

动态词典驱动连续命名实体识别

——融合密度采样与自适应权重的双重抗遗忘机制

顾苗苗^{ID}

上海理工大学健康科学与工程学院, 上海

收稿日期: 2025年8月30日; 录用日期: 2025年9月22日; 发布日期: 2025年9月30日

摘要

连续命名实体识别(Continual Named Entity Recognition, CNER)作为自然语言处理领域的新兴研究方向, 致力于解决模型在顺序学习新增实体类型时的核心挑战。与传统NER不同, CNER要求模型在仅访问当前阶段标注数据的情况下, 持续扩展其识别能力, 同时避免对已学实体类型的遗忘。然而, 这一过程面临着严峻的灾难性遗忘问题——当模型学习新实体类型时, 旧类型的识别性能会显著下降。这一现象在CNER任务中尤为突出, 因为在增量学习过程中, 历史步骤中的旧实体类型会被强制重新标注为“O”(非实体)类别。针对上述挑战, 本文提出了一种融合动态词典与自适应权重调整的创新框架。该框架通过双重机制协同应对语义漂移和灾难性遗忘问题: 首先, 基于特征空间可视化分析, 采用密度敏感的采样策略构建动态词典, 在训练过程中智能补充关键样本以维持特征空间的完整性; 其次, 设计了一种考虑实体分布特性和历史学习表现的动态权重策略, 有效平衡新旧知识的学习强度。在CoNLL2003、I2B2和OntoNotes5三个基准数据集上的实验结果表明, 该方法显著提升了模型在持续学习环境下的综合性能。在涵盖十个不同增量学习任务的测试中, 该方法实现了Macro-F1平均8.47%的提升, 其中对低频实体的识别改进尤为显著。消融研究验证了各组件对最终性能的贡献, 特征分析进一步揭示了该方法在维持特征空间稳定性方面的独特优势。

关键词

连续学习, 命名实体识别, 灾难性遗忘, 语义漂移

Dynamic Dictionary-Driven Continual Named Entity Recognition

—A Dual Anti-Forgetting Mechanism Integrating Density-Aware Sampling and Adaptive Weighting

Miaomiao Gu^{ID}

School of Health Science and Engineering, University of Shanghai for Science and Technology, Shanghai

Received: August 30, 2025; accepted: September 22, 2025; published: September 30, 2025

文章引用: 顾苗苗. 动态词典驱动连续命名实体识别[J]. 建模与仿真, 2025, 14(10): 39-53.
DOI: 10.12677/mos.2025.1410604

Abstract

Continual Named Entity Recognition (CNER), an emerging research direction in the field of natural language processing, is dedicated to addressing the core challenges that models face when sequentially learning new entity types. Unlike traditional NER, CNER requires models to continuously expand their recognition capabilities while only having access to the annotated data of the current stage, simultaneously avoiding forgetting previously learned entity types. However, this process is fraught with the severe problem of catastrophic forgetting—when models learn new entity types, the performance on old types significantly decreases. This phenomenon is particularly pronounced in CNER tasks because, during the incremental learning process, old entity types from historical steps are forcibly relabeled as the “O” (non-entity) category. To address the aforementioned challenges, this paper proposes an innovative framework that integrates a dynamic dictionary with adaptive weight adjustment. This framework addresses semantic drift and catastrophic forgetting through a dual mechanism: Firstly, based on feature space visualization analysis, a density-sensitive sampling strategy is employed to construct a dynamic dictionary, intelligently supplementing key samples during the training process to maintain the integrity of the feature space; Secondly, a dynamic weight strategy that considers the distribution characteristics of entities and historical learning performance is designed to effectively balance the learning intensity of new and old knowledge. Experimental results on the CoNLL2003, I2B2, and OntoNotes three benchmark datasets demonstrate that this method significantly enhances the model’s comprehensive performance in a continual learning environment. In tests covering ten different incremental learning tasks, this approach achieved an average Macro-F1 improvement of 8.47%, with particularly significant improvements in the recognition of low-frequency entities. Ablation studies confirm the contribution of each component to the final performance, and feature analysis further reveals the unique advantage of this method in maintaining the stability of the feature space.

Keywords

Continual Learning, Named Entity Recognition, Catastrophic Forgetting, Semantic Drift

Copyright © 2025 by author(s) and Hans Publishers Inc.

This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Open Access

1. 引言

命名实体识别(NER)是自然语言理解(NLU)中的一个重要研究领域。其目的是为序列中的每个标记分配多个实体类型或非实体类型[1]。最近,预训练语言模型(PLMs)的出现[2]将 NER 带入了一个新的领域。传统上,NER 在一个范例中操作,其中令牌被分类为一些固定的实体类型(如组织、个人等),NER 模型经历一次性学习过程。然而,更现实的场景要求 NER 模型在新实体类型出现时不断识别它们,而不需要完全重新训练。这种范式称为连续命名实体识别(CNER),由于其实际应用是有前途的,已经获得了大量的研究关注[3][4]。一个恰当的例子是像 Siri 和小艾这样的语音助手,它们经常需要提取新的实体类型(如 Genre、Actor)来理解新的用户意图(例如, GetMovie)。当前 CNER 领域已发展出多种有效方法来应对这两个问题。知识蒸馏框架作为基础方案,通过教师-学生模型架构传递旧知识[5],同时扩展模型容量学习新类型,如 AddNER 保留原有输出层并新增处理层,Extend NER 则直接扩展输出维度[3]。更先进的 CPFD 方法创新性地结合池化特征蒸馏和置信度伪标签策略,既平衡模型稳定性与可塑性,又有效处理非

实体类型的语义偏移问题[6]。权重调优与融合策略通过动态调整学习率和参数融合比例，以模型无关方式实现知识保留[7]。概念驱动方法则引入形式概念分析构建知识结构，设计多层次蒸馏方案[8]。此外，学习-复习框架采用两阶段策略，首先生成包含旧类型的合成样本，再进行二次蒸馏，有效缓解类型混淆问题[9]。这些方法通过不同技术路线，共同推动了 CNER 领域的发展。

CNER 核心解决两个问题，一个是学习新的实体，一个是保持对旧实体的识别能力。现有的方法普遍面临两大核心挑战：第一个挑战是语义偏移问题，其表现形式主要包含两个方面：非实体语义偏移和实体语义偏移。其中非实体语义偏移问题是 CNER 特有的任务设定，如图 1 所示，传统 NER 中明确的 O 标签语义被解构为三重内涵：真实非实体、已学旧实体和未学新实体的复杂混合体。这种语义混杂导致模型在每个增量阶段都面临重新定义决策边界的挑战。针对这一难题，研究者们提出了多层次解决方案：[3]开创性地将知识蒸馏引入 CNER，通过软目标传递保留旧实体知识；[9]设计的“学习-复习”框架通过合成样本扩充训练数据，明确区分不同语义成分；[4]从因果推理视角提取 Other-Class 中的因果效应；[6]则创新性地采用中位数熵阈值筛选高置信度预测。然而这些方法在伪标签质量控制方面仍存在明显局限，特别是对低频实体和未来实体干扰的处理效果有限。另一方面，实体语义偏移问题则表现为更深层的特征空间退化现象，如图 1 所示。在增量学习过程中，旧实体特征经历着不可逆的系统性畸变：簇心偏移、密度降低和边界模糊等效应不断侵蚀着关键的判别特征。对此，[8]开发的形式概念分析(FCA)方法通过构建概念格显式建模标签关系；[6]提出的池化特征蒸馏(PFD)则巧妙平衡了特征空间的稳定性与可塑性。但必须指出，这些方法仅能延缓而无法根治特征退化问题，特别是在多步增量后，早期实体的特征表示仍会出现严重劣化，且计算开销较大的问题也亟待优化。

