

UCSF

UC San Francisco Previously Published Works

Title

The FACTS model of speech motor control: Fusing state estimation and task-based control

Permalink

<https://escholarship.org/uc/item/6vd1t91m>

Journal

PLOS Computational Biology, 15(9)

ISSN

1553-734X

Authors

Parrell, Benjamin
Ramanarayanan, Vikram
Nagarajan, Srikantan
et al.

Publication Date

2019

DOI

10.1371/journal.pcbi.1007321

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

RESEARCH ARTICLE

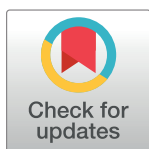
The FACTS model of speech motor control: Fusing state estimation and task-based control

Benjamin Parrell^{1☯*}, Vikram Ramanarayanan^{2,3☯}, Srikantan Nagarajan^{2,4}, John Houde²

1 Department of Communication Sciences and Disorders, University of Wisconsin–Madison, Madison, Wisconsin, United States of America, **2** Department of Otolaryngology - Head and Neck Surgery, University of California, San Francisco, San Francisco, California, United States of America, **3** Educational Testing Service R&D, San Francisco, California, United States of America, **4** Department of Radiology and Biomedical Imaging, University of California, San Francisco, San Francisco, California, United States of America

☯ These authors contributed equally to this work.

* bparrell@wisc.edu



OPEN ACCESS

Citation: Parrell B, Ramanarayanan V, Nagarajan S, Houde J (2019) The FACTS model of speech motor control: Fusing state estimation and task-based control. PLoS Comput Biol 15(9): e1007321. <https://doi.org/10.1371/journal.pcbi.1007321>

Editor: Adrian M Haith, Johns Hopkins University, UNITED STATES

Received: April 4, 2019

Accepted: August 2, 2019

Published: September 3, 2019

Copyright: © 2019 Parrell et al. This is an open access article distributed under the terms of the [Creative Commons Attribution License](https://creativecommons.org/licenses/by/4.0/), which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

Data Availability Statement: All simulation data is available in the Supporting Information files.

Funding: This work was funded by the following grants from the National Institutes of Health (www.nih.gov): R01DC017696 (JH,SN), R01DC013979 (SN, JH), R01NS100440 (SN, JN, MLGT), and R01DC017091 (SN, JH). The funders had no role in study design, data collection and analysis, decision to publish, or preparation of the manuscript.

Competing interests: The authors have declared that no competing interests exist.

Abstract

We present a new computational model of speech motor control: the Feedback-Aware Control of Tasks in Speech or *FACTS* model. *FACTS* employs a hierarchical state feedback control architecture to control simulated vocal tract and produce intelligible speech. The model includes higher-level control of speech tasks and lower-level control of speech articulators. The task controller is modeled as a dynamical system governing the creation of desired constrictions in the vocal tract, after Task Dynamics. Both the task and articulatory controllers rely on an internal estimate of the current state of the vocal tract to generate motor commands. This estimate is derived, based on efference copy of applied controls, from a forward model that predicts both the next vocal tract state as well as expected auditory and somatosensory feedback. A comparison between predicted feedback and actual feedback is then used to update the internal state prediction. *FACTS* is able to qualitatively replicate many characteristics of the human speech system: the model is robust to noise in both the sensory and motor pathways, is relatively unaffected by a loss of auditory feedback but is more significantly impacted by the loss of somatosensory feedback, and responds appropriately to externally-imposed alterations of auditory and somatosensory feedback. The model also replicates previously hypothesized trade-offs between reliance on auditory and somatosensory feedback and shows for the first time how this relationship may be mediated by acuity in each sensory domain. These results have important implications for our understanding of the speech motor control system in humans.

Author summary

Speaking is one of the most complex motor tasks humans perform, but it's neural and computational bases are not well understood. We present a new computational model that generates speech movements by comparing high-level language production goals

with an internal estimate of the current state of the vocal tract. This model reproduces many key human behaviors, including making appropriate responses to multiple types of external perturbations to sensory feedback, and makes a number of novel predictions about the speech motor system. These results have implications for our understanding of healthy speech as well as speech impairments caused by neurological disorders. They also suggest that the mechanisms of control are shared between speech and other motor domains.

Introduction

Producing speech is one of the most complex motor activities humans perform. To produce even a single word, the activity of over 100 muscles must be precisely coordinated in space and time. This precise spatiotemporal control is difficult to master, and is not fully adult-like until the late teenage years [1]. How the brain and central nervous system (CNS) controls this complex system remains an outstanding question in speech motor neuroscience.

Early models of speech relied on servo control [2]. In this type of *feedback control* schema, the current feedback from the *plant* (thing to be controlled—for speech, this would be the articulators of the vocal tract, as well as perhaps the phonatory and respiratory systems) is compared against desired feedback and any discrepancy between the current and desired feedback drives the generation of motor commands to move the plant towards the current production goal. A challenge for any feedback control model of speech is the short, rapid movements that characterize speech motor behavior, with durations in the range of 50–300 ms. This is potentially shorter than the delays in the sensory systems. For speech, measured latencies to respond to external perturbations of the system range from 20–50 ms for unexpected mechanical loads [3, 4] to around 150 ms for auditory perturbations [5, 6]. Therefore, the information about the state of the vocal tract conveyed by sensory feedback to the CNS is delayed in time. Such delays can cause serious problems for feedback control, leading to unstable movements and oscillations around goal states. Furthermore, speech production is possible even in the absence of auditory feedback, as seen in the ability of healthy speakers to produce speech when auditory feedback is masked by loud noise [7, 8]. All of the above factors strongly indicate that speech cannot be controlled purely based on feedback control.

Several alternative approaches have been suggested to address these problems with feedback control in speech production and other motor domains. One approach, the equilibrium point hypothesis [9–11], relegates feedback control to short-latency spinal or brainstem circuits operating on proprioceptive feedback, with high-level control based on pre-planned feedforward motor commands. Speech models of this type, such as the GEPPETO model, are able to reproduce many biomechanical aspects of speech but are not sensitive to auditory feedback [12–16]. Another approach is to combine feedback and feedforward controllers operating in parallel [17, 18]. This is the approach taken by the DIVA model [19–22], which combines feedforward control based on desired articulatory positions with auditory and somatosensory feedback controllers. In this way, DIVA is sensitive to sensory feedback (via the feedback controllers) but capable of producing fast movements despite delayed or absent sensory feedback (via the feedforward controller).

A third approach, widely used in motor control models outside of speech, relies on the concept of state feedback control [23–26]. In this approach, the plant is assumed to have a state that is sufficiently detailed to predict the future behavior of the plant, and a controller drives the state of the plant towards a goal state, thereby accomplishing a desired behavior. A key

concept in state feedback control is that the true state of the plant is not known to the controller; instead, it is only possible to estimate this state from efference copy of applied controls and sensory feedback. The internal state estimate is computed by first predicting the next plant state based on the applied controls. This state prediction is then used to generate predictions of expected feedback from the plant, and a comparison between predicted feedback and actual feedback is then used to correct the state prediction. Thus, in this process, the actual feedback from the plant only plays an indirect role in that it is only one of the inputs used to estimate the current state, making the system robust to feedback delays and noise.

We have earlier proposed a speech-specific instantiation of a state feedback control system [27]. The primary purpose of this earlier work was to establish the plausibility of the state feedback control architecture for speech production and suggest how such an architecture may be implemented in the human central nervous system. Computationally, our previous work built on models that have been developed in non-speech motor domains [24, 25]. Following this work, we implemented the state estimation process as a prototypical Kalman filter [28], which provides an optimal posterior estimate of the plant state given a prior (the efference-copy-based state prediction) and a set of observations (the sensory reafference), assuming certain conditions such as a linear system. We subsequently implemented a one-dimensional model of vocal pitch control based on this framework [29].

However, the speech production system is substantially more complex than our one-dimensional model of pitch. First, speech production requires the multi-dimensional control of redundant and interacting articulators (e.g., lips, tongue tip, tongue body, jaw, etc.). Second, speech production relies on the control of high-level task goals rather than direct control of the articulatory configuration of the plant (e.g., for speech, positions of the vocal tract articulators). For example, speakers are able to compensate immediately for a bite block which fixes the jaw in place, producing essentially normal vowels [30]. Additionally, speakers react to displacement of a speech articulator by making compensatory movements of other articulators: speakers lower the upper lip when the jaw is pulled downward during production of a bilabial [b] [3, 4], and raise the lower lip when the upper lip is displaced upwards during production of [p] [31]. Importantly, these actions are not reflexes, but are specific to the ongoing speech task. No upper lip movement is seen when the jaw is displaced during production of [z] (which does not require the lips to be close), nor is the lower lip movement increased if the upper lip is raised during production of [f] (where the upper lip is not involved). Together, these results strongly indicate that the goal of speech is not to achieve desired positions of each individual speech articulator, but must rather be to achieve some higher-level goal. While most models of speech motor production thus implement control at a higher speech-relevant level, the precise nature of these goals (vocal tract constrictions [32–34], auditory patterns [2, 12, 21], or both [14, 22]) remains an ongoing debate.

One prominent model that employs control of high-level speech tasks rather than direct control of articulatory variables is the Task Dynamic model [32, 35]. In Task Dynamics, the state of the plant (current positions and velocities of the speech articulators) is assumed to be available through proprioception. Importantly, this information is not used to directly generate an error or motor command. Rather, the current state of the plant is used to calculate values for various constrictions in the vocal tract (e.g., the distance between the upper and lower lip, the distance between the tongue tip and palate, etc.). It is these *constrictions*, rather than the positions of the individual articulators, that constitute the goals of speech production in Task Dynamics.

The model proposed here (Feedback Aware Control of Tasks in Speech, or *FACTS*) extends the idea of articulatory state estimation from the simple linear pitch control mechanism of our previous SFC model to the highly non-linear speech articulatory system. This presents three

primary challenges: first, moving from pitch control to articulatory control requires the implementation of control at a higher level of speech-relevant tasks, rather than at the simpler level of articulator positions. To address this issue, FACTS is built upon the Task Dynamics model, as described above. However, unlike the Task Dynamics model, which assumes the state of the vocal tract is directly available through proprioception, here we model the more realistic situation in which the vocal tract state must be estimated from an efference copy of applied motor commands as well as somatosensory and auditory feedback. The second challenge is that this estimation process is highly non-linear. This required that the implementation of the observer as a Kalman filter in SFC be altered, as this estimation process is only applicable to linear systems. Here, we implement state estimation as an Unscented Kalman Filter [36], which is able to account for the nonlinearities in the speech production system and incorporates internal state prediction, auditory feedback, and somatosensory feedback. Lastly, the highly non-linear mapping between articulatory positions and speech acoustics must be approximated in order to predict the auditory feedback during speech. Here, we learn the articulatory-to-acoustic mapping using Locally Weighted Projection Regression or LWPR [37]. Thus, in the proposed FACTS model, we combine the hierarchical architecture and task-based control of Task Dynamics with a non-linear state-estimation procedure to develop a model that is capable of rapid movements with or without sensory feedback, yet is still sensitive to external perturbations of both auditory and somatosensory feedback.

In the following sections, we describe the architecture of the model and the computational implementation of each model process. We then describe the results of a number of model simulations designed to test the ability of the model to simulate human speech behavior. First, we probe the behavior of the model when sensory feedback is limited to a single modality (somatosensation or audition), removed entirely, or corrupted by varying amounts of noise. Second, we test the ability of the model to respond to external auditory or mechanical perturbations. In both of these sections, we show that the model performs similarly to human speech behavior and makes new, testable predictions about the effects of sensory deprivation. Lastly, we explore how sensory acuity in both the auditory and somatosensory pathways affects the model's response to external auditory perturbations. Here, we show that the response to auditory perturbations depends not only on auditory acuity but on somatosensory acuity as well, demonstrating for the first time a potential computational principle that may underlie the demonstrated trade-off in response magnitude to auditory and somatosensory perturbations across human speakers.

