

BCE3S: Binary Cross-Entropy Based Tripartite Synergistic Learning for Long-tailed Recognition

Weijia Fan^{1, 2, 3, 4}, Qiufu Li^{2, 3*}, Jiajun Wen^{1, 3}, Xiaoyang Peng⁵

¹College of Computer Science and Software Engineering, Shenzhen University

²School of Artificial Intelligence, Shenzhen University

³Guangdong Provincial Key Laboratory of Intelligent Information Processing, Shenzhen University

⁴CV:HCI Lab, Karlsruhe Institute of Technology

⁵School of Software Engineering, Sun Yat-sen University

weijia.fan@partner.kit.edu, liquifu@szu.edu.cn, wenjiajun@szu.edu.cn, pengxy77@mail2.sysu.edu.cn

Abstract

For long-tailed recognition (LTR) tasks, high intra-class compactness and inter-class separability in both head and tail classes, as well as balanced separability among all the classifier vectors, are preferred. The existing LTR methods based on cross-entropy (CE) loss not only struggle to learn features with desirable properties but also couple imbalanced classifier vectors in the denominator of its Softmax, amplifying the imbalance effects in LTR. In this paper, for the LTR, we propose a binary cross-entropy (BCE)-based tripartite synergistic learning, termed BCE3S, which consists of three components: (1) BCE-based joint learning optimizes both the classifier and sample features, which achieves better compactness and separability among features than the CE-based joint learning, by decoupling the metrics between feature and the imbalanced classifier vectors in multiple Sigmoid; (2) BCE-based contrastive learning further improves the intra-class compactness of features; (3) BCE-based uniform learning balances the separability among classifier vectors and interactively enhances the feature properties by combining with the joint learning. The extensive experiments show that the LTR model trained by BCE3S not only achieves higher compactness and separability among sample features, but also balances the classifier’s separability, achieving SOTA performance on various long-tailed datasets such as CIFAR10-LT, CIFAR100-LT, ImageNet-LT, and iNaturalist2018.

Code — <https://github.com/wakinghours-github/BCE3S>

Introduction

Though deep models trained on balanced datasets can surpass humans in visual recognition, their performance is still unsatisfactory on the imbalanced datasets. In the long-tailed datasets, the sample number decline significantly from the head to tail classes, which are more aligned with the data distribution in the real world, and the training of deep models on such datasets is easily dominated by the head classes, leading to imbalanced features and classifiers from the head to the tail classes, which consequently decreases the overall performance of the final models (Kang et al. 2020; Alshammari et al. 2022; Xuan and Zhang 2024).

*Corresponding author.

Mode	Methods & accuracy on CIFAR100-LT with IF = 100	Formula
JL : CL : UL	<i>CIFAR10</i> CB Loss 39.60 <i>ICLR21</i> Logit Adj 50.50 <i>ICLR23</i> RIDE&H2T 51.38	CE or Focal
✓	<i>NeurIPS20</i> SSP 43.40 <i>ICLR21</i> RIDE 48.60 <i>NeurIPS24</i> DiffuLT 50.70 <i>CVPR24</i> DeiT-LT 55.60	CE
✓ ✓	<i>ICCV21</i> PaCo 43.80 <i>ICCV21</i> PaCo 52.00 <i>TPAMI24</i> ProCo 52.80 <i>CVPR23</i> GLMC 58.41	CE or MSE
✓ ✓ ✓	<i>NeurIPS22</i> ETF Cls. 45.30 <i>AISTATS23</i> NC-cRT 48.70 <i>ICML23</i> RBL 53.10 BCE3S 59.50	BCE

Table 1: Comparison of BCE3S with previous LTR methods in terms of learning mode and performance. “JL”, “CL”, and “UL” are short for joint learning between sample features and classifier vectors, contrastive learning among features, and classifier’s uniform separability learning. Our Only BCE3S integrates all the three learning modes and achieves the highest accuracy (i.e., 59.50%) on CIFAR100-LT with IF = 100. More results can be found in Table 3.

In model training with cross-entropy (CE) loss, the classifier and sample features are jointly learned, while, on the long-tailed datasets, this learning mode is likely to fail to learn both well sample features and classifier. Currently, the re-balancing techniques such as re-sampling (Wallace et al. 2011), re-weighting (Lin et al. 2017; Cui et al. 2019; Alshammari et al. 2022), and logit adjustment (Menon et al. 2021; Zhao et al. 2022; Li, Cheung, and Lu 2022) have been developed and applied to improve the CE loss to balance the models’ attention among different classes during the training. Combined with the joint learning, contrastive learning (Chen et al. 2020; He et al. 2020) has also been used to enhance the feature properties and improve LTR (Cui et al. 2021; Kang et al. 2021; Du et al. 2023; Xuan and Zhang 2024), by pulling the features of the same class closer and pushing those from different classes apart.

In LTR, in addition to enhancing sample features, it is also necessary to improve the discriminative capability of the classifiers learned on long-tailed datasets. Currently, there is no published work directly focusing on learning a well classifier for LTR. However, according to neural collapse (Papayan, Han, and Donoho 2020; Zhu et al. 2021), the optimal classifier that CE can learn is an equiangular tight frame (ETF) classifier where any two classifier vectors ex-

hibit uniform and maximum separability. Then, customized ETF classifiers (Yang et al. 2022; Kasarla et al. 2022; Liu et al. 2023) are designed before the training, while they may not align well with the final learned features.

To improve LTR, a reasonable approach is to design a uniform separability learning framework for the classifier vectors and integrate it with joint learning and contrastive learning, creating a tripartite synergistic learning (TSL) paradigm to rebalance the features and classifiers during the model training. However, in the LTR tasks, the commonly used CE loss couples imbalanced metrics from different classes on its Softmax’s denominator, making it challenging to effectively suppress the imbalance effect by combining the various CE-based learning modes. As BCE (binary CE) decouples the metrics from different classes by adopting multiple Sigmoid, it can flexibly adjust these metrics and effectively suppress the imbalance effect. In fact, Cui et al. (2019) have experimentally demonstrated that BCE (i.e., Sigmoid cross-entropy in the reference) achieves better LTR than CE, yet the potential of BCE in LTR has not been fully explored.

Therefore, in this paper, we design a BCE-based tripartite synergistic learning approach for LTR, termed BCE3S, which integrates: (1) BCE-based joint learning between sample features and classifiers, (2) BCE-based contrastive learning among features, and (3) BCE-based uniform separability learning for the classifier vectors. As Table 1 shows, with these three learning modes working in concert, BCE3S achieves the best LTR performance on CIFAR100-LT.

The main contributions of this paper are as follows:

1. We introduce the concept of tripartite synergistic learning (TSL) and propose BCE3S, a BCE-based LTR method integrating the joint learning between features and classifier, contrastive learning among features, and uniform learning for classifier vectors.
2. We explain in-depth that BCE-based uniform separability learning helps to train balanced classifier vectors, and analyze the advantage of BCE over CE in LTR.
3. We conduct extensive experiments to evaluate the performance of BCE3S on LTR and find that compared to CE-based methods, BCE3S achieves better intra-class compactness and inter-class separability of sample features, and the BCE-based uniform learning can effectively balance the classifier’s separability. BCE3S achieves SOTA LTR results on four long-tailed datasets.

Related Work

Joint learning for both feature and classifier

Various techniques have been proposed to improve the LTR performance of the popular CE-based joint learning.

The re-sampling technique (Wallace et al. 2011) adjusts the sample distributions via over-sampling the tail classes or under-sampling the head classes to mitigate the class imbalance of the long-tailed datasets, which complicates the model training. The class-balanced loss (Cui et al. 2019) introduces a re-weighting factor inversely proportional to the sample numbers, which is adaptive to the CE, BCE, and focal losses based joint learning. In (Kang et al. 2020), the

authors decoupled the learning of sample features and classifier, and they rectified the recognition decision boundaries on the head and tail classes by fine-tuning with different re-balancing strategies, including the classifier re-weighting. DisAlign (Zhang et al. 2021) developed an adaptive calibration function to re-weight the learning of classifier vectors. The methods of BaLMS (Ren et al. 2020) and Logit Adj. (Menon et al. 2021) re-balance the model training by adjusting the logits according to the sample numbers. Besides these techniques, distillation (Rangwani et al. 2024), diffusion (Shao et al. 2024), and collaborative learning (Xu et al. 2023a) have also enhanced the LTR performance.

The above works are mainly based on CE loss, while the published researches (Cui et al. 2019; Wightman, Touvron, and Jegou 2021; Touvron, Cord, and Jégou 2022) have shown that BCE also performs well in image recognition. Especially, the result in (Cui et al. 2019) has experimentally shown that BCE has more potential than CE in LTR tasks, and LiVT (Xu et al. 2023b) is a BCE-trained Transformer for LTR. However, these works do not in-depth explain the advantages of BCE over CE in LTR or further explore the application of BCE in this task.

Contrastive learning on sample features

Under the long-tailed scenario, PaCo (Cui et al. 2021) incorporates a set of learnable class centers into contrastive learning, re-balancing the feature learning from a perspective of optimization. KCL (Kang et al. 2021) applies the same number of samples for each class in combining its positive pairs to learn balanced feature space by contrastive learning. SSD (Li, Wang, and Wu 2021) employs a self-distillation framework to enhance the information exploitation in the tail class. BCL (Zhu et al. 2022) introduces class-averaging and class-complement to balance the gradient contribution in the contrastive learning. GLMC (Du et al. 2023) enhances the robustness of LTR model via mitigating the head class bias by designing a re-weighting loss. Xuan and Zhang (2024) propose decoupled supervised contrastive loss and patch-based self-distillation to alleviate the imbalance effect in LTR. ProCo (Du et al. 2024) models the feature space using von Mises-Fisher distribution to eliminate the limitation of contrastive learning in requiring a large amount of samples in LTR tasks. These contrastive learning methods are designed on Softmax, measuring the relative value of feature similarity and amplifying the imbalance effect in LTR.

Uniform separability learning for classifier

The works of neural collapse (Papayan, Han, and Donoho 2020; Zhu et al. 2021) show that when training a recognition model on a balanced dataset and the CE loss reaches its minimum, the classifier vectors form an equiangular tight frame (ETF), where uniform and maximum separability exists between any two classifier vectors. However, on an imbalanced dataset, if the imbalance factor (IF) is too large, the classifier vectors of the tail classes will converge and collapse to the same vector (Fang et al. 2021), losing their separability.

To keep the classifier’s uniform separability, Yang et al. (2022), Liu et al. (2023), and Kasarla et al. (2022) customize ETF classifiers before the LTR model training, which do not

participate in the training and maintain the uniform separability; RBL (Peifeng et al. 2023) adopts a learnable orthogonal matrix to rotate the ETF classifier, adjusting its direction based on the sample features. However, these customized ETF classifiers struggle to align with the learned features. In this paper, we propose BCE-based uniform learning, helping to learn balanced classifiers aligned with features.

Method

Preliminary

Let $\mathcal{D} = \bigcup_{k=1}^K \mathcal{D}_k$ be a dataset from K classes, where $\mathcal{D}_k$ contains the samples from the k -th class, and $n_k = |\mathcal{D}_k|$ denotes its sample number. Suppose $\mathcal{D}$ is an imbalanced and long-tailed dataset, and its K subsets have been sorted by the sample numbers in descending order, i.e., $n_k \geq n_\ell$ for $\forall k < \ell$, then n_1/n_K denotes the imbalance factor (IF).

In recognition task, for any sample $\mathbf{X} \in \mathcal{D}$, a deep model $\mathcal{M}(\cdot)$ first extracts its feature $\mathbf{x} = \mathcal{M}(\mathbf{X}) \in \mathbb{R}^d$, then, a linear classifier $\mathcal{C} = \{(\mathbf{w}_j, b_j)\}_{j=1}^K$ converts it into K metrics $\{\mathbf{w}_j^T \mathbf{x} + b_j\}_{j=1}^K$, which are applied to predict the sample's label, $\hat{k} = \arg \max_j \{\mathbf{w}_j^T \mathbf{x} + b_j\}_{j=1}^K$.

In training, for a sample from the k -th subset $\mathcal{D}_k, \forall k$, we denote it as $\mathbf{X}^{(k)}$ and its feature as $\mathbf{x}^{(k)} = \mathcal{M}(\mathbf{X}^{(k)})$. In the previous works (Ren et al. 2020; Li, Cheung, and Lu 2022; Alshammari et al. 2022; Du et al. 2023; Xuan and Zhang 2024; Wang et al. 2021; Chen et al. 2023), to learn the better model $\mathcal{M}$ and classifier $\mathcal{C} = \{(\mathbf{w}_j, b_j)\}_{j=1}^K$, Softmax was first applied to compute the probabilities $\left\{ \frac{\exp(\mathbf{w}_j^T \mathbf{x}^{(k)} + b_j)}{\sum_{\ell=1}^K \exp(\mathbf{w}_\ell^T \mathbf{x}^{(k)} + b_\ell)} \right\}_{j=1}^K$ that $\mathbf{X}^{(k)}$ belongs to each class, then, cross-entropy was used to compute the CE loss,

$$L_{ce}^{(sc)}(\mathbf{x}^{(k)}) = -\log \left(\frac{\exp(\mathbf{w}_k^T \mathbf{x}^{(k)} + b_k)}{\sum_{\ell=1}^K \exp(\mathbf{w}_\ell^T \mathbf{x}^{(k)} + b_\ell)} \right). \quad (1)$$

When applying CE loss to train models, the classifier vectors $\{\mathbf{w}_k\}$ and the sample features $\{\mathbf{x}^{(k)}\}$ are jointly trained and learned. The re-balancing techniques, such as re-sampling and re-weighting, have been used to improve CE loss and enhance its performance on long-tailed recognition (LTR).

