

knn

November 27, 2024

```
[47]: # This mounts your Google Drive to the Colab VM.
from google.colab import drive
drive.mount('/content/drive')

# TODO: Enter the foldername in your Drive where you have saved the unzipped
# assignment folder, e.g. 'cs231n/assignments/assignment1/'
FOLDERNAME = None
assert FOLDERNAME is not None, "[!] Enter the foldername."

# Now that we've mounted your Drive, this ensures that
# the Python interpreter of the Colab VM can load
# python files from within it.
import sys
sys.path.append('/content/drive/My Drive/{}'.format(FOLDERNAME))

# This downloads the CIFAR-10 dataset to your Drive
# if it doesn't already exist.
%cd /content/drive/My\ Drive/$FOLDERNAME/cs231n/datasets/
!bash get_datasets.sh
%cd /content/drive/My\ Drive/$FOLDERNAME
```

```
-----
ModuleNotFoundError                                Traceback (most recent call last)
/Users/kimsh/Downloads/course/fall24/csc480/hw4/assignment1/knn.ipynb Cell 1
↳ line 2
      <a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall24/csc480.
↳ hw4/assignment1/knn.ipynb#W0sZmlsZQ%3D%3D?line=0'>1</a> # This mounts your
↳ Google Drive to the Colab VM.
----> <a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall24/csc480.
↳ hw4/assignment1/knn.ipynb#W0sZmlsZQ%3D%3D?line=1'>2</a> from google.colab
↳ import drive
      <a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall24/csc480.
↳ hw4/assignment1/knn.ipynb#W0sZmlsZQ%3D%3D?line=2'>3</a> drive.mount('/content
↳ drive')
      <a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall24/csc480.
↳ hw4/assignment1/knn.ipynb#W0sZmlsZQ%3D%3D?line=4'>5</a> # TODO: Enter the
↳ foldername in your Drive where you have saved the unzipped
```

```
<a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall124/csc480.
↪hw4/assignment1/knn.ipynb#W0sZmlsZQ%3D%3D?line=5'>6</a> # assignment folder, .
↪g. 'cs231n/assignments/assignment1/'
```

`ModuleNotFoundError: No module named 'google.colab'`

1 k-Nearest Neighbor (kNN) exercise

Complete and hand in this completed worksheet (including its outputs and any supporting code outside of the worksheet) with your assignment submission. For more details see the [assignments page](#) on the course website.

The kNN classifier consists of two stages:

- During training, the classifier takes the training data and simply remembers it
- During testing, kNN classifies every test image by comparing to all training images and transferring the labels of the k most similar training examples
- The value of k is cross-validated

In this exercise you will implement these steps and understand the basic Image Classification pipeline, cross-validation, and gain proficiency in writing efficient, vectorized code.

```
[48]: # Run some setup code for this notebook.

import random
import numpy as np
from cs231n.data_utils import load_CIFAR10
import matplotlib.pyplot as plt

# This is a bit of magic to make matplotlib figures appear inline in the
↪notebook
# rather than in a new window.
%matplotlib inline
plt.rcParams['figure.figsize'] = (10.0, 8.0) # set default size of plots
plt.rcParams['image.interpolation'] = 'nearest'
plt.rcParams['image.cmap'] = 'gray'

# Some more magic so that the notebook will reload external python modules;
# see http://stackoverflow.com/questions/1907993/
↪autoreload-of-modules-in-ipython
%load_ext autoreload
%autoreload 2
```

The autoreload extension is already loaded. To reload it, use:

```
%reload_ext autoreload
```

```
[49]: # Load the raw CIFAR-10 data.
cifar10_dir = 'cs231n/datasets/cifar-10-batches-py'
```

```

# Cleaning up variables to prevent loading data multiple times (which may cause
↳memory issue)
try:
    del X_train, y_train
    del X_test, y_test
    print('Clear previously loaded data.')
except:
    pass

X_train, y_train, X_test, y_test = load_CIFAR10(cifar10_dir)

# As a sanity check, we print out the size of the training and test data.
print('Training data shape: ', X_train.shape)
print('Training labels shape: ', y_train.shape)
print('Test data shape: ', X_test.shape)
print('Test labels shape: ', y_test.shape)

```

```

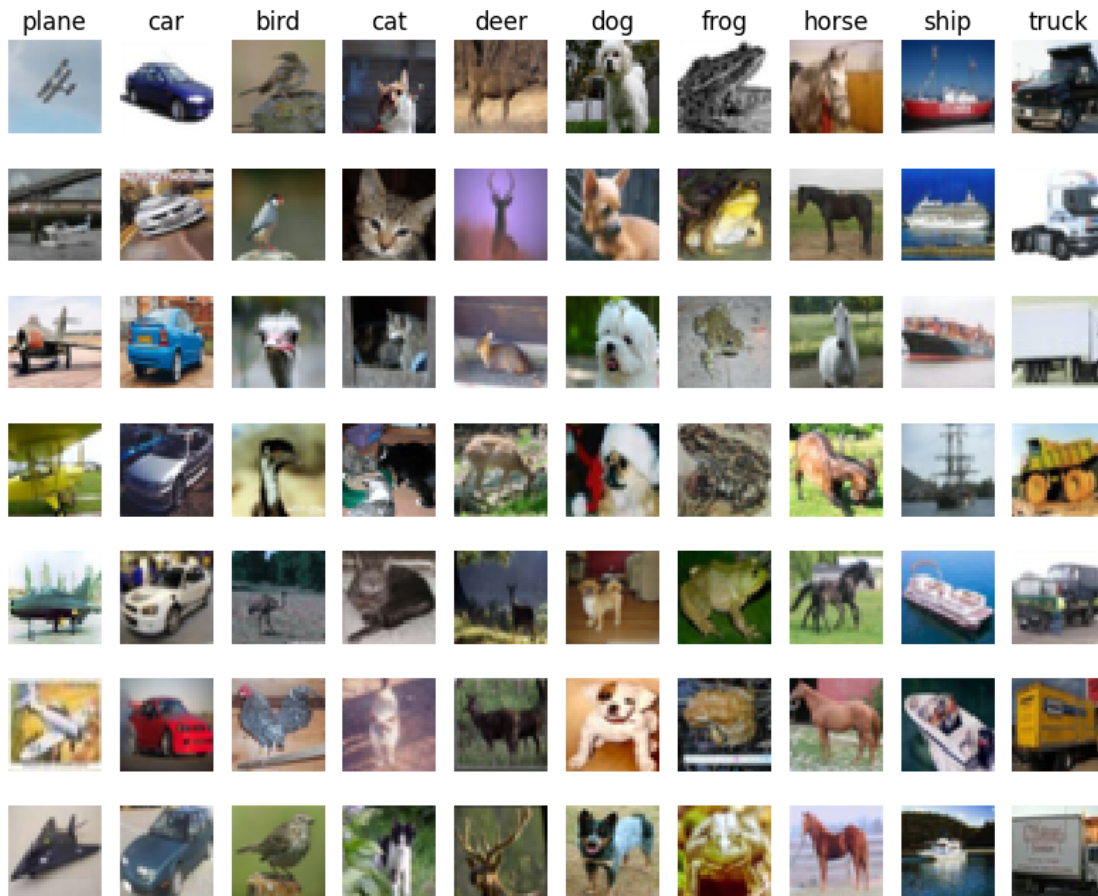
Clear previously loaded data.
Training data shape: (50000, 32, 32, 3)
Training labels shape: (50000,)
Test data shape: (10000, 32, 32, 3)
Test labels shape: (10000,)

```

```

[50]: # Visualize some examples from the dataset.
# We show a few examples of training images from each class.
classes = ['plane', 'car', 'bird', 'cat', 'deer', 'dog', 'frog', 'horse',
↳'ship', 'truck']
num_classes = len(classes)
samples_per_class = 7
for y, cls in enumerate(classes):
    idxs = np.flatnonzero(y_train == y)
    idxs = np.random.choice(idxs, samples_per_class, replace=False)
    for i, idx in enumerate(idxs):
        plt_idx = i * num_classes + y + 1
        plt.subplot(samples_per_class, num_classes, plt_idx)
        plt.imshow(X_train[idx].astype('uint8'))
        plt.axis('off')
        if i == 0:
            plt.title(cls)
plt.show()

```



```
[51]: # Subsample the data for more efficient code execution in this exercise
num_training = 5000
mask = list(range(num_training))
X_train = X_train[mask]
y_train = y_train[mask]

num_test = 500
mask = list(range(num_test))
X_test = X_test[mask]
y_test = y_test[mask]

# Reshape the image data into rows
X_train = np.reshape(X_train, (X_train.shape[0], -1))
X_test = np.reshape(X_test, (X_test.shape[0], -1))
print(X_train.shape, X_test.shape)
```

(5000, 3072) (500, 3072)

```
[52]: from cs231n.classifiers import KNearestNeighbor

# Create a kNN classifier instance.
# Remember that training a kNN classifier is a noop:
# the Classifier simply remembers the data and does no further processing
classifier = KNearestNeighbor()
classifier.train(X_train, y_train)
```

We would now like to classify the test data with the kNN classifier. Recall that we can break down this process into two steps:

1. First we must compute the distances between all test examples and all train examples.
2. Given these distances, for each test example we find the k nearest examples and have them vote for the label

Lets begin with computing the distance matrix between all training and test examples. For example, if there are N_{tr} training examples and N_{te} test examples, this stage should result in a $N_{te} \times N_{tr}$ matrix where each element (i,j) is the distance between the i-th test and j-th train example.

Note: For the three distance computations that we require you to implement in this notebook, you may not use the `np.linalg.norm()` function that numpy provides.

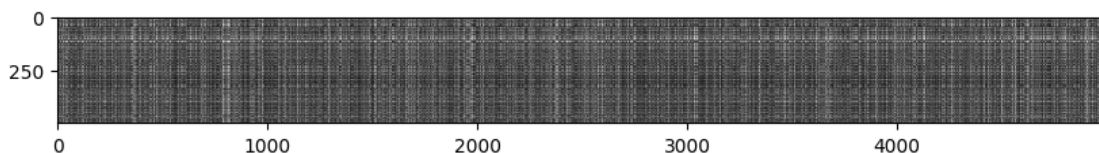
First, open `cs231n/classifiers/k_nearest_neighbor.py` and implement the function `compute_distances_two_loops` that uses a (very inefficient) double loop over all pairs of (test, train) examples and computes the distance matrix one element at a time.

```
[61]: # Open cs231n/classifiers/k_nearest_neighbor.py and implement
# compute_distances_two_loops.

# Test your implementation:
dists = classifier.compute_distances_two_loops(X_test)
print(dists.shape)
```

(500, 5000)

```
[62]: # We can visualize the distance matrix: each row is a single test example and
# its distances to training examples
plt.imshow(dists, interpolation='none')
plt.show()
```



Inline Question 1

Notice the structured patterns in the distance matrix, where some rows or columns are visibly

brighter. (Note that with the default color scheme black indicates low distances while white indicates high distances.)

- What in the data is the cause behind the distinctly bright rows?
- What causes the columns?

Your Answer : The bright lines indicate the dissimilarity between the test and training samples. The brighter the line, the more different test and training samples are from each other. A bright row means test image is very different from all the training images, and a bright column means a training image is very different from all the test images.

```
[68]: # Now implement the function predict_labels and run the code below:
# We use k = 1 (which is Nearest Neighbor).
y_test_pred = classifier.predict_labels(dists, k=1)

# Compute and print the fraction of correctly predicted examples
num_correct = np.sum(y_test_pred == y_test)
accuracy = float(num_correct) / num_test
print('Got %d / %d correct => accuracy: %f' % (num_correct, num_test, accuracy))
```

Got 137 / 500 correct => accuracy: 0.274000

You should expect to see approximately 27% accuracy. Now let's try out a larger k, say k = 5:

```
[69]: y_test_pred = classifier.predict_labels(dists, k=5)
num_correct = np.sum(y_test_pred == y_test)
accuracy = float(num_correct) / num_test
print('Got %d / %d correct => accuracy: %f' % (num_correct, num_test, accuracy))
```

Got 139 / 500 correct => accuracy: 0.278000

You should expect to see a slightly better performance than with k = 1.

Inline Question 2

We can also use other distance metrics such as L1 distance. For pixel values $p_{ij}^{(k)}$ at location (i, j) of some image I_k ,

the mean μ across all pixels over all images is

$$\mu = \frac{1}{nhw} \sum_{k=1}^n \sum_{i=1}^h \sum_{j=1}^w p_{ij}^{(k)}$$

And the pixel-wise mean μ_{ij} across all images is

$$\mu_{ij} = \frac{1}{n} \sum_{k=1}^n p_{ij}^{(k)}.$$

The general standard deviation σ and pixel-wise standard deviation σ_{ij} is defined similarly.

Which of the following preprocessing steps will not change the performance of a Nearest Neighbor classifier that uses L1 distance? Select all that apply. To clarify, both training and test examples are preprocessed in the same way.

1. Subtracting the mean μ ($\tilde{p}_{ij}^{(k)} = p_{ij}^{(k)} - \mu$.)
2. Subtracting the per pixel mean μ_{ij} ($\tilde{p}_{ij}^{(k)} = p_{ij}^{(k)} - \mu_{ij}$.)
3. Subtracting the mean μ and dividing by the standard deviation σ .
4. Subtracting the pixel-wise mean μ_{ij} and dividing by the pixel-wise standard deviation σ_{ij} .
5. Rotating the coordinate axes of the data, which means rotating all the images by the same angle. Empty regions in the image caused by rotation are padded with a same pixel value and no interpolation is performed.

Your Answer :

1, 2, 3

Your Explanation :

For 1, the 2 μ 's will cancel out as seen in the equation.

For 2, it is the same case as 1.

For 3, the distance comparisons will remain the same because each distance is scaled by the same amount.

```
[70]: # Now lets speed up distance matrix computation by using partial vectorization
# with one loop. Implement the function compute_distances_one_loop and run the
# code below:
dists_one = classifier.compute_distances_one_loop(X_test)

# To ensure that our vectorized implementation is correct, we make sure that it
# agrees with the naive implementation. There are many ways to decide whether
# two matrices are similar; one of the simplest is the Frobenius norm. In case
# you haven't seen it before, the Frobenius norm of two matrices is the square
# root of the squared sum of differences of all elements; in other words,
↳ reshape
# the matrices into vectors and compute the Euclidean distance between them.
difference = np.linalg.norm(dists - dists_one, ord='fro')
print('One loop difference was: %f' % (difference, ))
if difference < 0.001:
    print('Good! The distance matrices are the same')
else:
    print('Uh-oh! The distance matrices are different')
```

One loop difference was: 0.000000

Good! The distance matrices are the same

```
[71]: # Now implement the fully vectorized version inside compute_distances_no_loops
# and run the code
dists_two = classifier.compute_distances_no_loops(X_test)

# check that the distance matrix agrees with the one we computed before:
difference = np.linalg.norm(dists - dists_two, ord='fro')
print('No loop difference was: %f' % (difference, ))
if difference < 0.001:
```

```

    print('Good! The distance matrices are the same')
else:
    print('Uh-oh! The distance matrices are different')

```

No loop difference was: 0.000000

Good! The distance matrices are the same

```

[72]: # Let's compare how fast the implementations are
def time_function(f, *args):
    """
    Call a function f with args and return the time (in seconds) that it took
    to execute.
    """
    import time
    tic = time.time()
    f(*args)
    toc = time.time()
    return toc - tic

two_loop_time = time_function(classifier.compute_distances_two_loops, X_test)
print('Two loop version took %f seconds' % two_loop_time)

one_loop_time = time_function(classifier.compute_distances_one_loop, X_test)
print('One loop version took %f seconds' % one_loop_time)

no_loop_time = time_function(classifier.compute_distances_no_loops, X_test)
print('No loop version took %f seconds' % no_loop_time)

# You should see significantly faster performance with the fully vectorized
# implementation!

# NOTE: depending on what machine you're using,
# you might not see a speedup when you go from two loops to one loop,
# and might even see a slow-down.

```

Two loop version took 96.891194 seconds

One loop version took 96.526387 seconds

No loop version took 0.222191 seconds

1.0.1 Cross-validation

We have implemented the k-Nearest Neighbor classifier but we set the value $k = 5$ arbitrarily. We will now determine the best value of this hyperparameter with cross-validation.

```

[76]: num_folds = 5
      k_choices = [1, 3, 5, 8, 10, 12, 15, 20, 50, 100]

      X_train_folds = []

```

```

y_train_folds = []
#####
# TODO:
# Split up the training data into folds. After splitting, X_train_folds and
# y_train_folds should each be lists of length num_folds, where
# y_train_folds[i] is the label vector for the points in X_train_folds[i].
# Hint: Look up the numpy array_split function.
#####
# *****START OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****

X_train_folds = np.array_split(X_train, num_folds)
y_train_folds = np.array_split(y_train, num_folds)

# *****END OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****

# A dictionary holding the accuracies for different values of k that we find
# when running cross-validation. After running cross-validation,
# k_to_accuracies[k] should be a list of length num_folds giving the different
# accuracy values that we found when using that value of k.
k_to_accuracies = {}

#####
# TODO:
# Perform k-fold cross validation to find the best value of k. For each
# possible value of k, run the k-nearest-neighbor algorithm num_folds times,
# where in each case you use all but one of the folds as training data and the
# last fold as a validation set. Store the accuracies for all fold and all
# values of k in the k_to_accuracies dictionary.
#####
# *****START OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****

for k in k_choices:
    k_to_accuracies[k] = []
    for i in range(num_folds):
        X_train_temp = np.concatenate(np.compress(np.array(range(num_folds)) != i, X_train_folds, axis=0))
        y_train_temp = np.concatenate(np.compress(np.array(range(num_folds)) != i, y_train_folds, axis=0))
        classifier.train(X_train_temp, y_train_temp)
        y_pred_temp = classifier.predict(X_train_folds[i], k=k, num_loops=0)
        k_to_accuracies[k].append(np.sum(y_pred_temp == y_train_folds[i]) / len(y_pred_temp))

# *****END OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****

# Print out the computed accuracies

