

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Do cross-linguistic animacy-based constraints on plural marking reflect learning biases?

Permalink

<https://escholarship.org/uc/item/4nm5r85z>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Saldana, Carmen

Ivani, Jessica K.

Franzon, Dr. Francesca

Publication Date

2026

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Do cross-linguistic animacy-based constraints on plural marking reflect learning biases?

Carmen Saldana (carmen.saldana@ub.edu)^{1,2}, Jessica K. Ivani³ & Francesca Franzon²

¹Department of Catalan Philology and General Linguistics, Universitat de Barcelona

²Department of Translation and Language Sciences, Universitat Pompeu Fabra

³Department of Linguistics, UC Santa Barbara

Abstract

Nominal number marking is widespread across languages, yet number distinctions are often restricted to specific categories. Typological research suggests that such restrictions follow the Animacy Hierarchy Constraint (AHC), whereby higher-ranked categories in the animacy scale (human > animate > inanimate) are more likely to exhibit number contrasts than lower-ranked ones. However, large-scale cross-linguistic evidence for the AHC is limited, and a mechanistic explanation for its typological prevalence is missing. This study combines typological and experimental studies to investigate the impact of the AHC on number neutralisation in nominal paradigms, and assess whether the AHC mirrors cognitive biases at play during language learning, which could in turn explain cross-linguistic regularities. We analyse data from 509 diverse languages using Bayesian hierarchical models, demonstrating a robust monotonic increase in number neutralisation down the animacy hierarchy. To examine whether this cross-linguistic regularity reflects learning biases, we conduct an artificial language learning experiment manipulating animacy-conditioned number neutralisation patterns. Results show that systems conforming to the AHC are more learnable than AHC-violating systems. Together, these findings suggest that the AHC is supported by learning biases that contribute to its typological prevalence.

Keywords: number morphology; animacy; quantitative typology; learning biases; artificial language learning.

Introduction

Grammatical number distinctions are pervasive across the languages of the world. Most documented languages morphologically mark plurality in their pronominal and/or nominal systems. In personal pronouns, the majority of languages mark at least a SG/PL distinction (98% Saldana et al., 2025, $n = 1250$ languages). Beyond pronouns, most documented languages also productively mark (at least one type of) non-singularity either on nouns themselves or within the noun phrase (71% of documented languages in Grambank Skirgård et al., 2023).

Nevertheless, number distinctions are not exceptionless within languages that encode them: they are often restricted to particular forms and/or contexts. Typological surveys reveal that restrictions on number marking in nouns follow an *animacy hierarchy*: human > animate > inanimate (Corbett, 2000; Croft, 2002; Dixon, 1979; Silverstein, 1976; Smith-Stark, 1974). Higher-ranked categories on the hierarchy tend to exhibit more number distinctions. These accounts assume that when number distinctions are neutralised, (1) neutralisation targets categories lower in the hierarchy, and (2) a given

category is unlikely to show neutralisation unless all categories lower in animacy do so as well. This generalisation is captured by what Corbett, 2000 terms the Animacy Hierarchy Constraint (AHC hereafter): number distinctions in a language must affect a top segment of the hierarchy. This is also claimed to extend to optional number marking, such that optionality in number marking is expected primarily in the lower segments of the animacy hierarchy (Corbett, 2000; Santazilia, 2023).

This animacy hierarchy relevant to the AHC has often been considered to extend beyond common nouns and include pronouns. The resulting *extended* hierarchy comprises two further sub-hierarchies: a person hierarchy (1st person > 2nd person > 3rd person) and a referentiality hierarchy (pronoun > common noun). Recent typological work provides quantitative support for large-scale constraints on number neutralisation in pronominal systems that align with the person hierarchy, consistent with the AHC as applied to the top segment of the extended hierarchy (Saldana et al., 2025). By contrast, a large-scale typological analysis of the AHC that includes nominal systems is still lacking. More crucially, the mechanisms underlying this cross-linguistic regularity in nominal systems remain largely unexplored. In this study, we aim to fill these gaps. We assess the AHC using a typologically diverse sample of 509 languages and test whether the AHC on nominal systems mirrors learners' preferences in an artificial language learning experiment.

