

# Pairwise Confusion for Fine-Grained Visual Classification

Abhimanyu Dubey<sup>1</sup>, Otkrist Gupta<sup>1</sup>, Pei Guo<sup>2</sup>, Ramesh Raskar<sup>1</sup>, Ryan Farrell<sup>2</sup>,  
and Nikhil Naik<sup>1,3</sup>

<sup>1</sup> Massachusetts Institute of Technology, Cambridge MA 02139, USA  
`{dubeya,otkrist,raskar,naik}@mit.edu`

<sup>2</sup> Brigham Young University, Provo UT 84602, USA  
`peigu, farrell@cs.byu.edu`

<sup>3</sup> Harvard University, Cambridge MA 02139, USA  
`naik@fas.harvard.edu`

**Abstract.** Fine-Grained Visual Classification (FGVC) datasets contain small sample sizes, along with significant intra-class variation and inter-class similarity. While prior work has addressed intra-class variation using localization and segmentation techniques, inter-class similarity may also affect feature learning and reduce classification performance. In this work, we address this problem using a novel optimization procedure for the end-to-end neural network training on FGVC tasks. Our procedure, called Pairwise Confusion (PC) reduces overfitting by intentionally introducing confusion in the activations. With PC regularization, we obtain state-of-the-art performance on six of the most widely-used FGVC datasets and demonstrate improved localization ability. PC is easy to implement, does not need excessive hyperparameter tuning during training, and does not add significant overhead during test time.

## 1 Introduction

The Fine-Grained Visual Classification (FGVC) task focuses on differentiating between hard-to-distinguish object classes, such as species of birds, flowers, or animals; and identifying the makes or models of vehicles. FGVC datasets depart from conventional image classification in that they typically require expert knowledge, rather than crowdsourcing, for gathering annotations. FGVC datasets contain images with much higher visual similarity than those in large-scale visual classification (LSVC). Moreover, FGVC datasets have minute inter-class visual differences in addition to the variations in pose, lighting and viewpoint found in LSVC [1]. Additionally, FGVC datasets often exhibit long tails in the data distribution, since the difficulty of obtaining examples of different classes may vary. This combination of small, non-uniform datasets and subtle inter-class differences makes FGVC challenging even for powerful deep learning algorithms.

Most of the prior work in FGVC has focused on tackling the *intra-class* variation in pose, lighting, and viewpoint using localization techniques [1,2,3,4,5], and by augmenting training datasets with additional data from the Web [6,7].

However, we observe that prior work in FGVC does not pay much attention to the problems that may arise due to the *inter-class* visual similarity in the feature extraction pipeline. Similar to LSVC tasks, neural networks for FGVC tasks are typically trained with cross-entropy loss [1,7,8,9]. In LSVC datasets such as ImageNet [10], strongly discriminative learning using the cross-entropy loss is successful in part due to the significant inter-class variation (compared to intra-class variation), which enables deep networks to learn generalized discriminatory features with large amounts of data.

We posit that this formulation may not be ideal for FGVC, which shows smaller visual differences between classes and larger differences within each class than LSVC. For instance, if two samples in the training set have very similar visual content but different class labels, minimizing the cross-entropy loss will force the neural network to learn features that distinguish these two images with high confidence—potentially forcing the network to learn sample-specific artifacts for visually confusing classes in order to minimize training error. We suspect that this effect would be especially pronounced in FGVC, since there are fewer samples from which the network can learn generalizable class-specific features.

Based on this hypothesis, we propose that introducing *confusion* in output logit activations during training for an FGVC task will force the network to learn slightly less discriminative features, thereby preventing it from overfitting to sample-specific artifacts. Specifically, we aim to *confuse* the network, by minimizing the distance between the predicted probability distributions for random pairs of samples from the training set. To do so, we propose Pairwise Confusion (PC)<sup>4</sup>, a pairwise algorithm for training convolutional neural networks (CNNs) end-to-end for fine-grained visual classification.

In Pairwise Confusion, we construct a Siamese neural network trained with a novel loss function that attempts to bring class conditional probability distributions closer to each other. Using Pairwise Confusion with a standard network architecture like DenseNet [11] or ResNet [12] as a base network, we obtain state-of-the-art performance on six of the most widely-used fine-grained recognition datasets, improving over the previous-best published methods by 1.86% on average. In addition, PC-trained networks show better localization performance as compared to standard networks. Pairwise Confusion is simple to implement, has no added overhead in training or prediction time, and provides performance improvements both in FGVC tasks and other tasks that involve transfer learning with small amounts of training data.

## 2 Related Work

**Fine-Grained Visual Classification:** Early FGVC research focused on methods to train with limited labeled data and traditional image features. Yao et al. [13] combined strongly discriminative image patches with randomization techniques to prevent overfitting. Yao et al. [14] subsequently utilized template matching to avoid the need for a large number of annotations.

<sup>4</sup> Implementation available at <https://github.com/abhimanyudubey/confusion>.

**Table 1.** A comparison of fine-grained visual classification (FGVC) datasets with large-scale visual classification (LSVC) datasets. FGVC datasets are significantly smaller and noisier than LSVC datasets.

| Dataset            | num. classes | samples per class | Dataset        | num. classes | samples per class |
|--------------------|--------------|-------------------|----------------|--------------|-------------------|
| Flowers-102 [32]   | 102          | 10                | CIFAR-100 [38] | 100          | 500               |
| CUB-200-2011 [33]  | 200          | 29.97             | ImageNet [10]  | 1000         | 1200              |
| Cars [34]          | 196          | 41.55             | CIFAR-10 [38]  | 10           | 5000              |
| NABirds [35]       | 550          | 43.5              | SVHN [39]      | 10           | 7325.7            |
| Aircrafts [36]     | 100          | 100               |                |              |                   |
| Stanford Dogs [37] | 120          | 100               |                |              |                   |

Recently, improved localization of the target object in training images has been shown to be useful for FGVC [1,15,16,17]. Zhang et al. [15] utilize part-based Region-CNNs [18] to perform finer localization. Spatial Transformer Networks [2] show that learning a content-based affine transformation layer improves FGVC performance. Pose-normalized CNNs have also been shown to be effective at FGVC [19,20]. Model ensembling and boosting has also improved performance on FGVC [21]. Lin et al. [1] introduced Bilinear Pooling, which combines pairwise local feature sets and improves classification performance. Bilinear Pooling has been extended by Gao et al. [16] using a compact bilinear representation and Cui et al. [9] using a general Kernel-based pooling framework that captures higher-order interactions of features.

**Pairwise Learning:** Chopra et al. [22] introduced a Siamese neural network for handwriting recognition. Parikh and Grauman [23] developed a pairwise ranking scheme for relative attribute learning. Subsequently, pairwise neural network models have become common for attribute modeling [24,25,26,27].

**Learning from Label Confusion:** Our method aims to improve classification performance by introducing confusion within the output labels. Prior work in this area includes methods that utilize label noise (e.g., [28]) and data noise (e.g., [29]) in training. Krause et al. [6] utilized noisy training data for FGVC. Neelakantan et al. [30] added noise to the gradient during training to improve generalization performance in very deep networks. Szegedy et al. [31] introduced label-smoothing regularization for training deep Inception models.

In this paper, we bring together concepts from pairwise learning and label confusion and take a step towards solving the problems of overfitting and sample-specific artifacts when training neural networks for FGVC tasks.