Model overview

A schematic control diagram of the FACTS model is shown in Fig 1. Modules that build on Task Dynamics are shown in blue, and the articulatory state estimation process (or observer) is shown in red. Following Task Dynamics, speech tasks in the FACTS model are hypothesized to be desired constrictions in the vocal tract (e.g., close the lips for a [b]). Each of these speech tasks, or gestures, can be specified in terms of its constriction location (where in the vocal tract the constriction is formed) and its constriction degree (how narrow the constriction is). We model each gesture as a separate critically-damped second-order system [32]. Interestingly, similar dynamical behavior has been seen at a neural population level during the planning and execution of reaching movements in non-human primates [38, 39] and recently in human speech movements [40], suggesting that a dynamical systems model of task-level control may be an appropriate first approximation to the neural activity that controls movement production. However, the architecture of the model would also allow for tasks in other control spaces, such as auditory targets (c.f. [13, 19]), though an appropriate task feedback control law

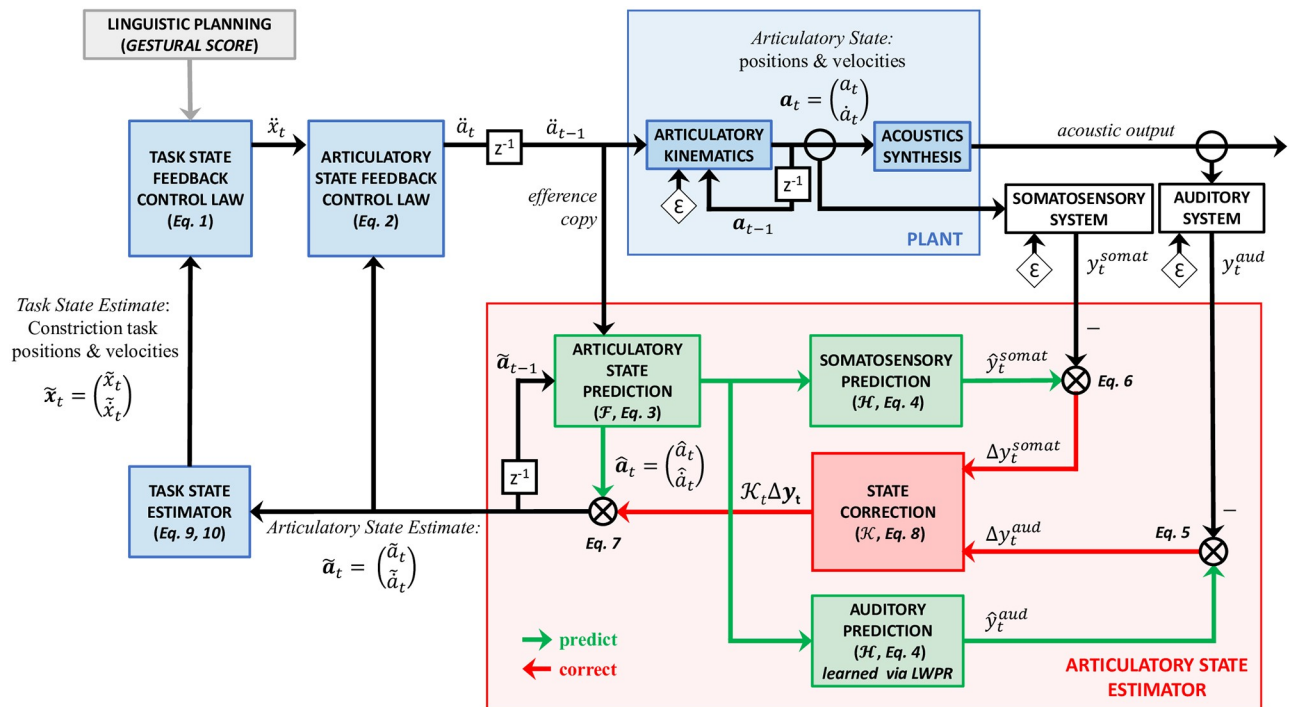


Fig 1. Architecture of the FACTS model. Boxes in blue represent processes outlined in the Task Dynamics model [32, 41]. The articulatory state estimator (shown in red) is implemented as an Unscented Kalman Filter, which estimates the current articulatory state of the plant by combining the predicted state generated by a forward model with auditory and somatosensory feedback. Additive noise is represented by ϵ . Time-step delays are represented by z^{-1} . Equation numbers correspond to equations found in the methods.

<https://doi.org/10.1371/journal.pcbi.1007321.g001>

for such targets would need to be developed (consistent with engineering control theory, we refer to the term “controller” as a “control law”). To what extent, if any, incorporation of auditory targets would impact or alter the results presented here is not immediately clear. This is a promising avenue for future research, as it may provide a way to bridge existing models which posit either constriction- or sensory-based targets. We leave such explorations for future work, but note here that the results presented here may apply only to the current formulation of FACTS with constriction-based targets.

FACTS uses as the Haskins Configurable Articulatory Synthesizer (or CASY) [41–43] as the model of the vocal tract **plant** being controlled. The relevant parameters of the CASY model required to move the tongue body to produce a vowel (and which fully describe the articulatory space for the majority of the simulations in this paper) are the Jaw Angle (JA, angle of the jaw relative to the temporomandibular joint), Condyle Angle (CA, the angle of the center of the tongue relative to the jaw, measured at the temporomandibular joint), and the Condyle Length (CL, distance of the center of the tongue from the temporomandibular joint along the Condyle Angle). The CASY model is shown in Fig 2.

The model begins by receiving the output from a linguistic planning module. Currently, this is implemented as a *gestural score* in the framework of Articulatory Phonology [33, 34]. These gestural scores list the control parameters (e.g., target constriction degree, constriction location, damping, etc.) for each gesture in a desired utterance as well as each gesture’s onset and offset times. For example, the word “mod” ([mad]) has four gestures: simultaneous activation of a gesture driving closure at the lips for [m], a gesture driving an opening of the velum for nasalization of [m], and a gesture driving a wide opening between the tongue and

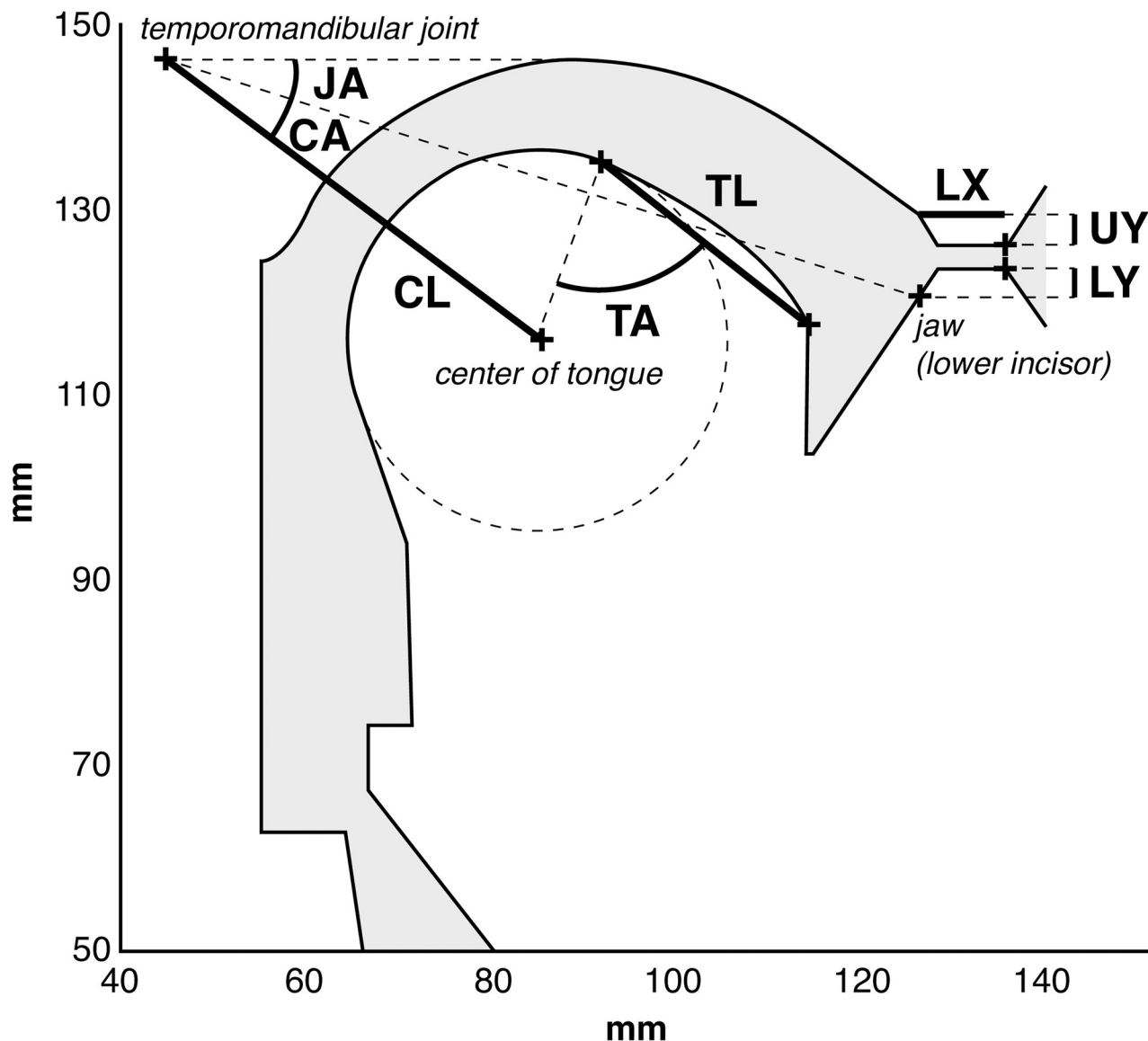


Fig 2. The CASY plant model. Articulatory variable relevant to the current paper are the Jaw Angle (JA), Condyle Angle (CA), and Condyle Length (CL). See text for a description of these variables. Diagram after [78].

<https://doi.org/10.1371/journal.pcbi.1007321.g002>

pharyngeal wall for the vowel [a]. These are followed by a gesture driving closure between the tongue tip and hard palate for [d] (Fig 3).

The **task state feedback control law** takes these gestural scores as input and generates a task-level command based on the current state of the ongoing constriction tasks. In this way, the task-level commands are dependent on the current task-level state. For example, if the lips are already closed during production of a /b/, a very different command needs to be generated than if the lips are far apart. These task-level commands are converted into motor commands that can drive changes in the positions of the speech articulators by the **articulatory state feedback control law**, using information about the current articulatory state of the vocal tract. The motor commands generate changes in the model vocal tract articulators (or **plant**), which are then used to generate an acoustic signal.

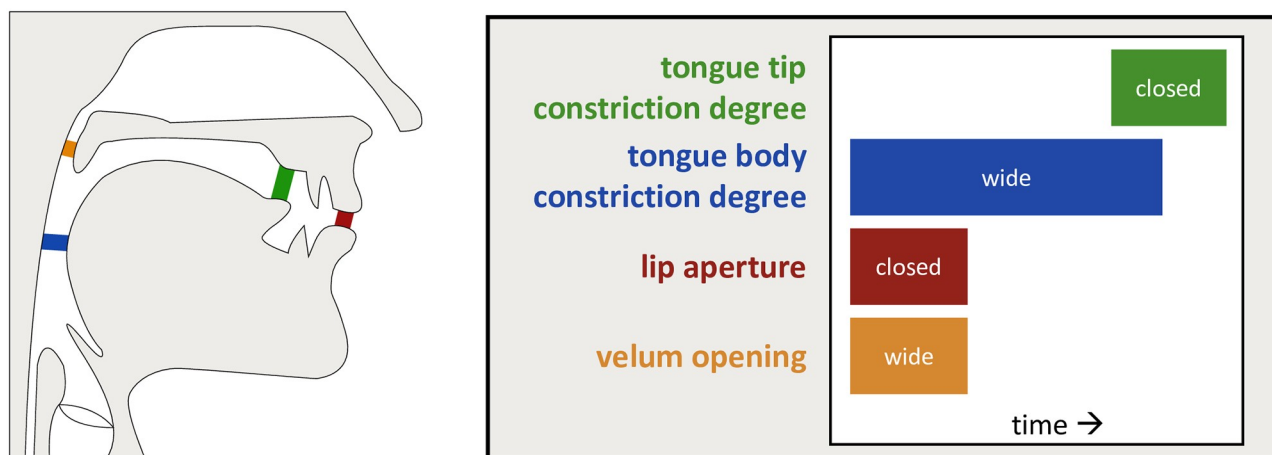


Fig 3. Example of task variables and gestural score for the word “mod”. A gestural score for the word “mod” [mad], which consists of a bilabial closure and velum opening gestures for [m], a wide constriction in the pharynx for [a], and a tongue tip closure gesture for [d]. Tasks are shown on the left, and a schematic of the gestural score on the right.

<https://doi.org/10.1371/journal.pcbi.1007321.g003>

The **articulatory state estimator** (sometimes called an observer in other control models) combines a copy of the outgoing motor command (or efference copy) with auditory and somatosensory feedback to generate an internal estimate of the articulatory state of the plant. First, the efference copy of the motor command is used (in combination with the previous articulatory state estimate) to generate a prediction of the articulatory state. This is then used by a **forward model** (learned here via LWPR) to generate auditory and somatosensory predictions, which are compared to incoming sensory signals to generate sensory errors. Subsequently, these sensory errors are used to correct the state prediction to generate the final state estimate.

The final articulatory state estimate is used by the articulatory state feedback control law to generate the next motor command, as well as being passed to the **task state estimator** to estimate the current task state, or values (positions) and first derivatives (velocities) of the speech tasks (note the Task State was called the Vocal Tract State in earlier presentations of the model [44, 45]). Finally, this estimated task-level state is passed to the task state feedback control law to generate the next task-level command.

