

正交空间约束的特定辐射源识别小样本类增量学习方法

孙露¹, 薛睿¹, 查浩然¹, 林云¹, 王巍²

(1. 哈尔滨工程大学信息与通信工程学院, 黑龙江 哈尔滨 150001; 2. 中国电子科技集团第三十六研究所, 浙江 嘉兴 314033)

摘要: 针对样本稀缺条件下的特定辐射源识别小样本类增量学习任务, 现有方法缺乏显式几何约束易导致新旧类混淆, 提出了一种基于正交空间约束的特定辐射源识别小样本类增量学习方法。首先, 引入一组相互正交的目标向量作为结构化先验知识, 并从理论上推导其数量的上下界。其次, 基于正交伪目标向量设计联合优化策略, 对新类潜在表示方向施加显式几何约束, 引导特征提取器在嵌入空间中预留可扩展的表示方向。最后, 提出了分类器权重校准策略, 通过量化新类样本在判别过程中的误分风险, 利用高风险样本增强判别边界。在 ADS-B 和 Wi-Fi 数据集上的实验结果表明, 所提方法优于其他基准方法, 尤其在增量训练样本仅为 1 的极端情况下, 展现出更优的识别性能。

关键词: 特定辐射源识别; 小样本类增量学习; 正交空间约束; 分类器权重校准

中图分类号: TN911.7

文献标志码: A

doi: 10.11959/j.issn.1000-436x.TXXB260021

Few-shot class-incremental learning for specific emitter identification with orthogonal space constraints

Sun Lu¹, Xue Rui¹, Zha Haoran¹, Lin Yun¹, Wang Wei²

1. College of Information and Communication Engineering, Harbin Engineering University, Harbin 150001, China

2. The 36th Research Institute of China Electronics Technology Group Corporation, Jiaxing 314033, China

Abstract: To address the challenge of few-shot class-incremental learning (FSCIL) for specific emitter identification (SEI) with scarce samples, and the issue that existing methods lacked explicit geometric constraints, leading to new class embeddings being easily confused with those of old classes, a method for FSCIL based on orthogonal space constraints was proposed. Firstly, a set of mutually orthogonal pseudo-target vectors was introduced as structured prior knowledge, and the lower and upper bounds of their quantity were theoretically derived. Next, a collaborative optimization strategy based on orthogonal pseudo-target vectors was proposed, so as to impose geometric constraints on new class representation directions and guide the feature extractor to reserve expandable representation directions in the embedding space. Finally, a classifier weight calibration strategy was designed to quantify the misclassification risk of new class samples during the decision process, using high-risk samples to enhance the decision boundaries. Experimental results on the ADS-B and Wi-Fi datasets show that, the proposed method outperforms other benchmark methods, especially in the extreme case where only one incremental training sample is available.

Key words: specific emitter identification, few-shot class-incremental learning, orthogonal space constraint, classifier weight calibration

收稿日期: 2026-01-07; 修回日期: 2026-04-08

通信作者: 薛睿, xuerui@hrbeu.edu.cn

基金项目: 国家自然科学基金资助项目(No.U23A20271); 中国船舶集团有限公司第七二二研究所创新基金资助项目(No.2023J-6)

Foundation Items: The National Natural Science Foundation of China (No.U23A20271), The Innovation Fund of China State Shipbuilding Corporation Limited 722 Research Institute (No.2023J-6)

0 引言

特定辐射源识别 (specific emitter identification, SEI) 技术通过提取无线发射器硬件固有的射频指纹 (radio frequency fingerprint, RFF) 实现对发射源的身份认证与识别^[1-4]。随着大量无线通信设备和终端节点持续接入网络, 传统面向固定类别集合的 SEI 模型面临类别持续扩展的应用需求^[5]。传统方法通常需要定期收集数据并对模型进行全局重新训练, 但随着新辐射源的不断引入, 内存消耗与训练时间急剧上升^[6]。此外, 受限于数据采集难度和标注成本, 新增辐射源通常难以获得充足的标注样本, 使 SEI 同时面临新类样本稀缺与旧类数据不可访问的双重约束。因此, 在双重约束下以较低代价持续扩展模型的识别范围, 已成为动态设备接入场景下 SEI 亟待解决的重要问题。

类增量学习 (class-incremental learning, CIL) 方法为动态设备接入场景下的 SEI 提供了有效的解决思路。已有研究尝试将 CIL 方法引入 SEI 任务, 通过正则化约束^[7-8]、网络结构扩展^[9]、生成式回放^[10-11]等策略, 在模型更新过程中仅依赖新类别数据完成参数更新, 从而在避免旧类别知识灾难性遗忘的同时, 实现对新类别的有效学习。然而, CIL 方法通常假设新类样本相对充足, 当新辐射源仅提供极少量标注样本时, 模型参数极易受小样本随机扰动的影响从而导致过拟合。小样本学习 (few-shot learning, FSL) 致力于在极少量标注样本条件下实现对新类别的学习, 相关研究主要通过数据增强^[12-13]、对比学习^[14]和元学习^[15]等方法提升模型的小样本适应能力。但 FSL 通常针对单次小样本学习任务展开研究, 未考虑新类别持续到来的动态演化过程。因此, 本文研究小样本类增量学习 (few-shot class-incremental learning, FSCIL) 问题, 使模型在新类别持续到来的场景下, 同时具备旧类别知识保持能力和对小样本新类别的快速适应能力。

此外, 部分研究在 FSCIL 基础上进一步引入开集识别^[16], 形成了小样本增量开集识别 (few-shot incremental open-set recognition, FIOSR)^[17-18]任务。虽然 FIOSR 任务的设定较 FSCIL 更宽, 但 FSCIL 关注的小样本类增量学习仍是 FIOSR 的核心环节, 二者在该部分方法设计上具有通用性。本文将 FIOSR 中的小样本类增量学习方法纳入比较范围, 为保证任务一致性, 比较时仅采用其闭集小样本类增

量学习部分, 并不涉及未知类检测与拒识能力。

现有 FSCIL 方法主要分为两类: 基于特征空间的方法和基于动态网络的方法。其中, 基于特征空间的方法通过将特征提取器与分类器解耦, 在基础训练阶段利用充足的基类样本学习具有较强判别性和迁移能力的特征提取器, 在增量阶段则冻结特征提取器, 仅利用新类别的少量样本对分类器参数或类别原型进行更新。

对于基础训练阶段特征提取器的学习, 相关方法可归纳为表征增强和空间预留方法两类。表征增强方法通过提升基础训练阶段特征空间的表征质量, 能够在一定程度上提升后续新类接入的可迁移性。文献[17]通过角度惩罚损失与原型损失联合优化特征提取器。文献[18]基于元学习思想, 从基类数据中采样大量伪任务进行模拟训练, 以增强特征提取器对小样本学习场景的适应能力。文献[19]利用自监督对比学习增强类内紧凑性并扩大类间距离, 以提升特征空间的可分性。文献[20]指出类间距离并非越大越好, 因此通过适度缩小类间距离, 使学习的特征在判别性与可迁移性之间取得更好的平衡。然而, 这类方法往往缺少对后续增量类的空间规划, 当基础类别特征与新类别特征相似时, 可能导致新旧类的混淆。为缓解上述问题, 空间预留方法尝试在基础训练阶段通过前瞻性地预留特征空间, 提高模型对增量任务中新类别的容纳能力。文献[21]在基础训练阶段预分配虚拟类别并进行联合优化, 以前瞻性地为后续增量类预留嵌入空间。文献[22]为每个类别分配可收缩的决策边界, 并在训练过程中进行自适应调整, 以优化不同阶段的决策空间布局。尽管这些方法在一定程度上提升了特征空间对后续增量类任务的容纳能力, 但仍停留于经验性空间预留, 缺乏对新类嵌入方向的显式几何约束, 因而难以在冻结特征空间中持续实现低冲突的新旧类分布。

对于增量阶段分类器的学习, 现有方法通过微调策略提升新类别在既有特征空间中的适应能力。文献[17]引入原型校准机制, 利用与新类相似度较高的基类原型对新类原型进行修正。文献[19]提出了分阶段训练策略, 通过融合不同的数据增强方法逐步优化分类器参数。然而, 受样本稀缺和分布偏差的影响, 位于决策边界附近的边缘样本更易发生误分, 现有方法更侧重于对整体分布进行统一修

正,在细粒度判别建模方面仍存在不足。

基于动态网络的方法通过构建可扩展的网络结构,在开放模型中的部分节点或参数参与新类别学习,在尽可能保持旧类性能的同时提升模型对新类别的适应能力。与此同时,为缓解小样本条件下的过拟合风险,部分研究进一步引入网络压缩机制以控制模型复杂度。文献[23]通过增设新的网络层提升模型对新类别的学习能力,并在训练完成后对扩展节点进行压缩。文献[24]通过固定模型的最优子网来保留旧任务性能,同时利用剩余节点学习新任务,并在模型容量不足时进一步扩展新的网络节点。然而,这类方法的模型复杂度和训练时间也随增量过程不断上升,相比之下,基于特征空间的方法具有结构简洁和增量阶段开销较低的优势。

综上所述,本文沿用基于特征空间的增量学习范式,重点解决两方面挑战:①特征空间可持续扩展性不足,现有方法虽然尝试通过特征空间预留缓解新旧类干扰,但对后续增量类嵌入方向缺乏显式几何约束,使冻结后的特征空间难以为后续新类持续提供低干扰的特征嵌入方向;②判别边界不确定,为提升新类在已有特征空间中的适应能力,现有方法大多侧重于新类原型整体优化,缺乏对误分风险的量化以及基于易误分样本的定向校准机制,导致小样本条件下构建的新类原型存在偏差,从而难以形成可靠的判别边界。

本文提出了一种基于正交空间约束的特定辐射源识别小样本类增量学习方法,主要工作如下。

(1) 针对特征空间可持续扩展性不足的问题,引入一组相互正交的伪目标向量作为结构化先验知识,并从理论上推导了伪目标数量的上下界。基于伪目标向量设计交叉熵损失、自监督对比损失和类中心分离损失,对后续增量类嵌入方向施加显式几何约束,从而提升特征空间对后续新类别的持续扩展能力。

(2) 针对判别边界不确定问题,提出了一种分类器权重校准策略,通过量化新类样本的误分风险,利用高风险样本对分类器权重进行调整,从而缓解小样本条件下新类原型估计偏差带来的判别边界不稳定问题,提升模型对持续类别扩展的识别性能。

