

面向开放世界持续学习的任务敏感提示驱动混合专家模型^{*}

李昱洁^{2,3}, 吴 晗¹, 孟 丹^{2,3}, 李天瑞⁴, 杨 新¹

¹(西南财经大学 计算机与人工智能学院, 四川 成都 611130)

²(西南财经大学 管理科学与工程学院, 四川 成都 611130)

³(智能金融教育部工程研究中心 (西南财经大学), 四川 成都 611130)

⁴(西南交通大学 计算机与人工智能学院, 四川 成都 614202)

通信作者: 杨新, E-mail: yangxin@swufe.edu.cn



摘 要: 开放世界持续学习 (OWCL) 旨在模拟现实环境中任务不断演化、类别动态变化且遇到未经训练的未知样本的情景. 一个好的开放世界持续学习模型不仅需要在学习新任务的同时保持对已学任务的记忆, 还需具备识别未知类别的能力, 进而实现持续且鲁棒的知识积累与泛化. 然而, 现有持续学习方法普遍建立在封闭世界假设之上, 无法有效应对开放类别带来的类别不确定性与任务间干扰, 尤其在知识稳定性与知识可塑性之间的权衡上表现出明显不足. 因此, 在开放世界持续学习问题的形式化定义基础上, 提出一种任务敏感提示驱动的混合专家模型 TP-MoE (task-aware prompt-driven mixture of experts), 以实现对任务语义的动态建模与专家模块的高效调度, 从而帮助模型进行知识传输和知识更新. 具体而言, TP-MoE 引入一种即插即用的任务提示聚合机制并改进门控机制用以专家网络路由, 在任务增量过程中持续融合历史与当前任务知识; 同时结合一种自适应开放边界阈值策略, 可根据新旧知识的迁移动态调整开放类别的判别边界, 从而提升开放类别检测能力与已知类别分类准确性. 实验结果表明, TP-MoE 在 Split-CIFAR100 和 Open-CORE25 基准数据集上对各类指标的测试均取得领先性能, 展现出良好的稳健性与泛化性, 为开放世界持续学习任务中的知识建模与任务调度提供了一种可扩展、可迁移的新框架.

关键词: 开放世界持续学习; 持续学习; 任务敏感; 混合专家模型; 知识传输

中图法分类号: TP18

中文引用格式: 李昱洁, 吴晗, 孟丹, 李天瑞, 杨新. 面向开放世界持续学习的任务敏感提示驱动混合专家模型. 软件学报, 2026, 37(4): 1531–1547. <http://www.jos.org.cn/1000-9825/7525.htm>

英文引用格式: Li YJ, Wu H, Meng D, Li TR, Yang X. Task-aware Prompt-driven Mixture-of-experts Model for Open-world Continual Learning. Ruan Jian Xue Bao/Journal of Software, 2026, 37(4): 1531–1547 (in Chinese). <http://www.jos.org.cn/1000-9825/7525.htm>

Task-aware Prompt-driven Mixture-of-experts Model for Open-world Continual Learning

LI Yu-Jie^{2,3}, WU Han¹, MENG Dan^{2,3}, LI Tian-Rui⁴, YANG Xin¹

¹(School of Computing and Artificial Intelligence, Southwestern University of Finance and Economics, Chengdu 611130, China)

²(School of Management Science and Engineering, Southwestern University of Finance and Economics, Chengdu 611130, China)

³(Engineering Research Center for Intelligent Finance (Southwestern University of Finance and Economics), Ministry of Education, Chengdu 611130, China)

⁴(School of Computing and Artificial Intelligence, Southwest Jiaotong University, Chengdu 614202, China)

Abstract: Open-world continual learning (OWCL) aims to simulate real-world scenarios where tasks evolve continuously, categories change dynamically, and unseen samples are encountered. A well-designed OWCL model is expected not only to retain knowledge of

* 基金项目: 国家自然科学基金 (62476228)

本文由“高可信自适应机器学习”专题特约编辑王熙照教授、张敏灵教授、曹付元教授推荐.

收稿时间: 2025-05-12; 修改时间: 2025-06-30, 2025-08-15; 采用时间: 2025-08-20; jos 在线出版时间: 2025-09-02

CNKI 网络首发时间: 2026-01-27

learned tasks while acquiring new tasks but also to recognize unknown categories, thus achieving continuous and robust knowledge accumulation and generalization. However, most existing continual learning methods are built upon the closed-world assumption and cannot effectively cope with the category uncertainty and inter-task interference introduced by open categories. In particular, they show clear limitations in balancing knowledge stability and plasticity. Therefore, based on the formal definition of the OWCL problem, this study proposes a task-aware prompt-driven mixture-of-experts model (TP-MoE), which realizes dynamic modeling of task semantics and efficient scheduling of expert modules, thus supporting knowledge transfer and knowledge update. Specifically, TP-MoE introduces a plug-and-play task prompt aggregation mechanism and improves the gating strategy for expert routing, enabling the continual integration of historical and current task knowledge during task increments. At the same time, an adaptive open-boundary thresholding strategy is incorporated, which dynamically adjusts the decision boundaries of open categories according to the transfer between new and old knowledge, thus enhancing both open-category detection capability and known-category classification accuracy. Experimental results demonstrate that TP-MoE achieves state-of-the-art performance across various metrics on the Split-CIFAR100 and Open-CORe25 benchmarks, exhibiting strong robustness and generalization. This study provides a scalable and transferable framework for knowledge modeling and task scheduling in open-world continual learning.

Key words: open-world continual learning (OWCL); continual learning; task-aware; mixture-of-experts (MoE) model; knowledge transfer

1 引言

持续学习 (continual learning), 又称终身学习或增量学习, 旨在使模型能够通过持续地学习不同的任务而不遗忘之前学习的任务, 并在不断学习新知识的同时保留已获得的知识^[1-3]. 现有的大多数持续学习方法通常建立在“封闭世界”的假设之上, 即各任务的数据分布保持不变, 每个任务的训练集和测试集中的类别集合保持一致^[4-6]. 然而, 现实应用环境往往远比这一假设复杂: 模型可能在测试阶段遇到从未见过的新类别 (即开放类别)^[7-9], 这就导致即便是性能最优的持续学习方法, 在面对测试时出现的未知类别时也会表现出显著的性能退化. 正是由于这些现实挑战, 推动了持续学习向更具实际应用价值的开放世界环境下拓展, 要求持续学习模型在适应性和鲁棒性方面具备更高能力, 以应对复杂动态的开放环境变化.

开放世界持续学习 (open-world continual learning, OWCL)^[6,9,10] 是一种兼具实用性与挑战性的新的机器学习范式. 在开放世界持续学习中, 模型需在开放环境中不断地适应潜在的动态任务序列, 其中包含的开放未知类别可能在测试阶段不可预测、随机地出现^[9-12]. 与传统封闭集假设下的持续学习模型不同, 开放世界持续学习强调“边学边用”, 这不仅要求开放世界持续学习模型能够识别从未见过的样本 (或未知类别), 还需要在不遗忘旧任务知识的前提下持续地学习和更新关于新出现的任务的知识. 如图 1 所示, 我们以任务 1 至任务 t 简要展示开放世界持续学习意图和本文研究动机. 在任务 1 中, 测试集中出现了训练集中没有出现过的未知类别, 面对这样的情况, 现有的持续学习模型 (对应输出 1) 无法正确地识别出开放类别, 而是强制地将其错误地分类至现有的类别中; 然而, 开放世界持续学习模型 (对应输出 2) 能够正确地识别出开放类别, 同时也能保证在已知类别 (即训练集中出现过的类别) 上的分类准确度. 随着任务不断地出现和学习, 到任务 t 时, 出现了之前任务 1 中的开放类别的有标记样本, 那么持续学习模型才能够正确地将其分类, 但是面对没有标记的未知样本依然会造成错误的判别; 然而开放世界持续学习模型不仅能够不遗忘已经学习过的知识, 还能够保持检测开放未知样本的良好能力, 减少错误分类的发生, 展示出更强的泛化能力.

因此, 开放世界持续学习面临两大核心挑战: 一是如何准确识别未知样本, 避免未知样本被错误归类到现有的已知类别中; 二是如何在持续学习的过程中保持对已学知识的稳定记忆和新学知识的增量学习, 即能够对已知样本做出良好的分类. 而需要强调的是, 开放世界持续学习中的“开放识别”与“已知分类”并非独立任务, 而是相互依赖、相互影响的: 未知样本的存在会加剧模型在稳定性 (stability) 与可塑性 (plasticity) 之间的权衡难度, 而新任务的逐步持续引入又进一步模糊了已知类别在嵌入空间中的决策边界, 使得开放识别也变得更为困难.

为了应对这两个挑战, 本文首先深入分析开放世界持续学习中“已知-未知”的相互作用机制, 尝试揭示实际应用中更加复杂的知识演化过程. 同时, 本文在实验部分详细划分了两种基础开放世界持续学习场景, 即类增量开放世界持续学习和域增量开放世界持续学习, 并在这两个场景下复现对比了大量基准模型. 然而, 现有的开放世界持

持续学习方法大多仍是“开放集识别方法+持续学习方法”的简单策略叠加, 忽视了未知样本本身所蕴含的可迁移知识和“开放识别”与“已知分类”的相互影响关系, 从而导致现有的开放世界持续学习方法在真实的开放动态环境中的适应性和泛化能力受限. 因此, 一个理想的开放世界持续学习方法应具备任务敏感的知识学习和知识迁移——既能在已经学习过的任务中的已知类别上保持良好效果, 也能有效地利用来自不同任务的未知样本的知识判别未知类别和已知类别.

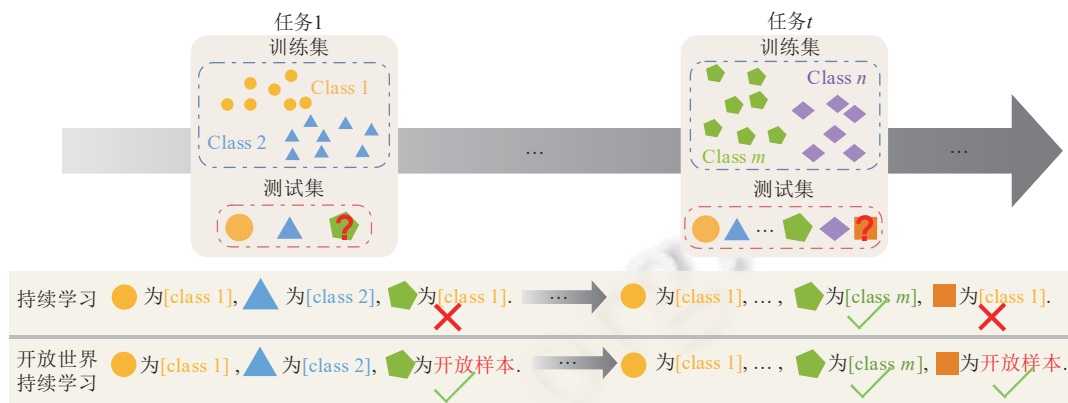


图1 开放世界持续学习问题的简明示意图

具体而言, 本文提出了一种面向开放世界持续学习的混合专家模型结构, 引入一种即插即用的任务敏感提示聚合策略, 用于持续地积累、融合来自不同任务的知识, 从而支持开放世界持续学习模型进行知识积累、知识传输甚至知识更新. 同时, 为了更精准地识别开放样本, 本文提出了一种自适应的阈值选择策略, 能够根据新知识的积累更新和旧知识的迁移实现自适应开放决策边界. 具体而言, 本文提出的 TP-MoE (task-aware prompt-driven mixture of experts) 模型基于混合专家模型改进了专家网络选择门控机制, 通过任务敏感的注意力提示机制编码与聚合任务共性及任务特性, 实现了开放世界持续任务中的知识积累、迁移和更新; 针对开放识别, 我们提出了一种自适应的开放边界阈值机制, 结合理论支撑的自适应阈值选择方法, 使模型能根据新知识不断调整其开放集识别策略. 并且, 本文在两个公开基准数据集上进行了充分的对比实验, 验证了所提模型在开放世界持续学习任务中的显著优势和稳健性.

