

ISD: Self-Supervised Learning by Iterative Similarity Distillation

Ajinkya Tejankar^{1,*} Soroush Abbasi Koohpayegani^{1,*} Vipin Pillai¹ Paolo Favaro² Hamed Pirsiavash^{1,3}

¹University of Maryland, Baltimore County ²University of Bern ³University of California, Davis

Abstract

Recently, contrastive learning has achieved great results in self-supervised learning, where the main idea is to pull two augmentations of an image (positive pairs) closer compared to other random images (negative pairs). We argue that not all negative images are equally negative. Hence, we introduce a self-supervised learning algorithm where we use a soft similarity for the negative images rather than a binary distinction between positive and negative pairs. We iteratively distill a slowly evolving teacher model to the student model by capturing the similarity of a query image to some random images and transferring that knowledge to the student. Specifically, our method should handle unbalanced and unlabeled data better than existing contrastive learning methods, because the randomly chosen negative set might include many samples that are semantically similar to the query image. In this case, our method labels them as highly similar while standard contrastive methods label them as negatives. Our method achieves comparable results to the state-of-the-art models. Our code is available here: <https://github.com/UMBCvision/ISD>.

1. Introduction

We can view the recent crop of SSL methods as iterative self-distillation where there is a teacher and a student. Both teacher and student improve simultaneously while the teacher is evolving more slowly (running average) compared to the student: (1) In the case of contrastive methods e.g. MoCo [20], we classify images to positive and negative pairs in the binary form. (2) In the case of clustering methods (DC-v2 [8], SwAV [8], SeLA [50]), the student predicts the quantized representations from the teacher. (3) In the case of BYOL [18], the student simply regresses the teacher’s embeddings vector. Here, we introduce a novel method using similarity based distillation to transfer the knowledge from the teacher to the student. We argue that our method is more regularized compared to prior work and

improves the quality of the features in transfer learning.

In the standard contrastive setting, e.g., MoCo [20], there is a binary distinction between positive and negative pairs, but in practice, many negative pairs may be from the same category as the positive one. Thus, forcing the model to classify them as negative is misleading. This can be more important when the unlabeled training data is unbalanced, for example, when a large portion of images are from a small number of categories. Such scenario can happen in applications like self-driving cars, where most of the data is just repetitive data captured from a high-way scene with a couple of cars in it. In such cases, the standard contrastive learning methods will try to learn features to distinguish two instances of the most frequent category that are in a negative pair, which may not be helpful for the down-stream task of understanding rare cases.

We are interested in relaxing the binary classification of contrastive learning with soft labeling, where the teacher network calculates the similarity of the query image with respect to a set of anchor points in the memory bank, converts that into a probability distribution over neighboring examples, and then transfers that knowledge to the student, so that the student also mimics the same neighborhood similarity. In the experiments, we show that our method is competitive with SOTA self-supervised methods on ImageNet and show an improved accuracy when trained on unbalanced, unlabeled data (for which we use a subset of ImageNet).

Our method is different from BYOL [18] in that we are comparing the query image with other random images rather than only with a different augmentation of the same query image. We believe our method can be seen as a more relaxed version of BYOL. Instead of imposing that the embedding of the query image should not change at all due to an augmentation (as done in BYOL), we are allowing the embedding to vary as long as its neighborhood similarity does not change. In other words, the augmentation should not change the similarity of the image compared to its neighboring images. This relaxation lets self-supervised learning focus on what matters most in learning rich features rather than forcing an unnecessary constraint of no change at all, which is difficult to achieve.

