

## Table of Contents

|                                                         |    |
|---------------------------------------------------------|----|
| 1. Introduction                                         | 2  |
| 2. Basic Terminologies and Background                   | 2  |
| 2.1 Definitions and Terminologies                       | 2  |
| 2.2 Gaussian Distribution                               | 2  |
| 3. Simple Linear Regression                             | 3  |
| 3.1 MLE Estimate for $\mathbf{w}$                       | 4  |
| 4. Bayesian Linear Regression                           | 6  |
| 4.1 MAP estimate of $\mathbf{w}$                        | 6  |
| 4.1.1 Bulk vs Sequential Learning: Updating Priors      | 8  |
| 4.1.2 Effect of Priors on Bayesian Linear Regression    | 8  |
| 4.2 Posterior Distribution for Prediction               | 10 |
| 4.3 Demonstrating the effect of predictive distribution | 11 |
| 5. Comparing Simple and Bayesian Linear Regression      | 12 |
| 6. Experimental Results                                 | 13 |
| 7. References                                           | 15 |

# 1. Introduction

In machine learning, linear regression is used for finding relationship between a target variable and one or more predictors. In this report, we compare and contrast two such methods of linear regression, namely the simple linear regression model and the Bayesian linear regression model.

Section 2 discusses the basic terminologies and background required for understanding linear regression while sections 3 and 4 take up simple linear regression and Bayesian linear regression respectively. Finally, section 5 summarizes the differences between the aforementioned two models.

## 2. Basic Terminologies and Background

### 2.1 Definitions and Terminologies

The following terminologies are used throughout the rest of the paper.

1. d-dimensional column vectors will be represented with lowercase bold typeface characters like **x**, **w** etc.
2. Matrices will be denoted by uppercase bold typeface characters like **A**, **B** etc.
3. The data is denoted by  $D$  and comprises of  $n$  points,  $[(\mathbf{x}_1, y_1), (\mathbf{x}_2, y_2), \dots, (\mathbf{x}_n, y_n)]$ . Here,  $\mathbf{x}$  represents the predictor vector while  $y$  is the target variable.
4. The predicted value is denoted by  $t$  and the parameter space is denoted by  $\Theta$ .

### 2.2 Gaussian Distribution

In probability theory, the gaussian distribution is a very common continuous probability distribution. A univariate gaussian distribution is defined by 2 parameters: mean ( $\mu$ ) and precision ( $p$ ). Its closed form is given by:

$$p(x) = \frac{\sqrt{p}}{\sqrt{2\pi}} \exp\left(-\frac{p}{2}(x - \mu)^2\right)$$

In shorthand, we write  $p(x)$  with mean  $\mu$  and precision  $p$  as:

$$N(x|\mu, p^{-1})$$

A multivariate gaussian distribution is a generalization of the univariate gaussian distribution to higher dimensions. A gaussian distribution in  $d$

variables is defined by 2 parameters: a mean (**m**) of size d and a precision matrix (**P**) of size  $d \times d$ . For independent dimensions, its closed form is given by:

$$p(x) = \frac{1}{2\pi^{\frac{d}{2}} |P^{-1}|^{\frac{1}{2}}} \exp\left(-\frac{1}{2}(x - m)^T P (x - m)\right)$$

In shorthand, we write  $p(x)$  as:

$$N(x|m, P^{-1})$$

Now, we simplify the above expression.

We know  $p(x) \propto \exp\left(-\frac{1}{2} \underbrace{(x-m)^T P (x-m)}_K\right)$

Simplifying  $K \rightarrow K = (x-m)^T P (x-m)$   
 $= (x^T - m^T) P (x-m)$   
 $= x^T P x - x^T P m - m^T P x + m^T P m$

Since each term in the above expression is a scalar,  
 $x^T P m = (x^T P m)^T = m^T P x$

$$\Rightarrow K = x^T P x - 2x^T P m + m^T P m$$

Quadratic  
in  $x$

Linear  
in  $x$

Constant in  $x$

Substituting this value of  $k$  and removing the constant term from proportionality, we get,

$$p(x) \propto \exp\left(-\frac{1}{2}(x^T P x - 2x^T P m)\right) \quad (2.1)$$

### 3. Simple Linear Regression

In linear regression, we try to predict a target function in  $x$  by modelling it to be linear in weights **w**. By using square error as the loss function and minimizing it, our target function comes out to be:

$$t = w^T x$$

Here,  $\mathbf{w}$  is the parameter vector to our learning model. The model needs to learn  $\mathbf{w}$  for it to predict values of  $y$  given  $\mathbf{x}$ .

