

超越硬共享: 使用专家混合监督的高效多任务语音文本建模

Hojun Jin*, Eunsoo Hong*, Ziwon Hyung, Sungjun Lim, Seungjin Lee, Keunseok Cho

Samsung Research, Korea

hojun.jin@samsung.com, eunsoo1.hong@samsung.com, ziwon.hyung@samsung.com, sungjun.lim@samsung.com, sjsr.lee@samsung.com, ks1.cho@samsung.com

Abstract

硬共享是在不同任务中共同使用模型的一个常用策略。然而,它通常导致任务干扰,妨碍模型性能。为了解决这个问题,我们提出了一种简单而有效的专家混合(S-MoE)。该专家混合模型不同,S-MoE通过利用特殊指定每个任务路由到其指定专家,从而消除了控制函数的需要。通过为每个任务分配一个专家的前缀,S-MoE克服了硬共享的限制。我们进一步将S-MoE用于语音到文本模型,使模型能够处理混合输入,同时结合自行语音识别(ASR)和语音翻译(ST)。实验结果证明了所提出的S-MoE的有效性,它同时用于识别和解码器,在字错误率(WER)方面取得了6.35%的提升。

Index Terms: automatic speech recognition, speech translation, multi-task learning, supervised mixture of experts

1. 介绍

语音文本(STT)模型通常在宽带(WB)语音上行,因为它能提供更丰富的语音特征。然而,窄带(NB)语音在移动通信中仍然是必不可少的,特别是在使用方面。由于采样率差异导致的带宽差异,NB和WB模型通常是分开的。此外,许多语音文本应用需要自行语音识别(ASR)和语音翻译(ST)。从部署的角度来看,由于移动和嵌入式设备上的资源限制,许多用于每个任务的模型是不切实际的。

多任务学习(MTL)是解决这一挑战的一个有前途的解决方案,它使模型能够在不同的输入和任务之间共享表示[1, 2, 3]。然而,使用硬共享的MTL方法通常遭遇任务干扰,其中一任务的性能降低其他任务的性能[4]。

为了解决这个问题,已经提出了多种方法,其中之一有效的方法是使用专家混合(MoE) [5] 的模块化架构。MoE因其能够根据不同任务分配不同的专家而受到关注。然而,MoE模型严重依赖于路由到特定专家的控制函数。

在本文中,我们提出了一种简单而有效的专家混合(S-MoE)的新方法,用于高效的多任务STT模型。我们的方法借鉴了MoE原理,但输入和输出组件均已明确的场景进行了定制。我们专注于简化的多组件任务,在特定情况下,监督路由提供了一种简单而有效的解决方案。与标准的MoE模型不同,我们的方法通过引入特殊的指定路由式地引入到路由的专家,从而消除了控制函数的需要。我们的方法简化了路由机制,同时保持模型在识别和推理过程中不增加额外的计算量,保持了效率。

本文的组织如下:

*These authors contributed equally.

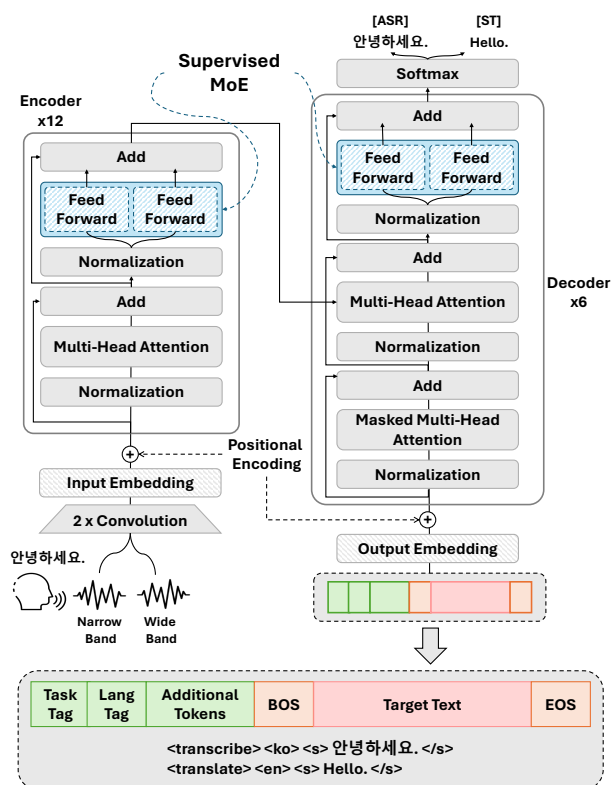


Figure 1: 嵌入在STT模型中的S-MoE。识别器中的S-MoE根据语音(窄带或宽带)输入向专家路由,而解码器中的S-MoE根据任务(自行语音识别或语音翻译)路由输入。