*Equal contribution

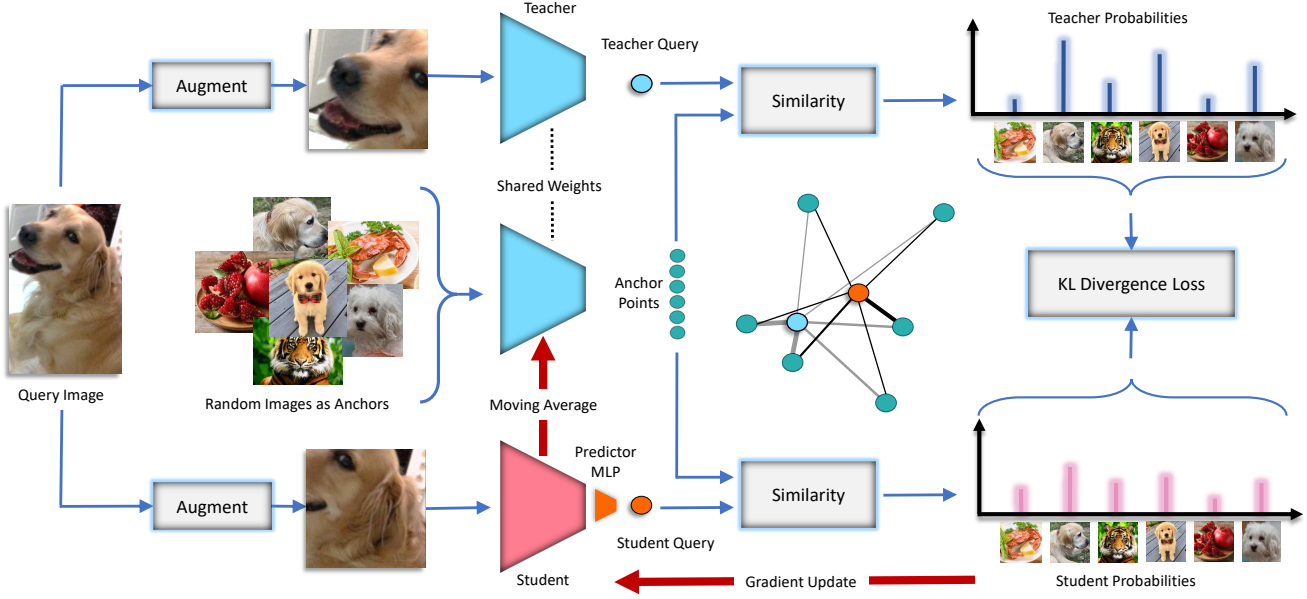


Figure 1. **Our method:** We initialize both teacher and student networks from scratch and update the teacher as running average of the student. We feed some random images to the teacher, and feed two different augmentations of a query image to both teacher and student. We capture the similarity of the query to the anchor points in the teacher’s embedding space and transfer that knowledge to the student. We update the student based on KL divergence loss and update the teacher to be a slow moving average of the student. This can be seen as a soft version of MoCo [20] which can handle negative images that are similar to the query image. Note that unlike contrastive learning and BYOL [18], we never compare two augmentations of the query images directly (positive pair).

Our distillation method is inspired by the CompRes method [1], which introduces an analogous similarity-based distillation method to compress a deeper self-supervised model to a smaller one and get better results compared to training the small model from scratch. Our method is different from [1] in that in our case, both teacher and student share the same architecture, we do self-supervised learning from scratch rather than compressing from another deeper model, and also the teacher evolves over time as a running average of the student rather than being frozen as in [1].

2. Method

We are interested in learning rich representations from unlabeled data. We have a teacher network and a student network. We initialize both models from scratch and update the teacher to be a slower version of the student: we use the momentum idea from MoCo in updating the teacher so that it is running average of the student. The method is described in Figure 1. Following the notation in [1], at each iteration, we pick a random query image and a bunch of random other images that we call anchor points. We augment those images and feed them to the teacher model to get their embeddings. Then, we augment the query again independent of the earlier augmentation and feed it to the

student model only. We calculate the similarity of the query point compared to the anchor points in the teacher’s embedding space and then optimize the student to mimic the same similarity for the anchor point at the student’s embedding space. Finally, we update the teacher with a momentum to be the running average of the student similar to MoCo and BYOL. Note that our method is closely related to CompRes method [1] which uses the similarity distillation for compressing a frozen larger model to a smaller one.

More formally, we assume a teacher model t and a student model s . Given a query image q , we augment it twice independently to get q_t and q_s . We also assume a set of n augmented random images $\{x_i\}_{i=1}^n$. We feed $\{x_i\}_{i=1}^n$ to the teacher model to get their embeddings $\{t(x_i)\}_{i=1}^n$ and call them anchor points. We also feed q_t to the teacher and q_s to the student to get $t(q_t)$ and $s(q_s)$ respectively. Then, we calculate the similarity of the query embedding $t(q_t)$ compared to all anchor points, divide by a temperature, and convert to a probability distribution using a SoftMax operator to get:

$$p_t(i) = -\log \frac{\exp(\text{sim}(t(q_t), t(x_i))/\tau_t)}{\sum_{j=1}^n \exp(\text{sim}(t(q_t), t(x_j))/\tau_t)}$$

where τ_t is the teacher’s temperature parameter and $\text{sim}(\cdot, \cdot)$ refers to the similarity between two vectors. In our exper-



Figure 2. **Positives vs. negatives:** We sample some query images randomly (left column), calculate their teacher probability distribution over all anchor points in the memory bank (size=128K) and rank them in descending order (right columns). The second left column is another augmented version of the query image that contrastive learning methods use for the positive pair. Our student learns to mimic the probability number written below each anchor image while contrastive learning method (e.g., MoCo) learn to predict the one-hot encoding written below the images. Note that there are lots of images in the top anchor points that are semantically similar to the query point that MoCo tries to discriminate them from the query while our method does not.

iments, we use cosine similarity which is standard in most recent contrastive learning methods.

Then, we calculate a similar probability distribution for the student’s query embedding to get:

$$p_s(i) = -\log \frac{\exp(\text{sim}(s(q_s), t(x_i))/\tau_s)}{\sum_{j=1}^n \exp(\text{sim}(s(q_s), t(x_j))/\tau_s)}$$

Where, τ_s is the student’s temperature parameter. Finally, we optimize the student only by minimizing the following loss:

$$L = KL(p_t || p_s)$$

and the teacher is updated using the following rule:

$$\theta_t \leftarrow m\theta_t + (1 - m)\theta_s$$

where θ refers to the parameters of a model and m is a the momentum hyperparameter that is set be close to one (0.99 in our ResNet50 experiments) as in MoCo. Since the teacher is not optimized by the loss directly, the loss can be simplified as cross entropy loss instead of KL divergence.

Note that unlike the positive pair in contrastive learning methods and BYOL, the query image is never compared to its own augmentation as it is not included in the

anchor points. We do this since when the features are mature, the similarity of the query to itself will be very large and will dominate the whole probability distribution. Note that one can convert our method to MoCo by including the query in anchor points and replacing the probability of the teacher p_t with a one-hot encoding vector in which the positive pair (query) corresponds to one and all other anchor points correspond to zero. Figure 2 shows some example teacher probabilities for both ISD and MoCo.

Our method can benefit from a large number of anchor points to cover the neighborhood of any query image, and also the anchor points are fed to the teacher only that evolves slowly. Hence, we use MoCo’s trick of a large memory bank (queue) for the anchor points. The queue size is 128K in our experiments which uses only 1.6% of the total memory and less than 1% of the total computation.

Different Temperature for student and teacher: Since the student is learning from the teacher, we can use a lower temperature for the teacher compared to the student to make the teacher more confident. In the extreme case, when the teacher uses zero temperature, its output will be a one-hot encoding over the anchor points which is a very sharp distribution. In the experiments we observe best results when the teacher has 10 times smaller temperature.