For this, we assume that for each given datapoint  $\mathbf{x}_i$ ,  $y_i$  was picked randomly from the distribution:

$$Y = w^T x + \epsilon$$

Where,  $\epsilon$  is the gaussian noise, with mean 0 and precision  $a$ .

Therefore, each  $y_i$  can be modelled as an independent random variable with a gaussian distribution like:

$$p(x_i, w) = N(w^T x, a^{-1})$$

### 3.1 MLE Estimate for $\mathbf{w}$

In simple linear regression, we follow the maximum likelihood approach for estimating  $\mathbf{w}$ , and therefore choose the  $\mathbf{w}$  that maximizes  $p(D|\mathbf{w})$  i.e.

$$w_{MLE} = \operatorname{argmax}_w (p(D|w))$$

Now,

$$p(D|w) = p(y_1, y_2, y_3, \dots, y_n | x_1, x_2, \dots, x_n, w)$$

Since each  $y_i$  is an independent observation,

$$p(D|w) = \prod_{i=1}^n p(y_i | x_1, x_2, \dots, x_n, w)$$

Also,  $y_i$  is conditionally independent of all  $x_j$ ,  $i \neq j$

$$\therefore p(D|w) = \prod_{i=1}^n p(y_i | x_i, w)$$

$$\text{Using eqn 3.1, } p(D|w) = \prod_{i=1}^n \frac{\sqrt{a}}{\sqrt{2\pi}} \exp\left(-\frac{a}{2} (y_i - w^T x_i)^2\right)$$

$$\Rightarrow P(D|\omega) = \frac{a^{n/2}}{(2\pi)^{n/2}} \exp\left(-\frac{a}{2} \sum_{i=1}^n (y_i - \omega^T x_i)^2\right)$$

↳ Since this is a scalar,  
 $\omega^T x_i = x_i^T \omega$

Let  $A = \begin{bmatrix} x_1^T \\ x_2^T \\ \vdots \end{bmatrix}$  &  $y = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \end{bmatrix}$

Simplifying, we get:  $p(D|\omega) = \frac{a^{n/2}}{(2\pi)^{n/2}} \exp\left[-\frac{a}{2} \underbrace{(y - A\omega)^T (y - A\omega)}_{K \text{ --- 3.2}}\right]$

Now,  $K = (y - A\omega)^T (y - A\omega) = y^T y - 2\omega^T A^T y + \omega^T A^T A \omega$

$$\Rightarrow P(D|\omega) \propto \exp\left[-\frac{1}{2} (\omega^T A^T A \omega - 2\omega^T A^T y)\right]$$

Comparing the eqn. with 2.1,  $P = A^T A$  &  $Pm = A^T y$   
 $\Rightarrow m = P^{-1} A^T y$

$$\Rightarrow P(D|\omega) = N(D | (A^T A)^{-1} A^T y, (A^T A)^{-1}) \text{ --- 3.3}$$

Since 3.3 is a gaussian distribution, its maximum occurs at the mean. Therefore, our MLE estimate for  $\omega$  is:

$$\omega = (A^T A)^{-1} A^T y \quad (3.4)$$

Here,  $A$  is called the design matrix. Consequently, our target variable  $t$  is predicted by:

$$t = ((A^T A)^{-1} A^T y)^T x \quad (3.5)$$

## 4. Bayesian Linear Regression

In the Bayesian viewpoint, we formulate linear regression using probability distributions rather than point estimates. Like simple linear regression, the output  $y$  is assumed to be drawn from the probability distribution:

$$y = N(w^T x, \sigma^2) \quad (4.1)$$

The aim of Bayesian Linear Regression is not to find the single “best” value of the model parameters, but rather to determine the posterior distribution for the model parameters. Therefore, not only is the response generated from a probability distribution, but the model parameters are assumed to come from a distribution as well. Here, we assume that  $\mathbf{w}$  comes from the distribution:

$$w = N(0, (\beta I)^{-1}) \quad (4.2)$$

Where,  $I$  is the identity matrix. The above distribution is the *prior* for  $\mathbf{w}$ .

