

# The Cognitive Modeling of Errors During the Japanese Phonological Awareness Formation Process

|       |                                                                                                                              |
|-------|------------------------------------------------------------------------------------------------------------------------------|
| メタデータ | 言語: eng<br>出版者:<br>公開日: 2020-08-21<br>キーワード (Ja):<br>キーワード (En):<br>作成者: Nishikawa, Jumpei, Morita, Junya<br>メールアドレス:<br>所属: |
| URL   | <a href="http://hdl.handle.net/10297/00027611">http://hdl.handle.net/10297/00027611</a>                                      |

# The Cognitive Modeling of Errors During the Japanese Phonological Awareness Formation Process

Jumpei Nishikawa (nishikawa.jumpei.16@shizuoka.ac.jp)<sup>1</sup>

Junya Morita (j-morita@inf.shizuoka.ac.jp)<sup>1</sup>

<sup>1</sup>Faculty of Informatics, Shizuoka University, 3-5-1, Johoku, Naka-ku, Hamamatsu, Shizuoka 432-8011, Japan

## Abstract

Both nature and nurture contribute to language development. In the case of phoneme segmentation, children have the natural ability to recognize a continuous sound in various units, but as they grow, they only selectively learn to recognize it as part of a series in the unit that is used in their mother tongue. This developmental process is supported by an ability called phonological awareness that allows children to become intentionally aware of units of phonology. It is known that erroneous pronunciation appears during the phonological awareness formation process. In this research, we aim to examine the factors that induce and reduce such errors. To do so, we modeled phonological awareness using the cognitive architecture ACT-R and performed simulations that manipulated ACT-R parameters that correspond to both nature and nurture factors. As a result, it was confirmed that errors due to a lack of phonological awareness can be modeled with the innate memory retrieval mechanism. We also observed that such errors were reduced when learning factors were added to the model. However, we could not simulate this learning process. In the future, we will study the interaction task that enables learning to reduce phonological errors and contribute to the acquisition of phonological awareness.

**Keywords:** cognitive modeling; phonological awareness; Japanese;

## Introduction

Language development is influenced by both innate and experiential factors. Newborn infants have an innate foundation for acquiring a variety of languages. Through experiences such as observing and imitating the behavior of familiar adults (Baron-Cohen, 1997; Tomasello, 1999), they gradually form a specific mother language structure, based on a general, innate language foundation. A prominent example of such convergence can be found in phoneme segmentation, which is one of the linguistic components. Infants initially have opportunities to divide a sound into various types of segments, like syllables or morae, though, as they grow up and master their language, they find it difficult to recognize phoneme segments that do not belong to their mother tongue.

In the fields of developmental psychology and speech-language pathology, this developmental process is partly attributed to an ability called phonological awareness, which enables one to intentionally pay attention to phonological aspects, such as phonemes and rhythms, of oral languages (Stahl & Murray, 1994). This ability is usually developed during the preschool period and relates to the acquisition of reading and writing skills. Researchers have detected several

phonological errors that seem to occur due to a lack of phonological awareness during the language development process (Kubozono, 1989). In addition, there are reported cases in which infants with atypical developmental traits, such as those on the autism spectrum, who experienced an overall delay in acquiring phonological units, improved these problems with support for enhancing their phonological awareness (Mugitani et al., 2019).

Although psychologists have revealed the age at which phonological awareness is formed and its role in language development, it has not been clarified that the internal processes, cognitive functions, and mechanisms behind the occurrence and suppression of errors relating this ability.

Based on this background, the current study aims to examine error occurrence and suppression factors that are related to the phonological awareness formation process. We focus on two language development factors, namely nature and nurture, and construct a model using the cognitive architecture ACT-R (Anderson, 2007). In addition, we examine the process and mechanism of the phonological awareness formation process by simulating Japanese language development.

The structure of this paper is as follows. The first section introduces the research related to this study. The model and the model-based simulation will then be shown. Finally, a summary of the current status and future issues is presented.

