

Zero-Shot Visual Recognition using Semantics-Preserving Adversarial Embedding Networks

Long Chen¹ Hanwang Zhang^{2,3} Jun Xiao^{1*} Wei Liu³ Shih-Fu Chang⁴

¹Zhejiang University ²Nanyang Technological University ³Tencent AI Lab ⁴Columbia University

{longc, junx}@zju.edu.cn; hanwangzhang@ntu.edu.sg; {wliu, sfchang}@ee.columbia.edu

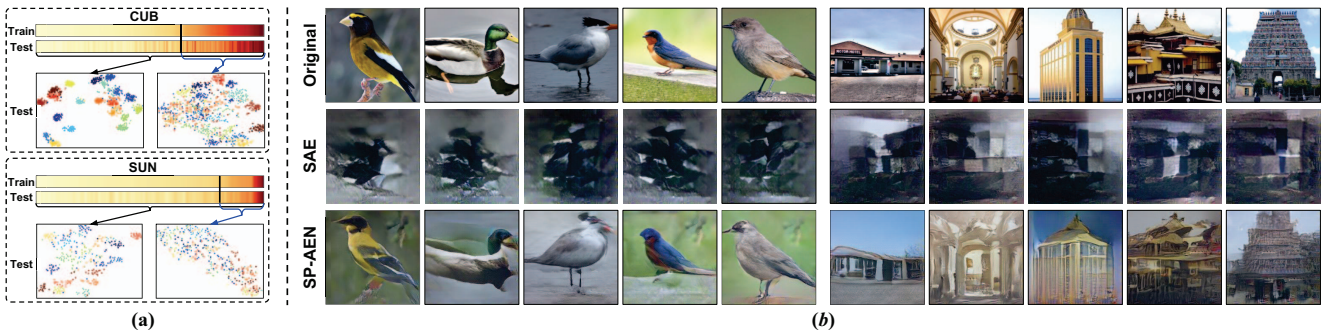


Figure 1: (a) Attribute variance heat maps of the 312 attributes in CUB birds [59] and the 102 attributes in SUN scenes [46] (lighter color indicates lower variance, *i.e.*, lower discriminability) and the t-SNE [34] visualizations of the test images represented by all attributes (left) and only the high-variance ones (right). Some of the low-variance attributes (the lighter part to the left of the cut-off line) discarded at training are still needed in discriminating unseen test classes. (b) Comparison of reconstructed images using SAE [25] and our proposed SP-AEN method, which is shown to retain sufficient semantics for photo-realistic reconstruction.

Abstract

We propose a novel framework called Semantics-Preserving Adversarial Embedding Network (SP-AEN) for zero-shot visual recognition (ZSL), where test images and their classes are both unseen during training. SP-AEN aims to tackle the inherent problem — semantic loss — in the prevailing family of embedding-based ZSL, where some semantics would be discarded during training if they are non-discriminative for training classes, but could become critical for recognizing test classes. Specifically, SP-AEN prevents the semantic loss by introducing an independent visual-to-semantic space embedder which disentangles the semantic space into two subspaces for the two arguably conflicting objectives: classification and reconstruction. Through adversarial learning of the two subspaces, SP-AEN can transfer the semantics from the reconstructive subspace to the discriminative one, accomplishing the improved zero-shot recognition of unseen classes. Comparing

with prior works, SP-AEN can not only improve classification but also generate photo-realistic images, demonstrating the effectiveness of semantic preservation. On four popular benchmarks: CUB, AWA, SUN and aPY, SP-AEN considerably outperforms other state-of-the-art methods by an absolute performance difference of 12.2%, 9.3%, 4.0%, and 3.6% in terms of harmonic mean values [62].

1. Introduction

Zero-shot visual recognition, or more generally, zero-shot learning (ZSL), recognizes novel classes that are unseen at training stage. The community has reached a consensus that ZSL is all about transferring knowledge from seen classes to unseen classes; Despite that there are fruitful ZSL methods, the transfer still follows the simple but intuitive mechanism: although “raccoon” is unseen, we can recognize it by checking if it satisfies the “raccoon signature”, *e.g.*, visual attributes “striped tail” [13, 27, 67], classeme “fox-like” [57, 30, 65, 52], or “raccoon” word vectors [47, 37]. These attributes can be modeled at training

*Corresponding Author

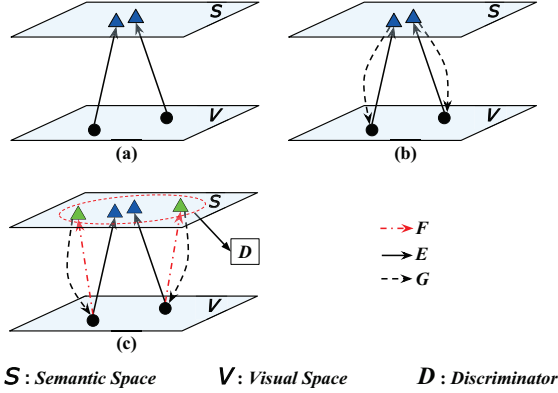


Figure 2: Three investigated ZSL paradigms. (a) Conventional visual-to-semantic mapping E trained on classification loss. (b) Semantic autoencoder [25], visual-to-semantic E and semantic-to-visual G trained on both classification and reconstruction losses. (c) The proposed SP-AEN, introducing an independent visual-to-semantic F and an adversarial-style discriminator D between the two sub-space embeddings (blue and green triangular).

stage and are expected to be sharable in both seen and unseen classes at test stage. After a decade of progress, the transfer has evolved from primitive attribute classifiers [27] to semantic embedding based framework [1, 14, 60], which is prevailing due to its simple and effective paradigm (cf. Figure 2 (a)): first, it maps images from visual space to semantic space where all the classes reside; then, ZSL is reduced to a simple nearest neighbor search — the image is assigned to the nearest class in embedding space.

