

Conversational AI

MARIA TELEKI

Will post
slides after
the talk!

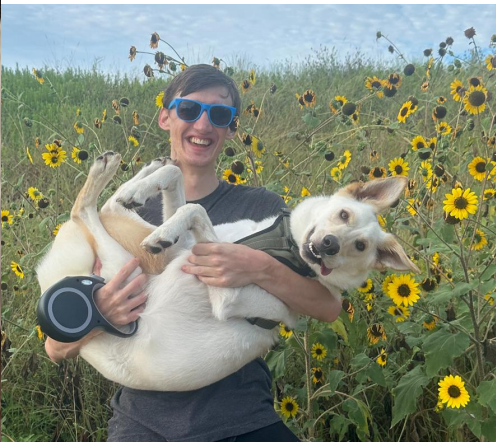


Texas A&M University
Department of Computer Science & Engineering





about me



Agenda

- 1. Speech language models**
- 2. Robustness of pipeline approaches**
- 3. Speaker variation**

First,

**Let's review a recent
survey of Speech
Language Models.**

The way we speak and write are different.

The way we **speak and **write** are different.**

So he calls me up, and he's like, 'I still love you,' and I'm like, I'm just, I mean, this is exhausting, you know – like we are never getting back together. Like, ever.

- TAYLOR SWIFT

The way we speak and write are different.

So he calls me up, and he's like, 'I still love you,' and I'm like, I'm just, I mean, this is exhausting, you know – like we are never getting back together. Like, ever.

- TAYLOR SWIFT

On The Landscape of Spoken Language Models: A Comprehensive Survey

Siddhant Arora^{1*} Kai-Wei Chang^{2*} Chung-Ming Chien^{3*} Yifan Peng^{1*} Haibin Wu^{2*#}
Yossi Adi⁴⁺ Emmanuel Dupoux⁵⁺ Hung-Yi Lee²⁺ Karen Livescu³⁺ Shinji Watanabe¹⁺

¹ Carnegie Mellon University, USA

² National Taiwan University, Taiwan

³ Toyota Technological Institute at Chicago, USA

⁴ Hebrew University of Jerusalem, Israel

⁵ ENS - PSL, EHESS, CNRS, France

ArXiv, highlighted at INTERSPEECH '25 keynote

Big picture context is we live in a **MULTIMODAL** world

LLMs = <i>Large Language Models</i>	SLMs = <i>Speech Language Models</i>	VLMs = <i>Vision Language Models</i>	RLMs = <i>LLMs for Recommendation (I made this name up)</i>	...
$p(\text{text} \text{text})$	$p(\text{text} \text{speech},\text{text})$	$p(\text{text} \text{image},\text{text})$	$p(\text{text} \text{collaborative-signal},\text{text})$	

And right now people are building models
for all these different types of modalities

SLMs can do tasks that LLMs can't do

Task	Examples of Instructions
Speech recognition	Recognize the speech and give me the transcription. (Tang et al., 2024) Repeat after me in English. (Grattafiori et al., 2024)
Speech translation	Translate the following sentence into English. (Grattafiori et al., 2024) Recognize the speech, and translate it into English (Chu et al., 2023)
Speaker recognition	How many speakers did you hear in this audio? Who are they? (Tang et al., 2024)
Emotion recognition	Describe the emotion of the speaker. (Tang et al., 2024) Can you identify the emotion? Categorise into: sad, angry, neutral, happy (Das et al., 2024)
Question answering	What happened to this person? (Wang et al., 2023b) Generate a factual answer to preceding question (Das et al., 2024) What medicine is mentioned? Briefly introduce that medicine. (Peng et al., 2024b)

Table 2: Examples of instructions for speech-related tasks used in SLM instruction tuning.

SLMs are trained in a super cool new way!!!!!!

Table 1: Typology of text and spoken LMs. We use a loose notation here, where *speech* and *text* are to be interpreted in context; for example, $p(\text{text}|\text{text})$ in post-trained text LMs corresponds to modeling some desired output text given an input text instruction or prompt. “Post-training” refers to any form of instruction-tuning and/or preference-based optimization of the SLM. Please see the sections below for details and references for the example models.

Type of LM	Training Strategy	Model distribution	Examples
pure text LM	pre-training	$p(\text{text})$	GPT, Llama
pure text LM	post-training	$p(\text{text} \text{text})$	ChatGPT, Llama-Instruct
pure speech LM	pre-training	$p(\text{speech})$	GSLM, AudioLM, TWIST
pure speech LM	post-training	$p(\text{speech} \text{speech})$	Align-SLM
speech+text LM	pre-training	$p(\text{text}, \text{speech})$	SpiRit-LM, Moshi (pre-trained)
speech+text LM	post-training	$p(\text{text}, \text{speech} \text{text}, \text{speech})$	Moshi (post-trained), Mini-Omni
speech-aware text LM	post-training	$p(\text{text} \text{speech}, \text{text})$	SALMONN, Qwen-Audio-Chat

“tokenizing speech” → speech becomes like text for modeling

WE ARE
HERE

Pipeline Approach vs. End-to-End Approach



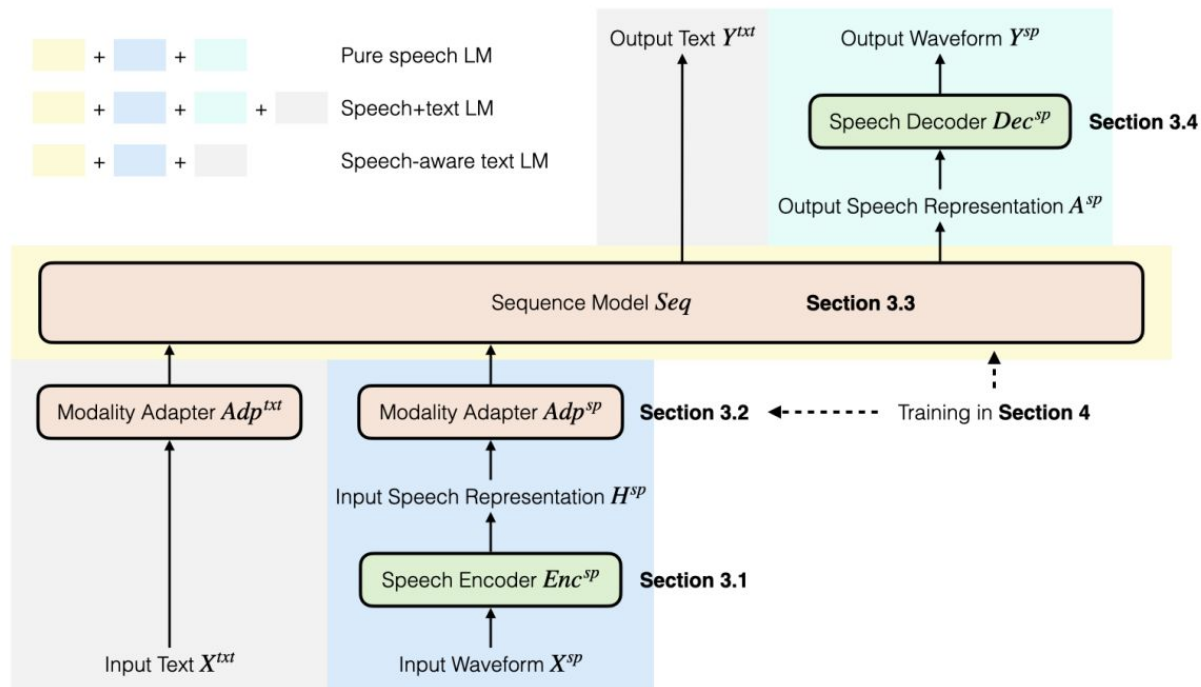
✓ preferred by industry for *controllability* → think custom vocabulary & ease of debugging

✓ currently in research stage but very promising → major barrier imo is downsampling problem (will explain)

The big meta-level question to pay attention to...

- What's the input?
- What's the output?
- How do you glue the parts together?
- Based on the input and output, **what is the SIGNAL being learned (by the model)?**

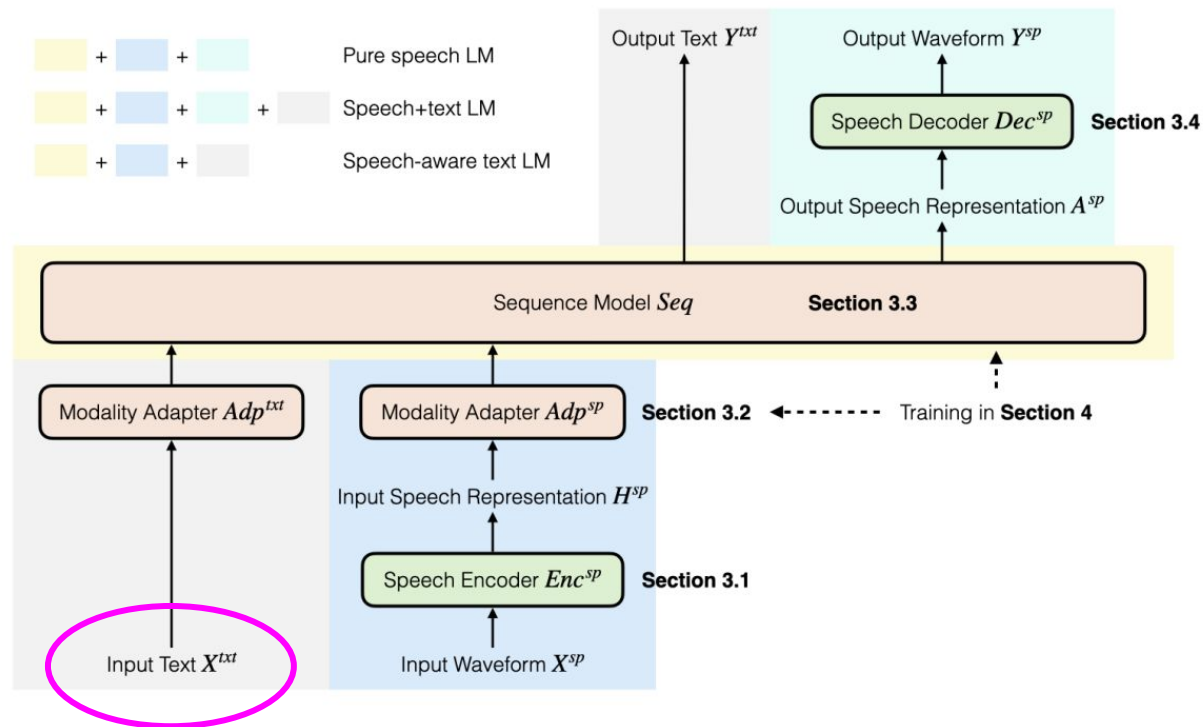
Let's talk SLM architecture



The paper is basically walking through this pic like it's a map – so that's what we're going to do! :)

Figure 1: Overview of SLM architecture. See Sections 3 and 4 for more detailed descriptions of the components and training methods, respectively.

Let's talk SLM architecture

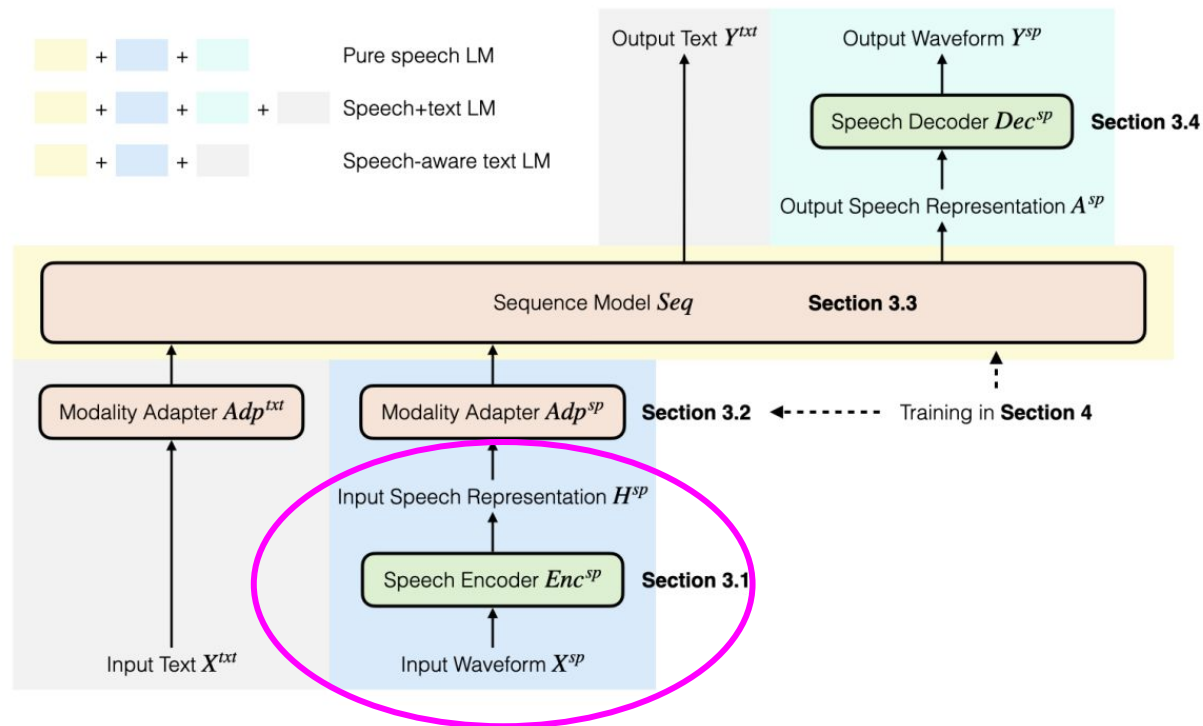


We know about
input text, noise



Figure 1: Overview of SLM architecture. See Sections 3 and 4 for more detailed descriptions of the components and training methods, respectively.

Let's talk SLM architecture



So now let's walk through the speech encoder!

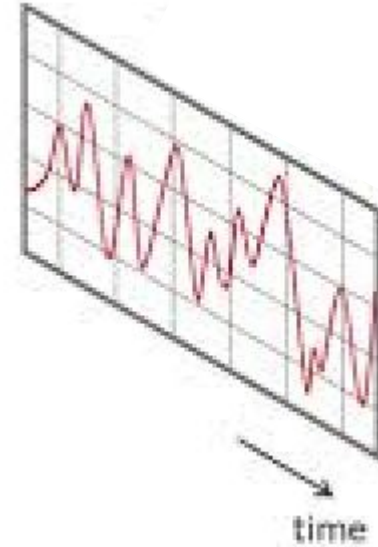
Figure 1: Overview of SLM architecture. See Sections 3 and 4 for more detailed descriptions of the components and training methods, respectively.

Let's start from the speech signal

What is it? It's air pressure over time.

Microphones pick up differences in air pressure – that's sound.

So we get a lil plot of that.

