

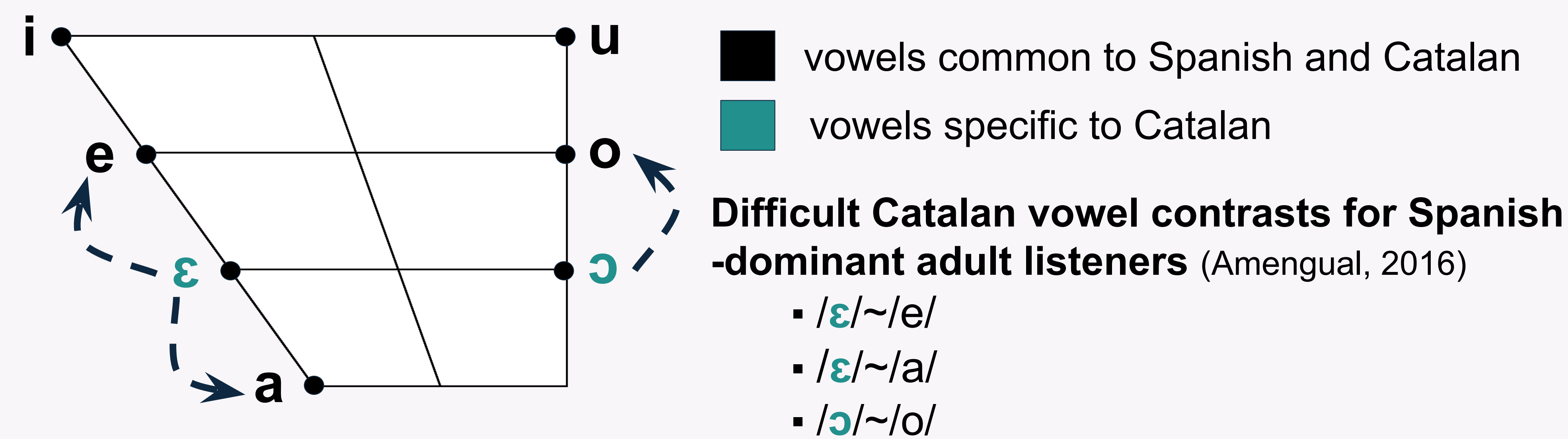
Small Neural Networks as Models of Cross-Linguistic Speech Perception

Annika Shankwitz

University of Maryland – ashankwi@umd.edu

TL;DR

Spanish has five vowels [i e a o u]; Catalan has seven vowels [i e ε a ɔ o u]



L1 Filter – Given a listener’s native language (L1), certain non-native speech sounds are harder to distinguish than others

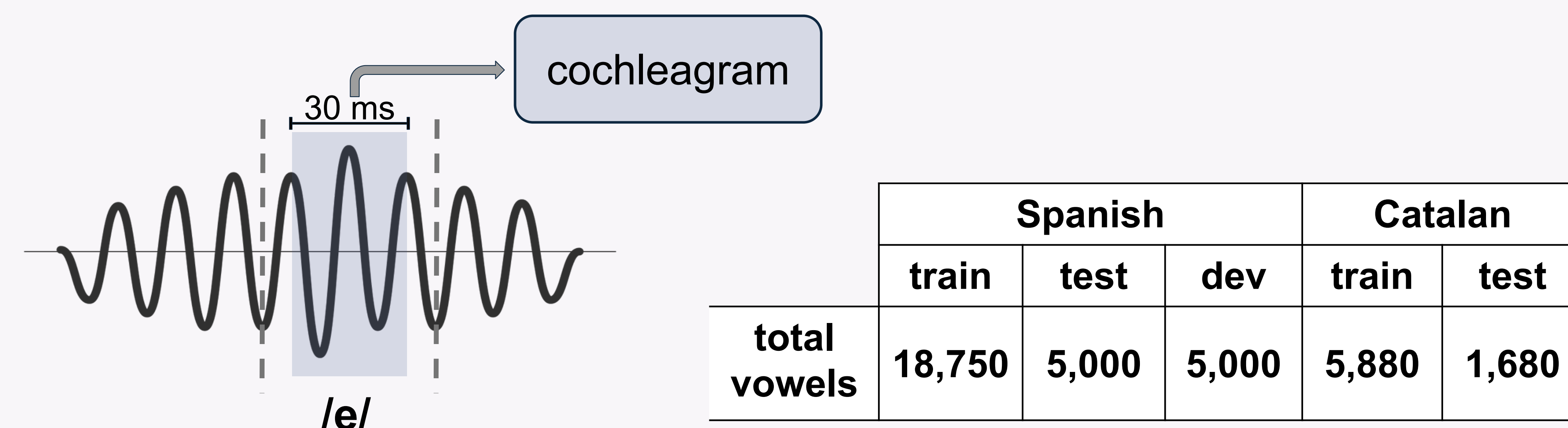
→ Unclear what combination of *model size* and *learning style* best mimics the L1 filter