Contrastive learning based on Softmax has also been introduced into LTR, comparing features in a projection space,

$$L_{ce}^{(ss)}(\mathbf{x}^{(k)}) = -\log \left(\frac{\exp\left(\frac{1}{\tau} \cos(\mathbf{z}^{(k)}, \mathbf{z}_*^{(k)})\right)}{\sum_{\ell=1}^N \exp\left(\frac{1}{\tau} \cos(\mathbf{z}^{(k)}, \mathbf{z}_*^{(\ell)})\right)} \right), \quad (2)$$

where $\mathbf{z}^{(k)} = \mathcal{P}(\mathbf{x}^{(k)})$, $\mathbf{z}_*^{(k)} = \mathcal{P}(\mathbf{x}_*^{(k)})$, and $\mathcal{P}$ is a non-linear projection operator (Chen et al. 2020). In Eq. (2), τ is a temperature factor, $\cos(\mathbf{z}, \mathbf{z}_*) = \frac{\langle \mathbf{z}, \mathbf{z}_* \rangle}{\|\mathbf{z}\| \|\mathbf{z}_*\|}$ denotes the cosine similarity of any two vectors $\mathbf{z}, \mathbf{z}_* \in \mathbb{R}^{d'}$, and $\{\mathbf{z}_*^{(k)}\}_{k=1}^K$ are temporary features from the K classes, which could be saved in memory bank during the training.

According to neural collapse (Papayan, Han, and Donoho 2020; Zhu et al. 2021), the optimal classifier that CE can learn is the one whose classifier vectors $\{\mathbf{w}_k\}_{k=1}^K$ form an ETF which has uniform and maximum separability, i.e., satisfying $\cos(\mathbf{w}_k, \mathbf{w}_{k'}) = -\frac{1}{K-1}$ and $\|\mathbf{w}_k\| = \|\mathbf{w}_{k'}\|$ for

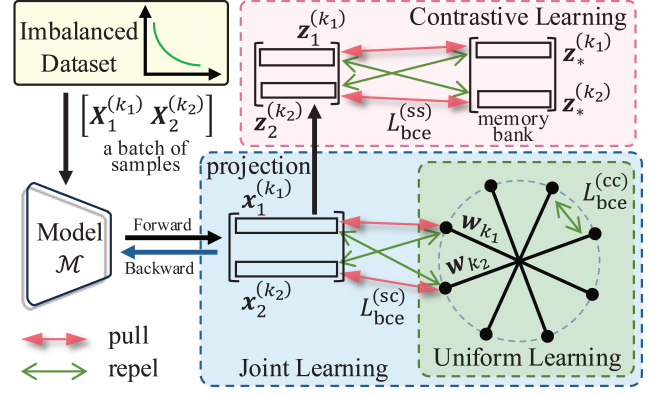


Figure 1: Pipeline of BCE3S, integrating all the three learning modes, i.e., joint learning (Eq. (3)) between sample feature and classifier, contrastive learning (Eq. (4)) among features, and classifier's uniform separability learning (Eq. (5)).

$\forall k \neq k'$. The customized ETF classifiers (Yang et al. 2022; Liu et al. 2023) have been used to design LTR methods.

BCE-based tripartite synergistic learning

We propose the BCE-based tripartite synergistic learning (TSL), i.e., BCE3S, to integrate the above three learning modes. BCE3S consists of BCE-based joint learning $L_{bce}^{(sc)}$, BCE-based contrastive learning $L_{bce}^{(ss)}$, and BCE-based uniform learning $L_{bce}^{(cc)}$. Fig. 1 shows its pipeline.

Joint learning. For a sample feature $\mathbf{x}^{(k)}$ from the k -th class, the joint learning $L_{bce}^{(sc)}$ measures its similarities with the normalized classifier vectors using BCE formula,

$$L_{bce}^{(sc)}(\mathbf{x}^{(k)}) = \log \left(1 + \exp(-\mathbf{w}_k^T \mathbf{x}^{(k)} - b_k) \right) + \sum_{\substack{j=1, j \neq k \\ p_j < r}}^K \log \left(1 + \exp(\mathbf{w}_j^T \mathbf{x}^{(k)} + b_j) \right), \quad (3)$$

where $\|\mathbf{w}_j\| = 1, \forall j$. In Eq. (3), the normalization of the classifier vectors could prevent their gradients from being dominated by the head class in training, which helps to learn more uniform classifier vectors in norm (Kang et al. 2020).

In $L_{bce}^{(sc)}$, $r \in (0, 1]$ is a re-sampling parameter for the negative sample-to-class metrics $\{\mathbf{w}_j^T \mathbf{x}^{(k)}\}_{j \neq k}$, and p_j is randomly sampled from the uniform distribution $U(0, 1)$.

Contrastive learning. For any sample feature $\mathbf{x}^{(k)}$ from class k , similar to (Chen et al. 2020), $L_{bce}^{(ss)}$ implements the contrastive learning in a projection space. A non-linear projection operator $\mathcal{P}$ is first applied to transform the sample feature into $\mathbf{z}^{(k)} = \mathcal{P}(\mathbf{x}^{(k)}) \in \mathbb{R}^{d'}$, and

$$L_{bce}^{(ss)}(\mathbf{x}^{(k)}) = \log \left(1 + \exp \left(-\frac{1}{\tau} \cos(\mathbf{z}^{(k)}, \mathbf{z}_*^{(k)}) \right) \right) + \sum_{\substack{j=1 \\ j \neq k}}^K \log \left(1 + \exp \left(\frac{1}{\tau} \cos(\mathbf{z}^{(k)}, \mathbf{z}_*^{(j)}) \right) \right), \quad (4)$$

where $\{z_*^{(j)} = \mathcal{P}(\mathbf{x}_*^{(j)})\}_{j=1}^K$ are projections of K sample features $\{\mathbf{x}_*^{(j)}\}_{j=1}^K$ stored in a memory bank.

Uniform separability learning for classifier. In the LTR tasks, the joint learning $L_{\text{bce}}^{(\text{sc})}$ tends to learn more discriminative classifier vectors for the head classes but less discriminative ones for the tail classes, and the contrastive learning $L_{\text{bce}}^{(\text{ss})}$ focuses to the sample features, which would aggravate the imbalanced effect. To balance the discriminative property of the classifier vectors for the different classes, we design a uniform separability learning for the classifier,

$$L_{\text{bce}}^{(\text{cc})}(\mathbf{w}_k) = \log(1 + \exp(-\mathbf{w}_k^T \mathbf{w}_k)) + \sum_{\substack{j=1 \\ j \neq k}}^K \log(1 + \exp(\mathbf{w}_k^T \mathbf{w}_j)), \quad (5)$$

where the positive term equals to $\log(1 + e^{-1})$ as $\|\mathbf{w}_k\|^2 = 1$, which is a constant and was omitted in our experiments. The uniform learning $L_{\text{bce}}^{(\text{cc})}$ tries to maximize the separability among the classifier vectors $\{\mathbf{w}_k\}_{k=1}^K$, which helps to learn classifier with uniform and maximum separability (i.e., ETF classifier) and aligning with the final sample features by working together with the joint learning in BCE3S.

Training pipeline. In the experiments, we comprehensively learn the sample features and classifier vectors using BCE3S. As Fig. 1 shows, in each iteration during the training, a batch of samples $[\mathbf{X}_i^{(k_i)}]_{i=1}^B$ are fed into the model $\mathcal{M}$ to extract their features $[\mathbf{x}_i^{(k_i)}]_{i=1}^B$, where B is batch size. A tripartite synergistic loss is computed to train the model,

$$L_{\text{bce}}^{(\text{tri})}([\mathbf{x}_i^{(k_i)}]) = \frac{1}{B} \sum_{i=1}^B L_{\text{bce}}^{(\text{sc})}(\mathbf{x}_i^{(k_i)}) + \frac{\lambda_{\text{ss}}}{B} \sum_{i=1}^B L_{\text{bce}}^{(\text{ss})}(\mathbf{x}_i^{(k_i)}) + \frac{\lambda_{\text{cc}}}{K} \sum_{k=1}^K L_{\text{bce}}^{(\text{cc})}(\mathbf{w}_k), \quad (6)$$

where λ_{ss} and λ_{cc} are weight factors.

Although we design BCE3S, it should be noted that other losses, such as CE and MSE, can also be integrated within the TSL framework. However, for the LTR tasks, BCE not only inherits the advantage of fast convergence from CE but also suppresses the imbalance effects.

Analysis for tripartite synergistic learning (TSL)

We analyze BCE3S in balancing the separability among the different classes and enhancing the feature's properties.

Classifier during the training. Batch algorithm and back propagation are commonly used in the model training. In a iteration, with a batch of features $[\mathbf{x}_i^{(k_i)}]_{i=1}^B$, the normalized classifier vector $\mathbf{w}_k$ is updated via its gradient, i.e.,

$$\mathbf{w}_k \leftarrow \mathbf{w}_k - \eta \frac{\partial L_{\text{bce}}^{(\text{tri})}([\mathbf{x}_i^{(k_i)}])}{\partial \mathbf{w}_k}, \quad \forall k, \quad (7)$$

where η is the learning rate, and

$$\frac{\partial L_{\text{bce}}^{(\text{tri})}([\mathbf{x}_i^{(k_i)}])}{\partial \mathbf{w}_k} = -\frac{1}{B} \left(\sum_{\substack{i=1 \\ k_i=k}}^B \underbrace{\sigma(-\mathbf{w}_k^T \mathbf{x}_i^{(k_i)}) \mathbf{x}_i^{(k_i)}}_{\text{pulling term}} - \sum_{\substack{i=1 \\ k_i \neq k}}^B \underbrace{\sigma(\mathbf{w}_k^T \mathbf{x}_i^{(k_i)}) \mathbf{x}_i^{(k_i)}}_{\text{repelling term}} \right) + \frac{1}{K} \sum_{\substack{j=1 \\ j \neq k}}^K \underbrace{\sigma(\mathbf{w}_k^T \mathbf{w}_j) \mathbf{w}_j}_{\text{interactive term}}, \quad (8)$$

where σ is Sigmoid. In Eq. (8), the pulling terms and repelling terms are derived from the joint learning $L_{\text{bce}}^{(\text{sc})}$, and they pull the classifier vector $\mathbf{w}_k$ using the sample features $\mathbf{x}_i^{(k_i=k)}$ from the same class and repel it using the ones $\mathbf{x}_i^{(k_i \neq k)}$ from the different classes. Due to the higher probability that the batch contains samples from the head classes, the classifier vectors of the head classes are prone to be repelled by the samples of other head classes in different directions, while the tail classifier vectors are not easily to be repelled by the samples from the other tail classes, which results in high separability among the head classifier vectors and the low ones among the tail classifier vectors. The tail classifier vectors are also easily repelled by the head class samples, which would lead to higher separability between classifier vectors of the head and tail classes, because the repulsion from the head class is more concentrated on the tail classifier vectors. Totally, the joint learning $L_{\text{bce}}^{(\text{sc})}$ will learn a imbalanced classifier on LTR, which will affect the feature learning and result in imbalanced features across classes.

In contrast, the uniform learning $L_{\text{bce}}^{(\text{cc})}$ leads to the interactive term in Eq. (8). For each $\mathbf{w}_k, \forall k$, in every batch, the interactive term provides $K - 1$ repelling forces from other classifier vectors, which directly, uniformly, and consistently maximize the separability among the K classifier vectors. Thereby, $L_{\text{bce}}^{(\text{cc})}$ helps to learn a balanced classifier on LTR. Combined with $L_{\text{bce}}^{(\text{sc})}$, the balanced classifier will align with the feature learning, and re-balance the features.

The contrastive learning $L_{\text{bce}}^{(\text{ss})}$ enhances the feature learning in the training, which does not directly update the classifier. However, when combined with joint learning $L_{\text{bce}}^{(\text{sc})}$, it can further exacerbate the imbalance of classifier separability through imbalanced sample features.

CE vs. BCE in LTR. On the LTR, both CE-based and BCE-based joint learning, $L_{\text{bce}}^{(\text{sc})}$ and $L_{\text{bce}}^{(\text{ss})}$, result in imbalanced classifier vectors among the head and tail classes, while they perform diversely on the sample features.

For a sample feature $\mathbf{x}^{(k)}$, it is updated by

$$\mathbf{x}^{(k)} \leftarrow \mathbf{x}^{(k)} - \eta \frac{\partial L_{\mu}^{(\text{sc})}(\mathbf{x}^{(k)})}{\partial \mathbf{x}^{(k)}}, \quad (9)$$

where $\mu \in \{\text{ce}, \text{bce}\}$ denotes CE- or BCE-based joint learning, and η is the learning rate. The gradients $\frac{\partial L_{\mu}^{(\text{sc})}(\mathbf{x}^{(k)})}{\partial \mathbf{x}^{(k)}}$ have the similar form in the CE or BCE,

$$-(1 - \text{Act}_{\mu}(\mathbf{w}_k^T \mathbf{x}^{(k)})) \mathbf{w}_k + \sum_{\substack{j=1 \\ j \neq k}}^K \underbrace{\text{Act}_{\mu}(\mathbf{w}_j^T \mathbf{x}^{(k)}) \mathbf{w}_j}_{\text{repelling}} \quad (10)$$

where, for $\forall j$,

$$\text{Act}_{\text{ce}}(\mathbf{w}_j^T \mathbf{x}^{(k)}) = \frac{\exp(\mathbf{w}_j^T \mathbf{x}^{(k)} + b_j)}{\sum_{\ell=1}^K \exp(\mathbf{w}_\ell^T \mathbf{x}^{(k)} + b_\ell)}, \quad (11)$$

$$\text{Act}_{\text{bce}}(\mathbf{w}_j^T \mathbf{x}^{(k)}) = \sigma(\mathbf{w}_j^T \mathbf{x}^{(k)}). \quad (12)$$

In the feature updating, $L_{\text{ce}}^{(\text{sc})}$ and $L_{\text{bce}}^{(\text{sc})}$ utilize the exponential inner products of the feature $\mathbf{x}^{(k)}$ with K classifier vectors $\{\mathbf{w}_j\}_{j=1}^K$ to compute one pulling term and $K - 1$ repelling terms, while the imbalanced classifier vectors $\{\mathbf{w}_j\}_{j=1}^K$ will slow down the feature learning.

For CE-based joint learning, $L_{\text{ce}}^{(\text{sc})}$, in computing each pulling or repelling term, the K imbalanced exponential inner products are coupling on the denominator of Softmax Act_{ce} , thereby, the imbalanced effect is **injected again** into the feature learning. In contrast, BCE decouples the metrics between the feature with K classifier vectors, and only one classifier vector is applied in computing one pulling or repelling term by Sigmoid Act_{bce} . Therefore, BCE will achieve more balanced features across classes.