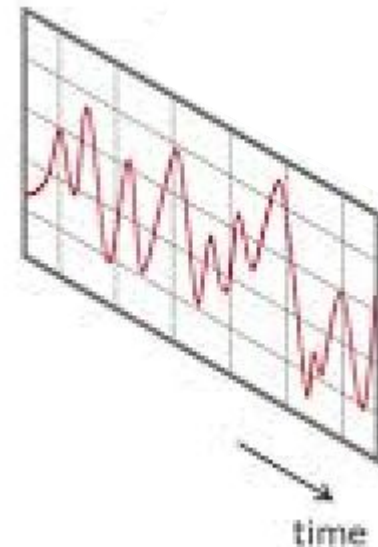


Let's start from the speech signal

What is it? It's air pressure over time.

Microphones pick up differences in air pressure – that's sound.

So we get a lil plot of that.



What's the problem? Why can't we just use that?

You have to super-duper oversample:

. Redundancy in Raw Waveform

- **Oversampling relative to perception:**

At 16 kHz, you get **16,000 samples per second**. But most speech information lives below ~4 kHz (Nyquist), and much of it is even lower (<2 kHz). → lots of extra samples.

- **Correlation across samples:**

Neighboring values in the waveform are highly correlated (they don't change randomly). For example, one period of a 200 Hz vowel spans ~80 samples — most of those points contain *the same information about the harmonic structure*.

- **Stationarity over frames:**

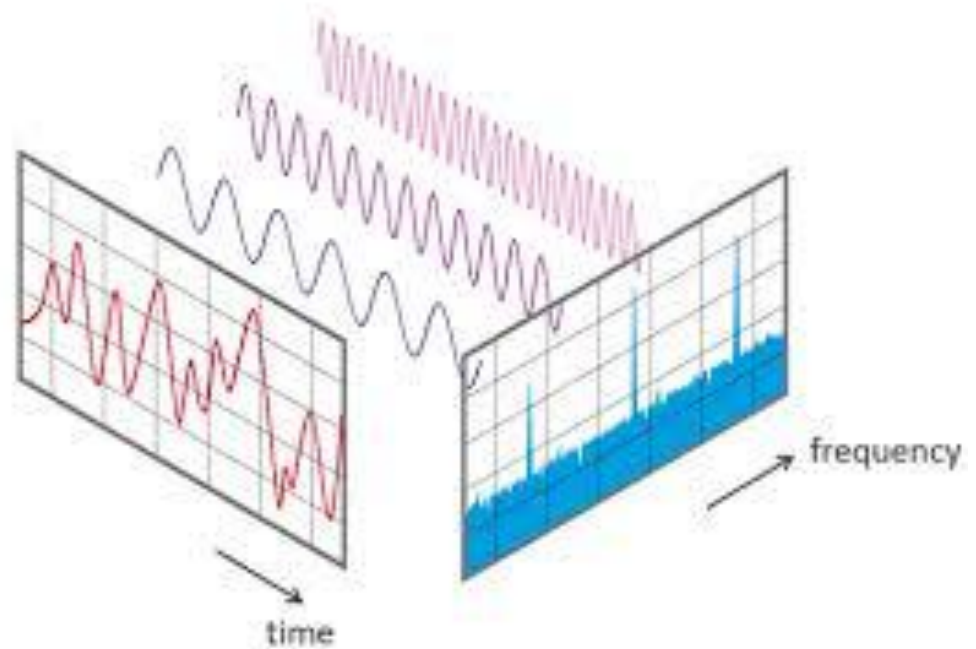
Speech characteristics (formants, pitch) are stable for 20–50 ms. That's **hundreds of samples** where the “meaning” doesn't change.

o raw data is very long, repetitive, and inefficient to model directly.

But, we have a trick!

We can do a fourier transform –

Just pull out the frequency components for that time chunk!

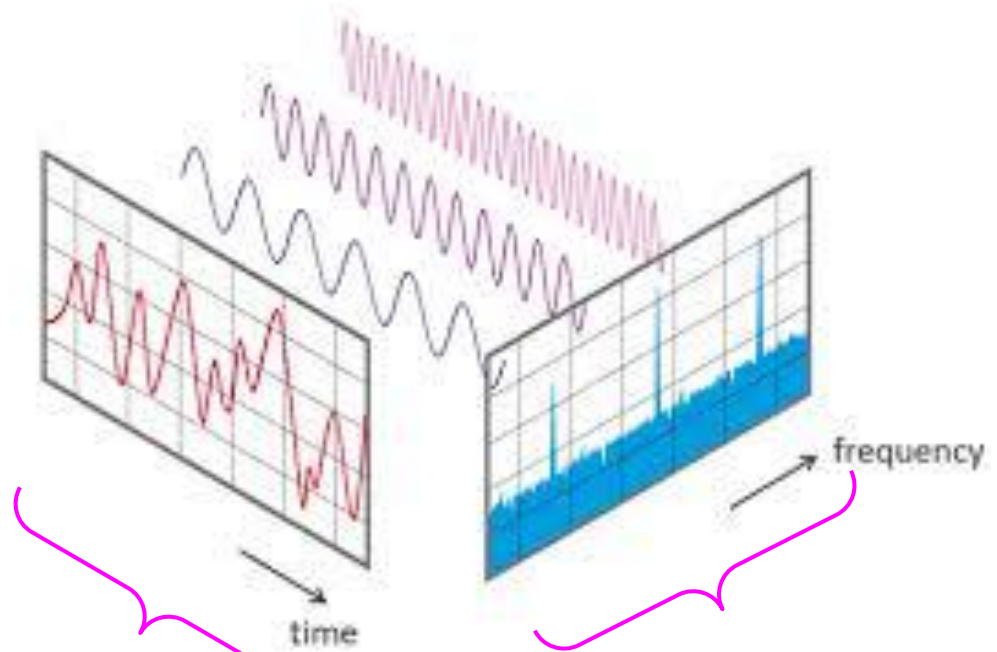


But, we have a trick!

We can do a fourier transform –

Just pull out the frequency components for that time chunk!

So now we can store 80 values per second *instead of 16,000 values per second*



Over 1 second, we're sampling 16,000 times to get this waveform

Or, we can just store 80 values for each frequency component

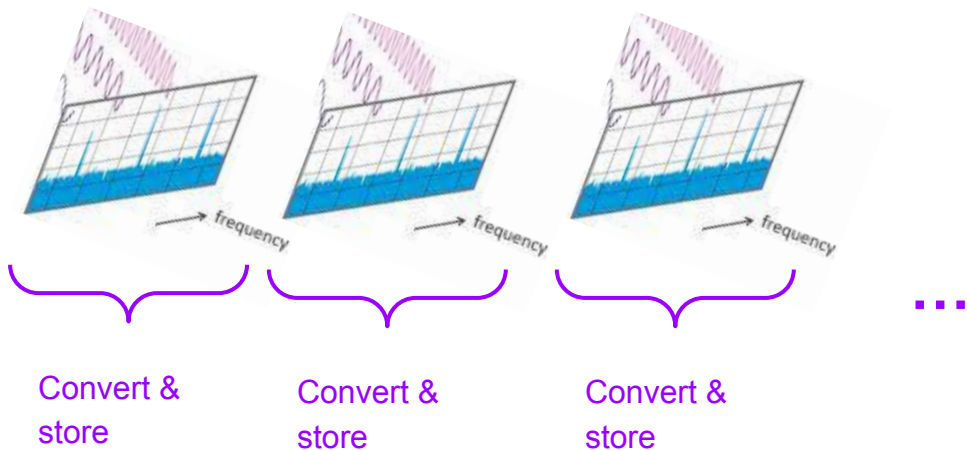
But, we have a trick!

We can do a fourier transform –

Just pull out the frequency components for that time chunk!

So now we can store 80 values per second *instead of 16,000 values per second*

So we store a “frequency snapshot” each second instead



But, we have a trick!

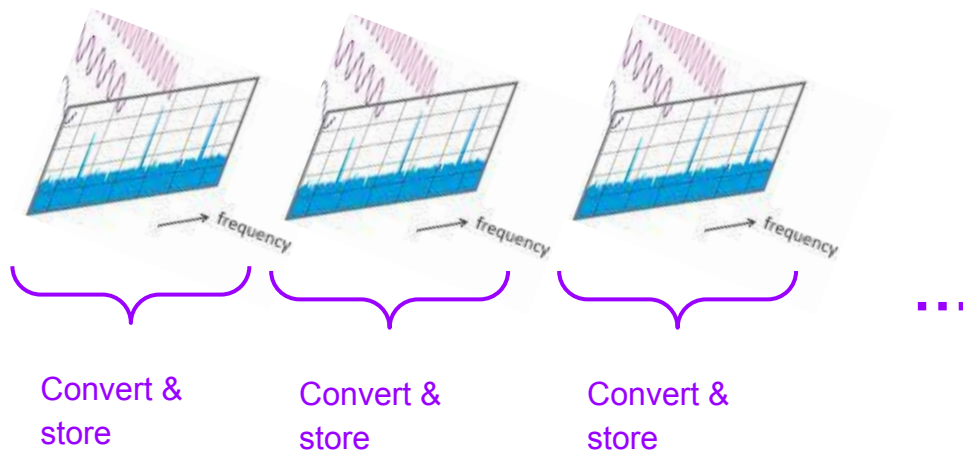
We can do a fourier transform –

Just pull out the frequency components for that time chunk!

So now we can store 80 values per second *instead of 16,000 values per second*

So we store a “frequency snapshot” each second instead

→ so we track how frequencies change *over time* while massively reducing data.

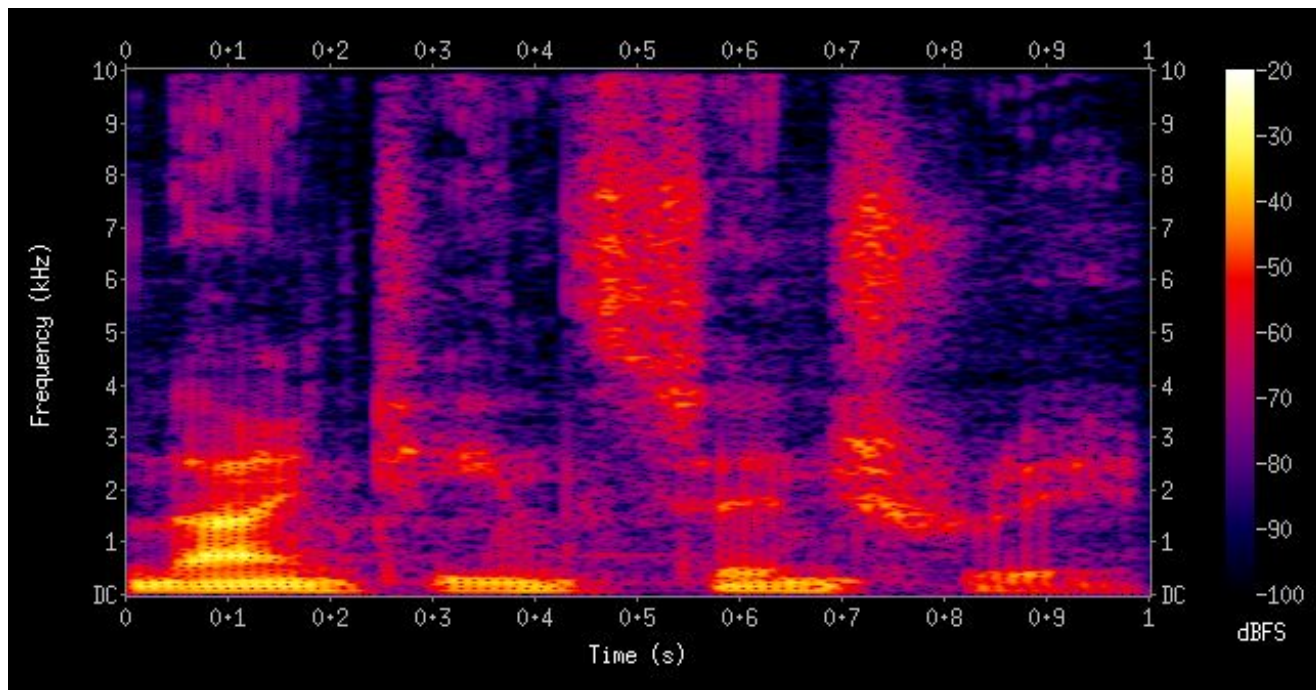


This technique is called the Short-Time Fourier Transform (STFT)!!!!!!!

It solves our over-sampling problem w/out losing important information <3

We end up with a **spectrogram** that looks like this!

Short-Time Fourier Transform (STFT) Result



In our example, this is the equivalent of **overlaying a couple seconds of the FTs** – so we're seeing how the *frequency components change over time*

Comparative Table of Function Approximation Methods

Aspect	Fourier Transform	Taylor Series	Neural Networks
Core Idea	Represent $f(x)$ as a sum of sinusoids: $f(x) = \sum_{k=-\infty}^{\infty} c_k e^{ikx}$, $c_k = \frac{1}{2\pi} \int f(x) e^{-ikx} dx$	Approximate $f(x)$ near a point a : $f(x) \approx \sum_{n=0}^N \frac{f^{(n)}(a)}{n!} (x-a)^n$	Learn approximation from data using nonlinear units: $f(x) \approx \sum_{i=1}^M w_i \sigma(\langle v_i, x \rangle + b_i)$ (1 hidden-layer NN)
Basis Functions	Fixed global sines/cosines.	Fixed monomials $(x-a)^n$.	Adaptive nonlinear features via activation functions.
Learning vs. Direct Computation	Coefficients c_k computed directly from integral formulas (no learning).	Coefficients $\frac{f^{(n)}(a)}{n!}$ computed directly from derivatives (no learning).	Coefficients w_i, v_i, b_i learned from data via optimization (gradient descent).
Domain of Usefulness	Oscillatory/periodic signals, frequency analysis.	Smooth analytic functions, valid locally around expansion point.	Arbitrary nonlinear, high-dimensional, non-analytic functions.
Local vs. Global	Global — each term affects all x .	Local — accurate only near a .	Both — architecture-dependent (RBFs local, deep nets capture global).
Interpretability	Coefficients $\leftrightarrow$ frequency content.	Coefficients $\leftrightarrow$ derivatives at expansion point.	Parameters (weights) not directly interpretable.
Convergence	Converges in L^2 for square-integrable f ; Gibbs phenomenon at discontinuities.	Converges within radius of convergence if f is analytic.	Universal Approximation Theorem: can approximate any continuous f on compact sets.
Computation	Fast Fourier Transform (FFT): $O(N \log N)$.	Easy if derivatives known; costly at high order.	Training costly (gradient descent), inference efficient.
Noise Sensitivity	Good for filtering (frequency cutoff).	Highly sensitive (derivatives amplify noise).	Can overfit to noise unless regularized.
Applications	Signal/image processing, PDEs, spectral analysis.	Physics models, perturbation methods, local expansions.	Pattern recognition, regression/classification, generative modeling.