A more detailed mathematical description of the model can be found in the methods.

Results

Here we present results showing the accuracy of the learned forward model and of various modeling experiments designed to test the ability of the model to qualitatively replicate human speech motor behavior under various conditions, including both normal speech as well as externally perturbed speech.

Forward model accuracy

Fig 4 visualizes a three dimensional subspace of the learned mapping from the 10-dimensional articulatory state space to the 3-dimensional space of formant frequencies (F1–F3). Specifically, we look at the mapping from the tongue condyle length and condyle angle to the first (see Fig 4A–4C) and second formants (see Fig 4D–4F), projected onto each two-dimensional plane. We also plot normalized histograms of the number of receptive fields that cover each region of the space (represented as a heatmap in Fig 4A and 4D and with a thick blue line in the other subplots). In each figure, the size of the circles is proportional to the absolute value of

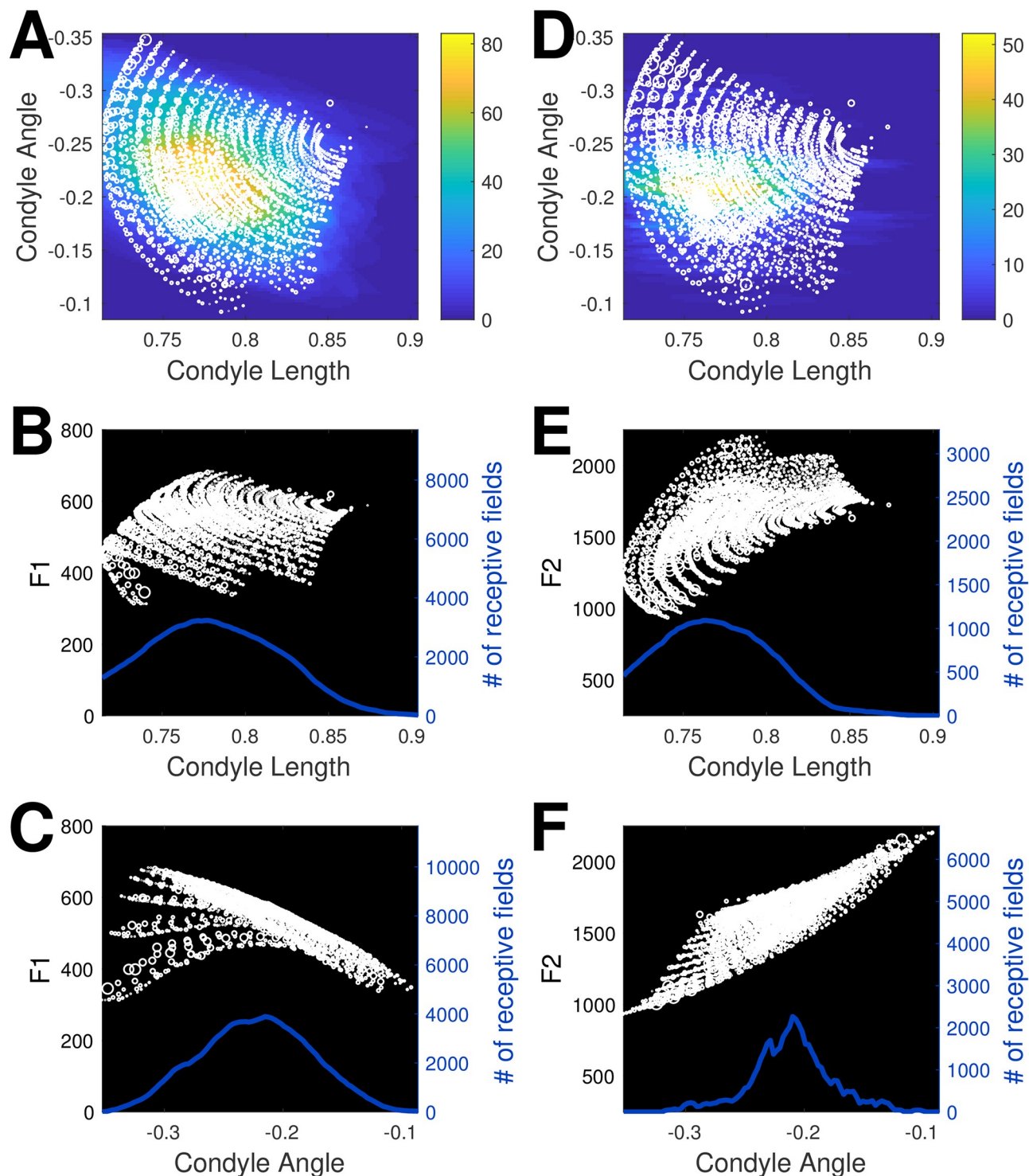


Fig 4. Learned LWPR transformations from CASY articulatory parameters to acoustics. Predicted formant values are shown as white circles. The width of each circle is proportional to the absolute value of the error between the actual formant values and the formant values predicted by the LWPR model. The density distributions reflect the number of receptive fields that cover each point (represented as colors in A, D and the height of the line in B, C, D, F). (A-C) show F1 values. (D-F) show F2 values.

<https://doi.org/10.1371/journal.pcbi.1007321.g004>

the error between the actual and predicted formant values. Overall, the fit of the model results in an average error magnitude of 4.2 Hz (std., 9.1 Hz) for F1 and 6.6 Hz (std., 19.1 Hz) for F2. For comparison, the range of the data was 310–681 Hz for F1 and 930–2197 Hz for F2. Fit error increases in regions of the space that are relatively sparsely covered by receptive fields. In addition, the higher frequency of smaller circles at the margins of the distribution (and therefore the edges of the articulatory space) suggest that we may need fewer receptive fields to cover these regions. Of course, this means that we do see some bigger circles in these regions where the functional mapping is not adequately represented by a small number of fields. Also note that we are only plotting the number of receptive fields that are employed to cover a given region of articulatory space, and this is *not* indicative of how much weight they carry in representing that region of space.

Model response to changes in sensory feedback availability and noise levels

How does the presence or absence of sensory feedback affect the speech motor control system? While there is no direct evidence to date on the effects of total loss of sensory information in human speech, some evidence comes from when sensory feedback from a single modality is attenuated or eliminated. Notably, the effects of removing auditory and somatosensory feedback differ. In terms of auditory feedback, speech production is relatively unaffected by its absence: speech is generally unaffected when auditory feedback is masked by loud masking noise [7, 8]. However, alterations to somatosensory feedback have larger effects: blocking oral tactile sensation through afferent nerve injections or oral anesthesia leads to substantial imprecision in speech articulation [46, 47].

Fig 5 presents simulations from the FACTS model testing the ability of the model to replicate the effects of removing sensory feedback seen in human speech. All simulations modeled the vowel sequence [ə a i]. 100 simulations were run for each of four conditions: normal feedback (5B), somatosensory feedback only (5C), auditory feedback only (5D), and no sensory feedback (5E). For clarity, only the trajectory of the tongue body in the CASY articulatory model is shown for each simulation.

In the normal feedback condition (Fig 5B), the tongue lowers from [ə] to [a], then raises and fronts from [a] to [i]. Note that there is some variability across simulation runs due to the noise in the motor and sensory systems. This variability is also found in human behavior and the ability of the state feedback control architecture to replicate this variability is a strength of this approach [25].

The effect of removing auditory feedback (Fig 5C) leads to a significant, though small, increase in the variability of the tongue body movement as measured by the tongue location at the movement endpoint (Fig 5I), though this effect was not seen in measures of Condyle Angle or Condyle Length variability (Fig 5G and 5H). Interestingly, while variability increased, prediction error slightly decreased in this condition (Fig 5F). Overall, these results are consistent with experimental results that demonstrate that speech is essentially unaffected, in the short term, by the loss of auditory information (though auditory feedback is important for pitch and amplitude regulation [48] as well as to maintain articulatory accuracy in the long term [48–50]).

Removing somatosensory feedback while maintaining auditory feedback (Fig 5D) leads to an increase in both variability across simulation runs as well as an increase in prediction error (Fig 5F–5I). This result is broadly consistent with the fact that reduction of tactile sensation via oral anaesthetic or nerve block leads to imprecise articulation for both consonants and vowels [47, 51] (though the acoustic effects of this imprecision may be less perceptible for vowels [47]). However, a caveat must be made that our current model does not include tactile

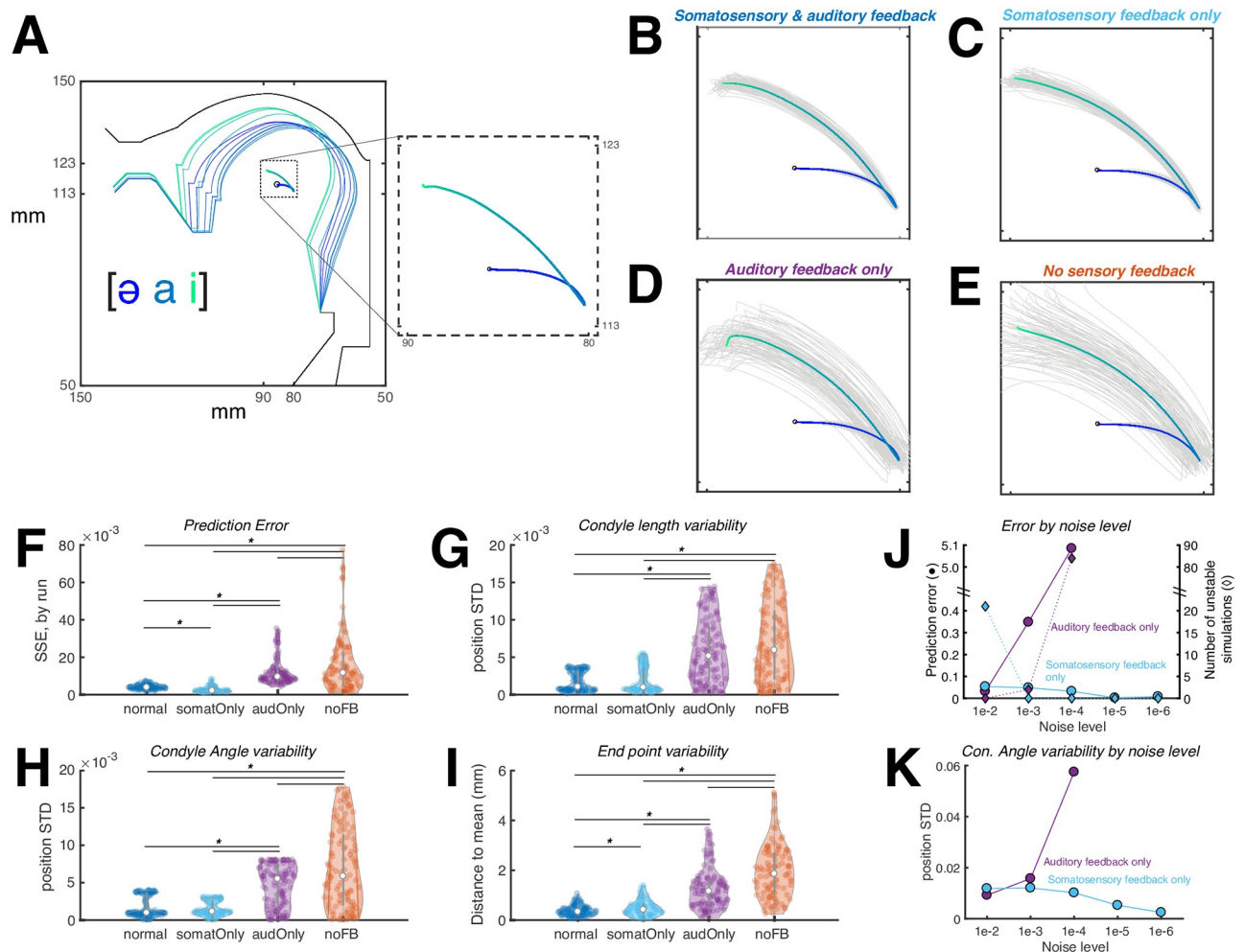


Fig 5. FACTS simulation of the vowel sequence [ə a i], varying the type of sensory feedback available to the model. (A) shows a sample simulation with movements of the CASY model as well as the trajectory of the tongue center. (B-E) each show tongue center trajectories from 100 simulations (gray) and their mean (blue-green gradient) with varying types of sensory feedback available. Variability is lower when sensory feedback is available, and lower when auditory feedback is absent compared to when somatosensory feedback is absent. (F) shows the prediction error in each condition. (G-H) show the produced variability in two articulatory parameters of the CASY plant model related to vowel production and (I) shows variability of the tongue center at the final simulation sample. (J) and (K) show prediction error and articulatory variability relative to sensory noise levels when only one feedback channel is available. Decreasing sensory noise leads to increased accuracy for somatosensation but decreased accuracy for audition. Colors in (F-K) correspond to the colors in the titles of (B-E).

