

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

A conflict-based model of speech error repairs in humans

Permalink

<https://escholarship.org/uc/item/2qg342vp>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 44(44)

Authors

Burgess, Jack L

Nozari, Nazbanou

Publication Date

2022

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

A conflict-based model of speech error repairs in humans

Jack Burgess (jackburg@andrew.cmu.edu)

Neuroscience Institute, 4400 Fifth Avenue
Pittsburgh, PA 15213 USA

Nazbanou Nozari (bnozari@andrew.cmu.edu)

Department of Psychology, 5000 Forbes Avenue
Pittsburgh, PA 15213 USA

Abstract

Fast and efficient correction of speech errors is essential to effective communication. Yet, despite several accounts of error *detection*, no computational account exists to explain how humans *repair* their speech errors. This paper proposes the first such model. We demonstrate that a simple automatic mechanism can form the basis of most repairs. We then demonstrate that augmenting the model with a conflict-based monitoring-control loop allows it to capture more nuanced findings in human speech error repair data.

Keywords: speech error; repair; computational model; language production

Introduction

The ability to detect and repair errors in one's own speech is a key aspect of successful communication. Several computational models exist to explain the *detection* process (Hartsuiker & Kolk, 2001; Hickok, 2012; Nozari et al., 2011; Tourville & Guenther, 2011; see Nozari, 2020, for a review), but, to date, no computational model has been implemented to explain the *repair* process. Empirical data on speech error repairs are relatively sparse, but they point to important characteristics that must be considered in a model of repair: (a) repairs can be performed extremely quickly. The gap between stopping an error and initiating a repair (e.g., “v-horizontal”; Levelt, 1983) can be as short as zero ms. (b) Repairs are not always accompanied by full consciousness over the error, an explicit intention to repair, or a clear knowledge of the target. This is most obvious in young children and individuals with aphasia (Clark, 1978; Nozari et al., 2011). These characteristics are incompatible with proposals that view the repair process as a deliberate, stop and restart-from-scratch process, and instead point to a fast and largely automatic process that quickly replaces the error with an available repair without a full restart (Nooteboom & Quené, 2019; Nozari et al., 2019).

This paper proposes the first computational account of such a process. The main idea is that the system leverages the co-activation of multiple responses, a natural property of the production system, to quickly replace an error with an available repair, by monitoring the activation dynamics for a brief period after response selection. *Simulation I* extends the two-step model of word production to a time-based model that has the right temporal properties for modeling such a process. *Simulation II* introduces the basic model of repair and tests its fundamental assumptions. *Simulation III* shows

the limitation of the basic model in capturing the adaptive nature of repairs. Finally, *Simulation IV* proposes the conflict-based repair model augmented with a monitoring-control loop.

Model description

The model is based on the two-step interactive activation model (Foygel & Dell, 2000). It is a neural network with three layers of representation (semantic features, lexical items, and phonemes), and it maps semantics onto phonology to produce a word (e.g., /kæt/; Fig. 1). Naming in the model has two steps. In the first step (semantic-lexical mapping), an input vector activates the semantic features of a concept (10 units per concept, each receiving 10 units of activation). Activation spreads in a bidirectional way in the network, and each node's activation is updated in parallel according to Eq. 1

$$A(j, t + 1) = (1 - d)A(j, t) + \sum_i w_{ij} \cdot A(i, t) + A(j, t)N(0, a) + N(0, i) \quad \text{Eq. 1}$$

where $A(j, t + 1)$, the activation of node j at time $t + 1$, is a sum of the activation of node j at time t after decaying at rate d , the input that node j receives from each node i multiplied by the respective connection weight w_{ij} , and two sources of noise (intrinsic noise and activation-based noise, drawn from normal distributions with means 0 and standard deviations i and a , respectively). After eight time steps, the most highly activated node in the lexical layer is selected. The second step (lexical-phonological mapping) starts by giving a jolt of 100 units to the selected lexical node. Activation spreads for eight more time steps in the manner described above. The process concludes by selecting the most highly activated nodes in the phoneme layer for each syllabic position (i.e., onset, vowel, and coda, for a CVC like “/kæt/”). Two key parameters are s and p , the strengths of the connections between semantic features and lexical items, and lexical items and phonemes, respectively. The model has been highly successful in capturing the various patterns of speech errors in neurotypical adults, children, and individuals with brain damage (Budd et al., 2011; Dell et al., 1997, 2013; Nozari et al., 2010).