According to Bayes' Theorem,

$$Posterior = \frac{Likelihood * Prior}{Normalization}$$

Since normalization is a constant factor,

$$Posterior \propto Likelihood * Prior \quad (4.3)$$

### 4.1 MAP estimate of $\mathbf{w}$

In Bayesian linear regression, we follow the *maximum a posteriori* approach for estimating  $\mathbf{w}$ , and therefore choose the  $\mathbf{w}$  that maximizes  $p(\mathbf{w} | D)$  i.e.

$$w_{MAP} = \underset{w}{argmax} (p(w|D))$$

According to eqn. 4.3,

$$p(\omega|D) \propto p(D|\omega) \times p(\omega)$$

Substituting the values of  $p(D|\omega)$  from 3.2 &  $p(\omega)$  from 4.2,

$$p(\omega|D) \propto \exp \left[ \underbrace{-\frac{a}{2} (y - A\omega)^T (y - A\omega) - \frac{b}{2} \omega^T \omega}_k \right]$$

$$\begin{aligned} k &\propto a(y - A\omega)^T (y - A\omega) - b\omega^T \omega \\ &= ay^T y - 2a\omega^T A^T y + \omega^T (A^T A + bI) \omega \end{aligned}$$

$$\Rightarrow p(\omega|D) \propto \exp \left[ -\frac{1}{2} \left[ \omega^T (A^T A + bI) \omega - 2a\omega^T A^T y \right] \right]$$

Comparing with 2.1, we get:

$$P_i = aA^T A + bI \quad \Delta \quad m_i = (A^T A + \frac{b}{a} I)^{-1} A^T y$$

$$\Delta \quad p(\omega|D) = N(\omega | m_i, P_i^{-1}) \quad \text{--- 4.4}$$

Since,  $p(\mathbf{w}|D)$  is a gaussian distribution, its maximum lies at its mean.

Therefore, our MAP estimate for  $\mathbf{w}$  is:

$$\mathbf{w} = (A^T A + \frac{b}{a} I)^{-1} A^T y \quad (4.5)$$

Here,  $\frac{b}{a}$  is called the regularization constant. This estimate is equivalent to simple linear regression with regularization and prevents the model from overfitting. However, in simple linear regression, the estimate for the target variable is a single point value given by:

$$t = \mathbf{w}^T \mathbf{x}$$

Whereas, in Bayesian Linear Regression we are interested in finding the predictive posterior distribution of  $t$ ,  $P(t|\mathbf{x}, D)$ .

#### 4.1.1 Bulk vs Sequential Learning: Updating Priors

In Bayesian linear regression too, like simple linear regression, sequential learning and bulk learning yield the same results. This can be shown by the following proof.

We know,

$$\begin{aligned}
 p(\omega|D) &\propto p(D|\omega) \times p(\omega) \\
 &= p(y_2|x_2, \omega) \times \underbrace{p(y_1|x_1, \omega) p(\omega)}_{\text{COMBINE}} \\
 &= p(y_2|x_2, \omega) \times p(\omega|y_1)
 \end{aligned}$$

$\nwarrow$  Likelihood of Data  $\nearrow$  non-informative prior  
 $\downarrow$  Likelihood of next point  $\nearrow$  prior after learning 1 pt.

$$\Rightarrow p(D|\omega) \times p(\omega) = p(y_n|x_n, \omega) \times p(\omega|y_1, \dots, y_{n-1})$$