The semantic transfer ability of this embedding-based ZSL framework is limited by the **semantic loss** problem. As shown in Figure 1, discarding the low-variance attributes (*i.e.*, less discriminative) is beneficial to classification at training; However, due to the semantic discrepancy between seen and unseen classes, these attributes would be discriminative at test time, resulting in a lossy semantic space that is problematic for unseen class recognition. The main reason is that although the class embedding has rich semantic meanings, it is still a lonely point in the semantic space, where the mappings of many images will inevitably collapse to it [36, 15]. One may consider the extreme case that all the class embeddings are one-hot label vectors, degenerating to the traditional supervised classification, therefore, no semantics can be transferred.

An arguably possible solution is to preserve semantics by reconstruction — the embedded semantic vector from one image should be able to map the image back, where any two semantic embeddings are expected to preserve sufficient semantics to be apart, otherwise the reconstruction would fail [24, 63, 69, 19]. However, reconstruction and classification are essentially two conflicting objectives: the

former aims to preserve as many image details as possible while the latter focuses on suppressing irrelevant ones. For example, using only “head” and “torso” attributes might be sufficient for “person” recognition while the color attributes “red” and “white” are indeed disturbing. To further illustrate this, as shown in Figure 2 (b), suppose $E: V \rightarrow S$ and $G: S \rightarrow V$ are two mapping transformations between the visual and semantic spaces. For classification, we want $x, x' \in V$ of the same class to be mapped to close semantic embeddings $s, s' \in S$, *i.e.*, $E(x) = s \approx s' = E(x')$; For reconstruction, we want $G(s) \approx x$ and $G(s') \approx x'$, which is difficult to be satisfied as $s \approx s'$. Therefore, joint training of the two objectives is ineffective to preserve semantics (*e.g.*, SAE [25]). For example, as illustrated in Figure 1 (b), if we want to achieve good classification performance, the reconstruction will fail generally.

To resolve this conflict, we propose a novel visual-semantic embedding framework: Semantics-Preserving Adversarial Embedding Network (SP-AEN). As illustrated in Figure 2 (c), we introduce a new mapping $F: V \rightarrow S$ and an adversarial objective [17] where the discriminator D and encoder F try to make $F(x)$ and $E(x)$ indistinguishable. There are two benefits of introducing F and D to help E preserve semantics: 1) **Semantic Transfer**. Even though the semantic loss is inevitable by E , we can avoid it using F by borrowing ingredients from $E(x)$ of other classes, and the discriminator D will eventually transfer semantics from $F(x)$ to $E(x)$ by tailoring the two semantic embeddings into the same distribution. For example, for a “bird” image where the attribute “spotty” in $E(\text{bird})$ is lost, we can retain it by using $E(\text{leopard})$ because “spotty” is a discriminative and preserved attribute in “leopard” images. 2) **Disentangled Classification and Reconstruction**. As the reconstruction is only imposed to F and G , E is disentangled to focus on classification. In this way, the conflict between classification and reconstruction is resolved because the constraint $G(E(x)) \approx x$ and $G(E(x')) \approx x'$ is relaxed to $G(F(x)) \approx x$ and $G(F(x')) \approx x'$, as $F(x)$ and $F(x')$ are not necessarily to be close with each other to comply with the discriminative objective as E . As shown in Figure 1 (b), compared to the reconstruction style in Figure 2 (b) [25], our visual-semantic embedding $G(F(x))$ can reconstruct photo-realistic images, suggesting that the semantic is better preserved.

We can deploy state-of-the-art network structures for SP-AEN in a flexible plug-and-play and end-to-end fine-tune fashion, *e.g.*, E may use the powerful model for classification [21], F and G may use the encoder and decoder of the image generation architecture [11]. The overall architecture is illustrated in Figure 3 and will be detailed in Section 4.1. We validate the effectiveness of SP-AEN on four popular benchmarks: CUB [59], AWA [27], SUN [46], and aPY [13], surpassing the state-of-the-art performances [62]

by 12.2%, 9.3%, 4.0%, and 3.6% in harmonic mean values, respectively. To the best of our knowledge, SP-AEN is the first ZSL model that empowers photo-realistic image generation from the semantic space. We hope that it will facilitate the ZSL community for better visual investigations of knowledge transfer.

2. Related Work

Zero-Shot Learning One main stream of ZSL is the attribute-based visual recognition [13, 27, 49, 43, 10, 22] where the attributes serve as an intermediate feature space that transfer semantics across classes, supporting zero-shot recognition of unseen classes. To scale up ZSL, embedding based methods are prevailing [14, 2, 3, 49, 61, 55, 25, 31]. These methods directly learn a mapping from the image visual space to a semantic space, represented by semantic vectors such as word vectors [37, 47, 42] or textual descriptions [29, 12, 9]. Our proposed SP-AEN is an embedding based ZSL that exploits the ranking based classification loss as [14]. However, to the best of our knowledge, SP-AEN is the first ZSL method that can reconstruct images from the semantic embeddings. The evaluation used in the experiments follows a similar setting for practical ZSL applications [8, 62]. Note that ZSL is also closely related to few-shot learning [18] and domain adaptation [40, 45], where both of them assume that a small number of training images given in the test classes; however, no image is exposed to test classes at training in ZSL.

Domain Shift and Hubness. Similar problems to the semantic loss have been reported in other terms. Domain shift [51, 15] is a generic problem that resides in all types of visual recognition, where the data from train and test are in different distributions. Hubness [36] states the phenomenon that the mapped semantic embeddings from images would be collapsed to hubs, which are near many other points without being similar to the class label in any meaningful way. We believe that semantic loss is one of the main reason for hubness, which can be alleviated by reconstruction [24, 63, 69, 19]. In this paper, we find that jointly training [25] reconstruction and classification is not effective to preserve semantics. Another way of countering semantic loss is to learn independent attribute classifiers [39], which is not applicable when attribute annotation is unavailable.

Generative Adversarial Network (GAN). The idea of GAN [17] is to train a generator that can fool a discriminator to confuse the distributions of the generated and true samples. In theory, this max-min training procedure can lead the generator to perfectly model the data distribution. SP-AEN is similar to the GAN applied in the feature-level [44, 58, 35, 53]. Recently, several ZSL models adopt generative model for data augmentation of unseen classes [38, 6]. However, they violate the ZSL assumption that the unseen class is prohibitively seen at training.

