

Adults regularize variation when linguistic cues suggest low input reliability

Yiran Chen & Kathryn Schuler*

Abstract. Children regularize inconsistent probabilistic patterns in linguistic input, yet they also acquire and match probabilistic sociolinguistic variation. What factors in the language input contribute to whether children will regularize or match the probabilistic patterns they are exposed to? Here, we test the hypothesis that low input reliability facilitates regularization. As a first step, we asked adult participants to acquire a variable plural marking pattern from a written (Exp 1) and a spoken (Exp 2) artificial language under different conditions, where they were led to believe input was more, or less, reliable. In both experiments, input reliability was manipulated through both information about the speaker (e.g., whether the speaker was likely to make mistakes) and linguistic cues (e.g., typos or pronunciation errors). Results showed that adults regularized the written language more only when they were told the speaker would make mistakes *and* the plural variants resembled typos (Exp 1), whereas they regularized the spoken language more when the plural variants resembled pronunciation errors regardless of the speaker's said reliability in the spoken language. We conclude that input reliability is an important factor that can modulate learners' regularization of probabilistic linguistic input, and that linguistic cues may play a more critical role than top-down knowledge about the speaker. The current study lays down an important foundation for future work exploring whether children are able to incorporate input reliability cues when learning probabilistic linguistic variation.

Keywords. acquisition; variation; regularization; input reliability; psycholinguistics

1. Introduction. Although children inevitably learn language from input, they do not always reproduce the patterns in their input veridically. Under some circumstances, especially when faced with unpredictable variation, children are known to change the language, making it more regular. Researchers have observed this regularization in natural language when children's language input contains inconsistencies, such as in emerging speech communities (e.g., Senghas & Coppola 2001, Kegl, Senghas & Coppola 1999, Kocab, Senghas & Snedeker 2016), in communities where pidgins and creoles are spoken (e.g., Hall 1966, Bickerton 1984, DeGraff 1999), and when acquiring language solely from non-native or late-learning parents (Ross 2001, Ross & Newport 1996, Singleton & Newport 2004). To illustrate, consider Simon, a Deaf child who acquired American Sign Language (ASL) solely from his late-learning parents (Singleton & Newport 2004). Singleton & Newport (2004) examined the family's production of ASL movement morphemes and found that, while Simon's parents produced the correct form only 65% of the time on average, Simon himself produced the correct form 88% of the time — on par with Deaf children his age learning from native-signing parents. The researchers argued that

* The authors would like to acknowledge the University of Pennsylvania and NSF#2043724 for funding. We thank Annalise Kendrick, Tula Childs, Cynthia Gu, Aja Altenhof and Lauren Kim for help with data collection and coding, the audience at Language and Cognition Lab at University of Pennsylvania as well as LSA 2022 for helpful discussion and feedback. Authors: Yiran Chen, University of Pennsylvania (cheny39@sas.upenn.edu), Kathryn Schuler, University of Pennsylvania (kschuler@sas.upenn.edu)

Simon was able to surpass his parents' performance by regularizing his input — boosting his parents' most frequent forms. Children's tendency to regularize inconsistent variation in their language input has been corroborated in laboratory experiments as well (Hudson Kam & Newport 2005, 2009, Austin, Schuler, Furlong & Newport 2021).

However, children do not always regularize probabilistic variation in language. In the field of developmental sociolinguistics, accumulating evidence suggests that children can indeed learn probabilistic variation, particularly if the variation reflects important features of their dialect (Labov 1989, Roberts 1997, 2002, Smith, Durham & Fortune 2007, Smith, Durham & Richards 2013, Miller 2013, Hendricks, Miller & Jackson 2018). For example, Labov (1989) found that 7-year-old David was able to probabilistically drop coronal stops in word final consonant clusters — a phenomenon known to as T/D deletion (e.g., dropping the “t” in “west”, or “d” in “sand”). Importantly, David's T/D deletion patterns obeyed the same constraints that condition the rates of dropping in adult speech: just like his parents, he dropped more often when the following segment was an obstruent compared to a vowel, and when the conversational context was more casual (Labov 1989).

Given robust evidence for both behaviors, one intriguing question follows: when do children regularize probabilistic variation in language, and when do they learn it?

Several hypotheses have been proposed to address this question. One commonly assumed possibility is that only conditioned variation gets reproduced, while unconditioned variation gets regularized. Variation is conditioned when certain factors systematically predict the occurrence of certain variants. For example, the aforementioned T/D deletion is considered conditioned variation, since the phonological environment and social context can make systematic predictions about the rate of deletion. Indeed, most sociolinguistic variables are thought to be examples of conditioned variation (Labov 1989, Roberts 1997, 2002, Smith et al. 2007, 2013, Miller 2013). Further support for the idea that conditioned variation is learned, not regularized, comes from artificial language learning experiments with children when the artificial language contained a conditioning factor, children regularized less (Wonnacott 2011, Hudson Kam 2015, Austin et al. 2021, Samara et al. 2017).

However, this hypothesis does not explain the range of behaviors we observe in the literature. First, there are a few documented circumstances in which children regularize conditioned variation (Hudson Kam 2015, Samara et al. 2017), but learn unconditioned variation (Hendricks et al. 2018). Second, adult learners seem to be readily able to learn and match both conditioned and unconditioned of variation (Austin et al. 2021, Hudson Kam & Newport 2005, 2009, Real & Griffiths 2009, Perfors & Burns 2010, Perfors 2012, Culbertson, Smolensky & Legendre 2012, Smith et al. 2017). More specifically, while children tend to regularize unconditioned variation in experiments, adults exposed to the same language tend to match the probabilities of the variants they are exposed to (Hudson Kam & Newport 2005, 2009, Austin et al. 2021). Lastly, if learners always regularize unconditioned variation and learn conditioned variation, then we would expect language change to emerge only under circumstances in which unconditioned variation is observed. However, this is not the case: conditioned variation in natural language can also undergo change (e.g., Roberts 1997, Kerswill & Williams 2000).

