

Response to Reviewer 3 comments about the article “*A Bayesian Statistical Method to Estimate the Climatology of Extreme Temperature under Multiple Scenarios: the ANKIALE Package*”

ROBIN, Y., VRAC, M., RIBES, A., BARBAUX, O. and NAVEAU, P.

October 23, 2025

Note In this document, the text in regular format corresponds to the reviewers questions. The answers from authors are given in the grey blocks.

1 Reviewer 3 (Richard Chandler)

1.1 Global comment

The main purpose of this paper seems to be to improve on a methodology for extreme event attribution that was proposed by two of the authors. There are two main improvements: one is the ability to consider multiple emissions scenarios simultaneously, while the other relates to the adoption of faster algorithms for the analysis. In scientific terms, the first of these is much more interesting. Regarding the algorithms: it’s useful to know that the software is faster and I’m aware that the implementation can be a lot of work, but I wouldn’t regard it as a headline message for the paper, especially it seems to come just from adopting a more modern (but now fairly standard) MCMC method.

Thank you for this summary and for highlighting these two important points. We would like to point out that the targeted journal (GMD) deals with theoretical developments as much as technical developments. We therefore believe that both messages are important.

I’m not aware of other work that attempts to incorporate information from multiple emissions scenarios simultaneously within the framework of a single statistical representation, as is done here: most analyses treat individual scenarios separately which, as the authors correctly point out, leads to potential inconsistencies when focusing on quantities that are not scenario-dependent. As a first attempt to highlight this issue therefore — albeit in an application where the role of the scenarios needs to be explained more clearly (see general comments below) — this paper is welcome.

Unfortunately however, I think it needs a very major rewrite before it can be considered publishable. As far as I can tell, the methodology is probably consistent with many approaches to attribution — although I think it may be unnecessarily complicated and the approach to fitting the GAMs is possibly not optimal (see below). However, the manuscript is poorly written and, as a result, I find it hard to understand exactly what is being done and why. The scientific content is therefore hard to evaluate.

We apologise that our initial submission were difficult to evaluate. Your comments were therefore very welcome and have greatly helped us to improve our article.

The main problem is that the paper doesn't contain all of the information needed to make sense of it. There are also problems with the organisation of the material; with terminology that is not properly defined or is used incorrectly; and with mathematical notation that is not always defined or explained. Conversely, the paper makes heavy weather of some things that are actually rather straightforward (for example, the estimation of GAMs with a common component across scenarios — see below).

Thank you for your comments. We have completely rewritten the methodology section, trying to simplify the notation as much as possible without losing the meaning. It now contains a simple general idea of the method, with the example of Paris to help build intuition. The mathematical details have been moved to the appendix and supplementary material. We hope that this new organization will make it easier to read and understand.

Comments and questions follow. Some of them are probably due to my lack of understanding, for which I apologise. I have done my best, within the 12 hours or so that I'm willing to spend on a review.

1.2 General comments

1. The paper is not self-contained. The authors state explicitly that they will only describe the changes relative to their earlier paper (denoted as RR20); but a reader should not have to look at another paper in order to understand this one. A brief but complete summary of the relevant points from RR20 is needed.

The methodology section has been completely rewritten, and the appendices and supplement have been completed. The paper now contains all relevant details of the methodology.

2. There is no clear statement of what the analysis is for. The abstract (line 1) says that it's about estimating the statistics of temperature extremes; but the paper itself is framed around the methodology of attribution studies. Reading through the paper for the first time, I didn't know where it was going or why. To give just one example: at some point the aim seems to be the estimation of a quantity called κ which is described as "the random variable of a multi-model synthesis" (line 128): I still don't know what this represents, and from here on I started to get completely lost. It would be helpful to include, early on, a clear statement of what you're trying to estimate and to explain, in broad terms, the information requirements (I don't mean the precise datasets used: rather, the need for real-world time series of TX3x and whatever else you need for the real world; and GCM simulations of [etc. etc.]). It might also be helpful to include a schematic diagram — again, early in the paper — showing how all of these quantities fit together in the steps of the analysis.

The abstract and introduction have been rewritten for clarity, as well as the methodology section. A figure illustrating the whole procedure has also been added.

3. I don't understand the role of the scenarios in the context of an attribution study. Scenarios relate to future climate (albeit starting in 2015 for the CMIP6 runs), whereas at-

tribution studies typically work with information on the past and present. It's possible that the scenarios allow more precise estimation of anthropogenic effects using the GCM outputs, but this isn't explained anywhere. It's also possible that I have missed the point — but in that case it hasn't been explained clearly enough.

In an attribution context, scenarios allow us to project a current (real) event or class of events into the future to see how it will behave. We then speak of “prospective attribution”. Using climate models also allows us to strengthen the estimation of the parameters that drive non-stationarity (in particular μ_1 and σ_1). We have rephrased this to better highlight and clarify these different points.

4. In some way related to the previous two points: throughout, the paper attempts to describe what was done, but there needs to be more justification for why it's being done.

