

## Lexical Storage and Retrieval of Prefixed Words

MARCUS TAFT AND KENNETH I. FORSTER

*Monash University*

Three experiments are described which support the hypothesis that in a lexical decision task, prefixed words are analyzed into their constituent morphemes before lexical access occurs. The results show that nonwords that are stems of prefixed words (e.g., *juvenate*) take longer to classify than nonwords which are not stems (e.g., *pertoire*), suggesting that the nonword stem is directly represented in the lexicon. Further, words which can occur both as a free and as a bound morpheme (e.g., *vent*) take longer to classify when the bound form is more frequent than the free form. Finally, prefixed nonwords took longer to classify when they contained a real stem (e.g., *dejuvenate*), compared with control items which did not (e.g., *depertoire*). A general model of word recognition is presented which incorporates the process of morphological decomposition.

In order to recognize that a visually presented sequence of letters forms a word, some kind of representation of the sensory input must be matched with an internal lexical representation of the word. One of the key issues in the study of this process of lexical access concerns the form in which the sensory signal is represented and the possible recodings that might be carried out in order for accessing to take place. For example, there has been considerable debate as to whether word recognition is based on features of the whole word or on individual features of components of the word (e.g., Smith, 1971; Wheeler, 1970). Interest has also been focused on the possibility that the orthographic stimulus is converted into phonological form prior to accessing (Rubenstein, Lewis, & Rubenstein, 1971; Baron, 1973; Forster & Chambers, 1973; Meyer, Schvaneveldt, & Ruddy, 1974).

The type of recoding which will be of prime consideration here is the possible morphological decomposition of an item. For example, the word *unlucky* is clearly composed

of the base word *luck*, the adjectival suffix *-y*, and the negative prefix *un-*. It is possible that the internal lexicon is organized so that *unlucky* is stored in conjunction with *luck* (along with *lucky*, *luckily*, *luckless*, and so on), and more particularly, that there is no separate lexical entry for the word *unlucky*, this word being constructed from the entry *luck* by the addition of the affixes *un-* and *-y*. Such a method of storing information entails that recognition of the word *unlucky* requires a prior morphological analysis of the word, that is, the prefix *un-* and the suffix *-y* must be stripped off before the lexical representation of *unlucky* (namely *luck*) can be accessed. Similarly, it would seem logical and economical for the word *cats* to be filed in the lexicon as *cat*, and thus be recognized as a word only after the *-s* has been stripped off. The results obtained by Kintsch (1972), Gibson and Guinet (1971), Snodgrass and Jarvella (1972), and Murrell and Morton (1974) are consistent with the notion that affixed words, or at least suffixed words, are stored in their base form in the lexicon. Such an economy is quite plausible, and is, in fact, employed in information retrieval systems (Knuth, 1973).

Requests for reprints should be addressed to M. Taft, Department of Psychology, Monash University, Clayton, Victoria, Australia 3168.

However, the notion of morphological decomposition is no longer so appealing when one is confronted with more extreme cases. For example, it would have to be argued that the word *unremittingly* is stored as the entry *mit* even though *mit* does not form a word on its own. If this were so then in order to access the word *unremittingly* one would have to strip off its affixes *un-*, *re-*, *-ing*, and *-ly* and then search the lexicon for the entry *mit*. On finding it, one can then ascertain from information stored in this entry whether *un + re + mit + ing + ly* is a valid combination or not (this information would be similar to the output of the filter system described by Halle, 1973). Note that the same lexical entry, namely, *mit*, would have to be accessed also for the purposes of recognizing *submit*, *commit*, *admit*, *permit*, *emit*, *transmit*, and possibly, *omit* (but not *limit*, *summit*, and *mitten*, since these are not composed of stem plus affix).

Although it may seem quite counter intuitive to claim that nonwords can be given lexical status, it seems very difficult to make a clear distinction between recognizing a word which is obviously composed of a base word plus one or more affixes (e.g., *unlucky*, *reorganize*) and a word composed of a nonword stem plus one or more affixes (e.g., *unremittingly*, *rejuvenate*). That is, it would be difficult to design the system so that morphological decomposition was applied in the cases that clearly involve a real-word stem, but was prevented from applying in the cases that do not involve a real word as a stem.