Given these facts, many researchers have embraced the position that, beyond the mere presence of a conditioning factor, the language input itself likely provides additional cues to indicate whether variation should be learned or regularized (Hudson Kam 2015, Johnson & White 2020, Shih 2016, Hendricks et al. 2018). Candidate cues could include the frequency of

the alternating contexts (Shih 2016), input quantity (Hendricks et al. 2018), and whether probabilistic pattern is shared across speakers (Chen & Schuler 2021).

In the current study, we examine another important factor that has been underexplored in this literature: the reliability of the input. If we consider again the circumstances in which children are shown to robustly regularize — children learning from pidgin speakers (Bickerton 1984) or from parents who are late-learning signers (Singleton & Newport 2004) — we find that the input is produced by non-native or late-learning speakers, who might provide subtle cues to indicate they are unreliable language models. Simon’s late-learning parents, for example, not only committed more errors than native signers, but also produced different types of errors (uncommon among native signers) and showed more disfluency than native signers (Ross & Newport 1996, Singleton & Newport 2004).

While these observations suggest that subtle cues to reliability exist in children’s input, several questions remain unanswered. First, are children sensitive to such cues about input reliability? Evidence from the word-learning literature suggests this is indeed the case: children prefer learning words from speakers that they deem more reliable (e.g., Koenig, Clément & Harris 2004, Koenig & Harris 2005, Jaswal & Neely 2006, Kinzler, Corriveau & Harris 2010, Harris & Corriveau 2011). Importantly, such preferences include a preference to learn new words from native rather than non-native speakers (Corriveau, Kinzler & Harris 2013). For example, Corriveau and colleagues (2013) familiarized 3-5 year olds with two speakers: one with a native accent and one with a nonnative accent. During the novel label learning task that followed, the children preferred to ask the native speaker for the label, and endorsed the labels provided by the native speaker over the ones provided by nonnative speakers (Corriveau et al. 2013). Further, work with infants has shown that babies as young as 5-6 months can distinguish nonnative speakers based on speech cues (Kinzler et al. 2007).

Second, when faced with cues suggesting that input reliability is low, do learners respond by regularizing more? Initial evidence suggests that this is true, at least for adult learners. Perfors (2016) demonstrated that adults regularized more when they believed the input to be unreliable. In one experiment, adults were asked to learn novel labels for objects consisting of stem-affix pairs. Each object had a consistent stem, but which affix occurred was inconsistent: one dominant affix appeared 60% of the time (regardless of stem), and each of four noise affixes occurred 10% of the time. When asked to label objects at test, adults regularized more — producing a single affix more than 60% of the time — if they had been explicitly told their language input came from a previous participant who might have made mistakes, particularly in a condition where the non-dominant affixes were designed to resemble typos of the dominant affix.

Taken together, the evidence reviewed above suggests (1) that children are sensitive to input reliability signaled by linguistic cues during word learning and (2) that adult learners regularize more when they perceive their input to be unreliable. However, in order to provide direct support for our hypothesis — that low input reliability contributes to children’s regularization behavior in pidgin creole or non-native parent circumstances — we need to show (1) that learners are sensitive to the input reliability during *rule learning* and (2) that learners’ regularization behavior is sensitive to subtle cues in the language input that likely exist in input to children. Recall that the Perfors’ (2016) study signaled low reliability by telling learners that the input was likely to contain mistakes and using noise affixes that resembled typos of the dominant form, neither of which are available cues in children’s input during natural language learning. Therefore, in the current study, we begin by replicating Perfors (2016) in a more complex rule-learning context

(Experiment 1) and then test whether learners indeed regularize more when input reliability is signaled by two cues more likely to be present in children’s natural language input: identity as nonnative speaker and variants’ resemblance to pronunciation errors (Experiment 2).

2. Experiment 1. To begin to ask whether low input reliability facilitates learners’ regularization in natural language learning, we begin by replicating the Perfors (2016) experiment in a more complex context that is common in child language acquisition. Specifically, we ask whether adult learners are more likely to regularize unpredictable variation if they are led to believe that their linguistic input may contain mistakes, in the context of learning a language with a variable plural marking pattern.

2.1. METHODS.

2.1.1. PARTICIPANTS. We recruited 134 adults (aged 18-46 years) via Prolific (www.prolific.co) to participate in Experiment 1. Eligible participants were native English speakers who lived in the United States, had unimpaired (or corrected) hearing and vision, and had an acceptance rate over 85% on the Prolific platform. Participants received \$8.35/hour for the approximately 20-minute experiment.

2.1.2. PROCEDURE. Participants learned an artificial language with inconsistent plural marking through 180 exposure trials. On each exposure trial, participants saw a written sentence and were asked to guess which of two animal images the sentence described (see Figure 1). After submitting a guess, participants were shown the correct answer as feedback. Of the 180 exposure trials, 60 were singular (10 for each noun) and 120 were plural (20 for each noun). To encourage participants to learn the plural marking more implicitly, the two image alternatives on each trial only differed in which noun was depicted, never in grammatical number. That is, participants selected between two singular images on singular trials (e.g., one duck vs. one cow), and two plural images on plural trials (e.g., two sheep vs. two pigs). Further, to prevent learners from forming an association between a noun and a specific image exemplar, all exemplar images of each animal type (10 singular and 20 plural) are unique.