## Related Research

In this section, we first introduce previous studies on language development and the formation of phonological awareness. Next, we introduce cognitive modeling and cognitive architecture as the method for understanding and explaining human cognitive processes. We also review research on language learning using cognitive models.

## Studies on Phonological Awareness

In the field of phonology, researchers have developed systems for classifying the sounds that humans perceive and utter as mental representations of sounds. The most famous system is the distinctive feature (Chomsky & Halle, 1968), which is the basic unit that distinguishes phonemes; it is typically defined as a binomial variable that takes the value of + or – by classifying the tongue and throat movements associated with vocalization. Table 1 provides the distinctive features for some phonemes.

Table 1: Distinctive feature examples

|                 | a | i | u | e | o | k | g | s | z | t | d | n |
|-----------------|---|---|---|---|---|---|---|---|---|---|---|---|
| consonantal     | - | - | - | - | - | + | + | + | + | + | + | + |
| syllabic        | + | + | + | + | + | - | - | - | - | - | - | - |
| sonorant        | + | + | + | + | + | - | - | - | - | - | - | + |
| high            | - | + | + | - | - | + | + | - | - | - | - | - |
| back            | + | - | + | - | + | + | + | - | - | - | - | - |
| low             | + | - | - | - | - | - | - | - | - | - | - | - |
| anterior        |   |   |   |   |   | - | - | + | + | + | + | + |
| coronal         |   |   |   |   |   | - | - | + | + | + | + | + |
| voice           |   |   |   |   |   | - | + | - | + | - | + | + |
| continuant      |   |   |   |   |   | - | - | + | + | - | - | - |
| nasal           |   |   |   |   |   | - | - | - | - | - | - | + |
| strident        |   |   |   |   |   | - | - | + | + | - | - | - |
| delayed release |   |   |   |   |   | - | - | - | - | - | - | - |
| round           | - | - | + | - | + | - | - | - | - | - | - | - |

Around the world, there are varieties of systems that combine these innate phonetic elements as units. According to Trubetzkoy (1969), such systems fall into two groups, namely “mora languages” and “syllable languages,” based on the smallest prosodic unit that is used in that language. Among the mora languages, Japanese defines the mora as a unit of duration (Bloch, 1950), and each mora is associated with a single *kana* (Japanese character), which consists of five vowels (*a, i, u, e, o*) and 59 combinations of consonants and vowels (*ka, ki, ku, ke, ko, sa, si, etc.*), and other special morae representing the duration or gemination of sounds. In other words, Japanese phonic elements have a clear connection with written symbols. We think that this characteristic offers an advantage for modeling phonological awareness in a symbolic cognitive architecture. Therefore, in this study, we focus on Japanese morae and try to reveal factors that relate to the acquisition process of this ability.

Several Japanese researchers have investigated the formation of phonological awareness. Hara (2001) presented a study where participants engaged in several phonological manipulation tasks. For example, in the mora deletion task, participants were required to respond with a mora sequence that involved deleting a single mora from an orally presented word (e.g., *i-ko* is the correct answer for a task that deletes *ta* from the orally presented word *ta-i-ko*, which means “a drum”). Similarly, in the mora reversal task, participants were required to respond with a reverse-ordered mora sequence after an oral word prompt (e.g., *ka-i-su* for *su-i-ka*, which means “a watermelon”). In her study, the performance of such tasks was increased according to participants’ written and reading language skills.

In a clinical setting, Japanese speech-language-hearing therapists (ST) reported cases where children confused specific types of morae. For example, it was reported that young children around 2 to 3 years old tend to confuse morae that include the consonants /r/ and /d/ or /s/ and /sj/, though Japanese adults can clearly distinguish these (Kobayashi, 2018). Moreover, there have been reports that some children with developmental disorders have difficulty distinguishing a mora that consists of only a vowel as well as other morae that include that vowel (e.g., confusions between *a* and *ka*, *sa*, or *ta*) (Ishida & Ishizaka, 2016). Typically developed chil-

dren also exhibited similar erroneous utterances in which they omitted consonants (e.g., pronouncing *a-p-pa* instead of *ra-p-pa*; trumpet) (Nakamura, Kojima, & Fujiwara, 2015). This type of confusion suggests the existence of a developmental process that migrates from the innate phonological system (Chomsky & Halle, 1968) to the language-specific mora system.

