

A closer look at latent representations of end-to-end TTS models

Martin Lenglet, Olivier Perrotin, Gérard Bailly, *Univ. Grenoble Alpes, CNRS, Grenoble INP, GIPSA-lab*

How are acoustic parameters encoded in a model?

Context

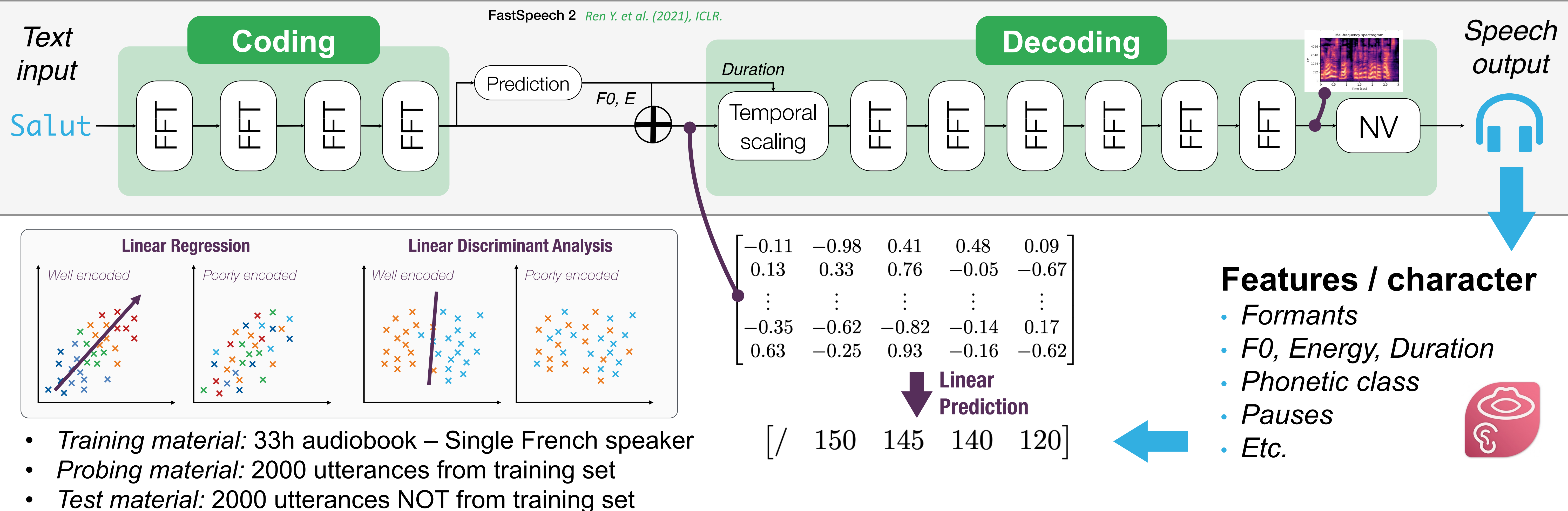
Hypothesis on neural audio modelling:

- Given the high-quality output of audio modelling in general, most speech information should be encoded in the model.

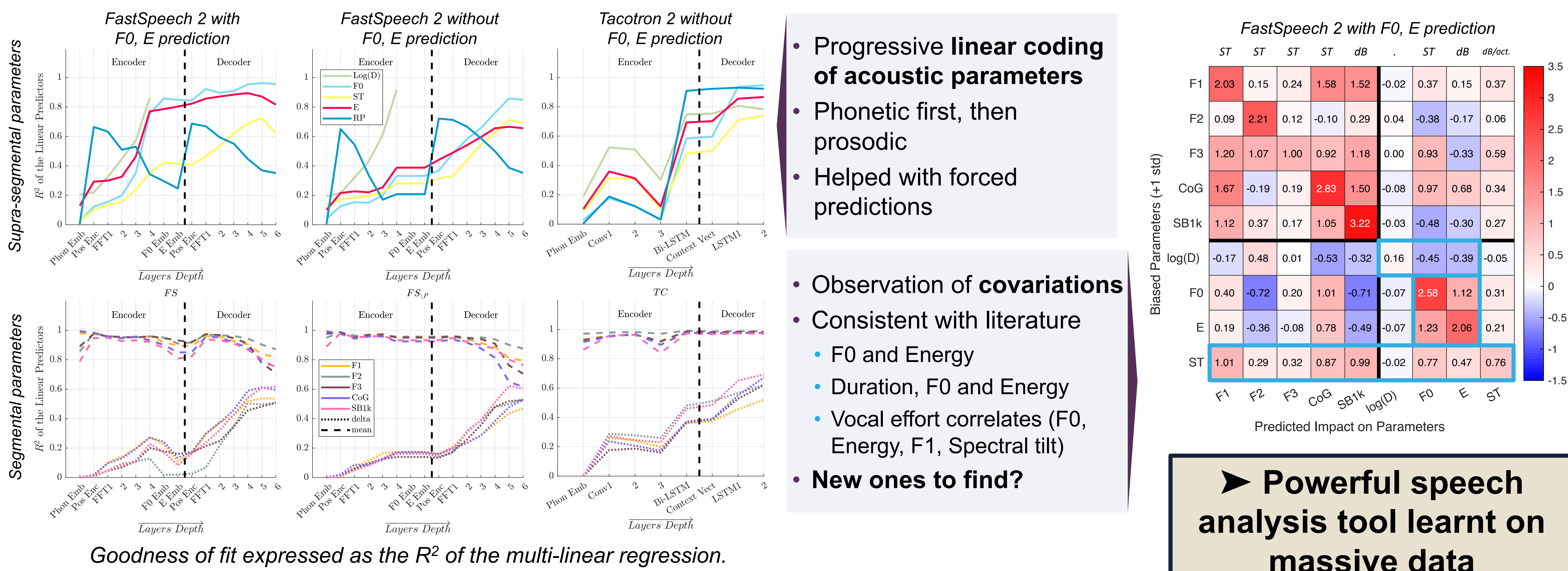
Previous studies on speech model probing:

- Phonetic on TTS: *Perquin et al. (2020), arXiv*
- Acoustic on TTS: *Tits et al. (2021), Informatics 8(4)*
- Language in SSL: *Vaidya et al. (2022), ICML*
- Articulatory in SSL: *Cho et al. (2023), ICASSP*

Method



Results



Can we control them, and how does the system behave?

Results

