

Listeners do not need an F1/F2 target to perceive vowel quality or regional accentedness

Kevin B. McGowan, Stella Takvoryan
 Department of Linguistics, University of Kentucky

BACKGROUND

Silent centers (SC): listeners can identify vowel quality in a CVC syllable even with 65% of tense vowels and 50% of lax vowels removed (Strange & Jenkins, 2013). However, listeners may still require vowel centers to hear social information. Three complementary ideas in the literature suggest that social information in vowel centers may be *essential*:

- Primacy of F1/F2 at the vowel midpoint**, sometimes taken along with duration, e.g., sociophonetics, sound change, second language acquisition, etc. (Kelley & Tucker, 2020; Labov et al., 1972; Nycz & Hall-Lew, 2013; Thomas, 2014)
- Hybrid silent centers** (Rakerd & Verbrugge, 1987; Verbrugge & Rakerd, 1986): pairing SC syllable edges from different talkers does not undermine vowel perception so argue vowel edges do not carry social information
- Vowel normalization** (Johnson, 2005; Johnson & Sjerps, 2021) assumes that variation is problematic for listeners so models typically operate on vowel centers where contextual variation is least (c.f. Barreda, 2025; Fruehwald, 2024)

METHODOLOGY

- Talkers:** Three non-Southern talkers from the Wildcat corpus (Van Engen et al., 2010) and two Southern talkers (KY)
- Stimuli:** BVT syllables with [i, ɪ, e, ɛ, æ, u, ʊ, o, ʌ, ɔ, a]; **middle 50% for lax vowels & middle 65% for tense vowels** (Strange et al., 1983) excised with a custom Praat script (see Figure 3)
- Procedure:** 2AFC; listeners heard a CVC word and answered either “what did you hear?” with a pair of words or “who did you hear?” and the maps in Figure 1. “What?” trials displayed a map congruent with the talker; “Who?” task trials displayed the word that was being spoken.