We have rephrased that.

5. The approach uses the GCM outputs to place a prior distribution on real-world quantities, based on an assumption that the GCMs are centred on reality (lines 184–185). This is known to be unrealistic: climate models are collectively biased relative to the real-world, and the biases are not independent. I am aware that the authors' assumption is often made in the literature. However, the paper should at least provide an honest acknowledgement that it is known to be (very) wrong — this is obvious to anyone who understands CMIP, and Reto Knutti and co-authors have demonstrated it empirically beyond doubt. There is also literature that aims to address the problem by relaxing the assumption e.g. to a notion of co-exchangeability defined by Rougier et al. (2013).

I can imagine an argument that the precise assumptions about the GCMs don't matter too much because they are only being used to set a prior. My response to that would be: if the prior doesn't influence the results then you don't need it because the real-world data are already sufficiently informative; but if the prior does influence the results then it needs to be defensible. The work of Knutti et al. shows that this prior is not even close to being defensible. At the very least, there should be some investigation of sensitivity to plausible alternative prior choices.

A minor point, related to this issue, is that the formulation doesn't date back to Ribes et al. (2017) as claimed on line 185 — it goes back much further.

We never assume that GCMs are centred on reality, an assumption that is obviously false, as you point out. Indeed, there is a problem with our notations, and probably an error in what we wrote in the initial submission. Here is the approach we take, which is the one used by ANKIALE. The prior is constructed as a synthesis of different climate models, using the following hypothesis: “*the models are statistically indistinguishable from truth*”, developed by Ribes et al. (2017). Let:

- $\theta_* \sim \mathcal{N}(\hat{\theta}_*, \Sigma_{\hat{\theta}_*})$ be the desired multi-model synthesis
- $\theta_{\mathcal{M}}$ the mean response of an infinite set of models, and $\Sigma_{\mathcal{M}}$ the climate modelling uncertainty (assumed to be equal for each model);
- $\check{\theta}$ the “truth” (not assumed to be equal to the multi-model mean).

The indistinguishability hypothesis is equivalent to saying that θ_m , θ_* and $\check{\theta}$ all come from the same distribution, i.e. $\theta_m, \theta_*, \check{\theta} \sim \mathcal{N}(\theta_{\mathcal{M}}, \Sigma_{\mathcal{M}})$. Note the difference with the “truth plus error” hypothesis (not used here), where the θ_m are centred on reality, i.e. $\theta_m \sim \mathcal{N}(\check{\theta}, \Sigma_{\mathcal{M}})$.

In the new version of our document, we have clarified the notation, and all the mathematical details are given in the supplementary material, Sect. S.1.2.

6. The proposed model seems needlessly complicated. I’m not 100% sure of this, because the presentation around Eq. 1 is unclear (e.g. in lines 106–107 we learn that this equation wasn’t actually used). I am also aware that the proposed methodology arises from a line of reasoning that is accepted in the attribution literature. Nonetheless:
 - As presented, the model seems overparameterised. For example, there are constant terms in the model for $X_t^{R,F}$ in equation (1), but also in the models for μ_t and $\log \sigma_t$. These constant terms can presumably be merged: the only reason for distinguishing between them is the workflow which splits the analysis into a first stage analysing the global and regional temperatures, and a second stage in which the results are fed into the GEV estimation. It’s not clear to me why the analysis should be split in this way: surely the Bayesian computational machinery allows you to do everything in one step? A clear justification for the two-step approach is needed.

- First, the two steps are not treated in the same way at all. Global forcings can be constrained analytically (52 parameters, using the Gaussian conditioning theorem), whereas the five parameters of the GEV law require MCMC. Mixing the constants would require using only an MCMC approach to constrain all of the parameters, i.e. in 57 dimensions, which is much more complex to implement.
- Next, ANKIALE idea, which was not sufficiently emphasised, is to offer a set of extensible tools. In this context, certain statistical models do not have a constant that can be unified. For example, the following model is used for precipitation P_t in attribution studies (see the work of, e.g., van der Wiel et al., 2017; van Oldenborgh et al., 2017; Uhe et al., 2018; Tradowsky et al., 2023):

$$\begin{cases} P_t \sim \text{GEV}(\mu_t, \sigma_t, \xi_t), \\ \mu_t = \mu_0 \exp(\alpha/\mu_0 X_t), \\ \sigma_t = \sigma_0 \exp(\alpha/\mu_0 X_t), \\ \xi_t \equiv \xi_0. \end{cases} \quad (1)$$

This model is not present in ANKIALE, but it is offered in SDFC (on which ANKIALE relies for inference) and can easily be added.

