

Recurrent Pixel Embedding for Instance Grouping

Shu Kong, Charless Fowlkes

Dept. of Computer Science, University of California, Irvine
{skong2, fowlkes}@ics.uci.edu

Abstract

We introduce a differentiable, end-to-end trainable framework for solving pixel-level grouping problems such as instance segmentation consisting of two novel components. First, we regress pixels into a hyper-spherical embedding space so that pixels from the same group have high cosine similarity while those from different groups have similarity below a specified margin. We analyze the choice of embedding dimension and margin, relating them to theoretical results on the problem of distributing points uniformly on the sphere. Second, to group instances, we utilize a variant of mean-shift clustering, implemented as a recurrent neural network parameterized by kernel bandwidth. This recurrent grouping module is differentiable, enjoys convergent dynamics and probabilistic interpretability. Backpropagating the group-weighted loss through this module allows learning to focus on correcting embedding errors that won't be resolved during subsequent clustering. Our framework, while conceptually simple and theoretically abundant, is also practically effective and computationally efficient. We demonstrate substantial improvements over state-of-the-art instance segmentation for object proposal generation, as well as demonstrating the benefits of grouping loss for classification tasks such as boundary detection and semantic segmentation.

1. Introduction

The successes of deep convolutional neural nets (CNNs) at image classification has spawned a flurry of work in computer vision on adapting these models to pixel-level image understanding tasks, such as boundary detection [1, 89, 63], semantic segmentation [59, 10, 45], optical flow [86, 20], and pose estimation [84, 7]. The key ideas that have enabled this adaption thus far are: (1) deconvolution schemes that allow for upsampling coarse pooled feature maps to make detailed predictions at the spatial resolution of individual pixels [89, 27], (2) skip connections and hyper-columns which concatenate representations across multi-resolution feature maps [31, 10], (3) atrous convolution which allows

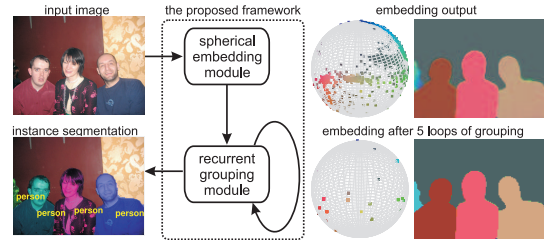


Figure 1: Our framework embeds pixels into a hyper-sphere where recurrent mean-shift dynamics groups pixels into a variable number of object instances. Here we visualize random projections of a 64-dim embeddings into 3-dimensions.

efficient computation with large receptive fields while maintaining spatial resolution [10, 45], and (4) fully convolutional operation which handles variable sized input images.

In contrast, there has been less innovation in the development of specialized loss functions for training. Pixel-level labeling tasks fall into the category of structured output prediction [4], where the model outputs a structured object (e.g., a whole image parse) rather than a scalar or categorical variable. However, most CNN pixel-labeling architectures are trained with loss functions that decompose into a simple (weighted) sum of classification or regression losses over individual pixel labels.

The need to address the output space structure is more apparent when considering problems where the set of output labels isn't fixed. Our motivating example is object instance segmentation, where the model generates a collection of segments corresponding to object instances. This problem can't be treated as k-way classification since the number of objects isn't known in advance. Further, the loss should be invariant to permutations of the instance labels within the same semantic category.

As a result, most recent successful methods for instance segmentation have adopted more heuristic approaches that first use an object detector to enumerate candidate instances and then perform pixel-level segmentation of each instance [56, 17, 54, 55, 2]. Alternately one can generate generic proposal segments and then label each one with a semantic detector [30, 12, 31, 16, 81, 33]. In either case the detection and segmentation steps can both be mapped to standard binary classification losses. While effective,

these approaches are somewhat unsatisfying since: (1) they rely on the object detector and non-maximum suppression heuristics to accurately “count” the number of instances, (2) they are difficult to train in an end-to-end manner since the interface between instance segmentation and detection is non-differentiable, and (3) they underperform in cluttered scenes as the assignment of pixels to detections is carried out independently for each detection¹.

Here we propose to directly tackle the instance grouping problem in a unified architecture by training a model that labels pixels with unit-length vectors that live in some fixed-dimension embedding space (Fig. 1). Unlike k -way classification where the target vectors for each pixel are specified in advance (i.e., one-hot vectors at the vertices of a $k-1$ dimensional simplex) we allow each instance to be labeled with an arbitrary embedding vector on the sphere. Our loss function simply enforces the constraint that the embedding vectors used to label different instances are far apart. Since neither the number of labels, nor the target label vectors are specified in advance, we can’t use standard soft-max thresholding to produce a discrete labeling. Instead, we utilize a variant of mean-shift clustering which can be viewed as a recurrent network whose fixed point identifies a small, discrete set of instance label vectors and concurrently labels each pixel with one of the vectors from this set.

This framework is largely agnostic to the underlying CNN architecture and can be applied to a range of low, mid and high-level visual tasks. Specifically, we carry out experiments showing how this method can be used for boundary detection, object proposal generation and semantic instance segmentation. Even when a task can be modeled by a binary pixel classification loss (e.g., boundary detection) we find that the grouping loss guides the model towards higher-quality feature representations that yield superior performance to classification loss alone. For the problem of instance segmentation, we demonstrate a notable boost in object proposal generation (improving the state-of-the-art average recall for 10 proposals per image from 0.56 to 0.77). To summarize our contributions: (1) we introduce a simple, easily interpreted end-to-end model for pixel-level instance labeling which is widely applicable and highly effective, (2) we provide theoretical analysis that offers guidelines on setting hyperparameters, and (3) benchmark results show substantial improvements over existing approaches.

2. Related Work

Common approaches to instance segmentation first generate region proposals or class-agnostic bounding boxes, segment the foreground objects within each proposal and classify the objects in the bounding box [90, 52, 30, 12,

¹This is less a problem for object proposals that are jointly estimated by bottom-up segmentation (e.g., MCG [70] and COB [63]). However, such generic proposal generation is not informed by the top-down semantics.

17, 55, 33]. [54] introduce a fully convolutional approach that includes bounding box proposal generation in end-to-end training. Recently, “box-free” methods [68, 69, 56, 36] avoid some limitations of box proposals (e.g. for wiry or articulated objects). They commonly use Faster RCNN [73] to produce “centeredness” score on each pixel and then predict binary instance masks and class labels. Other approaches have been explored for modeling segmentation and instance labeling jointly in a combinatorial framework (e.g., [40]) but typically don’t address end-to-end learning. Alternately, recurrent models that sequentially produce a list of instances [75, 72] offer another approach to address variable sized output structures in a unified manner.

The most closely related to ours is the associative embedding work of [66], which demonstrated strong results for grouping multi-person keypoints, and unpublished work from [22] on metric learning for instance segmentation. Our approach extends on these ideas substantially by integrating recurrent mean-shift to directly generate the final instances (rather than heuristic decoding or thresholding distance to seed proposals). There is also an important and interesting connection to work that segments instances using an embedding that is directly learned using a supervised regression loss rather than a pairwise associative loss. [79] train a regressor that predicts the distance to the contour centerline for boundary detection, while [3] predict the distance transform of the instance masks which is then post-processed with watershed transform to generate segments. [81] predict an embedding based on scene depth and direction towards the instance center (similar to Hough voting).

