

Interactive comment on “Inverse modelling of in situ soil water dynamics: accounting for heteroscedastic, autocorrelated, and non-Gaussian distributed residuals” by B. Scharnagl et al.

D. Kavetski

dmitri.kavetski@adelaide.edu.au

Received and published: 17 April 2015

This is an interesting and well-written paper that investigates the improvement of residual error models for environmental model calibration, in particular, the treatment of autocorrelation.

My main question concerns the application of eqn (20) - the main contribution of this work - in a "forecasting" context, where future observations are not available.

C1087

Another clarification question I have is whether the study used an independent "validation" period or are all results shown for the "calibration" period?

The paper refers to Schoups and Vrugt (2009) on page 2174 for details on how to compute prediction limits. However, Schoups and Vrugt used a standard AR(1) error model, which is easy to apply in prediction because it depends only on past observations.

On the other hand, in Scharnagl et al (under review), the new residual error model in eqn (20) defines the innovations "nu" as dependent on both past and future "epsilons", which in turn implies dependence of "nu" on future observations (if my reading of this is right, this should be noted more explicitly in the report!).

This dependence on future observations is not a major problem when applying the residual error model over a period with given observations. For example, when computing the predictive limits over the calibration period, eqn (20) can be "solved" for ϵ_t , and a randomly sampled value of "nu" used to compute the corresponding sample of streamflow at time t. I am assuming this is what was done here? I suggest this is an important detail currently missing in the presentation.

However, what happens if a future observation is not yet available? This is of course of major significance in any practical prediction application. How is eqn (20) used to construct a predictive distribution in this case?

Going further, it would be helpful to see a more detailed description of the data used in this work. For example, did the study use separate calibration and validation periods? Are Figures 2, 4 and 5 showing the predictions computed over the calibration period? I did not see this mentioned anywhere, yet this is obviously a very important point. Clearly the studies' conclusions would be more convincing if the methods were compared over an independent validation period not used for parameter calibration.