### 3 Method

FGVC datasets in computer vision are orders of magnitude smaller than LSVC datasets and contain greater imbalance across classes (see Table 1). Moreover, the samples of a class are not accurately representative of the complete variation in the visual class itself. The smaller dataset size can result in overfitting when

training deep neural architectures with large number of parameters—even with preliminary layers being frozen. In addition, the training data may not be completely representative of the real-world data, with issues such as more abundant sampling for certain classes. For example, in FGVC of birds, certain species from geographically accessible areas may be overrepresented in the training dataset. As a result, the neural network may learn to latch on to sample-specific artifacts in the image, instead of learning a versatile representation for the target object. We aim to solve both of these issues in FGVC (overfitting and sample-specific artifacts) by bringing the different class-conditional probability distributions closer together and *confusing* the deep network, subsequently reducing its prediction over-confidence, thus improving generalization performance.

Let us formalize the idea of “confusing” the conditional probability distributions. Consider the conditional probability distributions for two input images  $\mathbf{x}_1$  and  $\mathbf{x}_2$ , which can be given by  $p_\theta(\mathbf{y}|\mathbf{x}_1)$  and  $p_\theta(\mathbf{y}|\mathbf{x}_2)$  respectively. For a classification problem with  $N$  output classes, each of these distributions is an  $N$ -dimensional vector, with each element  $i$  denoting the belief of the classifier in class  $\mathbf{y}_i$  given input  $\mathbf{x}$ . If we wish to *confuse* the class outputs of the classifier for the pair  $\mathbf{x}_1$  and  $\mathbf{x}_2$ , we should learn parameters  $\theta$  that bring these conditional probability distributions “closer” under some distance metric, that is, make the predictions for  $\mathbf{x}_1$  and  $\mathbf{x}_2$  similar.

While KL-divergence might seem to be a reasonable choice to design a loss function for optimizing the distance between conditional probability distributions, in Section 3.1, we show that it is infeasible to train a neural network when using KL-divergence as a regularizer. Therefore, we introduce the Euclidean Distance between distributions as a metric for confusion in Sections 3.2 and 3.3 and describe neural network training with this metric in Section 3.4.

### 3.1 Symmetric KL-divergence or Jeffrey’s Divergence

The most prevalent method to measure dissimilarity of one probability distribution from another is to use the Kullback-Liebler (KL) divergence. However, the standard KL-divergence cannot serve our purpose owing to its asymmetric nature. This could be remedied by using the *symmetric* KL-divergence, defined for two probability distributions  $P, Q$  with mass functions  $p(\cdot), q(\cdot)$  (for events  $u \in \mathcal{U}$ ):

$$\mathbb{D}_J(P, Q) \triangleq \sum_{u \in \mathcal{U}} \left[ p(u) \cdot \log \frac{p(u)}{q(u)} + q(u) \cdot \log \frac{q(u)}{p(u)} \right] = \mathbb{D}_{\text{KL}}(P||Q) + \mathbb{D}_{\text{KL}}(Q||P) \quad (1)$$

This symmetrized version of KL-divergence, known as Jeffrey’s divergence [40], is a measure of the average relative entropy between two probability distributions [41]. For our model parameterized by  $\theta$ , for samples  $\mathbf{x}_1$  and  $\mathbf{x}_2$ , the Jeffrey’s divergence can be written as:

$$\mathbb{D}_J(p_\theta(\mathbf{y}|\mathbf{x}_1), p_\theta(\mathbf{y}|\mathbf{x}_2)) = \sum_{i=1}^N \left[ (p_\theta(\mathbf{y}_i|\mathbf{x}_1) - p_\theta(\mathbf{y}_i|\mathbf{x}_2)) \cdot \log \frac{p_\theta(\mathbf{y}_i|\mathbf{x}_1)}{p_\theta(\mathbf{y}_i|\mathbf{x}_2)} \right] \quad (2)$$

Jeffrey’s divergence satisfies all of our basic requirements of a symmetric divergence metric between probability distributions, and therefore could be included as a regularizing term while training with cross-entropy, to achieve our desired confusion. However, when we learn model parameters using stochastic gradient descent (SGD), it can be difficult to train, especially if our distributions  $P, Q$  have mass concentrated on different events. This can be seen in Equation 2. Consider Jeffrey’s divergence with  $N = 2$  classes, and that  $\mathbf{x}_1$  belongs to class 1, and  $\mathbf{x}_2$  belongs to class 2. If the model parameters  $\theta$  are such that it correctly identifies both  $\mathbf{x}_1$  and  $\mathbf{x}_2$  by training using cross-entropy loss,  $p_\theta(\mathbf{y}_1|x_1) = 1 - \delta_1$  and  $p_\theta(\mathbf{y}_2|x_2) = 1 - \delta_2$ , where  $0 < \delta_1, \delta_2 < \frac{1}{2}$  (since the classifier outputs correct predictions for the input images), we can show:

$$\mathbb{D}_J(p_\theta(\mathbf{y}|\mathbf{x}_1), p_\theta(\mathbf{y}|\mathbf{x}_2)) \geq (1 - \delta_1 - \delta_2) \cdot (2 \log(1 - \delta_1 - \delta_2) - \log(\delta_1 \delta_2)) \quad (3)$$

Please see the supplementary material for an expanded proof.

As training progresses with these labels, the cross-entropy loss will motivate the values of  $\delta_1$  and  $\delta_2$  to become closer to zero (but never equaling zero, since the probability outputs  $p_\theta(\mathbf{y}|\mathbf{x}_1), p_\theta(\mathbf{y}|\mathbf{x}_2)$  are the outputs from a softmax). As  $(\delta_1, \delta_2) \rightarrow (0^+, 0^+)$ , the second term  $-\log(\delta_1 \delta_2)$  on the R.H.S. of inequality (3) typically grows whereas  $(1 - \delta_1 - \delta_2)$  approaches 1, which makes  $\mathbb{D}_J(p_\theta(\mathbf{y}|\mathbf{x}_1), p_\theta(\mathbf{y}|\mathbf{x}_2))$  larger as the predictions get closer to the true labels. In practice, we see that training with  $\mathbb{D}_J(p_\theta(\mathbf{y}|\mathbf{x}_1), p_\theta(\mathbf{y}|\mathbf{x}_2))$  as a regularizer term diverges, unless a very small regularizing parameter is chosen, which removes the effect of regularization altogether.

A natural question that can arise from this analysis is that cross-entropy training itself involves optimizing KL-divergence between the target label distribution and the model’s predictions, however no such divergence occurs. This is because cross-entropy involves only one direction of the KL-divergence, and the target distribution has all the mass concentrated at one event (the correct label). Since  $(x \log x)|_{x=0} = 0$ , for predicted label vector  $\mathbf{y}'$  with correct label class  $c$ , this simplifies the cross-entropy error  $\mathcal{L}_{CE}(p_\theta(\mathbf{y}|\mathbf{x}), \mathbf{y}')$  to be:

$$\mathcal{L}_{CE}(p_\theta(\mathbf{y}|\mathbf{x}), \mathbf{y}') = - \sum_{i=1}^N \mathbf{y}'_i \log\left(\frac{p_\theta(\mathbf{y}_i|\mathbf{x})}{\mathbf{y}'_i}\right) = -\log(p_\theta(\mathbf{y}_c|\mathbf{x})) \geq 0 \quad (4)$$

This formulation does not diverge as the model trains, i.e.  $p_\theta(\mathbf{y}_c|\mathbf{x}) \rightarrow 1$ . In some cases where label noise is added to the label vector (such as label smoothing [28,42]), the label noise is a fixed constant and not approaching zero (as in the case of Jeffery’s divergence between model predictions) and is hence feasible to train. Thus, Jeffrey’s Divergence or symmetric KL-divergence, while a seemingly natural choice, cannot be used to train a neural network with SGD. This motivates us to look for an alternative metric to measure “confusion” between conditional probability distributions.