<https://doi.org/10.1371/journal.pcbi.1007321.g005>

sensation, only proprioceptive information. Unfortunately, it is impossible to block proprioceptive information from the tongue, as that afferent information is likely carried along the same nerve (hypoglossal nerve) as the efferent motor commands [52]. It is difficult to prove, then, exactly how a complete loss of proprioception would affect speech. Nonetheless, the current model results are consistent with studies that have shown severe dyskinesia in reaching movements after elimination of proprioception in non-human primates (see [53] for a review) and in human patients with severe proprioceptive deficits [54]. In summary, although the FACTS model currently includes only proprioceptive sensory information rather than both proprioceptive and tactile signals, these simulation results are consistent with a critical role for the somatosensory system in maintaining the fine accuracy of the speech motor control system.

While removal of only auditory feedback lead to only small increases in variability (in both FACTS simulations and human speech), our simulations show speech in the complete absence of sensory feedback (Fig 5E) shows much larger variability than the absence of either auditory or somatosensory feedback alone. This is consistent with human behavior [51], and occurs because without sensory feedback there is no way to detect and correct for the noise inherent in the motor system (shown by the large prediction errors and increased articulatory variability in Fig 5F).

The effects of changing the noise levels in the system can be seen in Fig 5J and 5K. For these simulations, only one type of feedback was used at a time: somatosensory (cyan) or auditory (purple). Noise levels (shown on the x axis) reflect both the sensory system noise and the internal estimate of that noise, which were set to be equal. Each data point reflects 100 stable simulations. Data for the acoustic-only simulations are not shown for noise levels below $1e-5$ as the model became highly unstable in these conditions due to inaccurate articulatory state estimates (the number of unstable or divergent simulations is shown in Fig 5J). For the somatosensory system, the prediction error and articulatory variability (shown here for the Condyle Angle) *decrease* as the noise decreases. However, for the auditory system, both prediction error and articulatory variability *increase* as the noise decreases. Because of the Kalman gain, decreased noise in a sensory or predictive signal leads not only to a more accurate signal, but also to a greater reliance on that signal compared to the internal state prediction. When the system relies more on the somatosensory signal, this results in a more accurate state estimate as the somatosensory signal directly reflects the state of the plant. When the system relies more on the auditory signal, however, this results in a less accurate state estimate as the auditory signal only indirectly reflects the state of the plant as a result of the nonlinear, many-to-one articulatory-to-acoustic mapping of the vocal tract. Thus, relying principally on the auditory signal to estimate the state of the speech articulators leads to inaccuracies in the final estimate and, subsequently, high trial-to-trial variability in movements generated from these estimates.

In sum, FACTS is able to replicate the variability seen in human speech, as well as qualitatively match the effects of both auditory and somatosensory masking on speech accuracy. While the variability of human speech in the absence of proprioceptive feedback remains untested, the FACTS simulation results make a strong prediction that could be empirically tested in future work if some manner of blocking or altering proprioceptive signals could be devised.

Model response to mechanical and auditory perturbations

When a downward mechanical load is applied to the jaw during the production of a consonant, speakers respond by increasing the movements of the other speech articulators in a task-specific manner to achieve full closure of the vocal tract [3, 4, 31]. For example, when the jaw is perturbed during production of a bilabial stops /b/ or /p/, the upper lip moves downward to a greater extent than normal to compensate for the lower jaw position. This upper lip lowering is not found for jaw perturbations during /f/ or /z/, indicating it is specific to sounds produced using the upper lip. Conversely, tongue muscle activity is larger following jaw perturbation for /z/, which involves a constriction made with the tongue tip, but not for /b/, for which the tongue is not actively involved.

The ability to sense and compensate for mechanical perturbations relies on the somatosensory system. We tested the ability of FACTS to reproduce the task-specific compensatory responses to jaw load application seen in human speakers by applying a small downward acceleration to the jaw (Jaw Angle parameter in CASY) starting midway through a consonant closure for stops produced with the lips (/b/) and tongue tip (/d/). The perturbation continued to

the end of the word. As shown in Fig 6A, the model produces greater lowering of the upper lip (as well as greater raising of the lower lip) when the jaw is fixed during production of /b/, but not during /d/, mirroring the observed response in human speech.

In addition to the task-specific response to mechanical perturbations, speakers will also adjust their speech in response to auditory perturbations [55]. For example, when the first vowel formant (F1) is artificially shifted upwards, speakers produce within-utterance compensatory responses by lowering their produced F1. The magnitude of these responses only partially compensates for the perturbation, unlike the complete responses produced for mechanical perturbations. While the exact reason for this partial compensation is not known, it has been hypothesized to relate to small feedback gains [19] or conflict with the somatosensory feedback system [14]. We explore the cause of this partial compensation below, but focus here on the ability of the model to replicate the observed behavior.

To test the ability of the FACTS model to reproduce the observed partial responses to auditory feedback perturbations, we simulated production of a steady-state [ə] vowel. After a brief stabilization period, we abruptly introduced a +100 Hz perturbation of F1 by adding 100 Hz to the perceived F1 signal in the auditory processing stage. This introduced a discrepancy between the produced F1 (shown in black in Fig 6B) and the perceived F1 (shown in blue in Fig 6B). Upon introduction of the perturbation, the model starts to produce a compensatory lowering of F1, eventually reaching a steady value below the unperturbed production. This compensation, like the response in human speakers, is only partial (roughly 20 Hz or 15% of the total perturbation).

Importantly, FACTS produces compensation for auditory perturbations despite having no auditory targets in the current model. Previously, such compensation has been seen as evidence in favor of the existence of auditory targets for speech [14]. In FACTS, auditory perturbations cause a change in the estimated state of the vocal tract on which the task-level and articulatory-level feedback controllers operate. This causes a change in motor behavior compared to the unperturbed condition, resulting in what appears to be compensation for the auditory perturbation. Our model results thus show that this compensation is possible without explicit auditory goals. Of course, these results do not argue that auditory goals do not exist. Rather, we show that they are not necessary for this particular behavior.

Model trade-offs between auditory and somatosensory acuity

The amount of compensation to an auditory perturbation has been found to vary substantially between individuals [55]. One explanation for the inter-individual variability is that the degree of compensation is related to the acuity of the auditory system. Indeed, some studies have found a relationship between magnitude of the compensatory response to auditory perturbation of vowel formants and auditory acuity for vowel formants [56] or other auditory parameters [57]. This point is not uncontroversial, however, as this relationship is not always present and the potential link between somatosensory acuity and response magnitude has not been established [58]. If we assume that acuity is inversely related to the amount of noise in the sensory system, this explanation fits with the UKF implementation of the state estimation procedure in FACTS, where the weight assigned to the auditory error is ultimately related to the estimate of the noise in the auditory system. In Fig 7B, we show that by varying the amount of noise in the auditory system (along with the internal estimate of that noise), we can drive differences in the amount of compensation the model produces to a +100 Hz perturbation of F1. When we double the auditory noise compared to baseline (top), the compensatory response is reduced. When we halve the auditory noise (bottom), the response increases.

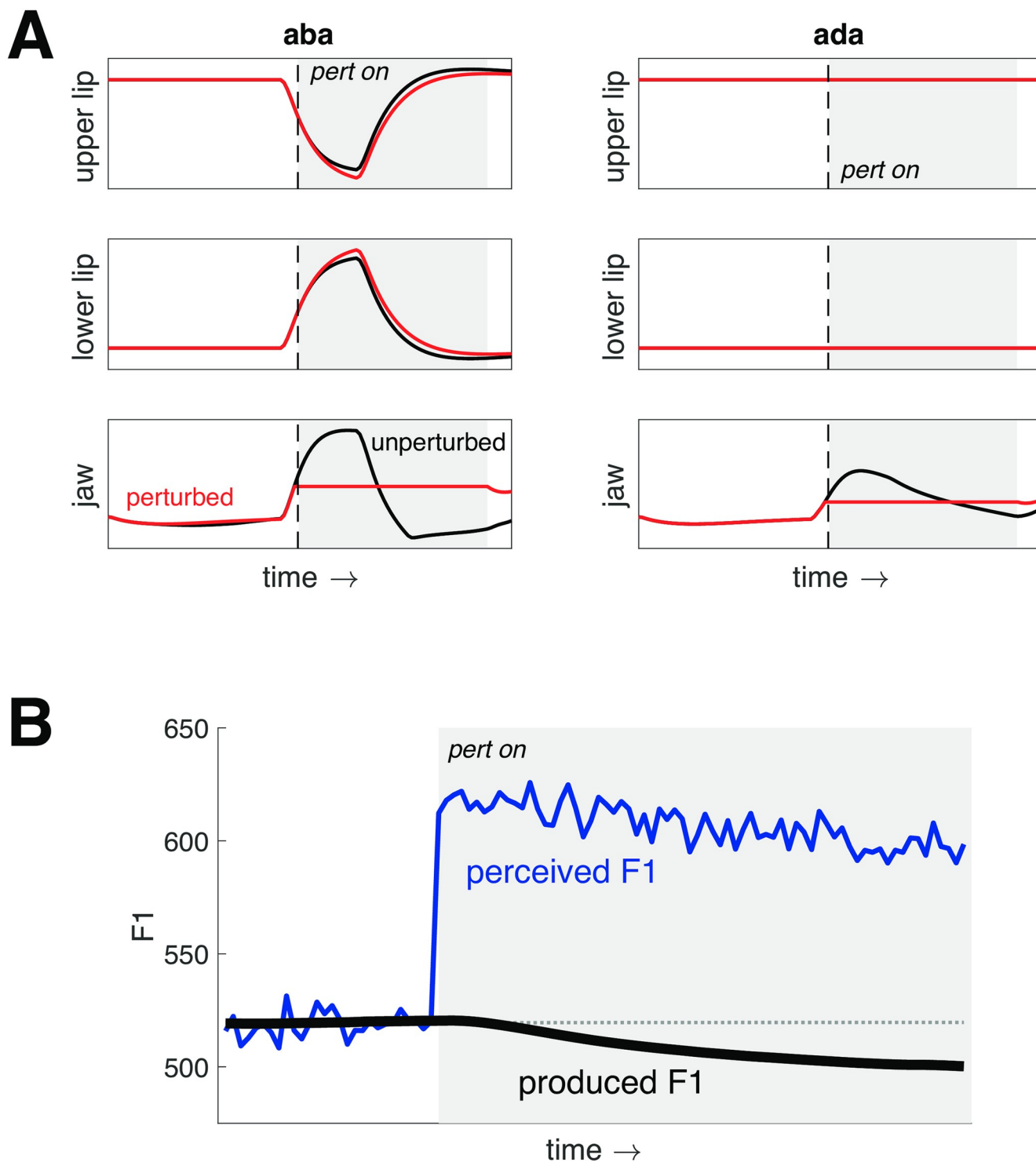


Fig 6. FACTS model simulations of mechanical and auditory perturbations. Times when the perturbations are applied are shown in gray. (A) shows the response of the model to fixing the jaw in place (simulating a downward force applied to the jaw) midway through the closure for second consonant in [aba] (left) and [ada] (right). Unperturbed trajectories are shown in black and perturbed trajectories in red. The upper and lower lips move more to compensate for the jaw perturbation only when the perturbation occurs on [b], mirroring the task-specific response seen in human speakers. (B) shows the response of the FACTS model to a +100 Hz auditory perturbation of F1 while producing [ə]. The produced F1 is shown in black and the perceived F1 is shown in blue. Note that the perceived F1 is corrupted by noise. The model responds to the perturbation by lowering F1 despite the lack of an explicit auditory target. The partial compensation to the perturbation produced by the model matches that observed in human speech.