- More fundamentally perhaps: I am aware of the fashion for regressing on global and regional mean temperatures in the attribution literature. I am also aware that this literature tends to focus on partitioning variation into natural and anthropogenically-forced components. If I understand correctly however, the partitioning here is derived (lines 140–147) from a regression on either the outputs of an energy balance model (EBM), or the net radiative forcings. It's not clear which option has been used to obtain the presented results (or, if it was an EBM, which one was used). However, the results in (for example) Figures S1 and S2 show a strong resemblance between the “naturally-attributed” component of global mean surface temperature and the natural component of the forcings (not shown in the paper, but I'm familiar with the forcing time series). I guess that the correlation between this component and the input forcings is at least 0.95, and possibly greater than 0.99 — indeed, it will be 1 if the forcings themselves have been used in the estimation procedure. Moreover, the estimated anthropogenic components of the global and regional temperatures will be smooth curves, just like the anthropogenic component of the forcings (or of the EBM outputs). I will therefore bet a moderate sum of money that one can obtain almost identical results to those reported in the paper by ignoring the global and regional mean temperatures altogether, and instead using the natural and anthropogenic components of the net radiative forcing in the models for the GEV parameters. This would simplify the analysis both conceptually and computationally (e.g. it would eliminate the need for GAMs).

It is perfectly true that the correlation between the natural part of the forcings and the counterfactual signal is close to 1. The aim here is to determine the response to the forcings, constrained by observations. One of the advantages is that it allows us to attribute the corresponding regional warming, as well as the warming of extremes, to a level of global warming, elements that are absent when other forcings, such as the net radiative forcing, are used directly.

We nevertheless attempted to replace the GAM decomposition with EBM (Leach et al., 2021; Smith et al., 2024). In the test performed, we have tested with EBM instead of forcings because there is no reason why global temperature should directly follow forcings, whether natural or anthropogenic. The results were unsatisfactory: for reasons that remain unknown, EBM are unable to reproduce the climate change seen in climate models.

The attribution community may consider this suggestion to be very oversimplistic. Nonetheless, I would be interested to see how it performs. With the ANKIALE software as described, this should not be hard to implement: just replace the estimates of $X_t^{*,N}$ and $X_t^{*,A}$ with the corresponding components of the net radiative forcing, and see what happens.

This approach would take us away from the subject of the article: to present improvements to an existing method, a tool allowing its use with an example application. We have instead proposed it as a possibility in the updated perspectives section.

7. The split of figures and tables between the main paper and supplement seems strange. For example, Figures S1–S3 are helpful in understanding the process and the results, but they have been relegated to the supplement. On the other hand, Tables 2 and 3 in the main paper take up a lot of space and will be of interest to relatively few potential readers: these could reasonably be moved to the supplement.

These figures were revised along with the methodology section and moved back into the main text. The tables were moved to the supplement.

1.3 Specific / detailed comments and queries

Line 1 “estimating the statistics of temperature extremes”: as noted above, it isn’t clear from this what you’re actually doing.

We have reworded.

Line 29 typo “combining”.

Thanks, corrected.

Line 30 a Bayesian framework?

Modified.

Lines 33–34 I would remove the sentence “However...inconsistencies”. It is redundant given the content of the subsequent paragraph, and only one of the issues is actually an inconsistency.

We have reworded.

Line 47 the time taken to “process the entire domain” presumably depends on the size of the domain! To put this in context, it would be helpful to indicate the approximate number of grid cells being referred to here.

More generally: if there are any real applications in which it’s necessary to carry out an analysis over a domain of this size, then a time scale of up to a week is perhaps still not fast enough. It is certainly too slow for the kinds of real-time attribution studies carried out by World Weather Attribution as described in lines 18–20, given that the time available for data analysis in such studies is only a fraction of the total time available. For those kinds of applications though, the focus is typically on specific events in much smaller regions. It’s certainly helpful to give an indication of computational cost; but it would be more relevant to indicate the cost for realistic applications of the methodology. Alternatively, give (e.g.) an indicative cost per thousand grid cells, so that readers can figure out for themselves the likely cost for their own applications.

We have rephrased this to specify that the entire procedure takes approximately 2.5 hours per grid point.

Line 51 what’s the relevance of the number of countries and their partial inclusion? What does “exact ratio” refer to?

This is simply information about the domain, and knowing what percentage of a country is included in it allows us to know to what extent the value presented here represents the country or not.

Line 64 “refer to” => “referred to”.

Thanks, corrected.

Lines 73–74 “we use observations from a weather station...: But you just said you were using ERA5. Perhaps the station data are used to illustrate the methodology and then, as a separate study, the ERA5 data are used to examine the entire European area. If this is the case, probably it would be helpful for the reader if the paper is reorganised clearly along these lines.

We now only use ERA5 throughout in order to simplify the presentation.

Lines 75–76 how does the use of a longer dataset allow you to “better verify the contribution”? Given more data, any reasonable method is going to do better. And, if you need an unusually long dataset to identify the benefits of the new approach, then the implication is presumably that the benefits aren’t obvious for datasets of more realistic size! (I’m playing devil’s advocate, but you need to think more carefully about what you’re trying to say).

You are right, and since there is a risk of consistency issues in the station series, we prefer to use only ERA5 in the revised version.

