

On Neighbourhood Cross Validation*

Simon N. Wood

School of Mathematics, University of Edinburgh, U.K.

`simon.wood@ed.ac.uk`

April 26, 2024

Abstract

It is shown how to efficiently and accurately compute and optimize a range of cross validation criteria for a wide range of models estimated by minimizing a quadratically penalized smooth loss. Example models include generalized additive models for location scale and shape and smooth additive quantile regression. Example losses include negative log likelihoods and smooth quantile losses. Example cross validation criteria include leave-out-neighbourhood cross validation for dealing with unmodelled short range autocorrelation as well as the more familiar leave-one-out cross validation. For a p coefficient model of n data, estimable at $O(np^2)$ computational cost, the general $O(n^2p^2)$ cost of ordinary cross validation is reduced to $O(np^2)$, computing the cross validation criterion to $O(p^3n^{-2})$ accuracy. This is achieved by directly approximating the model coefficient estimates under data subset omission, via efficiently computed single step Newton updates of the full data coefficient estimates. Optimization of the resulting cross validation criterion, with respect to multiple smoothing/precision parameters, can be achieved efficiently using quasi-Newton optimization, adapted to deal with the indefiniteness that occurs when the optimal value for a smoothing parameter tends to infinity. The link between cross validation and the jackknife can be exploited to achieve reasonably well calibrated uncertainty quantification for the model coefficients in non standard settings such as leaving-out-neighbourhoods under residual autocorrelation or quantile regression. Several practical examples are provided, focussing particularly on dealing with un-modelled auto-correlation.

1 Introduction

Cross validation has been widely used for model selection, evaluation and estimation of tuning parameters for a long time [e.g 44, 45], with a corresponding diversity of variants developed [see 1]. The basic idea of repeatedly omit-

ting a subset of the data during model fitting, and then assessing the quality of model predictions for the subset, gives rise to: the leave one out cross validation underpinning several smoothing parameter estimation criteria [e.g 6, 13, 15, 49]; leave out several cross-validation used to deal with auto-correlated data [e.g. 4, 39, 52, §6.2.2], see also [35]; and the k-fold cross validation [e.g. 18, §7.10] often used for validation and tuning parameter selection. In principle cross validation has the advantage of considerable generality. It can be applied to almost any loss function used for estimation, unlike the marginal likelihood based methods for smoothing parameter or variance parameter estimation [48, 50], for example. But a major problem with cross validation is computational cost: in general the model has to be refitted for each omitted subset (fold). For some special cases, such as univariate least squares spline smoothing, the cost can be reduced essentially to that of a single model fit [e.g 11, 19, 8]. But beyond the univariate least squares setting cost considerations lead to the use of approximations such as GCV [6, 13]. Splitting the data into only a small number of folds, as in k-fold cross validation, is another approach that reduces cost, but optimizing such scores with respect to hyperparameters is typically quite expensive, while the choice of folds introduces a degree of arbitrariness in the results. This paper shows how cross validation criteria can be computed efficiently and to high accuracy for a wide class of penalized regression models, optimized equally efficiently with respect to hyper-parameters, and exploited for uncertainty quantification. The methods apply equally to likelihood and non-likelihood based regular loss functions. Particular attention is paid to how this facilitates model estimation in the presence of unmodelled short-range residual auto-correlation – an issue of considerable applied importance.

Consider n data y_i . Suppose that we have a model for y_i , with parameters θ_i , which are in turn determined by p coefficients, β , and that the fit of the model is measured

*Submitted to *Statistical Science* 13th June 2023. Still waiting.

by a regular loss function

$$\sum_{i=1}^n \mathcal{D}(y_i, \theta_i).$$

For example $\mathcal{D}$ might be a negative log likelihood, a negative log pseudo or quasi likelihood, or a different loss such as the ELF loss [12] used in quantile regression. $\mathcal{D}$ may itself also depend on a small number of hyperparameters, but there will typically be a larger number of hyperparameters associated with penalties on β applied during estimation. Typically these penalties control the smoothness of spline terms in the model, or the dispersion of random effects.

For example, in the special case of a generalized additive model [17, 52] for some exponential family distributed y_i , $\mathcal{D}$ might be the corresponding deviance and θ_i would simply be $\mu_i = \mathbb{E}(y_i)$, parameterized in terms of the coefficients, β , of basis expansions for the model component smooth functions of covariates. These coefficients would then be estimated to minimize

$$\sum_{i=1}^n \mathcal{D}(y_i, \mu_i) + \sum_j \lambda_j \beta^T \mathbf{S}_j \beta \quad (1)$$

where the penalties $\beta^T \mathbf{S}_j \beta$ measure function wiggleness. The hyperparameters, λ_j , therefore control how smooth the estimated model functions will be. In the case in which the loss is a negative log likelihood and each parameter, θ_i , is modelled by its own smooth additive linear predictor, the class of models is the distributional regression models (generalized additive models for location scale and shape) of [57, 38, 30, 21, 22, 56, 54, 42] etc.

