

Advanced Voice Cloning

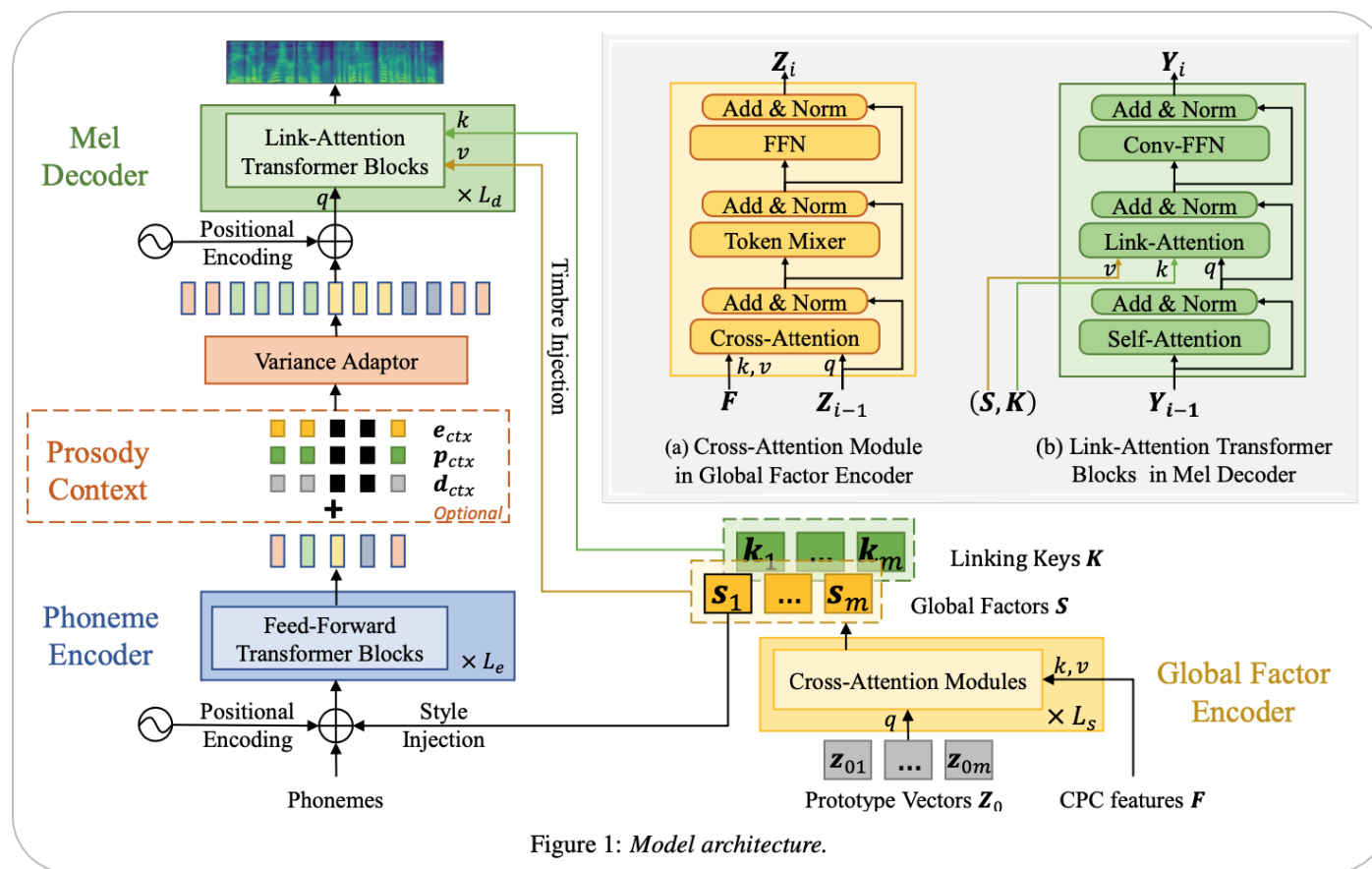
2022-11-13

候选方案汇总

- Zero-Shot 音色克隆
 - Global and Local Factors: RetrieverTTS
 - Fine-Grained Speaker Embedding
- 基于 Adapter 的音色克隆
 - 基于 Adapter 的 TTS 工作
 - 思考: Conditional Adapter
- SOTA: Scaling & Shifting Features

借鉴网络结构，应用时与Few-Shot 微调模型相结合，从而达到更优的效果

Zero-Shot 音色克隆: RetrieverTTS



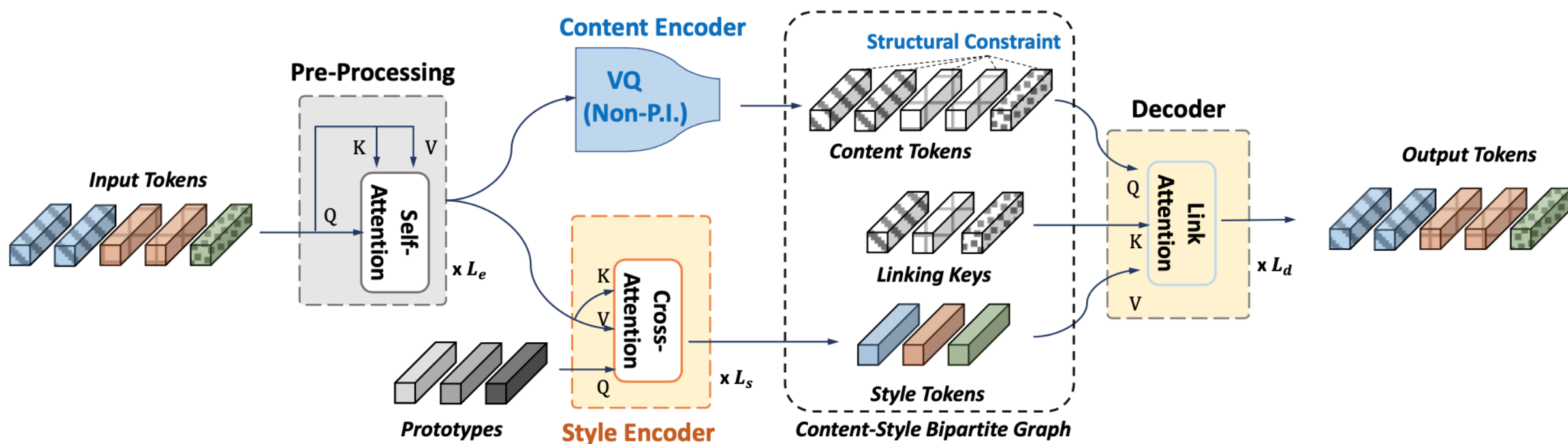
requires whole sentence generation. We also find that our system's SMOS score largely outperforms the state-of-the-art zero-shot voice adaptive TTS method, which leverages an explicit speaker-related meta-learning objective. This indicates that our system is powerful enough to handle the personalized TTS task, even though our design does not fully aim at this goal.