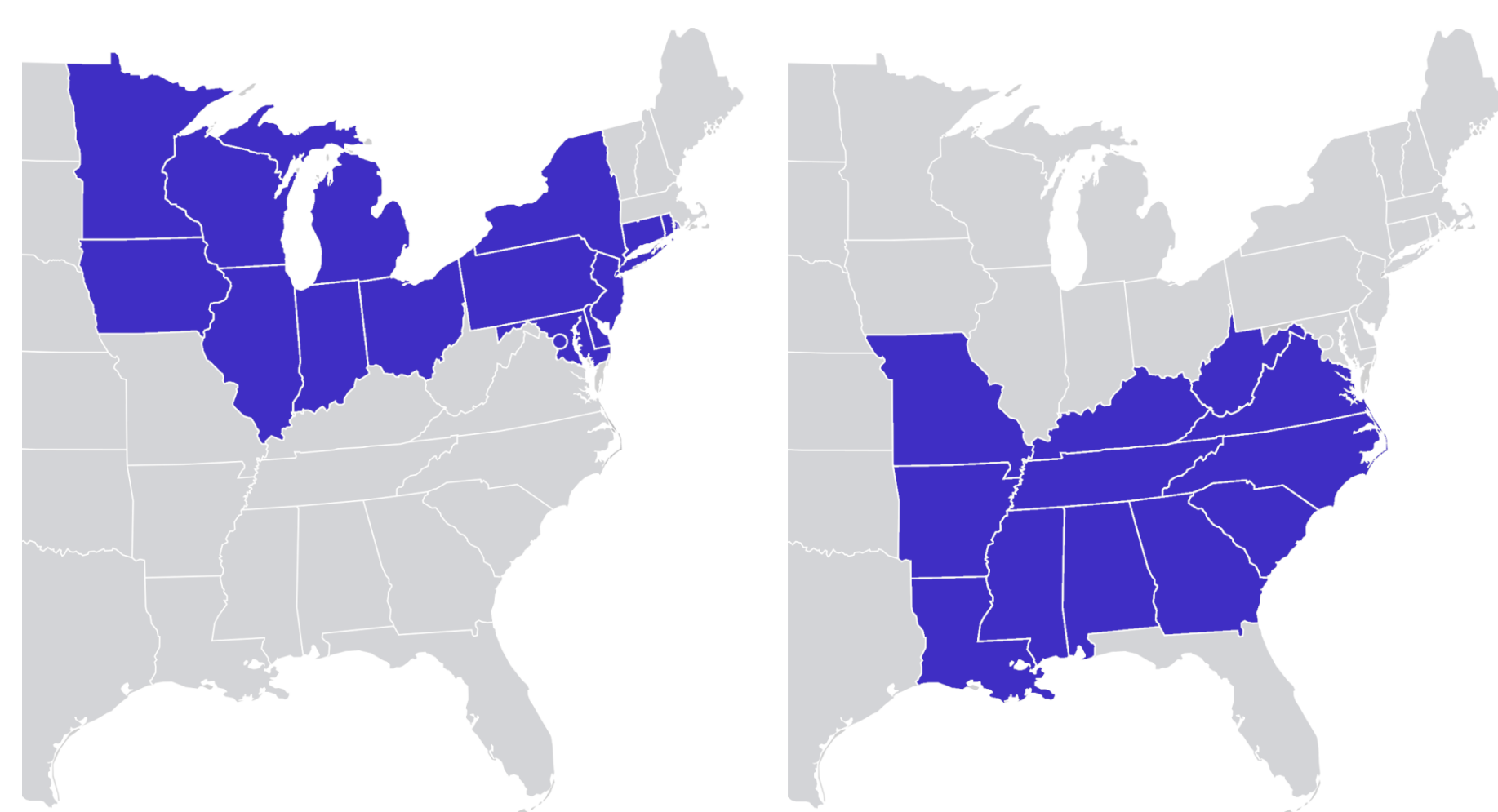


Figure 1: “Non-Southern” and “Southern” stimuli

- Participants:** 60 US participants recruited via Prolific
- Analysis:** BRMS logistic regression in R (Bürkner, 2017), NHST with bayestestR (Arel-Bundock et al., 2024; Makowski et al., 2019)
- Many studies have found that listeners perform poorly when asked to label regional accents (Campbell-Kibler, 2025; Clopper & Pisoni, 2004; Milroy & McClenaghan, 1977). Our simplified maps are intended to represent Clopper & Pisoni’s “dialect clusters”
- While it is clear that listeners do not need vowel centers to perceive vowel quality accurately, it is not yet known whether listeners can perceive, for example, regional accent without the vowel center.

PREDICTIONS

	If V center is necessary	If V center is not necessary
"Who do you hear?" task	Chance or below-chance accuracy	Above-chance accuracy
"What do you hear?" task	High accuracy (SC replication)	High accuracy (SC replication)

Figure 2: Predictions under two assumptions about social information

RESULTS

Summary: While Southern talkers were perceived less accurately overall, there is no significant difference between listeners’ ability to perceive either vowel quality or region in the Full Vowel and Silent Centers vowel manipulation conditions.