Lines 77–78 as noted above, I believe it is fairly common to use global or regional mean temperatures in attribution studies. Nonetheless, it would be helpful to clarify this — perhaps

as part of an introductory summary of the basic steps / logic of the approach (see general point 2 above).

Done.

Line 85 do the emissions scenarios really cover the historical period? I think the intended meaning is probably that historical estimates of emissions were used for the period 1850–2014, while scenarios based on alternative SSPs were used for the period 2015–2100.

Of course, we have reworded it.

Lines 89–90 is there a reference to support the claim that the current warming trajectory is close to the SSP2-4.5 scenario? Or is IPCC (2023) intended to be this reference? As written, it seems to be a reference just for the 2.8K estimate.

Sentence deleted.

Line 92 (also line 137) the GCM simulations of TX3x are broadly comparable with the ERA5 estimates because the latter are gridded. However, they are not comparable with the Paris station- based observations. How is this handled?

In itself, this is not a problem because the models are only used as priors. The question no longer arises since the station has been removed.

While on the subject of gridded data: each GCM has its own grid, and many of them are different from the grid used in the ERA5 dataset. How was this dealt with?

The parameters for each GCM are inferred on their own grids. It is during multi-model synthesis that the GCMs are regridded to the ERA5 grid using nearest neighbour interpolation. A fairly similar alternative approach was tested: nearest neighbour interpolation of the covariance matrices, and bilinear interpolation for the mean to smooth the effect of the GCMs' grid cells. The contribution was not very interesting. We have now added this information to the revised text.

Lines 86–97 see general comment 1 above.

See the response to the comment 1.

Line 104 this is where the paper starts to become hard to follow. There is no statement of what the covariate $X_t^{R,F}$ actually is: all we're told is that it's a proxy for something else. And how do the quantities in the final line of equation (1) relate to the regional covariate X_t^R mentioned immediately beforehand? And why are these quantities needed, given that the GEV parameters are modelled using the original quantity X_t^R ?

The answers to some of these questions emerge implicitly later in the paper. If I understand correctly, equation (1) and the subsequent paragraph do not correctly describe what was done. My best guess is that μ is represented

$$\mu_0 + \sum_{j \in \{G,R\}} (\mu_{N,j} x^{N,j} + \mu_{A,j}^{A,j})$$

with a corresponding expression for $\log \sigma$, where $x^{N,j}$ and $x^{A,j}$ represent respectively the ‘natural’ and ‘anthropogenic’ components of variation in the global or regional temperature series. Setting the $\{\mu_{A,j}\}$ and $\{\sigma_{A,j}\}$ to zero then allows an estimate of the GEV parameters in the absence of anthropogenic influence. If this is correct, then the presentation in the manuscript seems very over-complicated.

There is a misunderstanding of the term “add”, which did not imply an addition between covariates, but rather an addition of the covariate to the statistical model. This entire section has been rewritten and we do hope that the new version is now much clearer.

Line 105 typo “anthropogenic”.

Thanks, corrected.

Lines 106–107 “We also add...” this confirms that equation (1) does not describe what was done. See above.

By “add”, we meant this equation:

$$\left\{ \begin{array}{l} T_t \sim \text{GEV}(\mu_t, \sigma_t, \xi_t) \\ \mu_t := \mu_0 + \mu_1 X_t^{\text{R,F}} \\ \log(\sigma_t) := \sigma_0 + \sigma_1 X_t^{\text{R,F}} \\ \xi_t \equiv \xi_0 \\ X_t^{\text{R,F}} := X^{\text{R,0}} + X_t^{\text{R,N}} + X_t^{\text{R,A}} \\ X_t^{\text{G,F}} := X^{\text{G,0}} + X_t^{\text{G,N}} + X_t^{\text{G,A}} \end{array} \right.$$

It is now directly integrated into the text.

Lines 109–110 it is not correct that the model can be seen as a linear model. It’s not linear, and linear models have additive noise terms. Suggest deleting this sentence.

Sorry, it is a *generalized* linear model, from the exponential family. Corrected.

Line 112 “construct” => “estimate”? Also, the ‘factual / counterfactual’ vocabulary seems unnecessary and confusing. I know that it comes from the literature on causal inference, but the present paper doesn’t really deal with causality. If my summary in point 15 above is correct, then none of it is really needed — and other explanations can be simplified as well.

The “counterfactual” vocabulary is widely used in application examples (e.g. in intensity changes). The section has been rewritten, but we retain this formulation.

Line 116 it’s not clear to me that all elements of θ should be described as ‘parameters’. The X s would be more precisely described as latent or hidden variables. Arguably this is a minor point, but it does affect the readability of the paper.

The section has been entirely rewritten for clarity.

Line 128 what does κ represent? See general point 2 above.

This is the multi-model synthesis. The section has been rewritten.

Lines 131–132 this repeats lines 96–97 (but see general comment 1).

Thanks, modified.

Line 143 ‘Energy Balance Model’ appears twice. And which EBM are you using?