(3) 在ADS-B和Wi-Fi数据集上构建的FSCIL实验中进行了验证,实验结果验证了本文方法的有

效性,其整体性能优于现有先进方法。

1 问题定义

在将原始比特信息转换为射频信号的过程中,信号会受发射链路中各类硬件组件的非理想特性影响,包括本地振荡器的相位噪声与载波频率偏移、混频器的同相/正交失衡以及功率放大器的非线性失真等。这些硬件引入的畸变会在射频信号中留下独特的射频指纹,可用于区分不同的设备。基于深度学习的SEI模型通常由特征提取器 f_θ 与分类器 g_ϕ 构成。在模型训练期间,训练数据集可表示为 $D = \{(\mathbf{x}_i, y_i)\}_{i=1}^N$,其中 $\mathbf{x}_i \in \mathbb{R}^d$ 表示训练样本, $y_i \in C$ 表示对应的设备标签, N 表示训练样本的数量。信号 $\mathbf{x}_i$ 经特征提取器提取射频指纹特征,并经分类器计算对数概率后,通过Softmax激活函数得到归一化的概率分布,即:

$$p(y_i|\mathbf{x}_i) = \text{Softmax}(g_\phi(\mathbf{z}_i)) = \frac{\exp(g_\phi(\mathbf{z}_i)_{y_i})}{\sum_{j=1}^{|C|} \exp(g_\phi(\mathbf{z}_i)_j)} \quad (1)$$

其中, $\mathbf{z}_i = f_\theta(\mathbf{x}_i)$ 为信号 $\mathbf{x}_i$ 的特征表示向量。模型的训练目标是寻求一组模型参数,使数据集上的联合概率分布达到最大值,从而确保信号 $\mathbf{x}_i$ 能准确地映射到其对应的设备标签 y_i 上。这一目标可以表示为:

$$\max \prod_{i=1}^N p(y_i|\mathbf{x}_i, \theta, \phi) \quad (2)$$

FSCIL是基于一个统一模型从依次到达的小样本增量数据集中持续学习新类别,过程描述如图1所示。初始任务中的基类训练数据集 $D_0 = \{(\mathbf{x}_i, y_i)\}_{i=1}^{N_0}$,标签空间为 C_0 , D_0 中具有充足的训练样本。第 t 个增量任务中的新类训练数据集 $D_t = \{(\mathbf{x}_i, y_i)\}_{i=1}^{N_t}$,标签空间为 C_t , D_t 中只有少量训练样本,通常为1个或5个样本。不同任务的标签空间不重叠,即对于 $\forall m \neq n$ 满足 $C_m \cap C_n = \emptyset$ 。

在每个增量学习阶段,模型仅利用当前的新类训练数据集 D_t 学习新类别,并保持对旧类别 $\bigcup_{m=0}^{t-1} C_m$ 良好的识别性能,即需在包含所有已学习类别的测试数据集上最小化样本风险。

$$(\theta_t, \phi_t) = \arg \min_{\theta, \phi} \mathbb{E}_{(x,y) \sim D_{0:t}} [\ell(g_\phi(f_\theta(x)), y)] \quad (3)$$

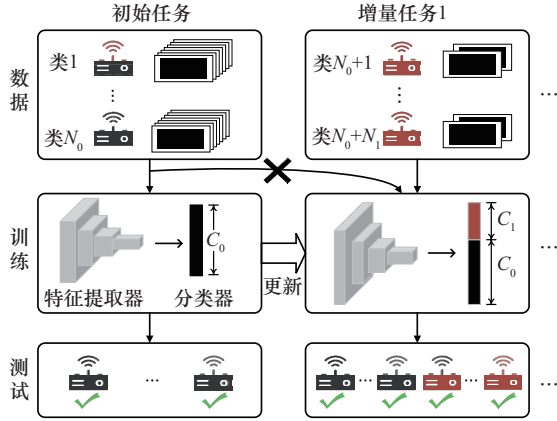


图1 FSCIL的过程描述

需要说明的是, 本文研究的是闭集测试条件下的特定辐射源识别小样本类增量学习问题, 测试样本均来自当前已学习类别集合, 不涉及未知类别的在线检测与拒识。

2 算法设计

2.1 算法框架

本文提出的基于正交空间约束的特定辐射源识别小样本类增量学习方法如图2所示, 由基类训练与小样本增量训练阶段构成。本文采用一维卷积神经网络作为特征提取器, 由6个一维卷积模块级联构成, 各模块均包含一维卷积、批归一化和一维最

大池化操作, 以提升训练稳定性和特征鲁棒性。网络逐层提取由局部到全局的信号特征, 并基于提取特征采用余弦相似度分类器进行分类判别。

阶段1 基类训练阶段学习可扩展的特征提取器 f_θ 和余弦相似度分类器的权重矩阵 W_0 (详见2.2节和2.3节)。首先, 构建相互正交的伪目标集合 $\mathcal{T}$ 。然后, 为每个基类分配一个伪目标作为分类器权重矩阵 $W_0 = [t_{h(1)}, \dots, t_{h(C_0)}]$, 未分配伪目标的基类作为后续增量类的预留方向。最后, 基于伪目标设计交叉熵损失、自监督对比损失和类中心分离损失, 以学习可扩展的特征提取器 f_θ 。

阶段2 小样本增量训练阶段固定特征提取器 f_θ , 仅学习分类器的新类权重矩阵 W_{new} (详见2.4节)。首先, 将分类器权重矩阵扩展为 $W = [W_{\text{old}}, W_{\text{new}}]$, 并采用类均值初始化新类权重。然后, 基于边际竞争约束和原型对齐约束对 W_{new} 进行校准微调。

2.2 伪目标构建

本文生成了一组类别无关且相互正交的伪目标向量, 其均匀分布在单位超球面上, 具有良好的方向分离性和覆盖性。定义一组随机向量 $\mathcal{T} = \{t_i\}_1^N \subset \mathbb{R}^d$, 通过最大化任意两个向量 $t_i, t_j \in \mathcal{T}$ 之间的夹角, 确保它们相互正交, 即:

$$\mathcal{L}_{\text{orth}} = \frac{1}{|\mathcal{T}|} \sum_{i=1}^{|\mathcal{T}|} \ln \sum_{j=1, j \neq i}^{|\mathcal{T}|} \exp\left(\frac{t_i \cdot t_j}{\tau_c}\right) \quad (4)$$

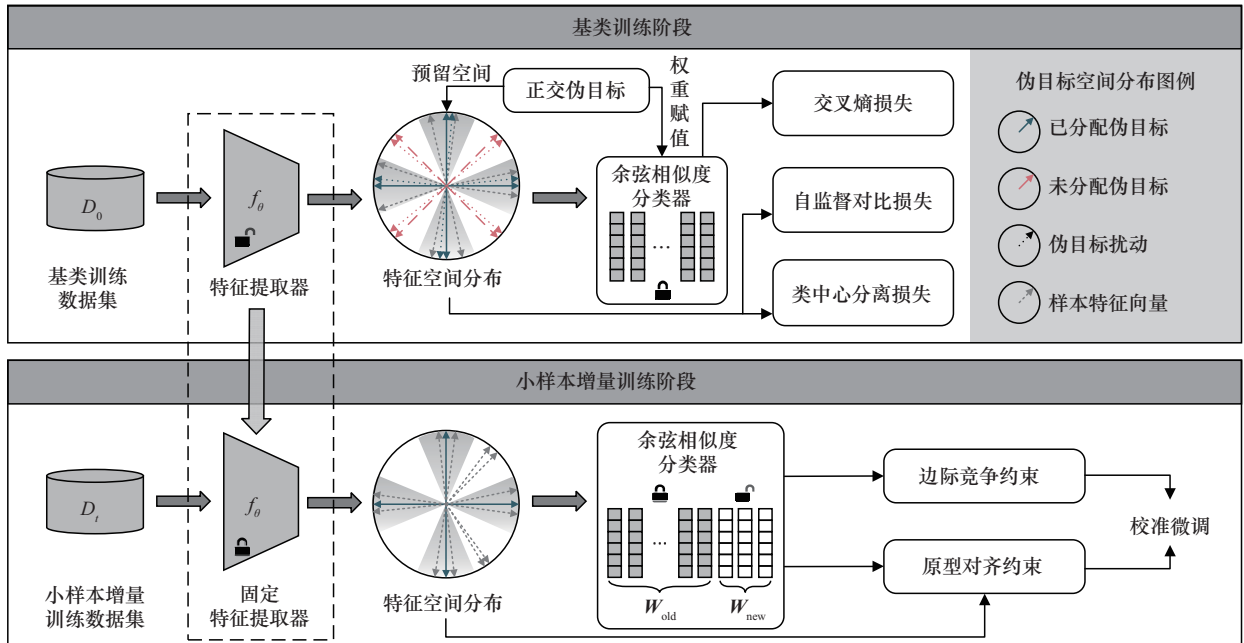


图2 基于正交空间约束的特定辐射源识别小样本类增量学习方法

其中, $|\mathcal{T}|$ 为集合 $\mathcal{T}$ 中向量数量, τ_c 为温度系数。为增强鲁棒性, 对每个伪目标施加均匀扰动, 生成扰动集合 $\tilde{\mathcal{T}}$ 为:

$$\tilde{\mathcal{T}} = \{\tilde{\mathbf{t}}_i\}_1^N \quad (5)$$

$$\tilde{\mathbf{t}}_i = \mathbf{t}_i + \boldsymbol{\varepsilon}_i \quad (6)$$

其中, $\boldsymbol{\varepsilon}_i \sim \mathcal{U}(-\lambda, \lambda)$, $\mathcal{U}$ 表示均匀分布, λ 表示扰动范围。通过一个固定的映射函数 $h: C_0 \rightarrow \{1, \dots, N\}$ 为每个基类 $c \in C_0$ 分配一个伪目标 $\mathbf{t}_{h(c)}$ 。因此, 伪目标集合可划分为两部分: 已分配伪目标集合 $\mathcal{T}_0 = \{\mathbf{t}_{h(c)} | c \in C_0\}$ 与未分配伪目标集合 $\mathcal{T}_u = \mathcal{T} \setminus \mathcal{T}_0$ 。其中, 已分配伪目标集合 $\mathcal{T}_0$ 是基类特征向量的优化方向, 而未分配伪目标集合 $\mathcal{T}_u$ 是为后续增量类预留的嵌入空间方向。

命题 1 伪目标数量 N 的理论范围应满足 $|C| \leq N \leq d + 1$, 其中 C 为全部类别的标签空间。上界由 d 维单位球面上最优均匀分离构型的最大可实现容量决定, 下界由后续增量类分离结构所需的最小几何支撑数量决定。

证明 为分析给定特征维度 d 下伪目标集合能达到的最大均匀分离容量, 定义伪目标集合 $\mathcal{T}$ 的全局分离度为:

$$\mathcal{D}(\mathcal{T}) = \sum_{i < j} (1 - \mathbf{t}_i^T \mathbf{t}_j) \quad (7)$$