The following experiments were designed to determine whether the concept of morphological decomposition, elaborated to its fullest, must be incorporated into models of word recognition.

#### EXPERIMENT I

If the nonword stems of derived words (e.g. *juvenate*) are stored in the lexicon, then the prediction can be made that such a stem will be more difficult to recognize as a nonword

than will a nonword that is not stored in the lexicon, that is, that is not the stem of a derived word (e.g., *luvenate*).

In cases like *luvenate*, the item would be recognized as a nonword after a search in the appropriate subset of the lexicon was found to be unsuccessful (Rubenstein, Garfield, & Millikan, 1970). But in cases like *juvenate* a lexical representation of the nonword would be found. However, there would have to be information in this lexical entry which stipulated that the item could not stand as a free morpheme, that is, was not a word on its own. Search would then have to continue on, in case there was an entry for this item which was a free morpheme, since this is a possibility. For example, if a reader has to recognize the item *vent* as a word, he might firstly find the nonword entry *vent*, which has been stored for the purposes of accessing *prevent*, *invent*, and so on, and thus could only identify *vent* as a word after further search discovered the free morpheme entry.

Therefore, on a task where subjects must decide whether the presented item is a word or not, that is, a lexical decision task, a nonword that is the stem of a derived word should take longer to respond to than a nonword that is not the stem of a derived word. This longer latency of response would result from the interruption of the search caused by finding an inappropriate lexical entry.

However, if items such as *juvenate* are found to take longer to recognize than items such as *luvenate* an explanation could be given without resort to morphological decomposition. It could merely be said that *juvenate* is more similar to a word than is *luvenate*. In order to avoid this problem, items such as *pertoire* were used in the following experiments in preference to items of the *luvenate* type. Both *juvenate* and *pertoire* form a complete word when the letter cluster *re* is added (i.e. *rejuvenate* and *repertoire*, respectively). In the former case, though, the cluster *re* forms a real prefix which contributes meaning to the word as a whole, and thus

*juvenate* rather than *rejuvenate* would be the lexical entry according to the morphological decomposition hypothesis. In the latter case, the cluster *re* contributes no meaning and is thus only a pseudoprefix. So, although *pertoire* is as similar as *juvenate* to a proper word, it cannot be considered as the stem of a word and, therefore, it is *repertoire* rather than *pertoire* which would be stored in the lexicon.

A nonword formed by the removal of a real prefix from a word will be termed a real stem (e.g., *juvenate*), and a nonword formed by the removal of a pseudoprefix will be termed a pseudo stem (e.g., *pertoire*). Thus, the hypothesis being tested in this experiment is that real stems should take longer to classify as nonwords than pseudo stems.

### Method

**Materials.** Twenty pairs of nonwords were constructed, one member of each pair being a real stem, one a pseudo stem. The members of each pair were as similar as possible in length, and the words from which they were generated were as similar as possible in frequency according to the Kučera–Francis count (Kučera & Francis, 1967). These items are given in the appendix. A total of 120 items were used, comprising 60 words (matched for length with the nonwords), 20 real stems, 20 pseudo stems, and 20 additional nonwords that were not relevant to the present experiment (these were also derived from real words by deleting initial letters). These items were presented in a semirandom order with practice effects distributed equally over conditions. There were 15 practice trials.

Items were assigned to the real stem condition if the words from which they were generated (i.e., the real prefixed words) met both of the following requirements: (a) if they were listed in the *Shorter Oxford Dictionary* as being derived from a prefix plus stem; (b) if this prefix contributed to the meaning of the word. Thus, the *re* of *rejuvenate* and *revive* means *again*, and hence *juvenate* and *vive* are classified as real stems.