For CMIP5 model, we use the EBM model defined by Held et al. (2010) and studied in CMIP5 by Geoffroy et al. (2013). In the new version, this is specified in the mathematical details of the supplementary material, Sect. S.1.1.

Line 145 residuals from what? (actually I don’t think they are residuals).

That is indeed poorly worded. We wanted to express that it was the spline smoothing of the regional temperature without natural forcings. The section has been rewritten.

Line 147 how is the variance of ε^* represented? Is it the same for all SSPs?

Here, ε^* appears as a multivariate normal distribution among all the smoothing hyperparameters. Since the variance may differ from one hyperparameter to another on the spline basis, it is therefore not necessarily the same for each scenario. All of this is explained in more details in Section S.1.1 of the supplementary material..

Lines 148–151 This presentation is needlessly complicated. If GAMs are needed at all (see general point 6), the model could be fitted directly using the `gam()` function in the `mgcv` package in R, e.g. using the `by()` argument to fit different smooths for the SSPs while retaining the common natural forcing term. The required setup is actually quite straightforward: it may even be available somewhere in python, although I appreciate that there’s nothing in python that gets close to the capabilities of `mgcv`. We probably don’t need to know the details of backfitting either: this is a completely standard approach to fitting GAMs.

The problem with this model using GAMs is that the natural term is shared, and there is no package that can take this feature into account when writing the model (or we have not found a way to do so). We have therefore explained how inference is performed.

Line 153 if it is necessary to retain the details on how the GAMs are fitted, then the expression in line 154 is a sum rather than an average as described: either the expression or the description is wrong. And in line 154, I assume both terms are summed: they should be enclosed in parentheses to make this clear, therefore.

Indeed, we forgot to normalise by the number of scenarios. This has been corrected.

Line 156 as noted above, it would be helpful to move Figure S2 into the main paper. Also however, there is no indication as to whether this was obtained using radiative forcings or EBM outputs.

This is a CMIP6 model, so we have used radiative forcings defined by Smith (2020). In the new version, this is specified in the mathematical details of the supplementary material, Sect. S.1.1.

Line 157 “We can see that...” how can we see that? What are we supposed to be looking at? How many model realisations are there? If more than one, this hasn’t been mentioned in the text (or how you have dealt with it) — it is noted in the caption to Table 1, but if affects the interpretation of results then it needs to be discussed in the text as well.

The section has been rewritten.

Line 158 “peaks are due to the volcanoes” they are troughs I think, not peaks.

Thanks, corrected.

Line 159 when referring to the “constant” signal, do you mean constant over time or constant across SSPs? If over time: why is this relevant? If over SSPs: why do the volcanic and solar effects differ between them as suggested at the end of the line?

It was constant over time; we simply noted that we had the expected result: a constant counterfactual world, apart from volcanoes and solar cycles.

Line 161 “we propose the following method...” I have no idea what is being done here, or why.

The section has been rewritten. This specific part is detailed in App. B1.

Line 162 “resampled” what does this mean? Resampled from what?

The section has been rewritten. This specific part is detailed in App. B1.

Line 164 what “method described above” is being referred to here?

The section has been rewritten. This specific part is detailed in App. B1 and Sect. S.1.1.

Line 165 “the four smoothed covariates” what are these? I thought the previous step (described in lines 138–155) was designed to estimate these covariates.

The section has been rewritten. This specific part is detailed in App. B1.

Lines 184–188 it is not correct to describe the formulation here in terms of models that are “statistically indistinguishable from reality”, which implies that reality and the models are exchangeable. The representation set out here is sometimes called the ‘truth plus error’ representation, I think.

Indeed, there is a problem with our notations, and probably an error in what we wrote in the initial submission. Here is the approach we take, which is the one used by ANKIALE. The prior is constructed as a synthesis of different climate models, using the following hypothesis: “*the models are statistically indistinguishable from truth*”, developed by Ribes et al. (2017). Let:

- $\theta_* \sim \mathcal{N}(\hat{\theta}_*, \Sigma_{\hat{\theta}_*})$ be the desired multi-model synthesis
- $\theta_{\mathcal{M}}$ the mean response of an infinite set of models, and $\Sigma_{\mathcal{M}}$ the climate modelling uncertainty (assumed to be equal for each model);
- $\check{\theta}$ the “truth” (not assumed to be equal to the multi-model mean).

The indistinguishability hypothesis is equivalent to saying that θ_m , θ_* and $\check{\theta}$ all come from the same distribution, i.e. $\theta_m, \theta_*, \check{\theta} \sim \mathcal{N}(\theta_{\mathcal{M}}, \Sigma_{\mathcal{M}})$. Note the difference with the “truth plus error” hypothesis, where the θ_m are centred on reality, i.e. $\theta_m \sim \mathcal{N}(\check{\theta}, \Sigma_{\mathcal{M}})$.

In the new version of our document, we have clarified the notation, and all the mathematical details are given in the supplementary material, Sect. S.1.2.

