

Dual Prompt Tuning based Contrastive Learning for Hierarchical Text Classification

Anonymous ACL submission

Abstract

Hierarchical text classification aims at categorizing texts into multi-tiered tree-like label hierarchy. Existing methods pay more attention to capture hierarchy-aware text feature by exploiting explicit parent-child relationships, while interactions between peer labels are rarely taken into account, resulting in severe label confusions within each layer. In this work, we propose a novel Dual Prompt Tuning (DPT) method, which emphasizes to identify discrimination among peer labels by performing contrastive learning on each hierarchical layer. We design an innovative hand-crafted prompt containing slots for both positive and negative label predictions to cooperate with contrastive learning. In addition, we introduce a label hierarchy self-sensing auxiliary task to ensure cross-layer label consistency. Extensive experiments demonstrate that DPT achieves significant improvements and outperforms the current state-of-the-art methods on BGC and RCV1-V2 benchmark datasets.

1 Introduction

As a sub-task of text classification, hierarchical text classification (HTC) has broader applications in the realistic scenes, such as intent recognition in dialogue system, commodity and book management (Cevahir and Murakami, 2016; Aly et al., 2019), where a large number of categories are organized into a hierarchical tree structure. The ultimate goal of HTC task is to classify texts or documents from the top to bottom level through the hierarchical label tree. Due to the challenges of large-scale, imbalanced and complex label hierarchy (Mao et al., 2019), simply transferring flat multi-label text classification algorithms to HTC often fails to achieve sufficient performance.

Full use of the hierarchical structure of labels is the key to achieving well-performing classification in HTC tasks, which enables the model to

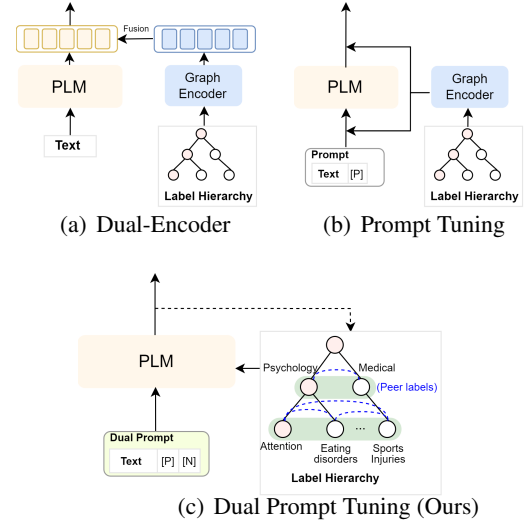


Figure 1: Architecture comparisons among existing methods and our proposed Dual Prompt Tuning.

predict labels that match the hierarchical relationship. Existing methods (Zhou et al., 2020; Deng et al., 2021; Chen et al., 2021; Zhu et al., 2023) apply dual-encoders framework to model the text and hierarchical structure separately, and then fusion them to obtain hierarchical label-wise text feature. Wang et al. (2022b) first proposes a hierarchy-aware prompt-tuning method, which incorporates the label hierarchy encoded by the Graph Attention Network into a soft prompt to bridge hierarchy and flat gap, as shown in Figure 1.

However, most studies pay close attention to exploit relations that explicitly displayed in the hierarchy, while interactions between "peer labels" which refer to a group of labels at the same hierarchical level are often neglected. PeerHTC (Song et al., 2023) recently tries to explore latent relevancy among peer labels with a complicated two-stage training procedure in which peer and adjacent level-wise feature are separately extracted by Graph Convolutional Neural Networks. Neverthe-

less, alleviating confusion within peer labels, especially fine-grained ones at lower level that share the same parent node, still remains challenging and highly valued.

To this end, we propose a novel **Dual Prompt Tuning (DPT)** method, aimed at alleviating label confusion between peer labels. We put forward Hierarchy-aware Peer-label Contrastive Learning (HierPCL) approach to extract discriminative pairwise representations. In detail, we create an original dual prompt template containing both positive and negative label slots, and then perform label-wise contrastive learning on the embeddings of both two slots. Dual prompt is multi-functional, targeted for predicting positive labels and recognizing incorrect but confused negatives at each hierarchical layer. Moreover, we design an adaptive hard negative sampling strategy and hierarchy-injected label representation method to further boost the performance of HierPCL.

Besides, we introduce a simple auxiliary Label Hierarchy Self-sensing task to keep our model in the best sense of holistic hierarchical structure. Instead of injecting label hierarchy into text semantics, we perform multi-task learning collaborating with label prediction to identify the correctness of each candidate label path. The basic idea of this task is to internalize structural hierarchy knowledge to ensure the cross-layer consistency of the final path prediction and improve the classification accuracy.

Our contributions are summarized as follows:

- We propose a Dual Prompt Tuning (DPT) method for HTC tasks to address label confusion between peers at each hierarchical layer, magnifying the power of prompts.
- We put forward Hierarchy-aware Peer-label Contrastive Learning approach based on DPT, which contributes to obtaining aligned and discriminative representation.
- We evaluate our proposed method on four popular benchmark datasets against the strong baselines. Experimental results demonstrate the advantage of our proposal.

2 Related Work

2.1 Hierarchical Text Classification

The HTC algorithms can generally be divided into *local* and *global* approaches (Zangari et al., 2023).

Local approaches construct multiple classifiers for different partitions of the label hierarchy tree usually in a "top-down" flow, for example, each node or each hierarchical level. Although there is a certain degree of connection between multiple classifiers, it is inevitable to lose the holistic structure information of label hierarchy. Global approaches use a single classifier to classify all labels with hierarchical dependencies simultaneously. Early works simplify the HTC task into a flat multi-label classification task, discarding all hierarchical features implicit in taxonomic label set. Later on, specialized hierarchy-aware methods are proposed. HiAGM (Zhou et al., 2020), HTCInfoMax (Deng et al., 2021), HiMatch (Chen et al., 2021), and HiTIN (Zhu et al., 2023) employ dual-encoders framework, which applies the text encoder and structure encoder to learn the representations of texts and labels respectively, and then fuse them to obtain enhanced text embeddings. Current state-of-the-art methods leverage the capabilities of deep learning techniques to improve HTC’s performance, i.e., sequence generative manners (Zhao et al., 2022; Ning et al., 2023; Huang et al., 2022) to mitigate label inconsistency phenomenon, data generation strategies (Wang et al., 2023) to rich text diversity, contrastive learning methods (Wang et al., 2022a; Ji et al., 2023) to enhance semantic expression and prompt-tuning paradigm (Wang et al., 2022b; Ji et al., 2023) to tap into the potential of pre-trained language models (PLMs) (Devlin et al., 2019; Brown et al., 2020; Raffel et al., 2023).