Image Generation. We seek algorithms that can generate perceptually realistic images [16, 23, 5, 32, 33, 56]. Besides pixel-level loss, these methods impose feature-level reconstruction loss for preserving perceptual similarity or adversarial loss to remove unreal artifacts. However, they are based on image-to-image transformation while we requires that the reconstruction is from the semantic embedding. Our reconstruction network relates to image generation from a bottleneck layer [11, 41, 66, 48].

3. Formulation

We start by formalizing the ZSL task and then introduce the training objectives of the proposed SP-AEN.

3.1. Preliminaries

Given a set of training set $\{x_i, l_i\}$, where $x_i \in \mathcal{V}$ is an image represented in the visual space, and $l_i \in \mathcal{L}_s$ is a class label in the *seen* class set, the goal of ZSL is to learn a classifier which can generalize to predict any image x at test stage to its correct label, which is not only in $\mathcal{L}_s$ but also in the *unseen* class set $\mathcal{L}_u$. As summarized in [62, 29], almost all types of ZSL methods can be unified into the embedding-based framework: we hope to find a visual-to-semantic mapping $E: \mathcal{V} \rightarrow \mathcal{S}$, where any class label l is embedded as $\mathbf{y}_l \in \mathbb{R}^d$ in the semantic space $\mathcal{S}$ (e.g., an attribute space). Therefore, the predicted label l^* can be obtained by following simple nearest neighbor search:

$$l^* = \max_{l \in \mathcal{L}} \mathbf{y}_l^T E(x). \quad (1)$$

In particular, if $l \in \mathcal{L}_u$, it is the *conventional* ZSL setting; if $l \in \mathcal{L}_s \cup \mathcal{L}_u$, it is the *generalized* ZSL setting, which is more practical for real applications. It is worth noting that Eq. (1) is not necessarily a linear model; in fact, as we will introduce in Section 4.1, it can also be highly non-linear in nature by using deep neural networks to implement E .

3.2. Classification Objective

As label prediction in Eq. (1) is fundamentally a ranking problem, we use a large-margin based ranking loss function for classification objective [14, 60], i.e., given a training pair (x, l) we want a higher dot-product similarity between $\mathbf{y}_l$ and $E(x)$ and a lower one for any wrongly labeled pair (x, l') , and the similarity margin between the correct one and the wrong one should be larger than a threshold:

$$L_{cls} = \sum_{l \neq l'} \max\{0, \gamma - \mathbf{y}_l^T E(x) + \mathbf{y}_{l'}^T E(x)\}. \quad (2)$$

Where $\gamma > 0$ is a hyperparameter for the margin. At each iteration in stochastic training, only one l' is randomly selected from the unpaired labels.

As mentioned in Section 1, the classification objective L_{cls} essentially forces the semantic embedding $E(x)$ of all

Where α and β are trade-off parameters for different objectives. The final goal is to solve:

$$E^* = \arg \min_{E, F, G} \max_D L(E, F, G, D). \quad (6)$$

As illustrated in Figure 3. By considering F as the encoder and G as the decoder, then the semantic embedding $F(x)$ can be considered as the bottleneck layer, regularized to match a supervised distribution $E(x)$. In this way, our SP-AEN is a supervised Adversarial Autoencoder (AAE) [35]. Therefore, SP-AEN has the potential flexibility to reform into other ZSL frameworks, *e.g.*, semi-supervised SP-AEN, by imposing another adversarial objective for $F(x)$ to match a prior embedding space. We leave this for our future work.

4. Implementation

4.1. Architecture

The overall architecture is detailed in Figure 3. It is an end-to-end network with the input of raw images and ground truth class embeddings. The embedder E is based on ResNet-101 [21] takes a cropped $224 \times 224 \times 3$ image as input and outputs a d -dimensional embedding vector, which is then fed into the classification loss function in Eq. (2). F is based on AlexNet [26] appended with two more fully-connected blocks that take the raw image as input and output a d -dimensional embedding vector, which is fed into the subsequent reconstruction network G . We adapt the architecture for G from [11] who has shown impressive results for generating images from a bottleneck layer. G contains five up-convolutional blocks with leaky ReLU [20] for transforming a vector into a 3-D feature map, which is eventually equal to the RGB color map. In particular, we append two fully-connected layers as the head of G that takes the embedding vector as input and then output a 4,096-dimensional vector that can match the input of [11]. D is a two-layer fully-connected layer plus a non-linear ReLU layer that takes the d -dimensional embedding vector as input.

4.2. Training Details

For all the experiments in this paper, the training images are resized with its shorter side to 256; ten 224×224 cropped image data augmentation trick is adopted with per-pixel mean subtraction [26]. For efficiency, we fixed the ResNet-101 in E , and initialized the AlexNet-like blocks in F with AlexNet and G with the pretrained generator [11]. Then, the rest of the modules was trained from scratch with MSRA random initializer [20]. The learning rate started from $1e^{-4}$ and is multiplied by 0.1 when the error is plateaus. We use grid search to select parameter α and β .

5. Experiments

5.1. Datasets

We used four popular benchmarks described as below. In particular, we followed the new split provided by [62] as the ILSVRC [50] 1K ImageNet classes, widely used as a pre-training source for CNN features, have already included the test classes in the conventional split of the benchmarks, hence violating the fundamental assumption of ZSL that the classes at test should be strictly unseen at training.

CUB [59]. It is the Caltech-UCSD-Birds 200-2011 dataset of 11,788 bird images from 200 fine-grained classes. Each image is annotated with 312 semantic attributes. The train split has 7,057 images across 150 classes (50 classes for validation); The test split has 1,764 images from the 150 seen classes and 2,967 images from the 50 unseen classes.

SUN [46]. It is a fine-grained scene data of 14,340 images across 717 scene classes. Each image is annotated with 102 semantic attributes. The train split has 10,320 images from 645 classes (65 classes for validation); The test split has 2,580 images from the 645 seen classes and 1,440 images from the 72 unseen classes.

