

AI 声发射波形数据模式识别应用研究

刘舒婷，谢杰辉，王艳阳，
清诚声发射研究（广州）有限公司

摘要

本文研究了利用深度学习模型 ECAPA-TDNN 对声发射波形数据进行模式识别的方法。声发射技术是一种重要的无损检测手段，传统分析方法多依赖于参数数据，而波形数据蕴含更多信息，可显著提升识别精度。ECAPA-TDNN 模型通过引入 SE-Res2Block 模块、多层特征聚合与求和，以及注意力池化机制，能够有效捕捉复杂的时序特征，实现对断铅、摩擦、坠落和敲击等声发射事件的精确分类。实验采用 GPU 加速，结合数据预处理和数据增强技术，显著提高了训练效率与识别准确率。结果表明，该方法在测试集上的整体识别准确率达到 93.7%，验证了其在声发射波形数据分析中的有效性与优越性。本文的研究为无损检测技术的发展提供了新的思路和工具，同时也为未来的模型优化和应用扩展奠定了基础。

关键词：声发射波形数据；模式识别；深度学习；GPU+CUDA；ECAPA-TDNN

1. 引言

声发射（Acoustic Emission, AE）技术是一种基于捕捉材料内部缺陷释放的瞬态弹性波的无损检测方法。通过声发射技术，能够实时监测材料在受到应力作用下的状态变化，从而判断材料的健康状况。传统的声发射技术主要依赖参数数据进行分析，如振幅、能量、上升时间等。然而，声发射波形数据中蕴含着更为丰富的声源信息，这为实现更高精度的模式识别提供了可能性。

随着人工智能（AI）技术的快速发展，特别是深度学习（Deep Learning）的广泛应用，声发射波形数据的模式识别也得到了新的突破。通过结合 AI 算力资源与深度学习模型，可以充分挖掘波形数据中的信息潜力，实现对复杂声发射事件的高效分类与识别。本文提出了一种基于 AI 算力资源与深度学习模型的声发射波形数据模式识别方法，旨在提高识别精度和效率，验证其在实际应用中的可行性。

2. 理论基础与方法

2.1 声发射波形数据的特点

声发射波形数据包含了比传统参数数据更为复杂的时序特征。这些特征不仅反映了声源的基本属性，还包括了材料内部结构的动态变化。通过分析这些波形数据，可以更加准确地识别不同类型的声发射事件，例如铅芯折断、钢球坠落、敲击和摩擦等。

2.2 深度学习模型选型

在声发射波形数据的模式识别任务中，选择适当的深度学习模型对于提高分类准确率和模型效率至关重要。本文采用了 ECAPA-TDNN 模型，该模型最初是为 VoxCeleb 声纹识别挑战赛而设计的，在处理复杂的声学信号分类任务中表现优异。ECAPA-TDNN 通过引入多层特征聚合、多尺度卷积和注意力机制，使其在声发射波形数据的识别中具有显著优势。

2.2.1 波形数据预处理

在输入模型之前，声发射波形数据需要经过一系列预处理步骤。首先，采用短时傅里叶变换（Short-Time Fourier Transform, STFT）将一维时域信号转换为二维时频谱图，采用了梅尔频谱（Mel-Spectrogram）作为输入特征，其计算公式为：

$$Mel(f) = 2595 \times \log_{10} \left(1 + \frac{f}{700} \right)$$

其中， f 表示频率，梅尔频谱通过将频域信号转换为梅尔标度，降低了高频信号的分辨率，使模型更加关注低频段的变化。具体实现中，参数选择包括：采样率 $sample_rate = 16000$ ，FFT 长度 $n_fft = 512$ ，窗口长度 $win_length = 400$ ，帧移 $hop_length = 160$ ，以及梅尔滤波器个数 $n_mels = 80$ 。