Finally, we note that these ideas are related to work on using embedding for solving pairwise clustering problems. For example, normalized cuts clusters embedding vectors given by the eigenvectors of the normalized graph Laplacian [77] and the spatial gradient of these embedding vectors was used in [1] as a feature for boundary detection. Rather than learning pairwise similarity from data and then embedding prior to clustering (e.g., [62]), we use a pairwise loss but learn the embedding directly. Our recurrent mean-shift grouping is reminiscent of other efforts that use unrolled implementations of iterative algorithms such as CRF inference [92] or bilateral filtering [39, 26]. Unlike general RNNs [6, 67] which are often difficult to train, our recurrent model has fixed parameters that assure interpretable convergent dynamics and meaningful gradients during learning.

3. Pairwise Loss for Pixel Embeddings

In this section we introduce and analyze the loss we use for learning pixel embeddings. This problem is broadly related to supervised distance metric learning [85, 47, 49] and clustering [48] but adapted to the specifics of instance labeling where the embedding vectors are treated as labels for a variable number of objects in each image.

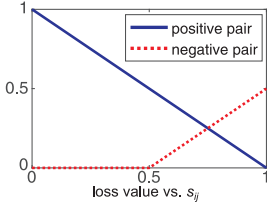


Figure 2: Pairwise loss $\ell(s_{ij})$ as a function of calibrated similarity score Eq. 1 with margin $\alpha = 0.5$. The gradient is constant, limiting the effect of noisy ground-truth labels (i.e., near an object boundary)

Our goal is to learn a mapping from an input image to a set of D -dimensional embedding vectors (one for each pixel). Let $\mathbf{x}_i, \mathbf{x}_j \in \mathbb{R}^D$ be the embeddings of pixels i and j respectively with corresponding labels y_i and y_j that denote ground-truth instance-level semantic labels (e.g., *car.1* and *car.2*). We measure the similarity of the embedding vectors using the cosine similarity, been scaled and offset to lie in the interval $[0, 1]$ for notational convenience:

$$s_{ij} = \frac{1}{2} \left(1 + \frac{\mathbf{x}_i^T \mathbf{x}_j}{\|\mathbf{x}_i\|_2 \|\mathbf{x}_j\|_2} \right) \quad (1)$$

In the discussion that follows we think of the similarity in terms of the inner product between the scaled embedding vectors (e.g., $\frac{\mathbf{x}_i}{\|\mathbf{x}_i\|}$) which live on the surface of a $(D - 1)$ dimensional sphere. Other common similarity metrics utilize Euclidean distance with a squared exponential kernel or sigmoid function [66, 22]. We prefer the cosine metric since it is invariant to the scale of the embedding vectors. This decouples the loss from other model design choices such as weight decay or regularization that implicitly limit the dynamic range of Euclidean distances.

Our goal is to learn an embedding so that pixels with the same label (positive pairs with $y_i = y_j$) have the same embedding (i.e., $s_{ij} = 1$). To avoid a trivial solution where all the embedding vectors are the same, we impose the additional constraint that pairs from different instances (negative pairs with $y_i \neq y_j$) are placed far apart. To provide additional flexibility, we include a weight w_i in the definition of the loss which specifies the importance of a given pixel. The total loss over all pairs and training images is:

$$\ell = \sum_{k=1}^M \sum_{i,j=1}^{N_k} \frac{w_i^k w_j^k}{N_k} \left(\mathbf{1}_{\{y_i=y_j\}} (1 - s_{ij}) + \mathbf{1}_{\{y_i \neq y_j\}} [s_{ij} - \alpha]_+ \right) \quad (2)$$

where N_k is the number of pixels in the k -th image (M images in total), and w_i^k is the pixel pair weight associated with pixel i in image k . The hyper-parameter α controls the maximum margin for negative pairs of pixels, incurring a penalty if the embeddings for pixels belonging to the same group have an angular separation of less than $\cos^{-1}(\alpha)$. Positive pairs pay a penalty if they have a similarity less than 1. Fig. 2 shows a graph of the loss function. [87] argue that the constant slope of the margin loss is more robust, e.g., than squared loss.

We carry out a simple theoretical analysis which provides a guide for setting the weights w_i and margin hyper-parameter α in the loss function. Proofs can be found in the supplementary material.

3.1. Instance-aware Pixel Weighting

We first examine the role of embedding dimension and instance size on the training loss.

Proposition 1 For n vectors $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$, the total intra-pixel similarity is bounded as $\sum_{i \neq j} \mathbf{x}_i^T \mathbf{x}_j \geq -\sum_{i=1}^n \|\mathbf{x}_i\|_2^2$. In particular, for n vectors on the hypersphere where $\|\mathbf{x}_i\|_2 = 1$, we have $\sum_{i \neq j} \mathbf{x}_i^T \mathbf{x}_j \geq -n$.

This proposition indicates that the total cosine similarity (and hence the loss) for a set of embedding vectors has a constant lower bound that does not depend on the dimension of the embedding space (a feature lacking in Euclidean embeddings). In particular, this type of analysis suggests a natural choice of pixel weighting w_i . Suppose a training example contains Q instances and $\mathcal{I}_q$ denotes the set of pixels belonging to a particular ground-truth instance q . We can write

$$\begin{aligned} \left\| \sum_{q=1}^Q \sum_{i \in \mathcal{I}_q} w_i \mathbf{x}_i \right\|^2 &= \sum_{q=1}^Q \left\| \sum_{i \in \mathcal{I}_q} w_i \mathbf{x}_i \right\|^2 + \\ &\quad \sum_{p \neq q} \left(\sum_{i \in \mathcal{I}_p} w_i \mathbf{x}_i \right)^T \left(\sum_{j \in \mathcal{I}_q} w_j \mathbf{x}_j \right) \end{aligned}$$

where the first term on the r.h.s. corresponds to contributions to the loss function for positive pairs while the second corresponds to contributions from negative pairs. Setting $w_i = \frac{1}{|\mathcal{I}_q|}$ for pixels i belonging to ground-truth instance q assures that each instance contributes equally to the loss independent of size. Furthermore, when the embedding dimension $D \geq Q$, we can simply embed the data so that the instance means $\mu_k = \frac{1}{|\mathcal{I}_q|} \sum_{i \in \mathcal{I}_q} \mathbf{x}_i$ are along orthogonal axes on the sphere. This zeros out the second term on the r.h.s., leaving only the first term which is bounded $0 \leq \sum_{q=1}^Q \left\| \frac{1}{|\mathcal{I}_q|} \sum_{i \in \mathcal{I}_q} \mathbf{x}_i \right\|^2 \leq Q$, and translates to corresponding upper and lower bounds on the loss that are independent of the number of pixels and embedding dimension (so long as $D \geq Q$).

Pairwise weighting schemes have been shown important empirically [22] and class imbalance can have a substantial effect on the performance of different architectures (see e.g., [57]). While other work has advocated online bootstrapping methods for hard-pixel mining or mini-batch selection [60, 46, 78, 88], our approach is much simpler. Guided by this result we simply use uniform random sampling of pixels during training, appropriately weighted by instance size in order to estimate the loss.