Sentences in the language followed the basic structure “gentif + {noun} + {plural marker}”, where “gentif” was a verb meaning “there is/are”; {noun} was one of six novel nouns (*mawg*, *geed*, *spad*, *daffin*, *flugat*, *clidam*), which were randomly mapped onto six different farm animals (duck, pig, cow, horse, chicken, or sheep); and {plural marker} was one of five novel plural markers, which occurred only on plural trials and were distributed according to the experimental condition (see Table 1). Importantly, plural marking in the language was inconsistent, such that learners could not predict which of the 5 plural markers would occur on a given trial. Instead, the plural marking pattern was probabilistic: one plural marker was *dominant*, occurring on 60% of all plural trials, and the remaining four were *non-dominant*, each occurring on 10% of all plural trials. After every 45 exposure trials, we tested participants’ learning with twelve test trials — two for each animal type. On each test trial, participants were presented with a novel plural image and asked to type a sentence to describe the image in the artificial language. We expected learners to accurately produce the basic plural sentence structure (gentif + noun + plural marker) and the correct noun. What was of crucial interest was how learners produced the inconsistent plural marking at test.



Figure 1. Sample singular and plural exposure trials (left) and sample test trial (right). The green frame represents the target of the trial.

2.1.3. CONDITIONS. Our main question is whether participants would be more likely to regularize unpredictable variation if they are led to believe that the input may contain unpredictable errors. Following Perfors (2016), we manipulated the perceived reliability of the input in two ways. First, we manipulated the cover story about the source of the language input such that participants in one condition would be more likely to expect the language input to contain mistakes. Participants in the *Experimenter* conditions were told that the input came from an experimenter, while participants in the *Participant* conditions were told that the input came from a previous participant, who was under time pressure and might therefore make mistakes.

Second, we manipulated the characteristics of the non-dominant markers such that participants in one condition would be more likely to interpret these markers as mistakes. Specifically, non-dominant markers in the *Distinct* condition were designed to be distinct from each other and the dominant marker. Non-dominant markers in the *Typo* condition, on the other hand, were designed to look like plausible typos of the dominant marker (i.e., one letter was omitted, reduplicated, or substituted by an adjacent letter on a QWERTY keyboard). To control for the possibility that a particular marker might be easier or harder to learn (or interpret as a typo), participants were randomly assigned to one of 5 dominant markers (*ka, po, su, ti, je*). In the *Distinct* condition, the remaining four markers served as the non-dominant markers; in the *Typo* condition, non-dominant markers were perturbations of whatever dominant marker was randomly assigned (see Table 1).

Dominant Marker (60%) Non dominant markers (10% each)

<i>Distinct & Typo</i>	<i>Distinct</i>	<i>Typo</i>
ka	po, su, ti, je	ja, kq, a, kka
po	ka, su, ti, je	lo, pi, o, ppo
su	ka, po, ti, je	wu, sy, u, ssu
ti	ka, po, su, je	ri, tu, i, tti
je	ka, po, su, ti	ne, jw, e, jje

Table 1. All plural markers in Experiment 1

To summarize, the experiment followed a 2×2 design, crossing Input Source (Experimenter or Participant) with Marker Type (Distinct or Typo). At the start of the experiment, participants

were randomly assigned to one of the four resulting conditions: Experimenter-Distinct, Experimenter-Typo, Participant-Distinct or Participant-Typo.

2.1.4. ANALYSIS. Our primary question is whether learners would be more likely to regularize the inconsistent plural markers if they believed the input might contain mistakes (i.e., in our Participant and/or Typo conditions). However, we wanted to first establish whether participants in all conditions learned the basic structure and vocabulary of the language equally well. To determine how well participants learned, we first coded participant test trials as either correct (1) or incorrect (0), depending on whether they produced the verb (“gentif”) followed by the correct noun. Then, we built a logistic mixed-effects model predicting whether a trial was correct with Input Source (Experimenter vs. Participant, simple coded), Marker Type (Distinct vs Typo, simple coded), and their interaction as fixed effects, with random intercepts for participants and random slopes of nouns for each participant.

Next, we analyzed learners’ regularization of the plural markers when they produced the correct verb and nouns. We used two different metrics to determine the extent to which participants regularized. First, we analyzed how often participants produced their most frequently used marker (max marker), allowing us to detect regularization toward a single marker (whether it was dominant in the participant’s input). To accomplish this, we built a logistic mixed-effect model, predicting each participant’s max marker by the Input Source (Experimenter or Participant), the Marker Type (Distinct or Typo), and their interaction (as fixed effects), with random intercepts for participants. Second, we analyzed the entropy of participant marker choices, allowing us to detect whether participants regularized in other ways, beyond simply boosting their use of a single form. Entropy provides a measure of how regular or predictable a participant’s marker productions are. For example, in our artificial language, an entropy of 0 would indicate a participant’s marker use was completely predictable (i.e., the participant used one plural marker consistently on every trial), while higher entropy would indicate their marker use was more unpredictable. To detect any entropy differences by condition, we ran a separate simple linear model predicting entropy by the same predictors in the previous model: the Input Source (Experimenter or Participant), the Marker Type (Distinct or Typo), and their interaction.

2.2. RESULTS AND DISCUSSION. Turning first to whether participants learned the basic structure and vocabulary of the language, we found that learners produced the verb followed by the correct noun on 64.97% of all test trials. There were no significant effects of Input Source ($p = 0.591$), Marker Type ($p = 0.386$), or their interaction ($p = 0.625$), suggesting that participants learned the language equally well in all conditions

Having established that participants’ learning did not differ across conditions, we turn next to participants’ regularization behavior. Our logistic mixed-effects model predicting participants’ max marker use revealed no significant effects of either Input Source (Est = 0.522, SE = 0.311, $p = 0.093$), or Marker Type (Est = 0.413, SE = 0.311, $p = 0.183$). In other words, participants did not regularize more when instructed that the input came from another participant who likely made mistakes, nor when the non-dominant markers in their input looked like typos. However, the interaction between Input Source and Marker Type was significant (Est = 1.377, SE = 0.622, $p = 0.026$). As shown in Figure 2a, participants only regularized more in the Participant-Typo condition, in which they were told the input was from former participants *and* the non-dominant markers resembled typos.

