

Do children use language structure to discover the recursive rules of counting?



Rose M. Schneider^{a,*}, Jessica Sullivan^b, Franc Marušič^d, Rok Žaucer^d,
Priyanka Biswas^c, Petra Mišmaš^d, Vesna Plesničar^d, David Barner^{a,c}

^a Psychology Department, University of California, San Diego, United States

^b Department of Psychology, Skidmore College, United States

^c Department of Linguistics, University of California, San Diego, United States

^d Center for Cognitive Science of Language, University of Nova Gorica, Slovenia

ARTICLE INFO

Keywords:

Cross-linguistic

Count list

Successor function

Natural number concepts

Number acquisition

Conceptual development

ABSTRACT

We test the hypothesis that children acquire knowledge of the successor function — a foundational principle stating that every natural number n has a successor $n + 1$ — by learning the productive linguistic rules that govern verbal counting. Previous studies report that speakers of languages with less complex count list morphology have greater counting and mathematical knowledge at earlier ages in comparison to speakers of more complex languages (e.g., Miller & Stigler, 1987). Here, we tested whether differences in count list transparency affected children's acquisition of the successor function in three languages with relatively transparent count lists (Cantonese, Slovenian, and English) and two languages with relatively opaque count lists (Hindi and Gujarati). We measured 3.5- to 6.5-year-old children's mastery of their count list's recursive structure with two tasks assessing productive counting, which we then related to a measure of successor function knowledge. While the more opaque languages were associated with lower counting proficiency and successor function task performance in comparison to the more transparent languages, a unique within-language analytic approach revealed a robust relationship between measures of productive counting and successor knowledge in almost every language. We conclude that learning productive rules of counting is a critical step in acquiring knowledge of recursive successor function across languages, and that the timeline for this learning varies as a function of count list transparency.

1. Introduction

Linguistic expressions of number - like *sixteen* and *seventy-two* - provide humans with a powerful ability to exactly quantify a limitless array of objects and other entities. This ability transcends our specific experience with numbers: Although most people have never counted to (or even thought about) the number *three million and thirty-two*, we can immediately recognize it as a possible number, and can judge without hesitation that adding “1” would generate *three million and thirty-three*. Somehow, when we learn to count in childhood, we use a finite training set to extract a set of rules that yield a potential infinity of numbers. In this sense, number words - like language more generally - make “infinite use of finite means” (Chomsky, 1965; von Humboldt, 1999), perhaps by virtue

* Corresponding author at: Psychology Department, University of California, San Diego, 9500 Gilman Drive, La Jolla, CA 92093-0109, United States.

E-mail address: roschnei@ucsd.edu (R.M. Schneider).

<https://doi.org/10.1016/j.cogpsych.2019.101263>

Received 13 June 2019; Received in revised form 25 October 2019; Accepted 9 December 2019

0010-0285/ © 2019 Elsevier Inc. All rights reserved.

of the fact that they are fundamentally linguistic symbols, acquired by children as part of the process of acquiring language. In the present study, we pursued this analogy between natural language and the acquisition of counting, and tested how children's learning of recursive counting rules is affected by the grammatical structure of counting cross-linguistically. Specifically, we asked how the rule-governed structure of number word morphology in transparent languages (Cantonese, Slovenian, and English) and more opaque languages (Hindi and Gujarati) affected children's discovery of the recursive successor function that governs counting.

Although children begin reciting some portion of the count list at around 2 years of age (Fuson, 1988), they initially appear to treat it as a closed routine, or "unbreakable chain", rather than as a productive system governed by rules. At this point, many children are unable to count beyond 10, and show no understanding of how counting can be used to label the cardinality of sets. For example, a child who is able to count to 10 may nevertheless produce a random number of items if asked to generate a set of *one* (Le Corre, Van de Walle, Brannon, & Carey, 2006; Wynn, 1990, 1992). Despite being able to recite a partial count list, children acquire meanings for the numerals *one*, *two*, and *three* in highly protracted stages over the course of about 18 months between the ages of 2.5 and 4 years, and rarely even attempt to count when asked to label sets or give a particular number of objects (Le Corre & Carey, 2007). Only at around the age of 3.5 to 4 years do US children begin to systematically use counting to label and generate sets, evidence that they have acquired some form of the Cardinal Principle (CP) - e.g., that the final word used in a count routine labels the cardinality of the set as a whole (Gelman & Gallistel, 1978).

This apparent discontinuity in children's understanding of number is not currently well understood, though there is growing evidence that it is importantly related to changes in children's readiness for subsequent numerical learning (Geary, 2018; Geary et al., 2018; Spaepen, Gunderson, Gibson, Goldin-Meadow, & Levine, 2018). On some accounts, acquisition of the CP indicates a moment of conceptual change in which children discover the semantic content of counting (Carey, 2004, 2009; Sarnecka & Carey, 2008; Le Corre & Carey, 2007; Wynn, 1990, 1992). On this view, children make a "wild induction" based on an analogical mapping between counting and cardinality: They notice that as one counts up from *one* to *two* and then from *two* to *three*, the cardinality of the sets labeled by these words grows in increments of exactly 1 (see also Gentner, 2010; Marchand & Barner, 2018; Wynn, 1992). On the basis of this isomorphism between the count list and cardinal meanings, children hypothesize that the meaning of the next numeral in their list (*four*) differs from the cardinality of the previous numeral by exactly 1 as well, and that more generally every number, n , has a successor defined as $n + 1$. This, as noted by Sarnecka and Carey (2008) amounts to acquiring implicit knowledge of the successor function, a central element of the Peano axioms, which provide a logical foundation for arithmetic, a subset of which are as follows:

1. 1 is a natural number.
2. If n is a natural number, then $S(n)$ is also a natural number.
3. For every natural number n , $S(n) \neq 1$.
4. If P is a property of natural numbers such that (a) 1 has property P , and (b) whenever a natural number has property P , so does its successor, then all natural numbers have property P , and every number has a natural successor.

Critically, this account posits that children acquire a recursive successor function at around the age of 3.5 or 4 years, allowing them to accurately label and generate sets of any size that is within their known count list. As evidence for this hypothesis, Sarnecka and Carey (2008) used a paradigm they called the "Unit Task." In this task, an experimenter placed either 4 or 5 items into a box, providing the appropriate cardinal label - e.g., "There are 4 frogs in the box." - and then added 1 or 2 additional items, while asking, "Are there 5 or 6 frogs in the box now?" Using this method, they found that only CP-knowers exhibited above-chance performance (around 66%), providing support for the claim that acquisition of the successor function is related to acquisition of the CP. However, while CP-knowers as a group performed above chance on this task, many CP-knowers appeared to fail completely, raising the question of whether acquisition of the successor function is what makes children CP knowers, as Sarnecka and Carey claimed, or whether instead this knowledge is acquired sometime later.

Subsequent work has also found that although being a CP-knower may be a necessary condition for learning about the successor function (Spaepen et al., 2018), many CP-knowers lack knowledge of the successor function, and appear to master it only after becoming exceptionally strong counters (Cheung, Rubenson, & Barner, 2017; Davidson, Eng, & Barner, 2012; Wagner, Kimura, Cheung, & Barner, 2015). For example, Cheung et al. (2017) found that US children only succeed at the Unit Task for the largest numbers in their count lists by around age 5.5, and that this coincides with the moment at which they begin to claim that, rather than being finite, numbers never end (for related evidence that children first judge numbers to be infinite at this age, see Evans, 1983; Hartnett & Gelman, 1998). This finding is consistent with a much earlier finding, by Secada, Fuson, and Hall (1983), that children as old as 5 or 6 struggle with a task almost identical to the Unit Task - which they use to assess "counting-on" (i.e., the ability to add a set of 5 to a set of 1 without recounting the set of five objects after it is labeled for them).

Critical to our study, Cheung et al. (2017) note that performance on the Unit Task is best predicted by how high a child can count. Although it is perhaps unsurprising that counting experience might be related to discovering the underlying logic of the count system, it remains unknown how it might help. Nothing about memorizing a finite list guarantees that children should impute a recursive rule to this list, or conclude that numbers are infinite; after all, other lists, like the ABCs or the months of the year, aren't generated by a recursive rule. One possibility is that, as proposed by Carey and colleagues, children posit a recursive function based *purely* on a mapping between cardinalities and the ordered count list, such that the inference that numbers never end is an inductive generalization based on this analogy (Carey, 2004, 2009; Gentner, 2010; see Marchand & Barner, 2018, for discussion). However, as Cheung et al. (2017) note, another possibility is that the recursive rule takes its origin in the morpho-syntactic structure of the count list itself (see also Barner, 2017; Hurford, 1987; Rule, Dechter, & Tenenbaum, 2015; Yang, 2016). For example, a child learning English might notice that after each decade term (*twenty*, *thirty*, *forty*, etc.) the next number in the count sequence can be generated by appending

Table 1

Examples of number words in Cantonese, Slovenian, Hindi, Gujarati, and English. Languages are arranged from most to least transparent. Bolded items refer to irregular/novel items in the count list.

	Numeral							
	2	5	10	20	25	50	52	102
Cantonese	yih	ñgh	sahp	yihsahp	yihsahpñgh	ñghsahp	ñghsahpyih	yātbāaklingyih
Slovenian	dva	pet	deset	dvajset	petindvajset	petdeset	dvainpetdeset	sto dva
English	two	five	ten	twenty	twenty-five	fifty	fifty-two	one hundred two
Gujarati	be	pāñch	das	vis	panchis	pachās	bāvan	eka so be
Hindi	do	pāñch	das	bis	pachchis	pachās	bāvan	ek sau do

the words *one* through *nine* in order - i.e., an additive decade + unit rule. In parallel, some children might also learn a multiplicative unit*decade rule for forming decade labels - e.g., seven*ty; eight*ty (though noticing such a rule is not strictly necessary to become a productive counter, since the decades can be readily memorized, and in some cases must be due to their irregularity). A child who has learned an additive decade + unit rule might use this to productively generate increasingly large number words which go beyond their input. Given this, high counters' better performance on measures like the Unit Task might result not merely from more number language input in general, but also from their knowledge of the productive morphological rules that allow numbers to be freely generated. Children's belief that numbers never end might originate in a rule that suggests that number *words* might never end.

Several studies provide preliminary evidence that young children learn the productive rules that govern counting prior to exhibiting errorless counting ability. First, when English-speaking children are asked to count as high as they can, their errors are non-random and often occur on decade transitions like *twenty-nine* and *thirty-nine* (Fuson, Richards, & Briars, 1982; Gould, 2017; Siegler & Robinson, 1982; Wright, 1994). If children simply memorized their count routine as an unstructured list like the alphabet, we might expect their errors to be randomly distributed, such that errors should be just as likely for numbers like 22, 25, and 29. However, the finding that children's frequent failure to recall decade terms - but not the words preceding them - suggests that their ability to count up to these words is not driven purely by memory, but instead by the application of a rule that combines the highest known decade label with the numbers 1–9. As a second form of evidence for children's acquisition of productive counting rules, several studies have compared how children learn to count in languages which have more or less transparent morphological rules governing counting. For example, in languages like Cantonese, in which the numbers 11–99 can be generated via rule-governed combinations of the verbal labels for 1–10, (see Table 1) children count higher and make fewer errors than same-aged children learning English, which has multiple exceptions in the teens and most decade labels (Miller, Smith, Zhu, & Zhang, 1995; Miller & Stigler, 1987). Speakers of Korean, which like Cantonese is highly transparent, also exhibit better performance on multidigit addition and subtraction problems and on identifying place-value names than age-matched English-speaking children (Fuson & Kwon, 1992). In a within-culture comparison, children learning Welsh (which has a highly regular count list) outperformed English-speaking peers in tasks assessing place-value comprehension (Dowker, Bala, & Lloyd, 2008). Finally, several studies have found that children learning less transparent languages (such as English, French, and Swedish) demonstrate a weaker understanding of the base-10 system in comparison to speakers of more transparent languages like Japanese, Korean, and Cantonese (Miura, Kim, Chang, & Okamoto, 1988; Miura & Okamoto, 1989).

Although there have been several points of connection made between the regularity of counting systems, how high children can count, and mathematical achievement, no previous work has provided direct evidence that such effects are actually due to differences in how readily children extract recursive counting rules (e.g., by examining individual differences within a particular culture). For example, while previous findings are consistent with a role for counting transparency, there are also known differences in the levels of counting, number, and mathematics exposure across many of these previously studied groups (Pan, Gauvain, Liu, & Cheng, 2006; Towse & Saxton, 1998). Consequently, it is unclear whether these findings reflect differences in linguistic structure, or in cross-cultural practices surrounding number and mathematics instruction. Further, these advantages in mathematics education outcomes extend well into elementary school and high school (Siegler et al., 2012; Watts, Duncan, Siegler, & Davis-Kean, 2014) at a time when computations depend mainly on written numerals, such that counting transparency should play a much weaker role, if any, in learning. Very generally, many cross-cultural differences including language, mathematics curriculum, and societal attitudes toward the importance of early math education may impact children's early counting fluency without necessarily implicating children's ability to detect recursive counting rules. If previously attested differences in mathematics learning result from the impact of counting transparency on the acquisition of counting rules, then children learning transparent counting systems should be faster to extract recursive rules from their count list, and should be faster to apply these rules to reasoning about simple addition facts, like those tested by Sarnecka and Carey (2008) Unit Task.

