

CSCI 6962/4140: Homework 4

Assigned Thursday October 24 2024. Due by 11:59pm Thursday November 7 2024.

Create a Jupyter notebook for this assignment, and use Python 3. Write documented, readable and clear code (e.g. use reasonable variable names). Submit this notebook making sure the answers to each question are legible and clearly labeled. You will be graded primarily based on the solutions and answers already present in the notebook, but it must also be runnable to reproduce your results.

In this homework, you will implement a one hidden-layer autoencoder by hand. You will follow a realistic workflow: you will implement a specific architecture, then explore the hyperparameter space to find choices that make this architecture work well, document the performance, and reflect on the results. **NB, this will take time: both to debug your backprop implementation, and to do your fitting and hyperparameter selection.**

1. [60 points]

Write a class `autoencoder` that implements a one hidden-layer autoencoder with loss function given by average binary cross-entropy error. The autoencoder will be trained, roughly, as follows:

```
ae = autoencoder(d, n1)

trainingLosses = np.array([0.0]*numepochs)
validationLosses = np.array([0.0]*numepochs)

for e in range(numepochs):
    minibatches = np.split(np.random.permutation(ntrain), numsplits)
    for indices in minibatches:
        Xbatch = Xtrain[:, indices]
        ae.forward(Xbatch)
        ae.backward()
        ae.adam_step()
    trainingLosses[e] = ae.loss(Xtrain)
    validationLosses[e] = ae.loss(Xvalid)
    print("Epoch [%d/%d]\n-----" % (e+1, numepochs))
    print("Training loss: %.2f" % trainingLosses[e])
    print("Validation loss: %.2f" % validationLosses[e])
```

To enable this training workflow, the class should implement the following functions:

(Constructor) A constructor `__init__(self, d, n1)`, where `d` is the dimensionality of the input vectors and `n1` is the dimensionality of the hidden layer.

The constructor should initialize the weights and biases for the two layers, `W1`, `W2`, `b1`, and `b2`. Initialize the biases to zero, and sample the initial weights using the Glorot initialization scheme¹

$$w_{ij}^{\ell} \sim \text{Unif} \left[-\sqrt{\frac{6}{n_{\ell-1} + n_{\ell}}}, \sqrt{\frac{6}{n_{\ell-1} + n_{\ell}}} \right].$$

Also initialize here the quantities necessary to implement the Adam algorithm described below: the timestep counter, and the storage for the gradient and curvature information for the parameters `W1`, `W2`, `b1`, and `b2`.

(Activations) Activation functions `h1(self, a)` and `h2(self, a)` for the hidden and output layers, respectively. They should accept matrix arguments and respectively return the pointwise evaluation of the ReLu and logistic sigmoid activation functions.

¹See <http://proceedings.mlr.press/v9/glorot10a/glorot10a.pdf> for more details.

(Activation derivatives) The derivatives of the activation functions, `dh1(self, a)` and `dh2(self, a)`. They should accept matrices and return the pointwise evaluation of the derivatives of the corresponding activation functions.

(Loss) The loss function `loss(self, Xbatch)` that computes the average binary cross-entropy reconstruction loss when the autoencoding of `Xbatch` is used to approximate `Xbatch`, given by

$$\ell(\mathbf{Y}, \mathbf{X}) = -\frac{1}{dm} \sum_{i=1}^d \sum_{j=1}^m [x_{ij} \ln(y_{ij}) + (1 - x_{ij}) \ln(1 - y_{ij})],$$

where $\mathbf{Y}$ is the output of the autoencoder given the input $\mathbf{X}$. This function should not change any of the attributes of the instance.

(Loss derivative) The derivative of the loss function given above with respect to the output of the autoencoder, `dloss(self, Ybatch)`. It should assume that `Ybatch` was obtained by passing a minibatch through the autoencoder, and use the internal state of the autoencoder to compute the derivative of the loss above with respect to the output `Ybatch`.

(Forward) An inference/forward-pass function `forward(self, Xbatch)`, where `Xbatch` is a $d \times m$ matrix whose columns comprise a minibatch. The inference function should compute the output of passing this minibatch through the autoencoder. In particular, the returned matrix should have the same size as `Xbatch`.

To facilitate their use in backpropagation, this function should store `Xbatch` and the activations `a1` and `a2` and outputs `o1` and `o2` as instance attributes.

(Backward) A backprop/backward-pass function `backward(self)`. This function should compute and store attributes `dw1`, `dw2`, `db1`, `db2` containing the gradients of the training objective with respect to the parameters of the autoencoder, using the values stored from an earlier call to `forward`.

(Adam Optimizer) A function `adam_step(self, alpha=0.001, rho1=0.9, rho2=0.999, delta=1e-8)` that implements a single step of the Adam optimization algorithm to update the autoencoder's parameters. This function assumes that `forward` then `backward` were called immediately before it to generate the quantities needed in the optimization process.

As discussed in class, the Adam algorithm stores exponentially weighted estimates of the gradients and curvature information for the variables being optimized, rescales them to debias them², and applies an update similar to AdaGrad. At the t -th timestep, Adam updates a given parameter $\mathbf{P}$ (a bias vector or weight matrix) according to the scheme

$$\begin{aligned} \mathbf{nuP} &\leftarrow \rho_1 \mathbf{nuP} + (1 - \rho_1) \mathbf{dP} \\ \mathbf{hP} &\leftarrow \rho_2 \mathbf{hP} + (1 - \rho_2) \mathbf{dP} \odot \mathbf{dP} \\ \mathbf{nuPhat} &\leftarrow (1 - \rho_1^t)^{-1} \mathbf{nuP} \\ \mathbf{hPhat} &\leftarrow (1 - \rho_2^t)^{-1} \mathbf{hP} \\ \mathbf{P} &\leftarrow \mathbf{P} - \alpha \cdot \mathbf{nuPhat} / (\sqrt{\mathbf{hPhat}} + \delta). \end{aligned}$$

The quantities `nuP` and `hP` respectively contain gradient and curvature information that must be stored as instance attributes. The timestep counter used in the optimization algorithm should be incremented inside this function.

2. [30 points] Load the Fashion MNIST dataset from Torchvision and scale and reshape it appropriately for input into your autoencoder class. Modify the training code to save the

²For simplicity, the debiasing step was not discussed in class. It helps convergence but is not necessary. See the original Adam paper for more discussion.

parameters of the best performing model (judged on the validation dataset) over the course of the training.

Train the autoencoder with `n1=100` until you are satisfied with the reconstruction results from the best performing model: you choose the minibatch sizes, number of epochs, and Adam hyperparameters (if you choose to modify these). For this portion, you may find it helpful to collaborate with other students to discuss and choose good hyperparameter choices.

Display the training and validation losses in the same plot, with appropriate labeling of the axes and data series. Place a red dot on both data series to indicate the epoch at which the best performing model occurred.

3. [10 points] Choose ten images randomly from the test set, one from each class, and display them in a two by ten grid, where the images in the first row are the original images and the images in the second row are the corresponding reconstructed images using the best performing model. Your submission pdf should contain this image and the following: the final loss on your test set, and your choices for the hyperparameters.