2.2 Prompt Tuning

Prompt Tuning (Schick and Schütze, 2021; Liu et al., 2023a) refers to tuning pre-trained language model by reconstructing downstream task into cloze test which bridges the gap in goals between fine-tuning and pre-training stages. It involves two key steps: (1) template construction which generates a template containing special tokens, and (2) label word verbalizer design which defines a function from token embedding to answer words. There are two types of template construction methods. Hard prompt methods (Shin et al., 2020; Gao et al., 2021a; Han et al., 2022) directly concatenate explicit discrete tokens with the original text and maintain them unchanged throughout entire training process, which do not introduce any parameters. Soft prompt methods (Qin and Eisner, 2021; Lester et al., 2021; Gu et al., 2022; Liu et al., 2023b) convert the template into a group of continuous vectors

as the template and update the vector parameters based on specific contextual semantics and task objectives during training. Prompt-based HPT (Wang et al., 2022b) adopts a soft prompt for HTC tasks, inserting a fixed number of learnable virtual label words to the input text. As for verbalizer, HierVerb (Ji et al., 2023) proposes Multi-verbalizer (Multi-Verb) framework which integrates the hierarchical information, bringing notable performance improvement under few-shot settings.

2.3 Contrastive Learning

Contrastive learning (He et al., 2020; Chen et al., 2020) aims to pull anchor sample close to its positive samples while push it apart from negative samples, which has been proven to elevate the alignment and uniformity of feature space. Contrastive learning has various forms, typically differing in the construction of positive and negative pairs and loss formulas. Under self-supervised settings, positive samples are usually obtained through data augmentations and repeating twice dropout mask (Gao et al., 2021b) operation. Under supervised settings, positives are other samples of the same category (Khosla et al., 2020). Negative samples are regularly selected from other samples or samples of other categories within a batch. Prototypical Contrastive Learning (Li et al., 2021) is proposed to enhance semantic discrimination and balance. Different from instance-level contrastive learning, it encourages instance to be closer to their assigned prototypes. The prototype of a class can set as the label semantics (Ma et al., 2022), the average of all embeddings of the same category (Xiao et al., 2021) and learnable parameters (Cui et al., 2022).

3 Methodology

In this section, we present a detailed description of our proposed DPT model to address HTC tasks. As shown in Figure 2, our model is based on prompt-tuning framework with Multi-verbalizer (Ji et al., 2023). The principal innovations of DPT are twofold, including (1) the implementation of Hierarchy-aware Peer-label Contrastive Learning to obtain rich discriminative features, and (2) the incorporation of Label Hierarchy Self-sensing auxiliary task to enhance encoder’s ability for an in-depth understanding of label hierarchy structure.

3.1 Preliminary

Given a set of inputs $D = \{t_1, t_2, \dots, t_N\}$ where $t_i = \{x_j\}_{j=1}^n$ denotes a text composed of n words,

and a predefined hierarchical label set $\mathcal{Y}$ which is commonly organized as a tree-like taxonomy structure $\mathcal{G}$, the goal of HTC is to select labels for t_i at each layer starting from the root label node of $\mathcal{G}$. Assuming L is the maximum depth of $\mathcal{G}$, the labels $\{y_1, y_2, \dots\}$ of an input text correspond to single or multiple paths of the label tree, each of which typically consists of no more than L continuous individual labels with hierarchical relationship within $\mathcal{G}$.

3.2 Dual Prompt Tuning

For the given text t_i , a prompt template is utilized to wrap the original text to generate a new form of model input. For example, t_i is converted to "[CLS] It was 1 level: [MASK] 2 level: [MASK] ... L level: [MASK]. t_i [SEP]." (Ji et al., 2023). Different from vanilla prompt tuning methods, we utilize a dual prompt template to reserve two types of slot positions, instead of only formulating positive label slots. For instance, a common dual prompt is formulated as follow:

$$T = \{[\text{CLS}] t_i [\text{SEP}] \text{ It belongs to } [\text{MASK}] \dots [\text{MASK}] \text{ rather than } [\text{MASK}] \dots [\text{MASK}] [\text{SEP}]\} \quad (1)$$

The number of [MASK] repetitions of positive or negative slots is equal to the depth of the label hierarchy L . In this paper, we define [MASK] tokens inserted in the prompt template at the position of label slots as "label mask tokens". In the above example, positive label mask tokens locate between "It belongs to" and "rather than", while negative label mask tokens are behind "rather than".

Consistent with other competing methods, we employ BERT (Devlin et al., 2019) as the backbone of our model to encode input texts and obtain all token embeddings:

$$V = \text{BERT}(T(t_i)) \quad (2)$$

Let $\{v_l^p\}_{l=1}^L$ and $\{v_l^n\}_{l=1}^L$ respectively represent the embeddings of l -th positive label mask token and l -th negative label mask token. We inherit HierVerb (Ji et al., 2023) to adopt a depth-oriented Multi-verbalizer framework mapping label token embeddings $\{v_l^p\}_{l=1}^L$ to label words. Probability distribution of t_i can be expressed as:

$$\begin{aligned} Z &= \{z_l\}_{l=1}^L \\ &= \{V_1(v_1^p), \dots, V_L(v_L^p)\} \end{aligned} \quad (3)$$

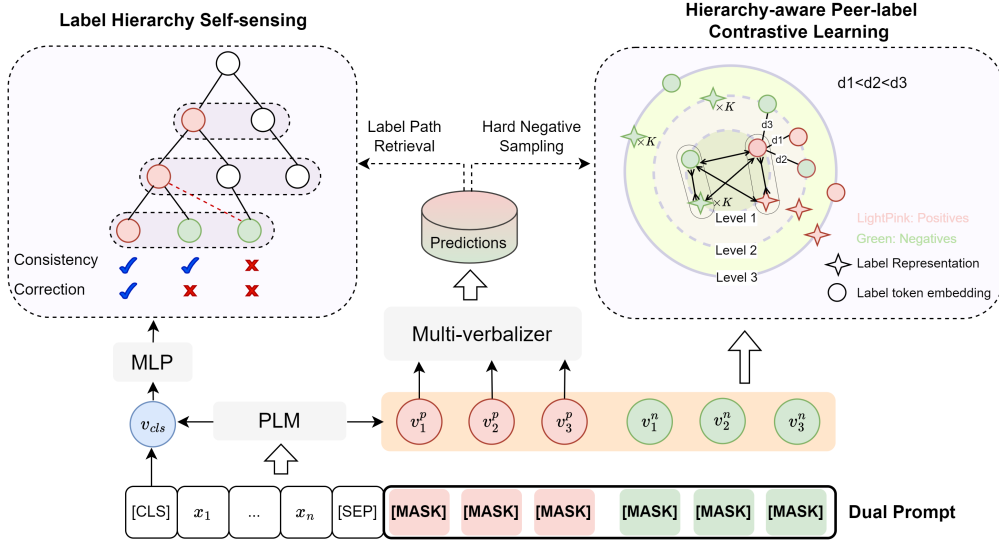


Figure 2: Model architecture of Dual Prompt Tuning (DPT) based on a multi-verbalizer framework. DPT is composed of two modules, including Hierarchy-aware Peer-label Contrastive Learning and Label Hierarchy Self-sensing task. Label token embeddings and label representations are encoded in the unified embedding space, and contrastive loss are calculated according to their affiliation and hierarchical relationship. Label Hierarchy Self-sensing task is used to simultaneously restrain path predictions with correctness and consistency. Note that the figure only depicts the Hierarchy-aware Peer-label Contrastive learning between the first and second levels.

where the l -th verbalizer V_l acts in predicting l -level labels. More details about Multi-verbalizer framework can be found in Ji et al. (2023).