The present work addresses these issues by directly assessing whether individual children have acquired productive counting rules, how knowledge of such rules is related to successor function knowledge, and how both differ across languages and cultures. First, in keeping with previous work, we tested the role of count-list transparency on acquisition of successor function knowledge by comparing children learning languages with relatively transparent count lists to children learning languages with much more opaque count lists. Second, we reasoned that if successor function knowledge is driven by learning recursive counting rules, then this should be true across all cultures, regardless of how opaque the count list, how long it takes for children to become competent counters, or how much input they receive. Third, to assess the value of this within-culture approach, we tested how other cultural factors might

impact learning by comparing children learning similarly transparent languages across cultures that differ with respect to previously attested outcomes in mathematics education. To accomplish these goals, we conducted two Experiments with data from five different groups that varied with respect to both counting transparency and cultural practices surrounding mathematics education.

In Experiment 1, we tested children learning Cantonese (in Hong Kong), Slovenian (in Slovenia), and English (in the US). As already noted, Cantonese is a fully regular count system, with the entirety of the count list from 11 to 99 generated using the verbal labels for 1–10 (see Table 1). For example, the label for 25 (*yihshapnigh*) can be described as the joint application of two distinct rules: the multiplicative unit*decade rule (two*decade, or *yih*sahp*) and the additive decade + unit rule (decade + five, or *yih*sahp* + *nigh*), each of which is fully regular (i.e., exceptionless) up to 99 in Cantonese. Given this, Cantonese-speaking children may be quick to learn the structure of the count list because there are no exceptions to this structure, allowing them to potentially notice rules after memorizing a relatively small subset of the count list. Slovenian is slightly less transparent than Cantonese; while numerals from 11 to 99 are all formed according to a unit + unit*decade structure, it features several irregular formulations (e.g., the teens more closely resemble English). Prior to grade school, however, Slovenian children appear to receive relatively little counting exposure and often cannot count past 10 before age 5, when US children count to nearly 50 (Almoammer et al., 2013; Marušič, Plesničar, Razboršek, Sullivan, & Barner, 2016). Thus, while Slovenian-speaking children do have the benefit of a fairly regular count system, they appear to have low exposure to counting routines, allowing us to weigh the relative effects of productive counting knowledge vs. number training. Critically, however, our main question was whether the acquisition of productive counting rules is related to successor function knowledge across opaque and transparent languages, and independent of cross-linguistic difference in learning timecourse.

In Experiment 2, we investigated the relationship between productive counting and successor function knowledge in two languages with highly opaque count lists: Hindi and Gujarati. While Hindi and Gujarati numbers are generated through a base-10 system, both exhibit many more irregularities, morphological variations, and unpredictable phonological changes in comparison to English (Berger, 1992; Bright, 1969). In fact, it has been argued that the phonological properties of Hindi and Gujarati are so irregular that the numbers 1–100 can only be learned through rote memorization, a claim we test in our study (Bright, 1969; Comrie, 2011; see Appendix for a table demonstrating irregularities in the Hindi count list). As an example of this complexity, whereas in Hindi the numbers 2, 5, and 10 are *do*, *pānch*, and *das*, respectively, the label for 20 is *bīs*, 25 is *pachchīs*, and 50 is *bāvan*. Such patterns are found throughout both languages. Given this, we hypothesized that children learning Hindi and Gujarati might have to memorize a larger segment of their count list before converging on productive counting rules. However, while we were interested in cross-linguistic differences, our main focus, as in Experiment 1, was to test within-group relations between children's mastery of the count list's structure and their knowledge of the successor function.

Since previous studies do not establish a single gold standard for evaluating children's acquisition of productive counting rules, in each of these Experiments we took an exploratory approach to evaluating this knowledge. In particular, we developed and pre-registered several measures that sought to differentiate between children with a fully memorized list versus those who count using rules. The first was a commonly used measure of children's counting mastery which involves testing how high they can count without error - what we call their "Initial Highest Count" (Almoammer et al., 2013; Barth, Starr, & Sullivan, 2009; Cheung et al., 2017; Davidson et al., 2012; Marušič et al., 2016; Fuson et al., 1982; Siegler & Robinson, 1982; Wagner, Chu, & Barner, 2019). However, a problem with this measure is that for some children it has the potential to underestimate knowledge of productive counting rules: Whereas some children — especially those who make errors on decade transitions like 29 — can count higher when prompted, other children who make errors nearby in the count list (e.g., 32) are unable to continue counting, suggesting that their list is likely memorized, and not generated by rules (Siegler & Robinson, 1982). Therefore, when used as a standalone measure, Initial Highest Count is ambiguous with respect to children's productivity.

Given these concerns, we provided children with prompts when they made errors or omissions in their count routine. In addition to analyzing children's initial errors when counting, we also measured their Final Highest Count, a potentially stronger indicator of their ability to generate new number labels, and then built a measure based on the difference between these two, which we expected would provide an especially strong measure of productivity since it reflects the ability to continue counting after an initial error when given a prompt. Finally, to assess knowledge of productive rules outside of the count routine, we tested children's ability to name the next number in the count list from arbitrary points, both within and beyond the range of their Initial Highest Count (e.g., "105, what comes next?"). We reasoned that only children with strong knowledge of how to productively generate number labels should perform well on this task. Taking this exploratory approach, we compared the relative ability of these metrics to predict children's acquisition of a recursive successor function as measured by Sarnecka and Carey (2008) Unit Task, with the goal of identifying a strong measure of productivity that explains knowledge across different languages and cultures.

2. Experiment 1

2.1. Method

The methods and analyses of this study were pre-registered prior to any data collection. The pre-registration can be found at <https://osf.io/2vzxr>. All methodological and analytical choices were as pre-registered, unless stated otherwise in-text. Throughout the method section, we use English examples; however, stimuli were always presented in the child's language.

2.1.1. Participants

We pre-registered a minimum n of 80 per language group to conduct analyses, and a maximum n of 150. We recruited 378 children aged 3.5 to 6.5 years from preschools, elementary schools, and the surrounding community in Hong Kong; Nova Gorica,

Slovenia; and San Diego, California, USA. Fifty of these children were tested, but excluded from analyses as pre-registered for missing data from more than 20% of trials ($n = 25$); not completing the Highest Count task ($n = 5$, including children who were unable to count to two);¹ missing Highest Count recording ($n = 4$); being out of age range ($n = 6$); experimenter error ($n = 6$); non-native primary language ($n = 2$); noted by experimenter for exclusion ($n = 1$); or for parental interference ($n = 1$). As pre-registered, an additional two participants were excluded only from analyses involving the Next Number and WPPSI tasks in our English dataset due to failure to complete the minimum number of test trials for those tasks.

After these exclusions, our final sample included 328 participants. A breakdown of demographic information by language is shown in Table 2.

2.1.2. Stimuli, design, and procedure

Children were tested individually by native speakers of Cantonese, Slovenian, or English in a room set apart from the classroom (in Hong Kong, Nova Gorica, and San Diego), or in a small lab testing space (in San Diego). Participants received the tasks in a fixed order (Abbreviated Give-N, Highest Count, Unit Task, Next Number, and WPPSI Picture Memory Test).

2.1.2.1. Abbreviated Give-N. This task was a conservative test of whether children understood the CP. The experimenter provided children with 10 plastic objects (e.g., buttons, bananas, apples, or bears), and a small plastic plate. After familiarizing the child with the purpose of the game, the experimenter asked them to put N items on the plate (trials included 6, 9, 7, and 5, in that order). After the child finished placing a set on the plate, the experimenter asked, “Is that N ? Can you count to make sure?” If the child answered in the negative they were permitted to fix the set. If children were able to correctly generate only three of the four requested sets, they were given a second try on the failed trial. Children were classified as CP-knowers if they correctly generated sets for all four numbers. Both CP- and subset knowers were included in analyses.

2.1.2.2. Highest Count. The experimenter introduced the task to the child by saying, “In this game I want you to count as high as you can. Can you start counting with one?” If the child did not begin counting after *one*, the experimenter repeated the prompt with a rising intonation. If the child made an error, the experimenter immediately stopped them by saying, “Wait a minute, what comes after N ?” This provided the child an opportunity to self-correct. If the child failed to correct the error the experimenter provided the next number by saying, “Actually, what comes after N is $N + 1$. Can you keep counting?” If the child did not continue, the experimenter repeated the previous three numbers (including the prompt) with a rising intonation to encourage the child to continue. If the child made an error immediately after being given a prompt, or was otherwise unable to continue, the experimenter stopped the task. The task was similarly ended if the child made more than three errors within a single decade, or more than three consecutive errors (i.e., counts with one prompt between each number). Otherwise, children were allowed to count to 140, and were then stopped and congratulated (“Wow! You counted to 140!”). Throughout the task, children were allowed up to 14 prompts (an average of one per decade), as we hypothesized that even children who were highly familiar with the base system might nevertheless struggle to recall decade transitions, which are often irregular. No child used all 14 prompts; the maximum number given was 12, with an average of 2.53 prompts across all languages. All children’s counting data were recorded on a voice recorder and independently coded and validated by two other researchers, which allowed us to apply the same coding criteria to all children.

2.1.2.3. Unit Task. To assess children’s understanding of the successor function, we used a modified version of the Unit Task (Sarnecka & Carey, 2008) presented on a tablet. The experimenter presented children with a scene depicting a picture of a frog on a lilypad, saying, “This is my friend, Froggie. Froggie is going to tell you about some fish she sees in the pond. You have to listen very carefully to Froggie, because she is going to ask you some questions. Let’s see what Froggie has to show us.”

Every trial had three phases. First, the child saw some number of fish move into the middle of the screen, and heard a pre-recorded female voice say in their native language, “Look! There is/are N fish in the pond!” The fish were visually presented for approximately 1.5 s. Next, a lilypad covered the fish so that they were no longer visible, and children were given a memory check: “How many fish are in the pond?” If the child failed this first memory check, the experimenter went back to the start of the trial, saying, “Let’s try that again!” If the child failed this second memory check, the experimenter told the child how many fish were in the pond, and proceeded with the remainder of the trial.

After the memory check, children heard, “Look!” and saw one fish swim in from the right side of the screen and remain directly to the right of the lilypad. Children then heard the critical question, “Are there $N + 1$ or $N + 2$ fish now?” Order of alternatives ($N + 1$ or $N + 2$) was counterbalanced across trials. If children failed to pick one of the presented alternatives, the experimenter provided the alternatives again verbally. Participants completed a training trial (with 1 fish; on this trial, participants received feedback) and then 12 test trials (numbers queried were 5, 7, 16, 24, 52, 71, 105, 107, 116, 224, 252, and 271). In contrast to previous work using this task (Cheung et al., 2017; Davidson et al., 2012; Sarnecka & Carey, 2008), we included items so large that they that could not have been produced in the Highest Count task (224, 252, and 271).

For both the Unit Task and the Next Number task (below), the correct response for a given N was $N + 1$. “I don’t know” responses were coded as incorrect. If a participant did not respond for a given N , that trial was excluded from analysis, but otherwise all numeric responses were included in analyses. Trials were additionally classified as being either within or outside of a child’s Initial

¹ Children who were unable to count to two in the Highest Count task were excluded because it was unclear whether such failure to count reflected a lack of knowledge or an unwillingness to participate.

Table 2

Demographic information for Cantonese, Slovenian, and US English samples.

	<i>n</i>	<i>n</i> CP-knower	<i>M</i> _{age} (<i>SD</i>)	<i>Median</i> _{age}
Cantonese	118 (55 female, 63 male)	98	5.16 (0.81)	5.07
Slovenian	99 (45 female, 54 male)	65	5.21 (0.75)	5.25
English (US)	111 (43 female, 68 male)	71	4.79 (0.85)	4.78

Highest Count.

2.1.2.4. Next Number task. The experimenter introduced the task by saying, “This is called ‘What Comes Next.’ In this game, I’m going to say a number, and you’ll tell me the one that comes next.” For every number, the experimenter prompted the child by saying, “N, what comes next?” If a child gave a response that was less than the initial prompt the experimenter reminded the child that the game was called, ‘What comes next,’ and allowed them to change their response. Children only received one such reminder. The numbers queried in this task were the same as in the Unit Task.

2.1.2.5. Picture memory task. This task was adapted for display on a tablet from the WPPSI-IV (Wechsler, 2012) picture memory task, and was included to assess children’s nonverbal working memory. Children were presented with pictures of familiar items (e.g., bell, block, and hat) and told to remember them with the prompt, “Look at this/these picture(s)!”. After either 3s (single target) or 5s (multiple targets), children saw a set containing the target and some number of distractor objects (e.g., chair, drum, and rainbow). Children were asked to touch the objects they had just seen with the prompt, “Point to the picture(s) I just showed you.” As is typical for this task, to prevent the use of verbal rehearsal strategies, children were stopped from saying the names of the objects and told that they had to be silent during the game.