Pseudo stems were derived from words which began with the same letters as a real prefix, but were either not listed in the dictionary as being derived from a prefix plus stem (e.g., *regulate*, *undulate*) or if they were so listed, the prefix no longer contributed meaning to the word (e.g., *devout*, *rebel*, *precinct*, *infant*). Thus, *gulate*, *dulate*, *vout*, *bel*, *cinct*, and *fant* would all be instances of pseudo stems.

**Procedure.** Stimulus items were typed and filmed on 16 mm movie film which was presented using a variable speed motion projector. Each test item was presented for 500 msec with an intertrial interval of approximately 4 sec. The starting of the projector on each trial served as a warning signal. Subjects were instructed to press a YES button if the item was a valid English word, otherwise to press a NO button. The time taken for this response to be made was measured by a millisecond timer. Subjects were told to respond as quickly as possible, but to avoid errors.

**Subjects.** A total of 30 volunteer undergraduate and graduate students served as subjects, and were paid for their participation.

### Results and Discussion

In this and in all subsequent experiments, the effects of isolated trials with exceptionally long or short latencies were minimized by establishing cutoff points two standard deviation units away from the mean for each subject, and setting any outlying values equal to the cutoff. Trials on which an error was made were omitted.

The mean decision times together with percentage error rates for each condition are shown in Table 1. It can be seen that as predicted, real stems took longer to classify than pseudo stems, this difference being significant<sup>1</sup>,  $\min F'(1, 29) = 4.44$ ,  $p < .05$ . Analysis of the errors made also showed a significant effect,  $\min F'(1, 25) = 6.98$ ,  $p < .02$ .

<sup>1</sup> Using the conservative estimate of the quasi *F* ratio suggested by Clark (1973).

TABLE 1  
MEAN LEXICAL DECISION TIMES AND  
PERCENTAGE ERROR RATES FOR REAL  
STEMS AND PSEUDO STEM (EXPERIMENT  
I)

|              | RT<br>(msec) | Error<br>(%) |
|--------------|--------------|--------------|
| Real stems   | 769          | 17.0         |
| Pseudo stems | 727          | 4.0          |

with more errors being made in the real stem condition. The results of this experiment indicate clearly that real stems are perceived as being more word-like than pseudo stems, and it seems very difficult to account for this fact without assuming that in some way, the stems of real prefixed words are directly represented in the lexicon. In terms of the morphological decomposition model outlined earlier, it would be assumed that the increased latency for real stem nonwords represents the time taken to check the contents of the lexical entry to determine whether the stem can stand alone. Evidently this check is not always carried out effectively, since the error rate of 17% in this condition was unusually high.

It may be objected that some of the real stem items have been incorrectly classified on the grounds that in the words from which they are generated, the stem does not have any clear meaning of its own which is modified by the prefix. For example, it might be suggested that the meanings of *embezzle*, *obligate*, and *insipid* cannot really be analyzed into separate components in the same way as the meanings of *overwhelm*, *unwieldy*, and *rejuvenate*. This problem will be taken up later, but for the moment it will suffice to point out that the force of such an objection would be to claim that these items are, in fact, unanalyzable into stem plus prefix, and hence should be placed in the pseudo stem condition. However, any such errors of classification would work against the hypothesis under considera-

tion, making it more difficult to detect any difference between the conditions (especially in view of the fact that the *min F'* test establishes generality over both subjects and items simultaneously).

## EXPERIMENT II

If it is true that stems of derived words are represented in the lexicon then a further consequence of this model of word storage is that there must exist items in the lexicon which not only are real words (free morphemes), but also are stems which cannot stand alone (bound morphemes). For example, while *vent* is a free morpheme (i.e., an outlet for air), it is also a bound morpheme as in *prevent*, *advent*, *invent*, and, possibly, *convent* and *event*. Since these two forms of *vent* have quite different functions, they can be considered as two quite separate entries, *vent*<sub>1</sub> (which can stand alone) and *vent*<sub>2</sub> (which cannot).