3.2. Margin Selection

To analyze the appropriate margin, let's first consider the problem of distributing labels for different instances as far apart as possible on a 3D sphere, sometimes referred to as Tammes's problem, or the hard-spheres problem [76]. This can be formalized as maximizing the smallest distance among n points on a sphere: $\max_{\mathbf{x}_i \in \mathbb{R}^3} \min_{i \neq j} \|\mathbf{x}_i - \mathbf{x}_j\|_2$. Asymptotic results in [28] provide the following proposition:

Proposition 2 *Given N vectors $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$ on a 2-sphere, i.e. $\mathbf{x}_i \in \mathbb{R}^3$, $\|\mathbf{x}_i\|_2 = 1, \forall i = 1 \dots n$, choosing $\alpha \leq 1 - \left(\frac{2\pi}{\sqrt{3}N}\right)$, guarantees that $[s_{ij} - \alpha]_+ \geq 0$ for some pair $i \neq j$. Choosing $\alpha > 1 - \frac{1}{4} \left(\left(\frac{8\pi}{\sqrt{3}N}\right)^{\frac{1}{2}} - CN^{-\frac{2}{3}} \right)^2$, guarantees the existence of an embedding with $[s_{ij} - \alpha]_+ = 0$ for all pairs $i \neq j$.*

Proposition 2 gives the maximum margin for a separation of n groups of pixels in a three dimensional embedding space (sphere). For example, if an image has at most $\{4, 5, 6, 7\}$ instances, α can be set as small as $\{0.093, 0.274, 0.395, 0.482\}$, respectively.

For points in a higher dimension embedding space, it is a non-trivial problem to establish a tight analytic bound for the margin α . Despite its simple description, distributing n points on a $(D - 1)$ -dimensional hypersphere is considered a serious mathematical challenge for which there is no general solutions [76, 61]. We adopt a safe (trivial) strategy. For n instances embedded in $n/2$ dimensions one can use value of $\alpha = 0.5$ which allows for zero loss by placing a pair of groups antipodally along each of the $n/2$ orthogonal axes. We adopt this setting for the majority of experiments in the paper where the embedding dimension is set to 64.

4. Recurrent Mean-Shift Grouping

While we can directly train a model to predict embeddings using the loss described in the previous section, it is not clear how to generate the final instance segmentation from the resulting (imperfect) embeddings. One can utilize heuristic post-processing [18] or utilize clustering algorithms that estimate the number of instances [56], but these are typically not differentiable and thus unsatisfying. Instead, we introduce a mean-shift grouping model (Fig. 3) which operates recurrently on the embedding space in order to condense the embedding vectors into a small number of instance labels.

Mean-shift and closely related algorithms [25, 13, 14, 15] use kernel density estimation to approximate the probability density from a set of samples and then perform clustering on the input data by assigning or moving each sample to the nearest mode (local maxima). From our perspective,

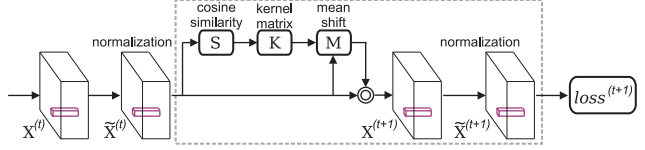


Figure 3: Recurrent mean shift grouping module is unrolled during training.

the advantages of this approach are (1) the final instance labels (modes) live in the same embedding space as the initial data, (2) the recurrent dynamics of the clustering process depend smoothly on the input, allowing for easy backpropagation, (3) the behavior depends on a single parameter, the kernel bandwidth, which is easily interpretable and can be related to the margin used for the embedding loss.

4.1. Mean Shift Clustering

A common choice for non-parametric density estimation is to use the isotropic multivariate normal kernel $K(\mathbf{x}, \mathbf{x}_i) = (2\pi)^{-D/2} \exp\left(-\frac{\delta^2}{2} \|\mathbf{x} - \mathbf{x}_i\|_2^2\right)$ and approximate the data density non-parametrically as $p(x) = \frac{1}{N} \sum K(x, x_i)$. Since our embedding vectors are unit norm, we instead use the von Mises-Fisher distribution, $K(\mathbf{x}, \mathbf{x}_i) \propto \exp(\delta \mathbf{x}^T \mathbf{x}_i)$, which is the natural extension of the multivariate normal to the hypersphere [24, 5, 64, 41]. The kernel bandwidth, δ determines the smoothness of the kernel density estimate. While it is straightforward to learn δ during training, we instead choose to tie it to the margin used in the embedding loss. Specifically we set $\frac{1}{\delta} = \frac{1-\alpha}{3}$ throughout our experiments so that the cluster separation (margin) in the learned embedding space is three standard deviations of the kernel.

We write the mean shift algorithm compactly in matrix form. Let $\mathbf{X} \in \mathbb{R}^{D \times N}$ denote the stacked embedding vectors of an N -pixel image. The kernel matrix is given by $\mathbf{K} = \exp(\delta \mathbf{X}^T \mathbf{X}) \in \mathbb{R}^{N \times N}$. Let $\mathbf{D} = \text{diag}(\mathbf{K}^T \mathbf{1})$ denote the diagonal matrix of total affinities, referred to as the degree when $\mathbf{K}$ is viewed as a weighted graph adjacency matrix. At each iteration, we compute the mean shift $\mathbf{M} = \mathbf{X} \mathbf{K} \mathbf{D}^{-1} - \mathbf{X}$, which is the difference vector between $\mathbf{X}$ and the kernel weighted average of $\mathbf{X}$. We then modify the embedding vectors by moving them in the mean shift direction with step size η :

$$\begin{aligned} \mathbf{X} &\leftarrow \mathbf{X} + \eta(\mathbf{X} \mathbf{K} \mathbf{D}^{-1} - \mathbf{X}) \\ &\leftarrow \mathbf{X}(\eta \mathbf{K} \mathbf{D}^{-1} + (1 - \eta)\mathbf{I}) \end{aligned} \quad (3)$$

Note that unlike standard mean-shift mode finding, we recompute $\mathbf{K}$ at each iteration. These update dynamics are termed the explicit- η method and were analyzed by [9]. When $\eta = 1$ and the kernel is Gaussian, this is also referred to as Gaussian Blurring Mean Shift (GBMS) and has been shown to have cubic convergence [9] under appropriate conditions. Unlike deep RNNs, the parameters of our

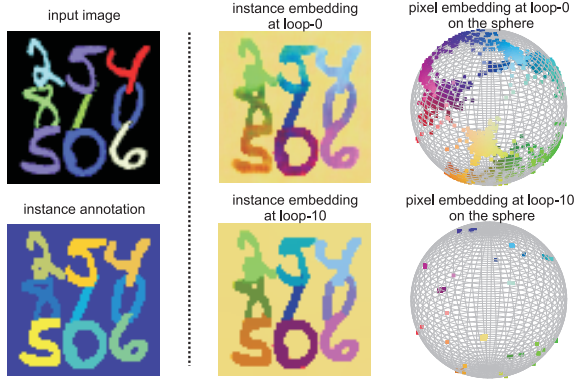


Figure 4: Demonstration of mean-shift grouping on a synthetic image and with ground-truth instance identities (left panel). Right panel: the pixel embedding visualization at 3-dimensional embedding sphere (upper row) and after 10 iterations of recurrent mean-shift grouping (bottom row).

recurrent module are not learned and the forward dynamics are convergent under general conditions. In practice, we do not observe issues with exploding or vanishing gradients during back-propagation through a (small) finite number of iterations².

Fig. 4 demonstrates a toy example of applying the method to perform digit instance segmentation on synthetic images from MNIST [53]. We learn 3-dimensional embedding in order to visualize the results before and after the mean shift grouping module. From the figure, we can see the mean shift grouping transforms the initial embedding vectors to yield a small set of instance labels which are distinct (for negative pairs) and compact (for positive pairs).