The typology of animacy effects on number neutralisation patterns

Materials and methods

Typological data We surveyed number neutralisation patterns in pronominal and nominal systems in the Tymber database (a typological database of nominal number; unpublished, but for a subset of the data see Ivani & Zakharko, 2019). The dataset comprises 509 typologically diverse languages from 101 different families, plus 24 isolates. Families with more than 5% representation in our data include Sino-Tibetan (19%), Indo-European (10%), Afro-Asiatic (6%) and Austronesian (6%). For pronominal systems, we gathered data for 1st, 2nd and 3rd person pronouns. Within 3rd person, for languages that morphologically encode animacy distinctions, we distinguished human, animate and inanimate pronouns. For nominal systems, we gathered information about human, ani-

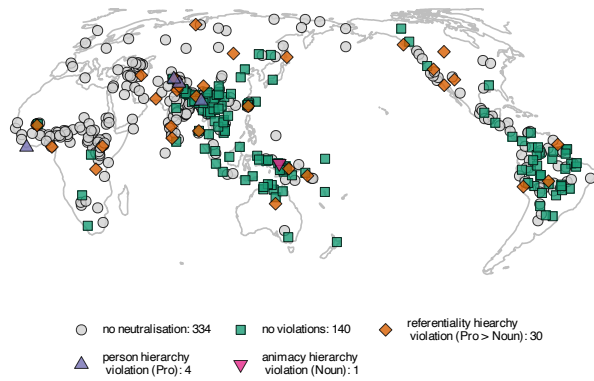


Figure 1: Geographical distribution of languages with number marking systems violating of the animacy hierarchy, the person hierarchy, and the referentiality hierarchy.

mate and inanimate common nouns (excluding proper nouns and kinship terms). Each category was coded for whether it systematically neutralises the SG/PL distinction. Figure 1 shows the geographical distribution of the languages in the database and indicates their conformity with or violation of the AHC (human > animate > inanimate), as well as the other two hierarchies involved in the extended AHC, the person hierarchy (1st person > 2nd person > 3rd person) and the referentiality hierarchy (pronoun > common noun). Number neutralisation is widespread in the sample: 44% of the languages ($n=174$) neutralise number in at least one pronominal or nominal category. Violations of the animacy hierarchy in nominal systems are extremely rare ($n=1$), as are violations of the person hierarchy ($n=4$). By contrast, violations of the referentiality hierarchy are comparatively more frequent ($n=30$; 17% of the 174).

Typological data analysis To test how number neutralisation follows the ordered animacy hierarchy and its extended version, we fit a Bayesian hierarchical binomial regression with a monotonic effect of category. The dependent variable (DV) is whether a given category systematically marks a singular–plural distinction (coded as 1) or neutralises it (coded as 0). Category is entered as an ordered fixed effect reflecting the hypothesised extended animacy hierarchy and is modelled as a monotonic effect. This constrains the effect to be directional across the ordered scale while allowing the sizes of the steps between adjacent categories to vary freely, thus testing whether position in the animacy hierarchy predicts number neutralisation without assuming equal distances between levels. The model includes random intercepts for language and by-family random intercepts and slopes for the monotonic effect of category, allowing both baseline rates of neutralisation and the strength of the hierarchical effect to vary across language families. We set a Student- t prior ($DF = 6, \mu = 0, \sigma = 1.5$) on regression coefficients and

variance components, and a uniform Dirichlet prior over the simplex parameters governing the relative size of the monotonic steps; for the random effects, we set a half-Cauchy prior with scale parameter 1 (McElreath, 2016).

We fit a further Bayesian hierarchical binomial regression model restricted to nominal categories, targeting differences in the probability of number neutralisation across human, non-human animate, and inanimate nouns. The DV in this model is whether the singular–plural distinction is neutralised (coded as 1) or not (coded as 0) for each animacy category. Category is included as a categorical fixed effect with three levels: ANIMATE as reference, HUMAN, and INANIMATE. Random effects include intercepts for language and family, and by-family slopes for the effect of animacy category.

Results

The model with a monotonic effect of category provides strong evidence that number neutralisation increases as we move down the extended animacy hierarchy, and that this increase is not evenly distributed across adjacent steps in the hierarchy. Figure 2 shows the posterior distributions of the global monotonic effect and the simplex components, illustrating the relative contribution of each transition along the hierarchy. The set of ordered categories we test is *1st > 2nd > 3rd human > 3rd animate > 3rd inanimate > human nouns > animate nouns > inanimate nouns*.

The overall monotonic effect of category is clearly negative ($\hat{\beta}_{mo} = -0.885$, 90% CI $[-1.159, -0.644]$, $P(\hat{\beta} < 0) = 1.00$), indicating that moving down the hierarchy is associated with a decrease in probability of keeping a SG/PL distinction. This confirms that the ordered extended animacy hierarchy captures a robust directional constraint on number marking.