本文第2节分析持续学习和开放世界持续学习的相关方法和研究现状. 第3节介绍本文所需的基础知识, 包括基于提示学习的持续学习和涉及混合专家模型的先验知识和基础架构. 第4节构建本文提出的 TP-MoE 模型. 第5节通过大量实验验证所提模型的有效性. 最后一节总结全文并提出后续研究方向.

2 开放世界持续学习相关研究及现状分析

持续学习旨在使模型能够在不需要反复重新训练所有数据的前提下, 持续增量地从数据流中学习新知识, 同时避免遗忘已学到的知识. 现有持续学习方法主要可分为3类^[13,14]: (1) 基于网络结构扩展的方法, (2) 基于正则约束的方法, (3) 基于回放机制的方法. 具体而言, 第1类基于网络结构扩展类方法^[15-19]通过拓展模型结构来适配新任务并保留旧知识, 例如 PackNet^[16]利用参数掩码隔离不同任务的权重, BNS^[17]利用强化学习构建任务专属学习器. 第2类基于正则约束的方法^[20,21]则通过惩罚关键参数的变化来缓解灾难性遗忘问题, 例如经典的持续学习基准框架 EWC^[18,22]利用费舍尔信息矩阵约束模型重要参数的更新, LwF^[23]通过正则约束使得模型在学习新任务时保持对旧任务输出的一致性, 实现无需访问旧数据的持续学习. 第3类基于回放机制的方法^[24-27]则通过保存部分旧任务样本或生成伪样本与新任务数据联合训练, 从而缓解灾难性遗忘, 例如 iCaRL^[27]使用的伪样本合并当前任务训练的回放策略.

尽管上述持续学习方法在持续地学习新知识和缓解灾难性遗忘方面取得了成效, 然而, 现实世界环境充满不

确定性, 在测试阶段经常会遇到从未见过的未知类别 (open/unseen classes), 使得建立在封闭世界假设之上的传统持续学习方法在应对开放样本时往往会产生知识误判、错误整合或更新失效等问题, 严重影响模型的泛化能力和适应能力.

为应对上述不足, 开放世界持续学习 (OWCL) 被提出^[6,12], 旨在使模型能够在面对新任务持续学习的同时, 具备识别并适应未知类别的能力. 作为一个新兴的研究领域, 开放世界持续学习更贴近现实应用需求, 因其模拟了模型在动态开放环境中“边学边识别”的过程. 近期研究对开放世界学习关注日益增强, 部分工作也尝试将开放识别能力引入持续学习框架中, 例如 Kim 等人^[10]强调新颖性检测的必要性, Liu 等人^[6]提出的 SOLA 框架实现任务自适应增量学习, Li 等人^[9]提出开放样本的知识传输以帮助模型避免在持续学习过程中由于测试时出现开放样本造成的误分类.

此外, 当前开放世界持续学习方法仍普遍存在两个关键不足: 一是缺乏统一的知识表征与整合机制, 难以有效地将历史任务中获得的知识迁移并适配到后续任务; 二是面对频繁出现的未知样本时, 模型易受其干扰, 进一步放大“开放风险”, 造成知识退化或错误迁移. 这两个不足不仅影响模型对未来任务的适应能力, 也严重阻碍了开放世界持续学习在现实应用中的落地可能性. 因此, 如何在持续学习过程中构建稳健的知识迁移机制, 以实现已知与未知样本的协同学习与表示, 是开放世界持续学习当前亟待解决的关键问题.

在此背景下, 研究开放世界持续学习具有重要现实意义与理论价值. 一方面, 开放世界持续学习为构建具备自主学习与适应能力的智能系统提供了基础框架, 能更好地支持深度学习模型在动态环境下的长期任务执行; 另一方面, 开放世界持续学习也推动了对知识迁移机制、不确定性处理和分布漂移响应能力等问题的系统化研究. 因此, 开发一种兼具开放识别能力、持续学习能力与知识迁移能力的开放世界持续学习模型, 是本文解决的主要问题.

3 先验知识

本节就开放世界持续学习的问题定义、所提方法涉及的相关概念和基本知识予以介绍.

3.1 问题定义: 开放世界持续学习

给定一系列 (可能无穷的) 连续的任务, 在开放世界持续学习的过程中, 如果模型在测试阶段遇到任何新出现的未知类别 (novel object), 则应将其检测为“未知 (unknown)”. 随着新任务的不断学习, 若模型在任何后续的新任务中接触到这些“未知样本”的真实标签 (ground-truth label), 这些“未知”将被转化为“已知 (known)”. 开放世界持续学习模型不仅要具备不遗忘已有知识的能力, 即克服灾难性遗忘问题, 还需要具备识别未知样本并不断更新来自未知样本的知识和知识迁移的能力.

简要地, 我们对所研究的问题, 即开放世界持续学习, 进行如下数学形式化定义.

问题定义: 开放世界持续学习要求学习一个包含 T 个任务的序列 $1, \dots, t, \dots, T$, 其中每个任务都拥有各自的训练集 $\mathcal{D}_t = \{(x_i^t, y_i^t)\}_{i=1}^n$, n 表示任务 t 的训练样本数量, $x_i^t \in X^t$ 是输入的训练样本, $y_i^t \in Y^t$ 是其对应的类别标签. 遵循现有的类增量持续学习设定, 我们给定 OWCL 中, 各任务的训练集中类别互不重叠, 即 $Y^n \cap Y^m = \emptyset, \forall n \neq m$, 其中 n, m 表示任务序号, $\bigcup_{t=1}^T Y^t = Y_{\text{tr}}$ 表示所有训练任务的总体类别集合. 并且, 最关键的一点在于测试阶段可能会出现从未在训练集中出现过的、未经训练过的未知类别样本, 即 $Y_{\text{tr}} \neq Y_{\text{te}}$.

总体来说, 开放世界持续学习的目标是构建一个统一的预测函数或分类器, 其能够: 1) 检测出未知/开放类别的测试样本; 2) 在不遗忘已学知识的前提下, 识别每个已知类别的测试样本. 需要特别注意的是, 在开放世界持续学习的设定下, 在训练任务 t 时, 先前任务的训练数据 ($\mathcal{D}_1^t, \dots, \mathcal{D}_{t-1}^t$) 均是不可获取和使用的.

如第 1 节所述, 在开放世界持续学习场景中, 模型需要在保持历史任务知识 (知识的稳定性) 与快速适应新任务知识 (知识的可塑性) 之间实现有效权衡. 为此, 我们从梯度优化的视角对模型训练过程进行了理论分解.

具体地, 第 t 个任务训练中的总梯度更新可以通用地表示为如下形式:

$$\nabla \mathcal{L}_t(\theta) = \nabla \mathcal{L}_{\text{old}}(\theta) + \nabla \mathcal{L}_{\text{new}}(\theta) \quad (1)$$

其中, $\nabla \mathcal{L}_{\text{old}}$ 表示与历史任务保持相关的梯度项, 例如由记忆回放、正则项或结构冻结产生的梯度; 而 $\nabla \mathcal{L}_{\text{new}}$ 对应于当前任务学习的损失梯度. 由于这两个梯度项往往存在方向冲突, 容易导致灾难性遗忘, 用于衡量其冲突程度的指标通常被如下定义:

$$\cos(\theta_{\text{conflict}}) = \frac{\nabla \mathcal{L}_{\text{old}} \cdot \nabla \mathcal{L}_{\text{new}}}{|\nabla \mathcal{L}_{\text{old}}| |\nabla \mathcal{L}_{\text{new}}|} \quad (2)$$

当 $\cos(\theta_{\text{conflict}}) < 0$ 时, 表明两类梯度方向相悖, 模型更新可能会导致旧知识遗忘.

分析发现, 现有基于正则约束的持续学习方法通常使用费舍尔信息矩阵对历史参数进行重要程度加权, 进而显式地构造稳定性约束; 然而, 开放世界持续学习由于未知样本存在造成任务边界难以构建, 需要通过任务敏感选择实现对隐式历史知识的非显式保留. 此外, 相比于使用显式投影的持续学习方法, 开放世界持续学习要求能够自适应地选择适合当前任务的参数子网络, 将历史任务与新任务的梯度分别投射到参数空间的不同子空间中, 实现参数隔离, 降低梯度冲突, 从而在无需显式正交投影的前提下达到知识的平衡稳定性与可塑性.

3.2 基于提示的持续学习方法

除第2节介绍的传统的3类持续学习方法, L2P^[28]提出了基于提示 (prompt-based) 的持续学习方法, 此类方法作为无回放 (rehearsal-free) 持续学习的有效 $p \in \mathbb{R}^{L_p \times d}$ 策略, 受到了越来越多研究的关注. 此类方法通常利用一个预训练模型主干网络 f_θ 进行特征提取, 其中主干网络的模型参数 θ 通常保持冻结状态, 即不参与训练更新. 为提升模型在新任务上的适应能力, 基于提示的持续学习方法引入提示 (prompt), 通过训练得到一组小规模、可学习的参数用以调控多头自注意力 (multi-head self-attention, MSA) 模块中的计算更新过程^[29,30]. 不同的模型使用不同的策略将提示融合至多头自注意力的查询矩阵 (query)、键矩阵 (key) 和值矩阵 (value) 中, 从而帮助持续学习模型有效地学习新任务, 不仅不会遗忘学习过的任务, 甚至能够使用过去学习到的知识帮助模型更好地处理新任务.