其中, 每个伪目标均满足单位范数约束 $\|\mathbf{t}_i\|_2 = 1$ 。该指标度量所有伪目标两两之间的整体分离程度, $\mathcal{D}(\mathcal{T})$ 越大, 表面整组伪目标在单位球面上分布得越分散。对于伪目标数量 $\mathcal{T} = \{\mathbf{t}_i\}_1^N$, 由恒等式:

$$\left\| \sum_{i=1}^N \mathbf{t}_i \right\|_2^2 = \sum_{i=1}^N \|\mathbf{t}_i\|_2^2 + 2 \sum_{i < j} \mathbf{t}_i^T \mathbf{t}_j = N + 2 \sum_{i < j} \mathbf{t}_i^T \mathbf{t}_j \quad (8)$$

可得:

$$\sum_{i < j} \mathbf{t}_i^T \mathbf{t}_j = \frac{1}{2} \left\| \sum_{i=1}^N \mathbf{t}_i \right\|_2^2 - \frac{N}{2} \quad (9)$$

代入全局分离度 $\mathcal{D}(\mathcal{T})$ 有:

$$\mathcal{D}(\mathcal{T}) = \frac{N(N-1)}{2} - \sum_{i < j} \mathbf{t}_i^T \mathbf{t}_j = \frac{N^2}{2} - \frac{1}{2} \left\| \sum_{i=1}^N \mathbf{t}_i \right\|_2^2 \quad (10)$$

由于范数平方项恒非负, 因此有 $\mathcal{D}(\mathcal{T}) \leq \frac{N^2}{2}$,

当且仅当 $\sum_{i=1}^N \mathbf{t}_i = \mathbf{0}$ 时取等号, 即当伪目标集合的质心位于原点时, 其全局分离度最大, 即伪目标在单位球面上达到全局平衡。

然而, 仅有全局平衡仍不足以唯一确定最均匀的极限排布, 满足该条件的点集仍可能存在局部聚集或成分分离不均匀的现象。为了刻画最理想的均匀分离情形, 进一步引入等角对称性约束, 即要求任意两个不同伪目标之间满足相同的夹角关系, 即:

$$\mathbf{t}_i^T \mathbf{t}_j = \alpha, \quad i \neq j \quad (11)$$

将该条件与 $\sum_{i=1}^N \mathbf{t}_i = \mathbf{0}$ 联立, 对任意 k 有:

$$\mathbf{t}_k^T \sum_{i=1}^N \mathbf{t}_i = \mathbf{t}_k^T \mathbf{0} = 0 \quad (12)$$

结合 $\|\mathbf{t}_k\|_2^2 = 1$ 和 $\mathbf{t}_k^T \mathbf{t}_i = \alpha (i \neq k)$, 可得:

$$\alpha = -\frac{1}{N-1} \quad (13)$$

这表明在最优均匀分离构型条件下, 任意两个不同伪目标之间的内积均为 $-\frac{1}{N-1}$ 。引入 Gram 矩阵将伪目标之间的两两内积关系写成统一的矩阵形式为:

$$\mathbf{G} = [\mathbf{t}_i^T \mathbf{t}_j]_{i,j=1}^N \quad (14)$$

其中, 对角元素均为 1, 非对角元素均为 $-\frac{1}{N-1}$ 。该矩阵的秩满足 $\text{rank}(\mathbf{G}) = N - 1$, 由于 Gram 矩阵的秩不能超过向量所在空间的维数 $d^{[25]}$, 从而可得伪目标数量的理论上界为:

$$N \leq d + 1 \quad (15)$$

为推导伪目标数量的理论下界, 额外引入如下假设: 除已分配给基类的 $|C_0|$ 个伪目标外, 剩余的 $N - |C_0|$ 个伪目标不仅作为未使用方向存在, 而且必须能够表征或支撑后续 $|C| - |C_0|$ 个增量类的分离结构。在该假设下, 剩余伪目标数量至少不应少于后续增量类数量, 即:

$$N - |C_0| \geq |C| - |C_0| \quad (16)$$

伪目标数量的理论下界为:

$$N \geq |C| \quad (17)$$

需要指出的是, 该理论下界是从预留几何支撑

数量不少于后续增量类数量的角度得到的保守必要条件。它并不要求后续增量类数量在增量阶段与剩余伪目标逐一显式对齐，只要求剩余伪目标数量足以支撑后续增量类整体的分离结构。

2.3 基类训练

基类训练阶段结合伪目标约束学习特征提取器，使已知类别特征向对应方向聚集，同时为后续增量类别预留潜在嵌入方向。

(1) 交叉熵损失。为了将基类特征向量聚集到对应的伪目标方向，本文将基类的伪目标作为余弦相似度分类器权重，即 $\mathbf{W}_0 = [\mathbf{t}_{h(1)}, \dots, \mathbf{t}_{h(C_0)}]$ ，并且分类器权重在训练过程中保持不变。交叉熵损失函数定义为：

$$\mathcal{L}_{ce} = - \sum_{(\mathbf{x}_i, y_i) \in D_0} \ln \frac{\exp(\mathbf{z}_i^T \mathbf{t}_{h(y_i)})}{\sum_{t_k \in \mathcal{T}_0} \exp(\mathbf{z}_i^T \mathbf{t}_k)} \quad (18)$$

其中， $\mathbf{z}_i = f_\theta(\mathbf{x}_i) \in \mathbb{R}^d$ 为信号 $\mathbf{x}_i$ 的特征向量。通过优化特征提取器的参数，在嵌入空间中形成结构良好的类簇分布，如图 3(a) 所示。

(2) 自监督对比损失。为了对预留的嵌入方向施加显式约束，本文定义了自监督对比学习的锚点和正负样本集，并设置两类锚点：一类为样本特征向量，另一类为未分配伪目标。采用样本特征向量作为锚点，能够将同类样本和对应的伪目标聚集在一起，增强类内紧凑性，并将其与其他类样本分离，增强类间可分性，如图 3(b) 中的三角形所示。采用未分配伪目标作为锚点，能够引导基类的特征向量远离预留的嵌入方向，从而为后续新增类别保留明确的嵌入空间，有效缓解增量学习中新旧类别的类间表征挤压问题，如图 3(b) 中的五角星所示。

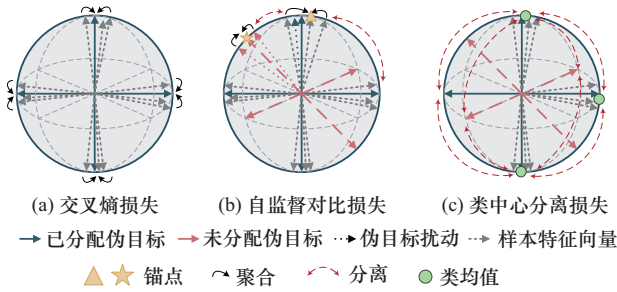


图 3 基础阶段损失函数优化目标

具体而言，当样本特征 $\mathbf{z}_i$ 作为锚点时，正集包含同类样本的特征向量，同时将该类对应的伪目标

向量及其扰动版本纳入正集中，即：

$$\tilde{P}_i = P_i \cup \mathbf{t}_{h(y_i)} \cup \tilde{\mathbf{t}}_{h(y_i)} \quad (19)$$

其中， $P_i = \{\mathbf{z}_j | (\mathbf{x}_j, y_j) \in D_0, y_j = y_i\}$ ，负集包含不同类样本的特征向量，即：

$$N_i = \{\mathbf{z}_j | (\mathbf{x}_j, y_j) \in D_0, y_j \neq y_i\} \quad (20)$$

当未分配伪目标 $\mathbf{t}_i \in \mathcal{T}_u$ 作为锚点时，正集包含自身扰动，以此强化该预留方向的表示，即：

$$\tilde{P}_i = \{\tilde{\mathbf{t}}_i\} \quad (21)$$

负集包含所有基类样本特征向量、已分配伪目标向量及其扰动，旨在推动基类样本远离该方向，即：

$$N_i = \left(\{\mathbf{z}_j | (\mathbf{x}_j, y_j) \in D_0\} \cup \mathcal{T}_0 \cup \tilde{\mathcal{T}}_0 \right) \setminus \{\mathbf{t}_i \cup \tilde{\mathbf{t}}_i\} \quad (22)$$

自监督对比损失函数定义为：

$$\mathcal{L}_s = - \sum_{a \in \mathcal{A}} \frac{1}{|\tilde{P}_a|} \sum_{p \in P_a} \ln \frac{\exp\left(\frac{\mathbf{a}^T \mathbf{p}}{\tau_s}\right)}{\sum_{n \in N_a} \exp\left(\frac{\mathbf{a}^T \mathbf{n}}{\tau_s}\right)} \quad (23)$$

其中， $\mathcal{A}$ 表示锚点集合， τ_s 表示温度系数， $\tilde{P}_a$ 和 N_a 分别表示锚点的正/负集合。

(3) 类中心分离损失。该损失从类中心层面施加约束，通过促使基类的类中心与未分配伪目标向量趋近正交，有效拉大不同类别中心间的距离，如图 3(c) 所示。具体而言，对于一个训练批次 B ，首先计算该批次中每个类别出现的特征均值，作为其临时的类中心，即：

$$\mu_c = \frac{1}{|B_c|} \sum_{i=1}^{|B_c|} f_\theta(\mathbf{x}_i), c \in C_0^B \quad (24)$$

其中， C_0^B 表示训练批次 B 中出现的类别集合， B_c 表示其中属于类别 c 的样本子集。对于基类中未出现在当前训练批次的类别，使用其预先分配的伪目标向量 $\mathbf{t}_{h(c)}$ 作为其中心代理。结合未分配伪目标 $\mathcal{T}_u$ 构造类中心集合，即：

$$M = \{U_B \cup \mathcal{T}_B \cup \mathcal{T}_u\} \quad (25)$$

$$U_B = \{\mu_c | c \in C_0^B\} \quad (26)$$

$$\mathcal{T}_B = \{\mathbf{t}_{h(c)} | c \in C_0 \setminus C_0^B\} \quad (27)$$

类中心分离损失定义为：

$$\mathcal{L}_c = \frac{1}{|M|} \sum_{m_i \in M} \ln \sum_{m_j \in M, j \neq i} \exp\left(\frac{m_i m_j}{\tau_c}\right) \quad (28)$$

其中, τ_c 为温度系数, 该损失函数通过惩罚任意两个类中心之间的高相似度, 驱动所有已见类中心和未分配伪目标在特征空间中趋向正交。最终, 基类训练阶段的总体损失由交叉熵损失、自监督对比损失和类中心分离损失共同构成, 且各损失项的权重相同。

$$\mathcal{L}_{\text{init}} = \mathcal{L}_{\text{ce}} + \mathcal{L}_s + \mathcal{L}_c \quad (29)$$