- 我们提出了一种S-MoE方法,通过使用指定每个任务路由,减少计算成本,从而避免MTL中的任务干扰。
- 我们进一步将S-MoE用于一个STT模型,使其能够处理混合输入,同时结合自行ASR和ST。
- 我们的方法同时用于识别和解码器,在WER上取得了6.35%的提升。

2. 方法

在本节中,我们提出了一种如图1所示的架构。第一步的路径在以下子图中介绍。

2.1. 监督家混合模型 (S-MoE)

通篇稍微调整后的 MoE [?] 中, S-MoE 模型的输出 y 表示为方程 1。标准 MoE 中相似, $E_i(x)$ 表示固定输入 x 的第 i 个专家的输出。S-MoE 的模型在于控制机制。标准 MoE 模型依赖于可控制 G 不同, 我们的方法采用固定的控制函数 G' , 无需额外的控制。如图 1 所示, 在整个工作过程中, 编码器和解码器中的专家数量 n 均固定为 2。

$$y = \sum_{i=0}^{n-1} G'(x)_i E_i(x) \quad (1)$$

2.1.1. 编码器 S-MoE 的控制功能

在理解 (WB) 信号, 基于窄带 (NB) 信号的模型向于表现不佳 [?, ?]。为了在单一模型中有效理解 NB 和 WB 信号, 我们在编码器中引入了一个 S-MoE 架构。在模型中, 前两个专家 (FFN) 作为每个输入信号的控制专家, 而编码器模型中的其余部件在 NB 和 WB 输入之间共享。模型信息, 每个输入先通过, 使得控制函数 $G'(x)$ 可以有选择地激活合适的专家。标准 MoE 类似, 如果 $G'(x)_i = 0$, 则不计算相应的专家 $E_i(x)$ 。因此, WB 信号由专家 E_0 处理, 而 NB 信号由专家 E_1 处理。编码器 S-MoE 的控制函数在公式 2 中被正式定义。

$$G'_{enc}(x)_i = \begin{cases} G'_{enc}(x)_0 = \begin{cases} 0, & \text{if } x \text{ is NB signals} \\ 1, & \text{if } x \text{ is WB signals} \end{cases} \\ G'_{enc}(x)_1 = \begin{cases} 0, & \text{if } x \text{ is WB signals} \\ 1, & \text{if } x \text{ is NB signals} \end{cases} \end{cases} \quad (2)$$

2.1.2. 解码器 S-MoE 的控制函数

类似地, 我们在解码器中使用了 S-MoE 架构, 以有效地理解 ASR 和 ST 任务。为了制定合适的解码专家, 任务被添加到文本输入的前面。基于目标任务, 控制函数 ASR 输入向专家 E_1 , ST 输入向专家 E_0 。解码器 S-MoE 的控制函数在方程 3 中给出。

$$G'_{dec}(x)_i = \begin{cases} G'_{dec}(x)_0 = \begin{cases} 0, & \text{if task of } x \text{ is ASR} \\ 1, & \text{if task of } x \text{ is ST} \end{cases} \\ G'_{dec}(x)_1 = \begin{cases} 0, & \text{if task of } x \text{ is ST} \\ 1, & \text{if task of } x \text{ is ASR} \end{cases} \end{cases} \quad (3)$$

如图 1 所示, 模型程序包含特殊 token 以引导模型。在每个序列的开头, 输入一个任意的 token (<transcribe> 或 <translate>), 其后是目标语言 token (<en> 或 <ko>)。其他特殊 token 也可能被包括在内, 但本研究中未使用。序列然后以 <beginning_of_sentence> token 开始, 后面是目标文本。图 2 展示了所提出的基于 S-MoE 的多任务模型的嵌入流程。