Lines 202–203 I don’t know what ‘weights’ are referred to here, but this sentence suggests that outlying GCMs are excluded from the analysis. No clear justification is given for doing this but, if what’s written is correct, it is very worrying because it suggests that the analysis uses only those models that agree with each other. This is very poor scientific practice, and will also lead to underestimation of uncertainty in the final results.

We simply noted that this model appears almost like an outlier with respect to the multi-model synthesis, which is why it is not widely included in the synthesis (note, however, that the intersection is not empty, just improbable). We also specified which model exhibits this behavior when rewriting the text.

Line 207 ‘Recall that we have...’ when I read the paper, I didn’t recognise that this information had been provided previously. All of the relevant information needs to be provided, clearly and unambiguously, at the outset.

We recall the notations of the observations, described in Sect. 2, where the data are grouped together to make reading easier.

Line 208 “Our goal is to estimate the distribution...” really? I was unaware of this on reading the paper. And there has still been no clear statement of what *kappa* actually is, or why we might want to learn about it.

The section has been rewritten.

Line 210 the ‘double conditioning’ notation in this equation does not exist in conventional practice, as far as I’m aware. I think the denominator is incorrect, as well: what’s being done here is Bayes’ theorem conditional on $X_t^{*,o}$ which should lead to

$$\mathbb{P}(\kappa | T_t^o, X_t^{*,o}) = \frac{\mathbb{P}(T_t^o | \kappa, X_t^{*,o}) \mathbb{P}(\kappa | X_t^{*,o})}{\mathbb{P}(T_t^o | X_t^{*,o})}$$

I suspect there is also an unstated assumption that $X_t^{*,o}$ is irrelevant for T_t^o once κ is known, so that $\mathbb{P}(T_t^o | \kappa, X_t^{*,o}) = \mathbb{P}(T_t^o | \kappa)$. Apart from the denominator, this then gives equation (5) in the paper.

By being more rigorous about the notation, your equation appears to be:

$$\mathbb{P}(\kappa|T_t^o, X_t^{*,o}) = \mathbb{P}[(\kappa|\{T_t = T_t^o\})|\{X_t = X_t^{*,o}\}]$$

Whereas ours is:

$$\mathbb{P}[\kappa|(\{T_t = T_t^o\} \cap \{X_t = X_t^{*,o}\})]$$

This explains the differences between your calculation and ours.

Lines 211–212 “In other words...” this makes very heavy weather of quite basic material, while at the same time failing to provide enough information as to what’s being done and why.

This is the translation of Eq. (5). It may be clearer with the correct equation.

Line 212 “are constrained” by what? Why?

See comment below.

Line 216 why do you want “the posterior of global and regional temperatures”?

We think that at this stage, important information has been missed. The section has been rewritten and we think that it should now contain all relevant information.

Line 218 in this equation, the left-hand side depends on t but the right-hand side doesn’t. I have no idea what is intended. What are the dimensions of the subsequent matrices? What are you doing?

The exact form of this matrices is now given in the Supplementaty material, Sect. S.1.3.1. The time axis was in matrix A, which contains the spline basis.

Line 224 again, what are the dimensions of these quantities?

See answer above.

Line 227 N^{SSP} isn’t defined.

We apologize for this error, the notations have been revised.

Line 228 what do you mean by ‘null values’? If you mean zeroes, say so.

It was just zero, so we put zeros.

Lines 234–235 what do you mean by “scenarios” and “observations” here? Are the “scenario” data from the climate models? If so, what’s the basis for this assumption? After all, the available scenarios (assuming you mean the SSPs that are included in the analysis) are arbitrary: the observations can’t be the average over all collections of scenarios.

If the scenario data here are not from the climate models, where are they from?

Scenarios here come from the multi-model synthesis (the prior), conditioned by a SSP. The observational constraints using the Gaussian conditioning theorem consist of writing the observations as a transformation of the prior by a projection matrix, plus an observed error term. Here, as we are dealing with several scenarios simultaneously, we have a particular configuration: we may have several scenarios over the observed period. We simply propose to take the average of the scenarios over this period, which we justify by the fact that the scenarios are almost the same over the period 2014-2025. We have rewritten the text to clarify our approach.

Line 237 how is it ‘more general’ to take the average of the scenarios? All you’re doing is to making one choice instead of another.

We now clearly present the two approaches with their advantages and disadvantages, and explain why we chose one over the other. In a nutshell:

- One approach where the average of the scenarios can be associated with the observations. This is justified by the fact that the climate scenarios are sufficiently similar over the period 2015–2024.
- One approach where we choose a scenario, which is therefore slightly less general because it assumes expertise about the future that we are borrowing.

Line 240 what does ‘this constraint’ refer to?

The constraint of κ by $X_t^{*,o}$.

Line 241 ‘the covariate Europe’?! At this point, I’m starting to wonder whether all of the authors checked the manuscript before it was submitted.

It is certainly a French idiom that apparently does not translate into English. The section has been rewritten.