Method	Ref	Batch Size	Epochs	Sym. Loss 2x FLOPS	Top-1 Linear	NN	20-NN
ResNet-50							
Supervised	-	256	100	-	76.2	71.4	74.8
SwAV [8]	[11]	4096	200	✓	69.1	-	-
SimCLR[9]	[9]	4096	1000	✓	69.3	-	-
MoCo-V2 [20]	[11]	256	200	✓	69.9	-	-
SimSiam [11]	[11]	256	200	✓	70.0	-	-
BYOL [18]	[11]	4096	200	✓	70.6	-	-
MoCo-V2 [10]	[10]	256	400	✓	71.0	-	-
MoCo-V2 [20]	[20]	256	800	✗	71.1	57.3	61.0
CompRes* [1]	[1]	256	1K+130	✗	71.9	63.3	66.8
BYOL [18]	[18]	4096	1000	✓	74.3	62.8	66.9
SwAV † [8]	[8]	4096	800	✓	75.3	-	-
MoCo-V2 [20]	[11]	256	200	✗	67.5	-	-
CO2 [45]	[45]	256	200	✗	68.0	-	-
BYOL-asym	-	256	200	✗	69.3	55.0	59.2
MSF ‡ [25]	[25]	256	200	✗	72.4	62.0	64.9
ISD	-	256	200	✗	69.8	59.2	62.0
ResNet-18							
Supervised	-	256	100	-	69.8	63.0	67.6
MoCo-V2 [20]	[11]	256	200	✗	51.0	37.7	42.1
BYOL-asym	-	256	200	✗	52.6	40.0	44.8
ISD	-	256	200	✗	53.8	41.5	46.6

Table 1. **Evaluation on full ImageNet:** We compare our method with other state-of-the-art SSL methods by evaluating the learned features on the full ImageNet. A single linear layer is trained on top of a frozen backbone. Note that methods using symmetric losses use $2\times$ computation per mini-batch. Thus, it is not fair to compare them with the asymmetric loss methods. Further, we find that given a similar computational budget, both asymmetric MoCo-V2 (400 epochs) and symmetric MoCo-V2 (800 epochs) have similar accuracies (71.0 vs 71.1). Under similar resource constraints, our method performs competitively with other state-of-the-art methods. * is compressed from ResNet-50x4. †: SwAV is not comparable as it uses multiple crops together. ‡: is our concurrent (future of this!) work.

3. Experiments

We describe various experiments and their results in this section. We compare our proposed self-supervised method with other state-of-the-art methods on ImageNet and transfer learning. We also demonstrate the advantage of our method compared to MoCo on unbalanced, unlabeled dataset.

Implementation details: For all experiments, we use PyTorch with SGD optimizer (momentum = 0.9, weight decay = $1e-4$, batch size = 256) except when stated otherwise. Details about the specific architecture, epochs for training, and learning rate are described for each experiment in its respective section. We follow the evaluation protocols in [1] for nearest neighbor (NN) and linear layer (Linear) evaluation. We use the ImageNet labels only in the setting of evaluating the learned features. To evaluate how SSL features transfer to new tasks, we perform Linear layer evaluation on different datasets including Food101 [6], SUN397 [47], CIFAR10 [27], CIFAR100 [27], Cars [26], Flowers [30], Pets [35], Caltech-101 [15] and DTD [13]. More details about the datasets and training can be found in the appendix. We follow [18] setting for transfer learning and reproduce BYOL results for fairness.

3.1. Self-supervised learning

BYOL-asym (baseline). ResNet-50 is recently used as a benchmark in the community. Unfortunately, we cannot run it for 1000 epochs because of resource constraints. Some methods [18, 8] are even slower as they forward the mini-batch through the model more than once. For instance, BYOL method forwards the images twice to calculate the symmetric loss, so 100 epochs of symmetric BYOL is equivalent to almost 200 epochs of asymmetric BYOL in terms of running time. As shown in [11], given a constant budget, there is no big difference between symmetric and asymmetric losses. Thus, for a fair comparison with our method and MoCo, we use asymmetric loss, a small batch size (256), momentum for the teacher is 0.99, and train for 200 epochs. We implement BYOL in PyTorch following [18]. We call this baseline as BYOL-asym since it’s asymmetric version of BYOL. For ResNet18, hidden units in the MLP for projection and prediction layers is 1024, and output embedding dimension is 128. For ResNet50, hidden units in the MLP for projection and prediction layers is 4096, and output embedding dimension is 512. We use cos learning rate scheduler with initial learning rate of 0.05.

Method	Ref.	Epochs	Food 101	CIFAR 10	CIFAR 100	SUN 397	Cars 196	DTD	Pets	Caltech 101	Flowers 102	Mean
ResNet-50												
Sup-IN	[18]		72.3	93.6	78.3	61.9	66.7	74.9	91.5	94.5	94.7	80.9
SimCLR [9]	[18]	1000	72.8	90.5	74.4	60.6	49.3	75.7	84.6	89.3	92.6	68.6
MoCo v2 [10]	-	800	72.5	92.2	74.6	59.6	50.5	74.4	84.6	90.0	90.5	76.5
BYOL [18]	rep.	1000	75.4	92.7	78.1	62.1	67.1	76.8	89.8	92.2	95.5	81.1
BYOL [18]	[18]	1000	75.3	91.3	78.4	62.2	67.8	75.5	90.4	94.2	96.1	81.2
BYOL-asym [18]	-	200	70.2	91.5	74.2	59.0	54.0	73.4	86.2	90.4	92.1	76.8
MoCo v2 [10]	-	200	70.4	91.0	73.5	57.5	47.7	73.9	81.3	88.7	91.1	75.0
MSF ‡ [25]	[25]	200	71.2	92.6	76.3	59.2	55.6	73.2	88.7	92.7	92.0	77.9
ISD	-	200	68.6	90.8	72.0	55.8	45.8	68.6	89.1	90.3	87.4	74.3
ResNet-18												
BYOL-asym [18]	-	200	55.0	83.4	59.3	48.2	26.6	65.4	74.1	82.7	82.3	64.1
MoCo v2 [10]	-	200	56.7	83.0	59.7	48.8	30.4	64.4	70.1	80.5	83.1	64.1
ISD	-	200	58.3	83.3	62.7	49.6	36.1	65.6	76.4	84.5	87.4	67.1