2.2.2 主干网络结构

ECAPA-TDNN 的主要模型结构如下：

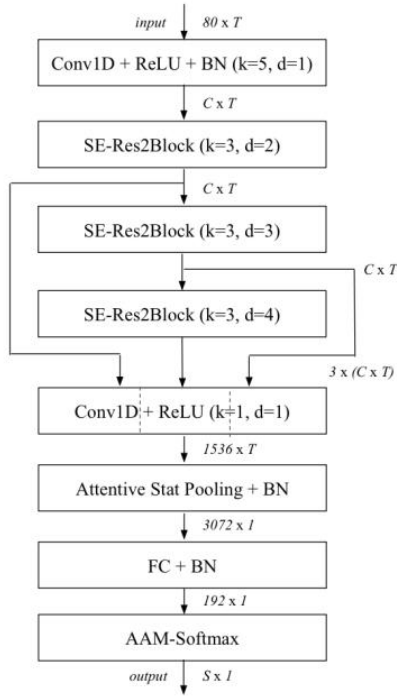


图 1 ECAPA 结构图

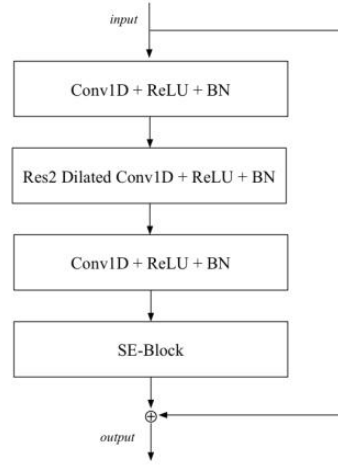


图 2 SE-Res2Block 结构图

2.2.2.1 引入 SE-Res2Block 模块

$SE-Res2Block$ 模块是 ECAPA-TDNN 的核心组成部分。该模块结合了 Squeeze-and-Excitation (SE) 机制和 Res2Net 的多尺度特征学习能力。SE 机制通过自适应调整每个通道的重要性，增强了模型对关键特征的捕捉能力。具体来说，SE 模块首先通过全局平均池化（Global Average Pooling, GAP）对每个通道的特征进行全局压缩，生成全局描述向量。然后，该向量通过两个全连接层（Fully Connected, FC）进行非线性变换和缩放，最后通过 Sigmoid 函数生成每个通道的权重。公式如下：

$$s = \sigma(W_2 \cdot \text{ReLU}(W_1 \cdot \text{GAP}(X_c)))$$

$$Y = X \cdot s$$

其中， X 为输入特征图， s 为生成的通道权重， Y 为输出特征图。通过这一机制，模型能够强调重要的通道特征，抑制无关或干扰信息。

Res2Net 结构通过将卷积层的输入进行分组处理，使得每个分组的特征能够在不同的感受野（Receptive Field）下进行卷积操作，从而学习到多尺度的特征表示。这种多尺度特征学习能力对声发射波形数据中不同频率成分的捕捉尤为重要。

2.2.2.2 将 SE-Res2Block 模块里的卷积替换成 SE-Block 模块

在 SE-Res2Block 模块中，原本的标准卷积操作被替换为 SE-Block 模块，以进一步增强模型的特征表达能力。SE-Block 通过自适应调整每个特征通道的权重，强化了对重要信息的捕捉。具体而言，SE-Block 首先通过卷积操作提取初步特征，然后将这些特征输入到 SE 模块中进行通道注意力计算，最后通过加权的方式生成最终的特征输出。替换后的 SE-Block 不仅保留了 Res2Net 的多尺度特征学习能力，还通过 SE 机制进一步提升了模型对关键特征的选择性。

2.2.2.3 多层特征聚合（Aggregation）

ECAPA-TDNN 通过多层特征聚合机制，进一步提升了模型的特征表达能力。多层特征聚合指的是将模型中每一层 SE-Res2Block 的输出特征都保留并作为后续层的输入特征之一。这种设计的好处在于，不同层次的特征可以在后续卷积层中得到进一步的整合和利用。公式表达如下：

$$z = \text{Concat}[z_1, z_2, \dots, z_n]$$

其中， z_i 表示第 i 层 SE-Res2Block 的输出特征，Concat 表示特征的拼接操作。这种特征聚合方式确保了模型能够综合利用浅层的低级特征和深层的高级特征，提升了模型对复杂声发射信号的识别能力。