Furthermore, several researchers have used the popular Japanese word game “*shiritori*” as a task for examining a development stage of phonological awareness. In this game, players take turns providing a noun whose first character (mora) is the same as the last character of the previously given noun. For example, after a player answers “*ri-n-go*” (meaning: apple), the other player continues with “*go-ri-ra*” (meaning: gorilla). The game is over when a player provides a word that ends with /N/, since no Japanese word can begin with this character (e.g., “*ri-n-go*” → “*go-ri-ra*” → “*ra-p-pa*” (trumpet) → “*pa-n*” (Bread) → game over). To avoid looping, the game also ends if a player repeats a word that was already provided as an answer in the game (e.g., “*ma-su-ku*” (mask) → “*ku-ru-ma*” (car) → “*ma-su-ku*” → ...).

Takahashi (1997) examined the conditions for being able to play *shiritori* through a cross-sectional developmental psychological experiment involving children with typical development. This research indicated that playing *shiritori* requires the ability to divide sounds into morae and the mental lexicon indexed by phonemes, and that the acquisition of kana characters is effective for indexing vocabulary by morae. These results suggest that playing *shiritori* requires phonological awareness, paying attention to a phoneme in the mother language (mora) in a sound, and that such ability is enhanced by presenting visual aids (kana character) that correspond to the sound. Furthermore, it has been shown that even if a child does not have the phonological awareness that is necessary to play *shiritori*, s/he can participate with help from adults, who can provide hints.

Kubozono (2000) also conducted an experiment that examined the *shiritori* process in a 4-year-old Japanese child in order to present the phonological awareness formation process. In his experiment, the participant was sometimes confused about the unit of mora. For example, she gave “*mo-n-shi-ro-cho-u*” (cabbage butterfly) as an answer for “*do-ra-e-mo-n*” (doraemon). She also noted “*yo-u-gu-ru-to*” (yogurt) after “*ta-i-yo-u*” (sun) was presented. Based on these two examples, the participant seems to have recognized the consonant-vowel-vowel sequence (“*mo-n*” and “*yo-u*”) as a single unit, although the Japanese mora system divides this into two morae, namely “consonant-vowel” and “vowel.” These reports suggest that the language-specific phonetic system is experientially acquired after childhood, and the misconfiguration of the system can be observed culturally while playing popular word games.

Based on the above-mentioned previous findings, the current research focuses on the ability to extract the mora at the end of a word from the continuous sounds that correspond

to words and the ability to search for words by initial mora, especially in phonological awareness. In addition, we apply *shiritori* as the task set based on Takahashi’s and Kubozono’s work. In addition, we refer to Chomsky and Halle (1968)’s definition and use a distinctive feature to express and characterize the mora as a symbol.

## Cognitive Modeling

One of the methods for understanding and explaining the mechanisms and processes that are related to human cognition is through the cognitive modeling approach in the field of cognitive science. In this approach, a model that approximates a person’s task execution (a cognitive model) is built as a computer program, and a simulation is done using the model. By observing the cognitive model’s behavior and internal states during the simulation, we can infer a person’s internal states and cognitive processes during task execution.

Cognitive architectures have been developed as a basis for cognitive modeling. Using a cognitive architecture, it is possible to construct a model that isolates the factors related to task achievement on a common basis. From among the various cognitive architectures that have been developed, we selected ACT-R (Anderson, 2007) for use in this study. Since ACT-R is based on psychological experiments on thinking and memory, it enables us to comprehensively grasp various phenomena related to human cognition. Furthermore, ACT-R is a production system with multiple modules. There are various parameters that define the module’s behavior, thus making it easier to model the individual. There are also modules that handle interaction with the external environment, thus making it possible to predict reaction times and associate them with experimental data.