基类训练阶段的伪代码如算法1所示。

算法1 基类训练阶段

输入 初始阶段训练集 D_0 , 基类集合 C_0 , 训练轮数 E_{base} , 学习率 l_{base} , 伪目标数量 N , 扰动范围 λ , 温度系数 τ_s 和 τ_c

输出 特征提取器参数 θ , 分类器权重矩阵 W_0

- 1) 随机初始化参数
- 2) 根据式(4)生成相互正交的伪目标向量 $T = \{\mathbf{t}_i\}_1^N \subset \mathbb{R}^d$ 及其扰动
- 3) 给每个基类 $c \in C_0$ 分配伪目标 $\mathbf{t}_{h(c)}$, 作为分类器权重矩阵 W_0 并不再更新
- 4) for $e = 1, 2, \dots, E_{\text{base}}$ do
- 5) for $B \subset D_0$ do
- 6) 对 $(\mathbf{x}_i, y_i) \in B$ 提取样本特征向量 $\mathbf{z}_i = f_{\theta}(\mathbf{x}_i)$
- 7) 根据式(18)计算交叉熵损失 $\mathcal{L}_{\text{ce}}$
- 8) 根据式(23)计算自监督对比损失 $\mathcal{L}_s$
- 9) 根据式(28)计算类中心分离损失 $\mathcal{L}_c$
- 10) 计算总损失 $\mathcal{L}_{\text{init}} = \mathcal{L}_{\text{ce}} + \mathcal{L}_s + \mathcal{L}_c$
- 11) $\theta \leftarrow \theta - l_{\text{base}} \nabla \mathcal{L}_{\text{init}}$
- 12) end for
- 13) end for

2.4 小样本增量训练

为了在学习新类别的同时, 有效避免因特征表示漂移而引发的灾难性遗忘, 并缓解因样本数量极少而导致的模型过拟合, 本文采用固定特征提取器参数的策略, 仅对分类器中新增类别的权重向量进行学习和更新。分类器权重矩阵可表示为新旧两部分的拼接, 即:

$$W = \text{Concat} \left(\underbrace{W_{0:t-1}}_{W_{\text{old}}}, \underbrace{[w_0^t, \dots, w_{C_t}^t]}_{W_{\text{new}}} \right) \quad (30)$$

新类权重通常使用其类均值进行初始化, 而在

小样本条件下, 由此估计得到的权重向量可能存在偏差, 易导致新类被误分为旧类。因此, 本文量化新类别样本在判别过程中存在的误判风险, 从而引导新类权重向更具判别性的方向校准。

具体而言, 对于每一个新类 $l \in C_t$, 计算其样本 $\mathbf{x}_i$ 在其他类别上的得分, 并从中选取得分最高的 Top- k 个类别作为竞争类。本文引入一个温度系数 τ_{fuse} , 对竞争类上的得分进行平滑处理, 即:

$$s_{\text{comp}}^i = \tau_{\text{fuse}} \ln \sum_{u \in \text{Top-}k} \exp\left(\frac{u}{\tau_{\text{fuse}}}\right) \quad (31)$$

其中, $u = \mathbf{z}_i^T \mathbf{w}$ 为样本 $\mathbf{x}_i$ 在竞争类上的得分, $\mathbf{z}_i = f_{\theta}(\mathbf{x}_i)$ 为样本 $\mathbf{x}_i$ 的特征向量, $\mathbf{w}$ 为竞争类对应的分类器权重向量。同时, 计算样本 $\mathbf{x}_i$ 在其真实类别 l 上的得分为:

$$s_{\text{self}}^i = \mathbf{z}_i^T \mathbf{w}_l^t \quad (32)$$

其中, $\mathbf{w}_l^t$ 表示 l 类的分类器权重向量。通过真实类别与竞争类别得分之间的相对关系量化误判风险。若样本真实类别上的得分 s_{self}^i 与竞争类别上的得分 s_{comp}^i 的差值小于预设间隔 q , 则该样本处于决策边界冲突区域, 因而将其判定为困难样本用于权重校准。困难样本集合表示为:

$$\mathcal{H} = \{i | s_{\text{comp}}^i - s_{\text{self}}^i + q > 0\} \quad (33)$$

基于困难样本集合构建边际竞争约束校准新类权重, 旨在通过优化扩大新类与竞争类之间的分类间隔, 即:

$$\mathcal{L}_h = \frac{1}{|\mathcal{H}|} \sum_{i \in \mathcal{H}} (s_{\text{comp}}^i - s_{\text{self}}^i + q) \quad (34)$$

此外, 为避免权重更新偏离语义中心, 额外引入原型对齐约束, 使校准的新类权重与其类原型方向保持一致, 即:

$$\mathcal{L}_a = 1 - \boldsymbol{\mu}_l^T \mathbf{w}_l^t \quad (35)$$

其中, $\boldsymbol{\mu}_l$ 表示 l 类的原型向量。综上所述, 新类分类器权重的总体优化目标由上述两项损失函数共同构成, 其中原型对齐约束损失权重 λ_a 设置为 1.6, 即:

$$\mathcal{L}_{\text{inc}} = \mathcal{L}_h + \lambda_a \mathcal{L}_a \quad (36)$$

小样本增量训练阶段的伪代码如算法2所示。

算法2 小样本增量训练阶段

输入 增量阶段训练集 D_t , 增量类集合 C_t , 训练轮数 E_{inc} , 学习率 l_{inc} , 冻结的特征提取器 f_{θ} , 旧

类分类器权重 $\mathbf{W}_{\text{old}}$, Top- k 数量, 预设间隔 q , 温度系数 τ_{fuse} , 损失权重 λ_a

输出 分类器权重 $\mathbf{W} = [\mathbf{W}_{\text{old}}, \mathbf{W}_{\text{new}}]$

- 1) 计算新类别 $l \in C_t$ 的类原型向量, 初始化新类分类器权重 $\mathbf{W}_{\text{new}}$
- 2) 拼接 $\mathbf{W}_{\text{old}}$ 与 $\mathbf{W}_{\text{new}}$ 得到扩展的分类器权重 $\mathbf{W} = [\mathbf{W}_{\text{old}}, \mathbf{W}_{\text{new}}]$
- 3) for $l \in C_t$
- 4) for $e = 1, 2, \dots, E_{\text{inc}}$
- 5) 根据式(31)~式(34)构建困难样本集合 $\mathcal{H}$, 并计算边际竞争损失 $\mathcal{L}_h$
- 6) 根据式(35)计算原型对齐约束 $\mathcal{L}_a$
- 7) 计算总损失 $\mathcal{L}_{\text{inc}} = \mathcal{L}_h + \lambda_a \mathcal{L}_a$
- 8) $\mathbf{W}_{\text{new}} \leftarrow \mathbf{W}_{\text{new}} - l_{\text{inc}} \nabla \mathcal{L}_{\text{inc}}$
- 9) end for
- 10) end for

2.5 复杂度分析

设特征嵌入维度为 d , 批大小为 B , 伪目标数量为 N , 当前已学习类别数为 $|C|$ 。特征提取器单次前向与反向传播的复杂度分别表示为 $\mathcal{F}_{\text{fwd}}$ 和 $\mathcal{F}_{\text{bwd}}$ 。

正交伪目标集合的构造通过迭代优化实现, 每次迭代的核心计算涉及 N 个伪目标间的两两相似度计算, 其复杂度约为 $O(N^2 d)$ 。若迭代次数为 I , 则总复杂度为 $O(IN^2 d)$ 。由于该过程仅在训练前执行一次且不参与模型推理, 因此不会引入在线推理负担。

在基类训练阶段, 计算开销主要由特征提取开销和正交约束开销组成。特征提取开销源于特征提取器的前向与反向传播, 复杂度为 $O(B(\mathcal{F}_{\text{fwd}} + \mathcal{F}_{\text{bwd}}))$ 。正交约束开销主要涉及 B 个特征向量与 N 个伪目标向量的相似度度量, 复杂度为 $O(BNd)$ 。因此, 基类训练阶段的总体复杂度近似为 $O(B(\mathcal{F}_{\text{fwd}} + \mathcal{F}_{\text{bwd}}) + BNd)$ 。

在小样本增量训练阶段, 特征提取器被固定, 仅更新分类器权重, 计算开销主要集中于前向传播过程及分类层的参数更新。前向传播的复杂度为 $O(B\mathcal{F}_{\text{fwd}})$, 对当前包含旧类与新类在内的 $|C|$ 个类别进行余弦相似度计算, 复杂度为 $O(B|C|d)$ 。由于本文提出的校准策略及 Top- k 竞争类筛选带来的额外计算复杂度远小于全类评分计算项, 且不改变

整体复杂度阶数。因此, 小样本增量训练阶段的复杂度可近似为 $O(B\mathcal{F}_{\text{fwd}} + B|C|d)$ 。

在推理阶段, 单样本仅需一次特征前向传播及对 $|C|$ 个类别的余弦相似度打分, 其复杂度为 $O(\mathcal{F}_{\text{fwd}} + |C|d)$, 与当前类别数呈线性相关。由于本文方法未引入额外的网络分支或复杂的辅助结构, 推理复杂度与采用相同主干网络的对比方法保持同阶。

3 仿真分析

3.1 实验数据和参数设置

为验证所提方法在 SEI 任务中的有效性, 本文基于真实 ADS-B 和 Wi-Fi 信号构建实验数据集。在基类训练阶段, 将每个类别 80% 的样本划分为训练集, 其余 20% 划分为测试集。在小样本增量训练阶段, 每个新类别仅选取 5 个样本用于学习, 其余样本作为测试集。所有实验均在 Python 3.8.19 环境下实现, 并基于 NVIDIA GeForce RTX 3090 GPU 进行训练与测试。

(1) ADS-B 数据集。本文采用公开可用的 ADS-B 数据集^[26]进行实验。原始信号由软件无线电接收机 USRP B210 在 1 090 MHz 频段以 8 MHz 采样率采集获得, 截取其前 1 024 个复数采样点作为单条样本。为保证各类别具有充足的样本数量, 本文首先筛选出样本数不少于 500 条的无线应答器, 在此基础上进一步选取出现频次最高的前 100 个无线应答器, 构建最终的 ADS-B 数据集。

(2) Wi-Fi 数据集。本文采用公开可用的 Wi-Fi 数据集^[27]进行实验。该数据集包含来自 174 个 Wi-Fi 发射器的信号样本, 原始信号由多台软件无线电接收机在 2 462 MHz 频段以 25 MHz 采样率采集获得, 每条信号样本长度为 256。考虑到不同接收机采集条件差异可能引入域偏移, 从而对模型性能评估造成干扰, 本文仅保留由同一接收机采集且样本数不少于 50 条的发射器, 最终筛选出 130 个 Wi-Fi 发射器构建实验数据集。

