

07-280 Notes

Pre-reading: Feature Engineering & Logistic Regression

Carnegie Mellon University
School of Computer Science

Contents

1	Feature Engineering	1
1.1	Intuition: Weakness of Linear Regression	1
1.2	Idea: Creating New Features in our Dataset	3
1.3	Aside: Another Perspective of Feature Engineering	4
1.4	Generalizing Example to Polynomial Features	4
1.5	Abstract Feature Transforms	5
1.6	Feature Engineering and Overfitting	6
2	Logistic Regression	7
2.1	Logistic Regression: a Linear Model for Classification	7

1 Feature Engineering

If your dataset is complex (i.e. input values have complex relationships with output values), simpler models may be unable to appropriately model/fit to said datasets. One powerful and interesting tool in machine learning is feature engineering, which allows simpler models to fit more complex distributions of data.

1.1 Intuition: Weakness of Linear Regression

Recall our linear regression model on $x \in \mathbb{R}$

$$h(x) = wx + b$$

If we were given a dataset $\mathcal{D} = \{(x^{(i)}, y^{(i)})\}_{i=1}^N$ that followed a linear trend, we could easily fit a linear model, as shown below, left.

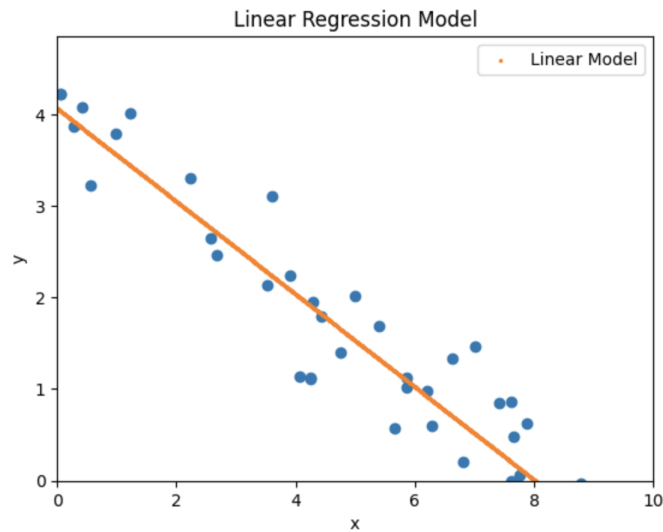


Figure 1: Linear Dataset with Linear Regression Model

The choice of the model for this dataset makes clear sense, and it fits the data well. However, suppose that our data is non-linear:

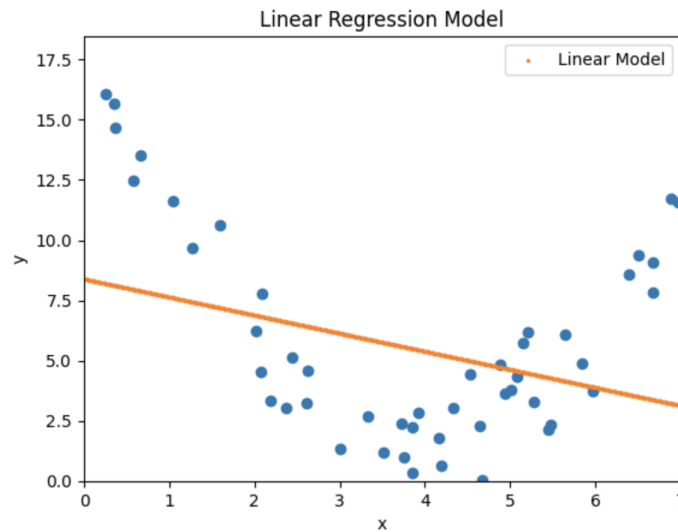


Figure 2: Linear Regression Model Fitted on Nonlinear Dataset

If we attempt to fit a linear regression model to this dataset, we can see that it doesn't necessarily model the dataset distribution well.