Children received three trials with feedback at the start of the task. If they selected the incorrect items on these trials, the experimenter showed them the target items again, saying, “I showed you these pictures, so you should choose these pictures.” Following WPPSI protocol, if children were younger than 4 years, they began the task with 1-item trials, while children older than 4 years began the task with 2-item trials. A response was only considered correct if the child correctly identified every target item. The task was terminated after 3 consecutive incorrect trials, and there were either 28 or 32 possible test trials, depending on the age of the child. As the task progressed, both the number of target items and the number of distractors increased (up to 7 targets and 12 distractors).

Children received one point for every correct trial. We summed all correct trials for each participant to obtain their raw working memory score.

2.2. Measures of productivity

2.2.1. Highest Count

This task yielded three measures of children’s counting knowledge, each of which might capture knowledge of productive counting rules. Initial Highest Count was the highest number counted to without errors. As noted in the Introduction, however, this measure alone may not yield a reliable measure of productivity since the source of children’s errors in a Highest Count task is often ambiguous. Our second measure, Final Highest Count, was defined as the highest number reached during the counting task with the aid of experimenter prompts. We reasoned that this measure should distinguish between children who understand how to productively generate the remainder of the numbers within a given decade, and those who do not. This measure is appealing as a measure of productivity because it captures the difference between a child who, having counted to, e.g., 29, can continue counting up to, e.g., 39, when prompted with the label *thirty* vs. one who cannot continue counting. Additionally, restricting the number of prompts for children who did not demonstrate knowledge of the count list’s structure meant that these prompts only provided meaningful Final Highest Count improvements if the child was able to use those prompts to progress substantially through the count list.

Our third measure was a binary classification based on the difference between these first two measures (Initial and Final Highest Count) in which we classified children as either “Resilient” or “Non-Resilient” counters, on the hypothesis that Resilient counters rely on knowledge of the base-system to count up from errors. Children were classified as Resilient if they were able to count at least two decades past any error (without making more than three errors in those two decades). This pre-registered criterion allowed children two prompts at each decade transition within those two decades, plus one additional mid-decade or decade-beginning error. Our logic in developing this classification was that, if children have some understanding of the structure of the count list, but perhaps have made an error due to an irregular decade label, they should be able to productively use experimenter prompts to continue counting beyond their Initial Highest Count. The two-decade criterion was chosen *a priori* because it provided stronger evidence that children’s counting behavior after a prompt was rule-governed (as opposed to a one-decade criterion) while also accommodating children who might have initial counts close to 100, where additional embedding may cause further errors. Children who were unable to meet this criterion for any error made during the Highest Count task were classified as Non-Resilient. In Experiments 1 and 2, only nine children used ten or more prompts, and 67 children used five or more prompts. Overall, Resilient and Non-Resilient counters tended to use prompts similarly, and to the same degree. This categorical classification contains some noise, however; for example, it does not distinguish between children who met these criteria after counting to 30 versus those who were able to count to 100.

Additionally, it may classify children who were able to count quite high but not able to continue after a prompt as Non-Resilient. Keeping these shortcomings of the classification in mind, we report Resilience as a broad diagnostic of productive counting knowledge.

2.2.2. Next Number

The Next Number task acted as a fourth potential measure of counting productivity. We reasoned that children who have productive counting rules should have an easier time labeling the next number in a sequence, especially for numbers beyond their Initial Highest Count. Also, unlike the other measures, this task required children to generate the next number without the benefit of the count routine's momentum. Children's Highest Contiguous Next Number was defined as the highest number for which they were able to generate a successor, provided all previously queried items were also correct. For example, if a child responded correctly for 5, 7, and 24, but incorrectly for 16, their Highest Contiguous Next Number would be 7. For children who made an error on 1, the training item in this task, their Highest Contiguous Next Number was 0.²

3. Results

3.1. Highest Count

A breakdown of counting profiles by language and Resilience is shown in Table 3. Consistent with prior work, we found that Cantonese-speaking children demonstrated overall greater counting proficiency than either US or Slovenian children in both their Initial and Final Highest Counts. As expected, we also found lower levels of counting proficiency in Slovenian-speaking children in comparison to the other two languages (Table 3). Children in all three languages used experimenter prompts to a similar degree. In grouping children by Resilience, we found that Resilient counters had higher Initial and Final Highest Counts than Non-Resilient counters in all languages. Nevertheless, we still found that Cantonese-speaking children were able to count higher than English- and Slovenian-speaking children, regardless of Resilience classification.

Consistent with our hypothesis that acquiring productive counting rules is facilitated by increased count list transparency or exposure, we found the greatest number of Resilient counters in Cantonese, with 51% of children identified as Resilient. In contrast, only 26% of Slovenian children were classified as Resilient. Finally, 38% of children were identified as Resilient in our US sample, which has higher rates of counting exposure in comparison to Slovenian, but a lower level of count list transparency.

A visualization of counting profiles is shown in Fig. 1. Consistent with our motivation for seeking alternative measures of counting productivity, we found that a majority of children (63% across all three languages) who stopped before 140 were nevertheless able to count beyond their Initial Highest Count when provided with prompts. Further, many children were classified as Resilient despite having fairly low Initial Highest Counts (i.e., they could count at least 2 decades beyond their first error). While Initial Highest Count was strongly correlated with Final Highest Count in all languages (Cantonese: $r = 0.86$, $p < .0001$; Slovenian: $r = 0.80$, $p < .0001$; English (US): $r = 0.82$, $p < .0001$), the frequency of children who were able to count beyond their initial errors, sometimes by many decades, indicates that this measure may not always fully capture knowledge of a productive counting rule.

3.2. Predictors of Unit Task performance

3.2.1. Within-language analyses

In this section, our main goals were to test whether productive counting knowledge is predictive of successor function knowledge within each language group, and to test which measure of productivity was the strongest predictor. To address these questions, we constructed four models for each language group, separately predicting Unit Task performance from our candidate measures of productivity: (1) Initial Highest Count; (2) Final Highest Count; (3) Counting Resilience; and (4) Highest Contiguous Next Number. For each model, we also included age, whether the number was within or outside the child's Initial Highest Count, and the numerical magnitude of the Unit Task trial,³ with subject as a random factor.⁴ We tested whether each candidate measure significantly predicted Unit Task performance by conducting a Likelihood Ratio Test between each of these four individual models and the base model. Next, we constructed our large models. In constructing large models for this and all other analyses, we used hierarchical model comparison to test whether candidate measures of productive counting knowledge explained unique or overlapping variance in children's performance. The base of each large model contained the candidate measure associated with the lowest AIC, to which we added predictors in order of increasing AIC. Candidate measures were retained on the basis of a significant χ^2 value.

As demonstrated in Table 4, we found that several candidate measures of counting productivity were related to acquisition of the

² We did not pre-register a Highest Contiguous Next Number for children who failed the first trial of the Next Number task, as we did not encounter any children who failed this trial during piloting. Failure on this initial trial was rare in all datasets: n Cantonese = 8; n Slovenian = 9; n US English = 8; n Hindi = 4; n Indian English = 5.

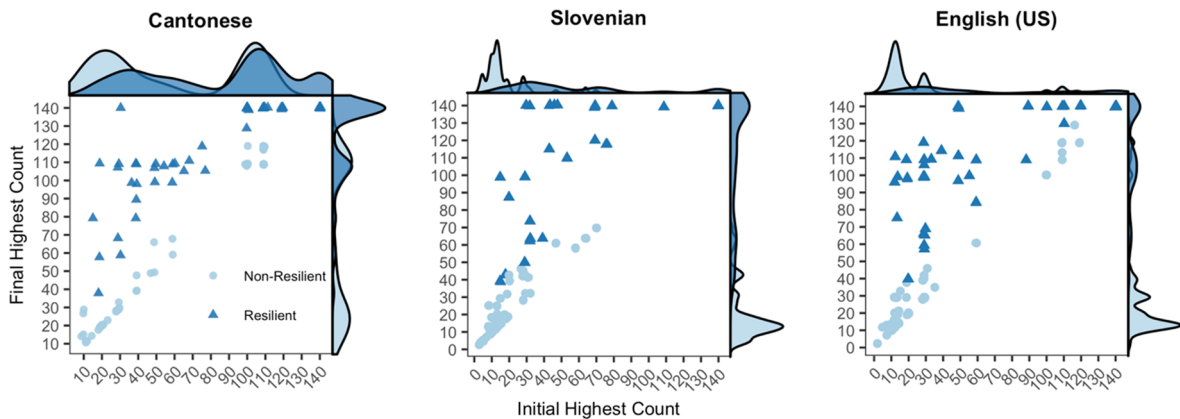
³ Note that we pre-registered the use of numerical magnitude in these models for Experiment 2, but not for Experiment 1. Based on the results for Experiment 2, we felt that the most informative analyses were those that took into account how large the number was on the Unit Task.

⁴ Models were generalized linear mixed effects models constructed in R using the 'lme4' package (Bates, Maechler, Bolker, & Walker, 2014) with the formula: Correct ~ [Resilience/Initial Highest Count/Final Highest Count/Highest Contiguous Next Number] + Trial Within/Outside Initial Highest Count + Item Magnitude + Age + (1|Subject). Continuous predictors were centered and scaled to facilitate model fit.

Table 3

Counting data by language. Initial and Final Highest Count are rounded.

	<i>n</i>	<i>M</i> IHC (<i>SD</i>)	<i>M</i> FHC (<i>SD</i>)	<i>M</i> Prompts (<i>SD</i>)
Cantonese				
Overall	118	73 (41.91)	94 (45.06)	2.63 (1.93)
Resilient	61	85 (40.00)	122 (25.44)	3.87 (2.12)
Non-Resilient	57	60 (40.24)	64 (42.43)	1.51 (0.63)
Slovenian				
Overall	99	27 (28.36)	44 (45.07)	2.31 (1.87)
Resilient	26	56 (38.34)	109 (36.08)	4.30 (2.40)
Non-Resilient	73	17 (13.80)	21 (16.04)	1.67 (1.06)
English (US)				
Overall	111	41 (41.53)	61 (49.19)	3.04 (2.77)
Resilient	42	64 (46.32)	110 (28.33)	5.68 (3.36)
Non-Resilient	69	27 (31.19)	31 (32.01)	1.74 (0.92)

**Fig. 1.** Initial and Final Highest Counts by language, grouped by Resilience. Points indicate the relation between a participant's Initial and Final Highest Counts. Points are jittered slightly to avoid overplotting. Density plots indicate the distribution of Initial (top) and Final (right) Highest Count by Resilience.**Table 4**

Within-language Unit Task models in Cantonese, Slovenian, and US English. IHC = Initial Highest Count; FHC = Final Highest Count; HCNN = Highest Contiguous Next number. Only final models shown for each language: predictors without coefficient estimates did not significantly improve the fit of that language's model in a Likelihood Ratio Test. Significance was calculated using the standard normal approximation to the *t* distribution (Barr, Levy, Scheepers, & Tily, 2013).

Predictors	Cantonese			Slovenian			English (US)		
	β	CI	<i>p</i>	β	CI	<i>p</i>	β	CI	<i>p</i>
(Intercept)	0.69	0.44–0.93	< 0.001	0.22	0.03–0.41	0.023	0.58	0.36–0.79	< 0.001
IHC	0.81	0.56–1.06	< 0.001	—	—	—	0.75	0.44–1.06	< 0.001
FHC	—	—	—	0.36	0.11–0.61	0.004	—	—	—
HCNN	—	—	—	0.35	0.12–0.57	0.002	0.43	0.15–0.71	0.003
Trial Within IHC	−0.05	−0.43–0.34	0.811	0.26	−0.11–0.64	0.173	0.09	−0.29–0.47	0.639
Item Magnitude	−0.52	−0.70–−0.33	< 0.001	−0.33	−0.49–−0.17	< 0.001	−0.48	−0.65–−0.32	< 0.001
Age	0	−0.23–0.23	0.993	0.12	−0.08–0.31	0.241	0.03	−0.20–0.26	0.792

successor function, though no single measure of productivity consistently emerged as the best predictor in each of the languages when considered individually. In Cantonese, Initial Highest Count was the strongest predictor of Unit Task performance ($\chi^2_{(1)} = 38.91$, $p < .0001$), with neither Final Highest Count ($\chi^2_{(1)} = 1.39$, $p = .24$) nor Highest Contiguous Next Number ($\chi^2_{(1)} = 1.86$, $p = .17$) improving the fit of this model. In Slovenian, we found that Highest Contiguous Next Number and Final Highest Count were both predictors of Unit Task performance in Slovenian: Final Highest Count significantly improved the fit of a model containing Highest Contiguous Next Number ($\chi^2_{(1)} = 8.03$, $p = .005$), but neither the addition of Initial Highest Count ($\chi^2_{(1)} = 0.02$, $p = .89$) nor Resilience ($\chi^2_{(1)} = 0.07$, $p = .80$) explained additional variance. Finally, in English, both Initial Highest Count and Highest Contiguous Next Number were predictors of Unit Task performance: Highest Contiguous Next Number significantly improved the fit of a model containing Initial Highest Count ($\chi^2_{(1)} = 8.65$, $p = .003$), while Final Highest Count did not

($\chi^2_{(1)} = 0.10, p = .75$). Thus, taking this approach, we found that different measures of counting ability predicted Unit Task performance in different languages. In our next analyses, we asked which of these measures was the best fit of the entire data set (including all three language groups).