<https://doi.org/10.1371/journal.pcbi.1007321.g006>

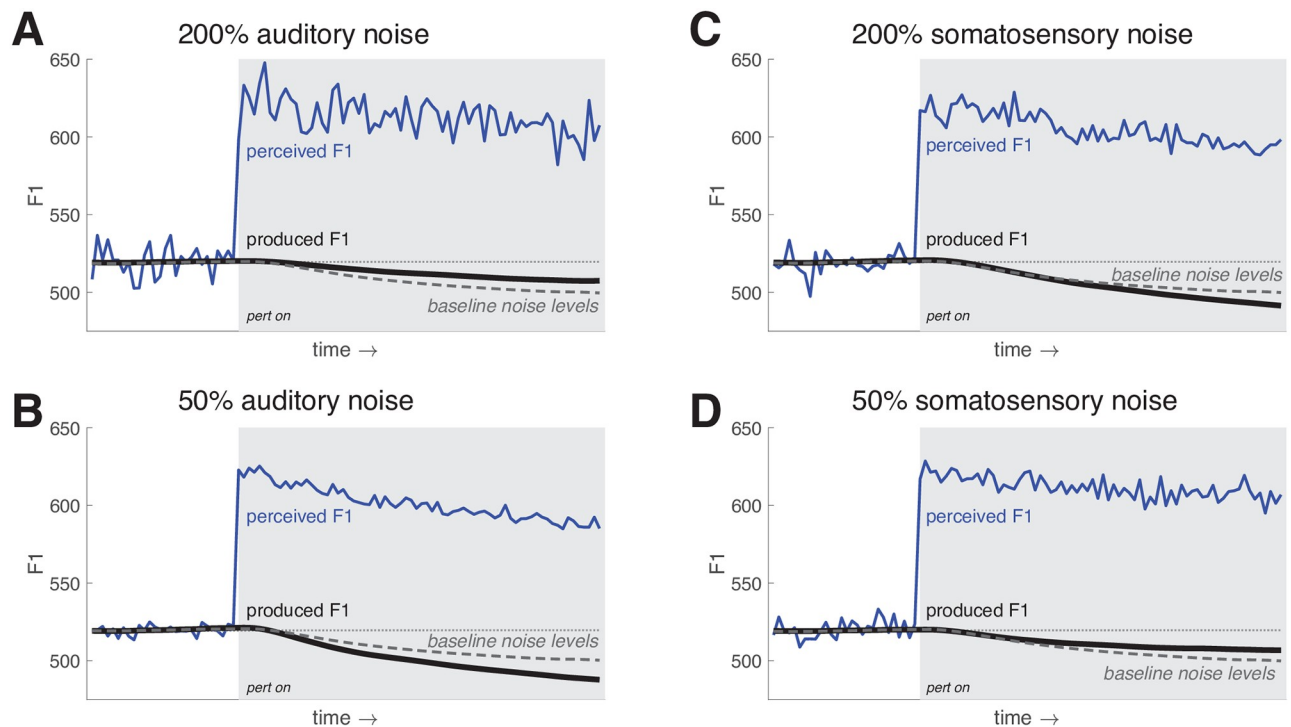


Fig 7. FACTS model simulations of the response to a +100 Hz perturbation of F1. (A) and (B) show the effects of altering the amount of noise in the auditory system in tandem with the observer's estimate of that noise. An increase in auditory noise (A) leads to a smaller perturbation response, while a decrease in auditory noise (B) leads to a larger response. (C) and (D) show the effects of altering the amount of noise in the somatosensory system in tandem with the observer's estimate of that noise. The pattern is the opposite of that for auditory noise. Here, an increase in somatosensory noise (A) leads to a larger perturbation response, while a decrease in somatosensory noise (B) leads to a smaller response. The baseline perturbation response is shown as a dashed grey line in each plot.

<https://doi.org/10.1371/journal.pcbi.1007321.g007>

Interestingly, the math underlying the UKF suggests that the magnitude of the response to an auditory error should be tied not only to the acuity of the auditory system, but to the acuity of the somatosensory system as well. This is because the weights assigned by the Kalman filter take the full noise covariance of all sensory systems into account. We verified this prediction empirically by running a second set of simulated responses to the same +100 Hz perturbation of F1, this time maintaining the level of auditory noise constant while varying only the level of somatosensory noise. The results can be seen in Fig 7C and 7D. When the level of somatosensory noise is increased, the response to the auditory perturbation increases. Conversely, when the level of somatosensory noise is reduced, the compensatory response is reduced as well. These results suggest that the compensatory response in human speakers should be related to the acuity of the somatosensory system as well as the auditory system, a hypothesis which we are currently testing experimentally. Broadly, however, these results agree with, and provide a testable hypothesis about the cause of, empirical findings that show a trading relationship across speakers in their response to auditory and somatosensory perturbations [59].

Discussion

The proposed FACTS model provides a novel way to understand the speech motor system. The model is an implementation of state feedback control that combines high-level control of speech tasks with a non-linear method for estimating the current articulatory state to drive speech motor behavior. We have shown that the model replicates many important

characteristics of human speech motor behavior: the model produces stable articulatory behavior, but with some trial-to-trial variability. This variability increases when somatosensory information is unavailable, but is largely unaffected by the loss of auditory feedback.

The model is also able to reproduce task-specific responses to external perturbations. For somatosensory perturbations, when a downward force is applied to the jaw during production of an oral consonant, there is an immediate task-specific compensatory response only in those articulators needed to produce the current task. This is seen in the increased movement of the upper and lower lips to compensate for the jaw perturbation during production of a bilabial /b/ but no alterations in lip movements when the jaw was perturbed during production of a tongue-tip consonant /d/. The ability of the model to respond to perturbations in a task-specific manner replicates a critical aspect of human speech behavior and is due to the inclusion of the task state feedback control law in the model [35].

For auditory perturbations, we showed that FACTS is able to produce compensatory responses to external perturbations of F1, even though there is no explicit auditory goal in the model. Rather, the auditory signal is used to inform the observer about the current state of the vocal tract articulators. We additionally showed that FACTS is able to produce the inter-individual variability in the magnitude of this compensatory response as well as the previously observed relationship between the magnitude of this response and auditory acuity.

We have also shown that FACTS makes some predictions about the speech motor system that go beyond what has been demonstrated experimentally to date. FACTS predicts that a complete loss of sensory feedback would lead to large increases in articulatory variability beyond those seen in the absence of auditory or somatosensory feedback alone. Additionally, FACTS predicts that the magnitude of compensation for auditory perturbations should be related not only to auditory acuity, but to somatosensory acuity as well. These concrete predictions can be experimentally tested to assess the validity of the FACTS model, testing which is ongoing in our labs.

It is important to note here that, while the FACTS model qualitatively replicates the patterns of variability seen in human movements when feedback is selectively blocked, alternative formulations of the model could potentially lead to a different pattern of results. In the current model, somatosensory and auditory feedback are combined to estimate the state of the speech articulators for low-level articulatory control. Given that somatosensory feedback is more directly informative about this state, it is perhaps unsurprising that removing auditory feedback results in smaller changes in production variability than removing somatosensory feedback. However, auditory feedback may be more directly informative about the task-level state, including cases where the task-level goals are articulatory [60] (as in the current version of the model) or, more obviously, where the task-level goals are themselves defined in the auditory dimension [19]. Indeed, a recent model for limb control has suggested that task-level sensory feedback (vision) is incorporated into a task-level state estimator, rather than being directly integrated with somatosensory feedback in the articulatory controller [61]. A similar use of auditory feedback in the task-level state estimator in FACTS, rather than in the articulatory-level estimator in the current version, may produce different patterns of variability when sensory feedback is blocked. We are currently working on developing such an alternative model to address this issue.

The current version of FACTS uses constriction targets as the goals for task-level control. There are a few considerations regarding this modelling choice that warrant some discussion. First, the ultimate goal of speech production in any theory, at an abstract level, must be to communicate a linguistic message through acoustics. Additionally, all speech movements will necessarily have deterministic acoustic consequences. However, this does not imply that auditory goals must be used at the level of control, which is implemented in FACTS based only on

constriction targets. Second, the current results should not be taken as arguing against the existence of auditory goals. Indeed, we believe that auditory goals may play an important role in speech production. While we have shown that auditory targets are not necessary for compensation to acoustic perturbations, they may well be necessary to explain other behaviors [14, 62]. Future work can test the ability of FACTS to explain these behaviors. Lastly, while the architecture of FACTS is compatible with auditory goals at the task level, the results of the current model may depend on the choice of task-level targets. Again, future work is planned to explore this issue.

One of the major drawbacks of the current implementation of FACTS is that the model of the plant only requires kinematic control of articulatory positions. While a kinematic approach is relatively widespread in the speech motor control field—including both DIVA and Task Dynamics—there is experimental evidence that the dynamic properties of the articulators, such as gravity and tissue elasticity, need to be accounted for [63–66]. Moreover, speakers will learn to compensate for perturbations of jaw protrusion that are dependent on jaw velocity [59, 67–69], indicating that speakers are able to generate motor commands that anticipate and cancel out the effects of those altered articulatory dynamics. While the FACTS model in its current implementation does not replicate this dynamical control of the speech articulators, the overall architecture of the model is compatible with control of dynamics rather than just kinematics [23]. Control of articulatory dynamics would require a dynamic model of the plant and the implementation of a new articulatory-level feedback control law that would output motor commands as forces, rather than (or potentially in addition to) articulatory accelerations. Coupled with parallel changes to the articulatory state prediction process, this would allow for FACTS to control a dynamical plant without any changes to the overall architecture of the model.

Another limitation of the current FACTS model is that it does not incorporate sensory delays in any way. Sensory delays are non-negligible in speech (roughly 30–50 ms for somatosensation and 80–120 ms for audition [19]). We are currently exploring methods to incorporate these delays into the model. One potential avenue is to use an extended state representation, where the state (articulatory and/or task) is represented as a matrix where each column represents a time sample [70]. Interestingly, this approach has shown that shorter-latency signals are assigned higher weights in the Kalman gain, even when they are inherently more noisy. This suggests another potential reason for why the speech system may rely more on somatosensation for online control than audition, since its latency is much shorter.

While a detailed discussion of the neural basis of the computations in FACTS is beyond the scope of the current paper, in order to demonstrate the plausibility of FACTS as a neural model of speech motor control, we briefly touch on potential neural substrates that may underlie a state-feedback control architecture in speech [23, 24, 27]. The cerebellum is widely considered to play a critical role as an internal forward model to predict future articulatory and sensory states [26, 71]. The process of state estimation may occur in the parietal cortex [24], and indeed inhibitory stimulation of the inferior parietal cortex with transcranial magnetic stimulation impairs sensorimotor learning in speech [72], consistent with a role in this process. However, state estimation for speech may also (or alternatively) reside in the ventral premotor cortex (vPMC) for speech, where the premotor cortices are well situated for integrating sensory information (received from sensory cortices via the arcuate fasciculus and the superior longitudinal fasciculus) with motor efference copy from primary motor cortex and cerebellum [27]. Another possible role for the vPMC might be in implementing the task state feedback control law [61]. Primary motor cortex (M1), with its descending control of the vocal tract musculature and bidirectional monosynaptic connections to primary sensory cortex, is the likely location of the articulatory feedback control law, converting task-level commands

from vPMC to articulatory motor commands. Importantly, the differential contributions of vPMC and M1 observed in the movement control literature is consistent with the hierarchical division of task and articulatory control into two distinct levels as specified in FACTS. Interestingly, recent work using electrocorticography has shown that areas in M1 code activation of task-specific muscle synergies similar to those proposed in Task Dynamics and FACTS [73]. This suggests that articulatory control may rely on muscle synergies or motor primitives, rather than the control of individual articulators or muscles [74].

We have currently implemented the state estimation process in FACTS as an Unscented Kalman Filter. We intend this to be purely a mathematically tractable approximation of the actual neural computational process. Interestingly, recent work suggests that a related approach to nonlinear Bayesian estimation, the Neural Particle Filter, may provide a more neurobiologically plausible basis for the state estimation process [75]. Our future extensions of FACTS will involve exploring implementing this type of filter.

In conclusion, the FACTS model uses a widely accepted domain-general approach to motor control, is able to replicate many important speech behaviors, and makes new predictions that can be experimentally tested. This model pushes forward our knowledge of the human speech motor control system, and we plan to further develop the model to incorporate other aspects of speech motor behavior, such as pitch control and sensorimotor learning, in future work.

Methods

Notation

We use the following mathematical notation to present the analyses described in this paper. Matrices are represented by bold uppercase letters (e.g., $\mathbf{X}$), while vectors are represented in italics without any bold case (either upper or lower case). We use the notation $\mathbf{X}^T$ to denote the matrix transpose of $\mathbf{X}$. Concatenations of vectors are represented using bold lowercase letters (e.g., $\mathbf{x} = [x \ \dot{x}]^T$). Scalar quantities are represented without bold and italics. Derivatives and estimates of vectors are represented with dot and tilde superscripts, respectively (i.e., $\dot{\mathbf{x}}$ and $\tilde{\mathbf{x}}$, respectively).

Task state feedback control law

In FACTS, we represent the state of the vocal tract tasks $\mathbf{x}_t = [x_t \ \dot{x}_t]^T$ at time t by a set of constriction task variables x_t (given the current gestural implementation of speech tasks, this is a set of constriction degrees such as lip aperture, tongue tip constriction degree, velic aperture, etc. and constriction locations, such as tongue tip constriction location) and their velocities $\dot{x}_t$. Given a gestural score generated using a linguistic gestural model [76, 77], the *task state feedback control law* (equivalent to the Forward Task Dynamics model in [32]) allows us to generate the dynamical evolution of $\mathbf{x}_t$ using the following simple second-order critically-damped differential equation:

$$\ddot{\mathbf{x}}_t = \mathbf{M}^{-1}(-\mathbf{B}\tilde{\dot{\mathbf{x}}}_t - \mathbf{C}(\tilde{\mathbf{x}}_t - \mathbf{x}_0)) \quad (1)$$

where $\mathbf{x}_0$ is the active task (or gestural) goal, $\mathbf{M}$, $\mathbf{B}$, and $\mathbf{C}$ are respectively the mass matrix, damping coefficient matrix, and stiffness coefficient matrix of the second-order dynamical system model. Essentially, the output of the task feedback controller, $\ddot{\mathbf{x}}_t$, can be seen as a desired change (or command) in task space. This is passed to the articulatory state feedback control law to generate appropriate motor commands that will move the plant to achieve the desired task-level change.

