

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Simulating the Elimination and Enhancement of the Plosivity Effect in Reading Aloud

Permalink

<https://escholarship.org/uc/item/7485289s>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 32(32)

ISSN

1069-7977

Authors

Liu, Qiang
Kawamoto, Alan

Publication Date

2010

Peer reviewed

Simulating the Elimination and Enhancement of the Plosivity Effect in Reading Aloud

Qiang Liu (qxliu@ucsc.edu) & Alan H. Kawamoto (ahk@ucsc.edu)

Department of Psychology, University of California at Santa Cruz
Santa Cruz, CA 95064 USA

Abstract

In this paper, we consider articulatory processes in a connectionist model of reading aloud to account for effects of manner of articulation of the initial segment in a variety of tasks. We first describe experimental results showing how flexibility in articulation can completely eliminate the *a priori* acoustic latency difference between plosives and non-plosives in some tasks, and exaggerate this difference in other tasks. We then simulate an expanded version of the Connectionist Incremental Articulation Model that incorporates Stevens' (1998) 3 phases involved in articulating a speech segment.

Keywords: Phonological Priming; Word Naming; Speech Production; Segment Duration; Jordan Network; Minimal Unit of Articulation; Response Criterion

The Precise Articulatory Sequence in the Production of Individual Sounds

Researchers have known that the onset of the articulatory response occurs long before the onset of the acoustic response (Bell-Berti & Harris, 1981). This asynchrony between the onset of articulatory and acoustic responses is due to the fact that individual sounds are produced in three distinct articulatory phases: (1) the movements of the articulators toward the formation of the oral constriction; (2) the flow of air behind or around the oral constriction; and (3) the release of the current constriction and movement toward the next constriction (Stevens, 1998).

The exact point at which the acoustic event is produced in the sequence above for initial segments that differed in manner can be distinguished according to their specific articulatory requirements. For non-plosive segments such as /m/, /l/, /r/, /s/, etc., acoustic energy can be generated shortly after their respective oral constriction are formed. However, for plosive segments such as /p/, /t/, /k/, etc., acoustic energy can be generated only after (a) the buildup of sufficient intra-oral pressure, and (b) the release of the current oral constriction. That is, acoustic onset for non-plosive segments occurs during the second phase, whereas for plosive segments, it occurs during the third.

The Plosivity Effect

This differential requirement for the production of the acoustic events of plosive and non-plosive segments is the basis for the Plosivity Effect. Because the acoustic event for plosive segments requires the buildup of the intra-oral pressure, the onset of acoustic energy (acoustic latency)

for responses beginning with plosive segments is typically 50 – 100 ms slower than responses beginning with non-plosive segments, although this difference can be as short as 20 ms in speeded naming tasks (see Kawamoto, Kello, Jones & Bame, 1998). Notably, the plosivity effect can be completely eliminated in some tasks (Kawamoto, Liu, Mura, & Sanchez, 2008), or enhanced in others (Liu, Kawamoto, & Grebe, 2009).

The Elimination of the Plosivity Effect

In a study examining the temporal relationship between the onsets of the articulatory and acoustic responses in the delayed naming task, participants were presented with the complete stimulus (a monosyllabic word) at the beginning of a trial and were told to respond as quickly and accurately as possible only after a signal to respond was given, which was given after a delay. Using stimuli that began with the segments /p/, /t/, /m/, and /n/, Kawamoto and colleagues found that what constituted the initiation of the response to participants was the onset of the acoustic response, despite the fact that for long delays, the onset of the articulatory response occurred before the signal to respond. In essence, participants were moving their articulators into the optimal position for the phonation of the acoustic response while they were waiting for the signal to respond. For sufficiently long delays, that meant holding the acoustic response in abeyance until the signal to respond was detected. For responses beginning with non-plosive segments (/m/ and /n/), that meant holding the response at the cusp of the second articulatory phase, whereas for responses beginning with plosive segments (/p/ and /t/), it was held at the cusp of the third articulatory phase.