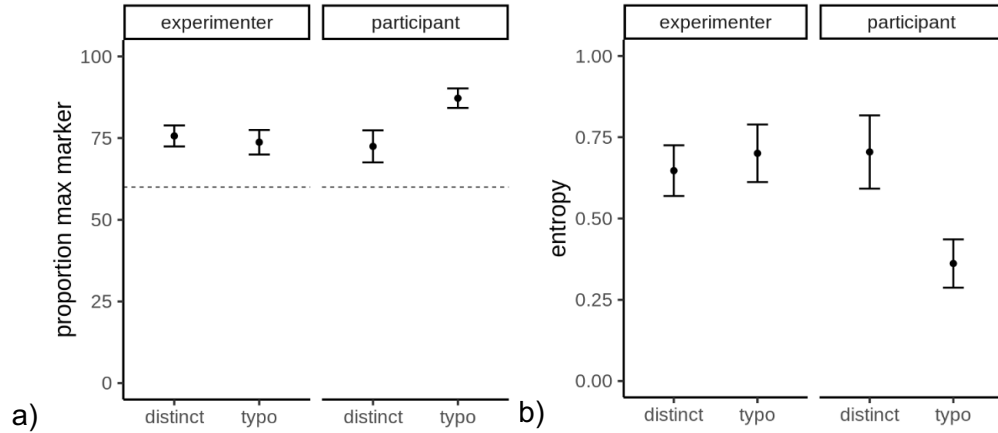


Figure 2. a) Participants' mean maximum marker proportion at test in Experiment 1. The dashed line is the proportion of the dominant marker in the input (60%). b) The mean entropy of participants' plural markers produced at test in Experiment 1. Error bars represent standard errors.

These results were reflected in our entropy analysis as well. Our linear model predicting entropy of participants' marker production at test revealed the same pattern: no main effect of Input Source or Marker Type, but a significant interaction between them ($Est = -0.396$, $SE = 0.179$, $p = 0.028$). We can see this interaction reflected in Figure 2b: participants in the Participant-Typo condition had significantly lower entropy, the condition in which they were told input was from a former participant *and* the non-dominant markers resembled typos.

Thus, like Perfors (2016), we found that adults regularized more when they were told input speakers were likely to make mistakes *and* when markers looked like plausible typos. However, unlike Perfors (2016), we did not find a significant main effect of Input Source. In the current study, information about the input speaker alone was not enough to influence adults' regularization behavior. We attribute this difference to the increased complexity in our study, in which adults were learning an inconsistent plural marking rule, rather than stem-affix pairs. As shown in Figure 2a (left), in this more complicated task, learners showed some boosting of their max marker above the input probability, even in the Experimenter conditions. In sum, our learners only regularized more when they received converging evidence of low input reliability, conveyed by both linguistic properties of the markers and information about the speaker.

3. Experiment 2. Experiment 1 showed that input reliability modulates learners' tendency to regularize, but only when the expectation that the speaker will make mistakes is confirmed by linguistic properties of the variants (e.g., typos). However, we are still several steps away from concluding that input reliability contributes to regularization in natural language learning. First, crucially, children do not learn natural languages in written form. Written compared to spoken or signed modalities might demand fundamentally different learning mechanisms. Second, children can only take advantage of cues that are available in their input, and neither the Experimenter v. Participant contrast nor typos are cues that children encounter in natural language. Bearing this in mind, in Experiment 2, we sought to test how input reliability modulates learners' regularization when learning a spoken language containing unpredictable variation. Importantly, this time, we manipulated input reliability with cues that are available in children's natural language input: whether the input was provided by a native or non-native speaker of the language (rather than Experimenter v. Participant) and whether the noise markers are plausible pronunciation errors (rather than typos). Specifically, we ask whether learners regularize more

when the variation in markers could be pronunciation errors, and when the input is known to contain mistakes (here because the speaker they are learning from is non-native).

3.1. METHOD.

3.1.1. PARTICIPANTS. We recruited another 79 adults (age 18-46) via Prolific (www.prolific.co) to participate in Experiment 2. Eligible participants were native English speakers who lived in the United States, had unimpaired (or corrected) hearing and vision, and had an acceptance rate over 85% on the Prolific platform. Participants received \$10.56/hour for the approximately 35-minute experiment.

3.1.2. PROCEDURE. The procedure was identical to that of Experiment 1, except the experiment was conducted with spoken rather than written language. This meant that, during exposure, participants heard the sentences rather than reading them, and were instructed to describe the pictures out loud rather than typing them at test.

3.1.3. CONDITIONS. Our main question in Experiment 2 was the same as in Experiment 1: would participants be more likely to regularize unpredictable variation if they believe the input may contain mistakes? Crucially, however, we wanted to make sure that our input reliability manipulations — Input Source and Marker Type — were conveyed to participants in ways that could plausibly occur in a child's natural language input. To accomplish this, we made two changes to our conditions. First, we changed the Input Source manipulation from Experimenter vs. Participant used in Experiment 1, to Native vs. Non-native. Recall that, in natural cases where regularization is reported, the input speakers are often nonnative or late-learning speakers of the language (e.g., Bickerton 1984, Singleton & Newport 2004). And, in experimental settings, children are shown to distinguish native and nonnative speakers as early as 5-6 month old (Kinzler et al. 2007) and disprefer the latter as models for language learning by age 3-5 years (Corriveau et al. 2013). To convey the fluency of the language model in our experiment, participants were told that the input speaker, Mari, was either a Native or a Non-native speaker of the language. To make our intention clear, participants in the Native condition were told Mari learned this language as a baby, spoke it for her entire life, and almost never made mistakes, while participants in the Non-native conditions were told that Mari had just learned the language a month ago and made a lot of mistakes. Following these backstories, participants were asked to rate how well they believed Mari spoke the language on a sliding scale from 0 - 100.