第二个挑战是持续学习的固有难题——灾难性遗忘[10]-[12]。在 CNER 中展现出新的特性。模型参数在优化新实体识别时，会不可逆地覆盖对旧实体识别至关重要的网络权重，这种现象因 CNER 中极端的类别不平衡而加剧。当前研究形成了三大解决路径：Extend NER 通过冻结部分参数和扩展输出层来保护旧知识[3]；CFNER 创新性地引入自适应权重机制[13]。这些方法虽然在一定程度上缓解了遗忘问题，但仍存在稳定性偏向，导致新实体学习不足，特别是对少样本实体的保护效果不佳，模型在实体数量差异显著时的表现仍有很大提升空间。

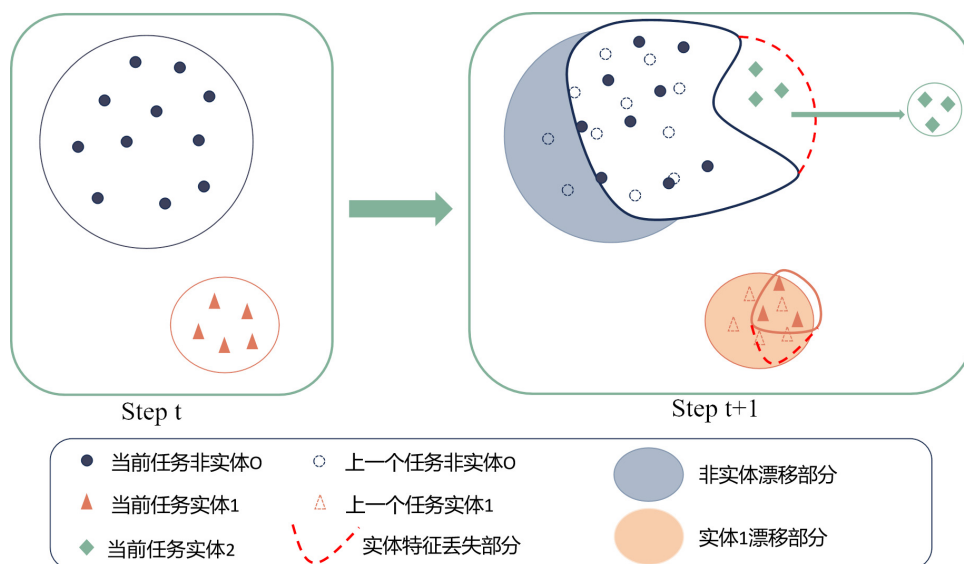


Figure 1. Semantic drift and feature attenuation in continual named entity recognition
图 1. 连续命名实体识别中的语义漂移与特征衰减

本文提出了一种名为动态特征保存与动态调整权重的新方法，该方法通过动态记忆库构建和权重调整机制协同缓解连续命名实体识别中的灾难性遗忘和特征偏移问题。利用历史模型信息在两个关键维度进行优化。首先，本文设计了基于特征空间分析的动态词典机制，通过在每轮训练后可视化模型特征分布(t-SNE 降维)，采用密度敏感采样策略选择关键样本加入动态词典，这些样本在后续训练中与新数据联合优化，有效修复特征覆盖盲区；其次，引入类型自适应的权重调节策略，根据记忆库与当前数据的类型比例动态调整损失权重。本文的核心贡献在于：

1) 提出基于密度采样的动态记忆库机制，通过记忆库补充关键样本，缓解实体语义漂移问题并直接对抗灾难性遗忘；

2) 设计基于实体历史表现和数量动态调整在交叉熵损失中的类别权重。

在 CoNLL2003、I2B2、OntoNotes5 这三个基准数据集的十种 CNER 设置中，DFPA 平均提升 Mi-F1 和 Ma-F1，最高单任务增益达 18.83%，显著超越现有最优方法。

2. 相关工作

2.1. 持续学习

持续学习在不降低先前任务性能的情况下学习连续任务[14]-[16]。本文将现有的持续学习方法分为基于记忆的、基于动态架构的和基于规则化的。基于记忆的方法[17]-[19]通过将保存或生成的旧样本整合到当前训练样本中来学习新任务。基于动态架构的方法[20]-[22]通过动态扩展模型架构以学习新任务。基于正则化的方法则通过对网络权重施加约束[1][23]-[25]、调控中间特征[26]或约束输出概率[27][28]来缓解灾难性遗忘。

2.2. 连续命名实体识别

传统 NER 专注于开发各种深度学习模型，旨在从非结构化文本中提取实体[29]。最近，PLM 已被广泛用于 NER，并已实现 SOTA 性能[2]。然而，大多数现有方法被设计为识别预定义实体类型的固定集合。作为响应，CNER 将持续学习范式与传统 NER 相结合[6][11][30]。

2019 年，Monaiikul 等人开创性地提出了 AddNER 和 Extend NER 框架，首次将知识蒸馏技术引入 CNER 领域。这两种架构通过教师 - 学生模型的知识传递机制，在扩展模型识别新实体类型能力的同时，有效保留了已有实体类型的识别能力，为后续研究奠定了基础[3]。

2021 年，夏雨等人提出的“学习 - 复习”框架代表了重要突破。该方法的创新性在于采用两阶段训练策略，首先生成包含旧实体类型的合成样本，然后进行二次蒸馏，显著缓解了新旧实体类型间的混淆问题，在 CoNLL-03 和 OntoNotes 数据集上取得了优于基线模型的表现[9][31][32]。

2023 年，张都臻等人开发的 CFPD 方法将研究推向新高度。该方法不仅创新性地引入池化特征蒸馏来平衡模型稳定性与可塑性，还设计了基于置信度的伪标签策略，有效缓解了非实体类型的语义偏移问题。但伪标签策略虽然减少了错误传播但未能完全消除旧模型的预测偏差[6]。

2023 年，张都臻等人提出了名为 RDP 的方法。该方法针对增量式命名实体识别(INER)中的两大挑战——灾难性遗忘和背景偏移问题，提出了任务关系蒸馏方案和原型伪标签策略。任务关系蒸馏通过跨任务关系蒸馏损失和内任务自熵损失，平衡了模型的稳定性和可塑性，有效缓解了灾难性遗忘；原型伪标签则通过纠正旧模型的预测错误，生成高质量伪标签，解决了背景偏移问题[13]。

2024 年，于雅涵等人提出的权重调优与融合策略展现了新思路。通过 WT 策略的学习率衰减计划和 WF 推理阶段的动态参数融合，该方法以模型无关的方式增强了知识保留能力，在十个 CNER 设置中均表现出稳定的性能提升。不过文章也承认该方法对实体类型的学习顺序较为敏感，且未在 BERT 之外的

大型语言模型上验证其普适性，这在一定程度上限制了方法的适用范围[7]。

2025 年，刘浩等人的概念驱动方法开辟了新方向。通过将形式概念分析引入 CNER，构建概念格来捕捉深层知识结构，并设计多层次蒸馏方案，该方法在保持旧知识方面取得了显著效果。然而该方法的一个明显局限是仅关注已学习过的旧类型而忽略了未来可能出现的实体类型，且完全监督的学习设置需要大量标注数据，在实际应用中可能面临标注成本高的挑战[8]。

3. 方法

针对连续命名实体识别(CNER)中的语义漂移与数据不平衡而导致的灾难性遗忘问题，本文提出动态词典与自适应权重调整相结合的双重优化框架。系统通过密度感知采样构建增量原型库，并利用自适应权重有效平衡新旧知识的学习强度，共同缓解灾难性遗忘现象。具体模型框架如图 2 所示。