Basically, when participants were afforded the opportunity (i.e., long delays) to not only form the appropriate oral constriction (phase 1), but also build the required intra-oral pressure (phase 2) for plosive initial segments, the plosivity effect disappeared completely. But when there was insufficient time for the first two phases to be completed for plosive initial segments (i.e., short delays), a plosivity effect could still be observed, although the magnitude of the plosivity effect diminishes as a function of delay.

The Enhancement of the Plosivity Effect

In a different study, which examined the minimal unit of phonological information needed to initiate articulation (i.e., minimal unit of articulation), participants again produced monosyllabic utterances that began with the

segments /p/, /t/, /m/, and /n/ under a variant of the delayed naming task, called the pair-wise priming task (Liu, Kawamoto, & Grebe, 2009).

Procedurally, the pair-wise priming task is identical to the delayed naming task except that instead of presenting the complete stimulus at the beginning of the trial, only the initial letter was presented (e.g., *m*__). The complete stimulus (e.g., *mood*) was presented only after a variable delay (i.e., stimulus onset asynchrony or SOA) of 300 ms or 600 ms.

The key issue was whether or not the initial segment was sufficient for a participant to initiate the articulatory response or even the acoustic response. Liu and colleagues (2009) found that participants could in fact initiate the articulatory response on the basis of the initial segment alone. In fact, in the 600 ms SOA condition, articulatory onset for both plosive and non-plosive segments on average occurred before the presentation of the complete stimulus.

The next issue is whether knowledge beyond the initial segment is required to initiate the acoustic response. If the acoustic response can be initiated on the basis on the initial segment alone in certain conditions, one can lengthen the acoustic event for these segments until the subsequent segment becomes available. Indeed, for nasals, a 10.08 ms increase in the mean acoustic segment duration was observed in the 600 ms SOA condition compared with the 300 ms SOA condition (see Figure 1). In fact, in certain trials, acoustic onset occurred prior to the presentation of the complete stimulus. Although the total number of trials where this was observed was limited to 6 out of 474 trials in total, they represent a clear existence proof that for non-plosive initial segments, the acoustic response can be initiated based on knowledge of the initial segment alone.

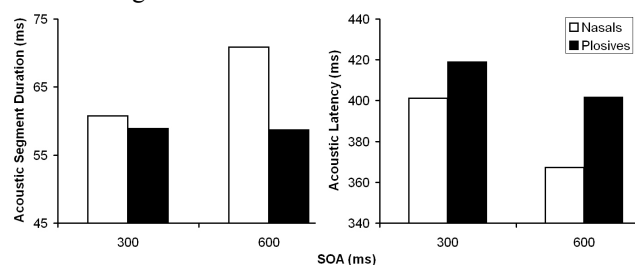


Figure 1: The acoustic segment durations (left) and acoustic latencies (right) of plosive and nasal initial segments across priming conditions reported by Liu et al., (2009).

However, the results for plosive initial segments differed from those for nasals. For plosives, acoustic onset corresponds to the explosive release of pent up pressure that require knowledge of the following segment (i.e., phase 3), and thus acoustic onset always occurs some time after the complete target is presented. Moreover, because this release of pressure occurs more or less in an all or none fashion, plosive segments are relatively resistant to acoustic lengthening. Indeed, acoustic duration for plosive

segments were almost identical across the different priming conditions (58.91 ms and 58.72 ms for the 300 ms and 600ms SOA, resp.). So, despite the fact that initiation of the articulatory response for plosive initial segments does not require knowledge of the subsequent segment, the initiation of the acoustic response is contingent on the following segment.

