

# Discussion: Quantification of uncertainties in the assessment of atmospheric release source with application to the autumn 2017 $^{106}\text{Ru}$ event

Joffrey Dumont Le Brazidec<sup>1,2</sup>, Marc Bocquet<sup>2</sup>, Olivier Saunier<sup>1</sup>, and Yelva Roustan<sup>2</sup>

<sup>1</sup>IRSN, PSE-SANTE, SESUC, BMCA, Fontenay-aux-Roses, France

<sup>2</sup>CEREA, Joint laboratory École des Ponts ParisTech and EDF R&D, Université Paris-Est, Marne-la-Vallée, France

**Correspondence:** Joffrey Dumont Le Brazidec (joffrey.dumont@enpc.fr)

## Report 2

We thank the reviewer for his/her constructive comments, revisions and technical comments, which allowed us to clarify some points of the paper.

The manuscript presents an evaluation of the impact on assumptions surrounding statistical model selection on posterior estimates of source properties of (unknown) radiological releases. This manuscript will be informative to researchers and operational users. I have tried to avoid repetition if the previously posted comment. One major drawback of the manuscript is that much of the motivation is based on a straw-man argument, i.e. the original Gaussian set-up is designed in the manuscript such that it will fail. Below are a number of suggestions for revisions to improve the manuscript, followed by technical comments.

- Paragraph starting line 110: This is a straw-man argument. The assumption from the start is that the error is larger for higher measurements. This may not always be valid, e.g. an incorrectly dispersed wide plume with a very high concentration. It may be that the error is much larger for the smaller measurements than the larger measurements. There are other approaches to improve the validity of Gaussian (or any) likelihoods through transformation of variables. For example, using a non-linear forward model. Caveats, justification and ‘typical errors’ needs explaining, preferably at the start of section 2.

We agree that the fourth criterion results from a "multiplicative" consideration of the error distribution, which is not necessarily true. However, we believe that this simplification is relevant in the case where a limited number of hyperparameters are available to model  $\mathbf{R}$ . We have specified in the text that this criterion relies on this underlying simplification.

- Much of the arguments surround having independent and identically distributed (iid) model-measurement error in the covariance. Many of the arguments throughout can be countered by the use of non-iid covariances, e.g.  $\mathbf{tR}$ , where the diagonal of  $\mathbf{R}$  is the measurement value and  $\mathbf{t}$  is a scaling – equivalent to having e.g. a 10% model-measurement error. This needs further discussions and better justification for the arguments proposed (e.g. non-negativity).

This is true, our argument is indeed based on the hypothesis of modelling  $\mathbf{R}$  by some hyperparameters. However, it should be noted that the variance hyperparameters of the matrix  $\mathbf{R}$  are estimated through a hierarchical Bayesian approach inside the MCMC procedure. In other words, it is not estimated independently of the MCMC in an empirical Bayesian approach. More precisely, if we note  $\mathbf{x}_S = (x_1, x_2, \ln \mathbf{q})$  the vector of variables describing the source except for the hyperparameter  $r$ , we try to estimate :

$$\begin{aligned} p(\mathbf{x}_S, r_1, \dots, r_N | \mathbf{y}) &\propto p(\mathbf{y} | \mathbf{x}_S, r_1, \dots, r_N) p(\mathbf{x}_S, r_1, \dots, r_N) \\ &\propto p(\mathbf{y} | \mathbf{x}_S, r_1, \dots, r_N) p(\mathbf{x}_S | r_1, \dots, r_N) p(r_1, \dots, r_N) \end{aligned} \quad (1)$$

and

$$p(\mathbf{x}_S, r_1, \dots, r_N | \mathbf{y}) = p(\mathbf{x}_S | \mathbf{y}, r_1, \dots, r_N) p(r_1, \dots, r_N | \mathbf{y}) \quad (2)$$

If we model  $\mathbf{R}$  with a term  $t$  of scaling model-measurement, we suppose that the values  $r_1, \dots, r_N$  depend on  $t$ , i.e., we suppose that these values depend on the difference between the observations  $\mathbf{y}$  and the predictions  $\mathbf{y}_S = \mathbf{H}\mathbf{x}_S$ . However, this difference depends on the definition of  $\mathbf{x}_S$ . In other words, with a definition of  $\mathbf{R}$  depending on  $t$  then  $r_1, \dots, r_N$  depend on  $\mathbf{x}_S$ . This would not be rigorously defined in a hierarchical Bayesian formalism.

