

Short Communication

Modulations of the auditory M100 in an imitation task

Matthias K. Franken^{*}, Peter Hagoort, Daniel J. Acheson

Max Planck Institute for Psycholinguistics, P.O. Box 310, 6500 AH Nijmegen, The Netherlands

Radboud University Nijmegen, Donders Institute for Brain, Cognition and Behaviour, P.O. Box 9101, 6500 HB Nijmegen, The Netherlands

ARTICLE INFO

Article history:

Accepted 4 January 2015

Available online 2 February 2015

Keywords:

Speech production

MEG

Forward model

Imitation

Individual differences

ABSTRACT

Models of speech production explain event-related suppression of the auditory cortical response as reflecting a comparison between auditory predictions and feedback. The present MEG study was designed to test two predictions from this framework: (1) whether the reduced auditory response varies as a function of the mismatch between prediction and feedback; (2) whether individual variation in this response is predictive of speech-motor adaptation.

Participants alternated between online imitation and listening tasks. In the imitation task, participants began each trial producing the same vowel (/e/) and subsequently listened to and imitated auditorily-presented vowels varying in acoustic distance from /e/.

Results replicated suppression, with a smaller M100 during speaking than listening. Although we did not find unequivocal support for the first prediction, participants with less M100 suppression were better at the imitation task. These results are consistent with the enhancement of M100 serving as an error signal to drive subsequent speech-motor adaptation.

© 2015 Elsevier Inc. All rights reserved.

1. Introduction

Feedback plays a crucial role in speech motor control as it signals a speaker whether a speech motor action was successful or not. In order to account for an extensive number of findings related to adaptation and feedback processing in motor control, a number of theories have posited a monitoring mechanism that utilizes forward models. Here, motor commands sent to speech articulators also send an 'efference copy' through forward models that predict the somatosensory and/or auditory consequences of those commands (Hickok, 2012; Houde & Nagarajan, 2011; Tian & Poeppel, 2010).

The workings of this internal forward model for speech production are thought to be reflected in a reduction of the auditory cortex response to self-produced speech relative to listening to recordings of the same speech. Magneto- and Electrophysiological studies have found a reduction of the M100, a well-known auditory component that occurs roughly 100 ms after audio onset (Naatanen & Picton, 1987). Theoretical models (Guenther, Ghosh, & Tourville, 2006; Hickok, 2012; Houde & Nagarajan, 2011) explain M100 suppression (M100S) as reflecting a comparison mechanism: if the internal forward model's prediction of the auditory

consequences of speech commands matches the actual auditory input, the cortical response is attenuated. When the prediction does not match the auditory feedback entirely, there is a reduction of M100S (i.e., the auditory cortex shows less suppression). This reduction of M100S then acts as an error signal, driving compensatory mechanisms that adapt motor output towards internal, auditory goals.

Over the last decade, a number of properties of M100S have emerged. Using magnetoencephalography (MEG), Houde and colleagues showed that masking the auditory feedback abolished M100S (Houde, Nagarajan, Sekihara, & Merzenich, 2002). Subsequent studies have shown that the amount of suppression can be modulated by properties of the feedback (Behroozmand & Larson, 2011; Heinks-Maldonado, Nagarajan, & Houde, 2006; Ventura, Nagarajan, & Houde, 2009). Such feedback can also reflect self-produced variation (Sitek et al., 2013). For instance, Niziolek, Nagarajan, and Houde (2013) found that the M100S correlates with the distance between a people's production and the centroid of their vowel space. Together, these studies support a view in which M100S reflects a match between predicted and actual auditory feedback. Additional support for this view comes from direct cortical recordings (Chang, Niziolek, Knight, Nagarajan, & Houde, 2013; Flinker et al., 2010), where it was also argued that the reduced M100S may be caused either by less neural suppression (i.e., less SIS) or via neural enhancement on top of stable SIS. This neural enhancement was hypothesized to reflect prediction error in the

^{*} Corresponding author at: Radboud University Nijmegen, Donders Institute for Brain, Cognition and Behaviour, P.O. Box 9101, 6500 HB Nijmegen, The Netherlands.
E-mail address: m.franken@donders.ru.nl (M.K. Franken).

auditory processing, whereas SIS is helpful in distinguishing self-produced speech from external speech. However, few studies have directly linked M100S to behavioral output (Chang et al., 2013). The current study was designed to more clearly establish this link by having people engage in an imitation task in which auditory feedback is critical to performance, and in which motor consequences of mismatch are likely to adapt in real time.

The goal of this study was to replicate M100S in an imitation task and to test two claims of the aforementioned theories of speech motor control. First, if the M100S indexes the match between a prediction and the incoming auditory signal, then the amount of suppression should relate to the degree of mismatch between the prediction and the auditory signal (Behroozmand & Larson, 2011; Heinks-Maldonado et al., 2006; Houde et al., 2002; Liu, Meshman, Behroozmand, & Larson, 2011). Second, if a reduction of M100S serves as an error signal to drive motor adaptation, then individual variation in the amount of suppression should be predictive of the individual variation in imitation aptitude. Specifically, people who show a smaller M100S should show larger adaptations to their speech as they more readily produce error signals that could drive imitation performance.