AWA [27]. It is the coarse-grained Animals with Attributes dataset of 30,475 images from 50 animals. Each class is annotated with 85 semantic attributes. The train split has 23,527 images from 40 classes (13 classes for validation); The test split has 5,882 images from the 40 seen classes and 7,913 images from the 10 unseen classes. In particular, we used the AWA2 released by [62] as the images from the original one are restricted due to photo copyright reasons.

aPY [13]. It is the coarse-grained Attribute Pascal and Yahoo dataset of 12,051 images from 32 generic object classes (*e.g.*, 20 Pascal classes and 12 popular Yahoo classes). Each image is annotated with 64 semantic attributes. The train split has 5,932 images from 20 classes (5 classes for validation); The test split has 1,483 images from the 20 seen classes and 7,924 from the 12 unseen classes.

For fair comparison and reproductivity, we used the class label embeddings provided by [62], each of which is an L2-normalized vector.

5.2. Settings and Evaluation Metrics

To evaluate the ZSL performances over all classes, we applied the following three settings. 1) $U \rightarrow U$: The test images and the prediction labels set are limited to the unseen classes; 2) $S \rightarrow T$: The test images are from the seen classes and the prediction labels set is the union of both seen and unseen classes; 3) $U \rightarrow T$: The test images are from the unseen classes and the prediction labels set is the union of both seen and unseen classes. Note that $U \rightarrow U$ and $U \rightarrow T$ are also known as the conventional and the generalized ZSL settings.

We followed [62] and used the per-class top-1 accuracy

	SUN (0.9851)				CUB (0.9575)				AWA (0.7459)				aPY (0.5847)			
Setting	U→U	U→T	S→T	H	U→U	U→T	S→T	H	U→U	U→T	S→T	H	U→U	U→T	S→T	H
DAP [27]	39.9	4.2	25.1	7.2	40.0	1.7	67.9	3.3	46.1	0.0	84.7	0.0	33.8	4.8	78.3	9.0
IAP [27]	19.4	1.0	37.8	1.8	24.0	0.2	72.8	0.4	35.9	0.9	87.6	1.8	36.6	5.7	65.6	10.4
SSE [68]	51.5	2.1	36.4	4.0	43.9	8.5	46.9	14.4	61.0	8.1	82.5	14.8	34.0	0.2	78.9	0.4
CONSE [43]	38.8	6.8	39.9	11.6	34.3	1.6	72.2	3.1	44.5	0.5	90.6	1.0	26.9	0.0	91.2	0.0
SYNC [7]	56.3	7.9	43.3	13.4	55.6	11.5	70.9	19.8	46.6	10.0	90.5	18.0	23.9	7.4	66.3	13.3
CMT [55]	39.9	8.1	21.8	11.8	34.6	7.2	49.8	12.6	37.9	0.5	90.0	1.0	28.0	1.4	85.2	2.8
CMT* [55]	—	8.7	28.0	13.3	—	4.7	60.1	8.7	—	8.7	89.0	15.9	—	10.9	74.2	19.0
LATEM [61]	55.3	14.7	28.8	19.5	49.3	15.2	57.3	24.0	55.8	11.5	77.3	20.0	35.2	0.1	73.0	0.2
DeViSE [14]	56.5	16.9	27.4	20.9	52.0	23.8	53.0	32.8	59.7	17.1	74.7	27.8	39.8	4.9	76.9	9.2
ALE [2]	58.1	21.8	33.1	26.3	54.9	23.7	62.8	34.4	62.5	14.0	81.8	23.9	39.7	4.6	73.7	8.7
SJE [3]	53.7	14.7	30.5	19.8	53.9	23.5	59.2	33.6	61.9	8.0	73.9	14.4	32.9	3.7	55.7	6.9
ESZSL [49]	54.5	11.0	27.9	15.8	53.9	12.6	63.8	21.0	58.6	5.9	77.8	11.0	38.3	2.4	70.1	4.6
SAE [25]	40.3	8.8	18.0	11.8	33.3	7.8	54.0	13.6	54.1	1.1	82.2	2.2	8.3	0.4	80.9	0.9
SP-AEN	59.2	24.9	38.6	30.3	55.4	34.7	70.6	46.6	58.5	23.3	90.9	37.1	24.1	13.7	63.4	22.6

Table 1: Performances (accuracy% and H%) of all the comparing methods under the three settings on the four datasets. Cosine similarity between the attribute variances of the disjoint train images and test images are given in brackets. As demonstrated in Figure 1, lower similarity indicates larger semantic loss.

as the evaluation metric, where the prediction using Eq. (1) is successful if the predicted class is the correct ground truth. We averaged the accuracies of all classes. For generalized ZSL setting, we also used the recently proposed harmonic mean (H) [62] of accuracies on seen classes $\mathcal{L}_s$ ($Acc_{S \rightarrow T}$) and unseen classes $\mathcal{L}_u$ ($Acc_{U \rightarrow T}$):

$$H = 2 \times Acc_{S \rightarrow T} \times Acc_{U \rightarrow T} / (Acc_{S \rightarrow T} + Acc_{U \rightarrow T}). \quad (7)$$

H offers a comprehensive metric in the practical ZSL case: in many real applications, the test image would belong to any class from both seen and unseen sets, and it is required that the accuracy should be high on both of them. Note that we slightly abuse the setting notations: $U \rightarrow U$, $S \rightarrow T$, $U \rightarrow T$ as the accuracy calculated in the corresponding setting.

5.3. Comparisons with State-of-The-Arts

Comparing Methods. We compared SP-AEN with a variety of ZSL methods as reported in [62]. These methods fall into two categories. 1) *embedding based*: DeVISE [14], ALE [2], SJE [3], ESZSL [49], LATEM [61], CMT/CMT* [55], SAE [25]. As SP-AEN does, this category maps images into the semantic embedding space where all class labels reside. Note that CMT* is CMT with novelty detection and hence is not applicable for $U \rightarrow U$ setting. To the best of our knowledge, SAE is the only ZSL method that uses reconstruction to tackle the semantic loss problem. 2) *attribute based*: DAP [27], IAP [27], SSE [68], CONSE [43], and SYNC [7]. These methods are based on an intermediate inference of attributes in ZSL. Note that this category cannot be applied in generic class embeddings without the attribute annotations.