Second, given the switch to spoken language, we manipulated input reliability by designing non-dominant markers that resembled possible pronunciation errors of the dominant marker (rather than typos as in Experiment 1). The entire inventory of markers is shown in Table 2. The markers in Distinct conditions were exactly the same as Experiment 1. In the Pronunciation Error conditions, in our best effort to ensure that the non-dominant markers could be perceived as plausible pronunciation errors of the dominant marker, we constructed the markers with the following rules: 1) The first variant differs from the dominant marker with one feature: voicing of the consonant. 2) The second differs from the dominant marker in that it turns the stop consonant into a fricative closest in place of articulation found in English sound inventory. 3) The third differs from the dominant marker in only one vowel feature - height, frontness, or roundedness. 4) The fourth differs from the target via labialization (for back-vowel targets) or palatalization (for front-vowel targets).

Dominant Marker (60%) Non dominant markers (10% each)

<i>Distinct & Speech Error</i>	<i>Distinct</i>	<i>Pronunciation Error</i>
ka	po, su, ti, je	ga, ha, ko, kwa
po	ka, su, ti, je	bo, fo, pa, pwo
su	ka, po, ti, je	zu, thu, syu, swo
ti	ka, po, su, je	di, tsi, te, tje
je	ka, po, su, ti	che, ye, ji, jie

Table 2. All plural markers in Experiment 2

In sum, as in Experiment 1, Experiment 2 has a 2×2 design, crossing Input Source with Marker Type. Therefore, the four conditions in Experiment 2 were Native Distinct, Native Pronunciation Error, Non-native Distinct, and Non-native Pronunciation Error.

3.2. ANALYSIS. The analysis was conducted as in Experiment 1, with an additional coding process. Two research assistants listened to recordings of each participants' production at the test and coded which marker the participant produced: the dominant marker, one of the noise markers, other or null. Inter-coder reliability was 90.05% for nouns and 94.59% for markers, and we included only trials on which the two transcribers agreed in subsequent analyses.

3.3. RESULTS AND DISCUSSION. First, we looked at participants' overall learning of the spoken language across different conditions. As in Experiment 1, accuracy on the basic structure and vocabulary of the language did not differ by Input Source ($p = 0.349$), Marker Type ($p = 0.995$) or their interaction ($p = 0.888$).

Recall that we asked participants to rate how well Mari speaks the language after reading the cover story. We analyzed this proficiency rating to determine whether learners indeed perceived Mari to have different language proficiency based on condition. Results showed that participants who were told Mary was a native speaker rated her proficiency consistently high (nearly all at or very close to 100), while those who were told that Mary was not a native speaker rated her proficiency much lower (Figure 3). Thus, participants paid attention to the cover story and had different beliefs about the input speaker's proficiency in the language in the two Input Source conditions.

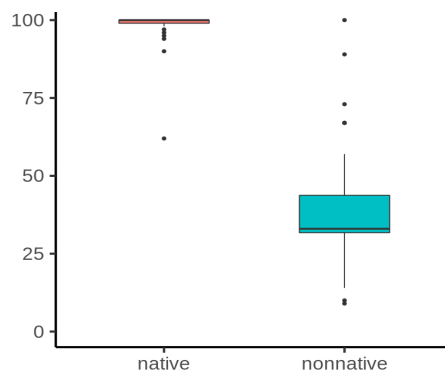


Figure 3. Participants' mean rating score of Mari's proficiency in the artificial language

Although participants had different beliefs about Mari’s proficiency, Input Source did not affect how they learned the inconsistent plural-marking rule: participants were not significantly more likely to use their max marker when told their input came from a non-native speaker (Est = -0.447, SE = 0.373, $p = 0.232$) (Figure 4a). Instead, we found a robust main effect of Marker Type (Est = 1.168, SE = 0.375, $p = 0.002$) — participants were more likely to regularize to their max marker in the Pronunciation Error conditions — with no significant interaction with Input Source (Est = 0.130, SE = 0.744, $p = 0.862$). The analysis of entropy matched the analysis on max marker use: participants’ marker use at test had lower entropy, and was thus more regular, when the non-dominant markers were plausible pronunciation errors in both Input Source conditions (Est = -0.354, SE = 0.106, $p = 0.001$, Figure 4b). Input Source did not have a significant effect on participants’ entropy (Est = 0.118, SE = 0.106, $p = 0.272$), nor did it interact with Marker Type (Est = -0.197, SE = 0.213, $p = 0.357$). In sum, participants in Experiment 2 were more likely to regularize when the markers sounded like errors, regardless of their language model (native or non-native).

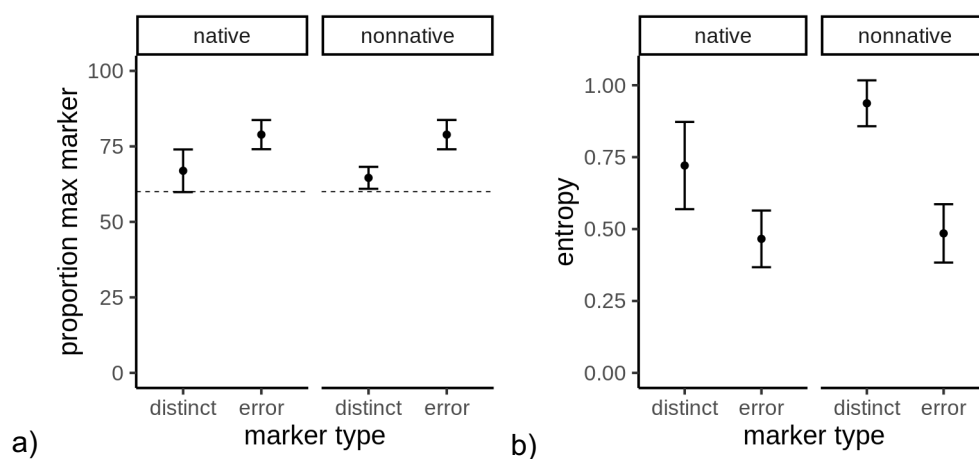


Figure 4. a) Participants’ mean maximum marker proportion at test in Experiment 2. The dashed line is the proportion of the dominant marker in the input (60%). b) Participants’ mean entropy of plural markers produced at test. Error bars represent standard errors.