3.3 Hierarchy-aware Peer-label Contrastive Learning

To extract hierarchy-aware discriminative feature on the basis of dual prompt tuning, the ideal embeddings at label mask tokens should satisfy the following intents: (1) Token embeddings of the positive label mask tokens should be close to representations of their positive labels, and far away from negative labels in the feature space. The same desire applies to token embeddings of negative label mask tokens. (2) The semantic similarity between high-level label and its ground truth sub-label is expected to be greater than that with other sub-labels, further greater than sub-labels of other nodes at the same hierarchical layer. Based on the above, we propose a Hierarchy-aware Peer-label Contrastive Learning (HierPCL) method to capture latent semantic relevancy between peer labels as well as parent-child labels.

Objective of HierPCL The objective function of HierPCL consists of three components: positive label contrastive learning, negative label contrastive learning and cross-hierarchical rank loss. The basic idea of HierPCL is to encourage the embeddings of the label mask tokens encoded by PLM

closer to the representations of their positive labels which should be filled in the label slots of the template.

(1) Positive label contrastive learning is performed on positive label mask tokens. The target positives are ground truth labels while the negatives are the sampled K negative labels and negative label mask token. The loss function is formulated as:

$$\mathcal{L}_{CL}^p = -\frac{1}{L} \sum_{l=1}^L \log \frac{\sum_{m=1}^M \exp(s(v_l^p, r_{l,m}^p)/\tau)}{\sum_{u \in \mathcal{A}^p} \exp(s(v_l^p, u)/\tau)} \quad (4)$$

where $r_{l,m}^p$ and $r_{l,k}^n$ respectively denote the representation vectors of m -th positive labels and the k -th negative label in l -th level of sentence t_i , and $s(\cdot)$ represents cosine similarity function. All participants above are denoted as $\mathcal{A}^p := \{\{r_{l,m}^p\}_{m=1}^M, \{r_{l,k}^n\}_{k=1}^K, v_l^p\}$.

(2) Negative label contrastive learning is performed on negative label mask tokens. Opposite to $\mathcal{L}_{CL}^p$, the target positives are negative labels of this instance while the negatives are ground truth labels and positive label mask token. Let $\mathcal{A}^n := \{\{r_{l,m}^p\}_{m=1}^M, \{r_{l,k}^n\}_{k=1}^K, v_l^p\}$, the loss function is formulated as:

$$\mathcal{L}_{CL}^n = -\frac{1}{L} \sum_{l=1}^L \log \frac{\sum_{k=1}^K \exp(s(v_l^n, r_{l,k}^n)/\tau)}{\sum_{u \in \mathcal{A}^n} \exp(s(v_l^n, u)/\tau)} \quad (5)$$

(3) Cross-hierarchical rank loss aims to align label token embeddings with the representations of their sub-labels. In other words, high-level labels tend to have higher semantic similarities with positive child-labels, compared to negative peer child-labels. The loss function is formulated as:

$$\mathcal{L}_R = \sum_{l=1}^{L-1} \left(\sum_{m=1}^M \sum_{k=1}^K \max(0, s(v_l^p, r_{l+1,k}^n) - s(v_l^p, r_{l+1,m}^p)) \right) + \sum_{r_{c,l+1}^n} \sum_{r_{o,l+1}^n} \max(0, s(v_l^p, r_{o,l+1}^n) - s(v_l^p, r_{c,l+1}^n)) \quad (6)$$

where $r_{c,l+1}^n$ denotes the representation vector of $(l+1)$ -th negative label which belongs to child-labels of l -th ground truth label, and $r_{o,l+1}^n$ denotes that doesn't belong to child-labels.

Finally, the objective function of HierPCL is formulated as follow:

$$\mathcal{L}_1 = \alpha \mathcal{L}_{CL}^p + (1 - \alpha) \mathcal{L}_{CL}^n + \beta \mathcal{L}_R \quad (7)$$

where α and β are hyper-parameters used for balancing the relative weights of three components.

Hard Negative Sampling The selection strategy of negative samples is critical to contrastive learning. Since the estimated label predictions can be considered a reliable source for generating hard negatives, we design an adaptive self-produced hard negative sampling strategy under the guidance of predictions during training. It adopts the top K hard negative labels according to confidence scores in descending order output by the Verbalizer at each hierarchical level. In our experiment, K is set to 10% of the number of labels, which achieves a balance between performance and memory consumption. The effects of different settings of K are described in Appendix B for detail.

Hierarchy-injected Label Representation

Label representation vectors are learned through the shared PLM without prior statistics. To avoid the side impact of the overlapping interaction between label names and text words, we assign a unique fabricated symbol for each label, like "L0", and add them to the vocabulary list used in model training. To incorporate hierarchy information to label representation, we flatten the parent-child hierarchy of the label to form a label sequence, as follow:

$$Q = \{[CLS] W [SEP] W^f [SEP] \{W^c\} [SEP]\} \quad (8)$$

where W^f means the parent label symbol of W , $\{W^c\}$ is on behalf of all children labels. "Root"

and "None" are used as fictitious tokens when parent label or child label doesn't exist. We use the embedding on W enriched with hierarchical dependencies as the label representation r .

3.4 Label Hierarchy Self-sensing Task

For the purpose of elevating the model's perceptual ability of label hierarchical structure, we introduce a Label Hierarchy Self-sensing task as an auxiliary task consisting of two sub-tasks: (1) determining whether the label nodes at each level can form a valid path, and (2) determining whether the ground truth path exists in the prediction. A simple base unit of feed-forward module is utilized on the top of $[CLS]$ token after PLM. The consistency and correction loss functions are respectively designed based on Binary Cross Entropy (BCE) (De Boer et al., 2005):

$$\mathcal{L}_{con} = \text{BCE}(\bar{y}_{con}, \bar{p}_{con}) \quad (9)$$

$$\mathcal{L}_{cor} = \text{BCE}(\bar{y}_{cor}, \bar{p}_{cor}) \quad (10)$$

where $\bar{p}_{con}$ represents the probability of that the predictions can form the label paths, and $\bar{p}_{cor}$ represents the probability of that predictions are the ground truth label paths. We retrieve all label paths from the label nodes of the sentence. $\bar{y}_{con} = 1$ if all predicted label nodes can exactly form label paths. $\bar{y}_{cor} = 1$ if all combined label paths are the ground truth labels of the sentence. Otherwise, the value of $\bar{y}_{con}$ or $\bar{y}_{cor}$ is 0.