To test these claims, we measured MEG during online imitation and listening tasks. In the speech imitation task, subjects were instructed to produce the vowel /e/ when a visual cue appeared. At the same moment, subjects heard a recording of themselves producing an auditory stimulus that was the same, close, or far from /e/, and they were asked to imitate this stimulus by adjusting their ongoing vowel production. By varying the acoustic distance between /e/ and the auditory target, we were able to investigate whether this acoustic distance modulated the magnitude of the M100S and people's subsequent speech-motor performance.

2. Results & discussion

2.1. Behavioral performance

In order to assess whether participants appropriately performed the imitation task, an initial analysis focused on people's speech output as a function of time across each imitation condition (see Fig. 1). Results show that for both the first formant (F1) and second formant (F2), participants started at similar values every trial, and the formant values subsequently diverged depending on the vowel. Results of a 9 (Time) $\times$ 5 (Vowel) repeated-measures ANOVA on F1 values (see Fig. 1a) showed significant main effects of both Time ($F(9, 270) = 23.97$; $p < 0.0001$) and Vowel ($F(4, 120) = 145.8$; $p < 0.0001$), as well as a significant Time $\times$ Vowel interaction ($F(36, 1080) = 106.9$; $p < 0.0001$). Similarly, for F2 (Fig. 1b) both main effects as well as the interaction were significant (Time: $F(9, 270) = 9.25$; $p < 0.0001$; Vowel: $F(4, 120) = 79.49$; $p < 0.0001$; interaction: $F(36, 1080) = 60.78$; $p < 0.0001$). With the exception of /i/ condition, all vowels differed from the /e/ condition. Note that although all five vowels were phonemically distinct in Dutch, an important cue to distinguish /i/ from /e/ is vowel duration, a parameter that was lost in this experiment (Adank, van Hout, & Smits, 2004). Importantly, however, these behavioral results demonstrate that participants did imitate the auditory stimuli.

In order to quantify whether the changes in F1 and F2 brought participants closer to the imitation target, we examined the Euclidian distance in F1–F2 space between subjects' speech output and their imitation target on a given trial (Fig. 2c). Results showed significant main effects of Time ($F(9, 270) = 42.1$; $p < 0.0001$) and Vowel ($F(4, 120) = 45.79$; $p < 0.0001$) as well as an interaction ($F(36, 1080) = 34.62$; $p < 0.0001$). Post-hoc tests revealed that for

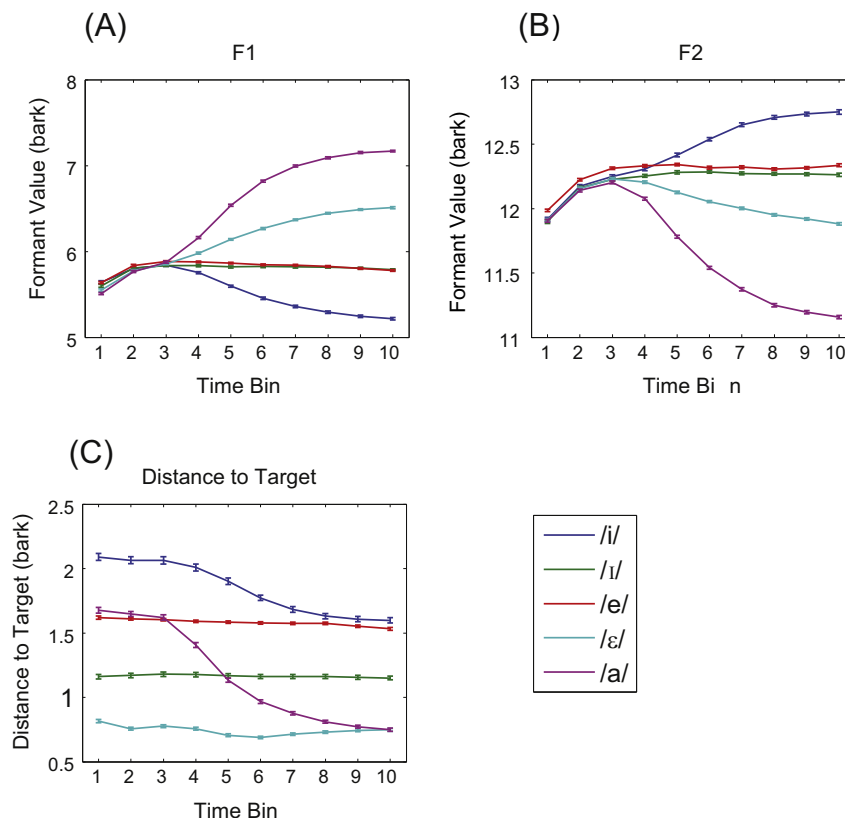


Fig. 1. Average behavioral results across participants. Formant values are expressed in bark (Zwicker, 1961). All error bars represent within-subject standard error of the mean. (A) F1 formant values averaged per time bin and vowel across subjects. Each time bin is 90 ms long. (B) F2 formant values averaged per time bin and condition across subjects. (C) Distance to the imitation target averaged per time bin and condition across subjects.

Download English Version:

<https://daneshyari.com/en/article/7284439>

Download Persian Version:

<https://daneshyari.com/article/7284439>

[Daneshyari.com](https://daneshyari.com)