

Graph Few-shot Class-incremental Learning

Zhen Tan

Arizona State University
ztan36@asu.edu

Ruocheng Guo

City University of Hong Kong
ruocheng.guo@cityu.edu.hk

Kaize Ding

Arizona State University
kaize.ding@asu.edu

Huan Liu

Arizona State University
huanliu@asu.edu

ABSTRACT

The ability to incrementally learn new classes is vital to all real-world artificial intelligence systems. A large portion of high-impact applications like social media, recommendation systems, E-commerce platforms, etc. can be represented by graph models. In this paper, we investigate the challenging yet practical problem, *Graph Few-shot Class-incremental* (Graph FCL) problem, where the graph model is tasked to classify both newly encountered classes and previously learned classes. Towards that purpose, we put forward a Graph Pseudo Incremental Learning paradigm by sampling tasks recurrently from the base classes, so as to produce an arbitrary number of training episodes for our model to practice the incremental learning skill. Furthermore, we design a Hierarchical-Attention-based Graph Meta-learning framework, HAG-Meta. We present a task-sensitive regularizer calculated from task-level attention and node class prototypes to mitigate overfitting onto either novel or base classes. To employ the topological knowledge, we add a node-level attention module to adjust the prototype representation. Our model not only achieves greater stability of old knowledge consolidation, but also acquires advantageous adaptability to new knowledge with very limited data samples. Extensive experiments on three real-world datasets, including Amazon-clothing, Reddit, and DBLP, show that our framework demonstrates remarkable advantages in comparison with the baseline and other related state-of-the-art methods.

CCS CONCEPTS

• Computing methodologies → Lifelong machine learning.

KEYWORDS

Graph Neural Networks, Incremental Learning, Few-shot Learning

1 INTRODUCTION

Graph-structured data, such as citation graphs [36], biomedical graphs [34], and social networks [9, 28], are nowadays ubiquitous in different real-world applications. Recently, a spectrum of Graph Neural Networks (GNNs) [1, 3, 9, 10, 16, 41, 45] has been proposed to model graph-structured data by transforming node features and propagating the embedded features along the graph structure. As a central task in graph machine learning, semi-supervised node classification aims to infer the missing labels for unlabeled nodes. By capturing the information carried by both labeled and unlabeled nodes as well as the relations between them, GNNs are able to achieve superior performance over other approaches.

However, most of the existing work on node classification primarily focuses on a single task, where the model is tasked to classify

the unlabeled nodes to a fixed set of classes [13, 16, 25, 51]. In practice, real-world graphs grow rapidly and novel node classes could emerge incrementally in different time periods. For example, E-commerce platforms like Amazon can be naturally modeled by graphs, where products are represented as nodes, and the interactions between products (e.g., viewed by the same customers) are represented as edges. As newly emerged product categories will be continually added to the platform, the underlying GNN models need to handle a sequence of incremental learning sessions for product categorization, where each session introduces a set of novel classes. It is worth mentioning that, different from those base classes (in the first learning session) provided with abundant labeled data, only a few labeled samples are available for those newly emerged classes in any incremental learning session. Ideally, the desired GNN model is supposed to be capable of accurately recognizing those novel classes introduced in a new session while preserving the performance on all the "seen" node classes in previous sessions. Such *Graph Few-shot Class-incremental Learning* (Graph FCL) problem has critical implications in both the academic and industrial communities. However, little effort has been devoted to this topic.

The main challenges of the Graph FCL problem center around the so-called *stability-plasticity* dilemma [27], which is a trade-off between the preservation of previously learned graph knowledge and the capability of acquiring new knowledge. Specifically, the novel classes in each new learning session have much fewer nodes compared to the base classes, resulting in a severe class-imbalance problem [12, 37, 44]. This may engender two potential problems: (1) on the one hand, if trained on all the data samples naively, the learned graph model could be substantially biased towards those base classes with significantly more nodes, resulting in the inertia to learn new node classes [37]. Moreover, to retain the existing knowledge, many of the methods from few-shot incremental learning [31, 46, 47] use a fixed feature encoder, which will not be updated after being pre-trained on base classes. Such a design will also exacerbate the difficulty of adapting the model to the new incremental learning tasks; (2) on the other hand, as the node classes from new tasks only have few-labeled samples, imposing the graph learning model to focus on new tasks will easily lead to overfitting to those new tasks and erase the existing knowledge for previously learned classes, which is known as *Catastrophic Forgetting* [8, 14, 18]. Considering the complex interactions between the nodes on graph-structured data, such learning errors would also be propagated on the graph and result in serious performance degradation. Thus, it is vital to explore and develop a new approach that can quickly adapt to the new class-incremental learning tasks while avoiding the forgetting of existing knowledge.

To address the aforementioned challenges of the studied problem, we first propose a new Graph Pseudo Incremental Learning (GPIL) training paradigm, which can facilitate the graph learning model to better adapt to new tasks. Initially, we split the original base dataset into base classes and pseudo novel classes with disjoint label spaces. Then, we pretrain the encoder on the base classes, and keep it trainable during the pseudo incremental learning process. For each episode during meta-training, all the few-shot node classification tasks are sampled from pseudo novel classes and base classes to mimic the incremental process in the evaluation phase. This way we obtain abundant meta-training episodes to learn a transferable model initialization for the incremental learning phase. We then propose a *Hierarchical Attention Graph Meta-learning* framework, HAG-Meta, which can effectively handle the stability-plasticity dilemma. Specifically, our framework uses a dynamically scaled cross-entropy loss regularizer where the scale factors [20, 24] are multiplied to each task-level loss to adjust their contribution for model training. Ideally, the scaling factors can help the model to down-weight the contribution of easy or insignificant tasks while focus on those hard or important tasks. Due to the fact that tasks in the Graph FCL problem have a hierarchical structure (nodes form classes, classes form tasks), we propose a hierarchical attention module that automatically captures the importance of different tasks and learns the scaling factors. On one hand, the *Task-Level Attention* (TLA) will estimate the importance of each task based on their aggregated prototypes and output the scaling factors to balance the contribution of different tasks. On the other hand, the *Node-Level Attention* (NLA) aims to learn prototypes that maintain a better balance between existing and novel knowledge within nodes and provide them to TLA. Being progressively trained in GPIL, the hierarchical attention module can gradually obtain the generalizability to produce the scaling factors for both encountered tasks and subsequent tasks. Training with this dynamically scaled regularizer, the proposed model will not only achieve better old knowledge consolidation but also acquire principled adaptability to new knowledge with merely limited data samples. The effectiveness of the proposed framework is validated with comprehensive experiments on three real-world datasets. The contribution of this work can be summarized as follows:

- **Problem:** We present a novel Graph Few-shot Incremental Learning problem and formulate it with node classification tasks, where the model is tasked to accomplish node classification on base classes and all few-shot novel classes encountered during incremental sessions.
- **Algorithm:** We propose Hierarchical Graph Attention modules tailored for the Graph FCL problem, and design a Graph Pseudo Incremental Learning paradigm to enable effective training to mimic the environment in the evaluation phase.
- **Evaluation:** Experiments on the Amazon-clothing, Reddit, and DBLP datasets show that the proposed framework significantly outperforms the baseline and other related state-of-the-art methods by a considerable margin.