The existence of the entry *vent*<sub>2</sub> may complicate recognition of the word *vent*, since, in a lexical search undertaken to classify *vent* as a word, the nonword entry *vent*<sub>2</sub> may be encountered before the free morpheme entry *vent*<sub>1</sub>. This interference will, of course, only occur if *vent*<sub>2</sub> comes before *vent*<sub>1</sub> in the subject's lexicon. If we assume that one of the factors governing the order in which entries are searched is frequency of occurrence (Rubenstein, Lewis & Rubenstein, 1971; Forster & Chambers, 1973) then this situation would arise if the nonword stem *vent* occurs more often than the word *vent*. In fact, this is the case, since the word *vent* has a frequency of 10, while the stem *vent* has a frequency of at least 83 (since *prevent* has a frequency of 83). Therefore *vent*<sub>2</sub> should be encountered before *vent*<sub>1</sub>, and this should interfere with the recognition of *vent* as a word. On the other hand, no such interference will occur in the recognition of a word such as *card*, since the frequency of the word *card* is 26, whereas the frequency of the stem *card* is only 1

(coming from *discard*). Thus the entry *card*<sub>2</sub> (for the bound form) comes after the entry *card*<sub>1</sub> (for the word), and hence there is no opportunity for interference to occur, since *card*<sub>2</sub> will never be encountered.

The prediction, then, is that items where the bound form is more frequent than the free form (B/F words, e.g., *vent*) will take longer to classify as words than items having only a free form (F words, e.g., *coin*), since in the latter condition there are no nonword entries to lead the subject astray. However, items where the bound form is less frequent than the free form (F/B words, e.g., *card*) should take no longer to recognize than F words, since in both cases, no interfering nonword entries are actually accessed.

### Method

**Materials.** Twenty B/F items (e.g., *vent*) were matched with 20 F items (*coin*) for length and for frequency. The frequency values given to the B/F items were the frequency values of their free forms. Twenty F/B items (*card*) were similarly matched with a different set of 20 F items (*fist*). A direct comparison of B/F and F/B words was not possible since it was extremely difficult to match them for frequency. These 80 words were presented in semirandom order with 60 distractor items (nonwords which were not the endings of any other words, e.g. *cint*, *shride*).

**Procedure.** The procedure was the same as in Experiment I, except that tachistoscopic presentation was employed. Items were typed on cards and presented on a two-field tachistoscope for 500 msec with an intertrial interval of approximately 5 sec. The experimenter said "ready" before each trial. Fourteen subjects were used and were paid for their participation.

### Results

Subject means for the two experimental conditions and their respective controls are presented in Table 2.

It can be seen that B/F words did indeed take longer to recognize than F words, this

difference being significant,  $\min F'(1, 34) = 4.95$ ,  $p < .05$ . Also as predicted, F/B words took no longer to recognize than F words,  $\min F'(1, 23) < 1$ . The error analyses revealed no differences at all between any of the conditions.

TABLE 2

MEAN LEXICAL DECISION TIMES AND  
PERCENTAGE ERROR RATES FOR B/F, F,  
F/B WORDS (EXPERIMENT II)

|           | RT<br>(msec) | Error<br>(%) |
|-----------|--------------|--------------|
| B/F Words | 637          | 2.5          |
| F Words   | 605          | 4.3          |
| F/B Words | 604          | 4.3          |
| F Words   | 612          | 3.9          |

These results confirm two quite independent theoretical assumptions. Firstly, the fact that interference occurs in the B/F condition adds force to the argument that the stems of derived words are stored as lexical items, since it would be very difficult to explain the results in any other way. Secondly, the fact that no interference occurs in the F/B condition argues strongly for the fundamental assumption that lexical entries are examined in serial order from high to low frequency of occurrence.

### EXPERIMENT III

In Experiment I, real stems (*juvenate*) were found to take longer to classify than pseudo stems (*pertoire*). One possible interpretation of this result is that the increase in latency was due to uncertainty as to whether the real stems could be used on their own. In actual fact, at least one of the real stem nonwords, namely, *whelm*, was still in use early this century, and it might be the case that other items occur sufficiently often to create confusion.

This problem can be overcome by adding an

inappropriate prefix. Thus, while it is conceivable that some subjects may be uncertain whether *whelm* can stand on its own, there would be no uncertainty about whether *diswhelm* is a proper word.