### 3.2 Euclidean Distance as Confusion

Since the conditional probability distribution over  $N$  classes is an element within  $\mathbb{R}^N$  on the unit simplex, we can consider the Euclidean distance to be a metric

of “confusion” between two conditional probability distributions. Analogous to the previous setting, we define the **Euclidean Confusion**  $\mathbb{D}_{\text{EC}}(\cdot, \cdot)$  for a pair of inputs  $\mathbf{x}_1, \mathbf{x}_2$  with model parameters  $\theta$  as:

$$\mathbb{D}_{\text{EC}}(p_\theta(\mathbf{y}|\mathbf{x}_1), p_\theta(\mathbf{y}|\mathbf{x}_2)) = \sum_{i=1}^N (p_\theta(\mathbf{y}_i|\mathbf{x}_1) - p_\theta(\mathbf{y}_i|\mathbf{x}_2))^2 = \|p_\theta(\mathbf{y}|\mathbf{x}_1) - p_\theta(\mathbf{y}|\mathbf{x}_2)\|_2^2 \quad (5)$$

Unlike Jeffrey’s Divergence, Euclidean Confusion does not diverge when used as a regularization term with cross-entropy. However, to verify this unconventional choice for a distance metric between probability distributions, we prove some properties that relate Euclidean Confusion to existing divergence measures.

**Lemma 1.** *On a finite probability space, the Euclidean Confusion  $\mathbb{D}_{\text{EC}}(P, Q)$  is a lower bound for the Jeffrey’s Divergence  $\mathbb{D}_{\text{J}}(P, Q)$  for probability measures  $P, Q$ .*

*Proof.* This follows from Pinsker’s Inequality and the relationship between  $\ell_1$  and  $\ell_2$  norms. Complete proof is provided in the supplementary material.

By Lemma 1, we can see that the Euclidean Confusion is a conservative estimate for Jeffrey’s divergence, the earlier proposed divergence measure. For finite probability spaces, the Total Variation Distance  $\mathbb{D}_{\text{TV}}(P, Q)^2 = \frac{1}{2}\|P - Q\|_1$  is also a measure of interest. However, due to its non-differentiable nature, it is unsuitable for our case. Nevertheless, we can relate the Euclidean Confusion and Total Variation Distance by the following result.

**Lemma 2.** *On a finite probability space, the Euclidean Confusion  $\mathbb{D}_{\text{EC}}(P, Q)$  is bounded by  $4\mathbb{D}_{\text{TV}}(P, Q)^2$  for probability measures  $P, Q$ .*

*Proof.* This follows directly from the relationship between  $\ell_1$  and  $\ell_2$  norms. Complete proof is provided in the supplementary material.

### 3.3 Euclidean Confusion for Point Sets

In a standard classification setting with  $N$  classes, we consider a training set with  $m = \sum_{i=1}^N m_i$  training examples, where  $m_i$  denotes the number of training samples for class  $i$ . For this setting, we can write the total Euclidean Confusion between points of classes  $i$  and  $j$  as the average of the Euclidean Confusion between all pairs of points belonging to those two classes. For simplicity of notation, let us denote the set of conditional probability distributions of all training points belonging to class  $i$  for a model parameterized by  $\theta$  as  $\mathcal{S}_i = \{p_\theta(\mathbf{y}|\mathbf{x}_1^i), p_\theta(\mathbf{y}|\mathbf{x}_2^i), \dots, p_\theta(\mathbf{y}|\mathbf{x}_{m_i}^i)\}$ . Then, for a model parameterized by  $\theta$ , the Euclidean Confusion is given by:

$$\mathbb{D}_{\text{EC}}(\mathcal{S}_i, \mathcal{S}_j; \theta) \triangleq \frac{1}{m_i m_j} \left( \sum_{u,v}^{m_i, m_j} \mathbb{D}_{\text{EC}}(p_\theta(\mathbf{y}|\mathbf{x}_u^i), p_\theta(\mathbf{y}|\mathbf{x}_v^j)) \right) \quad (6)$$

We can simplify this equation by assuming an equal number of points  $n$  per class:

$$\mathbb{D}_{\text{EC}}(\mathcal{S}_i, \mathcal{S}_j; \theta) = \frac{1}{n^2} \left( \sum_{u,v}^{n,n} \|p_\theta(\mathbf{y}|\mathbf{x}_u^i) - p_\theta(\mathbf{y}|\mathbf{x}_v^j)\|_2^2 \right) \quad (7)$$

This form of the Euclidean Confusion between the two sets of points gives us an interesting connection with another popular distance metric over probability distributions, known as the **Energy Distance** [43].

Introduced by Gabor Szekely [43], the **Energy Distance**  $\mathbb{D}_{\text{EN}}(F, G)$  between two cumulative probability distribution functions  $F$  and  $G$  with random vectors  $X$  and  $Y$  in  $\mathbb{R}^N$  can be given by

$$\mathbb{D}_{\text{EN}}(F, G)^2 \triangleq 2\mathbb{E}\|X - Y\| - \mathbb{E}\|X - X'\| - \mathbb{E}\|Y - Y'\| \geq 0 \quad (8)$$

where  $(X, X', Y, Y')$  are independent, and  $X \sim F, X' \sim F, Y \sim G, Y' \sim G$ . If we consider the sets  $\mathcal{S}_i$  and  $\mathcal{S}_j$ , with a uniform probability of selecting any of the  $n$  points in each of these sets, then we obtain the following results.

**Lemma 3.** *For sets  $\mathcal{S}_i, \mathcal{S}_j$  and  $\mathbb{D}_{\text{EC}}(\mathcal{S}_i, \mathcal{S}_j; \theta)$  as defined in Equation (7):*

$$\frac{1}{2}\mathbb{D}_{\text{EN}}(\mathcal{S}_i, \mathcal{S}_j; \theta)^2 \leq \mathbb{D}_{\text{EC}}(\mathcal{S}_i, \mathcal{S}_j; \theta)$$

where  $\mathbb{D}_{\text{EN}}(\mathcal{S}_i, \mathcal{S}_j; \theta)$  is the Energy Distance under Euclidean norm between  $\mathcal{S}_i$  and  $\mathcal{S}_j$  (parameterized by  $\theta$ ), and random vectors are selected with uniform probability in both  $\mathcal{S}_i$  and  $\mathcal{S}_j$ .

*Proof.* This follows from the definition of Energy Distance with uniform probability of sampling. Complete proof is provided in the supplementary material.