	MOS@long	SMOS
Ground-Truth	4.20 ± 0.14	4.45 ± 0.19
C.Tang et al. [4]	2.73 ± 0.25	—
Meta-StyleSpeech [12]	3.88 ± 0.20	2.91 ± 0.31
RetrieverTTS (Ours)	3.95 ± 0.17	3.70 ± 0.25

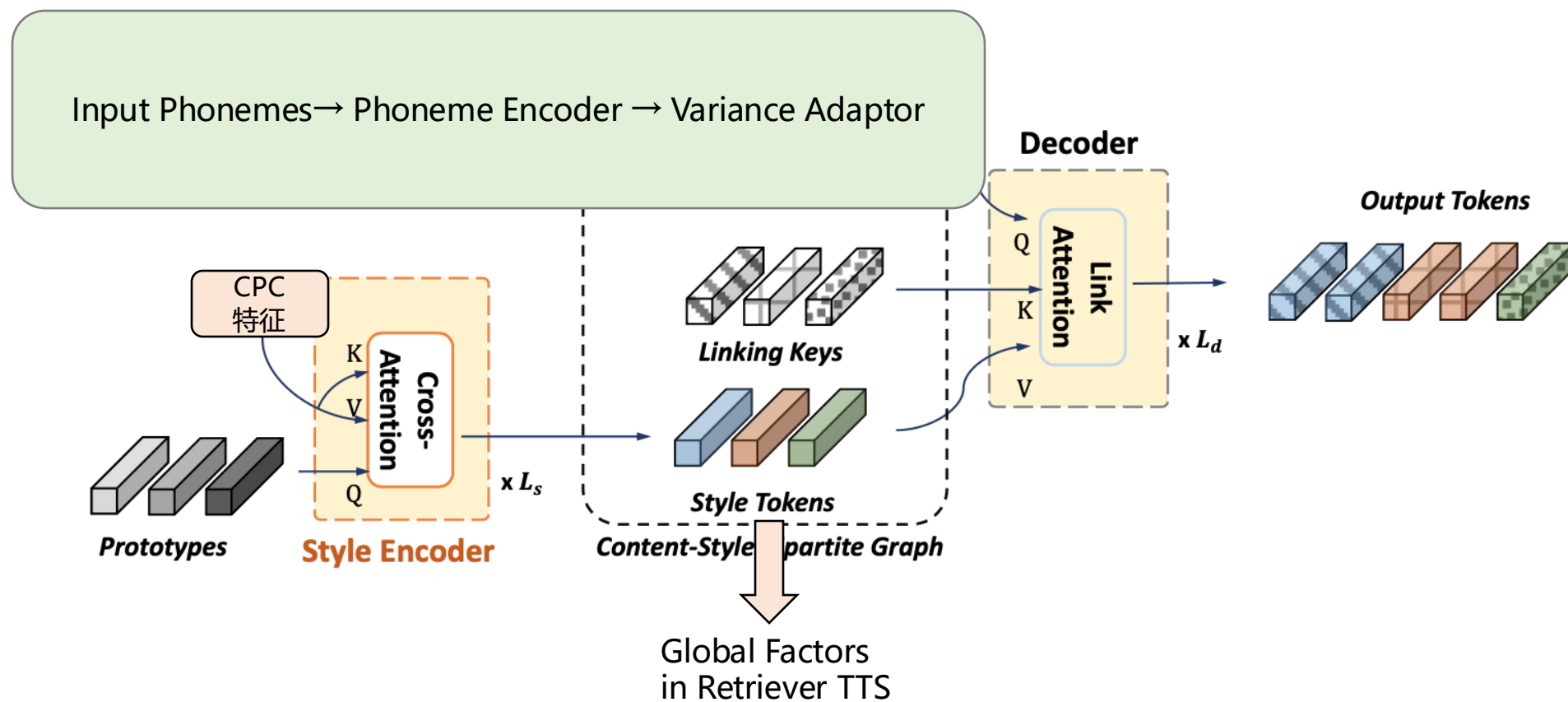
RetrieverTTS 用于音色克隆

- 复现模块
 - CPC 特征
 - Cross Attention
 - Token Mixer ?
 - Link Attention
- 去除模块 (暂时不实现语音编辑)
 - Prosody Context

