement. It would be ideal if we could give an account which amalgamated the insights of these four theorists and concomitantly discarded their particular failings. We should like to outline a possibility regarding such an amalgamation.

First, in Table 8, we will specify the characteristics of segmental errors that our hypothesis can observationally account for (our hypothesis being that segmental errors are in fact the mis-selection of constituents of the same type, and these types are the syllabic constituents as defined by the rules in Figure 1, Chapter 2, Section 5 and the sub-syllabic constituents as defined by the modification to these rules as defined in Figure 23, Chapter 2, Section 9 and Figure 36, Chapter 2, Section 10). Secondly, in Table 9, we will specify those further characteristics we may account for by integrating our hypothesis into Garrett's model, and the concomitant loss in empirical scope. Thirdly we will specify those segmental errors which remain unaccounted for (in Table 10), and suggest the direction that an account of some of these errors might take. Finally, we will suggest in the broadest terms how we might characterize the normal speech production mechanism such that this observationally adequate account becomes 'descriptively' adequate.²² To do this we must explain why it is that the processor selects syllabic constituents rather than whole morphemes when these are available to the processor, provided as units by the lexicon. That is, we will suggest a reason for the splitting apart of these 'holistic' forms into syllabic constituents.

Table 8²³Segmental Errors: Facts Accounted for By Our Hypothesis

- (i) they do not occur in affixes
- (ii) they involve the interaction of segments from identical syllabic constituents
- (iii) if sequences of segments are involved in an error these sequences constitute a syllabic constituent; they do not cross syllabic constituent boundaries
- (iv) sequences of segments may interact with single segments, both of them, the sequence and the single segment, being dominated by the same syllabic constituent node
- (v) segments may turn up in positions where no segments were intended; these segments always move into syllabic constituent or sub-constituent positions identical to the one they have vacated
- (vi) phonotactic constraints are not violated
- (vii) segments that interact in errors may be completely dissimilar in terms of the feature complexes which characterize these segments
- (viii) segments which interact in errors will often be alike

Table 9

Integration into Garrett's Model: Gains and Losses in
Empirical Coverage

A. Gains:

- (i) most (93% of ours) segmental errors occur in open class items
- (ii) most segmental errors (87% of ours) occur in 'adjacent open class' items²⁴
- (iii) segmental errors do not honour grammatical category
- (iv) segmental errors share these characteristics (i - iii) with morpheme exchanges - i.e. Garrett's 'stranding errors'

B. Losses:

- (i) closed class items may be affected (7% of ours) (usually deictic morphemes, adverbs, pronouns, or prepositions)

Table 10Errors 'Incorrectly' GeneratedA. Not Generated:

- (i) some within-word errors
- (ii) errors which seem to involve features - (some of Fromkin's Class L)

B. Incorrectly Generated:

- (i) a number of within-word errors

C. Overgenerated:

- (i) errors involving the sub-constituent obstruent position in codas and onsets

Although we extend Garrett's account of segmental exchanges to include perseverations and anticipations, by integrating our hypothesis into his model, we can predict the locus of only about 93% of segmental errors because we lose our account of errors which involve closed class items (mainly deictics and prepositions). These errors also obey the syllabic-position-constraint, hence would seem to involve the same type of processing as major class syllabic constituent errors. Garrett suggests that prepositions are major class (i.e. 'open class') items at the functional level but not at the positional level because they are involved in word exchanges at the functional level (like nouns, verbs, for instance) but they are not involved in morpheme exchanges (i.e. 'stranding errors') or segmental exchanges at the

positional level. With regard to the latter he is wrong. They can be involved in segmental exchanges. For instance (51) - (52):

(51) pass out

pat ous

(52) see through the wall

three sue...

They also take part in anticipations and perservations, (53) - (56) respectively.

(53) with a brush

wish a brush

(54) we'll put your trash
cans back

...bans back

(55) been away

been abay

(56) tag along

tag alog

These errors suggest that a thorough investigation is required to determine what items belong in the 'closed' classes and 'open' classes in production. This could lead to an accurate account of these errors.

The errors which pose the major stumbling block for our hypothesis are the within-word errors. As illustrated in Chapter 2 (Section 12.) most of these obey the syllabic-constituent-constraint but only if the word is taken as domain. We pointed out in the introduction to this chapter that the processor distinguishes roots from both derivational and inflectional affixes. It is clear that the

domain of syllable-based allophonic rules is not limited to roots. That is, affixes must be integrated into the syllable structure of words before these allophonic rules apply. It may be that the processor distinguishes derivational from inflectional affixes and that there is a level of processing in which derivational morphological rules apply and the word is the focus of attention. It could be at this level of processing that within-word errors occur, and possibly word blends. As we pointed out in our analysis of word blends (Chapter 2, Section 11) we could not determine clearly whether roots or words were the implicated syllabification domain but since most blends involved the juxtaposition of whole words (including affixes) it seemed reasonable to assume the word as domain.

To really come to grips with what is going on with regard to 'closed' class items, word blends and within-word errors would seem to require a richly developed theory of the morphological component of the production grammar. This is an area beyond the scope of this thesis but it would seem to be where one might find the solutions to these errors. Unfortunately we have no specific suggestions regarding the other problematical errors listed in Table 10.

We have now come to the point where it is incumbent upon us to attempt to explain why the processor is selecting

syllabic constituents for insertion into the positional level phrasal frame rather than whole morphemes. It seems to us that the major insight of the most recent version of the Shattuck-Hufnagel model is its recognition of the necessity for a prosodic structure frame. Such a structure seems mandatory for handling both segmental and suprasegmental facts. At the same time the major shortcoming of this model is the requirement that segments lose their serial position after directing the construction of this frame, only to be reinserted in the structure they have just caused to be built. We would like to suggest that such prosodic structures are not built - they are available for fleshing out. The bottom level of these empty prosodic frames consists of a string of empty syllabic sub-constituents, each one waiting to be fleshed out, or not, by the segments which will make up the utterance. We suggest then that in the generation of every utterance the processor must assign content to strings of syllable templates. We suggest that Garrett's positional level phrases consist of a string of such templates; how many syllables belong in the string is an empirical question. (Nooteboom (1969) has suggested that errors often involve segments within seven syllables of each other so perhaps these positional level 'phrases' - i.e. strings of syllable templates are 7 syllables long.)