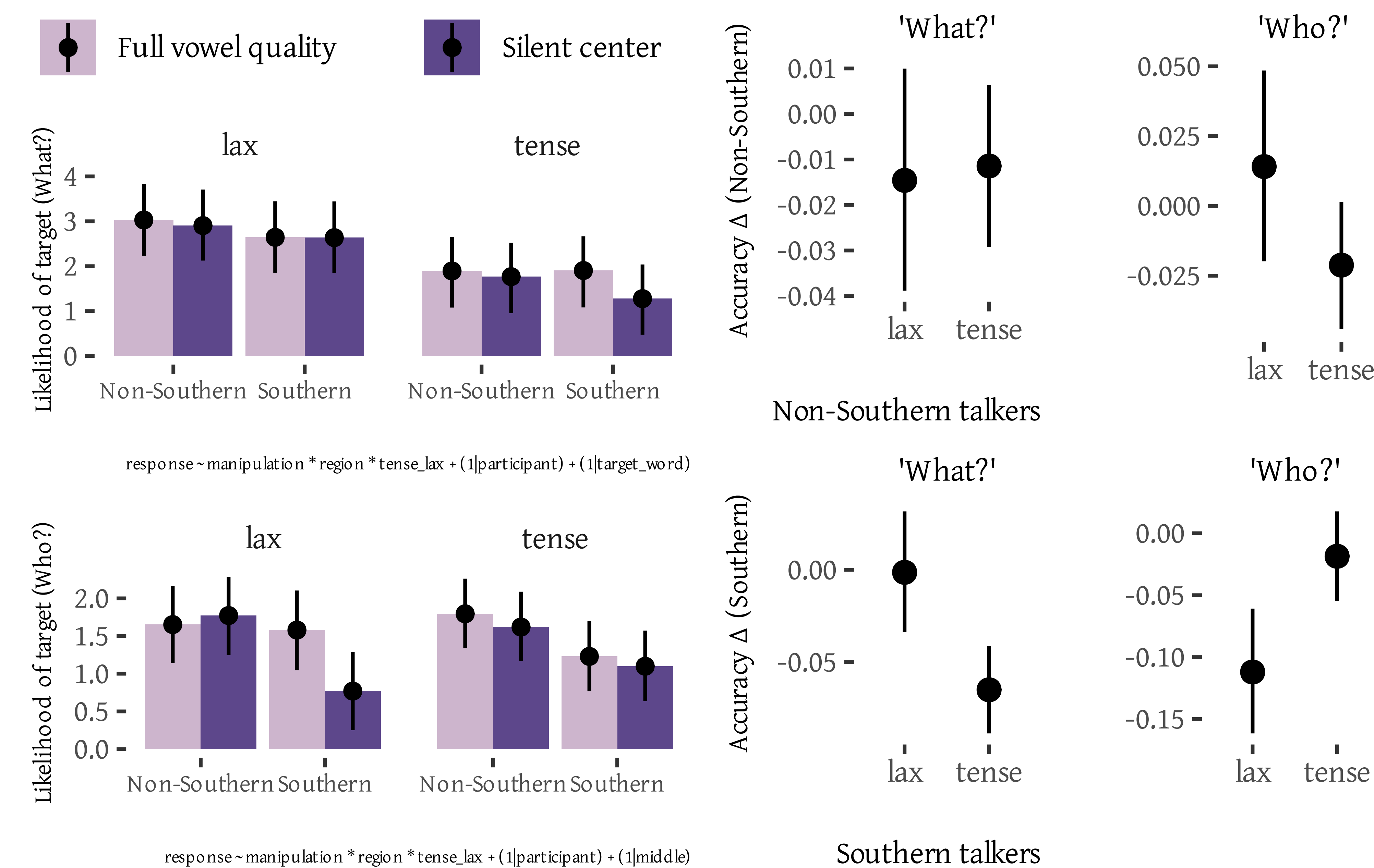


Figure 4: “What?” (top left) and “Who?” (bottom left) model predictions (95% HDI) and Accuracy differences for responses to Non-Southern (top row) and Southern (bottom row) talkers