现有的基于提示的持续学习方法, 通常通过为每个新任务学习新的 (一批) 提示增强样本嵌入, 以缓解灾难性遗忘 (catastrophic forgetting) 问题. 在测试推理阶段, 模型需从已有的提示集合中选择合适的单个或多个提示组合, 以应对来自各任务的测试样本. 例如, L2P^[28]提出了共享提示池 (prompt pool) 的机制, 即每次训练新的任务后都更新提示池; 在测试时, 使用最新训练得到的提示池, 通过键-值对匹配 (key-query matching) 机制进行提示检索, 为样本的特征嵌入进行知识增强. DualPrompt^[31]进一步细分提示池, 引入了任务无关 (G-prompt) 与任务特定 (E-prompt) 两类提示, 以分别捕获通用性与特异性信息. S-prompt^[32]则专注于学习任务特定提示, 采用了与 L2P 类似的提示选择和样本增强策略. CODA-P^[33]则选择扩展提示池, 并通过引入注意力权重矩阵对提示池中的不同提示进行加权融合, 以提升提示选择的灵活性, 增强样本嵌入信息. 最近的 HiDe-prompt^[34]进一步提出了持续学习目标的分层解方法, 通过独立优化每一子目标, 显著提升了整体性能, 达到了当前最优性能水平.

具体来说, 在使用提示对样本嵌入进行知识增强时, 提示参数通常被记作 p , 其中 $\mathbb{R}$ 表示嵌入空间, L_p 表示提示序列长度, d 表示嵌入维度. 提示被划分为两个部分 $p^K, p^V \in \mathbb{R}^{L_p/2 \times d}$, 分别融合合并至样本嵌入的键值向量上.

$$f_{\text{prompt}}(p, X^Q, X^K, X^V) \sim \text{MSA} \left(X^Q, \left[\begin{array}{c} p^K \\ X^K \end{array} \right], \left[\begin{array}{c} p^V \\ X^V \end{array} \right] \right) \quad (3)$$

其中, X^Q, X^K 和 X^V 分别为输入经过投影后的查询、键和值向量, $||$ 表示向量拼接操作, MSA 表示多头自注意力层操作, 其机制具体细节将在第3.3节中介绍. 拼接后的提示向量以嵌入增强前缀形式参与注意力计算, 从而引导模型进行任务特定的上下文建模.

3.3 注意力机制与混合专家模型

Transformer 架构在传统的多序列建模任务中取得了卓越的性能, 其核心能力源于注意力机制 (attention mechanism)^[29,30], 该机制允许模型在处理输入序列时动态捕捉全局依赖关系. 在近年来的模型扩展中, 专家混合结构 (mixture of experts, MoE)^[35-37]被广泛用于进一步提升模型容量, 其底层数学形式与注意力机制存在结构类似性, 为网络模型拓展至开放世界持续学习提供了可能.

目前, 最常见且被广泛应用的注意力形式为缩放点积注意力 (scaled dot-product attention)^[38].

设查询矩阵 $Q \in \mathbb{R}^{M \times d_Q}$, 键矩阵 $K \in \mathbb{R}^{N \times d_K}$, 值矩阵 $V \in \mathbb{R}^{N \times d_V}$, 其中 N 表示键-值对数量, M 表示查询向量数量. 那么, 缩放点积注意力计算为:

$$Attention(Q, K, V) = Softmax\left(\frac{Q \cdot K^T}{\sqrt{d_K}}\right) \quad (4)$$

其中, $Softmax$ 函数对 $Q \cdot K^T \in \mathbb{R}^{M \times N}$ 的每一行进行归一化, 输出为查询对值的加权聚合.

该注意力机制被应用于组成 Transformer 基本框架的多头自注意力层 $MSA(\cdot)$ 中, 每个头独立学习不同的注意力子空间. 令输入序列为 $X = [x_1, \dots, x_N]^T \in \mathbb{R}^{N \times d}$, 查询、键、值分别记为 X^Q, X^K, X^V . 设 l 为注意力头的数量, 那么 MSA 输出结构为:

$$MSA(X^Q, X^K, X^V) = Concat(h_1, \dots, h_l) W \quad (5)$$

$$h_i = Attention(X^Q \cdot W_i^Q, X^K \cdot W_i^K, X^V \cdot W_i^V), i \in [1, l] \quad (6)$$

其中, $W \in \mathbb{R}^{md_V \times d}$, $W_i^Q, W_i^K, W_i^V \in \mathbb{R}^{d_K \times d}$ 均表示可训练投影矩阵, 通常设置 $d_K = d_V = \frac{d}{m}$.

为进一步提升模型容量与计算效率, 专家混合结构 (MoE 结构)^[35] 被提出用于对不同专家网络进行稀疏激活. 该结构引入了门控函数 (gating function), 在给定输入的条件下对多个子网络输出进行加权聚合. 如图 2 所示, MoE 架构中所使用的专家混合层 (MoE 层) 通过引入多个子网络 (称为专家) 与一个门控网络实现稀疏激活机制. 在每个 MoE 层中, 输入首先经过门控网络进行评分, 动态选择其中一小部分专家 (如 Top- k) 参与前向计算. 被选中的专家对输入进行独立变换, 其输出经过门控权重加权后求和, 作为该层的最终输出. 该结构允许模型在保持计算成本不变的情况下显著提升参数规模与表达能力, 被广泛用于构建高效的大规模神经网络. 因此, MoE 层可以被看作是嵌入在循环神经网络结构 (recurrent neural network, RNN) 中的一种特殊编码层. 在 MoE 层中, 稀疏门控函数选择两个专家参与计算, 它们的输出将根据门控网络的输出进行调制.

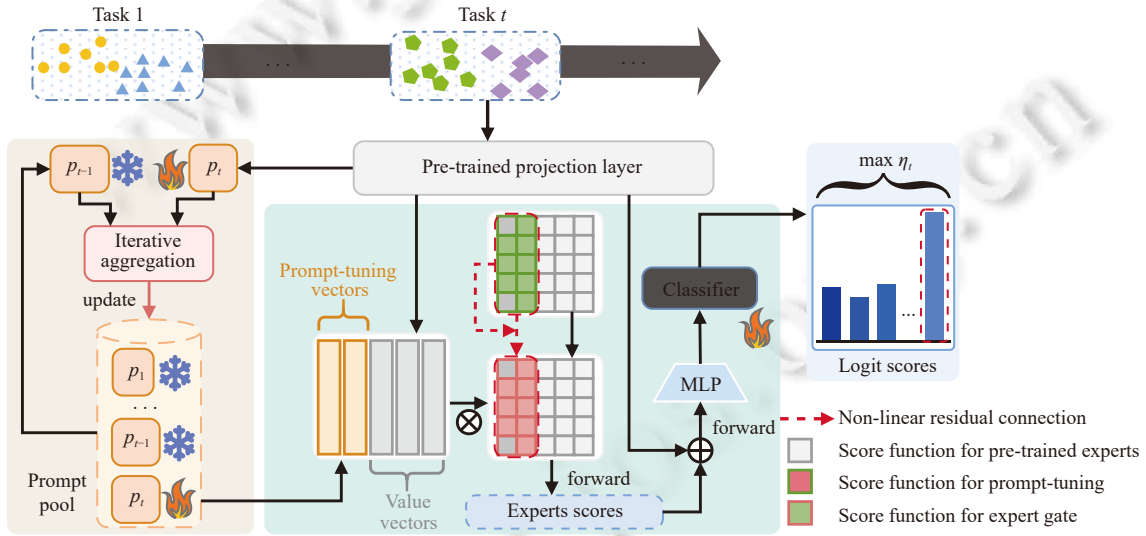


图 2 TP-MoE 模型架构简要示意图.

具体来说, MoE 层由给定的 N 个专家网络 $\{f_j: \mathbb{R}^d \rightarrow \mathbb{R}^N\}_{j=1}^N$ 和一个门控函数 $\{G: \mathbb{R}^d \rightarrow \mathbb{R}^N\}$ 组成. 当给定输入 $h \in \mathbb{R}^d$, MoE 层的输出为:

$$y = \sum_{j=1}^N G(h)_j \cdot f_j(h) = \sum_{j=1}^N \frac{\exp(s_j(h))}{\sum_{i=1}^N \exp(s_i(h))} \cdot f_j(h) \quad (7)$$

其中, $s_j(h)$ 为门控网络对第 j 个专家的打分, $G(h)_j$ 为第 j 个专家的 $Softmax$ 输出的权重分布.

可以发现, 该专家混合门控机制与注意力机制具有高度结构一致性, 二者本质上都通过 *Softmax* 权重将多个分支 (键值对或专家) 输出进行加权聚合. 因此, MoE 层可被视为注意力机制在模型容量层面的推广形式. 随着研究的深入, MoE 层逐渐被确立为扩展模型规模、提升参数利用率的基础模块.

因此, 在本研究中, 我们不仅受到现有的使用提示的持续学习研究的启发, 还通过探索自注意力机制和 MoE 结构的相似性, 在遵循了 MoE 结构提出的基础上, 将任务敏感注意力权重引入提示选择和融合策略, 设计了一个自适应确定开放阈值的判别策略; 然后, 一并引入至一个统一的预训练 MoE 结构中, 最终提出了一个具备开放类别检测能力和在动态复杂环境下进行持续学习的开放世界持续学习模型, 即 TP-MoE.

4 面向开放世界持续学习的 TP-MoE 模型

在本节中, 我们对所提出的 TP-MoE 模型进行详细的构建. 该模型以 ViT^[39] 为主干网络, 引入任务敏感的提示融合机制、非线性残差连接的稀疏激活混合专家门控模块, 直接作用于原主干网络注意力层, 来调节提示池中的不同提示对应不同专家网络的权重, 帮助模型优化分类头. 根据分类头的输出, 模型使用了一种基于 logits 分数的任务敏感开放决策边界构建方法, 使用当前任务下所有类别的最大 logits 分数作为阈值, 帮助模型能够在测试推理阶段准确地检测出未知样本.

4.1 任务敏感的提示融合策略

受现有基于提示的持续学习研究的启发, 我们首先通过更有效地融合任务特定的知识来改进单个任务下的分类效果. 为此, 给定任务 t 我们构建一个任务敏感的提示池 (task-aware prompt pool), 以整合来自不同任务的知识, 并通过交叉熵损失进行训练优化该提示池. 同时, 旧任务学习到的提示都会被储存并冻结, 用以进行知识积累, 缓解灾难性遗忘问题. 不同于现有的使用提示方法的持续学习模型, 针对任务敏感的提示池, 设计了一种融合旧知识的提示融合策略, 能够有效地融合任务敏感知识.

为了学习当前任务 t 下的提示 p_t , 我们在初始化时使用 $t-1$ 任务学习到的提示 $p_t \leftarrow p_{t-1}$, 并进一步结合所有旧任务 $1, \dots, t-1$ 学习到的提示的加权组合进行迭代优化:

$$p_t = \alpha \sum_{i=1}^{t-1} p_i + (1-\alpha) p_t \quad (8)$$

其中, α 为用于控制旧任务知识在新任务学习提示时的回放程度的超参数. 然而, 尽管该策略有助于新任务性能, 可能导致与旧任务表示的重叠, 从而影响单个任务下的分类效果表现. 为解决此问题, 我们受到经典正则约束的持续学习方法的启发, 引入旧任务 $1, \dots, t-1$ 表示的特征分布信息作为约束信息, 通过计算类别质心, 进一步设计正则约束项 $\mathcal{L}_{\text{regu}}(p_t)$:

$$\mathcal{L}_{\text{regu}}(p_t) = \sum_{h \in \mathcal{H}_t} \frac{1}{\sum_{i=1}^{t-1} |\mathcal{Y}_i|} \sum_{i=1}^{t-1} \sum_{c \in \mathcal{Y}_i} \log \frac{\exp(h \cdot \mu_c / \tau)}{\sum_{h' \in \mathcal{H}_t} \exp(h \cdot h' / \tau) + \sum_{i=1}^{t-1} \sum_{c \in \mathcal{Y}_i} \exp(h \cdot \mu_c / \tau)} \quad (9)$$

其中, $\mathcal{H}_t$ 是将当前任务的训练数据集 $\mathcal{D}_t$ 经过通用冻结投影层映射后的嵌入集合, τ 为温度系数超参数, 通常设为 0.4–0.8, μ_c 表示的是类别质心, 通过均值计算得到.