本文采用的实验参数设置如表 1 所示, 其中基类训练阶段的学习率设置为 0.01, 小样本增量训练阶段的学习率设置为 0.08。其原因在于, 基类训练需在多种损失约束下稳定地构建具有良好可扩展性的特征空间结构, 因此采用较小的学习率以避免特征表示的剧烈波动。小样本增量训练主要针对参数规模较小且目标明确的分类器权重进行优化, 用于

快速修正新旧类别间的判别偏移,因而采用较大的学习率以提高决策边界调整效率。

表1 仿真参数

仿真参数	基类训练阶段	小样本增量训练阶段
训练轮数/轮	100	50
训练批大小	128	128
优化器	SGD	SGD
学习率	0.01	0.08
早停轮数/轮	4	4

3.2 评估指标

将完成第 t 个任务后的Top-1准确率记为 A_t , A_t 值越高表明预测准确率越高。由于模型会不断更新,随着学习的任务增多准确率往往会下降。因此,最后一个阶段的准确率 A_T 是衡量所有类别整体准确率的一个合适指标。由于基类类别数量占比较高,准确率易被基类性能主导,从而难以真实反映新类识别水平。为缓解这一偏置,调和准确率已被广泛采用为FSCIL的一种可靠评估指标^[28],将完成第 t 个任务后的调和准确率记为 H_t 。

$$H_t = \frac{2A_{\text{old}}^t A_{\text{new}}^t}{A_{\text{old}}^t + A_{\text{new}}^t} \quad (37)$$

其中, A_{old}^t 表示所有旧类上的准确率,包括初始任务和先前所有增量任务上的所有类别, A_{new}^t 表示所有新类上的准确率,包括当前任务上的所有类别。此外,为评估模型在增量学习过程中的遗忘程度,将完成第 t 个任务后的遗忘率记为 F_t 。

$$F_t = \frac{1}{t} \sum_{k=0}^{t-1} \left(\max_{l \in \{0, \dots, t-1\}} A_{l,k} - A_{t,k} \right) \quad (38)$$

其中, $A_{l,k}$ 为完成第 l 个任务后在先前的第 k 个任务

上的准确率,且 $l > k$,遗忘率越低表示对先前任务的遗忘更少^[29]。然而,仅仅比较单个任务上的指标忽略了增量学习过程中性能的演变情况。因此,本文通过平均准确率、平均调和准确率和平均遗忘率来评估所有增量阶段的表现。

$$\bar{A} = \frac{1}{T} \sum_{t=1}^T A_t \quad (39)$$

$$\bar{H} = \frac{1}{T} \sum_{t=1}^T H_t \quad (40)$$

$$\bar{F} = \frac{1}{T} \sum_{t=1}^T F_t \quad (41)$$

3.3 实验结果

3.3.1 不同方法的比较分析

本节将本文方法与多种先进方法进行了对比,包括SEI领域的FSCIL方法OFSCIL^[17]、Meta-RFF^[18]和FSCIL-SEI^[19],以及图像分类领域的FSCIL方法Closer^[20]、ADBS^[22]和DSN^[23]。对于FI-OSR任务的相关方法OFSCIL和Meta-RFF,本文保留其小样本类增量学习部分用于比较。为了确保与基准方法的公平比较,本文在相同数据集上实现并重新训练了比较的方法。所有比较的方法都通过超参数网格搜索来确定最优参数,从而优化基线方法的性能。

(1) 基础设置下的性能比较

在ADS-B和Wi-Fi数据集中,初始训练阶段各有60个类别,随后每个增量阶段分别增加5个和10个类别进行训练,每类有5个训练样本。表2和表3展示了两种数据集上不同方法在准确率方面的性能比较,评估了基类任务和各增量任务完成后的Top-1准确率 A_t 、平均准确率 $\bar{A}$ 和平均调和准确率 $\bar{H}$ 。可以观察到,随着增量任务数量的增加,

表2 在ADS-B数据集上5类-5样本设置下的准确率对比

方法	A_t									$\bar{A}$	$\bar{H}$
	0	1	2	3	4	5	6	7	8		
FSCIL-SEI	99.73%	97.73%	96.02%	93.65%	90.42%	88.56%	87.19%	84.80%	83.00%	90.17%	86.02%
OFSCIL	99.98%	98.92%	98.32%	98.10%	96.95%	96.27%	95.69%	95.03%	94.50%	96.72%	92.43%
Meta-RFF	97.875%	96.558%	95.554%	94.867%	94.141%	93.412%	92.556%	92.355%	92.575%	94.00%	93.15%
DSN	99.94%	97.81%	97.13%	96.07%	94.52%	93.93%	92.56%	92.29%	91.88%	94.52%	88.52%
Closer	98.46%	97.67%	97.02%	95.17%	94.56%	94.53%	94.31%	92.28%	90.61%	94.52%	90.61%
ADBS	100%	98.98%	98.20%	97.57%	97.25%	95.13%	94.49%	94.55%	94.68%	96.36%	97.60%
本文方法	99.77%	98.62%	98.48%	98.15%	97.81%	97.82%	97.83%	97.42%	97.06%	97.90%	98.11%

表3 在Wi-Fi数据集上10类-5样本设置下的准确率对比

方法	A_t								$\bar{A}$	$\bar{H}$
	0	1	2	3	4	5	6	7		
FSCIL-SEI	93.92%	79.06%	70.66%	62.71%	60.49%	56.61%	55.15%	52.54%	62.46%	70.56%
OFSCIL	93.75%	90.88%	87.02%	87.35%	86.83%	87.16%	86.28%	84.39%	87.13%	84.41%
Meta-RFF	83.78%	80.48%	81.40%	81.30%	78.83%	78.07%	76.57%	75.08%	78.82%	76.40%
DSN	94.60%	90.60%	87.89%	88.24%	86.63%	86.79%	86.19%	84.92%	87.32%	82.65%
Closer	86.66%	83.76%	82.77%	82.20%	80.14%	79.27%	80.17%	79.69%	81.14%	79.68%
ADBS	95.61%	93.88%	90.01%	88.91%	85.82%	85.51%	84.77%	84.23%	87.59%	88.30%
本文方法	95.10%	94.02%	93.26%	93.06%	92.10%	90.28%	89.62%	88.62%	91.56%	88.70%

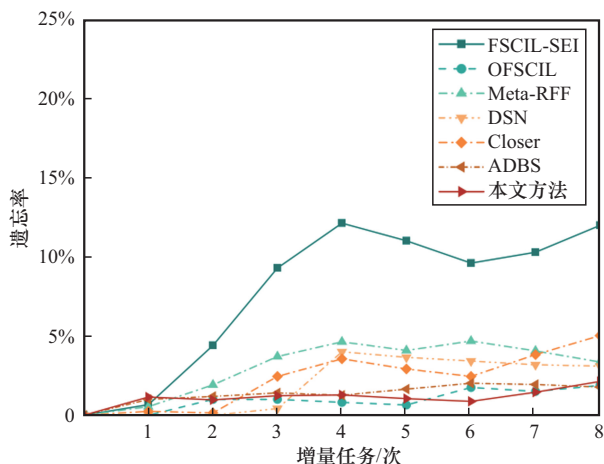
Top-1 准确率呈下降趋势。这主要由于增量学习过程中旧类不可避免地会被遗忘，而旧类在类别集合中的占比不断增大，从而使Top-1 准确率随增量任务推进而降低。具体而言，FSCIL-SEI的准确率出现显著下降，相比之下OFSCIL、DSN和ADBS的性能下降缓慢，而本文方法在各个增量阶段均保持了较高的准确率，并且在增量后期性能优势更加显著。Meta-RFF和Closer方法在Wi-Fi数据集上的初始任务学习效果不佳，这在一定程度上限制了其后续增量阶段的性能表现，进而导致其在后续新类别学习中的准确率持续偏低。本文方法在两种数据集上的平均准确率分别达到97.90%和91.56%，相较于最佳的对比方法分别提升1.18%和3.97%，表现出良好的性能优势。

为了探究各方法的抗灾难性遗忘能力，图4展示了在两种数据集上的遗忘率曲线。从图4可以观察到，随着增量任务数量的增加，各方法的遗忘率呈上升趋势，表明在持续引入新类的过程中普遍存在一定程度的灾难性遗忘现象。与现有方法相比，

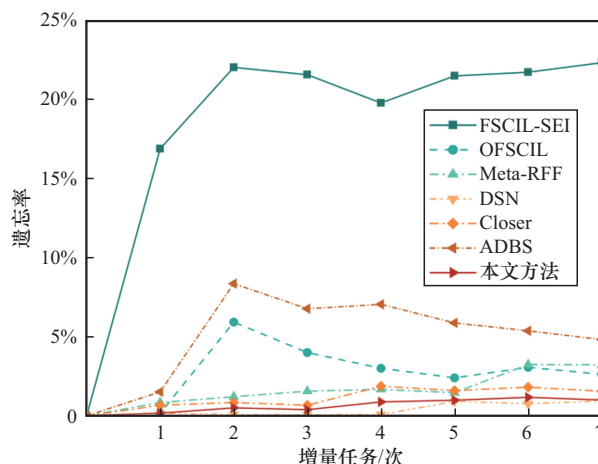
本文方法在整个增量过程中始终保持较低的遗忘率。此外，尽管本文方法的遗忘率略高于OFSCIL和DSN，但由于每次学习新类时能获得更高的初始准确率，即使遗忘率较大，最终的准确率仍高于OFSCIL和DSN。

(2) 不同初始类别数量下的性能比较

本文将初始类别数量设置为40、50、60和70，并在每次增量训练中ADS-B数据引入5类新类别，Wi-Fi数据引入10类新类别，直至完成全部类别的学习。不同初始类别数量设置下各方法的平均准确率比较结果如图5所示。从图5可以看出，FSCIL-SEI在Wi-Fi数据集上的平均准确率明显较低，这可能是由于该方法依赖于新类的原型增强更新分类器。Wi-Fi信号样本较短、有效判别信息相对有限，少量样本构建的类原型更易产生偏差，而数据增强过程可能进一步放大该偏差，并将其传递至分类边界，从而导致识别性能下降。总体来看，随着初始类别数量的增加，各种FSCIL方法的平均准确率均呈上升趋势，这是因为初始阶段包含的类别越多，