Lastly, we paid special attention to the marker participants used in the Pronunciation Error conditions. As we introduced before, we designed the four non-dominant markers to have different types of phonological similarity with the dominant marker. On average, participants used one of the non-dominant markers on 12.12% of test trials in the Native Pronunciation Error condition, and 17.17% in the Non-native Pronunciation Error condition. When they produced a non-dominant marker at test, which type of non-dominant markers were they most likely to use? As shown in Figure 5, when participants produced non-dominant markers at test, they most often produced the fricative change variant (e.g., ha when ka was the dominant marker) and the voicing change variant (e.g., ga when ka was the dominant marker). However, given that participants produced these non-dominant markers so infrequently, we caution against drawing any conclusions.

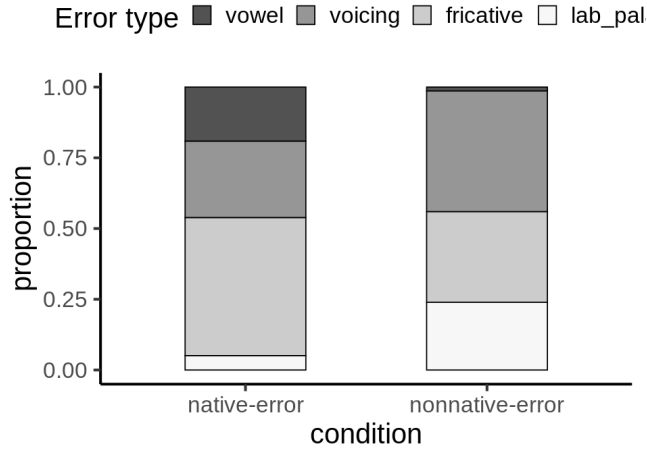


Figure 5. Proportion of each type of non-dominant marker out of all non-dominant markers produced at test in Experiment 2. Vowel stands for the variants that only differ from the dominant marker in vowel features. Voicing stands for the variants that only differ in one voicing feature. Fricative refers to the variants that substituted one stop consonant in the dominant marker with a fricative. Labialization/Palatalization refers to the variants that only differ from the dominant marker with an added labialization/palatalization.

4. General discussion and conclusion. Taken together, our two experiments show that adult learners’ regularization behavior during rule learning is indeed modulated by input reliability. In both experiments, adults learned an artificial language with an inconsistent plural marking rule, in which one dominant plural marker appeared 60% of the time and four other non-dominant markers each appeared 10% of the time. In Experiment 1, in which participants learned a written language, learners only regularized more if they were told their language model was likely to make mistakes *and* the non-dominant markers resembled mistakes (in the form of possible typos of the dominant marker). In Experiment 2, in which participants learned a spoken language, adult learners regularized more when the non-dominant markers resembled mistakes (here in the form of plausible pronunciation errors of the dominant marker) *regardless* of whether they were told their language model was likely to make mistakes (whether the input speaker was a native or nonnative speaker).

Thus, across these two experiments, we found that Input Source — knowledge about whether the language model was likely to make mistakes — was a relatively weak predictor of learners’ regularization behavior. Especially in Experiment 2, regularization did not depend on whether learners were told they were learning from a native or nonnative speaker. This is perhaps surprising given previous findings that listeners respond to nonnative speech differently on semantic (Gibson et al. 2017), syntactic (Hanulíková et al. 2012) and pragmatic (Fairchild, Mathis & Papafragou 2020) levels, and that learners (including children) show a preference for learning new words from native rather than nonnative speakers (e.g., Corriveau et al. 2013).

One possible explanation for this unexpected finding is simply that our native vs. nonnative manipulation was not as strong as in these prior studies. Indeed, the only cue to language fluency in our study was information provided in the backstory — we did not include any speech-level cues such as different accents or disfluencies (e.g., Gibson et al., 2017, Hanulíková et al., 2012, Corriveau et al 2013). However, this explanation would suggest that the backstories we provided did not cause learners to form different assumptions about the input speaker’s language fluency, which is not likely to be the case. First, in many studies, differential attitude towards the speaker was successfully elicited based only on a backstory (e.g., Fairchild et al. 2020). Second, in the

current study, participants *did* form different beliefs about the input speaker's proficiency in the language, as indicated by the dramatically different proficiency ratings for the speaker across conditions (see Figure 4).

How, then, should we account for this lack of difference? One possibility is that, as novice learners of the artificial language themselves, our participants were not confident enough to attribute the inconsistent variation they observed to mistakes made by a non-native speaker. While previous studies found that participants would indeed accommodate for or form differential learning attitudes towards nonnative speakers (Gibson et al. 2017, Hanulíková et al. 2012, Corriveau et al. 2013, Kinzler et al. 2011, Fairchild et al. 2020), these studies involved participants making decisions about nonnative speakers of their own native language — a language in which they are highly fluent experts.

Another possibility is that learners are more sensitive to local rather than global cues during rule learning (an arguably implicit learning task). In the current study, we provided information about the input speaker's language fluency only once, at the beginning of the experiment. While our rating test showed that participants used this information to form judgements about the speaker's language fluency immediately after, we do not know whether learners held this information in mind throughout the experiment, nor whether they used it to inform their rule learning. If they did not, perhaps a more effective manipulation would be to provide a local cue to the speaker's language fluency on every learning trial (e.g., speech-level cues). In the future, we hope to explore whether providing characteristics of the input speaker's language fluency *during learning* — such as disfluency and/or accent — would affect learners' regularization behavior.

