

1 Just some Regular Ol' Definitions

1.1 Regularization

1. **Why Regularize?:** Machine learning models tend to overfit if they memorize the training data too closely, failing to generalize well to unseen examples. Regularization techniques tackle this by penalizing complex models, pushing them towards simpler solutions that perform better on unseen data. Norms, like L0, L1, and L2, offer different ways to quantify and penalize model complexity.
2. **L2-norm:** Also known as the Euclidean norm, is a measure of the magnitude of a vector in Euclidean space. It is calculated as:

$$\|\mathbf{x}\|_2 = \sqrt{\sum_{i=1}^N x_i^2} \quad (1)$$

where N is the dimensionality of the vector $\mathbf{x}$.

3. **L1-norm:** a measure of the absolute magnitude of a vector. It is calculated as:

$$\|\mathbf{x}\|_1 = \sum_{i=1}^N |x_i| \quad (2)$$

again, where N is the dimensionality of the vector $\mathbf{x}$.

4. **L0-norm:** also known as the "counting norm", is a measure of the number of non-zero elements in a vector. It is calculated as:

$$\|x\|_0 = \sum_{i=1}^N \mathbb{I}(x_i \neq 0) \quad (3)$$

where $\mathbb{I}(x_i \neq 0)$ is the indicator function that is equal to 1 if $x_i \neq 0$ and 0 otherwise.

Properties:

1. Sparsity: L0 and L1 encourage sparsity (many zero parameters), while L2 prefers smaller but not necessarily zero values.
2. Optimization: L1 and L2 are differentiable, facilitating optimization, while L0 is not.

2 Regularization

Think back to linear regression. We would calculate our objective function that we want to minimize as follows: $J(\boldsymbol{\theta}) = \frac{1}{N} \|\mathbf{y} - X\boldsymbol{\theta}\|_2^2$.

Consider we add L2 regularization to this objective function.

This would give us: $J(\boldsymbol{\theta}) = \|\mathbf{y} - X\boldsymbol{\theta}\|_2^2 + \lambda \|\boldsymbol{\theta}\|_2^2$. Determine the closed-form solution of this objective function.

3 Multivariate Gaussian Distribution

3.1 Covariance Matrix

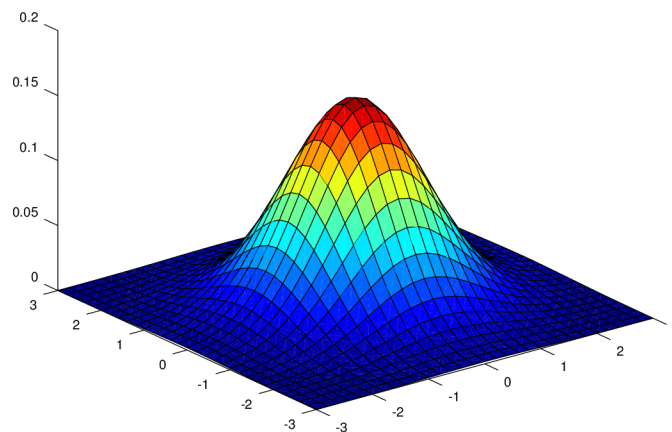
You may be familiar with Gaussian distributions for scalar random variables, i.e., one-dimensional, but we can also have N -dimensional random vector, $X = [X_1, X_2, \dots, X_N]^T$. For this, we use multivariate Gaussian distribution. Notice we cannot use the normal variance, σ^2 , anymore, so we use a covariance matrix. The covariance matrix, Σ , has the dimensions $N \times N$ and $\Sigma_{i,j} = \text{cov}[X_i, X_j] = \mathbb{E}[X_i X_j] - \mathbb{E}[X_i] \mathbb{E}[X_j]$

What do the diagonal elements represent?

What are some special properties of the covariance matrix?