4.2. End-to-end training for Instance Grouping

It's straightforward to compute the gradients of the recurrent mean shift grouping module output w.r.t $\mathbf{X}$ so our whole system is end-to-end trainable using standard back-propagation. Details about the gradient computation can be found in the supplementary material. To understand the benefit of end-to-end training, we visualize the embedding gradient with and without the grouping module (Fig. 5). Interestingly, we observe that the gradient backpropagated through mean shift results in large magnitude updates to the embedding in uncertain regions (e.g., instance boundaries) and small magnitude updates for those embedding vectors that will be clustered correctly by the mean-shift iteration.

While we could simply apply the pairwise embedding loss to the final output of the mean-shift grouping, in practice we accumulate the loss over all iterations (including the initial embedding regression). We unroll the recurrent grouping module into T loops, and accumulate the same

²Some intuition about stability of forward dynamics can be gained by noting that the eigenvalues of $\mathbf{KD}^{-1}$ lie in the interval $[0, 1]$. We have not been able to prove useful corresponding bounds on the spectrum of the Jacobian.

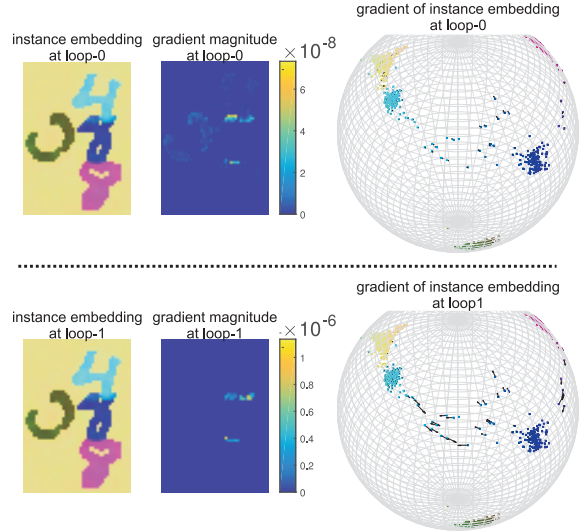


Figure 5: We compare the embedding vector gradients backpropagated through zero or one iteration of mean shift grouping. The length of arrows in the projection demonstrates the gradient magnitude (also depicted in maps as the second column). Backpropagating the loss through the grouping module serves to focus updates on embeddings of ambiguous pixels near boundaries while ignoring pixels with small errors that will be corrected by the subsequent grouping process.

loss function at the unrolled loop- t :

$$\ell^t = \sum_{k=1}^N \sum_{i,j \in S_k} \frac{w_i^k w_j^k}{|S_k|} \left(\mathbf{1}_{\{y_i=y_j\}} (1-s_{ij}^t) + \mathbf{1}_{\{y_i \neq y_j\}} [s_{ij}^t - \alpha]_+ \right)$$

We note that gradient magnitudes grow with the iteration depth T . Though this indicates potential issue of gradient explosion, we did not have such issues during training with with fixed unrolling depth.

5. Experiments

We now describe experiments in training our framework to deal a variety of pixel-labeling problems, including boundary detection, object proposal detection, semantic segmentation and instance-level semantic segmentation.

5.1. Tasks, Datasets and Implementation

We illustrate the advantages of the proposed modules on several large-scale datasets. First, to illustrate the ability of the instance-aware weighting and uniform sampling mechanism to handle imbalanced data and low embedding dimension, we use the BSDS500 [1] dataset to train a boundary detector for boundary detection ($> 90\%$ pixels are non-boundary pixels). We train with the standard split [1, 89], using 300 train-val images to train our model based on ResNet50 [34] and evaluate on the remaining 200 test images. Second, to explore instance segmentation and object proposal generation, we use PASCAL VOC 2012

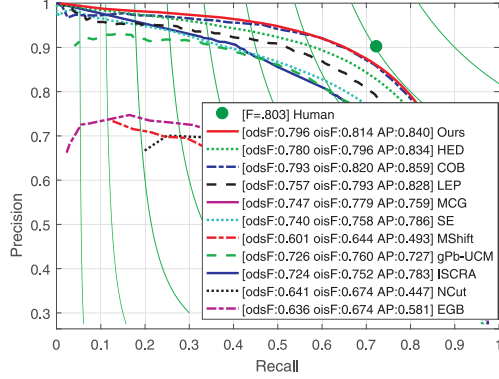


Figure 6: Boundary detection performance on BSDS500

dataset [21] with additional instance mask annotations provided by [29]. This provides 10,582 and 1,449 images for training and evaluation, respectively.

We implement our approach using the toolbox MatConvNet [83], and train using SGD on a single Titan X GPU³. To compute calibrated cosine similarity, we utilize an L2-normalization layer before matrix multiplication [44], which also contains random sampling with a hyper-parameter to control the ratio of pixels to be sampled for an image. In practice, we observe that performance does not depend strongly on this ratio and hence set it based on available (GPU) memory.

While our modules are architecture agnostic, we use the ResNet50 and ResNet101 models [34] pre-trained over ImageNet [19] as the backbone. Similar to [10], we increase the output resolution of ResNet by removing the top global 7×7 pooling layer and the last two 2×2 pooling layers, replacing them with atrous convolution with dilation rate 2 and 4, respectively to maintain a spatial sampling rate. Our model thus outputs predictions at $1/8$ the input resolution which are upsampled for benchmarking.

We augment the training set using random scaling by $s \in [0.5, 1.5]$, in-plane rotation by $[-10^\circ, 10^\circ]$ degrees, random left-right flips, random crops with 20-pixel margin and of size divisible by 8, and color jittering. When training the model, we fix the batch normalization in ResNet backbone, using the same constant global moments in both training and testing. Throughout training, we set batch size to one where the batch is a single input image. We use the “poly” learning rate policy [10] with a base learning rate of $2.5e-4$ scaled as a function of iteration by $(1 - \frac{iter}{maxiter})^{0.9}$.

5.2. Boundary Detection

For boundary detection, we first train a model to group the pixels into boundary or non-boundary groups. Similar to COB [63] and HED [89], we include multiple branches over ResBlock 2, 3, 4, 5 for training. Since the number of

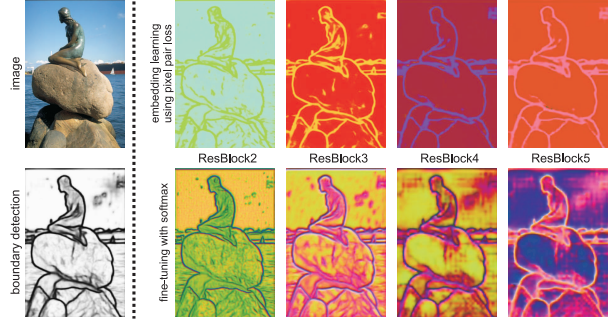


Figure 7: Visualization of boundary detection embeddings. We show the 3D embedding as RGB images (more examples in supplement). The upper and lower row in the right panel show embedding vectors at different layers from the model before and after fine-tuning using logistic loss. After fine-tuning, embeddings not only predict the boundary pixels, but also encode boundary orientation and signed distance to the boundary, similar to supervised embedding approaches [79, 81, 3]

instances labels is 2, we learn a simple 3-dimensional embedding space which has the advantage of easy visualization as an RGB image. Fig. 7 shows the resulting embeddings in the first row of each panel. Note that even though we didn’t utilize mean-shift grouping, the trained embedding already produces compact clusters. To compare quantitatively to the state-of-the-art, we learn a fusion layer that combines embeddings from multiple levels of the feature hierarchy fine-tuned with a logistic loss to make a binary prediction. Fig. 7 shows the results in the second row. Interestingly, we can see that the fine-tuned model embeddings encode not only boundary presence/absence but also the orientation and signed distance to nearby boundaries.