**Corollary 1.** *For sets  $\mathcal{S}_i, \mathcal{S}_j$  and  $\mathbb{D}_{\text{EC}}(\mathcal{S}_i, \mathcal{S}_j; \theta)$  as defined in Equation (7), we have:*

$$\mathbb{D}_{\text{EC}}(\mathcal{S}_i, \mathcal{S}_i; \theta) + \mathbb{D}_{\text{EC}}(\mathcal{S}_j, \mathcal{S}_j; \theta) \leq 2\mathbb{D}_{\text{EC}}(\mathcal{S}_i, \mathcal{S}_j; \theta)$$

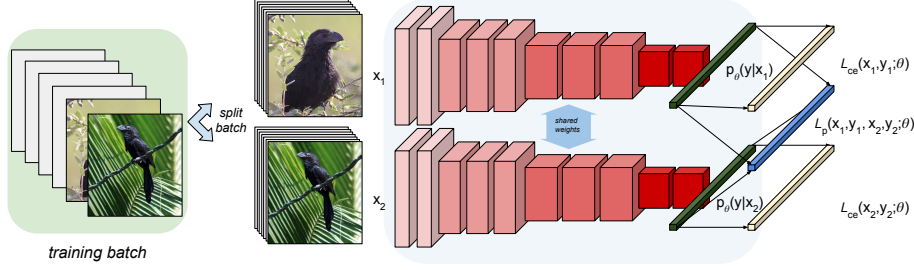
with equality only when  $\mathcal{S}_i = \mathcal{S}_j$ .

*Proof.* This follows from the fact that the Energy Distance  $\mathbb{D}_{\text{EN}}(\mathcal{S}_i, \mathcal{S}_j; \theta)$  is 0 only when  $\mathcal{S}_i = \mathcal{S}_j$ . The complete version of the proof is included in the supplement.

With these results, we restrict the behavior of Euclidean Confusion within two well-defined conventional probability distance measures, the Jeffrey’s divergence and Energy Distance. One might consider optimizing the Energy Distance directly, due to its similar formulation and the fact that we uniformly sample points during training with SGD. However, the Energy Distance additionally includes the two terms that account for the negative of the average all-pairs distances between points in  $\mathcal{S}_i$  and  $\mathcal{S}_j$  respectively, which we do not want to maximize, since we do not wish to push points within the same class further apart. Therefore, we proceed with our measure of Euclidean Confusion.

### 3.4 Learning with Gradient Descent

We proceed to learn parameters  $\theta^*$  for a neural network, with the following learning objective function for a pair of input points, motivated by the formulation



**Fig. 1.** CNN training pipeline for Pairwise Confusion (**PC**). We employ a Siamese-like architecture, with individual cross entropy calculations for each branch, followed by a joint energy-distance minimization loss. We split each incoming batch of samples into two mini-batches, and feed the network pairwise samples.

of Euclidean Confusion:

$$\theta^* = \arg \min_{\theta} \sum_{\substack{i=1, j \neq i \\ u, v}}^{N, N} \left[ \mathcal{L}_{\text{CE}}(p_{\theta}(\mathbf{y}|\mathbf{x}_u^i), \mathbf{y}_u^i) + \mathcal{L}_{\text{CE}}(p_{\theta}(\mathbf{y}|\mathbf{x}_v^j), \mathbf{y}_v^j) + \frac{\lambda}{n^2} \mathbb{D}_{\text{EC}}(p_{\theta}(\mathbf{y}|\mathbf{x}_v^j), p_{\theta}(\mathbf{y}|\mathbf{x}_u^i)) \right] \quad (9)$$

This objective function can be explained as: for each point in the training set, we randomly select another point from a different class and calculate the individual cross-entropy losses and Euclidean Confusion until all pairs have been exhausted. For each point in the training dataset, there are  $n \cdot (N - 1)$  valid choices for the other point, giving us a total of  $n^2 \cdot N \cdot (N - 1)$  possible pairs. In practice, we find that we do not need to exhaust all combinations for effective learning using gradient descent, and in fact we observe that convergence is achieved far before all observations are observed. We simplify our formulation instead by using the following procedure described in Algorithm 1.

**Training Procedure:** As described in Algorithm 1, our learning procedure is a slightly modified version of the standard SGD. We randomly permute the training set twice, and then for each pair of points in the training set, add Euclidean Confusion only if the samples belong to different classes. This form of sampling approximates the exhaustive Euclidean Confusion, with some points with regular gradient descent, which in practice does not alter the performance. Moreover, convergence is achieved after only a fraction of all the possible pairs are observed. Formally, we wish to model the conditional probability distribution  $p_{\theta}(\mathbf{y}|\mathbf{x})$  over the  $p$  classes for function  $f(\mathbf{x}; \theta) = p_{\theta}(\mathbf{y}|\mathbf{x})$  parameterized by model parameters  $\theta$ . Given our optimization procedure, we can rewrite the total loss for a pair of

**Algorithm 1** Training Using Euclidean Confusion

---

```

Training data  $D$ , Test data  $\hat{D}$ , parameters  $\theta$ , hyperparameters  $\hat{\theta}$ 
for  $epoch \in [0, \text{max\_epochs}]$  do
   $D_1 \leftarrow \text{shuffle}(D)$ 
   $D_2 \leftarrow \text{shuffle}(D)$ 
  for  $i \in [0, \text{num\_batches}]$  do
     $\mathcal{L}_{\text{batch}} = 0$ 
    for  $(d_1, d_2) \in \text{batch } i \text{ of } (D_1, D_2)$  do
       $\gamma \leftarrow 1$  if  $\text{label}(d_1) \neq \text{label}(d_2)$ , 0 otherwise
       $\mathcal{L}_{\text{pair}} \leftarrow \mathcal{L}_{\text{CE}}(d_1; \theta) + \mathcal{L}_{\text{CE}}(d_2; \theta) + \lambda \cdot \gamma \cdot \mathbb{D}_{\text{EC}}(d_1, d_2; \theta)$ 
       $\mathcal{L}_{\text{batch}} \leftarrow \mathcal{L}_{\text{batch}} + \mathcal{L}_{\text{pair}}$ 
    end for
     $\theta \leftarrow \text{Backprop}(\mathcal{L}_{\text{batch}}, \theta, \hat{\theta})$ 
  end for
   $\hat{\theta} \leftarrow \text{ParameterUpdate}(\text{epoch}, \hat{\theta})$ 
end for

```

---

points  $\mathbf{x}_1, \mathbf{x}_2$  with model parameters  $\theta$  as:

$$\mathcal{L}_{\text{pair}}(\mathbf{x}_1, \mathbf{x}_2, \mathbf{y}_1, \mathbf{y}_2; \theta) = \sum_{i=1}^2 [\mathcal{L}_{\text{CE}}(p_{\theta}(\mathbf{y}|\mathbf{x}_i), \mathbf{y}_i)] + \lambda \gamma(\mathbf{y}_1, \mathbf{y}_2) \mathbb{D}_{\text{EC}}(p_{\theta}(\mathbf{y}|\mathbf{x}_1), p_{\theta}(\mathbf{y}|\mathbf{x}_2)) \quad (10)$$

where,  $\gamma(\mathbf{y}_1, \mathbf{y}_2) = 1$  when  $\mathbf{y}_i \neq \mathbf{y}_j$ , and 0 otherwise. We denote **training** with this general architecture with the term *Pairwise Confusion* or **PC** for short. Specifically, we train a Siamese-like neural network [22] with shared weights, training each network individually using cross-entropy, and add the **Euclidean Confusion** loss between the conditional probability distributions obtained from each network (Figure 1). During training, we split an incoming batch of training samples into two parts, and evaluating cross-entropy on each sub-batch identically, followed by a pairwise loss term calculated for corresponding pairs of samples across batches. During testing, only one branch of the network is active, and generates output predictions for the input image.