While we found Input Source — whether the speaker was likely to make mistakes — to be a relatively weak cue in both experiments, we found Marker Type to be a strong predictor of regularization in both experiments. Recall that, in Experiment 1, participants regularized more when the non-dominant markers resembled typos, *only if* they were also told the language input came from a previous participant who may have made mistakes. In Experiment 2, participants regularized more if the non-dominant markers resembled pronunciation errors, *regardless of* whether they were told their language input came from a native or nonnative speaker. These results further support the notion that local linguistic properties of the variants — the fact that they can be interpreted as mistakes — are more salient cues to input reliability and are more likely to lead to changes in learners' regularization behavior. However, what remains less clear is through what mechanism having error-resembling variants increases regularization.

Admittedly, with the current design, we cannot differentiate many plausible accounts. One possibility, the *perception account*, is that participants were more likely to mishear the error-resembling variants and therefore not register them as being different from the dominant marker. Another, the *discarding account*, is that, though correctly perceived, the error-resembling variants were more likely to be “discarded” by the learners and not taken as part of the input at all. A third possibility, the *inference account*, is that, though registered in representation, participants may have inferred that the error-resembling markers were speaker mistakes and then “corrected” them to the dominant marker in their representation of the input. In the future, we hope to design finer-grained experiments to differentiate these possible mechanistic accounts.

In sum, our current study shows that subtle linguistic cues about input reliability can lead adult learners to be more likely to regularize an inconsistent plural marking rule. This provides support for the hypothesis that the reliability of the input may play a role in explaining when learners regularize variation and when they learn and match it. The current work provides a

foundation to examine whether input reliability also modulates regularization behavior in children.

References

- Austin, Alison C., Kathryn D. Schuler, Sarah Furlong & Elissa L. Newport. 2021. Learning a language from inconsistent input: Regularization in child and adult learners. *Language Learning and Development*. 1–29. <https://doi.org/10.1080/15475441.2021.1954927>.
- Bickerton, Derek. 1984. The language bioprogram hypothesis. *Behavioral and Brain Sciences* 7(2). 173–188. <https://doi.org/10.1017/S0140525X00044149>.
- Chen, Yiran & Kathryn D. Schuler. 2021. The role of shared variability on the acquisition of variable rule. Paper presented at the Boston University Conference on Language Development 46.
- Corriveau, Kathleen H., Katherine D. Kinzler & Paul L. Harris. 2013. Accuracy trumps accent in children's endorsement of object labels. *Developmental Psychology* 49(3). 470. <https://doi.org/10.1037/a0030604>.
- Culbertson, Jennifer, Paul Smolensky & Géraldine Legendre. 2012. Learning biases predict a word order universal. *Cognition* 122 (3). 306–329. <https://doi.org/10.1016/j.cognition.2011.10.017>.
- DeGraff, Michel. 1999. *Language creation and language change: Creolization, diachrony, and development*. Cambridge, MA: MIT Press.
- Fairchild, Sarah, Ariel Mathis & Anna Papafragou. 2020. Pragmatics and social meaning: Understanding under-informativeness in native and non-native speakers. *Cognition* 200. 104171. <https://doi.org/10.1016/j.cognition.2019.104171>.
- Gibson, Edward, Caitlin Tan, Richard Futrell, Kyle Mahowald, Lars Konieczny, Barbara Hemforth & Evelina Fedorenko. 2017. Don't underestimate the benefits of being misunderstood. *Psychological Science* 28(6). 703–712. <https://doi.org/10.1177/0956797617690277>.
- Hall, Robert A. 1967. Pidgin and creole languages. *Journal of the American Oriental Society* 87(2).
- Hanulíková, Adriana, Petra M. Van Alphen, Merel M. Van Goch & Andrea Weber. 2012. When one person's mistake is another's standard usage: The effect of foreign accent on syntactic processing. *Journal of Cognitive Neuroscience* 24(4). 878–887. https://doi.org/10.1162/jocn_a_00103.
- Harris, Paul L. & Kathleen H. Corriveau. 2011. Young children's selective trust in informants. *Philosophical Transactions of the Royal Society B: Biological Sciences* 366 (1567). 1179–1187. <https://doi.org/10.1098/rstb.2010.0321>.
- Hendricks, Alison Eisel, Karen Miller & Carrie N. Jackson. 2018. Regularizing unpredictable variation: Evidence from a natural language setting. *Language Learning and Development* 14 (1). 42–60. <https://doi.org/10.1080/15475441.2017.1340842>.
- Hudson Kam, Carla L. & Elissa L. Newport. 2009. Getting it right by getting it wrong: When learners change languages. *Cognitive Psychology* 59(1). 30–66. <https://doi.org/10.1016/j.cogpsych.2009.01.001>.
- Hudson Kam, Carla L. & Elissa L. Newport. 2005. Regularizing unpredictable variation: The roles of adult and child learners in language formation and change. *Language Learning and Development* 1(2), 151–195. <https://doi.org/10.1080/15475441.2005.9684215>.