由此, 单个任务下分类的总损失函数可由分类损失 (用于优化分类头参数) 和正则约束项计算得到:

$$\mathcal{L}_1(f_\theta, p_t) = \mathcal{L}_{\text{CE}}(f_\theta, p_t) + \lambda \cdot \mathcal{L}_{\text{regu}}(p_t) \quad (10)$$

其中, λ 为超参数, 用于平衡多个损失超参数.

在开放世界持续学习中, 模型在推理阶段是无法知道任务标识的, 因此, 在保证单个任务下的分类效果后, 模型需要进一步提升对任务标识识别的能力, 同时帮助开放世界持续学习模型更好地构建已知类别和未知类别的决策边界.

因此, 接下来为了识别任务标识, 我们构造任务敏感的辅助分类头 f_θ , 用于识别已知样本所属的任务标识, 从而保证在训练过程中持续适配新任务而不遗忘旧任务.

$$\mathcal{L}_2(f_{\delta_i}) = \frac{1}{\sum_{i=1}^{t-1} |\mathcal{Y}_i|} \sum_{i=1}^{t-1} \sum_{c \in \mathcal{Y}_i} \sum_{h' \in \mathcal{H}_{i,c}^{t,c}} -\log \frac{\exp(f_{\delta_i}(h')[i])}{\sum_{j=1}^t \exp(f_{\delta_i}(h')[i])} \quad (11)$$

其中, i 表示任务标识, 给定当前任务为 t , 那么 $i = [1, t]$; h' 是通过 f_{δ_i} 表示得到的样本的特征嵌入. 为了预测当前样本的特征嵌入表示 h' 所属任务 i 的对数似然损失, 我们通过交叉熵函数优化任务敏感分类头 f_{δ_i} , 每个任务 i 的类别数量将用于进行总损失归一化, 避免不同任务或类别数量影响梯度.

因此, 任务敏感的提示融合策略涉及 (1) 任务敏感的提示学习损失函数 $\mathcal{L}_1$ 以及 (2) 任务敏感的任务标识识别损失函数 $\mathcal{L}_2$. 该模块的最终优化目标即优化 $\mathcal{L}_1 + \mathcal{L}_2$.

4.2 非线性残差门控机制

本文将 MoE 结构拓展到开放世界持续学习的问题中, 并结合使用任务敏感提示增强后的样本嵌入, 进一步改进传统的线性门控函数的连接机制, 用于更有效地进行专家网络选择 (expert routing), 提升模型在面对不同任务的开放样本和已知样本的预测能力.

正如第 3 节中所提及的, MoE 结构引入了门控函数在给定输入的条件下对多个子网络输出进行加权聚合. 被选中的专家对输入进行独立变换, 其输出经过门控权重加权后求和, 作为该层的最终输出. 该结构允许模型在保持计算成本不变的情况下显著提升参数规模与表达能力.

具体地, 给定任务 t , 通过优化 $\mathcal{L}_1 + \mathcal{L}_2$ 得到了任务敏感的提示之后, 通过 MoE 中的注意力模块将提示与样本嵌入进行融合, 对样本进行知识增强. 在计算专家选择得分函数时, 我们引入了简单有效的非线性残差用以提升传统的门控函数.

$$\text{expert score}_{i,m} = \frac{x_i^T \cdot W_l^Q \cdot W_l^{K^T} \cdot p_i}{\sqrt{d_v}} + \alpha \cdot \sigma \left(\beta \cdot \frac{x_i^T \cdot W_l^Q \cdot W_l^{K^T} \cdot p_i}{\sqrt{d_v}} \right) \quad (12)$$

其中, x_i 表示任务 t 中的第 i 个样本; W_l^Q 和 W_l^K 分别表示注意力层中的第 l 个注意力头对应的查询矩阵和键矩阵, $\sqrt{d_v}$ 为缩放参数, α 和 β 均为可学习的标量参数. 并且, 我们对激活函数 σ 的选用也尝试了多种不同的激活函数进行评估, 相应的消融对比结果已展示在第 5 节中.

在训练过程中, MoE 模块的优化目标与主任务分类模块联合进行. 在每一次前向传播中, 模型首先根据公式 (12) 中定义的任务敏感非线性门控函数 $\text{expert score}_{i,m}$ 为每个输入分配对应的提示向量及其在注意力空间中的权重, 并通过 Top- k 稀疏激活策略来激活前 k 个专家网络. 在训练中, 为确保所有专家在训练过程中被充分利用, 我们在训练阶段还引入了辅助负载均衡损失 (auxiliary load balancing loss)^[35], 记作 $\mathcal{L}_3$. 该损失鼓励门控网络在不同样本间公平分配专家资源, 从而避免部分专家长期未被激活而退化. 辅助负载均衡损失函数可表示为:

$$\mathcal{L}_3 = \lambda_{\text{load}} \left\{ \frac{1}{N} \sum_{j=1}^N \left(f_j - \frac{1}{N} \sum_{k=1}^N f_k \right)^2 \right\} \quad (13)$$

其中, f_j 表示第 j 个专家在当前 batch 中被激活的相对频率, N 为专家总数, λ_{load} 为负载均衡损失的权重超参数.

因此, $\mathcal{L}_3$ 本质上是对所有专家激活频率的方差进行最小化, 目标是促使所有专家在训练中被均匀调用, 缓解稀疏激活机制可能带来的模式坍塌问题.

需要强调的是, 为保证稀疏门控机制的可训练性, 只有被选中的 Top- k 专家网络及其对应的门控路径参与反向传播, 其余未激活部分在该训练步骤中保持冻结状态. 这种局部反向传播策略确保了在保持高表达容量的同时显著降低梯度计算成本, 使 TP-MoE 能够在大规模开放世界持续学习任务中实现高效优化与动态泛化能力的统一.

同时, 本文所提出 TP-MoE 模型使用预训练 ViT 作为主干网络, 在此基础上, 我们提出的任务敏感的门控函数通过直接作用于注意力层来调节提示池中的不同提示对应不同专家网络的权重, 从而帮助模型更好地区分不同任务中不同提示的语义特征, 从模型架构的角度进一步提升提示融合的效果并显著提升样本使用效率. 通过引入非线性残差连接改进专家门控机制, 有效提升了模型在持续学习分类任务上的表现, 为模型自适应地确定开放检测阈值提供了保障, 进而增强了整体的开放世界持续学习能力.

4.3 自适应的开放检测阈值

在本节中, 我们详细探讨如何在已知样本与未知样本之间构建判别边界, 重点在于如何在开放世界持续学习过程中针对不同的任务均能够在测试推理阶段准确地检测出未知样本. 为应对这一挑战, 提出一种用于自适应地进行开放集检测阈值选择的新方法, 该方法不需要参与训练即能够在测试推理阶段高效地进行开放样本检测, 借助于已学习任务中获得的知識来构建一个紧凑的判别边界, 用于识别未知样本.

受到开放集识别现有工作的启发, 识别未知样本的关键在于充分利用对已知类别的学习, 我们提出的自适应开放决策边界构建方法的核心思想是自适应地选择 *Softmax* 熵上的阈值, 以任务敏感的提示改进后的 MoE 融合样本表征为基础, 利用对已知类别的知識设定合理的开放边界.

具体地, 本文提出了一种基于 logits 分数的任务敏感开放决策边界构建方法, 通过估计测试样本的分类输出 logits 得分来实现判断: 如果某个测试样本的最大 logits 得分不对应任何已学习任务类别, 则自然可以推断该样本属于开放集 (未知类别). 在完成任务 t 的训练后, 将经过提示和多专家输出增强的训练样本标识输入至冻结的预训练主干网络, 最终由一个可训练的分类器 f_v 输出该样本未缩放的 *Softmax* 熵得分.

对于任务 t 中的训练样本 x_i^t , 首先使用通用冻结投影层得到的样本嵌入表征 h_i^t ; 然后, 通过第 4.1 节中公式 (8) 得到当前任务 t 下学习到的一批提示 p_t , 对样本嵌入 h_i^t 进行 token 增强, 我们将最终的增强嵌入记作 $\hat{h}_i^t$; 因此接下来可以得到:

$$\text{Logit Score}_{i,t} = f_v(\hat{h}_i^t) \quad (14)$$

这样, 我们就能够得到该样本的最大 *Softmax* 熵得分, 并通过对任务 t 的所有样本的最大 *Softmax* 熵得分取平均, 来确定通过任务 t 学习到的开放集检测的阈值 η_t :

$$\eta_t = \varsigma \cdot \frac{1}{|\mathcal{X}_t|} \sum_{i=1}^{|\mathcal{X}_t|} \max_i (\text{Logit Score}_{i,t}) \quad (15)$$

其中, $|\mathcal{X}_t|$ 是任务 t 中训练样本的总数量, ς 为可调节的缩放因子超参数.

在推理阶段, 我们使用该阈值来判断一个样本是否属于开放集, 仅需将样本最终得到的最大 *Softmax* 熵得分与阈值 η_t 进行比较: 当样本最大 *Softmax* 熵得分小于阈值时, 则该样本被判定为未知对象, 模型将其标注为“unknown”; 反之, 则该样本被判定为已知对象, 模型则调用训练得到的分类头将其分类到已知类别中. 通过这种高效且即插即用的自适应策略, 模型能够有效支持开放世界持续学习场景下针对未知类别的检测, 避免产生错误分类. 同时, 由于样本经过任务敏感的提示融合和多专家网络输出融合, 能够有效地进行知识传输和更新.

5 实验分析

在本节中, 为了全面且充分地验证所提出的 TP-MoE 模型在开放世界持续学习任务下的表现, 我们对比了两个广为应用且具有挑战性的基准数据集, 即 Split-CIFAR100 和 Open-CORe25. 在这两个数据集的基础上, 本文完整地将开放世界持续学习在两种不同的任务设定下对所提模型进行了验证. 同时, 针对性地解决了两个主要研究问题, 验证 TP-MoE 的效果, 并通过一系列消融实验和参数稳健性检验进一步验证不同模块的有效性和模型的稳健性.