Small supervised neural networks trained on realistic input can mimic the L1 filter

→ Models trained on Spanish classify Catalan /ε/ as /e/ or /a/, and Catalan /ɔ/ as /o/

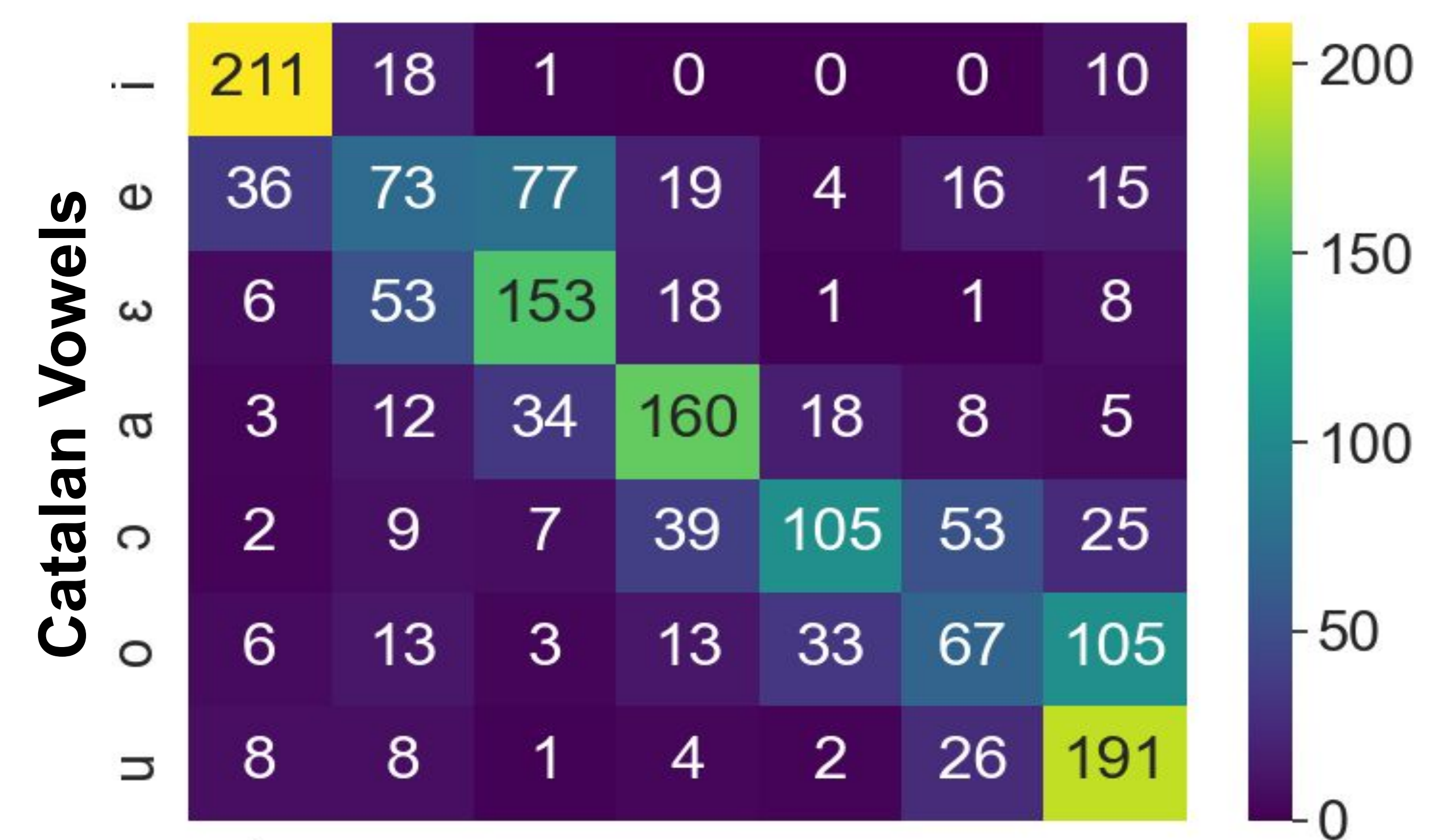