Finally, the loss function of the auxiliary task is formulated as:

$$\mathcal{L}_2 = \mathcal{L}_{con} + \mathcal{L}_{cor} \quad (11)$$

3.5 Multi-task Training

Overall, multi-task training objective is to minimize the weighted combination of classification loss, Peer-label contrastive learning loss, label hierarchy self-sensing loss and MLM loss retaining from the original BERT pre-training. The choice of classification loss function can be contingent upon given circumstances. For the sake of universality, the standard Binary Cross-Entropy function is employed:

$$\mathcal{L}_{CLS} = \sum_{l=1}^L \text{BCE}(y_l, z_l) \quad (12)$$

Final joint loss can be formulated as:

$$\mathcal{L} = \mathcal{L}_{MLM} + \mathcal{L}_{CLS} + \lambda_1 \mathcal{L}_1 + \lambda_2 \mathcal{L}_2 \quad (13)$$

where λ_1 and λ_2 are hyper-parameters.

Dataset	WoS	RCV1	BGC	NYT
L	2	4	4	8
$ \mathcal{Y} $	141	103	146	166
$\text{Avg}(\mathcal{Y}_i)$	2.0	3.24	3.01	7.6
#Train	30070	20833	58715	23345
#Val	7518	2316	14785	5834
#Test	9397	781265	18394	7292

Table 1: Statistics of HTC datasets. L , $|\mathcal{Y}|$ and $\text{Avg}(|\mathcal{Y}_i|)$ represent the maximum depth, total number of categories and the average number of labels respectively.

4 Experiment

4.1 Experiment Setup

Datasets We conduct experiments on 4 benchmark datasets: Web-of-Science (WoS) (Kowsari et al., 2017), NYTimes (NYT) (Sandhaus, 2008), RCV1-V2 (Lewis et al., 2004) and Blurb Genre Collection (BGC) ¹ (Aly et al., 2019). Note that taxonomic of WoS is single-path while the rest three datasets are for multi-path HTC. The detailed statistics of these datasets are list in Table 1.

Evaluation Metrics The performance of our model is evaluated by popular Micro-F1 and Macro-F1 metrics, together with path-constrained metrics proposed by Yu et al. (2022), including C-MicroF1 and C-MacroF1, in which a prediction is considered as correct only when all its ancestor nodes are predicted accurately. The hierarchical path consistency is taken into account as well as the accuracy of label nodes to the comprehensive metrics of HTC.

Implementation Details The backbone of DPT is initialized with bert-base-uncased². The batch size is set to 32 for BGC and 16 for other datasets. The AdamW optimizer is used with the learning rate of 2e-5 for WoS and 3e-5 for others. We apply the early stopping strategy with 5 patient epochs. For fair comparison, we perform the same data processing and splitting method as HPT. The reported results of our main experiments are the average score of 5 runs over different random seeds. Experimental settings of all hyper-parameters are described in Appendix A.

Baselines We compare our methods with following advanced HTC methods:

- **HGCLR** (Wang et al., 2022a) incorporates label hierarchy into text encoder through

hierarchy-guided contrastive learning between text and its generated positive samples with the most closest label paths.

- **Seq2Tree** (Yu et al., 2022) and **PAAM-HiA-T5** (Huang et al., 2022) treat HTC as the sequence generation task. Seq2Tree designs a constrained decoding strategy with dynamic vocabulary to ensure label consistency. PAAM-HiA-T5 proposes a multi-level sequential label generative T5 model with a path-adaptive attention mechanism to focus on label dependency prediction.
- **HPT** (Wang et al., 2022b) exploits the effects of prompt-tuning by a dynamic virtual template and a zero-bounded multi-label cross entropy loss, which achieves the state-of-the-art performances on most of datasets.
- **HiTIN** (Zhu et al., 2023) introduces the structural entropy to construct a coding tree for the label hierarchy and then build a novel structure Encoder to enhance text representations.

4.2 Main Result

Experimental results are shown in Tabel 2. Our model consistently outperforms previous advanced approaches across 3 datasets except for WoS. On WoS dataset with label depth of 2, our proposed DPT achieves comparable results with HPT but decreased performance compared to PAAM-HiA-T5 model. Our model establishes state-of-the-art results on RCV1-V2 and BGC datasets. It improves 0.5% and 0.76% absolute Micro-F1 and Macro-F1 on RCV1-V2 dataset comparing to the current best results. The performance boost of Micro-F1 and Macro-F1 on BGC reaches 0.53% and 1.52% over the SoTA HPT model. The notable advancements on Macro-F1 indicate that our model performs well on sparse labels. On NYT dataset, our model surpasses them on Micro-F1 by 0.14% but slightly lower than HPT on Macro-F1.

Without introducing any additional network parameters to extract the semantics of labels and their hierarchical structures, our model surprisingly outperforms previous methods of encoding label names and label hierarchies using GNNs. Compared to HGCLR which use instance-level contrastive learning with complex positive sample generation operation, DPT makes impressive performances progress over all datasets.

¹<https://www.inf.uni-hamburg.de/en/inst/ab/lt/resources/data/blurb-genre-collection.html>

²<https://huggingface.co/bert-base-uncased>

Model	WoS		RCV1-V2		BGC		NYT	
	Micro-F1	Macro-F1	Micro-F1	Macro-F1	Micro-F1	Macro-F1	Micro-F1	Macro-F1
BERT(Wang et al., 2022a)	85.63	79.07	85.65	67.02	-	-	78.24	66.08
BERT+HiAGM(Wang et al., 2022a)	86.04	80.19	85.58	67.93	-	-	78.64	66.76
BERT+HTCInfoMax(Wang et al., 2022a)	86.30	79.97	85.53	67.09	-	-	78.75	67.31
BERT+HiMatch(Chen et al., 2021)	86.70	81.06	86.33	68.66	78.89	63.19	76.79	63.89
HGCLR(Wang et al., 2022a)	87.11	81.20	86.49	68.31	-	-	78.86	67.96
Seq2Tree(Yu et al., 2022)	87.20	82.50	86.88	70.01	79.72	63.96	-	-
PAAM-HiA-T5(Huang et al., 2022)	90.36	81.64	87.22	70.02	-	-	77.52	65.97
HPT(Wang et al., 2022b)	87.16	81.93	87.26	69.53	81.32 [‡]	66.69 [‡]	80.42	70.42
HiTIN(Zhu et al., 2023)	87.19	81.57	86.71	69.95	-	-	79.65	69.31
DPT (Ours)	87.25	81.51	87.76	70.78	81.85	68.21	80.56	70.28

Table 2: Experimental results on four HTC datasets. The best results are in bold format. The result of BERT+HiMatch on NYT dataset is reported by Huang et al. (2022). The result of BERT+HiMatch on BGC is reported by Yu et al. (2022). [‡] means the results are reproduced upon the release project by ourselves.