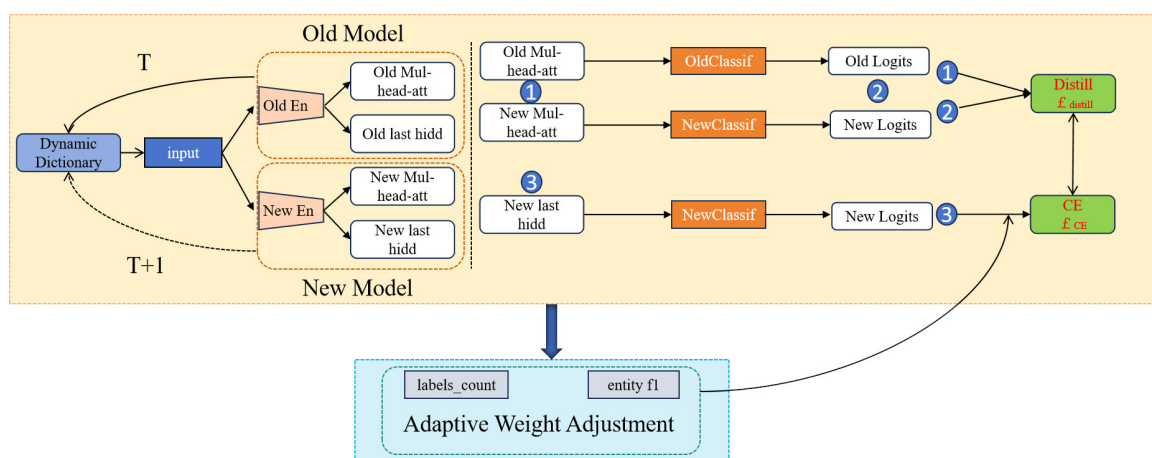


Figure 2. NER model architecture with dynamic dictionary and adaptive weight adjustment

图 2. 动态词典与自适应权重调整的 NER 模型架构

3.1. 动态词典与密度感知采样机制

在连续命名实体识别(Continual Named Entity Recognition, CNER)任务中，灾难性遗忘和语义漂移是影响模型性能的两大核心挑战。传统的经验回放(Experience Replay)方法通常采用随机采样的方式维护历史样本，但这种方式难以保证特征空间的代表性，尤其是在低密度边界区域，模型更容易遗忘历史知识。为此，本文提出了一种基于密度感知采样的动态词典(Dynamic Dictionary)构建方法，结合特征空间分析与自适应权重机制，显著提升了模型对旧知识的保持能力和对新知识的适应能力。

具体而言，在每一增量任务 t 结束后($t < T_{total}-1$)，首先利用当前训练好的模型对本轮数据集 D_t 进行特征提取。由于 CNER 的增量设定要求所有历史实体在当前任务中均被标记为 O-label，导致数据中 O 标签样本占比极高，这不仅会造成特征空间的主导效应，还会加剧梯度消失和决策边界模糊化等问题。为缓解这一现象，本文在特征采样阶段对 O 标签样本实施严格的随机下采样，仅保留 5% 的 O 标签特征用于后续分析。设 S_0 为采样后的 O 标签特征集合，具体采样过程可表示为：

$$S_0 \leftarrow \text{RandomSubset}(\{y = 0\}, p = 0.05)$$

其中， $y = 0$ 表示 O 标签， $p = 0.05$ 为采样比例。

在获得初步特征集合后，对所有特征进行 z-score 归一化处理，以消除不同维度间的尺度影响，确保后续聚类与降维分析的有效性。归一化后的特征 $\mathbf{h}'$ 计算如下：

$$\mathbf{h}' = \frac{\mathbf{h} - \mu_{\mathbf{h}}}{\sigma_{\mathbf{h}}}$$

其中, $\mu_{\mathbf{h}}$ 和 $\sigma_{\mathbf{h}}$ 分别为特征的均值和标准差。随后, 采用 t-SNE 算法将高维特征 $\mathbf{h}'$ 降维至二维空间, 以最大程度保留标签间的局部拓扑结构。t-SNE 降维的结果 $\mathbf{h}_{2D}$ 可表示为:

$$\mathbf{h}_{2D} = t-SNE(\mathbf{h}'; \text{perplexity} = 30)$$

通过可视化分析, 发现特征空间中存在明显的覆盖盲区, 尤其是在新旧类别的边界区域。为此, 本文设计了密度感知采样机制, 重点关注低密度区域的样本, 以提升动态词典的代表性和多样性。

具体地, 对于每个目标实体 $k \in Y_{old}$, 提取其 B-k/I-k 标签的样本集合 P_k , 并对其二维特征执行 K-means 聚类 ($K=10$), 以避免样本分布过度重叠。每个聚类簇 C_j 内, 进一步计算所有样本的平均 k 近邻距离 ($k=20$), 以衡量该簇的密度水平。密度计算公式如下:

$$\rho_j = \frac{1}{|C_j|} \sum_{i \in C_j} \frac{1}{20} \sum_{m=1}^{20} \|\mathbf{h}_i - \mathbf{h}_m^{(i)}\|$$

其中, $\mathbf{h}_i$ 为第 i 个样本的特征, $\mathbf{h}_m^{(i)}$ 为其第 m 个最近邻。基于逆密度采样策略, 从每个簇中优先选取密度最低的 10% 样本, 形成边界关键样本集 $\mathcal{S}_{\text{density}}$, 以最大化动态词典对特征空间边界的覆盖能力。采样公式如下:

$$\mathcal{S}_{\text{density}} = \left\{ \mathbf{h}_i \mid i \in \arg \underset{\text{top}}{\text{sort}}(\rho) \right\}$$

采样完成后, 将选中的 token 反向映射至原始句子, 并与历史动态词典合并, 得到当前任务的动态词典 $\mathcal{M}_t$ 。动态词典的更新规则如下:

$$\mathcal{M}_t = \begin{cases} \mathcal{S}_0 \cup \mathcal{S}_{\text{density}}, & t = 0 \\ \mathcal{M}_{t-1} \cup \mathcal{S}_{\text{density}}, & t \geq 1 \end{cases}$$

需要注意的是, 在最后一个任务阶段, 动态词典不需要再更新。

3.2. 自适应权重调整策略

在连续命名实体识别场景下, 类别不平衡问题会显著影响模型的增量学习效果。传统静态权重分配方法难以适应持续变化的实体分布。为了解决这一挑战, 本文提出了一种自适应权重优化方法, 该方法基于多维动态评估, 能够动态地调节训练过程中的损失函数, 以适应实体分布的持续变化。

本文的方法首先涉及到新实体的初始化。在检测到新增实体类型时, 采用一种启发式权重初始化策略, 该策略基于样本量, 通过以下公式计算:

$$we_{init} = \sqrt{\frac{N_{total}}{N_e + f}}$$

这里, N_{total} 是整个数据集中的样本总量, N_e 是新增实体类型的样本量, 而 f 是一个小的数值稳定项, 用于避免除以零的情况。

随后, 对权重进行动态归一化处理, 以确保权重分布的合理性并避免极端值的干扰。这一步骤通过以下公式实现:

$$w_{\text{norm}} = 0.8 \times \frac{(w - w_{\min})}{(w_{\max} - w_{\min})} + 0.6$$

在这个公式中, $w_{\min}$ 和 $w_{\max}$ 分别代表所有权重中的最小值和最大值, 通过这种方式, 可以将权重映射到一个预定的区间内, 从而实现归一化。

接着, 计算每个实体的权重, 这涉及到一个复合权重公式, 该公式综合考虑了实体的基础 F1 分数、全局计数和样本比率, 具体如下:

$$\text{weight} = (\text{base_f1}^{1.5}) \times \sqrt{\frac{\text{global_counts}}{\text{total_count} + 1e-5}} \times \left(1 + \log\left(\frac{\text{sample_ratio}}{1}\right)\right)$$