Table 2. **Linear transfer evaluation:** We linear classifiers on top of frozen features for various downstream datasets. Hyperparameters are tuned individually for each method and the results are reported on the hold-out test sets. Our ResNet-18 is significantly better than other state-of-the-art SSL methods. “rep.” refers to the reproduction with our framework for a fair comparison. ‡: our concurrent work.

Method	Epochs	Top-1		Top-5	
		1%	10%	1%	10%
<i>Entire network is fine-tuned.</i>					
Supervised		25.4	56.4	48.4	80.4
PIRL [28]	800	-	-	57.2	83.8
CO2 [45]	200	-	-	71.0	85.7
SimCLR [9]	1000	48.3	65.6	75.5	87.8
InvP [44]	800	-	-	78.2	88.7
BYOL [18]	1000	53.2	68.8	78.4	89.0
SwAV [†] [8]	800	53.9	70.2	78.5	89.9
<i>Only the linear layer is trained.</i>					
BYOL [‡] [18]	1000	55.7	68.6	80.0	88.6
CompRes* [1]	1K+130	59.7	67.0	82.3	87.5
MoCo v2 [10]	200	43.6	58.4	71.2	82.9
BYOL-asym [18]	200	47.9	61.3	74.6	84.7
ISD	200	53.4	63.0	78.8	85.9

Table 3. **Evaluation on limited labels ImageNet for ResNet-50:** We evaluate our model for the 1% and 10% ImageNet linear evaluation. Unlike other methods, we only train a single linear layer on top of the frozen backbone. We observe that our method is better than other state-of-the-art methods given similar computational budgets. * is compressed from ResNet-50x4

ResNet-18 experiments: Following CompRes [1], we train our self-supervised ResNet-18 model with the initial learning rate set to 0.01 and multiplied by 0.2 at epochs 140 and 180. We follow [18] and add a prediction layer for the student. The hidden and output dimensions of the prediction MLP layer are set to 512. We do not have any projection layer for ResNet18. We use the same set of augmentations used in [10, 9, 18]. We use the same temperature for teacher and student $\tau_s = \tau_t = 0.02$. The memory bank size is 128K

and momentum m for teacher encoder is 0.999. We choose these parameters based on ablations which can be found in the appendix. The results are shown in Tables 1 and 2. Our model outperforms both baselines on full ImageNet linear and transfer linear benchmarks.

It is important to note that our method can be seen as a soft version of MoCo. So, ISD outperforming MoCo, empirically supports our main motivation of improving representations by smoothing the contrastive learning: not considering all negatives equally negative.

ResNet-50 experiments: We train ResNet-50 with different settings than ResNet-18. We use same architecture and settings as ResNet-50 BYOL-asym. We use cos learning rate scheduler with initial learning rate of 0.05. We use temperature of $\tau_s = 0.1$ for the student and $\tau_t = 0.01$ for the teacher. Memory bank size is 128K, and momentum m for teacher encoder is 0.99. We study the effect of memory bank size in Figure 3 which shows that memory bank size of even 16K is on-par with 128K and above. Additionally, inspired by [40], we train our model with two different augmentation sets which we call it weak (random cropping and random horizontal flipping) and strong (same as [10, 9, 18]). The teacher view uses weak augmentation while the student view uses the strong augmentation. We evaluate the effect of different augmentation in Table 5. ResNet50 results are shown in Tables 1 and 2.

Tables 1 and 2 show the results on ImageNet and transfer learning settings respectively. Our method is comparable to SOTA SSL methods including BYOL in Linear and nearest neighbor evaluation on ImageNet. As mentioned earlier, we believe our method is more relaxed compared to BYOL as our method lets the embeddings of augmented images move as long as their similarity relationship with neighbors has

not changed. Table 3 shows our results when only limited labels are available in ImageNet dataset. Figure 5 shows random image samples from random clusters where each row corresponds to a cluster. Note that each row contains almost semantically similar images.

Evolution of teacher and student models: In Figure 4, for every 10 epoch of ResNet-18, we evaluate both teacher and student models for BYOL, MoCo, and ISD methods using nearest neighbor. For all methods, the teacher performs usually better than the student in the initial epochs when the learning rate is relatively large and then is very close to the student when it shrinks. This is interesting as we have not seen previous papers comparing the teacher with the student. This might happen since the teacher is a running average of the student so can be seen as an ensemble over many student networks similar to [41]. We believe this deserves more investigation as future work.

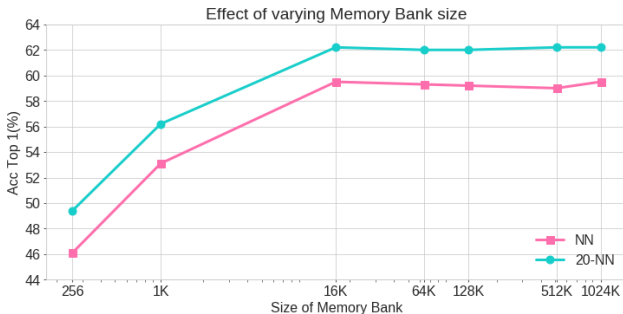


Figure 3. **Effect of Memory Bank Size:** We study the effect of memory bank size by varying from 256 to 1024K for ISD on ResNet-50 model.