This is seemingly a setback for linear regression models. The linearity constraint of linear regression models makes it a model that is easily interpretable, easy to optimize, and simple; however, it would seem that it doesn't make it a good model for more complex (e.g. non-linear) distributions of data. This is where the idea of feature engineering can come into play.

1.2 Idea: Creating New Features in our Dataset

Recall Figure 2. We can see that the dataset strongly resembles a quadratic curve. For example, if we had some kind of model that learned optimal parameters $[b, w_1, w_2]^\top$ for

$$h(x) = b + w_1x + w_2x^2$$

then we could create a better hypothesis for our non-linear dataset $\mathcal{D}$ because it can represent a quadratic data distribution, but how do we go about this?

While the quadratic formulation of $h(x)$ is non-linear with respect to the input data $x \in \mathbb{R}$, we can **note that $h(x)$ is still a linear function with respect to the model's parameters (b, w_1, w_2)** . This means that we can apply linear regression! However, instead of just having one feature $x \in \mathbb{R}$, we have essentially constructed a new dimension/feature in our dataset $x^2 \in \mathbb{R}$ that we are using as a feature to train on, in addition to the original feature, x .

Specifically, when given the non-linear dataset $\mathcal{D}$, we can construct (or engineer) new features to create a dataset $\mathcal{D}'$ that allows us to properly model this more complex distribution.

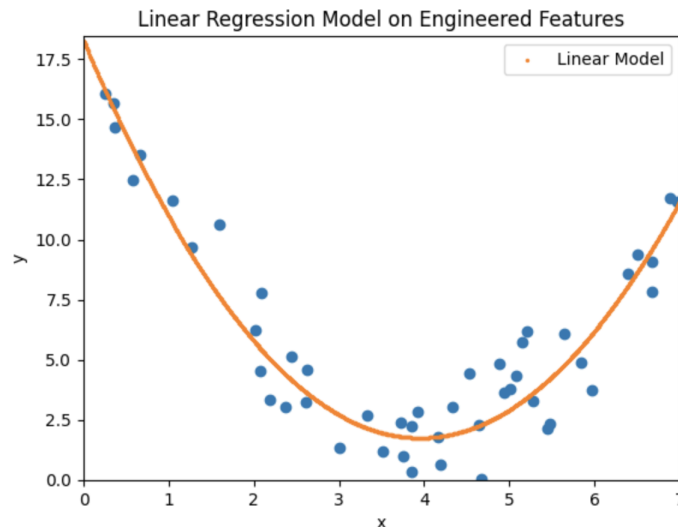
y	x		y	x	x^2
5	6		5	6	36
12	7		12	7	49
1	3		1	3	9
5	2		5	2	4
15	0		15	0	0
$\vdots$	$\vdots$		$\vdots$	$\vdots$	$\vdots$

Table 1: Original Dataset $\mathcal{D}$ (left) and Feature-Engineered Dataset $\mathcal{D}'$ (right)

Now, we can train a linear regression model with this new dataset $\mathcal{D}'$. By adding the new feature, x^2 , by hand into the dataset, we can allow our model to fit to a quadratic function. If we trained a linear regression model using the dataset $\mathcal{D}'$ where one feature is x and another feature is x^2 , then we would receive a model that matches our desired quadratic hypothesis

$$h(x) = b + w_1x + w_2x^2$$

and as a result, it would fit better to the dataset. Below is the result of training a linear regression model on dataset $\mathcal{D}'$:



This is extremely powerful. Note that this curve still represents the predictions of a linear regression model that takes in samples from $\mathcal{D}$, our original dataset, as input. However, when predicting on a sample, $x \in \mathbb{R}$, then we first have to engineer the feature x^2 , which would then allow our model, $h(x) = b + w_1x + w_2x^2$, to predict. Thus, we can see that when we train a model on a feature-engineered dataset, we have to make sure that we follow the same transformations when predicting with this model as well.