在这个公式中, base_f1 代表实体的基础 F1 分数, total_count 是实体的全局计数, global_counts 是所有实体的全局计数, 而 sample_ratio 是样本比率。通过结合实体识别的性能和样本分布特征来动态调整每个实体的权重。在训练过程中, 这些计算得到的权重被用于加权交叉熵损失函数, 以优化模型:

$$L = -\frac{1}{N} \sum (w_e \times y_i \times \log(p_i))$$

这里, p_i 表示模型对实体 i 的预测概率, w_e 是相应实体 e 的权重。通过这种方式, 本文的方法能够在反向传播过程中自动强化关键实体的梯度信号, 同时抑制不必要的更新。

本文提出的动态词典与密度感知采样机制, 能够在每一轮增量学习后动态维护特征空间的代表性样本, 显著提升了模型对历史知识的保持能力。同时, 动态权重机制有效缓解了类别不平衡带来的训练偏置, 使得模型在面对新旧知识冲突时能够实现更优的平衡。

4. 实验

4.1. 实验设置

本文严格遵循 CFNER 中提出的实验设置[13]。选用三个广泛应用的命名实体识别数据集进行评估, 分别为 CoNLL2003、I2B2 和 OntoNotes5。训练集被划分为不相交的切片, 每个切片对应一个连续学习步骤, 切片的划分和采样算法均与 CFNER 保持一致, 每个切片仅保留当前学习的实体类型标签, 其余标签被屏蔽为非实体类型。具体采样细节可参考 CFNER 附录 B。在 CNER 设置方面, 实体类型按照字母顺序依次引入, 数据切片顺序地用于模型训练。以 CoNLL2003 为例, 采用 FG-1-PG-1 和 FG-2-PG-1 两种设置; I2B2 和 OntoNotes5 数据集则采用 FG-1-PG-1、FG-2-PG-2、FG-8-PG-1 和 FG-8-PG-2 四种设置。评估时, 验证集仅保留当前学习的实体类型标签, 测试集则保留所有已学习过的实体类型标签, 其余均视为非实体类型。性能评估指标采用微观 F1 (Mi-F1) 和宏观 F1 (Ma-F1), 并以所有步骤的平均结果(包括第一步)作为最终性能。此外, 为了更全面地分析模型表现, 本文还引入了逐步性能对比曲线, 并通过配对 t 检验(显著性水平为 0.05)评估改进的统计学意义。本文在上述基础上引入了动态词典机制: 在每次训练结束后, 自动挑选并保存每个实体类型不超过 100 条的数据样本, 作为后续训练阶段的辅助知识。

4.2. 整体结果

CoNLL2003 数据集的实验对比结果如表 1 所示, Ours 为本文方法的结果。

Table 1. Comparison results on the CoNLL2003 dataset

表 1. Conll2003 数据集对比实验结果

数据集	方法	FG-1-PG-1		FG-2-PG-1	
		Mi-F1	Ma-F1	Mi-F1	Ma-F1
CoNLL2003	PODNet	36.74 ± 0.52	29.43 ± 0.28	59.12 ± 0.54	58.39 ± 0.99

续表

LUCIR	74.15 ± 0.43	70.48 ± 0.66	80.53 ± 0.31	77.33 ± 0.31
ST	76.17 ± 0.91	72.88 ± 1.12	76.65 ± 0.24	66.72 ± 0.11
Extend NER	76.36 ± 0.98	73.04 ± 1.80	76.66 ± 0.66	66.36 ± 0.64
CFNER	80.91 ± 0.29	79.11 ± 0.50	80.83 ± 0.36	75.20 ± 0.32
CFNER*	80.80 ± 0.80	78.93 ± 1.13	79.92 ± 0.36	72.83 ± 1.59
RDP	82.55 ± 0.26	80.64 ± 0.12	85.82 ± 0.36	83.59 ± 0.37
RDP*	81.76 ± 0.32	79.77 ± 0.37	85.54 ± 0.28	83.26 ± 0.29
CPFD	82.24 ± 0.63	79.94 ± 0.66	85.70 ± 0.19	83.49 ± 0.16
CPFD*	82.32 ± 0.21	80.03 ± 0.34	85.77 ± 0.37	83.26 ± 0.35
**(Ours)	84.34 ± 0.27	82.42 ± 0.33	86.38 ± 0.26	84.56 ± 0.23

注：*表示为作者复现的分数，下表同。

I2B2 数据集的实验对比结果如表 2 所示，Ours 为本文方法的结果。

Table 2. Comparison results on the I2B2 dataset
表 2. I2B2 数据集对比实验结果

数据集	方法	FG-1-PG-1		FG-2-PG-2		FG-8-PG-1		FG-8-PG-2	
		Mi-F1	Ma-F1	Mi-F1	Ma-F1	Mi-F1	Ma-F1	Mi-F1	Ma-F1
I2B2	PODNet	12.31 ± 0.35	17.14 ± 1.03	34.67 ± 2.65	24.62 ± 1.76	39.26 ± 1.38	27.23 ± 0.93	36.22 ± 12.9	26.08 ± 7.42
	LUCIR	43.86 ± 2.43	31.31 ± 1.62	64.32 ± 0.76	43.53 ± 0.59	57.86 ± 0.87	33.04 ± 0.39	68.54 ± 0.27	46.94 ± 0.63
	ST	31.98 ± 2.12	14.76 ± 1.31	55.44 ± 4.78	33.38 ± 3.13	49.51 ± 1.35	23.77 ± 1.01	48.94 ± 6.78	29.00 ± 3.04
	Extend NER	42.85 ± 2.86	24.05 ± 1.35	57.01 ± 4.14	35.29 ± 3.38	43.95 ± 2.01	23.12 ± 1.79	52.25 ± 5.36	30.93 ± 2.77
	CFNER	62.73 ± 3.62	36.26 ± 2.24	71.98 ± 0.50	49.09 ± 1.38	59.79 ± 1.70	37.30 ± 1.15	69.07 ± 0.89	51.09 ± 1.05
	CFNER*	64.48 ± 1.13	37.74 ± 1.16	73.27 ± 0.42	52.92 ± 1.04	59.65 ± 1.53	38.50 ± 1.18	67.81 ± 0.94	50.74 ± 1.02
	RDP	71.39 ± 1.01	44.00 ± 2.31	77.45 ± 0.55	53.48 ± 0.66	77.50 ± 1.26	62.99 ± 0.36	80.08 ± 0.40	63.72 ± 0.71
	RDP*	50.57 ± 8.02	38.67 ± 1.81	76.32 ± 1.63	54.97 ± 0.52	77.45 ± 2.53	62.55 ± 1.49	81.80 ± 0.62	65.85 ± 0.90
	CPFD	74.19 ± 0.95	48.34 ± 1.45	78.19 ± 0.58	56.04 ± 1.22	74.75 ± 1.35	56.19 ± 2.46	81.05 ± 0.87	65.04 ± 1.13
	CPFD*	73.97 ± 0.78	47.09 ± 1.08	79.24 ± 0.28	58.36 ± 1.05	74.58 ± 2.01	55.39 ± 1.73	81.19 ± 0.83	64.45 ± 1.50
	**(Ours)	78.66 ± 0.76	65.92 ± 0.77	82.37 ± 0.19	68.85 ± 0.20	85.95 ± 0.29	71.67 ± 0.65	86.22 ± 0.19	73.11 ± 0.16

OntoNotes5 数据集的实验对比结果如表 3 所示，Ours 为本文方法的结果。

Table 3. Comparison results on the Ontonotes5 dataset
表 3. I2B2 数据集对比实验结果

数据集	方法	FG-1-PG-1		FG-2-PG-2		FG-8-PG-1		FG-8-PG-2	
		Mi-F1	Ma-F1	Mi-F1	Ma-F1	Mi-F1	Ma-F1	Mi-F1	Ma-F1
Ontonotes5	PODNet	9.06 ± 0.56	8.36 ± 0.57	19.04 ± 1.08	16.93 ± 0.85	29.00 ± 0.86	20.54 ± 0.91	37.38 ± 0.26	25.85 ± 0.29
	LUCIR	28.18 ± 1.15	21.11 ± 0.84	56.40 ± 1.79	40.58 ± 1.11	66.46 ± 0.46	46.29 ± 0.38	76.17 ± 0.09	55.58 ± 0.55