Model	RCV1-V2		BGC		NYT	
	C-MiF1	C-MaF1	C-MiF1	C-MaF1	C-MiF1	C-MaF1
HPT	86.80	68.71	80.88	65.36	79.33	68.80
DPT (Ours)	87.47	70.20	81.43	66.97	79.77	68.70

Table 3: Evaluation results of label consistency. MiF1 and MaF1 are abbreviations for Micro-F1 and Macro-F1, respectively.

4.3 Results on Label Consistency

For HTC tasks, cross-layer label consistency is also an important metric, which signifies the fact that each layer label predicted by the model should conform to the hierarchical relationship. Table 3 illustrates the label consistency performance of the proposed DPT and the SoTA model HPT. Our model improves consistency of label hierarchy on RCV1-V2 and BGC by a large margin, respectively exceeding HPT by 1.49% and 1.61% on C-MacroF1 metric. Although our model focuses more on the interaction between peer labels at each layer, the knowledge of label hierarchy has also been internalized. The accuracy of both individual label nodes and label paths has been improved, indicating that our methods are reasonable and efficient.

4.4 Results on Imbalanced Hierarchy

To further clarify the superiority of our methods, we intent to explore the performance on imbalanced hierarchy. Following long-tailed learning setting, we sort the test set in descending order based on the quantity of class instances and evenly cluster the dataset into head, middle, and tail groups. The visualization Macro-F1 results are shown in Figure 3³. It's evident that DPT comprehensively outperforms

HPT in RCV1-V2 and BGC datasets and shows significant improvements on tail classes with few training samples, demonstrating the effectiveness of our methods in eliminating the impact resulting from imbalanced distribution.

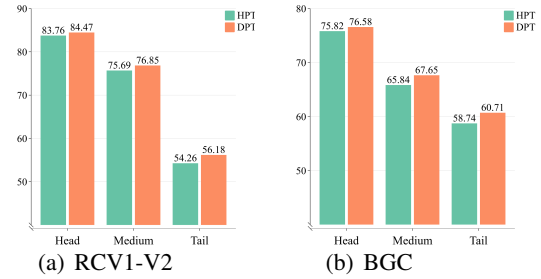


Figure 3: Macro-F1 score on head, medium and tail class groups

4.5 Ablation Study

To investigate the effects of each component of our proposed model, we implement different variants on BGC and RCV1-V2 dataset, and results are shown in Table 4 and Tabel 5 respectively. Upon only employing the HierPCL module, the performances in all metrics realize considerable enhancement and are superior to current state-of-the-art models, confirming its significant effectiveness and reliability. By removing negative contrastive part, the scores undergo sharp declines, which demonstrates that negative label contrastive learning plays a prominent role in HierPCL. As a strong contrast, we replace our self-produced hard negative sampling with random sampling, resulting in marked decrease in metrics, which validates the advantages of our negative sampling strategy. Cross-hierarchical Rank Loss in HierPCL and label

³We drew this picture in the website <https://www.chiplot.online/>.

hierarchy self-sensing auxiliary task also improve gains of our methods, further boosting the model performance especially in terms of Macro-F1 metric.

Ablation Models	BGC			
	MiF1	MaF1	C-MiF1	C-MaF1
Multi-Verb(Baseline)	81.38	66.74	80.92	65.58
DPT(Ours)	81.85	68.21	81.43	66.97
r.m. $\mathcal{L}_2$	81.71	68.09	81.27	66.56
r.m. $\mathcal{L}_{LC}^n$	81.62	67.35	81.35	66.11
r.m. $\mathcal{L}_R$	81.83	67.99	81.40	66.50
r.m. $\mathcal{L}_R$ & $\mathcal{L}_{LC}^n$	81.36	67.41	80.97	66.26
r.p. Random Sampling	81.48	67.46	81.03	66.59

Table 4: Ablation study on BGC dataset.

Ablation Models	RCV1-V2			
	MiF1	MaF1	C-MiF1	C-MaF1
Multi-Verb(Baseline)	87.19	69.16	86.68	68.31
DPT(Ours)	87.76	70.78	87.47	70.20
r.m. $\mathcal{L}_2$	87.34	70.28	87.01	69.45
r.m. $\mathcal{L}_{LC}^n$	87.14	70.05	86.61	69.32
r.m. $\mathcal{L}_R$	87.47	70.35	87.11	69.52
r.m. $\mathcal{L}_R$ & $\mathcal{L}_{LC}^n$	87.22	69.52	86.79	68.73
r.p. Random Sampling	87.26	69.55	86.94	68.93

Table 5: Ablation study on RCV1-V2 dataset.

4.6 Insight into case effects

In order to gain insight into the practical effects of our model, we conduct detailed case studies on the test set. We define 3 types of label errors from the perspective of multi-label classification, as follows:

- Misjudged, which means the label is mistakenly identified by the model as one of the ground truth labels of the instance while the instance actually belongs to another labels.
- Excess, which means the label is unnecessarily identifies as one label for the instances.
- Missing, which means the ground truth label which the model has failed to recall.

We separately calculate the distribution of label error types for baseline model (prompt-tuning with Multi-Verb framework) and our improved model on the test set. As shown in Figure 4, we find that optimization effects of our model in cases are manifested in recalling missing labels, correcting misjudged labels and removing excess labels, respectively accounting for 44.69%, 40.21% and 15.10% of the proportion. It demonstrates the strong power

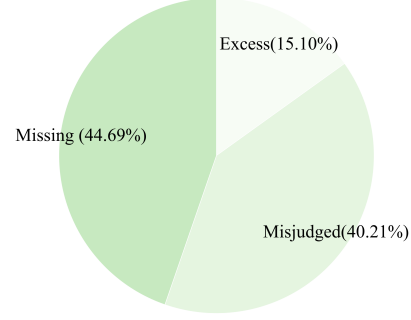


Figure 4: Proportion of error types corrected by our method.

of DPT to capture discrimination representation and then relieve label confusion. Some specific cases are illustrated in Appendix C.

5 Conclusion

In this paper, we present a novel Dual Prompt Tuning method for HTC tasks. Firstly, we propose a Hierarchy-aware Peer-label Contrastive Learning approach to alleviate confusion between peer labels. An original dual prompt template is created with slots for both positive and negative label, on which the contrastive learning is performed at each layer. Secondly, to further strengthen knowledge of label hierarchy structure, we design a Label Hierarchy Self-sensing auxiliary task to identify consistency and correctness of model predictions. Experimental results illustrate that our proposed DPT model achieves significant improvements on popular HTC datasets. Particularly, DPT exhibits outstanding efficacy in preserving label path consistency and addressing imbalanced hierarchy challenge. It excels in the accurate recognition of negative labels and contributes to obtaining hierarchy-aware discriminative features.

6 Limitations

In our work, hierarchical labels serve as supplementary instances. Both positive labels and the sampled K negative labels are input into PLM to calculate the representation vector, which brings additional memory consumption during model training. Besides, There is still room for improvement in hand-craft prompt design. It's worth exploring in depth. Combined with large language models to enhance discriminative ability for HTC is one of the development directions of our future work.