Without changing the model, we have allowed it to generalize to a new family of hypothesis functions by simply changing the dataset that we were given, $\mathcal{D}$, to a new dataset $\mathcal{D}'$. This is the primary strength that feature engineering brings to machine learning and is extremely prevalent throughout the entire field.

1.3 Aside: Another Perspective of Feature Engineering

Note that while our original dataset, $\mathcal{D}$, had only 1 feature. The linear regression model that we trained was trained on the 2 features from the engineered dataset $\mathcal{D}'$. As a result, because we are still working with 2D inputs, we can visualize the dataset that our model was trained on.

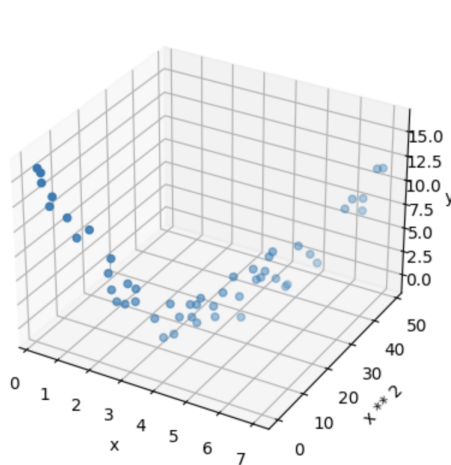


Figure 3: Feature-Engineered $\mathcal{D}'$ Plot

Because our model is trained on 2 features, instead of fitting a line to the dataset, we are now fitting a plane. We can note that this dataset $\mathcal{D}'$ can easily be fitted by a plane in this dimension. That is precisely what our newly trained model h is doing when it was trained on the engineered dataset $\mathcal{D}'$.

1.4 Generalizing Example to Polynomial Features

Note that in order to model a quadratic function, we simply engineered a new feature x^2 that allowed us to develop a model for $h(x) = b + w_1x + w_2x^2$. However, suppose that we now wanted to model a cubic function,

$$h(x) = b + w_1x + w_2x^2 + w_3x^3$$

or even a quartic function (degree 4 polynomial)

$$h(x) = b + w_1x + w_2x^2 + w_3x^3 + w_4x^4$$

Then, all we have to do is simply engineer the features for the higher-powered terms, and we will be able to successfully have our h model those polynomials. In general, if we wanted to model a polynomial of degree M on our original (1-feature) dataset, $\mathcal{D}$, all we have to do is add features to

construct a new dataset, as shown below. While we're at it, we'll insert a column of ones to account for the bias term, b .

y	x
$y^{(1)}$	$x^{(1)}$
$y^{(2)}$	$x^{(2)}$
$y^{(3)}$	$x^{(3)}$
$y^{(4)}$	$x^{(4)}$
$y^{(5)}$	$x^{(5)}$
$\vdots$	$\vdots$

y	$-$	x	x^2	x^3	$\dots$	x^M
$y^{(1)}$	1	$x^{(1)}$	$(x^{(1)})^2$	$(x^{(1)})^3$	$\dots$	$(x^{(1)})^M$
$y^{(2)}$	1	$x^{(2)}$	$(x^{(2)})^2$	$(x^{(2)})^3$	$\dots$	$(x^{(2)})^M$
$y^{(3)}$	1	$x^{(3)}$	$(x^{(3)})^2$	$(x^{(3)})^3$	$\dots$	$(x^{(3)})^M$
$y^{(4)}$	1	$x^{(4)}$	$(x^{(4)})^2$	$(x^{(4)})^3$	$\dots$	$(x^{(4)})^M$
$y^{(5)}$	1	$x^{(5)}$	$(x^{(5)})^2$	$(x^{(5)})^3$	$\dots$	$(x^{(5)})^M$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$

Table 2: Original Dataset $\mathcal{D}$ (left) and Feature-Engineered Dataset $\mathcal{D}'$ (right)

Notation alert: $x^{(i)}$ is the i -th sample, not an exponent!