There have been many studies on language acquisition using ACT-R. For instance, models related to the acquisition of irregular verbs in English learning (Taatsgen & Anderson, 2002) and models of infants’ noun learning (Van Rijn, Van Rijn, & Hendriks, 2010) have been constructed. In addition, brain dysfunction has also been modeled by ACT-R, and some studies have explained the errors in sentence comprehension in aphasia using ACT-R parameters (Mätzig, Vassith, Engelmann, Caplan, & Burchert, 2018).

In this study, we aim to model the language development process by mapping phonological awareness from the memory retrieval mechanism of the ACT-R declarative module. In particular, we focus on parameters that express knowledge similarity and the effect of learning, assuming that changes in these parameters through simulation correspond to innate, experiential language development factors.

## Model

In this section, we describe the model of errors in the Japanese phonological awareness formation process. The model executes *shiritori* to exhibit reported errors and demonstrate the factors that suppress such errors.

An overview of the model is shown in Figure 1. This model includes two agents (dashed line area) who keep a game of

*shiritori* going by taking turns providing words. Boxes in the agent area corresponds to each module of ACT-R. In the following, we show how the *shiritori* process is realized through the ACT-R module structure.

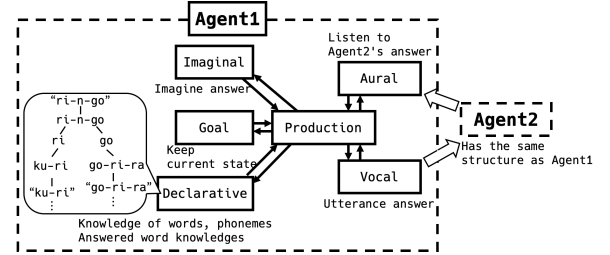


Figure 1: An overview of the model

## Module Structure

**Declarative Module** The declarative module of ACT-R contains knowledge that is required for task execution in the form of chunks. The model in this study retains three types of chunks that relate to word (vocabulary) and phonological knowledge, and the association between words and phonological knowledge (Table 2). The model also has chunks that store the words that have already been provided in the current *shiritori* trial. These chunks do not exist in the declarative module at the beginning of the trial; rather, they are generated and stored as the trial progresses using imaginal module.

Table 2: The model’s declarative memory

(a) Word knowledge

| word   | sound    |
|--------|----------|
| ringo  | “ringo”  |
| gorira | “gorira” |
| kuri   | “kuri”   |
| ...    | ...      |

(b) Phonological knowledge

| mora | sound |
|------|-------|
| /ri/ | “ri”  |
| /go/ | “go”  |
| /ku/ | “ku”  |
| ...  | ...   |

(c) The word–mora relationship

| word   | mora | position |
|--------|------|----------|
| ringo  | /ri/ | head     |
| gorira | /go/ | tail     |
| gorira | /go/ | head     |
| ...    | ...  | ...      |

**Production Module** The production module selects and applies rules, and operates the module, while using information and states that are held by other modules. In the model of this research, when word information is received as a partner’s answer, the novel word is retrieved and provided as an answer according to the rules of *shiritori* that were presented in the previous section.

Figure 2 summarizes this process. Using the word chunk (chunk type b in Table 2) acquired by the aural module, the model retrieves a chunk that connects the word and the ending mora (chunk type c in Table 2). Phonological knowledge