3.2.2. Cross-linguistic analyses

In our next set of analyses, we had two goals. First, we sought to ask which of the four measures of counting ability best predicted children's Unit Task performance across all three languages when included in a single model. Second, we sought to provide a basic characterization of how the samples in our study differed, and thus whether our data are compatible with past reports that find cross-cultural differences, and in particular an advantage for Mandarin- or Cantonese-speaking children. We found that the overall pattern of performance on the Unit Task was largely consistent with prior work finding an advantage for Chinese children, with higher mean performance for Cantonese ($M = 0.63, SD = 0.23$) in comparison to both Slovenian ($M = 0.55, SD = 0.22$) and English ($M = 0.59, SD = 0.25$).⁵ In order to address limitations of past work, however, we attempted to account for some individual differences in counting exposure by assuming that a child's Initial Highest Count provides some signal of counting training. Children's rote counting ability has been shown to be correlated with levels of parental number talk, such that children who can count higher without error are likely to have had more exposure to counting than children who have lower rote counts (Lefevre, Clarke, & Stringer, 2002; Manolitsis, Georgiou, & Tziraki, 2013; see also Saxe, Guberman, Gearhardt, & Gelman, 1987). While children's highest errorless count may be indicative of their exposure to counting training, it does not offer an unequivocal proxy, as children eventually impute productive counting rules which allow them to surpass a rote memorized list. In this way, Initial Highest Count may provide a nonlinear estimate of exposure if some children are faster to become productive counters than others. Therefore, our logic in including Initial Highest Count in our cross-linguistic analyses was to remain as conservative as possible regarding the nature of cross-linguistic differences on our other potential measures of productivity (Final Highest Count, Resilience, and Highest Contiguous Next Number). Thus, to the extent that Initial Highest Count captures some of the variance explained by knowledge of productive rules, it will lead us to produce conservative estimates for these other measures of productivity.

Additionally, unlike most previous studies, we included a working memory term to control for domain-general cognitive differences between groups. We built three models⁶ predicting Unit Task performance from (1) Counting Resilience; (2) Final Highest Count; and (3) Highest Contiguous Next Number. As in our within-language analyses, these models included effects of item magnitude, whether the item was within or outside the child's Initial Highest Count, age, and raw working memory score, with a random effect of subject.

As demonstrated in Table 5, Highest Contiguous Next Number was the single best counting productivity predictor of Unit Task performance cross-linguistically, improving the model fit in comparison to the base ($\chi^2_{(1)} = 22.70, p < .0001$). This model also revealed a significant main effect of language: when controlling for other factors, English-speaking children exhibited better performance relative to Cantonese-speaking children ($\beta = 0.43, p = .002$), and there were no other significant language-wise differences (Slovenian vs. Cantonese: $p = .22$; English vs. Slovenian: $p = .15$). These effects were significant even when controlling for Initial Highest Count, which we included in an effort to provide a conservative estimate of the effects of cross-linguistic differences and productivity. Additionally, this model indicated that Initial Highest Count predicted Unit Task performance, with higher Initial Highest Counts associated with better performance ($\beta = 0.55, p < .001$). Critical to our hypotheses, it is important to note that these effects of Highest Contiguous Next Number, language, and Initial Highest Count were significant when accounting for the effects of trial difficulty, age, and working memory. Thus, the results of our model yield a more nuanced picture of the relationship between count list morphology and numerical knowledge than a comparison of mean performance alone. These data suggest that when controlling for differences in age, working memory, and exposure to the count routine, the mean differences in Unit Task performance across languages may not be best explained by count list transparency, at least not in languages with relatively small differences in count list structure, such as English, Cantonese, and Slovenian.

3.3. Predictors of Next Number performance

3.3.1. Within-language analyses

In the preceding analyses, we found that Highest Contiguous Next Number emerged as the best overall predictor of successor knowledge in a model including all three languages (Cantonese, Slovenian, and English). Further, performance on the Next Number task was significantly related to Unit Task performance in all three languages when analyzed individually,⁷ and was one of the

⁵ Children's overall performance was lower here in comparison to previous work testing the relationship between counting ability and successor knowledge (Cheung et al., 2017) due to our inclusion of (a) both subset and CP knowers in the sample (in contrast to Cheung et al., which included only CP knowers); and (b) extremely large numbers (224, 252, and 271). Despite our inclusion of these large numbers and of subset knowers, we replicated Cheung and colleagues' finding that US children who have higher Initial Highest Counts are significantly more accurate on the Unit Task, even for these very large numbers. Additionally, highly competent counters (with Initial Highest Counts > 80) have near-ceiling performance (89%) on the Unit Task. This analysis is detailed in the Supplementary Online Materials.

⁶ Models were generalized linear mixed effects models with the formula: Correct ~ [Resilience/Final Highest Count/Highest Contiguous Next Number] + Language*Initial Highest Count + Trial Within/Outside Initial Highest Count + Item Magnitude + Age + WPPSI + (1|Subject). Continuous predictors were scaled and centered to facilitate model fit.

⁷ Highest Contiguous Next Number significantly improved the fit of the Cantonese Unit Task base model ($\chi^2_{(1)} = 11.83, p = .0006$), but did not explain unique variance when included with a model containing Initial Highest Count ($\chi^2_{(1)} = 1.86, p = .17$).

Table 5

Cross-linguistic Unit Task regression models with Cantonese (left) and Slovenian (right) selected as a reference group. IHC = Initial Highest Count; HCNN = Highest Contiguous Next Number.

Predictors	Comparison to Cantonese			Comparison to Slovenian		
	β	CI	p	β	CI	p
(Intercept)	0.31	0.11–0.51	0.003	0.50	0.25–0.76	< 0.001
HCNN	0.34	0.20–0.47	< 0.001	0.34	0.20–0.47	< 0.001
Cantonese	—	—	—	–0.19	–0.51–0.12	0.229
Slovenian	0.19	–0.12–0.51	0.229	—	—	—
English (US)	0.43	0.16–0.70	0.002	0.24	–0.08–0.55	0.142
IHC	0.54	0.34–0.75	< 0.001	0.55	0.24–0.86	0.001
Trial Within IHC	0.10	–0.12–0.32	0.388	0.10	–0.12–0.32	0.388
Item Magnitude	–0.44	–0.54–0.34	< 0.001	–0.44	–0.54–0.34	< 0.001
Age	0.10	–0.06–0.25	0.217	0.10	–0.06–0.25	0.217
WPPSI	0.07	–0.04–0.17	0.217	0.07	–0.04–0.17	0.223
Cantonese: IHC	—	—	—	–0.01	–0.34–0.32	0.966
Slovenian: IHC	0.01	–0.32–0.34	0.966	—	—	—
English (US): IHC	0.20	–0.08–0.47	0.162	0.19	–0.17–0.55	0.297

strongest predictors in English and Slovenian. These findings provide support for the proposal that Next Number knowledge is a critical precursor to acquiring the successor function (Barner, 2017; Cheung et al., 2017; Davidson et al., 2012), and also that learning the recursive rules by which number words are productively generated may support the induction that every natural number has a successor. In our next analyses, we explored the Next Number task, testing which candidate measure of counting productivity best predicts children's performance. To do this, we evaluated the relation between children's performance on the Next Number task and Highest Count task. These analyses closely mirror our Unit Task analyses: Within each language, we constructed three models⁸ predicting Next Number performance from (1) Counting Resilience; (2) Final Highest Count; and (3) Initial Highest Count.

These within-language models revealed that counting ability significantly predicted Next Number performance in all three languages, although once again the best predictor differed across languages (Table 6). Final Highest Count was the best predictor of Next Number performance in Slovenian ($\chi^2_{(1)} = 72.93, p < .0001$) and English ($\chi^2_{(1)} = 43.93, p < .0001$), and Initial Highest Count did not significantly improve the fit of these models ($ps > 0.05$). In Cantonese, however, Initial Highest Count was the strongest predictor of Next Number performance ($\chi^2_{(1)} = 68.63, p < .0001$), and Final Highest Count did not explain additional variance ($\chi^2_{(1)} = 0.37, p = .54$). In our next set of analyses, we tested which measure best predicted Next Number performance for the entire dataset.

3.3.2. Cross-linguistic analyses

Perhaps surprisingly, descriptive statistics indicated no evidence of an advantage for Cantonese-speakers in this task. Mean performance was around 50% for Cantonese ($M = 0.49, SD = 0.33$), Slovenian ($M = 0.46, SD = 0.32$), and English ($M = 0.49, SD = 0.35$). As above, we constructed models that predicted Next Number performance across languages: One predicting Next Number performance from Counting Resilience, and another predicting it from Final Highest Count. Importantly, these models tested whether these predictors remained significant when controlling for differences in overall exposure to counting by including an interaction between language and Initial Highest Count. These models also included effects of item magnitude, whether the item was within or outside the child's Initial Highest Count, age, and raw working memory score, with a random effect of subject.

As demonstrated in Table 7, Final Highest Count emerged as the strongest predictor of Next Number performance, and explained significant additional variance in comparison to our base model ($\chi^2_{(1)} = 38.96, p < .0001$); the addition of Resilience did not improve the fit of this model ($\chi^2_{(1)} = 0.27, p = 0.60$). Also, similar to our Unit Task models we found the surprising result that, when accounting for working memory and Initial Highest Count - a measure used by previous studies as a proxy for training with counting - Cantonese-speaking children's performance was actually significantly poorer than that of both English- ($\beta = -1.82, p < .001$) and Slovenian-speaking children ($\beta = 1.77, p < .001$), while there was no difference in performance between English and Slovenian children ($\beta = 0.04, p = .88$, Fig. 2).

4. Discussion

In three languages with relatively transparent counting systems, we investigated the relation between knowledge of count list structure and acquisition of the successor function. Consistent with previous reports, we found striking cross-cultural differences in counting proficiency and numerical knowledge that were broadly related to language transparency: Cantonese-speaking children were able to count much higher than English- and Slovenian-speaking children, and also had greater mean successor task

⁸ Models were generalized linear mixed effects models with the formula: Correct ~ [Resilience/Initial Highest Count/Final Highest Count] + Trial Within/Outside Initial Highest Count + Item Magnitude + Age + (1|Subject). Continuous predictors were scaled and centered to facilitate model fit.

Table 6

Within-language Next Number models for Cantonese, Slovenian, and US English. IHC = Initial Highest Count; FHC = Final Highest Count. Only final models shown for each language: predictors without coefficient estimates did not significantly improve the fit of that language's model in a Likelihood Ratio Test.

Predictors	Cantonese			Slovenian			English (US)		
	β	CI	p	β	CI	p	β	CI	p
(Intercept)	-0.47	-0.78--0.15	0.004	-0.31	-0.62--0.01	0.043	-0.42	-0.75--0.09	0.012
IHC	1.62	1.26--1.98	< 0.001	—	—	—	—	—	—
FHC	—	—	—	1.83	1.43--2.22	< 0.001	1.49	1.06--1.92	< 0.001
Trial Within IHC	0.75	0.26--1.24	0.003	1.19	0.69--1.70	< 0.001	1.64	1.15--2.12	< 0.001
Item Magnitude	-1.12	-1.40--0.84	< 0.001	-1.09	-1.37--0.81	< 0.001	-0.66	-0.89--0.42	< 0.001
Age	0.34	0.01--0.67	0.041	0.47	0.15--0.80	0.004	0.69	0.27--1.11	0.001

Table 7

Cross-linguistic Next Number regression models with Cantonese (left) and Slovenian (right) selected as a reference group. FHC = Final Highest Count; IHC = Initial Highest Count.