We suggest that a fully fleshed out prosodic structure is interpreted for articulation. Selkirk (1980) suggests that her prosodic structures may be the units of production and perception. We are borrowing this notion, but believe that basic templates are provided for filling in. Clearly this is not specific to English processing. A preliminary look at the errors in Meringer (Meringer and Mayer, 1895) suggest that they too are subject to the syllabic-position-constraint.

We have specified the structure of the English syllable template as consisting of three root dominated nodes, peak, onset and coda, the latter two containing sub-constituent structure. Clearly a number of languages will not include either the coda constituent or onset substructure in their template. The English facts lead us to speculate however, that no language will include a rhyme in their production grammar template. English even though a quantity sensitive system, shows no evidence of a rhyme constituent, hence we speculate that it may not show up in any language. We could speculate further that higher levels of prosodic structure might consist of just such empty templates waiting to be fleshed out - an area for future research.

This notion of 'fleshing out' means that a good deal of processing involves selection and subsequent interpretation,

rather than building on the basis of interpretation. One wonders whether this will hold up under deeper investigation. We have thus borrowed Shattuck-Hufnagel's notion of a static organizing frame discarding the necessity for scrambling. This in turn extends Garrett's empirical coverage.

In Garrett's slot-extended model, blends must occur before the construction of the segmental slots otherwise the scanner-copier could not match segments (or syllabic constituents) to segmental (or syllabic constituent) slots. They must occur after the construction of the phrase is directed by the functional level however, because there presumably are no subcategorization facts associated with the products of blends, and without these facts the phrases of the positional level can't be built. Hence blend errors in this new slot model must occur before slot construction. This is unlikely given that they involve the mis-selection of syllabic constituents. They must occur when such items are being selected - i.e. after 'stranding errors'. Garrett notes that these errors must occur relatively late, i.e. after the phrasal frame is constructed. By taking the Shattuck-Hufnagel slots into his model they have to be moved forward. If we assume that no prosodic structure building goes on then these errors need not be pushed forward and the fact that they, like segmental errors involve the

mis-selection of identical syllabic constituents can be accounted for.

We retain Crompton's insights regarding syllable routines without the problems. Syllables are not misordered, because they are not selected - they form part of the organizing frame as Shattuck-Hufnagel suggests. If as Crompton argues the syllable is the 'routine', or as Fujumura and Lovins term it, the 'concatenative unit', accessed in production, it need not be misordered. The routine IS the routine specified by the set of features dominated by the syllable node. The strings of templates are strings of 'potential' routines.

This is only the beginning but we seem to be heading in an acceptable direction. It would require very little experience to determine the characteristics of your language's production grammar template. This would provide part of the organizing frame for utterances and part of the organizing frame for lexical items.

We will speculate no further. The account seems at least plausible, and it incorporates the major insights of all the models without their failings.

NOTES TO CHAPTER 3

¹It was suggested to us that speech error theorists should consider speaker intuitions regarding speech errors, as part of the data they should account for. That is, the suggestion was, that for a speech model to reach a level beyond observational adequacy it should determine what speakers judged to be possible and impossible errors and account for these judgements. This is analogous to a competence theorist being required to account for speaker judgements regarding what is and what is not a possible sentence in some language. It is not at all clear to us in what way such judgements would be enlightening with regard to the units and/or processes of production. Firstly, presumably there is no 'grammar' which generates sets of possible errors whose output speakers could have intuitions about. Speakers have no privileged access to the rules and representations of their grammars, only the output of these rules. Secondly, competence grammarians use speaker intuitions regarding possible sentences to determine which sentences their grammar should generate and which ones it should not. If speakers were to consistently judge some error as 'impossible' that had in fact often occurred, the production theorist could not, on this basis, decide not to allow his model to generate that error. The production modeler cannot idealize away performance phenomena - this is his domain of interest.

It is not at all clear to us what would be suitable criteria for attributing descriptive adequacy to a production model, but it seems unlikely that speaker intuitions regarding possible/impossible errors will figure in these criteria.

²Presumably she didn't distinguish ambiguous part errors from whole errors. But it is interesting that she has such a low percentage of part errors. We have 45% (n=63) of our set of onset cluster errors that are unambiguously part errors. See Chapter 2, Section 9. See below for discussion.

³In fact she finds 18 but 8 of them are V + liquid and should be classified as peak errors. See the discussion above in Chapter 2, Section 6.3.

⁴As discussed in Chapter 2, Section 7, Kozhevnikov and Chistovich (1965) do define an 'articulatory syllable' which consists of all C's up to the preceding V and the following V. For instance 'thank you' would be made up of two 'articulatory syllables' 'tha' and 'nkyou'.

⁵In presenting Garrett's model I have looked at the '76, '80 and '82 versions and amalgamated them.

⁶In fact as segments were being assigned their syllabic role, the errors occurred.

⁷He uses observations from hesitation studies and experimentally induced speech errors in normals to support the model so derived. He also uses the data from pathological populations as a testing ground for his model.

⁸He suggests that this level might also include a third level, one which determines detailed phonetic form.

⁹He suggests that the functional level may be like an Aspects-type deep structure, but the comparison is less than exact because no serial order is imposed at this processing level. His model would fit better into a Stowell-type (see Stowell, 1981) framework where verbs are subcategorized for their arguments but the order of arguments is determined not by the verb's lexical entry, but by the universal principles of the government, binding and X-bar theories.