**CNN Architectures** We experiment with VGGNet [44], GoogLeNet [42], ResNets [12], and DenseNets [11] as base architectures for the Siamese network trained with **PC** to demonstrate that our method is insensitive to the choice of source architecture.

## 4 Experimental Details

We perform all experiments using Caffe [45] or PyTorch [46] over a cluster of NVIDIA Titan X, Tesla K40c and GTX 1080 GPUs. Our code and models are available at [github.com/abhimanyudubey/confusion](https://github.com/abhimanyudubey/confusion). Next, we provide brief descriptions of the various datasets used in our paper.

**Table 2.** Pairwise Confusion (**PC**) obtains state-of-the-art performance on six widely-used fine-grained visual classification datasets (A-F). Improvement over the baseline model is reported as ( $\Delta$ ). All results averaged over 5 trials.

| (A) CUB-200-2011         |              |          | (B) Cars                |              |          | (C) Aircrafts            |              |          |
|--------------------------|--------------|----------|-------------------------|--------------|----------|--------------------------|--------------|----------|
| Method                   | Top-1        | $\Delta$ | Method                  | Top-1        | $\Delta$ | Method                   | Top-1        | $\Delta$ |
| Gao <i>et al.</i> [16]   | 84.00        | -        | Wang <i>et al.</i> [17] | 85.70        | -        | Simon <i>et al.</i> [49] | 85.50        | -        |
| STN[2]                   | 84.10        | -        | Liu <i>et al.</i> [48]  | 86.80        | -        | Cui <i>et al.</i> [9]    | 86.90        | -        |
| Zhang <i>et al.</i> [47] | 84.50        | -        | Lin <i>et al.</i> [8]   | 92.00        | -        | LRBP [50]                | 87.30        | -        |
| Lin <i>et al.</i> [8]    | 85.80        | -        | Cui <i>et al.</i> [9]   | 92.40        | -        | Lin <i>et al.</i> [8]    | 88.50        | -        |
| Cui <i>et al.</i> [9]    | 86.20        | -        | ResNet-50               | 91.71        | -        | ResNet-50                | 81.19        | -        |
| ResNet-50                | 78.15        | (2.06)   | <b>PC-ResNet-50</b>     | <b>93.43</b> | (1.72)   | <b>PC-ResNet-50</b>      | 83.40        | (2.21)   |
| <b>PC-ResNet-50</b>      | 80.21        | -        | Bilinear CNN [1]        | 91.20        | -        | BilinearCNN [1]          | 84.10        | -        |
| Bilinear CNN [1]         | 84.10        | (1.48)   | <b>PC-Bilinear CNN</b>  | 92.45        | (1.25)   | <b>PC-BilinearCNN</b>    | 85.78        | (1.68)   |
| <b>PC-BilinearCNN</b>    | 85.58        | -        | DenseNet-161            | 91.83        | -        | DenseNet-161             | 86.30        | -        |
| DenseNet-161             | 84.21        | -        | <b>PC-DenseNet-161</b>  | 92.86        | (1.03)   | <b>PC-DenseNet-161</b>   | <b>89.24</b> | (2.94)   |
| <b>PC-DenseNet-161</b>   | <b>86.87</b> | (2.66)   |                         |              |          |                          |              |          |

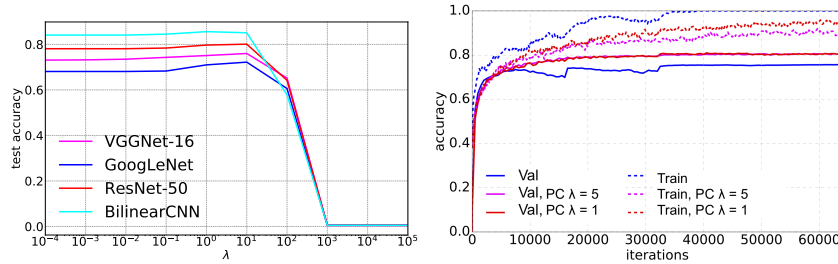
| (D) NABirds                |              |          | (E) Flowers-102        |              |          | (F) Stanford Dogs        |              |          |
|----------------------------|--------------|----------|------------------------|--------------|----------|--------------------------|--------------|----------|
| Method                     | Top-1        | $\Delta$ | Method                 | Top-1        | $\Delta$ | Method                   | Top-1        | $\Delta$ |
| Branson <i>et al.</i> [19] | 35.70        | -        | Det.+Seg. [51]         | 80.66        | -        | Zhang <i>et al.</i> [3]  | 80.43        | -        |
| Van <i>et al.</i> [35]     | 75.00        | -        | Overfeat [52]          | 86.80        | -        | Krause <i>et al.</i> [6] | 80.60        | -        |
| ResNet-50                  | 63.55        | -        | ResNet-50              | 92.46        | -        | ResNet-50                | 69.92        | -        |
| <b>PC-ResNet-50</b>        | 68.15        | (4.60)   | <b>PC-ResNet-50</b>    | 93.50        | (1.04)   | <b>PC-ResNet-50</b>      | 73.35        | (3.43)   |
| BilinearCNN [1]            | 80.90        | -        | BilinearCNN [1]        | 92.52        | -        | BilinearCNN [1]          | 82.13        | -        |
| <b>PC-BilinearCNN</b>      | 82.01        | (1.11)   | <b>PC-BilinearCNN</b>  | <b>93.65</b> | (1.13)   | <b>PC-BilinearCNN</b>    | 83.04        | (0.91)   |
| DenseNet-161               | 79.35        | -        | DenseNet-161           | 90.07        | -        | DenseNet-161             | 81.18        | -        |
| <b>PC-DenseNet-161</b>     | <b>82.79</b> | (3.44)   | <b>PC-DenseNet-161</b> | 91.39        | (1.32)   | <b>PC-DenseNet-161</b>   | <b>83.75</b> | (2.57)   |

#### 4.1 Fine-Grained Visual Classification (FGVC) datasets

1. **Wildlife Species Classification:** We experiment with several widely-used FGVC datasets. The Caltech-UCSD Birds (**CUB-200-2011**) dataset [33] has 5,994 training and 5,794 test images across 200 species of North-American birds. The **NABirds** dataset [35] contains 23,929 training and 24,633 test images across over 550 visual categories, encompassing 400 species of birds, including separate classes for male and female birds in some cases. The **Stanford Dogs** dataset [37] has 20,580 images across 120 breeds of dogs around the world. Finally, the **Flowers-102** dataset [32] consists of 1,020 training, 1,020 validation and 6,149 test images over 102 flower types.

2. **Vehicle Make/Model Classification:** We experiment with two common vehicle classification datasets. The **Stanford Cars** dataset [34] contains 8,144 training and 8,041 test images across 196 car classes. The classes represent variations in car make, model, and year. The **Aircraft** dataset is a set of 10,000 images across 100 classes denoting a fine-grained set of airplanes of different varieties [36].