3.2 pdf

A multivariate gaussian distribution would look something like this:



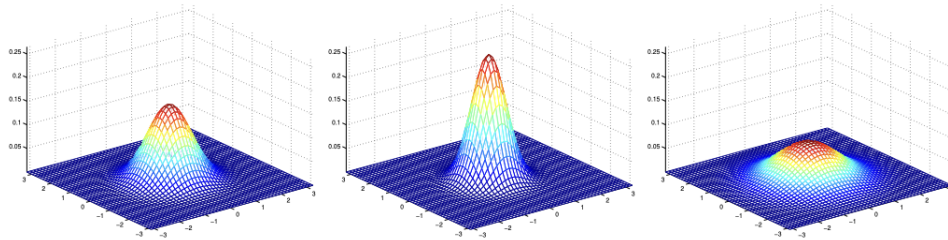
where the probability density function (pdf) is given by:

$$p(\mathbf{x} \mid \boldsymbol{\mu}, \Sigma) = \frac{1}{(2\pi)^{N/2} |\Sigma|^{1/2}} \exp\left(-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu})^T \Sigma^{-1} (\mathbf{x} - \boldsymbol{\mu})\right)$$

$\boldsymbol{\mu} \in \mathbb{R}^N$ is the mean and $\Sigma \in \mathbb{R}^{N \times N}$ is the covariance matrix. $|\Sigma|$ refers to the determinant of Σ .

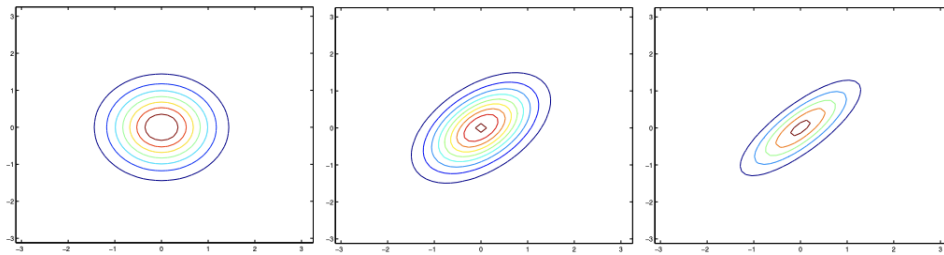
The following μ and Σ values correspond to these multivariate Gaussian examples:

$$\mu = \begin{bmatrix} 0 \\ 0 \end{bmatrix} \Sigma = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}, \quad \mu = \begin{bmatrix} 0 \\ 0 \end{bmatrix} \Sigma = \begin{bmatrix} 0.6 & 0 \\ 0 & 0.6 \end{bmatrix}, \quad \mu = \begin{bmatrix} 0 \\ 0 \end{bmatrix} \Sigma = \begin{bmatrix} 2 & 0 \\ 0 & 2 \end{bmatrix}$$



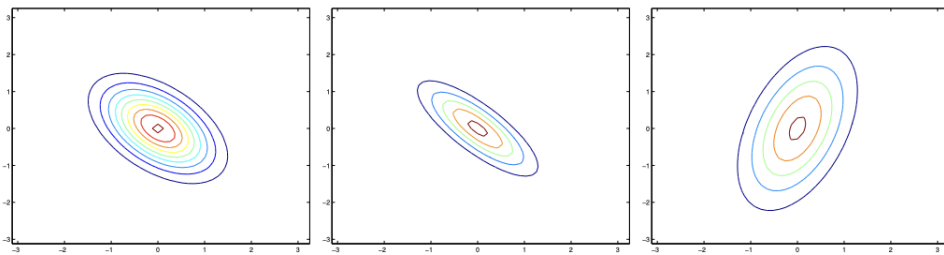
We can look at the contour lines to analyze Σ .

$$\Sigma = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}, \quad \Sigma = \begin{bmatrix} 1 & 0.5 \\ 0.5 & 1 \end{bmatrix}, \quad \Sigma = \begin{bmatrix} 1 & 0.8 \\ 0.8 & 1 \end{bmatrix} :$$



3.2.1 Covariance Matching

Given the gaussian plots below, use this [colab notebook](#) to determine what covariance matrix results in each plot. (Hint: Try starting off with 1's on the diagonal and try varying the other sliders.)