Comparative Table of Function Approximation Methods

Aspect	Fourier Transform	Taylor Series	Neural Networks
Core Idea	Represent $f(x)$ as a sum of sinusoids: $f(x) = \sum_{k=-\infty}^{\infty} c_k e^{ikx}$, $c_k = \frac{1}{2\pi} \int f(x) e^{-ikx} dx$	Approximate $f(x)$ near a point a : $f(x) \approx \sum_{n=0}^N \frac{f^{(n)}(a)}{n!} (x - a)^n$	Learn approximation from data using nonlinear units: $f(x) \approx \sum_{i=1}^M w_i \sigma(w_i \cdot x) + b$
Basis Functions	Fixed global sines/cosines.	Fixed monomials $(x - a)^n$.	Adaptive basis functions.
Learning vs. Direct Computation	Coefficients c_k computed directly from integral formulas (no learning).	Coefficients $\frac{f^{(n)}(a)}{n!}$ computed directly from derivatives (no learning).	Requires training data and optimization.
Domain of Usefulness	Oscillatory/periodic signals, frequency analysis.	Smooth analytic functions, valid locally around expansion point.	Very broad, can approximate almost any function.
Local vs. Global	Global approximation.	Highly accurate only near a .	Both — a single layer is local, deep nets capture global).
Interpretability	Highly interpretable (coefficients relate to frequency).	Points $\leftrightarrow$ derivatives at point.	Parameters (weights) not directly interpretable.
Convergence	Good for smooth periodic functions.	Within radius of analyticity.	Universal Approximation Theorem: can approximate any continuous f on compact sets.
Computational Cost	Fast for many applications.	Easy if derivatives known; costly at high order.	Training costly (gradient descent), inference efficient.
Noise Sensitivity	Good (averaging effect).	Highly sensitive (derivatives amplify noise).	Can overfit to noise unless regularized.
Applications	Signal/image processing, PDEs, spectral analysis.	Physics models, perturbation methods, local expansions.	Pattern recognition, regression/classification, generative modeling.

Speech moves through the air in waves of air pressure... so this “math tool” pretty realistically models how speech works & is cheap to use & solves our downsampling problem

There are lots of signals where we don’t really KNOW how they work... so this “math tool” is the best tool we can use

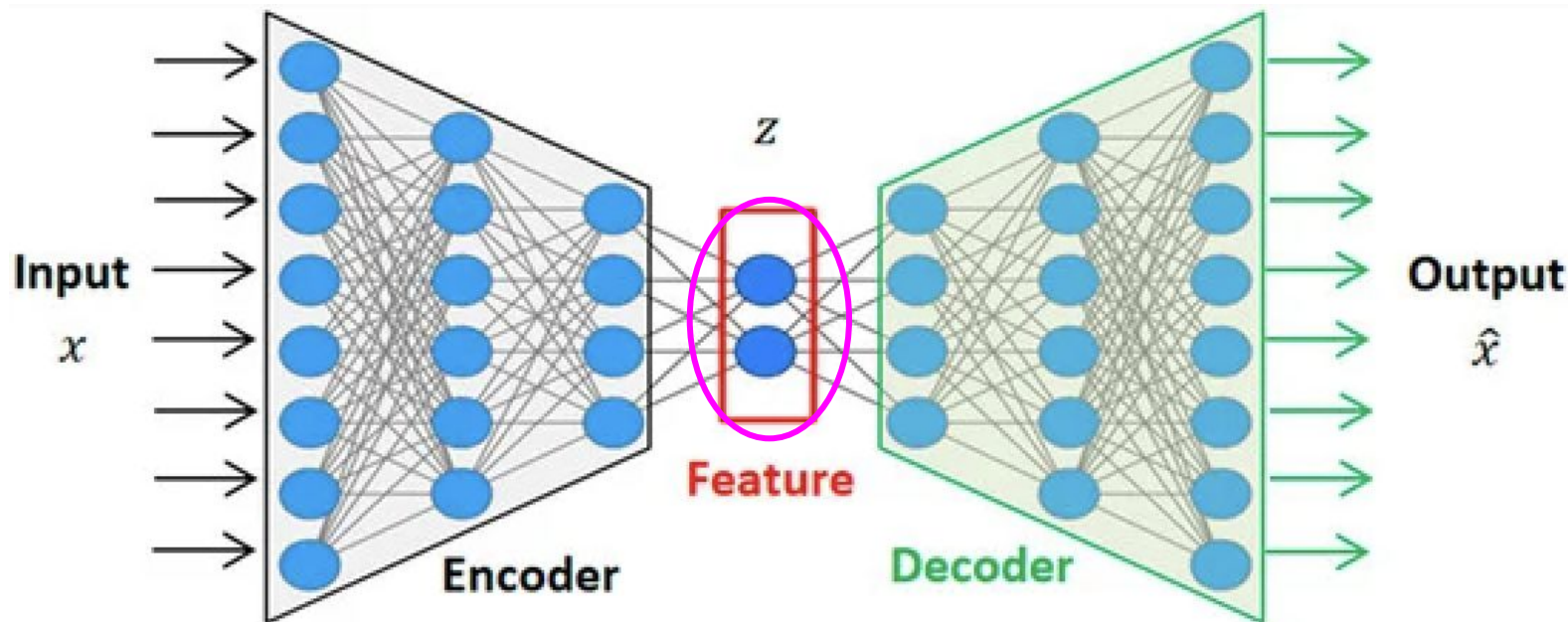
There are actually lots of ways to encode speech features –

3.1.1 Continuous Features

To extract informative representations from raw waveforms, a speech representation model—either a learned encoder or a digital signal processing (DSP) feature extractor—converts speech into continuous features. These continuous features may include:

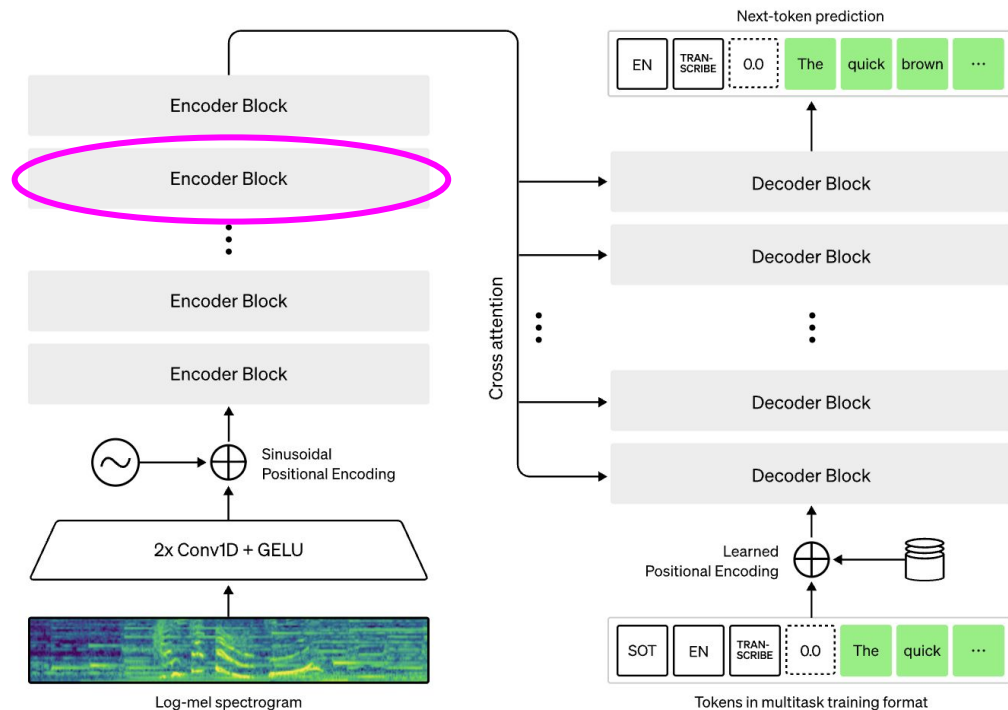
- ✓ 1. Traditional spectrogram features, such as mel filter bank features (Huang et al., 2001).
- 2. Hidden representations from self-supervised learning-based (SSL) speech encoders, such as wav2vec 2.0 (Baevski et al., 2020), HuBERT (Hsu et al., 2021), or WavLM (Chen et al., 2022).
- 3. Hidden representations from supervised pre-trained models, such as Whisper (Radford et al., 2023) or USM (Zhang et al., 2023b).
- 4. Hidden representations from neural audio codec models, such as SoundStream (Zeghidour et al., 2022) or EnCodec (Défossez et al., 2023).

2. Hidden representations from self-supervised learning-based (SSL) speech encoders, such as wav2vec 2.0 (Baevski et al., 2020), HuBERT (Hsu et al., 2021), or WavLM (Chen et al., 2022).



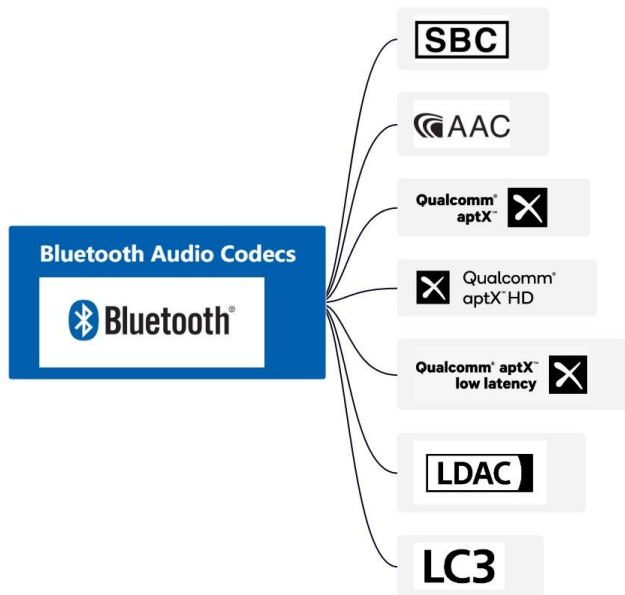
Take it from an autoencoder

3. Hidden representations from supervised pre-trained models, such as Whisper (Radford et al., 2023) or USM (Zhang et al., 2023b).



Take it from a Whisper layer

4. Hidden representations from neural audio codec models, such as SoundStream (Zeghidour et al., 2022) or EnCodec (Défossez et al., 2023).



Codec = compressed,
lossLESS audio
representation

(big topic rn)

Encoder

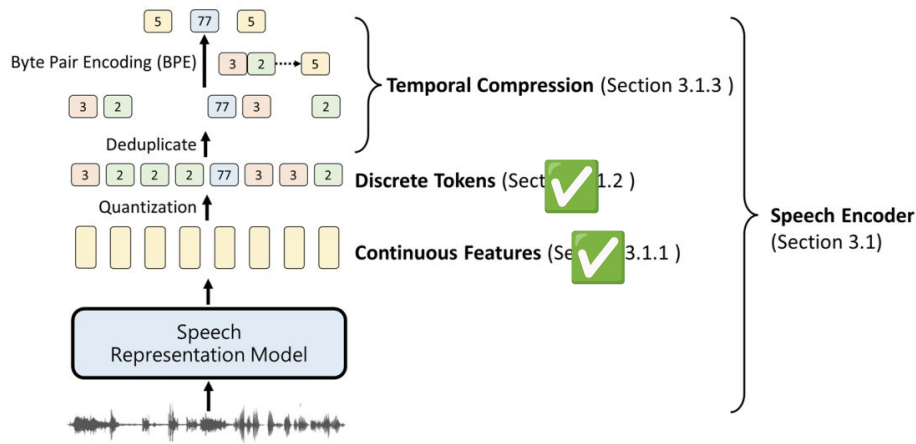


Figure 2: A general pipeline for speech encoders. Note that different encoders use different components of the pipeline. See Section 3.1 for more details.

Now we know how to go from
speech signal -> tokens

(more deets in the paper)

But we're not done – there's
ANOTHER
over-sampling/compression
issue...

Encoder

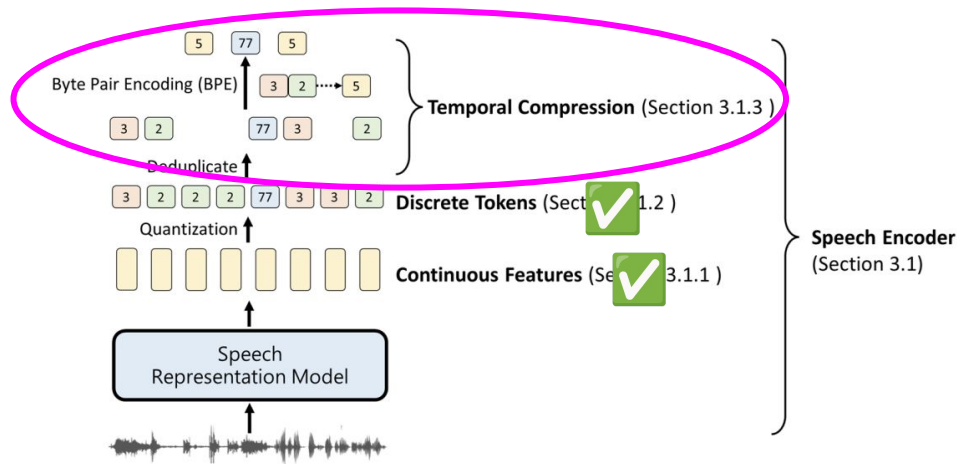


Figure 2: A general pipeline for speech encoders. Note that different encoders use different components of the pipeline. See Section 3.1 for more details.

So there are lots of tools to apply here to downsample AGAIN:

- BPE
- Multi-stream capture
- Duration prediction

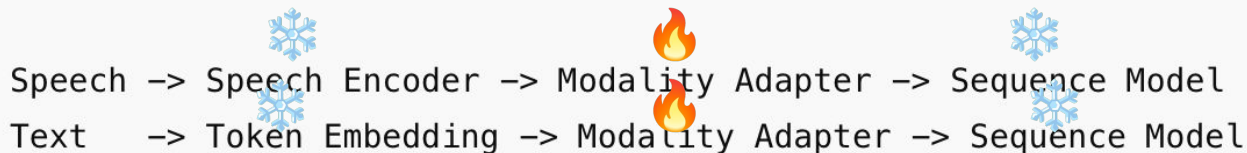
Big theme: downsampling is a big deal in audio/speech world!!!!!!!!!!