Furthermore, even if we agree that  $\mathbf{R} = r\mathbf{I}$  is a simplification and that better and more rigorous modelling are possible, this is a very classic choice in the literature.

- Section 3.2.3: It would be useful for many readers to provide a brief conceptual introduction to MCMC methods (i.e. asymptotically exact methods not reliant on closed form solutions or conjugacy).

A few sentences have been added.

- Section 3.4: This section needs expanding considerably. It is a paper with lot of content. At current, the summary provides an overview of the approach of the paper but no summary of the finding. The summary should summarise the results and findings to adequately inform the lazy reader on the paper's content.

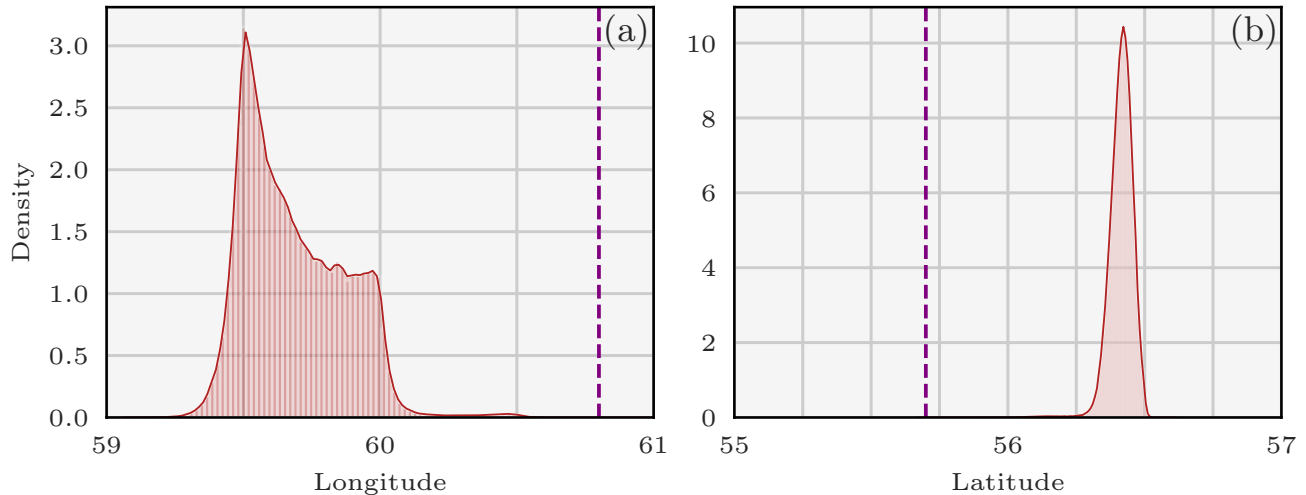
This is true and we have extended the conclusion accordingly. Thank you.

- I understand the following would require a lot of extra work, and so do not mandate it for publication. It would however, much improve the paper. Seeing as an ensemble is used, it would seem sensible to me to use a simulated dataset (a simulation using an ensemble member) to draw conclusions from the various experiment. It is not perfect, but would be useful to have a "truth".

Indeed, we agree that this may be interesting to use a synthetic case. However, we decided in this paper to focus on the development of methods rather than on the maximum exploitation of the results on the  $^{106}\text{Ru}$  case.

- It would be useful in the analysis and plot to also show the original case of a Gaussian likelihood. This is needed to prove the worth of using a non-Gaussian likelihood.

We believe that the relevance of a non-Gaussian likelihood can be proved in theory but not on a particular case (Liu et al., 2017). Inverse modelling with a Gaussian likelihood might provide pertinent results in some cases : the problem is that these results are based on the use of the few largest observations instead of the whole dataset. The information in some of the data is not exploited, which can lead to discrepancies between the reconstructions in some cases. The longitude and latitude distributions are reconstructed with a Gaussian likelihood in figure 1. We can note that the results differ significantly from the ones obtained using log-normal, log-Laplace, or log-Cauchy likelihoods.



**Figure 1.** Density of the coordinates describing the  $^{106}\text{Ru}$  source for a Gaussian likelihood of the Ruthenium source sampled using the parallel tempering method with the help of the observation sorting algorithm and the HRES meteorology..

Technical comments:

- Title: “modelling uncertainties” is ambiguous as it can refer to a statistical model or a transport model. The title also is not grammatically correct. A suggested improvement is “Quantification of uncertainties in atmospheric release source assessment applied to the autumn 2017 Ruthenium 106 source”.