```

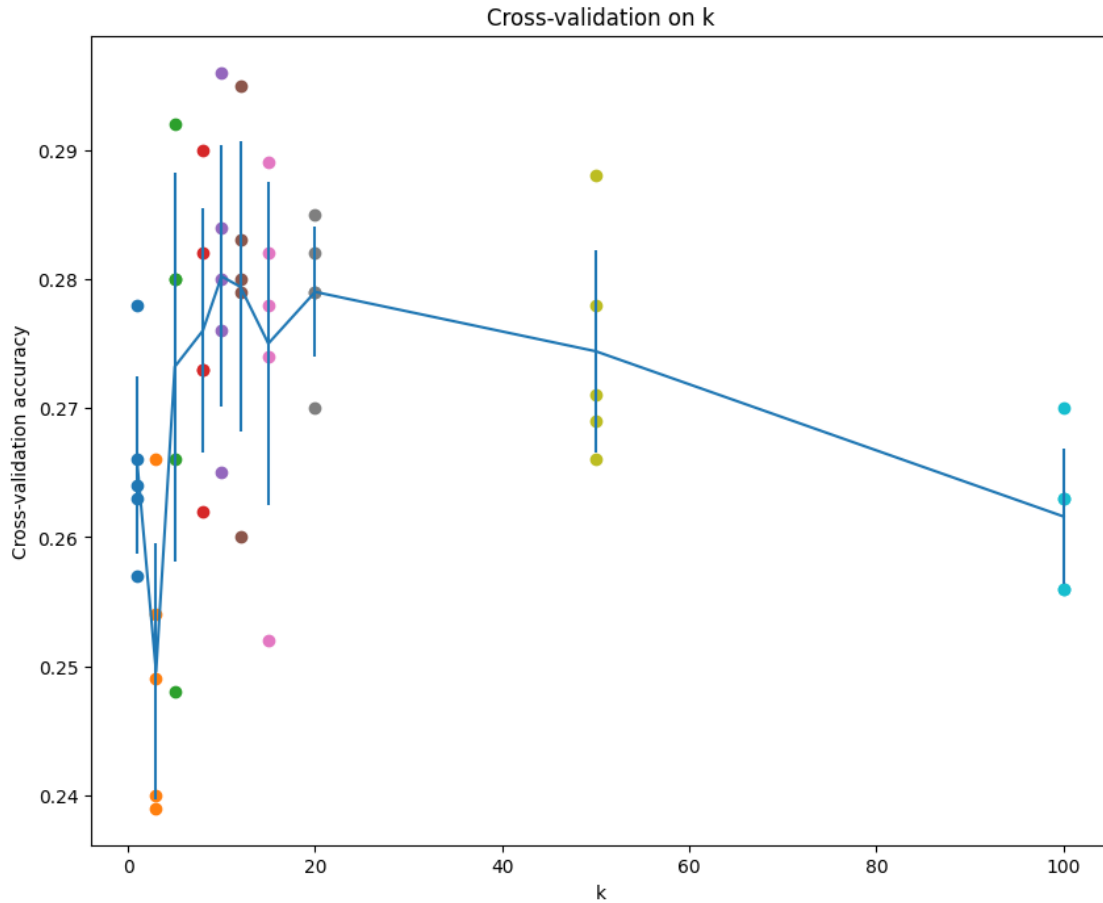
```
for k in sorted(k_to_accuracies):  
    for accuracy in k_to_accuracies[k]:  
        print('k = %d, accuracy = %f' % (k, accuracy))
```

```
k = 1, accuracy = 0.263000  
k = 1, accuracy = 0.257000  
k = 1, accuracy = 0.264000  
k = 1, accuracy = 0.278000  
k = 1, accuracy = 0.266000  
k = 3, accuracy = 0.239000  
k = 3, accuracy = 0.249000  
k = 3, accuracy = 0.240000  
k = 3, accuracy = 0.266000  
k = 3, accuracy = 0.254000  
k = 5, accuracy = 0.248000  
k = 5, accuracy = 0.266000  
k = 5, accuracy = 0.280000  
k = 5, accuracy = 0.292000  
k = 5, accuracy = 0.280000  
k = 8, accuracy = 0.262000  
k = 8, accuracy = 0.282000  
k = 8, accuracy = 0.273000  
k = 8, accuracy = 0.290000  
k = 8, accuracy = 0.273000  
k = 10, accuracy = 0.265000  
k = 10, accuracy = 0.296000  
k = 10, accuracy = 0.276000  
k = 10, accuracy = 0.284000  
k = 10, accuracy = 0.280000  
k = 12, accuracy = 0.260000  
k = 12, accuracy = 0.295000  
k = 12, accuracy = 0.279000  
k = 12, accuracy = 0.283000  
k = 12, accuracy = 0.280000  
k = 15, accuracy = 0.252000  
k = 15, accuracy = 0.289000  
k = 15, accuracy = 0.278000  
k = 15, accuracy = 0.282000  
k = 15, accuracy = 0.274000  
k = 20, accuracy = 0.270000  
k = 20, accuracy = 0.279000  
k = 20, accuracy = 0.279000  
k = 20, accuracy = 0.282000  
k = 20, accuracy = 0.285000  
k = 50, accuracy = 0.271000  
k = 50, accuracy = 0.288000  
k = 50, accuracy = 0.278000  
k = 50, accuracy = 0.269000
```

```
k = 50, accuracy = 0.266000
k = 100, accuracy = 0.256000
k = 100, accuracy = 0.270000
k = 100, accuracy = 0.263000
k = 100, accuracy = 0.256000
k = 100, accuracy = 0.263000
```

```
[77]: # plot the raw observations
      for k in k_choices:
          accuracies = k_to_accuracies[k]
          plt.scatter([k] * len(accuracies), accuracies)

      # plot the trend line with error bars that correspond to standard deviation
      accuracies_mean = np.array([np.mean(v) for k,v in sorted(k_to_accuracies.
          ↪items())])
      accuracies_std = np.array([np.std(v) for k,v in sorted(k_to_accuracies.
          ↪items())])
      plt.errorbar(k_choices, accuracies_mean, yerr=accuracies_std)
      plt.title('Cross-validation on k')
      plt.xlabel('k')
      plt.ylabel('Cross-validation accuracy')
      plt.show()
```



```
[88]: # Based on the cross-validation results above, choose the best value for k,
# retrain the classifier using all the training data, and test it on the test
# data. You should be able to get above 28% accuracy on the test data.
best_k = k_choices[np.argmax(accuracies_mean)]
print(accuracies_mean)
print(best_k)

classifier = KNearestNeighbor()
classifier.train(X_train, y_train)
y_test_pred = classifier.predict(X_test, k=best_k)

# Compute and display the accuracy
num_correct = np.sum(y_test_pred == y_test)
accuracy = float(num_correct) / num_test
print('Got %d / %d correct => accuracy: %f' % (num_correct, num_test, accuracy))

[0.2656 0.2496 0.2732 0.276 0.2802 0.2794 0.275 0.279 0.2744 0.2616]
10
Got 141 / 500 correct => accuracy: 0.282000
```

Inline Question 3

Which of the following statements about k -Nearest Neighbor (k -NN) are true in a classification setting, and for all k ? Select all that apply. 1. The decision boundary of the k -NN classifier is linear. 2. The training error of a 1-NN will always be lower than or equal to that of 5-NN. 3. The test error of a 1-NN will always be lower than that of a 5-NN. 4. The time needed to classify a test example with the k -NN classifier grows with the size of the training set. 5. None of the above.

Your Answer :

2 and 4

Your Explanation :

For 2, 1-NN will overfit KNN to the training set which guarantees more accuracy than 5-NN on the training set.

For 4, increasing the training set will increase the time KNN needs to classify a test example because it has to compare it to more points in the training set.

SVM

November 27, 2024

```
[1]: # This mounts your Google Drive to the Colab VM.
from google.colab import drive
drive.mount('/content/drive')

# TODO: Enter the foldername in your Drive where you have saved the unzipped
# assignment folder, e.g. 'cs231n/assignments/assignment1/'
FOLDERNAME = None
assert FOLDERNAME is not None, "[!] Enter the foldername."

# Now that we've mounted your Drive, this ensures that
# the Python interpreter of the Colab VM can load
# python files from within it.
import sys
sys.path.append('/content/drive/My Drive/{}'.format(FOLDERNAME))

# This downloads the CIFAR-10 dataset to your Drive
# if it doesn't already exist.
%cd /content/drive/My\ Drive/$FOLDERNAME/cs231n/datasets/
!bash get_datasets.sh
%cd /content/drive/My\ Drive/$FOLDERNAME
```

```
-----
ModuleNotFoundError                                Traceback (most recent call last)
/Users/kimsh/Downloads/course/fall24/csc480/hw4/assignment1/svm.ipynb Cell 1
↳ line 2
      <a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall24/csc480.
↳ hw4/assignment1/svm.ipynb#W0sZmlsZQ%3D%3D?line=0'>1</a> # This mounts your
↳ Google Drive to the Colab VM.
----> <a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall24/csc480.
↳ hw4/assignment1/svm.ipynb#W0sZmlsZQ%3D%3D?line=1'>2</a> from google.colab
↳ import drive
      <a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall24/csc480.
↳ hw4/assignment1/svm.ipynb#W0sZmlsZQ%3D%3D?line=2'>3</a> drive.mount('/content
↳ drive')
      <a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall24/csc480.
↳ hw4/assignment1/svm.ipynb#W0sZmlsZQ%3D%3D?line=4'>5</a> # TODO: Enter the
↳ foldername in your Drive where you have saved the unzipped
```

```
<a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall124/csc480.
↪hw4/assignment1/svm.ipynb#W0sZmlsZQ%3D%3D?line=5'>6</a> # assignment folder, .
↪g. 'cs231n/assignments/assignment1/'
```

`ModuleNotFoundError: No module named 'google.colab'`

1 Multiclass Support Vector Machine exercise

Complete and hand in this completed worksheet (including its outputs and any supporting code outside of the worksheet) with your assignment submission. For more details see the [assignments page](#) on the course website.

In this exercise you will:

- implement a fully-vectorized **loss function** for the SVM
- implement the fully-vectorized expression for its **analytic gradient**
- **check your implementation** using numerical gradient
- use a validation set to **tune the learning rate and regularization** strength
- **optimize** the loss function with **SGD**
- **visualize** the final learned weights

```
[2]: # Run some setup code for this notebook.
import random
import numpy as np
from cs231n.data_utils import load_CIFAR10
import matplotlib.pyplot as plt

# This is a bit of magic to make matplotlib figures appear inline in the
# notebook rather than in a new window.
%matplotlib inline
plt.rcParams['figure.figsize'] = (10.0, 8.0) # set default size of plots
plt.rcParams['image.interpolation'] = 'nearest'
plt.rcParams['image.cmap'] = 'gray'

# Some more magic so that the notebook will reload external python modules;
# see http://stackoverflow.com/questions/1907993/
↪autoreload-of-modules-in-ipython
%load_ext autoreload
%autoreload 2
```

1.1 CIFAR-10 Data Loading and Preprocessing

```
[3]: # Load the raw CIFAR-10 data.
cifar10_dir = 'cs231n/datasets/cifar-10-batches-py'

# Cleaning up variables to prevent loading data multiple times (which may cause ↵
↪memory issue)
```

```

try:
    del X_train, y_train
    del X_test, y_test
    print('Clear previously loaded data.')
except:
    pass

X_train, y_train, X_test, y_test = load_CIFAR10(cifar10_dir)

# As a sanity check, we print out the size of the training and test data.
print('Training data shape: ', X_train.shape)
print('Training labels shape: ', y_train.shape)
print('Test data shape: ', X_test.shape)
print('Test labels shape: ', y_test.shape)

```

```

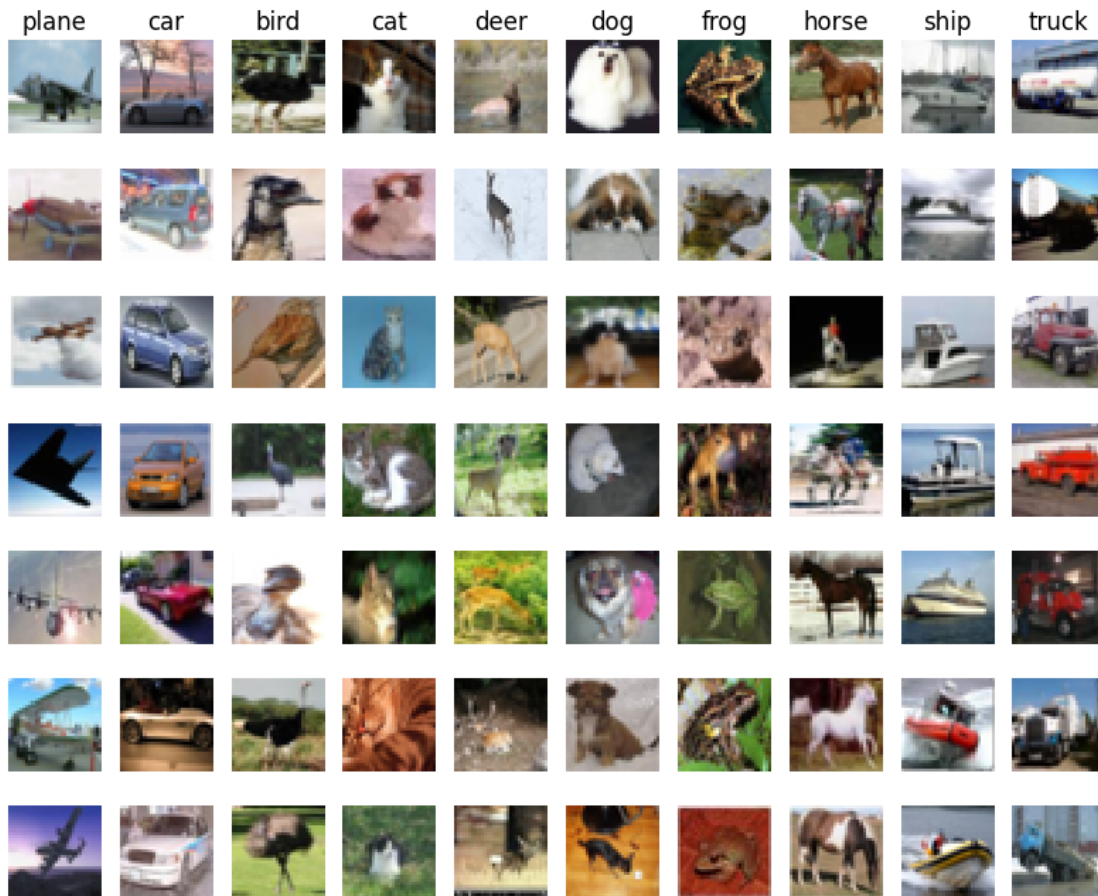
Training data shape: (50000, 32, 32, 3)
Training labels shape: (50000,)
Test data shape: (10000, 32, 32, 3)
Test labels shape: (10000,)

```

```

[4]: # Visualize some examples from the dataset.
# We show a few examples of training images from each class.
classes = ['plane', 'car', 'bird', 'cat', 'deer', 'dog', 'frog', 'horse', '
↳ 'ship', 'truck']
num_classes = len(classes)
samples_per_class = 7
for y, cls in enumerate(classes):
    idxs = np.flatnonzero(y_train == y)
    idxs = np.random.choice(idxs, samples_per_class, replace=False)
    for i, idx in enumerate(idxs):
        plt_idx = i * num_classes + y + 1
        plt.subplot(samples_per_class, num_classes, plt_idx)
        plt.imshow(X_train[idx].astype('uint8'))
        plt.axis('off')
        if i == 0:
            plt.title(cls)
plt.show()

```



```
[5]: # Split the data into train, val, and test sets. In addition we will
# create a small development set as a subset of the training data;
# we can use this for development so our code runs faster.
num_training = 49000
num_validation = 1000
num_test = 1000
num_dev = 500

# Our validation set will be num_validation points from the original
# training set.
mask = range(num_training, num_training + num_validation)
X_val = X_train[mask]
y_val = y_train[mask]

# Our training set will be the first num_train points from the original
# training set.
mask = range(num_training)
X_train = X_train[mask]
```

```

y_train = y_train[mask]

# We will also make a development set, which is a small subset of
# the training set.
mask = np.random.choice(num_training, num_dev, replace=False)
X_dev = X_train[mask]
y_dev = y_train[mask]

# We use the first num_test points of the original test set as our
# test set.
mask = range(num_test)
X_test = X_test[mask]
y_test = y_test[mask]

print('Train data shape: ', X_train.shape)
print('Train labels shape: ', y_train.shape)
print('Validation data shape: ', X_val.shape)
print('Validation labels shape: ', y_val.shape)
print('Test data shape: ', X_test.shape)
print('Test labels shape: ', y_test.shape)

```

```

Train data shape: (49000, 32, 32, 3)
Train labels shape: (49000,)
Validation data shape: (1000, 32, 32, 3)
Validation labels shape: (1000,)
Test data shape: (1000, 32, 32, 3)
Test labels shape: (1000,)

```

```

[6]: # Preprocessing: reshape the image data into rows
X_train = np.reshape(X_train, (X_train.shape[0], -1))
X_val = np.reshape(X_val, (X_val.shape[0], -1))
X_test = np.reshape(X_test, (X_test.shape[0], -1))
X_dev = np.reshape(X_dev, (X_dev.shape[0], -1))

# As a sanity check, print out the shapes of the data
print('Training data shape: ', X_train.shape)
print('Validation data shape: ', X_val.shape)
print('Test data shape: ', X_test.shape)
print('dev data shape: ', X_dev.shape)

```

```

Training data shape: (49000, 3072)
Validation data shape: (1000, 3072)
Test data shape: (1000, 3072)
dev data shape: (500, 3072)

```

```

[7]: # Preprocessing: subtract the mean image
# first: compute the image mean based on the training data
mean_image = np.mean(X_train, axis=0)

```

```

print(mean_image[:10]) # print a few of the elements
plt.figure(figsize=(4,4))
plt.imshow(mean_image.reshape((32,32,3)).astype('uint8')) # visualize the mean_
    ↳image
plt.show()

# second: subtract the mean image from train and test data
X_train -= mean_image
X_val -= mean_image
X_test -= mean_image
X_dev -= mean_image

# third: append the bias dimension of ones (i.e. bias trick) so that our SVM
# only has to worry about optimizing a single weight matrix W.
X_train = np.hstack([X_train, np.ones((X_train.shape[0], 1))])
X_val = np.hstack([X_val, np.ones((X_val.shape[0], 1))])
X_test = np.hstack([X_test, np.ones((X_test.shape[0], 1))])
X_dev = np.hstack([X_dev, np.ones((X_dev.shape[0], 1))])

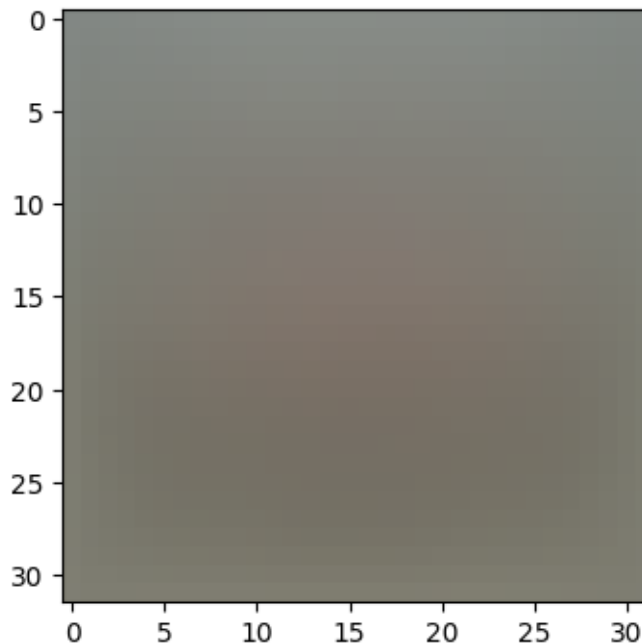
print(X_train.shape, X_val.shape, X_test.shape, X_dev.shape)

```

```

[130.64189796 135.98173469 132.47391837 130.05569388 135.34804082
 131.75402041 130.96055102 136.14328571 132.47636735 131.48467347]

```



```

(49000, 3073) (1000, 3073) (1000, 3073) (500, 3073)

```

1.2 SVM Classifier

Your code for this section will all be written inside `cs231n/classifiers/linear_svm.py`.