3 Main Approaches to TRAIN 🚂 SLMs

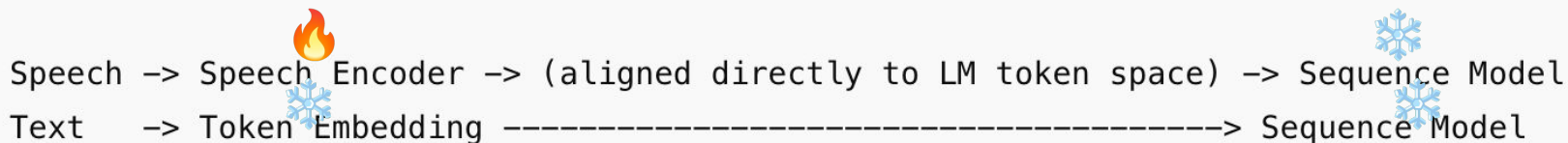
Train the (big) sequence model (aka the LLM) [most expensive \$\$\$]



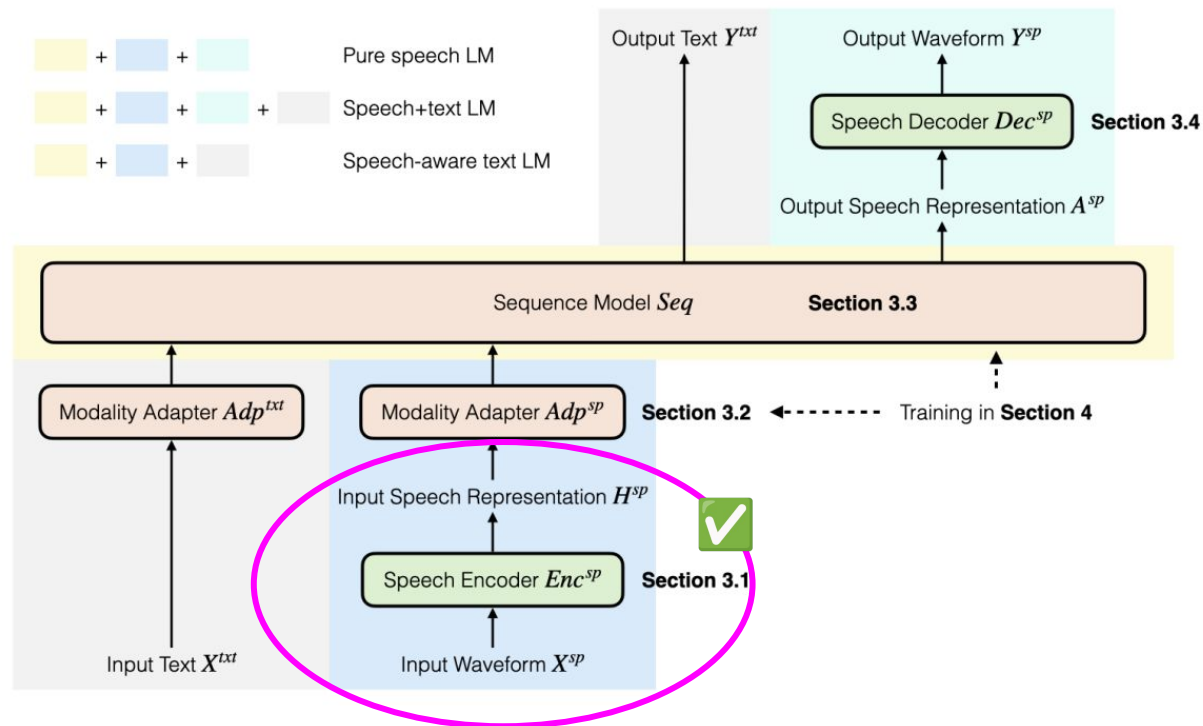
Train the modality adapters [least expensive \$]



Train the speech encoder [mid-expensive \$\$] [PS this is kinda like just training 1 modality adapter]



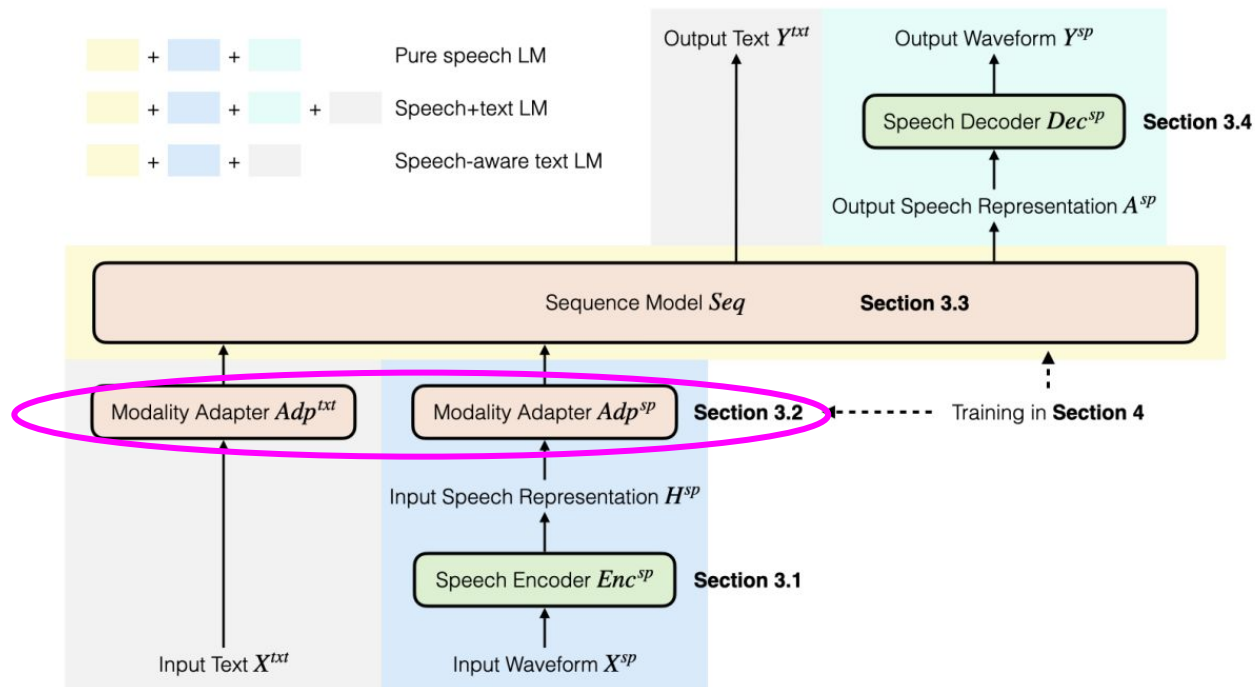
Let's talk SLM architecture



Ok yay so now we know about the speech encoder

Figure 1: Overview of SLM architecture. See Sections 3 and 4 for more detailed descriptions of the components and training methods, respectively.

Let's talk SLM architecture



Let's talk about
the modality
adapters

Figure 1: Overview of SLM architecture. See Sections 3 and 4 for more detailed descriptions of the components and training methods, respectively.

Modality Adapters

Train the modality adapters [least expensive \$]

Speech  -> Speech Encoder  -> Modality Adapter  -> Sequence Model 
Text  -> Token Embedding  -> Modality Adapter  -> Sequence Model 

In many SLMs (especially speech-aware text LMs), the speech encoder (Section 3.1) and the sequence model (Section 3.3) are initially developed separately and then combined. It is therefore necessary to somehow align the output of the speech encoder with the expectations of the sequence model, and this is the role of the modality adapter. The modality adapter is typically trained on downstream tasks or, in the case of speech+text LMs, as part of pre-training (see Section 4 for more details on training).

The whole point is to make the **SPEECH TOKENS** and **TEXT TOKENS** match up correctly!!!!

Lots of ways to do a Modality Adapter... Broad Overview

Comparative Table of Adapter Architectures in Multimodal Models

Adapter Architecture	Speech–Language	Vision–Language	Recommender–Language
Linear / MLP Projection	Project continuous speech embeddings (e.g., wav2vec2, HuBERT) into LLM token space.	Project vision encoder outputs (ViT/CNN/CLIP) into the same dimensionality as LLM embeddings.	Map item embeddings or categorical features into LLM-compatible vectors.
Quantization (Discrete Tokens)	VQ-VAE, k-means clustering to produce “pseudo-text” tokens from speech, directly fed into LLM.	Less common (though discrete VQ-image tokens exist, e.g., VQGAN), but most models use continuous projection instead.	Rare — item IDs can act as discrete tokens, sometimes directly embedded as vocab.
Cross-Attention Adapters	LLM attends to speech embeddings via cross-attention blocks (e.g., spoken QA grounding).	Popular: Flamingo inserts cross-attention layers where LLM queries image embeddings.	LLM attends over user history/item embeddings via cross-attention, aligning behavioral context with language tokens.
Prefix / Prompt Tuning	Encode speech embeddings into a sequence of pseudo-tokens prepended as “soft prompts.”	Encode image embeddings as prefix tokens before text input (BLIP-2 Q-Former, Kosmos-1).	Encode user/item history as prompt tokens describing context (P5, TALM).
Bottleneck Adapters	Small trainable modules (e.g., LSTM, Conformer bottlenecks) compress variable-length speech into manageable embeddings.	Smaller adapter layers inserted in the vision encoder or between encoder–decoder.	Trainable bottlenecks align high-dim item features with LLM hidden states.
Fusion / Multi-Modal Attention	Speech features and LLM hidden states jointly processed via fusion transformer blocks.	Multimodal transformer fusion (e.g., CLIP-LMs, PaLM-E) combining vision & text tokens.	Fusion of user/item embeddings with natural language tokens for personalized reasoning.

The whole point is to make the [SPEECH/VISION/REC] TOKENS and LLM TOKENS match up correctly!!!!

Lots of ways to do a Modality Adapter...

Connectionist Temporal Classification (CTC)-based compression: This method compresses H^{SP} (Eq. 1) according to the posterior distribution from a CTC-based speech recognizer (Gaido et al., 2021). CTC (Graves et al., 2006), a commonly used approach for ASR, assigns each time step a probability distribution over a set of label tokens, including a blank (“none of the above”) symbol. The time steps with high non-blank probabilities indicate segments that are likely to carry important linguistic information. CTC compression aggregates the frame-level labels, specifically by merging repeated non-blank labels and removing blanks. This approach produces a compressed representation intended to retain the relevant content of the original sequence while significantly reducing its length (Wu et al., 2023b; Tsunoo et al., 2024).

How CTC collapsing works



How do you make speech tokens & text tokens match up correctly?
Compression!

Lots of ways to do a Modality Adapter...

How do you
make speech
tokens & text
tokens match up
correctly?
Compression!

Q-Former: The Q-Former (Li et al., 2023) is an adapter that produces a fixed-length representation by encoding a speech representation sequence of arbitrary length into M embedding vectors, where M is a hyperparameter (Lu et al., 2024b).

Let the input speech representation sequence be:

$$X = \{x_1, x_2, \dots, x_{L'}\}, \quad x_i \in \mathbb{R}^{d'}, \quad (3)$$

where L' is the sequence length and d' is the dimension of the embeddings.

To achieve a fixed-length representation, Q-Former introduces M trainable query embeddings:

$$Q = \{q_1, q_2, \dots, q_M\}, \quad q_i \in \mathbb{R}^{d'}. \quad (4)$$

These queries interact with X via a cross-attention mechanism:

$$\text{Attn}(Q, X) = \text{softmax} \left(\frac{QW_Q(XW_K)^T}{\sqrt{d'}} \right) XW_V, \quad (5)$$

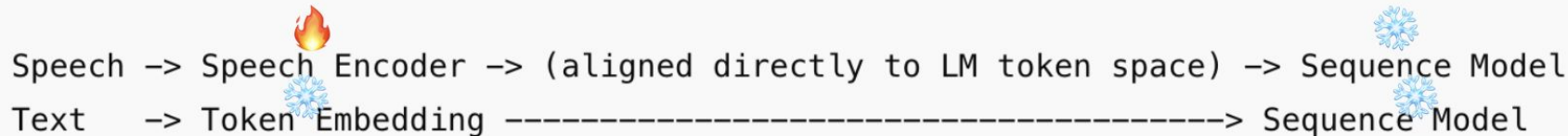
where W_Q , W_K , and W_V are learnable projection matrices. The result is a sequence of M embeddings.

In some approaches, instead of directly encoding the entire utterance into M vectors, a *window-level Q-Former* is applied (Yu et al., 2024; Pan et al., 2024; Tang et al., 2024) to retain temporal information. In the window-level Q-Former, the input embedding sequence is segmented, and the Q-Former is applied to each segment.

Lu et al. (2024a) compare the Q-Former with CNN-based modality adapters in a speech-aware text LM, finding that the Q-Former produces better performance on the Dynamic-SUPERB benchmark (Huang et al., 2024) (see Section 7 for more on this and other SLM benchmarks).

This approach
plugs more
directly into the
LLM (in the
attention
mechanism)

Train the speech encoder [mid-expensive \$\$] [PS this is kinda like just training 1 modality adapter]



Implicit alignment Speech and text modalities can be implicitly aligned through techniques such as the “modal-invariance trick” (Fathullah et al., 2024) or behavior alignment (Wang et al., 2023a). The idea is that the model should produce identical responses regardless of the input modality, provided the input conveys the same meaning. This approach often utilizes ASR datasets. The text transcript is input to a text LLM to generate a text response, while the corresponding speech recording is input into the SLM, which is trained to generate the same text response. Another idea found to be useful for implicit alignment is training spoken LLMs for audio captioning, where a spoken LLM takes audio as input and outputs its description. It has been observed that training a spoken LLM solely through audio captioning can generalize to tasks it has never seen during training (Lu et al., 2024a;b).

Explicit alignment Speech and text modalities can also be explicitly aligned by matching speech features to corresponding text embeddings, via optimization of appropriate distance/similarity measures. For example, Wav2Prompt (Deng et al., 2024) and DiVA (Held et al., 2024) align modalities by minimizing the L_2 distance between speech features and the token embeddings of their transcripts in a text LLM while keeping the text embeddings fixed.

Train the (big) sequence model (aka the LLM) [most expensive \$\$\$]



You just train the whole thing end-to-end, it's expensive
and there's not really any tricks

→ ***Another great reference!***

Summarizing Speech: A Comprehensive Survey

**Fabian Retkowski¹ Maïke Züfle¹ Andreas Sudmann² Dinah Pfau³
Shinji Watanabe⁴ Jan Niehues¹ Alexander Waibel^{1,4}**

¹KIT ²University Bonn ³Deutsches Museum ⁴CMU

{fabian.retkowski,maike.zuefle,jan.niehues,alex.waibel}@kit.edu
asudmann@uni-bonn.de d.pfau@deutsches-museum.de swatanab@andrew.cmu.edu