¹⁰Prepositions are considered major category items at this level as they are syntactic heads. They are considered closed class items at the phonological positional level.

¹¹Spell-out must be after the error occurs but in fact attachment need not be. It depends on what the morpheme is attached to. If attachment is to a particular morpheme then attachment must occur after the error. Otherwise the error would be (ii) It wait to pays. However if attachment is to a particular slot rather than to a particular morpheme attachment could come before the occurrence of the error. Spell-out, not attachment, must occur after the insertion of the verb.

¹²See below for discussion of lexical retrieval.

¹³See Liberman and Prince, 1977.

¹⁴The probability of each mis-selection, times the probability of the intersection of all these, times the probability of inserting each in the spot which will produce the two morphemes.

¹⁵In Garrett 1980 Garrett believed that this look-up involved grammatical category but in Garrett 1982 he has come to believe otherwise.

¹⁶Certain stranding errors, Garrett believes, do occur across phrases. He suggests that the locus of these errors is the functional level where word exchanges occur. An example he gives of such an exchange is (i) which knows you gotta m--means you gotta know. It's not clear why in some cases you get stranding at the functional level and in other cases you do not as in example (20) above.

¹⁷It is not clear whether this is due to a misreading of the features attached to each slot or whether it is due to a misreading of the feature constellations in the semantic sub-lexicon. Moreover it's not obvious how the two could be distinguished since the processor goes directly to the semantic sub-lexicon on the basis of the feature descriptions attached to the slots.

¹⁸Also like Fromkin, Crompton doesn't specify a syllabification algorithm.

¹⁹It's not clear how he would distinguish errors involving entire clusters from those involving only parts of clusters. Presumably it can misread an entire set of features or only a few features.

²⁰Crompton mistakenly believes that Garrett considers perseverations, anticipations and exchanges to all be errors of one type. In fact he (Garrett) distinguishes the latter from the former two. (See discussion Section 3 above).

²¹Assuming he would consider anticipations and perseverations as occurring at the same level as segmental exchanges.

²²As sketchily characterized in the Introduction to this Chapter.

²³See Chapter 2 for details as to how our hypothesis accounts for these facts.

²⁴See this Chapter Section 3 for definition of 'adjacent open class'.

CONCLUSION

The basic assumption that the grammar lies at the root of our capacity to use language was the motivating factor behind the investigation in this dissertation. Phonological theory in recent years has been primarily concerned with the nature of phonological representations and the syllable has figured prominently in the questions that have been raised. There have been numerous arguments for its inclusion as a unit in phonological theory, and suggestions as to how it is to be incorporated, and the kind of structure it is, have been diverse and numerous. Given the convincing arguments for assigning it a prominent status in phonological theory it was reasonable to expect that it would play an equally important role in production. Many of the constructs proposed by competence grammarians have shown up in speech error analysis. Hence we expected that speech errors might point towards the syllable as a processing unit. Moreover we hoped that speech error data might help us both to select some structure from among the many divergent notions of competence theorists, and to determine where in a production model that syllable structure was assigned. Indeed we discovered that speech errors provided insights in both these areas, as well as pointing to the necessity for its inclusion in a model of production.

As we saw in Chapter 1 there appears to be no characteristic of syllable structure that is universally accepted by all competence theorists, except perhaps the very basic requirement that it dominate some syllabic (or potentially) syllabic element. Similarly there is no agreement on when syllabification occurs nor on its domain of application, nor on the external boundaries of the syllable, nor on the basis for determining these external boundaries. The structure that we derived on the basis of sub-morphemic segmental speech errors was not isomorphic in structure to any of the competence proposals but was a composite of characteristics: all of them individually suggested by one or another theorist.

We argued for a trinary branching root syllable node, but the nodes dominated onset, peak, and coda - not the constituents dominated by the three branches of the Halle and Vergnaud or Fudge syllables. The onset and coda of our template were branching but they were ternary unlike the binary structures proposed by Fudge and Selkirk. We argued, like Fudge, that the peak was only one structural position which dominated vowels or vowel glide sequences. But unlike Fudge we assigned pre-consonantal 'r' to peaks, supporting Kahn's claim that 'r' is a glide in English.¹ Our nodes were labelled, but differently from those of the Cairns and Feinstein templates. We proposed empty positions but unlike

Fudge our segments are assigned to one position only. They do not move into 'head-like' positions if these head positions are available.

The external boundaries of our syllable were those determined by the maximization of onsets,² and we found no evidence of fuzzy boundaries - i.e. ambisyllabic segments. We argued that roots were the domain of syllabification at the level of processing at which between-word segmental errors occur but suggested that the word might be the implicated domain with respect to the locus at which within-word errors and word blends occur.

Our analysis supported those competence-theorists like Selkirk and Cairns and Feinstein who argue that syllable structure is lexical. There are sound arguments both from experimental studies and from malapropisms and word-stress errors that the production lexicon includes syllable structure. Our analysis of segmental errors provided more support for this hypothesis.

We see then that the structure that we found to characterize the units of production converges in a number of respects with aspects independently espoused by competence theorists. This is not surprising given the assumption that the grammar underlies the capacity for use.

✓

What was the source of our conclusions regarding the nature of the production template? A number of competence constructs have been supported as units in production on the basis of their moveability as units in errors. For instance Fromkin argues that both the segment and the feature are units of production on this basis. The syllable per se does not undergo such movement however so its viability as a processing unit could not be confirmed on these grounds. We noticed (as had others before us) that segments which interacted in errors seemed to involve segments from similar syllabic positions. In order to test this we required a precise algorithm for determining which segments belonged to which position. We adopted Kahn's slow speech rules, their domain of application (i.e. the word), and modified them by distinguishing three positions, onset, peak and coda. The analysis of over 500 errors confirmed that segments which interacted in errors were indeed, for the most part, from the same syllabic position. This included perseverations, anticipations, exchanges and the two string positions of single elements involved in 'movement' or 'addition' errors.