续表

ST	50.71 ± 0.79	33.24 ± 1.06	68.93 ± 1.67	50.63 ± 1.66	73.59 ± 0.66	49.41 ± 0.77	77.07 ± 0.62	53.32 ± 0.63
Extend NER	50.53 ± 0.86	32.84 ± 0.84	67.61 ± 1.53	49.26 ± 1.49	73.12 ± 0.93	49.55 ± 0.90	76.85 ± 0.77	54.37 ± 0.57
CFNER	58.94 ± 0.57	42.22 ± 1.10	72.59 ± 0.48	55.96 ± 0.69	78.92 ± 0.58	57.51 ± 1.32	80.68 ± 0.25	60.52 ± 0.84
CFNER*	58.22 ± 1.12	41.54 ± 1.08	72.39 ± 0.58	55.05 ± 0.57	80.18 ± 0.26	60.06 ± 1.29	81.60 ± 0.4	61.72 ± 0.53
RDP	68.28 ± 1.09	53.56 ± 0.39	74.38 ± 0.26	57.73 ± 0.54	79.89 ± 0.20	63.20 ± 0.58	83.30 ± 0.30	66.92 ± 1.26
RDP*	56.13 ± 2.56	45.38 ± 2.00	68.84 ± 0.77	54.48 ± 0.82	77.58 ± 1.34	60.89 ± 1.40	83.00 ± 0.20	66.01 ± 0.76
CPFD	66.73 ± 0.70	54.12 ± 0.30	74.33 ± 0.30	57.75 ± 0.35	81.87 ± 0.47	65.52 ± 1.05	83.38 ± 0.18	66.27 ± 0.75
CPFD*	66.57 ± 0.89	53.74 ± 0.51	74.32 ± 0.36	57.77 ± 0.74	82.72 ± 0.29	66.17 ± 0.91	84.00 ± 0.20	67.08 ± 0.50
** (Ours)	66.12 ± 0.65	59.39 ± 1.22	71.84 ± 0.41	64.46 ± 0.66	79.36 ± 0.30	70.53 ± 0.36	81.05 ± 0.29	72.22 ± 0.22

4.3. 案例分析

图 3 通过具体案例对比分析展示了本文提出模型与基线模型在命名实体识别任务上的预测性能差异。该图选取具有代表性的文本序列，并列呈现真实标注标签、本文模型预测结果及参考模型预测结果，清晰展示了各模型在实体边界识别准确度、类型标注一致性以及错误模式方面的对比情况。从可视化结果可以看出，本文模型在复杂实体识别和歧义消解方面表现出显著优势，其预测结果与真实标签的一致性更高，而基线模型则存在更多的实体漏检和类型误判问题。该案例分析为定量实验结果提供了直观佐证，进一步验证了本文所提方法的有效性和优越性。

Input Sentence	Results in last of the group matches at the World Grand Prix badminton finals on Friday :														
CPFD PL	[O]	[O]	[O]	[O]	[O]	[O]	[O]	[O]	[O]	[B-ORG]	[I-ORG]	[I-ORG]	[O]	[O]	[O]
Ours	[O]	[O]	[O]	[O]	[O]	[O]	[O]	[O]	[O]	[B-MISC]	[I-MISC]	[I-MISC]	[O]	[O]	[O]
Golden Labels	[O]	[O]	[O]	[O]	[O]	[O]	[O]	[O]	[O]	[B-MISC]	[I-MISC]	[I-MISC]	[O]	[O]	[O]
Input Sentence	He is a programmer and works as a Industrial Engineering and Manufacturing Technologist at Activision Blizzard .														
CPFD PL	[O]	[O]	[O]	[O]	[O]	[O]	[O]	[B-PROF]	[O]	[O]	[O]	[O]	[O]	[B-HOSP]	[I-HOSP]
Ours	[O]	[O]	[O]	[B-PROF]	[O]	[O]	[O]	[B-PROF]	[I-PROF]	[I-PROF]	[I-PROF]	[I-PROF]	[I-PROF]	[O]	[B-ORG]
Golden Labels	[O]	[O]	[O]	[B-PROF]	[O]	[O]	[O]	[B-PROF]	[I-PROF]	[I-PROF]	[I-PROF]	[I-PROF]	[I-PROF]	[O]	[B-ORG]

Figure 3. Comparative analysis of cases based on true labels and multi-model predictions
图 3. 基于真实标签与多模型预测的案例对比分析

4.4. 可视化分析

4.4.1. 特征可视化

图 4 通过特征空间可视化对比了在 I2B2 数据集前四轮训练中，基线方法与本文提出方法的特征表示学习效果。从可视化结果可以明显观察到，加入本文提出的方法后，模型学习到的特征表示呈现出更加清晰的类别边界和更紧凑的类内分布。具体而言，基线方法的特征分布相对分散，不同实体类型的特征存在较大重叠，这反映了模型在区分相似实体方面的困难；而引入本文的方法后，各类别特征形成了更加分离的聚类簇，表明模型能够学习到判别性更强的特征表示。这种改进在医学文本的复杂实体识别任务中尤为重要，因为医学术语往往具有高度的语义相似性。前四轮的渐进式改善进一步证明了本文方法在训练过程中的稳定性和有效性，为模型性能提升提供了直观的特征层面解释。