2 RELATED WORK

2.1 Class-incremental Learning

Incremental learning (IL) [31, 46], also known as continual learning or lifelong learning, has drawn growing attention recently. IL aims to train machine learning models to acquire new knowledge while preserving the utmost existing knowledge. In this work, we mainly focus on *Class-incremental Learning* (CIL) where novel classes emerge in subsequent sessions and the model is tasked to fulfill classification on all the classes it has encountered, rather than *Task-incremental Learning* (TIL), where usually a task identifier is available, so the model can have multiple classifiers and finish the final classification on classes in a single task [15]. The mainstream of CIL methods can be categorized into two families. The first family includes replay-based methods [22, 30, 32, 35], which maintain a subset of previous samples, and train models together with samples in the new session. The other family of methods is regularization-based methods [2, 18], where various regularizers are proposed to regularize the parameters of a neural network so that more important parameters concerning the previous task can be protected when models are trained on each new task. A common choice of regularizer is a *Knowledge Distillation* (KD) [11] based loss proposed in LwF [19]. iCaRL [30] is the first work that combines both replay and KD regularization methods, and puts forward the data imbalance problem between old classes and novel classes in CIL. A series of work focuses on this problem [44, 48]. [31, 46, 47] further extend the situation to the *Few-shot Class-incremental Learning* (FCL) setting on image domain. For the FCL setting, KD performance will degrade tremendously due to the extreme scarcity of samples in novel classes. Instead, to overcome catastrophic forgetting, those methods usually adopt a decoupling method, where the encoder is fixed after pre-training, and extra modules are involved for learning incremental classes during meta-training. However, directly applying those FCL methods to the graph domain can lead to drastic performance degradation. Nodes in a graph are not i.i.d. data as usually assumed for images. Their representations are learned via sampling and aggregation from their neighbors. Fixing the encoder, if nodes in novel classes in an impending session are densely linked with nodes in base classes, their representation can have evident overlap, and the boundaries between those classes will be blurred. Since no existing work is suitable for FCL on graphs, our paper aims at bridging this gap.

2.2 Graph Few-shot Learning

Graph Neural Network (GNN) [1, 3, 9, 10, 16, 41, 45] is a family of deep neural models tailored for graph-structured data, which has been widely used in various applications, such as recommendation [42], anomaly detection [4], and text classification [5]. Generally, GNNs exploit a recurrent neighborhood aggregation strategy to preserve the graph structure information and transform the nodes' attributes simultaneously. For instance, variants of GCN [1, 3, 10, 16], GraphSAGE [9], GAT [41], and GIN [45] put forward different aggregation schemes to try to enhance the representation power of GNN. However, all those conventional GNNs may easily fail to learn expressive node representation when the labeled nodes are extremely scarce. Recently, increasing research attention has been devoted to graph few-shot learning problems. Especially, the

episodic meta-learning paradigm [38] has become the most popular strategy for this problem, which transfers knowledge learned from many similar FSL tasks. Based on it, Meta-GNN [50] applies MAML [7] to tackle the low-resource learning problem on graph. Furthermore, RALE [21] uses GNNs to encode graph path-based hubs and capture the task-level dependency, to achieve knowledge transfer. GPN [6] adopts Prototypical Networks [33] to make the classification based on the distance between the node feature and the prototypes. AMM-GNN [43] leverages an attribute-level attention mechanism to characterize the feature distribution differences between different tasks and learns more meaningful transferable knowledge across tasks. However, all those methods cannot be generalized to the Class-incremental learning scenarios, where the model is tested not only on the novel classes in the current task but also on all the classes in previous tasks. Catastrophic forgetting will erase the knowledge specific for previously learned classes.

3 PROBLEM STATEMENT

Formally, a graph $G = (\mathcal{V}, \mathcal{E}, \mathbf{X})$, where $\mathcal{V}$, $\mathcal{E}$, and $\mathbf{X}$ denote the set of nodes, edges and node features respectively, can be alternatively represented by $G = \{\mathbf{A}, \mathbf{X}\}$, where $\mathbf{A}$ is the adjacency matrix. The Graph FCL task assumes the existence of a sequence of homogeneous datasets within a graph, i.e., $\mathcal{D} = \{\mathcal{D}^0, \mathcal{D}^1, \dots, \mathcal{D}^i, \dots, \mathcal{D}^T\}$. In any session i , $\mathcal{C}^i$, the label space of the dataset $\mathcal{D}^i$, has no overlapping with the label space of any other session, i.e., $\forall i, j \in \{0, \dots, T\}, i \neq j, \mathcal{C}^i \cap \mathcal{C}^j = \emptyset$. Then, the dataset in each learning session can be represented as $\mathcal{D}^i = \{\mathbf{A}_{\mathcal{C}^i}, \mathbf{X}_{\mathcal{C}^i}\}$, where $\mathbf{A}_{\mathcal{C}^i}$ denotes the attributes of nodes whose labels belong to the label space $\mathcal{C}^i$. Notably, in the first session, the dataset $\mathcal{D}^0$ is a relatively large dataset where a sufficient amount of data is available for normal semi-supervised node classification training. The classes in $\mathcal{D}^0$ are the *base classes*. Datasets in following sessions, $\mathcal{D}^i \in \mathcal{D}, i \neq 0$, are few-shot datasets, the classes in which are named as *novel classes*. Now, we present the formal definition of a Few-shot Class-incremental Node Classification task:

Definition 3.1. Few-shot Class-incremental Node Classification on Graphs: For a specific session i , given a graph $\mathcal{G} = (\mathbf{A}, \mathbf{X})$, and a set of support nodes with labels, $\mathcal{S}^i$, from the label space $\mathcal{C}^i$, the model is tasked to predict labels for the nodes in corresponding query set $\mathcal{Q}^i$. The label space of the query set $\mathcal{Q}^i$ includes the base set $\mathcal{C}^0$, all the novel sets in previous sessions $\{\mathcal{C}^1, \mathcal{C}^2, \dots, \mathcal{C}^{i-1}\}$, and the novel set encountered in the current session $\mathcal{C}^i$.

For each session i , we have such a Few-shot Class-incremental Node Classification task $\mathcal{T}^i$. In the corresponding support set $\mathcal{S}^i$, we denote the number of novel classes as N and the number of support nodes in each class as K . This task is named as an **N -way K -shot** incremental node classification task. Alternatively, learning through that sequence of datasets can be represented as a sequence of tasks: i.e., $T = \{\mathcal{T}^0, \mathcal{T}^1, \dots, \mathcal{T}^i, \dots, \mathcal{T}^T\}$. In essence, we want our graph model able to retain a decent performance when fulfilling the classification on both base and novel classes.

4 METHODOLOGY

In this section, we introduce a Hierarchical Attention Graph Meta-learning framework, **HAG-Meta** for solving Graph FCL problem.

We first describe our proposed training procedures in Section 4.1. Then, we present the model proposed in Section 4.2. The overview of HAG-Meta is shown in Figure 1. Pseudocode-style algorithm descriptions are given in Algorithm 1 and Appendix C.

4.1 Graph Pseudo Incremental Learning

Data Splits: To solve the Graph FCL problem, We split the dataset $\mathcal{D}$ into two folds, $\mathcal{D}_{base}$ and $\mathcal{D}_{novel}$, with disjoint categories. $\mathcal{D}_{base}$ is randomly split into $\mathcal{D}_{base/tr}$, $\mathcal{D}_{base/val}$, and $\mathcal{D}_{base/test}$ for pre-training. $\mathcal{D}_{novel}$ has three splits of $\mathcal{D}_{novel/tr}$, $\mathcal{D}_{novel/val}$, $\mathcal{D}_{novel/test}$, with disjoint categories. To simulate the Graph FCL problem, the target few-shot data $\{\mathcal{D}^1, \dots, \mathcal{D}^i, \dots, \mathcal{D}^T\}$ are sampled from $\mathcal{D}_{novel/test}$, and the corresponding nodes and edges are masked during pre-training and meta-training. The base data $\mathcal{D}^0$ consists of $\mathcal{D}_{base}$ and $\mathcal{D}_{novel/tr}$, where our proposed Graph Pseudo Incremental Learning is conducted. Details are given next.

Pre-training: We pre-train a GNN-based encoder g_θ on the split $\mathcal{D}_{base/tr}$, following the normal semi-supervised node classification process. The encoder is still trainable after pre-training.