This was particularly interesting given our assumptions (borrowed from Garrett) regarding the processing mechanism. On the assumption that grammars are closely related to language processing it followed that processing would be modular and that errors would involve the interaction of

elements of the same processing type. Some noteworthy characteristics of segmental errors could be explained if one were to assume that segmental errors were in fact syllabic constituent errors. First of all errors often involved the interaction of 'no segment' with 'some segment'. For example (1) - (3).

- | | |
|--------------------------|--------------------|
| (1) an ice cream cone | a [kays] ream cone |
| (2) pounds and francs | prounds and fancis |
| (3) brisket and potatoes | ...protatoes |

How could 'no segment' be analysed as a unit of the same type as 'k' and 'r' as seemed to be the case in (1) - (3)? If the processor were selecting syllabic constituents, not segments, this identification of 'something' with 'nothing' would be explained. Similarly, sequences of segments often interacted with single segments as in (4).

- | | |
|---------------------|---------------|
| (4) highway driving | dryway hiving |
|---------------------|---------------|

How could the processor mistake one segment for two? A plausible answer was again that as they both were elements of the same syllabic constituent type, they were being selected and mis-selected as such by the processor.

Thirdly only certain sequences of segments were regularly involved in errors - only those that belonged to one syllabic constituent. It is clear that above the

segmental level, sequences of segments that are involved in errors are always analysable as constituents of some sort. This was equally true of sub-morphemic sequences. They were onsets, peaks or codas, and interestingly, not rhymes.

We also found a number of errors which involved single segments from onset clusters. The pattern they exhibited in errors duplicated that of entire onsets, hence we concluded that they consisted of three sub-constituents. There was a scarcity of errors involving coda clusters so we could only speculate tentatively, assigning them a structure analogous to onsets. We proposed algorithms for generating errors, mirroring what we took to be the mis-selection of identical syllabic constituents and sub-constituents by the processor. These algorithms can generate about 98% of our between-word segmental errors.

Subsequent to our investigation of segmental errors we analysed a number of word blends. They too involved the mis-selection of identical constituents. Indeed this turned out to be the case for all blends that involved a simple juxtaposition of two words. Hence we could generate all of these.

It remained to be seen whether our hypothesis, namely, that segmental errors were syllabic constituent errors, could be integrated into some model of production in an explanatory

way. We examined three models of speech production which pinpointed the locus of segmental errors. We found that our hypothesis fitted into Garrett's model without any change to the model. However, recently Garrett (Garrett, 1982) suggested incorporating the segmental slots of the Shattuck-Hufnagel model into his. We found we had to minimally change his model to effect this integration. The change only extended his empirical coverage, so far as we could determine. Our hypothesis could be incorporated easily because his ~~lexical items~~ are syllabified and segmental errors are hypothesized to occur when lexical items (roots or stems) are mapped onto the segmental slots of the phrasal frame.

However, although our hypothesis was observationally accurate we had not motivated this manipulation of syllabic constituents within the normal functioning of the processing mechanism. And indeed, an explicitly elaborated formal proposal in this regard is well beyond the scope of this dissertation, although we suggest the direction in which such an account might be forthcoming. The Fromkin, Shattuck-Hufnagel and Garrett accounts of segmental errors all fail in this respect. All of the models insert segments rather than morphemes into the required positions in order to account for the fact that segments are found out of position in errors. None of them motivates this piece-meal insertion

of morphemes in terms of the normal mechanism. It is their ad hoc way of accounting for errors. We believe that the root of the trouble in all three models lies in their conception of this processing level as a syntactic level. We have no suggestions regarding the precise mapping between syntactic and prosodic structure (cf. Selkirk, 1985 and note 3) but since segmental errors seem clearly to be constrained by prosodic facts, namely syllabic constituency membership, it seems reasonable to conclude that the errors occur when prosodic structure is the computationally relevant type. Hence we propose that at this level the lowest level of prosodic structure (the syllable) is being fleshed out.

We suggested that the generation of utterances involved the insertion of lexical material into a prosodic frame. We believe that, at least with regard to the syllable level of prosodic structure, this frame consists of a series of empty syllable templates into which morphemes have to be inserted piece by piece - they can not be inserted as holistic units because their allophonic variants will depend on adjacent morphemes. If, as has been suggested, syllables are the concatenative units of production, or the routines used by the processor, then all items of an utterance must be assigned the appropriate syllable structure as determined by the adjacent morphemes. These syllables will flesh out empty feet templates and so on up through the hierarchy of prosodic

units. Or at least, so we speculate.³

Sussman (1984) speculates that canonical syllable forms are equivalent to neuronal templates "where each consonant and vowel position is associated with a specific cell assembly network" (p. 169). Might there be neuronal templates that correspond to all units of prosodic structure? Is the prosodic word, for instance, implicated in within-word errors and word blends?

It has often been noted that errors in aphasia are not unlike errors in normals. Will we find that aphasic errors are constrained by the syllabic-position-constraint?

Might a number of second language learner errors be the result of mapping the units of the second language into the native language prosodic templates?

These are all empirical questions; they can be formulated as a result of the rich and precisely formulated hypotheses of competence theorists. We expect that the extensive research currently, underway in the domain of prosody will lead to further insights into the nature of the production grammar prosodic units. It seems clear that hypotheses that come out of the competence domain provide the best working hypotheses with regard to the processing domain. The assumption that the rules and representations of grammars

will form a basic component in a theory of performance has indeed been an extremely fruitful assumption.

NOTES TO CONCLUSION

¹Our algorithm assigns vowel glide sequences directly to the peak node and consonants directly to the sub-constituents of onsets and codas. We believe that the CV tier might prove to be a viable mediating level to deal with 'complex segments' such as the vowel glide sequences and the tautomorphic obstruent sequences such as the 'ks' and 'axe' which we assigned to one structural node. This we believe is a clear area for further research. We must leave it unsolved. See Chapter 2, Section 6.4.