Quantitatively, we compare our model to COB [63], HED [89], LEP [65], UCM [1], ISCRA [74], NCuts [77], EGB [23], and the original mean shift (MShift) segmentation algorithm [15]. Fig. 6 shows standard benchmark precision-recall for all the methods, demonstrating our model achieves state-of-the-art performance. Our model has the same architecture of COB [63] except with a different loss and no explicit branches to compute boundary orientation. Note that it is possible to surpass human performance with several sophisticated techniques [43], we don’t pursue this as it is out the scope of this paper.

5.3. Object Proposal Detection

Object proposals are an integral part of current object detection and semantic segmentation pipelines [73, 33], as they provide a reduced search space of locations, scales and shapes for subsequent recognition. State-of-the-art methods usually involve training models that output large numbers of proposals, particularly those based on bounding boxes. Here we demonstrate that by training our framework with 64-dimensional embedding space on the object instance level annotations, we are able to produce very high qual-

³The code and trained models can be found at <https://github.com/aimerykong/Recurrent-Pixel-Embedding-for-Instance-Grouping>

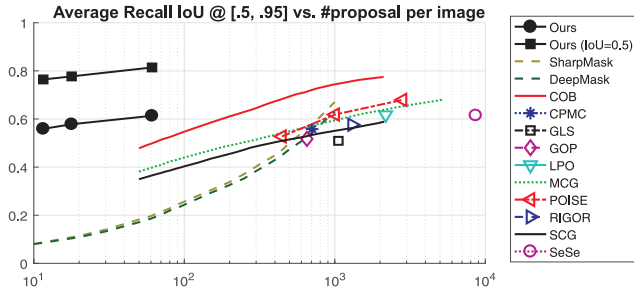


Figure 8: Segmented object proposal evaluation on PASCAL VOC 2012 validation set measured by Average Recall (AR) at IoU from 0.5 to 0.95 and step size as 0.5. We also plot the curve for our method at IoU=0.5.

#prop.	SCG [70]	MCG [70]	COB [63]	inst-DML [22]	Ours
10	-	-	-	0.558	0.769
60	0.624	0.652	0.738	0.667	0.814

Table 1: Object instance proposal evaluation on PASCAL VOC 2012 validation set measured by total Average Recall (AR) at IoU=0.50 and various number of proposals per image.

ity object proposals by grouping the pixels into instances. It is worth noting that due to the nature of our grouping module, far fewer number of proposals are produced with much higher quality. We compare against the most recent techniques including POISE [38], LPO [51], CPMC [8], GOP [50], SeSe [82], GLS [71], RIGOR [37].

Fig. 8 shows the Average Recall (AR) [35] with respect to the number of object proposals⁴. Our model performs remarkably well compared to other methods, achieving high average recall of ground-truth objects with two orders of magnitude fewer proposals. We also plot the curves for SharpMask [68] and DeepMask [69] using the proposals released by the authors. Despite only training on PASCAL, we outperform these models which were trained on the much larger COCO dataset [58]. In Table 1 we report the total average recall at IoU=0.5 for some recently proposed proposal detection methods, including unpublished work inst-DML [22] which is similar in spirit to our model but learns a Euclidean distance based metric to group pixels. We can clearly see that our method achieves significantly better results than existing methods.

5.4. Semantic Instance Detection

As a final test of our method, we also train it to produce semantic labels which are combined with our instance proposal method to recognize the detected proposals.

For semantic segmentation which is a k-way classification problem, we train a model using cross-entropy loss alongside our embedding loss. Similar to our proposal detection model, we use a 64-dimension embedding space in-

⁴Our basic model produces ~ 10 proposals per image. In order to plot a curve for our model for larger numbers of proposals, we run the mean shift grouping with multiple smaller bandwidth parameters, pool the results, and remove redundant proposals.

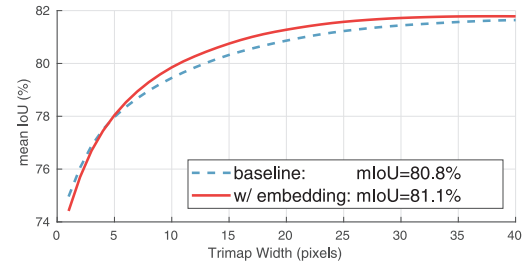


Figure 9: Semantic segmentation performance as a function of distance from ground-truth object boundaries comparing a baseline model trained with cross-entropy loss versus a model which also includes embedding loss.

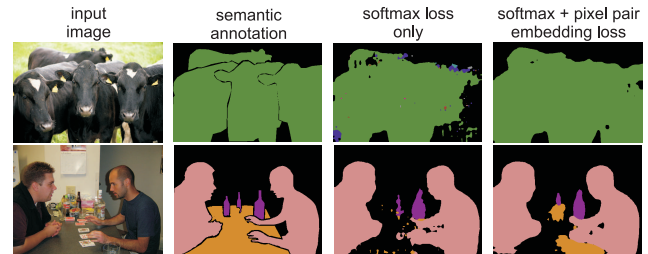


Figure 10: The proposed embedding loss improves semantic segmentation by forcing the pixel feature vectors to be similar within the segments. Randomly selected images from PASCAL VOC2012 val

side DeepLab model [11] as our model. While there are more complex methods in literature such as PSPNet [91] and which augment training with additional data (e.g., COCO [58] or JFT-300M dataset [80]) and utilize ensembles and post-processing, we focus on a simple experiment training the base model with/without the proposed pixel pair embedding loss to demonstrate the effectiveness.

In addition to reporting mean intersection over union (mIoU) over all classes, we also computed mIoU restricted to a narrow band of pixels around the ground-truth boundaries. This partition into figure/boundary/background is sometimes referred to as a tri-map in the matting literature and has been previously utilized in analyzing semantic segmentation performance [42, 10, 27]. Fig. 9 shows the mIoU as a function of the width of the tri-map boundary zone. This demonstrates that training with embedding loss yields performance gains over cross-entropy loss primarily far from ground-truth boundaries where it successfully fills in holes in the segments output (see also qualitative results in Fig. 10). This is similar in spirit to the model in [32], which considers local consistency to improve spatial precision. However, our uniform sampling allows for long-range interactions between pixels.