This generalization of feature engineering to M -th degree polynomials is known as a **polynomial feature transform** where the dataset is now augmented/engineered to include polynomial terms. We will take a look at some models trained on datasets with larger degree polynomial feature transforms; however, before looking at that, we can generalize our understanding of feature engineering even further.

1.5 Abstract Feature Transforms

So far, we have looked only at feature transformations that involve exponentiating some terms inside of our dataset. However, this is just a specific case of feature engineering. Instead of applying an exponent, we could have chosen any other function we desired. We could have applied a period function $\sin$ to our feature x to construct a new feature $\sin(x)$. We could have transformed our original features by including the distance from the origin, $\|\mathbf{x}\|_2$, as a feature as well.

In reality, we could choose any transformation that we would like to apply to our features, as long as we apply the same transformation to each sample in our dataset for training and at time of prediction.

We can represent this notion of feature transformation mathematically, if we are given a dataset $X \in \mathbb{R}^{N \times K}$ where we have N samples and K features per sample, then, we can define a feature transformation function, $\phi : \mathbb{R}^K \rightarrow \mathbb{R}^M$, where M is the number of features that are produced by the transformation ϕ . Now, when we want to train our model on our feature-transformed dataset, we can write the notation as follows, Φ , where:

$$X = \begin{bmatrix} -\mathbf{x}^{(1)\top} - \\ -\mathbf{x}^{(2)\top} - \\ -\mathbf{x}^{(3)\top} - \\ \dots \\ -\mathbf{x}^{(N)\top} - \end{bmatrix} \quad \Phi = \begin{bmatrix} -\phi(\mathbf{x}^{(1)})^\top - \\ -\phi(\mathbf{x}^{(2)})^\top - \\ -\phi(\mathbf{x}^{(3)})^\top - \\ \dots \\ -\phi(\mathbf{x}^{(N)})^\top - \end{bmatrix} \quad (\text{often just renamed } X)$$

Note, that now, the dataset matrix $\Phi \in \mathbb{R}^{N \times M}$ because we've applied the feature transform ϕ to each sample $\mathbf{x}^{(i)}$ in our dataset. We now have a generalized notion of feature transformations as this function ϕ which machine learning engineers (that's you!) can define to be anything.

Example 1: Relating back to our original example, when performing a polynomial feature transformation of degree 3, we were using the feature transform $\phi(\mathbf{x}) = [1, x, x^2, x^3]^\top$.

1.6 Feature Engineering and Overfitting

Before concluding this chapter on an introduction to feature engineering, one might be asking, “why not just have more and more features?”. In the examples above, we’ve shown how adding more features can allow us to create models that are able to fit to more complex data distributions, such as k -dimensional polynomial curves.

However, if adding more features allows our model to generalize to a larger family of functions, why not add as many features as we can? What downside is there to having too many features? Unfortunately, there are some downsides to applying high-dimensional and complex transformations to your dataset. To show this, we will look at what happens to a linear regression model when we apply polynomial feature transforms of degree k for increasing k on a linear dataset, $\mathcal{D}$.