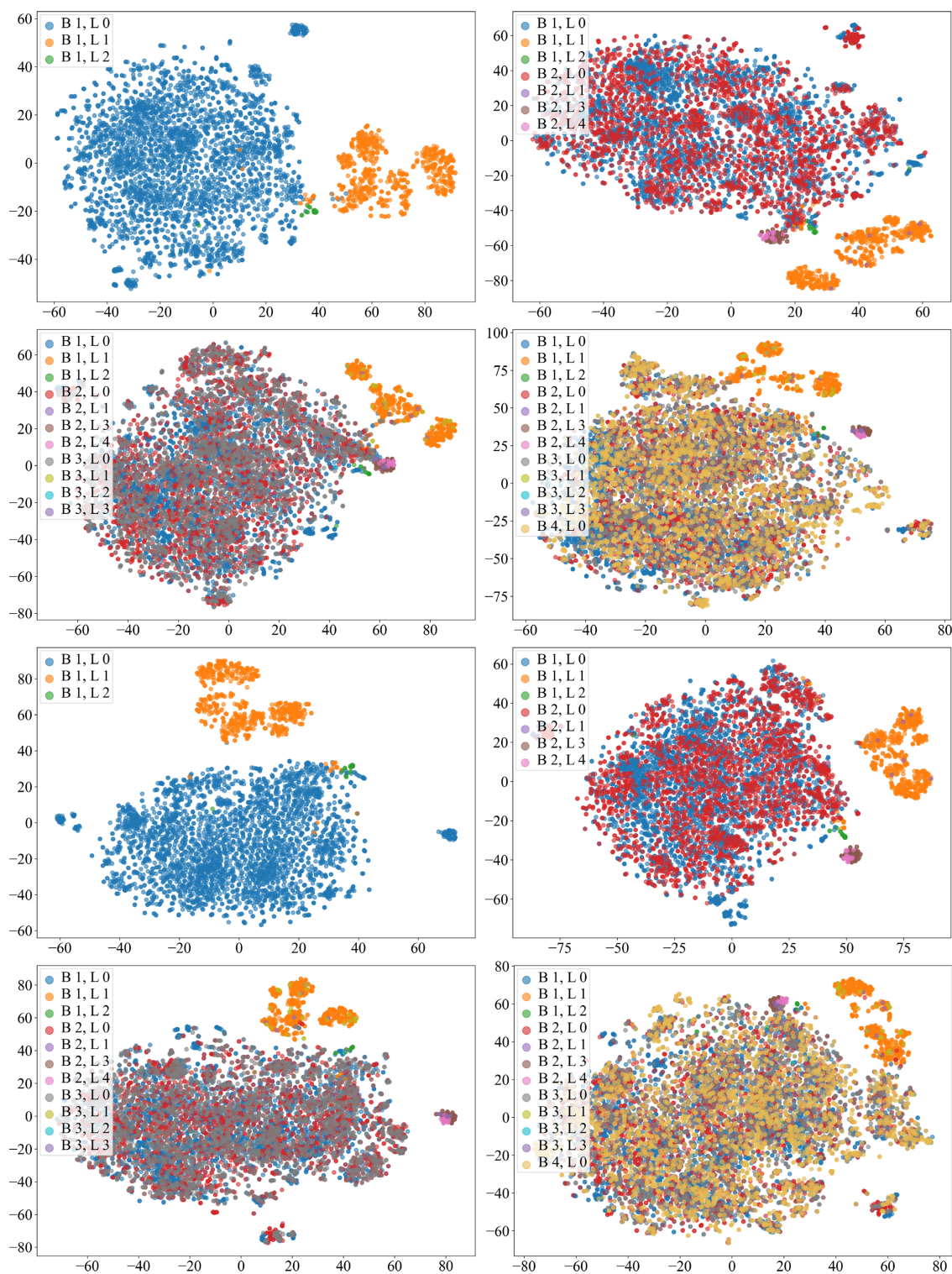


Figure 4. Feature visualization comparison on I2B2 dataset (first four rounds)

图 4. I2B2 数据集前四轮特征可视化对比分析

4.4.2. 结果分数可视化

图 5 展示了 I2B2 和 OntoNotes5 数据集在不同任务设置下各对比方法的 Ma-F1 分数变化趋势。从图

中可以清晰观察到, 本文提出的方法(Ours)在所有实验设置上均优于其他基线方法, 展现出卓越的性能稳定性和泛化能力。这种跨数据集的一致优势充分证明了本文方法在不同领域和任务设置下的强鲁棒性。性能提升主要归因于动态词典机制的有效性、自适应权重调整的优化效果等。对于 8-1、8-2 这样实验设置下的提升效果没有 1-1、2-2 提升明显, 是因为可进步空间缩小。

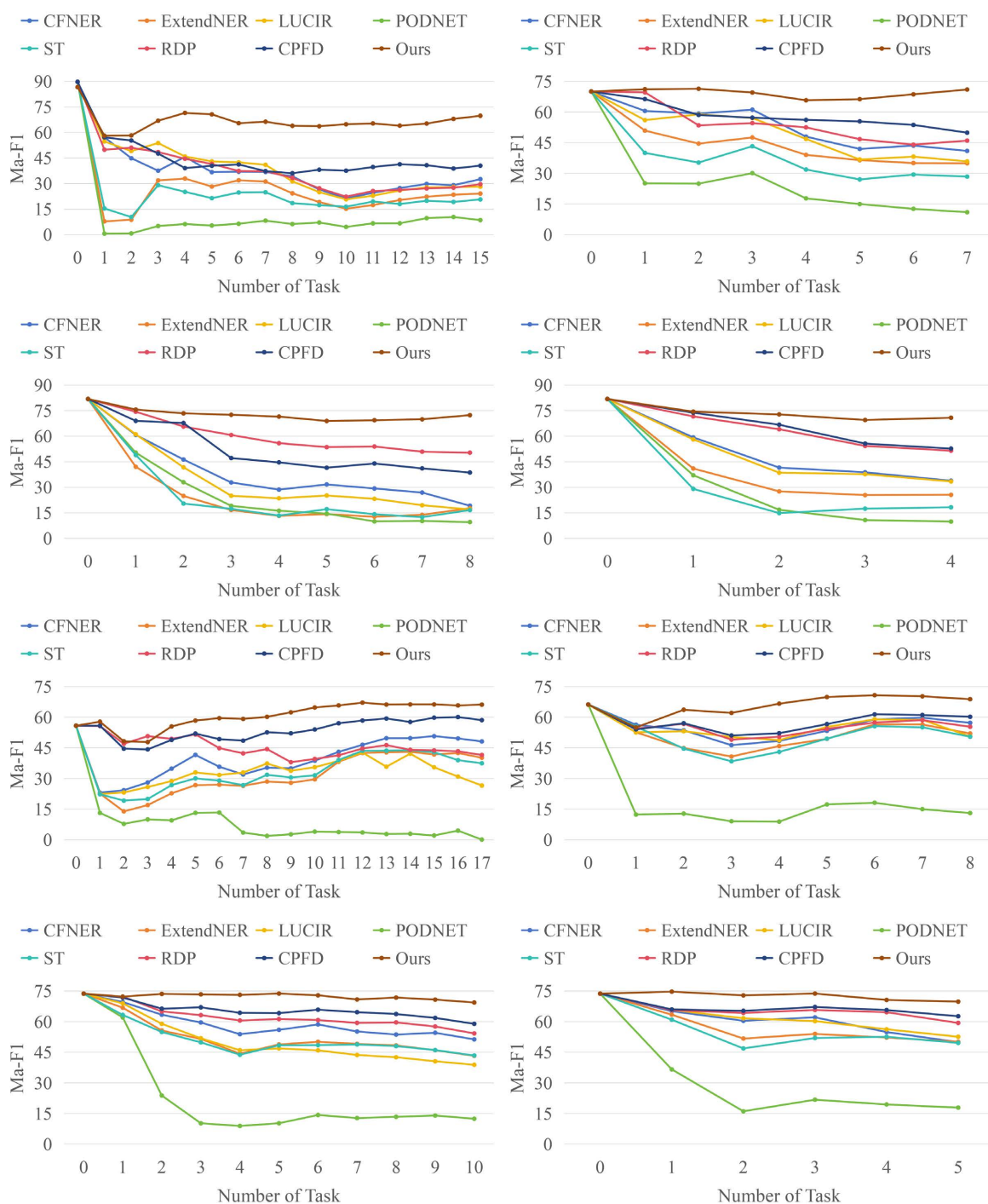


Figure 5. Ma-F1 score comparison across different task settings on I2B2 and Ontonotes5 datasets

图 5. I2B2 与 OntoNotes5 数据集多任务设置下的 Ma-F1 分数对比

4.5. 消融实验

CoNLL2003 数据集的消融实验结果如表 4 所示。

Table 4. Ablation study results on the CoNLL2003 dataset

表 4. CoNLL2003 数据集消融实验结果

实验设置	FG-1-PG-1		FG-2-PG-1	
	Mi-F1	Ma-F1	Mi-F1	Ma-F1
Baseline (CPFD)	82.32 ± 0.21	80.03 ± 0.34	85.77 ± 0.37	83.26 ± 0.35
+动态词典	83.46 ± 0.34	81.33 ± 0.38	86.68 ± 0.14	84.80 ± 0.12
+自适应权重	84.34 ± 0.27	82.42 ± 0.33	86.38 ± 0.26	84.56 ± 0.23

I2B2 数据集的消融实验结果如表 5 所示。

Table 5. Ablation study results on the I2B2 dataset

表 5. I2B2 数据集消融实验结果

实验设置	FG-1-PG-1		FG-2-PG-2		FG-8-PG-1		FG-8-PG-2	
	Mi-F1	Ma-F1	Mi-F1	Ma-F1	Mi-F1	Ma-F1	Mi-F1	Ma-F1
Baseline (CPFD)	73.97 ± 0.78	47.09 ± 1.08	79.24 ± 0.28	58.36 ± 1.05	74.58 ± 2.01	55.39 ± 1.73	81.19 ± 0.83	64.45 ± 1.50
+动态词典	78.41 ± 0.16	66.15 ± 0.31	80.73 ± 0.32	62.93 ± 0.75	84.62 ± 0.18	68.62 ± 0.56	85.19 ± 0.26	70.43 ± 0.98
+自适应权重	78.66 ± 0.76	65.92 ± 0.77	82.37 ± 0.19	68.85 ± 0.20	85.95 ± 0.29	71.67 ± 0.65	86.22 ± 0.19	73.11 ± 0.16