Meta-training: To learn an initialization with more transferable meta-knowledge within the graph, here, we propose the **Graph Pseudo Incremental Learning (GPIL)** paradigm, where a model would be trained on $\mathcal{D}_{base/tr}$ and $\mathcal{D}_{novel/tr}$. Similar to the episodic meta-learning strategy, during each session i , we randomly sample N novel categories from $\mathcal{D}_{novel/tr}$, and K nodes per category to form a novel support set $\mathcal{S}_{novel}^i = \{(\mathbf{x}_j^i, y_j^i)\}_{j=1}^{N \times K}$. The query set is composed of samples from the base categories, the novel categories from the previous sessions, and N novel categories in the current session. We sample K samples from all those categories to form a query set $\mathcal{Q}^i = \mathcal{Q}_{base}^i \cup \mathcal{Q}_{novel}^i$. The novel categories in each session will be merged into the base categories for the next session. We will cache those novel support nodes in $\mathcal{S}_{novel}^i$ for the classification of old classes in later sessions. So similarly, $\mathcal{S}^i = \mathcal{S}_{base}^i \cup \mathcal{S}_{novel}^i$. During each session i , the parameters are updated by the loss proposed in Section 4.2.1 for the classification of queries in $\mathcal{Q}^i$. During training, we reset the base categories and novel categories whenever the number of left novel categories is less than N . In this way, the proposed model can be trained on an arbitrary number of episodes despite the limitation of the number of novel categories in $\mathcal{D}_{novel/tr}$, which is crucial for training the attention models (See details in Section 4.2).

Evaluation: For each evaluation session, we randomly sample N novel categories from $\mathcal{D}_{novel/test}$ (or $\mathcal{D}_{novel/val}$ for validation). The proposed model is fine-tuned on the sample. Base samples for test are sampled from $\mathcal{D}_{base/test}$ (or $\mathcal{D}_{base/val}$ for validation) and $\mathcal{D}_{novel/train}$. For the next session, those novel categories are merged into base categories.

4.2 Model

The baseline model we deploy is based on Prototypical Network (PN) [33]. We replace the multilayer perceptron (MLP) encoder in the original PN with a GNN encoder g_θ . We call it Proto-GNN for convenience. Then, given a graph $\mathcal{G} = (\mathbf{A}, \mathbf{X})$, the latent features can be defined as:

$$\mathbf{Z} = g_\theta(\mathbf{A}, \mathbf{X}) \quad (1)$$

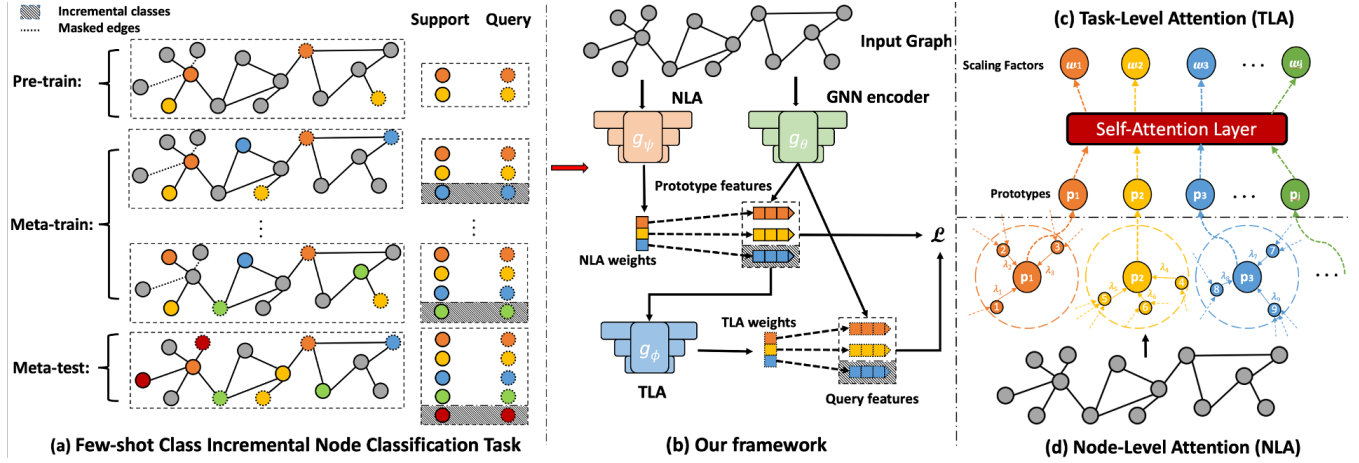


Figure 1: (a) Illustration of Few-shot Class-incremental Node Classification Task. (b) The illustration of our framework: HAG-Meta, as in Section 4.2. (c) Structure of the Task-Level Attention module, computed with self-attention layer, as in Section 4.2.1. (d) Structure of the Node-Level Attention module. We adopt GCN layers to generate the weight, as in Section 4.2.2.

As mentioned in [33], the vanilla way to compute a prototype for a category is the average over all the features of nodes in the category:

$$\mathbf{p}_k = \frac{1}{|S_k|} \sum_{j \in S_k} \mathbf{z}_j \quad (2)$$

Where S_k is the set of all the support nodes index in category k and $|S_k|$ is its cardinality. Then, a distance function d is used to generate the distribution over classes for a query node v_q based on a softmax over distances to the prototypes in the latent space:

$$p(y = k|q) = \frac{\exp(-d(\mathbf{z}_q, \mathbf{p}_k))}{\sum_{k'} \exp(-d(\mathbf{z}_q, \mathbf{p}_{k'}))} \quad (3)$$

For the choice of d , we use squared Euclidean distance, which has been shown as a simple and effective distance function [33].

To overcome the class-imbalance challenge in Section 1, we here propose our Hierarchical-Attention based Graph Attention module. The goal is to learn a strong regularizer that can dynamically scale the contribution for nodes in different tasks. Tasks in the Graph FCL problem naturally have a hierarchical structure (nodes form classes, classes form tasks). Thus, we propose a two-level hierarchical attention mechanism: *Task-Level Attention* and *Node-Level Attention* to estimate the importance of different tasks.

4.2.1 Task-Level Attention. To deal with the challenge that our model may overfit onto base or novel classes, we here propose the **Task-Level Attention (TLA)** to estimate the importance of classes learned in different tasks. Ultimately, we want to learn a series of scaling factors for the loss, which should be competent to automatically down-weight the contribution of easy or insignificant tasks during meta-training and rapidly focus the model on hard or important tasks. As described in the training procedure, during meta-training, the query nodes are from classes in all the previous training sessions. In this case, we find that a task-weighted loss will serve the purpose. For a PN, the prototypes matrix, $\mathbf{P}^i = \{\mathbf{p}_k^i\}$ ($\mathbf{p}_k^i \in \mathbb{R}^{|S_k| \times h}$, where h is the size of prototype features), of the support nodes, S^i , in a certain session i , serves as the classifier

for the queries, Q^i . Hence, we can make the hypothesis that the prototypes in sessions are representative enough to express the knowledge of the task $\mathcal{T}^i$. Based on the self-attention by [40], TLA aims at learning the attentions (scaling factors) $\mathbf{w}^{i,j} \in \mathbf{W}$, ($\forall j \in [1, i]$, $\mathbf{W} \in \mathbb{R}^{1 \times i}$) between the current task, $\mathcal{T}^i$, and all the tasks the model has been trained on, $\{\mathcal{T}^1, \mathcal{T}^2, \dots, \mathcal{T}^i\}$, including the current task. The desirable property of TLA is that the attention mechanism is inductive and permutation invariant, which suits the Graph FCL problem where novel tasks and classes come in sequence. The structure of the TLA model is shown in Figure 1 (c). We use g_ψ to denote the TLA generator. Because the number of classes in the base is much larger than that of each novel task, we first use MLPs to project the prototypes of all the tasks into the same size:

$$\mathbf{u}^j = \text{MLP}(\mathbf{p}^j), \forall j \in [1, i] \quad (4)$$

Where $\mathbf{u}^j$ is the projections of the prototypes $\mathbf{p}^j$ at session j . Then, the weights $\mathbf{W}$ can be computed as:

$$\mathbf{w}^{i,j} = \frac{\exp(\mathbf{u}^i \cdot \mathbf{u}^j)}{\sum_{j'=1}^i \exp(\mathbf{u}^i \cdot \mathbf{u}^{j'})} \quad (5)$$