As you can see, we have prefilled the function `svm_loss_naive` which uses for loops to evaluate the multiclass SVM loss function.

```
[26]: # Evaluate the naive implementation of the loss we provided for you:
from cs231n.classifiers.linear_svm import svm_loss_naive
import time

# generate a random SVM weight matrix of small numbers
W = np.random.randn(3073, 10) * 0.0001

loss, grad = svm_loss_naive(W, X_dev, y_dev, 0.000005)
print('loss: %f' % (loss, ))
```

loss: 8.447939

The `grad` returned from the function above is right now all zero. Derive and implement the gradient for the SVM cost function and implement it inline inside the function `svm_loss_naive`. You will find it helpful to interleave your new code inside the existing function.

To check that you have correctly implemented the gradient, you can numerically estimate the gradient of the loss function and compare the numeric estimate to the gradient that you computed. We have provided code that does this for you:

```
[35]: # Once you've implemented the gradient, recompute it with the code below
# and gradient check it with the function we provided for you

# Compute the loss and its gradient at W.
loss, grad = svm_loss_naive(W, X_dev, y_dev, 0.0)

# Numerically compute the gradient along several randomly chosen dimensions, and
# compare them with your analytically computed gradient. The numbers should
# match
# almost exactly along all dimensions.
from cs231n.gradient_check import grad_check_sparse
f = lambda w: svm_loss_naive(w, X_dev, y_dev, 0.0)[0]
grad_numerical = grad_check_sparse(f, W, grad)

# do the gradient check once again with regularization turned on
# you didn't forget the regularization gradient did you?
loss, grad = svm_loss_naive(W, X_dev, y_dev, 5e1)
f = lambda w: svm_loss_naive(w, X_dev, y_dev, 5e1)[0]
grad_numerical = grad_check_sparse(f, W, grad)
```

```
numerical: -0.105736 analytic: -0.105736, relative error: 6.510338e-09
numerical: -26.893066 analytic: -26.893066, relative error: 9.869183e-12
numerical: -18.346430 analytic: -18.419940, relative error: 1.999364e-03
numerical: 13.584597 analytic: 13.671742, relative error: 3.197254e-03
```

```

numerical: 5.378144 analytic: 5.332022, relative error: 4.306328e-03
numerical: 2.661139 analytic: 2.661139, relative error: 5.403595e-11
numerical: -9.221664 analytic: -9.221664, relative error: 7.374423e-11
numerical: 0.733091 analytic: 0.686527, relative error: 3.280027e-02
numerical: -6.772661 analytic: -6.772661, relative error: 4.340004e-12
numerical: -20.722216 analytic: -20.722216, relative error: 2.118431e-11
numerical: -2.239230 analytic: -2.357256, relative error: 2.567743e-02
numerical: 14.238183 analytic: 14.238183, relative error: 2.684619e-11
numerical: -3.124048 analytic: -3.124048, relative error: 5.859279e-11
numerical: 17.358027 analytic: 17.270592, relative error: 2.524922e-03
numerical: 2.007831 analytic: 1.977626, relative error: 7.578925e-03
numerical: -10.184294 analytic: -10.184294, relative error: 2.068810e-11
numerical: -3.526441 analytic: -3.541763, relative error: 2.167720e-03
numerical: 9.459157 analytic: 9.459157, relative error: 2.864724e-11
numerical: -6.940359 analytic: -7.025943, relative error: 6.127856e-03
numerical: -0.788321 analytic: -0.788321, relative error: 4.729160e-10

```

Inline Question 1

It is possible that once in a while a dimension in the gradcheck will not match exactly. What could such a discrepancy be caused by? Is it a reason for concern? What is a simple example in one dimension where a gradient check could fail? How would change the margin affect of the frequency of this happening? *Hint: the SVM loss function is not strictly speaking differentiable*

Your Answer : It could be a fundamental difference in the numeric versus analytic computation for loss and gradient that cause the discrepancy.

```

[44]: # Next implement the function svm_loss_vectorized; for now only compute the
      ↪ loss;
      # we will implement the gradient in a moment.
      tic = time.time()
      loss_naive, grad_naive = svm_loss_naive(W, X_dev, y_dev, 0.000005)
      toc = time.time()
      print('Naive loss: %e computed in %fs' % (loss_naive, toc - tic))

      from cs231n.classifiers.linear_svm import svm_loss_vectorized
      tic = time.time()
      loss_vectorized, _ = svm_loss_vectorized(W, X_dev, y_dev, 0.000005)
      toc = time.time()
      print('Vectorized loss: %e computed in %fs' % (loss_vectorized, toc - tic))

      # The losses should match but your vectorized implementation should be much
      ↪ faster.
      print('difference: %f' % (loss_naive - loss_vectorized))

```

```

Naive loss: 8.447939e+00 computed in 0.051316s
Vectorized loss: 8.447939e+00 computed in 0.002823s
difference: 0.000000

```

```
[45]: # Complete the implementation of svm_loss_vectorized, and compute the gradient
# of the loss function in a vectorized way.

# The naive implementation and the vectorized implementation should match, but
# the vectorized version should still be much faster.
tic = time.time()
_, grad_naive = svm_loss_naive(W, X_dev, y_dev, 0.000005)
toc = time.time()
print('Naive loss and gradient: computed in %fs' % (toc - tic))

tic = time.time()
_, grad_vectorized = svm_loss_vectorized(W, X_dev, y_dev, 0.000005)
toc = time.time()
print('Vectorized loss and gradient: computed in %fs' % (toc - tic))

# The loss is a single number, so it is easy to compare the values computed
# by the two implementations. The gradient on the other hand is a matrix, so
# we use the Frobenius norm to compare them.
difference = np.linalg.norm(grad_naive - grad_vectorized, ord='fro')
print('difference: %f' % difference)
```

```
Naive loss and gradient: computed in 0.060948s
Vectorized loss and gradient: computed in 0.003200s
difference: 0.000000
```

1.2.1 Stochastic Gradient Descent

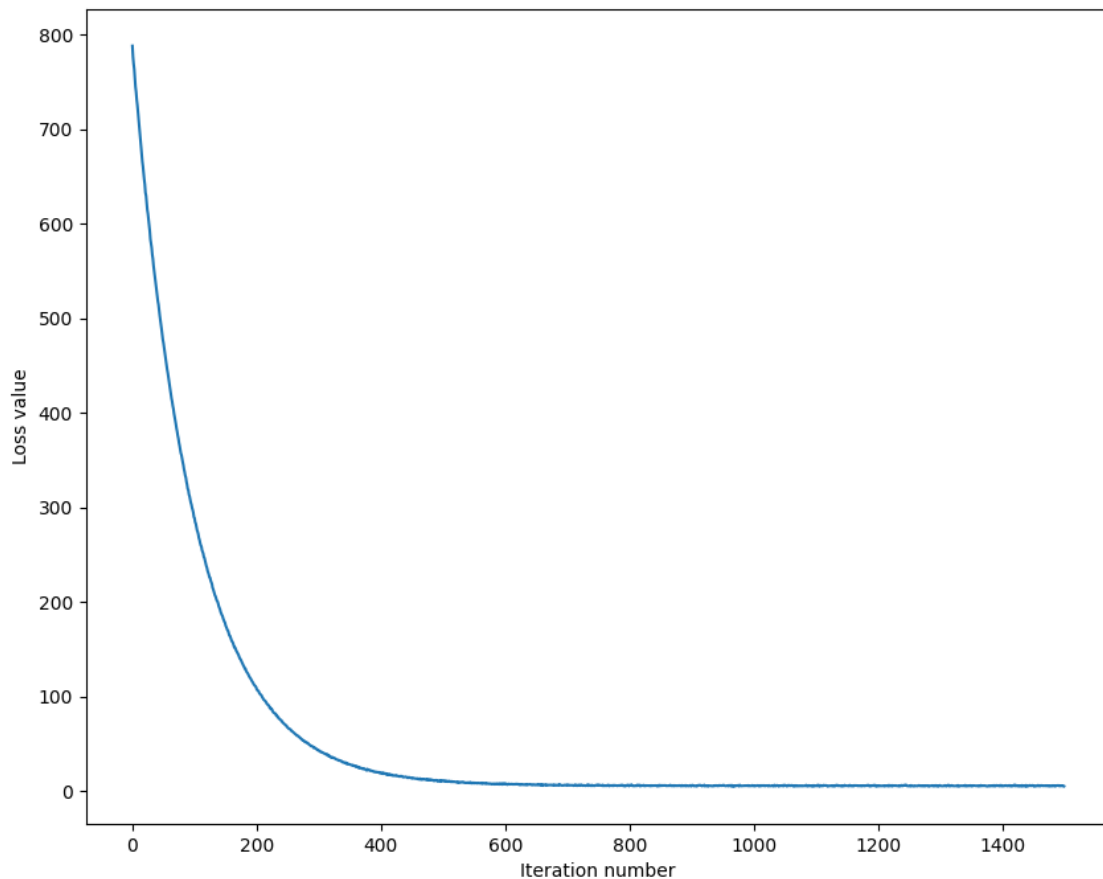
We now have vectorized and efficient expressions for the loss, the gradient and our gradient matches the numerical gradient. We are therefore ready to do SGD to minimize the loss. Your code for this part will be written inside `cs231n/classifiers/linear_classifier.py`.

```
[47]: # In the file linear_classifier.py, implement SGD in the function
# LinearClassifier.train() and then run it with the code below.
from cs231n.classifiers import LinearSVM
svm = LinearSVM()
tic = time.time()
loss_hist = svm.train(X_train, y_train, learning_rate=1e-7, reg=2.5e4,
                      num_iters=1500, verbose=True)
toc = time.time()
print('That took %fs' % (toc - tic))
```

```
iteration 0 / 1500: loss 787.942203
iteration 100 / 1500: loss 288.521461
iteration 200 / 1500: loss 107.473239
iteration 300 / 1500: loss 42.500261
iteration 400 / 1500: loss 18.575221
iteration 500 / 1500: loss 10.498323
iteration 600 / 1500: loss 7.789996
iteration 700 / 1500: loss 6.401790
```

```
iteration 800 / 1500: loss 5.852610
iteration 900 / 1500: loss 5.771643
iteration 1000 / 1500: loss 5.695896
iteration 1100 / 1500: loss 5.215097
iteration 1200 / 1500: loss 5.016683
iteration 1300 / 1500: loss 5.232318
iteration 1400 / 1500: loss 4.976558
That took 2.301672s
```

```
[48]: # A useful debugging strategy is to plot the loss as a function of
# iteration number:
plt.plot(loss_hist)
plt.xlabel('Iteration number')
plt.ylabel('Loss value')
plt.show()
```



```
[49]: # Write the LinearSVM.predict function and evaluate the performance on both the
# training and validation set
y_train_pred = svm.predict(X_train)
```

```
print('training accuracy: %f' % (np.mean(y_train == y_train_pred), ))
y_val_pred = svm.predict(X_val)
print('validation accuracy: %f' % (np.mean(y_val == y_val_pred), ))
```

```
training accuracy: 0.371776
validation accuracy: 0.380000
```

```
[58]: # Use the validation set to tune hyperparameters (regularization strength and
# learning rate). You should experiment with different ranges for the learning
# rates and regularization strengths; if you are careful you should be able to
# get a classification accuracy of about 0.39 (> 0.385) on the validation set.

# Note: you may see runtime/overflow warnings during hyper-parameter search.
# This may be caused by extreme values, and is not a bug.

# results is dictionary mapping tuples of the form
# (learning_rate, regularization_strength) to tuples of the form
# (training_accuracy, validation_accuracy). The accuracy is simply the fraction
# of data points that are correctly classified.
results = {}
best_val = -1 # The highest validation accuracy that we have seen so far.
best_svm = None # The LinearSVM object that achieved the highest validation
    ↪ rate.

#####
# TODO:
# Write code that chooses the best hyperparameters by tuning on the validation #
# set. For each combination of hyperparameters, train a linear SVM on the #
# training set, compute its accuracy on the training and validation sets, and #
# store these numbers in the results dictionary. In addition, store the best #
# validation accuracy in best_val and the LinearSVM object that achieves this #
# accuracy in best_svm.
#
# Hint: You should use a small value for num_iters as you develop your #
# validation code so that the SVMs don't take much time to train; once you are #
# confident that your validation code works, you should rerun the validation #
# code with a larger value for num_iters.
#####

# Provided as a reference. You may or may not want to change these
    ↪ hyperparameters
learning_rates = [1e-7, 5e-5]
regularization_strengths = [2.5e4, 5e4]

# *****START OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****

import itertools
```

```

for lr, reg in itertools.product(learning_rates, regularization_strengths):
    svm = LinearSVM()
    svm.train(X_train, y_train, lr, reg, num_iters=10000)

    y_train_pred, y_val_pred = svm.predict(X_train), svm.predict(X_val)
    results[(lr, reg)] = np.mean(y_train == y_train_pred), np.mean(y_val ==
↪y_val_pred)
    if results[(lr, reg)][1] > best_val:
        best_val = results[(lr, reg)][1]
        best_svm = svm

# *****END OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****

# Print out results.
for lr, reg in sorted(results):
    train_accuracy, val_accuracy = results[(lr, reg)]
    print('lr %e reg %e train accuracy: %f val accuracy: %f' % (
        lr, reg, train_accuracy, val_accuracy))

print('best validation accuracy achieved during cross-validation: %f' %
↪best_val)

```

```

lr 1.000000e-07 reg 2.500000e+04 train accuracy: 0.369592 val accuracy: 0.369000
lr 1.000000e-07 reg 5.000000e+04 train accuracy: 0.361653 val accuracy: 0.372000
lr 5.000000e-05 reg 2.500000e+04 train accuracy: 0.100265 val accuracy: 0.087000
lr 5.000000e-05 reg 5.000000e+04 train accuracy: 0.100265 val accuracy: 0.087000
best validation accuracy achieved during cross-validation: 0.372000

```

```

[59]: # Visualize the cross-validation results
import math
import pdb

# pdb.set_trace()

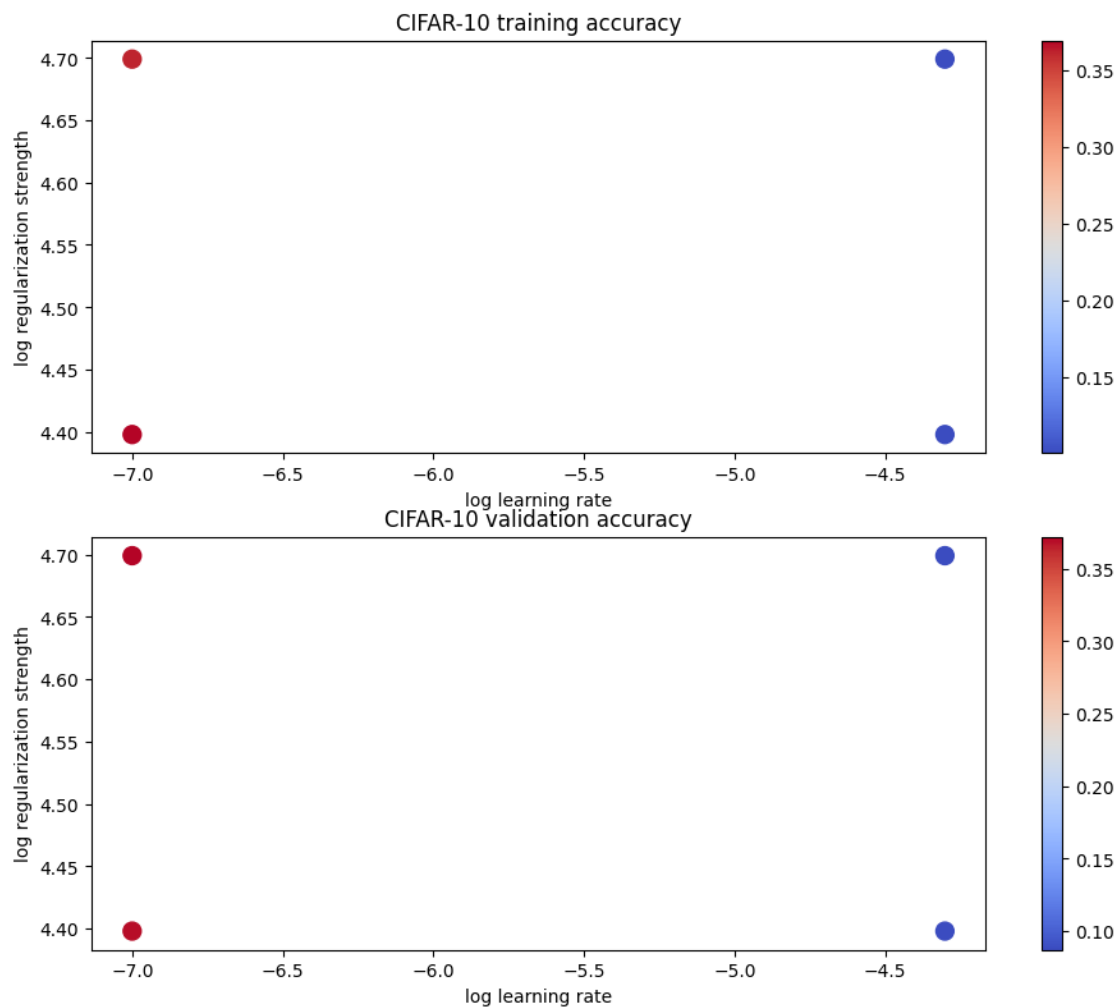
x_scatter = [math.log10(x[0]) for x in results]
y_scatter = [math.log10(x[1]) for x in results]

# plot training accuracy
marker_size = 100
colors = [results[x][0] for x in results]
plt.subplot(2, 1, 1)
plt.tight_layout(pad=3)
plt.scatter(x_scatter, y_scatter, marker_size, c=colors, cmap=plt.cm.coolwarm)
plt.colorbar()
plt.xlabel('log learning rate')
plt.ylabel('log regularization strength')

```

```
plt.title('CIFAR-10 training accuracy')

# plot validation accuracy
colors = [results[x][1] for x in results] # default size of markers is 20
plt.subplot(2, 1, 2)
plt.scatter(x_scatter, y_scatter, marker_size, c=colors, cmap=plt.cm.coolwarm)
plt.colorbar()
plt.xlabel('log learning rate')
plt.ylabel('log regularization strength')
plt.title('CIFAR-10 validation accuracy')
plt.show()
```



```
[55]: # Evaluate the best svm on test set
y_test_pred = best_svm.predict(X_test)
test_accuracy = np.mean(y_test == y_test_pred)
```

```
print('linear SVM on raw pixels final test set accuracy: %f' % test_accuracy)
```

linear SVM on raw pixels final test set accuracy: 0.351000

```
[56]: # Visualize the learned weights for each class.
# Depending on your choice of learning rate and regularization strength, these
# may
# or may not be nice to look at.
w = best_svm.W[:-1,:] # strip out the bias
w = w.reshape(32, 32, 3, 10)
w_min, w_max = np.min(w), np.max(w)
classes = ['plane', 'car', 'bird', 'cat', 'deer', 'dog', 'frog', 'horse', 'ship', 'truck']
for i in range(10):
    plt.subplot(2, 5, i + 1)

    # Rescale the weights to be between 0 and 255
    wimg = 255.0 * (w[:, :, :, i].squeeze() - w_min) / (w_max - w_min)
    plt.imshow(wimg.astype('uint8'))
    plt.axis('off')
    plt.title(classes[i])
```



Inline question 2

Describe what your visualized SVM weights look like, and offer a brief explanation for why they

look the way they do.