Similarly, BCE-based contractive learning and uniform separability learning, $L_{\text{bce}}^{(\text{ss})}$ and $L_{\text{bce}}^{(\text{cc})}$, could respectively result in better classifier and features than the CE-based ones, $L_{\text{ce}}^{(\text{ss})}$ and $L_{\text{ce}}^{(\text{cc})}$ (see supplementary for their formulas).

During the training, compared to CE, BCE adds at most $K - 1$ logarithmic operations, which is negligible. In the testing or inference, it does not incur any additional cost.

Experiments and Results

We evaluate BCE3S on four long-tailed datasets: CIFAR10-LT/CIFAR100-LT (Yue et al. 2016), ImageNet-LT (Liu et al. 2022a), and iNaturalist2018 (Van Horn et al. 2018). The CIFAR variants have IF of 100, 50, and 10, while ImageNet-LT and iNaturalist2018 have IF of 256 and 500, respectively. We train ResNets on their imbalance training set and evaluate them on the balance test set. Following (Kang et al. 2020; Alshammari et al. 2022; Du et al. 2024), we report total accuracy on the entire test set and three subsets: *Many* (> 100 samples per class), *Medium* (≤ 100 and ≥ 20 samples per class), and *Few* (< 20 samples per class). More details about datasets and training can be found in supplementary.

Ablation Study

LTR results. We conduct ablation study for the proposed BCE3S on CIFAR100-LT with IF = 100 using ResNet32, and Table 2 presents the results. We first take the CE-based joint learning $L_{\text{ce}}^{(\text{sc})}$ as baseline. When training the model using BCE-based joint learning $L_{\text{bce}}^{(\text{sc})}$, though it reduces the accuracy on the *Many* subset, it improves the accuracy on the *Medium* and *Few* subsets, and increasing the overall accuracy from 51.48% to 52.88%, indicating that BCE pay more attention to the tail classes and improves the total accuracy at the cost of reduced accuracy on the head classes.

Based on $L_{\text{ce}}^{(\text{sc})}$, we compare CE- and BCE-based contrastive learning, $L_{\text{ce}}^{(\text{ss})}$ and $L_{\text{bce}}^{(\text{ss})}$. As Table 2 shows, $L_{\text{bce}}^{(\text{ss})}$ improves accuracy across all subsets, increasing total accuracy

Methods						Many	Med.	Few	All
$L_{\text{ce}}^{(\text{sc})}$	$L_{\text{ce}}^{(\text{ss})}$	$L_{\text{ce}}^{(\text{cc})}$	$L_{\text{bce}}^{(\text{sc})}$	$L_{\text{bce}}^{(\text{ss})}$	$L_{\text{bce}}^{(\text{cc})}$				
✓						82.29	51.37	15.67	51.48
			✓			81.11	55.06	17.40	52.88
-	✓	-	-	-	-	84.09	54.80	17.53	53.87
✓				✓		84.17	55.20	17.90	54.15
	✓		✓			83.37	55.83	19.87	54.68
			✓	✓		82.74	56.57	20.63	54.95
-	✓	-	✓	-	-	82.31	51.83	17.20	52.11
✓					✓	81.23	53.69	18.17	52.67
		✓	✓			81.57	55.51	17.93	53.36
			✓		✓	81.03	56.51	19.20	53.90
-	✓	-	✓	-	-	83.97	54.54	18.87	54.14
✓				✓	✓	83.77	54.49	18.67	53.99
	✓	✓	✓			83.77	56.17	21.40	55.40
			✓	✓	✓	83.34	57.09	22.80	55.99

Table 2: Ablation study for the proposed BCE3S and the CE based counterparts on CIFAR100-LT (IF = 100). The formulas of CE-based contrastive learning and uniform learning, $L_{\text{ce}}^{(\text{ss})}$ and $L_{\text{ce}}^{(\text{cc})}$, are presented in supplementary.

from 53.87% to 54.15%, and further to 54.95% when combined with $L_{\text{bce}}^{(\text{sc})}$. For uniform learning, $L_{\text{bce}}^{(\text{cc})}$ enhances the performance on *Medium* and *Few* despite reducing that on *Many*, yielding better overall performance (52.67%), which increases to 53.90% when combined with $L_{\text{bce}}^{(\text{sc})}$; meanwhile, it always perform better than CE-based one $L_{\text{ce}}^{(\text{cc})}$.

We finally compare different tripartite synergistic learning approaches based on CE and BCE, and achieve the best total LTR accuracy (55.99%) using the proposed BCE3S.

Separability and compactness. To analyze the models trained with various learning methods, for any class k , we define intra-class compactness $\mathcal{E}_k^{(\text{x,com})}$, inter-class separability $\mathcal{E}_k^{(\text{x,sep})}$ among features, and separability $\mathcal{E}_k^{(\text{w,sep})}$ among the classifier vectors; as their names imply, $\mathcal{E}_k^{(\text{x,com})}$ measures the compactness of features within the k -th class, $\mathcal{E}_k^{(\text{x,sep})}$ measures the separability between the features from k -th class and that from other classes, and $\mathcal{E}_k^{(\text{w,sep})}$ measures the separability between the k -th classifier vector with other ones, seeing supplementary for their definitions. The higher values suggest the better separability or compactness of the classifier vectors or features. Fig. 2 shows the results of these metrics for CE- and BCE-based methods, as well as their mean and standard deviation across different classes.

As the first row in the figure shows, the four CE-based methods achieve similar intra-class compactness, with their average values all approximately 82, and a noticeable imbalance is shown from the head to tail classes. In contrast, for the BCE-based methods, the average compactness of $L_{\text{bce}}^{(\text{sc})}$ reaches to 86.02. With the addition of $L_{\text{bce}}^{(\text{ss})}$ and $L_{\text{bce}}^{(\text{cc})}$, the compactness gradually increase to 95.47, while the imbalance across different classes gradually decreases and standard deviation of the compactness decreases from 5.55 to 1.81. As the second row shows, compared to the CE-based methods, the BCE-based ones achieve higher inter-class separability for sample features, with the highest mean of separability increases from 49.71 ($L_{\text{ce}}^{(\text{sc})} + L_{\text{ce}}^{(\text{ss})}$) to 50.84

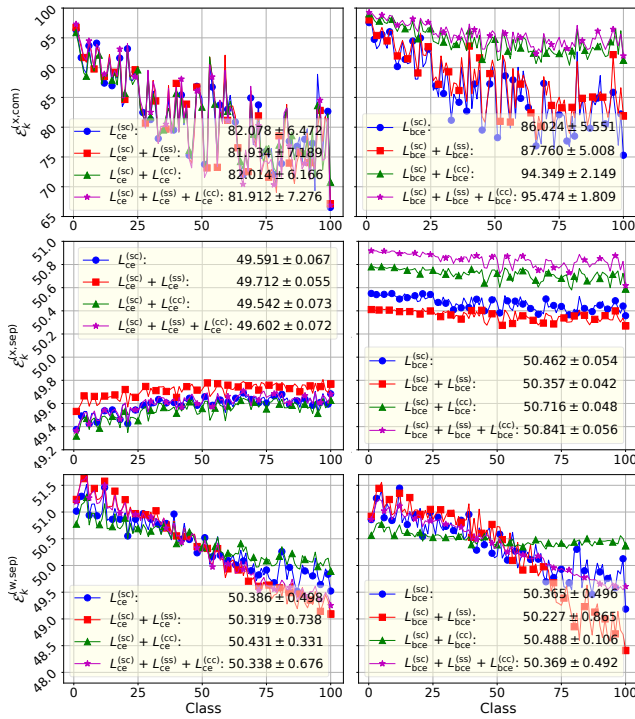


Figure 2: The intra-class compactness (top), inter-class separability (middle) of sample features, and separability (bottom) of classifier vectors on the training set of CIFAR100-LT (IF = 100), with the model trained using different CE- (left) and BCE-based (right) methods.

($L_{bce}^{(sc)} + L_{bce}^{(ss)} + L_{bce}^{(cc)}$). Totally, compared to CE-based methods, BCE3S enhances the feature properties and balances the property difference on the head and tail classes.

As the third row of Fig. 2 shows, the average separability of classifier vectors for the various methods is similar, close to 50.50, while they have significant differences across the different classes. The classifier separability of the two joint learning $L_{ce}^{(sc)}$ and $L_{bce}^{(sc)}$ has significant imbalance from the head to tail classes. When the CE- and BCE-based contrastive learning are added, the imbalance is respectively amplified, while the two kinds of uniform learning can reduce the imbalance effect. In particular, the BCE one significantly suppresses the separability imbalance, resulting in a standard deviation of only 0.106 for $L_{bce}^{(sc)} + L_{bce}^{(cc)}$.

These results indicate that BCE3S can not only enhance the feature properties, but also effectively improve the balance of features and classifier for the LTR models.

Feature distribution in t-SNE. To intuitively compare the feature properties between our BCE3S with CE-based methods, for ResNet32 trained on CIFAR10-LT with IF = 100, we visually show the feature distributions of the 10 classes on the test set in Fig. 3 using t-SNE. As the figure shows, the feature clusters of class 3 and 5 (i.e., “cat” and “dog”) of CE-based joint learning $L_{ce}^{(sc)}$ overlap with each other; though the intersection decreases with the addition of $L_{ce}^{(ss)}$ and $L_{ce}^{(cc)}$, the final feature clusters of the two classes are not completely separated, indicating unsatisfactory inter-class separability. Meanwhile, the features of the tail classes

Methods	CIFAR10-LT			CIFAR100-LT		
	100	50	10	100	50	10
CB-FocalCVPR’19	74.60	79.30	87.10	39.60	45.20	58.00
SSPNurIPS’20	77.80	82.10	88.50	43.40	47.10	58.90
TSCCVPR’22	79.70	82.90	88.70	43.80	47.40	59.00
ETF+DRNurIPS’22	76.50	-	87.70	45.30	-	-
NC-DRWAISTATS’23	81.90	-	89.80	48.60	-	63.10
RIDE (3 exp.)ICLR’21	-	-	-	48.60	51.40	59.80
BCLCVPR’22	84.50	87.20	91.10	51.90	56.40	64.60
NC.-cRTAISTATS’23	82.60	-	90.20	48.70	-	63.60
ResLT+H2T AAAI’24	81.77	84.99	-	49.60	54.39	-
Logit Adj.ICLR’21	84.30	87.10	90.90	50.50	54.90	64.00
BCL+DODAICLR’24	-	-	-	51.00	53.60	62.70
RIDE+H2TAAAI’24	-	-	-	51.38	55.54	-
DiffuLTNurIPS’24	84.70	86.90	90.70	51.50	56.30	63.80
PaCoICCV’21	-	-	-	52.00	56.00	64.20
DiffuLT+RIDENurIPS’24	85.30	87.30	90.90	52.40	56.90	64.20
ProCoTPAMI’24	85.90	88.20	91.90	52.80	57.10	65.50
RBLICML’23	84.70	87.60	-	53.10	57.20	-
DeiT-LT CVPR’24	87.50	89.80	-	55.60	60.50	-
GLMCCVPR’23	88.50	91.04	94.87	57.97	63.78	73.40
GLMC + MN	89.58	92.04	94.87	58.41	64.57	74.28
BCE3S	90.08	92.55	95.71	59.50	65.23	76.13

Table 3: LTR on CIFAR10-LT and CIFAR100-LT, IF = 100, 50, and 10. BCE3S consistently achieves the best results.

(the 8th and 9th classes) are spreading over relatively elongated regions, indicating unsatisfactory intra-class compactness. In contrast, for BCE, the ten feature clusters of $L_{bce}^{(sc)}$ are distributed in relatively independent regions, while with the addition of $L_{bce}^{(ss)}$ and $L_{bce}^{(cc)}$, the clusters of the tail classes becomes increasingly compact, indicating higher separability and compactness among the features.

Parameter study. BCE3S applied hyper-parameters λ_{ss} , λ_{cc} , and τ to balance the impact of three learning components, and we have conducted a comprehensive study using ResNet32 on CIFAR100-LT with IF = 100. The detailed experimental results for the parameter study can be found in the supplementary.

Comparison with SOTA

We compare our BCE3S with a series of SOTA LTR methods on CIFAR10-LT, CIFAR100-LT, ImageNet-LT, and iNaturalist2018 with different backbones. To further explore the potential of BCE3S, we boost its performance by combining it with re-balancing techniques (Alshammari et al. 2022; Ren et al. 2020). The detailed experimental settings and more results can be found in the supplementary.

On **CIFAR10-LT and CIFAR100-LT**, we adopted a similar training strategy used in MaxNorm (MN) of Alshammari et al. (2022). As Table 3 shows, our BCE3S achieves the best top-1 accuracy on both CIFAR10-LT and CIFAR100-LT with different IFs. For example, on CIFAR100-LT, BCE3S achieves accuracies of 76.13%, 65.23%, and 59.50% with the three IFs, which respectively surpass the previous best results by 1.85%, 0.66%, and 1.09%, being new SOTA. On CIFAR10-LT dataset, BCE3S has also achieved SOTA results. Those results highlight the potential of a model trained using BCE3S.