(a) 在ADS-B数据集上5类-5样本设置下的遗忘率比较



(b) 在Wi-Fi数据集上10类-5样本设置下的遗忘率比较

图4 不同数据集上的遗忘率比较

模型越能够学习到充分的类间结构信息,从而为后续小样本增量学习提供更好的表征基础。相比之下,本文方法在ADS-B和Wi-Fi数据集上的所有初始类别数量设置下均取得了最高的平均准确率。

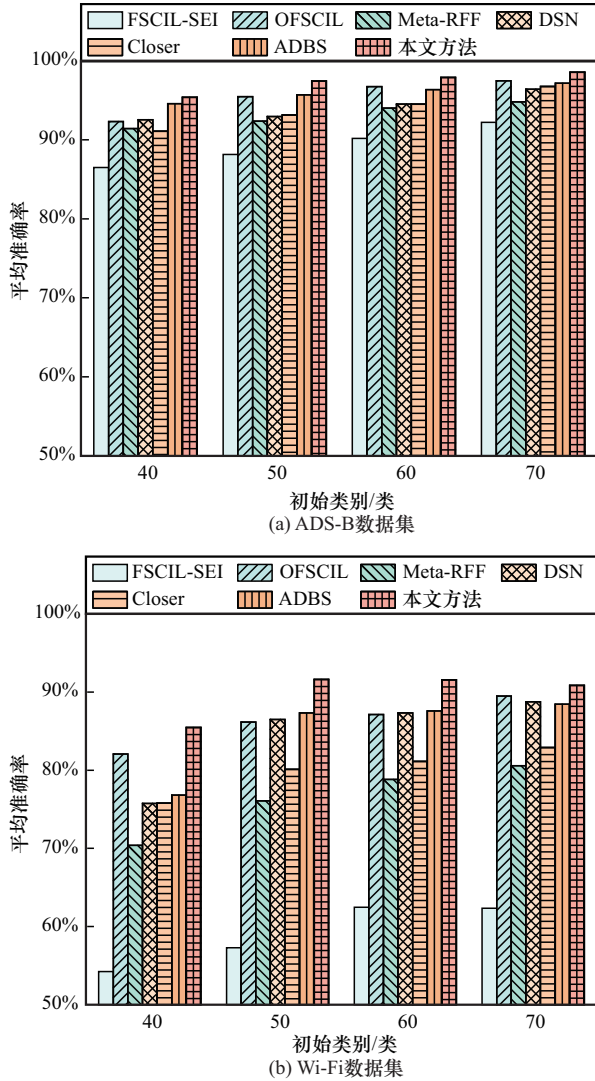


图5 不同初始类别数量设置下的平均准确率比较

(3) 不同增量类别数量下的性能比较

本文将初始类别数量设定为60个,并在后续增量阶段依次引入不同数量的新类别,直至完成全部类别的学习。图6给出了不同增量类别数量设置下各类FSCIL方法的平均准确率比较结果。从图6可以看出,本文方法在所有增量类别数量设置下的平均准确率均优于其他对比方法。此外,随着增量类别数量增加,各方法的平均准确率整体呈上升趋势。这是因为在单次增量中引入更多类别,有助于模型同时感知更丰富的类间关系,从而更充分地优

化分类边界。相比于将这些类别分散到多个任务中逐步学习,这种方式能够减少多阶段增量带来的误差累积,因此有利于提升整体性能。

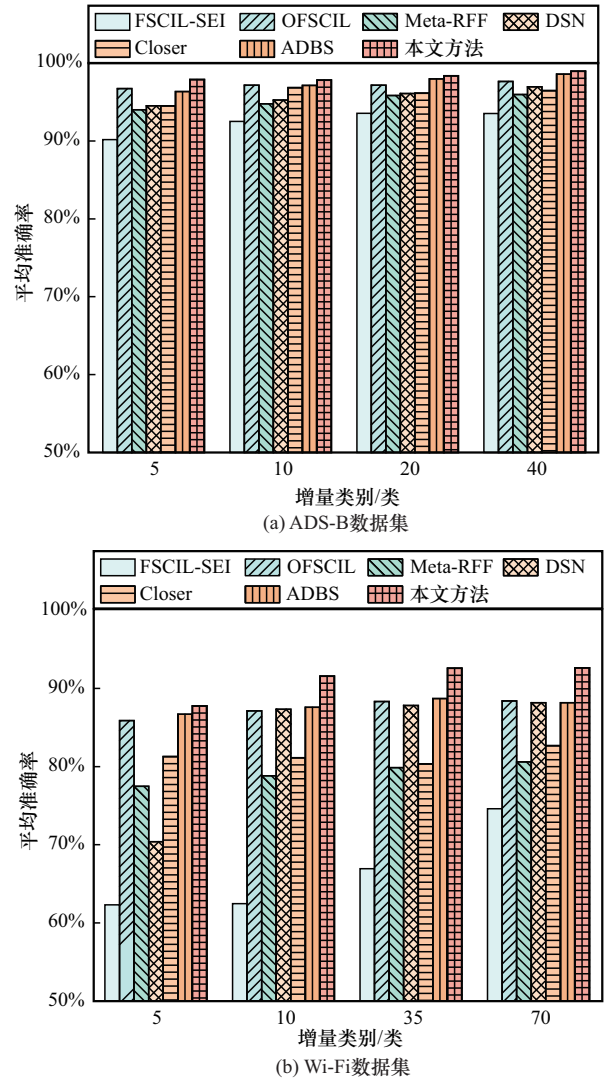


图6 不同增量类别数量设置下的平均准确率比较

(4) 单样本设置下的性能比较

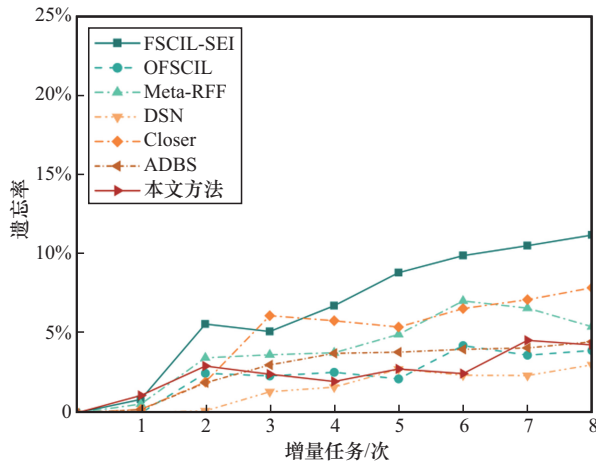
本文进一步考虑了每个增量类别仅包含一个训练样本的极端情况。表4和表5给出了不同FSCIL方法在准确率方面的比较结果。从表4和表5可以看出,随着增量阶段的推进,各方法的准确率整体呈下降趋势。尽管如此,本文方法的准确率始终保持在93%以上,在两种数据集的最终增量阶段分别较表现最优的对比方法提高了1.73%和2.92%。图7进一步展示了不同FSCIL方法在单样本设置下的遗忘率比较结果。可以看出,本文方法的遗忘率整体较低,尤其在样本长度较短的Wi-Fi数据集上

表 4 在 ADS-B 数据集上 5 类-1 样本设置下的准确率对比

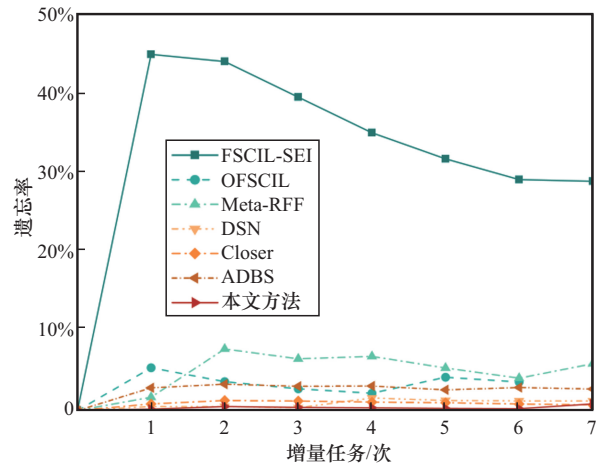
方法	A_t									$\bar{A}$	$\bar{H}$
	0	1	2	3	4	5	6	7	8		
FSCIL-SEI	99.73%	95.19%	92.07%	87.93%	84.23%	79.94%	76.13%	73.25%	71.30%	82.51%	63.37%
OFSCIL	99.98%	96.98%	95.41%	94.80%	93.69%	92.47%	90.96%	90.18%	88.98%	92.93%	84.05%
Meta-RFF	98.42%	95.69%	93.98%	93.48%	91.92%	88.84%	86.97%	85.54%	84.73%	90.14%	80.67%
DSN	99.94%	97.69%	94.98%	93.20%	92.06%	91.32%	89.90%	89.20%	88.94%	92.16%	82.99%
Closer	98.46%	96.19%	95.46%	91.85%	90.86%	89.41%	88.03%	85.82%	84.51%	90.27%	81.83%
ADBS	100%	98.25%	97.70%	96.93%	95.95%	94.03%	92.39%	92.28%	91.71%	94.90%	92.56%
本文方法	99.77%	97.58%	96.96%	96.40%	95.67%	94.77%	94.47%	93.47%	93.44%	95.35%	93.08%

表 5 在 Wi-Fi 数据集上 10 类-1 样本设置下的准确率对比

方法	A_t								$\bar{A}$	$\bar{H}$
	0	1	2	3	4	5	6	7		
FSCIL-SEI	93.41%	54.84%	39.33%	33.37%	34.75%	35.32%	37.66%	38.54%	39.12%	48.04%
OFSCIL	93.41%	92.45%	87.52%	87.79%	87.34%	86.88%	84.60%	83.23%	87.12%	84.01%
Meta-RFF	85.14%	82.62%	78.78%	78.61%	73.56%	71.65%	70.46%	67.77%	74.78%	68.64%
DSN	94.60%	91.88%	89.39%	88.91%	87.54%	87.25%	85.69%	84.85%	87.93%	83.67%
Closer	86.66%	84.90%	83.90%	84.43%	82.37%	82.02%	82.26%	82.39%	83.18%	82.63%
ADBS	95.61%	92.31%	90.01%	88.91%	85.11%	84.59%	83.18%	83.23%	86.76%	86.26%
本文方法	95.10%	93.73%	92.39%	91.71%	90.48%	90.00%	88.96%	87.77%	90.72%	87.38%



(a) 在 ADS-B 数据集上 5 类-1 样本设置下的遗忘率比较



(b) 在 Wi-Fi 数据集上 10 类-1 样本设置下的遗忘率比较

图 7 单样本设置下的遗忘率比较

取得了最低的遗忘率。上述结果表明, 本文方法能够在极端小样本增量条件下有效学习新知识并抑制遗忘, 体现出较强的适应能力。此外, 本文注意到 FSCIL-SEI 在 Wi-Fi 数据集上的遗忘率曲线呈现出相反的变化趋势。其原因可能在于, 在第一个增量任务中, 该方法对基类的识别准确率出现了大幅度下降, 从而导致初始阶段遗忘率较高。随着增量任务数量