To label detected instances with semantic labels, we use the semantic segmentation model described above to generate labels and then use a simple voting strategy to transfer these predictions to the instance proposals. In order to produce a final confidence score associated with each proposed

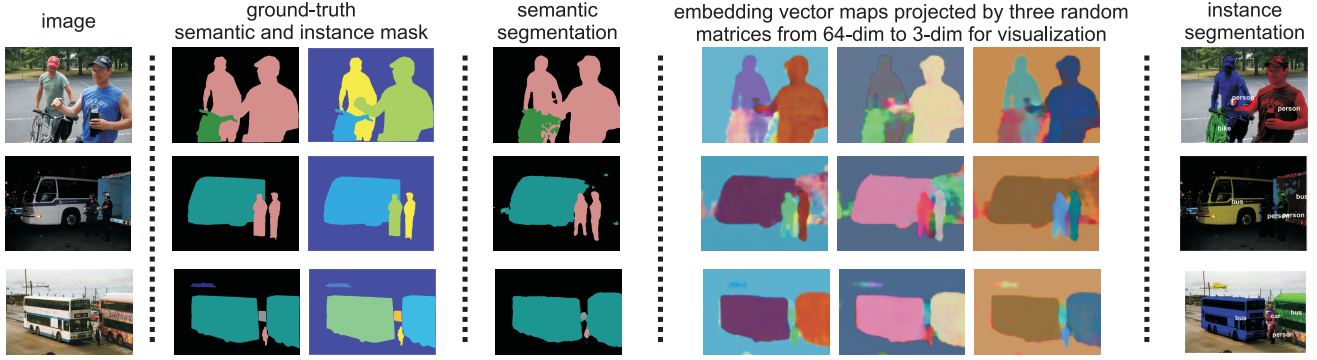


Figure 11: Visualization of generic/instance-level semantic segmentation on random PASCAL VOC 2012 validation images.

Method	plane	bike	bird	boat	bottle	bus	car	cat	chair	cow	table	dog	horse	motor	person	plant	sheep	sofa	train	tv	mean
SDS [30]	58.8	0.5	60.1	34.4	29.5	60.6	40.0	73.6	6.5	52.4	31.7	62.0	49.1	45.6	47.9	22.6	43.5	26.9	66.2	66.1	43.8
Chen et al. [12]	63.6	0.3	61.5	43.9	33.8	67.3	46.9	74.4	8.6	52.3	31.3	63.5	48.8	47.9	48.3	26.3	40.1	33.5	66.7	67.8	46.3
PFN [56]	76.4	15.6	74.2	54.1	26.3	73.8	31.4	92.1	17.4	73.7	48.1	82.2	81.7	72.0	48.4	23.7	57.7	64.4	88.9	72.3	58.7
MNC [17]	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	63.5
Li et al. [54]	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	65.7
R2-IOs [55]	87.0	6.1	90.3	67.9	48.4	86.2	68.3	90.3	24.5	84.2	29.6	91.0	71.2	79.9	60.4	42.4	67.4	61.7	94.3	82.1	66.7
Assoc. Embed. [66]	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	35.1
inst-DML [22]	69.7	1.2	78.2	53.8	42.2	80.1	57.4	88.8	16.0	73.2	57.9	88.4	78.9	80.0	68.0	28.0	61.5	61.3	87.5	70.4	62.1
Ours	85.9	10.0	74.3	54.6	43.7	81.3	64.1	86.1	17.5	77.5	57.0	89.2	77.8	83.7	67.9	31.2	62.5	63.3	88.6	74.2	64.5

Table 2: Instance-level segmentation comparison using APr metric at 0.5 IoU on the PASCAL VOC 2012 validation set.

object, we train a linear regressor to score each object instance based on its morphology (e.g., size, connectedness) and the consistency w.r.t. the semantic segmentation prediction. We note this is substantially simpler than approaches based, e.g. on Faster-RCNN [73], which use much richer convolutional features to classify segmented instances [33].

Comparison of instance detection performance are displayed in Table 2. We use a standard IoU threshold of 0.5 to identify true positives, unless a ground-truth instance has already been detected by a higher scoring proposal in which case it is a false positive. We report the average precision per-class mean over all classes (as in [29]). Our approach yields competitive performance on VOC validation despite our simple re-scoring. Among the competing methods, the one closest to our model is inst-DML [22], that learns Euclidean distance based metric with logistic loss. The inst-DML approach relies on generating pixel seeds to derive instance masks. The pixel seeds may fail to correctly detect thin structures which perhaps explains why this method performs $10\times$ worse than our method on the bike category. In contrast, our mean-shift grouping approach doesn’t make strong assumptions about the object shape or topology.

For visualization purposes, we generate three random projections of the 64-dimensional embedding and display them in the spatial domain as RGB images. Fig. 11 shows the embedding visualization, as well as predicted semantic segmentation and instance-level segmentation. From the visualization, we can see the instance-level semantic seg-

mentation outputs complete object instances even though semantic segmentation results are noisy, such as the bike in the first image in Fig. 11. The instance embedding provides important details that resolve both inter- and intra-class instance overlap which are not emphasized in the semantic segmentation loss.

6. Conclusion and Future Work

We have presented an end-to-end trainable framework for solving pixel-labeling vision problems based on two novel contributions: a pixel-pairwise loss based on spherical max-margin embedding and a variant of mean shift grouping embedded in a recurrent architecture. These two components mesh closely to provide a framework for robustly recognizing variable numbers of instances without requiring heuristic post-processing or hyperparameter tuning to account for widely varying instance size or class-imbalance. The approach is simple and amenable to theoretical analysis, and when coupled with standard architectures yields instance proposal generation which substantially outperforms state-of-the-art. Our experiments demonstrate the potential for instance embedding and open many opportunities for future work including learn-able variants of mean-shift grouping, extension to other pixel-level domains such as encoding surface shape, depth and figure-ground and multi-task embeddings.

Acknowledgments This project is supported by NSF grants IIS-1618806, IIS-1253538, DBI-1262547 and a hardware donation from NVIDIA.