Results. Table 1 summarizes the performances (accuracy% and H%) of all the comparing methods under the three settings on the four datasets. We have the following two key

observations: 1) Using the generalized ZSL setting metric ($U \rightarrow T$ and H), SP-AEN significantly outperforms the best competitors by around 4% to 12%. In particular, we can clearly see that the performance gap between SP-AEN and others becomes larger as the cosine similarity between the attribute variances of the disjoint train and test splits. As larger cosine similarity indicates smaller semantic loss, it demonstrates the effectiveness of SP-AEN in alleviating the semantic loss of ZSL. 2) Under the conventional ZSL setting ($U \rightarrow U$), in most cases, SP-AEN achieves the best performance. This is reasonable as the search space of label prediction is merely limited to the unseen sets, however, the semantic loss may cause the semantic mapping of unseen class images very similar to one of the seen classes, resulting in incorrect recognition.

Method	SUN	CUB	AWA	aPY
DirectMap	0.079	0.069	0.075	0.085
SAE	0.285	0.281	0.259	0.275
SplitBranch	0.070	0.058	0.059	0.076
SP-AEN	0.053	0.040	0.047	0.055

Table 3: The mean squared pixel-level loss between the input images and its reconstructed images of various reconstruction settings over four datasets.

5.4. Ablation Studies

5.4.1 Conflict between Classification & Reconstruction

To validate our key motivation for the design of SP-AEN: The conflict between classification and reconstruction, as illustrated in Figure 6. We propose three possible architectures that can achieve the semantic-to-visual reconstruction as SP-AEN: 1) **DirectMap**: For each input image, we use E to get its discriminative semantic embedding and then

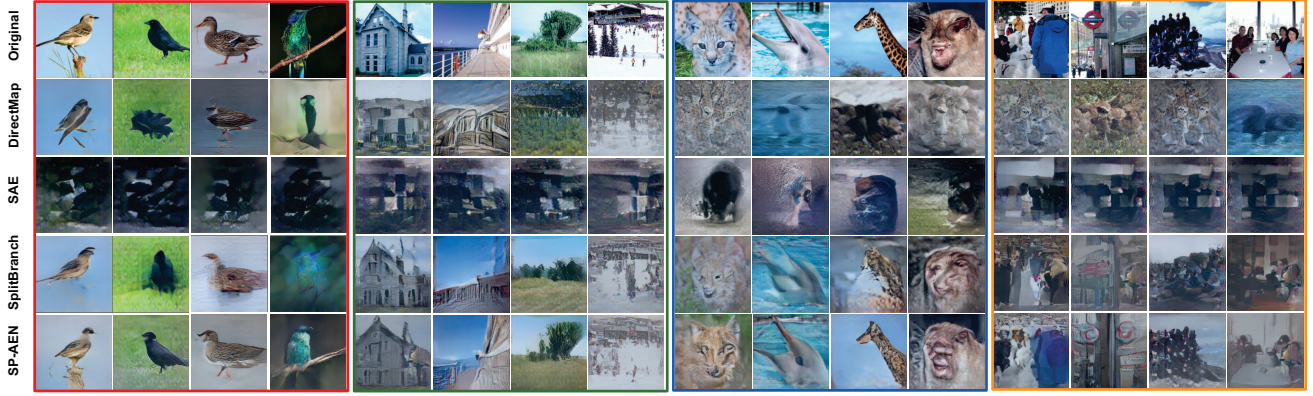


Figure 4: Example reconstruction results of various architectures on CUB, SUN, AWA and aPY respectively.

	SUN				CUB				AWA				aPY			
Setting	U→U	U→T	S→T	H	U→U	U→T	S→T	H	U→U	U→T	S→T	H	U→U	U→T	S→T	H
Cls. Only	56.8	17.2	29.0	21.6	52.2	23.5	55.0	32.9	60.2	17.5	76.7	28.5	35.8	5.5	72.9	10.2
Full Obj.	59.2	24.9	38.6	30.3	55.4	34.7	70.6	46.6	58.5	23.3	90.9	37.1	24.1	13.7	63.4	22.6

Table 2: Performances (accuracy% and H%) of all the comparing methods under the three settings on the four datasets.

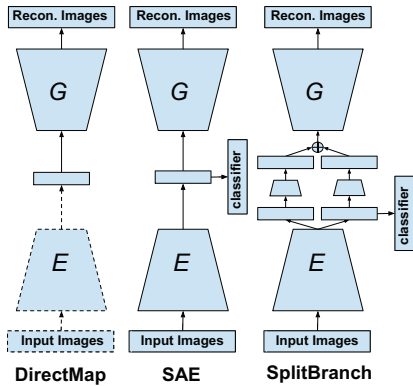


Figure 6: Three reconstruction architectures evaluated. Dashed line indicates fixed parameters at training.

use G to map it back to the image space. In this architecture, we fix E and train G . DirectMap is used to evaluate how much reconstructive semantics are preserved in the discriminatively trained semantic embedding. 2) **SAE** [25]: We adapt the SAE architecture using our image reconstructor G as the decoder and E as the encoder. The bottleneck layer as semantic embedding is used for classification. We jointly train E and G . 3) **SplitBranch**: We split the semantic output of encoder E into two branches, and only the first semantic embedding branch is used for classification. Two semantic embedding branches are concatenated after two respective fully-connected layers. The merged semantic representation is fed to decoder G to reconstruct image.

Reconstruction Results. Figure 4 shows some reconstructed images and Table 3 reports the reconstruction losses of the unseen images in the test splits of the four dataset-

s. We can have the following observations: 1) On CUB and SUN, the images reconstructed by DirectMap are close to SP-AEN, which have the highest quality. However, on AWA and aPY, the reconstruction quality of DirectMap significantly drops. Again, this is due to the semantic loss between train and test set, as the attribute variance cosine similarities of AWA and aPY are much larger than those of CUB and SUN. 2) If we jointly train the discriminative E and reconstructive G as in SAE, the reconstruction fails in all examples; If we jointly train them as in SplitBranch, we can observe significant quality improvement, closing to SP-AEN. However, we find that the weight for merging the semantic embedding of the classification branch is almost zero, meaning the contribution of the semantic embedding for reconstruction is minor. Therefore, the semantic transfer is ineffective. This motivates us to use adversarial loss in SP-AEN to allow semantic transfer and high-quality reconstruction at the same time.