From the standpoint of the morphological decomposition hypothesis, the adding of inappropriate prefixes changes the situation very little. In the real stem condition (e.g., *dejuvenate*), the prefix *de-* is identified and removed, and a search is begun for the entry *juvenate*. When this is found, the contents of the entry are examined to see whether *de-* is a legitimate prefix. When it is found that it is not, there must be an additional search in case *dejuvenate* is listed as a complete word, like *repertoire*. However, in the pseudo stem condition (e.g., *depertoire*), the prefix *de-* is also removed, and a search begun for *pertoire*. No entry is found, and after an additional search to check that *depertoire* is not listed, the item is established to be a nonword. Thus, real stem nonwords still involve an extra step of checking the legitimacy of the prefix, and hence should take longer to classify.

Since this experiment is an extension of the first experiment, it was felt desirable to employ a totally new set of items. Thus the words from which the items in the real stem and pseudo stem conditions were derived are totally different from those used earlier.

### Method

The apparatus and procedure were the same as for Experiment II. Twenty pairs of real stems (e.g., *juvenate*) and pseudo stems (e.g., *pertoire*) were selected using the same criteria as in Experiment I. (None of the pairs were the same as used in Experiment I.) Inappropriate prefixes were added to both real stems (e.g., *dejuvenate*) and pseudo stems (e.g., *depertoire*). These 40 experimental items (listed in the appendix) were presented in a semirandom order along with 50 distractor items, all of which were prefixed words (e.g., *demolish*, *incriminate*). A total of 15 subjects was used.

### Results

The means in the two conditions are shown in Table 3. As predicted, the real stem nonwords once again took longer to classify than the pseudo stem nonwords,  $\min F'(1, 33) = 5.60$ ,  $p < .05$ , and also produced more errors,  $\min F'(1, 36) = 6.80$ ,  $p < .02$ .

TABLE 3

MEAN LEXICAL DECISION TIMES AND PERCENTAGE ERROR RATES FOR REAL STEM NONWORDS AND PSEUDO STEM NONWORDS (EXPERIMENT III)

|                      | RT<br>(msec) | Error<br>(%) |
|----------------------|--------------|--------------|
| Real stem nonwords   | 836          | 18.7         |
| Pseudo stem nonwords | 748          | 3.3          |

These results strongly suggest that the increased latency observed for real stems in Experiment I could not have been due to uncertainty about the correct response, since no such uncertainty could have been present in this experiment. This conclusion is further supported by the fact that the error rates in the real stem condition are very similar in the two experiments (17.0% in Experiment I, 18.7% in Experiment III).

### GENERAL DISCUSSION

The results of the preceding experiments are all consistent with the assumption that a morphological analysis of words is attempted prior to lexical search. The essential features of the model of word recognition proposed as an explanation of the results are shown in Fig. 1.

In Experiment I, classification of real stems (*juvenate*) required the following steps: 1, 4, 5, 4, 7. Pseudo stems (*pertoire*) required the following sequence of operations: 1, 4, 7. The increased time required for real stems comes from the additional cycle 5, 4. In Experiment II, B/F words (*vent*) involved the sequence 1, 4, 5, 4, 5, 6, whereas F/B words

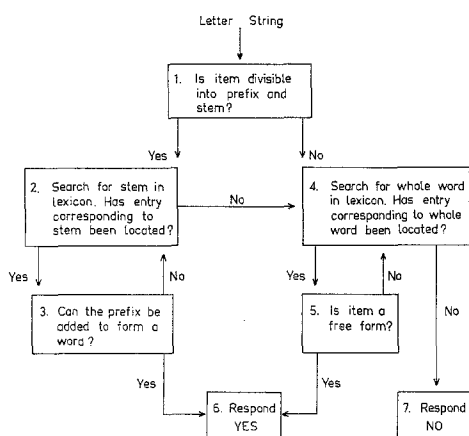


FIG. 1. Model of word recognition incorporating a morphological analysis.