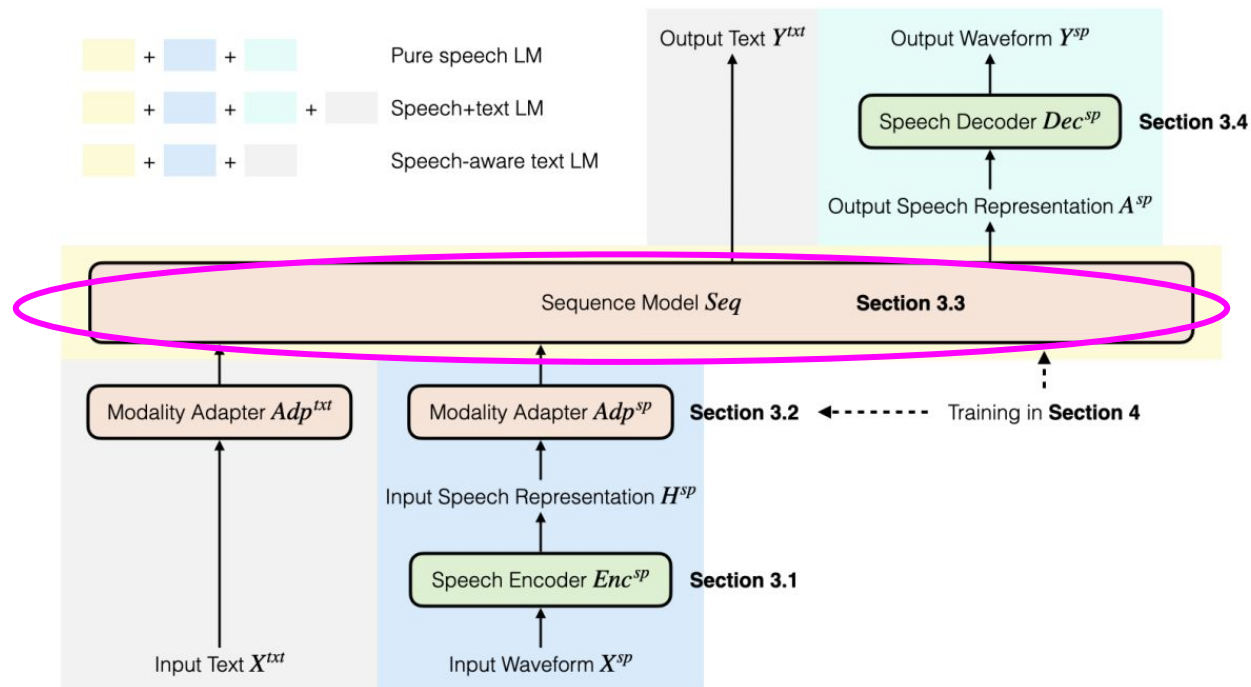
EMNLP '25

→ *Another great reference!*

Reference	Audio Encoder	Projector	LLM
Fathullah et al. (2024)	🔥 Conformer (Gulati et al., 2020)	🔥 Linear	❄️ LLaMA-2-7B-chat (Touvron et al., 2023)
Shang et al. (2024)	🔥 Conformer (Gulati et al., 2020)	🔥 Q-Former (Li et al., 2023)	≈ LLaMA-2-7B-chat (Touvron et al., 2023)
Microsoft et al. (2025)	🔥 Conformer (Gulati et al., 2020)	🔥 MLP	❄️ Phi-4-mini-instruct (Microsoft et al., 2025)
Kang and Roy (2024)	🔥 HuBERT-Large (Hsu et al., 2021)	🔥 Linear	❄️ MiniChat-3B (Zhang et al., 2024a)
Züfle et al. (2025)	❄️ HuBERT-Large (Hsu et al., 2021)	🔥 Q-Former (Li et al., 2023)	❄️ LLaMA3.1-8B-Instruct (Grattafiori et al., 2024)
He et al. (2025)	❄️ MERaLiON-Whisper (He et al., 2025)	🔥 MLP	≈ SEA-LION V3 (He et al., 2025)
Chu et al. (2024)	🔥 Whisper-large-v3 (Radford et al., 2023)	🔥 Linear	🔥 Qwen-7B (Bai et al., 2023)
Eom et al. (2025)	❄️ Whisper-large-v2 (Radford et al., 2023)	🔥 Q-Mamba (Eom et al., 2025)	🔥 Mamba-2.8B-Zephyr (xiuyul/mamba-2.8b-zephyr)

Table 2: Overview of Audio Encoder → Projector → LLM Architectures (🔥 trainable, ❄️ frozen, ≈ LoRA)

Let's talk SLM architecture

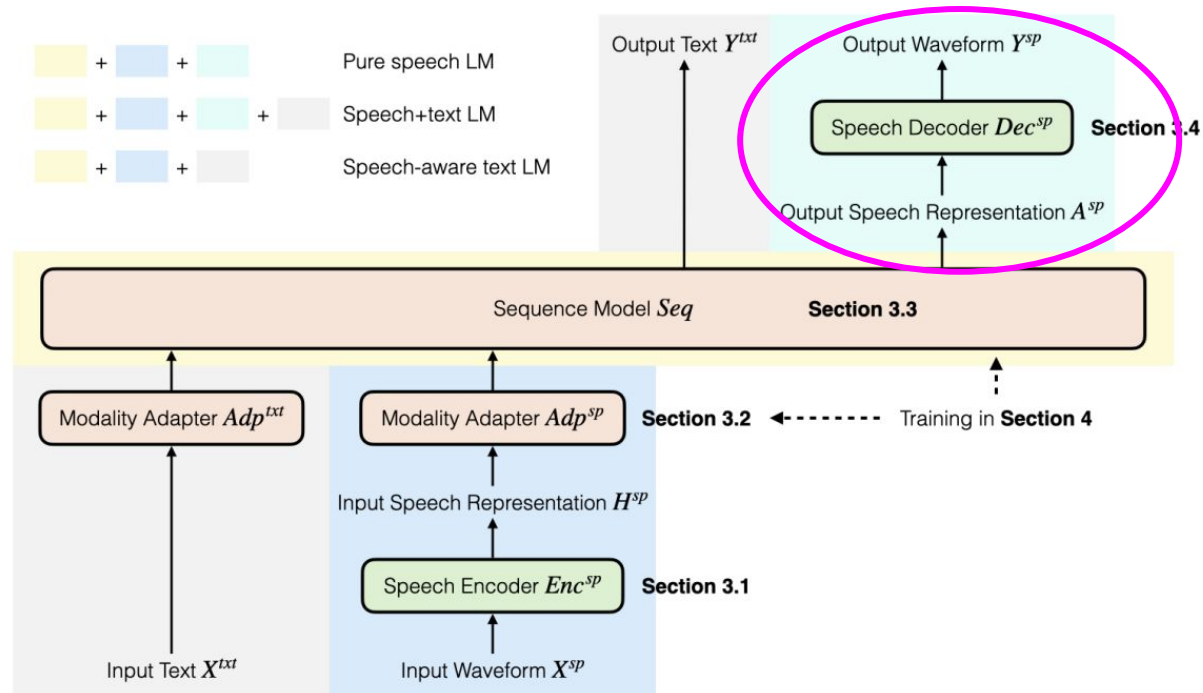


Let's talk about
the **sequence
model**

... not much to
say, **this is the
core LLM**

Figure 1: Overview of SLM architecture. See Sections 3 and 4 for more detailed descriptions of the components and training methods, respectively.

Let's talk SLM architecture



Let's talk about
the **speech
decoder**

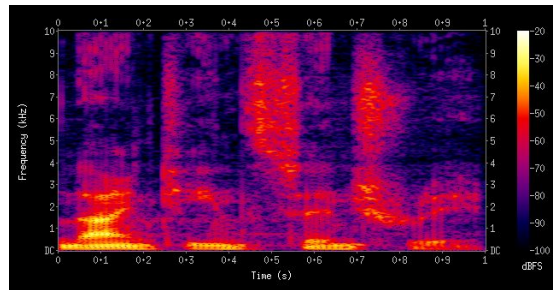
Figure 1: Overview of SLM architecture. See Sections 3 and 4 for more detailed descriptions of the components and training methods, respectively.

3.4 Speech Decoder

The speech decoder converts speech representations—whether continuous features, phonetic tokens, or audio codec tokens—back into waveforms. The speech decoder can take various forms:

1. Vocoder (Kong et al., 2020) for continuous features, similar to those used in traditional synthesis systems. For instance, in Spectron (Nachmani et al., 2024), a generated mel spectrogram is synthesized into audio using the WavFit vocoder (Koizumi et al., 2023).

They generate these →



2. Unit-based vocoder (Polyak et al., 2021) based on HiFi-GAN (Kong et al., 2020) for phonetic tokens. These vocoders take phonetic tokens as inputs and optionally combine them with additional information to improve synthesis quality. For example, when phonetic tokens are deduplicated, a duration modeling module is often included in the vocoder (Lakhotia et al., 2021).
3. Codec decoder (Guo et al., 2025). When the SLM generates audio codec tokens, these tokens can be input directly into the corresponding pre-trained audio neural codec decoder (without additional training) to get the waveform.

They generate codecs

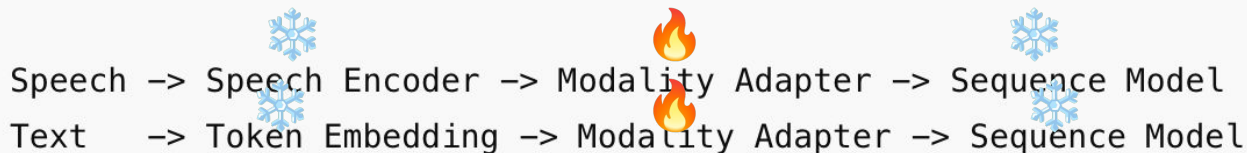
3 Main Approaches to TRAIN 🚂 SLMs

Train the (big) sequence model (aka the LLM) [most expensive]

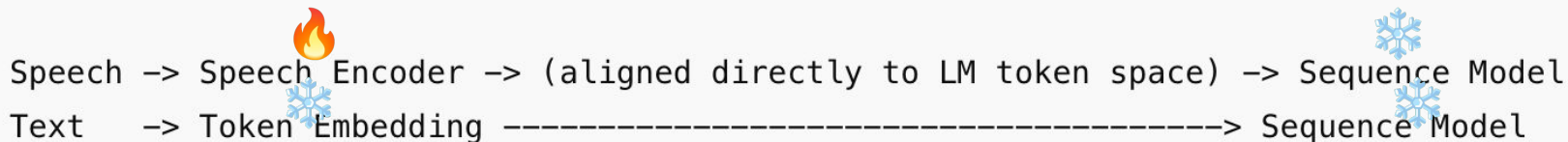


Let's revisit this...

Train the modality adapters [least expensive \$]



Train the speech encoder [mid-expensive \$\$] [PS this is kinda like just training 1 modality adapter]



They're doing the training with special speech tokens!!!!

$$p(\cdot | \text{speech}, \text{instruction}) \left\{ \begin{array}{l} p(\cdot | \text{speech}, \text{text instruction}), \\ p(\cdot | \text{speech}, \text{speech instruction}) \end{array} \right.$$

Either one! Just depending on setup

TASK-SPECIFIC TUNING

INSTRUCTION TUNING

Token / Specifier	Training Paradigm	Task	Example Input	Expected Output
<task:TRANSCRIBE>	Task-specific	ASR	[speech: "Hello, how are you today?"]	"Hello, how are you today?"
<task:SUMMARIZE>	Task-specific	Speech Summarization	[speech: lecture audio]	"The lecture explains Fourier transforms as frequency decompositions."
<task:TRANSLATE_EN>	Task-specific	Speech Translation	[speech: "Bonjour, comment ça va?"]	"Hello, how are you?"
<speech_instruction:CLASSIFY_EMOTION>	Instruction-tuning	Emotion Detection	[speech: angry utterance]	"angry"
<speech_instruction:CONVERSATION>	Instruction-tuning	Dialogue Response	[speech: "What time is it?"]	"It's 3:00 PM."
<speech_instruction:MULTIMODAL_QA>	Instruction-tuning	Multimodal Fusion (Speech + Text)	[speech: "Look at this picture."] + [text: "What color is the car?"]	"The car is red."

Task-specific = "menu of fixed commands" (rigid, controlled).
Instruction tuning = "free-form instructions" (flexible, generalizable).

So let's talk about how you design your loss aka objective function for training 🔥

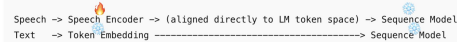
If you
want to
do
THIS

Set up your
math like
THIS

Train the (big) sequence model (aka the LLM) [most expensive \$\$\$]



Train the speech encoder [mid-expensive \$\$] [PS this is kinda like just training 1 modality adapter]



Train the modality adapters [least expensive \$]



**TASK-SPECIFIC
TUNING**

**INSTRUCTION
TUNING**

How CTC collapsing works



Stage

Objective / Loss

Formula

Pretraining (AR LM)

Autoregressive next-token prediction

$$\mathcal{L}_{\text{AR}} = -\sum_t \log p(y_t | y_{<t}, x)$$

Adapter / Alignment

L2 alignment

$$\mathcal{L}_{\text{align}} = \|f_{\text{speech}}(X^{\text{sp}}) - f_{\text{text}}(X^{\text{txt}})\|_2^2$$

Contrastive (CLIP-style)

$$\mathcal{L}_{\text{contrast}} = -\log \frac{\exp(\text{sim}(h^{\text{sp}}, h^{\text{txt}})/\tau)}{\sum_{h'} \exp(\text{sim}(h^{\text{sp}}, h')/\tau)}$$

Task-specific training

Cross-entropy w/ task token

$$\mathcal{L}_{\text{task}} = -\sum_t \log p(y_t | y_{<t}, X^{\text{sp}}, \langle \text{task} \rangle)$$

Instruction tuning

Cross-entropy w/ instruction

$$\mathcal{L}_{\text{instr}} = -\sum_t \log p(y_t | y_{<t}, X^{\text{sp}}, \text{instruction})$$

Specialized (optional)

CTC (speech alignment)

$$\mathcal{L}_{\text{CTC}} = -\log p(\text{transcript} | X^{\text{sp}})$$

Reconstruction

$$\mathcal{L}_{\text{recon}} = \|X^{\text{sp}} - \hat{X}^{\text{sp}}\|_2^2$$

Multi-task combo

Weighted mix

$$\mathcal{L} = \sum_i \lambda_i \mathcal{L}_i$$

You'll
probably
end up
combo-ing
multiple of
these
together

This is cool:
interruption
handling w/ a
duplex model

Duplex = 2 parallel
streams for user and SLM,
open at all times → robust to
interruptions, no assumption
of “turn-taking” really

Walkie-Talkie vs. Phone Call

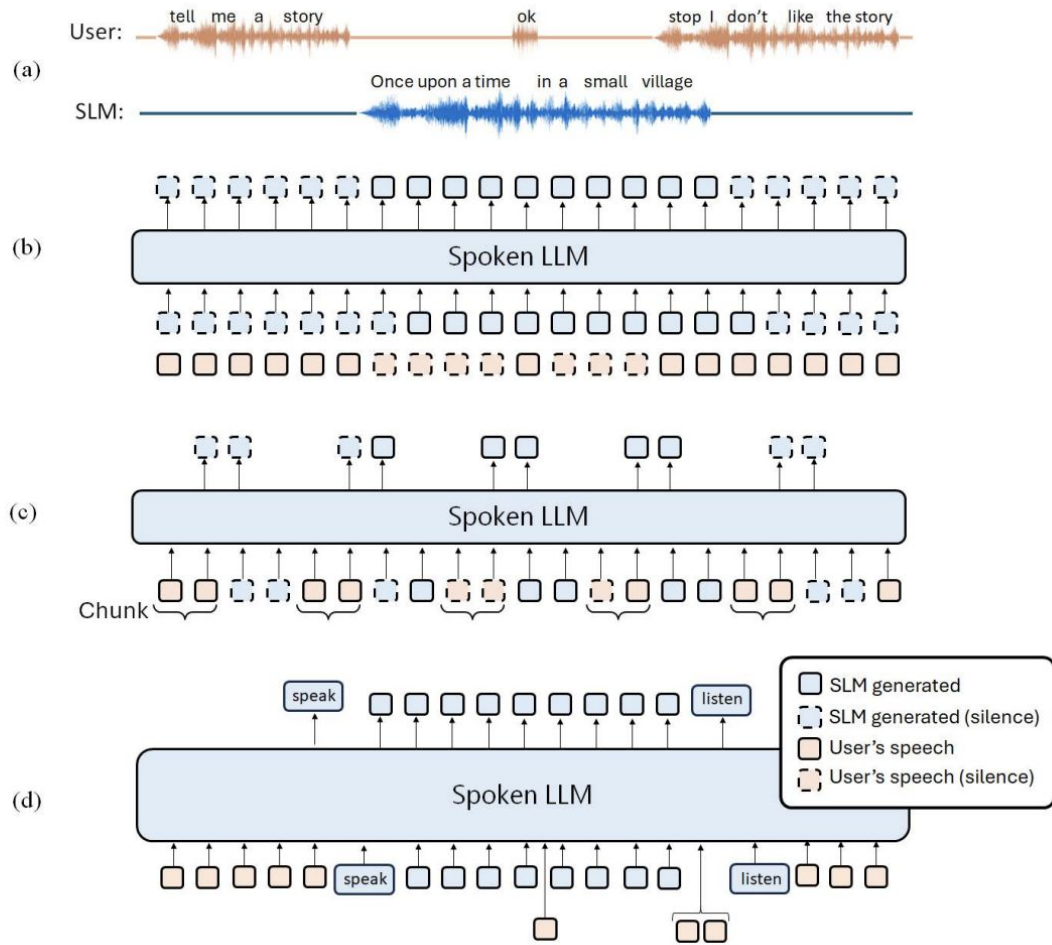


Figure 5: Full-duplex speech conversation. (a) An example of full-duplex speech conversation between a user and an SLM. (b) Dual-channel approach. (c) Time-multiplexing approach (with equal chunks). (d) Time-multiplexing approach (where the SLM controls the switching between listening and speaking modes).

