

Comment on Article by Rubio and Steel

Xinyi Xu *

Prior elicitation is an important and challenging problem in Bayesian analysis. When little prior knowledge is available for the model parameters (which is commonly the case when the model is high dimensional), a standard approach is to use “noninformative” or “weakly-informative” prior distributions. When the posterior distribution is proper and the sample size is reasonably large, the Bayesian estimates under these priors are usually close to those obtained by frequentist methods, and thus are viewed as “objective” estimates. However, this approach falls apart in some situations.

The authors of this paper show an interesting case where even for quite simple models such as the two-piece location-scale models, the widely used “noninformative” Jeffreys prior leads to improper posteriors and thus prevents valid Bayesian inference for the models. They cleverly propose two alternative classes of priors for the two-piece location-scale models and particularly recommend one of them, which focuses on the Arnold-Groeneveld (AG) measure of skewness. These AG priors have nice interpretations, lead to proper posteriors for all practically interesting subclasses of these models, and can be easily implemented by practitioners in many scientific and industrial fields. This work provides significant methodological and practical contributions to the literature. I’d like to discuss two aspects of prior elicitation that are reflected in this work.

1 The impact of model parametrization on prior elicitation

Although some priors such as the Jeffreys priors are invariant to model parametrization, many common priors are not, so different model parameterizations can lead to different prior choices. In this paper, the authors consider two different parameterizations of the two-piece location-scale models in Sections 2.1 and 2.2. In both model specifications, the model parameters are not directly interpretable, nor can people easily collect information on them. Therefore, improper “noninformative” priors are placed on the model parameters in pursuit of “objective” analysis. This is an all too common practice in Bayesian inference. However, when the models are parameterized with non-interpretable parameters, the “noninformative” priors on convenient model specifications are not necessarily noninformative; instead, they could implicitly contain strong undesirable information on important model features.

In the work of Rubio and Steel, for the Inverse Scale Factors (ISF) model, the Jeffreys

*The Ohio State University xinyi@stat.osu.edu