Ablation study: We varied the temperature for our method on ResNet-18 with 130 epochs and reported the results in Table 6. Here, $LR = 0.01$ and it is multiplied by 0.2 at 90 and 120 epochs. Also, for more fair comparison with BYOL on ResNet18, we varied the learning rate and chose the best one for BYOL. Table 7 shows the results of this experiment.

3.2. Self-Supervised Learning on Unbalanced Dataset

Most recent self-supervised learning methods are benchmarked by training on unlabeled ImageNet. However, we know that ImageNet has a particular bias of having an almost uniform distribution over the number of samples per category. We believe this bias does not exist in many real-world applications where the data is unbalanced: a few categories have a large number of samples while the rest of the data have a small number of samples. For instance, in self-driving car applications, it is really important to learn features for understanding rare scenes while most of the data is captured from repetitive safe highway scenes. Hence, we

believe it is important to design and evaluate self-supervised learning methods for such unbalanced data.

As mentioned earlier, since standard contrastive learning methods e.g., MoCo, consider all negative examples equally negative, when the query image is from a large category, it is possible to have multiple samples from the same category in the memory bank. Then, the contrastive loss in MoCo pushes their embeddings to be far apart as negative pairs. However, our method can handle such cases since our teacher assigns a soft label to the negative samples, so if an anchor example is very similar to the query, it will have high similarity and the student is optimized to reproduce such similarity.

To study our method on unbalanced data, we design a controlled setting to introduce the unbalanced data in the SSL training only and factor out its effect in the feature evaluation step. Hence, we sub-sample ImageNet data with 38 random categories where 8 categories are large (use all of almost 1300 images per category) and 30 categories are small (use only 100 images per category.) We train our SSL method and then evaluate by nearest neighbor (NN) classifier on the balanced validation data. To make sure that the feature evaluation is not affected by the unbalanced data, we keep both evaluation and the training data of NN search balanced, so for NN search, we use all ImageNet training images (almost 1300×38 images) for those 38 categories.

We repeat the sampling of 38 categories 10 times to come up with 10 datasets and report the results for our method and also MoCo in Table 4. To measure the effect of the unbalanced data, we report the accuracy on all 38 categories and also on those 30 small categories only separately. Our method performs consistently better than MoCo, but more interestingly, the gap the improvement is larger when we evaluate on the 30 small categories only. We believe this empirically proves our hypothesis that our method may be able to handle unbalanced data more effectively. For a fair comparison, we train both our model and MoCo for 400 epochs with a memory bank size of 8192 and cosine learning schedule.

4. Related Work

Self-supervised learning: The task of learning representations by solving a pretext task without using any supervised annotations is called self-supervised learning. Various pretext tasks like solving jigsaw puzzles [31], predicting rotations [17], counting the visual primitives [32], filling up a missing patch [37], predicting missing channels of input [52, 53], and contrastive learning [19, 20] are explored in the literature. We are proposing a novel pretext task based on iterative similarity based distillation.

Contrastive learning: The task of learning unsupervised representations by contrasting the representations of an image with other images is called contrastive learning

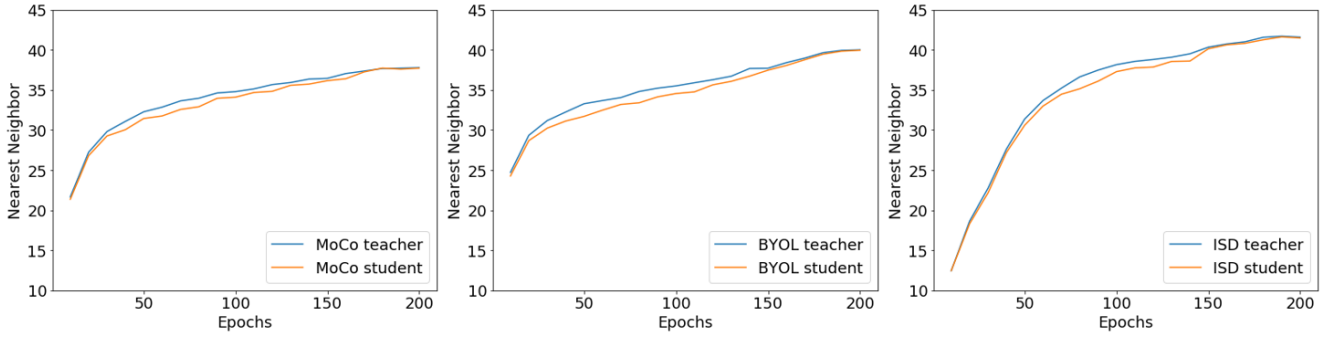


Figure 4. **Evolution of teacher and student models:** Comparing the teacher and student ResNet18 models using Nearest Neighbor while training for MoCo, BYOL, and ISD methods. Interestingly, the teacher performs better than the student before shrinking the learning rate. Most previous works use the student as the final model which seems to be sub-optimal. We believe this is due to ensembling effect similar to [41] and needs more investigation.



Figure 5. **Random Clusters:** We cluster ImageNet dataset into 1000 clusters using k-means and show random samples from random clusters. We did not do cherry-picking for this visualization. Each row corresponds to a cluster. Note that semantically similar images are clustered together. More results can be found in the appendix.

