

English dative preferences as statistical inference over verb-specific distributions

Lean Luo (leluo@ucdavis.edu), Emily Morgan

Department of Linguistics, UC Davis

Introduction. Alternating dative verbs in English vary in their preference for double object (DO) versus prepositional object (PO) forms, with more frequent verbs showing a stronger DO bias overall [1]. The reason for the pattern remains unclear: notably, it is found to be under-predicted by the *abstract constraints* argued to regulate dative alternation (e.g., shorter arguments tend to come first). We study the emergence of this pattern by modeling speakers that additionally exploit *verb-specific distributions* from their experience. Our work is built around two empirical assumptions (see [1-3]): *first*, speakers internalize item-specific distributions to a greater extent for more frequent verbs (frequency-dependent distributional learning); and *second*, a verb’s dative use can also be modulated by its non-dative occurrences in surface ditransitive forms, the majority of which are PO-patterning and found among infrequent dative verbs (e.g., *drove the kids to the school*, where the second argument is not a recipient as in datives).

Method and Results. We adopt an iterated learning framework: in each generation a speaker learns from the previous generation’s output and generates utterances as input for the next generation. Generation: Utterance tokens are generated for each dative verb v taken from the corpus in [1] according to its estimated college-age exposure frequency. A token i is sampled as DO or PO with its DO probability $p(DO)_{v,i}$, where $p(DO)_{v,i} = \text{logit}^{-1}(\alpha + u_v + \eta_{v,i})$. Here α is the baseline DO bias, u_v is the idiosyncratic DO bias for the verb, and $\eta_{v,i}$ represents the overall abstract constraint effect for the token. For each i , its $\eta_{v,i}$ is computed based on a dative utterance of v sampled from [1], using the effects of relevant constraints estimated in the same study. α is initialized at 0 (i.e., a 0.5 DO preference) and each u_v is sampled from $N(0, 0.15^2)$ in generation 1. Learning: In each subsequent generation, α and u_v are updated based on a mixed-effects logistic regression model fitted to the input data, with frequency-dependent distributional learning captured via partial pooling [4]. We define Model 1 as a baseline learner that updates using the α_D and u_v^D estimated from dative input with prior knowledge of $\eta_{v,i}$, which accounts for much of the token-level DO/PO variation and thus limits learning of verb-specific distributions. We then model a learner without such prior knowledge (Model 2), and also the ones that consider a sampled set of non-dative utterances in their updates (Models 3 and 4). See Table 1. Evaluation: We run 30 chains of 20 iterations, obtaining in each chain a set of model-predicted $p(DO)_v$ values and a slope of linear fit relating $p(DO)_v$ to verb frequency. Results show that alignment between model outputs and corpus-observed values from [1] increases across generations for Models 2-4, but not for Model 1 (Figure 1).

Discussion. Our findings argue for the role of verb-specific representations, and suggest that English dative preference patterns plausibly emerge from noisy inference over complex input: our models with better results all track verb distributions, specifically by attributing abstract constraint effects to verb-specific biases, or/and learning from verbs’ non-dative usages that share surface patterns with datives.

Table 1. Learning model summary

Model 1 (η given, dative-only)	Fit: $p(DO)_{v,i} = \text{logit}^{-1}(\alpha_D + u_v^D + \eta_{v,i})$
Model 2 (η unknown, dative-only)	Fit: $p(DO)_{v,i} = \text{logit}^{-1}(\alpha_D + u_v^D)$
Model 3 (η given, dative + non-dative)	Same as Model 1, but with non-dative observations included and weighted 1 versus 2 for dative observations
Model 4 (η unknown, dative + non-dative)	Same as Model 2, but with non-dative observations included and weighted 1 versus 2 for dative observations

Notes: (i) Dative observations are assigned a weight of 2 in Models 1 and 2 for consistency across models; (ii) the 2:1 weighting of dative versus non-dative observations in Models 3 and 4 follows is based on [2].

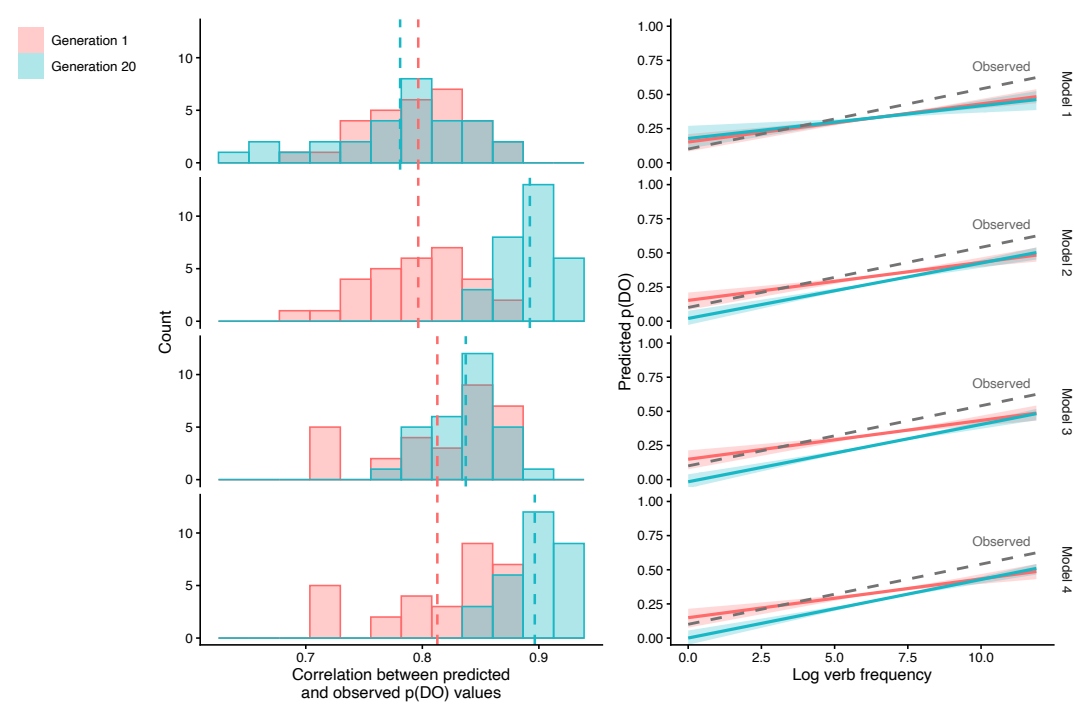


Figure 1. Left: Distributions of correlations between model-predicted and corpus-observed verb-level p(DO) values across chains; dashed vertical lines indicate means. **Right:** Model-predicted p(DO) as a function of log verb frequency; shaded bands indicate 95% confidence intervals (narrow in most cases) across chains, and the gray dashed line shows the linear fit to the corpus data.

References:

[1] Goodwin et al. (2025). Syntactic choice is shaped by fine-grained, item-specific knowledge. *Proc. CogSci.* // [2] Goodwin et al. (2026). *How do speakers learn verb biases from realistically complex data?* HSP2026. // [3] Bock & Loebell (1990). Framing sentences. *Cognition.* // [4] Kapatsinski (2021). Hierarchical inference in sound change. *Front. Psychol.*