The conflict-based model of error detection. An extension of the two-step model, the *conflict-based model of error detection*, has successfully shown that the information available during production can be used as a signal for

detecting speech errors (Nozari et al., 2011). This model is based on a key component of the two-step model: throughout the process, the spread of activation activates not only the target, but also related representations through their shared connections (e.g., “dog” and “cat”; Fig. 1). The conflict-based model posits that higher levels of conflict are associated with higher probabilities of an error, a signal that the monitoring system uses to detect errors, sometimes before they become overt. Predictions of the conflict-based account have found support in children, individuals with brain damage, and L2 speakers (Hanley et al., 2016; Nozari et al., 2011, 2019). The current paper will extend the conflict-based model to explain not only error *detection* but also the mechanism underlying error *repairs*.

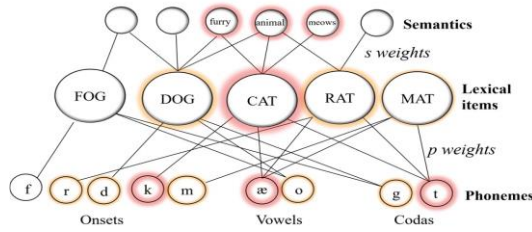


Figure 1. Schematic of the two-step interactive model of word production. “cat” is the target of the current trial.

Simulation I (a, b): the time-based model

The goal of Simulation I is to create a time-based version of the two-step model (Ia), and to determine that conflict generated during its processing can reliably distinguish between correct and error trials at the lexical layer (Ib).

Simulation Ia. Despite its success in explaining error patterns, the two-step model does not address the temporal dynamics of the production process. This is clear in Figure 2a. Activation starts at very high levels and quickly drops over time. In reality, neural activation starts at low levels and gradually builds up as input processing converges on a given representation (a process modeled as evidence accumulation in drift diffusion or similar models). Since the repair process unfolds over time, it is critical that the underlying model is more realistic in capturing the temporal dynamics of word production. This can be achieved by a re-tuning of the model parameters. In our simulations, we also clamped the activation of the input, as the data simulated under Simulations III and IV came from tasks in which visual stimuli for naming are visible to the participants during the entire production time. Table 1 demonstrates the parameters in the original and time-based versions of the model. Equation 2 shows the new activation rule,

$$A(j, t+1) = (1-d)A(j, t) + \sum_i w_{ij} \cdot A(i, t) + \sum_i w_{ij} \cdot fX_i + A(j, t)N(0, a) + N(0, i) \quad \text{Eq. 2}$$

where clamping is implemented as $\sum_i w_{ij} \cdot fX_i$, where X is the binary input feature vector ($X_i = 1$ signifies feature i is

present) and f is the input feature strength. Twenty batches of 10,000 trials were simulated using the original and time-based models.

Results of simulation Ia. As shown in Figure 2b, activation starts at zero and grows over time in the time-based model, with the activation of the target gaining over competitors with more time. The mean values and SDs for correct, semantic, and other errors were $96.93 \pm 0.13\%$, $2.76\% \pm 0.15$, and $0.31\% \pm 0.03\%$ for the original model, and $97.81\% \pm 0.16\%$, $2.04\% \pm 0.14\%$, and $0.15\% \pm 0.06\%$ for the time-based model, showing largely comparable response profiles across the two models, despite the changed temporal dynamics. Since this project targets lexical repairs and our proposed mechanism is focused on the lexical layer, all models past this point will only include the first step of mapping (semantic-lexical). “correct” and “error” will correspondingly refer to the target and semantic errors in the lexical layer (almost all of the errors at this stage in a normal production system).

Table 1: Original and time-based model parameters. See text for parameter descriptions.

Parameter	Original	Time-based
s weight	0.04	0.0005
p weight	0.04	0.0005
d	0.6	0.05
f	10	0.1
i	0.01	0.003
a	0.16	0.16
t	8	20

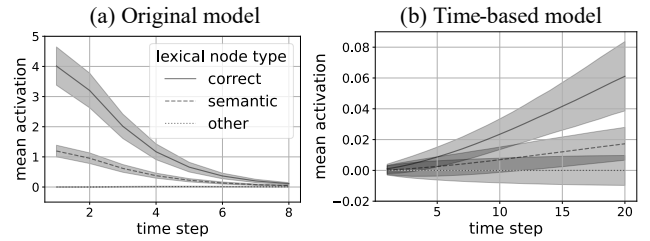


Figure 2. Trajectories of mean lexical activation ($\pm$ SD) over time for correct, semantic, and other error types in the original (a) and time-based model (b).