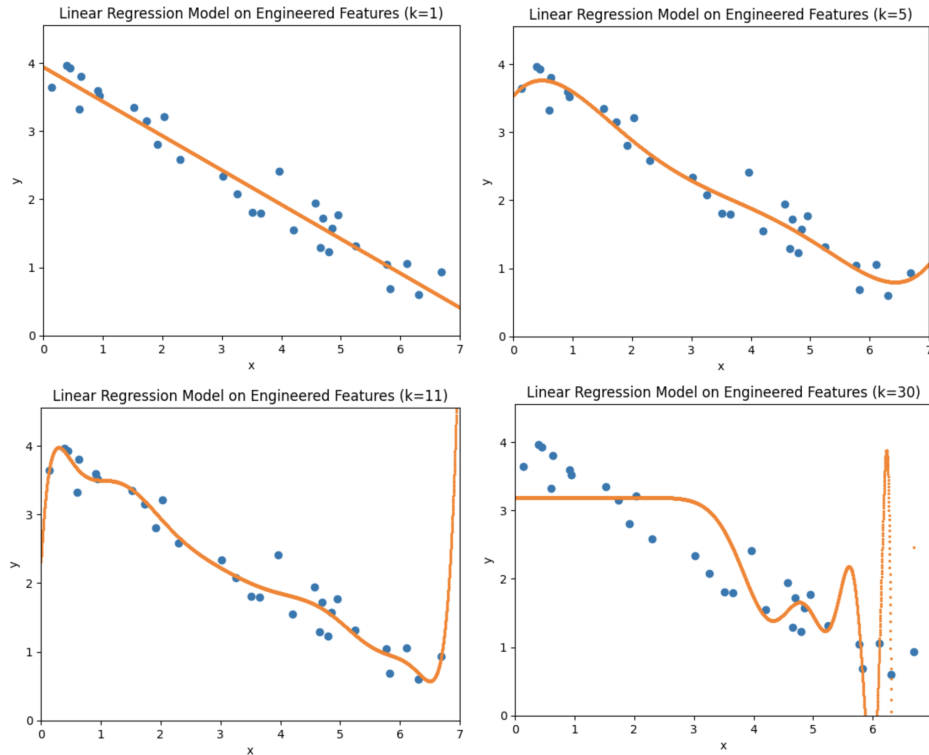


Figure 4: Linear regression with $k = \{1, 5, 11, 30\}$ polynomial feature transforms

As we can see, when we increase the number of features to be very large, the model starts to break down and begins to prioritize matching the exact distribution of the training data as opposed to capturing its general trend. This is a prime example of model overfitting and can commonly occur when the dimensionality of the data is too high, which is the case with complex feature transforms. As a result, we can note that it may not be best to always have very large feature transformations. Instead, we must make reasonable transformations that both improve the model’s ability to generalize well to the data distribution without allowing it to overfit.

There’s certainly more to come in this course on how to combat overfitting despite having very expressive models, such as linear regression with high-degree polynomial features and even neural networks.

2 Logistic Regression

On a different topic, we are going to take a look at classification from a new perspective. Previously, we have viewed classification as making an explicit decision on whether a given sample, $\mathbf{x}^{(i)}$, belongs to one of K classes. However, making a discrete decision on the class of $\mathbf{x}^{(i)}$ doesn't necessarily provide information on how strongly we believe $\mathbf{x}^{(i)}$ belongs to a given class. This idea will raise the notion of developing a *probability distribution* over classes when performing prediction.

Some classification models, such as decision trees, aren't necessarily linear functions with respect to their input, $\mathbf{x}^{(i)}$. Instead, they perform some non-linear function in order to decide their classes. This is where logistic regression models arise. Logistic regression models are an intuitive extension of our well-known linear regression models that allow us to perform classification while building upon a linear model.

2.1 Logistic Regression: a Linear Model for Classification

Now that we have covered this new intuition for performing classification in a probabilistic manner, we can look at a new model that performs this kind of behavior, the logistic regression model. We will inspire an understanding of this model with an example.

Suppose that we wanted to perform classification on a dataset $\mathcal{D} = \{x^{(i)}, y^{(i)}\}_{i=1}^N$ where $x^{(i)} \in \mathbb{R}$ and $y^{(i)} \in \{0, 1\}$ and that we wanted to create a linear classifier (i.e. the decision boundary is linear).

