

Modeling individual and item differences in atypicality inferences with ACT-R

Guifu Liu¹, Vera Demberg^{1,2} · CogSci 2026

Background

Atypicality inferences: Some events are implied by the context and thereby redundant if uttered:

(1) Mary went to a restaurant. She ate there!

To repair this redundancy, comprehenders **infer** that event to be **atypical** and **accommodate** with explanations.

(2) Mary doesn't usually eat at a restaurant because she only orders beer.

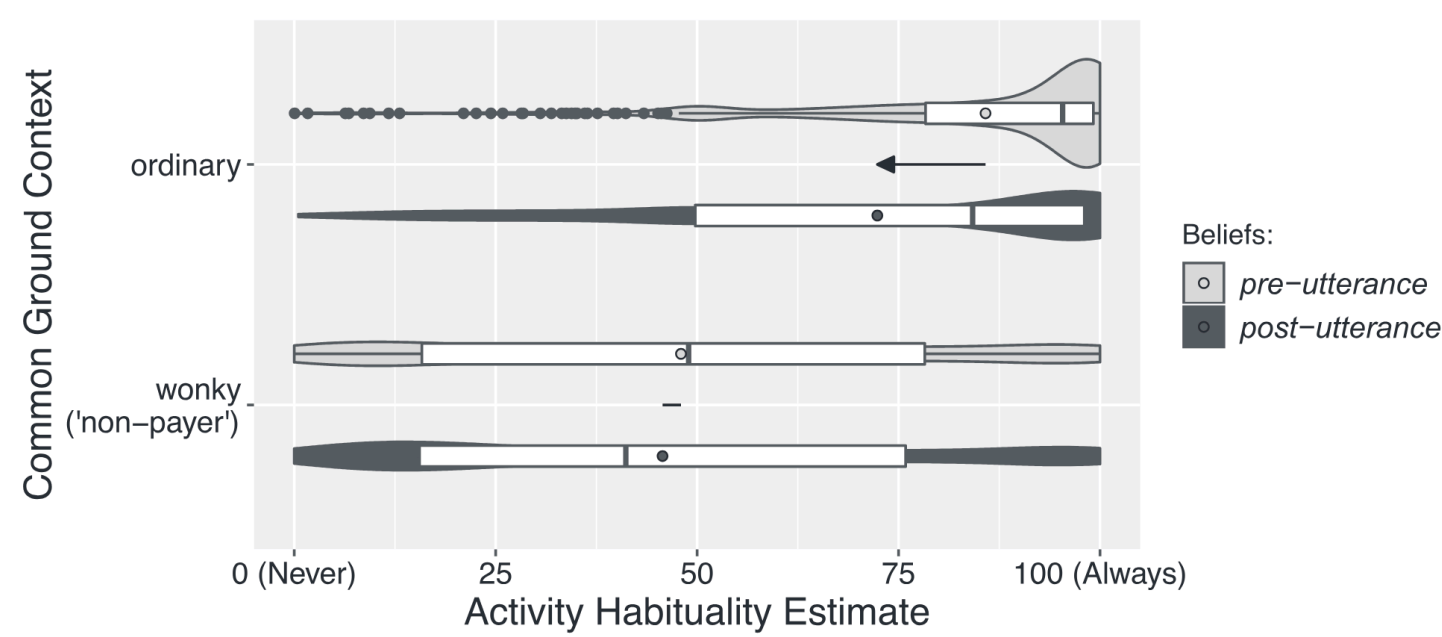


Figure 1. How often do you think Mary usually eats at a restaurant?

Individual differences: Not everyone draws such inference all the same.

- ↑ reasoning abilities measured by *Raven's Matrix* (Ryzhova et al., 2023), ↑ inference rate

Item differences: Inferences drawn from different contexts and events aren't all the same.

- Items have different priors for how habitual an event is (Kravtchenko and Demberg, 2022). *Picking topping when ordering pizza < Eating when going to a restaurant*

Takeaway

How to model *individual* and *item* differences in atypicality inferences?