SILENT CENTERS VISUALIZED

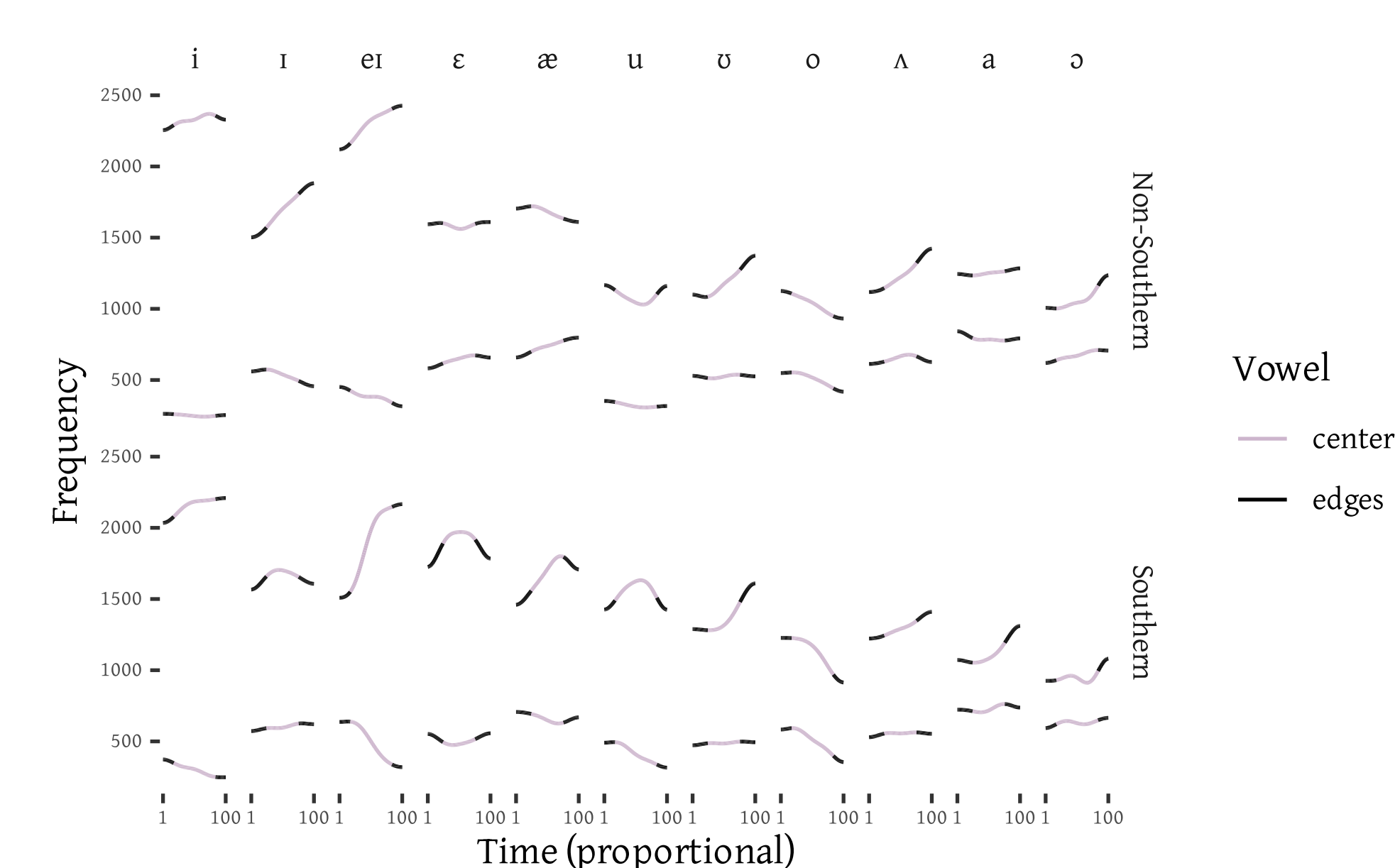


Figure 3: Vowel stimuli unnormed F1/F2 DCTs with excised portions indicated

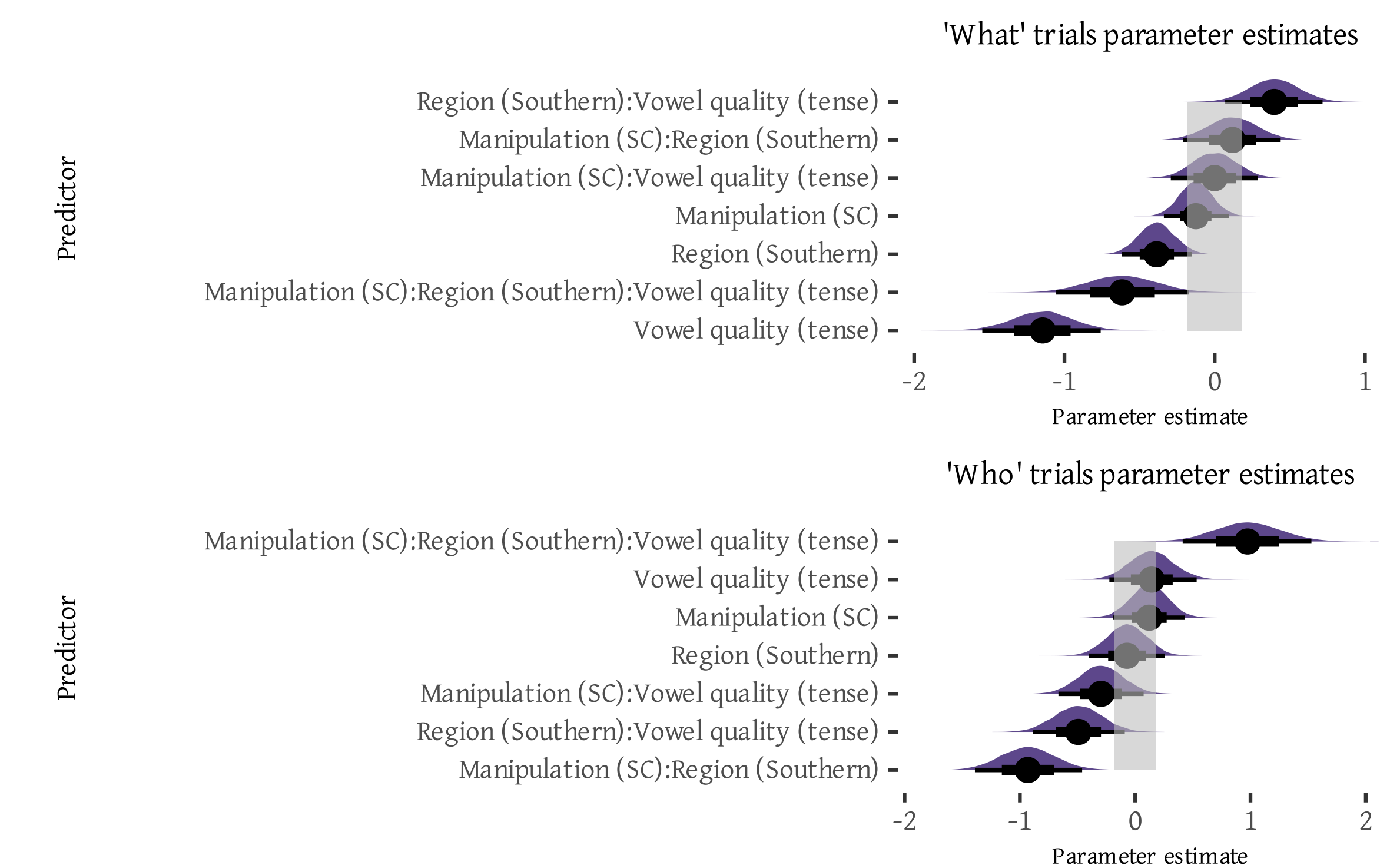


Figure 5: Model coefficient parameter estimates for “what?” and “who?” trials

DISCUSSION

- Listeners do not need the vowel center to perceive vowel quality (replicated):** Listener accuracy on the “what?” trials is a straightforward, successful replication of the silent centers effect (Strange et al., 1983; Strange & Jenkins, 2013)
- Listeners do not need the vowel center to perceive regional accentedness:** Contra hybrid silent-centers work that paired incongruous syllable edges (Rakerd & Verbrugge, 1987; Verbrugge & Rakerd, 1986), listeners to the “who?” trials can, indeed, perceive regional accent from SC vowels
- Tense and lax vowel qualities:** different vowel qualities encode regional variation differently, particularly along this dimension, in this study

CONCLUSIONS

- A single point measure to characterize the vowels of a talker, a community, or a time period are missing information that listeners use. Even multiple points per vowel are only measuring vowel centers if they begin at 20% (or later) and end at 80%
- Tense/Lax: it may be that listeners need a greater percentage of a more dynamic vowel quality to perceive regional variation or this may be due to varying levels of awareness of particular vowel qualities (Babel, 2025)
- These results are inconsistent with models of sound change or normalization that operate on a single F1/F2 measure (with or without duration) and more consistent with, e.g., Beddor (2009) (sound change) or Fruehwald (2025) (normalization)

REFERENCES

ACKNOWLEDGEMENTS

We are grateful to Josef Fruehwald, Jennifer Cramer, Kendal Smith, Austin Coleman, Kyler Laycock, and Shane O’Nan for their invaluable assistance with this project.