Zero-Shot 音色克隆: RetrieverTTS 与 Retriever

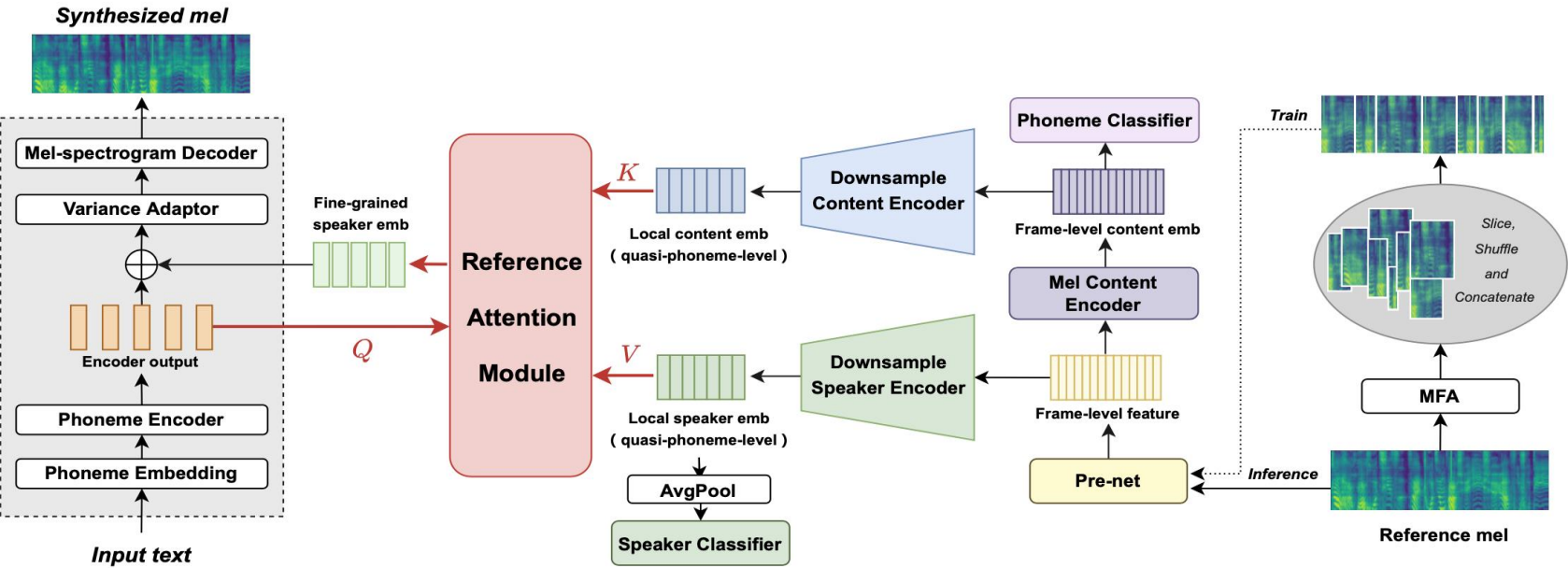


Zero-Shot 音色克隆: RetrieverTTS 与 Retriever

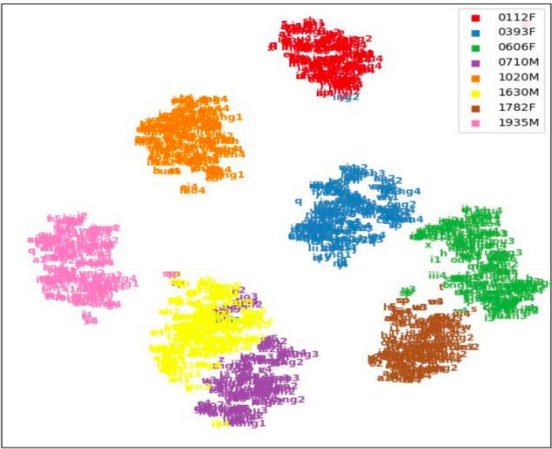


Zero-Shot 音色克隆

Content-Dependent Fine-Grained Speaker Embedding (CDFSE)

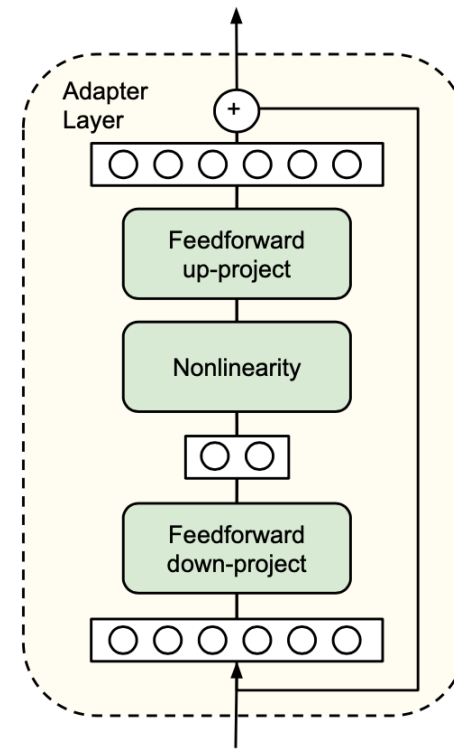
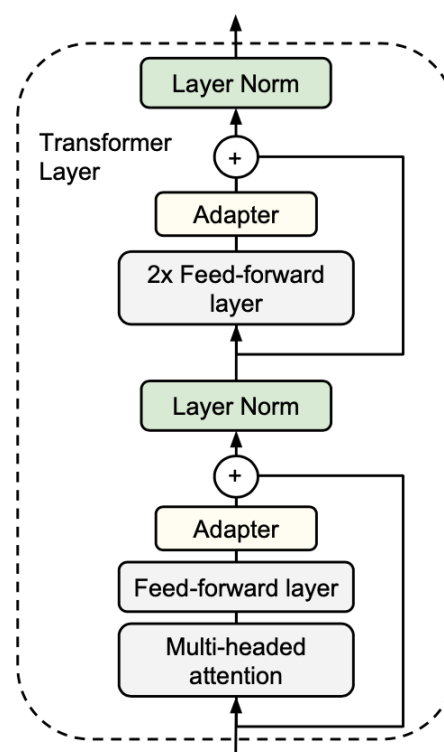
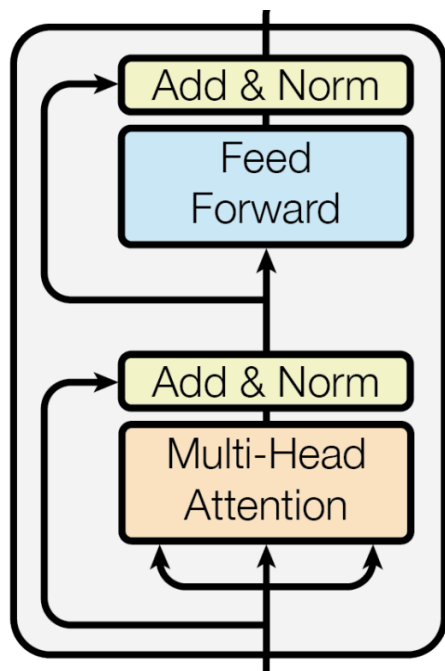