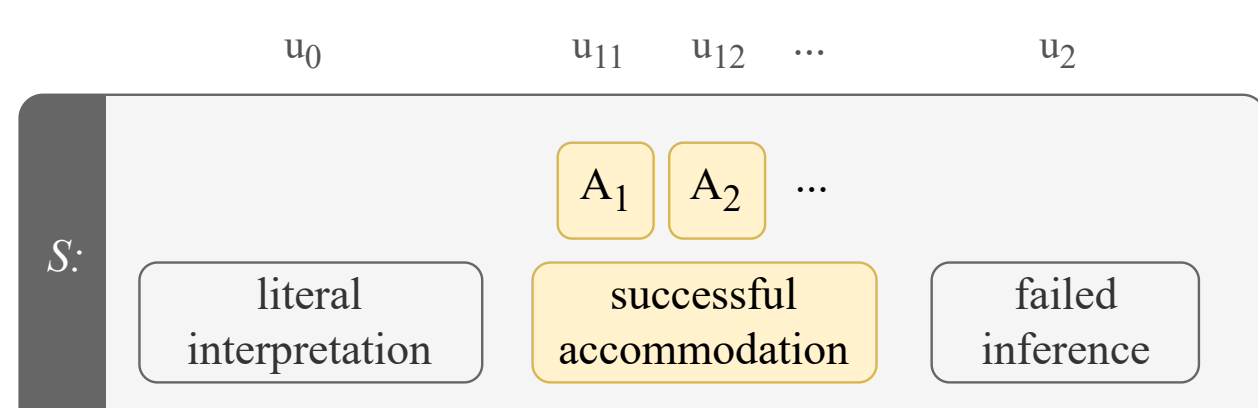
Hybrid model (*cognitive modeling + event plausibility*) can help explain pragmatic inference variance among humans.

ACT-R Cognitive Model

A comprehender model (Fig. 4) that can be **parameterized individually** (shown to correlate with Raven's) and **implements several interpretation strategies**.

- Persistence** (τ) and **Negative Feedback** (F_{neg}): Individual exploration tendencies to persist and disengaged in failed inference.
- Initial utility** u : Item predispositions to interpret literally or accommodate with explanations.

Strategy S : A comprehender can interpret literally, accommodate with valid or invalid explanations, or fail to make inference. Encodes the success likelihood and the strategy cost.