²We used 'permissible initial clusters' (borrowed from Kahn) as our criteria and found no segmental errors which suggested different criteria such as the sonority hierarchy. For instance we found no error like (i)

(i) the athlete bumped

the ablete...

where the sonority principle would have analysed the 'th' as an onset and 'permissible clusters' would have analysed 'th' as a coda. However, the word blend in (165) Chapter 2 supports the sonority hierarchy and the structure we assign to onsets on the basis of our corpus is in line with this hierarchy with respect to the obstruent position. The error referred to involved the juxtaposition of a nasal and a liquid, viz. move/run becomes mrun. Such clusters occur word medially in English in such forms as 'Lumley' and 'Monroe'. The assignment of the sibilants 's' and 'sh' to the 's' position violates sonority and supports Cairns' and Feinstein's arguments. (See Chapter 1, Section 3 (iv).)

³Selkirk's most recent work (Selkirk, 1985) argues for a syntax-phonology mapping which involves a metrical grid along the lines of Prince, 1983. We believe that this work may provide a most fruitful area for investigating in detail our speculations regarding a prosodic structure frame in the production process.

REFERENCES

- Aitchison, J., and M. Straf (1981) "Lexical Storage and Retrieval: A Developing Skill?," Linguistics 19, 751-795. Reprinted in A. Cutler (ed.) Slips of the Tongue and Language Production, Mouton, The Hague.
- Baars, B.J. (1980) "The Competing Plans Hypothesis: An Heuristic Viewpoint on the Causes of Errors in Speech," in H. Dechert and M. Raupach (eds.), Temporal Variables in Speech: Studies in Honour of Frieda Goldman-Eisler, Mouton, The Hague.
- Baars, B.J., and M. Motley (1976) "Spoonerisms as Sequencer Conflicts: Evidence from Artificially Elicited Errors," American Journal of Psychology 89, 467-484.
- Bell, Alan, and Joan B. Hooper (eds.) (1978) Syllables and Segments, North-Holland, New York.
- Boomer, D.S., and J.D. Laver (1968) "Slips of the Tongue," British Journal of Disorders of Communication 3.
- Broselow, E. (1976) The Phonology of Egyptian Arabic, Doctoral Dissertation, University of Massachusetts, Amherst, Massachusetts.
- Broselow, E. (1983) "Non-Obvious Transfer: On Predicting Epenthesis Errors," in S. Gass and L. Selinker (eds.), Language Transfer in Language Learning, Newbury House, Rowley, Massachusetts.
- Browman, C.P. (1978) "Tip of the Tongue and Slip of the Ear: Implications for Language Processing." UCLA Working Papers in Phonetics 42.
- Browman, C.P. (1980) "Perceptual Processing: Evidence from Slips of the Ear," in V.A. Fromkin (ed.), Errors in Linguistic Performance, Academic Press, New York.
- Brown, R., and McNeill, D. (1966) "The 'Tip of the Tongue' Phenomenon," Journal of Verbal Learning and Verbal Behaviour 5, 325-337.
- Buckingham, H.W. (1980) "On Correlating Aphasic Errors with Slips-of-the Tongue," Applied Psycholinguistics 1, 199-220.

Butterworth, B. (1981) "Speech Errors: Old Data in Search of New Theories," Linguistics 19-7/8, 627-662.

Caplan, David, and Christine Futter (1983) "A Case Study," paper presented at BABBLE, Niagara Falls, Ontario.

Chomsky, Noam (1965) Aspects of the Theory of Syntax, MIT Press, Cambridge, Massachusetts.

Chomsky, Noam (1970) "Remarks on Nominalization," in R. Jacobs and P. Rosenbaum (eds.) Readings in English Transformational Grammar, 184-21, Ginn, Waltham.

Chomsky, Noam (1976) Reflections on Language, Temple Smith, London.

Chomsky, Noam (1980) Rules and Representations, Columbia University Press, New York.

Chomsky, Noam (1982) Lectures on Government and Binding, Foris Publications, Dordrecht.

Chomsky, Noam, and Morris Halle (1968) The Sound Pattern of English, Harper & Row, New York.

Clements, George N., and Samuel J. Keyser (1983) CV Phonology - A Generative Theory of the Syllable, MIT Press, Cambridge, Massachusetts.

Cohen, A. (1980) Correcting of Speech Errors in a Shadowing Task, in V.A. Fromkin (ed.), Errors in Linguistic Performance: Slips of the Tongue, Ear, Pen, and Hand, Academic Press, New York.

Cole, R.A. (1973) Listening for Mispronunciations: A Message of What we Hear During Speech, Perception and Psychophysics 11, 153-156.

Crompton, A. (1981) Syllables and Segments in Speech Production, Linguistics 19, 663-716.

Cutler, Anne (1980a) "Syllable Omission Errors and Isochrony," in H.W. Dechert and M. Raupach (eds.), Temporal Variables in Speech, Mouton, The Hague.

Cutler, Anne (1980b) "Errors of Stress and Intonation," in V.A. Fromkin (ed.), Errors in Linguistic Performance: Slips of the Tongue, Ear, Pen, and Hand, Academic Press, New York.

Cutler, Anne (1981) "The Reliability of Speech Error Data," Linguistics, 561-582.

Cutler, Anne (ed.) (1982a) Slips of the Tongue and Language Production, Mouton, The Hague.

Cutler, Anne (1982b) Speech Errors: A Classified Bibliography, Indiana University Linguistics Club, Blommington, Indiana.

Cutler, Anne, and D. Fay (1978) Introductory Essay, in Rudolf Meringer and Karl Mayer: Versprechen und Verlesen. Re-issue, John Benjamins, Amsterdam.

Cutler, Anne, and S.D. Isard (1980) "The Production of Prosody," in B. Butterworth (ed.), Language Production, Academic Press, London.