***Your Answer :** It seems like the SVM weights are blurred versions of the images for their class, maybe as the average of all the images for each respective class which comes out to the final blurred image for every class. The SVM is trained to find distinct features for every class, and we can see that it generally does so in the final weights for every class.*

softmax

November 27, 2024

```
[ ]: # This mounts your Google Drive to the Colab VM.
from google.colab import drive
drive.mount('/content/drive')

# TODO: Enter the foldername in your Drive where you have saved the unzipped
# assignment folder, e.g. 'cs231n/assignments/assignment1/'
FOLDERNAME = None
assert FOLDERNAME is not None, "[!] Enter the foldername."

# Now that we've mounted your Drive, this ensures that
# the Python interpreter of the Colab VM can load
# python files from within it.
import sys
sys.path.append('/content/drive/My Drive/{}'.format(FOLDERNAME))

# This downloads the CIFAR-10 dataset to your Drive
# if it doesn't already exist.
%cd /content/drive/My\ Drive/$FOLDERNAME/cs231n/datasets/
!bash get_datasets.sh
%cd /content/drive/My\ Drive/$FOLDERNAME
```

1 Softmax exercise

Complete and hand in this completed worksheet (including its outputs and any supporting code outside of the worksheet) with your assignment submission. For more details see the [assignments page](#) on the course website.

This exercise is analogous to the SVM exercise. You will:

- implement a fully-vectorized **loss function** for the Softmax classifier
- implement the fully-vectorized expression for its **analytic gradient**
- **check your implementation** with numerical gradient
- use a validation set to **tune the learning rate and regularization strength**
- **optimize** the loss function with **SGD**
- **visualize** the final learned weights

```
[1]: import random
import numpy as np
```

```

from cs231n.data_utils import load_CIFAR10
import matplotlib.pyplot as plt

%matplotlib inline
plt.rcParams['figure.figsize'] = (10.0, 8.0) # set default size of plots
plt.rcParams['image.interpolation'] = 'nearest'
plt.rcParams['image.cmap'] = 'gray'

# for auto-reloading external modules
# see http://stackoverflow.com/questions/1907993/
↳ autoreload-of-modules-in-ipython
%load_ext autoreload
%autoreload 2

```

```

[2]: def get_CIFAR10_data(num_training=49000, num_validation=1000, num_test=1000,↳
↳ num_dev=500):
    """
    Load the CIFAR-10 dataset from disk and perform preprocessing to prepare
    it for the linear classifier. These are the same steps as we used for the
    SVM, but condensed to a single function.
    """
    # Load the raw CIFAR-10 data
    cifar10_dir = 'cs231n/datasets/cifar-10-batches-py'

    # Cleaning up variables to prevent loading data multiple times (which may↳
    ↳ cause memory issue)
    try:
        del X_train, y_train
        del X_test, y_test
        print('Clear previously loaded data.')
    except:
        pass

    X_train, y_train, X_test, y_test = load_CIFAR10(cifar10_dir)

    # subsample the data
    mask = list(range(num_training, num_training + num_validation))
    X_val = X_train[mask]
    y_val = y_train[mask]
    mask = list(range(num_training))
    X_train = X_train[mask]
    y_train = y_train[mask]
    mask = list(range(num_test))
    X_test = X_test[mask]
    y_test = y_test[mask]
    mask = np.random.choice(num_training, num_dev, replace=False)
    X_dev = X_train[mask]

```

```

y_dev = y_train[mask]

# Preprocessing: reshape the image data into rows
X_train = np.reshape(X_train, (X_train.shape[0], -1))
X_val = np.reshape(X_val, (X_val.shape[0], -1))
X_test = np.reshape(X_test, (X_test.shape[0], -1))
X_dev = np.reshape(X_dev, (X_dev.shape[0], -1))

# Normalize the data: subtract the mean image
mean_image = np.mean(X_train, axis = 0)
X_train -= mean_image
X_val -= mean_image
X_test -= mean_image
X_dev -= mean_image

# add bias dimension and transform into columns
X_train = np.hstack([X_train, np.ones((X_train.shape[0], 1))])
X_val = np.hstack([X_val, np.ones((X_val.shape[0], 1))])
X_test = np.hstack([X_test, np.ones((X_test.shape[0], 1))])
X_dev = np.hstack([X_dev, np.ones((X_dev.shape[0], 1))])

return X_train, y_train, X_val, y_val, X_test, y_test, X_dev, y_dev

# Invoke the above function to get our data.
X_train, y_train, X_val, y_val, X_test, y_test, X_dev, y_dev = ↳ get_CIFAR10_data()
print('Train data shape: ', X_train.shape)
print('Train labels shape: ', y_train.shape)
print('Validation data shape: ', X_val.shape)
print('Validation labels shape: ', y_val.shape)
print('Test data shape: ', X_test.shape)
print('Test labels shape: ', y_test.shape)
print('dev data shape: ', X_dev.shape)
print('dev labels shape: ', y_dev.shape)

```

```

Train data shape: (49000, 3073)
Train labels shape: (49000,)
Validation data shape: (1000, 3073)
Validation labels shape: (1000,)
Test data shape: (1000, 3073)
Test labels shape: (1000,)
dev data shape: (500, 3073)
dev labels shape: (500,)

```

1.1 Softmax Classifier

Your code for this section will all be written inside `cs231n/classifiers/softmax.py`.

```
[7]: # First implement the naive softmax loss function with nested loops.
# Open the file cs231n/classifiers/softmax.py and implement the
# softmax_loss_naive function.

from cs231n.classifiers.softmax import softmax_loss_naive
import time

# Generate a random softmax weight matrix and use it to compute the loss.
W = np.random.randn(3073, 10) * 0.0001
loss, grad = softmax_loss_naive(W, X_dev, y_dev, 0.0)

# As a rough sanity check, our loss should be something close to -log(0.1).
print('loss: %f' % loss)
print('sanity check: %f' % (-np.log(0.1)))
```

loss: 2.285304

sanity check: 2.302585

Inline Question 1

Why do we expect our loss to be close to $-\log(0.1)$? Explain briefly.**

Your Answer: We made 10 classes so the chance to randomly guess the correct class is 1 in 10, and our softmax loss function takes the negative log of the probability.

```
[8]: # Complete the implementation of softmax_loss_naive and implement a (naive)
# version of the gradient that uses nested loops.
loss, grad = softmax_loss_naive(W, X_dev, y_dev, 0.0)

# As we did for the SVM, use numeric gradient checking as a debugging tool.
# The numeric gradient should be close to the analytic gradient.
from cs231n.gradient_check import grad_check_sparse
f = lambda w: softmax_loss_naive(w, X_dev, y_dev, 0.0)[0]
grad_numerical = grad_check_sparse(f, W, grad, 10)

# similar to SVM case, do another gradient check with regularization
loss, grad = softmax_loss_naive(W, X_dev, y_dev, 5e1)
f = lambda w: softmax_loss_naive(w, X_dev, y_dev, 5e1)[0]
grad_numerical = grad_check_sparse(f, W, grad, 10)
```

numerical: -0.847350 analytic: -0.847350, relative error: 1.815383e-08

numerical: 0.536715 analytic: 0.536715, relative error: 3.631984e-08

numerical: 0.105854 analytic: 0.105853, relative error: 2.697530e-07

numerical: 0.212442 analytic: 0.212442, relative error: 2.231282e-07

numerical: -3.683886 analytic: -3.683886, relative error: 1.986645e-08

numerical: -1.687341 analytic: -1.687342, relative error: 3.446106e-08

numerical: -0.181585 analytic: -0.181586, relative error: 1.668307e-07

numerical: -1.390851 analytic: -1.390851, relative error: 2.939417e-08

numerical: 0.345366 analytic: 0.345365, relative error: 1.979911e-07

numerical: -0.659346 analytic: -0.659346, relative error: 1.212902e-07

```

numerical: -0.674677 analytic: -0.674677, relative error: 3.939843e-08
numerical: -1.405604 analytic: -1.405604, relative error: 1.606118e-08
numerical: 0.542502 analytic: 0.542502, relative error: 1.177172e-07
numerical: -0.291087 analytic: -0.291088, relative error: 1.150962e-07
numerical: 0.531040 analytic: 0.531040, relative error: 4.504543e-09
numerical: 2.869777 analytic: 2.869777, relative error: 2.236330e-08
numerical: 1.605881 analytic: 1.605881, relative error: 7.500883e-09
numerical: -1.832096 analytic: -1.832096, relative error: 9.675029e-09
numerical: 0.537810 analytic: 0.537810, relative error: 2.739524e-09
numerical: -0.454793 analytic: -0.454793, relative error: 8.666566e-08

```

```

[10]: # Now that we have a naive implementation of the softmax loss function and its
      ↪ gradient,
      # implement a vectorized version in softmax_loss_vectorized.
      # The two versions should compute the same results, but the vectorized version
      ↪ should be
      # much faster.
tic = time.time()
loss_naive, grad_naive = softmax_loss_naive(W, X_dev, y_dev, 0.000005)
toc = time.time()
print('naive loss: %e computed in %fs' % (loss_naive, toc - tic))

from cs231n.classifiers.softmax import softmax_loss_vectorized
tic = time.time()
loss_vectorized, grad_vectorized = softmax_loss_vectorized(W, X_dev, y_dev, 0.
    ↪ 000005)
toc = time.time()
print('vectorized loss: %e computed in %fs' % (loss_vectorized, toc - tic))

# As we did for the SVM, we use the Frobenius norm to compare the two versions
# of the gradient.
grad_difference = np.linalg.norm(grad_naive - grad_vectorized, ord='fro')
print('Loss difference: %f' % np.abs(loss_naive - loss_vectorized))
print('Gradient difference: %f' % grad_difference)

```

```

naive loss: 2.285304e+00 computed in 0.051379s
vectorized loss: 2.285304e+00 computed in 0.002026s
Loss difference: 0.000000
Gradient difference: 0.000000

```

```

[14]: # Use the validation set to tune hyperparameters (regularization strength and
      # learning rate). You should experiment with different ranges for the learning
      # rates and regularization strengths; if you are careful you should be able to
      # get a classification accuracy of over 0.35 on the validation set.

from cs231n.classifiers import Softmax
results = {}
best_val = -1

```

```

best_softmax = None

#####
# TODO:
# Use the validation set to set the learning rate and regularization strength. #
# This should be identical to the validation that you did for the SVM; save #
# the best trained softmax classifier in best_softmax. #
#####

# Provided as a reference. You may or may not want to change these
↳ hyperparameters
learning_rates = [1e-7, 5e-7, 5]
regularization_strengths = [2.5e4, 5e4, 5]

# *****START OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****

import itertools

for lr, reg in itertools.product(learning_rates, regularization_strengths):
    softmax = Softmax()
    softmax.train(X_train, y_train, lr, reg, num_iters=1_000)

    y_train_pred, y_val_pred = softmax.predict(X_train), softmax.predict(X_val)
    results[(lr, reg)] = np.mean(y_train == y_train_pred), np.mean(y_val ==
↳ y_val_pred)

    if results[(lr, reg)][1] > best_val:
        best_val = results[(lr, reg)][1]
        best_softmax = softmax

# *****END OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****

# Print out results.
for lr, reg in sorted(results):
    train_accuracy, val_accuracy = results[(lr, reg)]
    print('lr %e reg %e train accuracy: %f val accuracy: %f' % (
        lr, reg, train_accuracy, val_accuracy))

print('best validation accuracy achieved during cross-validation: %f' %
↳ best_val)

```

```

/Users/kimsh/Downloads/course/fall24/csc480/hw4/assignment1/cs231n/classifiers/s
oftmax.py:73: RuntimeWarning: invalid value encountered in divide
    probs /= probs.sum(axis=1, keepdims=True)
/Users/kimsh/Downloads/course/fall24/csc480/hw4/assignment1/cs231n/classifiers/s
oftmax.py:74: RuntimeWarning: divide by zero encountered in log
    loss = -np.log(probs[range(X.shape[0]), y]).sum()

```

```

lr 1.000000e-07 reg 5.000000e+00 train accuracy: 0.224429 val accuracy: 0.222000
lr 1.000000e-07 reg 2.500000e+04 train accuracy: 0.327653 val accuracy: 0.343000
lr 1.000000e-07 reg 5.000000e+04 train accuracy: 0.305224 val accuracy: 0.328000
lr 5.000000e-07 reg 5.000000e+00 train accuracy: 0.295122 val accuracy: 0.296000
lr 5.000000e-07 reg 2.500000e+04 train accuracy: 0.313633 val accuracy: 0.331000
lr 5.000000e-07 reg 5.000000e+04 train accuracy: 0.310776 val accuracy: 0.326000
lr 5.000000e+00 reg 5.000000e+00 train accuracy: 0.100265 val accuracy: 0.087000
lr 5.000000e+00 reg 2.500000e+04 train accuracy: 0.100265 val accuracy: 0.087000
lr 5.000000e+00 reg 5.000000e+04 train accuracy: 0.100265 val accuracy: 0.087000
best validation accuracy achieved during cross-validation: 0.343000

```

```

[15]: # evaluate on test set
# Evaluate the best softmax on test set
y_test_pred = best_softmax.predict(X_test)
test_accuracy = np.mean(y_test == y_test_pred)
print('softmax on raw pixels final test set accuracy: %f' % (test_accuracy, ))

```

softmax on raw pixels final test set accuracy: 0.336000

Inline Question 2 - True or False

Suppose the overall training loss is defined as the sum of the per-datapoint loss over all training examples. It is possible to add a new datapoint to a training set that would leave the SVM loss unchanged, but this is not the case with the Softmax classifier loss.

Your Answer :

Yes

Your Explanation :

It can be done for SVM loss because SVM only considers the support vectors, while it cannot be done for softmax because softmax accounts for all probabilities for so a new datapoint is guaranteed to affect the softmax loss.

```

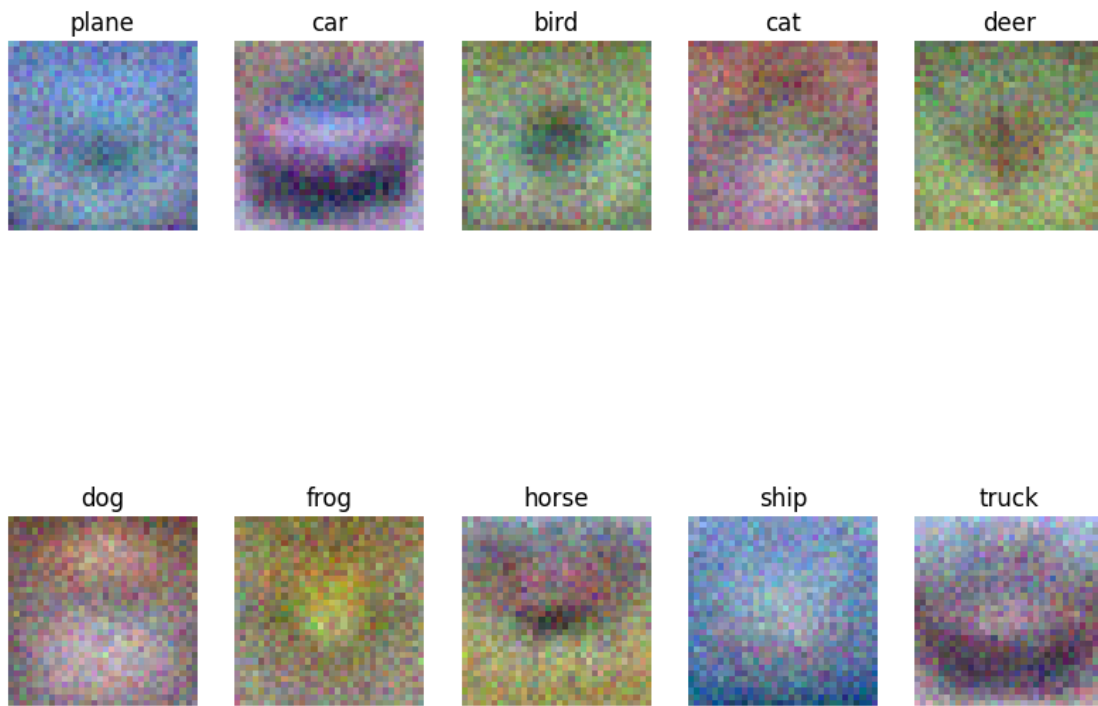
[16]: # Visualize the learned weights for each class
w = best_softmax.W[:-1,:] # strip out the bias
w = w.reshape(32, 32, 3, 10)

w_min, w_max = np.min(w), np.max(w)

classes = ['plane', 'car', 'bird', 'cat', 'deer', 'dog', 'frog', 'horse', '
↳ 'ship', 'truck']
for i in range(10):
    plt.subplot(2, 5, i + 1)

    # Rescale the weights to be between 0 and 255
    wimg = 255.0 * (w[:, :, :, i].squeeze() - w_min) / (w_max - w_min)
    plt.imshow(wimg.astype('uint8'))
    plt.axis('off')
    plt.title(classes[i])