Bulk learning Sequential learning

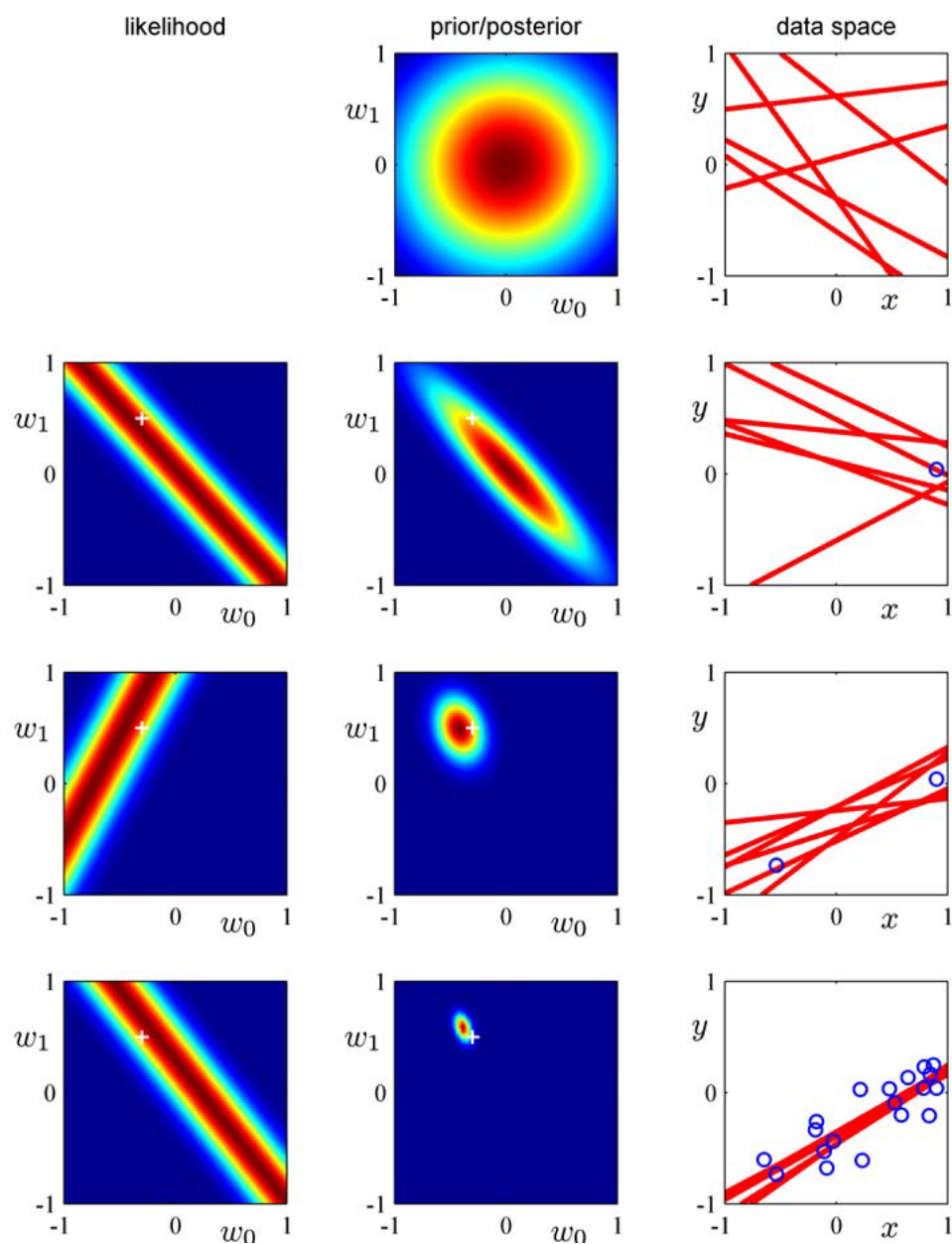
Therefore, we can treat sequential learning Bayesian Linear Regression as learning one point, where the prior for  $\mathbf{w}$  (which is now an informative prior) provides the rest of the information for calculating the posterior distribution of  $\mathbf{w}$ .