Metric	Model	seen speakers	unseen speakers
MOS	GSE	3.50 ± 0.16	3.56 ± 0.12
	CLS	3.51 ± 0.14	3.53 ± 0.11
	Attentron*	3.63 ± 0.16	3.57 ± 0.13
	CDFSE	3.59 ± 0.17	3.54 ± 0.12
SMOS	GSE	3.89 ± 0.14	3.08 ± 0.14
	CLS	3.79 ± 0.16	3.12 ± 0.14
	Attentron*	4.04 ± 0.17	3.29 ± 0.13
	CDFSE	4.11 ± 0.15	3.51 ± 0.14



基于 Adapter 的方法

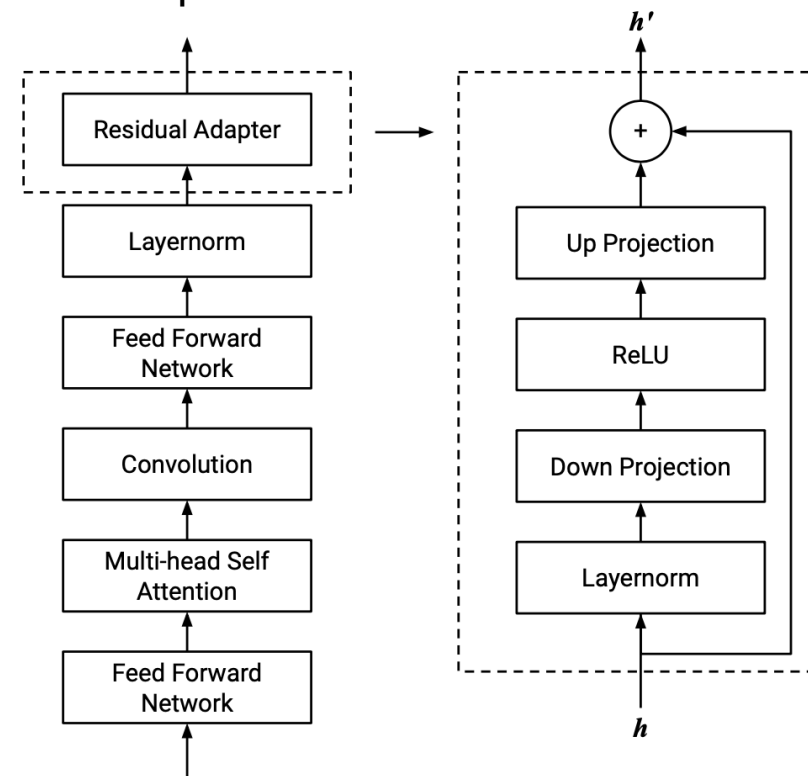
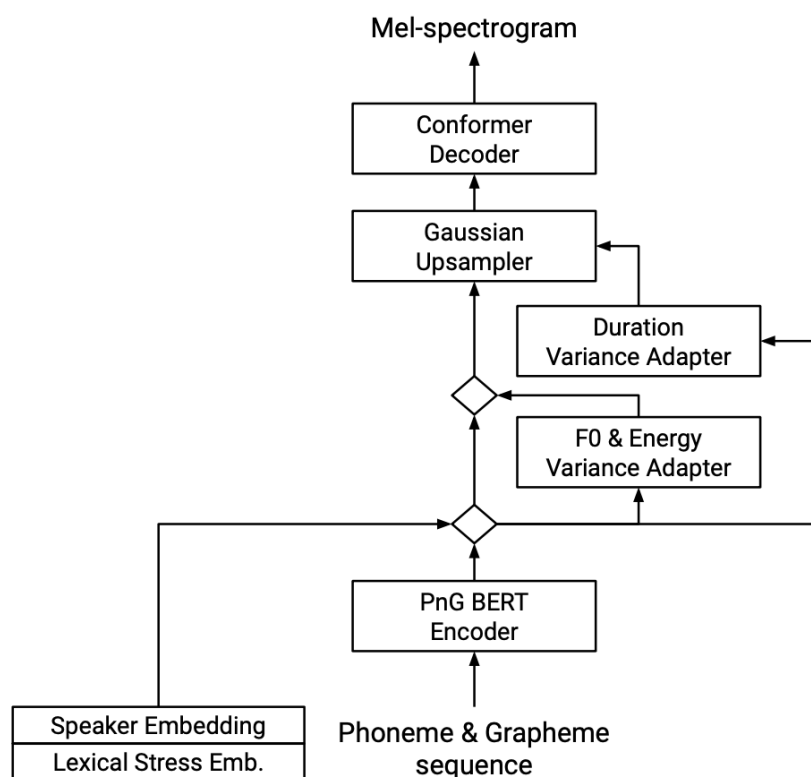
Adapter 的提出: Parameter-Efficient Transfer Learning for NLP