OntoNotes5 数据集的消融实验结果如表 6 所示。

Table 6. Ablation study results on the I2B2 dataset

表 6. I2B2 数据集消融实验结果

实验设置	FG-1-PG-1		FG-2-PG-2		FG-8-PG-1		FG-8-PG-2	
	Mi-F1	Ma-F1	Mi-F1	Ma-F1	Mi-F1	Ma-F1	Mi-F1	Ma-F1
Baseline (CPFD)	66.57 ± 0.89	53.74 ± 0.51	74.32 ± 0.36	57.77 ± 0.74	82.72 ± 0.29	66.17 ± 0.91	84.00 ± 0.20	67.08 ± 0.50
+动态词典	65.98 ± 0.75	59.99 ± 0.53	72.59 ± 0.29	62.63 ± 0.59	80.43 ± 0.70	69.66 ± 0.89	82.49 ± 0.20	69.28 ± 0.83
+自适应权重	66.12 ± 0.65	59.39 ± 1.22	71.84 ± 0.41	64.46 ± 0.66	79.36 ± 0.30	70.53 ± 0.36	81.05 ± 0.29	72.22 ± 0.22

4.6. 结果分析

为深入探究本方法对持续学习性能的提升本质，我们进行了细致的类别级性能分析。如表 1 与表 2

所示，与基线方法相比，我们提出的动态词典驱动框架所带来的宏观 F1 (Macro-F1)提升并非均匀分布，其增益主要来源于有效遏制了对已学旧类别的灾难性遗忘，以及对难样本新类别的充分学习。

在 CoNLL-2003 数据集的 1-1 任务增量学习场景中，我们的方法将平均 Macro-F1 从 80.30 提升至 83.08。具体而言，在最终任务中，所有旧类别(location, misc, organisation)的 F1 分数均获得稳定改善或保持高位，尤其是 misc 类别，其 F1 分数从 60.66 显著提升至 73.63，这清晰地表明我们的方法通过密度感知记忆与自适应加权机制，有效巩固了模型对历史知识的保持能力。

在更具挑战性的 I2B2 数据集 1-1 任务场景中，我们方法的优势更为显著，平均 Macro-F1 从 47.93 大幅提升至 66.81，相对提升近 40%。基线模型表现出严重的灾难性遗忘，众多早期类别(如 CITY、COUNTRY、IDNUM)的 F1 分数在后续任务中暴跌至 0，而我们的方法成功地维持了这些类别的有效识别能力。例如，CITY 和 COUNTRY 类别的 F1 分数被分别稳定在 50~60 和 60~65 的区间，而非归零。同时，对于后续任务中的难样本类别(如 MEDICALRECORD, ZIP)，我们的方法也实现了更优的学习效果，F1 分数分别达到 81.06 和 83.92。

综上所述，性能提升直接验证了动态词典作为“记忆锚点”在缓解遗忘上的有效性，以及自适应权重策略在促进难样本学习上的必要性。我们的方法确保了模型在整个增量学习过程中保持稳定且公平的类别级性能，从而实现了宏观 F1 指标的整体飞跃。

4.7. 局限性

本研究在连续命名实体识别任务中取得了重要进展，但仍需客观审视其存在的若干局限性。从方法论角度而言，当前动态词典的构建机制完全依赖于精确的监督信号，这一特性使其在弱监督或远程监督场景下的鲁棒性受到显著制约。计算效率方面，尽管相较于基线方法已实现优化，但每轮增量学习过程中进行的特征空间分析(包括 t-SNE 降维和 K-means 聚类等操作)仍不可避免地引入额外的计算负担，这在处理超大规模文本语料时可能构成系统瓶颈。在理论框架层面，现有的权重调整公式基于实体独立性假设构建，未能充分考虑实体间的潜在语义关联，这一局限在需要复杂语义推理的嵌套实体识别任务中表现得尤为明显。此外，当前实现方案对新引入实体类型的数量存在内在约束，当单步增量过程中新增实体类型超过特定阈值时，动态词典的内存占用将呈现显著增长趋势，这对资源受限的实际应用场景提出了挑战。这些局限性不仅界定了本研究的适用范围，同时也为后续研究指明了潜在改进方向，包括但不限于：开发更具鲁棒性的弱监督学习范式、设计更高效的特征空间更新算法，以及探索跨语言迁移的有效策略等。

5. 结论

本研究为连续命名实体识别(CNER)这一自然语言理解领域的新兴研究方向奠定了重要基础。针对 CNER 任务中存在的两大核心挑战——灾难性遗忘问题与非实体类型的语义漂移现象，本文提出了一套创新性的解决方案。首先，设计了一种基于密度感知采样的动态词典构建机制，通过智能选择历史数据中的关键样本来直接缓解模型遗忘问题。其次，开发了一种自适应权重调整策略，该策略综合考虑实体类别的历史学习表现和样本分布特征，动态优化交叉熵损失函数中的权重分配，有效平衡新旧类别之间的学习强度。在 CoNLL2003、I2B2 和 OntoNotes5 这三个基准数据集的十个不同 CNER 任务设置下进行的系统性实验表明，本文提出的方法在所有测试场景中都显著超越了现有最优方法，展现了卓越的泛化能力和鲁棒性。

致 谢

在本论文的研究与写作过程中，我得到了许多老师、同学和朋友的关心与帮助。在此，谨向所有给

予我指导和支持的人表示衷心的感谢。

首先，衷心感谢我的导师，在论文选题、研究方法、论文撰写等各个阶段都给予了我悉心的指导和宝贵的建议，使我能够顺利完成本论文。感谢实验室的各位同学和朋友，在数据处理、实验设计和学术讨论中给予了我极大的帮助和支持。你们的鼓励和陪伴让我在科研道路上不断前行。