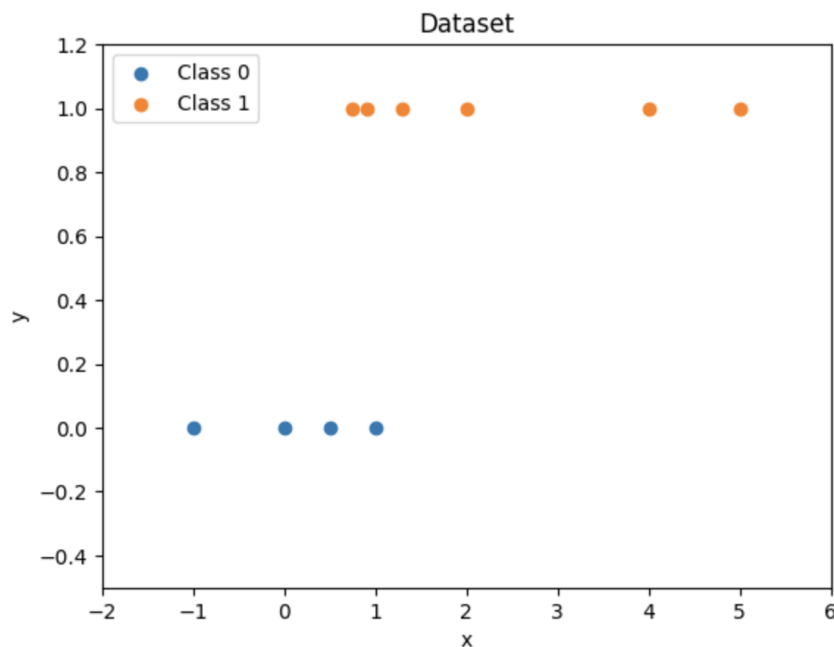


Figure 5: Dataset to Train Classifier On

One idea that we might have to implement this linear classifier is to use a simple linear regression model where $\hat{y} = wx + b$. Training such a model produces the following regression.

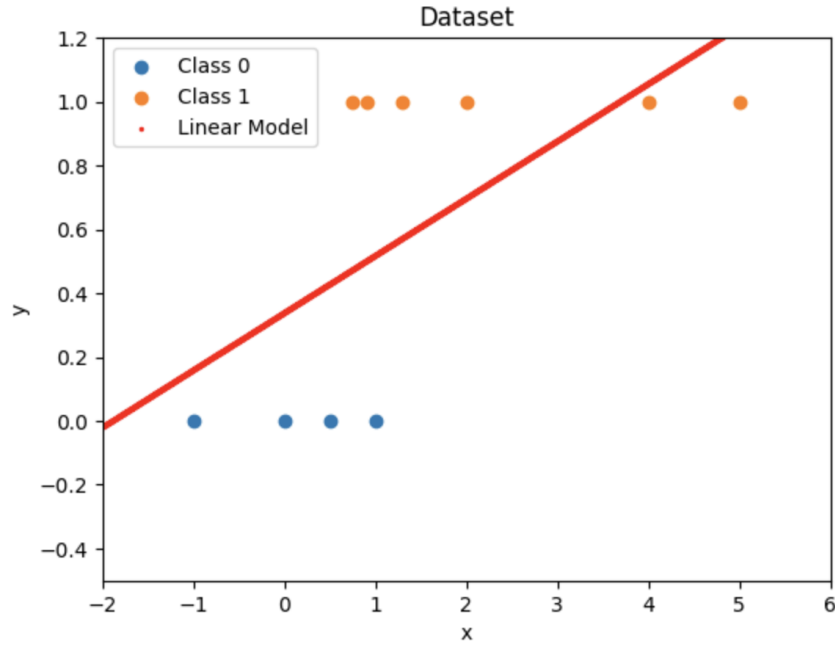


Figure 6: $\mathcal{D}$ trained on Linear Regression Model

We can see that our model learns some relationship between the input x and the class y , where the class will increase if the input increases. However, this doesn't necessarily constrain the output of a model to the appropriate class that we would like, either a 0 or a 1. As a result, we can consider placing some threshold on this linear regression. For example, at a value of 0.5, we can threshold our function to class 0 if the linear regression is less than 0.5 and to class 1 otherwise. This thresholding produces a new model formulation $\hat{y} = g_{threshold}(wx + b)$. Training such a model produces a fit as follows