The simplex parameters further show that the overall monotonic effect is unevenly distributed across adjacent categories, indicating that specific transitions along the hierarchy account for different shares of the effect. The largest contribution comes from the step between 2nd and 3rd person pronouns ($\hat{\zeta}_{2 \rightarrow 3hum} = 0.298$, 90% CI $[0.190, 0.407]$), with a substantial contribution also from the step between 1st and 2nd person pronouns ($\hat{\zeta}_{1 \rightarrow 2} = 0.152$, 90% CI $[0.025, 0.297]$). By contrast, the steps within the 3rd person domain are comparatively small, both between human and animate and between animate and inanimate pronouns ($\hat{\zeta}_{3hum \rightarrow 3anim} = 0.061$, 90% CI $[0.011, 0.114]$; $\hat{\zeta}_{3anim \rightarrow 3inanim} = 0.058$, 90% CI $[0.010, 0.113]$). Further down the hierarchy, the transitions from 3rd person inanimate pronouns into the nominal domain and subsequent steps within the nominal domain contribute more substantially ($\hat{\zeta}_{3inanim \rightarrow Nhum} = 0.159$, 90% CI $[0.100, 0.223]$; $\hat{\zeta}_{Nhum \rightarrow Nanim} = 0.135$, 90% CI $[0.083, 0.193]$; and $\hat{\zeta}_{Nanim \rightarrow Ninanim} = 0.136$, 90% CI $[0.083, 0.193]$).

Taken together, these results demonstrate that number neutralisation is governed by a strong monotonic effect, while also revealing structured heterogeneity within the hierarchy: different adjacent transitions make different contributions. In

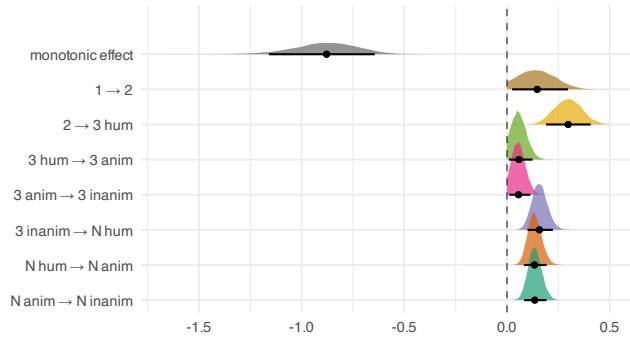


Figure 2: Posterior densities (log-odds scale) of the global monotonic effect and the simplex parameters from the regression model with a monotonic effect for category. The simplex components represent distances between adjacent categories along the hierarchy. Black dots indicate the posterior means; horizontal bars show 90% credible intervals.

particular, we find clear evidence for (i) a robust person hierarchy within pronouns, (ii) a referential split between pronominal and nominal categories, and (iii) a pronounced animacy hierarchy within the nominal domain, which appears stronger than within 3rd person pronouns. To directly test this latter animacy effect on nominals, we fitted a Bayesian regression model to nominal data only. The model estimates suggest a substantially higher probability of neutralisation for inanimate nouns compared to animate nouns ($\beta_{\text{inanimate}} = 3.136$ log-odds scale, 90% CI [1.774, 4.944], $P(\hat{\beta} > 0) = 1.00$) and a markedly lower probability of number neutralisation in human than in animate nouns ($\beta_{\text{human}} = -3.246$, 90% CI [-5.069, -1.908], $P(\hat{\beta} < 0) = 1.00$).

The learnability of number neutralisation patterns in nominal systems

Materials and Methods

We carry out an artificial language learning experiment to investigate how participants acquire systems of nominal number marking with different patterns of number neutralisation (i.e., absence of a number distinction) conditioned on animacy. Specifically, we teach and test participants on nominal systems that mark plurality on some but not all nouns, depending on their animacy category: HUMAN, ANIMATE, and INANIMATE. In conditions where nominal paradigms conform to the Animacy Hierarchy Constraint (AHC), number distinctions are present on nouns whose referents occupy higher positions on the animacy hierarchy (i.e., following HUMAN > ANIMATE > INANIMATE) and absent on nouns with referents lower on the hierarchy. In other conditions, participants are exposed to systems that violate the AHC. Our participants are English-speaking, whose native language systematically marks plurality and does not exhibit animacy-based restrictions on plural marking.

We hypothesise that the cross-linguistically recurrent conformity to the Animacy Hierarchy Constraint reflects cognitive biases in language learning and use, potentially stemming

from saliency asymmetries in number distinctions associated with the representation of animacy categories. Differences in the learnability of AHC-conforming and AHC-violating systems in our experiment provide a direct test of whether animacy-based restrictions on nominal number marking are supported by learners' biases. If conformity to the AHC facilitates the learning of number systems, these results are suggestive of a bias that goes beyond input-specific linguistic knowledge, which may contribute to a transmission bias of AHC-conforming patterns, thereby shaping the cross-linguistic tendencies observed in typological data.