[19]. Contrastive learning is essentially positive/negative classification where the positive and negative pairs of embeddings can be defined in various ways. In [46, 9, 20], the positive pairs are augmented views of the same image while the negative pairs are those of different images. In [21], the positive pairs are a patch and context embeddings from the same image while the negatives pairs are a patch and context embedding from different images. In [23, 4], the positives pairs are global and local features from the same image while the negative pairs are global and local features from different images. In [50, 8, 7], the positive pairs are members of the same cluster while the negative pairs are member of different clusters. We are different from these contrastive methods in that we do not consider all negatives equally: we calculate a soft labeling for the negatives using similarity of the data points. Also, we do not consider positive pairs directly: the comparison for the positive pairs

is done through the similarity distillation. A few methods have attempted to fix false negative problem of contrastive learning by debiasing the loss [12] and by sampling the local neighborhood as positives [44]. We’re different as we simply make the contrastive learning soft. [24] is a concurrent work that identifies some wrong negatives and cancels them in constrictive leaning.

Knowledge Distillation: The task of transferring the knowledge from one model to the other is called knowledge distillation [22, 3]. The knowledge from the teacher can be extracted and transferred in various ways. The knowledge in the activations of intermediate layers can be transferred through regression [39, 49, 51]. While in most works the teacher is a deeper model and the student is a shallower model, in [5, 16] both teacher and the student use the same architecture. Techniques from knowledge distillation can also be used in an unsupervised way to improve self-

Method	D_1	D_2	D_3	D_4	D_5	D_6	D_7	D_8	D_9	D_{10}	Mean
Evaluation On All 38 Categories											
MoCo	48.9	57.8	59.4	56.7	62.0	54.6	57.8	57.2	62.4	54.4	57.12
ISD	49.6	57.9	61.9	58.6	62.3	56.3	57.8	58.1	62.5	55.3	58.03
Diff	+0.7	+0.1	+2.5	+1.9	+0.3	+1.7	0	+0.9	+0.1	+0.9	+0.91
Evaluation Only on 30 Rare Categories											
MoCo	44.7	52.3	57.3	53.1	57.7	50.7	51.1	51.9	58.9	59.8	53.75
ISD	46.5	53.9	60.8	56.8	60.5	54.5	53.1	55.0	60.7	61.5	56.33
Diff	+1.7	+1.6	+3.5	+3.7	+2.8	+3.8	+2.0	+3.1	+1.8	+1.7	+2.57

Table 4. **Unbalanced dataset:** Nearest Neighbor (NN) results with ResNet-18 model for the unbalanced data when we consider all 38 categories and 30 small categories separately. We repeat the experiment 10 times with different random sets of 38 categories. NN is done on the validation set of ImageNet (which has uniform distribution) by searching the nearest neighbors among all ImageNet training data of those 38 categories (so the training data of NN also has uniform distribution). Hence, the whole evaluation is on balanced data to make sure we observe the effect of the unbalanced, “unlabeled” data only. “Diff” shows the improvement of our method over MoCo. Interestingly the improvement is bigger in the rare categories. This is aligned with our hypothesis that our method can handle unbalanced, unlabeled data better since it does not consider all negative images equally negative.

Method	Student Aug.	Teacher Aug.	NN	20-NN
ISD	weak	weak	40.4	43.5
ISD	strong	weak	22.9	26.3
ISD	strong	strong	58.0	61.2
ISD	weak	strong	59.2	62.0

Table 5. **Effect of augmentation strategies:** Effect of using weak or strong augmentations for ResNet-50 trained with 200 epochs.

τ	0.003	0.007	0.01	0.02	0.04	0.06
NN	37.2	37.8	37.7	39.7	35.3	32.5

Table 6. **Effect of temperature:** Effect of changing temperature for our method ISD on ResNet-18 model.

LR	0.01	0.05	0.10	0.20
NN	37.3	40.0	38.6	37.3

Table 7. **Effect of learning rates for BYOL:** Comparison of different learning rates for BYOL on ResNet-18 with 200 epochs and cosine learning rate scheduler.

supervised learning [1, 33, 48]. Instead of using knowledge distillation for model compression or reducing the generalization gap, we use it iteratively to evolve the teacher and student together to learn rich representations from scratch.

Similarity based knowledge distillation: While the above methods [22, 3, 39, 49, 51] only extract the information about a single data point from the teacher, similarity based distillation methods [36, 38, 34, 43, 2, 42, 1, 14] represent the knowledge from the teacher in terms of similarities between data points. CompRes [1] and SEED [14] are the closest related work to our method that uses similarity based distillation to compresses a large self-supervised model to a smaller one. Our method is different as we use

similarity based distillation to iteratively distill an evolving teacher to a student.

Consistency regularization: Consistency regularization is a method of regularization that seeks to make the output of a model consistent across small perturbations in either the input [29] or the model parameters [41]. Recently, BYOL [18] applied a variant of Mean Teachers [41] for self-supervised learning. Our method uses a variation of this idea through similarity based loss rather than regression loss defined on data points individually. [45] is probably the closest to ours that uses a loss similar to our as a regularizer in addition to the MoCo loss. Our method is different as we optimize our loss from scratch as the main objective without adding it to another method. Also, our method benefits from different temperatures for the teacher and student networks which is inspired by [14].

5. Conclusion

We introduce ISD, a novel self-supervised learning method. It is a variation of contrastive learning (e.g., MoCo) in which negative samples are not all treated equally. The similarity between images in the teacher’s embedding space determines how much each anchor image should be contrasted with. Our extensive experiments show that our method performs comparable to the state-of-the-art SSL methods on ImageNet, transfer learning tasks, and when the unlabeled data is unbalanced.

Acknowledgment: This material is based upon work partially supported by the United States Air Force under Contract No. FA8750-19-C-0098, funding from SAP SE, and also NSF grant numbers 1845216 and 1920079. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the authors and do not necessarily reflect the views of the United States Air Force, DARPA, or other funding agencies.