Given that the initiation of the acoustic response for plosive initial segments must wait until the next segment is known but it does not for non-plosive segments, the onset of the acoustic response for non-plosive initial segments can be initiated much earlier than plosive initial segments, particularly in the 600 ms SOA condition. This differential constraint is what was driving the 16.64 ms enhancement of the plosivity effect on acoustic latency observed across the different priming conditions (i.e., the difference between the 17.77 ms plosivity effect in the 300 ms SOA condition and the 34.42 ms plosivity effect in the 600 ms SOA condition as illustrated in Figure 1).

Individual Differences: A Case for multiple Response Criteria. Although the results of Liu and colleagues' (2009) study represent the strongest evidence to date that participants can initiate both the articulatory and acoustic responses before the full phonological code of a word is generated, not all participants behaved in this manner. In fact, there is a wide range of individual differences.

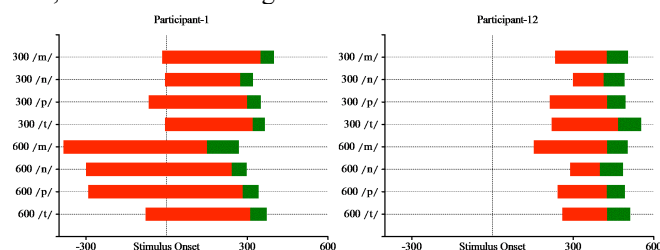


Figure 2: Data for two participants (1 and 12) reported by Liu et al., (2009). The articulatory onset to acoustic onset interval (AAI) is shown in red, and the acoustic segment duration is shown in green.

At one extreme (e.g., participant 1 in Figure 2), there are participants who initiated the articulatory response shortly after the prime is processed and long before the complete stimulus is presented. In a sense, data from these participants typified a response strategy where articulation is based on a segmental criterion (cf. Kawamoto et al., 1998). For these participants, articulation of the initial segment was necessarily lengthened because articulation of the initial segment was initiated before the subsequent segment was available. Since information for the next segment comes much later in the 600 ms SOA condition, these participants showed the largest AAI and acoustic segment duration effects (for nasals). When the difference in the SOA is taken into account, there is virtually no difference in when the articulatory response is initiated (articulatory latency)

Table 1: Examples of the input and output representations used in the current model. Input for the Articulatory Control structure (in bold), with the “-” denoting the Unspecified Segment unit (used only for priming), and the “\$” denoting the Metrical Slot unit (specified or not). The first and last sweep represents the neutral state.

The Complete Input Plan for <i>mood</i>																																	
Sweep	Onset1										Onset2						Vowel						Coda										
	s	p	m	n	t	f	l	r	-	\$	p	t	r	l	-	\$	a	e	i	u	ɪ	U	-	\$	d	p	t	f	l	r	-	\$	
1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
2-8	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	1	0	0	0	0	0	0	0
9	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0

Corresponding Target Output Values																				
Sweep	Articulatory Phase of Current Segment										Velum	Tongue Tip	Tongue Body	Lip Vertical	Lip Horizontal	Pressure	Glottis	Acoustic Energy		
1	0	0	0	0	0	0	0	0	0	0	0	0	0.1	0.5	0.2	0.8	0.8	0	0.5	0
2	1	0	0	0	0	0	0	0	0	0	0	0	0.1	0.45	0.2	0.1	0.7	0	0.4	0
3	0	1	0	0	0	0	0	0	0	0	0	0	0.9	0.4	0.15	0.1	0.5	0.4	0.2	0.5
4	0	0	1	0	0	0	1	0	0	0	0	0	0.5	0.3	0.12	0.15	0.45	0.4	0.2	0.6
5	0	0	0	0	0	0	0	1	0	0	0	0	0.1	0.1	0.1	0.4	0.3	0.45	0.1	0.7
6	0	0	0	0	0	0	0	0	1	1	0	0	0.1	0.5	0.12	0.5	0.4	0.55	0.4	0.5
7	0	0	0	0	0	0	0	0	0	0	1	0	0.1	0.9	0.15	0.7	0.7	0.8	0.5	0.1
8	0	0	0	0	0	0	0	0	0	0	0	1	0.1	0.5	0.2	0.8	0.8	0.2	0.5	0.8
9	0	0	0	0	0	0	0	0	0	0	0	0	0.1	0.5	0.2	0.8	0.8	0	0.5	0

relative to the presentation of the prime.