2.1.3. 推理阶段

输入音, 无论是 NB (8 kHz) 还是 WB (16 kHz), 首先被理解并输入到 Transformer 模型中。在模型中, 基于 Encoder S-MoE 的控制函数和 FFN。Encoder S-MoE 的控制函数根据输入/推理期输入信用的固定专家, 如在第 2.1.1 节中所述。这意味着对于 NB 和 WB 输入使用不同的 FFN, 从而使模型能有效捕捉特定输入表的

示。在解码器中, Decoder S-MoE 利用独立的 FFN 处理共享的 token 表示。每个推理批次由任一任务的文本组成, 以确保任意的变化。模型通过交替批次运行, 这意味着 ASR 和 ST 任务交替进行。

2.1.4. 推理阶段

在推理过程中, 模型可以同时生成 ASR 和 ST 输出。通过批量大小设置为 2, 相应地分配 ASR 和 ST 任务, 可以在一次推理步骤中同时得到 ASR 和 ST 结果。

3. 数据集

3.1. 数据集

3.1.1. 数据集

我们的数据集由来自公开可用的 AIHub 材料 [?, ?, ?, ?, ?, ?, ?, ?] 的 40,000 小时语音数据组成。对于自语音 (ASR), 我们使用提供的文本作为目标文本。对于语音 (ST), 如果有相应的英文文本, 我们使用机器翻译 (MT) 模型生成目标文本。如果提供了原始文本, 我们直接使用。因此, 每个语音样本都有配套的 ASR 和 ST 数据, 而无需多任务。

为了在不同条件下训练模型, 我们生成了数据的 NB 和 WB 数据。WB 数据 16 kHz, 而 NB 数据 8 kHz。我们采用以下方法: (1) 将 16 kHz 下采样到 8 kHz, (2) 使用 AMR-WB [?]、AMR-NB 和 G.711-NB [?] 解码器进行基于解码器的推理。这些数据用于 15 % NB/WB

3.1.2. 数据集

为了评估 ASR 和 ST 的性能, 我们使用了两个数据集以及公开可用的 Fleurs [?] 数据集。我们的数据集由 1000 个自日常生活的男女语音样本组成, 由两个作者撰写的文本。此外, 为了在不同条件下训练模型, 我们使用数据集中相同的推理方法生成数据集的 NB/WB 数据。

3.2. 模型

我们的基模型遵循一个基于 Transformer 的序列到序列架构 [?], 包含一个 12 层的编码器和 6 层解码器。我们采用正弦位置嵌入和具有 GLU 激活的正向化 Transformer 块, 使用 8 个注意力头。前两个层的维度为 2048, 嵌入维度为 512, 且使用 0.15 的置熵率。对于分词, 我们使用大小为 40,000 的字节对编码 (BPE) 分词器。输入音被重采样到 16 kHz, 并分成 80 毫秒的梅尔倒谱滤波器 (fbank) 特征, 其特征是使用 25 ms 窗口和 10 ms 步长提取的。我们使用最大音序列长度 3,000 和最大文本序列长度 120 的限制, 并去掉少于 30 秒的语音片段。模型训练 10 万次, 使用大小为 3,200 的有效批次, 采用余弦退火调度器进行训练率衰减。章 4.1 的解码器-解码器 S-MoE 模型通过章 ?? 的解码器 S-MoE 模型进行初始化, 并使用 NB/WB 微小子集 5 万次。

3.3. 度量

对于 ST, 我们使用 BLEU 指标 [?] 衡量模型质量。我们使用了 sacrebleu。对于 ASR, 我们报告 WER, 它是 ASR 研究中广泛使用的指标。

4. 结果

为了评估 S-MoE 在解码器模型中的有效性, 我们进行了任务推理: 从英语 (ko2en) 的语音 (ST)