8 Challenges and future work

Model architecture The optimal representation of speech within SLMs remains unclear. Speech representations in SLMs include both discrete and continuous varieties, derived from a wide range of encoders. This design choice can also influence other architectural choices in an SLM, for example depending on the information rate of the encoder and whether it encodes more phonetic or other types of information.

Another open question is determining the best method to combine speech and text, which applies to all aspects of SLM modeling and training. We have described various choices of modality adapters and approaches for interleaving speech and text. These have not been thoroughly compared, so the effect of each modeling choice is still unclear.

A final architectural challenge is that current SLMs are large and slow, making them impractical for real-time and on-device settings. To some extent this is because various compression algorithms (e.g., (Lai et al., 2021; Peng et al., 2023a; Ding et al., 2024)) and alternative architectures (e.g., Park et al. (2024)) have not been widely applied to SLMs. However, there is also an inherent efficiency challenge that arises when combining multiple pre-trained components, sometimes with different architectures and frame rates.

Second,

**Let's talk about the
robustness of
pipeline approaches.**

WE ARE
HERE

Pipeline Approach vs. End-to-End Approach



✓ preferred by industry for *controllability* → think custom vocabulary & ease of debugging

✓ currently in research stage but very promising → major barrier imo is downsampling problem (will explain)

Pipeline approaches are super susceptible to issues processing **disfluencies**:



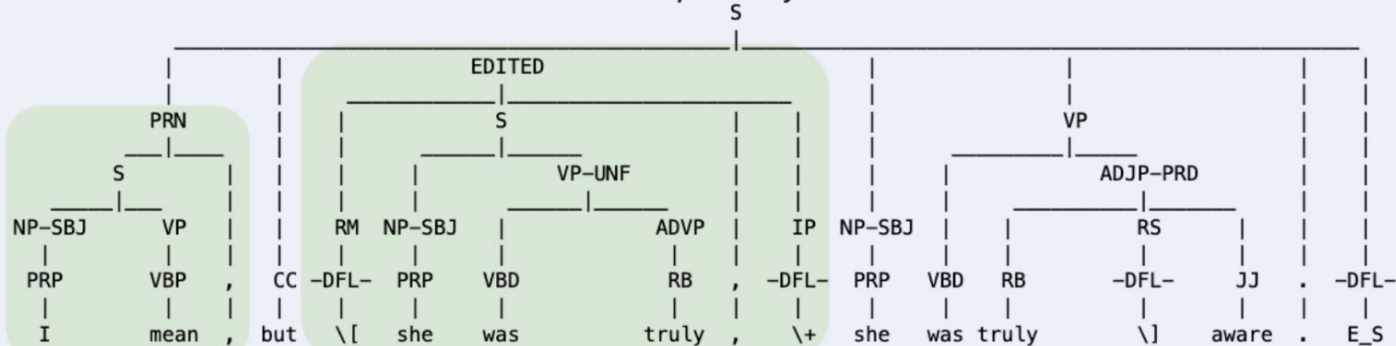
Fig. 1. Removing disfluencies such as INTJ (*uh*), EDITED (*gas station* is replaced with *grocery store*), and PRN (*you know*) ensures clean text input for downstream tasks. Our

Pipeline approaches are super susceptible to issues processing **disfluencies**:

t_{tree}

```
((S(PRN(S (NP-SBJ (PRP I))(VP (VBP mean)))(, ,))(CC but)(EDITED(RM(-DFL- \[])(S(NP-SBJ(PRP she))(VP-UNF(VBD wa  
s)(ADVP(RB truly)))(, ,)(IP(-DFL- \+)))(NP-SBJ(PRP she))(VP(VBD was)(ADJP-PRD(RB truly)(RS(-DFL- -\)))(JJ awar  
e)))(. .)(-DFL- E_S)))
```

Equivalently:



$t_{disfluent}$

"I mean, but she was truly, she was truly aware."

t_{tag}

[PRN, PRN, NONE, EDITED, EDITED, EDITED, NONE, NONE, NONE, NONE]

t_{fluent}

"But she was truly aware."

Quantifying the Impact of Disfluency on Spoken Content Summarization

Maria Teleki, Xiangjue Dong, James Caverlee

Texas A&M University

College Station, Texas, USA

{mariateleki, xj.dong, caverlee}@tamu.edu

LREC-COLING '24

Simple disfluencies can kill model performance.

Original

Hello and welcome to our podcast! Let's get right to it. Today we're going to be interviewing a very special guest, someone I know you guys have been excited about having on the show.

Repeats with N=3

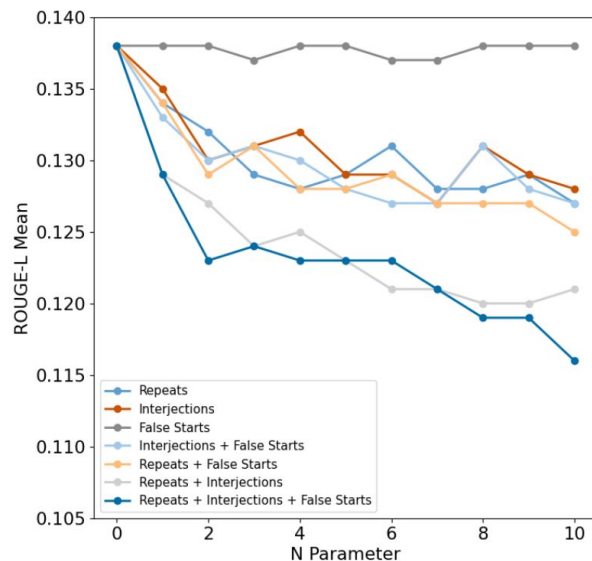
Hello and welcome to our podcast! Let's get **get get get** right to it. Today we're going to be interviewing a **a a a** very special guest, someone I know you guys have been excited about having on the show.

Interjections with N=3

Hello and welcome to our podcast! Let's get right **uh okay okay** to it. Today we're going to be interviewing a very special **um so I mean** guest, someone I know you guys have been excited about having on the show.

False Starts with N=3

Hello and welcome to our podcast! Let's get right to it. Today we're **today we're today we're today we're** going to be interviewing a very special guest, someone I know you guys have been excited about having on the show.



(a) ROUGE-L over increased N on **BART** model.

Decrease Rouge-L

Increase N

Ok but BART's old...

I agree! Here's our recent work evaluating:

- gpt-{4o,4o-mini,o4}
- llama-{1B,3B,8B,70B}
- qwen-{0.6B,1.7B,4B,8B}
- phi-4-mini
- mobileLLM{125M,350M,600M,1B}

DRES: BENCHMARKING LLMs FOR DISFLUENCY REMOVAL

Maria Teleki, Sai Janjur, Haoran Liu, Oliver Grabner, Ketan Verma, Thomas Docog, Xiangjue Dong, Lingfeng Shi, Cong Wang, Stephanie Birkelbach, Jason Kim, Yin Zhang, James Coverlee

ArXiv '25

Texas A&M University

Long story short, these models still struggle!