Predictors	Comparison to Cantonese			Comparison to Slovenian		
	β	CI	p	β	CI	p
(Intercept)	-1.51	-1.84--1.17	< 0.001	0.27	-0.16--0.69	0.217
FHC	1.05	0.73--1.38	< 0.001	1.05	0.73--1.38	< 0.001
Cantonese	—	—	—	-1.77	-2.31--1.24	< 0.001
Slovenian	1.77	1.24--2.31	< 0.001	—	—	—
English (US)	1.82	1.38--2.25	< 0.001	0.04	-0.49--0.57	0.880
IHC	0.65	0.28--1.02	0.001	0.88	0.28--1.48	0.004
Trial Within IHC	1.18	0.89--1.46	< 0.001	1.18	0.89--1.46	< 0.001
Item Magnitude	-0.92	-1.07--0.77	< 0.001	-0.92	-1.07--0.77	< 0.001
Age	0.61	0.34--0.88	< 0.001	0.61	0.34--0.88	< 0.001
WPPSI	0.15	-0.02--0.33	0.088	0.15	-0.02--0.33	0.088
Cantonese: IHC	—	—	—	-0.23	-0.81--0.34	0.428
Slovenian: IHC	0.23	-0.34--0.81	0.428	—	—	—
English (US): IHC	0.13	-0.31--0.56	0.571	-0.11	-0.71--0.49	0.730

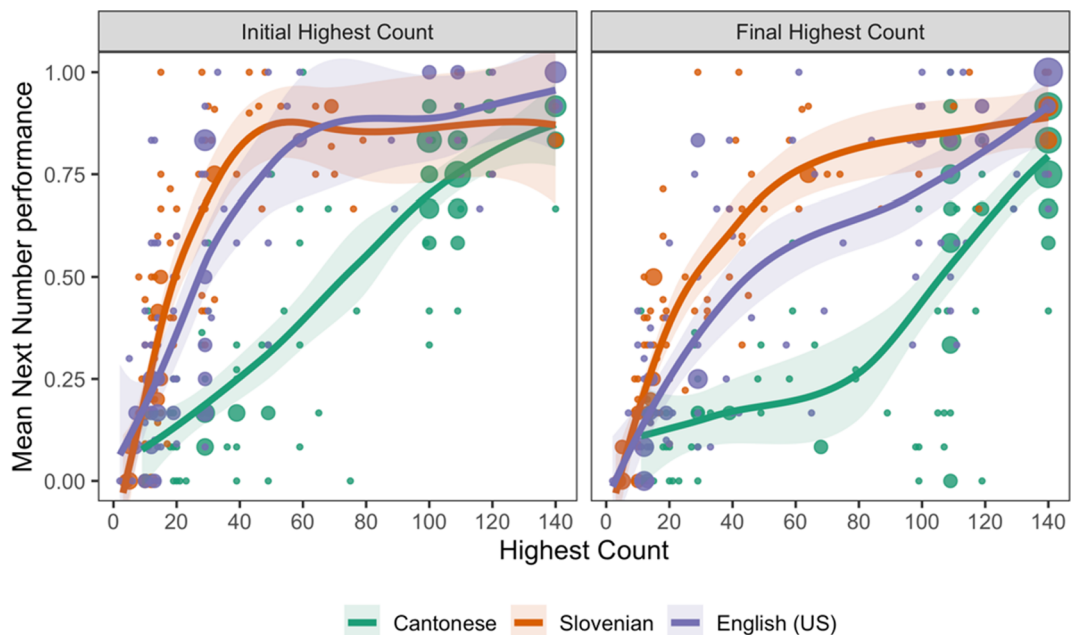


Fig. 2. Scatterplot relating Highest Count (Initial, left and Final, right) to mean Next Number performance by Language. Smooth curve fitted by locally weighted regression, and shaded areas indicate 95% confidence intervals. Size of points indicate frequency of highest count.

performance. Despite these cross-linguistic differences, however, our within-language analyses in Cantonese, Slovenian, and English each found significant relations between measures of counting productivity and performance on the Unit Task, our measure of successor function knowledge. Also, models including all three languages found that performance on the Next Number task, a measure of children's ability to count-up from arbitrary points in the count list, was the best overall predictor of performance on the Unit Task. Finally, we found that while Initial Highest Count predicted performance on both the Unit and Next Number tasks in all three languages, it was never *solely* predictive, and other productivity measures such as Next Number performance or Final Highest Count were better indicators of children's performance on these two tasks in English and Slovenian. In particular, the relationship between children's Next Number and Unit Task performance suggests that successor function knowledge arises in part from knowledge of productive counting rules.

A key element of our design was to assess the role of language structure on number knowledge within languages, rather than relying on cross-linguistic comparisons. The rationale for this was a concern that children learning different languages also typically belong to different cultures, with different practices surrounding early numeracy. Children learning Cantonese, for example, may outperform children learning English not because of the relatively small differences between their structures, but because they receive more intensive math training in early elementary school years (Towse & Saxton, 1998). This concern was vindicated by our analyses that included all three language groups. Although Cantonese-speaking children performed better than other groups on most measures, when factors like working memory, age, and a child's Initial Highest Count (a reasonable proxy for training exposure) were included in models, they were either not different from other groups, or performed significantly worse.

Still, as already noted, the grammatical differences in counting transparency between Cantonese, English, and Slovenian are relatively modest. For the most part, English and Slovenian have relatively transparent decade rules that recur from the 20s through to 100, with only decade labels and syntax presenting points of contrast with Cantonese. Other languages, however, like Hindi and Gujarati, are substantially less transparent, and feature multiple exceptions in every decade all the way to 100. We explore these languages in Experiment 2.

5. Experiment 2

In Experiment 1, we contrasted three languages which, though different with respect to base-system transparency, were nevertheless relatively similar in structure, and generally transparent in nature. In Experiment 2, we investigated two languages that are substantially less transparent in nature: Hindi and Gujarati, as well as a sample of English-speaking Indian children living in the same region of India. Based on Experiment 1, we predicted that in these languages counting productivity would again be predictive of Unit Task performance, but that far fewer children should be productive overall.

5.1. Method

As in Experiment 1, the methods and analyses of this study were pre-registered prior to any data collection. The pre-registration can be found at <https://osf.io/5zxrt>.

5.1.1. Participants

All participants were recruited from Hindi, Gujarati, and English medium schools in Vadodara, Gujarat, India. Our research group has conducted numerous studies in Vadodara for over a decade, and is familiar with local schools and consultants who are knowledgeable about local educational practices. Additionally, the research team included a PhD with expertise in Indian linguistics, as well as three undergraduate research assistants fluent in Hindi and Gujarati. Children's inclusion in a particular dataset (Hindi, Gujarati, or English), was determined by their primary language of instruction. Although relatively few children learn Hindi as a first language in Gujarat, Hindi medium schools exist in the area to serve children of recent immigrants to the region. English medium children spoke Gujarati, Hindi, or another language at home (based on parental report).

We pre-registered a minimum n of 80 per language group to conduct analyses, and a maximum n of 150. We recruited 288 children aged 3.5 to 6.5 years. As per our pre-registration, 47 of these children were excluded from all analyses for missing data from > 20% of trials ($n = 19$); not completing the Highest Count task ($n = 9$); failure to comprehend the task ($n = 8$); missing Highest Count recording ($n = 4$); being out of age range ($n = 3$); experimenter error ($n = 2$); language impairment ($n = 1$); or, for Gujarati and Hindi, native language other than the language of instruction ($n = 1$). In addition to the 47 who were excluded from all analyses, 24 were removed from a subset of analyses if they did not complete the pre-registered minimum number of test trials for that task. After these exclusions, our final sample included 241 participants (Table 8).

Table 8

Demographic information for Hindi, Gujarati, and English (US and Indian).

	n	n CP-knower	M_{age} (SD)	$Median_{age}$
Hindi	91 (41 female, 50 male)	44	5.63 (0.68)	5.81
Gujarati	80 (48 female, 25 male, 7 sex not recorded)	50	5.65 (0.41)	5.72
English (US)	111 (43 female, 68 male)	71	4.79 (0.85)	4.78
English (India)	70 (35 female, 32 male, 3 sex not recorded)	54	5.24 (0.81)	5.37

Table 9

Counting data by language. Initial and Final Highest Count are rounded.

	<i>n</i>	<i>M</i> IHC (<i>SD</i>)	<i>M</i> FHC (<i>SD</i>)	<i>M</i> Prompts (<i>SD</i>)
Hindi				
Overall	91	24 (13.63)	31 (22.99)	2.00 (1.37)
Resilient	8	38 (22.04)	81 (43.28)	4.00 (1.85)
Non-Resilient	83	22 (11.86)	26 (12.42)	1.80 (1.14)
Gujarati				
Overall	80	27 (17.75)	37 (25.87)	2.26 (1.25)
Resilient	11	52 (27.28)	86 (31.34)	3.36 (0.81)
Non-Resilient	69	23 (11.55)	29 (13.64)	2.08 (1.22)
English (India)				
Overall	70	48 (39.64)	75 (48.62)	3.10 (2.12)
Resilient	28	59 (37.58)	117 (22.60)	4.74 (2.31)
Non-Resilient	42	40 (39.57)	47 (40.15)	2.05 (1.08)

These exclusions had their greatest effect on the Indian English dataset, perhaps because this was the only group tested in their second language. Consequently we did not reach our minimum *n* defined in our pre-registration, and instead conducted our primary analyses, as pre-registered, using the US English dataset from Experiment 1. Because US children's performance may substantially differ from that of Hindi- and Gujarati-speaking children for many reasons other than language, we interpreted these results of these analyses with caution, and compared them to *post hoc* analyses which included the Indian English dataset for tasks in which we observed low exclusion rates (Highest Count and Next Number). Thus, our Indian English dataset helped to isolate the role of language in children's performance on these measures.

5.1.2. Stimuli, methods, and procedure

These were identical to Experiment 1, with the exception that audio prompts were translated into the language of instruction. Children were tested individually in a small room set apart from the classroom (in Vadodara and San Diego), or individually in a small lab testing space (in San Diego).

6. Results

6.1. Highest Count

Table 9 shows a breakdown of counting profiles by language and Resilience. Overall, counting proficiency seemed to be strongly related to count list transparency: Hindi- and Gujarati-speaking children had much lower Initial and Final Highest counts in comparison to English-speaking children. As in Experiment 1 these effects of language persisted when children were grouped by Resilience, but to a lesser degree; Resilient counters had higher Initial and Final Highest Counts than Non-Resilient counters in all three languages, although counts still tended to be higher for English-speaking children. Critically, we found that very few children were able to meet the criteria for Resilience in Hindi and Gujarati (9% and 14% of children in each language respectively). On the other hand, 40% of Indian English-speaking children were identified as Resilient, which is similar to the proportion observed in our US English sample (38%).

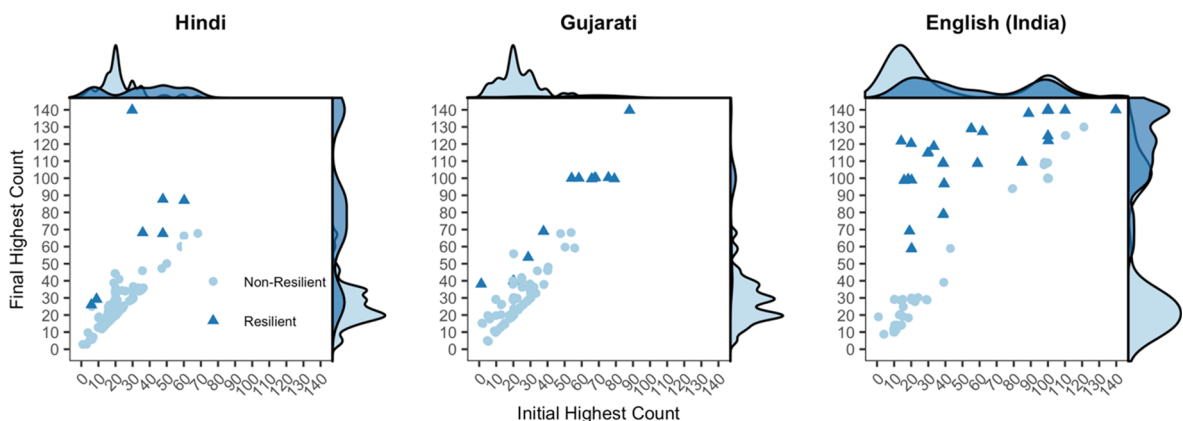


Fig. 3. Initial and Final Highest Counts by language, grouped by Resilience. Points indicate the relation between a participant's Initial and Final Highest Counts. Points are jittered slightly to avoid overplotting. Density plots indicate the distribution of Initial (top) and Final (right) by Resilience.

Table 10

Within-language Unit Task models for Hindi, Gujarati, and US English. Only final models shown for each language: predictors without coefficient estimates did not significantly improve the fit of that language's model in a Likelihood Ratio Test. IHC = Initial Highest Count; FHC = Final Highest Count; HCNN = Highest Contiguous Next Number.

Predictors	Hindi			Gujarati			English (US)		
	β	CI	<i>p</i>	β	CI	<i>p</i>	β	CI	<i>p</i>
(Intercept)	−0.48	−0.67–−0.29	< 0.001	−0.35	−0.54–−0.16	< 0.001	0.58	0.36–0.79	< 0.001
IHC	0.35	0.14–0.56	0.001	—	—	—	0.75	0.44–1.06	< 0.001
FHC	—	—	—	—	—	—	—	—	—
HCNN	0.33	0.12–0.53	0.002	—	—	—	0.43	0.15–0.71	0.003
Trial Within IHC	0.03	−0.35–0.40	0.894	0.54	0.16–0.92	0.005	0.09	−0.29–0.47	0.639
Item Magnitude	−0.39	−0.57–−0.22	< 0.001	−0.16	−0.33–0.01	0.063	−0.48	−0.65–−0.32	< 0.001
Age	0.30	0.13–0.48	0.001	0.19	0.02–0.35	0.026	0.03	−0.20–0.26	0.792

Once again, we found that a majority of children (62% across all three languages) were able to continue counting past their Initial Highest Count if they stopped prior to 140, although very few children counted far enough past their first error to be classified as Resilient in Hindi and Gujarati (Fig. 3). Thus, while Initial Highest Count was strongly correlated with Final Highest Count (English (India): $r = 0.77$, $p < .0001$; Hindi: $r = 0.77$, $p < .0001$; Gujarati: $r = 0.91$, $p < .0001$), we again found evidence that it may underestimate children's productive counting knowledge.