(*card*) involved only 1, 4, 5, 6, and hence B/F words took longer to classify as words. In Experiment III, the real stem nonwords with a prefix (*dejuvenate*) involved the sequence 1, 2, 3, 2, 4, 7, whereas the pseudo stem nonwords with a prefix (*depertoire*) involved the sequence 1, 2, 4, 7, and hence the real stem nonwords again took longer to classify.

There are several alternative models which should be mentioned. These all assume morphological analysis of some kind, and represent variations on the general theme. Firstly, it is possible that a search for the whole word is carried out to begin with and if no lexical entry is found, as would be the case with *dejuvenate*, *depertoire*, and *rejuvenate* (since *rejuvenate* would be stored as *juvenate*), morphological decomposition would be undertaken, that is, steps 1, 2, and 3. A "no" decision at step 2 would lead directly to a "no" response. This system would of course lead to longer recognition times for prefixed words (e.g., *rejuvenate*) than for nonprefixed words (e.g., *somersault*). There are two points that are inconsistent with this view, one logical and one empirical. Firstly, according to this model, prefixed words would be recognized faster if morphological decomposition was not involved, that is, if *rejuvenate* was stored

as *rejuvenate*, and therefore there is no economy at all in having morphological decomposition. Secondly, Taft, Forster, and Garrett (1974) found no difference in the recognition times of prefixed words (*rejuvenate*) and nonprefixed words (*somersault*). A similar result was obtained by R. H. Freeman (personal communication).

The second possible modification to the model is that the search for the whole word is conducted in parallel with morphological decomposition; steps 1, 2, 3 are undertaken at the same time as steps 4, 5. Again a "no" decision at step 2 would lead directly to a "no" response. No judgment can be made from the evidence at hand as to the validity of this parallel model as opposed to the serial model set out in Fig. 1.

A third possibility concerns the form in which prefixed words are stored. In order to explain the results of the experiments reported in this paper, essentially one must explain why *juvenate* and *rejuvenate* are more similar to each other than *pertoire* and *repertoire*, and why *dejuvenate* and *rejuvenate* are more similar than *depertoire* and *repertoire*. The model proposed does this first by assuming that *rejuvenate* is actually stored as *juvenate*, but *repertoire* is not stored as *pertoire*; and second, by assuming that any prefixes on the test item are discarded temporarily while the search takes place. But there are other possibilities that should be considered. For example, *rejuvenate* might be stored as a structure *re(juvenate)*<sup>2</sup>, while *repertoire* is simply stored as *repertoire*. The test item *juvenate*, then, shares a structural element with a lexical entry, while *pertoire* shares no structural elements with a lexical entry. Similarly, the test item *dejuvenate* will share a structural element with a lexical entry, but *depertoire* will share none.

On the present evidence, this reformulation of the model cannot be distinguished from

<sup>2</sup> Actually, the form would probably be *re(juven(ate))* since *-ate* functions as a suffix. However, this is not critical for the argument.

the model proposed earlier. Both theories assume that morphological decomposition is involved in the storage and retrieval of lexical items, and differ only in assumptions about how the prefixes are represented in the lexical entry. The theory that states that the lexical entry for *rejuvenate* is *juvenate* would claim that *admit*, *remit*, *submit*, *permit*, *commit*, *transmit*, and *emit* are all accessed via the same lexical entry, namely *mit*, whereas the theory that states that the entry for *rejuvenate* is *re(juvenate)* would claim that *admit*, *remit*, and so on, all have separate lexical entries. Further research would be required to ascertain which view is correct.

The most critical issue in the design of the current series of experiments concerns the procedure for determining morphological structure. Ideally, one should rely on an independent authority for information about morphology, but the problem here is that dictionaries such as the *Shorter Oxford Dictionary* list words such as *devout*, *rebel*, *precinct*, and *infant* as prefixed words. This presumably reflects historical facts about these words rather than current usage, since the prefixes no longer appear to contribute to the meaning of the whole word. Because of this problem, we have ultimately relied on our own intuitions about morphological structure (using the criteria listed earlier), and the fact that significant results were obtained provides indirect support for those intuitions. However, it is quite possible that others will not share these intuitions, and will claim, for instance, that there is no basis for asserting that *embezzle* contains the prefix *em-* while its control item *emigrate* does not (although the latter does, of course, contain the prefix *e-*). One could only answer this objection by noting the fact that *bezzle* took 772 msec to classify, while *igrate* took only 673 msec (see appendix).

