

The visual processing of prefixed ‘zero forms’: Evidence from fMRI

Aditi Lahiri (aditi.lahiri@ling-phil.ox.ac.uk)¹, Charles Redmon^{1,2}, Frans Plank¹, Hilary Wynne¹

¹Language and Brain Laboratory, University of Oxford,

²Department of Language and Linguistics, University of Essex

Some of the most consequential word-building operations in language leave no trace on the surface at all, e.g. English *chain*_N ~ *chain*_V. Words formed from these bases (e.g. *unchain*) often carry layered internal complexities that are not necessarily transparent in the words’ surface forms. Neurolinguistic research has indicated that the brain may engage in substantial structural analysis to carry out silent transformation in morphologically complex words: fMRI studies (Meinzer et al., 2009; Pliatsikas et al., 2014) have isolated the involvement of the left inferior frontal gyrus (LIFG) in the processing of morphologically complex words. Specifically, the BOLD-signal response obtained for suffixed words with additional derivational depth (‘two-step’ forms, e.g. *bridging*_V < *bridge*_V < *bridge*_N) has been found to be significantly stronger than simpler (‘one-step’) suffixed words (e.g. *freezing*_V < *freeze*_V). What remains unknown is whether this effect is equivalent in prefixed words.

Using fMRI, we tested whether the brain tracks the hidden, multi-step derivational structure of English prefixed *un-* verbs”. In words like *unchain*, the verbal prefix *-un* is not the same prefix as the adjectival *un-*, which attaches to nouns (and adjectives) to form adjectives (e.g. *unsure*). Verbal *un-* can only attach to verbs, thus requiring the meaning to denote a reversible change of state. Thus, in almost all cases, verbal *un-* forces the system to posit a zero morpheme, making it an ideal affix to employ for an investigation of derivational depth.

Scanning was carried out on a 3T Siemens Magnetom Prisma scanner. 20 monolingual British English speakers (8F/12 M) completed a visual lexical decision task in a *go/no-go* design. Our test items consisted of (a) *un-*prefixed 1-step verbs containing stems that were basic verbs (e.g. *unfold*), (b) *un-*prefixed 2-step verbs containing stems that were basic nouns (e.g. *unchain*). Stem status was confirmed via a native speaker judgement task (N=25). For control items, we included verbs beginning with the false prefix *-in* (e.g. *incept*), and verbs beginning with the prefix *re-* (e.g. *regroup*). Group-level mixed-effects analyses for *two-step* > *one-step* test items revealed **four significant clusters**: in the right lingual gyrus, left inferior frontal gyrus (pars opercularis), left frontal orbital cortex, and medial superior frontal/paracingulate cortex (see Figure 1, Table 1). Crucially, LIFG activation was stronger and more widespread for *two-step* items, particularly in comparisons with *false-prefixed* words (see Figure 2a). In contrast, *one-step* items showed weaker or no significant effects, suggesting greater engagement of left frontal language-related processes for two-step forms (Figure 2b). Our results indicate an increase in processing demand as a function of degree of morphological complexity in prefixed words.

Table 1: Example stimuli

un- 1 step	un- 2 step	false prefix	re- prefix
<i>undo</i>	<i>unbag</i>	<i>invert</i>	<i>rebuild</i>
<i>unmute</i>	<i>uncork</i>	<i>inflate</i>	<i>recheck</i>
<i>unwind</i>	<i>unbox</i>	<i>include</i>	<i>reclaim</i>
<i>unfreeze</i>	<i>unchain</i>	<i>invite</i>	<i>redraw</i>

Table 2: Stimuli factors

Condition	Stems		Whole Word	
	Zipf score	Length	Zipf score	Length
un- 1 step	4.02	4.26	2.38	6.26
un- 2 step	3.98	4.48	2.12	6.43
false prefix	3.35	NA	3.35	6.40
re- prefix	4.75	4.26	2.53	6.26

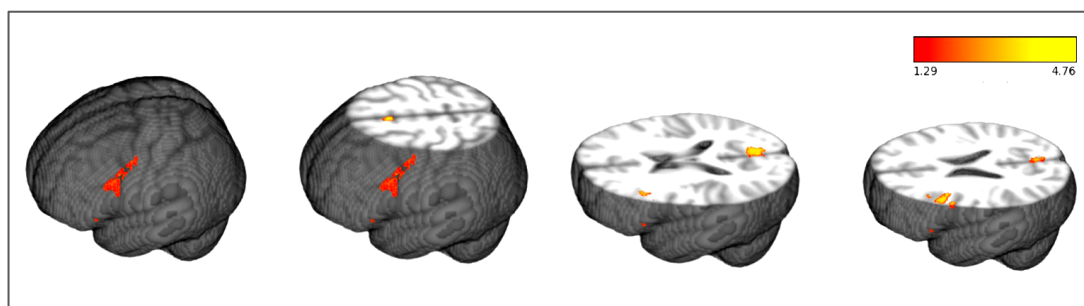


Figure 1: Activity for the processing of 2-step > 1-step *un-* verbs, rendered on the standard MNI template

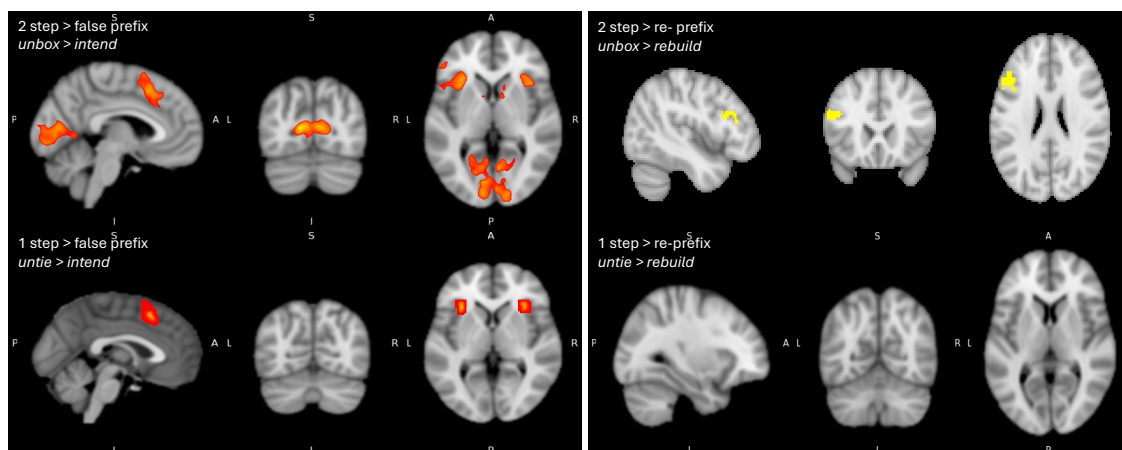


Figure 2a: (left) Activity for the processing of test items vs false prefixed controls, **Figure 2b: (right)** Activity for the processing of test items vs re-prefixed controls

References

Marcus Meinzer, Aditi Lahiri, Tobias Flaisch, Ronny Hannemann, and Carsten Eulitz. Opaque for the reader but transparent for the brain: Neural signatures of morphological complexity. *Neuropsychologia*, 47(8-9), 1964-1971. (2009).

Christos Pliatsikas, Linda Wheeldon, Aditi Lahiri, and Peter C. Hansen. Processing of zero-derived words in English: An fMRI investigation. *Neuropsychologia*, 53, 47-53. (2014).