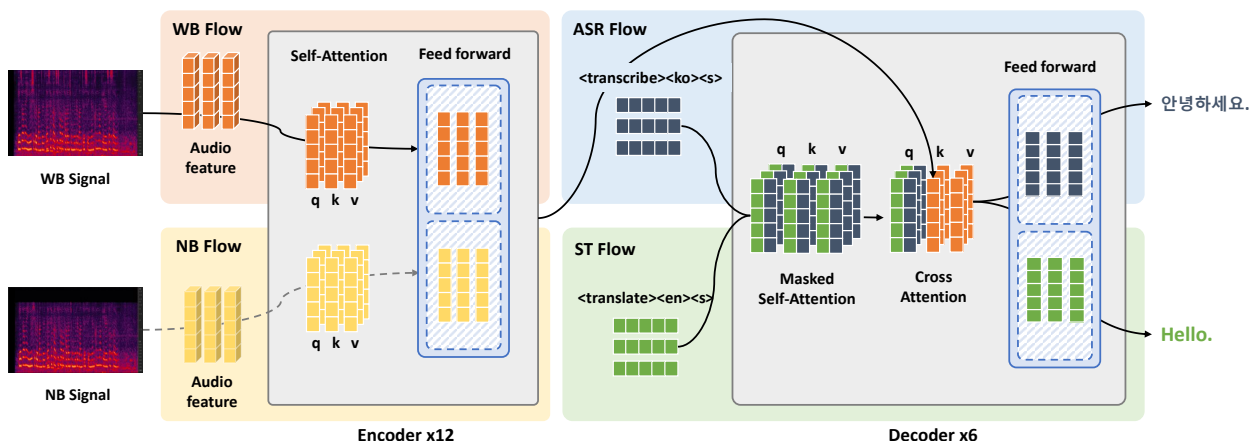


Figure 2: *S-MoE Transformer* 的嵌入流。在□□/推理期□,□□器□□□一□前□□,而解□器□的□□前□□都被利用。

Table 1: 解□器 *S-MoE* □ *ASR/ST* 的影□

Model	Param		Male TC		Female TC		Fleurs	
	Trainable	Active	WER	BLEU	WER	BLEU	WER	BLEU
Whisper-base	74M	74M	8.73	13.6	13.40	11.9	18.03	3.3
Whisper-small	244M	244M	5.15	26.5	7.40	25.8	10.27	12.6
Whisper-large-v3	1,550M	1,550M	4.46	33.7	6.87	32.8	5.32	20.1
Ours (Base-Model)	107M	107M	3.93	34.3	5.36	33.2	5.46	21.1
Ours (DecFFNx2)	119M	119M	3.83	34.1	5.36	33.5	5.46	21.5
Ours (DecS-MoE)	119M	107M	3.53	34.8	4.89	33.9	5.23	21.5

Table 2: □□器-解□器 *S-MoE* □ *ASR/ST* 的影□

Model	Codec	Male TC		Female TC		Fleurs	
		WER	BLEU	WER	BLEU	WER	BLEU
Base-Model	NB	4.11	33.8	6.60	32.5	5.91	19.9
	G.711 NB	4.18	33.9	6.77	32.4	6.50	19.9
	AMR NB	4.78	32.8	11.30	29.8	13.69	17.9
	WB	4.51	33.8	6.16	32.0	8.71	19.1
	AMR WB	3.98	33.9	5.46	32.9	5.84	19.8
EncDecS-MoE	NB	4.08	34.1	6.22	32.9	5.66	20.4
	G.711 NB	3.92	34.2	5.96	32.8	6.15	20.0
	AMR NB	4.65	34.0	10.43	30.5	12.96	18.2
	WB	4.17	34.2	6.98	32.3	8.30	19.1
	AMR WB	3.88	34.2	5.18	33.5	5.34	20.2

和□□自□□音□□(ko-ASR)。□果□□在表 1 中。第 3.2 □中描述的模型作□□比的主要基□(Base-Model)。

□着□家□量的增加,*S-MoE*的可□□□□□性增□。根据 [?] ,□管□□略有增加,但 *MoE* □□仍保留了□算□□,因□□基于路由机制□激活了部分□家□□。在我□的□□中,我□□□家□量□置□2,以支持 *ASR* 和 *ST* 的多任□□□。在表格中,□*S-MoE* □用于解□器的模型□□*DecS-MoE*。□了匹配 *DecS-MoE* 中的□□增加,我□□□□了一□□□,□□□的解□器 *FFN* □度加倍 (*DecFFNx2*)。