For each task, the weight $\mathbf{W}$ is determined by the number of classes in the task, and then normalized by the number of classes in that task (different classes in the same task share the same weight):

$$\mathbf{W}_C = \frac{\tilde{\mathbf{W}}}{|C^j|} \quad (6)$$

where $|C^j|$ is the number of classes in session j , and $\tilde{\mathbf{W}} \in \mathbb{R}^{1 \times |C|}$ is the expanded weight vector of $\mathbf{W}$. $|C|$ is the number of all the classes having been seen. And $\mathbf{W}_C$ is the target scaling factor of classes within the task. With all the factors $\mathbf{W}_C$ computed, finally we introduce our **TLA Loss**, which is the Cross-Entropy Loss (CEL) scaled by the TLA scaling factors $\mathbf{W}_C$:

$$\mathcal{L}_{TLA} = \sum_{k \in C} w_k \cdot [y_k \cdot \log(\hat{y}_k) + (1 - y_k) \log(1 - \hat{y}_k)] \quad (7)$$

The final loss is:

$$\begin{aligned}\mathcal{L} &= \mathcal{L}_{CEL} + \mathcal{L}_{TLA} \\ &= \sum_{k \in C} (1 + w_k) \cdot [y_k \cdot \log(\hat{y}_k) + (1 - y_k) \log(1 - \hat{y}_k)]\end{aligned}\quad (8)$$

$\mathcal{L}_{CEL}$ will function as the initialization of task contribution weight, and $\mathcal{L}_{TLA}$ will turn into a regularizer, which adjusts the contribution according to the importance of tasks.

4.2.2 Node-Level Attention. While TLA can weigh the importance of different graph tasks, it cannot fully capture the knowledge within the graph structure, which may lead to the inaccurate importance measured for tasks. To incorporate TLA with the graph knowledge, we propose to use a **Node-Level Attention (NLA)**, $\Lambda = \{\lambda_j\}, j \in \mathcal{S}^i$, to adjust the representation of the prototype features learned from the GNN encoder g_θ in each session i , where λ_j is the attention weight of the novel support nodes v_j to the prototype. We want the NLA to maintain the balance of existing knowledge and novel knowledge for each node. The schematic diagram of NLA can be found in Figure 1 (d). We propose to use the Graph Convolutional Network (GCN) [16] to calculate the final NLA. The propagation rule in the l th layer of our GCN can be represented as:

$$\mathbf{h}_j^l = \sigma'(\mathbf{R}^l(\mathbf{h}_j^{l-1} + \sum_{j' \in \mathcal{N}_j} \mathbf{h}_{j'}^{l-1} / \sqrt{d_j d_{j'}})) \quad (9)$$

where, at the l th GCN layer, $\mathbf{h}_j^l$ is the latent NLA representation of node v_j , σ' is an activation function, $\mathbf{R}^l$ is the learned weight matrix, $\mathcal{N}_j$ is a set of nodes adjacent to node v_j , d_j and $d_{j'}$ are the node degrees of node v_j and node $v_{j'}$ respectively. We set $\mathbf{h}_j^0 = \mathbf{X}_j$. Then, we use an MLP to project the latent NLA in the last (L)th GCN layer into scalar:

$$\lambda_j = \text{MLP}(\mathbf{h}_j^L) \quad (10)$$

Next, we apply the centrality adjustment method proposed in [26]:

$$\tilde{\lambda}_j = \sigma(\log(\text{degree}(v_j) + \epsilon) \cdot \lambda_j) \quad (11)$$

where σ is the sigmoid function, ϵ is a small constant. Finally we use softmax function to normalize the NLA:

$$\lambda_j = \frac{\exp(\tilde{\lambda}_j)}{\sum_{j' \in \mathcal{S}^i} \exp(\tilde{\lambda}_{j'})} \quad (12)$$

Then, we modify the original strategy, Eq.(2), to calculate prototype in PN with NLA:

$$\mathbf{p}_k = \sum_{j \in \mathcal{S}_k} \lambda_j \mathbf{z}_i \quad (13)$$

With the adjusted prototypes, we can then use Eq.(3) to get the final label. An overview of the incremental training procedure of each session is given in Appendix C.

5 EXPERIMENTS

In this section, we present the evaluation of our framework: HAG-Meta. We first introduce the used datasets and compared methods. Then, we show the result and analysis of the comparative study. Furthermore, we conduct comprehensive ablation experiments to validate the effectiveness of individual components in the proposed framework and study their characteristics. Also, we compare the

Algorithm 1 HAG-META

Input: Dataset $\mathcal{D}$, number of sessions M for GPIL, number of target evaluation sessions T for evaluation, random initialized GNN model g_θ , random initialized TLA weightier g_ϕ , and NLA weightier g_ψ .

Output: Trained models: g_θ , g_ϕ , and g_ψ .

```
// Data split
1: Split dataset  $\mathcal{D}$  into  $\mathcal{D}_{base/tr}$ ,  $\mathcal{D}_{base/val}$ ,  $\mathcal{D}_{base/test}$ ,  $\mathcal{D}_{novel/tr}$ ,  $\mathcal{D}_{novel/val}$ , and  $\mathcal{D}_{novel/test}$ .
// GNN back bone Pre-training
2: Pre-train  $g_\theta$  on  $\mathcal{D}_{base/tr}$  in normal supervised learning.
// Meta training (GPIL)
3: while  $i \leq M$  do
4:   Sample a Graph FCL task from  $\mathcal{D}_{novel/tr}$  and  $\mathcal{D}_{base/tr}$  according to Section 4.1:  $\mathcal{T}_{tr}^i = \{\mathcal{S}^i, \mathcal{Q}^i\}$ .
5:   Do one incremental training session:
      $\mathcal{A}^i, g_\theta, g_\phi, g_\psi = \text{INCREMENTALSESSION}(\mathcal{T}_{tr}^i, g_\theta, g_\phi, g_\psi)$ 
6: end while
// Evaluation
7: while  $j \leq T$  do
8:   Sample a Graph FCL task from  $\mathcal{D}_{novel/test}$  and  $\mathcal{D}_{base/test}$  according to Section 4.1:  $\mathcal{T}_{test}^j = \{\mathcal{S}^j, \mathcal{Q}^j\}$ .
9:   Do one incremental training session:
      $\mathcal{A}^j, g_\theta, g_\phi, g_\psi = \text{INCREMENTALSESSION}(\mathcal{T}_{test}^j, g_\theta, g_\phi, g_\psi)$ 
10: end while
```

result of our model with that of the best baseline under different N -way K -shot settings. Finally, to illustrate the advantage of our model, we visualize the learned embeddings.

5.1 Experiment Settings

Evaluation Datasets. We conduct our experiments ¹ on three widely used graph Few-shot learning datasets: Amazon-Clothing [23], DBLP [36], and Reddit [9]. More details about the datasets can be found in Appendix A. The statistic are shown in Table 1.

Compared Methods. In this paper, we compare our HAG-Meta framework with the following methods:

- Prototypical Networks on graphs: As discussed in Section 4.2, we implement Proto-GNN with two different encoders: GCN [16] and GAT [41], to reveal the inability of GNN to deal with the Graph FCL problem. We denote them as Proto-GCN and Proto-GAT.
- State-of-the-art Graph Few-shot learning methods: Meta-GNN [50] and GPN [6].
- Continual learning methods on graph: ER-GNN [49], and classic iCaRL [30] with the encoder substituted by a GNN. These models do not consider the Few-shot learning setting.
- Few-shot Class-incremental Learning: CEC [47]. It is one of the state-of-the-art methods but is primarily for the image domain, so we replace the encoder with a GNN encoder.