6.2. Predictors of Unit Task performance

6.2.1. Within-language analyses

We next tested whether candidate measures of productive counting were significantly predictive of Unit Task performance in Hindi and Gujarati. We predicted that although children learning these languages may become productive counters later in development, the same relation between productive counting and successor function knowledge should nevertheless exist. The individual measures, model specifications, and hierarchical model comparisons used to assess the relationship between productive counting and successor function knowledge were identical to Experiment 1. Once again, our candidate measures of productivity were (1) Counting Resilience; (2) Final Highest Count; (3) Initial Highest Count; and (4) Highest Contiguous Next Number.

We found evidence of a link between knowledge of productive counting rules and the acquisition of the successor function in Hindi, but not in Gujarati, perhaps because so few Gujarati-speaking children exhibited knowledge of either productive rules or the successor function (see Table 10). In Hindi, Initial Highest Count and Highest Contiguous Next Number were the strongest predictors of Unit Task performance: the addition of Highest Contiguous Next Number significantly improved the fit of a model containing only Initial Highest Count ($\chi^2_{(1)} = 9.91$, $p = .002$), but the addition of Final Highest Count to a model containing both Initial Highest Count and Highest Contiguous Next Number did not ($\chi^2_{(1)} = 0.93$, $p = .34$). However, none of our candidate measures of productivity were related to successor knowledge in Gujarati, though Unit Task performance was predicted by age ($\beta = 0.19$, $p < .03$), and performance was better for items within a participant's Initial Highest Count range ($\beta = 0.54$, $p < .005$).

6.2.2. Cross-linguistic analyses

In our cross-linguistic analyses we tested whether measures of counting productivity were related to Unit Task performance when all languages were included in a single model. Consistent with the hypothesis that learning a more morphologically complex count list may impede extraction of productive counting rules and the successor function, we found lower mean performance for Hindi ($M = 0.40$, $SD = 0.22$) and Gujarati ($M = 0.45$, $SD = 0.18$)⁹ children in comparison to US English ($M = 0.59$, $SD = 0.25$). However, as in Experiment 1, we worried that cross-cultural factors other than language might explain these differences. The following analyses explore this possibility in two steps. First, we built three separate models predicting Unit Task performance from (1) Counting Productivity, (2) Final Highest Count, and (3) Highest Contiguous Next Number in US English, Hindi, and Gujarati. Second, we conducted *post hoc* tests that included subsets of data from children learning Indian English, who are culturally more similar to the Hindi and Gujarati children than US children, but have learned the more regular English count system. Model specifications, measures, and comparison process were identical to Experiment 1.

In the model including Hindi, Gujarati, and US English data, Highest Contiguous Next Number emerged as the strongest predictor of Unit Task performance, and significantly improved model fit compared to our base model ($\chi^2_{(1)} = 13.18$, $p = .0003$; see Table 11). In contrast to Experiment 1, however, mean differences in Unit Task performance by language persisted even when controlling for between-group differences: Hindi- and Gujarati-speaking children had significantly lower Unit Task scores compared to English-

⁹ Hindi- and Gujarati-speaking children's overall below-chance performance on the Unit Task is largely driven by their performance on large numbers (> 100); for these larger numbers, many children gave an answer that was not an available alternative, which was always considered incorrect. Children's performance was not significantly different from chance for items under 100 (Hindi: $t(90) = -1.36$, $p = .18$; Gujarati: $t(79) = 1.62$, $p = .11$), however, indicating that they understood how to use these alternatives.

Table 11

Cross-linguistic Unit Task models with US English as a reference group. HCNN = Highest Contiguous Next Number; IHC = Initial Highest Count.

Predictors	Comparison to English (US)		
	β	CI	p
(Intercept)	0.43	0.23–0.63	< 0.001
HCNN	0.26	0.12–0.40	< 0.001
Hindi	−0.70	−1.02–−0.39	< 0.001
Gujarati	−0.72	−1.02–−0.43	< 0.001
IHC	0.45	0.27–0.63	< 0.001
Within IHC	0.16	−0.06–0.39	0.148
Item Magnitude	−0.37	−0.47–−0.27	< 0.001
Age	0.27	0.11–0.44	0.001
WPPSI	0.10	−0.01–0.21	0.071
Hindi:IHC	0.45	0.04–0.87	0.032
Gujarati:IHC	−0.42	−0.75–−0.09	0.012

speaking US children (Hindi: $\beta = -0.70$, $p < .0001$; Gujarati: $\beta = -0.72$, $p < .0001$). This much lower performance on the Unit Task for Hindi- and Gujarati-speaking children in comparison to English-speaking children suggests that acquiring the successor function may be more difficult in languages in which the recursion of the count list is less easily discoverable due to less transparent morphology. However, this conclusion is tempered by the fact that only US data were available for this particular analysis. To further probe whether this difference might be due to language in particular, our next analyses, which focused on the Next Number task, included *post hoc* tests with Indian English data.

6.3. Predictors of Next Number performance

6.3.1. Within-language analyses

The results of our Unit Task analyses indicated that although productive counting is related to successor knowledge in some children learning opaque count lists, this connection is somewhat more fragile. One potential factor limiting our ability to detect effects was the relative infrequency of children who demonstrated productive counting knowledge in these languages: In Hindi only 8 out of 91 children were classified as Resilient, and in Gujarati, only 11 out of 80. Thus, we surely lacked power to reliably detect relations between productivity and other outcomes. Nevertheless, in both Hindi and our cross-linguistic Unit Task analyses, Highest Contiguous Next Number again significantly predicted successor knowledge. As in Experiment 1, we next explored the best predictors of Next Number performance. We again constructed three within-language models as in Experiment 1, predicting Next Number performance from (1) Counting Resilience; (2) Final Highest Count; and (3) Initial Highest Count. Because we observed higher rates of comprehension in our Indian English sample for this task relative to other tasks, we include the results of their within-language analyses here.

Despite the lower levels of counting proficiency in Hindi and Gujarati, we again found that counting ability was significantly predictive of Next Number performance in all languages (see Table 12). In Hindi, although Initial Highest Count, Final Highest Count, and Resilience were each significantly related to Next Number performance, overall Initial Highest Count was the strongest predictor

Table 12

Within-language Next Number models for Hindi, Gujarati, and Indian English. IHC = Initial Highest Count; FHC = Final Highest Count. Only final models shown for each language: predictors without coefficient estimates did not significantly improve the fit of that language's model in a Likelihood Ratio Test.

Predictors	Hindi			Gujarati			English (India)		
	β	CI	p	β	CI	p	β	CI	p
(Intercept)	−1.88	−2.34–−1.43	< 0.001	−1.24	−1.69–−0.78	< 0.001	−0.76	−1.13–−0.39	< 0.001
IHC	1.41	0.93–1.88	< 0.001	—	—	—	0.64	0.13–1.15	0.014
FHC	—	—	—	1.78	1.13–2.42	< 0.001	0.85	0.25–1.44	0.005
Resilient	—	—	—	−1.75	−3.46–−0.04	0.045	—	—	—
Trial Within IHC	1.68	1.14–2.22	< 0.001	2.00	1.43–2.57	< 0.001	1.11	0.52–1.70	< 0.001
Item Magnitude	−0.92	−1.24–−0.61	< 0.001	−1.25	−1.59–−0.91	< 0.001	−1.37	−1.72–−1.02	< 0.001
Age	0.53	0.09–0.98	0.020	0.37	−0.02–0.75	0.064	0.18	−0.30–0.67	0.458

Table 13

Cross-linguistic Next Number models with US English (left) and Indian English (right) selected as a reference group. FHC = Final Highest Count; IHC = Initial Highest Count.

Predictors	Comparison to English (US)			Comparison to English (India)		
	β	CI	p	β	CI	p
(Intercept)	−0.63	−1.06–−0.19	0.005	−1.33	−1.88–−0.77	< 0.001
FHC	0.98	0.58–1.38	< 0.001	0.66	0.21–1.11	0.004
Hindi	−0.53	−1.24–0.17	0.136	0.30	−0.46–1.07	0.438
Gujarati	−0.53	−1.18–0.12	0.112	0.29	−0.42–1.01	0.422
IHC	0.42	0.00–0.84	0.047	0.25	−0.18–0.68	0.262
Trial Within IHC	1.73	1.42–2.04	< 0.001	1.66	1.33–2.00	< 0.001
Item Magnitude	−0.91	−1.07–−0.74	< 0.001	−1.21	−1.41–−1.01	< 0.001
Age	0.80	0.43–1.16	< 0.001	0.57	0.16–0.99	0.006
WPPSI	0.15	−0.08–0.37	0.210	0.13	−0.11–0.37	0.293
Hindi:IHC	1.48	0.58–2.38	0.001	2.05	1.19–2.90	< 0.001
Gujarati:IHC	0.45	−0.28–1.18	0.227	1.00	0.29–1.71	0.006

($\chi^2_{(1)} = 33.41, p < .0001$), and neither Final Highest Count ($\chi^2_{(1)} = 2.77, p = .10$) nor Resilience ($\chi^2_{(1)} = 0.37, p = .54$) significantly improved model fit. In Gujarati, Final Highest Count and Resilience were the best predictors of performance: Resilience significantly improved the fit of a model containing Final Highest Count ($\chi^2_{(1)} = 4.12, p = .04$), and Initial Highest Count did not add to this model ($\chi^2_{(1)} = 0.08, p = .35$). Finally, in Indian English, both Initial and Final Highest count produced the best fit to the data in comparison to a model containing Final Highest Count alone. ($\chi^2_{(1)} = 5.88, p = .02$).

6.3.2. Cross-linguistic analyses

Although we found that counting ability predicted Next Number performance in our within-language analyses, mean performance on the Next Number task was lower in Hindi ($M = 0.32, SD = 0.30$) and Gujarati ($M = 0.38, SD = 0.26$) in comparison to both US English ($M = 0.49, SD = 0.35$) and Indian English ($M = 0.45, SD = 0.29$), which did not significantly differ from one another ($t(177) = 0.88, p = .38$). As pre-registered, we constructed cross-linguistic models using our US English dataset. Because we observed

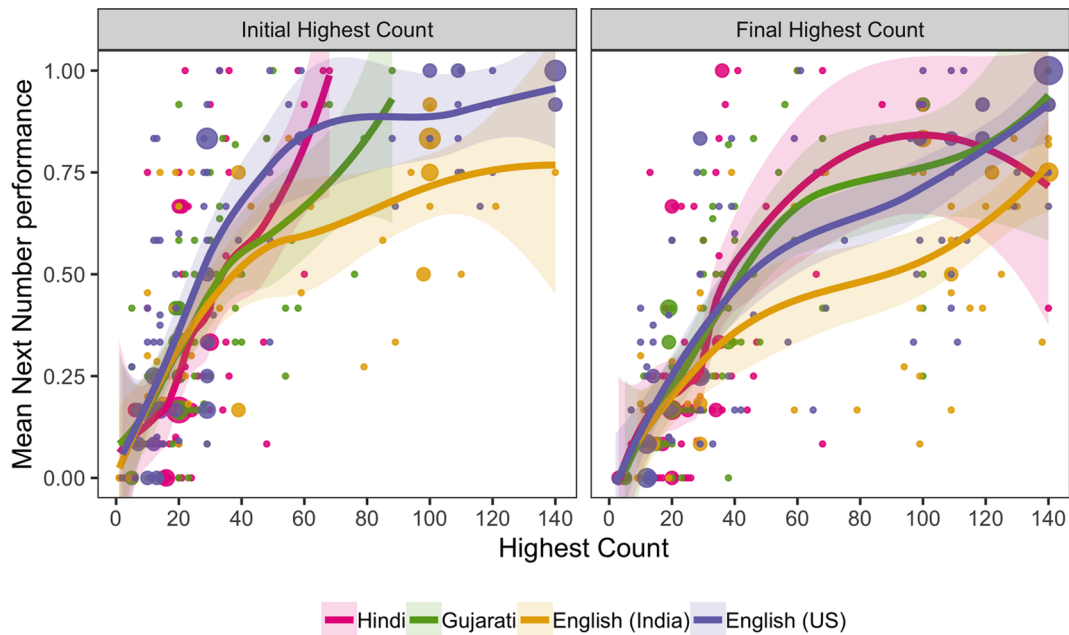


Fig. 4. Scatterplot relating Highest Count (Initial, left and Final, right) to mean Next Number performance by language. Smooth curve fitted by locally weighted regression, and shaded areas indicate 95% confidence intervals. Size of points indicate frequency of highest count.

a much higher rate of comprehension for the Next Number task in our Indian English sample, however, we also report *post hoc* analyses including these data in an attempt to isolate the effects of language transparency versus other cultural factors. In these analyses, we again controlled for between-group differences by including an interaction between language and Initial Highest Count, as well as a nonverbal working memory term. Once again, the model specifications and comparison process were identical to Experiment 1. Using this base model we constructed two generalized linear mixed effects models predicting Next Number performance from (1) Counting Resilience; and (2) Final Highest Count.

