

## **An investigation into whether English listeners are 'stress deaf'**

Becky Taylor ([becky.taylor@york.ac.uk](mailto:becky.taylor@york.ac.uk))

*University of York*

Sam Hellmuth ([sam.hellmuth@york.ac.uk](mailto:sam.hellmuth@york.ac.uk))

*University of York*

Word-accent, which is realised in English as stress, plays an important role in spoken language comprehension cross-linguistically (Cutler, Dahan, & van Donselaar, 1997). However, English speakers have been shown to have difficulty acquiring word-accent in a second language (L2): English learners of Polish perform only slightly better than French learners of Polish in a stress identification task (Kijak, 2009) and English learners of Japanese have difficulty learning which words take which word-accent pattern in Japanese (Taylor, 2011). English listeners thus appear to be 'stress deaf', to adopt the term coined in influential research by Dupoux and Peperkamp (DP) on the perceptual behaviour of French listeners and learners (Dupoux, Pallier, Sebastian & Mehler, 1997; Dupoux, Peperkamp, & Sebastián-Gallés, 2001; Dupoux, Sebastián-Gallés, Navarette, & Peperkamp, 2008).

DP attribute 'stress deafness' to the degree of unpredictability in the word-accent system of a particular language (Peperkamp, Vendelin, & Dupoux, 2010), supporting a lexical statistical model of linguistic knowledge: if word-accent is predictable, learners do not encode it in lexical representations in their first language (L1), nor, crucially, when learning an L2. If this account is correct, English learners should not be stress deaf, as the research above suggests them to be, since English stress is only partially predictable. Kijak (2009) attributes English listeners' apparent stress deafness to the low functional load of stress in English: fewer than two dozen minimal pairs can be found in English which contrast in stress alone (Cutler & Pasveer, 2006). However, Kijak also points out that the differing phonetic correlates of stress in a listener's L1 may affect their ability to perceive the position of L2 stress. In English, the primary phonetic cue to word-accent employed by listeners in word recognition has been shown to be vowel quality differences between stressed/unstressed vowels (Cooper, Cutler, & Wales, 2002; Cutler, Wales, Cooper, & Janssen, 2007). Since unstressed vowels are not reduced in Polish, the English listeners in Kijak's study may simply have failed to perceive stress since the targets displayed no vowel reduction.

In this paper we will present the results of a study (currently in progress) which i) replicates DP's methodology in order to confirm where English stands on their 'stress deafness' continuum, but also ii) introduces the phonetic realisation of word accent as an additional independent variable. Use of DP's robust testing paradigm is needed, since most of the existing data available for the performance of English listeners comes from studies in which they served as a control group, and in which they have performed at ceiling level (Altmann & Kabak, 2010), just as DP's French listeners did, using a purely acoustic strategy, in AX discrimination tasks.

In a perception study, participants (36 naïve English listeners who are speakers of British English) 'learn' pairs of nonsense words from a purported 'new language' which differ minimally in one parameter: word-accent position in Test Condition, ['nama]~[na'ma], or word-medial consonant in Control Condition, [□maba]~[□maga]. A word sequence memory task is then used to elicit responses from participants that tap into their abstract phonological processing, rather than a more surface acoustic memory strategy. To introduce our additional independent variable, participants are divided into three groups. The first group is tested on word accent realised with melodic cues only (Japanese-like word accent), the second group on melodic and dynamic cues (Spanish-like word accent) and the third group on melodic/dynamic cues with the unstressed vowel reduced (Dutch-like word accent).

We hypothesise that native English listeners do encode word-accent in lexical representations (as DP predict) but that these representations are phonetically rich, and thus encode the phonetic cues to word-accent which are used in L1 lexical recall (here, segmental, rather than suprasegmental). If correct these results support a hybrid model of lexical representations (Pierrehumbert, 2003; 2006), combining probabilistic knowledge with phonetic detail.