Experimental Conditions We run six critical conditions in which we manipulate the patterns of sg/pl neutralisation in nominal paradigms, conditioned on animacy. Singular is always unmarked (stem only), and plurality can be marked via affixation, or unmarked, resulting in number neutralisation.

Two conditions conform to the AHC, where nouns with referents higher up in the animacy hierarchy are marked for number (A), and those lower down in the hierarchy are not (B): in one condition, only human and non-human animate nouns are marked for number (condition AAB hereafter), and in another, only humans are marked for number (ABB hereafter). Two further conditions violate the AHC and plural is marked in categories which are not contiguous in the hierarchy, i.e., are discontinuous: in one, human and inanimate nouns are marked for number and number is neutralised in animate nouns (ABA hereafter); in the other, number marking is only marked in animate nouns and neutralised in human and inanimate nouns (BAB hereafter). The remaining two conditions are also AHC-violating but are contiguous: BAA and BBA, where referents higher up in the animacy hierarchy not marked for number, and those lower down in the hierarchy are. Participants are randomly assigned to one condition (i.e., between-participants design). Table 1 summarises these six critical conditions. We run a further control condition with systematic plurality marking across animacy categories (i.e., AAA).

condition	AHC	contiguity	NUMBER-marked nouns		
			HUMAN	ANIMATE	INANIMATE
AAB	conforming	contiguous	✓	✓	✗
ABB	conforming	contiguous	✓	✗	✗
BBA	violating	contiguous	✗	✗	✓
BAA	violating	contiguous	✗	✓	✓
ABA	violating	discontinuous	✓	✗	✓
BAB	violating	discontinuous	✗	✓	✗

Table 1: Summary of experimental conditions. The names of the conditions express which of the presented nouns are marked for plural (A) or are not marked (B), for the three categories of HUMAN, (non-human) ANIMATE and INANIMATE nouns. The AHC column indicates whether the nominal number marking system conforms to the Animacy Hierarchy Constraint (i.e., sg/pl distinctions should be found for element higher up in the animacy hierarchy).

Artificial lexicon The miniature lexicon consists of 12 nouns (four for each of the three animacy categories) and

one plural marker. The nouns are semi-nonce nouns (for English speakers). In the human category, the nouns *kukonu*, *daktawo*, *eldare*, *tinepi*, *bebiya* refer to chef, doctor, elder, teenager and toddler; in the non-human animate category, the nouns *pigaro*, *sipete*, *gotani*, *dogowu*, *katema* refer to pig, sheep, goat, dog and cat; and in the inanimate category, the nouns *penrua*, *paktoi*, *balise*, *botanu*, *magelo* refer to pencil, backpack, ball, bottle and mug. The plural marker is a single monosyllabic form, randomly selected for each participant from the array {*nu*, *pu*, *du*, *su*, *tu*}.

Each entity referred to in the lexicon is represented by four different images: two illustrate singular referents (showing one instance of the entity), while the other two depict plural referents (one with a group of 3, the other with a group of 5).

Experimental procedure Participants are asked to learn the mapping between meanings and forms via feedback learning, whereby they receive feedback after each trial and thus the testing trials serve as training themselves. This allows us to track both their overall accuracy and how accuracy changes over time. The experiment comprises a main critical testing phase which we refer to a learning phase, followed by a generalisation phase.

In the learning phase, we train and test participants on the expression of plural referents within nominal paradigms with the aforementioned different patterns of number neutralisation. In each trial, participants see the image for a given referent and are asked to select the corresponding word form out of an array of two, one marked with a plural suffix (i.e., stem+suffix) and one unmarked (i.e., a bare stem). They receive feedback after each selection and a bonus of £0.01 for each correct response. Participants go through 8 blocks of 6 trials each (i.e., total of 48 trials). In each block, they see one plural and one singular referent for each animacy category (human, animate and inanimate).

In the generalisation phase, participants are shown novel referents (not introduced before during the experiment). The new referents correspond to three new entities, one per animacy category, and each of them appears as a plural referent and as a singular referent. Trials are of the same type as in the learning phase. They see each referent in plural and singular once per block (6 trials per block), and they go through two blocks (12 trials total).

Participants Participants (N = 441) are adult English speakers recruited via Prolific (filters: English as a first and fluent language) for a study lasting a median of 8 minutes. They are compensated a base rate of £1.35, and can get up to £0.60 in bonuses according to performance (£0.01 per correct trial). Note that we exclude data from participants whose translations of the number markers in the post-experimental survey do not mention number-related meanings, and those who report to have taken written notes during the experiment. After exclusions (N = 107), we have between 45 and 56 participants for each of the six critical conditions and the control condition

(total N = 334).