最后，感谢我的家人对我学习和生活的理解与支持，是你们给予我坚强的后盾和无限的动力。

由于篇幅有限，未能一一列举所有帮助过我的人，在此一并表示诚挚的谢意！

参考文献

- [1] Ma, X.Z. and Hovy, E. (2016) End-to-End Sequence Labeling via Bi-Directional LSTM-CNNsCRF. *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, Berlin, 7-12 August 2016, 1064-1074.
- [2] Devlin, J., Chang, M.W., Lee, K. and Toutanova, K. (2019) BERT: Pre-Training of Deep Bidirectional Transformers for Language Understanding. *2019 Proceedings of NAACL-HLT*, Minneapolis, 2-7 June 2019, 4171-4186.
- [3] Monaikul, N., Castellucci, G., Filice, S. and Rokhlenko, O. (2021) Continual Learning for Named Entity Recognition. *Proceedings of the AAAI Conference on Artificial Intelligence*, **35**, 13570-13577. <https://doi.org/10.1609/aaai.v35i15.17600>
- [4] Zheng, J., Liang, Z., Chen, H. and Ma, Q. (2022) Distilling Causal Effect from Miscellaneous Other-Class for Continual Named Entity Recognition. *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, Abu Dhabi, 7-11 December 2022, 3602-3615. <https://doi.org/10.18653/v1/2022.emnlp-main.236>
- [5] Hinton, G., Vinyals, O., Dean, J., et al. (2015) Distilling the Knowledge in a Neural Network. arXiv: 1503.02531.
- [6] Zhang, D., Cong, W., Dong, J., Yu, Y., Chen, X., Zhang, Y., et al. (2023) Continual Named Entity Recognition without Catastrophic Forgetting. *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, Singapore, 6-10 December 2023, 8186-8197. <https://doi.org/10.18653/v1/2023.emnlp-main.509>
- [7] Yu, Y., Zhang, D., Chen, X. and Chu, C. (2024) Flexible Weight Tuning and Weight Fusion Strategies for Continual Named Entity Recognition. *Findings of the Association for Computational Linguistics ACL 2024*, Bangkok, 11-16 August 2024, 1351-1358. <https://doi.org/10.18653/v1/2024.findings-acl.79>
- [8] Liu, H., Xin, X., Peng, W., Song, J. and Sun, J. (2025) Concept-Driven Knowledge Distillation and Pseudo Label Generation for Continual Named Entity Recognition. *Expert Systems with Applications*, **270**, Article ID: 126546. <https://doi.org/10.1016/j.eswa.2025.126546>
- [9] Xia, Y., Wang, Q., Lyu, Y., Zhu, Y., Wu, W., Li, S., et al. (2022) Learn and Review: Enhancing Continual Named Entity Recognition via Reviewing Synthetic Samples. *Findings of the Association for Computational Linguistics: ACL 2022*, Dublin, 22-27 May 2022, 2291-2230. <https://doi.org/10.18653/v1/2022.findings-acl.179>
- [10] French, R. (1999) Catastrophic Forgetting in Connectionist Networks. *Trends in Cognitive Sciences*, **3**, 128-135. [https://doi.org/10.1016/s1364-6613\(99\)01294-2](https://doi.org/10.1016/s1364-6613(99)01294-2)
- [11] McCloskey, M. and Cohen, N.J. (1989) Catastrophic Interference in Connectionist Networks: The Sequential Learning Problem. In: *Psychology of Learning and Motivation*, Elsevier, 109-165. [https://doi.org/10.1016/s0079-7421\(08\)60536-8](https://doi.org/10.1016/s0079-7421(08)60536-8)
- [12] Robins, A. (1995) Catastrophic Forgetting, Rehearsal and Pseudorehearsal. *Connection Science*, **7**, 123-146. <https://doi.org/10.1080/09540099550039318>
- [13] Zhang, D., Li, H., Cong, W., Xu, R., Dong, J. and Chen, X. (2023) Task Relation Distillation and Prototypical Pseudo Label for Incremental Named Entity Recognition. *Proceedings of the 32nd ACM International Conference on Information and Knowledge Management*, Birmingham, 21-25 October 2023, 3319-3329. <https://doi.org/10.1145/3583780.3615075>
- [14] Chen, Z.Y. and Liu, B. (2018) Lifelong Machine Learning. *Synthesis Lectures on Artificial Intelligence and Machine Learning*, **12**, 1-207.
- [15] Dong, J., Wang, L., Fang, Z., Sun, G., Xu, S., Wang, X., et al. (2022) Federated Class-Incremental Learning. *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, New Orleans, 18-24 June 2022, 10154-10163. <https://doi.org/10.1109/cvpr52688.2022.00992>
- [16] Dong, J., Zhang, D., Cong, Y., Cong, W., Ding, H. and Dai, D. (2023) Federated Incremental Semantic Segmentation. *2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, Vancouver, 17-24 June 2023, 3934-3943. <https://doi.org/10.1109/cvpr52729.2023.00383>
- [17] Lopez-Paz, D. and Ranzato, M.A. (2017) Gradient Episodic Memory for Continual Learning. *Proceedings of the 31st*

- International Conference on Neural Information Processing Systems*, Long Beach, 4-9 December 2017, 6470-6479.
- [18] Rebuffi, S.A., Kolesnikov, A., Sperl, G. and Lampert, C.H. (2017) iCaRL: Incremental Classifier and Representation Learning. *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition*, Honolulu, 21-26 July 2017, 2001-2010.
 - [19] Shin, H., Lee, J.K., Kim, J. and Kim, J. (2017) Continual Learning with Deep Generative Replay. *Proceedings of the 31st International Conference on Neural Information Processing Systems*, Long Beach, 4-9 December 2017, 2994-3003.
 - [20] Mallya, A. and Lazebnik, S. (2018) PackNet: Adding Multiple Tasks to a Single Network by Iterative Pruning. 2018 *IEEE/CVF Conference on Computer Vision and Pattern Recognition*, Salt Lake City, 18-23 June 2018, 7765-7773. <https://doi.org/10.1109/cvpr.2018.00810>
 - [21] Rosenfeld, A. and Tsotsos, J.K. (2020) Incremental Learning through Deep Adaptation. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, **42**, 651-663. <https://doi.org/10.1109/tpami.2018.2884462>
 - [22] Yoon, J., Yang, E., Lee, J. and Hwang, S.J. (2018) Lifelong Learning with Dynamically Expandable Networks. arXiv: 1708.01547.
 - [23] Aljundi, R., Babiloni, F., Elhoseiny, M., Rohrbach, M. and Tuytelaars, T. (2018) Memory Aware Synapses: Learning What (not) to Forget. In: Ferrari, V., Hebert, M., Sminchisescu, C. and Weiss, Y., Eds., *Computer Vision—ECCV 2018*, Springer, 144-161. https://doi.org/10.1007/978-3-030-01219-9_9
 - [24] Kirkpatrick, J., Pascanu, R., Rabinowitz, N., Veness, J., Desjardins, G., Rusu, A.A., et al. (2017) Overcoming Catastrophic Forgetting in Neural Networks. *Proceedings of the National Academy of Sciences of the United States of America*, **114**, 3521-3526. <https://doi.org/10.1073/pnas.1611835114>
 - [25] Zenke, F., Poole, B. and Ganguli, S. (2017) Continual Learning through Synaptic Intelligence. *International Conference on Machine Learning*, Sydney, 6-11 August 2017, 3987-3995.
 - [26] Hou, S., Pan, X., Loy, C.C., Wang, Z. and Lin, D. (2019) Learning a Unified Classifier Incrementally via Rebalancing. 2019 *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, Long Beach, 15-20 June 2019, 831-839. <https://doi.org/10.1109/cvpr.2019.00092>
 - [27] Li, Z. and Hoiem, D. (2018) Learning without Forgetting. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, **40**, 2935-2947. <https://doi.org/10.1109/tpami.2017.2773081>
 - [28] Dong, J., Liang, W., Cong, Y. and Sun, G. (2023) Heterogeneous Forgetting Compensation for Class-Incremental Learning. 2023 *IEEE/CVF International Conference on Computer Vision (ICCV)*, Paris, 1-6 October 2023, 11708-11717. <https://doi.org/10.1109/iccv51070.2023.01078>
 - [29] Li, J., Sun, A., Han, J. and Li, C. (2022) A Survey on Deep Learning for Named Entity Recognition. *IEEE Transactions on Knowledge and Data Engineering*, **34**, 50-70. <https://doi.org/10.1109/tkde.2020.2981314>
 - [30] Ma, R., Chen, X., Lin, Z., Zhou, X., Wang, J., Gui, T., et al. (2023) Learning “O” Helps for Learning More: Handling the Unlabeled Entity Problem for Class-Incremental NER. *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, Toronto, 9-14 July 2023, 5959-5979. <https://doi.org/10.18653/v1/2023.acl-long.328>
 - [31] Tjong, E.F., Sang, K. and De Meulder, F. (2003) Introduction to the CoNLL-2003 Shared Task: Language-Independent Named Entity Recognition. *Proceedings of the Seventh Conference on Natural Language Learning at HLT-NAACL 2003*, Edmonton, 31 May 2003, 142-147.
 - [32] Hovy, E., Marcus, M., Palmer, M., Ramshaw, L. and Weischedel, R. (2006) OntoNotes: The 90% Solution. *Proceedings of the Human Language Technology Conference of the NAACL, Companion Volume: Short Papers on XX—NAACL’06*, New York, 4-9 June 2006, 57-60. <https://doi.org/10.3115/1614049.1614064>