Although the model does include a dynamical formulation of the evolution of speech tasks (following [32, 35]), this is not intended to model the dynamics of the vocal tract plant itself. Rather, the individual speech tasks are modelled as (abstract) dynamical systems.

Articulatory state feedback control law

The desired task-level state change generated by the task feedback control law, $\ddot{x}_t$, is passed to an articulatory feedback control law. In our implementation of this control law, we use Eq 2 (after [32]) to perform an inverse kinematics mapping from the task accelerations $\ddot{x}_t$ to the model articulator accelerations $\ddot{a}_t$, a process which is also dependent on the current estimate of the articulator positions $\tilde{a}_t$ and velocities $\dot{\tilde{a}}_t$. $\mathbf{J}(\tilde{a})$ is the Jacobian matrix of the forward kinematics model relating changes in articulatory states to changes in task states, $\dot{\mathbf{J}}(\tilde{a}, \dot{\tilde{a}})$ is the result of differentiating the elements of $\mathbf{J}(\tilde{a})$ with respect to time, and $\mathbf{J}(\tilde{a})^*$ is a weighted Jacobian pseudoinverse of $\mathbf{J}(\tilde{a})$.

$$\ddot{a}_t = \mathbf{J}(\tilde{a}_t)^* \ddot{x}_t - \dot{\mathbf{J}}(\tilde{a}_t, \dot{\tilde{a}}_t) \dot{\tilde{a}}_t \quad (2)$$

Plant

In order to generate articulatory movements in CASY, we use Runge-Kutta integration to combine the previous articulatory state of the plant ($[a_{t-1} \dot{a}_{t-1}]^T$) with the output of the inverse kinematics computation ($\ddot{a}_{t-1}$, the input to the plant, which we refer to as the **motor command**). This allows us to compute the model articulator positions and velocities for the next time-step ($[a_t \dot{a}_t]^T$), which effectively “moves” the articulatory vocal tract model. Then, a tube-based *synthesis model* converts the model articulator and constriction task values into the output acoustics (parameterized by the vector y_t^{aud}). In order to model noise in the neural system, zero-mean Gaussian white noise ε is added to the motor command ($\ddot{a}_{t-1}$) received by the plant as well as to the somatosensory (y_t^{somat}) and auditory (y_t^{aud}) signals passed from the plant to the articulatory state estimator. Currently, noise levels (standard deviation of Gaussian noise) are tuned by hand for each of these signals (see below for details). Together, the CASY model and the acoustic synthesis process constitute the plant. The model vocal tract in the current implementation of the FACTS model is the Haskins Configurable Articulatory Synthesizer (or CASY) [41–43].

Articulatory state estimator

The articulatory state estimator generates an estimate of the articulatory state of the plant needed to generate state-dependent motor commands. The final state estimate ($\hat{a}_t$) generated by the observer is a combination of an articulatory state prediction ($\tilde{a}_t$) generated from an efference copy of outgoing motor commands, combined with information about the state of the plant derived from the somatosensory and auditory systems (y_t). This combination of internal prediction and sensory information is accomplished through the use of an Unscented Kalman Filter (UKF) [36], which extends the linear Kalman Filter [28] used in most non-speech motor control models [23, 25] to nonlinear systems like the speech production system.

First, the state prediction is generated using a **forward model** ($\mathcal{F}$) that predicts the evolution of the plant based on an estimate of the previous state of the plant ($\tilde{a}_{t-1}$) and an efference copy of the previously issued motor command ($\ddot{a}_{t-1}$). Based on this predicted state, another forward model ($\mathcal{H}$) generates the predicted sensory output $\hat{y}_t = [\hat{y}_t^{somat} \hat{y}_t^{aud}]^T$ (comprising somatosensory and auditory signals $\hat{y}_t^{somat}$ and $\hat{y}_t^{aud}$, respectively) that would be generated by

the plant in the predicted state. Currently, auditory signals are modelled as the first three formant values (F1-F3; 3 dimensions), and somatosensory signals are modelled as the position and velocities of the speech articulators in the CASY model (20 dimensions).

$$\hat{\mathbf{a}}_t = \mathcal{F}(\tilde{\mathbf{a}}_{t-1}, \ddot{\mathbf{a}}_{t-1}) \quad (3)$$

$$\hat{\mathbf{y}}_t = \begin{bmatrix} \mathcal{H}^{somat}(\hat{\mathbf{a}}_t^{somat}) \\ \mathcal{H}^{aud}(\hat{\mathbf{a}}_t^{aud}) \end{bmatrix} \quad (4)$$

These predicted sensory signals are then compared with the incoming signals from the somatosensory (y_t^{somat}) and auditory (y_t^{aud}) systems, generating the sensory prediction error (comprising both somatosensory and auditory components) $\Delta \mathbf{y}_t = [\Delta y_t^{somat} \Delta y_t^{aud}]^T$:

$$\Delta y_t^{aud} = \hat{y}_t^{aud} - y_t^{aud} \quad (5)$$

$$\Delta y_t^{somat} = \hat{y}_t^{somat} - y_t^{somat} \quad (6)$$

These sensory prediction errors are used to correct the initial articulatory state prediction, giving a final articulatory state estimate $\tilde{\mathbf{a}}_t$:

$$\tilde{\mathbf{a}}_t = \hat{\mathbf{a}}_t + \mathcal{K}_t \Delta \mathbf{y}_t \quad (7)$$

where $\mathcal{K}_t$ is the Kalman Gain, which effectively specifies the weights given to the sensory signals in informing the final state estimate. Details of how we generate $\mathcal{F}$, $\mathcal{H}$, and $\mathcal{K}$ are given in the following sections.

Forward models for state and sensory prediction. One of the challenges in estimating the state of the plant is that both the process model $\mathcal{F}$ (that provides a functional mapping from $[\tilde{\mathbf{a}}_{t-1} \tilde{\dot{\mathbf{a}}}_{t-1} \tilde{\ddot{\mathbf{a}}}_{t-1}]^T$ to $\tilde{\mathbf{a}}_t$) as well as the observation model $\mathcal{H}$ (that maps from $\tilde{\mathbf{a}}_t$ to $\hat{\mathbf{y}}_t$) are unknown. Currently, we implement the process model $\mathcal{F}$ by replicating the integration of $\ddot{\mathbf{a}}_t$ used to drive changes in the CASY model, which ignores any potential dynamical effects in the plant. However, the underlying architecture (the forward model $\mathcal{F}$) is sufficiently general that non-zero articulatory dynamics could be accounted for in predicting $\hat{\mathbf{a}}$ as well.

Implementing the observation model is more challenging due to the nonlinear relationship between articulator positions and formant values. In order to solve this problem, we *learn* the observation model functional mappings from articulatory positions to acoustics ($\hat{y}_t^{aud} = \mathcal{H}(\hat{\mathbf{a}}_t)$) required for Unscented Kalman Filtering using Locally Weighted Projection Regression, or LWPR, a computationally efficient machine learning technique [37]. While we do not here explicitly relate this machine learning process to human learning, such maps could theoretically be learned during early speech acquisition, such as babbling [22]. Currently, we learn only the auditory prediction component of $\mathcal{H}$. Since the dimensions of the somatosensory prediction are identical to those of the predicted articulatory state, the former are generated from the latter via an identity function ($\hat{y}_t^{somat} = \hat{\mathbf{a}}_t$).

State correction using an Unscented Kalman Filter. Errors between predicted and afferent sensory signals are used to correct the initial efference-copy-based state prediction through the use of a Kalman gain $\mathcal{K}$, which effectively dictates how each dimension of the sensory error signal $\Delta \mathbf{y}_t$ affects each dimension of the state prediction $\hat{\mathbf{a}}_t$. Our previous SFC model of vocal pitch control implemented a Kalman filter to estimate the weights in the Kalman gain [29], which provides the optimal state estimate under certain strict conditions, including that the system being estimated is linear [28]. However, the traditional Kalman filter approach is only

applicable to linear systems, and the speech production mechanism, even in the simplified CASY model used in FACTS, is highly non-linear.

Our goal in generating a state estimate is to combine the state prediction and sensory feedback to generate an optimal or near-optimal solution. To accomplish this, we use an Unscented Kalman Filter (UKF) [36], which extends the linear Kalman Filter to non-linear systems. While the UKF has not been proven to provide optimal solutions to the state estimation problem, it consistently provides more accurate estimates than other non-linear extensions of the Kalman filter, such as the Extended Kalman Filter [36].

In FACTS, the weights of the Kalman gain are computed as a function of the estimated covariation between the articulatory state and sensory signals, given by the posterior covariance matrices $\mathbf{P}_{\mathbf{a}_t \mathbf{y}_t}$ and $\mathbf{P}_{\mathbf{y}_t \mathbf{y}_t}$ as follows:

$$\mathcal{K}_t = \mathbf{P}_{\mathbf{a}_t \mathbf{y}_t} \mathbf{P}_{\mathbf{y}_t \mathbf{y}_t}^{-1} \quad (8)$$

In order to generate the posterior means (the state and sensory predictions) and covariances (used to calculate $\mathcal{K}$) in an unscented Kalman Filter (UKF) [36], multiple prior points (called sigma points, $\mathcal{X}$) are used. These prior points are computed based on the mean and covariance of the prior state. Each of these points is then projected through the nonlinear forward model function ($\mathcal{F}$ or $\mathcal{H}$), after which the posterior mean and covariance can be calculated from the transformed points. This process is called the unscented transform (Fig 8). This is used both to predict the future state of the system (process model, $\mathcal{F}$) as well as the expected sensory feedback (observation model, $\mathcal{H}$).

First, the sigma points ($\mathcal{X}$) are generated:

$$\mathcal{X}_{t-1} = [\hat{\mathbf{s}}_{t-1} \pm \sqrt{(L + \lambda) \mathbf{P}_{t-1}}] \quad (9)$$

where $\hat{\mathbf{s}}_{t-1} = [\hat{\mathbf{a}}_{t-1}^T \hat{\mathbf{v}}_{t-1}^T \hat{\mathbf{n}}_{t-1}^T]^T$, and $\mathbf{v}$ and $\mathbf{n}$ are estimates of the process noise (noise in the plant articulatory system) and observation noise (noise in the sensory systems), respectively, L is the dimension of the articulatory state $\mathbf{a}$, λ is a scaling factor, and $\mathbf{P}$ is the noise covariance of $\mathbf{a}$, $\mathbf{v}$, and $\mathbf{n}$. In our current implementation, the level of noise for $\mathbf{v}$ and $\mathbf{n}$ are set to be equal to the level of Gaussian noise added to the plant and sensory systems.

The observer then estimates how the motor command $\ddot{a}_t$ would effect the speech articulators by using the integration model ($\mathcal{F}$) to generate the state prediction $\hat{\mathbf{a}}_t = [\hat{a}_t \hat{a}_t']^T$. First, all sigma points reflecting the articulatory state $\mathcal{X}^a$ and process noise $\mathcal{X}^v$ are passed through $\mathcal{F}$:

$$\mathcal{X}_{t|t-1}^a = \mathcal{F}[\mathcal{X}_{t-1}^a, \ddot{a}_{t-1}, \mathcal{X}_{t-1}^v] \quad (10)$$

and the estimated articulatory state is calculated as the weighted sum of the sigma points where the weights (W) are inversely related to the distance of the sigma point from the center of the distribution.

$$\hat{\mathbf{a}}_t = \sum_{i=0}^{2L} W_i \mathcal{X}_{i,t|t-1}^a \quad (11)$$

The expected sensory state ($\hat{y}_t$) is then derived based on the predicted articulatory state in a similar manner, first by projecting the articulatory $\mathcal{X}^a$ and observation noise $\mathcal{X}^n$ sigma points

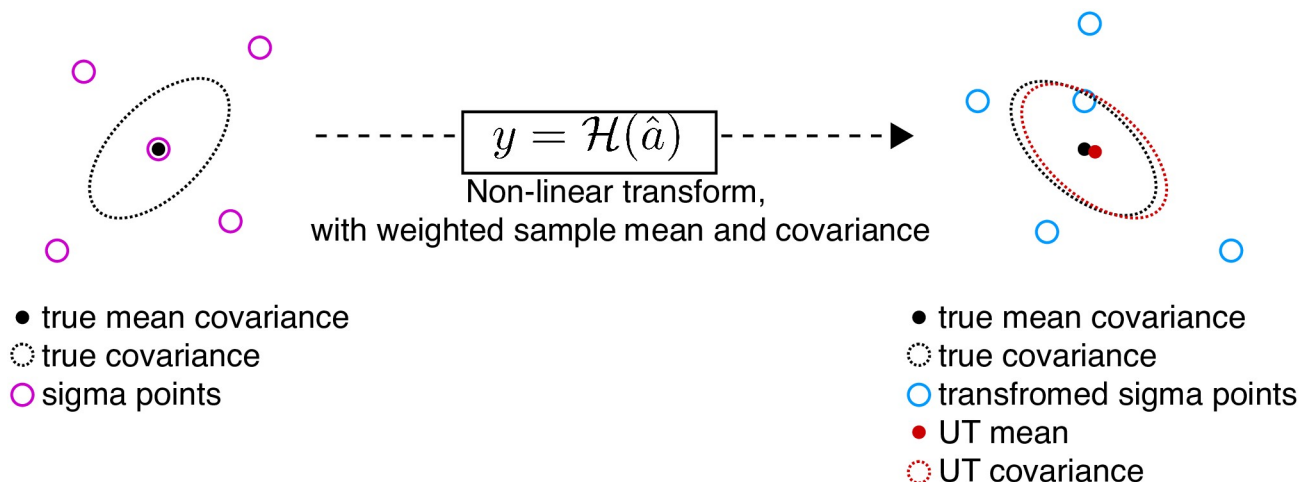


Fig 8. Cartoon showing the unscented transform (UT). The final estimate of the mean and covariance provide a better fit for the true values than would be achieved with the transformation of only a single point at the pre-transformation mean.