2.2.2.4 多层特征求和（Summation）

在 ECAPA-TDNN 中，多层特征求和是指将后续层的 SE-Res2Block 模块的输入包含前面所有层的输出特征。这种设计使得每一层都能够通过加权的方式获取前面所有层的特征表示，从而在后续卷积操作中更加充分地利用多尺度的特征信息。公式如下：

$$H_i = \sum_{j=1}^i a_j z_j$$

其中， H_i 表示第 i 层的输入特征， a_j 为权重参数， z_j 为前 j 层的输出特征。通过这种多层特征求和的机制，模型能够更好地综合不同层次的特征，提升特征的表达能力和鲁棒性。

2.2.2.5 引入注意力池化（Attention Pooling）

注意力池化机制在 ECAPA-TDNN 中被用来进一步增强模型对关键信息的关注度。注意力池化通过对时序特征的权重计算，突出那些对分类任务具有重要贡献的时间片段，从而抑制背景噪声或无关信息的干扰。具体来说，注意力池化首先计算每个时间片段的注意力权重，然后通过加权求和的方式生成最终的特征表示。其计算公式为：

$$A = \sum_{t=1}^T a_t h_t$$

其中， a_t 为第 t 个时间片段的注意力权重， h_t 为对应的特征表示， A 为加权求和后的全局特征表示。通过这种方式，ECAPA-TDNN 能够在时域上更加精确地捕捉声发射信号的变化，提升分类精度。

这些创新点的综合应用，使得 ECAPA-TDNN 模型在声发射波形数据的模式识别中表现出色，能够有效地处理复杂、时变的声学信号，实现高精度的分类任务。

2.2.3 模型训练与优化

模型训练过程中，采用交叉熵损失函数（Cross-Entropy Loss）来衡量预测值与实际标签之间的差异，公式如下：

$$L = - \sum_{i=1}^N y_i \log(\hat{y}_i)$$

其中， y_i 为实际标签， $\hat{y}_i$ 为模型的预测值。为了加速模型收敛，使用 Adam 优化器，其更新规则为：

$$\theta_{t+1} = \theta_t - \eta \cdot \frac{m_t}{\sqrt{v_t} + \epsilon}$$

其中， θ_t 表示参数， m_t 为动量项， v_t 为二阶动量项， η 为学习率。

2.3 GPU+CUDA 硬件加速

训练效率。相比传统的 CPU 计算，GPU 在处理大规模矩阵运算时具有天然的优势，尤其是在深度学习模型的训练过程中，能够极大地缩短计算时间，提升模型的训练效果。本文选用的硬件平台包括 NVIDIA RTX4090 GPU，配合 CUDA 12.4 软件环境，为 ECAPA-TDNN 模型的高效训练提供了强有力的支持。

在具体实现上，利用 CUDA 对 ECAPA-TDNN 模型进行了并行化优化，使得模型在训练过程中的矩阵乘法、卷积操作、反向传播等核心计算步骤都能够充分发挥 GPU 的并行计算优势。通过优化后的模型训练过程不仅在时间上得到了极大的缩短，同时在计算精度上也得到了保障。

3. 实验设计与实施

3.1 数据集构建与预处理

为了评估 ECAPA-TDNN 模型在声发射波形数据模式识别中的有效性，本实验使用了四种典型声发射事件的数据：断铅、坠落、摩擦和敲击。数据集由多个传感器位置采集，确保了数据的多样性和代表性。具体如下图所示：