5.1 实验设置与实验数据

第一, 考虑到开放世界持续学习中任务类增量的场景, 记为类增量开放世界持续学习 (class-incremental OWCL, CI-OWCL). 在此场景下, 我们将 Split-CIFAR100 数据集划分为 10 个任务 (每个任务包含 10 个类别), 测试时每个任务的类别无重叠, 即新类别是随着新任务不断出现的, 且在每次训练当前任务时, 其余任务的训练集均不可用; 测试时, 每次都使用所有任务的完整测试集进行评估, 这样既保证了在前 9 个任务中会不断出现未知样本, 又保证了未知样本会随着新任务的学习而不断得到对应的标签数据. 在此设定下, 开放世界持续学习模型不仅需要类增量任务上保持良好的分类效果, 还需要准确地识别没有学习过的开放样本, 同时, 还需要根据不断学习的新类别

更新关于开放样本的知识, 从而做出正确的分类.

第二, 考虑了分布变化的开放世界持续学习场景, 记作域增量开放世界持续学习 (domain-incremental OWCL, DI-OWCL). 在此场景中, 我们在一个经典且具有挑战性的目标识别数据集 Open-CORe50 的基础上, 构建了针对开放世界持续学习的 Open-CORe25 数据集. 该数据集从原有的 50 个已知类别随机地划分出 25 个类别, 作为仅在测试阶段使用的未知类别而不参与训练, 从而保证了开放世界持续学习中最基本的设定; 并且, 同一类别的新训练样本会随着任务的推进按顺序出现, 并伴随分布变化 (例如不同的拍摄条件和背景). 通过这种设定将分布的变化引入为未知的特征变化, 诸如此类的分布变化会给模型带来分布外知识, 极大地挑战了现有开放世界持续学习模型的能力.

最后, 值得一提的是, 我们也考虑了将类增量与域增量相结合的混合型开放世界持续学习场景, 例如在类别不断增长的同时, 特征分布也随环境变化而演化的复杂现实场景. 然而, 为了更清晰地验证本文 TP-MoE 模型在开放世界持续学习中的核心机制有效性, 并确保与现有所复现方法具备公平的可比性, 后续实验选择聚焦于类增量 (CI-OWCL) 和域增量 (DI-OWCL) 这两个基础但具代表性的任务设定, 方便对 TP-MoE 模型的泛化能力和模块组合策略进行系统评估, 也为后续更复杂场景的研究提供基础支撑.

如表 1 所示, 在 CI-OWCL 场景下的 Split-CIFAR100 数据集中, 共包含 90 个已知类别和 10 个未知类别, 总计 10 个任务和 55 000 个样本, 每轮训练 (即每个任务下) 包含 5 000 个样本. 验证集包括来自已知类别的 9 000 个样本和来自未知类别的 1 000 个样本. 在 DI-OWCL 场景下所使用的 Open-CORe25 数据集中, 包含 25 个已知类别和 25 个未知类别, 共 5 个任务, 总样本量为 59 936, 每轮训练的样本数分别如表 1 所示. 验证集包含 7 528 个来自已知类别的样本和 7 466 个来自未知类别的样本.

表 1 数据集描述

Dataset	已知类别数量	开放类别数量	任务数量	训练样本	各个任务下训练样本数	验证时使用的已知样本数量	验证时使用的未知样本数量
Split-CIFAR100	90	10	10	55 000	5 000	9 000	1 000
Open-CORe25	25	25	5	59 936	14 989	7 528	7 466
					14 986		
					14 995		
					14 966		
					14 975		

在实验实现方面, TP-MoE 方法基于 PyTorch 库实现, 所有实验均在单块 NVIDIA RTX 3090 GPU 上进行. 所使用的预训练骨干网络为 ViT-B/16, 该模型采用自监督方式在 ImageNet-21K 上完成预训练, 主干网络的预训练层全部冻结不参与模型优化. 在每个任务的测试阶段, 我们将尚未学习过的类别的测试样本视为开放类. 为了全面地验证所提模型的效果, 本文遵循现有持续学习社区的常用设定并广泛地与基准模型进行对比. 在训练阶段, 对所有数据集均采用了数据预处理操作, 包括随机缩放后裁剪为 224×224 像素, 并进行随机水平翻转. 在验证推理阶段, 除 Split-CIFAR100 外的所有数据集图像均统一缩放其短边至 256 像素, 再中心裁剪为 224×224 像素; Split-CIFAR100 数据集的图像从其原始的 32×32 像素上采样至 224×224 像素, 以确保与预训练 ViT-B/16 骨干网络的结构兼容性.

5.2 评价指标及基准模型

我们使用 3 个关键指标对模型在所有任务上的性能进行全面评估, 分别为 ACC_t 、 AUC_t 和 FPR_t .

具体而言, ACC_t 表示在所有 t 个任务中已知类别上的平均最终分类准确率, 用于衡量模型在避免遗忘方面的表现; AUC_t 表示所有已学的 t 个任务的 ROC 曲线下的平均面积, 用以评估模型在区分已知与未知实例方面的可靠性, 即开放检测能力; FPR_t (表示所有任务下的平均假阳性率) 用于量化开放集检测中的错误率, 表示模型将未知样本误判为已知类别的频率, 即误拒率.

值得注意的是, 由于目前开放世界持续学习模型非常有限, 现有的大多数研究通常将实验验证分为开放检测能力验证和持续学习分类能力验证, 因此, 所使用的对比模型也大多是持续学习模型与开放检测方法的组合. 因此, 我们使用了 5 个大规模预训练与提示调优结合的最新持续学习方法以及 4 个经典的分布外样本检测算法进行

组合, 将 TP-MoE 与这 20 个组合的基线模型进行比较. 另外, 在实验中还重点考虑了 2 个最新的开放世界持续学习模型, 也将基准 MoE 模型纳入了对比. 下面简要介绍所涉及的持续学习方法、分布外样本检测算法和开放世界持续学习模型.

- L2P^[28]: 该方法基于提示的微调机制 (prompt-based fine-tuning), 高效地将预训练模型适配到持续学习任务中, 充分利用了大规模预训练知识.

- DualPrompt^[31]: 该方法提出双重提示调优策略, 分别面向任务特定提示与任务无关提示, 以增强在持续学习场景中的泛化能力.

- CODA-P^[33]: CODA-P 结合对比学习与动态提示机制, 在适应新任务的同时保持已学知识, 有效提升持续学习性能.

- RanPAC^[40]: 这是一种内存高效的持续学习方法, 采用随机部分注意力机制 (randomized partial attention) 来捕捉与任务相关的知识, 同时缓解灾难性遗忘.

- ADAM^[41]: ADAM 重新审视基于注意力机制的持续学习范式, 提出一种利用动态注意力来更好地保持旧任务知识的结构.

- OpenMax^[4]: 这是一种 OOD 检测方法, 通过调整 *Softmax* 输出层识别未知类别, 实现开放集识别 (open-set recognition).

- MaxLogits^[42]: 该方法使用最大 logit 值作为指标来检测分布外样本, 为持续学习提供一种简单可扩展的 OOD 识别技术.

- Entropy^[43]: 该方法利用模型输出分布的熵值作为指标, 用于区分分布内与分布外数据, 进行 OOD 检测.

- EnergyBased^[44]: 该 OOD 方法基于神经网络输出计算能量分数, 以区分分布内样本与未见类别, 提升未知类别的识别能力.

- MORE^[12]: MORE 是一种 OWCL 方法, 提出了一种内存高效的正则化技术, 以更好地应对任务切换并避免遗忘.

- Pro-KT^[9]: Pro-KT 是一种基于原型的持续学习方法, 利用原型结构促进任务间的知识迁移, 有效应对开放世界学习场景.

5.3 主要实验结果与分析

为了评估 TP-MoE 的有效性, 我们针对性地解决了以下两个研究问题.

研究问题 1: TP-MoE 能否在开放世界持续学习的场景下比现有的模型在已知类别的分类上获得更好的结果?

研究问题 2: TP-MoE 能否在开放世界持续学习的场景下比现有的模型在未知类别的检测上获得更好的结果?

为了验证这两个问题的结果, 我们对比了所提出模型 TP-MoE 与第 5.2 节提到的现有主流的持续学习方法在两个基准数据集上使用不同的场景设定下的性能表现. 评估指标包括分类准确率 (ACC_t)、开放检测能力 (AUC_t) 和误拒率 (FPR_t), 对模型在所有任务上的性能进行全面评估.

从表 2 所示的实验结果中可以看出, 在 Split-CIFAR100 数据集上, TP-MoE 显著优于所有对比方法. ACC_t 为 0.876 ± 0.023 , 相比表现次佳的 Pro-KT (0.779 ± 0.010) 高出约 9.7%; AUC_t 达到 0.915 ± 0.016 , 显著领先于其他方法的平均水平, 反映出模型在识别未知类别时具备更强的辨别能力. 此外, TP-MoE 在 FPR_t 上也实现了相对较低的 0.280 ± 0.101 , 相比大多数传统方法 (如 L2P、DualPrompt、CODA-P) 普遍超过 0.9 的高误拒率, 有效减少了对已知类别样本的误拒现象, 表明其在开放环境中保持较高保守性的同时, 仍能维持准确识别. 在更加复杂的 Open-CORE25 数据集上, TP-MoE 依然保持优越的性能, 其 ACC_t 为 0.797 ± 0.087 , 在所有方法中排名第一. AUC_t 达到 0.789 ± 0.245 , 显著高于 Pro-KT (0.675 ± 0.125) 和 MORE (0.641 ± 0.077) 等代表方法, 说明模型在面对大规模、复杂场景下的开放类别具有良好的感知与区分能力. FPR_t 指标则为 0.320 ± 0.074 , 同样低于所有基线方法, 进一步验证了 TP-MoE 在处理开放集样本时具备更低的误拒倾向.

并且, 由于 MoE 模块本身并不具备持续学习与开放识别机制, 我们构造了一个仅使用 MoE 结构的 baseline

(即去除任务敏感的提示融合策略及开放阈值机制), 从实验结果中可以看出, 该方法在持续学习精度和开放检测能力上均表现极差, 说明其缺乏对任务和类别变化的建模能力, 进一步强调了本文所提 TP-MoE 框架中任务敏感的提示融合策略、路由机制与任务感知开放阈值策略的必要性.