#### 4.1.2 Effect of Priors on Bayesian Linear Regression

Figure 1 demonstrates sequential learning in Bayesian Linear Regression for approximating a straight-line function,  $y = -0.3 + 0.5x$ . The first column gives the likelihood distribution of weights given the next observation ( $P(y_i|x_i, \mathbf{w})$ ), and the second column represents the prior distribution of  $\mathbf{w}$  ( $P(\mathbf{w}|y_1, \dots, y_{i-1})$ ). Finally, for plotting the last column,  $w_0$  and  $w_1$  are picked from the posterior distribution of  $\mathbf{w}$  (obtained by multiplying the likelihood and the prior distribution in that row) and the corresponding probable target functions are plotted on the xy-plane. Note that the target functions are probable because  $w_0$  and  $w_1$  are sampled from a gaussian distribution.

We can see, as the bayesian model learns, the prior becomes more and more informative as it confines itself to closer to the actual point with each new data point. Row 4 shows the result of learning with 20 data points, where the posterior distribution has become very compact.

In the limit of infinite points, the posterior distribution would become a delta function centered on the true parameter values. At this point, the prior for  $\mathbf{w}$  is a single value, and its distribution's precision becomes  $\infty$ . This leads the target function for bayesian linear regression to be exactly similar to that of simple linear regression.



**Figure 1:** Illustration of sequential Bayesian learning for a simple linear model of the form  $y(x, w) = w_0 + w_1 x$ .

#### 4.2 Posterior Distribution for Prediction

As stated earlier, in Bayesian linear regression, the target output is not a single point but a distribution. This requires us to calculate the predictive distribution of the target variable ( $P(t|x, D)$ ).

$$\begin{aligned}
 P(t|x, D) &= \int p(t|x, D, \omega) p(\omega|x, D) d\omega \\
 &= \int N(t|\omega^T x, \sigma^2) N(\omega|m_1, P_1^{-1}) d\omega \\
 &\propto \int \exp \left[ \underbrace{-\frac{a}{2}(t - \omega^T x)^2 - \frac{1}{2}(\omega - m_1)^T P_1^{-1}(\omega - m_1)}_{-2k} \right] d\omega
 \end{aligned}$$

$$k = \omega^T \underbrace{(a x x^T + P_1^{-1})}_{P_2} \omega - 2\omega^T \underbrace{(x y a + P_1^{-1} m_1)}_{P_2 m_2} + a t^2$$

$$\Rightarrow k = \omega^T P_2 \omega - 2\omega^T P_2 m_2 + a t^2 + m_2^T P_2 m_2 - m_2^T P_2 m_2$$

$$\Rightarrow k = (\omega - m_2)^T P_2 (\omega - m_2) - m_2^T P_2 m_2 + a t^2$$

↳ Gaussian in  $\omega$

$$\text{Finally, } p(t|x, D) \propto \int N(\omega|m_2, P_2^{-1}) \exp \left[ \underbrace{-\frac{1}{2}(a t^2 - m_2^T P_2 m_2)}_{\text{Const in } \omega} \right] d\omega$$

$$\propto \exp \left[ \underbrace{-\frac{1}{2}(a t^2 - m_2^T P_2 m_2)}_k \right]$$

Simplifying  $k$  as done in previous proofs, we get  $\rightarrow$

$$p(t|x, D) \propto \exp \left[ -\frac{1}{2} \left[ (a - a^2 x^T P_2^{-1} x) t^2 - 2(a x^T P_2^{-1} P_1 m_1) t \right] \right]$$

$\therefore P(t|x, D)$  is gaussian in  $t$ , with

$$P_3 = a(1 - a x^T P_2^{-1} x) \quad \& \quad m_3 = \frac{1}{P_3} a x^T P_2^{-1} P_1 m_1$$

On substituting the values of  $m_1, m_2, P_1, P_2$  &  $P_3$ , we get  $\Rightarrow$

$$m_3 = m_1^T x \quad \& \quad \frac{1}{P_3} = \sigma_3^2 = \frac{1}{a} (1 + x^T (A^T A + \frac{b}{a} I)^{-1} x)$$

$$\& \quad p(t|x, D) = N(t | m_3, \sigma_3^2)$$

And hence, our estimate for  $t$  has a mean at:

$$t = ((A^T A + \frac{b}{a} I)^{-1} A^T y)^T x \quad (4.6)$$

Note that the equation above corresponds to the MAP estimate of  $\mathbf{w}$ . This implies, that the predicted value has the maximum probability to be at our MAP estimate of  $\mathbf{w}$  (which should intuitively be so). Finally, the variance of the distribution is given by:

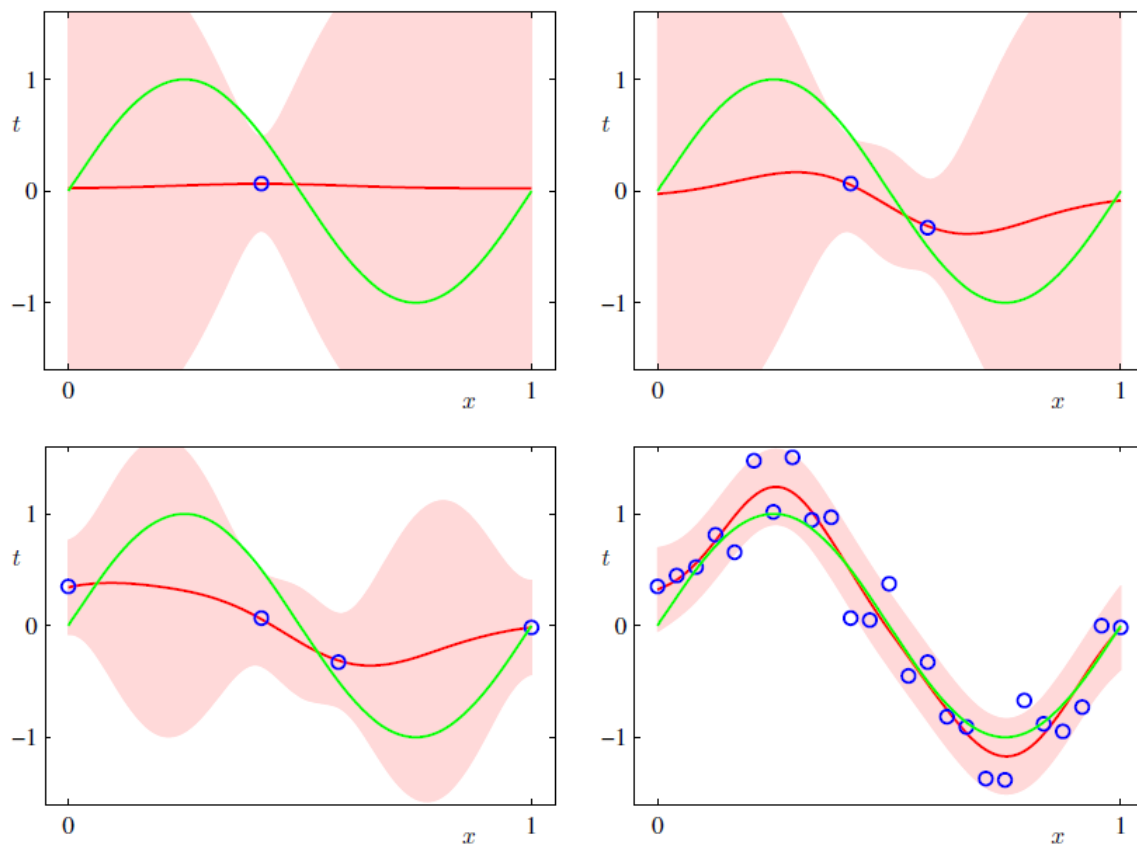
$$var = \frac{1}{a} (1 + x^T (A^T A + \frac{b}{a} I)^{-1} x) \quad (4.7)$$

### 4.3 Demonstrating the effect of predictive distribution

As an Illustration, Figure 2 fits a Bayesian linear regression model to datasets of various sizes and plots the corresponding posterior distributions for each  $x$ . Here, the green curves correspond to the function  $\sin(2\pi x)$  from which the data points were generated. Data sets of size 1, 2, 4 and 25 are shown in the four plots by the blue circles. The red curve denotes the mean of the corresponding gaussian predictive distribution and the pink shaded region

spans 1 standard deviation to either side of the mean for a 9-dimensional model.

Note that the predictive uncertainty depends on  $x$  and is smallest in the neighbourhood of the data points. Also note that the level of uncertainty decreases as more data points are observed. Therefore, for points that lie far away from a cluster of data-points, the uncertainty rises, and Bayesian linear regression allows us to effectively say “I don’t know” for the prediction of  $t$  at such points as the standard deviation is too high.