的增加, 早期较大的遗忘被分摊到更多增量阶段中, 整体遗忘率反而呈现下降趋势。

3.3.2 消融实验

为验证各模块的有效性, 本文在 ADS-B 数据集上进行了消融实验, 该实验包含 4 个核心设计: 伪目标策略、自监督对比损失 $\mathcal{L}_s$ 、类中心分离损失 $\mathcal{L}_c$ 和校准策略。实验设置初始基类为 60 类, 连

续进行8次增量,每次增加5个新类,每类仅提供5个训练样本。表6展示了各模块对性能的影响,包括在完成最后一个增量任务 $T=8$ 后,模型在基类任务及各增量任务上的识别准确率,以及平均准确率 $\bar{A}$ 、平均调和准确率 $\bar{H}$ 和平均遗忘率 $\bar{F}$ 这3个综合指标。

(1) 伪目标策略:将基类任务的识别准确率 A_{base}^T 从7.25%提升至56.42%,并将平均遗忘率降低了12.56%,证明其能够有效为未来增量任务预留特征空间,显著缓解了新类与基类间的混淆。

(2) 自监督对比损失:将基类任务的识别准确率 A_{base}^T 从56.42%进一步提升至99.88%,并将平均遗忘率进一步降低了5.23%,证明其有效增强了特征空间中的类内聚合性与类间分离性。

(3) 类中心分离损失:使所有增量任务的识别准确率均有明显提升,其中最后一个增量任务的识别准确率从82.50%提升至98.00%,证明其通过约束新类特征空间与基类特征空间趋于正交,有效缓解了将新类误分为基类的风险。

(4) 校准策略:进一步提升了各增量任务的准确率,并将最后一个增量任务的准确率从98.00%提升至98.75%,证明其有效优化了判别边界。

3.3.3 超参数敏感性分析

模型的性能主要受以下参数的影响:伪目标数量 N 、扰动范围 λ 、Top- k 数量和温度系数 τ_{fuse} 。图8展示了在ADS-B数据集上不同超参数对模型性能的影响,具体分析如下。

(1) 当伪目标数量从160个增至240个时,性能呈先升后降趋势,在伪目标数量为200个时达到最优。当伪目标数量较少时,由于未分配伪目标在单位超球面上的方向覆盖不足,新类难以匹配合适的伪目标向量,准确率较低;当伪目标数量较多时,由于伪目标之间的最小夹角被压缩,类间混淆

加剧了性能回落。

(2) 当扰动范围从0增至0.040时,性能呈现先升后降趋势,在扰动范围为0.010时达到最优。适度增加扰动有助于增强模型的鲁棒性,提升新类别的学习能力。然而过大的扰动会掩盖伪目标原有的类别信息,导致新旧类的空间可分性变差,从而导致性能下降。

(3) 当Top- k 数量从20个增加到60个时,性能保持稳定,最佳性能出现在60个时。实验表明模型对Top- k 数量变化具有较强的鲁棒性,不会因Top- k 数量的变化而显著影响性能,因此可选择较少的Top- k 值,既不会大幅影响性能,又可以降低计算复杂度。

(4) 当温度系数从0.001增加到0.500时,性能保持稳定,最佳性能出现在0.010时。温度系数在增量校准中用于调整新旧类别之间的表征平衡,实验结果表明,温度系数对模型有较好的鲁棒性。

3.3.4 识别效果可视化展示

图9展示了ADS-B数据集上不同方法的混淆矩阵。实验设置初始基类为60类,随后进行8次增量,每次增加5个新类,每类仅提供5个训练样本。图9中显示完成所有增量任务后,模型在100个类上的混淆矩阵。可以看出,OFSCIL与Closer在增量学习中存在新旧类混淆的问题,容易将新类误判为基类。本文方法通过引入伪目标策略,有效区分了新旧类别的特征空间,显著缓解了此类混淆现象。如图9(c)所示,本文方法的预测结果集中分布于对角线,几乎消除了跨类误判,表明其实现了显著更优的性能。

3.3.5 复杂度比较

在相同主干网络与硬件环境下,表7对比了在基类训练、增量训练和推理3个阶段的时间消耗数据。其中,基类训练阶段统计了60个基类的训练

表6 消融实验结果

伪目标	基类训练		校准策略	A_T										$\bar{A}$	$\bar{H}$	$\bar{F}$
	$\mathcal{L}_s$	$\mathcal{L}_c$		A_{base}^T	A_{inc1}^T	A_{inc2}^T	A_{inc3}^T	A_{inc4}^T	A_{inc5}^T	A_{inc6}^T	A_{inc7}^T	A_{inc8}^T				
×	×	×	×	7.25%	98.50%	95.25%	99.25%	98.75%	99.00%	99.50%	99.75%	98.75%	49.74%	62.54%	18.64%	
√	×	×	×	56.42%	97.50%	97.00%	99.00%	97.00%	99.75%	99.75%	98.75%	100.00%	82.45%	89.41%	6.08%	
√	√	×	×	99.88%	79.50%	87.25%	88.00%	64.75%	86.75%	87.25%	93.5%	82.50%	95.99%	90.55%	0.85%	
√	√	√	×	97.21%	89.75%	94.5%	95.25%	93.25%	97.00%	97.00%	94.75%	98.00%	97.38%	97.21%	1.61%	
√	√	√	√	97.44%	93.00%	95.75%	96.75%	95.75%	97.00%	98.25%	96.75%	98.75%	97.90%	98.11%	1.27%	

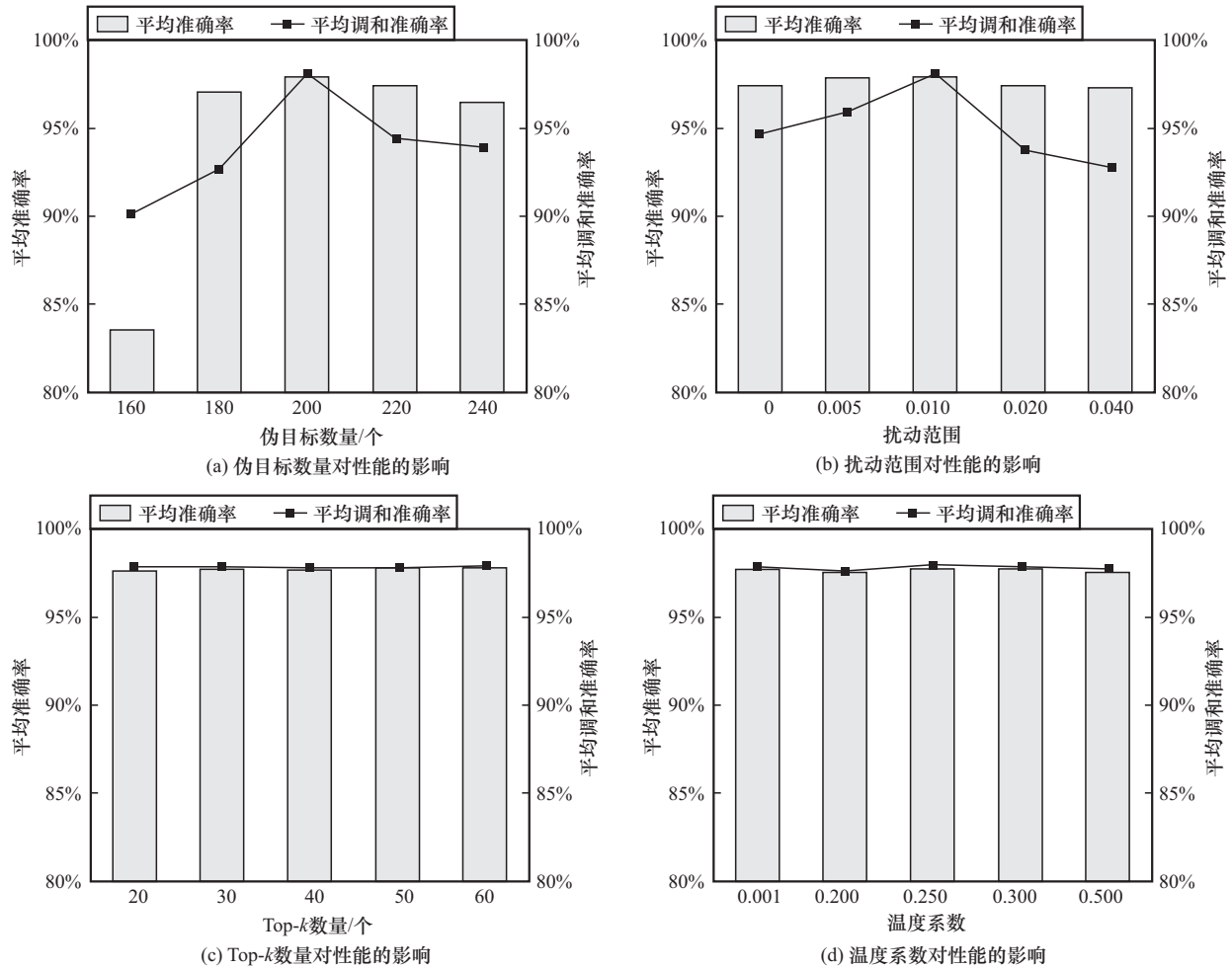


图8 不同超参数设置的性能对比

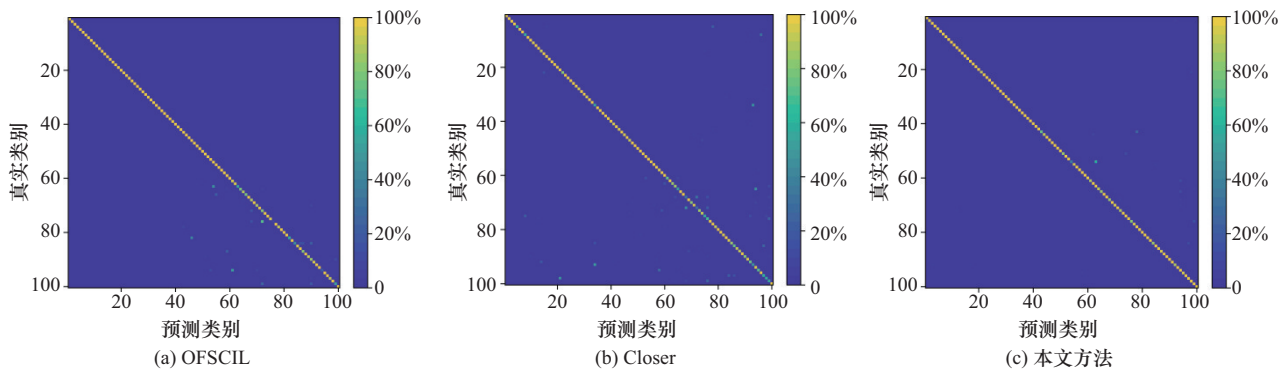


图9 不同方法的混淆矩阵对比

时间；增量训练阶段统计了单次增量5个新类、连续增量8次后的平均训练时间；推理阶段展示了在100类规模下的单样本推理时间。

本文方法的平均增量训练时间为19.384 s，相比之下，OFSCIL仅通过相似基类原型对新类原型进行一次迁移校准，不需要额外优化过程，因此能够有效减少增量训练时间。虽然本文方法带来了