Predictions Conformity to the AHC is expected to facilitate the learning and/or the generalisation of nominal number-marking systems. We derive two sets of predictions corresponding to the two experimental phases: learning during the critical testing phase, and generalisation to novel referents.

During the learning phase, participants' trial-level accuracy is predicted to depend on both the structure of number neutralisation and its conformity to the Animacy Hierarchy Constraint. Accuracy is expected to decrease as the number of animacy categories lacking a number distinction increases, such that systems with fewer neutralised categories (AAB/ABA/BAA) yield higher accuracy than systems with more extensive neutralisation (ABB/BAB/BBA). Crucially, controlling for this structural factor, accuracy is predicted to be higher for AHC-conforming systems than for AHC-violating systems.

In the generalisation phase, participants' accuracy in extending learned number-marking patterns to novel referents is predicted to depend on the same structural properties. Specifically, generalisation accuracy is expected to be higher for AHC-conforming systems than for AHC-violating systems, controlling for the number of animacy categories with a number distinction.

Data analysis We run two different Bayesian binomial mixed-effects models to separately test the predictions outlined above regarding learning and generalisation. Our dependent variable (henceforth, DV) across the two models is whether or not participants select the correct button in each trial (coded as 1 if correct, 0 if incorrect). In a first model, we analyse the data from the learning phase of the six critical conditions. As fixed effects, we include categorical predictors for the number of animacy categories with number neutralisation (one in AAB/ABA/BAA, and two in ABB/BAB/BBA) and AHC status (conforming or violating), and a continuous predictor of testing block (centered); we include all interactions amongst these fixed effects. Treatment coding is used, with AHC-conforming systems and two-neutralisation systems as reference levels. As random effects, we include by-participant intercepts and slopes for the effect of block. In a second model, we analyse the generalisation data of the same conditions with the same model structure with the only exception of non-centered block—i.e., we set the first block of generalisation trials as the intercept.

Transparency and Openness The preregistered hypotheses, design and analysis plan are accessible at <https://osf.io/2zydc/>. Note that (a) this study was registered as a replication of a previous study with Spanish speakers with an improved experimental design, and (b) we include two further experimental conditions which were not preregistered (i.e., AHC-violating contiguous: BAA and BBA) and strengthen our results. Experiments were coded using the *jsPsych* JavaScript libraries (De Leeuw, 2015). For data analysis, we use *R*'s *brms* package (Bürkner, 2018) as an interface