Davidson-Nielsen, N. (1975) "A Phonological Analysis of English sp, st, sk, with Special Reference to Speech Error Evidence," Journal of the IPA 5, 3-25.

Dell, G.S., and P.A. Reich (1980) "Toward a Unified Model of Slips of the Tongue," in V.A. Fromkin (ed.), Errors in Linguistic Performance: Slips of the Tongue, Ear, Pen, and Hand, Academic Press, New York.

Dell, G.S., and P.A. Reich (1981) "Stages in Speech Production: An Analysis of Speech Error Data," Journal of Verbal Learning and Verbal Behavior 20, 611-629.

Donegan, Patricia J., and David Stampe (1978) "The Syllable in Phonological and Prosodic Structure," in A. Bell and J. Hooper (eds.).

Dretske, F. (1974) "Explanation in Linguistics," in D. Cohen (ed.), Explaining Linguistic Phenomena, Hemisphere Publishing Co., Washington, D.C.

Eek, A., (1975) "Observations on the Durations of Some Word Structures II," Estonian Papers in Phonetics 4, 7-53.

Ellis, Andrew W. (ed.) (1982) Normality and Pathology in Cognitive Functions, Academic Press, New York.

Fay, D.A. (1980) "Transformational Errors," in V.A. Fromkin (ed.), Errors in Linguistic Performance: Slips of the Tongue, Ear, Pen, and Hand, Academic Press, New York.

- Fay, D.A., and Anne Cutler (1977) "Malapropisms and the Structure of the Mental Lexicon," Linguistic Inquiry 8, 5: 05-520.
- Fodor, J. (1975) The Language of Thought, Thomas Crowell Company, New York.
- Foley, J. (1970) "Phonological Distinctive Features," Folia Linguistica 4, 87-92.
- Forster, K.J. (1976) "Accessing the Mental Lexicon," in Walker and Wales, (eds.), New Approaches to Language Mechanisms, North Holland, Amsterdam.
- Fromkin, Victoria A. (1971) "The Non-Anomalous Nature of Anomalous Utterances," Language 47, 27-52.
- Fromkin, Victoria A. (1973) Speech Errors as Linguistic Evidence, Mouton, Paris.
- Fromkin, Victoria A. (1977) "Putting the EmPHasis on the Wrong SyLLABle," in L.M. Hyman (ed.), Studies in Stress and Accent, Southern California Occasional Papers in Linguistics 4, 15-26.
- Fry, D.B. (1973) "The Linguistic Evidence of Speech Errors," in V.A. Fromkin (ed.), Speech Errors as Linguistic Evidence.
- Fudge, E. (1969) "Syllables," Journal of Linguistics 5, 253-286.
- Fujimura, Osamu, and Julie B. Lovins (1978) "Syllables as Concatenative Phonetic Units," in Syllables and Segments, Allan Bell and Joan B. Hooper (eds.).
- Fujimura, Osamu, and Julie B. Lovins (1982) Syllables as Concatenative Phonetic Units, Indiana University Linguistics Club, Bloomington, Indiana.
- Garnes, S., and Z. Bond (1975) "Slips of the Ear: Errors in Perception of Casual Speech," Proceedings of the Eleventh Regional Meeting of the Chicago Linguistics Society, CLS, 214-225.
- Garnes, S., and Z. Bond (1980) "A Slip of the Ear: A Snip of the Ear? A Slip of the Year?," in V.A. Fromkin (ed.), Errors in Linguistic Performance: Slips of the Tongue, Ear, Pen, and Hand, Academic Press, New York.

- Garrett, Merrill (1975) "The Analysis of Sentence Production," in Bower (ed.), Psychology of Learning and Motivation, Vol. 9, Academic Press, New York.
- Garrett, Merrill (1976) "Syntactic Processes in Sentence Production," in Walker and Wales (eds.), New Approaches to Language Mechanisms, North Holland, Amsterdam.
- Garrett, Merrill (1980a) "Levels of Processing in Sentence Production," in B. Butterworth (ed.), Language Production: Volume 1: Speech and Talk, 197-220, Academic Press, London.
- Garrett, Merrill (1980b) "The Limits of Accommodation: Arguments for Independent Processing Levels in Sentence Production," in V.A. Fromkin (ed.), Errors in Linguistic Performance, Academic Press, New York.
- Garrett, Merrill (1982) "Production of Speech: Observations from Normal and Pathological Use," in Andrew W. Ellis (ed.), Normality and Pathology in Cognitive Functions, Academic Press, New York.
- Gay, Thomas (1978) "Articulatory Units: Segments or Syllables?," in Allan Bell and Joan B. Hooper, (eds.), Syllables and Segments.
- Goldsmith, John (1976a) Autosegmental Phonology, Doctoral Dissertation, MIT, Cambridge, Massachusetts.
- Goldsmith, John (1976b) "An Overview of Autosegmental Phonology," Linguistic Analysis 2, 23-68.
- Goldstein, L. (1980) "Categorical Features in Speech Perception and Production," Journal of the Acoustical Society of America 67, 1336-1348.
- Goldstein, L. (1980) "Bias and Asymmetry in Speech Perception," in V.A. Fromkin (ed.), Errors in Linguistic Performance: Slips of the Tongue, Ear, Pen, and Hand, Academic Press, New York.
- Goldstein, Myron (1968) "Some Slips of the Tongue," Psychological Reports 22.
- Greenberg, J. (ed.) (1963) Universals of Language, MIT, Cambridge, Massachusetts.
- Grunwald, Henry A. (ed.) (1984) "Questing for Quarks," Discover, Volume 5, No. 5.