Line 263 use of the NUTS is sensible, but there’s no need to provide details of, e.g. the generation of proposal distributions or anything else in the paragraph. The explanation provided isn’t enough for readers who are not unfamiliar with the algorithm, and they don’t need to understand how NUTS works in any case. All a reader needs to know is that this is a more modern MCMC method that can often overcome the difficulties described in the previous paragraph.

The section has been rewritten, and this information has been moved to App. B.3.2.

Line 280 what are “the laws”?

Laws, distributions, probability distributions. It is true that we have used these words as synonyms. We have rephrased.

Line 281 again, ‘recall’ implies that the reader has been told already.

The equations describing GEV have been moved to Appendix B.1, which we can refer to for the readers information.

Line 284 what’s an ‘intensity’? Don’t you just have a value of T here?

Yes, that is the language used in attribution. We have rephrased it.

Line 291 I don't think 'deduce' is the right word here. Rather, you can construct measures of the change.

Modified.

Line 295 Figure S4 is another that belongs in the main paper.

The section has been rewritten.

Lines 297 'we calculate...' why? *NB* the answer to this may be obvious to someone who understands what the authors are trying to do, but I lost track of that long ago.

The section has been rewritten.

Lines 301–303 why are you calculating the Wasserstein distance? Few readers will have a good intuition for what a given Wasserstein distance looks like, in terms of the distributions being compared. And most will be interested primarily in whether the distributions differ to an extent that would change any substantive conclusions of interest.

The section has been rewritten.

Line 312 'slightly ahead' in what sense?

The section has been rewritten.

Line 320 is the 'simultaneous analysis of several thousand grid points' often required? Given that you're just doing one grid point at a time (I think), attribution studies would usually focus on specific events that affect a subset of the grid points. More importantly though: in these kinds of applications, the focus is usually on weather that is simultaneously extreme over a large area. As far as I can see, the present paper doesn't address that problem: some discussion of this is needed, as part of the context for the work.

A discussion has been added.

Line 326 I wonder whether this kind of "package manual" material is appropriate for inclusion in a journal article? This query is perhaps for the journal editor.

This article focuses as much on methodology as it does on tools, which is why we chose GMD. We therefore believe that it has its place here.

Line 377 what does 'since 1940' mean in 'the maximum observed in 2024 since 1940'?

This is the maximal temperature observed between 1940 and 2024. We have reformulated.

Line 389 typo 'conter-factual'.

Thanks, corrected.

Line 451 what would be the challenges in extending the methodology to other variables such as wind and precipitation? It seems to me that it would be basically the same: it's all just based on the GEV distribution for annual maxima.

To our knowledge, no one has yet performed a multi-model synthesis followed by constraint with observations on variables other than temperature. In our previous paper (Robin and Ribes, 2020), the entire Sect. 3.3 sought to verify the best compromise (in terms of statistical model choice) in GCMs, and to verify the quality of the final fit to observations. This procedure would obviously need to be repeated when dealing with new variables. For information, the statistical model of Eq. 1, often used for precipitation, assumes a strong relationship between μ and σ (if one increases, so does the other), but we have been able to verify that this assumption is generally false in GCMs. Extending it to other variables is therefore a slightly more complex task than your proposal.

Caption to Table 2 what are these values averaged over? Grid cells, I assume — it would be helpful to clarify this. See also the earlier suggestion for moving this table (and Table 3) to the supplement.

Yes, the values are the average for the area, and the tables are in the supplement. It is now clarified.

Algorithm 1 this is a completely standard algorithm, and there is no need to include it (but if you do include it, the cross-references at the end are broken — similarly for Algorithm 2 on the next page).

It is certainly standard, but not necessarily known by the climate community. After rewriting the equations, we no longer needed it and removed it. The cross-references were indeed broken, which was corrected by updating the Copernicus template.

Appendix A I don't find Appendix A very helpful. The material on GAMs in Appendix A1 is mostly fairly standard, but seems more complicated than necessary; while I don't understand Appendix A2 because I still don't know what the purpose is.

All of this has been completely reworded in sections S.1.1. and S.1.3.1.

Finally: out of interest, before submitting my report I looked at the comments from the other reviewers. There is quite a lot of overlap with my comments above, particularly from Anonymous Reviewer 1 who also found the paper hard to follow. To be completely clear therefore: in this report, everything except this paragraph was written independently and without looking at the other review reports.

Thank you for your valuable comments and your honesty. As repeated along our responses, our article has been deeply modified (especially the method section) to clarify our approach and goals.