Simulation Ib. Next, we tested whether conflict in the time-based model could reasonably distinguish between error and correct trials. Two measures of conflict were tested. (1) *Point conflict* compares the activations of the two most activated nodes at the end of the semantic-lexical mapping process, i.e., the lexical selection point (Eq. 3).

$$\text{Point conflict} = \ln \left(\frac{1}{|A(\text{cat}, T) - A(\text{dog}, T)|} \right) \quad \text{Eq. 3}$$

(2) The second conflict measure, *Integral conflict*, measures the activation difference between the target and semantic competitor nodes throughout the entire process (Eq. 4).

$$\text{Integral conflict} = \ln \left(\frac{1}{\int |A(\text{cat}, t) - A(\text{dog}, t)| dt} \right) \quad \text{Eq. 4}$$

These two measures were calculated for each of 50,000 trials simulated with the time-based model with parameters in Table 1. For each measure, values were divided into two distributions (correct, error) based on the lexical response at time t_{20} (Fig. 3). The sensitivity of the two measures of conflict in distinguishing between error and correct trials was compared using Cohen's d , calculated as the difference between average conflict in error (m_e) and correct trials (m_c) scaled by the pooled standard deviation (s_p) of the two distributions (Eq. 5).

$$d = \frac{m_e - m_c}{s_p}, \quad s_p = \sqrt{\frac{(n_c - 1)s_c^2 + (n_e - 1)s_e^2}{n_c + n_e}} \quad \text{Eq. 5}$$

where n_c and n_e respectively denote the number of error and correct trials and s_c and s_e respectively denote the standard deviation of conflict in correct and error trials.

Results of simulation Ib. Figure 3 shows the results of Simulation Ib. Cohen's d for *Point conflict* (Fig. 3a) and *Integral conflict* (Fig. 3b) was 2.6 and 1.5, respectively, suggesting the former to be a more sensitive measure for distinguishing errors from correct responses. For this reason, we used the point conflict measure for gauging conflict in the later sections of this study.

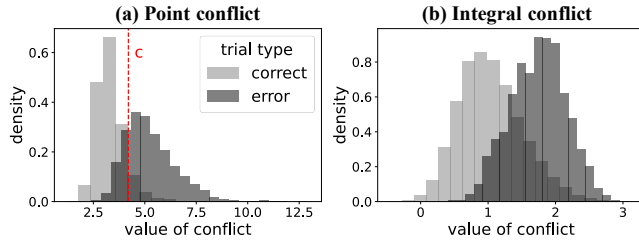


Figure 3: Density histograms of conflict measures for correct vs. error trials in the time-based model (scaled so area under each curve is 1). Point conflict (a) and Integral conflict (b). The vertical line shows a possible placement of the criterion or c , as a cutoff point above which a trial is determined to be an error.

Discussion. Simulation I showed that the time-based version of the two-step model preserves its original qualities, including conflict as a strong signal discriminating between correct and error trials, while better capturing activation dynamics. It thus provides a suitable basic model for implementing the repair process in Simulation II.

Simulation II: the basic model of repair

Simulation II implements the basic conflict-based model of repairs. The proposed mechanism is a “*respond and check*” process: at a given time step (t_{20}) the model produces a response by selecting the most highly activated node (this is similar to a deadline for response generation). Processing, however, is allowed to continue for a little while longer (until

t_{25}), at which time the model “checks” the response by re-examining which node has the highest activation. If it is the same as the produced response, no action is taken. If, on the other hand, a different node has a higher activation at this point, the model automatically replaces the old response with the new response, i.e., it generates a repair.

The rationale for the proposed mechanism is based both on the action monitoring literature and the dynamics of the time-based model. The former points to evidence for continued processing after response generation as one of the key monitoring mechanisms for action regulation (e.g., Yeung et al., 2004). The latter is based upon the differential dynamics of activation in correct and error trials. This is illustrated in Figure 4, as two examples. In the majority of the *correct trials* (e.g., Fig. 4a), the correct representation gains a clear advantage over the error representation early on and maintains that advantage. This leads to the generally low conflict demonstrated in Simulation I. Thus, there is a high probability that the selected response at t_{20} and the check at t_{25} are the same, i.e., no repair. *Error trials*, on the other hand, often entail noisy activation of both responses up to the response point (high conflict), which may lead to the selection of the error representation as response at t_{20} (Fig. 4b). Over time, however, continued processing increases the signal-to-noise ratio, leading to the correct representation gaining advantage over the error representation. We propose the detection of this new response triggers the repair process.