- finetune 时只更新 Adapter + LayerNorm 部分的参数

基于 Adapter 的方法

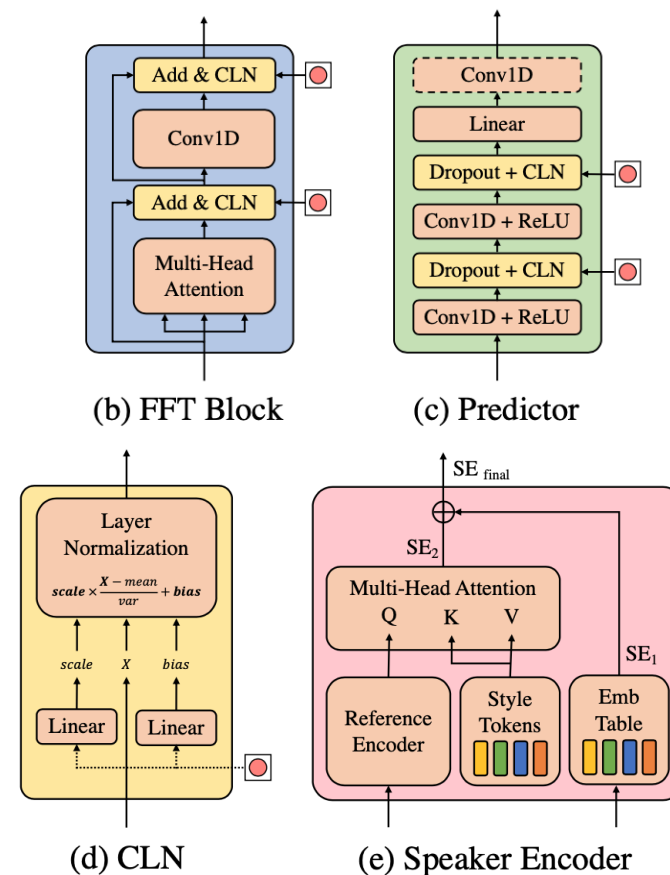
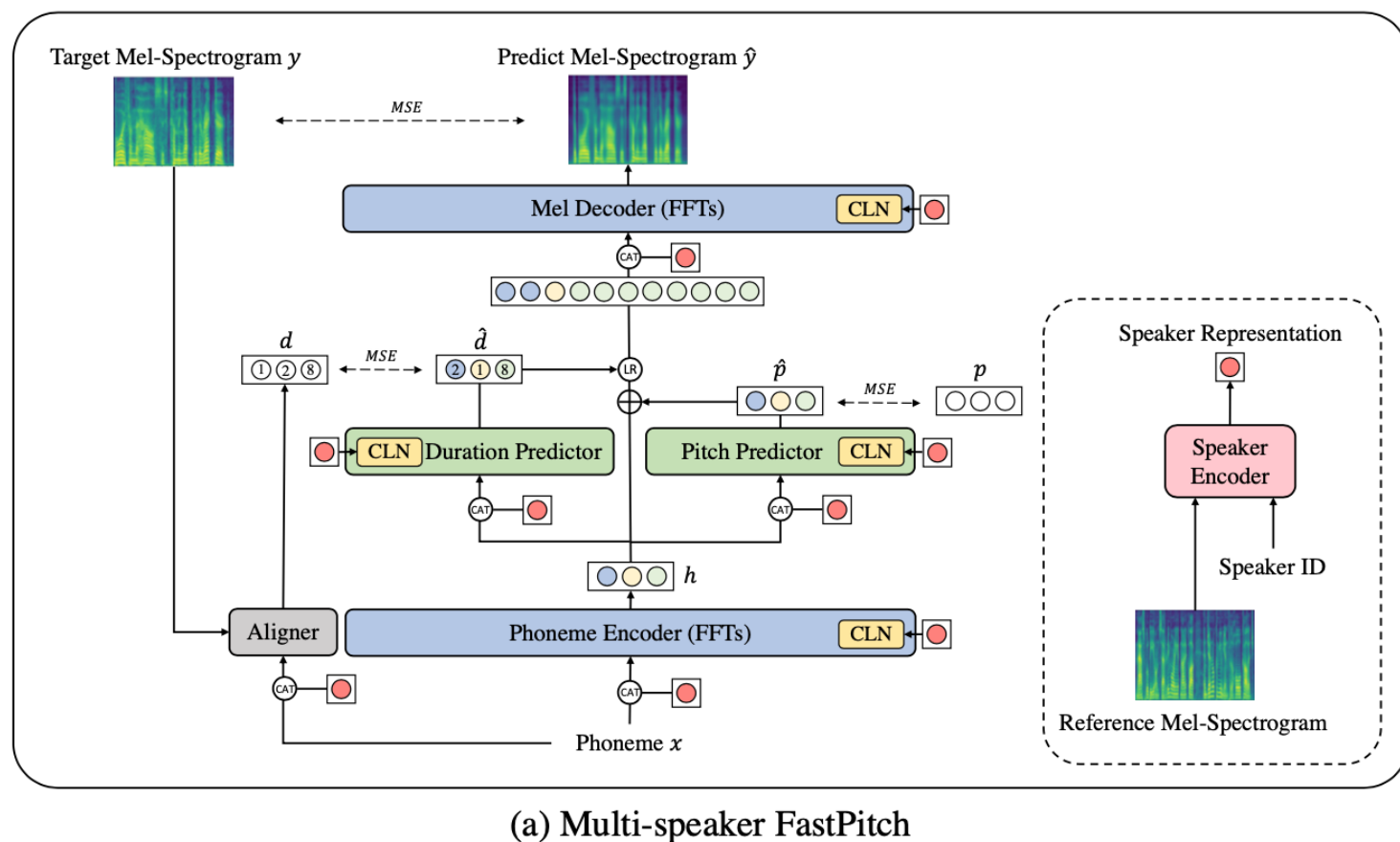
论文 1: Residual Adapters for Few-Shot Text-to-Speech Speaker Adaptation