- Hudson Kam, Carla L. 2015. The impact of conditioning variables on the acquisition of variation in adult and child learners. *Language* 91(4). 906–937. <https://doi.org/10.1353/lan.2015.0051>.
- Jaswal, Vikram K. & Leslie A. Neely. 2006. Adults don't always know best: Preschoolers use past reliability over age when learning new words. *Psychological Science* 17(9). 757. <https://doi.org/10.1111/j.1467-9280.2006.01778.x>.
- Johnson, Elizabeth K. & Katherine S. White. 2020. Developmental sociolinguistics: Children's acquisition of language variation. *Wiley Interdisciplinary Reviews: Cognitive Science* 11(1). e1515. <https://doi.org/10.1002/wcs.1515>.
- Kegl, Judy, Ann Senghas & Marie Coppola. 1999. Creation through contact: Sign language emergence and sign language change in Nicaragua. In Michel DeGraff (eds.) *Language creation and language change: Creolization, diachrony, and development*. 179–237. Cambridge, MA: MIT Press.
- Kerswill, Paul & Ann Williams. 2000. Creating a new town koine: Children and language change in Milton Keynes. *Language in Society* 29 (1). 65–115. <https://doi.org/10.1017/S0047404500001020>.
- Kinzler, Katherine D., Emmanuel Dupoux & Elizabeth S. Spelke. 2007. The native language of social cognition. *Proceedings of the National Academy of Sciences* 104(30). 12577–12580. <https://doi.org/10.1073/pnas.0705345104>.
- Kinzler, Katherine D, Kathleen H. Corriveau & Paul L. Harris. 2011. Children's selective trust in native-accented speakers. *Developmental Science* 14(1). 106–111. <https://doi.org/10.1111/j.1467-7687.2010.00965.x>.
- Kocab, Annemarie, Ann Senghas & Jesse Snedeker. 2016. The emergence of temporal language in Nicaraguan Sign Language. *Cognition* 156. 147–163. <https://doi.org/10.1016/j.cognition.2016.08.005>.
- Koenig, Melissa A., Fabrice Clément & Paul L. Harris. 2004. Trust in testimony: Children's use of true and false statements. *Psychological Science* 15(10), 694–698. <https://doi.org/10.1111/j.0956-7976.2004.00742.x>.
- Koenig, Melissa A. & Paul L. Harris. 2005. Preschoolers mistrust ignorant and inaccurate speakers. *Child Development* 76(6). 1261–1277. <https://doi.org/10.1111/j.1467-8624.2005.00849.x>.
- Labov, William. 1989. The child as linguistic historian. *Language Variation and Change* 1(1). 85–97. <https://doi.org/10.1017/S0954394500000120>.
- Miller, Karen. 2013. Acquisition of variable rules: /s/-lenition in the speech of Chilean Spanish-speaking children and their caregivers. *Language Variation and Change* 25(3). 311–340. <https://doi.org/10.1017/S095439451300015X>.
- Prefors, Amy & Nicholas Burns. 2010. Adult language learners under cognitive load do not over-regularize like children. *Proceedings of the Annual Meeting of the Cognitive Science Society*. 32(32).
- Prefors, Amy. 2012. When do memory limitations lead to regularization? An experimental and computational investigation. *Journal of Memory and Language* 67(4). 486–506. <https://doi.org/10.1016/j.jml.2012.07.009>.
- Perfors, Amy. 2016. Adult regularization of inconsistent input depends on pragmatic factors. *Language Learning and Development* 12(2). 138–155. <https://doi.org/10.1080/15475441.2015.1052449>.

- Real, Florencia & Thomas L. Griffiths. 2009. The evolution of frequency distributions: Relating regularization to inductive biases through iterated learning. *Cognition* 111(3). 317–328. <https://doi.org/10.1016/j.cognition.2009.02.012>.
- Roberts, Julie. 1997. Hitting a moving target: Acquisition of sound change in progress by Philadelphia children. *Language Variation and Change* 9(2). 249–266. <https://doi.org/10.1017/S0954394500001897>.
- Roberts, Julie. 2002. Child language variation. In J. K. Chambers, Peter Trudgill & Natalie Schilling-Estes (eds.), *The handbook of language variation and change*, 328–333. Wiley Online Library. <https://doi.org/10.1002/9780470756591.ch13>.
- Ross, Danielle S. & Elissa L. Newport. 1996. The development of language from non-native linguistic input. *Proceedings of the 20th annual Boston University Conference on Language Development* 2. 634–645.
- Ross, Danielle S. 2001. Disentangling the nature-nurture interaction in the language acquisition process: Evidence from deaf children of hearing parents exposed to non-native input. Rochester, NY: University of Rochester dissertation.
- Samara, Anna, Kenny Smith, Helen Brown & Elizabeth Wonnacott. 2017. Acquiring variation in an artificial language: Children and adults are sensitive to socially conditioned linguistic variation. *Cognitive Psychology* 94. 85–114. <https://doi.org/10.1016/j.cogpsych.2017.02.004>.
- Senghas, Ann & Marie Coppola. 2001. Children creating language: How Nicaraguan Sign Language acquired a spatial grammar. *Psychological Science* 12(4). 323–328. <https://doi.org/10.1111/1467-9280.00359>.
- Shin, Naomi L. 2016. Acquiring constraints on morphosyntactic variation: Children’s Spanish subject pronoun expression. *Journal of Child Language* 43(4). 914–947. <https://doi.org/10.1017/S0305000915000380>.
- Singleton, Jenny L. & Elissa L. Newport. 2004. When learners surpass their models: The acquisition of American sign language from inconsistent input. *Cognitive Psychology* 49(4). 370–407. <https://doi.org/10.1016/j.cogpsych.2004.05.001>.
- Smith, Jennifer, Mercedes Durham & Liane Fortune. 2007. “Mam, my trousers is fa’in doon!”: Community, caregiver, and child in the acquisition of variation in a Scottish dialect. *Language Variation and Change* 19(1). 63–99. <https://doi.org/10.1017/S0954394507070044>.
- Smith, Jennifer, Mercedes Durham & Hazel Richards. 2013. The social and linguistic in the acquisition of sociolinguistic norms: Caregivers, children, and variation. *Linguistics* 51(2). 285–324. <https://doi.org/10.1515/ling-2013-0012>.
- Smith, Kenny, Amy Perfors, Olga Fehér, Anna Samara, Kate Swoboda & Elizabeth Wonnacott. 2017. Language learning, language use and the evolution of linguistic variation. *Philosophical Transactions of the Royal Society B: Biological Sciences* 372(1711). 20160051. <https://doi.org/10.1098/rstb.2016.0051>.
- Wonnacott, Elizabeth. 2011. Balancing generalization and lexical conservatism: An artificial language study with child learners. *Journal of Memory and Language* 65(1). 1–14. <https://doi.org/10.1016/j.jml.2011.03.001>.