- Halle, Morris, and J. Vergnaud (1980) "Three Dimensional Phonology," Journal of Linguistic Research 1.1, 83-105.
- Halle, Morris, and George N. Clements (1983) Problem Book in Phonology, MIT Press, Cambridge, Massachusetts.
- Harris, J.W. (1983) Syllable Structure and Stress in Spanish: a Nonlinear Approach, Linguistic Inquiry Monograph No. 8, MIT Press.
- Haugen, E. (1956) "The Syllable in Linguistic Description," in M. Halle, H.G. Lunt, and H. McClean, eds., For Roman Jakobson, 213-221, Mouton, The Hague.
- Hayes, Bruce, (1981) A Metrical Theory of Stress Rules, Indiana University Linguistics Club, Bloomington, Indiana.
- Hoard, James E. (1971) "Aspiration, Tenseness and Syllabification in English," Language 47, 133-140.
- Hoard, James E. (1978) "Syllabification in Northwest Indian Languages with Remarks on the Nature of Syllabic Stops and Affricates," in A. Bell and J. Hooper, eds.
- Hockett, Charles F. (1955) A Manual of Phonology, IJAL Memoir 11.
- Hooper, Joan B. (1972) "The Syllable in Phonological Theory," Language 48, 525-540.
- Hooper, Joan B. (1975) "The ArchiSegment in Natural Generative Phonology," Language 51, 536-560.
- Hooper, Joan B. (1976) An Introduction to Natural Generative Phonology, Academic Press, New York.
- Horowitz, L.M., P.C. Chilian, and K.P. Dunnigan (1969) "Word Fragments and Their Reintegrative Powers," Journal of Experimental Psychology 80, 392-394.
- Horowitz, L.M., M.A. White, and D.W. Attwood (1968) "Word Fragments as Aids to Recall: The Organisation of a Word," Journal of Experimental Psychology 76, 219-226.
- Jakobson, R. (1968) Child Language, Aphasia, and Phonological Universals, Mouton, The Hague.

- Jespersen, O. (1909) A Modern English Grammar on Historical Principles, Carl Winer, Heidelberg.
- Kahn, D. (1976) Syllable-Based Generalizations in English Phonology, Indiana University Linguistics Club, Bloomington, Indiana.
- Kent, R.D., and J.C. Rosenbek (1982) "Prosodic Disturbance and Neurologic Lesion," Brain and Language 15, 259-291.
- Kiparsky, P. (1979) "Metrical Structure Assignment is Cyclic," Linguistic Inquiry 10.3, 521-541.
- Klatt, K. (1966) "Lexical Representations for Speech Production and Perception," in T. Myers et al. (eds.).
- Kohler, K. (1966) "Is the Syllable a Phonological Universal?," Journal of Linguistics 2, 107-208.
- Kozhevnikov, V.A., and L.A. Chistovich (1965) Speech: Articulation and Perception, (revised), Joint Publications Research Service 30, 543. Washington: U.S. Department of Commerce.
- Kupin, J.J. (1982) Tongue Twisters as a Source of Information about Speech Production. Indiana University Linguistics Club, Bloomington, Indiana.
- Kurylowicz, J. (1948) "Contribution a la theorie de la syllable," Biuletyn polskiego towarzystwa jezykoznawczego 8, 80-114.
- Lehiste, I. (1973) "Rhythmic Units and Syntactic Units in Production and Perception," Journal of the Acoustical Society of America 54, 1228-1234.
- Lehiste, I. (1977) "Isochrony Reconsidered," Journal of Phonetics 5, 253-263.
- Liberman, Mark (1975) The Intonational System of English, Doctoral Dissertation, MIT, Cambridge, Mass.
- Liberman, Mark, and Alan Prince (1977) "On Stress and Linguistic Rhythm," Linguistic Inquiry 8, 2, 249-336.
- Mackay, Donald G. (1969) "Forward and Backward Masking in Motor Systems," Kybernetik 6, 57-64.

- MacKay, Donald G. (1970) "Spoonerisms: The Structure of Errors in the Serial Order of Speech," Neuropsychologia 8, 323-350.
- MacKay, Donald G. (1972) "The Structure of Words and Syllables: Evidence from Errors in Speech," Cognitive Psychology 3, 210-227.
- MacKay, Donald G. (1978) "Speech Errors Inside the Syllable," in A. Bell and J.B. Hooper (eds.), Syllables and Segments, 201-212, North-Holland, New York.
- MacKay, Donald G. (1979) "Lexical Insertion, Inflection and Derivation: Creative Processes in Word Production," Journal of Psycholinguistic Research 8, 477-498.
- Marr, David (1984) Vision: A Computational Investigation into the Human Representation and Processing of Visual Information, W.H. Freeman, New York.
- Marslen-Wilson, W., and A. Welch (1978) "Processing Interactions and Lexical Access During Word Recognition in Continuous Speech," Cognitive Psychology 10, 29-63
- McCarthy, John, (1979) "On Stress and Syllabification," Linguistic Inquiry 10, 443-465.
- McCarthy, John, (1982) Formal Problems in Semitic Phonology and Morphology, Indiana University Linguistics Club, Bloomington, Indiana.
- Mehler, J., J. Segui, and U. Frauenfelder (1981) "The Role of the Syllable in Language Acquisition and Perception," in T. Myers et al. (eds.).
- Menn, L. (1978) "Phonological Units in Beginning Speech," in A. Bell and J.B. Hooper (eds.), Syllables and Segments, 157-171, North-Holland, New York.
- Meringer, R., and K. Mayer (1895/1978) Versprechen und Verlesen: Eine Psychologisch-Linguistische Studie, Goschen, Stuttgart. (Re-issued in 1978, John Benjamins, Amsterdam.)
- Mohanan, K.P., (1979) "On Syllabicity," MIT Working Papers in Linguistics, MIT Press, Cambridge, Massachusetts.
- Myers, T., J. Laver, and J. Anderson (eds.) (1983) The Cognitive Representation of Speech, North-Holland, New York.