References

- Rami Aly, Steffen Remus, and Chris Biemann. 2019. [Hierarchical multi-label classification of text with capsule networks](#). In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics: Student Research Workshop*, pages 323–330, Florence, Italy. Association for Computational Linguistics.
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel Ziegler, Jeffrey Wu, Clemens Winter, Chris Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. 2020. [Language models are few-shot learners](#). In *Advances in Neural Information Processing Systems*, volume 33, pages 1877–1901. Curran Associates, Inc.
- Ali Cevahir and Koji Murakami. 2016. [Large-scale multi-class and hierarchical product categorization for an E-commerce giant](#). In *Proceedings of COLING 2016, the 26th International Conference on Computational Linguistics: Technical Papers*, pages 525–535, Osaka, Japan. The COLING 2016 Organizing Committee.
- Haibin Chen, Qianli Ma, Zhenxi Lin, and Jiangyue Yan. 2021. [Hierarchy-aware label semantics matching network for hierarchical text classification](#). In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 4370–4379, Online. Association for Computational Linguistics.
- Ting Chen, Simon Kornblith, Mohammad Norouzi, and Geoffrey Hinton. 2020. [A simple framework for contrastive learning of visual representations](#). In *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pages 1597–1607. PMLR.
- Ganqu Cui, Shengding Hu, Ning Ding, Longtao Huang, and Zhiyuan Liu. 2022. [Prototypical verbalizer for prompt-based few-shot tuning](#). In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 7014–7024, Dublin, Ireland. Association for Computational Linguistics.
- Pieter-Tjerk De Boer, Dirk P Kroese, Shie Mannor, and Reuven Y Rubinfeld. 2005. A tutorial on the cross-entropy method. *Annals of operations research*, 134:19–67.
- Zhongfen Deng, Hao Peng, Dongxiao He, Jianxin Li, and Philip Yu. 2021. [HTCInfoMax: A global model for hierarchical text classification via information maximization](#). In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 3259–3265, Online. Association for Computational Linguistics.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. [BERT: Pre-training of deep bidirectional transformers for language understanding](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis, Minnesota. Association for Computational Linguistics.
- Tianyu Gao, Adam Fisch, and Danqi Chen. 2021a. [Making pre-trained language models better few-shot learners](#). In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 3816–3830, Online. Association for Computational Linguistics.
- Tianyu Gao, Xingcheng Yao, and Danqi Chen. 2021b. [SimCSE: Simple contrastive learning of sentence embeddings](#). In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 6894–6910, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Yuxian Gu, Xu Han, Zhiyuan Liu, and Minlie Huang. 2022. [PPT: Pre-trained prompt tuning for few-shot learning](#). In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 8410–8423, Dublin, Ireland. Association for Computational Linguistics.
- Xu Han, Weilin Zhao, Ning Ding, Zhiyuan Liu, and Maosong Sun. 2022. [Ptr: Prompt tuning with rules for text classification](#). *AI Open*, 3:182–192.
- Kaiming He, Haoqi Fan, Yuxin Wu, Saining Xie, and Ross Girshick. 2020. Momentum contrast for unsupervised visual representation learning. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 9729–9738.
- Wei Huang, Chen Liu, Bo Xiao, Yihua Zhao, Zhaoming Pan, Zhimin Zhang, Xinyun Yang, and Guiquan Liu. 2022. [Exploring label hierarchy in a generative way for hierarchical text classification](#). In *Proceedings of the 29th International Conference on Computational Linguistics*, pages 1116–1127, Gyeongju, Republic of Korea. International Committee on Computational Linguistics.
- Ke Ji, Yixin Lian, Jingsheng Gao, and Baoyuan Wang. 2023. [Hierarchical verbalizer for few-shot hierarchical text classification](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 2918–2933, Toronto, Canada. Association for Computational Linguistics.