```



[]:

two_layer_net

November 27, 2024

```
[ ]: # This mounts your Google Drive to the Colab VM.
from google.colab import drive
drive.mount('/content/drive')

# TODO: Enter the foldername in your Drive where you have saved the unzipped
# assignment folder, e.g. 'cs231n/assignments/assignment1/'
FOLDERNAME = None
assert FOLDERNAME is not None, "[!] Enter the foldername."

# Now that we've mounted your Drive, this ensures that
# the Python interpreter of the Colab VM can load
# python files from within it.
import sys
sys.path.append('/content/drive/My Drive/{}'.format(FOLDERNAME))

# This downloads the CIFAR-10 dataset to your Drive
# if it doesn't already exist.
%cd /content/drive/My\ Drive/$FOLDERNAME/cs231n/datasets/
!bash get_datasets.sh
%cd /content/drive/My\ Drive/$FOLDERNAME
```

1 Fully-Connected Neural Nets

In this exercise we will implement fully-connected networks using a modular approach. For each layer we will implement a **forward** and a **backward** function. The **forward** function will receive inputs, weights, and other parameters and will return both an output and a **cache** object storing data needed for the backward pass, like this:

```
def layer_forward(x, w):
    """ Receive inputs x and weights w """
    # Do some computations ...
    z = # ... some intermediate value
    # Do some more computations ...
    out = # the output

    cache = (x, w, z, out) # Values we need to compute gradients

    return out, cache
```

The backward pass will receive upstream derivatives and the `cache` object, and will return gradients with respect to the inputs and weights, like this:

```
def layer_backward(dout, cache):
    """
    Receive dout (derivative of loss with respect to outputs) and cache,
    and compute derivative with respect to inputs.
    """
    # Unpack cache values
    x, w, z, out = cache

    # Use values in cache to compute derivatives
    dx = # Derivative of loss with respect to x
    dw = # Derivative of loss with respect to w

    return dx, dw
```

After implementing a bunch of layers this way, we will be able to easily combine them to build classifiers with different architectures.

```
[1]: # As usual, a bit of setup
from __future__ import print_function
import time
import numpy as np
import matplotlib.pyplot as plt
from cs231n.classifiers.fc_net import *
from cs231n.data_utils import get_CIFAR10_data
from cs231n.gradient_check import eval_numerical_gradient, \
    eval_numerical_gradient_array
from cs231n.solver import Solver

%matplotlib inline
plt.rcParams['figure.figsize'] = (10.0, 8.0) # set default size of plots
plt.rcParams['image.interpolation'] = 'nearest'
plt.rcParams['image.cmap'] = 'gray'

# for auto-reloading external modules
# see http://stackoverflow.com/questions/1907993/
# autoreload-of-modules-in-ipython
%load_ext autoreload
%autoreload 2

def rel_error(x, y):
    """ returns relative error """
    return np.max(np.abs(x - y) / (np.maximum(1e-8, np.abs(x) + np.abs(y))))

[2]: # Load the (preprocessed) CIFAR10 data.
```

```
data = get_CIFAR10_data()
for k, v in list(data.items()):
    print('%s: ' % k, v.shape)
```

```
('X_train: ', (49000, 3, 32, 32))
('y_train: ', (49000,))
('X_val: ', (1000, 3, 32, 32))
('y_val: ', (1000,))
('X_test: ', (1000, 3, 32, 32))
('y_test: ', (1000,))
```

2 Affine layer: forward

Open the file `cs231n/layers.py` and implement the `affine_forward` function.

Once you are done you can test your implementaion by running the following:

```
[5]: # Test the affine_forward function

num_inputs = 2
input_shape = (4, 5, 6)
output_dim = 3

input_size = num_inputs * np.prod(input_shape)
weight_size = output_dim * np.prod(input_shape)

x = np.linspace(-0.1, 0.5, num=input_size).reshape(num_inputs, *input_shape)
w = np.linspace(-0.2, 0.3, num=weight_size).reshape(np.prod(input_shape),
    ↪output_dim)
b = np.linspace(-0.3, 0.1, num=output_dim)

out, _ = affine_forward(x, w, b)
correct_out = np.array([[ 1.49834967,  1.70660132,  1.91485297],
                        [ 3.25553199,  3.5141327,  3.77273342]])

# Compare your output with ours. The error should be around e-9 or less.
print('Testing affine_forward function:')
print('difference: ', rel_error(out, correct_out))
```

```
Testing affine_forward function:
difference: 9.769849468192957e-10
```

3 Affine layer: backward

Now implement the `affine_backward` function and test your implementation using numeric gradient checking.

```
[8]: # Test the affine_backward function
np.random.seed(231)
x = np.random.randn(10, 2, 3)
w = np.random.randn(6, 5)
b = np.random.randn(5)
dout = np.random.randn(10, 5)

dx_num = eval_numerical_gradient_array(lambda x: affine_forward(x, w, b)[0], x,
    ↪dout)
dw_num = eval_numerical_gradient_array(lambda w: affine_forward(x, w, b)[0], w,
    ↪dout)
db_num = eval_numerical_gradient_array(lambda b: affine_forward(x, w, b)[0], b,
    ↪dout)

_, cache = affine_forward(x, w, b)
dx, dw, db = affine_backward(dout, cache)

# The error should be around e-10 or less
print('Testing affine_backward function:')
print('dx error: ', rel_error(dx_num, dx))
print('dw error: ', rel_error(dw_num, dw))
print('db error: ', rel_error(db_num, db))
```

```
Testing affine_backward function:
dx error: 9.579607670311995e-11
dw error: 1.2604803862099753e-10
db error: 7.72167374950876e-11
```

4 ReLU activation: forward

Implement the forward pass for the ReLU activation function in the `relu_forward` function and test your implementation using the following:

```
[9]: # Test the relu_forward function

x = np.linspace(-0.5, 0.5, num=12).reshape(3, 4)

out, _ = relu_forward(x)
correct_out = np.array([[ 0.,          0.,          0.,          0.],
                        [ 0.,          0.,          0.04545455, 0.13636364],
                        [ 0.22727273, 0.31818182, 0.40909091, 0.5]])

# Compare your output with ours. The error should be on the order of e-8
print('Testing relu_forward function:')
print('difference: ', rel_error(out, correct_out))
```

```
Testing relu_forward function:
difference: 4.999999798022158e-08
```

5 ReLU activation: backward

Now implement the backward pass for the ReLU activation function in the `relu_backward` function and test your implementation using numeric gradient checking:

```
[10]: np.random.seed(231)
x = np.random.randn(10, 10)
dout = np.random.randn(*x.shape)

dx_num = eval_numerical_gradient_array(lambda x: relu_forward(x)[0], x, dout)

_, cache = relu_forward(x)
dx = relu_backward(dout, cache)

# The error should be on the order of e-12
print('Testing relu_backward function:')
print('dx error: ', rel_error(dx_num, dx))
```

```
Testing relu_backward function:
dx error:  3.2756349136310288e-12
```

5.1 Inline Question 1:

We've only asked you to implement ReLU, but there are a number of different activation functions that one could use in neural networks, each with its pros and cons. In particular, an issue commonly seen with activation functions is getting zero (or close to zero) gradient flow during backpropagation. Which of the following activation functions have this problem? If you consider these functions in the one dimensional case, what types of input would lead to this behaviour? 1. Sigmoid 2. ReLU 3. Leaky ReLU

5.2 Answer:

Sigmoid and ReLU have this problem.

The sigmoid function would be very close to 1 or 0 with a large positive or a large negative value. The derivative formula for sigmoid would then yield a gradient very close to 0.

If ReLU receives any negative value, the gradient will be 0 because of $\text{ReLU}(x) = \max(0, x)$.

6 “Sandwich” layers

There are some common patterns of layers that are frequently used in neural nets. For example, affine layers are frequently followed by a ReLU nonlinearity. To make these common patterns easy, we define several convenience layers in the file `cs231n/layer_utils.py`.

For now take a look at the `affine_relu_forward` and `affine_relu_backward` functions, and run the following to numerically gradient check the backward pass:

```
[12]: from cs231n.layer_utils import affine_relu_forward, affine_relu_backward
np.random.seed(231)
```

```

x = np.random.randn(2, 3, 4)
w = np.random.randn(12, 10)
b = np.random.randn(10)
dout = np.random.randn(2, 10)

out, cache = affine_relu_forward(x, w, b)
dx, dw, db = affine_relu_backward(dout, cache)

dx_num = eval_numerical_gradient_array(lambda x: affine_relu_forward(x, w,
    ↪b)[0], x, dout)
dw_num = eval_numerical_gradient_array(lambda w: affine_relu_forward(x, w,
    ↪b)[0], w, dout)
db_num = eval_numerical_gradient_array(lambda b: affine_relu_forward(x, w,
    ↪b)[0], b, dout)

# Relative error should be around e-10 or less
print('Testing affine_relu_forward and affine_relu_backward:')
print('dx error: ', rel_error(dx_num, dx))
print('dw error: ', rel_error(dw_num, dw))
print('db error: ', rel_error(db_num, db))

```

Testing affine_relu_forward and affine_relu_backward:

```

dx error:  1.70644420803027e-10
dw error:  8.162088973990222e-11
db error:  7.82672402145899e-12

```

7 Loss layers: Softmax and SVM

Now implement the loss and gradient for softmax and SVM in the `softmax_loss` and `svm_loss` function in `cs231n/layers.py`. These should be similar to what you implemented in `cs231n/classifiers/softmax.py` and `cs231n/classifiers/linear_svm.py`.

You can make sure that the implementations are correct by running the following:

```

[14]: np.random.seed(231)
num_classes, num_inputs = 10, 50
x = 0.001 * np.random.randn(num_inputs, num_classes)
y = np.random.randint(num_classes, size=num_inputs)

dx_num = eval_numerical_gradient(lambda x: svm_loss(x, y)[0], x, verbose=False)
loss, dx = svm_loss(x, y)

# Test svm_loss function. Loss should be around 9 and dx error should be around
    ↪the order of e-9
print('Testing svm_loss:')
print('loss: ', loss)
print('dx error: ', rel_error(dx_num, dx))

```

```

dx_num = eval_numerical_gradient(lambda x: softmax_loss(x, y)[0], x,
    verbose=False)
loss, dx = softmax_loss(x, y)

# Test softmax_loss function. Loss should be close to 2.3 and dx error should
    be around e-8
print('\nTesting softmax_loss:')
print('loss: ', loss)
print('dx error: ', rel_error(dx_num, dx))

```

```

Testing svm_loss:
loss: 8.999602749096233
dx error: 1.4021566006651672e-09

```

```

Testing softmax_loss:
loss: 2.3025458445007376
dx error: 8.234144091578429e-09

```

8 Two-layer network

Open the file `cs231n/classifiers/fc_net.py` and complete the implementation of the `TwoLayerNet` class. Read through it to make sure you understand the API. You can run the cell below to test your implementation.

```

[18]: np.random.seed(231)
N, D, H, C = 3, 5, 50, 7
X = np.random.randn(N, D)
y = np.random.randint(C, size=N)

std = 1e-3
model = TwoLayerNet(input_dim=D, hidden_dim=H, num_classes=C, weight_scale=std)

print('Testing initialization ... ')
W1_std = abs(model.params['W1'].std() - std)
b1 = model.params['b1']
W2_std = abs(model.params['W2'].std() - std)
b2 = model.params['b2']
assert W1_std < std / 10, 'First layer weights do not seem right'
assert np.all(b1 == 0), 'First layer biases do not seem right'
assert W2_std < std / 10, 'Second layer weights do not seem right'
assert np.all(b2 == 0), 'Second layer biases do not seem right'

print('Testing test-time forward pass ... ')
model.params['W1'] = np.linspace(-0.7, 0.3, num=D*H).reshape(D, H)
model.params['b1'] = np.linspace(-0.1, 0.9, num=H)
model.params['W2'] = np.linspace(-0.3, 0.4, num=H*C).reshape(H, C)

```

```

model.params['b2'] = np.linspace(-0.9, 0.1, num=C)
X = np.linspace(-5.5, 4.5, num=N*D).reshape(D, N).T
scores = model.loss(X)
correct_scores = np.asarray(
    [[11.53165108, 12.2917344, 13.05181771, 13.81190102, 14.57198434, 15.
↪33206765, 16.09215096],
    [12.05769098, 12.74614105, 13.43459113, 14.1230412, 14.81149128, 15.
↪49994135, 16.18839143],
    [12.58373087, 13.20054771, 13.81736455, 14.43418138, 15.05099822, 15.
↪66781506, 16.2846319 ]])
scores_diff = np.abs(scores - correct_scores).sum()
assert scores_diff < 1e-6, 'Problem with test-time forward pass'

print('Testing training loss (no regularization)')
y = np.asarray([0, 5, 1])
loss, grads = model.loss(X, y)
correct_loss = 3.4702243556
assert abs(loss - correct_loss) < 1e-10, 'Problem with training-time loss'

model.reg = 1.0
loss, grads = model.loss(X, y)
correct_loss = 26.5948426952
assert abs(loss - correct_loss) < 1e-10, 'Problem with regularization loss'

# Errors should be around e-7 or less
for reg in [0.0, 0.7]:
    print('Running numeric gradient check with reg = ', reg)
    model.reg = reg
    loss, grads = model.loss(X, y)

    for name in sorted(grads):
        f = lambda _: model.loss(X, y)[0]
        grad_num = eval_numerical_gradient(f, model.params[name], verbose=False)
        print('%s relative error: %.2e' % (name, rel_error(grad_num, grads[name])))

```

```

Testing initialization ...
Testing test-time forward pass ...
Testing training loss (no regularization)
Running numeric gradient check with reg = 0.0
W1 relative error: 1.22e-08
W2 relative error: 3.43e-10
b1 relative error: 9.83e-09
b2 relative error: 2.53e-10
Running numeric gradient check with reg = 0.7
W1 relative error: 2.53e-07
W2 relative error: 1.37e-07
b1 relative error: 1.56e-08

```

b2 relative error: 9.09e-10

9 Solver

Open the file `cs231n/solver.py` and read through it to familiarize yourself with the API. After doing so, use a `Solver` instance to train a `TwoLayerNet` that achieves about 36% accuracy on the validation set.

```
[23]: input_size = 32 * 32 * 3
      hidden_size = 50
      num_classes = 10
      model = TwoLayerNet(input_size, hidden_size, num_classes)
      solver = None

      #####
      # TODO: Use a Solver instance to train a TwoLayerNet that achieves about 36% #
      # accuracy on the validation set.                                           #
      #####
      # *****START OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****

      solver = Solver(model, data, optim_config={'learning_rate': 1e-4})
      solver.train()

      # *****END OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****
      #####
      #                               END OF YOUR CODE                               #
      #####
```

```
(Iteration 1 / 4900) loss: 2.302102
(Epoch 0 / 10) train acc: 0.118000; val_acc: 0.107000
(Iteration 11 / 4900) loss: 2.300908
(Iteration 21 / 4900) loss: 2.301796
(Iteration 31 / 4900) loss: 2.292424
(Iteration 41 / 4900) loss: 2.290181
(Iteration 51 / 4900) loss: 2.291646
(Iteration 61 / 4900) loss: 2.287809
(Iteration 71 / 4900) loss: 2.275967
(Iteration 81 / 4900) loss: 2.282732
(Iteration 91 / 4900) loss: 2.276917
(Iteration 101 / 4900) loss: 2.263951
(Iteration 111 / 4900) loss: 2.251224
(Iteration 121 / 4900) loss: 2.254937
(Iteration 131 / 4900) loss: 2.221195
(Iteration 141 / 4900) loss: 2.253322
(Iteration 151 / 4900) loss: 2.217544
(Iteration 161 / 4900) loss: 2.198788
(Iteration 171 / 4900) loss: 2.217726
(Iteration 181 / 4900) loss: 2.222377
```

(Iteration 191 / 4900) loss: 2.141118
(Iteration 201 / 4900) loss: 2.155688
(Iteration 211 / 4900) loss: 2.175805
(Iteration 221 / 4900) loss: 2.149188
(Iteration 231 / 4900) loss: 2.167391
(Iteration 241 / 4900) loss: 2.126053
(Iteration 251 / 4900) loss: 2.148920
(Iteration 261 / 4900) loss: 2.085525
(Iteration 271 / 4900) loss: 2.120765
(Iteration 281 / 4900) loss: 2.014044
(Iteration 291 / 4900) loss: 2.055537
(Iteration 301 / 4900) loss: 2.106982
(Iteration 311 / 4900) loss: 2.147217
(Iteration 321 / 4900) loss: 2.040379
(Iteration 331 / 4900) loss: 2.092310
(Iteration 341 / 4900) loss: 1.983126
(Iteration 351 / 4900) loss: 2.040591
(Iteration 361 / 4900) loss: 1.984567
(Iteration 371 / 4900) loss: 1.962470
(Iteration 381 / 4900) loss: 2.022835
(Iteration 391 / 4900) loss: 2.037144
(Iteration 401 / 4900) loss: 1.896024
(Iteration 411 / 4900) loss: 2.013811
(Iteration 421 / 4900) loss: 1.937920
(Iteration 431 / 4900) loss: 2.000016
(Iteration 441 / 4900) loss: 2.041284
(Iteration 451 / 4900) loss: 1.988159
(Iteration 461 / 4900) loss: 1.985817
(Iteration 471 / 4900) loss: 1.910620
(Iteration 481 / 4900) loss: 1.882914
(Epoch 1 / 10) train acc: 0.270000; val_acc: 0.298000
(Iteration 491 / 4900) loss: 1.930549
(Iteration 501 / 4900) loss: 1.861093
(Iteration 511 / 4900) loss: 1.889417
(Iteration 521 / 4900) loss: 1.987276
(Iteration 531 / 4900) loss: 1.993881
(Iteration 541 / 4900) loss: 1.959689
(Iteration 551 / 4900) loss: 1.794709
(Iteration 561 / 4900) loss: 1.991907
(Iteration 571 / 4900) loss: 2.019825
(Iteration 581 / 4900) loss: 1.875447
(Iteration 591 / 4900) loss: 1.906721
(Iteration 601 / 4900) loss: 1.855424
(Iteration 611 / 4900) loss: 1.958489
(Iteration 621 / 4900) loss: 1.956931
(Iteration 631 / 4900) loss: 1.880399
(Iteration 641 / 4900) loss: 1.863503
(Iteration 651 / 4900) loss: 1.821545