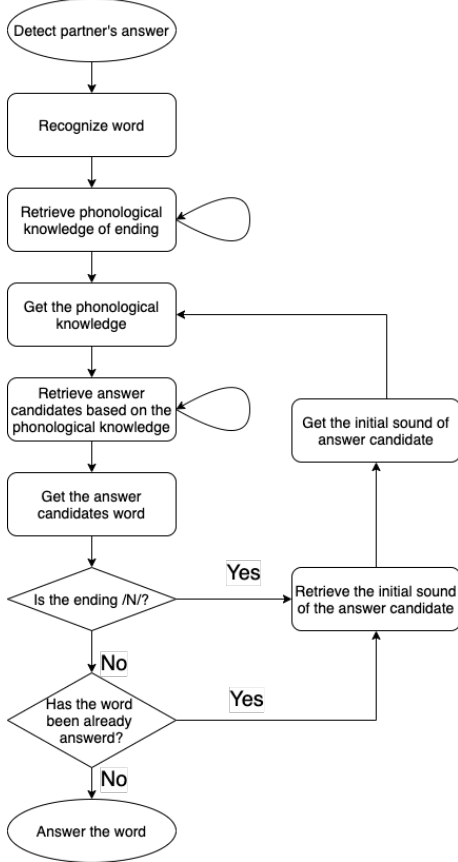


Figure 2: The answer process

(chunk type b in Table 2) is activated using this chunk. Activated phonological knowledge guides the retrieval of a chunk where the same mora is connected with a new word in an initial position (chunk type c in Table 2). When successful, the word knowledge in the retrieved chunk is stored in the goal module as an answer candidate.

After this, the model checks that the stored answer candidate is valid according to the rules of *shiritori*, such as not including /N/ as the end of a mora and not having been answered before in the current *shiritori* trial. If the current candidate violates these rules, the model re-searches for the answer candidate.

When the candidate word is confirmed as valid, the model stores it in the declarative module as an answered word and outputs the word through the speech module. During one *shiritori* task trial, the two agents alternately execute this procedure up to the given time limit.

## Model Parameters

An important parameter in this study is related to chunk retrieval in the declarative module (activation). An activation is assigned to each chunk that is stored in the declarative module and affects the success or failure of the retrieval and the time required for the retrieval. The activity value  $A_i$  that is assigned to chunk  $i$  is defined as the addition of multiple terms

as in Eq. 1:

$$A_i = B_i + S_i + P_i + \epsilon_i \quad (1)$$

If more than one chunk matches the search request from the production module, the chunk with the highest activation is selected. In our model, among the activation elements, we focus on the similarity  $P_i$  and the base level  $B_i$ .

**Similarity** The similarity term  $P_i$  of the activation assigned to chunk  $i$  is computed using Eq. 2:

$$P_i = \sum_k PM_{ki} \quad (2)$$

This value is computed as the summation of the weighted degree of similarity  $M_{ki}$  for each retrieval request  $k$  to the chunk  $i$ .  $M_{ki}$  usually takes a negative value, and  $P$  serves as a penalty in the effect of similarity retrieval. Introducing similarity effects into the model's knowledge allows for the use of a mechanism called partial matching. For production module search requests, it is possible to search for chunks that do not have an exact match, thus allowing for flexible selection and the reproduction of certain errors.

In this study, for each phonological knowledge combination, we set  $M_{ki}$  based on the distinctive feature (Chomsky & Halle, 1968). That is, we assumed that the similarity between two phonemes can be defined as the overlap of distinctive features. In the simulation that is reported in the next section, we treat this parameter as the innate factor of language development, with the expectation of frequently observing errors involving the confusion of similar morae in the early phase of development.

**Base Level** The base level is the basic element of the activity value that corresponds to learning and forgetting. It is represented by Eq. 3:

$$B_i = \ln \left( \sum_{j=1}^n t_j^{-d} \right) + \beta_i \quad (3)$$

The value is computed from the number of presentations for chunk  $i$  ( $n$ ) and the elapsed time since the chunk was referenced ( $t_j$ ).  $d$  indicates the decay rate, and  $\beta_i$  is the offset parameter that can be modulated according to the simulation's aim. In this study, the base level is introduced into the model to examine how the experiential factor influences the development of phonological awareness. It is assumed that the innate factor's (similarity) relative importance decreases as base level activation increases and that errors related to phonological awareness are suppressed.

## Learning Interactions

Children with inadequate phonological awareness need assistance from adults to play *shiritori* (Takahashi, 1997). In this study, we introduce assistance that encourages undeveloped agents to continue the *shiritori* task. We also assume asymmetric interactions, such as those between children and parents, setting the similarity parameter for only one of the two

agents. For the other agent, we incorporate the process of presenting the earlier answer again when the wrong answer has been received (for the wrong answer of “*o-ka-si*” (sweets) to “*ri-n-go*”, the agent presented “*ri-n-go*” again). Since ACT-R learning is calculated according to the frequency of using a chunk, it is expected that presenting correct answers increases activation of the correct association between words and morae.