5.4.2 Effectiveness of D and G

Since the score of seen classes is usually larger than that of unseen classes, a calibrated stacking rule [8] to solve it by subtracting a bias for seen classes to solve it:

$$l^* = \max_{l \in \mathcal{L}_u \cup \mathcal{L}_s} \mathbf{y}_l^T E(x) - \gamma [l \in \mathcal{L}_s].$$

Where the indicator function $[\cdot]$ indicates whether or not l is a seen class and $\gamma \in \mathbb{R}$ is a calibration factor. This calibrated stacking rule represents a middle ground between aggressively classifying each data point into seen classes and conservatively classifying every data point into unseen

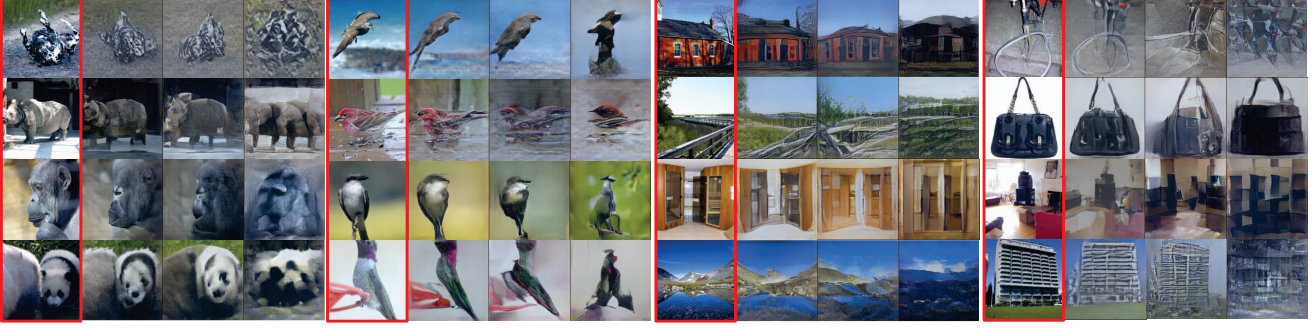


Figure 5: Reconstruction results of different α in AWA, CUB, SUN and aPY respectively. The left image in the red bounding box represents the original images, and decreases weights α in Eq. 5 from left to right.

classes. By varying the calibration factor γ , we can compute a series of classification accuracies ($Acc_{U \rightarrow T}$ and $Acc_{S \rightarrow T}$) and plot the Seen-Unseen accuracy Curve (SUC). The Area Under Seen-Unseen Accuracy Curve (AUSUC) is always used as a performance metric to show the balance capability between this two conflicting objectives $Acc_{U \rightarrow T}$ and $Acc_{S \rightarrow T}$ in the generalized ZSL.

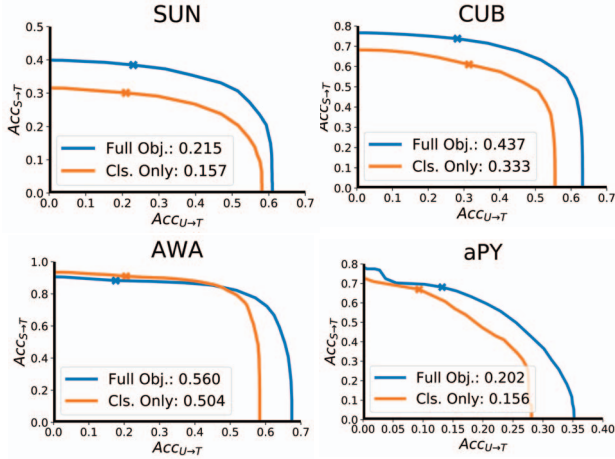


Figure 7: The Seen-Unseen accuracy Curve (SUC) [8] of SUN, CUB, AWA and aPY respectively. The blue line represents SP-AEN with whole loss objectives and the orange line represents SP-AEN with only classification loss objective. The cross denotes direct stacking when calibration factor $\gamma = 0$.

Table 2 reports ablative results of SP-AEN without the reconstruction G and discriminator D (Cls. Only). We can observe that by using adversarial training, we can significantly improve the H value by over 10% on all datasets. Figure 7 shows the AUSUC [8] of SP-AEN model with full objective (Full Obj.) and with only classification objective (Cls. Only)¹. We can observe that SP-AEN (Full Obj.) is

¹The performance results on Table 2 and Figure 7 are based on different hyper parameters setting. For results in Table 2, α and β are set to 10 and 50

consistently larger than model(Cls. Only) over all datasets. Both two types of metrics consistently demonstrates that SP-AEN enables effective semantic transfer. Figure 5 illustrates shows that by lowering the trade-off of G compared to D , the reconstruction quality will reduce.

6. Conclusions

We proposed a novel embedding based ZSL framework called Semantics-Preserving Adversarial Embedding Network (SP-AEN) to tackle the semantic loss problem in ZSL, which was rarely addressed by prior works. SP-AEN solves this problem by a novel visual reconstruction paradigm: 1) Introducing an independent visual-to-semantic mapping and then the reconstruction from the semantic space to the visual space would not affect the classification objective, whose contradiction to the reconstruction objective is extensively validated in this paper. 2) Semantic transfer can be achieved by adversarial learning between the two independent semantic embeddings. The first step preserves the semantics via reconstruction, while the second step enables semantics transfer across classes. We validated the effectiveness of SP-AEN through extensive comparative and ablative experiments on four ZSL benchmarks.

Our future works may focus on 1) incorporating generative models into SP-AEN so as to hallucinating photo-realistic images for unseen or even synthesized classes, and 2) developing new ZSL frameworks such as semi-supervised SP-AEN by imposing a prior semantic space.

Acknowledgement

This work was supported by National Key Research and Development Program of China (2017YFB0203001), National Natural Science Foundation of China (61572431), Zhejiang Natural Science Foundation (LZ17F020001), Key R&D Program of Zhejiang Province (2018C01006) and Joint Research Program of ZJU & Hikvision Research Institute.

5, and for results in Figure 7, α and β are set to 10 and 50.