**Figure 2:** Examples of predictive distribution for a Bayesian linear regression model to approximate  $\sin(2\pi x)$ .

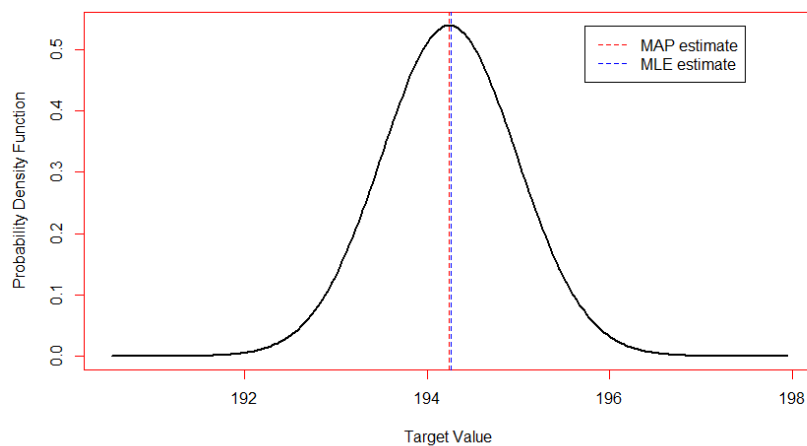
## 5. Comparing Simple and Bayesian Linear Regression

| Simple LR                                                                                        | Bayesian LR                                                                                                 | Reference |
|--------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------|-----------|
| <ul style="list-style-type: none"> <li>Gives a point estimate of the target variable.</li> </ul> | <ul style="list-style-type: none"> <li>Gives a probability distribution for the target variable.</li> </ul> | 3, 4, 4.3 |

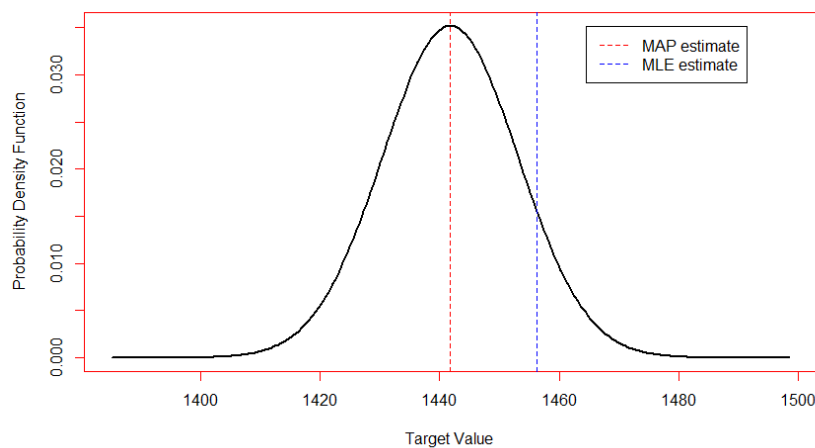
|                                                                                                                                              |                                                                                                                                                                                                                                            |                      |
|----------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------|
| <ul style="list-style-type: none"> <li>• Uses MLE to estimate the values of parameters.</li> </ul>                                           | <ul style="list-style-type: none"> <li>• Uses MAP to estimate the distribution of parameters.</li> </ul>                                                                                                                                   | <b>3.1, 4.1</b>      |
| <ul style="list-style-type: none"> <li>• It is prone to overfitting in complex models.</li> </ul>                                            | <ul style="list-style-type: none"> <li>• Prevents overfitting in complex models by implicitly adding a regularization parameter.</li> </ul>                                                                                                | <b>4.1</b>           |
| <ul style="list-style-type: none"> <li>• Complexity of the model needs to be limited according to number of datapoints available.</li> </ul> | <ul style="list-style-type: none"> <li>• Due to the availability of priors, the complexity of the model can be decided only by the complexity of the problem being solved. An increase in data washes out the effect of priors.</li> </ul> | <b>4.1.1, 4.1.2</b>  |
| <ul style="list-style-type: none"> <li>• It does not report any loss of confidence when extrapolating data points.</li> </ul>                | <ul style="list-style-type: none"> <li>• As data points are extrapolated, loss of confidence is given by standard deviation of the predicted distribution.</li> </ul>                                                                      | <b>3.1, 4.3, 4.4</b> |