These datasets contain (i) large visual diversity in each class [32,33,37], (ii) visually similar, often confusing samples belonging to different classes, and (iii) a large variation in the number of samples present per class, leading to greater class imbalance than LSVC datasets like **ImageNet** [10]. Additionally, some of these datasets have densely annotated part information available, which we do not utilize in our experiments.



**Fig. 2.** (left) Variation of test accuracy on CUB-200-2011 with logarithmic variation in hyperparameter  $\lambda$ . (right) Convergence plot of GoogLeNet on CUB-200-2011.

## 5 Results

### 5.1 Fine-Grained Visual Classification

We first describe our results on the six FGVC datasets from Table 2. In all experiments, we average results over 5 trials per experiment—after choosing the best value of hyperparameter  $\lambda$ . Please see the supplementary material for mean and standard deviation values for all experiments.

- 1. Fine-tuning from Baseline Models:** We fine-tune from three baseline models using the PC optimization procedure: ResNet-50 [12], Bilinear CNN [1], and DenseNet-161 [11]. As Tables 2-(A-F) show, PC obtains substantial improvement across all datasets and models. For instance, a baseline DenseNet-161 architecture obtains an average accuracy of 84.21%, but **PC**-DenseNet-161 obtains an accuracy of 86.87%, an improvement of **2.66%**. On NABirds, we obtain improvements of **4.60%** and **3.42%** over baseline ResNet-50 and DenseNet-161 architectures.
- 2. Combining PC with Specialized FGVC models:** Recent work in FGVC has proposed several novel CNN designs that take part-localization into account, such as bilinear pooling techniques [16,1,9] and spatial transformer networks [2]. We train a Bilinear CNN [1] with PC, and obtain an average improvement of 1.7% on the 6 datasets.

We note two important aspects of our analysis: (1) we do not compare with ensembling and data augmentation techniques such as Boosted CNNs [21] and Krause, *et al.* [6] since prior evidence indicates that these techniques invariably improve performance, and (2) we evaluate a single-crop, single-model evaluation without any part- or object-annotations, and perform competitively with methods that use both augmentations.

**Choice of Hyperparameter  $\lambda$ :** Since our formulation requires the selection of a hyperparameter  $\lambda$ , it is important to study the sensitivity of classification performance to the choice of  $\lambda$ . We conduct this experiment for four different models: GoogLeNet [42], ResNet-50 [12] and VGGNet-16 [44] and Bilinear-CNN [1] on the CUB-200-2011 dataset. PC’s performance is **not very sensitive to the choice of  $\lambda$**  (Figure 2 and Supplementary Tables S1-S5). For all six

datasets, the  $\lambda$  value is typically between the range [10,20]. On Bilinear CNN, setting  $\lambda = 10$  for all datasets gives average performance within 0.08% compared to the reported values in Table 2. In general, PC obtains optimum performance in the range of  $0.05N$  and  $0.15N$ , where  $N$  is the number of classes.

## 5.2 Additional Experiments

Since our method aims to improve classification performance in FGVC tasks by introducing confusion in output logit activations, we would expect to see a larger improvement in datasets with higher inter-class similarity and intra-class variation. To test this hypothesis, we conduct two additional experiments.

In the first experiment, we construct two subsets of ImageNet-1K [10]. The first dataset, **ImageNet-Dogs** is a subset consisting only of species of dogs (117 classes and 116K images). The second dataset, **ImageNet-Random** contains randomly selected classes from ImageNet-1K. Both datasets contain equal number of classes (117) and images (116K), but ImageNet-Dogs has much higher inter-class similarity and intra-class variation, as compared to ImageNet-Random. To test repeatability, we construct 3 instances of ImageNet-Random, by randomly choosing a different subset of ImageNet with 117 classes each time. For both experiments, we randomly construct a 80-20 train-val split from the training data to find optimal  $\lambda$  by cross-validation, and report the performance on the unseen ImageNet validation set of the subset of chosen classes. In Table 3, we compare the performance of training from scratch with- and without-PC across three models: GoogLeNet, ResNet-50, and DenseNet-161. As expected, PC obtains a larger gain in classification accuracy (1.45%) on ImageNet-Dogs as compared to the ImageNet-Random dataset ( $0.54\% \pm 0.28$ ).

In the second experiment, we utilize the CIFAR-10 and CIFAR-100 datasets, which contain the same number of total images. CIFAR-100 has  $10\times$  the number of classes and 10% of images per class as CIFAR-10 and contains larger inter-class similarity and intra-class variation. We train networks on both datasets from scratch using default train-test splits (Table 3). As expected, we obtain larger average gains of 1.77% on CIFAR-100, as compared to 0.20% on CIFAR-10. Additionally, when training with  $\lambda = 10$  on the entire ImageNet dataset, we obtain a top-1 accuracy of 76.28% (compared to a baseline of 76.15%), which is a smaller improvement, which is in line with what we would expect for a large-scale image classification problem with large inter-class variation.

Moreover, while training with PC, we observe that the rate of convergence is always similar to or faster than training without PC. For example, a GoogLeNet trained on CUB-200-2011 (Figure 2(right) above) shows that PC converges to higher validation accuracy faster than normal training using identical learning rate schedule and batch size. Note that the training accuracy is reduced when training with PC, due to the regularization effect. In sum, classification problems that have large intra-class variation and high inter-class similarity benefit from optimization with pairwise confusion. The improvement is even more prominent when training data is limited.

**Table 3.** Experiments with ImageNet and CIFAR show that datasets with large intra-class variation and high inter-class similarity benefit from optimization with Pairwise Confusion. Only the mean accuracy over 3 Imagenet-Random experiments is shown.

| Network           | ImageNet-Random |       | ImageNet-Dogs |       | CIFAR-10 |       | CIFAR-100 |       |
|-------------------|-----------------|-------|---------------|-------|----------|-------|-----------|-------|
|                   | Baseline        | PC    | Baseline      | PC    | Baseline | PC    | Baseline  | PC    |
| GoogLeNet [42]    | 71.85           | 72.09 | 62.35         | 64.17 | 86.63    | 87.02 | 73.35     | 76.02 |
| ResNet-50 [12]    | 82.01           | 82.65 | 73.81         | 75.92 | 93.17    | 93.46 | 72.16     | 73.14 |
| DenseNet-161 [11] | 78.34           | 79.10 | 70.15         | 71.44 | 95.15    | 95.08 | 78.60     | 79.56 |

**Table 4.** Pairwise Confusion (**PC**) improves localization performance in fine-grained visual classification tasks. On the CUB-200-2011 dataset, PC obtains an average improvement of 3.4% in Mean Intersection-over-Union (IoU) for Grad-CAM bounding boxes for each of the five baseline models.