References

- [1] Soroush Abbasi Koohpayegani, Ajinkya Tejankar, and Hamed Pirsiavash. Compress: Self-supervised learning by compressing representations. *Advances in Neural Information Processing Systems*, 33, 2020.
- [2] Sungsoo Ahn, Shell Xu Hu, Andreas Damianou, Neil D Lawrence, and Zhenwen Dai. Variational information distillation for knowledge transfer. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 9163–9171, 2019.
- [3] Jimmy Ba and Rich Caruana. Do deep nets really need to be deep? In *Advances in neural information processing systems*, pages 2654–2662, 2014.
- [4] Philip Bachman, R Devon Hjelm, and William Buchwalter. Learning representations by maximizing mutual information across views. In *Advances in Neural Information Processing Systems*, pages 15535–15545, 2019.
- [5] Hessam Bagherinezhad, Maxwell Horton, Mohammad Rastegari, and Ali Farhadi. Label refinery: Improving imagenet classification through label progression. *arXiv preprint arXiv:1805.02641*, 2018.
- [6] Lukas Bossard, Matthieu Guillaumin, and Luc Van Gool. Food-101 – mining discriminative components with random forests. In *European Conference on Computer Vision*, 2014.
- [7] Mathilde Caron, Piotr Bojanowski, Armand Joulin, and Matthijs Douze. Deep clustering for unsupervised learning of visual features. In *Proceedings of the European Conference on Computer Vision (ECCV)*, pages 132–149, 2018.
- [8] Mathilde Caron, Ishan Misra, Julien Mairal, Priya Goyal, Piotr Bojanowski, and Armand Joulin. Unsupervised learning of visual features by contrasting cluster assignments. *arXiv preprint arXiv:2006.09882*, 2020.
- [9] Ting Chen, Simon Kornblith, Mohammad Norouzi, and Geoffrey Hinton. A simple framework for contrastive learning of visual representations. *arXiv preprint arXiv:2002.05709*, 2020.
- [10] Xinlei Chen, Haoqi Fan, Ross Girshick, and Kaiming He. Improved baselines with momentum contrastive learning. *arXiv preprint arXiv:2003.04297*, 2020.
- [11] Xinlei Chen and Kaiming He. Exploring simple siamese representation learning, 2020.
- [12] Ching-Yao Chuang, Joshua Robinson, Yen-Chen Lin, Antonio Torralba, and Stefanie Jegelka. Debaised contrastive learning. In *Advances in Neural Information Processing Systems*, 2020.
- [13] Mircea Cimpoi, Subhansu Maji, Iasonas Kokkinos, Sammy Mohamed, and Andrea Vedaldi. Describing textures in the wild. In *Computer Vision and Pattern Recognition*, 2014.
- [14] Zhiyuan Fang, Jianfeng Wang, Lijuan Wang, Lei Zhang, Yezhou Yang, and Zicheng Liu. Seed: Self-supervised distillation for visual representation. In *International Conference on Learning Representations*, 2021.
- [15] Li Fei-Fei, Rob Fergus, and Pietro Perona. Learning generative visual models from few training examples: An incremental bayesian approach tested on 101 object categories. *Computer Vision and Pattern Recognition Workshop*, 2004.
- [16] Tommaso Furlanello, Zachary Chase Lipton, Michael Tschannen, Laurent Itti, and Anima Anandkumar. Born-again neural networks. In Jennifer G. Dy and Andreas Krause, editors, *Proceedings of the 35th International Conference on Machine Learning, ICML 2018, Stockholm, Sweden, July 10-15, 2018*, volume 80 of *Proceedings of Machine Learning Research*, pages 1602–1611. PMLR, 2018.
- [17] Spyros Gidaris, Praveer Singh, and Nikos Komodakis. Unsupervised representation learning by predicting image rotations. In *International Conference on Learning Representations*, 2018.
- [18] Jean-Bastien Grill, Florian Strub, Florent Altché, Corentin Tallec, Pierre H Richemond, Elena Buchatskaya, Carl Doersch, Bernardo Avila Pires, Zhaohan Daniel Guo, Mohammad Gheshlaghi Azar, et al. Bootstrap your own latent: A new approach to self-supervised learning. *arXiv preprint arXiv:2006.07733*, 2020.
- [19] Raia Hadsell, Sumit Chopra, and Yann LeCun. Dimensionality reduction by learning an invariant mapping. In *2006 IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR'06)*, volume 2, pages 1735–1742. IEEE, 2006.
- [20] Kaiming He, Haoqi Fan, Yuxin Wu, Saining Xie, and Ross Girshick. Momentum contrast for unsupervised visual representation learning. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 9729–9738, 2020.
- [21] Olivier J Hénaff, Aravind Srinivas, Jeffrey De Fauw, Ali Razavi, Carl Doersch, SM Eslami, and Aaron van den Oord. Data-efficient image recognition with contrastive predictive coding. *arXiv preprint arXiv:1905.09272*, 2019.
- [22] Geoffrey Hinton, Oriol Vinyals, and Jeff Dean. Distilling the knowledge in a neural network. *arXiv preprint arXiv:1503.02531*, 2015.
- [23] R Devon Hjelm, Alex Fedorov, Samuel Lavoie-Marchildon, Karan Grewal, Phil Bachman, Adam Trischler, and Yoshua Bengio. Learning deep representations by mutual information estimation and maximization. In *International Conference on Learning Representations*, 2019.
- [24] Tri Huynh, Simon Kornblith, Matthew R. Walter, Michael Maire, and Maryam Khademi. Boosting contrastive self-supervised learning with false negative cancellation, 2020.
- [25] Soroush Abbasi Koohpayegani, Ajinkya Tejankar, and Hamed Pirsiavash. Mean shift for self-supervised learning, 2021.
- [26] Jonathan Krause, Michael Stark, Jia Deng, and Li Fei-Fei. 3D object representations for fine-grained categorization. In *Workshop on 3D Representation and Recognition*, Sydney, Australia, 2013.
- [27] Alex Krizhevsky. Learning multiple layers of features from tiny images. Technical report, University of Toronto, 2009.
- [28] Ishan Misra and Laurens van der Maaten. Self-supervised learning of pretext-invariant representations. *arXiv preprint arXiv:1912.01991*, 2019.
- [29] Takeru Miyato, Shin-ichi Maeda, Masanori Koyama, and Shin Ishii. Virtual adversarial training: a regularization