M	k	$\mathcal{E}_P$	$\mathcal{E}_R$	$\mathcal{Z}_I$	$\mathcal{Z}_P$	$\mathcal{E}_P$	$\mathcal{E}_R$	$\mathcal{Z}_I$	$\mathcal{Z}_P$	$\mathcal{E}_R$	$\mathcal{Z}_P$
gpt-4o											
f	0	74.19 _{0.63}	77.91 _{0.63}	73.18 _{0.58}	83.86 _{0.15}	71.25 _{0.57}	52.52 _{0.91}	70.52 _{0.32}	75.56 _{0.99}	68.17 _{0.80}	86.13 _{0.40}
f	1	76.13 _{0.63}	76.94 _{0.68}	78.02 _{0.85}	85.11 _{0.77}	77.37 _{0.64}	62.71 _{0.98}	68.85 _{0.71}	74.35 _{0.65}	66.02 _{0.68}	86.05 _{0.47}
f	3	73.99 _{0.31}	81.33 _{0.10}	70.75 _{0.11}	77.83 _{0.20}	72.93 _{0.25}	49.35 _{0.73}	68.70 _{0.40}	73.20 _{0.59}	67.68 _{0.32}	86.64 _{0.10}
f	5	73.06 _{0.41}	81.24 _{0.69}	68.75 _{0.19}	76.92 _{0.55}	71.49 _{0.57}	42.39 _{0.38}	69.03 _{0.17}	76.93 _{0.18}	65.86 _{0.68}	84.77 _{0.10}
s	0	78.17 _{0.64}	78.44 _{0.81}	78.30 _{0.27}	82.92 _{0.67}	77.55 _{0.76}	69.80 _{0.86}	72.69 _{0.59}	75.61 _{0.75}	70.48 _{0.23}	85.20 _{0.23}
s	1	82.38 _{0.11}	83.61 _{0.01}	81.53 _{0.77}	83.77 _{0.41}	79.74 _{0.46}	79.65 _{0.40}	77.69 _{0.66}	81.84 _{0.00}	74.34 _{0.41}	84.56 _{0.47}
s	3	83.72 _{0.10}	85.64 _{0.39}	82.22 _{0.21}	83.42 _{0.08}	81.94 _{0.48}	78.26 _{0.57}	77.01 _{0.81}	82.10 _{0.10}	72.88 _{0.36}	83.03 _{0.15}
s	5	84.52 _{0.11}	88.09 _{0.34}	81.56 _{0.11}	81.15 _{0.20}	82.97 _{0.19}	77.14 _{0.06}	77.76 _{0.29}	83.14 _{0.27}	73.29 _{0.41}	83.84 _{0.27}
medium-o4-mini											
f	0	48.18 _{0.17}	33.04 _{0.81}	95.43 _{0.52}	96.90 _{0.46}	93.66 _{0.74}	96.51 _{0.17}	55.81 _{0.28}	41.40 _{0.29}	92.69 _{0.72}	95.49 _{0.10}
f	1	40.71 _{0.79}	26.19 _{0.57}	97.49 _{0.41}	98.19 _{0.50}	96.40 _{0.27}	98.68 _{0.44}	46.03 _{0.09}	31.17 _{0.83}	96.57 _{0.61}	97.42 _{0.44}
f	3	39.65 _{0.14}	25.20 _{0.42}	97.74 _{0.21}	98.67 _{0.94}	96.84 _{0.11}	98.40 _{0.24}	42.91 _{0.30}	28.11 _{0.49}	97.27 _{0.58}	98.02 _{0.46}
f	5	38.68 _{0.10}	24.43 _{0.01}	98.07 _{0.02}	98.89 _{0.27}	97.45 _{0.10}	98.00 _{0.20}	41.71 _{0.10}	27.03 _{0.11}	97.62 _{0.10}	98.15 _{0.11}
llama-3B-Instruct											
f	0	35.27 _{0.78}	54.71 _{0.27}	36.72 _{0.21}	44.18 _{0.36}	35.97 _{0.80}	23.35 _{0.79}	58.45 _{0.06}	49.69 _{0.98}	78.76 _{0.41}	84.46 _{0.19}
f	1	33.35 _{0.57}	28.67 _{0.55}	72.58 _{0.99}	75.84 _{0.84}	72.69 _{0.86}	66.53 _{0.43}	50.45 _{0.19}	41.22 _{0.27}	81.70 _{0.83}	86.01 _{0.16}
f	3	32.20 _{0.40}	26.65 _{0.59}	77.77 _{0.82}	80.83 _{0.69}	77.41 _{0.79}	73.35 _{0.29}	49.24 _{0.10}	46.74 _{0.42}	69.55 _{0.68}	76.19 _{0.41}
f	5	35.60 _{0.12}	34.99 _{0.79}	68.37 _{0.82}	73.13 _{0.46}	68.51 _{0.22}	59.15 _{0.18}	48.74 _{0.14}	50.54 _{0.79}	60.98 _{0.34}	67.52 _{0.10}
s	0	61.88 _{0.45}	68.31 _{0.65}	57.49 _{0.19}	56.38 _{0.73}	70.81 _{0.12}	27.84 _{0.32}	67.52 _{0.46}	65.81 _{0.48}	70.20 _{0.88}	77.96 _{0.24}
s	1	32.98 _{0.12}	29.09 _{0.45}	40.70 _{0.36}	43.15 _{0.13}	41.44 _{0.48}	34.47 _{0.08}	48.69 _{0.45}	45.74 _{0.13}	54.89 _{0.09}	57.20 _{0.45}
s	3	38.26 _{0.12}	31.37 _{0.39}	51.78 _{0.11}	53.39 _{0.11}	56.34 _{0.10}	38.24 _{0.11}	53.78 _{0.10}	57.00 _{0.14}	57.00 _{0.14}	54.27 _{0.10}
s	5	39.34 _{0.17}	30.24 _{0.29}	59.38 _{0.11}	62.51 _{0.10}	63.46 _{0.22}	44.34 _{0.10}	60.43 _{0.16}	71.60 _{0.15}	63.22 _{0.10}	60.13 _{0.15}
llama-70B-Instruct											
f	0	45.48 _{0.17}	31.93 _{0.82}	87.43 _{0.28}	88.57 _{0.10}	88.83 _{0.67}	81.12 _{0.27}	67.83 _{0.30}	63.90 _{0.75}	78.48 _{0.18}	81.14 _{0.82}
f	1	37.46 _{0.14}	27.09 _{0.68}	80.38 _{0.28}	81.35 _{0.81}	81.58 _{0.29}	75.18 _{0.80}	62.85 _{0.12}	58.99 _{0.77}	74.55 _{0.19}	78.90 _{0.41}
f	3	30.32 _{0.18}	18.92 _{0.19}	99.94 _{0.60}	100.00 _{0.00}	99.94 _{0.60}	99.85 _{0.01}	58.37 _{0.11}	48.95 _{0.06}	82.83 _{0.13}	84.17 _{0.30}
f	5	30.33 _{0.11}	18.92 _{0.19}	100.00 _{0.00}	100.00 _{0.00}	100.00 _{0.00}	100.00 _{0.00}	53.37 _{0.17}	44.39 _{0.14}	82.87 _{0.13}	83.42 _{0.17}
s	0	68.30 _{0.45}	63.97 _{0.94}	74.22 _{0.34}	80.84 _{0.79}	74.66 _{0.71}	60.62 _{0.11}	76.14 _{0.14}	77.87 _{0.12}	75.10 _{0.13}	73.61 _{0.81}
s	1	68.90 _{0.18}	67.91 _{0.62}	70.65 _{0.19}	75.98 _{0.23}	66.07 _{0.10}	69.04 _{0.12}	68.31 _{0.12}	75.97 _{0.14}	63.25 _{0.10}	61.85 _{0.19}
s	3	65.50 _{0.12}	66.80 _{0.36}	65.06 _{0.17}	71.32 _{0.10}	60.78 _{0.10}	62.12 _{0.11}	63.20 _{0.10}	72.30 _{0.10}	57.72 _{0.17}	53.18 _{0.10}
s	5	66.65 _{0.10}	67.05 _{0.67}	66.98 _{0.11}	74.17 _{0.45}	63.13 _{0.11}	63.23 _{0.11}	65.99 _{0.41}	76.46 _{0.85}	59.26 _{0.10}	54.21 _{0.12}
Qwen3-0.6B											
f	0	18.04 _{0.14}	43.29 _{0.59}	16.46 _{0.14}	22.51 _{0.29}	14.26 _{0.23}	11.56 _{0.47}	10.25 _{0.49}	86.02 _{0.34}	5.91 _{0.31}	10.36 _{0.49}
f	1	22.36 _{0.02}	29.02 _{0.01}	30.24 _{0.96}	35.09 _{0.05}	28.20 _{0.43}	27.18 _{0.81}	10.12 _{0.14}	79.60 _{0.78}	13.39 _{0.53}	16.82 _{0.74}
f	3	21.17 _{0.13}	38.69 _{0.16}	20.34 _{0.12}	24.92 _{0.19}	19.29 _{0.81}	27.88 _{0.71}	14.89 _{0.20}	89.20 _{0.39}	5.99 _{0.74}	8.82 _{0.22}
f	5	19.87 _{0.10}	40.95 _{0.17}	19.67 _{0.14}	24.09 _{0.15}	17.87 _{0.14}	16.06 _{0.42}	8.64 _{0.10}	85.42 _{0.17}	6.66 _{0.10}	10.17 _{0.12}
s	0	43.74 _{0.17}	44.90 _{0.86}	44.05 _{0.35}	60.78 _{0.78}	33.34 _{0.41}	41.94 _{0.15}	36.00 _{0.14}	71.41 _{0.04}	24.79 _{0.10}	39.13 _{0.48}
s	1	51.97 _{0.16}	47.23 _{0.62}	59.72 _{0.29}	70.39 _{0.57}	55.32 _{0.03}	63.93 _{0.20}	35.78 _{0.46}	65.77 _{0.26}	25.10 _{0.20}	33.64 _{0.67}
s	3	48.90 _{0.11}	45.89 _{0.84}	54.00 _{0.11}	62.76 _{0.25}	52.82 _{0.14}	42.36 _{0.84}	31.26 _{0.35}	67.30 _{0.40}	20.75 _{0.80}	18.55 _{0.10}
s	5	48.38 _{0.11}	48.74 _{0.75}	50.29 _{0.45}	58.66 _{0.18}	50.05 _{0.90}	36.34 _{0.73}	34.46 _{0.47}	81.04 _{0.07}	22.31 _{0.50}	23.00 _{0.10}
Qwen3-4B											
f	0	66.39 _{0.58}	75.58 _{0.31}	64.58 _{0.96}	62.47 _{0.52}	70.30 _{0.26}	49.66 _{0.18}	71.04 _{0.07}	69.94 _{0.12}	76.17 _{0.49}	75.24 _{0.17}
f	1	59.89 _{0.18}	79.63 _{0.43}	54.23 _{0.33}	55.44 _{0.35}	59.91 _{0.34}	55.60 _{0.48}	68.80 _{0.10}	74.54 _{0.30}	66.69 _{0.17}	67.52 _{0.37}
f	3	62.29 _{0.17}	80.84 _{0.18}	57.13 _{0.11}	54.85 _{0.10}	61.10 _{0.19}	39.32 _{0.14}	71.77 _{0.10}	76.82 _{0.16}	76.82 _{0.16}	77.84 _{0.10}
s	0	68.48 _{0.17}	67.96 _{0.96}	69.95 _{0.75}	72.43 _{0.78}	62.02 _{0.69}	71.46 _{0.10}	78.32 _{0.62}	66.17 _{0.11}	73.77 _{0.13}	73.77 _{0.13}
s	1	66.66 _{0.18}	80.77 _{0.61}	61.90 _{0.61}	72.82 _{0.60}	54.30 _{0.31}	57.46 _{0.22}	70.19 _{0.10}	74.11 _{0.12}	72.25 _{0.17}	73.73 _{0.81}
s	3	64.43 _{0.10}	80.56 _{0.33}	61.46 _{0.11}	66.70 _{0.10}	60.78 _{0.10}	44.52 _{0.09}	66.78 _{0.17}	74.06 _{0.17}	67.12 _{0.12}	67.81 _{0.81}
s	5	63.37 _{0.12}	83.74 _{0.14}	51.46 _{0.10}	64.45 _{0.10}	44.56 _{0.12}	44.35 _{0.12}	66.64 _{0.12}	73.42 _{0.10}	56.22 _{0.12}	63.24 _{0.46}
Phi-4-mini-reasoning											
f	0	27.65 _{0.02}	27.41 _{0.18}	58.22 _{0.36}	63.15 _{0.14}	54.33 _{0.71}	59.29 _{0.68}	31.72 _{0.17}	20.17 _{0.19}	85.91 _{0.38}	88.66 _{0.35}
f	1	26.03 _{0.12}	25.09 _{0.17}	62.63 _{0.63}	65.80 _{0.07}	55.60 _{0.63}	63.79 _{0.63}	30.84 _{0.24}	20.22 _{0.16}	86.03 _{0.48}	89.89 _{0.86}
f	3	27.00 _{0.05}	33.71 _{0.59}	35.84 _{0.18}	42.49 _{0.42}	32.25 _{0.35}	33.07 _{0.60}	34.79 _{0.31}	23.06 _{0.48}	79.18 _{0.29}	82.80 _{0.13}
f	5	25.64 _{0.34}	32.12 _{0.62}	32.86 _{0.16}	39.64 _{0.38}	28.62 _{0.47}	31.76 _{0.94}	33.42 _{0.12}	24.37 _{0.17}	75.20 _{0.24}	79.71 _{0.81}
s	0	36.91 _{0.51}	25.42 _{0.02}	70.73 _{0.71}	75.57 _{0.07}	72.50 _{0.88}	58.86 _{0.62}	36.93 _{0.71}	26.11 _{0.11}	62.88 _{0.07}	66.11 _{0.69}
s	1	39.14 _{0.38}	28.45 _{0.06}	66.77 _{0.60}	67.84 _{0.03}	71.55 _{0.34}	49.42 _{0.18}	36.84 _{0.24}	24.86 _{0.15}	75.31 _{0.61}	76.75 _{0.96}
s	3	44.57 _{0.19}	34.74 _{0.09}	67.76 _{0.06}	67.76 _{0.06}	69.93 _{0.06}	48.85 _{0.34}	37.15 _{0.23}	26.41 _{0.19}	72.46 _{0.19}	72.46 _{0.19}
s	5	49.37 _{0.12}	42.15 _{0.19}	62.92 _{0.10}	66.16 _{0.10}	67.10 _{0.10}	47.13 _{0.10}	38.44 _{0.10}	25.47 _{0.10}	73.07 _{0.10}	75.58 _{0.10}
MobileLLM-125M											
f	0	33.23 _{0.10}	20.58 _{0.99}	90.18 _{0.11}	91.42 _{0.11}	91.95 _{0.72}	83.76 _{0.02}	34.25 _{0.11}	21.85 _{0.11}	82.74 _{0.07}	82.78 _{0.14}
f	1	33.90 _{0.10}	21.13 _{0.45}	88.38 _{0.08}	91.02 _{0.08}	90.79 _{0.67}	83.77 _{0.14}	35.16 _{0.17}	24.75 _{0.12}	63.83 _{0.27}	68.40 _{0.10}
f	3	34.40 _{0.10}	22.20 _{0.35}	80.11 _{0.18}	84.82 _{0.07}	80.91 _{0.77}	70.45 _{0.13}	36.35 _{0.13}	28.03 _{0.03}	54.38 _{0.77}	60.63 _{0.12}
f	5	34.99 _{0.12}	22.64 _{0.11}	81.49 _{0.27}	83.64 _{0.29}	84.12 _{0.10}	66.45 _{0.12}	37.84 _{0.19}	31.96 _{0.19}	49.07 _{0.26}	57.55 _{0.24}
MobileLLM-600M											
f	0	33.96 _{0.10}	22.03 _{0.99}	77.69 _{0.08}	78.53 _{0.02}	78.33 _{0.42}	71.91 _{0.35}	33.38 _{0.22}	22.40 _{0.16}	68.78 _{0.19}	69.21 _{0.11}
f	1	36.13 _{0.17}	27.10 _{0.36}	80.91 _{0.08}	84.16 _{0.08}	81.41 _{0.67}	75.91 _{0.14}	33.32 _{0.12}	22.81 _{0.16}	63.68 _{0.27}	71.55 _{0.07}
f	3	36.18 _{0.18}	29.21 _{0.49}	48.84 _{0.47}	57.54 _{0.12}	48.52 _{0.44}	40.83 _{0.30}	33.55 _{0.29}	23.91 _{0.19}	59.55 _{0.19}	67.17 _{0.10}
f	5	36.98 _{0.18}	30.83 _{0.83}	49.57 _{0.98}	57.93 _{0.19}	46.94 _{0.36}	36.70 _{0.35}	34.57 _{0.15}	25		

So we make a set
of **concrete
recommendations**
to help these
models perform
better on
disfluent, spoken
data!

3. RESULTS & RECOMMENDATIONS

We analyze the disfluency removal behavior of LLMs and provide recommendations (R1-R9).

Open-Source vs. Proprietary. Looking to Table 3, proprietary models (gpt-4o, gpt-4o-mini) achieve the highest scores, with margins of 10–15 points over the best open-source alternatives. We attribute this to **training** exposure to Whisper-transcribed speech data [24]. **(R1) Proprietary models are currently the most reliable for production systems, while open-source models require targeted augmentation with spoken data.**

Segmentation (s) vs. Full Input (f). Segmenting transcripts consistently improves both mean performance and stability, e.g., gpt-4o improves from $\mathcal{E}_F=76.13$ (f) to 82.38 (s) at $k=1$. This supports prior **evidence** of long-context degradation in LLMs [25, 26]. **(R2) Segmentation is an effective preprocessing step that should be applied.**

Few-Shot Sensitivity (k). Increasing k does not uniformly improve results. Small models (e.g., MobileLLM) gain slightly, but others show degradation (e.g. Llama-3B/8B/70B) when more examples are provided. **(R3) Few-shot prompting should be used with caution, as some model families misinterpret exemplars and over-edit fluent text.**

Disfluency Category Performance. $\mathcal{Z}$ -Scores show that EDITED nodes are handled well, but INTJ and PRN nodes are frequently missed, despite **prior** work suggesting these are the easiest to detect [19, 17]. **(R4) Future modeling should focus on under-served categories (INTJ, PRN) to improve robustness across all disfluency types.**

Over-Deletion Failures. Several models (e.g., Llama-8B, o4-mini) achieve near perfect recall but at the cost of very low precision, deleting fluent tokens. Segmentation often mitigates this collapse mode. **(R5) Segment-level evaluation**

helps reduce over-deletion risk.

Under-Deletion Failures. Some models (e.g., Qwen series) exhibit the opposite trend of over-deletion, achieving high precision but low recall (purple). These models preserve most fluent tokens but fail to remove many true disfluencies, especially in INTJ and PRN categories. This reflects conservative editing strategies and **limited** exposure to conversational disfluency distributions. **(R6) Models prone to under-deletion require additional filtering or targeted fine-tuning to ensure sufficient disfluency coverage.**

Reasoning-Oriented Models. Models tuned for reasoning (o4-mini, Phi-4) perform **poorly**, showing high recall but extreme over-deletion (blue). **(R7) Reasoning capability does not translate to disfluency removal; specialized evaluation remains necessary.**

Impact of Model Size. Model scaling generally improves disfluency removal, with Qwen, GPT, and Llama families showing upward trends. However, gains are nonlinear – e.g., Qwen3-1.7B underperforms both smaller and larger variants – likely due to training **data** or optimization differences rather than capacity limits. **(R8) Model choice should be guided by empirical benchmarks on target domains and disfluency categories rather than size alone.**

Fine-Tuning and Generalization. Looking to Tables 4 and 5, fine-tuning improves performance to near SOTA levels (e.g., gpt-4o-mini_{f_t} achieves $\mathcal{E}_P=96.6$), but evaluation on GSM8K, MMLU, and CoQA shows degraded performance on unrelated tasks. **(R9) Fine-tuning is suitable for dedicated disfluency pipelines, but not for general-purpose conversational models.**

Ok, so that's LLMs – what about ASR systems?