- Nooteboom, S. (1969) "The Tongue Slips into Patterns," in V.A. Fromkin (ed.), (1973).
- Nooteboom, S. (1980) "Speaking and Unspeaking: Detection and Correction of Phonological and Lexical Errors in Spontaneous Speech," in V.A. Fromkin (ed.), Errors in Linguistic Performance: Slips of the Tongue, Ear, Pen, and Hand, Academic Press, New York.
- Nooteboom, S. (1981) "Lexical Retrieval from Fragments of Spoken Words: Beginnings Versus Endings," Journal of Phonetics 9, 4007-424.
- Pike, Kenneth L. (1945) The Intonation of American English, University of Michigan Press, Ann Arbor.
- Pike, Kenneth L. (1947) Phonemics, Longmans Canada Ltd., Don Mills.
- Pike, Kenneth L. (1967) Languages in Relation to a Unified Theory of the Structure of Human Behavior, 2nd ed, The Hague, Mouton.
- Prince, Alan (1980) "A Metrical Theory for Estonian Quantity," Linguistic Inquiry 11, 3.
- Prince, Alan (1983) "Relating to the Grid," Linguistic Inquiry 14, 19-100.
- Reed, I., O. Miyaoka, S. Jacobson, P. Afcan, and M. Krauss, (1977) Yup'ik Eskimo Grammar, Alaska Native Language Center and Yup'ik Language Workshop, University of Alaska.
- Roca, I. (1982) "Explorations on Extrametricality," unpublished manuscript, MIT.
- Safir, Ken (ed.) (1979) MIT Working Papers in Linguistics, MIT Press, Cambridge, Massachusetts.
- Saussure (1915 (1959)) Course in General Linguistics, Philosophical Library, New York.
- Selkirk, Elizabeth O. (1978) "On Prosodic Structure and its Relation to Syntactic Structure." Paper presented at Sloan Workshop on the Mental Representation of Phonology.

- Selkirk, Elizabeth O. (1980) "The Role of Prosodic Categories in English Word Stress," Linguistic Inquiry 11, 3.
- Selkirk, Elizabeth O. (1981) On Prosodic Structure and its Relation to Syntactic Structure, Indiana University Linguistics Club, Bloomington, Indiana.
- Selkirk, Elizabeth O. (1985) Phonology and Syntax The Relation Between Sound and Structure, MIT Press, Cambridge, Mass.
- Shattuck, Stefanie (1975) Speech Errors and Sentence Production, Doctoral Dissertation, MIT, Cambridge, Massachusetts.
- Shattuck-Hufnagel, Stefanie (1979) "Speech Errors as Evidence for a Serial-Ordering Mechanism in Sentence Production," in W.E. Cooper and E.C.T. Walker - (eds.), Sentence Processing: Psycholinguistic Studies Presented to Merrill Garrett, 295-342, Lawrence Erlbaum, Hillsdale, N.J.
- Shattuck-Hufnagel, Stefanie (1984) "A 'slot-and-filler' Model of Serial Ordering of Speech Production." Paper presented at Symposium on Motor and Sensory Language Processes, Montreal, April, 1984.
- Shattuck-Hufnagel, Stefanie, and D. Klatt (1979) "The Limited Use of Distinctive Features and Markedness in Speech Production: Evidence from Speech Error Data," Journal of Verbal Learning and Verbal Behavior 18, 41-55.
- Shattuck-Hufnagel, Stefanie, and D. Klatt (1980) "How Single Phoneme Error Data Rule Out Two Modes of Error Generation," in V.A. Fromkin (ed.), Errors in Linguistic Performance: Slips of the Tongue, Ear, Pen, and Hand, Academic Press, New York.
- Sherzer (1976) "Play Languages: Implications for Sociolinguistics," in B. Kirschenblatt-Gimblett, ed., Speech Play, University of Pennsylvania Press, Pennsylvania.
- Sievers, E. (1893) Grundzuge der Phonetik. Breitkopf und Hartel, Leipzig.
- Sommerstein, A. (1974) "On Phonotactically Motivated Rules," Journal of Linguistics 10, 71-94.
- Stanley, R. (1967) "Redundancy Rules in Phonology," Language 43, 393-436.

- Stowell, Timothy A. (1981) Origins of Phrase Structure,
Doctoral Dissertation, MIT, Cambridge, Massachusetts.
- Sussman, Harvey (1984) "A Neuronal Model for Syllable
Representation," Brain and Language 22, 167-177.
- Taft, M., and K.I. Forster (1976) "Lexical Storage and
Retrieval of Polymorphemic and Polysyllabic Words,"
Journal of Verbal Learning and Verbal Behavior 15,
607-620.
- Tent, J., and J.F. Clark (1980) An Experimental Investigation
into the Perception of Slips of the Tongue, Journal of
Phonetics 7, 317-325.
- Thrainsson, H. (1978) "On the Phonology of Icelandic
Preaspiration," Nordic Journal of Linguistics 1.1.,
3-54.
- Treiman, Rebecca (1983) "The Structure of Spoken Syllables:
Evidence from Novel Word Games," Cognition 15, 49-74.
- Treiman, Rebecca (1985) "Onsets and Rimes as Units of Spoken
Syllables: Evidence from Children," Journal of
Experimental Child Psychology 39, 161-181.
- Trubetzkoy, N.S. (1958) Grundzuge der Phonologie, Vandenhoeck
and Ruprecht Gottingen.
- van den Broecke, M., and L. Goldstein (1980) "Consonant
Features in Speech Errors," in V.A. Fromkin (ed.).
- Vennemann, T. (1972) "On the Theory of Syllabic Phonology,"
Linguistische Berichte 18.
- Vergnaud, J.-R. (1977) "Formal Properties of Phonological
Rules," in R. Butts and J. Hintikka, (eds.), Basic
Problems in Methodology and Linguistics, Reidel,
Dordrecht.
- Wells, R. (1973) "Predicting Slips of the Tongue," Yale
Scientific Magazine, XXVI, 9-30.
- Zwicky, A.M. (1980) Mistakes, Advocate Publishing Group,
Reynoldsburg, Ohio.