At the same time, there are also participants who, despite the fact that the prime provided sufficient information to initiate articulation, chose to wait until more information becomes available. These participants (e.g., Participant 12 in Figure 3) typified a response strategy where articulation is based on the whole word criterion (cf. Kawamoto et al., 1998). For these participants there was little or no difference in when and how the response was executed across priming conditions.

The Incremental Articulation Account: An Expanded Implementation

In this section, we describe an expanded implementation of the Incremental Articulation Account (Kawamoto et al., 1998). The core assumptions of the incremental account are that (1) the minimal unit of articulation is the segment, (2) multiple response criteria can be used, and (3) when articulation is initiated on a segment by segment basis, the articulation of the current segment/s may be lengthened to accommodate the time needed to process the subsequent segments.

The goal of the initial implementation of this model was to account for anticipatory coarticulation and segment duration effects (Kawamoto and Liu, 2007). The goal of the current implementation is to expand the generality of this model by demonstrating how the elimination and enhancement of the plosivity effect can be accounted for when the precise sequence of articulatory events involved in articulation is considered.

The Representations Used

The input representation for the current implementation, as in Kawamoto and Liu (2007), is a slot based local representation scheme that specifies the segmental content, syllabic frame, and encoding status of a metrical slot (see Table 1). The output and state representations in the current implementation correspond

to the current articulatory phase, the syllabic position of the segment being produced, and a small set of articulatory dimensions: the velar opening, the positions of the tongue tip and tongue body, the vertical and horizontal lip separations, the degree of intra-oral pressure, the constriction of the glottis, and acoustic intensity. Phase 1 of a word medial segment overlaps in the current output representation with phase 3 of the previous segment to convey the fact that these two phases are one and the same (see Table 1).

Model Architecture

The current model consists of two linked networks — a phonological network that provides the input to an articulatory network (see Figure 3).

The Phonological Network

The phonological network consists of three distinct components: a Phonological Buffer, a Response Criterion layer, and a Buffer Control structure. The Phonological Buffer functions simply as a temporary storage mechanism for the phonological code generated by the preceding processes such as phonological encoding (not modeled here). The Response Criterion layer denotes the particular set of response criteria used. In the current implementation there are three sets of units: (1) the Instruction units, (2) the Lengthening Criterion units, and (3) the Articulation Criterion units.

The Instruction units were designed to mimic the effect of instructions on the precise articulatory juncture that participants decide to hold the response in abeyance (i.e., {1, 0, 0} denotes articulatory onset, {0, 1, 0} denotes the acoustic onset, and {0, 0, 1} denotes vocalic onset). Their specific function in this implementation is to simulate the results of the delayed naming task. The Lengthening Criterion units dictate the manner in which a verbal response will be lengthened (i.e., {1, 0} for articulatory lengthening, and {0, 1} for acoustic lengthening).

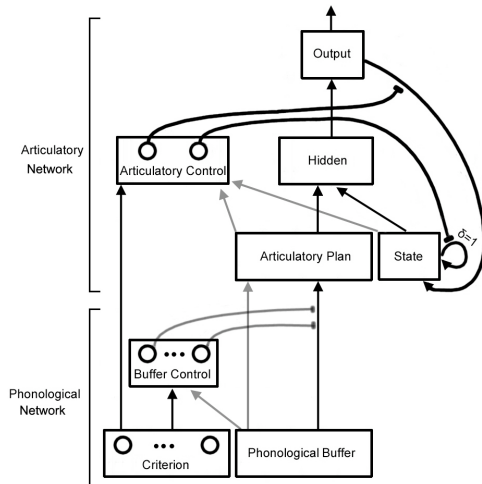


Figure 3: The basic architecture of the current model. The black arrows denote full connectivity with the exception that the connections between Output-to-State and State-to-State are one-to-one. The gray arrows denote that only some inputs are fed into the next layer (see discussion below).