$$h' = h + \text{ReLU}(\text{LayerNorm}(h) W_{\text{down}}) W_{\text{up}}$$

基于 Adapter 的方法

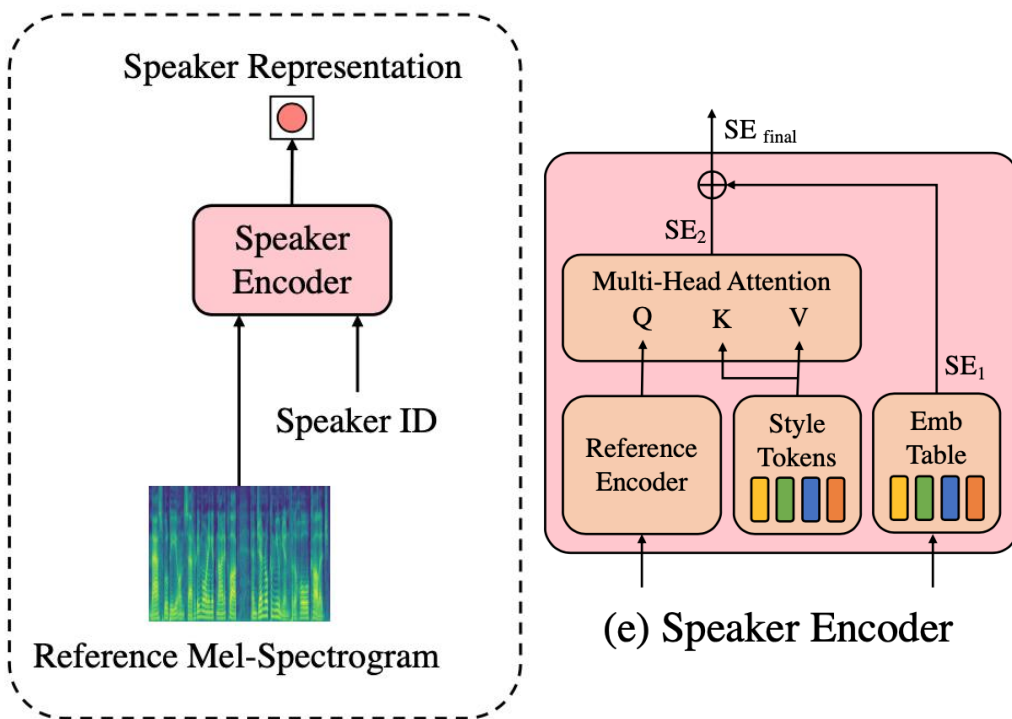
论文 2: Adapter-Based Extension of Multi-Speaker Text-to-Speech Model for New Speakers



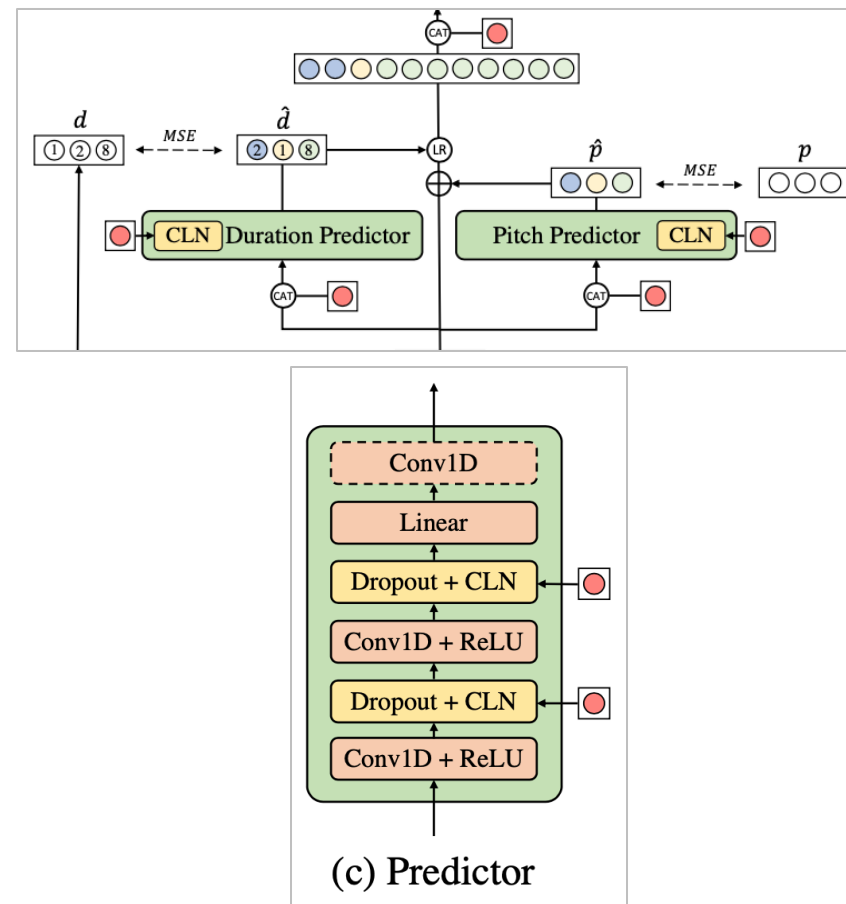
基于 Adapter 的方法

论文 2: Adapter-Based Extension of Multi-Speaker Text-to-Speech Model for New Speakers

借鉴点一



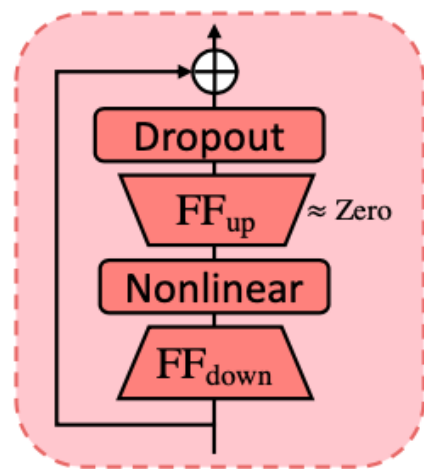
借鉴点二



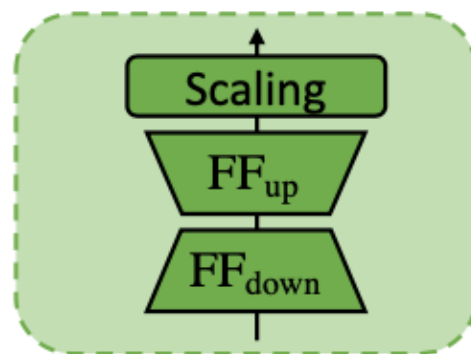
基于 Adapter 的方法

论文 2: Adapter-Based Extension of Multi-Speaker Text-to-Speech Model for New Speakers