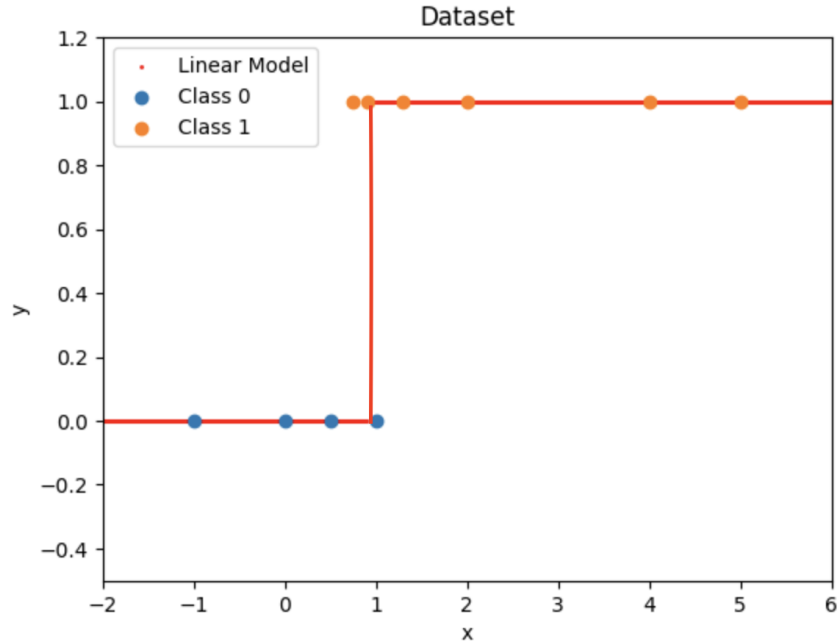


Figure 7: $\mathcal{D}$ trained on thresholding model

While our model now produces only valid values of $\hat{y}$, we still have an issue described earlier in these

notes. The model is extremely certain about its predictions! It transitions rapidly from one class to the other right around input values $x = 1$. However, looking at the overlap in our data, we probably want our model to state that for $x = 1$, $P(Y = 1 | x) \approx 0.5$. In addition to this hard-classification property of the threshold model, it has a derivative of zero essentially everywhere, making it very problematic to optimize over.

As a result, we can try working with a different model $h(x)$. Note, that in order to create a valid classifier, we simply applied the threshold function $g_{threshold}$ onto a linear combination of the features $wx + b$. We can apply the same idea but with a function that smoothly maps all real input values $(-\infty, \infty)$ to the range $(0, 1)$ (i.e., a probability distribution for binary classification). One such common function is a specific sigmoid (s-shaped) function, the **logistic function**, which is often just called the **sigmoid function** and denoted as σ :

$$g_{logistic}(z) = \sigma(z) = \frac{1}{1 + e^{-z}}$$

A plot of this function is in Figure 8. Prove to yourself that this function outputs values in $[0, 1]$ and outputs 0.5 when the input is 0.