References

- [1] P. Arbelaez, M. Maire, C. Fowlkes, and J. Malik. Contour detection and hierarchical image segmentation. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 33(5):898–916, 2011.
- [2] A. Arnab and P. H. Torr. Pixelwise instance segmentation with a dynamically instantiated network. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, 2017.
- [3] M. Bai and R. Urtasun. Deep watershed transform for instance segmentation. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 2858–2866. IEEE, 2017.
- [4] G. BakIr. *Predicting structured data*. MIT press, 2007.
- [5] A. Banerjee, I. S. Dhillon, J. Ghosh, and S. Sra. Clustering on the unit hypersphere using von mises-fisher distributions. *Journal of Machine Learning Research*, 6(Sep):1345–1382, 2005.
- [6] Y. Bengio, P. Simard, and P. Frasconi. Learning long-term dependencies with gradient descent is difficult. *IEEE transactions on neural networks*, 5(2):157–166, 1994.
- [7] Z. Cao, T. Simon, S.-E. Wei, and Y. Sheikh. Realtime multi-person 2d pose estimation using part affinity fields. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, 2017.
- [8] J. Carreira and C. Sminchisescu. Cpmc: Automatic object segmentation using constrained parametric min-cuts. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 34(7):1312–1328, 2012.
- [9] M. A. Carreira-Perpinán. Generalised blurring mean-shift algorithms for nonparametric clustering. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 1–8. IEEE, 2008.
- [10] L.-C. Chen, G. Papandreou, I. Kokkinos, K. Murphy, and A. L. Yuille. Deeplab: Semantic image segmentation with deep convolutional nets, atrous convolution, and fully connected crfs. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 40(4):834–848, 2018.
- [11] L.-C. Chen, G. Papandreou, F. Schroff, and H. Adam. Rethinking atrous convolution for semantic image segmentation. *arXiv preprint arXiv:1706.05587*, 2017.
- [12] Y.-T. Chen, X. Liu, and M.-H. Yang. Multi-instance object segmentation with occlusion handling. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 3470–3478, 2015.
- [13] Y. Cheng. Mean shift, mode seeking, and clustering. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 17(8):790–799, 1995.
- [14] D. Comaniciu and P. Meer. Mean shift analysis and applications. In *The Proceedings of the IEEE International Conference on Computer Vision (ICCV)*, volume 2, pages 1197–1203. IEEE, 1999.
- [15] D. Comaniciu and P. Meer. Mean shift: A robust approach toward feature space analysis. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 24(5):603–619, 2002.
- [16] J. Dai, K. He, and J. Sun. Convolutional feature masking for joint object and stuff segmentation. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 3992–4000, 2015.
- [17] J. Dai, K. He, and J. Sun. Instance-aware semantic segmentation via multi-task network cascades. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 3150–3158, 2016.
- [18] B. De Brabandere, D. Neven, and L. Van Gool. Semantic instance segmentation with a discriminative loss function. *arXiv preprint arXiv:1708.02551*, 2017.
- [19] J. Deng, W. Dong, R. Socher, L.-J. Li, K. Li, and L. Fei-Fei. Imagenet: A large-scale hierarchical image database. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 248–255. IEEE, 2009.
- [20] A. Dosovitskiy, P. Fischer, E. Ilg, P. Hausser, C. Hazirbas, V. Golkov, P. van der Smagt, D. Cremers, and T. Brox. FlowNet: Learning optical flow with convolutional networks. In *IEEE International Conference on Computer Vision (ICCV)*, pages 2758–2766, 2015.
- [21] M. Everingham, L. Van Gool, C. K. Williams, J. Winn, and A. Zisserman. The pascal visual object classes (voc) challenge. *International Journal of Computer Vision (IJCV)*, 88(2):303–338, 2010.
- [22] A. Fathi, Z. Wojna, V. Rathod, P. Wang, H. O. Song, S. Guadarrama, and K. P. Murphy. Semantic instance segmentation via deep metric learning. *arXiv preprint arXiv:1703.10277*, 2017.
- [23] P. F. Felzenszwalb and D. P. Huttenlocher. Efficient graph-based image segmentation. *International Journal of Computer Vision (IJCV)*, 59(2):167–181, 2004.
- [24] R. Fisher. Dispersion on a sphere. In *Proceedings of the Royal Society of London A: Mathematical, Physical and Engineering Sciences*, volume 217, pages 295–305. The Royal Society, 1953.
- [25] K. Fukunaga and L. Hostetler. The estimation of the gradient of a density function, with applications in pattern recognition. *IEEE Transactions on information theory*, 21(1):32–40, 1975.
- [26] R. Gadde, V. Jampani, M. Kiefel, D. Kappler, and P. V. Gehler. Superpixel convolutional networks using bilateral inceptions. In *European Conference on Computer Vision (ECCV)*, pages 597–613. Springer, 2016.
- [27] G. Ghiasi and C. C. Fowlkes. Laplacian pyramid reconstruction and refinement for semantic segmentation. In *European Conference on Computer Vision (ECCV)*, pages 519–534. Springer, 2016.
- [28] W. Habicht and B. Van der Waerden. Lagerung von punkten auf der kugel. *Mathematische Annalen*, 123(1):223–234, 1951.
- [29] B. Hariharan, P. Arbeláez, L. Bourdev, S. Maji, and J. Malik. Semantic contours from inverse detectors. In *IEEE International Conference on Computer Vision (ICCV)*, pages 991–998. IEEE, 2011.
- [30] B. Hariharan, P. Arbeláez, R. Girshick, and J. Malik. Simultaneous detection and segmentation. In *European Conference on Computer Vision (ECCV)*, pages 297–312. Springer, 2014.

- [31] B. Hariharan, P. Arbeláez, R. Girshick, and J. Malik. Hypercolumns for object segmentation and fine-grained localization. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 447–456, 2015.
- [32] A. W. Harley, K. G. Derpanis, and I. Kokkinos. Segmentation-aware convolutional networks using local attention masks. In *IEEE International Conference on Computer Vision (ICCV)*, volume 2, page 7, 2017.
- [33] K. He, G. Gkioxari, P. Dollár, and R. Girshick. Mask r-cnn. In *IEEE International Conference on Computer Vision (ICCV)*, 2017.
- [34] K. He, X. Zhang, S. Ren, and J. Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 770–778, 2016.
- [35] J. Hosang, R. Benenson, P. Dollár, and B. Schiele. What makes for effective detection proposals? *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 38(4):814–830, 2016.
- [36] H. Hu, S. Lan, Y. Jiang, Z. Cao, and F. Sha. Fastmask: Segment multi-scale object candidates in one shot. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016.
- [37] A. Humayun, F. Li, and J. M. Rehg. Rigor: Reusing inference in graph cuts for generating object regions. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 336–343, 2014.
- [38] A. Humayun, F. Li, and J. M. Rehg. The middle child problem: Revisiting parametric min-cut and seeds for object proposals. In *Proceedings of the IEEE International Conference on Computer Vision (ICCV)*, pages 1600–1608, 2015.
- [39] V. Jampani, M. Kiefel, and P. V. Gehler. Learning sparse high dimensional filters: Image filtering, dense crfs and bilateral neural networks. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 4452–4461, 2016.
- [40] A. Kirillov, E. Levinkov, B. Andres, B. Savchynskyy, and C. Rother. Instancecut: from edges to instances with multicut. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, 2017.
- [41] T. Kobayashi and N. Otsu. Von mises-fisher mean shift for clustering on a hypersphere. In *International Conference on Pattern Recognition (ICPR)*, pages 2130–2133. IEEE, 2010.
- [42] P. Kohli, P. H. Torr, et al. Robust higher order potentials for enforcing label consistency. *International Journal of Computer Vision (IJCV)*, 82(3):302–324, 2009.
- [43] I. Kokkinos. Pushing the boundaries of boundary detection using deep learning. *arXiv preprint arXiv:1511.07386*, 2015.
- [44] S. Kong and C. Fowlkes. Low-rank bilinear pooling for fine-grained classification. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, 2017.
- [45] S. Kong and C. Fowlkes. Recurrent scene parsing with perspective understanding in the loop. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, 2018.
- [46] S. Kong, X. Shen, Z. Lin, R. Mech, and C. Fowlkes. Photo aesthetics ranking network with attributes and content adaptation. In *European Conference on Computer Vision (ECCV)*, pages 662–679. Springer, 2016.
- [47] S. Kong and D. Wang. A dictionary learning approach for classification: Separating the particularity and the commonality. *European Conference on Computer Vision (ECCV)*, pages 186–199, 2012.
- [48] S. Kong and D. Wang. A multi-task learning strategy for unsupervised clustering via explicitly separating the commonality. In *International Conference on Pattern Recognition (ICPR)*, pages 771–774. IEEE, 2012.
- [49] S. Kong and D. Wang. Learning exemplar-represented manifolds in latent space for classification. In *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, pages 240–255. Springer, 2013.
- [50] P. Krähenbühl and V. Koltun. Geodesic object proposals. In *European Conference on Computer Vision (ECCV)*, pages 725–739. Springer, 2014.
- [51] P. Krahenbuhl and V. Koltun. Learning to propose objects. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 1574–1582, 2015.
- [52] L. Ladický, P. Sturges, K. Alahari, C. Russell, and P. H. Torr. What, where and how many? combining object detectors and crfs. In *European Conference on Computer Vision (ECCV)*, pages 424–437. Springer, 2010.
- [53] Y. LeCun, L. Bottou, Y. Bengio, and P. Haffner. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324, 1998.
- [54] Y. Li, H. Qi, J. Dai, X. Ji, and Y. Wei. Fully convolutional instance-aware semantic segmentation. In *IEEE Conf. on Computer Vision and Pattern Recognition (CVPR)*, 2016.
- [55] X. Liang, Y. Wei, X. Shen, Z. Jie, J. Feng, L. Lin, and S. Yan. Reversible recursive instance-level object segmentation. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 633–641, 2016.
- [56] X. Liang, Y. Wei, X. Shen, J. Yang, L. Lin, and S. Yan. Proposal-free network for instance-level object segmentation. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 2015.
- [57] T.-Y. Lin, P. Goyal, R. Girshick, K. He, and P. Dollár. Focal loss for dense object detection. In *IEEE International Conference on Computer Vision (ICCV)*, 2017.
- [58] T.-Y. Lin, M. Maire, S. Belongie, J. Hays, P. Perona, D. Ramanan, P. Dollár, and C. L. Zitnick. Microsoft coco: Common objects in context. In *European Conference on Computer Vision (ECCV)*, pages 740–755. Springer, 2014.
- [59] J. Long, E. Shelhamer, and T. Darrell. Fully convolutional networks for semantic segmentation. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 3431–3440, 2015.
- [60] I. Loshchilov and F. Hutter. Online batch selection for faster training of neural networks. In *International Conference on Learning Representations Workshop (ICLR Workshop)*, 2015.
- [61] L. Lovisolo and E. Da Silva. Uniform distribution of points on a hyper-sphere with applications to vector bit-plane en-