(Iteration 661 / 4900) loss: 1.871613
(Iteration 671 / 4900) loss: 2.016463
(Iteration 681 / 4900) loss: 1.872955
(Iteration 691 / 4900) loss: 1.839872
(Iteration 701 / 4900) loss: 1.885274
(Iteration 711 / 4900) loss: 1.829607
(Iteration 721 / 4900) loss: 1.954651
(Iteration 731 / 4900) loss: 1.880553
(Iteration 741 / 4900) loss: 1.799373
(Iteration 751 / 4900) loss: 1.828438
(Iteration 761 / 4900) loss: 1.871462
(Iteration 771 / 4900) loss: 1.880179
(Iteration 781 / 4900) loss: 1.783497
(Iteration 791 / 4900) loss: 1.867526
(Iteration 801 / 4900) loss: 1.888777
(Iteration 811 / 4900) loss: 1.888137
(Iteration 821 / 4900) loss: 1.932811
(Iteration 831 / 4900) loss: 1.768927
(Iteration 841 / 4900) loss: 1.860391
(Iteration 851 / 4900) loss: 1.771999
(Iteration 861 / 4900) loss: 1.901336
(Iteration 871 / 4900) loss: 1.855556
(Iteration 881 / 4900) loss: 1.770226
(Iteration 891 / 4900) loss: 1.762578
(Iteration 901 / 4900) loss: 1.881031
(Iteration 911 / 4900) loss: 1.976724
(Iteration 921 / 4900) loss: 1.716640
(Iteration 931 / 4900) loss: 1.871940
(Iteration 941 / 4900) loss: 1.849494
(Iteration 951 / 4900) loss: 1.688414
(Iteration 961 / 4900) loss: 1.753983
(Iteration 971 / 4900) loss: 1.706138
(Epoch 2 / 10) train acc: 0.359000; val_acc: 0.355000
(Iteration 981 / 4900) loss: 1.817515
(Iteration 991 / 4900) loss: 1.720335
(Iteration 1001 / 4900) loss: 1.816449
(Iteration 1011 / 4900) loss: 1.925526
(Iteration 1021 / 4900) loss: 1.829712
(Iteration 1031 / 4900) loss: 1.754395
(Iteration 1041 / 4900) loss: 1.903269
(Iteration 1051 / 4900) loss: 1.849956
(Iteration 1061 / 4900) loss: 1.744765
(Iteration 1071 / 4900) loss: 1.613304
(Iteration 1081 / 4900) loss: 1.839856
(Iteration 1091 / 4900) loss: 1.874830
(Iteration 1101 / 4900) loss: 1.830405
(Iteration 1111 / 4900) loss: 1.874367
(Iteration 1121 / 4900) loss: 1.770244

(Iteration 1131 / 4900) loss: 1.645877
(Iteration 1141 / 4900) loss: 1.596040
(Iteration 1151 / 4900) loss: 1.759439
(Iteration 1161 / 4900) loss: 1.953563
(Iteration 1171 / 4900) loss: 1.691994
(Iteration 1181 / 4900) loss: 1.881505
(Iteration 1191 / 4900) loss: 1.815949
(Iteration 1201 / 4900) loss: 1.776264
(Iteration 1211 / 4900) loss: 1.615266
(Iteration 1221 / 4900) loss: 1.769983
(Iteration 1231 / 4900) loss: 1.865118
(Iteration 1241 / 4900) loss: 1.834564
(Iteration 1251 / 4900) loss: 1.741519
(Iteration 1261 / 4900) loss: 1.781505
(Iteration 1271 / 4900) loss: 1.753204
(Iteration 1281 / 4900) loss: 1.892204
(Iteration 1291 / 4900) loss: 1.824438
(Iteration 1301 / 4900) loss: 1.715834
(Iteration 1311 / 4900) loss: 1.671238
(Iteration 1321 / 4900) loss: 1.605737
(Iteration 1331 / 4900) loss: 1.925910
(Iteration 1341 / 4900) loss: 1.688087
(Iteration 1351 / 4900) loss: 1.632900
(Iteration 1361 / 4900) loss: 1.639811
(Iteration 1371 / 4900) loss: 1.744500
(Iteration 1381 / 4900) loss: 1.837114
(Iteration 1391 / 4900) loss: 1.766576
(Iteration 1401 / 4900) loss: 1.713915
(Iteration 1411 / 4900) loss: 1.929609
(Iteration 1421 / 4900) loss: 1.825110
(Iteration 1431 / 4900) loss: 1.826942
(Iteration 1441 / 4900) loss: 1.622059
(Iteration 1451 / 4900) loss: 1.751760
(Iteration 1461 / 4900) loss: 1.788722
(Epoch 3 / 10) train acc: 0.378000; val_acc: 0.404000
(Iteration 1471 / 4900) loss: 1.677464
(Iteration 1481 / 4900) loss: 1.721759
(Iteration 1491 / 4900) loss: 1.649635
(Iteration 1501 / 4900) loss: 1.766489
(Iteration 1511 / 4900) loss: 1.721497
(Iteration 1521 / 4900) loss: 1.611361
(Iteration 1531 / 4900) loss: 1.723699
(Iteration 1541 / 4900) loss: 1.808270
(Iteration 1551 / 4900) loss: 1.664691
(Iteration 1561 / 4900) loss: 1.710850
(Iteration 1571 / 4900) loss: 1.651666
(Iteration 1581 / 4900) loss: 1.659262
(Iteration 1591 / 4900) loss: 1.717320

(Iteration 1601 / 4900) loss: 1.574300
(Iteration 1611 / 4900) loss: 1.675878
(Iteration 1621 / 4900) loss: 1.469033
(Iteration 1631 / 4900) loss: 1.725571
(Iteration 1641 / 4900) loss: 1.760362
(Iteration 1651 / 4900) loss: 1.682168
(Iteration 1661 / 4900) loss: 1.548418
(Iteration 1671 / 4900) loss: 1.576488
(Iteration 1681 / 4900) loss: 1.742529
(Iteration 1691 / 4900) loss: 1.557709
(Iteration 1701 / 4900) loss: 1.654732
(Iteration 1711 / 4900) loss: 1.579859
(Iteration 1721 / 4900) loss: 1.626347
(Iteration 1731 / 4900) loss: 1.697002
(Iteration 1741 / 4900) loss: 1.758854
(Iteration 1751 / 4900) loss: 1.756046
(Iteration 1761 / 4900) loss: 1.674506
(Iteration 1771 / 4900) loss: 1.689144
(Iteration 1781 / 4900) loss: 1.721371
(Iteration 1791 / 4900) loss: 1.837325
(Iteration 1801 / 4900) loss: 1.625401
(Iteration 1811 / 4900) loss: 1.548145
(Iteration 1821 / 4900) loss: 1.648151
(Iteration 1831 / 4900) loss: 1.628719
(Iteration 1841 / 4900) loss: 1.574192
(Iteration 1851 / 4900) loss: 1.691064
(Iteration 1861 / 4900) loss: 1.704529
(Iteration 1871 / 4900) loss: 1.746679
(Iteration 1881 / 4900) loss: 1.710360
(Iteration 1891 / 4900) loss: 1.638895
(Iteration 1901 / 4900) loss: 1.667065
(Iteration 1911 / 4900) loss: 1.706432
(Iteration 1921 / 4900) loss: 2.039567
(Iteration 1931 / 4900) loss: 1.514035
(Iteration 1941 / 4900) loss: 1.361862
(Iteration 1951 / 4900) loss: 1.660818
(Epoch 4 / 10) train acc: 0.414000; val_acc: 0.436000
(Iteration 1961 / 4900) loss: 1.678254
(Iteration 1971 / 4900) loss: 1.745978
(Iteration 1981 / 4900) loss: 1.652642
(Iteration 1991 / 4900) loss: 1.577038
(Iteration 2001 / 4900) loss: 1.645475
(Iteration 2011 / 4900) loss: 1.596167
(Iteration 2021 / 4900) loss: 1.688684
(Iteration 2031 / 4900) loss: 1.708798
(Iteration 2041 / 4900) loss: 1.804021
(Iteration 2051 / 4900) loss: 1.622153
(Iteration 2061 / 4900) loss: 1.408181

(Iteration 2071 / 4900) loss: 1.527459
(Iteration 2081 / 4900) loss: 1.717300
(Iteration 2091 / 4900) loss: 1.763393
(Iteration 2101 / 4900) loss: 1.674782
(Iteration 2111 / 4900) loss: 1.528108
(Iteration 2121 / 4900) loss: 1.655178
(Iteration 2131 / 4900) loss: 1.564793
(Iteration 2141 / 4900) loss: 1.769674
(Iteration 2151 / 4900) loss: 1.591140
(Iteration 2161 / 4900) loss: 1.748279
(Iteration 2171 / 4900) loss: 1.691106
(Iteration 2181 / 4900) loss: 1.653623
(Iteration 2191 / 4900) loss: 1.612167
(Iteration 2201 / 4900) loss: 1.550677
(Iteration 2211 / 4900) loss: 1.692134
(Iteration 2221 / 4900) loss: 1.565274
(Iteration 2231 / 4900) loss: 1.629941
(Iteration 2241 / 4900) loss: 1.708645
(Iteration 2251 / 4900) loss: 1.636191
(Iteration 2261 / 4900) loss: 1.626554
(Iteration 2271 / 4900) loss: 1.682221
(Iteration 2281 / 4900) loss: 1.710923
(Iteration 2291 / 4900) loss: 1.641846
(Iteration 2301 / 4900) loss: 1.658430
(Iteration 2311 / 4900) loss: 1.516800
(Iteration 2321 / 4900) loss: 1.768901
(Iteration 2331 / 4900) loss: 1.647812
(Iteration 2341 / 4900) loss: 1.615122
(Iteration 2351 / 4900) loss: 1.607570
(Iteration 2361 / 4900) loss: 1.818294
(Iteration 2371 / 4900) loss: 1.679956
(Iteration 2381 / 4900) loss: 1.594127
(Iteration 2391 / 4900) loss: 1.698212
(Iteration 2401 / 4900) loss: 1.764747
(Iteration 2411 / 4900) loss: 1.753230
(Iteration 2421 / 4900) loss: 1.543571
(Iteration 2431 / 4900) loss: 1.622546
(Iteration 2441 / 4900) loss: 1.608447
(Epoch 5 / 10) train acc: 0.418000; val_acc: 0.442000
(Iteration 2451 / 4900) loss: 1.697712
(Iteration 2461 / 4900) loss: 1.640289
(Iteration 2471 / 4900) loss: 1.692039
(Iteration 2481 / 4900) loss: 1.601466
(Iteration 2491 / 4900) loss: 1.762369
(Iteration 2501 / 4900) loss: 1.480687
(Iteration 2511 / 4900) loss: 1.748283
(Iteration 2521 / 4900) loss: 1.559910
(Iteration 2531 / 4900) loss: 1.640704

(Iteration 2541 / 4900) loss: 1.640517
(Iteration 2551 / 4900) loss: 1.561414
(Iteration 2561 / 4900) loss: 1.579191
(Iteration 2571 / 4900) loss: 1.494693
(Iteration 2581 / 4900) loss: 1.529990
(Iteration 2591 / 4900) loss: 1.669293
(Iteration 2601 / 4900) loss: 1.633461
(Iteration 2611 / 4900) loss: 1.607789
(Iteration 2621 / 4900) loss: 1.615504
(Iteration 2631 / 4900) loss: 1.598174
(Iteration 2641 / 4900) loss: 1.579691
(Iteration 2651 / 4900) loss: 1.500795
(Iteration 2661 / 4900) loss: 1.664102
(Iteration 2671 / 4900) loss: 1.643728
(Iteration 2681 / 4900) loss: 1.378938
(Iteration 2691 / 4900) loss: 1.600440
(Iteration 2701 / 4900) loss: 1.768620
(Iteration 2711 / 4900) loss: 1.617305
(Iteration 2721 / 4900) loss: 1.670802
(Iteration 2731 / 4900) loss: 1.638704
(Iteration 2741 / 4900) loss: 1.517932
(Iteration 2751 / 4900) loss: 1.531886
(Iteration 2761 / 4900) loss: 1.460224
(Iteration 2771 / 4900) loss: 1.697834
(Iteration 2781 / 4900) loss: 1.573468
(Iteration 2791 / 4900) loss: 1.627062
(Iteration 2801 / 4900) loss: 1.706540
(Iteration 2811 / 4900) loss: 1.710123
(Iteration 2821 / 4900) loss: 1.634137
(Iteration 2831 / 4900) loss: 1.425701
(Iteration 2841 / 4900) loss: 1.464077
(Iteration 2851 / 4900) loss: 1.583102
(Iteration 2861 / 4900) loss: 1.544762
(Iteration 2871 / 4900) loss: 1.408978
(Iteration 2881 / 4900) loss: 1.801534
(Iteration 2891 / 4900) loss: 1.612375
(Iteration 2901 / 4900) loss: 1.634389
(Iteration 2911 / 4900) loss: 1.444833
(Iteration 2921 / 4900) loss: 1.602934
(Iteration 2931 / 4900) loss: 1.681439
(Epoch 6 / 10) train acc: 0.448000; val_acc: 0.454000
(Iteration 2941 / 4900) loss: 1.454648
(Iteration 2951 / 4900) loss: 1.544542
(Iteration 2961 / 4900) loss: 1.692165
(Iteration 2971 / 4900) loss: 1.543428
(Iteration 2981 / 4900) loss: 1.668374
(Iteration 2991 / 4900) loss: 1.529515
(Iteration 3001 / 4900) loss: 1.632125

(Iteration 3011 / 4900) loss: 1.505268
(Iteration 3021 / 4900) loss: 1.478447
(Iteration 3031 / 4900) loss: 1.537118
(Iteration 3041 / 4900) loss: 1.429109
(Iteration 3051 / 4900) loss: 1.471354
(Iteration 3061 / 4900) loss: 1.657847
(Iteration 3071 / 4900) loss: 1.410303
(Iteration 3081 / 4900) loss: 1.573949
(Iteration 3091 / 4900) loss: 1.575670
(Iteration 3101 / 4900) loss: 1.638497
(Iteration 3111 / 4900) loss: 1.649795
(Iteration 3121 / 4900) loss: 1.615946
(Iteration 3131 / 4900) loss: 1.723065
(Iteration 3141 / 4900) loss: 1.573093
(Iteration 3151 / 4900) loss: 1.590917
(Iteration 3161 / 4900) loss: 1.756164
(Iteration 3171 / 4900) loss: 1.515779
(Iteration 3181 / 4900) loss: 1.695209
(Iteration 3191 / 4900) loss: 1.640636
(Iteration 3201 / 4900) loss: 1.564176
(Iteration 3211 / 4900) loss: 1.490512
(Iteration 3221 / 4900) loss: 1.383455
(Iteration 3231 / 4900) loss: 1.502551
(Iteration 3241 / 4900) loss: 1.490310
(Iteration 3251 / 4900) loss: 1.628431
(Iteration 3261 / 4900) loss: 1.518840
(Iteration 3271 / 4900) loss: 1.809047
(Iteration 3281 / 4900) loss: 1.617555
(Iteration 3291 / 4900) loss: 1.479001
(Iteration 3301 / 4900) loss: 1.546494
(Iteration 3311 / 4900) loss: 1.525875
(Iteration 3321 / 4900) loss: 1.505993
(Iteration 3331 / 4900) loss: 1.523785
(Iteration 3341 / 4900) loss: 1.478531
(Iteration 3351 / 4900) loss: 1.485320
(Iteration 3361 / 4900) loss: 1.494951
(Iteration 3371 / 4900) loss: 1.495819
(Iteration 3381 / 4900) loss: 1.560849
(Iteration 3391 / 4900) loss: 1.573796
(Iteration 3401 / 4900) loss: 1.516782
(Iteration 3411 / 4900) loss: 1.594974
(Iteration 3421 / 4900) loss: 1.515978
(Epoch 7 / 10) train acc: 0.472000; val_acc: 0.460000
(Iteration 3431 / 4900) loss: 1.651273
(Iteration 3441 / 4900) loss: 1.470304
(Iteration 3451 / 4900) loss: 1.464455
(Iteration 3461 / 4900) loss: 1.470630
(Iteration 3471 / 4900) loss: 1.536187

(Iteration 3481 / 4900) loss: 1.450538
(Iteration 3491 / 4900) loss: 1.666869
(Iteration 3501 / 4900) loss: 1.527344
(Iteration 3511 / 4900) loss: 1.596538
(Iteration 3521 / 4900) loss: 1.535595
(Iteration 3531 / 4900) loss: 1.420611
(Iteration 3541 / 4900) loss: 1.587542
(Iteration 3551 / 4900) loss: 1.522389
(Iteration 3561 / 4900) loss: 1.466910
(Iteration 3571 / 4900) loss: 1.501862
(Iteration 3581 / 4900) loss: 1.433496
(Iteration 3591 / 4900) loss: 1.571482
(Iteration 3601 / 4900) loss: 1.593070
(Iteration 3611 / 4900) loss: 1.461990
(Iteration 3621 / 4900) loss: 1.431752
(Iteration 3631 / 4900) loss: 1.432136
(Iteration 3641 / 4900) loss: 1.609786
(Iteration 3651 / 4900) loss: 1.571082
(Iteration 3661 / 4900) loss: 1.426692
(Iteration 3671 / 4900) loss: 1.445888
(Iteration 3681 / 4900) loss: 1.438834
(Iteration 3691 / 4900) loss: 1.569620
(Iteration 3701 / 4900) loss: 1.569926
(Iteration 3711 / 4900) loss: 1.432827
(Iteration 3721 / 4900) loss: 1.589051
(Iteration 3731 / 4900) loss: 1.343916
(Iteration 3741 / 4900) loss: 1.445398
(Iteration 3751 / 4900) loss: 1.550233
(Iteration 3761 / 4900) loss: 1.523298
(Iteration 3771 / 4900) loss: 1.550087
(Iteration 3781 / 4900) loss: 1.436836
(Iteration 3791 / 4900) loss: 1.493062
(Iteration 3801 / 4900) loss: 1.402204
(Iteration 3811 / 4900) loss: 1.522213
(Iteration 3821 / 4900) loss: 1.662094
(Iteration 3831 / 4900) loss: 1.708454
(Iteration 3841 / 4900) loss: 1.591888
(Iteration 3851 / 4900) loss: 1.513676
(Iteration 3861 / 4900) loss: 1.479884
(Iteration 3871 / 4900) loss: 1.758486
(Iteration 3881 / 4900) loss: 1.598733
(Iteration 3891 / 4900) loss: 1.376188
(Iteration 3901 / 4900) loss: 1.344747
(Iteration 3911 / 4900) loss: 1.411316
(Epoch 8 / 10) train acc: 0.472000; val_acc: 0.471000
(Iteration 3921 / 4900) loss: 1.411492
(Iteration 3931 / 4900) loss: 1.508765
(Iteration 3941 / 4900) loss: 1.624093