Instruction and Lengthening Criterion units are fed directly into the Articulatory Control structure (see below), while the Articulatory Criterion units are fed into the Buffer Control structure. The Buffer Control structure, which also takes input from the unspecified segment and Metrical Slot units in the Phonological Buffer (i.e., “-“ and “\$” units), monitors the content of the Phonological Buffer and gates the flow of information to the articulatory network according to the specific articulation criterion used.

Phonological Gating. The gating of information from the Phonological Buffer to the articulatory network is carried out by the Buffer Control units through Sigma-Pi connections from the Phonological Buffer to the Articulatory Plan (Buffer-to-Plan connections). All of the Buffer-to-Plan connections will be turned “off” initially for the whole word criterion (i.e., $\{0, 1\}$). They are turned “on” only when the segments in each of the specified metrical slots are known (i.e., when all of the Unspecified Segment units, “-“, are 0 or off). If the segment criterion is used (i.e., $\{1, 0\}$), all of the Buffer-to-Plan connections will be turned “on” from the very beginning. Thus, the flow of information from buffer to plan will always be unrestricted. In this respect, under the segment criterion, an incomplete syllabic code can be fed into the articulator network and incrementally updated to reflect ongoing processing, whereas under the whole word criterion, only the complete syllabic code is fed to the articulatory network.

The Articulatory Network

The articulatory network consists of two primary components, an Articulatory Control structure that is coupled to a Jordan net (Jordan, 1986).

Jordan Net. The basic architecture of the Jordan net in the current implementation remains relatively unchanged from that in Kawamoto and Liu (2007). As in Kawamoto and Liu (2007), the connections between Output-to-State and State-to-State are one-to-one. The decay parameter is set to equal 1 (i.e., $\delta=1$; no decay). Due to the actions of the Sigma-Pi connections from the control structure (see below), only the output of a single sweep will be buffered in the State units at any given time. Accordingly, the values of the State and Output Units are always offset by 1 sweep.

Articulatory Control Structure. The control structure is a simple feed forward network that acts as a monitoring and gating mechanism that takes input from the Instruction units, Articulation Criterion units, Lengthening Criterion units, and partial input from the Plan and State units.

The control structure checks the articulatory phase of the segment currently being produced with the availability of the following segment and gates the progress of the Jordan net according to the instruction or lengthening criterion specified by the criterion units. If information about the following segment is available, the output of the Jordan net will naturally move from producing the articulatory sequence of the current segment to those of the next. If not, the model will simply prolong the production of the articulatory phase/s according to the lengthening criterion specified. This control is accomplished by taking as inputs the Unspecified Segment units and the Metrical Slot units (i.e., the “-“ and “\$” units) from the Plan units and the Current Articulatory Phase units from the State units. The Unspecified Segment units denote whether or not the identity of a particular segment is unknown (1=“on” or unknown), the Metrical Slot units denote whether or not a particular metrical slot is specified in the word frame (1=“on” or specified), and the Current Articulatory Phase units denote the syllabic position and articulatory phase of the segment produced in the previous sweep. Together, these units feed into 2 control units, one for the Sigma-Pi connections to the Output-to-State connections and the other for the Sigma-Pi connections to the State-to-State connections, and turn them “on” or “off” accordingly. The specific operation of the Sigma-Pi connections under different situations is discussed in the simulation below.

Training Network

The Jordan net was trained on the Tlearn simulator (Plunkett & Elman, 1997). The training corpus for the Jordan net consisted of the input and output sequences corresponding to a small lexicon consisting of words such as *mood*, *neat*, *pit*, *pale*, etc. Training was carried out for 5000 epochs, with a learning rate = 0.05 and momentum = 0. Because the behavior of the Articulatory Control structure is just that of a simple table lookup, its operation was hardwired in the current simulation. The specific conditions that trigger its operation are discussed below.