As shown in Table 13, Final Highest Count was the strongest single predictor of Next Number performance in Hindi, Gujarati, and US English, significantly improving the fit of the base model ($\chi^2_{(1)} = 22.68, p < .0001$). A follow-up analysis which substituted our Indian for US English dataset as the comparison group replicated this effect ($\chi^2_{(1)} = 8.36, p = .004$). There was no effect of language in either the US or Indian English models (Fig. 4), suggesting that although counting mastery may vary across languages due to morphological complexity, the best predictor of performance on the Next Number task is nevertheless children's knowledge of productive counting rules, reflected here by their Final Highest Count.

7. Discussion

In Experiment 2, we investigated children learning Hindi and Gujarati, two languages that are substantially less transparent in nature than any of the languages studied in Experiment 1, to explore whether acquiring the successor function was made more difficult by a more complex count list. We again found substantial differences in counting ability and successor knowledge between transparent and opaque language groups, and these differences were greater than those reported in Experiment 1, even in a within-culture comparison. Using a within-language approach, however, we found that productive counting knowledge was related to Unit Task performance, despite lower numbers of productive counters in Hindi and Gujarati. Although this evidence was less robust than in Experiment 1, likely due to overall lower levels of counting ability in Hindi and Gujarati, we again found that knowledge of productive counting was significantly related to successor knowledge in Hindi. While the best predictors of Unit Task performance differed across our within-language models, once again our combined cross-linguistic models revealed Next Number performance, an indicator of children's ability to count up from an arbitrary point in the count list, as the strongest indicator of children's successor knowledge. Similarly, although there was some variability in the counting measures related to Next Number performance in within-language models, Final Highest Count emerged as the best predictor on this task in our cross-linguistic models. These results mirror our findings from Experiment 1; once again, we find that two strong indicators of children's productive counting knowledge are significantly related to their performance on the Unit and Next Number tasks.

In contrast to Experiment 1, however, our cross-linguistic models also revealed strong effects of count list transparency. Children learning less complex count lists seemed to benefit from this transparency in acquiring the successor function, even when controlling for factors such as age, working memory, and individual differences in Initial Highest Count, such that English-speaking US children performed better on the Unit Task relative to Hindi and Gujarati children. Thus, unlike in Experiment 1 where mean cross-linguistic differences disappeared when accounting for other factors, the significantly lower performance of Hindi and Gujarati children in comparison to English-speaking US children on the Unit Task suggests that extremely opaque count lists may confer a substantial disadvantage in acquiring other numerical knowledge. While these findings are tempered by the fact that we were unable to make a within-culture comparison for the Unit Task using our Indian English sample, *post hoc* comparisons indicated that these children's performance on the Highest Count and Next Number tasks was comparable to US English children, despite being tested in their second language.

8. General discussion

Using a large cross-linguistic dataset drawn from five languages across four cultures, we tested the hypothesis that children acquire the successor function through learning the productive morphological rules of their language's count list, while also exploring which candidate measures of productivity best predict successor function knowledge. We found large mean differences in counting ability and successor knowledge between languages with relatively transparent counting systems (Cantonese, Slovenian, English) compared to those with more opaque counting structure (Hindi, Gujarati). Also, despite these cross-linguistic differences, we found that productive counting knowledge was strongly related to successor function knowledge within almost every language. This was true even in a highly opaque language like Hindi where, despite the complexity posed by their count lists, a small number of children were nevertheless able to learn a productive rule, and exhibited stronger successor function knowledge.

In addition to these main findings, our study also revealed several important secondary results. First, although measures of counting productivity were related to successor function knowledge across languages, there was important variability in which measures of counting productivity best predicted successor function knowledge. Second, despite these interesting differences, when all languages were considered together, children's successor function knowledge in both Experiments 1 and 2 was best predicted by performance on the Next Number task - i.e., the ability to name the next number in the count list without counting up from 1. Third, somewhat surprisingly, we found that although, like in past reports, Cantonese-speaking children outperformed English-speaking children along several measures of number knowledge, these differences disappeared or were reversed when other factors, such as working memory and amount of counting exposure, were considered. Similarly, although Slovenian is more transparent than English,

Slovenian children exhibited much more limited counting abilities than English-speaking children. These results lead us to conclude that some previously reported cross-linguistic differences may be not be due to language alone, but may also be importantly affected by cultural differences in math practices. At the very least our results suggest that plain comparisons of counting ability across different languages should not be interpreted without also considering other factors that might explain differences.

This work began with two observations. First, previous work reported that children who learn relatively transparent languages, like Cantonese, make fewer counting errors than children learning English, and may be quicker to acquire early mathematical abilities. Second, previous studies found that children's acquisition of the successor function, as measured by the Unit Task, is strongly predicted by how high they can count. Together, these two observations led us to hypothesize that exposure to counting might lead children to move beyond a memorized list to derive productive rules that allow them to count indefinitely high, and that these rules might be the basis for deriving a recursive successor function: Knowledge that counting *words* are governed by a rule-like structure might lead children to infer that numbers themselves are generated by a recursive '+ 1' rule. To investigate this question, we took a novel approach. Although past work has found connections between the regularity of counting systems, counting proficiency, and mathematical achievement, none of this work has provided direct evidence that such effects are actually due to differences in how easily children are able to extract productive counting rules through exploring individual differences within a particular culture. For example, although previous findings (Fuson & Kwon, 1992; Miller, Kelly, & Zhou, 2005; Miller & Stigler, 1987; Miura et al., 1988; Miura & Okamoto, 1989) are consistent with the idea that such advantages are due to count list transparency, it is possible that they may also reflect differences in the levels of counting, number, and mathematics exposure across these groups (Pan et al., 2006; Towse & Saxton, 1998). Very generally, many cross-cultural differences including language, mathematics curriculum, and societal attitudes toward the importance of early math education may impact children's early counting fluency without necessarily implicating children's ability to detect recursive counting rules (Lefevre et al., 2002). If previously attested differences in mathematics learning result from the impact of counting transparency on the acquisition of counting rules, then children learning transparent counting systems should be faster to extract recursive rules from their count list, and should be faster to apply these rules to reasoning about simple addition facts, like those tested by Sarnecka and Carey (2008) Unit Task.

Our data suggest that the transparency of a child's count list likely does affect how readily they extract productive counting rules, but that this is not the only factor, and that training amount may overwhelm differences in transparency when grammatical differences in count structure are small. First, we found that when languages exhibited smaller differences in count list structure between them, as in the case of Slovenian, English, and Cantonese, these differences were not the best predictors of performance. For example, whereas English exhibits exceptions in the teens (*eleven, twelve, thirteen*) as well as on decade labels (*twenty, thirty, fifty*), Slovenian only has exceptions in the teens and one decade label (*twenty*), but nevertheless Slovenian children performed worse than US children on most tasks. One obvious account of why this might be is that Slovenian children likely receive much less exposure to counting in their preschool years, as evidenced by their very low Initial Highest Counts relative to English and Cantonese. Compatible with this, previous studies find that whereas US 5-year-olds typically have an initial highest count up to about 40 or 50 on average, Slovenian children at the same age can only count to about 10 before making their first error (Almoammer et al., 2013; Marušič et al., 2016). Our data also show that Cantonese children don't exhibit a general advantage over US children, particularly when we account for differences in working memory and Initial Highest Count. On the other hand, languages with more extreme morphological complexities were associated with lower performance on these tasks in comparison to more transparent languages, despite similar Initial Highest Count performance. We found that very few learners of Hindi and Gujarati were able to count very far beyond their Initial Highest Count, even when given a prompt, and that these groups performed significantly worse on the Unit Task than English-speaking US children.¹⁰ These data collectively support the idea that counting transparency likely plays a role in children's ability to extract productive counting rules when differences in transparency are significant, and when other factors, like training amount, don't compensate for differences in counting structure. Taken together, our results suggest that when making cross-cultural comparisons, we can't assume that differences between cultures are straightforwardly predicted by differences in grammatical structure, since other factors, like amount of input, surely also play a role.

One important lesson from these studies is that a common measure of children's counting ability - i.e., their Initial Highest Count - often underestimates children's counting ability, and provides a less powerful predictor of other abilities than do measures that are more sensitive to counting productivity. Our data show that, even within a language, two children who have an Initial Highest Count of 30 may have qualitatively different understanding of counting and the rules that govern it. One such child may have rote memorized all numbers up to ~30, without having extracted any rules to describe the count structure. Such children were frequent in Hindi and Gujarati, and surprisingly, in Cantonese, where children were often able to count quite high without error, yet still performed poorly on the Next Number task. Another type of child who counted up to ~30, however, may have noticed the recurrence of the numbers 1–9 in each of the first two decades and extracted a rule, stopping at 30 due to a random error - a common pattern in English, Cantonese, and Slovenian. Our methods allowed us to differentiate these two types of children by providing a prompt and asking whether they could continue: Children who have memorized up to 30 should have no idea what to do next, whereas children who have a rule may be able to recover and count up. Not only did we find this to be the case - that a large percentage of children who

¹⁰ A teacher survey found that teachers expected their Hindi and Gujarati students to be able to count as high as students in English medium schools, and also that Hindi, Gujarati, and English medium school teachers spend similar amounts of time on number and counting instruction in the classroom. Further, Hindi and Gujarati children often recited their count list as a memorized routine, indicating a high level of (rote) training. Despite their relatively high exposure to the count list, however, we still found very few Resilient counters in these two languages, and perhaps because of this, observed much lower mean performance on both the Unit and Next Number tasks.

initially counted to a relatively small number in fact had a productive counting rule (e.g., about 40% of English-speaking children) - we also found that this difference between children who could count-up vs. those who could not (i.e., what we termed “Resilience”) was predictive of performance on both the Next Number task and the Unit Task. Although we found that children’s Initial Highest Count is correlated with other measures of productivity, and can even be used to meaningfully predict their Unit Task performance, as in Cheung et al. (2017), our data provide strong evidence that this measure cannot alone identify productivity, and is generally not the strongest measure of children’s knowledge of rules.

An additional goal of the current work was exploratory, and sought to test several hypothesized measures of productivity in addition to Initial Highest Count to identify one which was robustly predictive across five diverse language groups. Both within several languages and cross-linguistically we found that Next Number performance was one of the strongest predictors of whether a child was successful on the Unit Task. Although these converging findings indicate the strength of this measure for assessing productivity, we also found that additional measures of productivity (such as Initial or Final Highest Count) explained Unit Task performance in several within-language analyses. In one sense, it is unsurprising that we should find that other productivity measures might jointly predict successor knowledge along with Next Number performance in some languages, as these measures were designed to capture the same construct. However, we also noticed some degree of variability across languages with respect to the relation between our different predictors. For example, in Slovenian, where children were comparatively weaker rote counters, we found that measures which incorporated either some amount of support (such as Final Highest Count) or did not require rote counting (such as the Next Number task) were much stronger predictors of successor knowledge than Initial Highest Count. In contrast we found that for US children, who have higher levels of rote counting training overall, Initial Highest Count and Next Number performance were the best predictors. Such differences are to be expected if children attain productivity via different routes in different languages - e.g., if children become productive with less rote memorization in Slovenian than in English because of its slightly greater transparency, or if children receive massive amounts of rote training in early childhood.

Across a diverse set of language groups, our findings suggest that when a language provides the basis for acquiring a productive rule for describing a count list’s base system, children can learn this rule and use it to generate very large numbers. In the Introduction, we speculated that acquiring such a productive rule might provide the basis for learning not only how to count, but also for discovering the recursive nature of the integers themselves (i.e., the concepts that are denoted by number words). Previous accounts of number word learning hypothesized that children might acquire knowledge of the successor function - and thus of integer concepts - using a form of analogical mapping defined over small number words (Carey, 2004, 2009; Gentner, 2010; Wynn, 1992). On this view, a child who has learned a handful of number words might notice an analogy “between the magnitudinal relationships of their own representations of numerosities, and the positional relationships of the number words” (Wynn, 1992, p.250), such that “she is in the position to make the crucial induction: For any word on the list whose quantificational meaning is known, the next word on the list refers to a set with another individual added” (Carey, 2004, p.67). Originally, this idea was proposed as an account of how children might become CP-knowers, since learning this type of analogical mapping would allow children to use a counting procedure to accurately give sets of any size within their count list. Although we now know that children fail to acquire successor function knowledge at this stage (Cheung et al., 2017; Davidson et al., 2012; Spaepen et al., 2018), it remains possible that this model might still explain how older children learn to interpret their count list. However, a critical problem with this idea is that, for children who have a finite count list, an inductive inference that applies to “all” numbers need not take the form of a recursive function that can generate an infinite number of numbers. Whereas the Peano-Dedekind axioms state that *every* number has a successor, the analogical mapping hypothesis described by Carey and others generates a much weaker inductive inference - i.e., that for all numbers, the successor of N in the count list has a cardinal value of $N + 1$. Whereas a child who has acquired a productive morphological rule for generating indefinitely many number words might take “all” numbers to be unbounded - and thus requiring a recursive rule - a child who knows only 30–40 number words and who believes that numbers are finite — as most 3- and 4-year-olds do (Cheung et al., 2017; Evans, 1983; Gelman, 1980; Hartnett & Gelman, 1998) — might have no basis for inducing a fully recursive successor function (for a related point, see Rips, Asmuth, & Bloomfield, 2006). Said otherwise, if children acquire the successor function by making an analogy between the structure of the count list and the relations between integer concepts, then learning that count words are generated by recursive rules might provide a basis for inferring that the integers, themselves, are governed by recursive rules.