一定的额外计算成本，但这是因为在增量训练阶段进一步采用了权重校准策略，以缓解小样本条件下的误判风险。实验结果表明，这一额外代价换来了更高的识别精度。本文方法的单样本推理时间为0.062 ms，与其他对比方法基本持平，表明本文方法在推理阶段不会增加额外的计算负担。此外，在当前主干网络架构下，模型的浮点运算次数

为 312.77×10^6 。结合较低的推理时延,本文方法的时间开销完全处于典型嵌入式设备的算力支撑与实时性容忍范围之内,能够良好地满足实际工业应用中效率的要求。

表7 训练与推理时间

方法	基类训练时间/s	平均增量训练时间/s	推理时间/ms
FSCIL-SEI	317.116	28.991	0.067
OFSCIL	211.039	4.82×10^{-4}	0.054
Meta-RFF	649.293	1.373	0.072
DSN	158.603	90.049	0.068
Closer	185.650	0.018	0.064
ADBS	128.272	8.144	0.067
本文方法	205.092	19.384	0.062

4 结束语

本文面向设备持续接入场景下的小样本类增量识别问题,提出了一种基于正交空间约束的特定辐射源识别小样本类增量学习方法。首先,设计了一组相互正交的伪目标向量以构建可扩展的特征嵌入空间,并基于理论分析给出了伪目标数量的上下界。其次,提出了结合交叉熵损失、自监督对比损失和类中心分离损失的协同优化策略,以引导特征提取器为后续新类预留充分的嵌入方向。最后,设计了分类器权重校准策略,通过显式约束分类器权重增强决策边界的可分性。基于公开的ADS-B和Wi-Fi数据集对本文方法进行了评估,实验结果表明,在样本稀缺条件下,本文方法在新类学习与旧类知识保持之间取得了较优平衡。未来将进一步面向跨域条件下的小样本持续类别扩展问题,研究适应不同信道环境与复杂数据分布的优化策略,在增强跨域准确识别能力的同时保持对小样本新类别的有效学习能力,从而提升模型的鲁棒性与持续适应能力。

参考文献:

[1] Hamamreh J M, Furqan H M, Arslan H. Classifications and applications of physical layer security techniques for confidentiality: a comprehensive survey[J]. IEEE Communications Surveys & Tutorials, 2019, 21(2): 1773-1828.

[2] Zha H R, Wang H H, Feng Z M, et al. LT-SEI: long-tailed specific emitter identification based on decoupled representation learning in low-resource scenarios[J]. IEEE Transactions on Intelligent Transportation

Systems, 2024, 25(1): 929-943.

[3] 何遵文,侯帅,张万成,等.通信特定辐射源识别的多特征融合分类方法[J].通信学报,2021,42(2):103-112.

He Z W, Hou S, Zhang W C, et al. Multi-feature fusion classification method for communication specific emitter identification[J]. Journal on Communications, 2021, 42(2): 103-112.

[4] 查浩然,刘畅,王巨震,等.面向无人机辐射源个体识别的域适应模型设计[J].信号处理,2024,40(4):650-660.

Zha H R, Liu C, Wang J Z, et al. Design of domain-adaptation model for specific emitter identification of UAV signal[J]. Journal of Signal Processing, 2024, 40(4): 650-660.

[5] 张立民,谭凯文,闫文君,等.基于持续学习和联合特征提取的特定辐射源识别[J].电子与信息学报,2023,45(1):308-316.

Zhang L M, Tan K W, Yan W J, et al. Specific emitter identification based on continuous learning and joint feature extraction[J]. Journal of Electronics & Information Technology, 2023, 45(1): 308-316.

[6] Ji P L, Peng Y, Lu G, et al. Incremental learning-based radio frequency fingerprint identification via exemplar replay[C]//Proceedings of the 2023 IEEE 23rd International Conference on Communication Technology (ICCT). Piscataway: IEEE Press, 2023: 1-5.

[7] Li D Z, Qi J, Hong S H, et al. A class-incremental approach with self-training and prototype augmentation for specific emitter identification[J]. IEEE Transactions on Information Forensics and Security, 2024, 19: 1714-1727.

[8] Sun L, Xue R, Zha H R, et al. AFD-IL: a long-term incremental learning approach with adaptive feature distillation for specific emitter identification[J]. IEEE Transactions on Cognitive Communications and Networking, 2026, 12: 2769-2781.

[9] Liu Y X, Wang J, Li J Q, et al. Class-incremental learning for wireless device identification in IoT[J]. IEEE Internet of Things Journal, 2021, 8(23): 17227-17235.

[10] 彭鹏,曹帅,贾勇,等.结合特征分布优化的辐射源类增量识别方法[J].电讯技术,2026,4:629-626.

Peng P, Cao S, Jia Y, et al. Incremental identification of radiation source classes combined with feature distribution optimization [J]. Telecommunication Engineering, 2026, 4: 629-626.

[11] Li D Z, Chen Z D, Shao M Y, et al. Non-exemplar class-incremental learning via prototype correction and hierarchical regularization for specific emitter identification[J]. IEEE Transactions on Intelligent Transportation Systems, 2025, 26(8): 12632-12646.

[12] Sun M H, Teng J Z, Liu X Y, et al. Few-shot specific emitter identification: a knowledge, data, and model-driven fusion framework[J]. IEEE Transactions on Information Forensics and Security, 2025, 20: 3247-3259.

[13] Cai Z R, Ma W X, Wang X R, et al. The performance analysis of time series data augmentation technology for small sample communication device recognition[J]. IEEE Transactions on Reliability, 2023, 72(2): 574-585.

[14] Liu C, Fu X, Wang Y, et al. Overcoming data limitations: a few-shot specific emitter identification method using self-supervised learning and adversarial augmentation[J]. IEEE Transactions on Information Fo-

- rensics and Security, 2024, 19: 500-513.
- [15] Xie C X, Zhang L M, Zhong Z G. Few-shot specific emitter identification based on variational mode decomposition and meta-learning[J]. Wireless Communications and Mobile Computing, 2022, 2022: 4481416.
- [16] Liu K S, Yan W J, Zhang L M, et al. An open-set communication-specific emitter identification method based on adaptive Weibull hierarchical decision[J]. IEEE Internet of Things Journal, 2026, 13(6): 12028-12038.
- [17] Zhang T, Shen X Y, Shang Z H, et al. An open-set few-shot class incremental learning framework for specific emitter identification[J]. IEEE Transactions on Cognitive Communications and Networking, 2026, 12: 602-616.
- [18] Li T T, Wen Z Y, Jin C D, et al. Meta-RFF: meta-task adaptive-based few-shot open-set incremental learning for RF fingerprint recognition[J]. IEEE Transactions on Cognitive Communications and Networking, 2026, 12: 2115-2130.
- [19] Li D Z, Chen X W, Hong S H, et al. FSCIL-SEI: few-shot class-incremental learning approach for specific emitter identification[J]. IEEE Transactions on Instrumentation and Measurement, 2025, 74: 2504914.
- [20] Oh J, Baik S, Lee K M. CLOSER: towards better representation learning for few-shot class-incremental learning[C]//Computer Vision-ECCV 2024. Berlin: Springer, 2024: 18-35.
- [21] Zhou D W, Wang F Y, Ye H J, et al. Forward compatible few-shot class-incremental learning[C]//Proceedings of the 2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). Piscataway: IEEE Press, 2022: 9036-9046.
- [22] Li L H, Tan Y Z, Yang S Y, et al. Adaptive decision boundary for few-shot class-incremental learning[J]. Proceedings of the AAAI Conference on Artificial Intelligence, 2025, 39(17): 18359-18367.
- [23] Yang B Y, Lin M B, Zhang Y X, et al. Dynamic support network for few-shot class incremental learning[J]. IEEE Transactions on Pattern Analysis and Machine Intelligence, 2023, 45(3): 2945-2951.
- [24] Shi W Q, Teng F, Lei Y K, et al. DPDS: a systematic framework for few-shot specific emitter incremental identification[J]. IEEE Internet of Things Journal, 2025, 12(7): 8725-8741.
- [25] Tropp J A, Dhillon I S, Heath R W, et al. Designing structured tight frames via an alternating projection method[J]. IEEE Transactions on Information Theory, 2005, 51(1): 188-209.
- [26] Liu Y X, Wang J, Niu S T, et al. ADS-B signals records for non-cryptographic identification and incremental learning[R]. IEEE Data-Port, 2021.
- [27] Hanna S, Karunaratne S, Cabric D. WiSig: a large-scale Wi-Fi signal dataset for receiver and channel agnostic RF fingerprinting[J]. IEEE Access, 2022, 10: 22808-22818.
- [28] Zhao L L, Lu J, Xu Y L, et al. Few-shot class-incremental learning via class-aware bilateral distillation[C]//Proceedings of the 2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). Piscataway: IEEE Press, 2023: 11838-11847.
- [29] Chaudhry A, Dokania P K, Ajanthan T, et al. Riemannian walk for incremental learning: understanding forgetting and intransigence[C]//Computer Vision-ECCV 2018. Berlin: Springer, 2018: 556-572.

[作者简介]



孙露 (2000-), 女, 哈尔滨工程大学博士生, 主要研究方向为信号处理、特定辐射源识别和增量学习。



薛睿 (1980-), 男, 博士, 哈尔滨工程大学教授、博士生导师, 主要研究方向为无线电移动通信系统、卫星通信系统、卫星导航与定位等。



查浩然 (1996-), 男, 哈尔滨工程大学博士生, 主要研究方向为信号处理、机器学习和数据分析。



林云 (1980-), 男, 博士, 哈尔滨工程大学教授、博士生导师, 主要研究方向为无线网络中的机器学习与数据分析、信号处理与分析等。



王巍 (1980-), 男, 博士, 中国电子科技集团第三十六研究所研究员, 主要研究方向为无线电信号处理、分析与识别和深度学习算法等。