Generically the direct cross validation approach to estimating hyperparameters is as follows. For $k = 1 \dots m$ choose subsets $\alpha(k) \subset \{1, \dots, n\}$ and $\delta(k) \subset \{1, \dots, n\}$. Usually $\delta(k) \subseteq \alpha(k)$. Let $\theta_i^{-\alpha(k)}$ denote the estimate of θ_i when data with indices in $\alpha(k)$ are omitted from the estimation data. Then the *neighbourhood* cross validation (NCV) criterion

$$V = \sum_{k=1}^m \sum_{i \in \delta(k)} \mathcal{D}(y_i, \theta_i^{-\alpha(k)}) \quad (2)$$

is optimized with respect to the hyperparameters. Leave out one cross validation corresponds to $m = n$ and $\alpha(k) = \delta(k) = k$. K-fold cross validation has $m = K$, $\alpha(k) = \delta(k)$ and each $\alpha(k)$ an approximately equally sized non-overlapping subset of $\{1, \dots, n\}$, s.t. $\cup_{k=1}^K \alpha(k) = \{1, \dots, n\}$. A particularly interesting case occurs when $\text{nei}(i)$ denotes some *neighbourhood* of i determined, for example, by spatial or temporal proximity

and $\alpha(k) = \text{nei}(k)$ with $m = n$ and $\delta(k) = k$. If it is reasonable to assume that there is short range residual correlation between point k and points in $\alpha(k)$, but not between k and $j \notin \alpha(k)$, then NCV provides a means to choose hyper-parameters without the danger of overfit that such short range (positive) autocorrelation otherwise causes [e.g. 4]. Since the residual autocorrelation is not being directly modelled, then if $\mathcal{D}$ is nominally a negative log likelihood, this approach is treating it as a pseudo-likelihood. Uncertainty quantification has to allow for this.

The various versions of cross-validation covered by (2) have a long history, but their use for routine hyperparameter estimation has been limited by computational cost. Other than in linear special cases, the cost of evaluating (2) is $O(mnp^2)$ for a regression model with $O(np^2)$ estimation cost, and often $m = n$. The aim of this paper is to provide methods by which hyperparameters can efficiently be estimated by optimization of accurate approximations to (2), reducing the cost to $O(np^2)$. Uncertainty quantification is also discussed: the link between cross validation and the jackknife being useful in the cases of a non-likelihood based loss, and in the presence of residual auto-correlation.

The paper is structured as follows. Section 2 covers efficient and accurate computation of the NCV criterion and its derivatives w.r.t. smoothing or precision parameters. Enhanced computational and statistical robustness is considered in section 3. Section 4 covers optimization of the criterion and section 5 discusses variance estimation and the link to the jackknife. Section 6 provides simulation examples of performance under auto-correlation and with quantile regression, while section 7 provides brief example applications to electricity load prediction, spatial modelling of extreme rainfall and spatio-temporal modelling of forest health data.

2 Computing the neighbourhood cross-validation criterion

Let $\hat{\beta}$ denote the model coefficient values optimizing the penalized loss, and let $\mathbf{g}_i$ denote the gradient of the unpenalized loss for y_i with respect to β at $\hat{\beta}$. Let $\mathbf{g}_{\alpha(i)} = \sum_{j \in \alpha(i)} \mathbf{g}_j$. Furthermore let $\mathbf{H}_\lambda$ denote the Hessian of the penalized loss with respect to β at $\hat{\beta}$, while $\mathbf{H}_{\alpha(i), \alpha(i)}$ is the Hessian of the unpenalized loss for y_j such that $j \in \alpha(i)$. Note that the rank of $\mathbf{H}_{\alpha(i), \alpha(i)}$ is the cardinality of $\alpha(i)$. The Hessian of the penalized loss based on the complement of $\alpha(i)$, at $\hat{\beta}$, is therefore $\mathbf{H}_{\lambda, \alpha(i)} = \mathbf{H}_\lambda - \mathbf{H}_{\alpha(i), \alpha(i)}$. The change in $\hat{\beta}$ on omission of $\mathbf{y}_{\alpha(i)}$ can then be approximated by taking one step of

Newton's method

$$\Delta^{-\alpha(i)} = \mathbf{H}_{\lambda, \alpha(i)}^{-1} \mathbf{g}_{\alpha(i)}. \quad (3)$$