References

- [1] Z. Akata, F. Perronnin, Z. Harchaoui, and C. Schmid. Label-embedding for attribute-based classification. In *CVPR*, 2013. 2
- [2] Z. Akata, F. Perronnin, Z. Harchaoui, and C. Schmid. Label-embedding for image classification. *TPAMI*, 2016. 3, 6
- [3] Z. Akata, S. Reed, D. Walter, H. Lee, and B. Schiele. Evaluation of output embeddings for fine-grained image classification. In *CVPR*, 2015. 3, 6
- [4] M. Arjovsky, S. Chintala, and L. Bottou. Wasserstein gan. In *arXiv*, 2017. 4
- [5] J. Bruna, P. Sprechmann, and Y. LeCun. Super-resolution with deep convolutional sufficient statistics. In *ICLR*, 2016. 3
- [6] M. Bucher, S. Herbin, and F. Jurie. Generating visual representations for zero-shot classification. In *ICCVW*, 2017. 3
- [7] S. Changpinyo, W.-L. Chao, B. Gong, and F. Sha. Synthesized classifiers for zero-shot learning. In *CVPR*, 2016. 6
- [8] W.-L. Chao, S. Changpinyo, B. Gong, and F. Sha. An empirical study and analysis of generalized zero-shot learning for object recognition in the wild. In *ECCV*, 2016. 3, 7, 8
- [9] L. Chen, H. Zhang, J. Xiao, L. Nie, J. Shao, W. Liu, and T.-S. Chua. Sca-cnn: Spatial and channel-wise attention in convolutional networks for image captioning. In *CVPR*, 2017. 3
- [10] B. Demirel, R. Gokberk Cinbis, and N. Ikizler-Cinbis. Attributes2classname: A discriminative model for attribute-based unsupervised zero-shot learning. In *ICCV*, 2017. 3
- [11] A. Dosovitskiy and T. Brox. Generating images with perceptual similarity metrics based on deep networks. In *NIPS*, 2016. 2, 3, 4, 5
- [12] M. Elhoseiny, B. Saleh, and A. Elgammal. Write a classifier: Zero-shot learning using purely textual descriptions. In *ICCV*, 2013. 3
- [13] A. Farhadi, I. Endres, D. Hoiem, and D. Forsyth. Describing objects by their attributes. In *CVPR*, 2009. 1, 2, 3, 5
- [14] A. Frome, G. S. Corrado, J. Shlens, S. Bengio, J. Dean, T. Mikolov, et al. Devise: A deep visual-semantic embedding model. In *NIPS*, 2013. 2, 3, 6
- [15] Y. Fu, T. M. Hospedales, T. Xiang, and S. Gong. Transductive multi-view zero-shot learning. *TPAMI*, 2015. 2, 3
- [16] L. Gatys, A. S. Ecker, and M. Bethge. Texture synthesis using convolutional neural networks. In *NIPS*, 2015. 3
- [17] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio. Generative adversarial nets. In *NIPS*, 2014. 2, 3, 4
- [18] B. Hariharan and R. Girshick. Low-shot visual recognition by shrinking and hallucinating features. In *ICCV*, 2017. 3
- [19] D. He, Y. Xia, T. Qin, L. Wang, N. Yu, T. Liu, and W.-Y. Ma. Dual learning for machine translation. In *NIPS*, 2016. 2, 3
- [20] K. He, X. Zhang, S. Ren, and J. Sun. Delving deep into rectifiers: Surpassing human-level performance on imagenet classification. In *ICCV*, 2015. 5
- [21] K. He, X. Zhang, S. Ren, and J. Sun. Deep residual learning for image recognition. In *CVPR*, 2016. 2, 4, 5
- [22] H. Jiang, R. Wang, S. Shan, Y. Yang, and X. Chen. Learning discriminative latent attributes for zero-shot classification. In *ICCV*, 2017. 3
- [23] J. Johnson, A. Alahi, and L. Fei-Fei. Perceptual losses for real-time style transfer and super-resolution. In *ECCV*, 2016. 3, 4
- [24] T. Kim, M. Cha, H. Kim, J. Lee, and J. Kim. Learning to discover cross-domain relations with generative adversarial networks. In *ICML*, 2017. 2, 3
- [25] E. Kodirov, T. Xiang, and S. Gong. Semantic autoencoder for zero-shot learning. In *CVPR*, 2017. 1, 2, 3, 4, 6, 7
- [26] A. Krizhevsky, I. Sutskever, and G. E. Hinton. Imagenet classification with deep convolutional neural networks. In *NIPS*, 2012. 4, 5
- [27] C. H. Lampert, H. Nickisch, and S. Harmeling. Learning to detect unseen object classes by between-class attribute transfer. In *CVPR*, 2009. 1, 2, 3, 5, 6
- [28] C. Ledig, L. Theis, F. Huszár, J. Caballero, A. Cunningham, A. Acosta, A. Aitken, A. Tejani, J. Totz, Z. Wang, et al. Photo-realistic single image super-resolution using a generative adversarial network. In *CVPR*, 2017. 4
- [29] J. Lei Ba, K. Swersky, S. Fidler, et al. Predicting deep zero-shot convolutional neural networks using textual descriptions. In *ICCV*, 2015. 3
- [30] L.-J. Li, H. Su, L. Fei-Fei, and E. P. Xing. Object bank: A high-level image representation for scene classification & semantic feature sparsification. In *NIPS*, 2010. 1
- [31] Y. Li, D. Wang, H. Hu, Y. Lin, and Y. Zhuang. Zero-shot recognition using dual visual-semantic mapping paths. In *CVPR*, 2017. 3
- [32] L. Ma, X. Jia, Q. Sun, B. Schiele, T. Tuytelaars, and L. Van Gool. Pose guided person image generation. In *NIPS*, 2017. 3
- [33] L. Ma, Q. Sun, S. Georgoulis, L. Van Gool, B. Schiele, and M. Fritz. Disentangled person image generation. In *CVPR*, 2018. 3
- [34] L. v. d. Maaten and G. Hinton. Visualizing data using t-sne. *JMLR*, 2008. 1
- [35] A. Makhzani, J. Shlens, N. Jaitly, and I. J. Goodfellow. Adversarial autoencoders. In *ICLRW*, 2016. 3, 5
- [36] A. G. MarcoBaroni. Hubness and pollution: Delving into cross-space mapping for zero-shot learning. In *ACL*, 2016. 2, 3
- [37] T. Mikolov, I. Sutskever, K. Chen, G. S. Corrado, and J. Dean. Distributed representations of words and phrases and their compositionality. In *NIPS*, 2013. 1, 3
- [38] A. Mishra, M. Reddy, A. Mittal, and H. A. Murthy. A generative model for zero shot learning using conditional variational autoencoders. In *arXiv*, 2017. 3
- [39] P. Morgado and N. Vasconcelos. Semantically consistent regularization for zero-shot recognition. In *CVPR*, 2017. 3
- [40] S. Motiian, M. Piccirilli, D. A. Adjeroh, and G. Doretto. Unified deep supervised domain adaptation and generalization. In *ICCV*, 2017. 3