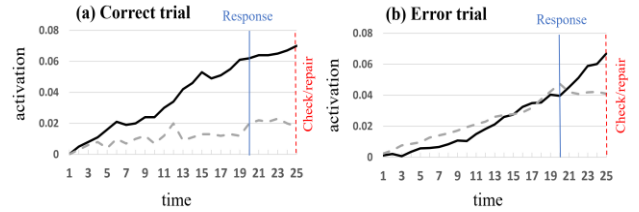


Figure 4. Example of activation dynamics in correct and error trials. Response is the most active node at t_{20} . Check happens at t_{25} . A repair is made if the response at $t_{20} \neq t_{25}$.

The simple and automatic nature of this process makes for an elegant model, but it also assumes that anytime the response and the check are different, a repair is undertaken. The efficiency of this claim must be verified, because if the mechanism also turns many correct responses into errors, then it is fairly useless as a repair mechanism. Simulation II addresses this issue within a signal detection framework. Using 50,000 simulations with the time-based model, we implemented the “respond and check” mechanism. The most activated nodes at t_{20} (response) and t_{25} (potential repair) were determined on each trial. Four trial types ensued from (t_{20} , t_{25}) response pairing: a hit (error, correct), a miss (error, error), a correct rejection (correct, correct), and a false alarm or FA (correct, error). A good repair model is expected to have a high hit rate, together with a low FA rate.

Results. The simulation returned a hit rate of 76% and a FA rate of 1.8%. While this FA does not look that high, recall that there are many more correct trials than error trials, and thus a FA > 1% implies many correct responses turned into errors, which is neither efficient, nor common. A reasonable FA rate must, however, be determined in light of what is reasonable as a hit rate. Corpus analyses of repairs report that hit rates often hover between 50-65% (Nooteboom, 2005). We must thus inspect what the model's FA rate is at hit rates corresponding to these values. But how to set different hit rates for the model? The conflict theory provides an answer: we set a criterion, a value of conflict above which the monitor determines that a trial is likely to be an error and therefore implements a repair if response at $t_{20} \neq t_{25}$ (Fig. 2a). Figure 5 shows the relationship between the values of criterion and hit rate/FA rates for criteria between 2.75 and 6.5 in increments of 0.75. The smaller the criterion, the higher the hit rate, but also the higher the FA rate. A criterion of 4.25 generates 62% hit rate (the upper bound in the corpus analyses), with a 0.7% FA rate, showing that for usual hit rates (< 65%), the model's FA rate is reasonably low.

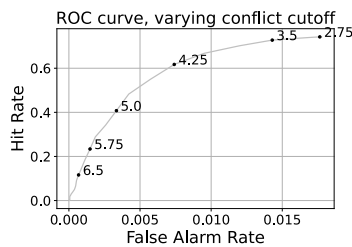


Figure 5. ROC of conflict-based repairs generated in Simulation II. The values on the curve show the conflict-based criteria for implementing a repair.

Discussion. Using a simple “respond and check” mechanism with conflict thresholding, the basic conflict-based repair model successfully captures the repair pattern in human data.

Simulation III: handling adaptive repairs

The basic repair model in Simulation II captures the fast and automatic nature of repairs well, but it is unlikely that this is all there is to repairs. Detecting the need for repairs is a strong signal that additional control resources are needed to maintain optimal performance. In humans, this is achieved by a monitoring-control loop, which constantly assesses the need, and deploys control accordingly (e.g., Botvinick et al., 2001; Yeung et al., 2004). The basic repair model has no such mechanism and is thus expected to fail in capturing the adaptive nature of repairs. Simulation III was designed to test this shortcoming using a well-replicated empirical finding in human speech error repair data, namely that an increase in the probability of errors is often accompanied by a greater proportion of repairs on such errors (Levelt, 1983, 1989; Nooteboom & Quené, 2015; Nozari et al., 2019). In other words, upon detecting the higher probability of committing a speech error, speakers adapt their behavior (e.g., by

becoming more vigilant) to repair more of those errors, and thus keep communication informative.

Empirical data are taken from Nozari et al. (2019). Participants watched short cartoon clips in a “haunted hotel” scenario, in which objects in a hotel room performed different actions on a second object, and described the event in sentences like “*The yellow curtain jumped over the brown window*” in real time. The second noun phrase, NP2 (e.g., the brown window), elicited significantly more errors than the first noun phrase, NP1, which was semantically related to it (e.g., the yellow curtain; Fig. 6a). Critically, participants also repaired a significantly greater proportion of their errors on NP2 than on NP1 (Fig. 6b).

