

Semantic attraction is selective: structure-sensitive interpretative mechanisms in verb-argument comprehension

Xingyu (Novak) Shi (xingyu.shi@ling-phil.ox.ac.uk)¹, Colin Phillips^{1, 2}

¹University of Oxford, ²University of Maryland, College Park

Semantic attraction effects in which illicit verb-argument bindings lead to “Illusions of Plausibility” [1-3] may reflect general structure-insensitive interpretative mechanisms, or more restricted mechanisms [4-7]. Here we show via 3 experiments, including speeded judgement (Exp 1), speeded cloze (Exp 2), and a hybrid forced choice paradigm (Exp 3) that semantic attraction is broad in terms of the range of verb-argument combinations that are considered (Exp 1), but also narrow in terms of the settings where it occurs (Exp 2 & 3). Together, the evidence suggests that verb-argument processing is mostly structure-sensitive, becoming vulnerable only in last-resort settings where normal processes fail.

Exp 1 tested specific verb-attractor relations that induce semantic attraction: frequent verb-attractor pairs (e.g., “tip the waitress”), plausible-but-infrequent (e.g., “singer”), and partial feature-matching (e.g., “drinker”), all compared to an implausible baseline (e.g., “chair”). A speeded judgement task crossing *Attractor* (four levels) and *Sentence Plausibility* (plausible or not) was conducted online ($N_{\text{participant}} = 58$, $N_{\text{item}} = 48$) (Fig 1). *Results* revealed comparable attraction effects for implausible sentences with frequent and infrequent-but-plausible attractors (both $p < .01$), in contrast to the partial-plausible ones ($p = .97$) (Fig 2). These suggest that the key factor to semantic attraction is plausibility rather than co-occurrence frequency or feature-matching. **Exp 2** asked whether semantic attraction occurs in cloze production, where verb confrontation is removed (see sample; $N_{\text{participant}} = 96$, $N_{\text{item}} = 42$). *Speeded cloze results* revealed little evidence for attraction in production, evidenced by little difference between the two attractor conditions in (i) attraction rates (Fig 3a) ($p = .53$), (ii) response latency (Fig 3b) ($p = .78$), and (iii) response entropy (Fig 3c) ($p = .96$). This suggests that semantic attraction in judgement (Exp1) depends on direct verb confrontation, with accurate encodings of the sentence representations prior to the verb. **Exp 3** asked whether the error-prone process identified in Exp1 can be avoided in some situations, specifically when comprehenders already have good prior predictions that match the incoming verbs. We implemented a novel hybrid forced choice paradigm online [8-9] ($N_{\text{participant}} = 86$, $N_{\text{item}} = 28$), where comprehenders were asked to choose between a good verb and a bad (illusory) verb at the end of sentence instead of making an acceptability judgement (Table 1). Choice trials were interleaved with production filler trials to encourage prediction. If comprehenders are protected by having good predictions, we would expect little attraction when an expected verb was present alongside the bad (illusory) verb, but no protection when the ‘good’ verb is unexpected. *Results* supported this predicted pattern (Fig 4), with Bayesian GLM revealing little attraction in the expected verb condition ($\beta = -.19$, [-1.10, .74], $P(\beta < 0) = 0.67$), but strong attraction in the unexpected verb condition ($\beta = -.65$, [-1.11, -.19], $P(\beta < 0) = 1$). RTs for correct trials further supported the strong protective effect from an expected verb by revealing a large speed advantage (Fig 5, main effect of *Verb Type* from a crossed Bayesian model: $\beta = .26$ [.19, .33], $P(\beta > 0) = 1$). **Together**, we show that semantic attraction likely results from a late checking process: only by checking can plausible-but-infrequent verb-attractor combinations be identified (Exp 1); attraction rarely shows up in production (Exp 2), and attraction can be blocked when comprehenders already have good predictions (Exp 3). These findings suggest selective attraction, supporting structure-sensitive interpretative mechanisms in verb-argument comprehension.

The pub owner was at the pub. She noticed the **bartender (vs. beer)** that the customer near the

chair
drinker
singer
waitress

Attractor	Plausibility (M)	RoBERTa Probability (M)
Implausible	1.93	<1%
Partial plausible (some features matched)	4.03	<1%
Plausible-infrequent	5.95	<1%
Plausible-frequent	6.30	33.3%
Targets (implausible/plausible)	2.0 / 6.27	/

Fig 1. Design for the speeded judgement task (Exp1). Verb-attractor plausibility ratings (on a 1-7 scale) were obtained from an independent rating task (N = 48).

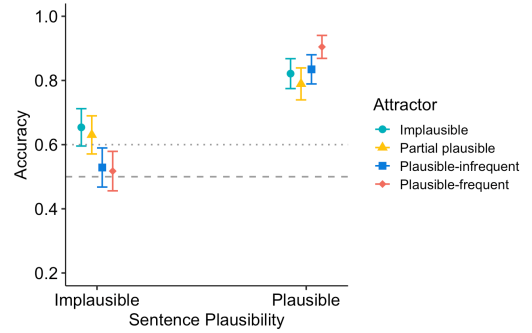


Fig 2. Accuracy results of Exp1.