(Iteration 3951 / 4900) loss: 1.539264
(Iteration 3961 / 4900) loss: 1.549394
(Iteration 3971 / 4900) loss: 1.424416
(Iteration 3981 / 4900) loss: 1.657262
(Iteration 3991 / 4900) loss: 1.520391
(Iteration 4001 / 4900) loss: 1.557669
(Iteration 4011 / 4900) loss: 1.607343
(Iteration 4021 / 4900) loss: 1.628428
(Iteration 4031 / 4900) loss: 1.434492
(Iteration 4041 / 4900) loss: 1.572193
(Iteration 4051 / 4900) loss: 1.574850
(Iteration 4061 / 4900) loss: 1.551645
(Iteration 4071 / 4900) loss: 1.539578
(Iteration 4081 / 4900) loss: 1.384278
(Iteration 4091 / 4900) loss: 1.519940
(Iteration 4101 / 4900) loss: 1.399808
(Iteration 4111 / 4900) loss: 1.416478
(Iteration 4121 / 4900) loss: 1.513898
(Iteration 4131 / 4900) loss: 1.519585
(Iteration 4141 / 4900) loss: 1.854212
(Iteration 4151 / 4900) loss: 1.558884
(Iteration 4161 / 4900) loss: 1.482525
(Iteration 4171 / 4900) loss: 1.439783
(Iteration 4181 / 4900) loss: 1.548880
(Iteration 4191 / 4900) loss: 1.426986
(Iteration 4201 / 4900) loss: 1.569824
(Iteration 4211 / 4900) loss: 1.525191
(Iteration 4221 / 4900) loss: 1.615340
(Iteration 4231 / 4900) loss: 1.456324
(Iteration 4241 / 4900) loss: 1.611235
(Iteration 4251 / 4900) loss: 1.530546
(Iteration 4261 / 4900) loss: 1.521066
(Iteration 4271 / 4900) loss: 1.457349
(Iteration 4281 / 4900) loss: 1.597607
(Iteration 4291 / 4900) loss: 1.487572
(Iteration 4301 / 4900) loss: 1.508705
(Iteration 4311 / 4900) loss: 1.486502
(Iteration 4321 / 4900) loss: 1.598110
(Iteration 4331 / 4900) loss: 1.505792
(Iteration 4341 / 4900) loss: 1.532898
(Iteration 4351 / 4900) loss: 1.537084
(Iteration 4361 / 4900) loss: 1.444370
(Iteration 4371 / 4900) loss: 1.461940
(Iteration 4381 / 4900) loss: 1.527139
(Iteration 4391 / 4900) loss: 1.445382
(Iteration 4401 / 4900) loss: 1.441289
(Epoch 9 / 10) train acc: 0.465000; val_acc: 0.483000
(Iteration 4411 / 4900) loss: 1.481944

(Iteration 4421 / 4900) loss: 1.277449
(Iteration 4431 / 4900) loss: 1.471724
(Iteration 4441 / 4900) loss: 1.469580
(Iteration 4451 / 4900) loss: 1.619321
(Iteration 4461 / 4900) loss: 1.518573
(Iteration 4471 / 4900) loss: 1.537314
(Iteration 4481 / 4900) loss: 1.619195
(Iteration 4491 / 4900) loss: 1.491149
(Iteration 4501 / 4900) loss: 1.553490
(Iteration 4511 / 4900) loss: 1.353703
(Iteration 4521 / 4900) loss: 1.468547
(Iteration 4531 / 4900) loss: 1.578264
(Iteration 4541 / 4900) loss: 1.375723
(Iteration 4551 / 4900) loss: 1.599778
(Iteration 4561 / 4900) loss: 1.394066
(Iteration 4571 / 4900) loss: 1.599768
(Iteration 4581 / 4900) loss: 1.499827
(Iteration 4591 / 4900) loss: 1.288340
(Iteration 4601 / 4900) loss: 1.477590
(Iteration 4611 / 4900) loss: 1.349604
(Iteration 4621 / 4900) loss: 1.513949
(Iteration 4631 / 4900) loss: 1.466576
(Iteration 4641 / 4900) loss: 1.572612
(Iteration 4651 / 4900) loss: 1.500355
(Iteration 4661 / 4900) loss: 1.460220
(Iteration 4671 / 4900) loss: 1.365463
(Iteration 4681 / 4900) loss: 1.440637
(Iteration 4691 / 4900) loss: 1.579253
(Iteration 4701 / 4900) loss: 1.367725
(Iteration 4711 / 4900) loss: 1.309591
(Iteration 4721 / 4900) loss: 1.371692
(Iteration 4731 / 4900) loss: 1.418223
(Iteration 4741 / 4900) loss: 1.318427
(Iteration 4751 / 4900) loss: 1.405286
(Iteration 4761 / 4900) loss: 1.497614
(Iteration 4771 / 4900) loss: 1.577086
(Iteration 4781 / 4900) loss: 1.399157
(Iteration 4791 / 4900) loss: 1.487526
(Iteration 4801 / 4900) loss: 1.448662
(Iteration 4811 / 4900) loss: 1.418818
(Iteration 4821 / 4900) loss: 1.413949
(Iteration 4831 / 4900) loss: 1.475848
(Iteration 4841 / 4900) loss: 1.645652
(Iteration 4851 / 4900) loss: 1.394688
(Iteration 4861 / 4900) loss: 1.610560
(Iteration 4871 / 4900) loss: 1.526104
(Iteration 4881 / 4900) loss: 1.383582
(Iteration 4891 / 4900) loss: 1.339698

(Epoch 10 / 10) train acc: 0.471000; val_acc: 0.492000

10 Debug the training

With the default parameters we provided above, you should get a validation accuracy of about 0.36 on the validation set. This isn't very good.

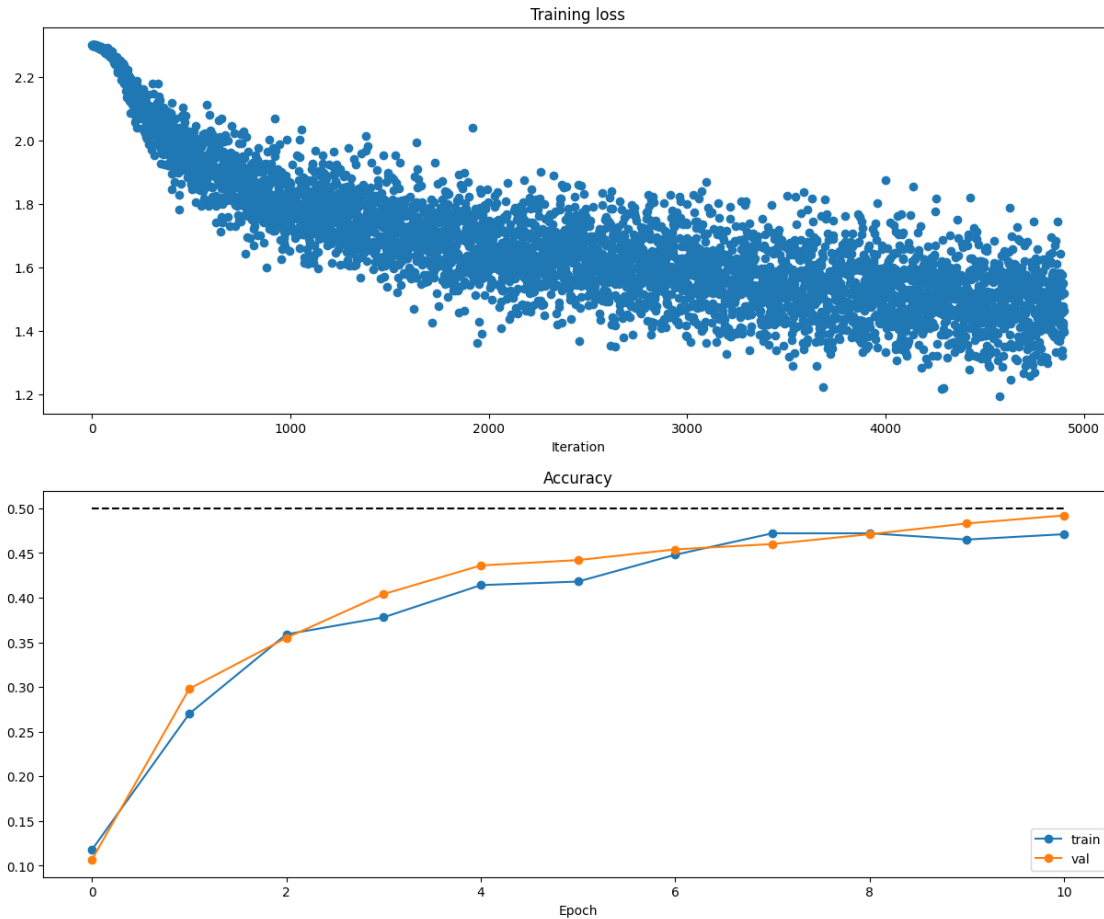
One strategy for getting insight into what's wrong is to plot the loss function and the accuracies on the training and validation sets during optimization.

Another strategy is to visualize the weights that were learned in the first layer of the network. In most neural networks trained on visual data, the first layer weights typically show some visible structure when visualized.

[24]: *# Run this cell to visualize training loss and train / val accuracy*

```
plt.subplot(2, 1, 1)
plt.title('Training loss')
plt.plot(solver.loss_history, 'o')
plt.xlabel('Iteration')

plt.subplot(2, 1, 2)
plt.title('Accuracy')
plt.plot(solver.train_acc_history, '-o', label='train')
plt.plot(solver.val_acc_history, '-o', label='val')
plt.plot([0.5] * len(solver.val_acc_history), 'k--')
plt.xlabel('Epoch')
plt.legend(loc='lower right')
plt.gcf().set_size_inches(15, 12)
plt.show()
```

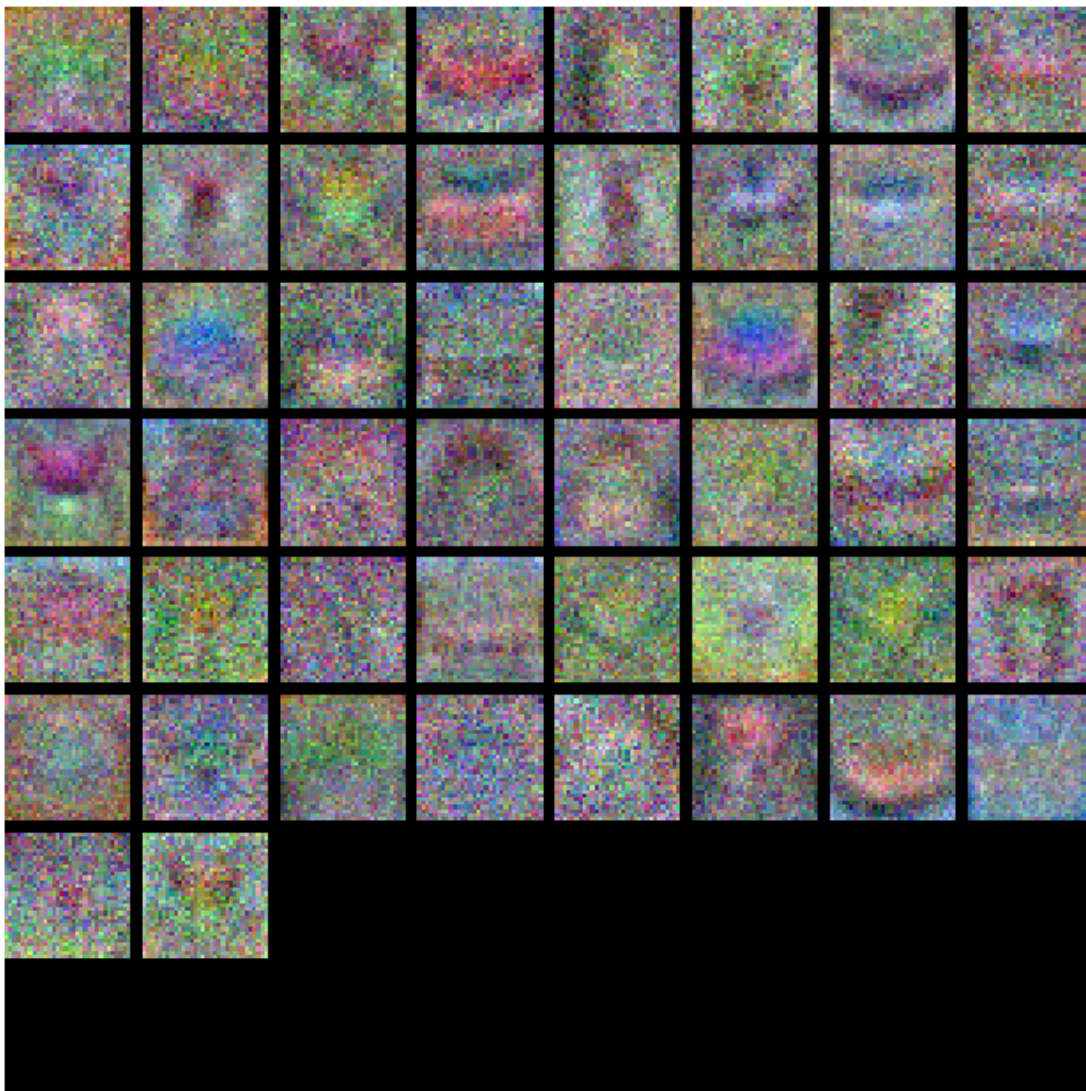


```
[25]: from cs231n.vis_utils import visualize_grid

# Visualize the weights of the network

def show_net_weights(net):
    W1 = net.params['W1']
    W1 = W1.reshape(3, 32, 32, -1).transpose(3, 1, 2, 0)
    plt.imshow(visualize_grid(W1, padding=3).astype('uint8'))
    plt.gca().axis('off')
    plt.show()

show_net_weights(model)
```



11 Tune your hyperparameters

What's wrong?. Looking at the visualizations above, we see that the loss is decreasing more or less linearly, which seems to suggest that the learning rate may be too low. Moreover, there is no gap between the training and validation accuracy, suggesting that the model we used has low capacity, and that we should increase its size. On the other hand, with a very large model we would expect to see more overfitting, which would manifest itself as a very large gap between the training and validation accuracy.

Tuning. Tuning the hyperparameters and developing intuition for how they affect the final performance is a large part of using Neural Networks, so we want you to get a lot of practice. Below, you should experiment with different values of the various hyperparameters, including hidden layer size, learning rate, number of training epochs, and regularization strength. You might also consider

tuning the learning rate decay, but you should be able to get good performance using the default value.

Approximate results. You should be aim to achieve a classification accuracy of greater than 48% on the validation set. Our best network gets over 52% on the validation set.

Experiment: Your goal in this exercise is to get as good of a result on CIFAR-10 as you can (52% could serve as a reference), with a fully-connected Neural Network. Feel free implement your own techniques (e.g. PCA to reduce dimensionality, or adding dropout, or adding features to the solver, etc.).

```
[28]: best_model = None

#####
# TODO: Tune hyperparameters using the validation set. Store your best trained
#
# model in best_model.
#
#
# To help debug your network, it may help to use visualizations similar to the
# ones we used above; these visualizations will have significant qualitative
# differences from the ones we saw above for the poorly tuned network.
#
# Tweaking hyperparameters by hand can be fun, but you might find it useful to
# write code to sweep through possible combinations of hyperparameters
# automatically like we did on the previous exercises.
#
#####
# *****START OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****

import itertools

res = {}
best = -1
learning_rates = np.geomspace(3e-4, 3e-2, 3)
regularization_strengths = np.geomspace(1e-6, 1e-2, 5)

for lr, reg in itertools.product(learning_rates, regularization_strengths):
    model = TwoLayerNet(hidden_dim=128, reg=reg)
```

```

    solver = Solver(model, data, optim_config={'learning_rate': lr},
    num_epochs=10, verbose=False)
    solver.train()

    res[(lr, reg)] = solver.best_val_acc
    if res[(lr, reg)] > best:
        best = res[(lr, reg)]
        best_model = model

for lr, reg in sorted(res):
    val_accuracy = res[(lr, reg)]
    print('lr %e reg %e val accuracy: %f' % (lr, reg, val_accuracy))

print('best validation accuracy achieved during CV: %f' % best)

# *****END OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****
#####
#                                     END OF YOUR CODE                                #
#####

lr 3.000000e-04 reg 1.000000e-06 val accuracy: 0.512000
lr 3.000000e-04 reg 1.000000e-05 val accuracy: 0.523000
lr 3.000000e-04 reg 1.000000e-04 val accuracy: 0.533000
lr 3.000000e-04 reg 1.000000e-03 val accuracy: 0.507000
lr 3.000000e-04 reg 1.000000e-02 val accuracy: 0.521000
lr 3.000000e-03 reg 1.000000e-06 val accuracy: 0.360000
lr 3.000000e-03 reg 1.000000e-05 val accuracy: 0.378000
lr 3.000000e-03 reg 1.000000e-04 val accuracy: 0.333000
lr 3.000000e-03 reg 1.000000e-03 val accuracy: 0.306000
lr 3.000000e-03 reg 1.000000e-02 val accuracy: 0.195000
lr 3.000000e-02 reg 1.000000e-06 val accuracy: 0.140000
lr 3.000000e-02 reg 1.000000e-05 val accuracy: 0.174000
lr 3.000000e-02 reg 1.000000e-04 val accuracy: 0.156000
lr 3.000000e-02 reg 1.000000e-03 val accuracy: 0.156000
lr 3.000000e-02 reg 1.000000e-02 val accuracy: 0.117000
best validation accuracy achieved during cross-validation: 0.533000

```

12 Test your model!

Run your best model on the validation and test sets. You should achieve above 48% accuracy on the validation set and the test set.

```

[29]: y_val_pred = np.argmax(best_model.loss(data['X_val']), axis=1)
      print('Validation set accuracy: ', (y_val_pred == data['y_val']).mean())

```

Validation set accuracy: 0.533

```
[30]: y_test_pred = np.argmax(best_model.loss(data['X_test']), axis=1)
      print('Test set accuracy: ', (y_test_pred == data['y_test']).mean())
```

Test set accuracy: 0.519

12.1 Inline Question 2:

Now that you have trained a Neural Network classifier, you may find that your testing accuracy is much lower than the training accuracy. In what ways can we decrease this gap? Select all that apply.

1. Train on a larger dataset.
2. Add more hidden units.
3. Increase the regularization strength.
4. None of the above.

Your Answer :

We can do 1 and 3.

Your Explanation :

The more data there is to train on, the more generalized and accurate the model gets beyond the training/validation set assuming that the extra training set data is varied and provides more variety than what is already in the training set.

Increasing regularization strength means penalizing large weights more which reduces the effect of outliers. This will be good for learning specific features of a class because the model wants to generalize features.