3.3 Multivariate Gaussian: What affects the “spread” of the distribution across the diagonal?

The covariance between variables X_1 and X_2 affects the spread of the distribution across the diagonal. A distribution that is more concentrated along the diagonal would have larger covariance between the two variables. A negative covariance between two variables would reverse the direction of the diagonal.

4 Maximum Likelihood Estimation (MLE)

Maximum likelihood estimation (MLE) is a method of estimating the parameters of a statistical model by maximizing the likelihood function.

$$\hat{\theta}_{MLE} = \underset{\theta}{\operatorname{argmax}} \mathcal{L}(\theta; \mathcal{D})$$

$\mathcal{D}$ is our data sample that we observe, $\mathcal{D} = \{y^{(i)}\}_{i=1}^N$.

$\mathcal{L}(\theta)$ is the likelihood of observing our sampled data points for a specific parameter θ , i.e., it is the joint density of the observed data sample given the parameters, $p(y^{(1)}, y^{(2)}, \dots, y^{(N)} \mid \theta)$. $\mathcal{L}$ is a real-valued function of θ so we want to pick the values for our parameters θ that maximizes the likelihood function. With more data points (larger N), we can arrive at parameters that are closer to their true values.

Oftentimes, we take the log of the $\mathcal{L}$ because maximizing the log function is easier. And, maximizing the log function would also maximize the likelihood function.

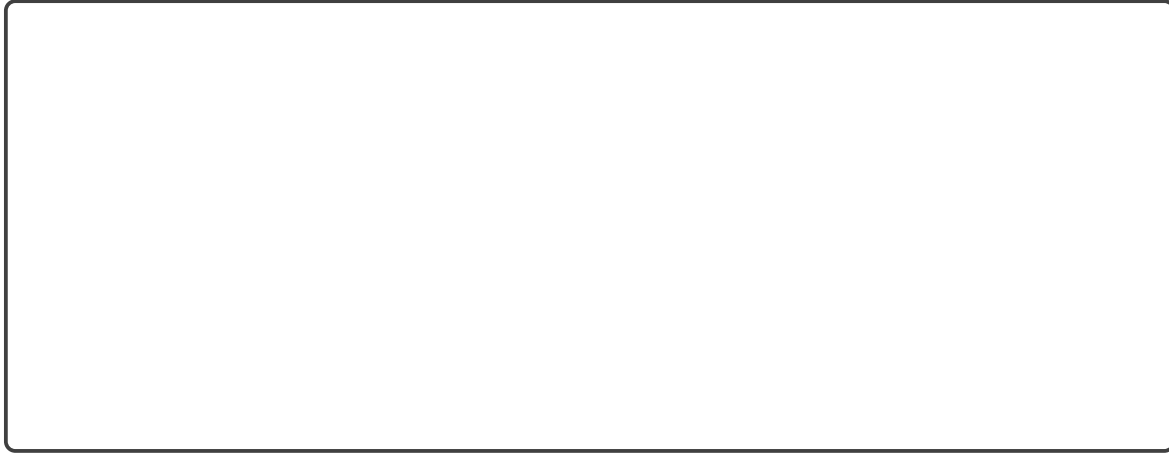
Very Important!!!: We assume that the data is i.i.d. here

4.1 Multivariate Gaussian

Assume you observe $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N$, where each $\mathbf{x}^{(i)} \in \mathbb{R}^M$ is drawn from $\mathcal{N}(\boldsymbol{\mu}, \Sigma)$

Derive the log-likelihood function first.

Now, using this result, derive $\hat{\boldsymbol{\mu}}$. Note that $\frac{\partial}{\partial \mathbf{v}} \mathbf{v}^T A \mathbf{v} = 2A\mathbf{v}$ if A is symmetric



5 Colab Demo [Link](#)