Consider the accuracy of the approximation, for a regular loss (meeting the Fisher regularity conditions, for example). To fix ideas suppose that B-spline based sieve estimators are used for the model smooth covariate effects, so that p may increase with n , typically at a slow rate such as $p \propto n^{1/5}$ [see 5]. Further assume that the sample size is increasing such that the rate of increase within the support of each basis function is the same. Then $\mathbf{H}_{\lambda, \alpha(i)}^{-1} = O(p/n)$, while $\mathbf{g}_{\alpha(i)} = \sum_{j \in \alpha(i)} \mathbf{g}_j = O(1)$, provided that the size of $\alpha(i)$ is not increasing with n . Note that as p increases the number of non-zero elements of $\mathbf{g}_{\alpha(i)}$ is $O(1)$, so $\Delta = O(p/n)$, and $\|\Delta\| = O(p^{3/2}/n)$ (Euclidean norm).

Newton's method is quadratically convergent, meaning that sufficiently close to the optimum, β^*

$$\|\beta^{k+1} - \beta^*\| \leq M \|\beta^k - \beta^*\|^2$$

for some finite positive constant M [see 31, Thm 3.5 and A.2]. Let $\epsilon^k = \beta^k - \beta^*$, and $\Delta^k = \beta^{k+1} - \beta^k = \epsilon^{k+1} - \epsilon^k$. Then $\|\Delta^k\| \leq \|\epsilon^k\| + \|\epsilon^{k+1}\|$, and hence $\|\Delta^k\| \leq \|\epsilon^k\| + M \|\epsilon^k\|^2$. Similarly $\|\epsilon^k\| \leq \|\Delta^k\| + M \|\epsilon^k\|^2$, so that $\|\Delta^k\| \geq \|\epsilon^k\| - M \|\epsilon^k\|^2$. So $\|\epsilon^k\|$ and $\|\Delta^k\|$ are of the same order. Hence if ϵ^0 is the error in $\hat{\beta}$ before adding Δ , then $\|\epsilon^0\| = O(p^{3/2}/n)$. The error after one Newton step is hence $\|\epsilon^1\| = O(p^3/n^2)$, which is in turn the error in V/n . For cubic splines, for example, $p = O(n^{1/5})$ is typically used ($p = O(n^{1/9})$ gives the optimal convergence rates for regression splines, but the higher rate is needed under penalization).

In itself the idea of taking single Newton steps is of course not new, and is what is done anytime a second order Taylor expansion is used to derive an easier to compute approximation. However, previous work on cross validation has not succeeded in retaining the low approximation error while also reducing the leading order computational cost to that of a single model fit [see e.g. 43, for a discussion of the issues in an ML context]. Without such a reduction leave-one-out cross validation is not computationally competitive with marginal likelihood or AIC type methods, and more general cross validation methods will have impractical cost as estimation criteria.

The difficulty is that naive computation with (3) would be $O(np^3)$, an improvement on the $O(n^2p^2)$ of full refit but still exceedingly expensive relative to the $O(np^2)$ of marginal likelihood methods. However, the naive computation turns out to be unnecessary. One approach would use the conjugate gradient method [e.g. 47, §5], to solve $(\mathbf{H}_\lambda - \mathbf{H}_{\alpha(i), \alpha(i)}) \Delta^{-\alpha(i)} = \mathbf{g}_{\alpha(i)}$ using $\mathbf{H}_\lambda$ as preconditioner. This only involves a few $O(p^2)$ steps (typically one or two in practice), but assumes that $\mathbf{H}_{\lambda, \alpha(i)}$

is positive definite. In practice it almost always is, but there is no finite sample size guarantee of this, so conjugate gradients do not offer a satisfactory general solution.

There are two aspects to the theoretical possibility of an indefinite Hessian. How to diagnose indefiniteness at low cost, and what to do when it is encountered. The pragmatic approach to the latter problem is simply to solve (3) using a method that does not require a positive definite Hessian, in order to ensure the continuity of V required for stable optimization. However, if this is done, it is imperative to at least have a means to detect that the problem has occurred, so that results can be carefully checked. A reasonable strategy starts from the Cholesky factorization of the full data penalized Hessian, $\mathbf{R}_0^\top \mathbf{R}_0 = \mathbf{H}_\lambda$. Updating of $\mathbf{R}_0$ when $\mathbf{H}_\lambda$ is subject to low rank updates is $O(p^2)$, and the update process reveals whether or not the updated Hessian is positive definite [14, §6.5.4]. Hence if $\mathbf{H}_{\lambda, \alpha(i)}$ is detected to be positive definite then its (triangular) Cholesky factor can be used to solve (3) at $O(p^2)$ cost. If an indefinite $\mathbf{H}_{\lambda, \alpha(i)}$ is detected then an iterative method suitable for symmetric indefinite matrices can be substituted. For example a pre-conditioned MINRES algorithm [36, 47, §6.4]. Alternatively the Woodbury identity [16, §18.2d] can be used. If $\mathbf{H}_{\lambda, \alpha(i)} = \mathbf{R}_0^\top \mathbf{R}_0 - \mathbf{U}\mathbf{U}^\top$ where $\mathbf{U}$ is $p \times k$ then