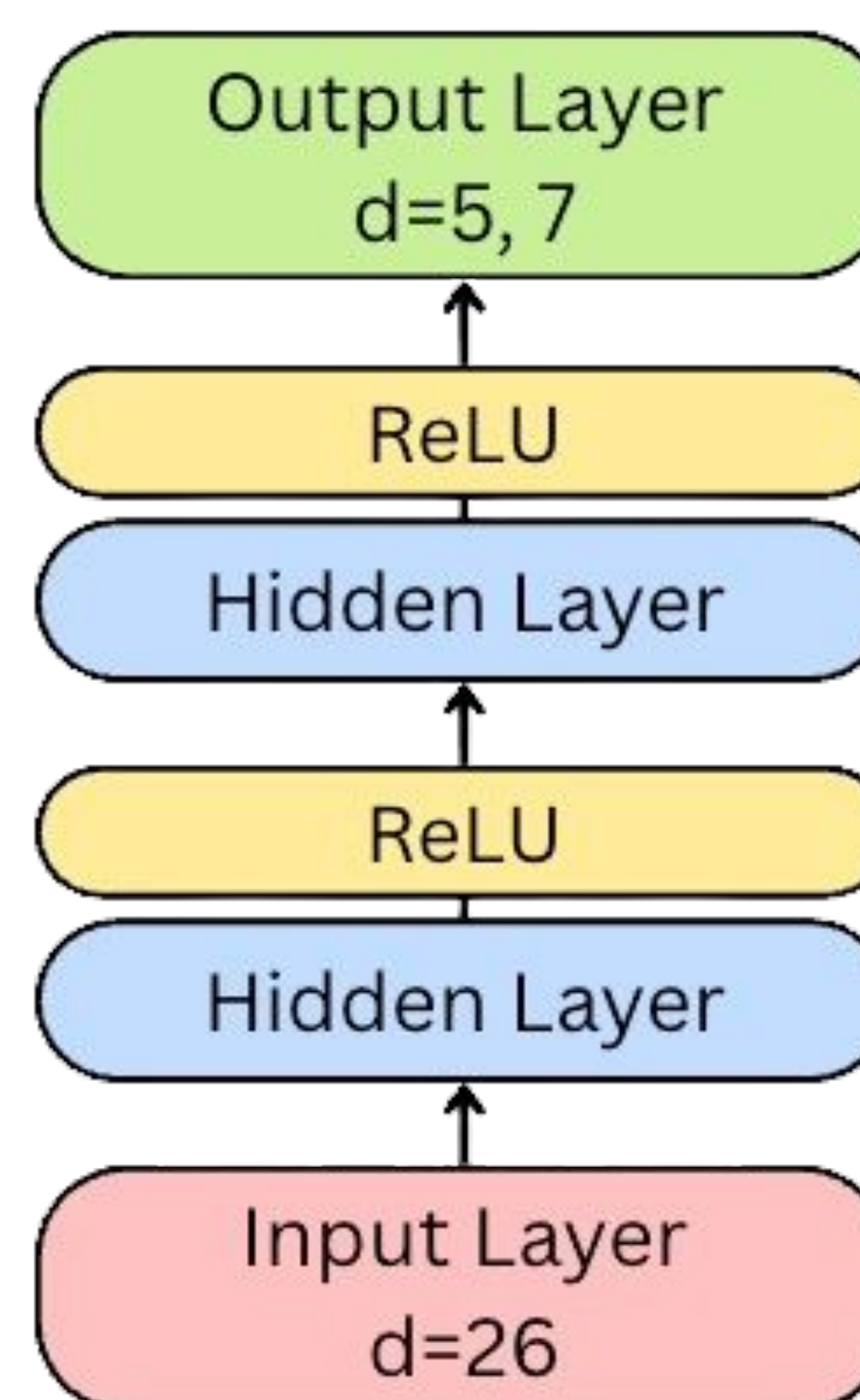
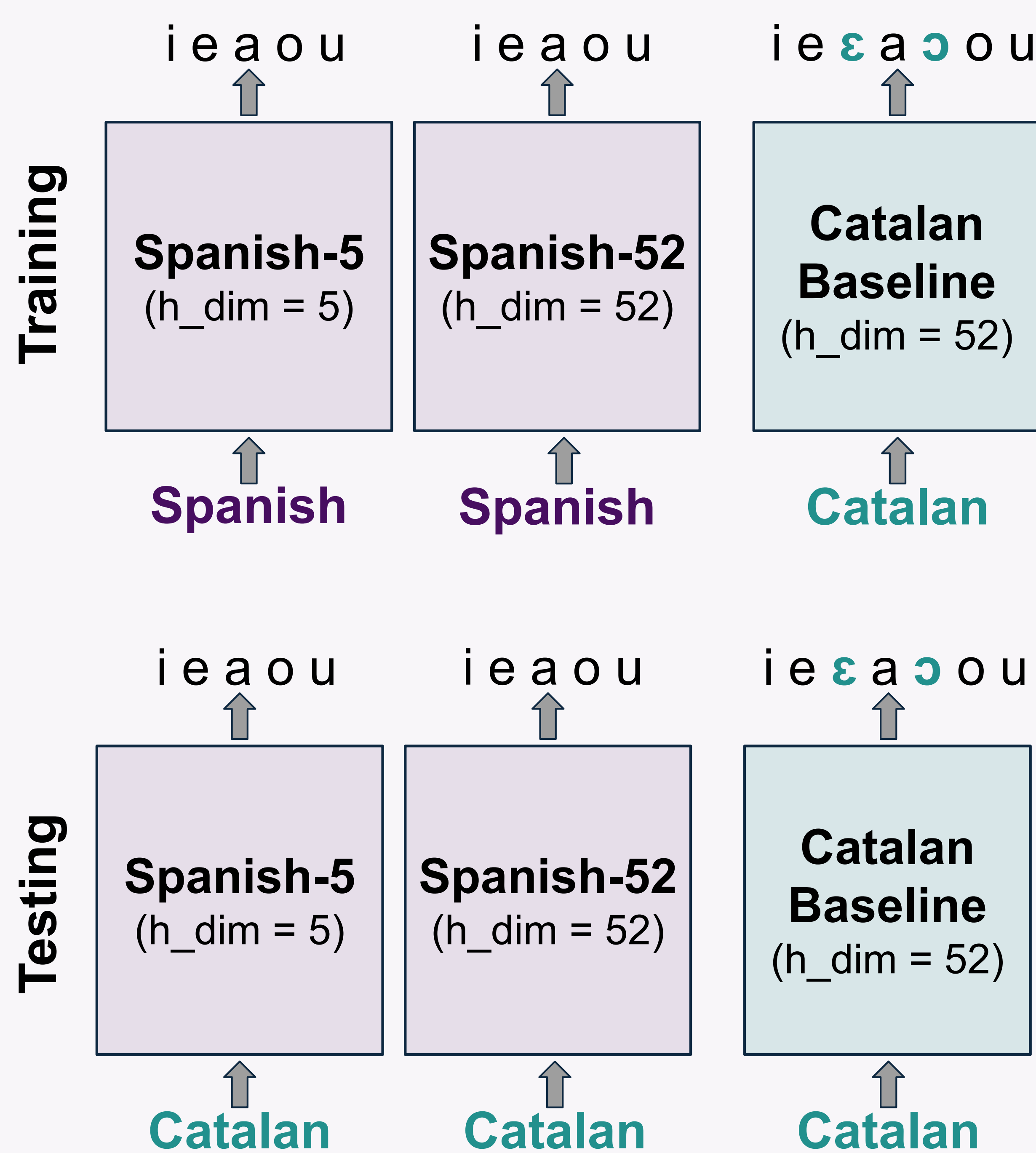
Data and Preprocessing