□了更□泛的比□,我□在□□中包括了 OpenAI 的 Whisper [?] 模型。Whisper 是一□□泛使用的多□言□音模型,使用多□□言□□行□□。我□的模型□□□一□□方向□行了□化,所以直接比□可能不完全公平,但 *Whis-*

per □□□ko-ASR 和 ko2en *ST* 性能提供了一□强有力的□照点。

如表 1 所示,*DecS-MoE* 在所有□□集上的 *ST* 和 *ASR* 方面均持□□于基□模型。我□□察到,在男性□□集上, *BLEU* 提高了 0.5 点, *WER* 降低了 0.4 % *BLEU* 0.7 *WER* 0.47 % *Fleurs* *DecS-MoE* *BLEU* 0.4 *WER* 0.2 % *DecS-MoE* *WER* 0.37 % *BLEU* 0.53 *WER* 8.06 % *BLEU* 1.77 % *S-MoE*

在可□□□□□量相同的情□下,*DecS-MoE* 甚至比 *DecFFNx2* 表□更佳。在男性□□集上, *BLEU* 提高了 0.7 分,而 *WER* 降低了 0.3 % *DecS-MoE* *BLEU* 0.4 *WER* 0.47 % *Fleurs* *ST* *WER* 0.23 % *WER* 0.33 % *BLEU* 0.37 *WER* 7.33 % *BLEU*

1.22 % FFN
DecS-MoE

即使使用 Whisper 模型, DecS-MoE 在所有 ASR/ST 数据集上也取得了卓越的性能。尽管其可训练参数量仅为 Whisper-large(1,550M) 的十分之一, DecS-MoE(119M) 的平均 BLEU 提升了 1.2, 同时 WER 减少了 1 %。这表明 WER 相对提升了 21.98 %, BLEU 提升了 3.99 %。

4.1. 编码器-解码器 S-MoE

我们提出了一组编码器 S-MoE 来处理不同的语音信号。具体而言, 编码器由两个专家组成, 一个用于 NB 信号, 一个用于 WB 信号。在 ?? 中, 我们详细了解了编码器 S-MoE 在结合 ASR 和 ST 任务时的效果。因此, 我们使用 S-MoE 用于编码器和解码器(EncDecS-MoE)。我们共训练了两个专家, EncDecS-MoE 模型有 144M 可训练参数, 而基线模型有 107M。然而, 在推理过程中, 我们模型使用的主参数量相同(107M)。例如, 在推理用于 ST 的 NB 语音输入时, EncDecS-MoE 模型通过编码器路由到 NB 专家, 解码器路由到 ST 专家, 只激活相关的专家以高效计算。

表 2 中的所有模型都在第 3.1.1 中描述的 NB/WB 数据集上进行了微调。由于 WB 信号相比, NB 语音信号的使用更有限, 因此 NB 不需要太多的训练数据。为了在 NB 和 WB 场景中都能有效工作, 我们对基线模型进行了微调, 模型最初在 WB 数据集上进行了训练。

如表 2 所示, EncDecS-MoE 在几乎所有情况下都表现出一致的性能提升。在 NB 场景中, EncDecS-MoE 的平均 BLEU 得分 28.57, 而基准模型 28.10, 平均 WER 6.67 % 7.09 % WER 相对提高了 6.35 % BLEU 得分提高了 1.63 % WB 场景下, 我们提出的模型在 WER 上相对于平均性能有 2.39 % 的相对提升, 在 BLEU 上有 1.15 % 的相对提升。NB 信号捕获率低, 可能在推理 WB 和 NB 数据的同一模型中产生负面影响。然而, 结果表明在编码器中施加独立的 FFNs 可以缓解这一问题。虽然使用独立的专家略微增加模型大小, 但活参数的量与基准模型相同。因此, 在编码器和解码器中使用 S-MoE 提升了性能, 而不影响推理速度。

5. 结论

我们提出了 S-MoE 架构, 旨在解决多任务中 STT 任务的任务干扰。通过使用指示性令牌代替传统的控制功能, S-MoE 在提高各项任务性能的同时, 保持了高效的推理和推理。我们的方法使得同一模型能够推理具有混合输入任务的 ASR 和 ST 任务, 使其成为资源受限环境的理想选择。未来的工作集中在扩展 S-MoE 以支持更多语音任务和多语言功能, 从而进一步增强其在现实世界中的适用性。

6. References