## 6. Experimental Results

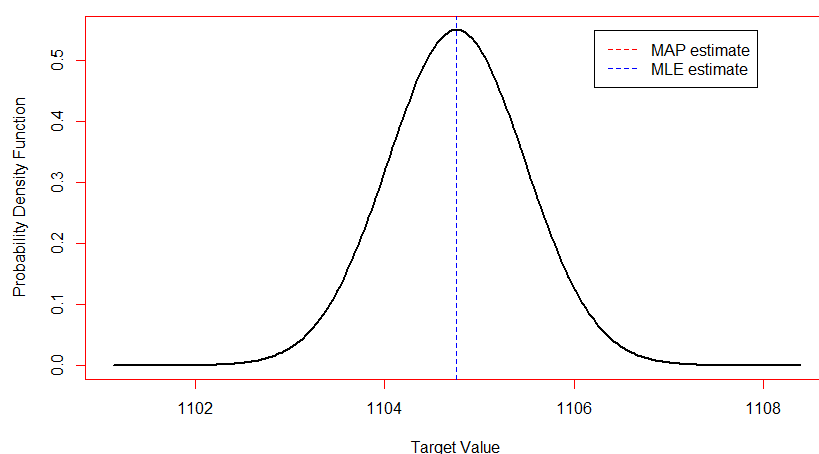
For comparing linear and Bayesian regression, we created two regression models which predict the number of calories burnt given the age, height, weight, duration of exercise, heart rate and the body temperature of a person. As mentioned before, simple linear regression gives a single point estimate which can be interpreted as the MLE estimate given the data while the Bayesian regression model gives the distribution (mean and standard deviation) of the target variable. We can use the mean as the MAP estimate.



**Figure 3:** Feeding points that are near the training data while testing leads to low standard deviation in MAP approach.



**Figure 4:** Extrapolation in MAP and MLE when only 10 data points are presented. High standard deviation in MAP indicates that the estimate is probably not a good one. MLE gives an answer without any indication of standard deviation



**Figure 5:** On giving 15,000 training points, MAP and MLE estimates by the two models perfectly coincide due to the priors being washed out.

For demonstrating the predictive uncertainty of the Bayesian regression model, we input test data points which were close enough to those in the training data. The Bayesian model output was almost same as the simple model output and the standard deviation output by the Bayesian model was very low (close to 0.2). But when we input a data point which was far away from the cluster of data points, the standard deviation of the output given by the Bayesian model was quite high while the simple model just gave a single point estimate through extrapolation. The above result illustrates one of the biggest downsides of using simple linear regression. Since the prediction made by the simple model does not report any loss of confidence when extrapolating data points, the estimate has more chances of being wrong.

Bayesian regression model also performs better when the size of the dataset is small because of the availability of priors. Simple linear regression may overfit the data while Bayesian model incorporates regularization and expresses the estimate as a distribution of possible values. As the number of data points increases, the likelihood washes out the priors, and thus, when data was increased to 15000 points, the output of the Bayesian model converged to that of simple linear regression. This property is illustrated in the below table.

## 7. References

- [1] M. Monk, "mathematicalmonk," Youtube, 2011. [Online]. Available: <https://www.youtube.com/user/mathematicalmonk>.
- [2] W. Koehrsen, "Towards Data Science," Medium, 14 April 2018. [Online]. Available: <https://towardsdatascience.com/introduction-to-bayesian-linear-regression-e66e60791ea7>.
- [3] C. Bishop, "Linear Models for Regression," in *Pattern Recognition and Machine Learning*, Springer, 2006, pp. 137-173.
- [4] C. B. Do, "The Multivariate Gaussian Distribution," *Lecture on Machine Learning*, 2008.
- [5] N. Goel, Introduction to ML, 2018.
- [6] N. Goel, Polynomial Curve Fitting, 2018.