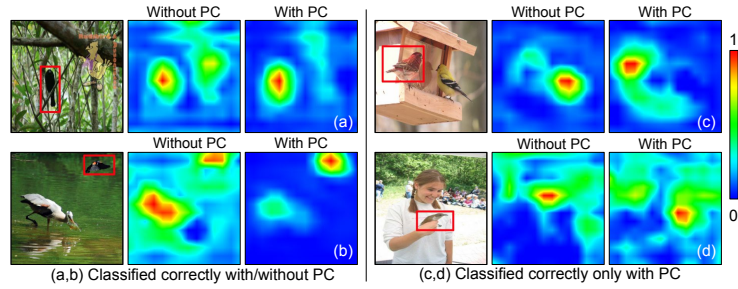
| Method               | GoogLeNet | VGG-16 | ResNet-50 | DenseNet-161 | Bilinear-CNN |
|----------------------|-----------|--------|-----------|--------------|--------------|
| Mean IoU (Baseline)  | 0.29      | 0.31   | 0.32      | 0.34         | 0.37         |
| Mean IoU (PC) - Ours | 0.35      | 0.34   | 0.35      | 0.37         | 0.39         |

### 5.3 Improvement in Localization Ability

Recent techniques for improving classification performance in fine-grained recognition are based on summarizing and extracting dense localization information in images [1,2]. Since our technique increases classification accuracy, we wish to understand if the improvement is a result of enhanced CNN localization abilities due to PC. To measure the regions the CNN localizes on, we utilize Gradient-Weighted Class Activation Mapping (Grad-CAM) [53], a method that provides a heatmap of visual saliency as produced by the network. We perform both quantitative and qualitative studies of localization ability of PC-trained models.

**Overlap in Localized Regions:** To quantify the improvement in localization due to PC, we construct bounding boxes around object regions obtained from Grad-CAM, by thresholding the heatmap values at 0.5, and choosing the largest box returned. We then calculate the mean IoU (intersection-over-union) of the bounding box with the provided object bounding boxes for the CUB-200-2011 dataset. We compare the mean IoU across several models, with and without PC. As summarized in Table 4, we observe an average 3.4% improvement across five different networks, implying better localization accuracy.

**Change in Class-Activation Mapping:** To qualitatively study the improvement in localization due to PC, we obtain samples from the CUB-200-2011 dataset and visualize the localization regions returned from Grad-CAM for both the baseline and PC-trained VGG-16 model. As shown in Figure 3, PC models provide tighter, more accurate localization around the target object, whereas sometimes the baseline model has localization driven by image artifacts. Figure 3-(a) has an example of the types of distractions that are often present in FGVC images (the cartoon bird on the right). We see that the baseline VGG-16 network pays



**Fig. 3.** Pairwise Confusion (PC) obtains improved localization performance, as demonstrated here with Grad-CAM heatmaps of the CUB-200-2011 dataset images (left) with a VGGNet-16 model trained without PC (middle) and with PC (right). The objects in (a) and (b) are correctly classified by both networks, and (c) and (d) are correctly classified by PC, but not the baseline network (VGG-16). For all cases, we consistently observe a tighter and more accurate localization with PC, whereas the baseline VGG-16 network often latches on to artifacts, even while making correct predictions.

significant attention to the distraction, despite making the correct prediction. With PC, we find that the attention is limited almost exclusively to the correct object, as desired. Similarly for Figure 3-(b), we see that the baseline method latches on to the incorrect bird category, which is corrected by the addition of PC. In Figures 3-(c-d), we see that the baseline classifier makes incorrect decisions due to poor localization, mistakes that are resolved by PC.

## 6 Conclusion

In this work, we introduce Pairwise Confusion (PC), an optimization procedure to improve generalizability in fine-grained visual classification (FGVC) tasks by encouraging confusion in output activations. PC improves FGVC performance for a wide class of convolutional architectures while fine-tuning. Our experiments indicate that PC-trained networks show improved localization performance which contributes to the gains in classification accuracy. PC is easy to implement, does not need excessive tuning during training, and does not add significant overhead during test time, in contrast to methods that introduce complex localization-based pooling steps that are often difficult to implement and train. Therefore, our technique should be beneficial to a wide variety of specialized neural network models for applications that demand for fine-grained visual classification or learning from limited labeled data.

**Acknowledgements:** We would like to thank Dr. Ashok Gupta for his guidance on bird recognition, and Dr. Sumeet Agarwal, Spandan Madan and Ishaan Grover for their feedback at various stages of this work. RF and PG were supported in part by the National Science Foundation under Grant No. IIS1651832, and AD, OG, RR and NN acknowledge the generous support of the MIT Media Lab Consortium.

## References

1. Lin, T.Y., RoyChowdhury, A., Maji, S.: Bilinear cnn models for fine-grained visual recognition. *IEEE International Conference on Computer Vision* (2015) 1449–1457
2. Jaderberg, M., Simonyan, K., Zisserman, A., Kavukcuoglu, K.: Spatial transformer networks. *Advances in Neural Information Processing Systems* (2015) 2017–2025
3. Zhang, Y., Wei, X.S., Wu, J., Cai, J., Lu, J., Nguyen, V.A., Do, M.N.: Weakly supervised fine-grained categorization with part-based image representation. *IEEE Transactions on Image Processing* **25**(4) (2016) 1713–1725
4. Krause, J., Jin, H., Yang, J., Fei-Fei, L.: Fine-grained recognition without part annotations. *IEEE Conference on Computer Vision and Pattern Recognition* (2015) 5546–5555
5. Zhang, N., Shelhamer, E., Gao, Y., Darrell, T.: Fine-grained pose prediction, normalization, and recognition. *International Conference on Learning Representations Workshops* (2015)
6. Krause, J., Sapp, B., Howard, A., Zhou, H., Toshev, A., Duerig, T., Philbin, J., Fei-Fei, L.: The unreasonable effectiveness of noisy data for fine-grained recognition. *European Conference on Computer Vision* (2016) 301–320
7. Cui, Y., Zhou, F., Lin, Y., Belongie, S.: Fine-grained categorization and dataset bootstrapping using deep metric learning with humans in the loop. *IEEE Conference on Computer Vision and Pattern Recognition* (2016)
8. Lin, T.Y., Maji, S.: Improved bilinear pooling with cnns. *arXiv preprint arXiv:1707.06772* (2017)
9. Cui, Y., Zhou, F., Wang, J., Liu, X., Lin, Y., Belongie, S.: Kernel pooling for convolutional neural networks. *IEEE Conference on Computer Vision and Pattern Recognition* (2017)
10. Deng, J., Dong, W., Socher, R., Li, L.J., Li, K., Fei-Fei, L.: Imagenet: A large-scale hierarchical image database. *IEEE Conference on Computer Vision and Pattern Recognition* (2009) 248–255
11. Huang, G., Liu, Z., van der Maaten, L., Weinberger, K.Q.: Densely connected convolutional networks. *IEEE Conference on Computer Vision and Pattern Recognition* (2017)
12. He, K., Zhang, X., Ren, S., Sun, J.: Deep residual learning for image recognition. *IEEE Conference on Computer Vision and Pattern Recognition* (2016) 770–778
13. Yao, B., Khosla, A., Fei-Fei, L.: Combining randomization and discrimination for fine-grained image categorization. *IEEE Conference on Computer Vision and Pattern Recognition* (2011) 1577–1584
14. Yao, B., Bradschi, G., Fei-Fei, L.: A codebook-free and annotation-free approach for fine-grained image categorization. *IEEE Conference on Computer Vision and Pattern Recognition* (2012) 3466–3473
15. Zhang, N., Donahue, J., Girshick, R., Darrell, T.: Part-based r-cnns for fine-grained category detection. *European Conference on Computer Vision* (2014) 834–849
16. Gao, Y., Beijbom, O., Zhang, N., Darrell, T.: Compact bilinear pooling. *IEEE Conference on Computer Vision and Pattern Recognition* (2016) 317–326
17. Wang, Y., Choi, J., Morariu, V., Davis, L.S.: Mining discriminative triplets of patches for fine-grained classification. *IEEE Conference on Computer Vision and Pattern Recognition* (June 2016)
18. Ren, S., He, K., Girshick, R., Sun, J.: Faster r-cnn: Towards real-time object detection with region proposal networks. *Advances in neural information processing systems* (2015) 91–99