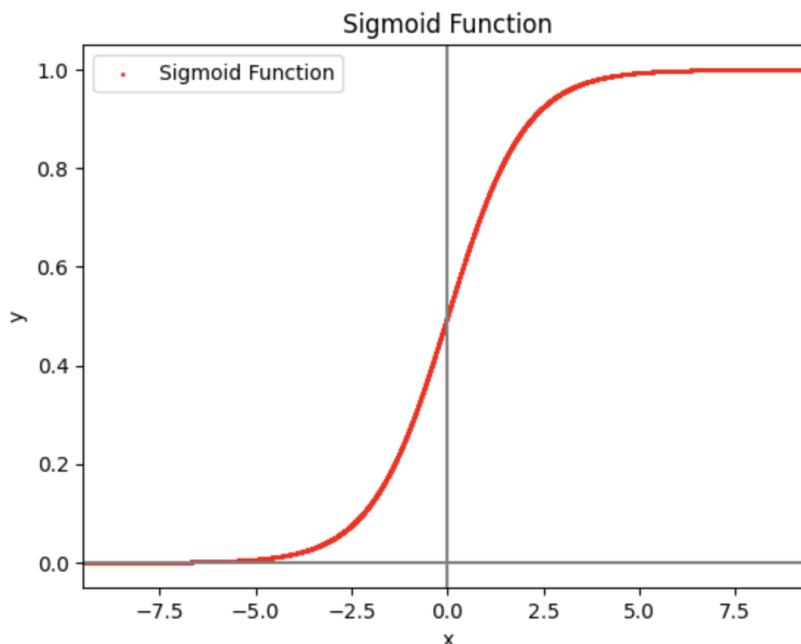


Figure 8: Plot of logistic (sigmoid) function

Now, with the logistic function, $\sigma(z)$, we can apply it to our linear transformation of our input, $wx + b$, to get a new model

$$h(x) = \hat{y} = \sigma(wx + b)$$

This is the logistic regression model for 1-dimensional inputs. If we wanted to generalize this model to samples $\mathbf{x} \in \mathbb{R}^M$, where we have M features (assume the bias term is included in this $\mathbf{x}$), then a more general definition of the logistic regression model would be:

$$h(x) = \hat{y} = \sigma(\theta^\top \mathbf{x})$$

Note that this model is used for performing binary classification. Because the output of the model $h(x)$ is a value in $(0, 1)$, we can re-think this value as a probability where $h(x) = P(Y = 1 | x)$. We can read this as our logistic regression model producing the probability of a sample, $\mathbf{x}$, belonging to class 1.

Finally, after defining a new logistic regression model, $h(x)$, we can optimize over it on our toy dataset $\mathcal{D}$, see Figure 9.

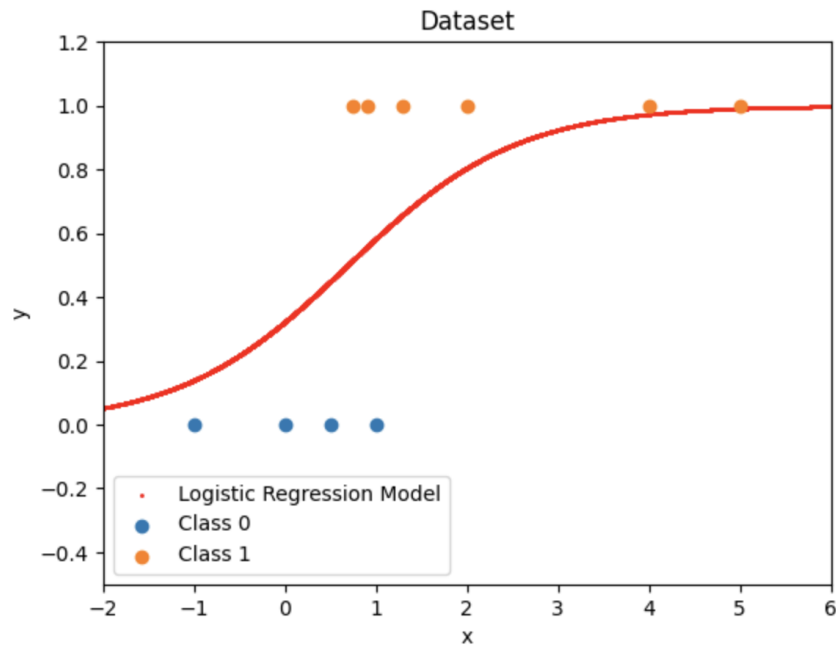


Figure 9: $\mathcal{D}$ trained on logistic regression model

We can now see that our logistic regression model is a prediction function that accounts for uncertainty in our data and indicates a probability for a sample being in class 1. While this example may seem simple, the logistic regression model is an extremely powerful model. It is used as a building block for many larger machine learning models, can be learned efficiently, and is great for simple classification problems.

In lecture, we will investigate this model further by looking at how we can estimate the error of this model on a dataset, how we can use that information to train our model, and how we can apply this model to multi-class classification problems.