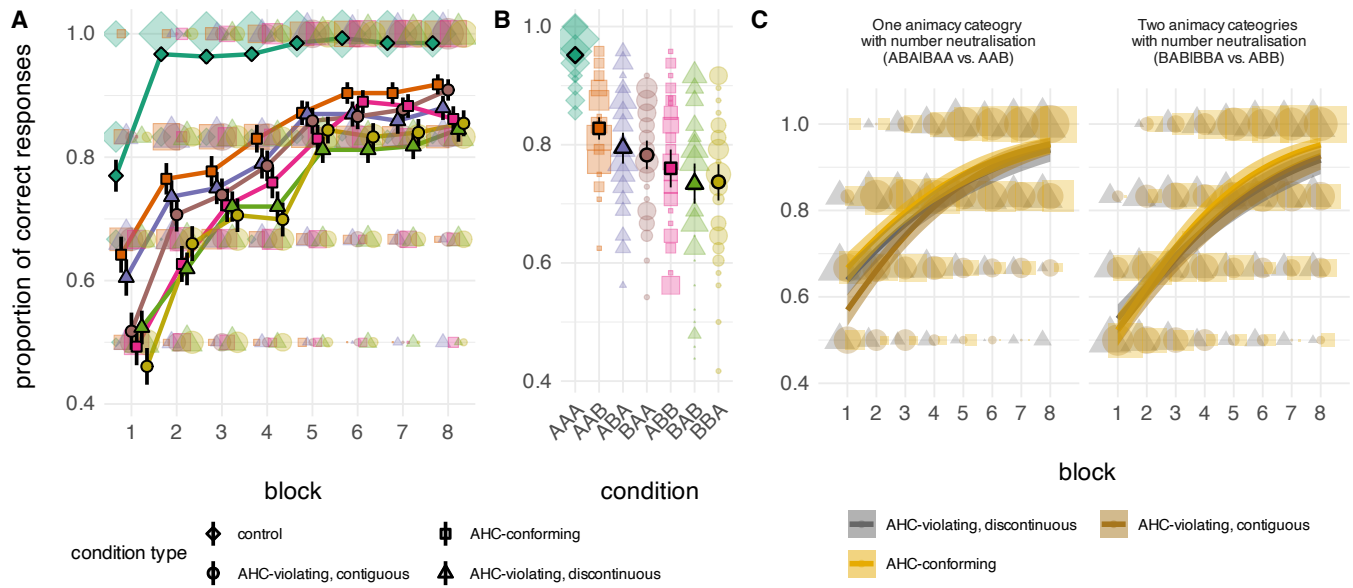


Figure 3: Data from the learning phase. **[A]** Accuracy by testing block for each of the six critical conditions and the control condition. Conditions are colour coded, and shapes represent the the conditions' type (control, AHC-violating discontinuous, AHC-violating contiguous, and AHC-conforming). Black-outlined shapes represent the accuracy means for each block and their standard errors. Shaded background shapes represent participants' individual scores, and larger ones represent more individuals for that given score. **[B]** Overall accuracy by condition. **[C]** Models' predicted accuracy means by testing block and condition type (AHC-violating discontinuous, AHC-violating contiguous, and AHC-conforming), subset by the number of animacy categories with number neutralisation (one in AAB|ABA|BAA, and two in ABB|BAB|BBA). Shaded shapes represent participants' data, and larger shapes represent more individuals; thick lines represent the model's predicted accuracy means conditioned on condition type and block. The shaded area shows the 90% credible intervals (CI).

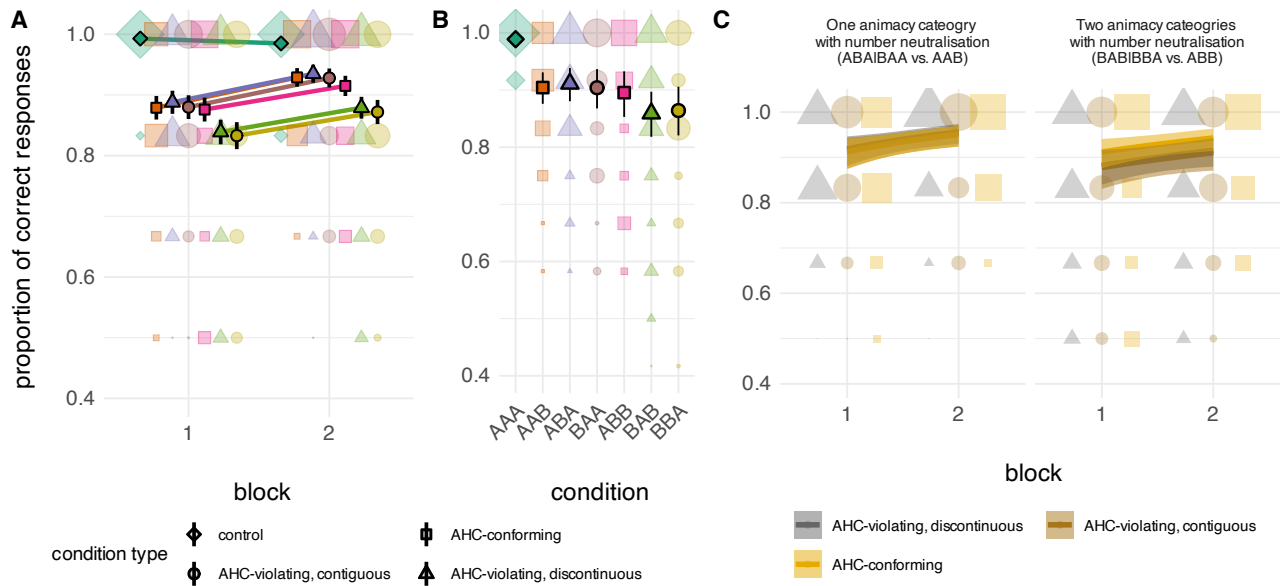


Figure 4: Data from the generalisation phase. **[A]** Accuracy by testing block for each of the six critical conditions and the control condition. Conditions are colour coded, and shapes represent the the conditions' type (control, AHC-conforming, AHC-violating contiguous, and AHC-violating discontinuous). Black-outlined shapes represent the accuracy means for each block and their standard errors. Shaded background shapes represent participants' individual scores, and larger ones represent more individuals for that given score. **[B]** Overall accuracy by condition. **[C]** Model's predicted accuracy means conditioned on condition type and block, subset by the number of animacy categories with number neutralisation (one in AAB|ABA|BAA, and two in ABB|BAB|BBA). The shaded area shows the 90%CI.

to Stan (Carpenter et al., 2017).