The greater error rate on NP2 stems from the semantic blocking effect (e.g., Schnur et al., 2009), a finding that producing a target word makes it more difficult to produce a semantically related word later. A common explanation for this is an error-based learning account, according to which, upon producing NP1, the connections between the subset of semantic features shared between NP1 and the related NP2 become stronger for NP1 and weaker for NP2 (Oppenheim et al., 2010; Oppenheim & Nozari, 2021). This differential adjustment of weights leads to a greater likelihood of producing NP1 as an error when NP2 next becomes the target. Our goal was not to implement the full learning model in the repair model. Rather, it was to simulate the end-point of learning, as implemented in Oppenheim and Nozari (2021), to capture the mechanism underlying the higher error rates on NP2. Simulation III was thus run with the default values in Table 1 for simulating NP1, and adjusted post-learning weights for simulating NP2 (i.e., positive and negative changes to the selective weights targeted by error-based learning, with $\delta = 4.25E^{-5}$). A criterion of 4.25 from Simulation II was used to adjust the hit rate/FA ratio. We ran 20 batches of 10,000 simulations for each NP.

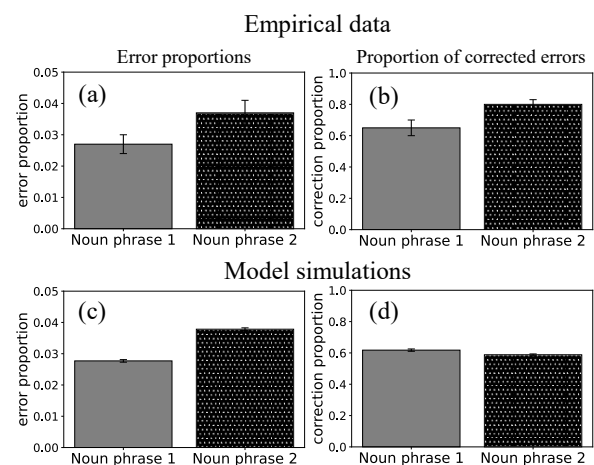


Figure 6. Results of Simulation III. Mean error proportions $\pm$ SE (a) and proportion of corrected error $\pm$ SE (b) in empirical data from Nozari et al. (2019). Corresponding data on error proportions (c) and proportion of corrected error (d) from Simulation III.

Results. Figure 6 shows the results of Simulation III. With simulated error rates (Fig. 6c) matched to the empirical data (Fig. 6a), and the criterion of 4.25, the model correctly estimates the proportion of corrected errors on NP1 to be around 60% (Fig. 6d), as observed in the empirical data (Fig. 6b). However, it fails to capture the increase in the proportion of corrected errors in the more error-prone (NP2) situation. If anything, the model estimates a slightly lower repair rate on NP2 (59%; SE = 0.75%) than NP1 (62%, SE = 0.62%; Fig. 6d), a pattern opposite of that observed in the empirical data.

Discussion. While the basic conflict-based model captures the repair pattern in the baseline (NP1) condition well, it fails to show the enhanced repair rates in the more difficult (NP2) condition observed in the empirical data. Simulation IV proposes an augmented model to address this problem.

Simulation IV: the augmented model of repair

Results of Simulation III showed the limitation of the basic model in simulating the adaptive nature of repairs. This is not surprising, since the model has no regulatory mechanism. Simulation IV augments the model with a feedback loop. This loop consists of a monitor that keeps track of the amount of conflict and scales the input proportional to that amount.

One way to construct such a loop is to use the value of conflict in each trial to scale the input accordingly. This model would represent a system with no memory. More reasonable, however, is a loop that retains some information about the relative difficulty and likelihood of error in certain situations. For the data set used in Simulation III, for example, the speakers often experience greater difficulty on the production of NP2 than NP1, learning that this position is an error-prone situation over the course of the experiment. We model this process by maintaining a running average of conflict on NP1 and NP2 separately, updated each time a new NP is produced. This constitutes the monitoring part of the loop. Control is implemented by scaling the input proportional to the running average of conflict in each condition (NP1 vs. NP2). Cognitively, this translates into focusing attention on the concept of the to-be-produced word, in order to produce the correct label, a process likely to increase accuracy in language production.