Testing the Network

Simulation of certain key results of the delayed naming experiment reported by Kawamoto et al., (2008) and the pair-wise priming task reported by Liu et al., (2009) were carried out. For the delayed naming task, our interest focused on simulating the elimination of the plosivity effect at long delays. For the pair-wise priming task, the result of interest to us was the enhancement of the plosivity effect driven by participants using the segment response criterion. To demonstrate the malleability of the plosivity effect as well as the differential lengthening of the articulatory and acoustic responses, the outputs of the Articulatory Network were computed offline from the weight matrices of the Jordan net and the Articulatory Control Structure.

Delayed Naming. The test sequences for delayed naming were simply the training sequences lengthened to 17 sweeps for test sequences beginning with plosive segments and 18 for non-plosive sequences. For both test sets, the neutral state was presented for the first 4 sweeps to represent the time that it takes for the phonological code to be generated. The complete syllabic code is presented in the 5th sweep and remained so until the 14th sweep for plosive test sequences and the 15th sweep for non-plosive sequences, after which both sets of test sequences returned to neutral state on the final sweep. Input from the Instruction units to the Articulatory Control units were set at {0, 1, 0} (indicating that the acoustic response will be held in abeyance) from sweeps 2-10 to indicate the time for the signal to respond to be detected, after which the input was switched to {0, 0, 0} for the remainder of each test sequence set.

The outputs of these test sequences clearly show that for sequences beginning with plosive and non-plosive segments, the articulatory phase prior to the generation of

acoustic energy is lengthened until the signal to respond is detected (i.e., sweep 10). Specifically, the Sigma-Pi connections to the Output-to-State connections were turned “on” and the connections to the State-to-State connections were turned “off” on the 5th sweep for non-plosive sequences and on the 6th sweep for plosive sequences. These actions turn the Jordan net into a feed forward network that simply updates the output of the earlier sweep. In essence, the articulation of the first articulatory phase for non-plosive sequences was lengthened until sweep 10, whereas for plosive sequences it was the second articulatory phase. When the Sigma-Pi units switch back to their default state on sweep 10, the articulatory network turned back into the Jordan net and produced the next articulation phase (phase 2 for non-plosives and phase 3 for plosives) in sweep 10 and for the remainder of the response in subsequent sweeps.

Since acoustic energy is generated in Phase 2 for non-plosive initial segments, and phase 3 for plosive initial segments, the output of the test sequences demonstrate that, due to the input from the Instruction units, acoustic energy for these test sequences are generated in synchronicity in sweep 8 — an elimination of the plosivity effect.

Pair-wise Priming. Two different sets of test sequences were used for the pair-wise priming task to simulate the behavior of participants using different response criteria: (1) the whole word, and (2) the segment criteria. The sequence lengths for these input sets were 17 and 15, respectively. For both test sets, the first 4 sweeps represent the neutral state. On sweeps 5-9, a fragmentary syllabic code consisting of information for the initial segment only (e.g., *m*___ or *p*___) with the unspecified segment represented by the “-” unit in the appropriate metrical slots. From sweeps 10 to the penultimate sweep,

Table 2: Example of the test sequences *m*___ → *mood* for the pair-wise priming task. The values of the Current Articulatory Phase Units from the Output Layer, the input from the Lengthening Criterion units, and the actions of the Articulatory Control units are to show the sequence of events under the **Segment Criterion**.