¹Codes are available at <https://github.com/Zhen-Tan-dmml/GFCIL>

	Amazon Clothing	DBLP	Reddit
# Nodes	24,919	40,672	232,965
# Edges	91,680	288,270	11,606,919
# Features	9,034	7,202	602
# Labels	77	137	41
Pre-train	20	37	11
Meta-train	30	50	10
Evaluation	27	50	20

Table 1: Statistics of the experimental datasets.

5.2 Implementation Details

In Table 1, we list the specific data split strategy for each dataset, following Section 4.1. Here, $|C_{base}|$ is the number of classes for pre-training, $|C_{novel/tr}|$ the number of classes for meta-training, and $|C_{novel/test}|$ the number of target few-shot classes for evaluation. E.g., for the Amazon-Clothing dataset, we pre-train our GNN encoder on 20 categories of nodes, which are viewed as base categories. Another 30 categories are used for meta-train for providing pseudo novel categories. For evaluation, the model will be fine-tuned on a sequence of tasks consisting of nodes in the remaining 27 target categories. We stop the encoder pre-training when its validation accuracy stops improving for more than 10 epochs. For model implementation, please refer to Appendix B for detail.

Evaluation Protocol: We evaluate the model after each session with the test set $\mathcal{D}^i$ (sampled from $\mathcal{D}_{novel/test}$). To reduce fluctuation, all the accuracy (Acc.) scores reported are averaged over 10 random seeds. We also calculate a performance dropping rate (PD) that measures the absolute accuracy drops in the last session w.r.t. the accuracy before the first evaluation session, i.e., $PD = \mathcal{A}^0 - \mathcal{A}^T$ where $\mathcal{A}^0$ is the classification accuracy in the last Meta-train session and $\mathcal{A}^T$ is the accuracy in the last session. To make the result more explicit, we define a new term: Relative Performance Dropping rate (RPD), which is the PD normalized by the initial accuracy, i.e., $RPD = \frac{PD}{\mathcal{A}^0}$.

5.3 Comparative Results

In this section, we present the comparison between our framework and the other four categories of baseline methods described in Section 5.1. To the best of our knowledge, we are the first to investigate the Graph Few-shot Class-incremental Learning Problem. To fairly compare those methods, all methods except the basic GAT model share the same pre-trained GNN encoder as the proposed framework HAG-Meta, a 2-layer GCN. Also, when experimented on each dataset, they share the same random seeds for data split, leading to identical evaluation data. We justify the advantage of the proposed framework from the following aspects:

Performance Degradation in Graph Few-shot Learning (GFSL) methods: A general observation is that, for those GFSL methods, their accuracy decreases substantially as new sessions emerge, especially for the first several sessions. Even though through meta-training, they have gained generalizability to a certain extent, the performance degrades constantly as the number of classes involved increases. This implies that the existing GFSL methods cannot maintain discriminative boundaries between all the base classes and novel classes. Without consolidated knowledge of classes in base and early sessions, traditional GFSL methods suffer grievously from Catastrophic Forgetting, as it adapts to classes in the latest episode.

Limitation in existing Incremental Learning methods. ER-GNN is one of the pioneers to task graph neural network models with a sequence of tasks. It adopts several Experience Replay methods to try to consolidate existing knowledge. However, the accuracy of ER-GNN diminishes tremendously when it is applied to a few-shot setting. Besides, iCaRL combines a piece of memory and knowledge distillation to consolidate existing knowledge. However, the limited number of samples in novel classes will affect the knowledge distillation process, leading to its overfitting onto old classes. Furthermore, the CEC method is one of the state-of-the-art few-shot Class-incremental learning methods for the image domain. It assumes that the data is i.i.d. distributed, which makes it overlook the topological relationships among tasks and nodes. It adopts a decoupling strategy, where the encoder is fixed after pre-training to retain existing knowledge. So when it is applied to graphs, the representation of nodes in base classes is fixed through all following sessions. But in graphs, the representation of novel classes nodes tightly depends on their neighboring nodes, which might belong to the base classes. This leads to indiscriminative boundaries between the base and the novel classes and unsatisfactory accuracy.

Advantages of the proposed HAG-Meta: Generally, HAG-Meta outperforms all baseline methods by a large margin, in terms of accuracy, PD, and RPD. Compared to the GFSL methods, the proposed penalty term $\mathcal{L}_{TLA}$ can effectively prevent catastrophic forgetting by regularization. In contrast with ER-GNN and iCaRL, our GPIL paradigm provides sufficient episodes to train the model such that it can learn an appropriate initialization to adapt to few-shot novel data. Also, we discard the decoupling strategy of the mainstream incremental learning methods on the image domain, like CEC. In contrast, we add node level attention, NLA, to learn the task and class dependency within the graph topological structure. The results in Table 2 verify that HAG-Meta is robust against increasing numbers of sessions with novel classes in the Graph FCL problem.

5.4 Ablation Study

In this section, we conduct more experiments to investigate the effectiveness of different components in our framework. We present the results of experiments on Amazon-Clothing datasets, under the 3-way 5-shot setting (similar results can be observed on the other datasets and settings). The results are shown in Table 3. For each method, the models share the same data splits for evaluation.

Specifically, the variant Proto-GNN stands for the Prototypical Graph Neural Networks baseline described in Section 4.2. NLA and TLA represent Node-Level Attention and Task-level Attention, respectively. A method without the TLA means the loss is the vanilla cross-entropy loss, without the $\mathcal{L}_{TLA}$ regularizer, see Eq.(8). A method without the NLA means there is no prototype representation adjustment. All the prototype representations are the vanilla average of the encoder output, see Eq.(2). For methods without GPIL, the model is directly fine-tuned on datasets with target novel classes in different sessions.

According to the results shown in Table 3, we can observe that each component in our proposed framework, HAG-Meta, is effective when tackling the Graph FCL problem:

- GPIL: Comparison between methods with and without GPIL shows that GPIL can improve the accuracy. With more episodes

Amazon-Clothing dataset (3-way 5-shot)													
Method	Acc. in each session (%) $\uparrow$										PD $\downarrow$	RPD $\downarrow$	improvement (PD/RPD)
	0	1	2	3	4	5	6	7	8	9			
Proto-GAT	60.22	44.32	41.87	38.28	35.15	32.25	30.67	28.54	26.54	25.43	34.79	57.77	(+1.43/+18.13)
Proto-GCN	60.52	43.11	42.41	39.80	36.91	33.84	31.47	29.67	27.18	25.76	34.76	57.44	(+1.40/+17.80)
Meta-GNN	79.62	65.29	63.46	58.49	56.36	54.13	51.76	49.50	46.32	43.07	36.55	45.90	(+3.19/+6.26)
GPN	80.76	66.98	64.46	62.63	60.03	55.48	52.35	50.76	48.15	45.83	34.93	43.25	(+1.57/+3.61)
ER-GNN	81.37	73.62	70.84	63.63	61.86	58.24	54.29	51.58	48.86	46.26	35.11	43.14	(+1.75/+3.50)
iCaRL	79.43	73.92	66.87	63.19	60.28	56.04	53.48	50.33	47.75	45.82	33.61	42.31	(+0.25/+2.67)
CEC	82.05	73.28	70.46	64.85	62.19	60.29	54.86	52.27	50.69	48.22	34.28	41.77	(+0.92/+2.13)
Ours	84.15	75.32	71.35	67.32	64.03	61.42	56.23	54.63	52.65	50.79	33.36	39.64	