Of course, it may be true that some people do not analyze *embezzle* as *em + bezzle*, just as some people may not have noticed that the meaning of *breakfast* is related in an interesting

way to the morphemes *break* and *fast*. Presumably, once the morphological structure is noticed, the form in which a word is stored will change. Thus, we do not claim that the experiments described would yield the same results in a less literate population than the one we have used, although we would claim that a more suitable set of items for that population would yield the same phenomena.

One approach that might be profitable would be to use the intuitions of speakers other than the investigators, although this presupposes that they will be as careful in their classifications as the investigators. Alternatively, one could rely on clear linguistic evidence about morphological structure, but unfortunately, there do not appear to be any tests that will show that *embezzle* is a prefixed word, and that *precinct* or *rebel* is not (see, for example, Bolinger, 1948). In fact, it might be claimed, instead, that the experimental data, per se, produced in these experiments constitute the best test.

The most interesting question raised by the current experiments concerns the purpose of morphological decomposition. Three explanations can be proposed. The first is that it is more economical to store the stem for a number of different words just once (e.g., *mit*). The problem with this explanation is that the storage capacity of the brain is usually assumed to be so vast that such an economy would be quite trivial. Second, organization by stems allows for semantically related words to appear near each other, even though the lexicon is organized orthographically or phonologically, for example, *rejuvenate* and *juvenile* could appear as adjacent entries even in an alphabetical listing if the prefix is removed. Third, Knuth (1973) has suggested that by stripping off the prefix, *re-* for example, one can use an alphabetical storage of words without having to list a very large number of words under the same description. Thus it is possible that a system using morphological decomposition would access prefixed words faster than a system which left the word



intact, since the entry for *rejuvenate* could be located without having to search through all the words beginning with *re*.

#### APPENDIX

Listed below are the items used in each experiment, together with the mean lexical decision time for the item.

##### *Experiment I*

The items are arranged in pairs with the real stem followed by its pseudo stem control the letters in parentheses being the letters that have been removed.

semble 751, sassin 630 (as); vive 800, lish 708 (re); bezzle 772, igrate 673 (em); cursion 799, tremity 738 (ex); wieldy 813, dulate 824 (un); trieve 756, gulate 683 (re); fect 659, digo 792 (in); sipid 647, kling 758 (in); juvenate 897, pertoire 713 (re); nihilate 794, tagonize 693 (an); ghastr 856, viary 804 (a); whelm 972, tures 697 (over); plored 796, itates 849 (im); fess 756, phet 606 (pro); flation 906, tellect 761 (in); herent 721, conuts 814 (co); lect 696, mond 701 (dia); sults 697, nings 697 (in); ligate 784, ituary 661 (ob); pudent 741, becile 705 (im).

##### *Experiment II*

Each B/F word is presented below paired with its F word control.

vent 725, coin 564; tribute 694, nervous 575; count 562, proud 553; vision 633, double 589; pulse 625, shout 640; tract 645, nurse 582; port 662, neat 612; tense 590, trick 597; quest 625, spray 639; pending 725, picking 642; verse 603, shade 548; lease 561, flock 608; pose 655, drum 596; hind 623, lash 612; duct, 738, dune 746; scribe 642, stripe 647; tent 610, nest 571; tempt 611, torch 567; crease 643, crutch 661; fuse 681, hook 614.

Each F/B word is presented below paired with its F word control.

grade 589, fruit 585; strain 634, branch 582;

habit 563, giant 662; patch 598, slice 674; locate 647, vessel 587; card 611, fist 617; text 609, moon 615; bark 569, bull 623; claim 597, fight 552; serve 579, reach 632; flame 582, bunch 699; chant 631, chunk 636; cover 552, carry 668; dense 617, dread 612; pile 624, pack 639; treat 592, storm 586; tour 630, dirt 533; view 586, dark 556; firm 620, deep 616; quaint 679, scream 593.