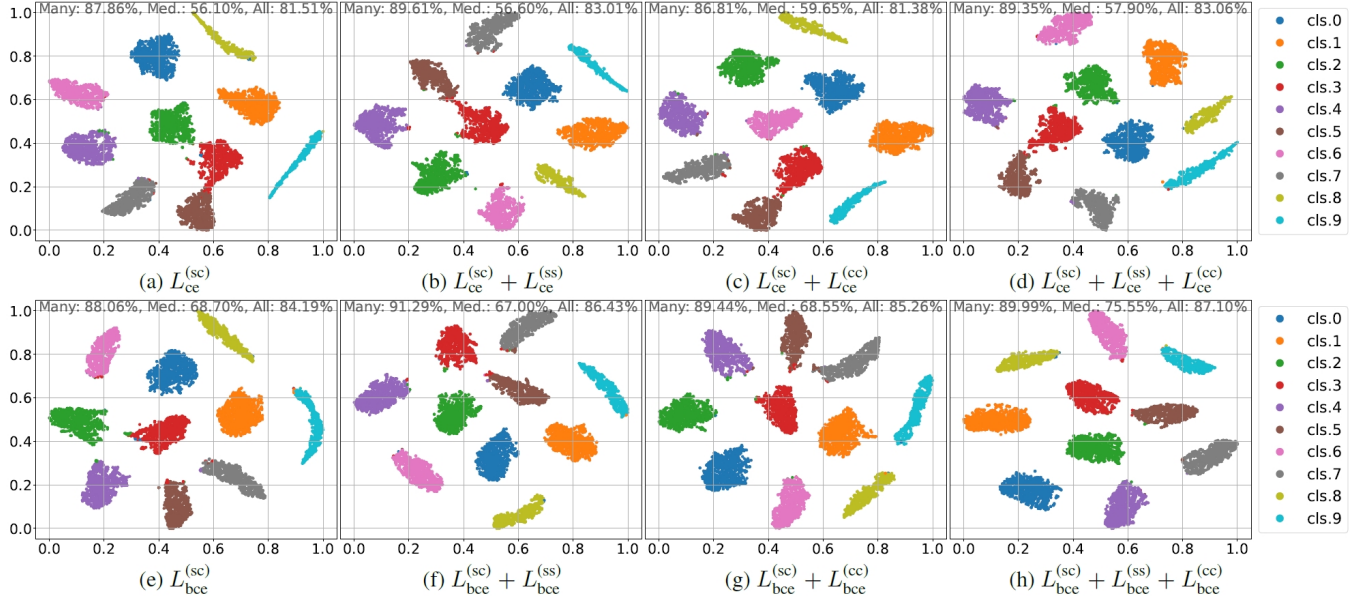


Figure 3: Feature distribution on the CIFAR10-LT test set with CE (top) and BCE (bottom) learning methods. Compared to CE methods, features extracted using BCE-based joint learning $L_{bce}^{(sc)}$ show improved intra-class compactness and inter-class separability. The contrastive learning $L_{bce}^{(ss)}$ and uniform learning $L_{bce}^{(cc)}$ further enhance these properties.

Methods	$\mathcal{M}$	ImageNet-LT			
		Many	Med.	Few	All
CB-FocalCVPR'19	R50	39.60	32.70	16.80	33.20
RISDAAAAI'22		-	-	-	49.30
LDAMNeurIPS'19		60.40	46.90	30.70	49.80
KCLICLR'21		61.80	49.40	30.90	51.50
TSCCVPR'22		63.50	49.40	30.40	52.40
GCLCVPR'22		62.24	48.62	52.12	54.51
GCL+H2TAAAI'24		62.36	48.75	52.15	54.62
BCLCVPR'22		65.30	53.50	36.30	55.60
DiffuLTNeurIPS'24		63.20	55.40	39.20	56.20
BCL+DODAICLR'24		66.90	54.10	37.40	56.90
DiffuLT+RIDENeurIPS'24		64.10	<u>55.80</u>	<u>39.90</u>	56.90
DSCLAAAI'24		68.50	55.20	35.40	57.70
ProCOTPAMI'24		<u>68.20</u>	55.10	38.10	<u>57.80</u>
BCE3S		68.14	55.90	40.56	57.85

Table 4: LTR results on ImageNet-LT using ResNet50 (R50). BCE3S achieves the best results on the test set.

On **ImageNet-LT**, we evaluate BCE3S using ResNet50. As Table 4 shows, BCE3S achieves the best performance on Medium, and Few, with accuracy of 55.90% and 40.56%, and obtains the best overall accuracy of 57.85%.

On **iNaturalist2018**, as Table 5 shows, we apply BCE3S to train ResNet50 using a strategy similar to that of Du et al. (2024). When trained for 180 epochs, BCE3S achieves a competitive accuracy of 73.99% on the overall test set, outperforming most previous methods. When extended to 400 epochs, BCE3S achieves the best performance on test subsets of Many and Med., and get an optimal overall accuracy of 75.91%, surpassing ProCo and DeiT-LT, though DeiT-LT uses ViT-B trained for 1,000 epochs.

Methods	iNaturalist 2018			
	Many	Med.	Few	All
KCLICLR'21	-	-	-	68.60
TSCCVPR'22	72.60	70.60	67.80	69.70
DRICLR'20	72.88	71.15	69.24	70.49
GCLCVPR'22	66.43	71.66	72.47	71.47
GCL+H2TAAAI'24	67.74	71.92	72.22	71.62
BCLCVPR'22	69.40	72.40	71.80	71.80
DR+H2TAAAI'24	71.73	72.32	71.30	71.81
DSCLAAAI'24	74.20	72.90	70.30	72.00
RIDEICLR'21	76.52	74.23	70.45	72.80
RIDE+H2TAAAI'24	75.69	74.22	71.36	73.11
PaCo (400 ep.)ICCV'21	70.30	73.20	73.60	73.20
ProCo (90 ep.)TPAMI'24	-	-	-	73.50
BCL+DODAICLR'24	71.20	73.20	73.40	73.70
BCE3S (180 ep.)	<u>77.16</u>	74.45	72.78	73.99
DeiT-LT(ViT-B, 1K ep.)CVPR'24	<u>70.30</u>	<u>75.20</u>	<u>76.20</u>	<u>75.10</u>
ProCo(400 ep.)TPAMI'24	74.00	<u>76.00</u>	76.80	<u>75.80</u>
BCE3S (400 ep.)	79.10	76.08	74.28	75.91

Table 5: LTR results on iNaturalist2018 using ResNet50.

Conclusion

For long-tailed recognition (LTR) on imbalanced datasets, this paper proposes the BCE-based tripartite synergistic learning method, i.e., BCE3S, by integrating the joint learning between sample features and classifier vectors, contrastive learning among the features, and uniform separability learning for the classifier vectors. Compared with CE-based methods, BCE3S suppresses the imbalance effect among the classifier vectors of the head and tail classes, improves the models' focus on the feature learning in the tail classes, and enhances the intra-class compactness and inter-class separability of the features. The extensive experiments demonstrate that our BCE3S achieves the optimal LTR performance on various long-tailed datasets.

Acknowledgements

The research was supported by National Natural Science Foundation of China under Grant No. 62576217 and 8226113862, Natural Science Foundation of Ningxia Province under Grant No. 2025AAC020002, Guangdong Provincial Key Laboratory under Grant No. 2023B1212060076, and Scientific Foundation for Youth Scholars of Shenzhen University under Grant No. 868-000001032180.

References

- Alshammari, S.; Wang, Y.-X.; Ramanan, D.; and Kong, S. 2022. Long-tailed recognition via weight balancing. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 6897–6907.
- Cao, K.; Wei, C.; Gaidon, A.; Arechiga, N.; and Ma, T. 2019. Learning imbalanced datasets with label-distribution-aware margin loss. *Advances in neural information processing systems*, 32.
- Chen, T.; Kornblith, S.; Norouzi, M.; and Hinton, G. 2020. A Simple Framework for Contrastive Learning of Visual Representations. arXiv:2002.05709.
- Chen, X.; Zhou, Y.; Wu, D.; Yang, C.; Li, B.; Hu, Q.; and Wang, W. 2023. Area: adaptive reweighting via effective area for long-tailed classification. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 19277–19287.
- Chen, X.; Zhou, Y.; Wu, D.; Zhang, W.; Zhou, Y.; Li, B.; and Wang, W. 2022. Imagine by reasoning: A reasoning-based implicit semantic data augmentation for long-tailed classification. In *Proceedings of the AAAI conference on artificial intelligence*, volume 36, 356–364.
- Cui, J.; Zhong, Z.; Liu, S.; Yu, B.; and Jia, J. 2021. Parametric contrastive learning. In *Proceedings of the IEEE/CVF international conference on computer vision*, 715–724.
- Cui, Y.; Jia, M.; Lin, T.-Y.; Song, Y.; and Belongie, S. 2019. Class-balanced loss based on effective number of samples. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 9268–9277.
- Du, C.; Wang, Y.; Song, S.; and Huang, G. 2024. Probabilistic contrastive learning for long-tailed visual recognition. *IEEE Transactions on Pattern Analysis and Machine Intelligence*.
- Du, F.; Yang, P.; Jia, Q.; Nan, F.; Chen, X.; and Yang, Y. 2023. Global and local mixture consistency cumulative learning for long-tailed visual recognitions. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 15814–15823.
- Fang, C.; He, H.; Long, Q.; and Su, W. J. 2021. Exploring deep neural networks via layer-peeled model: Minority collapse in imbalanced training. *Proceedings of the National Academy of Sciences*, 118(43): e2103091118.
- He, K.; Fan, H.; Wu, Y.; Xie, S.; and Girshick, R. 2020. Momentum contrast for unsupervised visual representation learning. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 9729–9738.
- Kang, B.; Li, Y.; Xie, S.; Yuan, Z.; and Feng, J. 2021. Exploring Balanced Feature Spaces for Representation Learning. In *International Conference on Learning Representations*.
- Kang, B.; Xie, S.; Rohrbach, M.; Yan, Z.; Gordo, A.; Feng, J.; and Kalantidis, Y. 2020. Decoupling Representation and Classifier for Long-Tailed Recognition. In *International Conference on Learning Representations*.
- Kasarla, T.; Burghouts, G.; Van Spengler, M.; Van Der Pol, E.; Cucchiara, R.; and Mettes, P. 2022. Maximum class separation as inductive bias in one matrix. *Advances in neural information processing systems*, 35: 19553–19566.
- Li, M.; Cheung, Y.-m.; and Lu, Y. 2022. Long-tailed visual recognition via gaussian clouded logit adjustment. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 6929–6938.
- Li, M.; Hu, Z.; Lu, Y.; Lan, W.; ming Cheung, Y.; and Huang, H. 2024. Feature Fusion from Head to Tail for Long-Tailed Visual Recognition. In *Proceedings of the AAAI Conference on Artificial Intelligence*, 13581–13589.
- Li, T.; Cao, P.; Yuan, Y.; Fan, L.; Yang, Y.; Feris, R. S.; Indyk, P.; and Katabi, D. 2022. Targeted supervised contrastive learning for long-tailed recognition. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 6918–6928.
- Li, T.; Wang, L.; and Wu, G. 2021. Self supervision to distillation for long-tailed visual recognition. In *Proceedings of the IEEE/CVF international conference on computer vision*, 630–639.
- Lin, T.-Y.; Goyal, P.; Girshick, R.; He, K.; and Dollár, P. 2017. Focal loss for dense object detection. In *Proceedings of the IEEE international conference on computer vision*, 2980–2988.
- Liu, X.; Zhang, J.; Hu, T.; Cao, H.; Yao, Y.; and Pan, L. 2023. Inducing neural collapse in deep long-tailed learning. In *International Conference on Artificial Intelligence and Statistics*, 11534–11544. PMLR.
- Liu, Z.; Miao, Z.; Zhan, X.; Wang, J.; Gong, B.; and Stella, X. Y. 2022a. Open long-tailed recognition in a dynamic world. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 46(3): 1836–1851.
- Liu, Z.; Miao, Z.; Zhan, X.; Wang, J.; Gong, B.; and Stella, X. Y. 2022b. Open long-tailed recognition in a dynamic world. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 46(3): 1836–1851.
- Menon, A. K.; Jayasumana, S.; Rawat, A. S.; Jain, H.; Veit, A.; and Kumar, S. 2021. Long-tail learning via logit adjustment. In *International Conference on Learning Representations*.
- Papayan, V.; Han, X. Y.; and Donoho, D. L. 2020. Prevalence of neural collapse during the terminal phase of deep learning training. *Proceedings of the National Academy of Sciences*, 117(40): 24652–24663.
- Peifeng, G.; Xu, Q.; Wen, P.; Yang, Z.; Shao, H.; and Huang, Q. 2023. Feature Directions Matter: Long-Tailed Learning