Sample stimuli for the speeded cloze task (Exp2). The pub owner was at the pub. She mixed the cocktail that the customer [waiting beside *the waitress vs. the chair*] vs. (none) had _____.

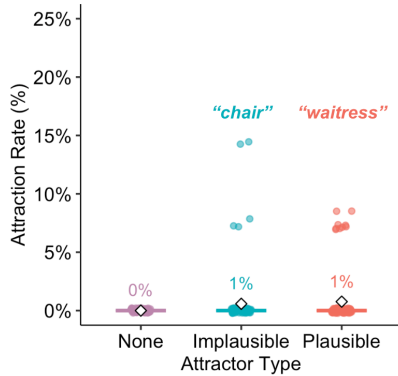


Fig 3a. By-subject (each dot) attraction rate (error responses compatible with attractors) for the speeded cloze task (Exp 2).

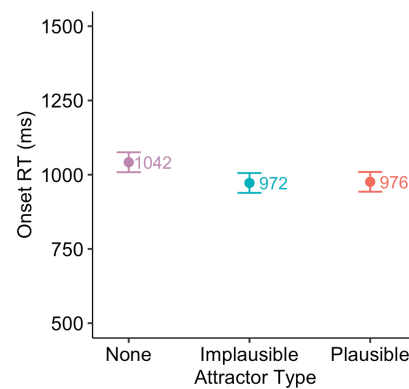


Fig 3b. Avg speech onset latency across attractors (Exp 2).

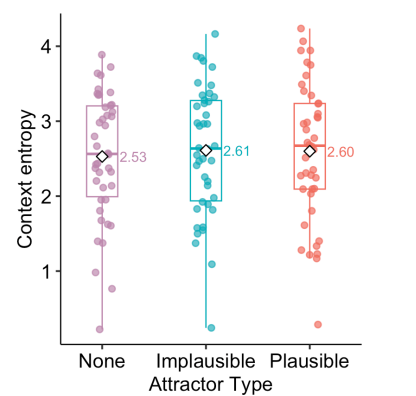


Fig 3c. By-item (each dot) sentential entropy (indexing cloze distributions) across attractors.

Table 1. Critical conditions and item characteristics for the hybrid forced choice paradigm (Exp 3)

Sentence prior to choice	Conditions	Verb choices (cloze; $M_{\text{plausibility ratings}}^*$)	
		Good (Verb Type)	Illusory
The pub owner was at the pub. She mixed the cocktail that the customer waiting beside <u>the waitress</u> vs. <u>the chair</u> had	<u>Expected verb</u> , <u>plausible</u> vs. implausible attractor	<u>ordered</u> (45 %; 6.9)	<u>tipped</u> (0 %; 2.3)
	<u>Unexpected-plausible verb</u> , <u>plausible</u> vs. implausible attractor	<u>tried</u> (0 %; 5.9)	<u>tipped</u>

*Plausibility ratings were obtained from another independent rating task (N = 30). Cloze values were from Exp 2.

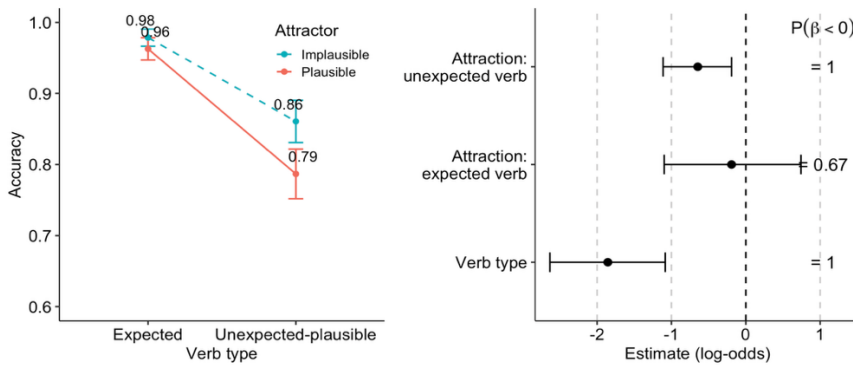


Fig 4. Accuracy results of the choice task (left) and nested Bayesian GLM estimates (right).

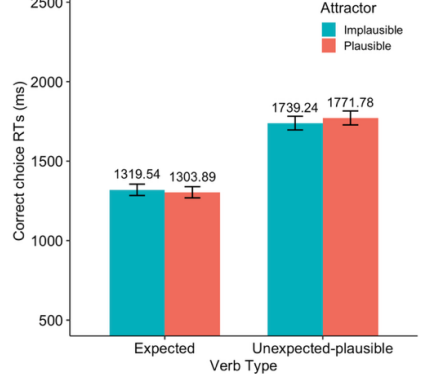


Fig 5. RT results of the choice task.

References. [1] Cunnings & Sturt, 2018. *JML*. [2] Fujita & Cunnings, 2022. *JEP: LMC*. [3] Laurinavichyute & von der Malsburg, 2022. *CogSci*. [4] Lewis & Vasissth, 2005. *CogSci*. [5] Wagers et al. 2009. *JML*. [6] Lago et al. 2021. *Open Mind*. [7] Keshev et al. 2025. *CogSci*. [8] Staub, 2009. *JML*. [9] Lee, 2025. *UMD PhD diss*.