

Robust Ranking Using Heavy-tailed Prior Distributions

H. He¹, T. Kenney¹ and H. Gu^{1*}

¹ *Department of Mathematics and Statistics, Dalhousie University:
hao.he@dal.ca, tkenney@mathstat.dal.ca, hgu@dal.ca*

**Presenting author*

Keywords. *Ranking; Prior Distributions; Posterior Mean.*

Ranking is the problem of ordering a collection of random variables, based on observation. The aim is to determine which variable has the higher mean. This can be complicated when the standard errors associated with different data points are different.

Ranking of these variables usually depends fundamentally on assumptions of how the true values are distributed — that is, the prior distribution. Posterior mean is a popular criterion for this sort of ranking, and has been applied to a number of real-life problems of this type. For selection purposes, using posterior mean maximises the expected value of the means of selected variables. It is possible to adapt this approach to use a different loss function that could more accurately reflect the relevant criteria for ranking and selection. Indeed there are examples of application of more general loss functions to this problem.

In most research on this topic, little work has gone into the question of robustness to choice of prior. In many cases, a parametric form for the prior is simply chosen (seemingly arbitrarily) and applied. Little work has been done on questioning whether the prior fits the data well, and if it does not, what the consequences are. This is different from a lot of model misspecification cases, because ranking is mostly focussed on the tail. Even if a certain prior fits most of the data very well, if it fits the tail badly there can be serious implications for ranking based on the posterior distribution.

In a simple example of this, we simulated a sample with a heavy-tailed prior, and attempted to rank the samples by posterior mean, using a normal prior. In this simulation, the posterior mean ranking discounted the points with large standard error, even if the observed values were also very large. The reason for this was that the observed values were too extreme under the normal distribution, so the posterior distribution too heavily discounted the observed values.

We study a range of prior distributions — ranging from light-tailed to heavy-tailed, both for simulation of MLE and estimation of posterior mean. We compare the results across a range of measures — Firstly, we look at the overall MSE of the posterior means, to confirm that using the correct distribution gives the best results, and that the difference between using too heavy a tail and too light a tail is relatively small. We then look at the MSE for the top 5% and the top 1%. Here we see that using too heavy a tail causes less harm than using too light a tail. The key point is that if the tail is too heavy, the results will end up close to the MLE, so there is a limit to how bad the results can be. Using a tail that is too light, there is no limit to how bad the results might be. We also examine the average of the means of selected variables, selecting the top 1% and top 5% by posterior mean under each prior. Again, this shows that using too heavy a prior has less potential harm than using too light a prior.

We also examine parameter estimation for the prior distribution. In terms of overall estimation of posterior means, the parameters should be as close to the best fitting values as possible. However, for estimating the posterior means of the largest points, in cases where the true prior distribution is more heavy-tailed than the prior distribution used, we show that it can be advantageous to use parameter estimates that are not the best-fitting ones. For example, using a larger variance than the true prior distribution can reduce the MSE for the points with largest true mean.