

COMS 4771: Machine Learning

Homework 4 solutions

Extra credit 1. By the definition of $\alpha = \mathbf{y} - \mathbf{X}\hat{\mathbf{w}}_\lambda$, since $n\lambda\hat{\mathbf{w}}_\lambda = \mathbf{X}^\top(\mathbf{y} - \mathbf{X}\hat{\mathbf{w}}_\lambda)$, we have $n\lambda\hat{\mathbf{w}}_\lambda = \mathbf{X}^\top\alpha$. Thus, $n\lambda\mathbf{X}\hat{\mathbf{w}}_\lambda = \mathbf{X}\mathbf{X}^\top\alpha$, and so $\frac{1}{n\lambda}\mathbf{X}\mathbf{X}^\top\alpha = \mathbf{X}\hat{\mathbf{w}}_\lambda$, and hence $\frac{1}{n\lambda}\mathbf{X}\mathbf{X}^\top\alpha + \alpha = \mathbf{X}\hat{\mathbf{w}}_\lambda + \mathbf{y} - \mathbf{X}\hat{\mathbf{w}}_\lambda = \mathbf{y}$.

Extra credit 2. We can rewrite the equation obtained in extra credit 1 as $(\frac{1}{n\lambda}\mathbf{X}\mathbf{X}^\top + \mathbf{I})\alpha = \mathbf{y}$, and hence $\alpha = (\frac{1}{n\lambda}\mathbf{X}\mathbf{X}^\top + \mathbf{I})^{-1}\mathbf{y}$. Then, since $n\lambda\hat{\mathbf{w}}_\lambda = \mathbf{X}^\top\alpha$, we conclude that

$$\hat{\mathbf{w}}_\lambda = \frac{1}{n\lambda}\mathbf{X}^\top \left(\frac{1}{n\lambda}\mathbf{X}\mathbf{X}^\top + \mathbf{I}_n \right)^{-1} \mathbf{y} = \mathbf{X}^\top (\mathbf{X}\mathbf{X}^\top + n\lambda\mathbf{I}_n)^{-1} \mathbf{y}.$$

In the last equation we use the fact that $(\frac{1}{n\lambda}\mathbf{X}\mathbf{X}^\top + \mathbf{I}_n)^{-1} = n\lambda(\mathbf{X}\mathbf{X}^\top + n\lambda\mathbf{I}_n)^{-1}$.

Part 1. We need to compute

$$\langle \hat{\mathbf{w}}_\lambda, \phi(\mathbf{x}) \rangle = \hat{\mathbf{w}}_\lambda^\top \phi(\mathbf{x}) = (\mathbf{X}^\top (\mathbf{X}\mathbf{X}^\top + n\lambda\mathbf{I}_n)^{-1} \mathbf{y})^\top \phi(\mathbf{x}) = \mathbf{y}^\top (\mathbf{X}\mathbf{X}^\top + n\lambda\mathbf{I}_n)^{-1} \mathbf{X}\phi(\mathbf{x}).$$

Here, we used the facts that $(\mathbf{A}\mathbf{B})^\top = \mathbf{B}^\top \mathbf{A}^\top$, and $(\mathbf{X}\mathbf{X}^\top + n\lambda\mathbf{I}_n)^{-1}$ is symmetric and so it equals its transpose. We now show that we can compute the matrix $(\mathbf{X}\mathbf{X}^\top + n\lambda\mathbf{I}_n)^{-1}$ and the vector $\mathbf{X}\phi(\mathbf{x})$ using the kernel function k , and thus we can compute the prediction $\langle \hat{\mathbf{w}}_\lambda, \phi(\mathbf{x}) \rangle$.

First, we observe that $\mathbf{X}\mathbf{X}^\top$ is an $n \times n$ matrix whose i, j -th is $\langle \phi(\mathbf{x}^{(i)}), \phi(\mathbf{x}^{(j)}) \rangle = k(\mathbf{x}^{(i)}, \mathbf{x}^{(j)})$. Hence, the matrix $\mathbf{X}\mathbf{X}^\top$ can be computed entry by entry using the kernel function. Once $\mathbf{X}\mathbf{X}^\top$ is computed, $(\mathbf{X}^\top (\mathbf{X}\mathbf{X}^\top + n\lambda\mathbf{I}_n)^{-1})$ can be computed. Next, we observe that $\mathbf{X}\phi(\mathbf{x})$ is a column vector of size n whose i -th entry is given by $\langle \phi(\mathbf{x}^{(i)}), \phi(\mathbf{x}) \rangle = k(\mathbf{x}^{(i)}, \mathbf{x})$. Thus $\mathbf{X}\phi(\mathbf{x})$ can be computed entry by entry. Finally, we can compute the prediction by computing the matrix product $\mathbf{y}^\top (\mathbf{X}\mathbf{X}^\top + n\lambda\mathbf{I}_n)^{-1} \mathbf{X}\phi(\mathbf{x})$. The pseudocode is given below.

- 1: Input: Training set S , regularization parameter λ , kernel function k , and test point $\mathbf{x}$.
- 2: Compute the matrix $\mathbf{X}\mathbf{X}^\top$ entry by entry setting the i, j -th entry to be $k(\mathbf{x}^{(i)}, \mathbf{x}^{(j)})$ for $i, j \in [n] \times [n]$.
- 3: Compute the matrix $(\mathbf{X}\mathbf{X}^\top + n\lambda\mathbf{I}_n)^{-1}$ numerically.
- 4: Compute the vector $\mathbf{X}\phi(\mathbf{x})$ entry by entry setting the i -th entry to be $k(\mathbf{x}^{(i)}, \mathbf{x})$, for $i \in [n]$.
- 5: Compute the product $\mathbf{y}^\top (\mathbf{X}\mathbf{X}^\top + n\lambda\mathbf{I}_n)^{-1} \mathbf{X}\phi(\mathbf{x})$ and output it.

Extra credit 3. As described in part 1, the matrix $\mathbf{X}\mathbf{X}^\top$ can be computed entrywise using the kernel function: the i, j -th entry is $k(\mathbf{x}^{(i)}, \mathbf{x}^{(j)})$. Once the matrix $\mathbf{X}\mathbf{X}^\top$ is computed, we can compute its eigendecomposition, $\mathbf{X}\mathbf{X}^\top = \mathbf{U}\mathbf{\Lambda}\mathbf{U}^\top$, where the columns of $\mathbf{U}$ are the eigenvectors of $\mathbf{X}\mathbf{X}^\top$, and $\mathbf{\Lambda}$ is a diagonal matrix with the eigenvalues along the diagonal. Now we can construct

the diagonal matrix $\mathbf{S}_\ell$ by choosing the top ℓ eigenvalues, taking their square roots, and arranging them on the diagonal of $\mathbf{S}_\ell$. The matrix $\mathbf{U}_\ell$ is formed column-wise by choosing the corresponding ℓ eigenvectors.

Part 2. $\mathbf{S}_\ell$ is an $\ell \times \ell$ matrix and $\mathbf{U}_\ell$ is an $n \times \ell$ matrix. As in part 1, we can compute $\mathbf{X}\mathbf{X}^\top$ entrywise using the kernel function: the i, j -th entry is $k(\mathbf{x}^{(i)}, \mathbf{x}^{(j)})$. We already have the matrices $\mathbf{S}_\ell$ and $\mathbf{U}_\ell$, and then the rank ℓ PCA representation of the training set is given by the columns of $\mathbf{S}_\ell^{-1}\mathbf{U}_\ell^\top\mathbf{X}\mathbf{X}^\top$.