- [41] A. Nguyen, J. Yosinski, Y. Bengio, A. Dosovitskiy, and J. Clune. Plug & play generative networks: Conditional iterative generation of images in latent space. In *CVPR*, 2017. 3
- [42] Y. Niu, Z. Lu, S. Huang, X. Gao, and J.-R. Wen. Feaboost: Joint feature and label refinement for semantic segmentation. In *AAAI*, 2017. 3
- [43] M. Norouzi, T. Mikolov, S. Bengio, Y. Singer, J. Shlens, A. Frome, G. S. Corrado, and J. Dean. Zero-shot learning by convex combination of semantic embeddings. In *ICLR*, 2014. 3, 6
- [44] A. Odena, C. Olah, and J. Shlens. Conditional image synthesis with auxiliary classifier gans. In *arXiv*, 2016. 3
- [45] P. Panareda Busto and J. Gall. Open set domain adaptation. In *ICCV*, 2017. 3
- [46] G. Patterson and J. Hays. Sun attribute database: Discovering, annotating, and recognizing scene attributes. In *CVPR*, 2012. 1, 2, 5
- [47] J. Pennington, R. Socher, and C. Manning. Glove: Global vectors for word representation. In *EMNLP*, 2014. 1, 3
- [48] S. Reed, Z. Akata, X. Yan, L. Logeswaran, B. Schiele, and H. Lee. Generative adversarial text to image synthesis. In *ICML*, 2016. 3
- [49] B. Romera-Paredes and P. Torr. An embarrassingly simple approach to zero-shot learning. In *ICML*, 2015. 3, 6
- [50] O. Russakovsky, J. Deng, H. Su, J. Krause, S. Satheesh, S. Ma, Z. Huang, A. Karpathy, A. Khosla, M. Bernstein, et al. Imagenet large scale visual recognition challenge. *I-JCV*, 2015. 5
- [51] K. Saenko, B. Kulis, M. Fritz, and T. Darrell. Adapting visual category models to new domains. In *ECCV*, 2010. 3
- [52] X. Shang, T. Ren, J. Guo, H. Zhang, and T.-S. Chua. Video visual relation detection. In *ACM MM*, 2017. 1
- [53] A. Shrivastava, T. Pfister, O. Tuzel, J. Susskind, W. Wang, and R. Webb. Learning from simulated and unsupervised images through adversarial training. In *arXiv*, 2016. 3
- [54] K. Simonyan and A. Zisserman. Very deep convolutional networks for large-scale image recognition. In *ICLR*, 2015. 4
- [55] R. Socher, M. Ganjoo, C. D. Manning, and A. Ng. Zero-shot learning through cross-modal transfer. In *NIPS*, 2013. 3, 6
- [56] Q. Sun, L. Ma, S. J. Oh, L. Van Gool, B. Schiele, and M. Fritz. Natural and effective obfuscation by head inpainting. In *CVPR*, 2018. 3
- [57] L. Torresani, M. Szummer, and A. Fitzgibbon. Efficient object category recognition using classemes. In *ECCV*, 2010. 1
- [58] E. Tzeng, J. Hoffman, K. Saenko, and T. Darrell. Adversarial discriminative domain adaptation. In *CVPR*, 2017. 3
- [59] P. Welinder, S. Branson, T. Mita, C. Wah, F. Schroff, S. Belongie, and P. Perona. Caltech-ucsd birds 200. 2010. 1, 2, 5
- [60] J. Weston, S. Bengio, and N. Usunier. Large scale image annotation: learning to rank with joint word-image embeddings. *Machine learning*, 2010. 2, 3
- [61] Y. Xian, Z. Akata, G. Sharma, Q. Nguyen, M. Hein, and B. Schiele. Latent embeddings for zero-shot classification. In *CVPR*, 2016. 3, 6
- [62] Y. Xian, B. Schiele, and Z. Akata. Zero-shot learning-the good, the bad and the ugly. In *CVPR*, 2017. 1, 2, 3, 5, 6
- [63] Z. Yi, H. Zhang, P. T. Gong, et al. Dualgan: Unsupervised dual learning for image-to-image translation. In *ICCV*, 2017. 2, 3
- [64] M. D. Zeiler and R. Fergus. Visualizing and understanding convolutional networks. In *ECCV*, 2014. 4
- [65] H. Zhang, Z. Kyaw, S.-F. Chang, and T.-S. Chua. Visual translation embedding network for visual relation detection. In *CVPR*, 2017. 1
- [66] H. Zhang, T. Xu, H. Li, S. Zhang, X. Huang, X. Wang, and D. Metaxas. Stackgan: Text to photo-realistic image synthesis with stacked generative adversarial networks. In *ICCV*, 2017. 3
- [67] H. Zhang, Z.-J. Zha, Y. Yang, S. Yan, Y. Gao, and T.-S. Chua. Attribute-augmented semantic hierarchy: towards bridging semantic gap and intention gap in image retrieval. In *ACM MM*, 2013. 1
- [68] Z. Zhang and V. Saligrama. Zero-shot learning via semantic similarity embedding. In *ICCV*, 2015. 6
- [69] J.-Y. Zhu, T. Park, P. Isola, and A. A. Efros. Unpaired image-to-image translation using cycle-consistent adversarial networks. In *ICCV*, 2017. 2, 3