- coding. *IEE Proceedings-Vision, Image and Signal Processing*, 148(3):187–193, 2001.
- [62] M. Maire, T. Narihira, and S. X. Yu. Affinity cnn: Learning pixel-centric pairwise relations for figure/ground embedding. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 174–182, 2016.
- [63] K.-K. Maninis, J. Pont-Tuset, P. Arbelaez, and L. Van Gool. Convolutional oriented boundaries: From image segmentation to high-level tasks. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 2017.
- [64] K. V. Mardia and P. E. Jupp. *Directional statistics*, volume 494. John Wiley & Sons, 2009.
- [65] L. Najman and M. Schmitt. Geodesic saliency of watershed contours and hierarchical segmentation. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 18(12):1163–1173, 1996.
- [66] A. Newell and J. Deng. Associative embedding: End-to-end learning for joint detection and grouping. In *Advances in Neural Information Processing Systems (NIPS)*, 2017.
- [67] R. Pascanu, T. Mikolov, and Y. Bengio. On the difficulty of training recurrent neural networks. In *International Conference on Machine Learning (ICML)*, pages 1310–1318, 2013.
- [68] P. O. Pinheiro, R. Collobert, and P. Dollár. Learning to segment object candidates. In *Advances in Neural Information Processing Systems (NIPS)*, pages 1990–1998, 2015.
- [69] P. O. Pinheiro, T.-Y. Lin, R. Collobert, and P. Dollár. Learning to refine object segments. In *European Conference on Computer Vision (ECCV)*, pages 75–91. Springer, 2016.
- [70] J. Pont-Tuset, P. Arbelaez, J. T. Barron, F. Marques, and J. Malik. Multiscale combinatorial grouping for image segmentation and object proposal generation. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 39(1):128–140, 2017.
- [71] P. Rantalankila, J. Kannala, and E. Rahtu. Generating object segmentation proposals using global and local search. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 2417–2424, 2014.
- [72] M. Ren and R. S. Zemel. End-to-end instance segmentation with recurrent attention. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, 2017.
- [73] S. Ren, K. He, R. Girshick, and J. Sun. Faster r-cnn: Towards real-time object detection with region proposal networks. In *Advances in Neural Information Processing Systems (NIPS)*, pages 91–99, 2015.
- [74] Z. Ren and G. Shakhnarovich. Image segmentation by cascaded region agglomeration. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 2011–2018, 2013.
- [75] B. Romera-Paredes and P. H. S. Torr. Recurrent instance segmentation. In *European Conference on Computer Vision (ECCV)*, pages 312–329. Springer, 2016.
- [76] E. B. Saff and A. B. Kuijlaars. Distributing many points on a sphere. *The mathematical intelligencer*, 19(1):5–11, 1997.
- [77] J. Shi and J. Malik. Normalized cuts and image segmentation. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 22(8):888–905, 2000.
- [78] A. Shrivastava, A. Gupta, and R. Girshick. Training region-based object detectors with online hard example mining. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 761–769, 2016.
- [79] A. Sironi, V. Lepetit, and P. Fua. Multiscale centerline detection by learning a scale-space distance transform. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, pages 2697–2704, 2014.
- [80] C. Sun, A. Shrivastava, S. Singh, and A. Gupta. Revisiting unreasonable effectiveness of data in deep learning era. In *IEEE International Conference on Computer Vision (ICCV)*, 2017.
- [81] J. Uhrig, M. Cordts, U. Franke, and T. Brox. Pixel-level encoding and depth layering for instance-level semantic labeling. In *German Conference on Pattern Recognition*, pages 14–25. Springer, 2016.
- [82] J. R. Uijlings, K. E. Van De Sande, T. Gevers, and A. W. Smeulders. Selective search for object recognition. *International Journal of Computer Vision (IJCV)*, 104(2):154–171, 2013.
- [83] A. Vedaldi and K. Lenc. Matconvnet: Convolutional neural networks for matlab. In *Proceedings of the 23rd ACM international conference on Multimedia*, pages 689–692. ACM, 2015.
- [84] S.-E. Wei, V. Ramakrishna, T. Kanade, and Y. Sheikh. Convolutional pose machines. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, 2016.
- [85] K. Q. Weinberger and L. K. Saul. Distance metric learning for large margin nearest neighbor classification. *Journal of Machine Learning Research*, 10(Feb):207–244, 2009.
- [86] P. Weinzaepfel, J. Revaud, Z. Harchaoui, and C. Schmid. Deepflow: Large displacement optical flow with deep matching. In *IEEE International Conference on Computer Vision (ICCV)*, pages 1385–1392, 2013.
- [87] C.-Y. Wu, R. Manmatha, A. Smola, and P. K. Sampling matters in deep embedding learning. In *The Proceedings of the IEEE International Conference on Computer Vision (ICCV)*, 2017.
- [88] Z. Wu, C. Shen, and A. v. d. Hengel. Bridging category-level and instance-level semantic image segmentation. *arXiv preprint arXiv:1605.06885*, 2016.
- [89] S. Xie and Z. Tu. Holistically-nested edge detection. In *Proceedings of the IEEE International Conference on Computer Vision (ICCV)*, pages 1395–1403, 2015.
- [90] Y. Yang, S. Hallman, D. Ramanan, and C. C. Fowlkes. Layered object models for image segmentation. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 34(9):1731–1743, 2012.
- [91] H. Zhao, J. Shi, X. Qi, X. Wang, and J. Jia. Pyramid scene parsing network. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition (CVPR)*, 2016.
- [92] S. Zheng, S. Jayasumana, B. Romera-Paredes, V. Vineet, Z. Su, D. Du, C. Huang, and P. H. Torr. Conditional random fields as recurrent neural networks. In *Proceedings of the IEEE International Conference on Computer Vision (ICCV)*, pages 1529–1537, 2015.