703	Prannay Khosla, Piotr Teterwak, Chen Wang, Aaron	<i>American Chapter of the Association for Computa-</i>	759
704	Sarna, Yonglong Tian, Phillip Isola, Aaron	<i>tional Linguistics: Human Language Technologies,</i>	760
705	Maschinot, Ce Liu, and Dilip Krishnan. 2020. Su-	pages 5203–5212, Online. Association for Computa-	761
706	perervised contrastive learning. <i>Advances in neural</i>	<i>tional Linguistics.</i>	762
707	<i>information processing systems</i> , 33:18661–18673.		
708	Kamran Kowsari, Donald E. Brown, Mojtaba Hei-	Colin Raffel, Noam Shazeer, Adam Roberts, Katherine	763
709	darysafa, Kiana Jafari Meimandi, Matthew S. Gerber,	Lee, Sharan Narang, Michael Matena, Yanqi Zhou,	764
710	and Laura E. Barnes. 2017. Hdltext: Hierarchical	Wei Li, and Peter J. Liu. 2023. Exploring the limits	765
711	deep learning for text classification . In <i>2017 16th</i>	of transfer learning with a unified text-to-text trans-	766
712	<i>IEEE International Conference on Machine Learn-</i>	former .	767
713	<i>ing and Applications (ICMLA)</i> , pages 364–371.		
714	Brian Lester, Rami Al-Rfou, and Noah Constant. 2021.	Evan. Sandhaus. 2008. The new york times annotated	768
715	The power of scale for parameter-efficient prompt	corpus ldc2008t19 . Philadelphia: Linguistic Data	769
716	tuning . In <i>Proceedings of the 2021 Conference on</i>	<i>Consortium.</i>	770
717	<i>Empirical Methods in Natural Language Processing</i> ,		
718	pages 3045–3059, Online and Punta Cana, Domini-	Timo Schick and Hinrich Schütze. 2021. Exploiting	771
719	can Republic. Association for Computational Lin-	cloze-questions for few-shot text classification and	772
720	guistics.	natural language inference . In <i>Proceedings of the</i>	773
721	David D. Lewis, Yiming Yang, Tony G. Rose, and Fan	<i>16th Conference of the European Chapter of the Asso-</i>	774
722	Li. 2004. Rcv1: A new benchmark collection for text	<i>ciation for Computational Linguistics: Main Volume,</i>	775
723	categorization research. 5:361–397.	pages 255–269, Online. Association for Computa-	776
724	Junnan Li, Pan Zhou, Caiming Xiong, and Steven Hoi.	<i>tional Linguistics.</i>	777
725	2021. Prototypical contrastive learning of unsuper-		
726	vised representations . In <i>International Conference</i>	Taylor Shin, Yasaman Razeghi, Robert L. Logan IV, Eric	778
727	<i>on Learning Representations</i> .	Wallace, and Sameer Singh. 2020. AutoPrompt: Elic-	779
728	Pengfei Liu, Weizhe Yuan, Jinlan Fu, Zhengbao Jiang,	iting Knowledge from Language Models with Auto-	780
729	Hiroaki Hayashi, and Graham Neubig. 2023a. Pre-	matically Generated Prompts . In <i>Proceedings of the</i>	781
730	train, prompt, and predict: A systematic survey of	<i>2020 Conference on Empirical Methods in Natural</i>	782
731	prompting methods in natural language processing.	<i>Language Processing (EMNLP)</i> , pages 4222–4235,	783
732	<i>ACM Computing Surveys</i> , 55(9):1–35.	Online. Association for Computational Linguistics.	784
733	Xiao Liu, Yanan Zheng, Zhengxiao Du, Ming Ding,	Junru Song, Feifei Wang, and Yang Yang. 2023. Peer-	785
734	Yujie Qian, Zhilin Yang, and Jie Tang. 2023b. Gpt	label assisted hierarchical text classification . In <i>Pro-</i>	786
735	understands, too . <i>AI Open</i> .	<i>ceedings of the 61st Annual Meeting of the Associa-</i>	787
736	Jie Ma, Miguel Ballesteros, Srikanth Doss, Rishita	<i>tion for Computational Linguistics (Volume 1: Long</i>	788
737	Anubhai, Sunil Mallya, Yaser Al-Onaizan, and Dan	<i>Papers)</i> , pages 3747–3758, Toronto, Canada. Associ-	789
738	Roth. 2022. Label semantics for few shot named	<i>ation for Computational Linguistics.</i>	790
739	entity recognition . In <i>Findings of the Association for</i>		
740	<i>Computational Linguistics: ACL 2022</i> , pages 1956–	Yue Wang, Dan Qiao, Juntao Li, Jinxiong Chang,	791
741	1971, Dublin, Ireland. Association for Computational	Qishen Zhang, Zhongyi Liu, Guannan Zhang, and	792
742	Linguistics.	Min Zhang. 2023. Towards better hierarchical text	793
743	Yuning Mao, Jingjing Tian, Jiawei Han, and Xiang Ren.	classification with data generation . In <i>Findings of</i>	794
744	2019. Hierarchical text classification with reinforced	<i>the Association for Computational Linguistics: ACL</i>	795
745	label assignment . In <i>Proceedings of the 2019 Confer-</i>	<i>2023</i> , pages 7722–7739, Toronto, Canada. Associa-	796
746	<i>ence on Empirical Methods in Natural Language Pro-</i>	<i>tion for Computational Linguistics.</i>	797
747	<i>cessing and the 9th International Joint Conference</i>		
748	<i>on Natural Language Processing (EMNLP-IJCNLP)</i> ,	Zihan Wang, Peiyi Wang, Lianzhe Huang, Xin Sun,	798
749	pages 445–455, Hong Kong, China. Association for	and Houfeng Wang. 2022a. Incorporating hierarchy	799
750	Computational Linguistics.	into text encoder: a contrastive learning approach	800
751	Bo Ning, Deji Zhao, Xinjian Zhang, Chao Wang, and	for hierarchical text classification . In <i>Proceedings</i>	801
752	Shuangyong Song. 2023. Ump-mg: A uni-directed	<i>of the 60th Annual Meeting of the Association for</i>	802
753	message-passing multi-label generation model for	<i>Computational Linguistics (Volume 1: Long Papers)</i> ,	803
754	hierarchical text classification. <i>Data Science and</i>	pages 7109–7119, Dublin, Ireland. Association for	804
755	<i>Engineering</i> , 8(2):112–123.	<i>Computational Linguistics.</i>	805
756	Guanghui Qin and Jason Eisner. 2021. Learning how	Zihan Wang, Peiyi Wang, Tianyu Liu, Binghuai Lin,	806
757	to ask: Querying LMs with mixtures of soft prompts .	Yunbo Cao, Zhifang Sui, and Houfeng Wang. 2022b.	807
758	In <i>Proceedings of the 2021 Conference of the North</i>	HPT: Hierarchy-aware prompt tuning for hierarchical	808
		text classification . In <i>Proceedings of the 2022 Con-</i>	809
		<i>ference on Empirical Methods in Natural Language</i>	810
		<i>Processing</i> , pages 3740–3751, Abu Dhabi, United	811
		Arab Emirates. Association for Computational Lin-	812
		guistics.	813
		Lin Xiao, Xiangliang Zhang, Liping Jing, Chi Huang,	814
		and Mingyang Song. 2021. Does head label help for	815

long-tailed multi-label text classification. *Proceedings of the AAAI Conference on Artificial Intelligence*, 35:14103–14111.

Chao Yu, Yi Shen, and Yue Mao. 2022. [Constrained sequence-to-tree generation for hierarchical text classification](#). In *Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval*, pages 1865–1869.

Alessandro Zangari, Matteo Marcuzzo, Michele Schiavinato, Matteo Rizzo, Andrea Gasparetto, Andrea Albarelli, et al. 2023. Hierarchical text classification: a review of current research. *EXPERT SYSTEMS WITH APPLICATIONS*, 224.

Deji Zhao, Bo Ning, Shuangyong Song, Chao Wang, Xiangyan Chen, Xiaoguang Yu, and Bo Zou. 2022. [Tosa: A top-down tree structure awareness model for hierarchical text classification](#). In *Web and Big Data: 6th International Joint Conference, APWeb-WAIM 2022, Nanjing, China, November 25–27, 2022, Proceedings, Part II*, page 23–37, Berlin, Heidelberg. Springer-Verlag.

Jie Zhou, Chunping Ma, Dingkun Long, Guangwei Xu, Ning Ding, Haoyu Zhang, Pengjun Xie, and Gongshen Liu. 2020. [Hierarchy-aware global model for hierarchical text classification](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 1106–1117, Online. Association for Computational Linguistics.

He Zhu, Chong Zhang, Junjie Huang, Junran Wu, and Ke Xu. 2023. [HiTIN: Hierarchy-aware tree isomorphism network for hierarchical text classification](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 7809–7821, Toronto, Canada. Association for Computational Linguistics.

A Hyper-parameters Setting

We list the hyper-parameter settings of all datasets in Table 6 for reproducibility.

Params	WoS	RCV1	BGC	NYT
α	0.2	0.6	0.6	0.2
β	0.1	0.2	0.1	0.1
λ_1	0.5	1.0	0.5	0.1
λ_2	0.6	0.5	0.2	0.2

Table 6: Hyper-parameter settings.

B Performance of Different Negative Sampling Ratio

To explain the rationality of selecting 10% negative labels from label sets for HierPCL, we compare the effects of negative sampling ratio of 10% and 100%. From Table 7 and Table 4, it’s obvious that

10% negative labels are sufficient, retaining the vast majority of accuracy on RCV1-V2 and even surpassing the performance of using all negatives on BGC dataset, which indicates that blindly increasing the number of negative samples is not always effective. We will explore the impact of positive and negative label ratios in future work.

ratio	MiF1	MaF1	C-MiF1	C-MaF1
0.1	87.76	70.78	87.47	70.20
1	87.81	70.01	87.50	69.24

Table 7: Results of different negative sampling ratio on RCV1-V2 dataset.

ratio	MiF1	MaF1	C-MiF1	C-MaF1
0.1	81.85	68.21	81.43	66.97
1	81.75	67.95	81.41	66.62

Table 8: Results of different negative sampling ratio on BGC dataset.