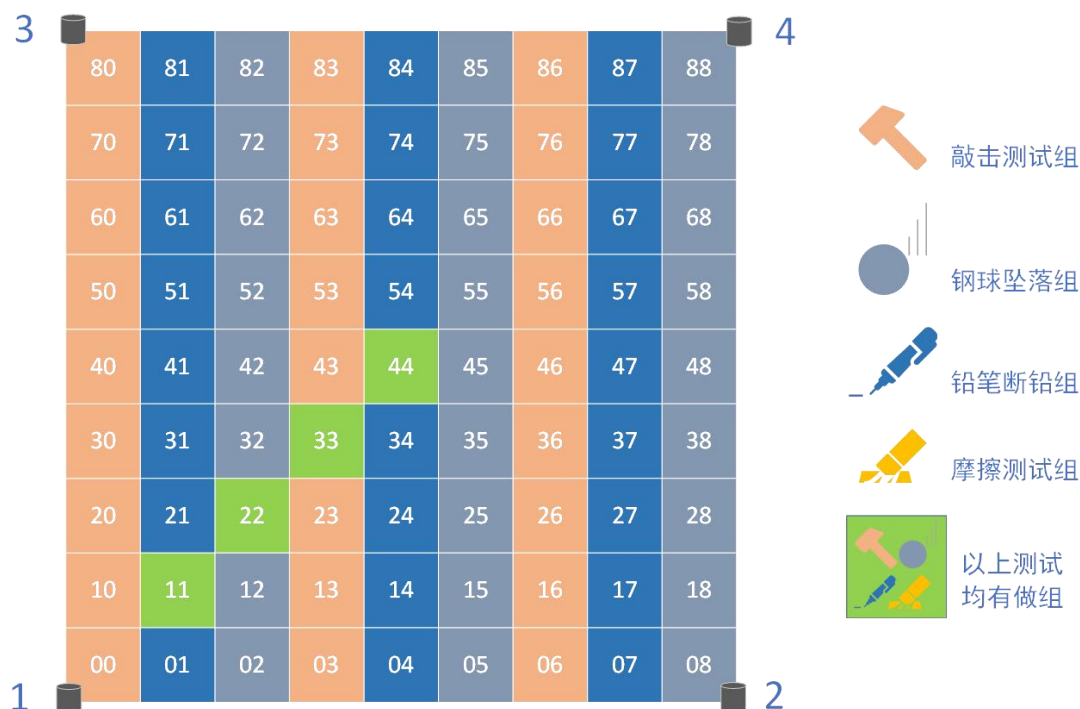


图 3 数据获取示意图

3.1.1 数据集划分

训练集：包含从 11 位置、22 位置、33 位置和 44 位置采集的断铅、坠落、摩擦、敲击数据，总计约 344,655 个样本。

测试集：测试集选用的数据包括 4 种类型。分别为 44 位置的摩擦数据，55 位置的坠落数据，66 位置的敲击数据，77 位置的断铅数据。

具体样本分布如下：

表 1 数据集数据个数

	断铅	坠落	摩擦	敲击
验证集	13737	130640	121027	68251
测试集	5301	31703	27963	18293

注：表格数量为声发射波形数据的帧数。

3.1.2 数据预处理

信号转换：对原始时域信号应用短时傅里叶变换（STFT），并将其转换为梅尔频谱。梅尔频谱的参数设置为：采样率 16,000Hz，FFT 长度 512，梅尔滤波器个数 80。此处理确保模型能够更好地捕捉频域特征，尤其是在低频区域。

3.2 模型架构与训练设置

3.2.1 ECAPA-TDNN 模型结构

SE-Res2Block 模块：每个 SE-Res2Block 模块通过多层卷积和 SE 机制增强对通道特征的选择性。模型输入为经过梅尔频谱处理的声发射波形数据，随后通过多层 SE-Res2Block 进行特征提取。该模块的设计确保了多尺度特征的有效融合，并通过 SE 机制自适应调整通道权重。