19. Branson, S., Van Horn, G., Belongie, S., Perona, P.: Bird species categorization using pose normalized deep convolutional nets. *British Machine Vision Conference* (2014)
20. Zhang, N., Farrell, R., Darrell, T.: Pose pooling kernels for sub-category recognition. *IEEE Computer Vision and Pattern Recognition* (2012) 3665–3672
21. Moghimi, M., Saberian, M., Yang, J., Li, L.J., Vasconcelos, N., Belongie, S.: Boosted convolutional neural networks. (2016)
22. Chopra, S., Hadsell, R., LeCun, Y.: Learning a similarity metric discriminatively, with application to face verification. *IEEE Conference on Computer Vision and Pattern Recognition* (2005) 539–546
23. Parikh, D., Grauman, K.: Relative attributes. *IEEE International Conference on Computer Vision* (2011) 503–510
24. Dubey, A., Agarwal, S.: Modeling image virality with pairwise spatial transformer networks. *arXiv preprint arXiv:1709.07914* (2017)
25. Souri, Y., Noury, E., Adeli, E.: Deep relative attributes. *Asian Conference on Computer Vision* (2016) 118–133
26. Dubey, A., Naik, N., Parikh, D., Raskar, R., Hidalgo, C.A.: Deep learning the city: Quantifying urban perception at a global scale. *European Conference on Computer Vision* (2016) 196–212
27. Singh, K.K., Lee, Y.J.: End-to-end localization and ranking for relative attributes. *European Conference on Computer Vision* (2016) 753–769
28. Reed, S., Lee, H., Anguelov, D., Szegedy, C., Erhan, D., Rabinovich, A.: Training deep neural networks on noisy labels with bootstrapping. *arXiv preprint arXiv:1412.6596* (2014)
29. Xiao, T., Xia, T., Yang, Y., Huang, C., Wang, X.: Learning from massive noisy labeled data for image classification. *IEEE Conference on Computer Vision and Pattern Recognition* (2015) 2691–2699
30. Neelakantan, A., Vilnis, L., Le, Q.V., Sutskever, I., Kaiser, L., Kurach, K., Martens, J.: Adding gradient noise improves learning for very deep networks. *arXiv preprint arXiv:1511.06807* (2015)
31. Szegedy, C., Vanhoucke, V., Ioffe, S., Shlens, J., Wojna, Z.: Rethinking the inception architecture for computer vision. *IEEE Conference on Computer Vision and Pattern Recognition* (2016) 2818–2826
32. Nilsback, M.E., Zisserman, A.: Automated flower classification over a large number of classes. *Indian Conference on Computer Vision, Graphics & Image Processing* (2008) 722–729
33. Wah, C., Branson, S., Welinder, P., Perona, P., Belongie, S.: The caltech-ucsd birds-200-2011 dataset. (2011)
34. Krause, J., Stark, M., Deng, J., Fei-Fei, L.: 3d object representations for fine-grained categorization. *IEEE International Conference on Computer Vision Workshops* (2013) 554–561
35. Van Horn, G., Branson, S., Farrell, R., Haber, S., Barry, J., Ipeirotis, P., Perona, P., Belongie, S.: Building a bird recognition app and large scale dataset with citizen scientists: The fine print in fine-grained dataset collection. *IEEE Conference on Computer Vision and Pattern Recognition* (2015) 595–604
36. Maji, S., Rahtu, E., Kannala, J., Blaschko, M., Vedaldi, A.: Fine-grained visual classification of aircraft. *arXiv preprint arXiv:1306.5151* (2013)
37. Khosla, A., Jayadevaprakash, N., Yao, B., Li, F.F.: Novel dataset for fine-grained image categorization: Stanford dogs. *IEEE International Conference on Computer Vision Workshops on Fine-Grained Visual Categorization* (2011) 1

38. Krizhevsky, A., Nair, V., Hinton, G.: The cifar-10 dataset otkrist (2014)
39. Netzer, Y., Wang, T., Coates, A., Bissacco, A., Wu, B., Ng, A.Y.: Reading digits in natural images with unsupervised feature learning. NIPS workshop on deep learning and unsupervised feature learning (2) (2011) 5
40. Jeffreys, H.: The theory of probability. OUP Oxford (1998)
41. Kullback, S., Leibler, R.A.: On information and sufficiency. The annals of mathematical statistics **22**(1) (1951) 79–86
42. Szegedy, C., Liu, W., Jia, Y., Sermanet, P., Reed, S., Anguelov, D., Erhan, D., Vanhoucke, V., Rabinovich, A.: Going deeper with convolutions. IEEE Conference on Computer Vision and Pattern Recognition (2015) 1–9
43. Székely, G.J., Rizzo, M.L.: Energy statistics: A class of statistics based on distances. Journal of statistical planning and inference **143**(8) (2013) 1249–1272
44. Simonyan, K., Zisserman, A.: Very deep convolutional networks for large-scale image recognition. arXiv preprint arXiv:1409.1556 (2014)
45. Jia, Y., Shelhamer, E., Donahue, J., Karayev, S., Long, J., Girshick, R., Guadarrama, S., Darrell, T.: Caffe: Convolutional architecture for fast feature embedding. ACM international conference on Multimedia (2014) 675–678
46. Paskze, A., Chintala, S.: Tensors and Dynamic neural networks in Python with strong GPU acceleration. <https://github.com/pytorch> Accessed: [January 1, 2017].
47. Zhang, X., Xiong, H., Zhou, W., Lin, W., Tian, Q.: Picking deep filter responses for fine-grained image recognition. IEEE Conference on Computer Vision and Pattern Recognition (2016) 1134–1142
48. Liu, M., Yu, C., Ling, H., Lei, J.: Hierarchical joint cnn-based models for fine-grained cars recognition. International Conference on Cloud Computing and Security (2016) 337–347
49. Simon, M., Gao, Y., Darrell, T., Denzler, J., Rodner, E.: Generalized orderless pooling performs implicit salient matching. International Conference on Computer Vision (ICCV) (2017)
50. Kong, S., Fowlkes, C.: Low-rank bilinear pooling for fine-grained classification. IEEE Conference on Computer Vision and Pattern Recognition (2017) 7025–7034
51. Angelova, A., Zhu, S.: Efficient object detection and segmentation for fine-grained recognition. IEEE Conference on Computer Vision and Pattern Recognition (2013) 811–818
52. Sharif Razavian, A., Azizpour, H., Sullivan, J., Carlsson, S.: CNN features off-the-shelf: An astounding baseline for recognition. IEEE Conference on Computer Vision and Pattern Recognition Workshops (June 2014)
53. Selvaraju, R.R., Das, A., Vedantam, R., Cogswell, M., Parikh, D., Batra, D.: Grad-cam: Why did you say that? visual explanations from deep networks via gradient-based localization. arXiv preprint arXiv:1610.02391 (2016)