The simplest way to implement a monitoring-control loop is to have a linear function to model the input gain based on average conflict. The linear function has two problems though: (a) theoretically, it implies that every trial will get some boost, i.e., at least some level of additional control is constantly recruited. It also means that there is no upper bound to how much control should be recruited. This would be extremely energy-consuming. (b) Empirically, a linear function fails to improve the pattern observed in Fig. 6d, because it cannot overcome the differential weight changes that caused the disadvantage for the correct response on NP2 to begin with. More reasonable is a function that scales input within a certain range, i.e., with lower and upper bounds on a linear function. This is best captured by a sigmoidal function. We thus modeled the input gain as a logistic function (Eq. 6).

$$\text{inputGain} = \frac{a}{1 + e^{-b(\bar{x}-c)}} \quad \text{Eq. 6}$$

where $\bar{x}$ is the running average of conflict within a condition, a determines the activation bounds, b determines how graded control implementation should be, and c determines a “set point” against which fluctuations of conflict are measured (Fig. 8a). It is reasonable to assume that such a set point is derived from speakers’ vast experience with language production, i.e., is the current situation easier or more difficult than what I am used to? Thus to determine this set point in the model, we ran the model for 200,000 trials over each NP1 and NP2, and averaged conflict over all trials to simulate a range of experience. The obtained value, 3.42, was used for parameter c in Simulation IV. No a priori information is available to determine the values of a and b . Therefore, we ran a parameter space search to determine how different values of these two parameters change the behavior of the augmented conflict-based model of repair.

The final model

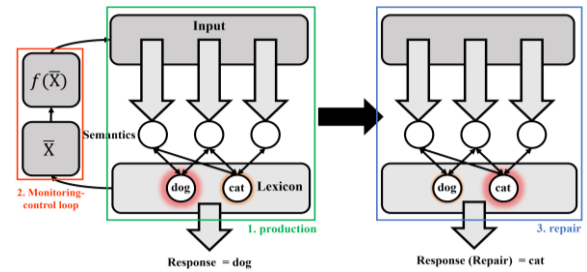


Figure 7. A schematic of the augmented conflict-based model of repair. The example shows a trial in which the error “dog” has been repaired to the target “cat”.

Figure 7 presents the augmented model of repairs. The model has three stages: *response production*, *monitoring-control loop*, and *repair*, with the steps outlined below.

1. Response production

- Input for “cat” is activated at the semantic level.
- Lexical nodes of “cat” and “dog” are activated.
- The most active lexical node is produced (response).

2. Monitoring-control loop

- Conflict (X) is measured at the time of response (t_{20}).
- Input at t_{21} receives a boost, which is scaled as a function of the running average of conflict ($f(\bar{X})$). The gain is near-zero for conditions with conflict levels lower than the set-point, thus control is effectively only recruited on demanding trials.

3. Repair (or not)

- Processing continues for a short period after response generation with the scaled input.
- The most active node in the lexical layer at t_{25} is selected as a potential repair.
- If original response = potential repair $\rightarrow$ no repair. If original response $\neq$ potential repair $\rightarrow$ change response to potential repair, i.e., the model repairs the

original response. A criterion can be imposed at this stage to adjust the hit rate. This model was used in Simulation IV.

Results. Figure 8b shows the results of the parameter search, and Figure 8c illustrates the simulation results using $(a, b) = (40, 70)$, which provide a good match to the empirical data. Two points are noteworthy about Fig. 8b. First, a large part of the parameter space (values above 0) produces the effect of interest, i.e., the increase in the proportion of repairs in the more error-prone condition. Second, only values of b above 70 produce effect sizes comparable to the empirical data. These values correspond to a steep logistic function, in which small values of conflict lead to near-zero gains on the input, while conflict values above the set-point deploy full control. This, in turn, means that control is rarely recruited in situations where the difficulty (indexed by conflict) is estimated to be around or lower than that generally experienced by the speakers, but the detection of a difficult situation can quickly engage high levels of attention.

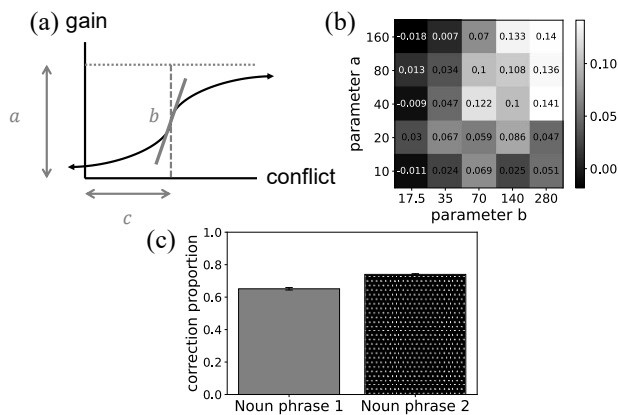


Figure 8. The gain function (a), and the results of the parameter search for its parameters a and b . The dependent variable in the heat map is the increase in proportion of repairs from NP1 to NP2 using the augmented repair model (b). Results of simulation IV with $a = 40$ and $b = 70$ (c).