DBLP dataset (5-way 5-shot)													
Method	Acc. in each session (%) $\uparrow$										PD $\downarrow$	RPD $\downarrow$	improvement (PD/RPD)
	0	1	2	3	4	5	6	7	8	9			
Proto-GAT	38.75	32.04	23.64	19.65	19.43	18.64	18.86	18.59	18.35	18.23	20.52	52.96	(+3.41/+21.62)
Proto-GCN	39.02	33.11	24.27	20.14	20.03	19.12	18.76	18.32	18.50	18.66	20.36	52.17	(+3.25/+20.83)
Meta-GNN	40.75	34.28	33.68	32.18	30.45	28.45	24.34	23.18	23.07	22.23	18.52	45.44	(+1.41/+14.10)
GPN	41.39	34.53	32.18	32.00	31.25	30.67	29.16	27.64	25.34	23.13	18.26	44.12	(+1.15/+12.78)
ER-GNN	46.74	41.70	38.55	36.82	35.03	34.28	33.32	31.59	30.67	28.86	17.88	38.25	(+0.77/+6.91)
iCaRL	45.76	40.03	37.92	36.56	34.81	33.22	32.10	30.57	29.34	28.24	17.52	38.29	(+0.41/+6.95)
CEC	46.45	40.25	38.28	36.67	35.26	34.48	32.24	31.68	30.84	28.76	17.69	38.08	(+0.58/+6.74)
Ours	54.59	47.89	46.45	45.05	43.83	41.38	39.67	39.07	38.15	37.48	17.11	31.34	

Reddit dataset (2-way 3-shot)													
Method	Acc. in each session (%) $\uparrow$										PD $\downarrow$	RPD $\downarrow$	improvement (PD/RPD)
	0	1	2	3	4	5	6	7	8	9			
Proto-GAT	48.23	42.07	37.52	33.43	32.38	31.05	28.32	26.42	24.39	22.47	25.76	53.41	(+5.14/+19.43)
Proto-GCN	48.04	42.77	37.35	35.04	33.62	31.21	28.04	25.86	24.42	22.61	25.43	52.94	(+4.81/+18.96)
Meta-GNN	53.14	48.56	45.63	42.52	40.42	38.20	35.12	32.21	30.68	29.45	23.69	44.58	(+5.07/+10.60)
GPN	55.28	51.84	46.36	43.71	41.18	39.07	37.82	35.04	32.48	31.65	23.63	42.75	(+3.01/+8.77)
ER-GNN	52.86	47.29	45.28	43.56	41.08	40.02	38.42	36.77	33.26	30.47	22.39	42.36	(+1.77/+8.38)
iCaRL	54.62	50.58	48.72	46.23	44.84	42.16	40.29	38.65	35.74	33.28	21.34	39.07	(+0.72/+5.09)
CEC	57.68	53.13	50.63	48.37	46.76	43.13	41.28	39.68	37.42	36.61	21.07	36.53	(+0.45/+2.55)
Ours	60.68	53.26	52.68	50.82	49.37	47.25	45.86	43.16	42.28	40.06	20.62	33.98	

Table 2: Comparative Results on the three datasets under different N-way K-shot settings.

Method	Acc. in each session (%) $\uparrow$											PD $\downarrow$	RPD $\downarrow$
	0	1	2	3	4	5	6	7	8	9			
Proto-GNN	60.52	43.11	42.41	39.80	36.91	33.84	31.47	29.67	27.18	25.76	34.76	57.44	
Proto-GNN + GPIL	78.24	68.77	62.48	55.20	51.48	45.65	42.89	37.42	34.62	32.52	45.72	58.44	
Proto-GNN + NLA	76.68	65.33	62.04	60.64	56.88	53.37	50.65	46.96	43.26	40.40	36.28	47.31	
Proto-GNN + TLA	75.54	63.87	60.78	58.27	56.24	53.41	51.78	48.61	44.18	41.35	34.19	45.26	
Proto-GNN + NLA + GPIL	80.75	74.16	70.48	64.88	61.56	57.71	54.65	50.25	48.89	44.5	36.25	44.89	
Proto-GNN + TLA + GPIL	79.94	73.67	70.01	65.21	62.33	58.24	54.14	51.33	48.65	45.32	34.62	43.31	
Proto-GNN + NLA + TLA	77.03	72.24	68.75	63.02	60.38	56.41	53.96	49.23	45.48	42.93	34.10	44.27	
Proto-GNN + NLA + TLA + GPIL	84.15	75.32	71.35	67.32	64.03	61.42	56.23	54.63	52.65	50.79	33.36	39.64	

Table 3: Ablation results on Amazon-Clothing dataset (3-way 5-shot).

of training. GPIL equips the hierarchical-attention module with better initialization to capture the importance among nodes and tasks.

- TLA and NLA: Generally, the scores show that both attention components help improve accuracy. Specifically, comparing methods with solely NLA or TLA, we can find that, at first several sessions, methods with NLA have higher scores, but

their scores decrease at a sharper rate compared to methods with only TLA. This shows that NLA is better in terms of capturing the latent information within the graph data from limited novel class samples, but it suffers from catastrophic forgetting. In contrast, methods with TLA can maintain a higher score at later sessions, which echoes the purpose of our design, that the TLA loss can alleviate the forgetting

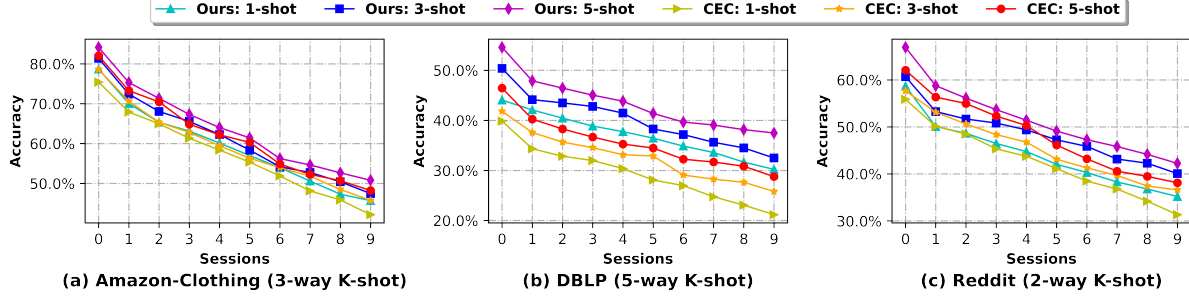


Figure 2: Parameter analysis on three real-world datasets: (a) Amazon-Clothing (b) DBLP (c) Reddit. For each dataset we experiment under three N-way k-shot settings.

problem by focusing more on important hard-to-learn tasks than insignificant easy-to-learn tasks.

- Our framework, HAG-Meta, containing NLA, TLA, and GPIL, as shown in the last line in Table 3, achieves the best performance. Through effective training in GPIL, the hierarchical attention components, TLA and NLA, have learned an initialization that can capture the topological information within the input graph to adapt to few-shot novel classes and mitigate forgetting of the existing knowledge simultaneously.

5.5 Parameter Analysis

In this section, we present extensive experiments to analyze the sensitivity of our HAG-Meta to the number of node classes (N -way) and support(query) set size (K -shot). For better comparison, we include the accuracy for both our HAG-Meta and the best comparable model CEC under different N -way K -shot settings for all three real-world datasets. As the result shown in Figure 2, our model outperforms CEC in every setting we test. Plus, with more supervisory signals, the classification accuracies of both models are higher with larger support and query set.

5.6 Visualization

To illustrate the effectiveness of our framework, we use the t-SNE [39] method to visualize the embedding after the second session of evaluation for the large-scale Reddit dataset under the 2-way 5-shot setting. As shown in Figure 3, each color signifies a class. It's evident that our model is capable to produce projection embedding that elicits much more discriminative decision boundaries than our best baseline model CEC.

6 CONCLUSION

In this work, we propose a new problem, the Graph Few-shot Class-incremental Learning (Graph FCL) problem, aiming at obtaining a graph model that can adapt to new tasks with a restricted number of labeled training samples in novel classes, and simultaneously keep a good performance on old tasks. We formalize it with the node classification task and propose a novel framework called HAG-Meta. We present a novel Graph Pseudo Incremental Learning paradigm, which allows our model to learn a generalizable initialization for the evaluation phase. Then we propose a hierarchical-attention-based module to solve the class-imbalance problem in Graph FCL.