Comparing ASR Systems in the Context of Speech Disfluencies

Maria Teleki¹, Xiangjue Dong¹, Soohwan Kim¹, James Caverlee¹

¹Texas A&M University, USA

{mariateleki, xj.dong, cocomox26, caverlee}@tamu.edu

INTERSPEECH '24

→ **TLDR: Whisper!**

Last,

**Let's talk about
speaker variation!**

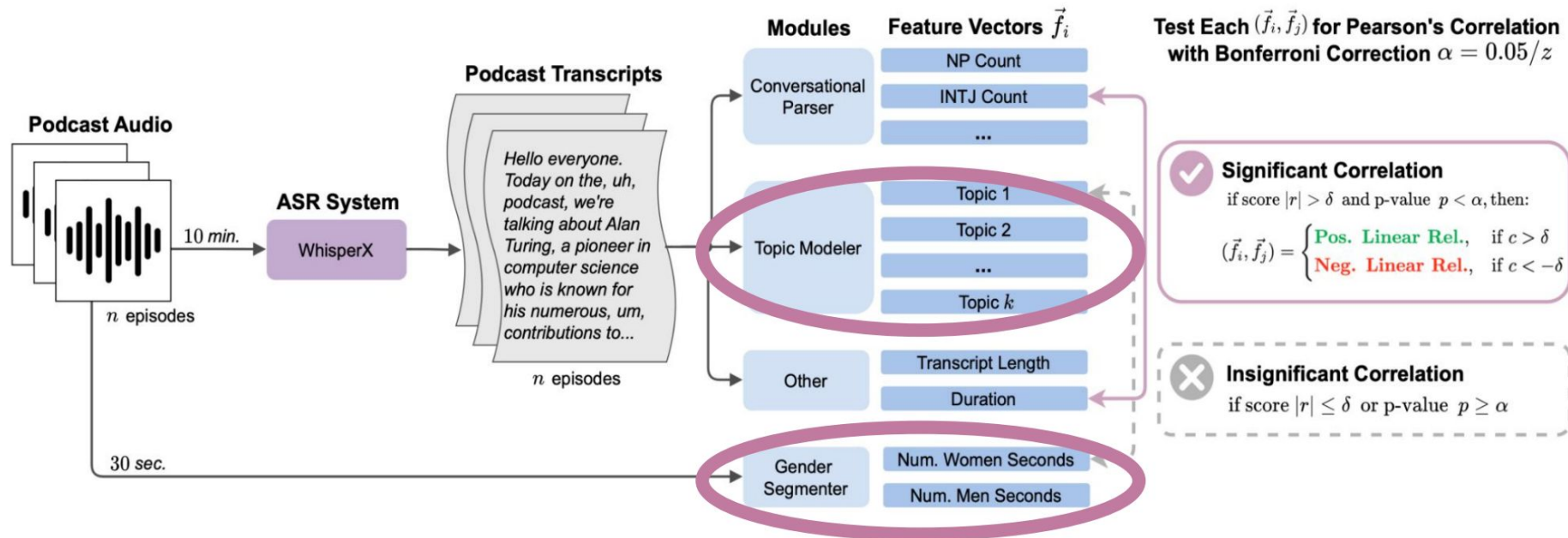
Masculine Defaults via Gendered Discourse in Podcasts and Large Language Models

Maria Teleki, Xiangjue Dong, Haoran Liu, James Caverlee

Texas A&M University
{mariateleki, xj.dong, liuhr99, caverlee}@tamu.edu

ICWSM '25 + IC2S2 '25, SiCon@ACL '25

We propose the **Gendered* Discourse Correlation Framework (GDCF)** to monitor & identify discourse terms at scale.



*** we use binary gender due to speech tool limitations - this is an important direction for future work.**

We find that women & men's speaking patterns are **different**.

Topic N	Gender	<i>r</i>	Topic N Word List	Topic N Categories	Topic N Gender
Topic 3	Women	0.15	women, woman, men, baby, pregnant, girls, male, doctor, health, birth	Content - Pregnancy	Women
	Men	-0.14			
Topic 10	Women	0.10	energy, body, feel, mind, space, yoga, love, beautiful, feeling, meditation	Content - Yoga	Women
	Men	-0.12			
Topic 49	Women	-0.21	game, know, think, team, going, mean, play, year, one, good	Content - Sports	Men
	Men	0.17			
Topic 71	Women	0.14	christmas, sex, girl, hair, love, get, date, girls, let, wear	Content - Dating	Women
	Men	-0.14			
Topic 54	Women	–	get, like, know, right, people, going, podcast, make, want, one	Discourse	Men
	Men	0.12			
Topic 60	Women	-0.27	going, know, think, get, got, one, really, good, well, yeah	Discourse	Men
	Men	0.20			
Topic 62	Women	0.33	like, know, really, going, people, want, think, get, things, life	Discourse	Women
	Men	-0.28			

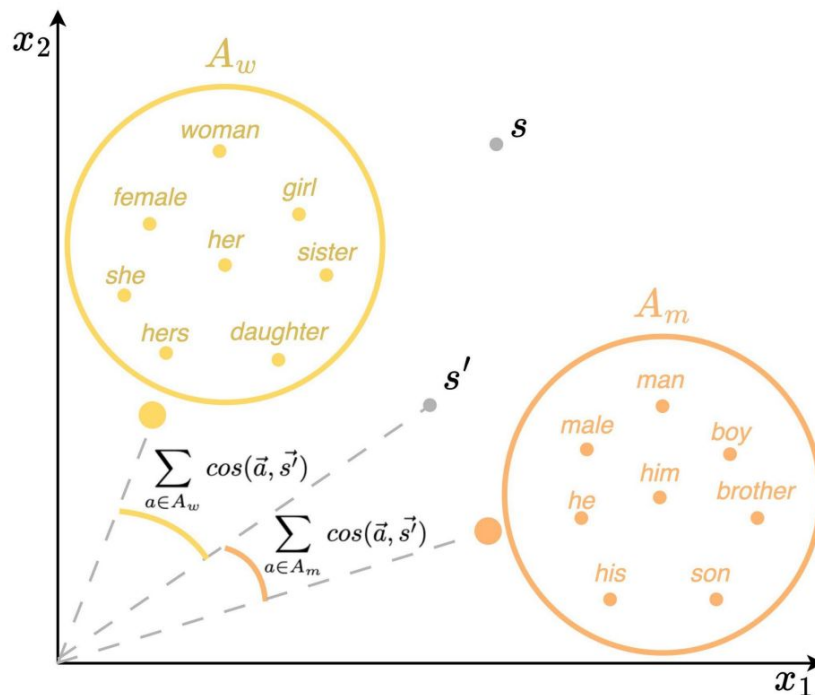
Men:

*s = And I was **going**, hey, it's cold outside...*

Women:

*s' = And I was **like**, hey, it's cold outside...*

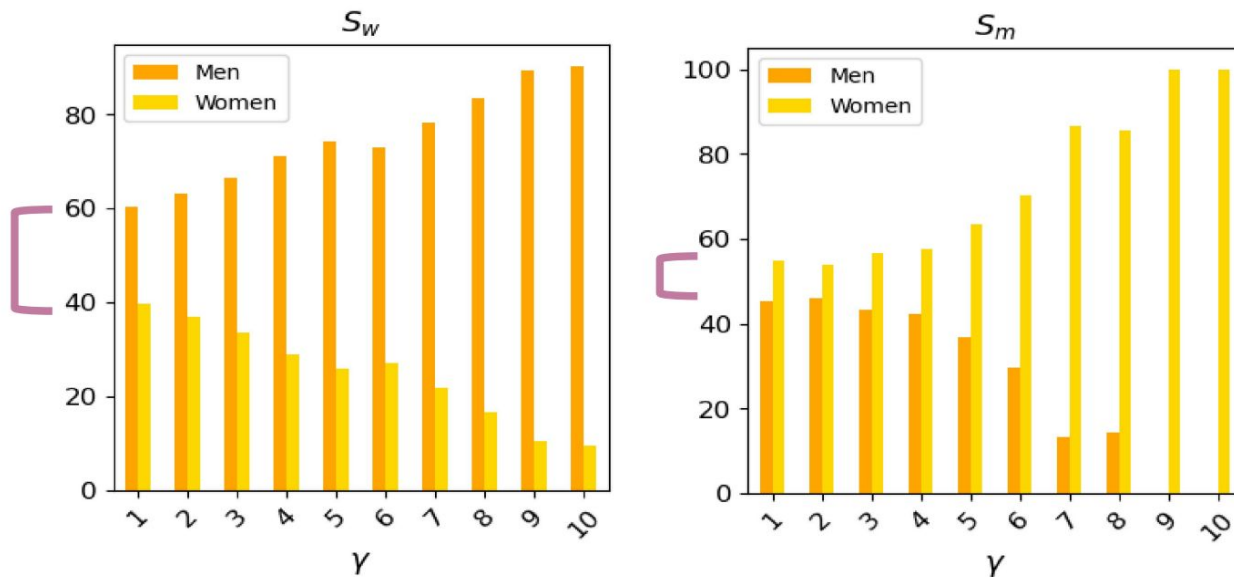
Is this difference reflected in LLM embeddings?



Men: $s = \text{And I was } \text{going}, \text{ hey, it's cold outside...}$

Women: $s' = \text{And I was } \text{like}, \text{ hey, it's cold outside...}$

Yes. Women have a **less stable/robust** embedding representation than men.



Men:

$s = \text{And I was going, hey, it's cold outside...}$

Women:

$s' = \text{And I was like, hey, it's cold outside...}$

The **use** of these masculine discourse terms is associated with **economic rewards**.

Topic N	Topic M	r	Topic N Word List	Topic N Categories	Topic M Word List	Topic M Categories
Topic 11	Topic 54	0.11	data, new, technology, public, bill, theory, science, system, security, article	Content - Technology/Political \$\$\$	get, like, know, right, people, going, podcast, make, want, one	Discourse (Men)
	Topic 62	-0.20			like, know, really, going, people, want, think, get, things, life	Discourse (Women)
Topic 12	Topic 54	0.24	business, money, company, market, buy, right, million, companies, pay, sell	Content - Business \$\$\$	get, like, know, right, people, going, podcast, make, want, one	Discourse (Men)
Topic 79	Topic 60	0.18	game, games, play, playing, like, played, nintendo, video, fun, switch	Content - Video Games	going, know, think, get, got, one, really, good, well, yeah	Discourse (Men)
	Topic 62	-0.13			like, know, really, going, people, want, think, get, things, life	Discourse (Women)

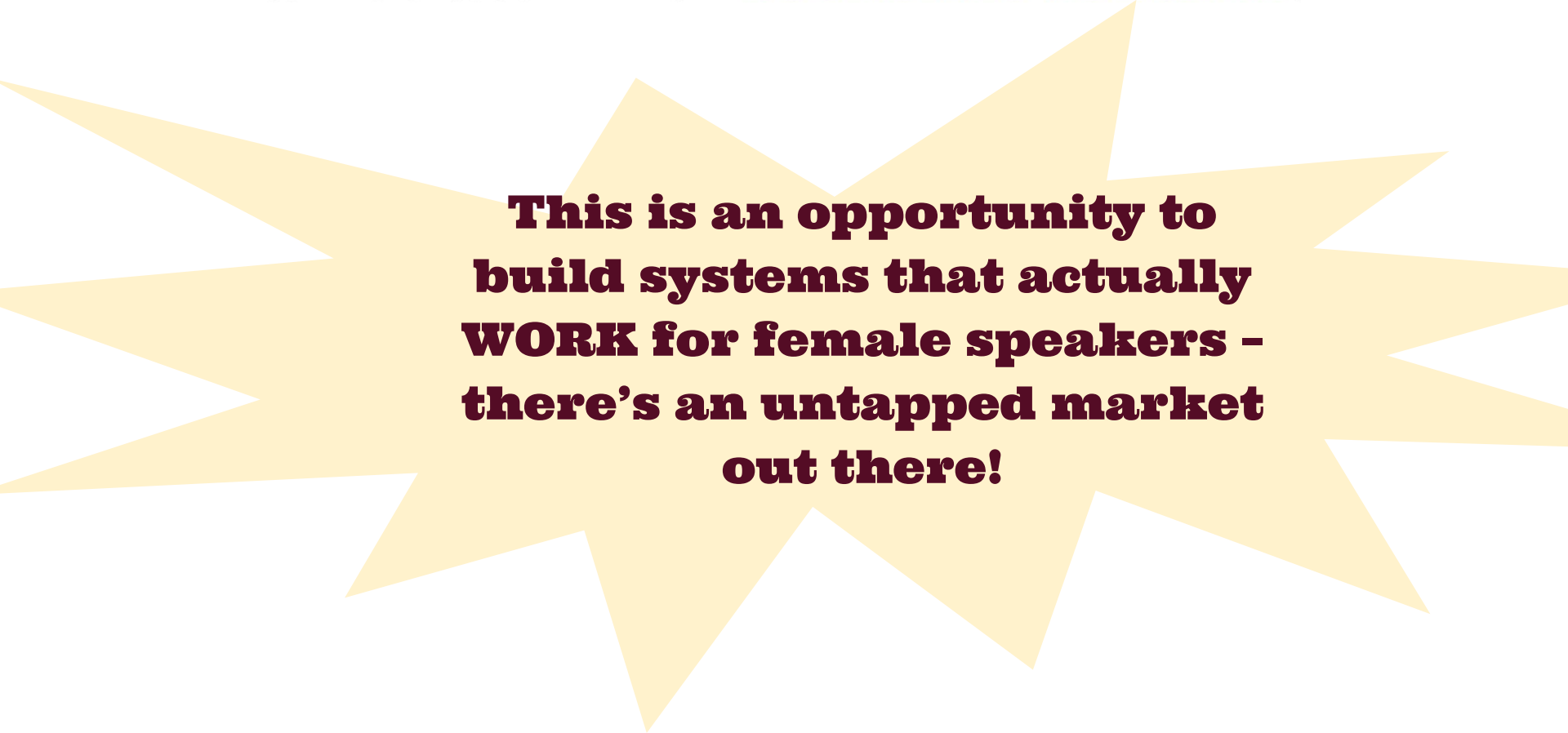
Men:

$s = \text{And I was } \textbf{going}, \text{ hey, it's cold outside...}$

Women:

$s' = \text{And I was } \textbf{like}, \text{ hey, it's cold outside...}$

The **use of these masculine discourse terms is associated with **economic rewards**.**



**This is an opportunity to
build systems that actually
WORK for female speakers -
there's an untapped market
out there!**

A quick plug for our recent work →



Maria Teleki ✓ • You

CS PhD Student @ TAMU | Speech AI/RecSys

1d • 🔄

...

In our new work, 🎵 CHOIR: Collaborative Harmonization fOr Inference Robustness, we show that different LLM personas often get different benchmark questions right! CHOIR leverages this diversity to boost performance across benchmarks. 🇮🇹

You always hear about the "bias-accuracy tradeoff," meaning that ⬇️ model bias, ⬇️ system performance, so ⬇️ \$\$\$ for a company. So much of the conversation around bias and diversity has focused on how to incentivize companies to debias their models (e.g., through new legislation).

🌟 To me, this work is super exciting because we take a totally different perspective: we show that ⬆️ diverse perspectives, ⬆️ system performance, so ⬆️ \$\$\$ for a company! With this work, we argue that <<< 🇺🇸 diverse perspectives are absolutely necessary >>> from an economic standpoint.

📄 https://lnkd.in/gr_Mmvsv

👤 [Xiangjue Dong](#) (1st author), [Cong \("Nicole"\) Wang](#), [Millennium Bismay](#), and [James Caverlee](#)

We've also got some upcoming work on social-media-as-a-signal, stay tuned!



Collaborators



James Caverlee
(my advisor)



Xiangjue Dong



Haoran Liu



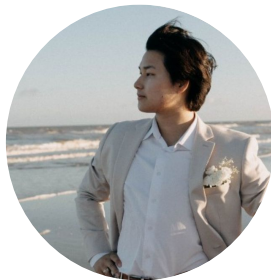
Sai Tejas
Janjur



Thomas
Docog



Oliver Grabner



Soohwan Kim



Stephanie
Birkelbach



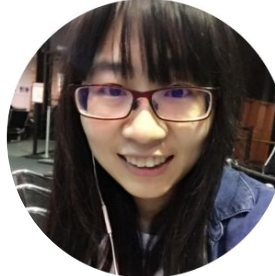
Rohan
Chaudhury



Ketan Verma



Chengkai Liu



Yin Zhang

+ Jason Kim, Lingfeng Shi, Cong Wang, and shoutout to work-in-progress collaborators too :)

Conversational AI

MARIA TELEKI

Will post
slides after
the talk!



Texas A&M University
Department of Computer Science & Engineering