<https://doi.org/10.1371/journal.pcbi.1007321.g008>

through the articulatory-to-sensory transform $\mathcal{H}$.

$$\mathcal{Y}_{i,t-1} = \mathcal{H}(\mathcal{X}_{i,t-1}^a, \mathcal{X}_{t-1}^m) \quad (12)$$

$$\hat{\mathbf{y}}_t = \sum_{i=0}^{2L} W_i \mathcal{Y}_{i,t-1} \quad (13)$$

Lastly, the posterior covariance matrices $\mathbf{P}_{\mathbf{x}_t \mathbf{y}_t}$ and $\mathbf{P}_{\mathbf{y}_t \mathbf{y}_t}$ necessary to generate the Kalman Gain $\mathcal{K}$ as well as the Kalman Gain itself are calculated in the following manner:

$$\mathbf{P}_{\mathbf{x}_t \mathbf{y}_t} = \sum_{i=0}^{2L} W_i [\mathcal{X}_{i,t-1} - \hat{\mathbf{a}}_t][\mathcal{Y}_{i,t-1} - \hat{\mathbf{y}}_t]^T \quad (14)$$

$$\mathbf{P}_{\mathbf{y}_t \mathbf{y}_t} = \sum_{i=0}^{2L} W_i [\mathcal{Y}_{i,t-1} - \hat{\mathbf{y}}_t][\mathcal{Y}_{i,t-1} - \hat{\mathbf{y}}_t]^T \quad (15)$$

$$\mathcal{K}_t = \mathbf{P}_{\mathbf{x}_t \mathbf{y}_t} \mathbf{P}_{\mathbf{y}_t \mathbf{y}_t}^{-1} \quad (16)$$

Task state estimator

Finally, we estimate the vocal tract state estimate at the next time step by passing the articulatory state estimate into a task state estimator, which in our current implementation is a forward kinematics model (see Eq 2) [32]. $\mathbf{J}(\tilde{\mathbf{a}})$, the Jacobian matrix relating changes in articulatory states to changes in task states, is the same as in Eq 2.

$$\tilde{\mathbf{x}}_t = f(\tilde{\mathbf{a}}_t) \quad (17)$$

$$\tilde{\dot{\mathbf{x}}}_t = \mathbf{J}(\tilde{\mathbf{a}}_t) \tilde{\dot{\mathbf{a}}}_t \quad (18)$$

This task state estimate is then passed to the task feedback controller to generate the next task-level command $\tilde{\mathbf{x}}_t$ using Eq 1.

Model parameter settings

There are a number of tunable parameters in the FACTS model. These include: 1) the noise added to $\ddot{a}$ in the plant, y_{aud} in the auditory system, and y_{soma} in the somatosensory system; 2) the internal estimates of the process ($\ddot{a}$) and observation y noise; and 3) initial values for the process, observation, and state covariance matrices used in the Unscented Kalman Filter. Internal estimates of the process and observation noise were set to be equal to the true noise levels. Noise levels were selected from a range from $1e-1$ to $1e-8$, scaled by the norm of each signal (equivalent to a SNR range of 10 to $1e8$), to achieve the following goals: 1) stable system behavior in the absence of external perturbations, 2) the ability of the model to react to external auditory and somatosensory perturbations, 3) and a partial compensation for external auditory perturbations in line with observed human behavior. The final noise values used were $1e-4$ for the plant/process noise, $1e-2$ for the auditory noise, and $1e-6$ for the somatosensory noise. The discrepancy in the values for the noise between the two sensory domains is proportional to the difference in magnitude between the two signals (300–3100 Hz for the auditory signal, 0–1.2 mm or mm/s for the articulatory position and velocity signals). Process and observation covariance matrices were initialized as identity matrices scaled by the process and observation noise, respectively. The state covariance matrix was initialized as an identity matrix scaled by $1e-2$. A relatively wide range of noise values produced similar behavior: the effects of changing the auditory and somatosensory noise levels are discussed in the results section.

Supporting information

S1 Data. The supporting information contains data from all simulations shown in this paper in MATLAB format.
(ZIP)

Author Contributions

Conceptualization: Benjamin Parrell, Vikram Ramanarayanan, Srikantan Nagarajan, John Houde.

Funding acquisition: Srikantan Nagarajan, John Houde.

Investigation: Benjamin Parrell, Vikram Ramanarayanan.

Methodology: Benjamin Parrell, Vikram Ramanarayanan.

Software: Benjamin Parrell, Vikram Ramanarayanan.

Visualization: Benjamin Parrell, Vikram Ramanarayanan.

Writing – original draft: Benjamin Parrell, Vikram Ramanarayanan.

Writing – review & editing: Benjamin Parrell, Vikram Ramanarayanan, Srikantan Nagarajan, John Houde.

References

1. Walsh B, Smith A. Articulatory movements in adolescents: evidence for protracted development of speech motor control processes. *J Speech Lang Hear Res.* 2002; 45(6):1119–33. [https://doi.org/10.1044/1092-4388\(2002/090\)](https://doi.org/10.1044/1092-4388(2002/090)) PMID: 12546482
2. Fairbanks G. Systematic Research In Experimental Phonetics.* 1. A Theory Of The Speech Mechanism As A Servosystem. *Journal of Speech and Hearing Disorders.* 1954; 19(2):133–139. <https://doi.org/10.1044/jshd.1902.133> PMID: 13212816

3. Abbs JH, Gracco VL. Control of complex motor gestures: orofacial muscle responses to load perturbations of lip during speech. *J Neurophysiol.* 1984; 51(4):705–23. <https://doi.org/10.1152/jn.1984.51.4.705> PMID: [6716120](#)
4. Kelso JA, Tuller B, Vatikiotis-Bateson E, Fowler CA. Functionally specific articulatory cooperation following jaw perturbations during speech: evidence for coordinative structures. *J Exp Psychol Hum Percept Perform.* 1984; 10(6):812–32. <https://doi.org/10.1037/0096-1523.10.6.812> PMID: [6239907](#)
5. Parrell B, Agnew Z, Nagarajan S, Houde JF, Ivry RB. Impaired Feedforward Control and Enhanced Feedback Control of Speech in Patients with Cerebellar Degeneration. *J Neurosci.* 2017; 37(38):9249–9258. <https://doi.org/10.1523/JNEUROSCI.3363-16.2017> PMID: [28842410](#)
6. Cai S, Beal DS, Ghosh SS, Tiede MK, Guenther FH, Perkell JS. Weak responses to auditory feedback perturbation during articulation in persons who stutter: evidence for abnormal auditory-motor transformation. *PLoS One.* 2012; 7(7):e41830. <https://doi.org/10.1371/journal.pone.0041830> PMID: [22911857](#)
7. Lane H, Tranel B. The Lombard Sign and the Role of Hearing in Speech. *Journal of Speech, Language, and Hearing Research.* 1971; 14(4):677–709. <https://doi.org/10.1044/jshr.1404.677>
8. Lombard E. Le signe de l'elevation de la voix. *Ann Maladies de L'Oreille et du Larynx.* 1911; 37:2.
9. Feldman A. Once more on the Equilibrium Point Hypothesis (A) for motor control. *Journal of Motor Behavior.* 1986; 18:17–54. <https://doi.org/10.1080/00222895.1986.10735369> PMID: [15136283](#)
10. Feldman A, Adamovich S, Ostry D, Flanagan J. The origin of electromyograms—Explanations based on the Equilibrium Point Hypothesis. Winters J, Woo S, editors. New York: Springer Verlag; 1990.
11. Feldman AG, Levin MF. The Equilibrium-Point Hypothesis—Past, Present and Future. In: *Progress in Motor Control. Advances in Experimental Medicine and Biology.* Springer, Boston, MA; 2009. p. 699–726.
12. Perrier P, Ma L, Payan Y. Modeling the production of VCV sequences via the inversion of a biomechanical model of the tongue. In: *Proceeding of the INTERSPEECH: Interspeech'2005 - Eurospeech, 9th European Conference on Speech Communication and Technology, Lisbon, Portugal, September 4-8, 2005; 2005.* p. 1041–1044.
13. Perrier P, Ostry D, Laboissière R. The Equilibrium Point Hypothesis and Its Application to Speech Motor Control. *Journal of Speech and Hearing Research.* 1996; 39:365–378. <https://doi.org/10.1044/jshr.3902.365> PMID: [8729923](#)
14. Perrier P, Fuchs SF. Motor Equivalence in Speech Production. In: Redford M, editor. *The Handbook of Speech Production.* Hoboken, NJ: Wiley-Blackwell; 2015. p. 225–247.
15. Sanguineti V, Laboissière R, Ostry DJ. A dynamic biomechanical model for neural control of speech production. *The Journal of the Acoustical Society of America.* 1998; 103(3):1615–1627. <https://doi.org/10.1121/1.421296> PMID: [9514026](#)
16. Laboissière R, Ostry DJ, Feldman AG. The control of multi-muscle systems: human jaw and hyoid movements. *Biological Cybernetics.* 1996; 74(4):373–384. <https://doi.org/10.1007/BF00194930> PMID: [8936389](#)
17. Kawato M, Gomi H. A computational model of four regions of the cerebellum based on feedback-error learning. *Biological Cybernetics.* 1992; 68(2):95–103. <https://doi.org/10.1007/BF00201431> PMID: [1486143](#)
18. Arbib MA. Perceptual Structures and Distributed Motor Control. In: Brookhart JM, Mountcastle VB, Brooks V, editors. *Handbook of Physiology, Supplement 2: Handbook of Physiology, The Nervous System, Motor Control;* 1981. p. 1449–1480.
19. Guenther FH. *Neural control of speech.* Cambridge, MA: The MIT Press; 2015.
20. Guenther FH. Speech Sound Acquisition, Coarticulation, and Rate Effects in a Neural Network Model of Speech Production. *Psychological Review.* 1995; 102:594–621. <https://doi.org/10.1037/0033-295X.102.3.594> PMID: [7624456](#)
21. Guenther FH, Hampson M, Johnson D. A theoretical investigation of reference frames for the planning of speech movements. *Psychological Review.* 1998; 105:611–633. <https://doi.org/10.1037/0033-295X.105.4.611-633> PMID: [9830375](#)
22. Tourville JA, Guenther FH. The DIVA model: A neural theory of speech acquisition and production. *Lang Cogn Process.* 2011; 26(7):952–981. <https://doi.org/10.1080/01690960903498424> PMID: [23667281](#)
23. Scott SH. The computational and neural basis of voluntary motor control and planning. *Trends Cogn Sci.* 2012; 16(11):541–9. <https://doi.org/10.1016/j.tics.2012.09.008> PMID: [23031541](#)
24. Shadmehr R, Krakauer JW. A computational neuroanatomy for motor control. *Exp Brain Res.* 2008; 185(3):359–81. <https://doi.org/10.1007/s00221-008-1280-5> PMID: [18251019](#)
25. Todorov E, Jordan MI. Optimal feedback control as a theory of motor coordination. *Nat Neurosci.* 2002; 5(11):1226–35. <https://doi.org/10.1038/nn963> PMID: [12404008](#)