##### *Experiment III*

Items are arranged in pairs with the real stem nonword followed by its pseudo stem nonword control, the letters in parentheses being the letters that have been replaced.

relineate 784, recimate 735 (de); besist 732, bescue 715 (re); displicate 824, dispetence 778 (com); conspector 763, contensity 811 (in); desume 698, demier 657 (pre); incocious 797, indacious 715 (pre); disfection 861, distegrity 773 (in); transcedent 1258, transcipice 796 (pre); discavation 735, disasperate 877 (ex); overgress 801, overquiem 755 (re); preterminable 829, pretermittent 826 (in); resert 842, refant 667 (in); prevacuated 878, prelocution 815 (e); conquisitive 775, condustrious 720 (in); perjection 795, pergenuity 712 (in); incapitate 953, inlinquent 667 (de); explenish 734, exverence 749 (re); incuperate 952, inceptacle 784 (re); prepugnant 875, preprimand 742 (re); devade 835, depoch 675 (e).

#### REFERENCES

- BARON, J. Phonemic stage not necessary for reading. *Quarterly Journal of Experimental Psychology*, 1973, **25**, 241-246.
- BOLINGER, D. L. On defining the morpheme. *Word*, 1948, **4**, 18-23.
- CLARK, H. H. The language-as-fixed-effect fallacy: A critique of language statistics in psychological research. *Journal of Verbal Learning and Verbal Behavior*, 1973, **12**, 335-359.
- FORSTER, K. I., & CHAMBERS, S. M. Lexical access and naming time. *Journal of Verbal Learning and Verbal Behavior*, 1973, **12**, 627-635.
- GIBSON, E. J., & GUINET, L. Perception of inflections in brief visual presentations of words. *Journal of Verbal Learning and Verbal Behavior*, 1971, **10**, 182-189.

- HALLE, M. Prolegomena to a theory of word formation. *Linguistic Inquiry*, 1973, 4, 3-16.
- KINTSCH, W. Abstract nouns: Imagery versus lexical complexity. *Journal of Verbal Learning and Verbal Behavior*, 1972, 11, 59-65.
- KNUTH, D. E. *The art of computer programming. Vol. 3: Sorting and searching*. Reading, Mass.: Addison-Wesley, 1973.
- KUČERA, H., & FRANCIS, W. N. *Computational analysis of present-day American English*. Providence, R.I.: Brown University Press, 1967.
- MEYER, D. E., SCHVANEVELDT, R. W., & RUDDY, M. G. Functions of graphemic and phonemic codes in visual word-recognition. *Memory & Cognition*, 1974, 2, 309-321.
- MURRELL, G. A., & MORTON, J. Word recognition and morphemic structure. *Journal of Experimental Psychology*, 1974, 102, 963-968.
- RUBENSTEIN, H., GARFIELD, L., & MILLIKAN, J. A. Homographic entries in the internal lexicon. *Journal of Verbal Learning and Verbal Behavior*, 1970, 9, 487-494.
- RUBENSTEIN, H., LEWIS, S. S., & RUBENSTEIN, M. A. Evidence for phonemic recoding in visual word recognition. *Journal of Verbal Learning and Verbal Behavior*, 1971, 10, 645-657.
- SMITH, F. *Understanding reading: A psycholinguistic analysis of reading and learning to read*. New York: Holt, Rinehart & Winston, 1971.
- SNODGRASS, J. G., & JARVELLA, R. J. Some linguistic determinants of word classification times. *Psychonomic Science*, 1972, 27, 220-222.
- TAFT, M., FORSTER, K. I., & GARRETT, M. F. *Lexical storage of derived words*. Paper presented at the 1st Experimental Psychology Conference held at Monash University, May 1974.
- WHEELER, D. D. Processes of word recognition. *Cognitive Psychology*, 1970, 1, 58-85.

(Received April 18, 1975)