The Pair-wise priming input example for <i>m</i> → <i>mood</i>																																
Onset1										Onset2						Vowel						Coda										
Sweep	s	p	m	n	t	f	l	r	-	\$	p	t	r	l	-	\$	a	e	i	u	I	U	-	\$	d	p	t	f	l	r	-	\$
1-4	0	0	0	0	0	0	0	0	0	0 0	0	0	0	0	0 0		0	0	0	0	0	0	0 0		0	0	0	0	0	0	0 0	
5-9	0	0	1	0	0	0	0	0	0 1		0	0	0	0	0 0		0	0	0	0	0	0	1 0		0	0	0	0	0	0	1 0	
10-14	0	0	1	0	0	0	0	0	0 1		0	0	0	0	0 0		0	0	0	1 0	0	0 1		1 0	0	0	0	0	0	0 1		
15	0	0	0	0	0	0	0	0	0 0		0	0	0	0	0 0		0	0	0	0	0	0	0 0		0	0	0	0	0	0	0 0	
Articulatory Control Units										Corresponding Values of the Current Articulatory Phase Units (Within the Output Layer)																						
Sweep	Lengthening Criterion		State-to-State	Output-to-State	Onset1				Onset2				Vowel				Coda															
1	0	0	0	1	.02	.00	.01	.01	.00	.00	.02	.02	.02	.01	.01	.01																
2-4	1	0	0	1	.01	.01	.02	.01	.02	.02	.02	.02	.01	.02	.01	.00	.02															
5-7	1	0	1	0	.98	.01	.01	.00	.00	.02	.00	.01	.00	.02	.00	.01																
8	0	1	0	1	.01	1.00	.01	.01	.01	.02	.01	.01	.01	.02	.02	.00																
9	0	1	1	0	.01	1.00	.01	.01	.01	.02	.01	.01	.01	.02	.02	.00																
10	0	1	0	1	.00	.01	.98	.01	.02	.01	1.00	.02	.00	.01	.01	.00																
11	0	1	0	1	.01	.02	.01	.01	.00	.01	.00	1.00	.01	.02	.02	.00																
12	0	1	0	1	.02	.00	.01	.00	.01	.00	.01	.02	.99	.99	.00	.01																
13	0	1	0	1	.01	.01	.01	.01	.02	.01	.01	.01	.02	.01	1.00	.02																
14	0	1	0	1	.00	.00	.01	.02	.00	.00	.00	.00	.01	.01	.01	.99																
15	0	0	0	1	.00	.00	.00	.01	.01	.01	.02	.01	.02	.00	.00	.01																

the test sequences contained the complete syllabic code (e.g., *mood* or *pit*); and neutral state on the final sweep (see Table 2).

Inputs to the criterion units, specifically, the Articulation Criterion and Lengthening Criterion units, for the whole word criterion test set were set to {0, 1} and {0, 0}, respectively. For the segment criterion test set, the input was set to {1, 0} and {1, 0} on sweeps 1-7, and switched to {1, 0} and {0, 1} for the remaining sweeps.

The output of the whole word test sets show that the first articulatory phase of the initial segment for both plosive and non-plosive test sequences was not produced until sweep 10, and subsequent articulatory phases were produced in each successive sweep until the entire response was completed on sweep 16. This is because the input from the Articulation Criterion units (i.e., {0, 1}) to the Buffer Control units coupled with input from the Unspecified Segment units (i.e., “-” units) in the Phonological Buffer turned “on” the Sigma-Pi connection from the Buffer Control units to the Buffer-to-Plan connections, which effectively shut off input from the Phonological Buffer units to the Plan units until the full phonological code became available on sweep 10. Once the Sigma-Pi connections to the Buffer-to-Plan connections were turned “off”, the entire syllabic code was then fed into the articulatory network, and the syllabic response is produced uninterrupted. Because the acoustic event for the initial segment was produced on the 11th sweep for non-plosive sequences, and on the 12th sweep for plosive sequences, a plosivity effect of 1 sweep was observed.