26. Wolpert DM, Miall RC. Forward Models for Physiological Motor Control. *Neural Netw.* 1996; 9(8):1265–1279. [https://doi.org/10.1016/S0893-6080\(96\)00035-4](https://doi.org/10.1016/S0893-6080(96)00035-4) PMID: 12662535
27. Houde JF, Nagarajan SS. Speech production as state feedback control. *Front Hum Neurosci.* 2011; 5:82. <https://doi.org/10.3389/fnhum.2011.00082> PMID: 22046152
28. Kalman RE. A New Approach to Linear Filtering and Prediction Problems. *Journal of Basic Engineering.* 1960; 82(1):35–45. <https://doi.org/10.1115/1.3662552>
29. Houde JF, Niziolek C, Kort N, Agnew Z, Nagarajan SS. Simulating a state feedback model of speaking. In: 10th International Seminar on Speech Production; 2014. p. 202–205.
30. Fowler CA, Turvey MT. Immediate compensation in bite-block speech. *Phonetica.* 1981; 37(5-6):306–326. <https://doi.org/10.1159/000260000> PMID: 7280033
31. Shaiman S, Gracco VL. Task-specific sensorimotor interactions in speech production. *Experimental Brain Research.* 2002-10-01; 146(4):411–418. <https://doi.org/10.1007/s00221-002-1195-5> PMID: 12355269
32. Saltzman EL, Munhall KG. A dynamical approach to gestural patterning in speech production. *Ecological Psychology.* 1989; 1(4):333–382. https://doi.org/10.1207/s15326969eco0104_2
33. Browman CP, Goldstein L. Articulatory phonology: An overview. *Phonetica.* 1992; 49(3-4):155–180. <https://doi.org/10.1159/000261913> PMID: 1488456
34. Browman CP, Goldstein L. Dynamics and articulatory phonology. In: Port R, van Gelder T, editors. *Mind as motion: Explorations in the dynamics of cognition.* Boston: MIT Press; 1995. p. 175–194.
35. Saltzman E. Task dynamic coordination of the speech articulators: A preliminary model. *Experimental Brain Research Series.* 1986; 15:129–144.
36. Wan EA, Van Der Merwe R. The unscented Kalman filter for nonlinear estimation. In: *Adaptive Systems for Signal Processing, Communications, and Control Symposium 2000. AS-SPCC. The IEEE 2000.* Ieee; 2000. p. 153–158.
37. Mitrovic D, Klanke S, Vijayakumar S. Adaptive optimal feedback control with learned internal dynamics models. In: *From Motor Learning to Interaction Learning in Robots.* Springer; 2010. p. 65–84.
38. Shenoy KV, Sahani M, Churchland MM. Cortical control of arm movements: a dynamical systems perspective. *Annu Rev Neurosci.* 2013; 36:337–59. <https://doi.org/10.1146/annurev-neuro-062111-150509> PMID: 23725001
39. Churchland MM, Cunningham JP, Kaufman MT, Foster JD, Nuyujukian P, Ryu SI, et al. Neural population dynamics during reaching. *Nature.* 2012; 487(7405):51–6. <https://doi.org/10.1038/nature11129> PMID: 22722855
40. Stavisky SD, Willett FR, Murphy BA, Rezaii P, Memberg WD, Miller JP, et al. Neural ensemble dynamics in dorsal motor cortex during speech in people with paralysis. *bioRxiv.* 2018.
41. Saltzman E, Nam H, Krivokapic J, Goldstein L. A task-dynamic toolkit for modeling the effects of prosodic structure on articulation. In: *Proceedings of the 4th International Conference on Speech Prosody (Speech Prosody 2008), Campinas, Brazil; 2008.*
42. Nam H, Mitra V, Tiede M, Hasegawa-Johnson M, Espy-Wilson C, Saltzman E, et al. A procedure for estimating gestural scores from speech acoustics. *The Journal of the Acoustical Society of America.* 2012; 132(6):3980–3989. <https://doi.org/10.1121/1.4763545> PMID: 23231127
43. Rubin P, Saltzman E, Goldstein L, McGowan R, Tiede M, Browman C. CASY and extensions to the task-dynamic model. In: *1st ETRW on Speech Production Modeling: From Control Strategies to Acoustics; 4th Speech Production Seminar: Models and Data, Autrans, France; 1996.*
44. Ramanarayanan V, Parrell B, Goldstein L, Nagarajan S, Houde J. A New Model of Speech Motor Control Based on Task Dynamics and State Feedback. In: *INTERSPEECH; 2016.* p. 3564–3568.
45. Parrell B, Ramanarayanan V, Nagarajan S, Houde JF. FACTS: A hierarchical task-based control model of speech incorporating sensory feedback. In: *Interspeech 2018; 2018.*
46. Ringel RL, Steer MD. Some Effects of Tactile and Auditory Alterations on Speech Output. *Journal of Speech, Language, and Hearing Research.* 1963; 6(4):369–378. <https://doi.org/10.1044/jshr.0604.369>
47. Scott CM, Ringel RL. Articulation without oral sensory control. *Journal of Speech, Language, and Hearing Research.* 1971; 14(4):804–818. <https://doi.org/10.1044/jshr.1404.804>
48. Lane H, Webster JW. Speech deterioration in postlingually deafened adults. *The Journal of the Acoustical Society of America.* 1991; 89(2):859–866. <https://doi.org/10.1121/1.1894647> PMID: 2016436
49. Cowie R, Douglas-Cowie E. Postlingually acquired deafness: speech deterioration and the wider consequences. vol. 62. *Walter de Gruyter; 1992.*
50. Perkell JS. Movement goals and feedback and feedforward control mechanisms in speech production. *Journal of Neurolinguistics.* 2012; 25(5):382–407. <https://doi.org/10.1016/j.jneuroling.2010.02.011> PMID: 22661828

51. Putnam AHB, Ringel RL. A Cineradiographic Study of Articulation in Two Talkers with Temporarily Induced Oral Sensory Deprivation. *Journal of Speech, Language, and Hearing Research*. 1976; 19(2):247–266. <https://doi.org/10.1044/jshr.1902.247>
52. Borden GJ. The Effect of Mandibular Nerve Block Upon the Speech of Four-Year-Old Boys. *Language and Speech*. 1976; 19(2):173–178. <https://doi.org/10.1177/002383097601900208> PMID: 1018564
53. Desmurget M, Grafton S. Feedback or Feedforward Control: End of a dichotomy. Taking action: Cognitive neuroscience perspectives on intentional acts. 2003; p. 289–338.
54. Gordon J, Ghilardi MF, Ghez C. Impairments of reaching movements in patients without proprioception. I. Spatial errors. *Journal of Neurophysiology*. 1995; 73(1):347–360. <https://doi.org/10.1152/jn.1995.73.1.347> PMID: 7714577
55. Purcell DW, Munhall KG. Compensation following real-time manipulation of formants in isolated vowels. *J Acoust Soc Am*. 2006; 119(4):2288–97. <https://doi.org/10.1121/1.2173514> PMID: 16642842
56. Villacorta VM, Perkell JS, Guenther FH. Sensorimotor adaptation to feedback perturbations of vowel acoustics and its relation to perception. *J Acoust Soc Am*. 2007; 122(4):2306–19. <https://doi.org/10.1121/1.2773966> PMID: 17902866
57. Martin CD, Niziolek CA, Duñabeitia JA, Perez A, Hernandez D, Carreiras M, et al. Online Adaptation to Altered Auditory Feedback Is Predicted by Auditory Acuity and Not by Domain-General Executive Control Resources. *Frontiers in Human Neuroscience*. 2018; 12. <https://doi.org/10.3389/fnhum.2018.00091>
58. Feng Y, Gracco VL, Max L. Integration of auditory and somatosensory error signals in the neural control of speech movements. *Journal of Neurophysiology*. 2011; 106(2):667–79. <https://doi.org/10.1152/jn.00638.2010> PMID: 21562187
59. Lametti DR, Nasir SM, Ostry DJ. Sensory preference in speech production revealed by simultaneous alteration of auditory and somatosensory feedback. *J Neurosci*. 2012; 32(27):9351–8. <https://doi.org/10.1523/JNEUROSCI.0404-12.2012> PMID: 22764242
60. Iskarous K. Vowel constrictions are recoverable from formants. *Journal of Phonetics*. 2010; 38(3):375–387. <https://doi.org/10.1016/j.wocn.2010.03.002> PMID: 20871808
61. Haar S, Donchin O. A revised computational neuroanatomy for motor control. *bioRxiv*. 2019.
62. Nieto-Castanon A, Guenther FH, Perkell J, Curtin H. A modeling investigation of articulatory variability and acoustic stability during American English /r/production. *Journal of the Acoustical Society of America*. 2005; 117(5):3196–3212. <https://doi.org/10.1121/1.1893271> PMID: 15957787
63. Shiller DM, Ostry DJ, Gribble PL. Effects of Gravitational Load on Jaw Movements in Speech. *Journal of Neuroscience*. 1999; 19(20):9073–9080. <https://doi.org/10.1523/JNEUROSCI.19-20-09073.1999> PMID: 10516324
64. Shiller DM, Ostry DJ, Gribble PL, Laboissière R. Compensation for the Effects of Head Acceleration on Jaw Movement in Speech. *Journal of Neuroscience*. 2001; 21(16):6447–6456. <https://doi.org/10.1523/JNEUROSCI.21-16-06447.2001> PMID: 11487669
65. Ostry DJ, Gribble PL, Gracco VL. Coarticulation of jaw movements in speech production: is context sensitivity in speech kinematics centrally planned? *The Journal of Neuroscience*. 1996; 16(4):1570–1579. <https://doi.org/10.1523/JNEUROSCI.16-04-01570.1996> PMID: 8778306
66. Perrier P, Payan Y, Zandipour M, Perkell J. Influence of tongue biomechanics on speech movements during the production of velar stop consonants: A modeling study. *Journal of the Acoustical Society of America*. 2003; 114(3):1582–1599. <https://doi.org/10.1121/1.1587737> PMID: 14514212
67. Tremblay S, Shiller DM, Ostry DJ. Somatosensory basis of speech production. *Nature*. 2003; 423(6942):866. <https://doi.org/10.1038/nature01710> PMID: 12815431
68. Tremblay S, Ostry D. The Achievement of Somatosensory Targets as an Independent Goal of Speech Production -Special Status of Vowel-to-Vowel Transitions. In: Divenyi P, Greenberg S, Meyer G, editors. *Dynamics of Speech Production and Perception*. Amsterdam, The Netherlands: IOS Press; 2006. p. 33–43.
69. Nasir SM, Ostry DJ. Somatosensory Precision in Speech Production. *Current Biology*. 2006; 16(19):1918–1923. <https://doi.org/10.1016/j.cub.2006.07.069> PMID: 17027488
70. Crevecoeur F, Munoz DP, Scott SH. Dynamic Multisensory Integration: Somatosensory Speed Trumps Visual Accuracy during Feedback Control. *Journal of Neuroscience*. 2016; 36(33):8598–8611. <https://doi.org/10.1523/JNEUROSCI.0184-16.2016> PMID: 27535908
71. Herzfeld DJ, Shadmehr R. Cerebellum estimates the sensory state of the body. *Trends Cogn Sci*. 2014; 18(2):66–7. <https://doi.org/10.1016/j.tics.2013.10.015> PMID: 24263038
72. Shum M, Shiller DM, Baum SR, Gracco VL. Sensorimotor integration for speech motor learning involves the inferior parietal cortex. *Eur J Neurosci*. 2011; 34(11):1817–22. <https://doi.org/10.1111/j.1460-9568.2011.07889.x> PMID: 22098364

73. Chartier J, Anumanchipalli GK, Johnson K, Chang EF. Encoding of Articulatory Kinematic Trajectories in Human Speech Sensorimotor Cortex. *Neuron*. 2018; 98(5):1042–1054.e4. <https://doi.org/10.1016/j.neuron.2018.04.031> PMID: 29779940
74. Ramanarayanan V, Goldstein L, Narayanan SS. Spatio-temporal articulatory movement primitives during speech production: Extraction, interpretation, and validation. *The Journal of the Acoustical Society of America*. 2013; 134(2):1378–1394. <https://doi.org/10.1121/1.4812765> PMID: 23927134
75. Kutschireiter A, Surace SC, Sprekeler H, Pfister JP. Nonlinear Bayesian filtering and learning: a neuronal dynamics for perception. *Scientific Reports*. 2017; 7(1):8722. <https://doi.org/10.1038/s41598-017-06519-y> PMID: 28821729
76. Nam H, Goldstein L, Saltzman E. Self-organization of syllable structure: a coupled oscillator model. In: Pellegrino F, Marisco E, Chitoran I, Coupé C, editors. *Approaches to phonological complexity*. Berlin/ New York: Mouton de Gruyter; 2009. p. 299–328.
77. Goldstein L, Nam H, Saltzman E, Chitoran I. Coupled oscillator planning model of speech timing and syllable structure. In: Fant G, Fujisaki H, Shen J, editors. *Frontiers in phonetics and speech science*. Beijing: The Commercial Press; 2009. p. 239–249.
78. Lammert A, Goldstein L, Narayanan S, Iskarous K. Statistical Methods for Estimation of Direct and Differential Kinematics of the Vocal Tract. *Speech Commun*. 2013; 55(1):147–161. <https://doi.org/10.1016/j.specom.2012.08.001> PMID: 24052685