Results

Learning Figure 3a-b shows the participants' accuracy across control and critical conditions. A visual inspection suggests that the control condition is indeed the easiest, and that systems become harder to learn the more animacy categories with number neutralisation. It also suggests that AHC-conforming systems are more learnable than AHC-violating ones, and that there is no difference between contiguous and discontinuous AHC-violating systems.

Results from the Bayesian hierarchical binomial model confirm these observations. Figure 3c shows the models' mean estimates by block, condition type, and number of animacy categories with number neutralisation. Accuracy in the reference condition (AHC-conforming systems with two number-neutralised animacy categories) was credibly above chance at the intercept ($\hat{\beta} = 1.54$, 90% CI [1.33, 1.76], $P(\hat{\beta} > 0) = 1.00$). Accuracy increased reliably across blocks ($\hat{\beta} = 0.41$, 90% CI [0.35, 0.48], $P(\hat{\beta} > 0) = 1.00$). In addition, systems with a single number-neutralised animacy category were learned more accurately than systems with two such categories ($\hat{\beta} = 0.32$, 90% CI [0.02, 0.62], $P(\hat{\beta} > 0) = 0.96$). However, learning slopes in these conditions were shallower over time ($\hat{\beta} = -0.08$, 90% CI [-0.18, 0.01], $P(\hat{\beta} < 0) = 0.92$) due to the higher initial accuracy scores.

Turning to the effect of condition type, both AHC-violating contiguous and AHC-violating discontinuous systems show evidence of lower overall accuracy than AHC-conforming systems (contiguous: $\hat{\beta} = -0.23$, 90% CI [-0.53, 0.08], $P(\hat{\beta} < 0) = 0.89$; discontinuous: $\hat{\beta} = -0.24$, 90% CI [-0.53, 0.04], $P(\hat{\beta} < 0) = 0.92$). Learning trajectories also differed across condition types: for AHC-violating discontinuous systems, we find shallower learning curves relative to AHC-conforming systems ($\hat{\beta} = -0.10$, 90% CI [-0.19, -0.01], $P(\hat{\beta} < 0) = 0.97$). However, the corresponding interaction for AHC-violating contiguous systems is modulated by the number of neutralised animacy categories: in systems with a single neutralised animacy category, the reduced learning rate associated with AHC-violating systems was attenuated ($\hat{\beta} = 0.12$, 90% CI [-0.01, 0.25], $P(\hat{\beta} > 0) = 0.93$).

Generalisation Figure 4a-b shows the participants' accuracy in the generalisation phase across control and critical conditions. A visual inspection suggests that participants are most accurate in their generalisations in the control condition. Accuracy scores are high, and we do not observe large differences across conditions with one or two animacy categories with number neutralisation, or across AHC-conforming and AHC-violating patterns. Accuracy scores for all conditions with one neutralisation seem comparably high, and so are accuracy scores for the ABB condition. The only slight difference is observed within the conditions with two neutralisations: AHC-conforming systems seem to be more generalisable.

Results from the Bayesian hierarchical binomial model confirm these observations. Figure 4c shows participants' accuracy scores and the model's mean estimates for the generalisation phase. Accuracy in the reference condition (AHC-conforming systems with two number-neutralised animacy categories) was high, with the intercept credibly above chance ($\hat{\beta} = 2.34$, 90% CI [1.96, 2.76], $P(\hat{\beta} > 0) = 1.00$). Accuracy increased across blocks ($\hat{\beta} = 0.46$, 90% CI [0.03, 0.90], $P(\hat{\beta} > 0) = 0.96$).

Overall, there was no reliable difference in accuracy between AHC-conforming and AHC-violating systems. However, at the reference level, both AHC-violating contiguous and discontinuous systems showed a slight tendency toward lower accuracy than AHC-conforming systems, but evidence in support of these differences is weak (contiguous: $\hat{\beta} = -0.37$, 90% CI [-0.91, 0.16], $P(\hat{\beta} < 0) = 0.87$; discontinuous: $\hat{\beta} = -0.33$, 90% CI [-0.84, 0.18], $P(\hat{\beta} < 0) = 0.85$).

Discussion

Across our typological and experimental studies, we provide convergent support for the Animacy Hierarchy Constraint. We show that plural number marking follows a robust downward monotonic pattern across the extended animacy hierarchy cross-linguistically, and that AHC-conforming systems are easier to learn than AHC-violating ones.