Finding 1: ACT-R Model parameters can simulate human variance

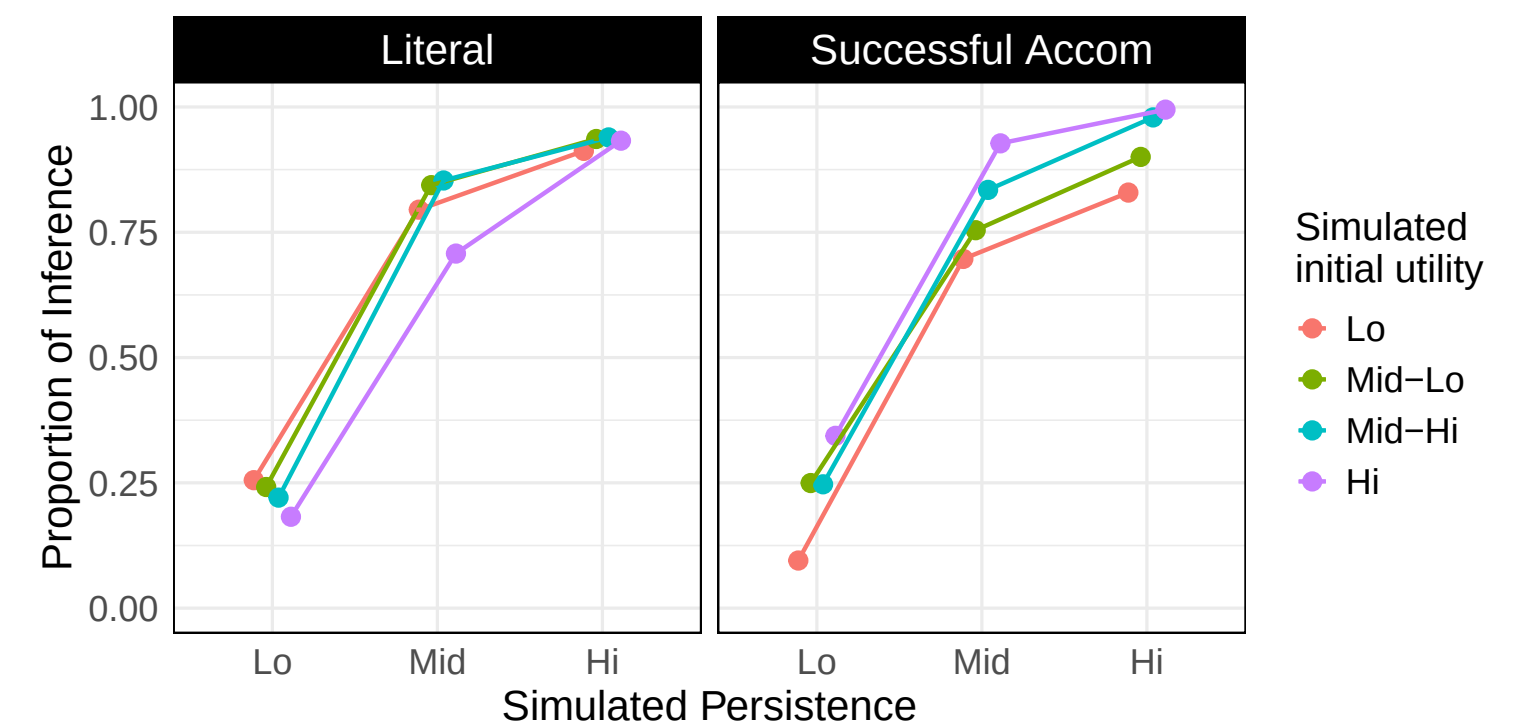


Figure 2. Higher persistence and higher initial utilities for accommodations (Right) ~ Higher inference

Finding 2: Estimated parameters can help predict an individual's performance

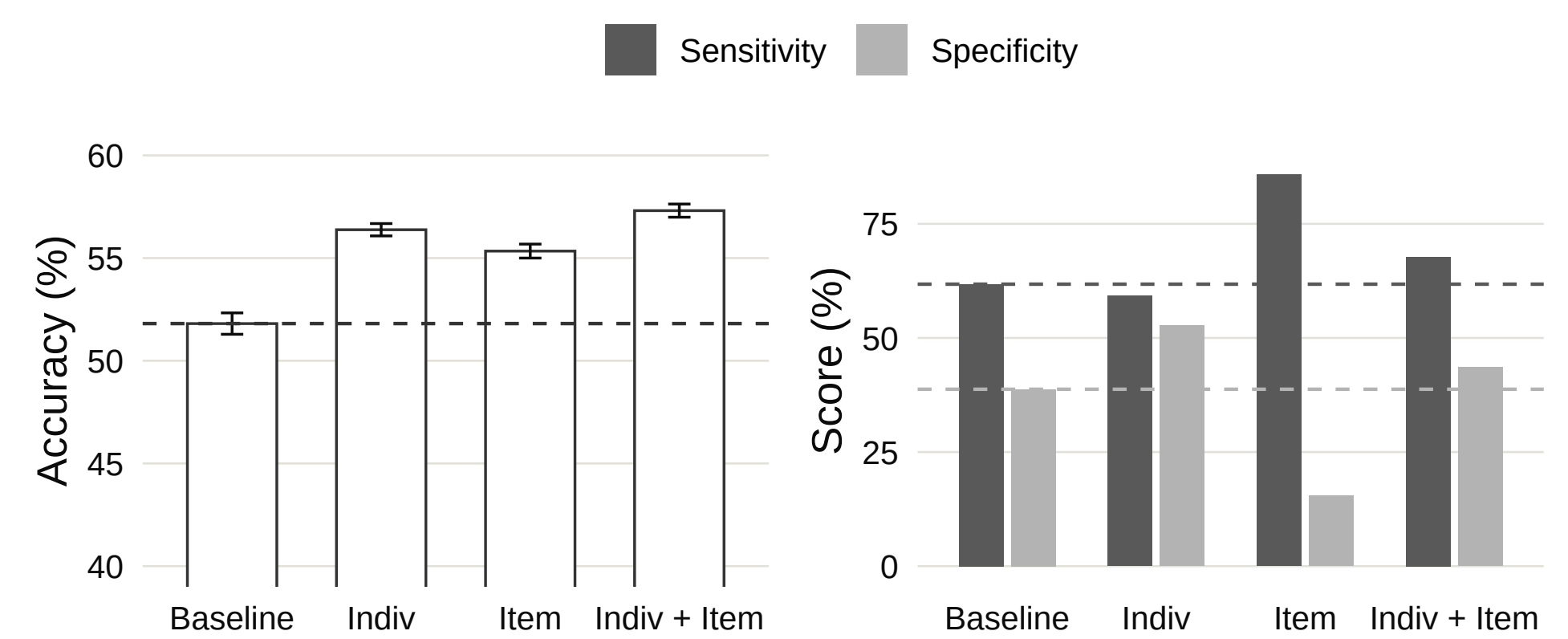


Figure 3. Ind / Item improves accuracy. Item helps detect inferences (+8.58%) but hurts non-inferences (-9.16%)

Comprehenders with low pragmatic profile consistently make no inference, and **item-level variation provides little help**.