We have changed the title following the suggestion of the reviewer. Thank you.

- Abstract, line 5: “improve on these distributions” is vague. ‘Better quantify’ or ‘improve estimates of these distributions’ would be better.

We have corrected the sentence following this suggestion.

- Line 7: A space is not needed before a colon in English.

This has been corrected. Thank you.

- Line 8: Clarify ‘model errors’ (I assume transport errors?)

We have added « physical » to clarify.

- Line 10: ‘several suited distributions for the errors are advised’ the passive voice reads as though you are advising the distributions. Better “we suggest several suitable distributions for the errors” or “are suggested” if sticking with the passive.

Active voice is now used. Thank you for this advice.

- Line 17: sources or a source; remove ‘many’; ‘Therefore’ doesn’t follow – delete.

This has been corrected, Thank you for these corrections.

- Line 31: ‘finds its origin in’ to ‘originates from’

This has been corrected. Thank you.

- Line 38: This is an oversimplification of the weighting strategy. See, for example, importance resampling or Ensemble Kalman filter methods.

We use a linear weighting strategy in order to be able to include the weights as variables of the ensemble to be sampled.

- Line 46: Why index vector  $x$  elements with  $x_1$ ,  $x_2$  and then  $\ln(q)$ ? Perhaps  $x_3 = \ln(q)$  would be clearer.

Classic coordinates names are  $(x, y)$  but  $y$  name is already used for describing the observations vector.  $\ln q$  being a vector, this might be confusing to use  $x_3$ .

- Line 47:  $R$  is better described as the covariance matrix containing the model-measurement errors.

This has been corrected.

- Line 52: Introduce for the reader what the prior is (i.e. the probability distribution of prior knowledge before considering data).

This has been added.

- Line 57: reconstructed posterior distribution

This has been corrected. Thank you

- Line 60: ‘transformation’ is better than ‘parameterisation’

This has been corrected.

- Line 61: ‘are the results’ would be better as ‘are the results of a simulation’

This has been corrected.

- Line 62: ‘and are therefore depending’ to ‘and depend on’  
This has been corrected. Thank you a lot.
- Line 70: This isn’t an expansion.  
We mean that the computation time is high.
- Line 103: A cost function is a non-probabilistic concept and so better to refer to as simply the negative log-likelihood.  
The authors do agree but the term of "cost" or "loss" is actually widespread in the statistic, data assimilation, or inverse problem literature. The use of the term "cost" rather than "negative log-likelihood" allows us to save a lot of space considering the number of occurrences of the term. Furthermore, the cost can correspond to the negative of the log-posterior distribution.
- Equation 4: There should be no divide by 2 in the first term.  
This comment has been removed by the reviewer.
- Line 112: A space not full stop is needed between units  
These have been corrected. Thank you.
- Line 115: Capital G on Gaussian.  
This has been corrected.
- Line 120: Unless there has been a transform (e.g.  $\ln(y)$ ). Square bracket is facing the wrong way.  
We do not understand. Observations can be positive or equal to zero.
- Line 123: Space between units.  
This has been corrected. Thank you.
- Line 156: ‘Large multiple’  
This has been corrected.
- Line 178: ‘as this paper’  
This has been corrected.
- Line 278: What are the upper/lower bounds of the uniform distribution?  
This has been added. Thank you a lot.
- Line 280: What are the shape parameters of the log-gamma distribution?  
This has been added.

- Line 348: ‘Harmful’ is an incorrect choice of word here. You can just delete it;  
This has been removed.
- Line 348-350: This sentence isn’t correct. Observations don’t have a high likelihood. Please rephrase.  
This has been corrected.
- Line 351: Change ‘totally legitimate’ to ‘valid’  
This has been corrected.
- Line 353-355: This sentence does not make sense. I’m unsure of its meaning, please revise.  
This has been corrected. Thank you.
- Line 368: 4c and 5a  
This has been corrected (4b and 5a).
- Equation A1 and A2: Second term is incorrect, not divide by 2 but multiplied by 2.  
This comment has been removed by the reviewer.

Thank you very much for all these technical comments.

## References

Liu, Y., Haussaire, J.-M., Bocquet, M., Roustan, Y., Saunier, O., and Mathieu, A.: Uncertainty quantification of pollutant source retrieval: comparison of Bayesian methods with application to the Chernobyl and Fukushima Daiichi accidental releases of radionuclides, *Quarterly Journal of the Royal Meteorological Society*, 143, 2886–2901, <https://doi.org/10.1002/qj.3138>, 2017.