## Simulation

In this section, we describe three simulations using the model.

### Error Occurrence Factors

**Simulation Settings** First, we test whether it is possible to model phonological errors using the ACT-R knowledge similarity function. Specifically, we perform simulations under the varying conditions of the ACT-R parameters (:mp), corresponding to the similarity weights  $P$  in Eq. 2.

The value of  $P$  was set to the following eight conditions: 1, 10, 20, 30, 40, 50, 60, and 70. Word knowledge in the model was selected from those listed in Amano and Kobayashi (2008)’s Japanese word database. Based on the rules of *shiritori*, we took out 20,544 nouns, excluding homonym duplications, and words consisting of only one mora, such as “*ro*” (furnace) and “*wa*” (ring). Next, referring to the child’s associative vocabulary survey (for Japanese Language & Linguistics, 1981), 2,054 words were randomly taken out and used as model knowledge. For phonological knowledge, 103 pieces of knowledge were defined based on Japanese morae.

In this research, we use a unit called a “chain,” which is the number of *shiritori* continuations. A trial in the simulation is terminated when the model achieves 100 chains or after 3,600 seconds have passed from the beginning of the trial. A total of 100 trials were simulated for each condition. Phonological knowledge similarity was calculated as the cosine similarity with the distinctive feature column of one mora as a vector, based on Chomsky’s table of distinctive features (Chomsky & Halle, 1968).

**Results and Discussion** Figure 3 shows the types and numbers of errors that appeared during the simulation. In this graph, the horizontal axis indicates a phonological knowledge pair (morae) arranged in order from left to right based on similarity. The vertical axis shows the number of errors that occurred during the simulation using the pair of morae. For example, an incorrect answer of “*ri-n-go*” → “*o-ka-si*” counts as a pair of “*go - o*.”

In the graph, the total number of errors is smaller in the condition where the value of  $P$  is larger, and the majority of errors is concentrated on the left side of the figure. From these results, it can be confirmed that phonological awareness errors are actually invoked by incorporating knowledge similarity. In addition, it can be observed that a larger similarity weight reduces the number of errors, and that errors are less likely to occur in low-similarity pairs. This is thought to correspond to the phenomenon that was confirmed in the

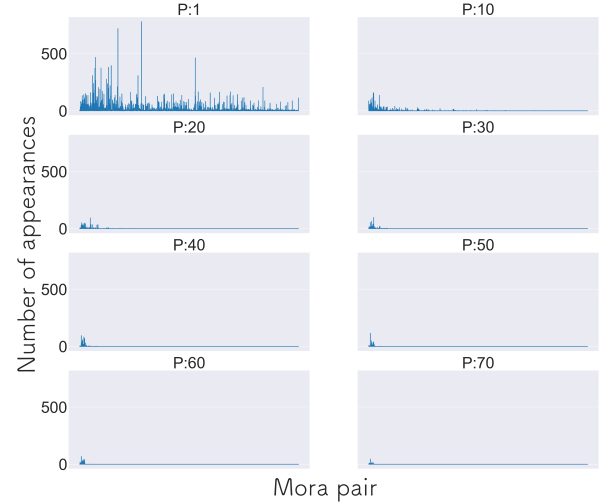


Figure 3: A comparison of similarity differences and errors

child development process where the “*ra*” sound is confused with a “*da*” sound (Kobayashi, 2018).

### Error Suppression Factors

**Simulation Settings** Following the simulations that were described in the previous section, we will now examine the factors that suppress phonological errors. Here, we hypothesize that the impact of similarity is relatively reduced by the effect of learning. To test this hypothesis, we simulate a condition in which the ACT-R parameter (:blc) corresponding to the value of the offset  $\beta_i$  in Eq. 3 is changed to the following six conditions: 0.1, 1, 5, 10, 15, and 20, while keeping the same setting in the other parameters, as in the previous section’s simulation.