In particular, in our experiment, we found that during learning, participants reliably improved over time across conditions, and systems with fewer number-neutralised animacy categories were learned more accurately overall. Crucially, AHC-violating systems were learned less accurately than their AHC-conforming counterparts. During generalisation, however, these asymmetries were less robust: participants systematically extended AHC-conforming patterns to novel referents more successfully than AHC-violating patterns in systems with two number-neutralised animacy categories; in systems with a single number-neutralised animacy category, which are overall easier to learn, we find no difference between AHC-conforming and violating patterns. We suggest this is due to the high success scores achieved in the last block of learning, which yield ceiling results in the generalisation phase, and differences are only observed for those least learnable (i.e., with deviations from systematic plural marking across more categories).

Our findings contribute to the growing body of work providing behavioural support for a hypothesised causal link between cognition and cross-linguistic regularities (Culbertson et al., 2012; Saldana et al., 2022). The convergence between typological patterns and experimental learning outcomes supports accounts whereby structural properties of languages are shaped by learning biases as they apply through repeated cycles of acquisition and transmission. However, further work is required to assess whether the observed learning biases can be extrapolated to different number systems and different linguistic populations, with more or less familiarisation with number marking and animacy effects on its marking.

Acknowledgements

This research has been supported by the Beatriu de Pinós programme (grant agreement No. 2022 BP 00101, held by Carmen Saldana) co-funded by the *Agència de Gestió d'Ajuts Universitaris i de Recerca* (AGAUR, Generalitat de Catalunya) and the European Union's Horizon 2020 research and innovation MSCA programme under agreement No 801370.

Ethics

The study was approved by the Institutional Committee for Ethical Review of Projects (CIREP) at Universitat Pompeu Fabra (Reference Nr. 345, PI: Carmen Saldana).

References

- Bürkner, P.-C. (2018). Advanced Bayesian multilevel modeling with the R package brms. *The R Journal*, 10(1), 395–411. <https://doi.org/10.32614/RJ-2018-017>
- Carpenter, B., Gelman, A., Hoffman, M. D., Lee, D., Goodrich, B., Betancourt, M., Brubaker, M., Guo, J., Li, P., & Riddell, A. (2017). Stan: A probabilistic programming language. *Journal of Statistical Software*, 76(1).
- Corbett, G. G. (2000). *Number*. Cambridge University Press.
- Croft, W. (2002). *Typology and universals*. Cambridge University Press.
- Culbertson, J., Smolensky, P., & Legendre, G. (2012). Learning biases predict a word order universal. *Cognition*, 122(3), 306–329. <https://doi.org/10.1016/j.cognition.2011.10.017>
- De Leeuw, J. R. (2015). jsPsych: A JavaScript library for creating behavioral experiments in a Web browser. *Behavior research methods*, 47(1), 1–12. <https://doi.org/10.3758/s13428-014-0458-y>
- Dixon, R. M. W. (1979). Ergativity. *Language*, 55, 59–138.
- Ivani, J., & Zakharko, T. (2019). Tymber, Database on nominal number marking constructions. <https://github.com/jivani/tymber/>
- McElreath, R. (2016). *Statistical rethinking: A Bayesian course with examples in R and Stan*. CRC press.
- Saldana, C., Cathcart, C., Ivani, J. K., Bickel, B., & Maldonado, M. (2025). Cross-linguistic regularities in pronominal number neutralisation are not driven by learning biases: Limits of the typological prevalence hypothesis for morphosyntactic categories. *OSF Preprint*. https://doi.org/10.31219/osf.io/qzjnp_v1
- Saldana, C., Herce, B., & Bickel, B. (2022). More or less unnatural: Semantic similarity shapes the learnability and cross-linguistic distribution of unnatural syncretism in morphological paradigms. *Open Mind*. https://doi.org/10.1162/opmi_a_00062
- Santazilia, E. (2023). *Animacy and inflectional morphology across languages* (Vol. 19). Brill.
- Silverstein, M. (1976). Hierarchy of features and ergativity. In R. Dixon (Ed.), *Grammatical categories in Australian languages* (pp. 112–171). Australian Institute of Aboriginal Studies.
- Skirgård, H., Haynie, H. J., Blasi, D. E., Hammarström, H., Collins, J., Latache, J. J., Lesage, J., Weber, T., Witzlack-Makarevich, A., Passmore, S., et al. (2023). Grambank reveals the importance of genealogical constraints on linguistic diversity and highlights the impact of language loss. *Science Advances*, 9(16), eadg6175. <https://doi.org/https://doi.org/10.1126/sciadv.adg6175>
- Smith-Stark, T. C. (1974). The plurality split. *Chicago Linguistic Society*, (10), 657–661.