Modeling Item and Individual Differences

These model parameters are estimated and **not learned**:

- Persistence is estimated by reasoning test scores (*Raven's Matrix*).
- Initial utility is estimated by habitual ratings from humans and event plausibility from LLM surprisal.

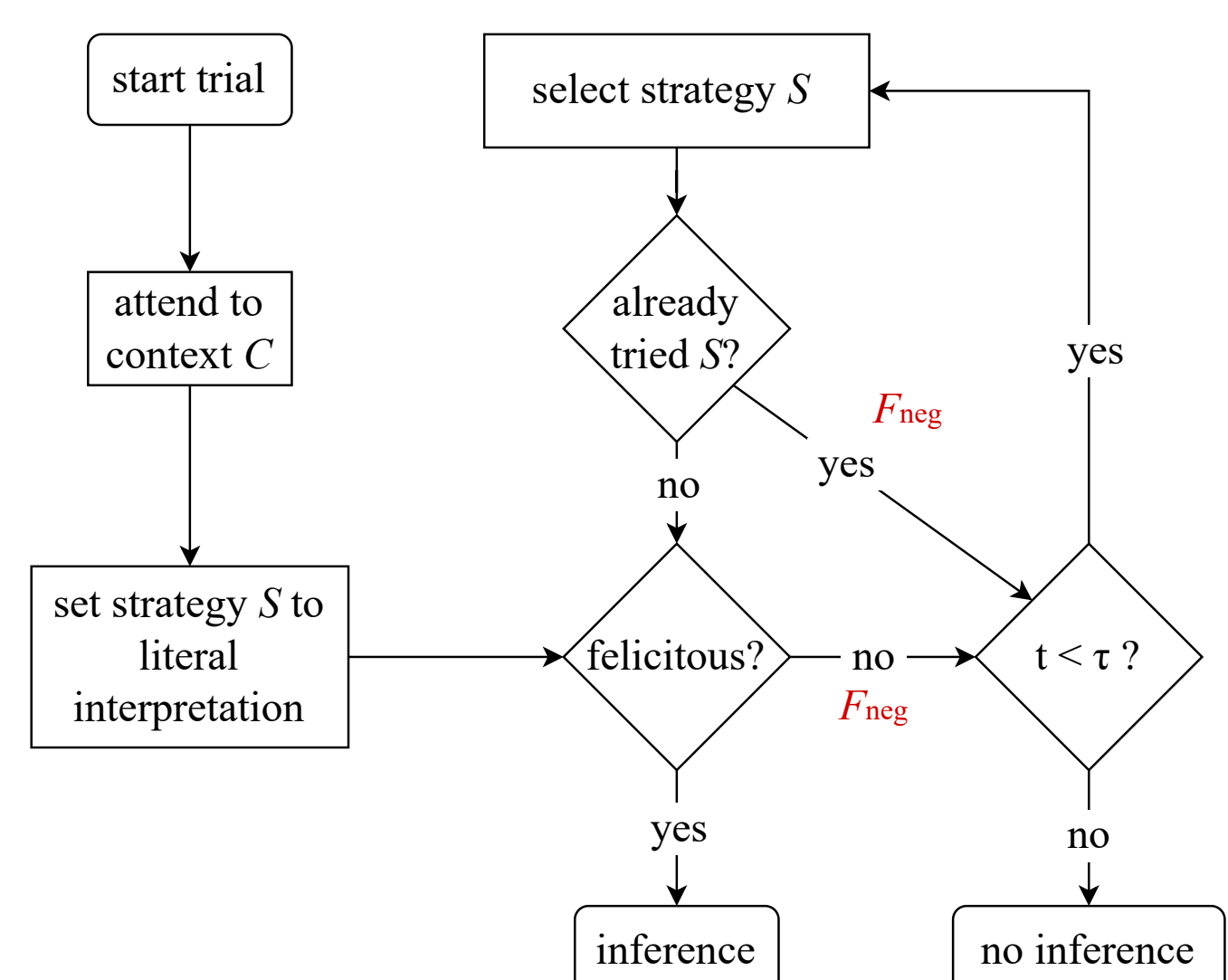


Figure 4. ACT-R comprehender model adapted from Duff et al., 2026