- Spanish Corpora: Multilingual Librispeech (918 hrs)
- Catalan Corpora: Crowdsourced high-quality Catalan speech data set (9.25 hrs)

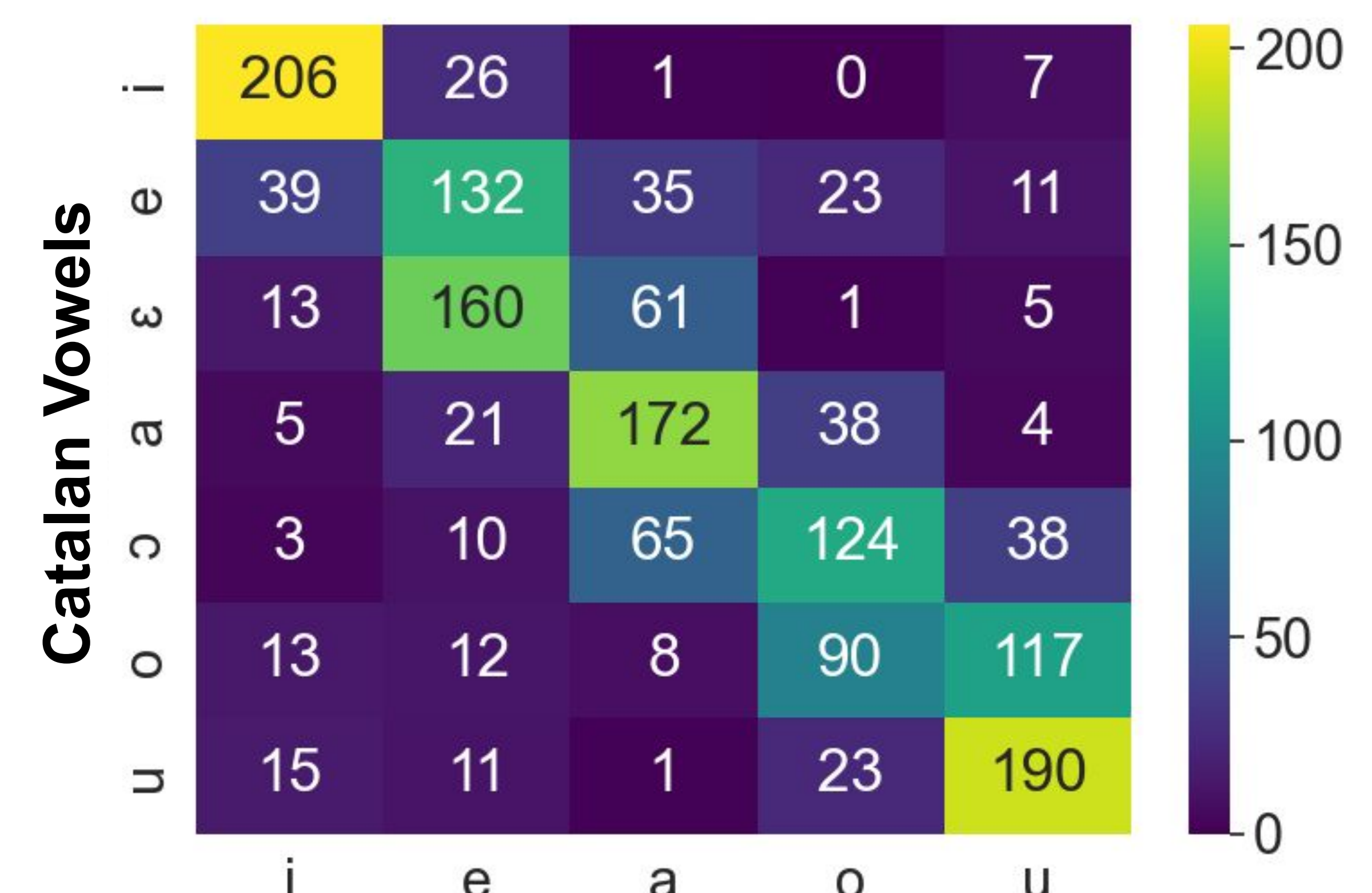


Model Architecture and Training

Four-layer feedforward neural networks (two hidden layers)



Catalan Baseline's Predictions



Spanish-5's Predictions



Spanish-52's Predictions

	/ε/ as /e/	/ε/ as /a/	/ɔ/ as /o/
Spanish-5	67%	25%	52%
Spanish-52	72%	20%	70%
Baseline	22%	8%	22%

/ε/ → /e/ or /a/ ✓

/ɔ/ → /o/ ✓

Using just 30 ms of speech, small supervised neural networks can mimic the L1 filter, and thus show promise as compact models of cross-linguistic speech perception

Spanish models should classify
/ε/ as /e/ or /a/ and /ɔ/ as /o/

References

- Mark Amengual. 2016. The perception and production of language-specific mid-vowel contrasts: Shifting the focus to the bilingual individual in early language input conditions. *International Journal of Bilingualism*, 20(2):133–152.
- Oddur Kjartansson, Alexander Gutkin, Alena Butryna, Isin Demirsahin, and Clara Rivera. 2020. Open-source high quality speech datasets for Basque, Catalan and Galician. In *Proceedings of the 1st Joint Workshop on Spoken Language Technologies for Under-resourced languages and Collaboration and Computing for Under-Resourced Languages*, pages 21–27, Marseille, France. European Language Resources association.
- Vineel Pratap, Qiantong Xu, Anuroop Sriram, Gabriel Synnaeve, and Ronan Collobert. 2020. MLS: A large-scale multilingual dataset for speech research. In *Interspeech 2020*, pages 2757–2761. ISCA.