

Separating Intervention from Interference in Wh-Dependencies

Arthur Stepanov (arthur.stepanov@ung.si)

Center for Cognitive Modeling of Language, University of Nova Gorica

Featural Relativized Minimality (fRM) and cue-based retrieval models both predict similarity-sensitive difficulty in syntactic dependencies but assign it to different mechanisms. In fRM, wh-dependencies are degraded when a feature-similar element structurally intervenes between the filler and its gap [1]. Cue-based models instead predict interference when retrieval cues partially match another memory representation [2]. Structural position is therefore diagnostic: fRM requires c-command intervention, whereas retrieval interference can arise from a linearly preceding but non-intervening wh-element.

We tested this contrast in two experiments on English wh-dependencies. The 2×3 design crossed dependency (extraction/no extraction) with wh-position (no wh/linear *how*/structural *how*). In the linear condition, *how* preceded the dependency site but did not structurally intervene. In the structural condition, *how* occupied the structural intervention configuration (Table 1). No-extraction sentences provided construction-specific baselines.

In offline Experiment 1, 40 native speakers rated 48 targets and 48 fillers on a 1–7 scale. Cumulative-link mixed models showed high ratings across no-extraction conditions and in extraction with no wh or linear *how*. Extraction with structural *how* was sharply degraded: the extraction $\times$ structural-*how* interaction was highly significant, whereas the extraction $\times$ linear-*how* interaction was not (Figure 1). Offline acceptability therefore depended on structural intervention, not linear wh-precedence.

In online Experiment 2, 53 native speakers completed a Maze task with the same targets and 72 fillers [3]. Full factorial linear mixed models at lexically aligned regions revealed a dependency $\times$ linear-*how* interaction at the critical verb: linear *how* facilitated no-extraction sentences but produced a small slowdown under extraction. At the critical preposition, both *how* conditions were slower than no wh, but neither *how* cost interacted with dependency; the additional extraction-specific linear-*how* effect was equivalent within $\pm 10\%$ bounds. In extraction trials, structural *how* produced a large selective slowdown one word later (33% vs. no wh; 28% vs. linear *how*), while linear *how* did not differ from no wh. At post-critical 2 both *how* conditions remained slower, and by post-critical 3 they converged (Figure 2).

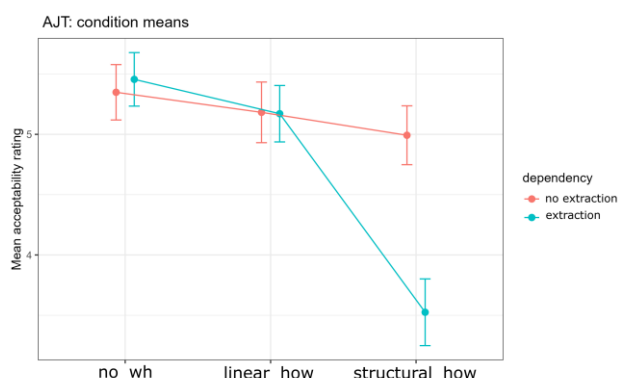
Thus, a merely linear wh-element produced an early dependency-sensitive effect, but the large intervention penalty required the structurally relevant configuration. The results distinguish a modest processing effect of linear wh-similarity from a later, much larger structure-dependent cost, while showing that the shared cost at the critical preposition was construction-general rather than extraction-specific.

Table 1. Target stimuli with linear vs. structural manipulation of wh-intervener

Extraction	pre-critical regions													critical P	post-critical 1	post-critical 2	post-critical 3
no-wh	What	did	the	lawyer,	after	choosing		to	proceed,	decide		to	argue	about	at	the	hearing?
linear-how	What	did	the	lawyer,	after	choosing	how	to	proceed,	decide		to	argue	about	at	the	hearing?
structural-how	What	did	the	lawyer,	after	choosing		to	proceed,	decide	how	to	argue	about	at	the	hearing?

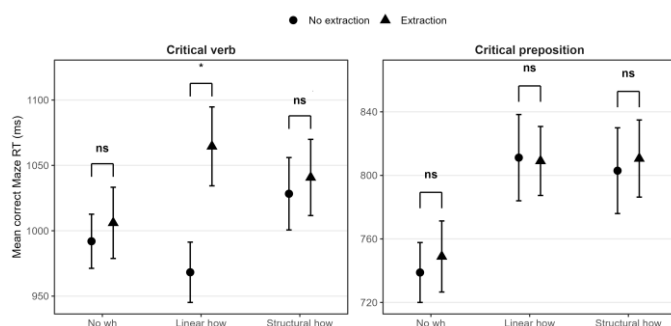
No extraction																		
no-wh	Did	the	lawyer,	after	choosing		to	proceed,	decide		to	argue	about	the	contract	at	the	hearing?
linear-how	Did	the	lawyer,	after	choosing	how	to	proceed,	decide		to	argue	about	the	contract	at	the	hearing?
structural-how	Did	the	lawyer,	after	choosing		to	proceed,	decide	how	to	argue	about	the	contract	at	the	hearing?

Figure 1. Experiment 1 (acceptability judgment task) results



	Est.	SE	z value	Pr(> z)
Dependency: extraction	0.191	0.159	1.197	0.231
wh_position: linear_how	-0.223	0.158	-1.411	0.158
wh_position: structural_how	-0.541	0.154	-3.507	0.0004
extraction*linear_how	-0.171	0.227	-0.752	0.452
extraction*structural_how	-2.118	0.230	-9.206	<0.001

Figure 2. Experiment 2 (Maze task): dependency interaction and extraction RTs



A. Dependency × wh-position at aligned regions

Critical verb

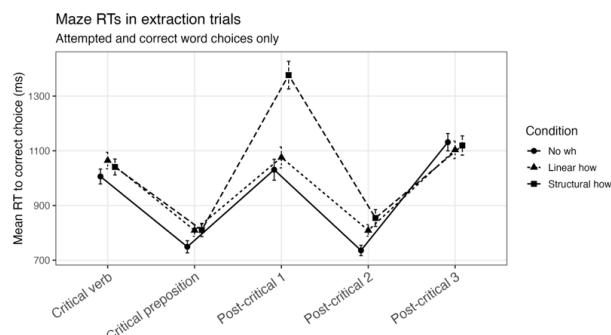
dependency × linear how: $\beta = .092$, $p = .009$

dependency × structural how: $\beta = .031$, $p = .399$

Critical preposition

dependency × linear how: $\beta = .016$, $p = .671$; TOST $p = .019$

dependency × structural how: $\beta = .026$, $p = .506$



B. Extraction-only RT profile

Post-critical 1 (extraction)

linear how – no wh: $\beta = .036$, $p = .337$

structural how – no wh: $\beta = .285$, $p < .001$

structural – linear: $\beta = .249$, $p < .001$

Post-critical 2 (extraction)

linear and structural how > no wh; structural – linear: n.s.

[1] Rizzi, L. (2004). Locality and left periphery. In A. Belletti (Ed.), *Structures and Beyond: The Cartography of Syntactic Structures*, Vol. 3. Oxford University Press.

[2] Lewis, R. L., & Vasisht, S. (2005). An activation-based model of sentence processing as skilled memory retrieval. *Cogn. Sci.* 29, 375–419.

[3] Boyce, V., Futrell, R., & Levy, R. P. (2020). Maze made easy. *J. Mem. Lang.* 111, 104082.