For the segment criterion test sets, the first articulatory phase was produced on sweeps 5-7 for both the plosive and non-plosive sequences because the initial segment only became available on sweep 5. On sweep 8, the input from the Lengthening Criterion units was switched to {0, 1}, at which point the action of the Articulatory Control units reverts the Sigma-Pi connections to their default state, thus, allowing the next articulatory (phase 2) to be produced. However, on sweep 9, based on the input from the Current Articulatory Phase units within the State units (values offset from those within the Output Layer, shown in Table 2, by 1 sweep), the Articulatory Control units again turned “on” the Sigma-Pi connections to the Output-to-State connections and turned “off” the Sigma-Pi connections to the State-to-State connections. Thus, phase 2 was repeated on sweep 9 (see Table 2). When the entire syllabic code became available on sweep 10, the Sigma-Pi connections again revert back to its default state allowing the remainder of the syllabic response to be produced in subsequent sweeps. Accordingly, the acoustic event for non-plosive sequences (phase 2) was produced on sweep 8, whereas, for plosive sequences (phase 3), it was produced on sweep 10, resulting in a plosivity effect of 2 sweeps.

Although the output of these two test sets taken together produced a mean plosivity effect of 1.50 sweeps

— a net plosivity enhancement of 0.50 sweeps, the exact magnitude of the enhancement ultimately depends on the proportion of participants using the segment criterion and the whole word criterion.

Conclusion

The results of the current simulations demonstrate that a single network using a sub-syllabic minimal unit and different response criteria can account for both the elimination and enhancement of the plosivity effect, as well as the differential lengthening of the articulatory and acoustic responses when the precise sequence of articulatory events involved in the generation of the individual sounds is taken into consideration. This approach can easily be extended to other latency and duration effects, both articulatory and acoustic, that can arise from a variety of processing difficulties in speech production. Moreover, with slight modifications, the current network can be easily coupled to existing models, such as the one described by Dell, Juliano, and Govindjee (1993) to account for a wider range of empirical data (e.g., latency data). Such an extension provides a way to explore the intricate coordination between different processing stages, and how potential asynchronies in the flow of information may reveal themselves in the interplay between different dependent measures such as articulatory latency, acoustic latency, and the duration of various components of a verbal response.

Acknowledgments

This work was partially supported by a grant from the UCSC faculty senate.

References

- Bell-Berti, F., & Harris, K. S. (1981). A temporal model of speech production. *Phonetica*, 38, 9-20.
- Dell, G. S., Juliano, C., & Govindjee, A. (1993). Structure and content in language production: A theory of frame constraints in phonological speech errors. *Cognitive Sciences*, 17, 149-195.
- Jordan, M. (1986). Serial order: A parallel distributed processing approach. Technical Report 8604. San Diego: Institute for Cognitive Science. University of California.
- Kawamoto, A. H., Kello, C. T., Jones, R. M., & Bame, K. (1998). Initial phoneme versus whole word criterion to initiate pronunciation: Evidence based on response latency and initial phoneme duration. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 24, 862-885.
- Kawamoto, A. H., & Liu, Q. (2007). A unified account of segment duration and coarticulation effects in speech production. In D. S. McNamara & J. G. Trafton (Eds.) *Proceedings of the 29th Annual Meeting of the Cognitive Science Society* (pp. 1151-1156). Austin, TX: Cognitive Science Society.
- Kawamoto, A. H., Liu, Q., Mura, K., & Sanchez, A. (2008). Articulatory preparation in the delayed naming task. *Journal of Memory and Language*, 58, 347-365.
- Liu, Q., Kawamoto, A. H., & Grebe, P. R. (2008, May). Incremental articulation: Evidence from onset priming. Poster presented at the 21st Annual Conference of the Association for the Psychological Science. San Francisco, CA.
- Plunkett, K., & Elman, J. L. (1997) *Exercises in Rethinking Innateness: A Handbook for Connectionist Simulations*. Cambridge, MA: MIT Press
- Stevens, K. N. (1998). *Acoustic Phonetics*. Cambridge: MIT.