Bibliography

- Geoffroy, O., D. Saint-Martin, G. Bellon, A. Voldoire, D. J. L. Olivié, and S. Tytéca (Mar. 2013). “Transient Climate Response in a Two-Layer Energy-Balance Model. Part II: Representation of the Efficacy of Deep-Ocean Heat Uptake and Validation for CMIP5 AOGCMs”. In: *J. Clim.* 26.6, pp. 1859–1876. ISSN: 0894-8755, 1520-0442. DOI: 10.1175/JCLI-D-12-00196.1.
- Held, I. M., M. Winton, K. Takahashi, T. Delworth, F. Zeng, and G. K. Vallis (May 2010). “Probing the Fast and Slow Components of Global Warming by Returning Abruptly to Preindustrial Forcing”. In: *J. Clim.* 23.9, pp. 2418–2427. ISSN: 0894-8755, 1520-0442. DOI: 10.1175/2009JCLI3466.1.
- Leach, N. J., S. Jenkins, Z. Nicholls, C. J. Smith, J. Lynch, M. Cain, T. Walsh, B. Wu, J. Tsutsui, and M. R. Allen (May 2021). “FaIRv2.0.0: A Generalized Impulse Response Model for Climate Uncertainty and Future Scenario Exploration”. In: *Geosci. Model Dev.* 14.5, pp. 3007–3036. ISSN: 1991-959X. DOI: 10.5194/gmd-14-3007-2021.
- Ribes, A., F. W. Zwiers, J.-M. Azaïs, and P. Naveau (2017). “A New Statistical Approach to Climate Change Detection and Attribution”. In: *Clim Dyn* 48.1, pp. 367–386. ISSN: 1432-0894. DOI: 10.1007/s00382-016-3079-6.
- Robin, Y. and A. Ribes (2020). “Nonstationary Extreme Value Analysis for Event Attribution Combining Climate Models and Observations”. In: *Adv. Stat. Clim. Meteorol. Oceanogr.* 6.2, pp. 205–221. ISSN: 2364-3579. DOI: 10.5194/ascmo-6-205-2020.
- Rougier, J., M. Goldstein, and L. House (Sept. 2013). “Second-Order Exchangeability Analysis for Multimodel Ensembles”. In: *J Am Stat Assoc* 108.503, pp. 852–863. ISSN: 0162-1459. DOI: 10.1080/01621459.2013.802963.
- Smith, C. (Aug. 2020). *Effective Radiative Forcing Time Series from the Shared Socioeconomic Pathways*. Zenodo. DOI: 10.5281/zenodo.3973015.
- Smith, C., D. P. Cummins, H.-B. Fredriksen, Z. Nicholls, M. Meinshausen, M. Allen, S. Jenkins, N. Leach, C. Mathison, and A.-I. Partanen (Dec. 2024). “Fair-Calibrate v1.4.1: Calibration, Constraining, and Validation of the FaIR Simple Climate Model for Reliable Future Climate Projections”. In: *Geosci. Model Dev.* 17.23, pp. 8569–8592. ISSN: 1991-959X. DOI: 10.5194/gmd-17-8569-2024.
- Tradowsky, J. S., S. Y. Philip, F. Kreienkamp, S. F. Kew, P. Lorenz, J. Arrighi, T. Bettmann, S. Caluwaerts, S. C. Chan, L. De Cruz, H. de Vries, N. Demuth, A. Ferrone, E. M. Fischer, H. J. Fowler, K. Goergen, D. Heinrich, Y. Henrichs, F. Kaspar, G. Lenderink, E. Nilson, F. E. L. Otto, F. Ragone, S. I. Seneviratne, R. K. Singh, A. Skålevåg, P. Termonia, L. Thalheimer, M. van Aalst, J. Van den Bergh, H. Van de Vyver, S. Vannitsem, G. J. van Oldenborgh, B. Van Schaeybroeck, R. Vautard, D. Vonk, and N. Wanders (June 2023). “Attribution of the Heavy Rainfall Events Leading to Severe Flooding in Western Europe during July 2021”. In: *Clim Change* 176.7, p. 90. ISSN: 1573-1480. DOI: 10.1007/s10584-023-03502-7.
- Uhe, P., S. Philip, S. Kew, K. Shah, J. Kimutai, E. Mwangi, G. J. van Oldenborgh, R. Singh, J. Arrighi, E. Jjemba, H. Cullen, and F. Otto (2018). “Attributing Drivers of the 2016 Kenyan Drought”. In: *Int. J. Climatol.* 38.S1, e554–e568. ISSN: 1097-0088. DOI: 10.1002/joc.5389.
- van der Wiel, K., S. B. Kapnick, G. J. van Oldenborgh, K. Whan, S. Philip, G. A. Vecchi, R. K. Singh, J. Arrighi, and H. Cullen (Feb. 2017). “Rapid Attribution of the August 2016 Flood-Inducing Extreme Precipitation in South Louisiana to Climate Change”. In: *Hydrol. Earth Syst. Sci.* 21.2, pp. 897–921. ISSN: 1027-5606. DOI: 10.5194/hess-21-897-2017.
- van Oldenborgh, G. J., K. van der Wiel, A. Sebastian, R. Singh, J. Arrighi, F. Otto, K. Haustein, S. Li, G. Vecchi, and H. Cullen (Dec. 2017). “Attribution of Extreme Rainfall from Hurricane Harvey, August 2017”. In: *Environ. Res. Lett.* 12.12, p. 124009. ISSN: 1748-9326. DOI: 10.1088/1748-9326/aa9ef2.