- via Rotated Balanced Representation. In Krause, A.; Brunskill, E.; Cho, K.; Engelhardt, B.; Sabato, S.; and Scarlett, J., eds., *Proceedings of the 40th International Conference on Machine Learning*, volume 202 of *Proceedings of Machine Learning Research*, 27542–27563. PMLR.
- Rangwani, H.; Mondal, P.; Mishra, M.; Asokan, A. R.; and Babu, R. V. 2024. DeiT-LT: Distillation Strikes Back for Vision Transformer Training on Long-Tailed Datasets. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 23396–23406.
- Ren, J.; Yu, C.; Ma, X.; Zhao, H.; Yi, S.; et al. 2020. Balanced meta-softmax for long-tailed visual recognition. *Advances in neural information processing systems*, 33: 4175–4186.
- Russakovsky, O.; Deng, J.; Su, H.; Krause, J.; Satheesh, S.; Ma, S.; Huang, Z.; Karpathy, A.; Khosla, A.; Bernstein, M.; et al. 2015. Imagenet large scale visual recognition challenge. *International journal of computer vision*, 115: 211–252.
- Shao, J.; Zhu, K.; Zhang, H.; and Wu, J. 2024. DiffuLT: Diffusion for Long-tail Recognition Without External Knowledge. *Advances in neural information processing systems*, 37.
- Touvron, H.; Cord, M.; and Jégou, H. 2022. DeiT iii: Revenge of the vit. In *European conference on computer vision*, 516–533. Springer.
- Van Horn, G.; Mac Aodha, O.; Song, Y.; Cui, Y.; Sun, C.; Shepard, A.; Adam, H.; Perona, P.; and Belongie, S. 2018. The inaturalist species classification and detection dataset. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, 8769–8778.
- Wallace, B. C.; Small, K.; Brodley, C. E.; and Trikalinos, T. A. 2011. Class imbalance, redux. In *2011 IEEE 11th international conference on data mining*, 754–763. Ieee.
- Wang, B.; Wang, P.; Xu, W.; Wang, X.; Zhang, Y.; Wang, K.; and Wang, Y. 2024. Kill Two Birds with One Stone: Rethinking Data Augmentation for Deep Long-tailed Learning. In *The Twelfth International Conference on Learning Representations*.
- Wang, X.; Lian, L.; Miao, Z.; Liu, Z.; and Yu, S. 2021. Long-tailed Recognition by Routing Diverse Distribution-Aware Experts. In *International Conference on Learning Representations*.
- Wightman, R.; Touvron, H.; and Jégou, H. 2021. ResNet strikes back: An improved training procedure in timm. In *NeurIPS 2021 Workshop on ImageNet: Past, Present, and Future*.
- Xu, Z.; Chai, Z.; Xu, C.; Yuan, C.; and Yang, H. 2023a. Towards Effective Collaborative Learning in Long-Tailed Recognition. *IEEE Transactions on Multimedia*.
- Xu, Z.; Liu, R.; Yang, S.; Chai, Z.; and Yuan, C. 2023b. Learning imbalanced data with vision transformers. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 15793–15803.
- Xuan, S.; and Zhang, S. 2024. Decoupled Contrastive Learning for Long-Tailed Recognition. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, 6396–6403.
- Yang, Y.; Chen, S.; Li, X.; Xie, L.; Lin, Z.; and Tao, D. 2022. Inducing neural collapse in imbalanced learning: Do we really need a learnable classifier at the end of deep neural network? *Advances in neural information processing systems*, 35: 37991–38002.
- Yue, C.; Long, M.; Wang, J.; Han, Z.; and Wen, Q. 2016. Deep quantization network for efficient image retrieval. In *Proceedings of the AAAI Conference on Artificial Intelligence*, 3457–3463.
- Zhang, S.; Li, Z.; Yan, S.; He, X.; and Sun, J. 2021. Distribution alignment: A unified framework for long-tail visual recognition. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2361–2370.
- Zhao, Y.; Chen, W.; Tan, X.; Huang, K.; and Zhu, J. 2022. Adaptive logit adjustment loss for long-tailed visual recognition. In *Proceedings of the AAAI conference on artificial intelligence*, volume 36, 3472–3480.
- Zhu, J.; Wang, Z.; Chen, J.; Chen, Y.-P. P.; and Jiang, Y.-G. 2022. Balanced contrastive learning for long-tailed visual recognition. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 6908–6917.
- Zhu, Z.; Ding, T.; Zhou, J.; Li, X.; You, C.; Sulam, J.; and Qu, Q. 2021. A geometric analysis of neural collapse with unconstrained features. *Advances in Neural Information Processing Systems*, 34: 29820–29834.

Appendix A: Separability and compactness of classifier and sample features

Let $\mathcal{D} = \bigcup_{k=1}^K \mathcal{D}_k$ be a dataset from K classes, where

$$\mathcal{D}_k = \left\{ \mathbf{X}_i^{(k)} \right\}_{i=1}^{n_k}$$

contains n_k samples from the k -th class. In recognition task, for each sample $\mathbf{X}_i^{(k)}$, an encoder $\mathcal{M}(\cdot)$ first extracts its feature $\mathbf{x}_i^{(k)} = \mathcal{M}(\mathbf{X}_i^{(k)}) \in \mathbb{R}^d$, then, a linear classifier $\mathcal{C} = \{(\mathbf{w}_j, b_j)\}_{j=1}^K$ converts it into K metrics $\{\mathbf{w}_j^T \mathbf{x}_i^{(k)} + b_j\}_{j=1}^K$, which are applied to predict the sample's label,

$$\hat{k} = \arg \max_j \{\mathbf{w}_j^T \mathbf{x}_i^{(k)} + b_j\}_{j=1}^K.$$

In the ablation experiments, to compare the learned classifier and features, we define three metrics, separability $\mathcal{E}_k^{(w, \text{sep})}$ among classifier vectors, intra-class compactness $\mathcal{E}_k^{(x, \text{com})}$, and inter-class separability $\mathcal{E}_k^{(x, \text{sep})}$ among sample features, for any class k ,

$$\mathcal{E}_k^{(w, \text{sep})} = \frac{1}{K-1} \sum_{\substack{j=1 \\ j \neq k}}^K \left(\frac{1 - \cos(\hat{\mathbf{w}}_k, \hat{\mathbf{w}}_j)}{2} \right) \times 100\% \quad (13)$$

$$\mathcal{E}_k^{(x, \text{com})} = \frac{1}{n_k^2 - n_k} \sum_{i=1}^{n_k} \sum_{\substack{i'=1 \\ i' \neq i}}^{n_k} \left(\frac{\cos(\mathbf{x}_i^{(k)}, \mathbf{x}_{i'}^{(k)}) + 1}{2} \right) \times 100\%, \quad (14)$$

$$\mathcal{E}_k^{(x, \text{sep})} = \frac{1}{n_k K - n_k} \sum_{i=1}^{n_k} \sum_{\substack{j=1 \\ j \neq k}}^K \left(\frac{1 - \cos(\hat{\mathbf{x}}_i^{(k)}, \bar{\mathbf{x}}^{(j)})}{2} \right) \times 100\%, \quad (15)$$

where, for $\forall k$,

$$\hat{\mathbf{w}}_k = \mathbf{w}_k - \bar{\mathbf{w}}, \quad (16)$$

$$\bar{\mathbf{w}} = \frac{1}{K} \sum_{j=1}^K \mathbf{w}_j, \quad (17)$$

$$\bar{\mathbf{x}}^{(k)} = \frac{1}{n_k} \sum_{i=1}^{n_k} \mathbf{x}_i^{(k)}, \quad (18)$$

$$\hat{\mathbf{x}}_i^{(k)} = \mathbf{x}_i^{(k)} - \bar{\mathbf{x}}^{(k)}, \quad (19)$$

and $\mathbf{x}_i^{(k)}$ is the feature of the i -th sample from the k -th class.

Appendix B: Neural collapse

The neural collapse was first found by Papayan, Han, and Donoho (2020), which refers to four properties about the sample features $\{\mathbf{x}_i^{(k)}\}$ and the classifier vectors $\{\mathbf{w}_k\}$ at the terminal phase of training.

- **NC1**, within-class variability collapse. Each feature $\mathbf{x}_i^{(k)}$ collapse to its class center $\bar{\mathbf{x}}^{(k)} = \frac{1}{n_k} \sum_{i'=1}^{n_k} \mathbf{x}_{i'}^{(k)}$, indicating the *maximal intra-class compactness*

- **NC2**, convergence to simplex equiangular tight frame. The set of class centers $\{\bar{\mathbf{x}}^{(k)}\}_{k=1}^K$ form a simplex equiangular tight frame (ETF), with equal and maximized cosine distance between every pair of feature means, i.e., the *maximal inter-class separability*.

- **NC3**, convergence to self-duality. The class center $\bar{\mathbf{x}}^{(k)}$ is ideally aligned with the classifier vector $\mathbf{w}_k, \forall k \in [K]$.

- **NC4**, simplification to nearest class center. The classifier is equivalent to a nearest class center decision.

The theory of neural collapse (Papayan, Han, and Donoho 2020; Zhu et al. 2021; Fang et al. 2021) indicates that the centered classifier vectors form an ETF framework at terminal training, i.e.,

$$\cos(\hat{\mathbf{w}}_k, \hat{\mathbf{w}}_{k'}) = \frac{K}{K-1} \delta_{k,k'} - \frac{1}{K-1}, \quad \text{and} \quad (20)$$

$$\|\mathbf{w}_k\| = \|\mathbf{w}_{k'}\|, \quad \forall k, k', \quad (21)$$

where $\delta_{k,k'}$ denotes the Kronecker delta function, which equals 1 when $k = k'$ and 0 otherwise.

From Fig. 2 and Fig. 6, on CIFAR100-LT one can find the separability of any classifier vectors is close to 50.50 which is approximately equal to $\frac{1+\frac{1}{K-1}}{2} \times 100 = \frac{500}{99}$.

Appendix C: Tripartite synergistic learning

In our BCE3S, we apply symbols ‘‘sc’’, ‘‘ss’’, and ‘‘cc’’ to indicate the joint learning, contrastive learning, and uniform learning, as they measures the metrics of pairs of sample-to-class, sample-to-sample, and class-to-class, respectively. Besides BCE3S, one can also design CE-based tripartite synergistic learning which also consists of three components: CE-based joint learning $L_{\text{ce}}^{(\text{sc})}$, CE-based contrastive learning $L_{\text{ce}}^{(\text{ss})}$, and CE-based uniform learning $L_{\text{ce}}^{(\text{cc})}$.

CE-based joint learning. During training, for any sample $\mathbf{X}^{(k)}$ from the k -th class, the model $\mathcal{M}$ extracts its feature as $\mathbf{x}^{(k)} = \mathcal{M}(\mathbf{X}^{(k)})$, and the classifier $\mathcal{C} = \{(\mathbf{w}_j, b_j)\}_{j=1}^K$ converts the feature into K metrics $\{\mathbf{w}_j^T \mathbf{x} + b_j\}_{j=1}^K$, then Softmax is used to compute class probabilities $\left\{ \frac{\exp(\mathbf{w}_j^T \mathbf{x} + b_j)}{\sum_{\ell=1}^K \exp(\mathbf{w}_\ell^T \mathbf{x} + b_\ell)} \right\}_{j=1}^K$ and cross-entropy is used to compute the loss for joint learning between sample feature and classifier,

$$L_{\text{ce}}^{(\text{sc})}(\mathbf{x}^{(k)}) = -\log \left(\frac{\exp(\mathbf{w}_k^T \mathbf{x}^{(k)} + b_k)}{\sum_{j=1}^K \exp(\mathbf{w}_j^T \mathbf{x}^{(k)} + b_j)} \right), \quad (22)$$

where $\|\mathbf{w}_j\| = 1$ for $\forall j$.

CE-based contrastive learning. Following SimCLR (Chen et al. 2020) and GLMtC (Du et al. 2023), we implement contrastive learning in a projection space. For any sample feature $\mathbf{x}^{(k)}$ from class k , we transform the sample feature into $\mathbf{z}^{(k)} = \mathcal{P}(\mathbf{x}^{(k)}) \in \mathbb{R}^{d'}$ using a non-linear projector $\mathcal{P}$, and

$$L_{\text{ce}}^{(\text{ss})}(\mathbf{x}^{(k)}) = -\log \left(\frac{\exp\left(\frac{1}{\tau} \cos(\mathbf{z}^{(k)}, \mathbf{z}_*^{(k)})\right)}{\sum_{j=1}^K \exp\left(\frac{1}{\tau} \cos(\mathbf{z}^{(j)}, \mathbf{z}_*^{(k)})\right)} \right), \quad (23)$$

where $\{z_*^{(j)} = \mathcal{P}(x_*^{(j)})\}_{j=1}^K$ are projections of K sample features $\{x_*^{(j)}\}_{j=1}^K$ stored in a memory bank, and τ is a temperature factor. Similar to SimCLR (Chen et al. 2020), the non-linear projector $\mathcal{P}$ consists of a two-layer MLP with a non-linear activation function. The CE-based contrastive learning is similar to that in GLMC (Du et al. 2023).

CE-based uniform learning. Softmax and cross-entropy can also be used to design uniform learning for the classifier,

$$L_{ce}^{(sc)}(w_k) = -\log\left(\frac{\exp(w_k^T w_k)}{\sum_{j=1}^K \exp(w_k^T w_j)}\right), \quad (24)$$

where the numerator in the Softmax is natural constant e .

CE3S. During the training, a batch of samples $[X_i^{(k_i)}]_{i=1}^B$ are fed into the model $\mathcal{M}$, and their features $[x_i^{(k_i)}]_{i=1}^B$ are extracted, where B is the batch size. The model training can be driven by the CE3S loss,

$$L_{ce}^{(tri)}([x_i^{(k_i)}]) = \frac{1}{B} \sum_{i=1}^B L_{ce}^{(sc)}(x_i^{(k_i)}) + \frac{\lambda_{ss}}{B} \sum_{i=1}^B L_{ce}^{(ss)}(x_i^{(k_i)}) + \frac{\lambda_{cc}}{K} \sum_{k=1}^K L_{ce}^{(cc)}(w_k). \quad (25)$$

If the models were trained using a two-stage training pipeline, we will take the whole CE3S in the first stage and only fine-tune the classifier vectors by using a class-balanced CE-based joint learning, similar to that in the work of Cui et al. (2019),

$$L_{ce-cb}^{(sc)}([x_i^{(k_i)}]) = \frac{1}{B} \sum_{i=1}^B \frac{1-\beta}{1-\beta^{n_{k_i}}} L_{ce}^{(sc)}(x_i^{(k_i)}), \quad (26)$$

where β is a re-weighting parameter. In our experiments with CE3S or BCE3S, we set $\beta = 0.9999$.

Appendix D: CE vs. BCE in LTR

In the main body of paper, we have in-depth analyzed and compared BCE-based and CE-based joint learning with respect to feature learning. We here compare their contrastive learning and uniform learning in terms of the learning and features and classifier.

Contrastive learning for feature learning. With BCE- and CE-based contrastive learning $L_{bce}^{(ss)}$ and $L_{ce}^{(ss)}$, the sample features are also updated in the projection space through their gradients,

$$z^{(k)} \leftarrow z^{(k)} - \eta \frac{\partial L_{\mu}^{(ss)}(z^{(k)})}{\partial z^{(k)}}, \quad (27)$$

where $\mu \in \{ce, bce\}$ denotes CE- or BCE-based contrastive learning, and η is the learning rate. The gradients $\frac{\partial L_{\mu}^{(ss)}(z^{(k)})}{\partial z^{(k)}}$

have the similar form in CE or BCE,

$$\begin{aligned} \frac{\partial L_{\mu}^{(ss)}(z^{(k)})}{\partial z^{(k)}} = & -\underbrace{\frac{1}{\tau} (1 - \text{Act}_{\mu}(\cos(z_*^{(k)}, z^{(k)})))}_{\text{pulling}} z_*^{(k)} \\ & + \sum_{\substack{j=1 \\ j \neq k}}^K \underbrace{\frac{1}{\tau} \text{Act}_{\mu}(\cos(z_*^{(j)}, z^{(k)}))}_{\text{repelling}} z_*^{(j)}, \end{aligned} \quad (28)$$

where, for $\forall j$,

$$\text{Act}_{ce}(\cos(z_*^{(j)}, z^{(k)})) = \frac{\exp(\cos(z_*^{(j)}, z^{(k)}))}{\sum_{\ell=1}^K \exp(\cos(z_*^{(\ell)}, z^{(k)}))}, \quad (29)$$

$$\text{Act}_{bce}(\cos(z_*^{(j)}, z^{(k)})) = \frac{1}{1 + e^{-\cos(z_*^{(j)}, z^{(k)})}}. \quad (30)$$

In the memory bank, the projected feature $z_*^{(k)}$ is updated by the last feature $x_i^{(k_i=k)}$ from the same class in the on-line batch; if there is no data from the same class k , the projected feature $z_*^{(k)}$ is not updated in this iteration using the batch. As the probability of samples from different classes in the long-tailed dataset appearing in the batch is different, so the learning degree of the projected features $\{z_*^{(k)}\}_{k=1}^K$ in the memory bank is not the same, resulting in strong imbalance. The imbalanced K projected features are coupled in the denominator of Softmax Act_{ce} by CE-based contrastive learning, while they are decoupled via Sigmoid Act_{bce} by BCE-based contrastive learning. Therefore, $L_{bce}^{(ss)}$ will achieve more balance features than $L_{ce}^{(ss)}$.