In addition, we examine whether learning occurs as a result of the execution of the tasks that were set in this study. We set the condition that the offset  $\beta_i$  is set to 0 and observe the time-series change through the execution of the task.

**Results and Discussion** Figure 4 shows the change in the rate of correct answers due to the manipulation of the offset value  $\beta_i$  in Eq. 3. The horizontal and vertical axes, respectively, represent the value of  $\beta_i$  and the rate of correct answers defined as the ratio of words suitable for the *shiritori* rule from among all the words that were provided during the task. The graph shows an increase in correct answers as the base level constant  $\beta_i$  increases, suggesting that the effect of learning  $B_i$  increased and the similarity effect  $P_i$  decreased relatively with respect to the ACT-R activation calculation (Eq. 1). In other words, the result indicates that the learning effect can suppress the phonological errors that the similarity effect causes.

Figure 5 shows the change in the rate of correct answers given during the execution of the task. The horizontal axis shows the time spent completing the task, divided into four intervals, while the vertical axis shows the percentage of cor-

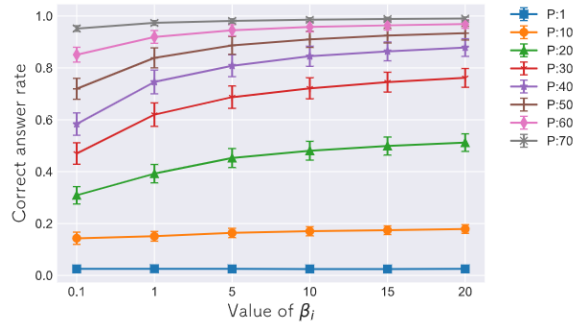


Figure 4: The change in the correct answer rate, as indicated by  $\beta_i$  (error bar is  $\pm \frac{SD}{5}$ )

rect answers in each interval. In the graph, there is no condition showing that the rate of correct answers is increasing, in some conditions, it is decreasing. This indicates that learning does not occur in the task that was set in this study.

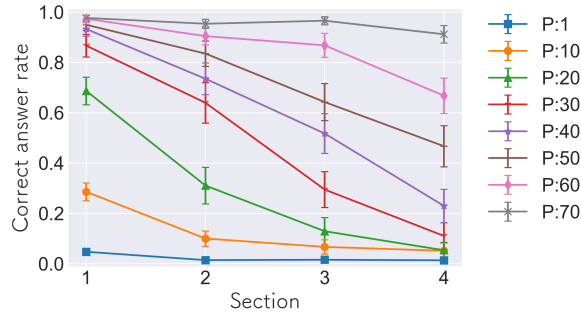


Figure 5: Changes in the rate of correct answers throughout the task (error bar is  $\pm \frac{SD}{5}$ )

In this regard, it is suggested that implemented assistance (repeating the previous answer) has no effect on learning, and we need to explore other forms of support to enhance phonological awareness in the *shiritori* task. In Takahashi (1997)'s study, the most effective assistance came in the form of offering hints that activate semantic knowledge pertaining to the correct answer.

In addition, through individual case observation, the number of times that *shiritori* has continued stands at about 10, and it is possible that a sufficient number of presentations of knowledge for learning may not have been secured. Furthermore, due to the research process' effect with regard to avoiding words that were already provided as answers and words that end with /N/, a problem can be assumed in which only the same words are searched and it is impossible to proceed to the next answer.

## Summary and Future Work

The purpose of this study was to investigate the factors that contribute to the occurrence and suppression of errors during the formation of phonological awareness. For this purpose, we modeled phonological awareness using the cognitive ar-

chitecture ACT-R and performed simulations that manipulated ACT-R parameters corresponding to both innate and experiential factors. As a result, it is suggested that phonological awareness and the errors that occur during its formation process can be modeled through the innate memory retrieval mechanism. Furthermore, the effect of learning may reduce errors by suppressing innate factors. However, the results did not show that learning occurred while performing the tasks set in this study.