Discussion. Augmenting the basic conflict-based model of repair with a monitoring-control loop enables it to simulate more nuanced patterns observed in human speech error repair data, including the adaptive nature of repairs.

General Discussion

The main idea expressed in this paper was that many speech errors can be repaired by quickly replacing a response with an already-activated alternative, instead of restarting the production process. In four simulations, we gradually built upon the basic structure of the two-step model of word production to arrive at an augmented conflict-based model, which, to our knowledge, is the first computational model of speech error repairs in humans. The simple and automatic basic mechanism of this model allows for fast and subconscious detection of lexical errors, even before they

become overt. The addition of the monitoring-control loop further enables the model to capture the nuances in speech error data that reflect the adaptable nature of human behavior.

Aside from capturing the data directly tested in the above simulations, the model provides a theoretical explanation for additional data not directly tested here. The basic subconscious process explains reports of subconscious repairs in children in speech (Clark, 1978) and adults in other modalities of language production and their rapid timing (Pinet & Nozari, 2021). The model also explains a behavior commonly observed in aphasia; patients “grope” for the correct response, by producing a quick string of related repairs. In segmental repairs, this is called *conduite d’approche* (Kohn, 1984), but it is also seen in lexical repairs, for example, “orange, peach, no, apple”, produced in response to the picture of a “watermelon” (Nozari, 2019). This behavior reflects the basic premise of the conflict-based repair model: change the current response if another one overtakes it quickly in activation. In healthy mature systems, the change is often from an error to a correct response, but in damaged systems, where the signal-to-noise ratio is low, it is possible to choose another error as repair, and repeatedly so, because spreading activation does not cleanly converge on the correct response.

One alternative to the conflict-based model is forward models, e.g., DIVA (Guenther, 1994; Tourville & Guenther, 2011). In DIVA, predicted perceptual outcomes of a speech motor command are compared to the actual perceptual outcomes of speech, and discrepancies lead to the generation of an error signal, which can later be used for correction. DIVA is a powerful model for explaining articulatory-phonetic adjustments to speech over time (i.e., on the next trial), but its dependence on overt perceptual outcomes makes it difficult for the model to explain fast lexical repairs initiated pre-verbally. A similar model, the Hierarchical State Feedback Control (HSFC; Hickok, 2012) detects errors through a discrepancy in the activation of motor and corresponding perceptual representations, but does not propose a mechanism for repairs.

Limitations and future directions

(1) This model only implemented lexical repairs, but the basic process is applicable to phonological repairs as well, which should be explored in future work. (2) The finding of an increased proportion of repairs as a function of increased error rates has been reported under different circumstances, only one of which was tested in Simulation III. Future explorations should include other manipulations of error rates (e.g., by changing planning time constraints) and examining how the model’s repair behavior changes as a function of changes to the error rate. (3) Related to the second limitation, more appropriate gain functions may exist, which must be explored based on simulations of a wider range of data.

In conclusion, this work has taken the first step to propose a computational model that successfully captures speech error repair behavior in humans, which can be used as a framework for future testing and improvement.