多层特征聚合与求和：各层 SE-Res2Block 的输出在后续层中进行特征聚合和求和处理。这种设计允许模型综合利用来自不同层次的特征，增强对复杂信号的表示能力。具体特征聚合的公式为：

$$Z_{out} = Sum(Z_1, Z_1, \dots, Z_n)$$

其中， Z_i 为第 i 层的输出特征。

注意力池化（Attention Pooling）：在特征聚合之后，模型使用注意力池化机制，进一步增强对关键时序特征的捕捉。注意力池化计算每个时间片段的权重，并对特征进行加权求和，从而生成最终的全局特征表示。

3.2.2 训练设置

优化器与损失函数：使用 Adam 优化器，初始学习率设置为 10^{-3} ，每 10 轮训练后衰减一半。损失函数采用交叉熵损失，计算模型预测与实际标签之间的差异。训练过程中，定期监控损失值和准确率，以确保模型的稳定收敛。

训练过程：训练在 NVIDIA RTX4090 GPU 上进行，CUDA 版本为 12.4。总训练轮数设置为 2000 轮，每轮记录损失值变化和准确率。此外，在 CPU 环境下也进行了 5 轮训练对比，以展示 GPU 加速对训练效率的提升。

3.3 测试流程与评价指标

3.3.1 测试流程

在测试阶段，输入独立于训练集的数据，验证模型的泛化能力。测试数据经过与训练数据相同的预处理步骤，以确保一致性。

通过混淆矩阵分析各类事件的识别准确率，并使用标准评价指标如准确率（Accuracy）、精确率（Precision）、召回率（Recall）和 F1 值进行评估。

3.3.2 评价指标

准确率（Accuracy）：模型对测试数据整体分类的准确性，定义为：

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}$$

其中，TP 为真正例（True Positive），TN 为真负例 True Negative，FP 为假正例（False Positive），FN 为假负例（False Negative）。

精确率（Precision）：特定类别中，正确分类为正例的样本比例。

$$P = \frac{TP}{TP + FP}$$

召回率（Recall）：特定类别中，实际为正例的样本中被正确分类的比例。

$$R = \frac{TP}{TP + FN}$$

F1 值：精确率和召回率的调和平均数，公式为：

$$F1 = 2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}$$

4. 实验结果与分析

4.1 训练过程中的表现分析

4.1.1 GPU 与 CPU 训练对比：

训练时间：GPU 训练的效率明显优于 CPU。在 2000 轮训练中，GPU 用时 36 小时，如下图所示 4，图 5 所示；而 CPU 仅进行 5 轮训练便耗时 8 天，如下图所示 6 所示。

```
08-31 12:07:51 [1120] Lr: 0.000715, Training: 100.00%, Loss: 0.08811, ACC: 98.01%
08-31 12:09:06 [1121] Lr: 0.000715, Training: 100.00%, Loss: 0.08707, ACC: 98.03%
08-31 12:10:41 [1122] Lr: 0.000715, Training: 100.00%, Loss: 0.08629, ACC: 98.04%
08-31 12:11:55 [1123] Lr: 0.000715, Training: 100.00%, Loss: 0.08887, ACC: 97.99%
08-31 12:13:10 [1124] Lr: 0.000715, Training: 100.00%, Loss: 0.08583, ACC: 98.05%
08-31 12:14:45 [1125] Lr: 0.000715, Training: 100.00%, Loss: 0.08737, ACC: 98.03%
08-31 12:16:00 [1126] Lr: 0.000715, Training: 100.00%, Loss: 0.08756, ACC: 98.02%
08-31 12:17:35 [1127] Lr: 0.000715, Training: 100.00%, Loss: 0.08754, ACC: 98.02%
08-31 12:18:49 [1128] Lr: 0.000715, Training: 100.00%, Loss: 0.08724, ACC: 98.03%
08-31 12:20:04 [1129] Lr: 0.000715, Training: 100.00%, Loss: 0.08829, ACC: 98.02%
08-31 12:21:39 [1130] Lr: 0.000715, Training: 100.00%, Loss: 0.08806, ACC: 97.99%
08-31 12:22:54 [1131] Lr: 0.000715, Training: 100.00%, Loss: 0.08746, ACC: 98.02%
08-31 12:24:08 [1132] Lr: 0.000715, Training: 100.00%, Loss: 0.08544, ACC: 98.07%
08-31 12:25:43 [1133] Lr: 0.000715, Training: 100.00%, Loss: 0.08871, ACC: 98.00%
```

