

Table A: Test accuracy (%) on four generic long-tailed image recognition datasets.

Methods	CIFAR-10-LT			CIFAR-100-LT			ImageNet-100-LT			Places-LT		
	Old	New	All	Old	New	All	Old	New	All	Old	New	All
ABC [†]	77.7	20.2	48.9	48.5	20.8	43.0	-	-	-	-	-	-
DARP [†]	77.6	35.2	56.4	54.5	24.0	48.4	-	-	-	-	-	-
OpenLDN [†]	74.8	49.0	61.9	50.4	32.9	46.9	-	-	-	-	-	-
TRSSL [†]	78.4	66.8	72.6	58.7	35.8	54.1	-	-	-	-	-	-
ORCA [†]	77.5	55.6	66.6	55.0	30.8	50.1	69.3	29.0	49.2	21.5	6.9	14.2
SimGCD	75.1	41.5	58.3	59.8	24.6	52.8	81.1	33.4	57.8	31.4	18.2	24.8
GCD	78.5	<u>71.7</u>	<u>75.1</u>	<u>65.5</u>	<u>49.0</u>	<u>62.2</u>	81.0	76.8	78.9	29.8	<u>22.7</u>	<u>26.2</u>
OpenCon [†]	87.2	47.2	67.2	64.2	40.9	59.6	<u>83.3</u>	42.8	63.0	30.6	12.4	21.6
BaCon-O	83.3	78.0	80.7	66.5	69.6	67.1	84.2	80.0	82.1	30.7	25.6	28.1
BaCon-S	85.6	78.7	82.2	66.8	71.0	67.6	83.9	80.6	82.3	31.1	28.4	29.9

Table B: Test accuracy (%) and balancedness (Std↓) on CIFAR-100-LT.

Methods	Known					Novel					Overall
	Many↑	Med↑	Few↑	Std↓	All↑	Many↑	Med↑	Few↑	Std↓	All↑	
ABC [†]	77.4	51.6	16.3	25.0	48.5	20.8	32.0	5.9	10.7	20.8	43.0
DARP [†]	74.8	55.2	33.4	16.9	54.5	29.9	30.1	10.0	9.4	24.0	48.4
OpenLDN [†]	77.4	49.1	24.8	21.5	50.4	<u>57.3</u>	24.8	14.4	17.9	32.9	46.9
TRSSL [†]	78.8	73.0	24.1	24.3	58.7	32.7	50.9	18.8	13.1	35.8	54.1
ORCA [†]	73.6	<u>60.0</u>	31.0	17.8	55.0	47.5	28.4	17.2	12.5	30.8	50.1
SimGCD	71.1	72.8	34.6	17.6	59.8	32.9	25.8	15.5	7.1	24.6	52.8
GCD	75.9	69.4	<u>52.9</u>	<u>9.7</u>	<u>65.5</u>	41.9	<u>54.2</u>	<u>49.8</u>	<u>5.1</u>	<u>49.0</u>	<u>62.2</u>
OpenCon [†]	86.8	69.7	35.8	21.2	64.2	55.8	48.9	15.3	17.7	40.9	59.6
BaCon-O	72.9	64.7	62.0	4.6	66.5	72.6	70.6	65.3	3.1	69.6	67.1
BaCon-S	73.5	65.6	61.2	5.1	66.8	74.7	73.1	64.5	4.5	71.0	67.6

Table C: Test accuracy (%) and balancedness (Std↓) on CIFAR-100-LT.

Methods	Known					Novel					Overall
	Many↑	Med↑	Few↑	Std↓	All↑	Many↑	Med↑	Few↑	Std↓	All↑	
SimGCD	71.1	72.8	34.6	17.6	59.8	32.9	25.8	15.5	7.1	24.6	52.8
+BCL	73.3	68.5	41.9	13.8	61.4	34.9	28.2	19.1	6.5	27.5	54.6
GCD	75.9	69.4	52.9	9.7	65.5	<u>41.9</u>	54.2	49.8	<u>5.1</u>	<u>49.0</u>	62.2
+BCL	76.6	67.4	<u>55.7</u>	<u>8.5</u>	<u>66.6</u>	39.2	<u>55.1</u>	<u>49.8</u>	6.6	48.7	<u>63.3</u>
BaCon-O	72.9	64.7	62.0	4.6	66.5	72.6	70.6	65.3	3.1	69.6	67.1
BaCon-S	73.5	65.6	61.2	5.1	66.8	74.7	73.1	64.5	4.5	71.0	67.6

Table D: Accuracy (%) and balancedness (Std↓) of the pseudo-labeling branch on CIFAR-10.

Epoch	Branch	Known					Novel					Overall
		Many↑	Med↑	Few↑	Std↓	All↑	Many↑	Med↑	Few↑	Std↓	All↑	
50	Pseudo-labeling	85.3	93.1	85.1	3.7	88.3	65.0	69.6	51.1	7.8	61.3	74.8
	Contrastive-learning	96.0	84.3	73.2	9.3	82.2	69.7	73.6	66.3	2.9	69.9	76.0
100	Pseudo-labeling	86.9	92.2	88.1	2.2	89.5	61.5	68.9	53.4	6.4	61.2	75.3
	Contrastive-learning	89.3	92.2	83.7	3.6	88.2	60.4	64.9	59.9	2.2	62.0	75.1

Table E: Accuracy of SSL methods with contrastive learning.

Methods	CIFAR-10-LT			CIFAR-100-LT		
	Old	New	All	Old	New	All
Opencos [†]	71.3	14.2	42.8	57.1	<u>19.7</u>	49.6
Simmatch [†]	<u>78.6</u>	<u>16.7</u>	<u>47.7</u>	<u>61.7</u>	10.4	<u>51.5</u>
BaCon-S	85.6	78.7	82.2	66.8	71.0	67.6

Table F: Accuracy on Herbarium.

Methods	Old	New	All
ORCA	14.7	4.9	9.8
GCD	41.4	24.2	32.8
SimGCD	51.4	<u>27.3</u>	<u>39.4</u>
BaCon-S	47.0	38.6	42.8

Table G: Test accuracy (%) on different novel classes.

Dataset	Old	New	All
Known+Novel_val	83.9	80.6	82.3
Known+Novel_test_A	80.4	85.8	83.1
Known+Novel_test_B	83.3	80.2	81.8
Known+Novel_test_C	82.1	81.2	81.7
Known+Novel_test_D	84.6	79.5	82.0
Known+Novel_test_E	82.2	82.7	82.4
Average (Acc/Std)	82.5 / 1.4	81.9 / 2.2	82.4 / 0.5