表 2 TP-MoE 与基准模型的对比结果

类型	方法	Split-CIFAR100 (CI-OWCL场景)			Open-CORe25 (DI-OWCL 场景)		
		$ACC_t(\uparrow)$	$AUC_t(\uparrow)$	$FPR_t(\downarrow)$	$ACC_t(\uparrow)$	$AUC_t(\uparrow)$	$FPR_t(\downarrow)$
持续学习与开放检测的组合方法	L2P	OpenMax	0.436±0.048	0.133±0.051	0.937±0.046	0.424±0.063	0.244±0.050
		MaxLogits	0.445±0.041	0.103±0.047	0.931±0.045	0.412±0.179	0.259±0.053
		Entropy	0.445±0.028	0.120±0.038	0.932±0.041	0.424±0.063	0.257±0.063
		EnergyBased	0.445±0.038	0.102±0.043	0.925±0.040	0.425±0.065	0.256±0.057
	DualPrompt	OpenMax	0.423±0.031	0.139±0.037	0.928±0.045	0.402±0.029	0.270±0.024
		MaxLogits	0.425±0.034	0.114±0.039	0.933±0.042	0.408±0.113	0.261±0.017
		Entropy	0.425±0.045	0.204±0.032	0.936±0.048	0.408±0.112	0.267±0.026
		EnergyBased	0.425±0.046	0.114±0.049	0.942±0.039	0.408±0.112	0.260±0.030
	CODA-P	OpenMax	0.439±0.036	0.125±0.042	0.936±0.046	0.413±0.032	0.278±0.045
		MaxLogits	0.439±0.047	0.090±0.039	0.942±0.047	0.415±0.033	0.183±0.028
		Entropy	0.439±0.036	0.106±0.053	0.936±0.047	0.415±0.033	0.256±0.056
		EnergyBased	0.439±0.031	0.089±0.053	0.946±0.039	0.415±0.033	0.169±0.022
	ADAM	OpenMax	0.440±0.044	0.112±0.042	0.876±0.043	0.433±0.025	0.283±0.038
		MaxLogits	0.440±0.042	0.099±0.042	0.882±0.038	0.433±0.026	0.238±0.029
		Entropy	0.440±0.042	0.221±0.043	0.876±0.047	0.433±0.025	0.197±0.016
		EnergyBased	0.440±0.035	0.386±0.042	0.869±0.041	0.433±0.027	0.345±0.023
	RanPAC	OpenMax	0.467±0.045	0.091±0.044	0.879±0.042	0.472±0.022	0.201±0.020
		MaxLogits	0.468±0.042	0.092±0.042	0.878±0.039	0.472±0.022	0.212±0.017
		Entropy	0.468±0.049	0.091±0.041	0.871±0.031	0.472±0.022	0.219±0.021
		EnergyBased	0.468±0.046	0.227±0.038	0.869±0.041	0.475±0.021	0.411±0.025
开放世界持续学习方法	MoE	0.212±0.072	0.200±0.030	0.995±0.002	0.120±0.035	0.102±0.031	0.991±0.004
	MORE	0.716±0.011	0.717±0.127	0.492±0.016	<u>0.641±0.077</u>	0.641±0.077	0.517±0.011
	Pro-KT	<u>0.779±0.010</u>	<u>0.745±0.010</u>	<u>0.397±0.026</u>	0.635±0.011	<u>0.675±0.125</u>	<u>0.451±0.015</u>
	TP-MoE	0.876±0.023	0.915±0.016	0.280±0.101	0.797±0.087	0.789±0.245	0.320±0.074

总体来看, TP-MoE 在两个基准数据集上都实现了在准确率、开放集检测能力和误拒控制上的均衡优化, 展现出良好的泛化性能与稳健性. 相比于现有方法通常存在“精度提升伴随高误拒”或“开放检测能力弱”的问题, TP-MoE 在 3 项指标上的优化能力体现了其在开放世界持续学习场景中强大的知识积累和知识迁移能力.

5.4 消融实验

为了深入验证 TP-MoE 模型中关键结构设计的有效性, 我们设计了两组消融实验, 分别针对任务敏感的提示融合模块与非线性残差门控机制进行评估. 这两部分模块共同构成了 TP-MoE 的消融核心, 其设计旨在增强模型对任务切换的适应能力与对类别边界的知识保持能力, 进而缓解开放世界持续学习中常见的灾难性遗忘与泛化退化问题.

表 3 展示了在移除任务敏感提示融合模块、保留残差门控的设置下, TP-MoE 在 Split-CIFAR100 数据集上 10 个任务中的表现. 从表中可以观察到, 模型在各项指标上均表现出良好的稳定性与判别能力, 尤其是在前半部分任务 (任务 1–任务 5) 中, 模型准确率 ACC_t 基本维持在 90% 以上, 其中任务 1 达到最高的 98.70%, 显示出模型在学习初期任务时的强泛化能力和高质量表示学习效果.

此外, 模型在任务切换过程中仍能保持较高的 AUC_t (如任务 1 达到 90.73%), 说明即使在不依赖任务敏感提示的前提下, 模型仍然具备出色的类别判别能力与边界保留能力. 这说明模型的专家选择机制在残差门控调控下仍能一定程度上进行合理路由, 有效维持知识迁移的稳定性. 在反映遗忘程度的指标 FPR_t 上, 模型在整个任务序

列中也展现出较强的抗遗忘能力, 进一步表明 TP-MoE 架构在无显式任务标识输入的条件下仍能保持类间区分的鲁棒性. 值得注意的是, 随着任务数量增加和类别间干扰加剧, 模型的准确率略有波动, 但整体仍保持在相对较高水平 (最低为 87.13%), 而且 AUC_t 并未出现明显下降, 甚至在任务 10 达到了 91.30% 的峰值, 体现出模型后期依然具备稳定的分类能力和抗干扰性能.

综合来看, 该组消融实验结果验证了 TP-MoE 中残差门控机制对于稳定知识迁移和维持判别能力的积极作用, 同时也进一步凸显任务敏感提示模块在提升专家动态选择精度与任务适应能力方面的重要价值.

表 4 展示了在保留提示融合、移除非线性残差门控机制的条件下模型的表现. 尽管该机制仍在一定程度上保留了对任务信息的非线性交互能力, 且在任务 1 的 FPR_t 降至 0.0952, 但前两个指标上整体表现相比表 3 有所下降. 任务间准确率呈现更明显的下滑趋势, 任务 10 的 ACC_t 降至 78.84%, 且 FPR_t 略微升高至 0.2781, 显示出模型在任务持续推进过程中对新旧知识切换的适应性变差, 容易出现类间混淆, 从而进一步说明非线性残差门控机制在模型中起到的关键作用.

表 3 TP-MoE 在消融掉任务敏感的提示融合策略后
在数据集 Split-CIFAR100 各个任务上的表现

消融模块1: 任务敏感的提示融合	$ACC_t (\uparrow) (\%)$	$AUC_t (\uparrow) (\%)$	$FPR_t (\downarrow)$
任务1	98.70	90.73	0.120
任务2	95.55	87.64	0.152
任务3	93.47	88.67	0.150
任务4	92.18	87.24	0.178
任务5	90.46	87.52	0.211
任务6	89.23	89.17	0.223
任务7	88.61	89.27	0.239
任务8	87.89	89.87	0.248
任务9	88.18	89.69	0.272
任务10	87.13	91.30	0.272

表 4 TP-MoE 在消融掉非线性残差门控机制后在数
据集 Split-CIFAR100 各个任务上的表现

消融模块2: 非线性残差门控	$ACC_t (\uparrow) (\%)$	$AUC_t (\uparrow) (\%)$	$FPR_t (\downarrow)$
任务1	96.00	86.98	0.0952
任务2	92.05	79.43	0.1435
任务3	88.57	81.02	0.1403
任务4	86.73	79.85	0.1678
任务5	84.14	79.93	0.1996
任务6	82.75	81.06	0.2149
任务7	81.61	81.00	0.2371
任务8	80.43	81.82	0.2478
任务9	80.33	81.45	0.2781
任务10	78.84	82.74	0.2781

综合两组实验结果可得, 任务敏感的提示融合机制和非线性残差门控机制在 TP-MoE 中起到了至关重要的作用, 其有效引导了专家模块的动态路由与任务偏好的结构适应, 从而提升了整体学习性能与迁移稳定性. 残差门控机制对长期学习稳定性的贡献也十分显著. 这进一步验证了我们提出的 TP-MoE 架构在结构设计上的合理性, 即通过提示驱动的专家激活与非线性交互协同作用, 实现对开放世界持续学习场景中多任务的有效泛化.

5.5 参数敏感性分析

为进一步验证 TP-MoE 模型在不同参数配置下的稳定性与适应能力, 本文在 Split-CIFAR100 与 Open-CORE25 两个开放世界持续学习数据集上开展了系统性的超参数敏感性分析. 具体而言, 我们考察了 3 类核心超参数对模型性能的影响: 其一是提示长度与提示个数, 通过调节提示维度和并行提示数量, 分析其对表示能力和任务间迁移性能的影响; 其二是非线性门控机制中使用的激活函数类型, 包括 ReLU、GELU、Sigmoid 和 tanh 等, 以评估不同激活函数对专家选择稳定性与路由精度的影响; 其三是门控阈值的设定策略, 比较了静态设定与动态自适应机制在控制专家稀疏激活过程中的表现, 尤其是在任务边界模糊或概念漂移频繁的场景下的鲁棒性. 为此, 在本节中, 我们全面评估了 TP-MoE 在关键结构参数变化下的泛化能力与性能波动, 为后续模型部署与扩展提供了实证依据. 相关结果如图 3、表 5 和图 4 所示.

图 3 展示了不同提示长度 ($\text{prompt length} \in \{5, 10, 15, 20\}$) 与提示个数 ($\text{prompt number} \in \{1, 3, 5, 7\}$) 组合下, TP-MoE 在两个数据集上的平均准确率表现. 从图 3(a) 的 Split-CIFAR100 数据集上使用不同的提示长度和提示个数对 TP-MoE 的性能影响可以看出, 模型对提示个数较为敏感, 当提示个数从 1 增加至 3 或 5 时性能迅速提升, 平均准确率由 84.67% 上升至 87.35% 和 87.50%, 随后在 7 个提示时趋于饱和. 同时, 提示长度的变化对性能略有影

响, 但整体相对稳定, 长度设为 15 或 20 时通常可获得更优解, 反映出较长提示向量在复杂任务场景下具有更强的特征学习能力.