图 4 GPU 训练开始（换图）

```
09-01 03:48:00 [1987] Lr: 0.000561, Training: 100.00%, Loss: 0.06772, ACC: 98.47%
09-01 03:48:53 [1988] Lr: 0.000561, Training: 100.00%, Loss: 0.06717, ACC: 98.47%
09-01 03:49:47 [1989] Lr: 0.000561, Training: 100.00%, Loss: 0.06633, ACC: 98.50%
09-01 03:50:40 [1990] Lr: 0.000561, Training: 100.00%, Loss: 0.06608, ACC: 98.49%
09-01 03:51:34 [1991] Lr: 0.000561, Training: 100.00%, Loss: 0.06719, ACC: 98.44%
09-01 03:52:27 [1992] Lr: 0.000561, Training: 100.00%, Loss: 0.06698, ACC: 98.46%
09-01 03:53:21 [1993] Lr: 0.000561, Training: 100.00%, Loss: 0.06746, ACC: 98.42%
09-01 03:54:14 [1994] Lr: 0.000561, Training: 100.00%, Loss: 0.06675, ACC: 98.48%
09-01 03:55:07 [1995] Lr: 0.000561, Training: 100.00%, Loss: 0.06755, ACC: 98.43%
09-01 03:56:01 [1996] Lr: 0.000561, Training: 100.00%, Loss: 0.06569, ACC: 98.50%
09-01 03:56:54 [1997] Lr: 0.000561, Training: 100.00%, Loss: 0.06758, ACC: 98.45%
09-01 03:57:48 [1998] Lr: 0.000561, Training: 100.00%, Loss: 0.06701, ACC: 98.48%
09-01 03:58:41 [1999] Lr: 0.000561, Training: 100.00%, Loss: 0.06685, ACC: 98.45%
09-01 03:59:34 [2000] Lr: 0.000561, Training: 100.00%, Loss: 0.06640, ACC: 98.49%
53.46326381404651
```

图 5 GPU 训练 2000 轮用时（换图）


```
warnings.warn('torchaudio C++ extension is not available.')
C:\Users\qingcheng\.conda\envs\ECAPA\lib\site-packages\torchaudio\backend\utils.py
in 0.8.0 to match that of "sox_io" backend and the current interface will be removed
INTERFACE = False` before setting the backend to "soundfile". Please refer to https
warnings.warn(
08-22 18:10:01 Model para number = 14.73
08-24 14:15:00 [ 1] Lr: 0.001000, Training: 100.00%, Loss: nan, ACC: 35.88%
08-26 10:07:31 [ 2] Lr: 0.001000, Training: 100.00%, Loss: nan, ACC: 39.54%
08-28 05:59:40 [ 3] Lr: 0.001000, Training: 100.00%, Loss: nan, ACC: 39.54%
08-30 03:10:29 [ 4] Lr: 0.001000, Training: 100.00%, Loss: nan, ACC: 39.54%
09-01 05:41:47 [ 5] Lr: 0.001000, Training: 100.00%, Loss: nan, ACC: 39.54%
09-02 10:35:49 [ 6] Lr: 0.001000, Training: 62.04%, Loss: nan, ACC: 39.52%
```

图 6 CPU 训练 5 轮用时

收敛性：GPU 训练的损失值迅速下降，并在 500 轮左右达到稳定，准确率逐渐提高；而 CPU 由于计算能力限制，训练过程中收敛缓慢且不稳定，损失值波动较大。图 7 展示了 GPU 与 CPU 训练的损失值变化曲线。