References

- Botvinick, M. M., Braver, T. S., Barch, D. M., Carter, C. S., & Cohen, J. D. (2001). Conflict monitoring and cognitive control. *Psychological Review*, 108(3), 624–652.
- Budd, M.-J., Hanley, J. R., & Griffiths, Y. (2011). Simulating children's retrieval errors in picture-naming: A test of Foygel and Dell's (2000) semantic/phonological model of speech production. *Journal of Memory and Language*, 64(1), 74–87.
- Clark, E. V. (1978). Awareness of language: Some evidence from what children say and do. In A. Sinclair, R. J. Jarvella, & W. J. M. Levelt (Eds.), *The Child's Conception of Language*. Springer Berlin Heidelberg.
- Dell, G. S. (1986). A spreading-activation theory of retrieval in sentence production. *Psychological Review*, 93(3), 283–321.
- Dell, G. S., Schwartz, M. F., Martin, N., Saffran, E. M., & Gagnon, D. A. (1997). Lexical access in aphasic and nonaphasic speakers. *Psychological Review*, 104(4), 801–838.
- Dell, G. S., Schwartz, M. F., Nozari, N., Faseyitan, O., & Branch Coslett, H. (2013). Voxel-based lesion-parameter mapping: Identifying the neural correlates of a computational model of word production. *Cognition*, 128(3), 380–396.
- Foygel, D., & Dell, G. S. (2000). Models of impaired lexical access in speech production. *Journal of Memory and Language*, 43(2), 182–216.
- Guenther, F. H. (1994). A neural network model of speech acquisition and motor equivalent speech production. *Biological Cybernetics*, 72(1), 43–53.
- Hanley, J. R., Cortis, C., Budd, M.-J., & Nozari, N. (2016). Did I say dog or cat? A study of semantic error detection and correction in children. *Journal of Experimental Child Psychology*, 142, 36–47.
- Hartsuiker, R. J., & Kolk, H. H. J. (2001). Error monitoring in speech production: A computational test of the perceptual loop theory. *Cognitive Psychology*, 42(2), 113–157.
- Hickok, G. (2012). Computational neuroanatomy of speech production. *Nature Reviews. Neuroscience*, 13(2), 135–145.
- Kohn, S. E. (1984). The nature of the phonological disorder in conduction aphasia. *Brain and Language*, 23(1), 97–115.
- Levelt, W. J. M. (1983). Monitoring and self-repair in speech. *Cognition*, 14(1), 41–104.
- Levelt, W. J. M. (1989). *Speaking: From intention to articulation*. The MIT Press.
- McDonagh, D., Fisher, A. V., & Nozari, N. (2020). Do taxonomic and associative relations affect word production in the same way? *Proceedings of the Annual Conference of the Cognitive Science Society*.
- Nooteboom, S. G. (2005). Listening to oneself: Monitoring speech production. In R. J. Hartsuiker, R. Bastiaanse, A. Postma, & F. Wijnen (Eds.), *Phonological encoding and monitoring in normal and pathological speech*. Psychology Press.
- Nooteboom, S., & Quené, H. (2015). Word onsets and speech errors. Explaining relative frequencies of segmental substitutions. *Journal of Memory and Language*, 78, 33–46.
- Nooteboom, S., & Quené, H. (2019). Temporal aspects of self-monitoring for speech errors. *Journal of Memory and Language*, 105, 43–59.
- Nozari, N. (2019). The dual origin of semantic errors in access deficit: Activation vs. inhibition deficit. *Cognitive Neuropsychology*, 36(1-2), 31–53.
- Nozari, N. (2020). A comprehension- or a production-based monitor? Response to Roelofs (2020). *Journal of Cognition*, 3(1), 19.
- Nozari, N., Dell, G. S., & Schwartz, M. F. (2011). Is comprehension necessary for error detection? A conflict-based account of monitoring in speech production. *Cognitive Psychology*, 63(1), 1–33.
- Nozari, N., Kittredge, A. K., Dell, G. S., & Schwartz, M. F. (2010). Naming and repetition in aphasia: Steps, routes, and frequency effects. *Journal of Memory and Language*, 63(4), 541–559.
- Nozari, N., Martin, C. D., & McCloskey, N. (2019). Is repairing speech errors an automatic or a controlled process? Insights from the relationship between error and repair probabilities in English and Spanish. *Language, Cognition and Neuroscience*, 34(9), 1230–1245.
- Oppenheim, G. M., Dell, G. S., & Schwartz, M. F. (2010). The dark side of incremental learning: A model of cumulative semantic interference during lexical access in speech production. *Cognition*, 114(2), 227–252.
- Oppenheim, G., & Nozari, N. (accepted). Behavioral interference or facilitation does not distinguish between competitive and noncompetitive accounts of lexical selection in word production. In T. Fitch, C. Lamm, H. Leder, & K. Tessmar (Eds.), *Proceedings of the 43rd Annual Conference of the Cognitive Science Society*. Cognitive Science Society.
- Pinet, S., & Nozari, N. (in press). Correction without consciousness in complex tasks: Evidence from typing. *Journal of Cognition*.
- Schnur, T. T., Schwartz, M. F., Kimberg, D. Y., Hirshorn, E., Coslett, H. B., & Thompson-Schill, S. L. (2009). Localizing interference during naming: Convergent neuroimaging and neuropsychological evidence for the function of Broca's area. *Proceedings of the National Academy of Sciences*, 106(1), 322–327.
- Tourville, J. A., & Guenther, F. H. (2011). The DIVA model: A neural theory of speech acquisition and production. *Language and Cognitive Processes*, 26(7), 952–981.
- Yeung, N., Botvinick, M. M., & Cohen, J. D. (2004). The neural basis of error detection: Conflict monitoring and the error-related negativity. *Psychological Review*, 111(4), 931–959.