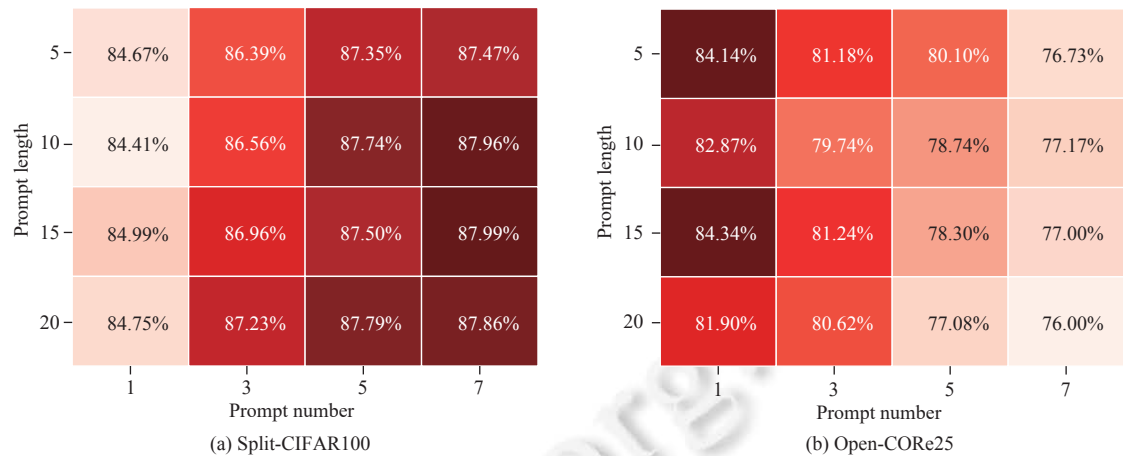


图 3 不同提示长度与不同的提示个数对 TP-MoE 的性能变化热力图

表 5 不同激活函数对 TP-MoE 的性能影响

激活函数	Split-CIFAR100			Open-CORe25		
	$ACC_t(\uparrow)(\%)$	$AUC_t(\uparrow)$	$FPR_t(\downarrow)$	$ACC_t(\uparrow)(\%)$	$AUC_t(\uparrow)$	$FPR_t(\downarrow)$
Sigmoid	87.69	0.917	0.273	77.05	0.764	0.301
ReLU	86.95	0.910	0.278	78.96	0.748	0.289
GELU	87.04	0.909	0.271	79.52	0.761	0.286
tanh	87.61	0.915	0.280	79.78	0.789	0.320

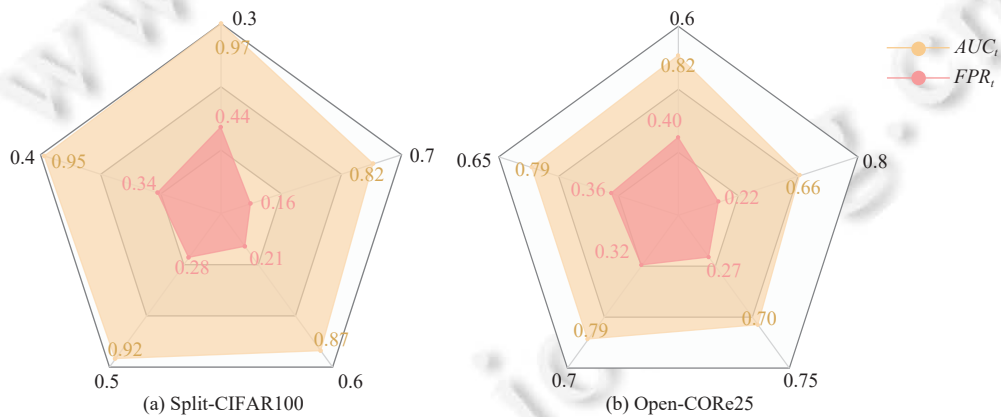


图 4 阈值对 TP-MoE 的开放检测效果影响雷达图

而在图 3(b) 的 Open-CORe25 数据集上使用不同的提示长度和提示个数对 TP-MoE 的性能影响图中, 提示长度和提示个数的组合对模型性能的影响更加显著. 尤其在提示个数增多时, 模型精度逐步下降, 提示个数为 7 时最低的模型精度为 76.00% (提示长度为 20), 说明过多的提示可能在特定领域的任务表征中引发干扰和泛化退化. 相较之下, 较短的提示长度 (5 或 10) 配合少量提示个数 (1 或 3) 能获得更稳健的性能, 表明提示机制在不同模态下应进行适配调整.

表 5 系统比较了不同激活函数 (Sigmoid, ReLU, GELU, tanh) 对 TP-MoE 模型性能的影响. Split-CIFAR100 中

的结果显示, 各激活函数下模型的准确率 ACC_i 和 AUC_i 指标均处于较高水平 (均超过 87% 和 0.91), 其中 $\tanh$ 表现略优, ACC_i 达到 87.47%, FPR_i 维持在 0.281 以下, 说明模型对激活函数的鲁棒性较强; 同时, Sigmoid 与 GELU 也展现出良好的泛化能力, 表现差异不大. Open-CORe25 实验结果中, 模型对激活函数更加敏感, $\tanh$ 虽在准确率上略高 (79.78%), 但其 FPR_i 达到 0.32, 显著高于 ReLU (0.29) 与 GELU (0.29), 提示在开放音频任务中过度平滑的激活函数可能导致过拟合边界. 总体来看, TP-MoE 对激活函数选择较为稳健.

图 4 展示了不同阈值设定对 TP-MoE 模型开放检测能力的影响. 图 4(a) 为 Split-CIFAR100 数据集上使用不同阈值设定下 TP-MoE 的开放检测效果, 图 4(b) 为 Open-CORe25 数据集上使用不同阈值设定下 TP-MoE 的开放检测效果.

从图 4(a) 可以观察到, 当阈值设定为 0.3 或 0.4 时, AUC_i 达到峰值 (0.97 和 0.95), 说明低阈值更有利于模型识别开放类样本. 而随着阈值上升, 模型更趋于保守, FPR_i 虽降低 (如 0.21@0.6), 但同时伴随显著的性能衰减 (AUC_i 降至 0.87@0.6), 表明存在明显的性能-风险权衡. 在图 4(b) 中, 类似趋势也得以体现, 最佳检测性能出现在阈值为 0.6 附近, 整体表现相对平稳. 总体看来, TP-MoE 在不同阈值设定下表现出良好的渐变响应与策略稳定性, 验证其在开放空间中具备良好的泛化能力.

综上, TP-MoE 模型在多项关键超参数 (提示结构、激活函数、阈值控制) 上的表现总体稳健, 且对不同数据集和任务设定均具备良好的自适应能力, 展现出较高的容错性与泛化性. 上述结果为未来 TP-MoE 的迁移扩展与部署应用提供了重要参考依据.

6 总 结

本文围绕开放世界持续学习中面临的分布漂移、类别开放不确定与判别泛化难的问题, 本文对开放世界持续学习进行了形式化描述, 并提出了一种任务敏感提示驱动的混合专家模型 (TP-MoE). 该模型通过引入任务敏感提示融合机制与非线性稀疏激活的专家结构, 实现了对不同任务语义的显式建模与知识路由调度, 有效保证了模型在处理开放世界持续学习任务时对已知类别的增量学习能力和对未知样本的检测能力. 本文在两个开放世界持续学习基准数据集上与 23 个组合基准模型进行了对比实验, 实验结果表明, TP-MoE 在多个指标上均显著优于现有方法. 消融实验验证了提示融合与门控机制的有效性, 而参数稳健性分析进一步揭示了模型对提示结构、激活函数与阈值变化的良好适应能力, 表现出高度可控性与迁移泛化性. 总体而言, TP-MoE 为开放世界持续学习中的无回放任务感知建模提供了新思路, 展示出在多模态、动态开放环境下进行自适应性的知识传输、知识判别调度和知识更新的能力.

References

- [1] Ke ZX, Liu B, Huang XC. Continual learning of a mixed sequence of similar and dissimilar tasks. In: Proc. of the 34th Int'l Conf. on Neural Information Processing Systems. Vancouver: Curran Associates Inc., 2020. 1553.
- [2] Chu ZX, Li RP, Rathbun S, Li S. Continual causal inference with incremental observational data. In: Proc. of the 39th IEEE Int'l Conf. on Data Engineering. Anaheim: IEEE, 2023. 3430–3439. [doi: 10.1109/ICDE55515.2023.00263]
- [3] Wang HM, Zhang YY, Yang YF, Zheng YF, Wong KF. Acquiring new knowledge without losing old ones for effective continual dialogue policy learning. IEEE Trans. on Knowledge and Data Engineering, 2024, 36(12): 7569–7584. [doi: 10.1109/TKDE.2023.3344727]
- [4] Bendale A, Boulton T. Towards open world recognition. In: Proc. of the 2015 IEEE Conf. on Computer Vision and Pattern Recognition. Boston: IEEE, 2015. 1893–1902. [doi: 10.1109/CVPR.2015.7298799]
- [5] Joseph KJ, Khan S, Khan FS, Balasubramanian VN. Towards open world object detection. In: Proc. of the 2021 IEEE/CVF Conf. on Computer Vision and Pattern Recognition. Nashville: IEEE, 2021. 5826–5836. [doi: 10.1109/CVPR46437.2021.00577]
- [6] Liu B, Mazumder S, Robertson E, Grigsby S. AI autonomy: Self-initiated open-world continual learning and adaptation. AI Magazine, 2023, 44(2): 185–199. [doi: 10.1002/aaai.12087]
- [7] Guo YD, Liu B, Zhao DY. Dealing with cross-task class discrimination in online continual learning. In: Proc. of the 2023 IEEE/CVF Conf. on Computer Vision and Pattern Recognition. Vancouver: IEEE, 2023. 11878–11887. [doi: 10.1109/CVPR52729.2023.01143]