Uniform learning for classifier vectors. In the training with the uniform learning, $L_{ce}^{(cc)}$ and $L_{bce}^{(cc)}$, the classifier vectors are updated via their gradients,

$$w_k \leftarrow w_k - \eta \frac{\partial L_{\mu}^{(cc)}(w_k)}{\partial w_k}, \quad \forall k, \quad (31)$$

where $\mu \in \{ce, bce\}$ denotes CE- or BCE-based uniform learning, and η is the learning rate. The gradients $\frac{\partial L_{\mu}^{(cc)}(w_k)}{\partial w_k}$ have the similar form in CE or BCE,

$$-\underbrace{(1 - \text{Act}_{\mu}(w_k^T w_k))}_{\text{pulling}} w_k + \sum_{\substack{j=1 \\ j \neq k}}^K \underbrace{\text{Act}_{\mu}(w_j^T w_k) w_j}_{\text{repelling}}, \quad (32)$$

where, for $\forall j$,

$$\text{Act}_{ce}(w_j^T w_k) = \frac{\exp(w_j^T w_k)}{\sum_{\ell=1}^K \exp(w_{\ell}^T w_k)}, \quad (33)$$

$$\text{Act}_{bce}(w_j^T w_k) = \frac{1}{1 + e^{-w_j^T w_k}}. \quad (34)$$

In updating classifier vectors, both $L_{ce}^{(cc)}$ and $L_{bce}^{(cc)}$ utilize exponential similarities between a classifier vector and all K classifier vectors to compute one pulling term and $K-1$ repelling terms. For CE-based uniform learning ($L_{ce}^{(cc)}$), the

computation of each pulling or repelling term involves K imbalanced exponential similarities coupled in the Softmax denominator (Act_{ce}), thereby injecting the imbalance effect into the uniform learning process. In contrast, BCE decouples the metrics between the K classifier vectors, employing only a single similarity to compute each pulling or repelling term through the Sigmoid function Act_{bce} . Consequently, $L_{\text{bce}}^{(\text{cc})}$ achieves more uniform separability across the K classifier vectors.

Appendix E: Datasets and training details

We evaluate the proposed BCE-based tripartite synergistic learning (BCE3S) on four popular LTR datasets, including CIFAR10-LT (Yue et al. 2016), CIFAR100-LT (Yue et al. 2016), ImageNet-LT (Liu et al. 2022a), and iNaturalist2018 (Van Horn et al. 2018).

CIFAR10-LT and CIFAR100-LT. Following the work of Yue et al. (2016), the long-tailed CIFAR10 and CIFAR100 are produced by sampling the training samples of the original datasets, using an exponential decay imbalance mode across classes, and the long-tailed datasets are referred to as CIFAR10-LT and CIFAR100-LT, respectively. For each of them, we produced their three different variants using three different imbalance factors (IF), 100, 50, and 10.

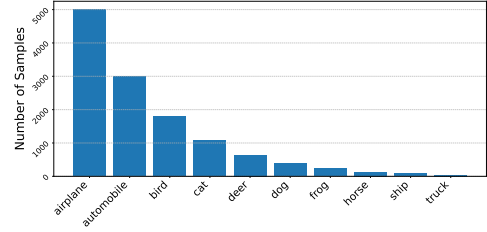
Fig. 4a illustrates the sample numbers of the ten classes in training set of CIFAR10-LT with IF = 100. The dataset comprises samples from 10 classes, i.e., airplane, automobile, bird, cat, deer, dog, frog, horse, ship, and truck. The first head class (airplane) contains 5,000 samples, while the last tail class (truck) has only 50 samples. The smallest class of CIFAR10-LT contains 50 samples at least, thereby no few subset exists.

Fig. 4b illustrates the distribution of the sample numbers of the 100 classes in the training set of CIFAR100-LT with IF of 100. The dataset comprises 100 distinct classes, with sample numbers ranging from 500 to merely 5, following a smooth exponential decay pattern across all classes.

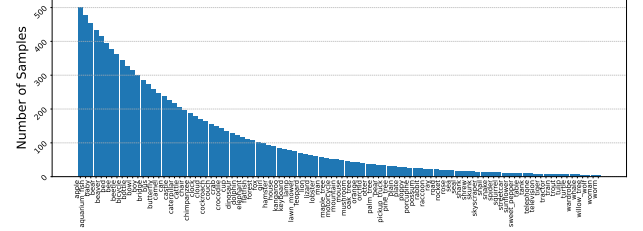
ImageNet-LT (Liu et al. 2022a) is produced by sampling a subset of ImageNet (Russakovsky et al. 2015) using the Pareto distribution with a power factor of $\alpha_p = 6$, which comprises 115,846 images from 1,000 classes and the sample number from the head class to the tail class decreases from 1,280 to 5, resulting in IF = 256. Fig. 4c illustrates the sample distribution in different classes in its training set; we sort the classes in descending order according to the sample number.

iNaturalist18. This large-scale dataset, collected from the iNaturalist species identification platform, encompasses 437,513 images across 8,142 species classes, which exhibits a natural long-tailed distribution, with sample numbers ranging from 1,000 for the most common species to merely 2 for the rarest ones, resulting in an IF of 500. This extreme class imbalance authentically reflects real-world bio-diversity observation patterns, where certain species are frequently encountered while others are exceptionally rare. Fig. 4d illustrates the sample distribution in its training set.

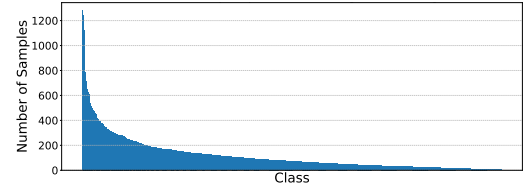
Training Details. We take ResNets as the classification models, and train them on each of the imbalanced datasets



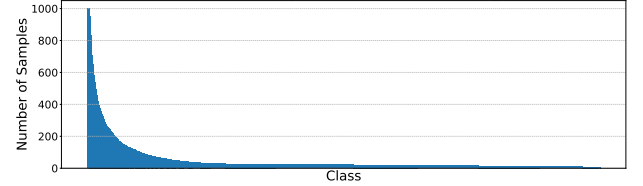
(a) Sample distribution in the training set of CIFAR10-LT with IF = 100.



(b) Sample distribution in the training set of CIFAR100-LT with IF = 100.



(c) Sample distribution in the training set of ImageNet-LT.



(d) Sample distribution in the training set of iNaturalist2018.

Figure 4: Dataset descriptions for different long-tailed datasets.

from scratch. For each training, we employ SGD optimizer with a momentum of 0.9 and decay learning rate with cosine scheduler. For the training on CIFAR10-LT and CIFAR100-LT, we adopt a customized ResNet32, and set the initial learning rate as 0.01 and weight decay rate as $5e-3$, which are the same with that used in MaxNorm (Alshammari et al. 2022). For a fair comparison, we try to reproduce the results in MaxNorm (Alshammari et al. 2022), and we got the recognition accuracy reported in the paper using the official code by training the ResNet32 with 320 epochs. Therefore, we evaluate our BCE3S with a training of 320 epochs.

On ImageNet-LT, following DSCL (Xuan and Zhang 2024), we train ResNet50 and ResNeXt50 for 200 epochs with the initial learning rate of 0.05, weight decay rate of $5e-4$, and batch size of 128. On iNaturalist2018, following ProCo (Du et al. 2024), we train ResNet50 for 180, 400

epochs, respectively, with initial learning rate of 0.1, weight decay of $1e-4$, and batch size of 256.

For each model trained on the imbalanced training set, it is evaluated on the corresponding balanced test set. Following (Kang et al. 2020; Alshammari et al. 2022; Liu et al. 2022b; Du et al. 2024), besides the total recognition accuracy on the whole test set, we also report the accuracy on three kinds of subsets, i.e., Many, Medium, and Few, while the numbers of corresponding training samples in each class are larger than 100, varying from 20 to 100, and less than 20, respectively.

BCE3S with re-balancing methods. Similar to previous approaches (Kang et al. 2020; Alshammari et al. 2022; Xuan and Zhang 2024), BCE3S’s performance in long-tailed recognition can be further improved through a two-stage training strategy. In the first stage, the model is trained using the complete BCE3S method. In the second stage, we fix the feature extractor parameters and only fine-tune the classifier using a class-balanced binary cross-entropy loss (Cui et al. 2019), defined as:

$$L_{\text{bce-cb}}^{(\text{sc})}([\mathbf{x}_i^{(k_i)}]) = \frac{1}{B} \sum_{i=1}^B \frac{1 - \beta}{1 - \beta^{n_{k_i}}} L_{\text{bce}}^{(\text{sc})}(\mathbf{x}_i^{(k_i)}), \quad (35)$$

where β is a hyperparameter that controls the degree of re-weighting based on class frequency. BCE3S’s performance can be boosted by existing re-balancing methods, which demonstrates the compatibility and potential of BCE3S.

Appendix F: Normalized sample feature and classifier

Similar to the previous work (Kang et al. 2020), we adopt normalized classifier in the BCE-based joint learning $L_{\text{bce}}^{(\text{sc})}$. We here compare the different normalization modes on the sample feature or classifier vectors. We train ResNet32 using $L_{\text{bce}}^{(\text{sc})}$ on CIFAR100-LT with IF = 100. Table 6 shows the results. One can find that if normalization is not applied or applied simultaneously on sample features and classifiers, the final accuracy of the model is similar, 49.89% and 49.93%, respectively. When the normalization is only applied on the sample features, the model training fails. If normalization is applied on the classifier, the model accuracy reaches a maximum of 51.56%. Therefore, in BCE3S, we only normalize the classifier vector for the joint learning.

Normalization		Many	Med.	Few	All
feature	Classifier				
		81.20	52.51	10.30	49.89
✓		2.86	0.00	0.00	1.00
	✓	81.83	52.20	15.50	51.56
✓	✓	79.60	49.11	16.27	49.93

Table 6: Comparison of normalized sample features and classifier vectors. All experiments employ BCE-based joint learning, $L_{\text{bce}}^{(\text{sc})}$ with $r = 1.0$ on CIFAR100-LT dataset with IF = 100.

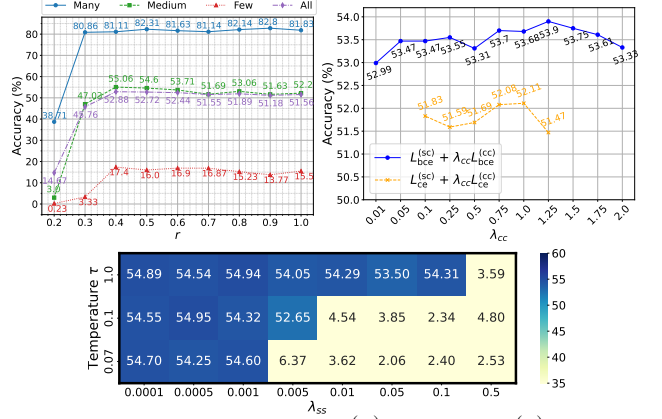


Figure 5: Parameter study for $L_{\text{bce}}^{(\text{sc})}$ (top left), $L_{\text{bce}}^{(\text{ss})}$ (bottom) and $L_{\text{bce}}^{(\text{cc})}$ (top right), on CIFAR100-LT with IF = 100.

Appendix G: More results for ablation experiments

Parameter study. BCE3S contains various hyperparameters in its three BCE components, and we conduct the parameter study on CIFAR100-LT with IF = 100. For the re-sampling parameter r in Eq. (3), we train ResNet32 using only the BCE joint learning $L_{\text{bce}}^{(\text{sc})}$ with different r , while the results are shown in the top left of Fig 5. As r decreases from 1.0 to 0.4, the accuracy on the Few and Medium subsets increase from 15.5% and 52.2% to 17.4% and 55.06%, respectively; though the accuracy on the Many subset decreases, the total accuracy on the whole set has increased from 51.56% to 52.88%. When r is below 0.4, the accuracy drops significantly. These results show that the hyper-parameter r helps to balance the performances on the different classes in the LTR tasks.

For λ_{ss} and temperature τ in the contrastive learning $L_{\text{bce}}^{(\text{ss})}$, we perform their parameter studies together using $L_{\text{bce}}^{(\text{sc})} + \lambda_{\text{ss}} L_{\text{bce}}^{(\text{ss})}$. When increasing λ_{ss} or decreasing τ , the contrastive learning is enhancing, which amplifies the imbalance effect of the long-tailed datasets and easily leads to training failure, as the bottom of Fig. 5 shows.

For the uniform learning $L_{\text{bce}}^{(\text{cc})}$, a too small λ_{cc} fails to enhance the classifier separability, while a too large one would optimize the classifier too quickly, making it difficult to be further tuned with the updating features, ultimately leading to a decreased LTR performance. As the top right of Fig. 5 shows, the best LTR corresponds to $\lambda_{\text{cc}} = 1.25$ when using $L_{\text{bce}}^{(\text{sc})} + \lambda_{\text{cc}} L_{\text{bce}}^{(\text{cc})}$, significantly higher than that of CE ones.