[]:

features

November 27, 2024

```
[1]: # This mounts your Google Drive to the Colab VM.
from google.colab import drive
drive.mount('/content/drive')

# TODO: Enter the foldername in your Drive where you have saved the unzipped
# assignment folder, e.g. 'cs231n/assignments/assignment1/'
FOLDERNAME = None
assert FOLDERNAME is not None, "[!] Enter the foldername."

# Now that we've mounted your Drive, this ensures that
# the Python interpreter of the Colab VM can load
# python files from within it.
import sys
sys.path.append('/content/drive/My Drive/{}'.format(FOLDERNAME))

# This downloads the CIFAR-10 dataset to your Drive
# if it doesn't already exist.
%cd /content/drive/My\ Drive/$FOLDERNAME/cs231n/datasets/
!bash get_datasets.sh
%cd /content/drive/My\ Drive/$FOLDERNAME
```

```
-----
ModuleNotFoundError                                Traceback (most recent call last)
/Users/kimsh/Downloads/course/fall24/csc480/hw4/assignment1/features.ipynb Cell,
  ↪1 line 2
      <a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall24/csc480.
  ↪hw4/assignment1/features.ipynb#W0sZmlsZQ%3D%3D?line=0'>1</a> # This mounts
  ↪your Google Drive to the Colab VM.
----> <a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall24/csc480.
  ↪hw4/assignment1/features.ipynb#W0sZmlsZQ%3D%3D?line=1'>2</a> from google.colab
  ↪import drive
      <a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall24/csc480.
  ↪hw4/assignment1/features.ipynb#W0sZmlsZQ%3D%3D?line=2'>3</a> drive.mount('/
  ↪content/drive')
      <a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall24/csc480.
  ↪hw4/assignment1/features.ipynb#W0sZmlsZQ%3D%3D?line=4'>5</a> # TODO: Enter th
  ↪foldername in your Drive where you have saved the unzipped
```

```
<a href='vscode-notebook-cell:/Users/kimsh/Downloads/course/fall124/csc480.
hw4/assignment1/features.ipynb#W0sZmlsZQ%3D%3D?line=5'>6</a> # assignment_
folder, e.g. 'cs231n/assignments/assignment1/'
```

ModuleNotFoundError: No module named 'google.colab'

1 Image features exercise

Complete and hand in this completed worksheet (including its outputs and any supporting code outside of the worksheet) with your assignment submission. For more details see the [assignments page](#) on the course website.

We have seen that we can achieve reasonable performance on an image classification task by training a linear classifier on the pixels of the input image. In this exercise we will show that we can improve our classification performance by training linear classifiers not on raw pixels but on features that are computed from the raw pixels.

All of your work for this exercise will be done in this notebook.

```
[2]: import random
import numpy as np
from cs231n.data_utils import load_CIFAR10
import matplotlib.pyplot as plt

%matplotlib inline
plt.rcParams['figure.figsize'] = (10.0, 8.0) # set default size of plots
plt.rcParams['image.interpolation'] = 'nearest'
plt.rcParams['image.cmap'] = 'gray'

# for auto-reloading external modules
# see http://stackoverflow.com/questions/1907993/
# autoreload-of-modules-in-ipython
%load_ext autoreload
%autoreload 2
```

1.1 Load data

Similar to previous exercises, we will load CIFAR-10 data from disk.

```
[3]: from cs231n.features import color_histogram_hsv, hog_feature

def get_CIFAR10_data(num_training=49000, num_validation=1000, num_test=1000):
    # Load the raw CIFAR-10 data
    cifar10_dir = 'cs231n/datasets/cifar-10-batches-py'

    # Cleaning up variables to prevent loading data multiple times (which may
    # cause memory issue)
```

```

try:
    del X_train, y_train
    del X_test, y_test
    print('Clear previously loaded data.')
except:
    pass

X_train, y_train, X_test, y_test = load_CIFAR10(cifar10_dir)

# Subsample the data
mask = list(range(num_training, num_training + num_validation))
X_val = X_train[mask]
y_val = y_train[mask]
mask = list(range(num_training))
X_train = X_train[mask]
y_train = y_train[mask]
mask = list(range(num_test))
X_test = X_test[mask]
y_test = y_test[mask]

return X_train, y_train, X_val, y_val, X_test, y_test

X_train, y_train, X_val, y_val, X_test, y_test = get_CIFAR10_data()

```

1.2 Extract Features

For each image we will compute a Histogram of Oriented Gradients (HOG) as well as a color histogram using the hue channel in HSV color space. We form our final feature vector for each image by concatenating the HOG and color histogram feature vectors.

Roughly speaking, HOG should capture the texture of the image while ignoring color information, and the color histogram represents the color of the input image while ignoring texture. As a result, we expect that using both together ought to work better than using either alone. Verifying this assumption would be a good thing to try for your own interest.

The `hog_feature` and `color_histogram_hsv` functions both operate on a single image and return a feature vector for that image. The `extract_features` function takes a set of images and a list of feature functions and evaluates each feature function on each image, storing the results in a matrix where each column is the concatenation of all feature vectors for a single image.

```

[4]: from cs231n.features import *

num_color_bins = 10 # Number of bins in the color histogram
feature_fns = [hog_feature, lambda img: color_histogram_hsv(img,
    ↪nbin=num_color_bins)]
X_train_feats = extract_features(X_train, feature_fns, verbose=True)
X_val_feats = extract_features(X_val, feature_fns)
X_test_feats = extract_features(X_test, feature_fns)

```

```

# Preprocessing: Subtract the mean feature
mean_feat = np.mean(X_train_feats, axis=0, keepdims=True)
X_train_feats -= mean_feat
X_val_feats -= mean_feat
X_test_feats -= mean_feat

# Preprocessing: Divide by standard deviation. This ensures that each feature
# has roughly the same scale.
std_feat = np.std(X_train_feats, axis=0, keepdims=True)
X_train_feats /= std_feat
X_val_feats /= std_feat
X_test_feats /= std_feat

# Preprocessing: Add a bias dimension
X_train_feats = np.hstack([X_train_feats, np.ones((X_train_feats.shape[0], 1))])
X_val_feats = np.hstack([X_val_feats, np.ones((X_val_feats.shape[0], 1))])
X_test_feats = np.hstack([X_test_feats, np.ones((X_test_feats.shape[0], 1))])

```

```

Done extracting features for 1000 / 49000 images
Done extracting features for 2000 / 49000 images
Done extracting features for 3000 / 49000 images
Done extracting features for 4000 / 49000 images
Done extracting features for 5000 / 49000 images
Done extracting features for 6000 / 49000 images
Done extracting features for 7000 / 49000 images
Done extracting features for 8000 / 49000 images
Done extracting features for 9000 / 49000 images
Done extracting features for 10000 / 49000 images
Done extracting features for 11000 / 49000 images
Done extracting features for 12000 / 49000 images
Done extracting features for 13000 / 49000 images
Done extracting features for 14000 / 49000 images
Done extracting features for 15000 / 49000 images
Done extracting features for 16000 / 49000 images
Done extracting features for 17000 / 49000 images
Done extracting features for 18000 / 49000 images
Done extracting features for 19000 / 49000 images
Done extracting features for 20000 / 49000 images
Done extracting features for 21000 / 49000 images
Done extracting features for 22000 / 49000 images
Done extracting features for 23000 / 49000 images
Done extracting features for 24000 / 49000 images
Done extracting features for 25000 / 49000 images
Done extracting features for 26000 / 49000 images
Done extracting features for 27000 / 49000 images
Done extracting features for 28000 / 49000 images
Done extracting features for 29000 / 49000 images

```

```

Done extracting features for 30000 / 49000 images
Done extracting features for 31000 / 49000 images
Done extracting features for 32000 / 49000 images
Done extracting features for 33000 / 49000 images
Done extracting features for 34000 / 49000 images
Done extracting features for 35000 / 49000 images
Done extracting features for 36000 / 49000 images
Done extracting features for 37000 / 49000 images
Done extracting features for 38000 / 49000 images
Done extracting features for 39000 / 49000 images
Done extracting features for 40000 / 49000 images
Done extracting features for 41000 / 49000 images
Done extracting features for 42000 / 49000 images
Done extracting features for 43000 / 49000 images
Done extracting features for 44000 / 49000 images
Done extracting features for 45000 / 49000 images
Done extracting features for 46000 / 49000 images
Done extracting features for 47000 / 49000 images
Done extracting features for 48000 / 49000 images
Done extracting features for 49000 / 49000 images

```

1.3 Train SVM on features

Using the multiclass SVM code developed earlier in the assignment, train SVMs on top of the features extracted above; this should achieve better results than training SVMs directly on top of raw pixels.

```

[7]: # Use the validation set to tune the learning rate and regularization strength

from cs231n.classifiers.linear_classifier import LinearSVM

learning_rates = [1e-9, 1e-8, 1e-7]
regularization_strengths = [5e4, 5e5, 5e6]

results = {}
best_val = -1
best_svm = None

#####
# TODO:                                     #
# Use the validation set to set the learning rate and regularization strength. #
# This should be identical to the validation that you did for the SVM; save   #
# the best trained classifier in best_svm. You might also want to play        #
# with different numbers of bins in the color histogram. If you are careful  #
# you should be able to get accuracy of near 0.44 on the validation set.      #
#####
# *****START OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****

```

```

import itertools

for lr, reg in itertools.product(learning_rates, regularization_strengths):
    svm = LinearSVM()
    svm.train(X_train_feats, y_train, lr, reg, num_iters=10_000)
    y_train_pred, y_val_pred = svm.predict(X_train_feats), svm.
    ↪predict(X_val_feats)
    results[(lr, reg)] = np.mean(y_train == y_train_pred), np.mean(y_val ==
    ↪y_val_pred)

    if results[(lr, reg)][1] > best_val:
        best_val = results[(lr, reg)][1]
        best_svm = svm

# *****END OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****

# Print out results.
for lr, reg in sorted(results):
    train_accuracy, val_accuracy = results[(lr, reg)]
    print('lr %e reg %e train accuracy: %f val accuracy: %f' % (
        lr, reg, train_accuracy, val_accuracy))

print('best validation accuracy achieved: %f' % best_val)

```

```

lr 1.000000e-09 reg 5.000000e+04 train accuracy: 0.091673 val accuracy: 0.101000
lr 1.000000e-09 reg 5.000000e+05 train accuracy: 0.416122 val accuracy: 0.425000
lr 1.000000e-09 reg 5.000000e+06 train accuracy: 0.417041 val accuracy: 0.417000
lr 1.000000e-08 reg 5.000000e+04 train accuracy: 0.416449 val accuracy: 0.422000
lr 1.000000e-08 reg 5.000000e+05 train accuracy: 0.411490 val accuracy: 0.417000
lr 1.000000e-08 reg 5.000000e+06 train accuracy: 0.401388 val accuracy: 0.398000
lr 1.000000e-07 reg 5.000000e+04 train accuracy: 0.415408 val accuracy: 0.424000
lr 1.000000e-07 reg 5.000000e+05 train accuracy: 0.396122 val accuracy: 0.383000
lr 1.000000e-07 reg 5.000000e+06 train accuracy: 0.334469 val accuracy: 0.347000
best validation accuracy achieved: 0.425000

```

```

[8]: # Evaluate your trained SVM on the test set: you should be able to get at least
    ↪0.40
y_test_pred = best_svm.predict(X_test_feats)
test_accuracy = np.mean(y_test == y_test_pred)
print(test_accuracy)

```

0.416

```

[9]: # An important way to gain intuition about how an algorithm works is to
# visualize the mistakes that it makes. In this visualization, we show examples
# of images that are misclassified by our current system. The first column
# shows images that our system labeled as "plane" but whose true label is

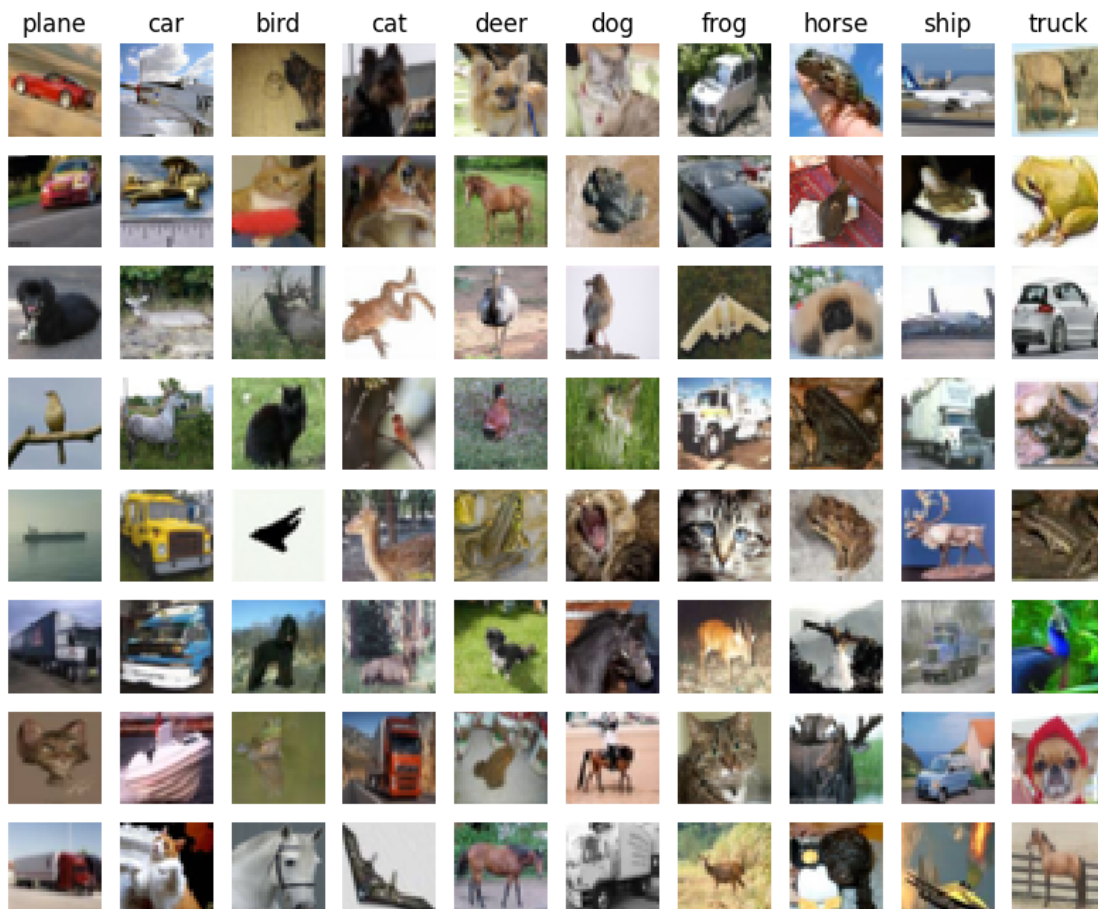
```

```

# something other than "plane".

examples_per_class = 8
classes = ['plane', 'car', 'bird', 'cat', 'deer', 'dog', 'frog', 'horse', 'ship', 'truck']
for cls, cls_name in enumerate(classes):
    idxs = np.where((y_test != cls) & (y_test_pred == cls))[0]
    idxs = np.random.choice(idxs, examples_per_class, replace=False)
    for i, idx in enumerate(idxs):
        plt.subplot(examples_per_class, len(classes), i * len(classes) + cls + 1)
        plt.imshow(X_test[idx].astype('uint8'))
        plt.axis('off')
        if i == 0:
            plt.title(cls_name)
plt.show()

```



1.3.1 Inline question 1:

Describe the misclassification results that you see. Do they make sense?

Your Answer :

They make sense because you can reasonably (or for the most part) see features of the misclassification in every mis-classified image.

1.4 Neural Network on image features

Earlier in this assignment we saw that training a two-layer neural network on raw pixels achieved better classification performance than linear classifiers on raw pixels. In this notebook we have seen that linear classifiers on image features outperform linear classifiers on raw pixels.

For completeness, we should also try training a neural network on image features. This approach should outperform all previous approaches: you should easily be able to achieve over 55% classification accuracy on the test set; our best model achieves about 60% classification accuracy.

```
[10]: # Preprocessing: Remove the bias dimension
# Make sure to run this cell only ONCE
print(X_train_feats.shape)
X_train_feats = X_train_feats[:, :-1]
X_val_feats = X_val_feats[:, :-1]
X_test_feats = X_test_feats[:, :-1]

print(X_train_feats.shape)
```

```
(49000, 155)
```

```
(49000, 154)
```

```
[17]: from cs231n.classifiers.fc_net import TwoLayerNet
from cs231n.solver import Solver

input_dim = X_train_feats.shape[1]
hidden_dim = 500
num_classes = 10

data = {
    'X_train': X_train_feats,
    'y_train': y_train,
    'X_val': X_val_feats,
    'y_val': y_val,
    'X_test': X_test_feats,
    'y_test': y_test,
}

net = TwoLayerNet(input_dim, hidden_dim, num_classes)
best_net = None
```

```
#####
# TODO: Train a two-layer neural network on image features. You may want to #
# cross-validate various parameters as in previous sections. Store your best #
# model in the best_net variable. #
#####
# *****START OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****

import itertools

res = {}
best = -1
learning_rates = np.geomspace(3e-4, 3e-2, 3)
regularization_strengths = np.geomspace(1e-6, 1e-2, 5)
data = {
    'X_train': X_train_feats,
    'X_val': X_val_feats,
    'X_test': X_test_feats,
    'y_train': y_train,
    'y_val': y_val,
    'y_test': y_test
}

for lr, reg in itertools.product(learning_rates, regularization_strengths):
    model = net
    solver = Solver(model, data, optim_config={'learning_rate': lr},
    num_epochs=15, verbose=False)
    solver.train()

    res[(lr, reg)] = solver.best_val_acc
    if res[(lr, reg)] > best:
        best = res[(lr, reg)]
        best_net = model

for lr, reg in sorted(res):
    val_accuracy = res[(lr, reg)]
    print('lr %e reg %e val accuracy: %f' % (lr, reg, val_accuracy))

print('best validation accuracy achieved during CV: %f' % best)

# *****END OF YOUR CODE (DO NOT DELETE/MODIFY THIS LINE)*****
```

```
lr 3.000000e-04 reg 1.000000e-06 val accuracy: 0.275000
lr 3.000000e-04 reg 1.000000e-05 val accuracy: 0.300000
lr 3.000000e-04 reg 1.000000e-04 val accuracy: 0.320000
lr 3.000000e-04 reg 1.000000e-03 val accuracy: 0.319000
lr 3.000000e-04 reg 1.000000e-02 val accuracy: 0.325000
lr 3.000000e-03 reg 1.000000e-06 val accuracy: 0.416000
lr 3.000000e-03 reg 1.000000e-05 val accuracy: 0.504000
```

```
lr 3.000000e-03 reg 1.000000e-04 val accuracy: 0.517000
lr 3.000000e-03 reg 1.000000e-03 val accuracy: 0.517000
lr 3.000000e-03 reg 1.000000e-02 val accuracy: 0.521000
lr 3.000000e-02 reg 1.000000e-06 val accuracy: 0.593000
lr 3.000000e-02 reg 1.000000e-05 val accuracy: 0.609000
lr 3.000000e-02 reg 1.000000e-04 val accuracy: 0.619000
lr 3.000000e-02 reg 1.000000e-03 val accuracy: 0.618000
lr 3.000000e-02 reg 1.000000e-02 val accuracy: 0.618000
best validation accuracy achieved during CV: 0.619000
```

```
[18]: # Run your best neural net classifier on the test set. You should be able
      # to get more than 55% accuracy.
```

```
y_test_pred = np.argmax(best_net.loss(data['X_test']), axis=1)
test_acc = (y_test_pred == data['y_test']).mean()
print(test_acc)
```

0.58