- [8] Fei GL, Wang S, Liu B. Learning cumulatively to become more knowledgeable. In: Proc. of the 22nd ACM SIGKDD Int'l Conf. on Knowledge Discovery and Data Mining. San Francisco: ACM Press, 2016. 1565–1574. [doi: [10.1145/2939672.2939835](https://doi.org/10.1145/2939672.2939835)]
- [9] Li YJ, Yang X, Wang H, Wang XK, Li TR. Learning to prompt knowledge transfer for open-world continual learning. In: Proc. of the 38th AAAI Conf. on Artificial Intelligence. Vancouver: Association for the Advancement of Artificial Intelligence, 2024. 13700–13708. [doi: [10.1609/aaai.v38i12.29275](https://doi.org/10.1609/aaai.v38i12.29275)]
- [10] Kim G, Xiao CN, Konishi T, Ke ZX, Liu B. Open-world continual learning: Unifying novelty detection and continual learning. Artificial Intelligence, 2025, 338: 104237. [doi: [10.1016/j.artint.2024.104237](https://doi.org/10.1016/j.artint.2024.104237)]
- [11] Dang SH, Cao ZJ, Cui ZY, Pi YM, Liu NY. Open set incremental learning for automatic target recognition. IEEE Trans. on Geoscience and Remote Sensing, 2019, 57(7): 4445–4456. [doi: [10.1109/TGRS.2019.2891266](https://doi.org/10.1109/TGRS.2019.2891266)]
- [12] Kim G, Xiao CN, Konishi T, Ke ZX, Liu B. A theoretical study on solving continual learning. In: Proc. of the 36th Int'l Conf. on Neural Information Processing Systems. New Orleans: Curran Associates Inc., 2022. 366.
- [13] De Lange M, Aljundi R, Masana M, Parisot S, Jia X, Leonardis A, Slabaugh G, Tuytelaars T. A continual learning survey: Defying forgetting in classification tasks. IEEE Trans. on Pattern Analysis and Machine Intelligence, 2022, 44(7): 3366–3385. [doi: [10.1109/TPAMI.2021.3057446](https://doi.org/10.1109/TPAMI.2021.3057446)]
- [14] Kar S, Castellucci G, Filice S, Malmasi S, Rokhlenko O. Preventing catastrophic forgetting in continual learning of new natural language tasks. In: Proc. of the 28th ACM SIGKDD Conf. on Knowledge Discovery and Data Mining. Washington: ACM Press, 2022. 3137–3145. [doi: [10.1145/3534678.3539169](https://doi.org/10.1145/3534678.3539169)]
- [15] Gao Q, Luo ZP, Klabjan D, Zhang FL. Efficient architecture search for continual learning. IEEE Trans. on Neural Networks and Learning Systems, 2023, 34(11): 8555–8565. [doi: [10.1109/TNNLS.2022.3151511](https://doi.org/10.1109/TNNLS.2022.3151511)]
- [16] Mallya A, Lazebnik S. PackNet: Adding multiple tasks to a single network by iterative pruning. In: Proc. of the 2018 IEEE/CVF Conf. on Computer Vision and Pattern Recognition. Salt Lake City: IEEE, 2018. 7765–7773. [doi: [10.1109/CVPR.2018.00810](https://doi.org/10.1109/CVPR.2018.00810)]
- [17] Qin Q, Peng H, Hu WP, Zhao DY, Liu B. BNS: Building network structures dynamically for continual learning. In: Proc. of the 35th Int'l Conf. on Neural Information Processing Systems. La Jolla: Curran Associates Inc., 2021. 1576.
- [18] Kirkpatrick J, Pascanu R, Rabinowitz N, Veness J, Desjardins G, Rusu AA, Milan K, Quan J, Ramalho T, Grabska-Barwinska A, Hassabis D, Clopath C, Kumaran D, Hadsell R. Overcoming catastrophic forgetting in neural networks. Proc. of the National Academy of Sciences of the United States of America, 2017, 114(13): 3521–3526. [doi: [10.1073/pnas.1611835114](https://doi.org/10.1073/pnas.1611835114)]
- [19] Wu Y, Chen YP, Wang LJ, Ye YC, Liu ZC, Guo YD, Fu Y. Large scale incremental learning. In: Proc. of the 2019 IEEE/CVF Conf. on Computer Vision and Pattern Recognition. Long Beach: IEEE, 2019. 374–382. [doi: [10.1109/CVPR.2019.00046](https://doi.org/10.1109/CVPR.2019.00046)]
- [20] Wang LY, Lei B, Li Q, Su H, Zhu J, Zhong Y. Triple-memory networks: A brain-inspired method for continual learning. IEEE Trans. on Neural Networks and Learning Systems, 2022, 33(5): 1925–1934. [doi: [10.1109/TNNLS.2021.3111019](https://doi.org/10.1109/TNNLS.2021.3111019)]
- [21] Wu YQ, Xie RB, Zhu YC, Zhuang FZ, Zhang X, Lin LY, He Qing. Personalized prompt for sequential recommendation. IEEE Trans. on Knowledge and Data Engineering, 2024 36(7): 3376–3389.
- [22] Chaudhry A, Ranzato MA, Rohrbach M, Elhoseiny M. Efficient lifelong learning with A-GEM. In: Proc. of the 7th Int'l Conf. on Learning Representations. New Orleans: OpenReview.net, 2019. 1–20.
- [23] Li ZZ, Hoiem D. Learning without forgetting. IEEE Trans. on Pattern Analysis and Machine Intelligence, 2018, 40(12): 2935–2947. [doi: [10.1109/TPAMI.2017.2773081](https://doi.org/10.1109/TPAMI.2017.2773081)]
- [24] Yang R, Wang S, Zhang H, Xu SY, Guo YH, Ye XT, Hou B, Jiao LC. Knowledge decomposition and replay: A novel cross-modal image-text retrieval continual learning method. In: Proc. of the 31st ACM Int'l Conf. on Multimedia. Ottawa: ACM Press, 2023. 6510–6519. [doi: [10.1145/3581783.3612207](https://doi.org/10.1145/3581783.3612207)]
- [25] Lao QC, Jiang X, Havaei M, Bengio Y. A two-stream continual learning system with variational domain-agnostic feature replay. IEEE Trans. on Neural Networks and Learning Systems, 2022, 33(9): 4466–4478. [doi: [10.1109/TNNLS.2021.3057453](https://doi.org/10.1109/TNNLS.2021.3057453)]
- [26] Ke ZX, Liu B, Ma NZ, Xu H, Shu L. Achieving forgetting prevention and knowledge transfer in continual learning. In: Proc. of the 35th Int'l Conf. on Neural Information Processing Systems. La Jolla: Curran Associates Inc., 2021. 1719.
- [27] Rebuffi SA, Kolesnikov A, Sperl G, Lampert CH. iCaRL: Incremental classifier and representation learning. In: Proc. of the 2017 IEEE/CVF Conf. on Computer Vision and Pattern Recognition. Honolulu: IEEE, 2017. 5533–5542. [doi: [10.1109/CVPR.2017.587](https://doi.org/10.1109/CVPR.2017.587)]
- [28] Wang ZF, Zhang ZZ, Lee CY, Zhang H, Sun RX, Ren XQ, Su GL, Perot V, Dy J, Pfister T. Learning to prompt for continual learning. In: Proc. of the 2022 IEEE/CVF Conf. on Computer Vision and Pattern Recognition. New Orleans: IEEE, 2022. 139–149. [doi: [10.1109/CVPR52688.2022.00024](https://doi.org/10.1109/CVPR52688.2022.00024)]
- [29] Le M, Nguyen A, Nguyen H, Nguyen T, Pham T, Van Ngo L, Ho N. Mixture of experts meets prompt-based continual learning. In: Proc. of the 38th Int'l Conf. on Neural Information Processing Systems. Vancouver: Curran Associates Inc., 2024. 3781.

- [30] Vaswani A, Bengio S, Brevdo E, Chollet F, Gomez A, Gouws S, Jones L, Kaiser Ł, Kalchbrenner N, Parmar N, Sepassi R, Shazeer N, Uszkoreit J. Tensor2Tensor for neural machine translation. In: Proc. of the 13th Conf. of the Association for Machine Translation in the Americas (Vol. 1: Research Track). Boston: Association for Machine Translation in the Americas, 2018. 193–199.
- [31] Wang ZF, Zhang ZZ, Ebrahimi S, Sun RX, Zhang H, Lee CY, Ren XQ, Su GL, Perot V, Dy J, Pfister T. DualPrompt: Complementary prompting for rehearsal-free continual learning. In: Proc. of the 17th European Conf. on Computer Vision. Tel Aviv: Springer, 2022. 631–648. [doi: [10.1007/978-3-031-19809-0_36](https://doi.org/10.1007/978-3-031-19809-0_36)]
- [32] Wang YB, Huang ZW, Hong XP. S-prompts learning with pre-trained Transformers: An Occam's razor for domain incremental learning. In: Proc. of the 36th Int'l Conf. on Neural Information Processing Systems. New Orleans: Curran Associates Inc., 2022. 411.
- [33] Smith JS, Karlinsky L, Gutta V, Cascante-Bonilla P, Kim D, Arbelle A, Panda R, Feris R, Kira Z. CODA-prompt: Continual decomposed attention-based prompting for rehearsal-free continual learning. In: Proc. of the 2023 IEEE/CVF Conf. on Computer Vision and Pattern Recognition. Vancouver: IEEE, 2023. 11909–11919. [doi: [10.1109/CVPR52729.2023.01146](https://doi.org/10.1109/CVPR52729.2023.01146)]
- [34] Wang LY, Xie JY, Zhang XX, Huang MY, Su H, Zhu J. Hierarchical decomposition of prompt-based continual learning: Rethinking obscured sub-optimality. In: Proc. of the 37th Int'l Conf. on Neural Information Processing Systems. New Orleans: Curran Associates Inc., 2023. 3022.
- [35] Shazeer N, Mirhoseini A, Maziarz K, Davis A, Le QV, Hinton GE, Dean J. Outrageously large neural networks: The sparsely-gated mixture-of-experts layer. In: Proc. of the 5th Int'l Conf. on Learning Representations. Toulon: OpenReview.net, 2017. 1–19.
- [36] Lepikhin D, Lee H, Xu YZ, Chen DH, Firat O, Huang YP, Krikun M, Shazeer N, Chen ZF. GShard: Scaling giant models with conditional computation and automatic sharding. In: Proc. of the 9th Int'l Conf. on Learning Representations. Appleton: OpenReview.net, 2020. 1–23.
- [37] Fedus W, Zoph B, Shazeer N. Switch Transformers: Scaling to trillion parameter models with simple and efficient sparsity. The Journal of Machine Learning Research, 2022, 23(1): 120.
- [38] Vaswani A, Shazeer N, Parmar N, Uszkoreit J, Jones L, Gomez AN, Kaiser Ł, Polosukhin I. Attention is all you need. In: Proc. of the 31st Int'l Conf. on Neural Information Processing Systems. Long Beach: Curran Associates Inc., 2017. 6000–6010.
- [39] Dosovitskiy A, Beyer L, Kolesnikov A, Weissenborn D, Zhai XH, Unterthiner T, Dehghani M, Minderer M, Heigold G, Gelly S, Uszkoreit J, Houlsby N. An image is worth 16x16 words: Transformers for image recognition at scale. In: Proc. of the 9th Int'l Conf. on Learning Representations. Appleton: OpenReview.net, 2021. 1–21.
- [40] McDonnell MD, Gong D, Parvaneh A, Abbasnejad E, van den Hengel A. RanPAC: Random projections and pre-trained models for continual learning. In: Proc. of the 37th Int'l Conf. on Neural Information Processing Systems. New Orleans: Curran Associates Inc., 2023. 526.
- [41] Zhou DW, Cai ZW, Ye HJ, Zhan DC, Liu ZW. Revisiting class-incremental learning with pre-trained models: Generalizability and adaptivity are all you need. Int'l Journal of Computer Vision, 2025, 133(3): 1012–1032. [doi: [10.1007/s11263-024-02218-0](https://doi.org/10.1007/s11263-024-02218-0)]
- [42] Hendrycks D, Basart S, Mazeika M, Zou A, Kwon J, Mostajabi M, Steinhardt J, Song D. Scaling out-of-distribution detection for real-world settings. In: Proc. of the 39th Int'l Conf. on Machine Learning. Baltimore: PMLR, 2022. 8759–8773.
- [43] Chan RB, Rottmann M, Gottschalk H. Entropy maximization and meta classification for out-of-distribution detection in semantic segmentation. In: Proc. of the 2021 IEEE/CVF Int'l Conf. on Computer Vision. Montreal: IEEE, 2021. 5108–5117. [doi: [10.1109/ICCV48922.2021.00508](https://doi.org/10.1109/ICCV48922.2021.00508)]
- [44] Liu WT, Wang XY, Owens JD, Li YX. Energy-based out-of-distribution detection. In: Proc. of the 34th Int'l Conf. on Neural Information Processing Systems. Vancouver: Curran Associates Inc., 2020. 1802.

作者简介

李昱洁, 博士生, CCF 学生会会员, 主要研究领域为深度学习, 持续学习, 开放世界持续学习, 智能金融, 数据挖掘。

吴晗, 本科生, CCF 学生会会员, 主要研究领域为深度学习, 持续学习, 开放世界持续学习。

孟丹, 博士, 教授, 博士生导师, 主要研究领域为智能金融, 智能决策, 不确定性信息处理, 数据挖掘。

李天瑞, 博士, 教授, 博士生导师, CCF 杰出会员, 主要研究领域为智能信息处理, 数据挖掘, 云计算, 大数据, 粗糙集与粒计算。

杨新, 博士, 教授, 博士生导师, CCF 杰出会员, 主要研究领域为机器学习, 联邦学习, 持续学习, 数据挖掘, 金融科技。