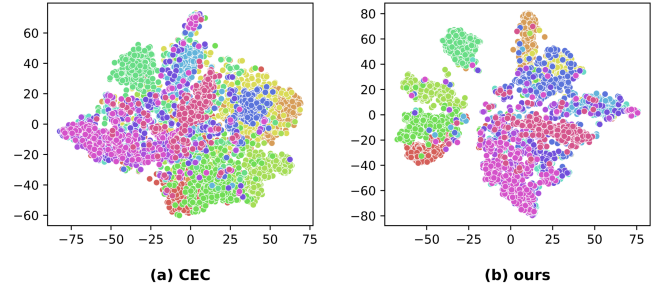


Figure 3: t-SNE embedding visualization: (a) CEC (b) ours.

Primarily, the Task-Level Attention will be trained to estimate the importance between different tasks for backpropagation, while the Node-Level Attention incorporates the Task-Level Attention with the ability to capture the knowledge within the graph structure. A dynamically scaled loss regularizer is then computed from the task importance and automatically adjusts the contribution of different tasks for training. We conduct experiments on real-world datasets to demonstrate the effectiveness and advantage of our framework.

Despite the promising results, the Graph FCL problem is far from being solved. In particular, many other graph settings are worth considering, such as the dynamic graph scenario, where the graph structure is continually evolving in different learning sessions. Besides, more sophisticated methods to design the regularizer, like from the perspective of causality, are worth considering. Also, we plan to study the scenario where the model cannot explicitly store any of the training data [30], e.g. for privacy issues. A potential direction could be to investigate the efficient methods to store or reproduce embedded features [17, 29].

REFERENCES

- [1] Shaosheng Cao, Wei Lu, and Qionghai Xu. 2016. Deep neural networks for learning graph representations. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 30.
- [2] Arslan Chaudhry, Puneet K Dokania, Thalaiyasingam Ajanthan, and Philip HS Torr. 2018. Riemannian walk for incremental learning: Understanding forgetting and intransigence. In *Proceedings of the European Conference on Computer Vision (ECCV)*. 532–547.
- [3] Michaël Defferrard, Xavier Bresson, and Pierre Vandergheynst. 2016. Convolutional neural networks on graphs with fast localized spectral filtering. *arXiv preprint arXiv:1606.09375* (2016).
- [4] Kaize Ding, Jundong Li, Rohit Bhanushali, and Huan Liu. 2019. Deep anomaly detection on attributed networks. In *SDM*.

- [5] Kaize Ding, Jianling Wang, Jundong Li, Dingcheng Li, and Huan Liu. 2020. Be more with less: Hypergraph attention networks for inductive text classification. In *EMNLP*.
- [6] Kaize Ding, Jianling Wang, Jundong Li, Kai Shu, Chenghao Liu, and Huan Liu. 2020. Graph prototypical networks for few-shot learning on attributed networks. In *Proceedings of the 29th ACM International Conference on Information & Knowledge Management*. 295–304.
- [7] Chelsea Finn, Pieter Abbeel, and Sergey Levine. 2017. Model-Agnostic Meta-Learning for Fast Adaptation of Deep Networks. arXiv:1703.03400 [cs.LG]
- [8] Ian J Goodfellow, Mehdi Mirza, Da Xiao, Aaron Courville, and Yoshua Bengio. 2013. An empirical investigation of catastrophic forgetting in gradient-based neural networks. *arXiv preprint arXiv:1312.6211* (2013).
- [9] William L. Hamilton, Rex Ying, and Jure Leskovec. 2018. Inductive Representation Learning on Large Graphs. arXiv:1706.02216 [cs.SI]
- [10] Mikael Henaff, Joan Bruna, and Yann LeCun. 2015. Deep convolutional networks on graph-structured data. *arXiv preprint arXiv:1506.05163* (2015).
- [11] Geoffrey Hinton, Oriol Vinyals, and Jeff Dean. 2015. Distilling the Knowledge in a Neural Network. arXiv:1503.02531 [stat.ML]
- [12] Saihui Hou, Xinyu Pan, Chen Change Loy, Zilei Wang, and Dahua Lin. 2019. Learning a unified classifier incrementally via rebalancing. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 831–839.
- [13] Fenyu Hu, Yanqiao Zhu, Shu Wu, Liang Wang, and Tieniu Tan. 2019. Hierarchical graph convolutional networks for semi-supervised node classification. *arXiv preprint arXiv:1902.06667* (2019).
- [14] Ronald Kemker, Marc McClure, Angelina Abitino, Tyler Hayes, and Christopher Kanan. 2018. Measuring catastrophic forgetting in neural networks. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 32.
- [15] Jong-Yeong Kim and Dong-Wan Choi. 2021. Split-and-Bridge: Adaptable Class Incremental Learning within a Single Neural Network. arXiv:2107.01349 [cs.LG]
- [16] Thomas N Kipf and Max Welling. 2016. Semi-supervised classification with graph convolutional networks. *arXiv preprint arXiv:1609.02907* (2016).
- [17] Thomas N. Kipf and Max Welling. 2016. Variational Graph Auto-Encoders. arXiv:1611.07308 [stat.ML]
- [18] James Kirkpatrick, Razvan Pascanu, Neil Rabinowitz, Joel Veness, Guillaume Desjardins, Andrei A Rusu, Kieran Milan, John Quan, Tiago Ramalho, Agnieszka Grabska-Barwinska, et al. 2017. Overcoming catastrophic forgetting in neural networks. *Proceedings of the national academy of sciences* 114, 13 (2017), 3521–3526.
- [19] Zhizhong Li and Derek Hoiem. 2017. Learning without Forgetting. arXiv:1606.09282 [cs.CV]
- [20] Tsung-Yi Lin, Priya Goyal, Ross Girshick, Kaiming He, and Piotr Dollár. 2017. Focal loss for dense object detection. In *Proceedings of the IEEE international conference on computer vision*. 2980–2988.
- [21] Zemin Liu, Yuan Fang, Chenghao Liu, and Steven CH Hoi. 2021. Relative and absolute location embedding for few-shot node classification on graph. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 35. 4267–4275.
- [22] David Lopez-Paz and Marc Aurelio Ranzato. 2017. Gradient episodic memory for continual learning. *Advances in neural information processing systems* 30 (2017), 6467–6476.
- [23] Julian McAuley, Rahul Pandey, and Jure Leskovec. 2015. Inferring Networks of Substitutable and Complementary Products. arXiv:1506.08839 [cs.SI]
- [24] Aditya Krishna Menon, Sadeep Jayasumana, Ankit Singh Rawat, Himanshu Jain, Andreas Veit, and Sanjiv Kumar. 2020. Long-tail learning via logit adjustment. *arXiv preprint arXiv:2007.07314* (2020).
- [25] Kenta Oono and Taiji Suzuki. 2019. Graph neural networks exponentially lose expressive power for node classification. *arXiv preprint arXiv:1905.10947* (2019).
- [26] Lawrence Page, Sergey Brin, Rajeev Motwani, and Terry Winograd. 1999. *The PageRank citation ranking: Bringing order to the web*. Technical Report. Stanford InfoLab.
- [27] Brian L Pulito, T Raju Damarla, and S Nariani. 1990. A two-dimensional shift invariant image classification neural network which overcomes the stability/plasticity dilemma. In *1990 IJCNN International Joint Conference on Neural Networks*. IEEE, 825–833.
- [28] Guo-Jun Qi, Charu Aggarwal, Qi Tian, Heng Ji, and Thomas Huang. 2011. Exploring context and content links in social media: A latent space method. *IEEE Transactions on Pattern Analysis and Machine Intelligence* 34, 5 (2011), 850–862.
- [29] Amal Rannen, Rahaf Aljundi, Matthew B. Blaschko, and Tinne Tuytelaars. 2017. Encoder Based Lifelong Learning. *2017 IEEE International Conference on Computer Vision (ICCV)* (Oct 2017). <https://doi.org/10.1109/iccv.2017.148>
- [30] Sylvestre-Alvise Rebuffi, Alexander Kolesnikov, Georg Sperl, and Christoph H. Lampert. 2017. iCaRL: Incremental Classifier and Representation Learning. arXiv:1611.07725 [cs.CV]
- [31] Mengye Ren, Renjie Liao, Ethan Fetaya, and Richard S. Zemel. 2019. Incremental Few-Shot Learning with Attention Attractor Networks. arXiv:1810.07218 [cs.LG]
- [32] David Rolnick, Arun Ahuja, Jonathan Schwarz, Timothy P. Lillicrap, and Greg Wayne. 2019. Experience Replay for Continual Learning. arXiv:1811.11682 [cs.LG]
- [33] Jake Snell, Kevin Swersky, and Richard S. Zemel. 2017. Prototypical Networks for Few-shot Learning. arXiv:1703.05175 [cs.LG]
- [34] Aravind Subramanian, Pablo Tamayo, Vamsi K Mootha, Sayan Mukherjee, Benjamin L Ebert, Michael A Gillette, Amanda Paulovich, Scott L Pomeroy, Todd R Golub, Eric S Lander, et al. 2005. Gene set enrichment analysis: a knowledge-based approach for interpreting genome-wide expression profiles. *Proceedings of the National Academy of Sciences* 102, 43 (2005), 15545–15550.
- [35] Binh Tang and David S. Matteson. 2021. Graph-Based Continual Learning. arXiv:2007.04813 [stat.ML]
- [36] Jie Tang, Jing Zhang, Limin Yao, Juanzi Li, Li Zhang, and Zhong Su. 2008. Arnetminer: extraction and mining of academic social networks. In *Proceedings of the 14th ACM SIGKDD international conference on Knowledge discovery and data mining*. 990–998.
- [37] Xiaoyu Tao, Xiaopeng Hong, Xinyuan Chang, Songlin Dong, Xing Wei, and Yihong Gong. 2020. Few-Shot Class-Incremental Learning. arXiv:2004.10956 [cs.CV]
- [38] Sebastian Thrun and Lorien Pratt. 2012. *Learning to learn*. Springer Science & Business Media.
- [39] Laurens Van der Maaten and Geoffrey Hinton. 2008. Visualizing data using t-SNE. *Journal of machine learning research* 9, 11 (2008).
- [40] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Lukasz Kaiser, and Illia Polosukhin. 2017. Attention is all you need. *arXiv preprint arXiv:1706.03762* (2017).
- [41] Petar Veličković, Guillem Cucurull, Arantxa Casanova, Adriana Romero, Pietro Lio, and Yoshua Bengio. 2017. Graph attention networks. *arXiv preprint arXiv:1710.10903* (2017).
- [42] Jianling Wang, Kaize Ding, Liangjie Hong, Huan Liu, and James Caverlee. 2020. Next-item recommendation with sequential hypergraphs. In *SIGIR*.
- [43] Ning Wang, Minnan Luo, Kaize Ding, Lingling Zhang, Jundong Li, and Qinghua Zheng. 2020. Graph Few-shot Learning with Attribute Matching. In *CIKM*.
- [44] Yue Wu, Yinpeng Chen, Lijuan Wang, Yuancheng Ye, Zicheng Liu, Yandong Guo, and Yun Fu. 2019. Large Scale Incremental Learning. arXiv:1905.13260 [cs.CV]
- [45] Keyulu Xu, Weihua Hu, Jure Leskovec, and Stefanie Jegelka. 2018. How powerful are graph neural networks? *arXiv preprint arXiv:1810.00826* (2018).
- [46] Sung Whan Yoon, Do-Yeon Kim, Jun Seo, and Jaekyun Moon. 2020. XtarNet: Learning to Extract Task-Adaptive Representation for Incremental Few-Shot Learning. arXiv:2003.08561 [cs.LG]
- [47] Chi Zhang, Nan Song, Guosheng Lin, Yun Zheng, Pan Pan, and Yinghui Xu. 2021. Few-Shot Incremental Learning with Continually Evolved Classifiers. arXiv:2104.03047 [cs.CV]
- [48] Bowen Zhao, Xi Xiao, Guojun Gan, Bin Zhang, and Shu-Tao Xia. 2020. Maintaining discrimination and fairness in class incremental learning. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 13208–13217.
- [49] Fan Zhou and Chengtai Cao. 2021. Overcoming Catastrophic Forgetting in Graph Neural Networks with Experience Replay. arXiv:2003.09908 [cs.LG]
- [50] Fan Zhou, Chengtai Cao, Kunpeng Zhang, Goce Trajcevski, Ting Zhong, and Ji Geng. 2019. Meta-gnn: On few-shot node classification in graph meta-learning. In *Proceedings of the 28th ACM International Conference on Information and Knowledge Management*. 2357–2360.
- [51] Chenyi Zhuang and Qiang Ma. 2018. Dual graph convolutional networks for graph-based semi-supervised classification. In *Proceedings of the 2018 World Wide Web Conference*. 499–508.