This study leaves a number of issues to be addressed. In its simulations, the authors arbitrarily determined values for each parameter. Although we were able to observe a variation in the correct answer rate, the threshold and maximum values need to be verified using further simulation. In addition, this study only showed that the correct answer rate was improved by setting the parameters in advance; it did not show that learning occurred or that the correct answer rate was improved through task performance. It is therefore necessary to further examine the interaction in which the learning effect can be expected to operate. Consideration should also be given to non-phonological aids, such as letters, pictures, and word meanings. In addition, evaluating the constructed model is essential. To ensure that the model's task-related learning process explains the child's language development, it is crucial to correlate it with experimental data from humans. We believe that after setting the learning task, it is possible to examine this by comparing it with the existing human experimental data (Hara, 2001).

## Acknowledgment

This research was supported by JSPS KAKENHI Grant Numbers 18H05068, 20H05560, 20H04996.

## References

- Amano, S., & Kobayashi, T. (2008). *Kihongo database : gogibetsu tangoshimmitsudo*. Gakken plus.
- Anderson, J. R. (2007). How can the human mind occur in the physical universe?
- Baron-Cohen, S. (1997). *Mindblindness: An essay on autism and theory of mind*. MIT press.
- Bloch, B. (1950). Studies in colloquial Japanese IV phonemics. *Language*, 26(1), 86–125.
- Chomsky, N., & Halle, M. (1968). *The sound pattern of English*. Harper & Row New York.
- for Japanese Language, N. I., & Linguistics. (1981). *Yoji / jido no rensogoihyo*. Tokyo Shoseki.
- Hara, K. (2001). Kenjoji ni okeru oninishiki no hattatsu. *The Japanese journal of communication disorders*, 18(1), 10–18.
- Ishida, H., & Ishizaka, I. (Eds.). (2016). *Gengochokakushi no tameno gengohattatsushogaigaku*. Ishiyaku shuppan.
- Kobayashi, H. (2018). *Oninishiki no keisei to kotoba no hattatsu - "kotoba ga osoi" wo kangaeru-*. Kodama shuppan.
- Kubozono, H. (1989). The mora and syllable structure in Japanese: Evidence from speech errors. *Language and Speech*, 32(3), 249–278.

- Kubozono, H. (2000). Kodomo no shiritori to mora no kakutoku. *Bulletin of the Faculty of Letters, Kobe University*(27), 587–602.
- Mätzig, P., Vasishth, S., Engelmann, F., Caplan, D., & Burchert, F. (2018). A computational investigation of sources of variability in sentence comprehension difficulty in aphasia. *Topics in cognitive science*, 10(1), 161–174.
- Mugitani, R., Homae, F., Hiroya, S., Satou, Y., Shirose, A., Tanaka, A., ... Tachiiri, H. (2019). *Kodomo no onsei* (No. 21). Corona Publishing Co., Ltd.
- Nakamura, T., Kojima, C., & Fujiwara, Y. (2015). Kenjohat-tatsu ni okeru onin purosesu no henka. *Journal of rehabilitation sciences, Seirei Christopher University*, 10, 1–13.
- Stahl, S. A., & Murray, B. A. (1994). Defining phonological awareness and its relationship to early reading. *Journal of educational Psychology*, 86(2), 221.
- Taatgen, N. A., & Anderson, J. R. (2002). Why do children learn to say “broke” ? a model of learning the past tense without feedback. *Cognition*, 86(2), 123–155.
- Takahashi, N. (1997). Yoji no kotobaasobi no hattatsu : “shiritori” o kanon isuru joken no bunseki. *The Japanese Journal of Developmental Psychology*, 8(1), 42–52.
- Tomasello, M. (1999). *The cultural origins of human cognition*. Harvard University Press.
- Trubetzkoy, N. S. (1969). Principles of phonology.
- Van Rij, J., Van Rijn, H., & Hendriks, P. (2010). Cognitive architectures and language acquisition: A case study in pronoun comprehension. *Journal of Child Language*, 37(3), 731–766.