- method for supervised and semi-supervised learning. *IEEE transactions on pattern analysis and machine intelligence*, 41(8):1979–1993, 2018.
- [30] Maria-Elena Nilsback and Andrew Zisserman. Automated flower classification over a large number of classes. In *Indian Conference on Computer Vision, Graphics and Image Processing*, 2008.
 - [31] Mehdi Noroozi and Paolo Favaro. Unsupervised learning of visual representations by solving jigsaw puzzles. In *European Conference on Computer Vision*, pages 69–84. Springer, 2016.
 - [32] Mehdi Noroozi, Hamed Pirsiavash, and Paolo Favaro. Representation learning by learning to count. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 5898–5906, 2017.
 - [33] Mehdi Noroozi, Ananth Vinjimoor, Paolo Favaro, and Hamed Pirsiavash. Boosting self-supervised learning via knowledge transfer. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 9359–9367, 2018.
 - [34] Wonpyo Park, Dongju Kim, Yan Lu, and Minsu Cho. Relational knowledge distillation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 3967–3976, 2019.
 - [35] O. M. Parkhi, A. Vedaldi, A. Zisserman, and C. V. Jawahar. Cats and dogs. In *Computer Vision and Pattern Recognition*, 2012.
 - [36] Nikolaos Passalis and Anastasios Tefas. Learning deep representations with probabilistic knowledge transfer. In *Proceedings of the European Conference on Computer Vision (ECCV)*, pages 268–284, 2018.
 - [37] Deepak Pathak, Philipp Krahenbuhl, Jeff Donahue, Trevor Darrell, and Alexei A Efros. Context encoders: Feature learning by inpainting. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 2536–2544, 2016.
 - [38] Baoyun Peng, Xiao Jin, Jiaheng Liu, Dongsheng Li, Yichao Wu, Yu Liu, Shunfeng Zhou, and Zhaoning Zhang. Correlation congruence for knowledge distillation. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 5007–5016, 2019.
 - [39] Adriana Romero, Nicolas Ballas, Samira Ebrahimi Kahou, Antoine Chassang, Carlo Gatta, and Yoshua Bengio. Fit-nets: Hints for thin deep nets. In Yoshua Bengio and Yann LeCun, editors, *3rd International Conference on Learning Representations, ICLR 2015, San Diego, CA, USA, May 7-9, 2015, Conference Track Proceedings*, 2015.
 - [40] Kihyuk Sohn, David Berthelot, Nicholas Carlini, Zizhao Zhang, Han Zhang, Colin A Raffel, Ekin Dogus Cubuk, Alexey Kurakin, and Chun-Liang Li. Fixmatch: Simplifying semi-supervised learning with consistency and confidence. *Advances in Neural Information Processing Systems*, 33, 2020.
 - [41] Antti Tarvainen and Harri Valpola. Mean teachers are better role models: Weight-averaged consistency targets improve semi-supervised deep learning results. In *Advances in neural information processing systems*, pages 1195–1204, 2017.
 - [42] Yonglong Tian, Dilip Krishnan, and Phillip Isola. Contrastive representation distillation. In *International Conference on Learning Representations*, 2020.
 - [43] Frederick Tung and Greg Mori. Similarity-preserving knowledge distillation. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 1365–1374, 2019.
 - [44] Feng Wang, Huaping Liu, Di Guo, and Sun Fuchun. Unsupervised representation learning by invariance propagation. In *Advances in Neural Information Processing Systems*, 2020.
 - [45] Chen Wei, Huiyu Wang, Wei Shen, and Alan Yuille. Co2: Consistent contrast for unsupervised visual representation learning. *arXiv preprint arXiv:2010.02217*, 2020.
 - [46] Zhirong Wu, Yuanjun Xiong, Stella X Yu, and Dahua Lin. Unsupervised feature learning via non-parametric instance discrimination. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 3733–3742, 2018.
 - [47] Jianxiong Xiao, James Hays, Krista A. Ehinger, Aude Oliva, and Antonio Torralba. Sun database: Large-scale scene recognition from abbey to zoo. In *Computer Vision and Pattern Recognition*, 2010.
 - [48] Xueting Yan, Ishan Misra, Abhinav Gupta, Deepti Ghadiyaram, and Dhruv Mahajan. Clusterfit: Improving generalization of visual representations. In *CVPR*, 2020.
 - [49] Junho Yim, Donggyu Joo, Jihoon Bae, and Junmo Kim. A gift from knowledge distillation: Fast optimization, network minimization and transfer learning. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 4133–4141, 2017.
 - [50] Asano YM., Rupprecht C., and Vedaldi A. Self-labelling via simultaneous clustering and representation learning. In *International Conference on Learning Representations*, 2020.
 - [51] Sergey Zagoruyko and Nikos Komodakis. Paying more attention to attention: Improving the performance of convolutional neural networks via attention transfer. In *ICLR*, 2017.
 - [52] Richard Zhang, Phillip Isola, and Alexei A Efros. Colorful image colorization. In *European conference on computer vision*, pages 649–666. Springer, 2016.
 - [53] Richard Zhang, Phillip Isola, and Alexei A Efros. Split-brain autoencoders: Unsupervised learning by cross-channel prediction. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 1058–1067, 2017.