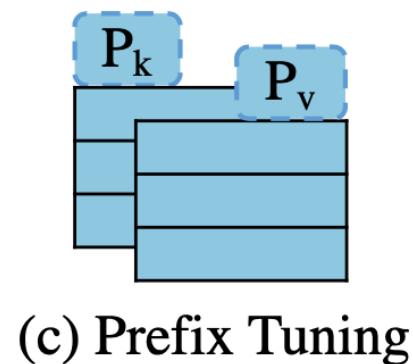
- 统一框架下的 Adapter 及变体



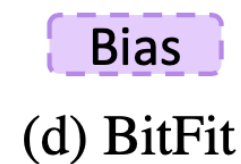
(a) Adapter



(b) LoRA



(c) Prefix Tuning

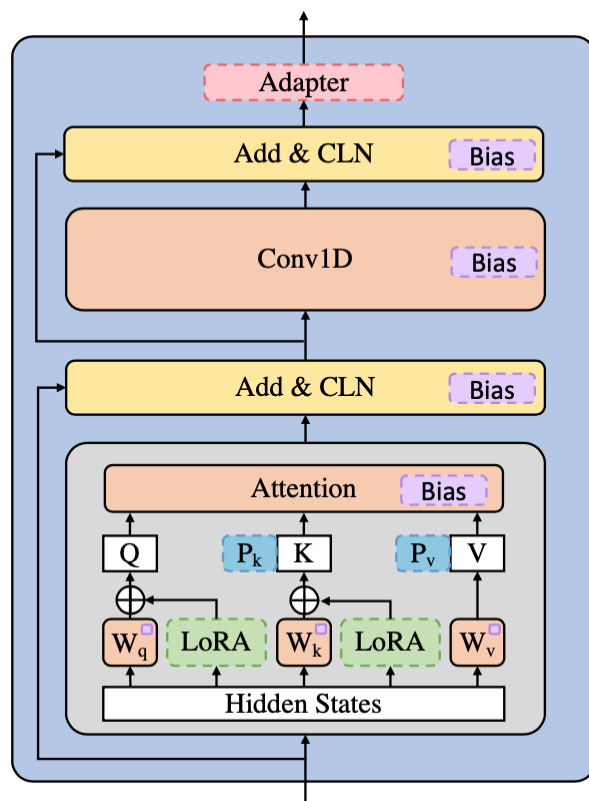


(d) BitFit

基于 Adapter 的方法

论文 2: Adapter-Based Extension of Multi-Speaker Text-to-Speech Model for New Speakers

- 统一框架下的 Adapter 及变体



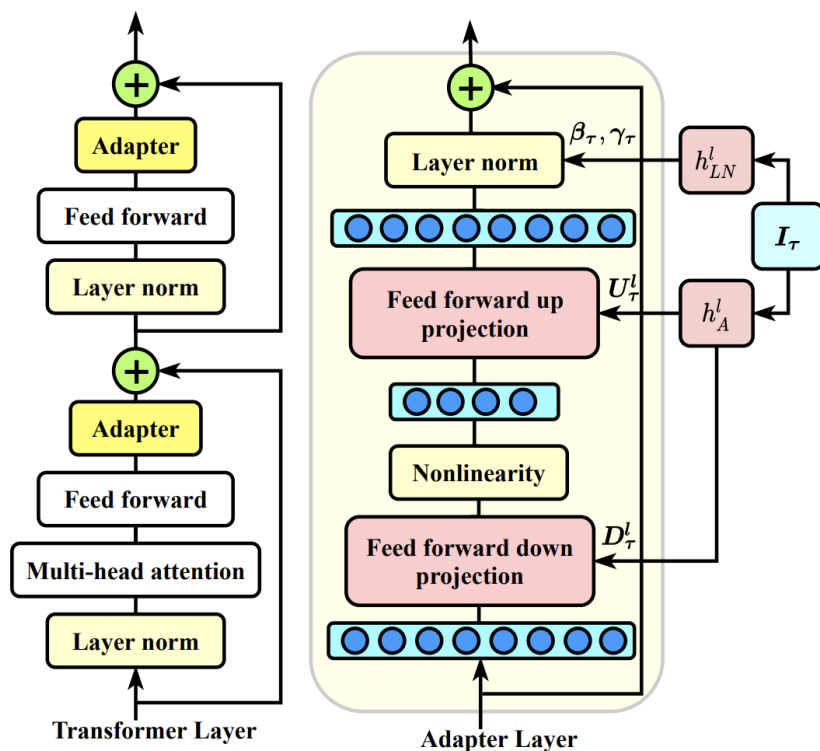
Method	SECS $\uparrow$	CFSD $\downarrow$	MSE _P $\downarrow$	MSE _D $\downarrow$	Params
<i>Different parameter-efficient methods</i>					
BitFit	0.452	56.4	71.0	19.1	2.2M
PrefixTuning (FFTs)	0.067	83.2	75.6	22.7	0.6M
LoRA (FFTs)	0.141	62.7	77.7	22.3	2.8M
Adapter (FFTs)	0.568	30.0	65.9	21.3	2.4M
<i>Different speaker-related modules</i>					
Adapter (FFTs/Predictors/Aligner)	0.575	28.3	62.7	16.9	3.5M
+ speaker embedding	0.586	27.2	63.6	16.2	3.5M
+ speaker embedding + CLN	0.540	46.9	66.8	17.2	7.8M
speaker embedding + CLN [13]	0.513	53.5	68.0	21.7	4.3M
Full fine-tuning	0.604	31.0	73.8	19.7	53.4M

Method	MOS $\uparrow$			SMOS $\uparrow$		
	LibriTTS	VCTK	HiFi-TTS	LibriTTS	VCTK	HiFi-TTS
Recording	4.16 $\pm$ 0.05	4.08 $\pm$ 0.05	4.12 $\pm$ 0.05	3.90 $\pm$ 0.07	3.80 $\pm$ 0.07	3.71 $\pm$ 0.08
Adapter	4.10 $\pm$ 0.05	3.89 $\pm$ 0.06	3.87 $\pm$ 0.06	3.44 $\pm$ 0.08	3.28 $\pm$ 0.08	3.27 $\pm$ 0.08
Full fine-tuning	3.96 $\pm$ 0.06	3.88 $\pm$ 0.06	3.75 $\pm$ 0.07	3.40 $\pm$ 0.08	3.33 $\pm$ 0.06	3.34 $\pm$ 0.08