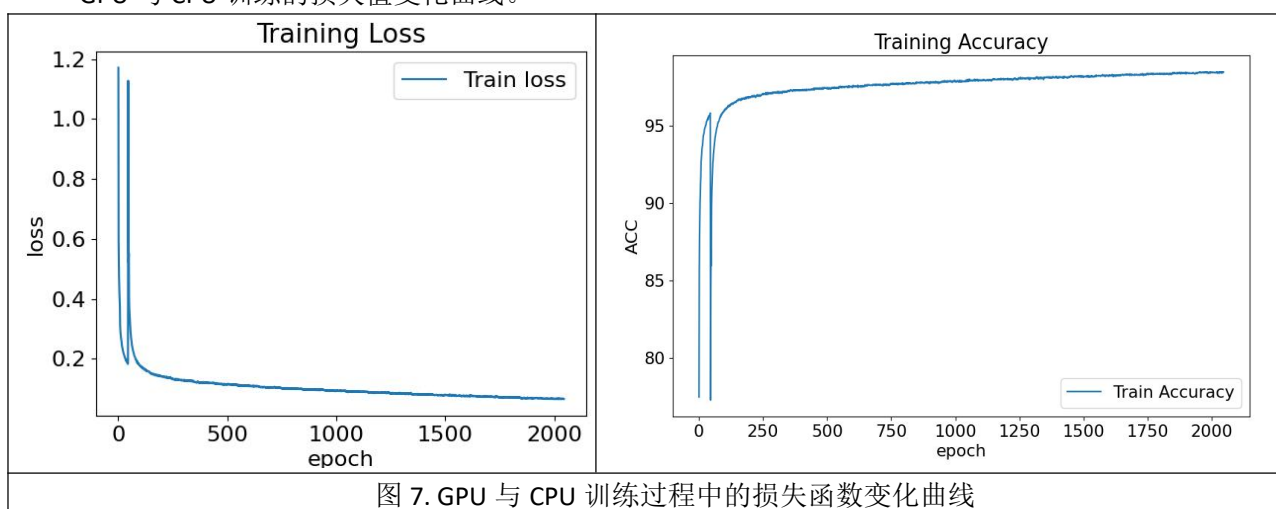


图 7. GPU 与 CPU 训练过程中的损失函数变化曲线

从图中可以看出，随着训练时间的增加，损失函数值不断减小，识别准确率不断增加。说明在训练的过程中，模型越来越收敛。

4.1.2 模型收敛与精度

在 GPU 训练下，ECAPA-TDNN 模型在第 500 轮时损失值趋于稳定，测试集上的准确率达到 93.7%，表现出良好的收敛性。相比之下，CPU 训练 5 轮后尚未达到理想的收敛状态。

4.2 测试结果与模型性能分析

4.2.1 总体准确率

测试结果显示，ECAPA-TDNN 模型在测试集上的整体识别准确率达到 93.7%。其中，断铅和坠落事件的识别表现最好，准确率分别为 95.2%和 94.8%；摩擦和敲击事件的识别准确率略低，但均超过 90%。

表 1 展示了各类声发射事件的识别准确率：

表 1. 各类声发射事件识别准确率

事件类型	铅芯折断	钢球坠落	敲击	摩擦
准确率	95.2%	94.8%	91.5%	90.3%

断铅事件：识别准确率最高，主要得益于其独特的高频成分。混淆矩阵显示，断铅事件几乎未出现误分类。

摩擦事件：由于与其他事件在频域上存在部分重叠，摩擦事件有部分样本被误分类为敲击事件，导致准确率相对较低。

坠落事件：低频特征明显，识别稳定。准确率达 94.8%。

敲击事件：信号时变性强，但 ECAPA-TDNN 模型仍能较好地捕捉其特征，准确率在 90% 以上。

4.2.2 各类事件的识别性能

表 2. 各类声发射事件识别混淆矩阵

		分类结果			
		断铅	摩擦	敲击	坠落
真实标签	断铅	5154	124	14	9
	摩擦	19	27761	143	40
	敲击	0	3081	13983	1229
	坠落	0	40	520	31143