Separability and compactness. In the ablation experiments, besides the CE and BCE-based tripartite synergistic learning, we also evaluate the combination of CE-based joint learning $L_{\text{ce}}^{(\text{sc})}$ with BCE-based contrastive learning $L_{\text{bce}}^{(\text{ss})}$ and uniform learning $L_{\text{bce}}^{(\text{cc})}$. The LTR accuracy of various learning methods on CIFAR100-LT with IF = 100 have been presented in Table 2, and we here further compare the classifier separability and features’ compactness and separability for all the models.

Fig. 6 shows the separability and compactness of the different models on the training set. We have compared the

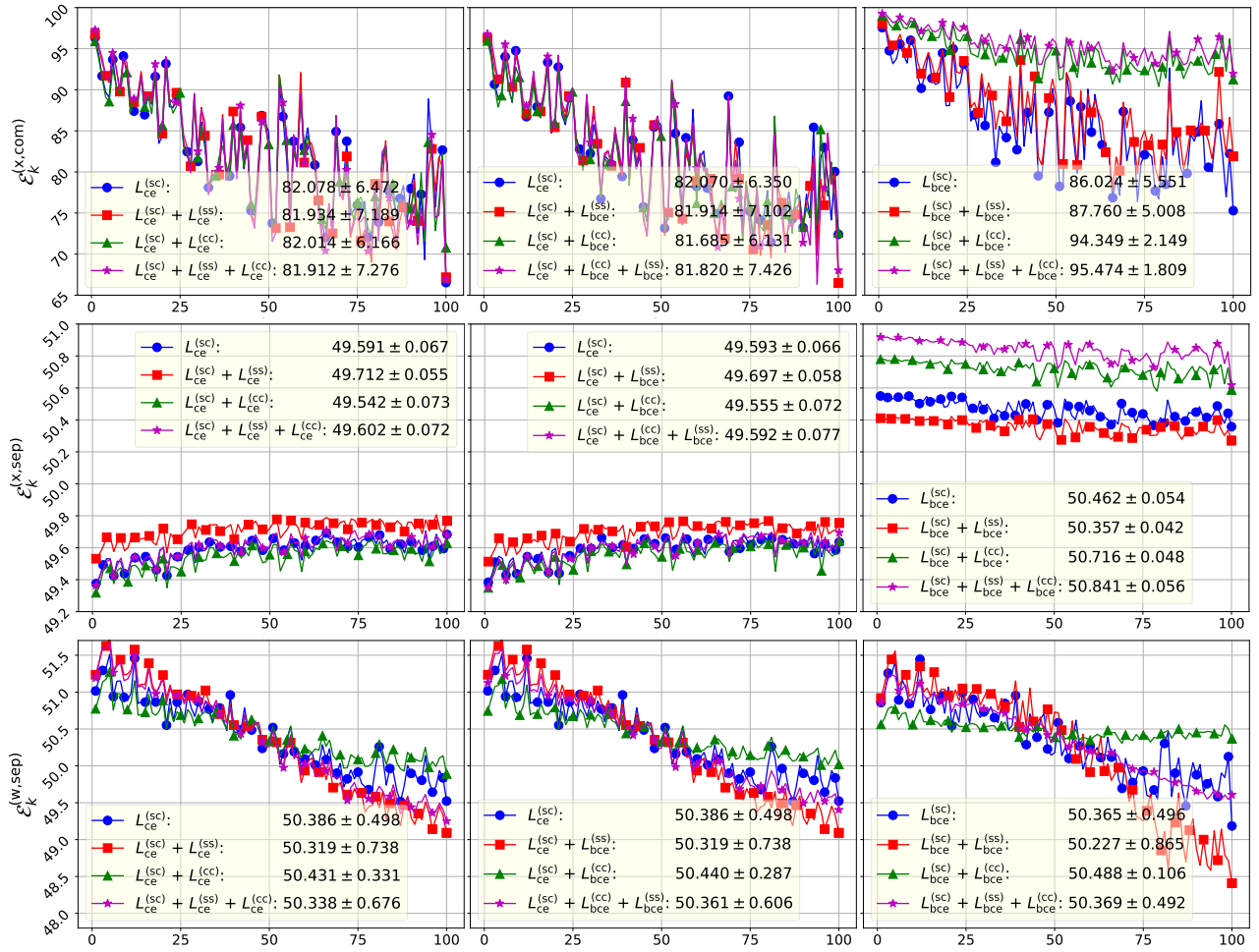


Figure 6: The intra-class compactness (top), inter-class separability (middle) of sample features, and separability (bottom) of classifier vectors on the training set of CIFAR100-LT (IF = 100), with ResNet32 trained using the methods of CE-based learning (left), CE with BCE-based learning (middle), and BCE-based learning (right).

results of CE- and BCE-based tripartite synergistic learning in the experimental section of the main paper. We find and explain that both CE and BCE-based joint learning, $L_{ce}^{(sc)}$ and $L_{bce}^{(sc)}$, will learn imbalanced classifiers, but their learned features have significant differences. Because BCE’s joint learning $L_{bce}^{(sc)}$ decouples the metric between any sample feature with K classifier vectors $\{\mathbf{w}_k\}_{k=1}^K$, it avoids the secondary injection of classifier’s imbalance effects into the feature learning. Therefore, compared to $L_{ce}^{(sc)}$, $L_{bce}^{(sc)}$ learns better features with higher compactness and separability. Meanwhile, BCE-based contrastive learning $L_{bce}^{(ss)}$ can enhance the compactness of features, and BCE-based uniform learning $L_{bce}^{(cc)}$ effectively balances the separability of classifiers. However, the two CE-based learning modes, $L_{ce}^{(ss)}$ and $L_{ce}^{(cc)}$, have no significant effect in this regard.

Here, we further found that although BCE-based contrastive learning $L_{bce}^{(ss)}$ and uniform learning $L_{bce}^{(cc)}$ can improve the properties of features and classifiers, they do not show these advantages when combined with CE-based joint learn-

ing $L_{ce}^{(sc)}$. As the left and middle columns in Fig. 6 show, when $L_{ce}^{(ss)}$ and $L_{ce}^{(cc)}$ are combined with $L_{ce}^{(sc)}$, the final classifier separability and features’ compactness and separability are almost identical to that obtained by combining $L_{ce}^{(ss)}$, $L_{ce}^{(cc)}$ with $L_{ce}^{(sc)}$. These results suggest that CE-based joint learning has inherent limitations in LTR tasks, which is difficult to overcome completely by external re-balancing techniques. We believe that this defect of CE comes from its Softmax function, which couples K metrics related to the imbalanced classifiers on its denominator.

Fig. 7 shows the features’ compactness and diversity of models trained with different learning methods on the test set of CIFAR100-LT with IF = 100. We draw similar conclusions from these results as we did on the training set: BCE-based joint learning $L_{bce}^{(sc)}$ can learn better sample features than CE-based one $L_{ce}^{(sc)}$, and its contrastive learning $L_{bce}^{(ss)}$ and uniform learning $L_{bce}^{(cc)}$ can further effectively improve the feature compactness, but they cannot improve the feature properties of CE-based joint learning $L_{ce}^{(sc)}$. These re-

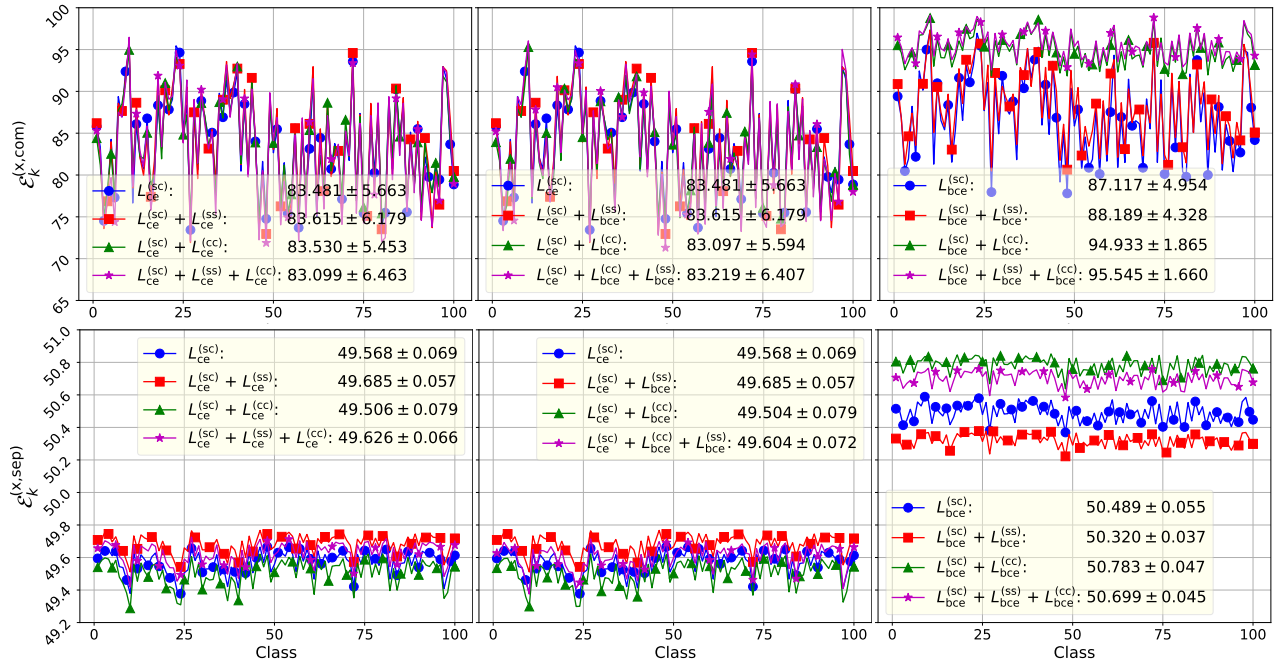


Figure 7: The intra-class compactness (top) and inter-class separability (bottom) of features on the test set of CIFAR100-LT (IF = 100), with ResNet32 trained using the methods of CE-based learning (left), CE with BCE-based learning (middle), and BCE-based learning (right).

sults also reflect that BCE3S has well generalization.

Separability matrix of classifiers. We also compute the separability matrix $S = [s_{jk}]$ for the classifier, where

$$s_{jk} = \begin{cases} 1 & j = k, \\ \frac{1 - \cos(\hat{w}_j, \hat{w}_k)}{2} & j \neq k. \end{cases} \quad (36)$$

Fig. 8 shows the separability matrix of classifiers in models trained by different CE and BCE-based methods on on CIFAR100-LT with IF = 100. In each matrix diagram, the upper left area reflects the classifier separability among head classes, the lower left and upper right areas reflect the classifier separability between head and tail classes, and the lower right area reflects the classifier separability among tail classes. Except for the elements on the diagonal, the brighter points indicate the higher separability between the two classifier vectors, while the darker points indicate the lower separability. From the first column of the figure, one can find that in each subfigure, the lower left and upper right areas are the brightest, followed by the upper left area, while the lower right area has more obvious black points, indicating that the BCE and CE-based joint learning, $L_{ce}^{(sc)}$ and $L_{bce}^{(sc)}$, achieve the highest classifier separability between the head and tail classes, followed by the classifier separability among the head classes, while the classifier separability among the tail classes is the worst.

From the second column in Fig. 8, one can find that when the two kind of contrastive learning are respectively added, the imbalance of the classifier separability of the models are further exacerbated, with the lower left and upper right area of each subfigure appearing brighter and the lower right area

appearing darker. The third column shows the separability matrix of the classifier trained by the combination of joint learning and uniform learning. One can find that the combination of $L_{bce}^{(sc)}$ and $L_{bce}^{(cc)}$ effectively improves the separability of the classifier, and the color of the entire separability matrix diagram is more balanced among its different parts. However, combined with CE-based joint learning $L_{ce}^{(sc)}$, neither BCE-based uniform learning $L_{bce}^{(cc)}$ nor CE-based one $L_{ce}^{(cc)}$ effectively balances the classifiers, and the colors of the separability matrix on the two diagonal directions still exhibit significant differences. In the fourth column, when contrastive learning is added to the model training again, the color of the separability matrix corresponding to CE and BCE in different parts becomes more uneven or re-uneven, indicating an imbalanced classifier.

Feature distribution in t-SNE. To further compare the features of CE and BCE-based learning modes, we trained models on CIFAR10-LT with IF of 100, then, using t-SNE, we visually show their feature distributions of all 10 classes on the test set, as Fig. 9 shows. In the figure, the first row shows the features of model trained by various learning combination of three CE-based learning modes, i.e., $L_{ce}^{(sc)}$, $L_{ce}^{(ss)}$, and $L_{ce}^{(cc)}$; the second row shows that of CE-based joint learning with BCE-based contrastive learning and uniform learning, i.e., $L_{ce}^{(sc)}$ with $L_{bce}^{(ss)}$, $L_{bce}^{(cc)}$; the third row shows that of BCE-based three learning modes, i.e., $L_{bce}^{(sc)}$, $L_{bce}^{(ss)}$, and $L_{bce}^{(cc)}$. Similarly, the sample features learned by BCE-based joint learning $L_{bce}^{(sc)}$ exhibit better compactness and separability than that learned by CE one $L_{ce}^{(sc)}$, and $L_{bce}^{(ss)}$ and $L_{bce}^{(cc)}$.