A DATASETS DESCRIPTION

In this section, we provide a detailed description of all those three datasets we experiment on. We follow the pre-processing procedures in [6].

- Amazon-Clothing [23] is a product network built with the products in “Clothing, Shoes, and Jewelry” on Amazon. In this dataset, each product is considered as a node and its description is used to construct the node attributes. We use the substitutable relationship (“also viewed”) to create links between products.
- DBLP [36] is a citation network between academic papers where each node represents a paper, and the links are the citation relations among different papers. The paper abstracts are used to construct node attributes. The class label of a node is defined as the paper venue.
- Reddit [9] is a large-scale post-to-post graph constructed with data sampled from Reddit, within which posts are represented by nodes and two posts are connected if they are commented by the same user. Each post is labeled with a community ID.

B IMPLEMENTATION DETAIL

For the model implementation, we implement the proposed framework in PyTorch. Specifically, the graph encoder g_θ consists of two GCN layers [16] with dimension size 32 and 16, respectively. Both of them are activated with the ReLU function. Regarding the TLA,

We use a 3-layer MLP to map the base prototypes into the same size as the novel prototypes in each session, namely, N . For the NLA, it consists of one fully connected layer and two GCN aggregation layers [16]. For each aggregation layer, we use ReLU function as the activation function.

The framework is trained with Adam optimizers whose learning rates are set to be 0.005 initially with a weight decay of 0.0005. And the coefficients for computing running averages of gradient and square are set to be $\beta_1 = 0.9$, $\beta_2 = 0.999$. For each dataset, we meta-train the model over 1000 episodes with an early-stopping strategy.

C PSEUDO-CODE STYLE DESCRIPTION OF EACH INCREMENTAL SESSION

Algorithm 2 INCREMENTALSESSION

Input: Sampled task $\mathcal{T} = \{\mathcal{S}, \mathcal{Q}\}$, g_θ , g_ϕ , and g_ψ .

Output: Accuracy $\mathcal{A}$, trained models: g_θ , g_ϕ , and g_ψ .

// One session of Graph FCL

- 1: Compute representation of nodes in $\mathcal{S}$ and $\mathcal{Q}$ with g_θ .
 - 2: Compute Λ and $\mathbf{p}$, with g_ψ .
 - 3: Compute $\mathcal{W}_C$, and labels for nodes in $\mathcal{Q}^i$ with g_ϕ .
 - 4: Compute TLA Loss and total Loss, $\mathcal{L}_{TLA}$ and $\mathcal{L}$, with Eq.(7) and Eq.(8) for backprop.
 - 5: Compute labels for nodes in $\mathcal{Q}$ and corresponding accuracy $\mathcal{A}$.
-