Much remains to be discovered about the process by which children acquire successor function knowledge. For example, while it is plausible that rules governing counting might be analogically extended to cardinal values, more evidence of this is needed. Many alternative hypotheses are possible, including the possibility that explicit arithmetic training - e.g., on problems like $2 + 1$, $3 + 1$, $12 + 1$, etc., forms the basis for an inductive inference supporting the successor function, and happens to co-occur developmentally with greater counting abilities (see Barner, 2017; Secada et al., 1983). Alternatively, performance on the Unit Task may simply be easier once children have acquired a productive counting rule, allowing them to deploy working memory resources previously devoted to tracking their position in the count list to the problem of reasoning about the corresponding set operations. Although all of our cross-linguistic models controlled for working memory, it remains possible that more subtle measures of how working memory is deployed during the Unit Task might find differences in overall load when children use a memorized list vs. one that is governed by rules. Future studies should explore these questions, while also investigated a broader range of languages, using both correlational and experimental designs.

Acknowledgements

We would like to thank Pierina Cheung, Kaiqi Guo, Calvin Lee, Sneha Patel, Shruti Shah, Soham Shah, Ashlie Pankonin, Sara Lee, Samuel Lucero, Hortensia Flores, Sonora Grimstead, Kaitlyn Seifert, Samuel Beech, and Sarah Valadez for their assistance with data

collection. We are grateful to our testing sites in Hong Kong, Nova Gorica, Vadodara, and San Diego, and to all of our participants for their time. This work was supported by the National Science Foundation [BCS-1749518, DGE-1321851] and the Javna Agencija za Raziskovalno Dejavnost RS [P6-0382].

Appendix A

From Comrie (2011): Demonstration of Hindi numeral irregularity between 1 and 99.

	0	1	2	3	4	5	6	7	8	9
-		ek	do	tīn	cār	pāñc	chah	sāt	āṭh	nau
10	das	gyārah	bārah	terah	caudah	pandrah	solah	satrah	aṭhārah	unnīs
20	bīs	ikkīs	bāīs	teīs	caubīs	paccīs	chabbīs	sattāīs	aṭṭāīs	untīs
30	tīs	ikattīs	battīs	tairtīs	caumtīs	pairntīs	chattīs	sairntīs	artīs	untālīs
40	cālīs	iktālīs	bayālīs	tairntālīs	cavālīs	pairntālīs	chiyālīs	sairntālīs	artālīs	uncās
50	pacās	ikyāvan	bāvan	tirpan	cauvan	pacpan	chappan	sattāvan	aṭṭhāvan	unsath
60	sāth	iksath	bāsath	tirsath	caumsath	paimsath	chiyāsath	sarsath	arsath	unhattar
70	sattar	ikhattar	bahattar	tihattar	cauhattar	pachattar	chihattar	sathattar	aṭhhattar	unyāsī
80	assī	ikyāsī	bayāsī	tirāsī	caurāsī	pacāsī	chiyāsī	sattāsī	aṭṭhāsī	navāsī
90	nave	ikyānve	bānve	tirānve	caurānve	pacānve	chiyānve	sattānve	aṭṭhānve	ninyānve

Appendix B. Supplementary material

Supplementary data to this article can be found online at <https://doi.org/10.1016/j.cogpsych.2019.101263>.

References

- Almoammer, A., Sullivan, J., Donlan, C., Marušič, F., O'Donnell, T., & Barner, D. (2013). Grammatical morphology as a source of early number word meanings. *Proceedings of the National Academy of Sciences*, 110(46), 18448–18453.
- Barner, D. (2017). Language, procedures, and the non-perceptual origin of number word meanings. *Journal of child language*, 44(3), 553–590.
- Barr, D. J., Levy, R., Scheepers, C., & Tily, H. J. (2013). Random effects structure for confirmatory hypothesis testing: Keep it maximal. *Journal of memory and language*, 68(3), 255–278.
- Barth, H., Starr, A., & Sullivan, J. (2009). Children's mappings of large number words to numerosities. *Cognitive Development*, 24(3), 248–264.
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2014). Fitting linear mixed-effects models using lme4. arXiv preprint arXiv:1406.5823.
- Berger, H. (1992). Modern Indo-aryan. *Indo-European Numerals*, 57, 243–287.
- Bright, W. (1969). *Hindi numerals*. Washington, D.C.: Distributed by ERIC Clearinghouse.
- Carey, S. (2004). Bootstrapping & the origin of concepts. *Daedalus*, 133(1), 59–68.
- Carey, S. (2009). Where our number concepts come from. *The Journal of philosophy*, 106(4), 220.
- Cheung, P., Rubenson, M., & Barner, D. (2017). To infinity and beyond: Children generalize the successor function to all possible numbers years after learning to count. *Cognitive psychology*, 92, 22–36.
- Chomsky, N. (1965). *Aspects of the Theory of Syntax*. Cambridge, MA: MIT Press.
- Comrie, B. (2011). Typology of numeral systems. *Numeral types and changes worldwide. Trends in Linguistics. Studies and monographs*, 118.
- Davidson, K., Eng, K., & Barner, D. (2012). Does learning to count involve a semantic induction? *Cognition*, 123(1), 162–173.
- Dowker, A., Bala, S., & Lloyd, D. (2008). Linguistic influences on mathematical development: How important is the transparency of the counting system? *Philosophical Psychology*, 21(4), 523–538.
- Evans, D. W. (1983). *Understanding zero and infinity in the early school years. Unpublished doctoral dissertation*. University of Pennsylvania.
- Fuson, K. C. (1988). *Children's counting and concepts of number*. New York: Springer-Verlag.
- Fuson, K. C., & Kwon, Y. (1992). Korean children's understanding of multidigit addition and subtraction. *Child Development*, 63(2), 491–506.
- Fuson, K. C., Richards, J., & Briars, D. J. (1982). The acquisition and elaboration of the number word sequence. *Children's logical and mathematical cognition* (pp. 33–92). New York, NY: Springer.
- Geary, D. C. (2018). Growth of symbolic number knowledge accelerates after children understand cardinality. *Cognition*, 177, 69–78.
- Geary, D. C., vanMarle, K., Chu, F. W., Rouder, J., Hoard, M. K., & Nugent, L. (2018). Early conceptual understanding of cardinality predicts superior school-entry number-system knowledge. *Psychological science*, 29(2), 191–205.
- Gelman, R., & Gallistel, C. R. (1978). *The child's concept of number*. Cambridge, MA: Harvard.
- Gelman, R. (1980). What young children know about numbers. *Educational Psychologist*, 15, 54–68.
- Gentner, D. (2010). Bootstrapping the mind: Analogical processes and symbol systems. *Cognitive Science*, 34(5), 752–775.
- Gould, P. (2017). Mapping the acquisition of the number word sequence in the first year of school. *Mathematics Education Research Journal*, 29(1), 93–112.
- Hartnett, P., & Gelman, R. (1998). Early understandings of numbers: Paths or barriers to the construction of new understandings? *Learning and Instruction*, 8(4), 341–374.
- Hurford, J. R. (1987). *Language and number: The emergence of a cognitive system*. Oxford: Basil Blackwell.
- Le Corre, M., & Carey, S. (2007). One, two, three, four, nothing more: An investigation of the conceptual sources of the verbal counting principles. *Cognition*, 105(2), 395–438.
- Le Corre, M., Van de Walle, G., Brannon, E. M., & Carey, S. (2006). Re-visiting the competence/performance debate in the acquisition of the counting principles. *Cognitive psychology*, 52(2), 130–169.
- Lefevre, J. A., Clarke, T., & Stringer, A. P. (2002). Influences of language and parental involvement on the development of counting skills: Comparisons of French- and English-speaking Canadian children. *Early Child Development and Care*, 172(3), 283–300.
- Manolitsis, G., Georgiou, G. K., & Tziraki, N. (2013). Examining the effects of home literacy and numeracy environment on early reading and math acquisition. *Early Childhood Research Quarterly*, 28(4), 692–703.
- Marchand, E., & Barner, D. (2018). Analogical mapping in numerical development. *Language and Culture in Mathematical Cognition* (pp. 31–47). Academic Press.
- Marušič, F., Plesničar, V., Razboršek, T., Sullivan, J., & Barner, D. (2016). Does grammatical structure accelerate number word learning? Evidence from learners of dual and non-dual dialects of Slovenian. *PLoS one*, 11(8), e0159208.
- Miller, K. F., Kelly, M., & Zhou, X. (2005). Learning mathematics in China and the United States: Cross-cultural insights into the nature and course of preschool mathematical development. In J. I. D. Campbell (Ed.). *Handbook of mathematical cognition* (pp. 163–177). New York, NY, US: Psychology Press.
- Miller, K. F., Smith, C. M., Zhu, J., & Zhang, H. (1995). Preschool origins of cross-national differences in mathematical competence: The role of number-naming

- systems. *Psychological Science*, 6(1), 56–60.
- Miller, K. F., & Stigler, J. W. (1987). Counting in Chinese: Cultural variation in a basic cognitive skill. *Cognitive Development*, 2(3), 279–305.
- Miura, I. T., Kim, C. C., Chang, C. M., & Okamoto, Y. (1988). Effects of language characteristics on children's cognitive representation of number: Cross-national comparisons. *Child Development*, 1445–1450.
- Miura, I. T., & Okamoto, Y. (1989). Comparisons of U.S. and Japanese First Graders' cognitive representation of number and understanding of place value. *Journal of Educational Psychology*, 81(1), 109–114.
- Pan, Y., Gauvain, M., Liu, Z., & Cheng, L. (2006). American and Chinese parental involvement in young children's mathematics learning. *Cognitive Development*, 21(1), 17–35.
- Rips, L. J., Asmuth, J., & Bloomfield, A. (2006). Giving the boot to the bootstrap: How not to learn the natural numbers. *Cognition*, 101(3), B51–B60.
- Rule, J., Dechter, E., & Tenenbaum, J. B. (2015). Representing and learning a large system of number concepts with latent predicate networks. *CogSci* (pp. 2051–2056).
- Saxe, G. B., Guberman, S. R., Gearhart, M., Gelman, R., Massey, C. M., & Rogoff, B. (1987). Social processes in early number development. *Monographs of the Society for Research in Child Development* i-162.
- Sarnecka, B. W., & Carey, S. (2008). How counting represents number: What children must learn and when they learn it. *Cognition*, 108(3), 662–674.
- Secada, W. G., Fuson, K. C., & Hall, J. W. (1983). The transition from counting-all to counting-on in addition. *Journal for Research in Mathematics Education*, 47–57.
- Siegler, R. S., Duncan, G. J., Davis-Kean, P. E., Duckworth, K., Claessens, A., Engel, M., ... Chen, M. (2012). Early predictors of high school mathematics achievement. *Psychological science*, 23(7), 691–697.
- Siegler, R. S., & Robinson, M. (1982). The development of numerical understandings. *Advances in child development and behavior* (pp. 241–312). JAI.
- Spaepen, E., Gunderson, E. A., Gibson, D. G., Goldin-Meadow, S., & Levine, S. C. (2018). Meaning before order: Cardinal principle knowledge predicts improvement in understanding the successor principle and exact ordering. *Cognition*, 180, 59–81.
- Towse, J., & Saxton, M. (1998). *Mathematics across national boundaries: Cultural and linguistic perspectives on numerical competence*.
- von Humboldt, W. F. (1999). *On language: On the diversity of human language construction and its influence on the mental development of the human species*. Cambridge University Press.
- Wagner, K., Kimura, K., Cheung, P., & Barner, D. (2015). Why is number word learning hard? Evidence from bilingual learners. *Cognitive Psychology*, 83, 1–21.
- Wagner, K., Chu, J., & Barner, D. (2019). Do children's number words begin noisy? *Developmental Science*, 22(1), e12752.
- Watts, T. W., Duncan, G. J., Siegler, R. S., & Davis-Kean, P. E. (2014). What's past is prologue: Relations between early mathematics knowledge and high school achievement. *Educational Researcher*, 43(7), 352–360.
- Wechsler, D., & Corporation, Psychological (2012). *WPPSI-IV: Wechsler preschool and primary scale of intelligence – (4th ed.)*. Bloomington, MN: Pearson, Psychological Corporation.
- Wright, R. J. (1994). A study of the numerical development of 5-year-olds and 6-year-olds. *Educational Studies in Mathematics*, 26(1), 25–44.
- Wynn, K. (1990). Children's understanding of counting. *Cognition*, 36(2), 155–193.
- Wynn, K. (1992). Children's acquisition of the number words and the counting system. *Cognitive psychology*, 24(2), 220–251.
- Yang, C. (2016). The linguistic origin of the next number. *Manuscript*.