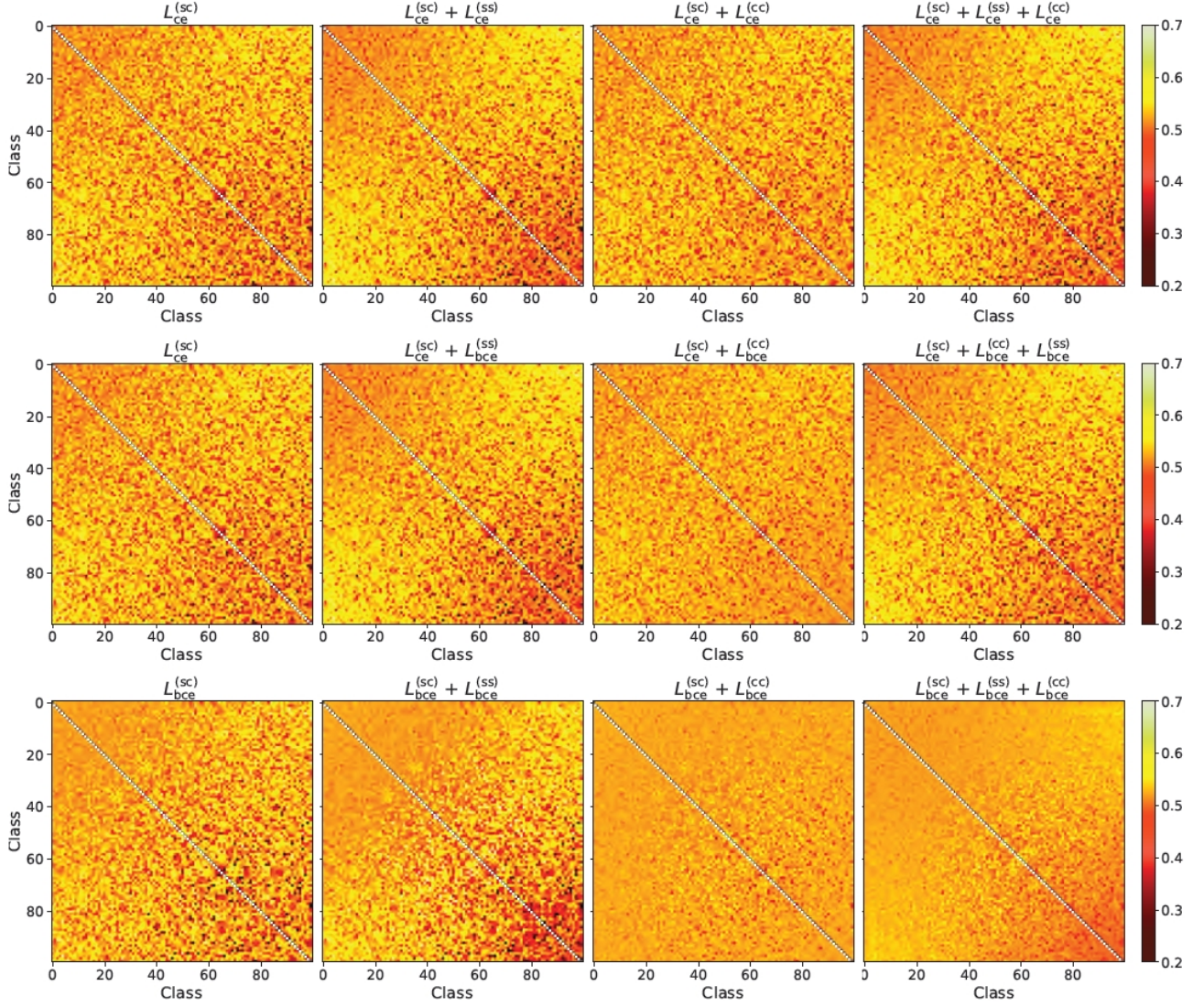


Figure 8: The visualizations of similarity matrices of classifier trained by different learning configurations: CE-based learning (top), CE with BCE-based learning (middle), and BCE-based learning (bottom).

can further effectively improve the properties of features. In contrast, the compactness and separability of sample features learned by CE-based joint learning $L_{ce}^{(sc)}$ are poor, and their corresponding contrastive learning $L_{ce}^{(ss)}$ and uniform learning $L_{ce}^{(cc)}$ cannot effectively improve the feature properties, which were only slightly improved by BCE-based contrastive learning $L_{bce}^{(ss)}$ and uniform learning $L_{bce}^{(cc)}$.

Appendix H: More results about Comparison BCE with SOTA

On **ImageNet-LT**, we apply BCE3S to train ResNet50 and ResNeXt50 by adopting a similar strategy used in DSCL (Xuan and Zhang 2024). We compare BCE3S with various state-of-the-art (SOTA) methods, including CB-Focal (Cui et al. 2019), RISDA (Chen et al. 2022), LDAM (Cao et al. 2019), KCL (Kang et al. 2021),

TSC (Li et al. 2022), GCL (Li, Cheung, and Lu 2022), GCL+H2T (Li et al. 2024), BCL (Zhu et al. 2022), DiffuLT (Shao et al. 2024), BCL+DODA (Wang et al. 2024), DiffuLT+RIDE (Shao et al. 2024), DSCL (Xuan and Zhang 2024), ProCo (Du et al. 2024), ETF+DR (Yang et al. 2022)+DR (Kang et al. 2020), FCL (Kang et al. 2021), RBL (Peifeng et al. 2023), NC-DRW (Liu et al. 2023), NC-cRT (Liu et al. 2023), Meta Softmax (Ren et al. 2020), PaCo (Cui et al. 2021), GLMC (Du et al. 2023), SSD (Li, Wang, and Wu 2021), GLMC+MN (Alshammari et al. 2022), and GLMC+MS (Ren et al. 2020). We have presented the LTR results of ResNet50 in the main body of the paper, and we here present the results of ResNeXt50.

Table 7 provides a detailed comparison of these methods using ResNeXt50 on ImageNet-LT. Although BCE3S achieved only one optimal result on the three test subsets

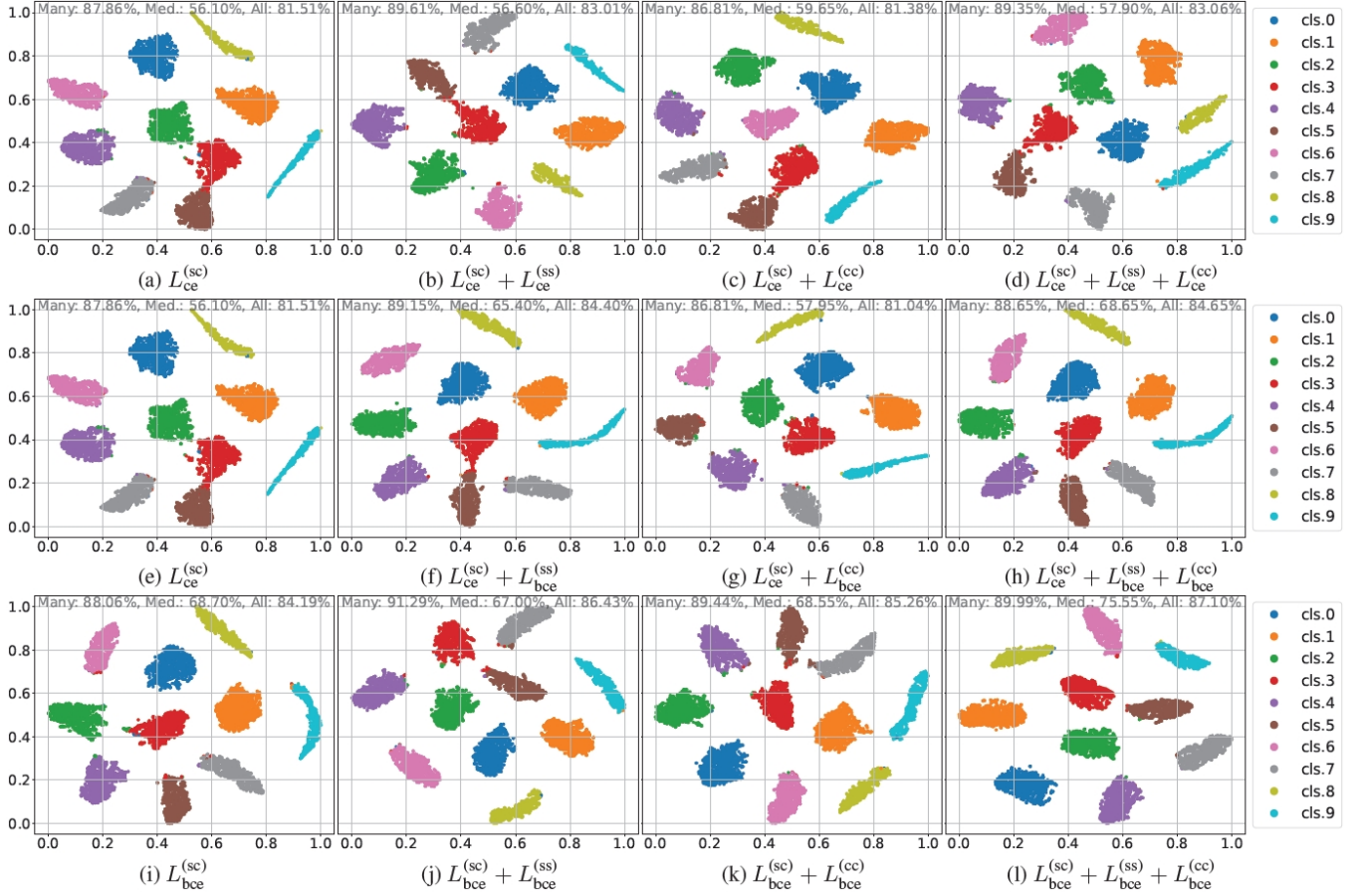


Figure 9: Feature distribution on the test set of CIFAR10-LT (IF = 100) with different learning configurations: CE-based learning (top), CE with BCE-based learning (middle), and BCE-based learning (bottom).

(Many, Medium, and Few), it outperformed all competing methods in terms of overall accuracy on the whole test set. Specifically, BCE3S achieved 58.54% with ResNeXt50, surpassing prior SOTA results like GLMC (Du et al. 2023), and ProCo (Du et al. 2024).

Appendix I: Further analysis on ImageNet-LT

Compactness and separability. Fig. 10 presents three metrics for ResNet50 trained on ImageNet-LT across all 1000 classes. For intra-class feature compactness, BCE3S demonstrates superior performance with an average of 96.889 and standard deviation of 0.790, compared to CE3S which achieves 85.868 with 2.172 standard deviation. This indicates BCE3S produces significantly more uniform compactness across both head and tail classes. Similarly, for inter-class feature separability, BCE3S exhibits enhanced uniformity across the class distribution, with 50.144 compared CE’s 50.099. Regarding classifier separability, despite the increased complexity from the large number of classes, BCE3S still outperforms CE3S consistently. These results provide compelling evidence that BCE effectively addresses the intrinsic limitations of CE in long-tailed recognition tasks, particularly in maintaining balanced feature represen-

Methods	$\mathcal{M}$	ImageNet-LT			
		Many	Med.	Few	All
ETF+DRNeurIPS’22	RX50	-	-	-	44.70
FCLICLR’21		61.40	47.00	28.20	49.80
KCLICLR’21		62.40	49.00	29.50	51.50
TSCVPR’22		63.50	49.70	30.40	52.40
RBLICML’23		64.80	49.60	34.20	53.30
NC-DRWAISTATS’23		67.10	49.70	29.00	53.60
NC-CRTAISTATS’23		65.60	51.20	35.40	54.20
Meta Soft.NeurIPS’20		65.80	53.20	34.10	55.40
PaCoICCV’21		64.40	<u>55.70</u>	33.70	56.00
GLMCCVPR’23		70.10	52.40	30.40	56.30
SSDICCV’21		66.80	53.10	35.40	56.60
GLMC+MNCVPR’22		60.80	55.90	45.50	56.70
BCLCVPR’22		67.90	54.20	36.60	57.10
GLMC + MS NeurIPS’20		64.76	55.67	<u>42.19</u>	57.21
ProCOTPAMI’24		-	-	-	58.00
BCE3S		71.13	54.98	36.78	58.54

Table 7: Benchmark results of top-1 accuracy (%) on ImageNet-LT. ResNeXt50 (RX50), trained by BCE3S, achieves the best results on the test set.

tations across the entire class distribution.

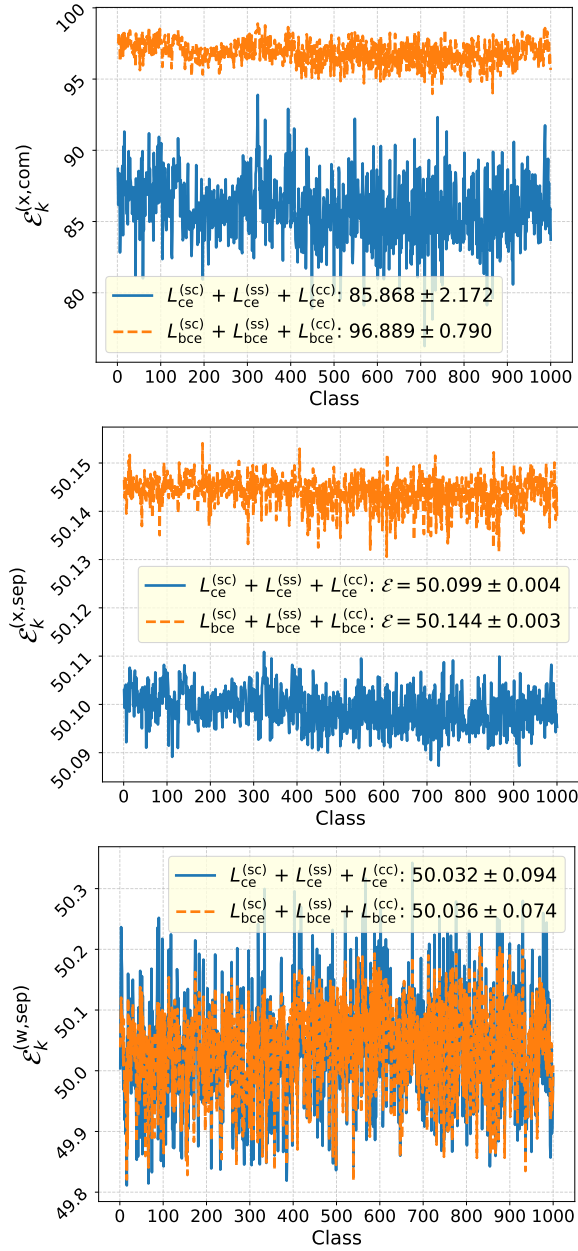


Figure 10: The intra-class compactness (left), inter-class separability (middle) of features and inter-class separability of classifier on the train set of ImageNet-LT with ResNet50.