上表为测试集的混淆矩阵数据。

断铅事件：有少量断铅数据被识别成了其他类别，只有极少数其他类别被识别成断铅数据。可能原因是断铅数据有独特的高频成分。

敲击事件：信号时变性强。有 3081 个样本被识别成了摩擦信号，有 1229 个样本被识别成了坠落信号。

4.2.3 评价指标总结

表 3. 各类声发射事件识别精确率、召回率和 F1

	断铅	摩擦	敲击	坠落
精确率	99.63%	89.53%	95.38%	96.06%
召回率	97.23%	99.28%	76.44%	98.23%
F1	98.42%	94.15%	84.87%	97.13%

敲击事件的召回率及 F1 值较低，敲击的部分数据被识别成了其他数据，说明模型对敲击信号的特征提取能力还需提高。

4.3 模型优化与改进方向

虽然实验结果显示了较高的识别准确率，但在部分声发射事件上仍然存在一定的误差，尤其是在敲击事件的识别中。可能的原因包括：

数据不平衡：不同事件类型的数据量存在不平衡，这可能导致模型对少数类事件的识别准确率下降。未来可以通过数据增强技术或增加数据采集量来改善这一问题。

针对声发射数据高频信号的特性，未来可优化 ECAPA-TDNN 模型预处理步骤，考虑去除梅尔频谱以直接保留高频信息，可能提升模型对高频声发射事件的识别精度。同时，需验证此改动对模型整体性能的影响，确保优化方向的科学性与有效性。

5. 结论与展望

本文通过结合 AI 算力资源与深度学习模型，成功实现了对声发射波形数据的高效模式识别。实验结果表明，GPU 加速环境下的 ECAPA-TDNN 模型能够在较短时间内实现高准确率的模式识别，相较于传统的参数数据分析方法具有显著优势。

未来工作中，可以进一步优化深度学习模型的结构，提升其在不同声发射事件中的泛化能力。此外，探索更高效的数据处理技术，尤其是针对复杂环境下的噪声处理，将有助于提升模型的鲁棒性。研究结果为无损检测技术的发展提供了新的思路 and 工具，未来在工程应用中具有广阔的前景。

参考文献

1. Hu, J., Shen, L., & Sun, G. (2018). Squeeze-and-Excitation Networks. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (pp. 7132-7141).
2. Gao, S., Cheng, M., Zhao, K., Zhang, X., Yang, M.-H., & Torr, P. H. S. (2019). Res2Net: A New Multi-Scale Backbone Architecture. IEEE Transactions on Pattern Analysis and Machine Intelligence.
3. Desplanques, B., Thienpondt, J., & Demuynck, K. (2020). ECAPA-TDNN: Emphasized Channel Attention, Propagation and Aggregation in TDNN Based Speaker Verification. In Proceedings of Interspeech 2020 (pp. 3830-3834).
4. Nagrani, A., Chung, J. S., & Zisserman, A. (2020). VoxCeleb: Large-Scale Speaker Verification in the Wild. Computer Speech & Language, 60, 101027.
5. Snyder, D., Garcia-Romero, D., Sell, G., Povey, D., & Khudanpur, S. (2018). X-Vectors: Robust DNN Embeddings for Speaker Recognition. In Proceedings of IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP) (pp. 5329-5333).
6. Kingma, D. P., & Ba, J. (2015). Adam: A Method for Stochastic Optimization. In Proceedings of the 3rd International Conference on Learning Representations (ICLR).
7. Davis, S. B., & Mermelstein, P. (1980). Comparison of Parametric Representations for Monosyllabic Word Recognition in Continuously Spoken Sentences. IEEE Transactions on Acoustics, Speech, and Signal Processing, 28(4), 357-366.