基于 Adapter 的方法

Conditional Adapter

参考论文: Parameter-efficient Multi-task Fine-tuning for Transformers via **Shared Hypernetworks**



$$A_{\tau}^l(x) = LN_{\tau}^l \left(U_{\tau}^l (\text{GeLU}(D_{\tau}^l(x))) \right) + x$$

$$(U_{\tau}^l, D_{\tau}^l) := h_A^l(I_{\tau}) = (W^{U^l}, W^{D^l}) I_{\tau}$$

$$W^{U^l} \in \mathbb{R}^{(d \times h) \times t}$$

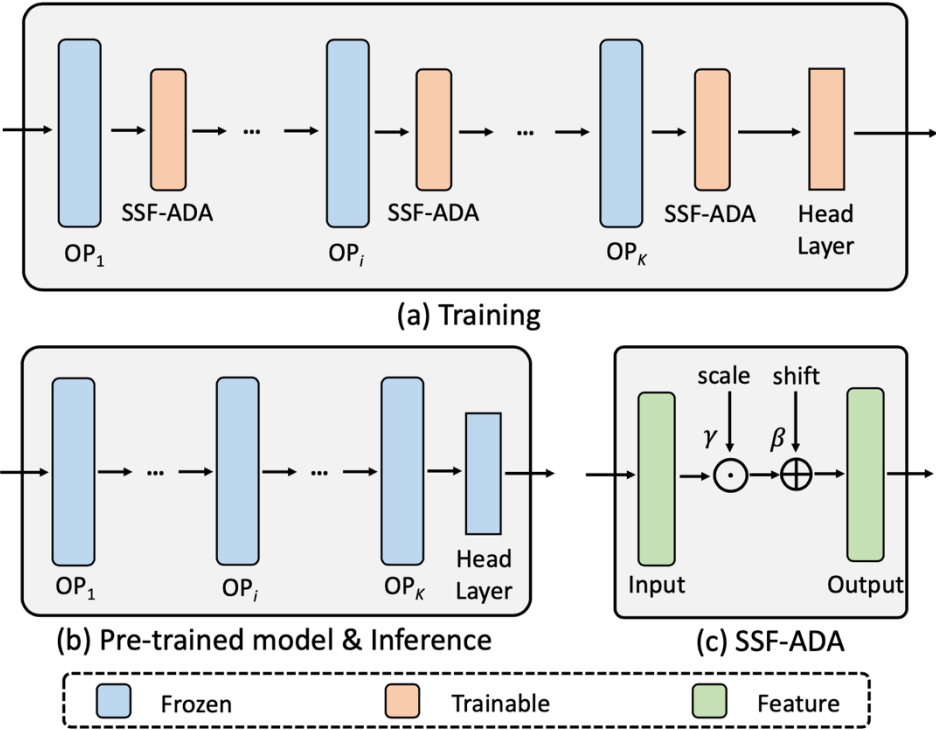
$$W^{D^l} \in \mathbb{R}^{(h \times d) \times t}$$

$$(\gamma_{\tau}^l, \beta_{\tau}^l) := h_{LN}^l(I_{\tau}) = (W^{\gamma^l}, W^{\beta^l}) I_{\tau}$$

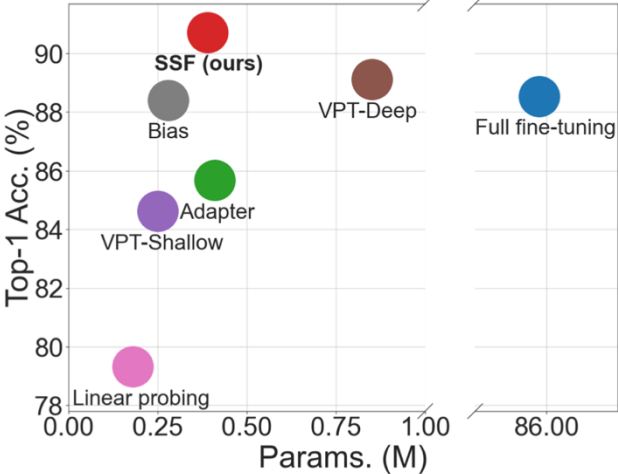
$$W^{\gamma^l} \in \mathbb{R}^{h \times t}$$

$$W^{\beta^l} \in \mathbb{R}^{h \times t}$$

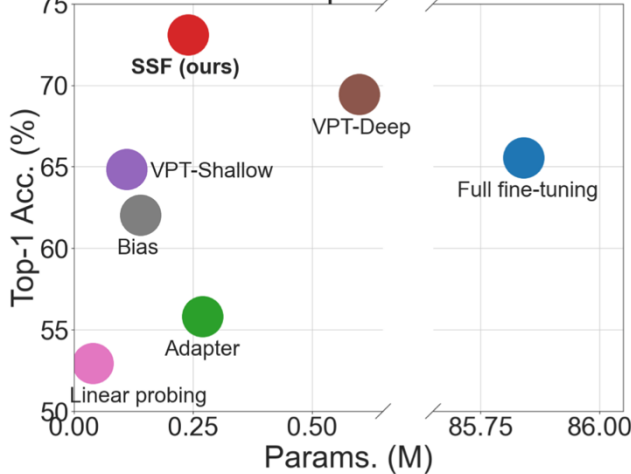
SOTA: Scaling & Shifting Features (SSF)



Fine-Grained Visual Classification Datasets



Visual Task Adaptation Benchmark



Method \ Dataset	CUB-200 -2011	NABirds	Oxford Flowers	Stanford Dogs	Stanford Cars	Mean	Params. (M)
Full fine-tuning	87.3	82.7	98.8	89.4	84.5	88.54	85.98
Linear probing	85.3	75.9	97.9	86.2	51.3	79.32	0.18
Adapter [36]	87.1	84.3	98.5	89.8	68.6	85.67	0.41
Bias [90]	88.4	84.2	98.8	91.2	79.4	88.41	0.28
VPT-Shallow [44]	86.7	78.8	98.4	90.7	68.7	84.62	0.25
VPT-Deep [44]	88.5	84.2	99.0	90.2	83.6	89.11	0.85
SSF (ours)	89.5	85.7	99.6	89.6	89.2	90.72	0.39

方案思考与总结

候选方案

- 按照优先级顺序排列
 - ❑ RetrieverTTS
 - ❑ (Conditional) Adapter-Based Method
 - ❑ Scaling & Shifting Features (SSF)
 - ❑ Content-Dependent Fine-Grained Speaker Embedding (CDFSE)