C Case Study

DPT performs Peer-label Contrastive Learning at each level, which enhances the model’s representation and discrimination abilities. Introduction of cross-hierarchical rank loss and label hierarchy self-sensing auxiliary task improve label consistency. To look into the practical effects, we conduct adequate case studies on the BGC dataset. Compared to the baseline model, main improvements of DPT are reflected in recalling missing labels, correcting misjudged labels, removing excess labels, and further correcting label inconsistencies. Table 9 provides some examples. Note that labels output by DPT model in the Table 9 are consistent with the ground truth labels.

DPT (Ours)	Teen & Young Adult-Teen & Young Adult Mystery & Suspense, Teen & Young Adult-Teen & Young Adult Fiction, Teen & Young Adult- Teen & Young Adult Social Issues
Multi-Verb (Baseline)	Teen & Young Adult-Teen & Young Adult Mystery & Suspense, Teen & Young Adult-Teen & Young Adult Fiction
Text	Secret, Silent Screams: For fans of Gillian Flynn, Caroline Cooney, and R.L. Stine comes Secret, Silent Screams from four-time Edgar Allen Poe Young Adult Mystery Award winner Joan Lowery Nixon. Is Barry's death the latest tragedy in a string of suicides at Farrington Park High School? Or is it murder? Marti is sure her friend Barry didn't take his own life, but no one will believe her except Police Officer Prescott. But opening an investigation takes time, and Marti is determined to find her friend's killer soon. Because even now he could be planning his next crime... "Enthralling suspense... satisfying[...]... [and an] intricate plot." —Publishers WeeklyFrom the Paperback edition.
DPT (Ours)	Fiction-Romance-Contemporary Romance, Fiction-Romance- Suspense Romance , Fiction- Women's Fiction
Multi-Verb (baseline)	Fiction-Romance-Contemporary Romance
Text	Summer in Eclipse Bay: A special message from Jayne Ann KrentzDear Reader:Summer has arrived in Eclipse Bay and things are definitely heating up between the Hartes and the Madisons. It seems that the mysterious new gallery owner, Octavia Brightwell, is thinking about having a scandalous fling with that rogue Nick Harte before she leaves town. As far as Nick is concerned, a short-term affair sounds perfect. But it isn't going to be easy.One big obstacle is Mitchell Madison. For reasons of his own, Mitchell has taken it upon himself to play guardian to Octavia. He's made it clear that if Nick fools around with her, there will be a price to pay. And then there's Nick's young son, Carson, who has his own agenda where Octavia is concerned. He doesn't want his father messing up his plans.Summer in Eclipse Bay is going to be eventful this year. Some long-buried secrets from the infamous Harte-Madison feud are about to surface. The past and the present are on a collision course. I hope you'll join me to watch the fireworks.Happy reading . . . Jayne Ann Krentz.
DPT (Ours)	Fiction- Graphic Novels & Manga
Multi-Verb (Baseline)	Fiction- Mystery & Suspense
Text	Corto Maltese: Beyond The Windy Isles: The second volume in the definitive English language edition of Hugo Pratt's masterpiece, Corto Maltese, presented in the original oversized B&W format and with new translations made from Pratt's original Italian scripts. "Mushroom Heads" begins in Maracaibo, Venezuela, where Corto Maltese and Professor Steiner lead an expedition on the trail of the legendary El Dorado, financed by the antiquarian Levi Colombia. In "Banana Conga," Corto has his first and nearly fatal encounter with the beautiful yet dangerous mercenary Venexiana Stevenson. Within this framework of adventure, Hugo Pratt weaves themes dealing with the exploitation of indigenous people, the noble struggle to gain freedom and independence, and how cowardice can poison men of all classes. The action, set in 1917, takes Corto Maltese from the Mosquito Coast to Barbados to a deadly struggle among Jivaro head-hunters in the Peruvian Amazon
DPTL (Ours)	Fiction- Women's Fiction
Multi-Verb (Baseline)	Fiction- Romance
Text	Nappily Ever After: SOON TO BE A NETFLIX ORIGINAL FILM STARRING SANAA LATHANWhat happens when you toss tradition out the window and really start living for yourself? Venus Johnston has a great job, a beautiful home, and a loving live-in boyfriend named Clint, who happens to be a drop-dead gorgeous doctor. She also has a weekly beauty-parlor date with Tina, who keeps Venus's long, processed hair slick and straight. But when Clint—who's been reluctant to commit over the past four years—brings home a puppy instead of an engagement ring, Venus decides to give it all up. She trades in her long hair for a dramatically short, natural cut and sends Clint packing. It's a bold declaration of independence—one that has effects she never could have imagined. Reactions from friends and coworkers range from concern to contempt to outright condemnation. And when Clint moves on and starts dating a voluptuous, long-haired beauty, Venus is forced to question what she really wants out of life. With wit, resilience, and a lot of determination, she finally learns what true happiness is—on her own terms. Told with style, savvy, and humor, Nappily Ever After is a novel that marks the debut of a fresh new voice in fiction.
DPT (Ours)	Nonfiction-Religion & Philosophy-Philosophy
Multi-Verb (Baseline)	Nonfiction-Religion & Philosophy-Philosophy, Nonfiction-Religion & Philosophy- Religion
Text	Malice: Despite our tendencies to separate the mind and body, good and evil, Flahault argues that both stem from the same source within us. This knot, inherent to the human condition, is the tension between our desire for absolute self-affirmation and the fact that each of us can only exist through mediation by others. The dependence on others weighs heavy on our shoulders, hampering our very existence.Malice, then, is not merely a result of our biological constitution, but is also a response to our feelings. These can often resemble those of Milton's and Shelley's monsters, stories the author calls upon to understand features of the nature of evil that reason alone cannot grasp.From the Preface: 'By combining several disciplines—philosophy, anthropology and literary criticism, as well as psychoanalysis—Flahault scrutinizes the origin of malevolence and reveals that, contrary to the view presented by moral philosophy, it is within us that the roots of wickedness are to be found... Taking issue with the widely accepted view that monotheism constitutes moral progress, he argues that by instigating a dualism between good and evil, monotheism has in fact foreclosed the possibility of acknowledging the ambivalence of our fascination with the limitless and infinity.' Chantal Mouffe.
DPT (Ours)	Humor
Multi-Verb (Baseline)	Humor- Graphic Novels & Manga (Label inconsistency)
Text	Garfield Caution: Wide Load: Indulge the Bulge Garfield believes that a full belly is a happy belly—and he intends to keep his stomach ecstatic. Fans of the fat cat will gleefully fill up on this latest smorgasbord of fun!

Table 9: Case Studies on BGC dataset.