$$\mathbf{H}_{\lambda, \alpha(i)}^{-1} = (\mathbf{R}_0^\top \mathbf{R}_0 - \mathbf{U}\mathbf{U}^\top)^{-1} = \mathbf{R}_0^{-1} \{ \mathbf{I}_p - \mathbf{R}_0^{-\top} \mathbf{U} (\mathbf{U}^\top \mathbf{R}_0^{-1} \mathbf{R}_0^{-\top} \mathbf{U} - \mathbf{I}_k)^{-1} \mathbf{U}^\top \mathbf{R}_0^{-1} \} \mathbf{R}_0^{-\top} \quad (4)$$

which has $O(kp^2)$ cost for $k < p$.

2.1 Hessian factor downdate algorithms

In practice the down-dating of Cholesky factors is performed by repeated application of rank 1 down dates, using the Givens and hyperbolic rotation based algorithms given in Golub and van Loan [14, §6.5.4]. This is straightforward in simple generalized regression settings where the unpenalized Hessian is $\mathbf{H} = \mathbf{X}^\top \mathbf{W} \mathbf{X}$ for some model matrix, $\mathbf{X}$, and diagonal weight matrix, $\mathbf{W}$. Each rank 1 update of $\mathbf{H}$ is of the form $-\mathbf{X}_i^\top W_{ii} \mathbf{X}_i$, where $\mathbf{X}_i$ is the i th row of $\mathbf{X}$ corresponding to the omission of the i th observation. If any rank one update fails, the Cholesky factor before the update attempt is restored, and the failed update added to an array of failed updates. Computation with this modified Hessian can then use (4) with the partially updated Cholesky factor as $\mathbf{R}_0$ and the failed updates making up $-\mathbf{U}\mathbf{U}^\top$. Obviously any negative W_{ii} will result in an update that can not fail.

A broad class of models to which neighbourhood cross validation can be applied are the distributional regression models in which multiple parameters of a negative

log likelihood are each dependent on a separate linear predictor [38, 42, 57, 56, 21, 22, 54]. Some sharing of terms between predictors is also possible. Generically there are K linear predictors $\boldsymbol{\eta}^k = \mathbf{X}[J_k]\boldsymbol{\beta}[J_k]$ where J_k is a vector of indices of the model matrix columns and coefficients involved in the k th predictor. The log likelihood for y_i is then of the form $l(y_i, h_1(\eta_i^1), h_2(\eta_i^2), \dots, h_K(\eta_i^K))$ where the h_j are inverse link functions. Writing the second partial derivative of l_i w.r.t. η_i^j and η_i^k as l_i^{jk} then the Hessian is computed as

```

H = 0 $p \times p$ 
for  $j = 1 : K$  do
  for  $k = 1 : K$  do
    H[ $J_j, J_k$ ]+ =  $-\mathbf{X}[J_j]^\top \text{diag}(l^{jk}) \mathbf{X}[J_k]$ 
  end for
end for

```

Obviously in practice the symmetry of $\mathbf{H}$ is exploited for computational efficiency.

The down dates of the Hessian are less obvious in this case. The following algorithm breaks the computation of the dropping of the i th point down into rank 1 up and down dates, assuming the existence of a function `cholup(R, u, up)` which updates $\mathbf{R}$ to the Cholesky factor of $\mathbf{R}^\top \mathbf{R} + \{2\mathbb{I}(\text{up} = \text{TRUE}) - 1\} \mathbf{u} \mathbf{u}^\top$, if this is possible ($\mathbf{R}$ package `mgcv` provides such a function, for example).

```

R = R0; b = 0
for  $q = 1 : 2$  do  $\triangleright$  2 pass ensures up- before down-date
  for  $j = 1 : K$  do
    for  $k = (j + 1) : K$  do
       $\alpha = l_i^{jk}$ 
      if  $q = 1$  then  $b[j] + = \alpha$ ;  $b[k] + = \alpha$ .
      if  $(q = 1, \alpha > 0)$  or  $(q = 2, \alpha < 0)$  then
        v = 0
         $\mathbf{v}[J_j] = \sqrt{|\alpha|} \mathbf{X}[i, J_j]$ 
         $\mathbf{v}[J_k] + = \sqrt{|\alpha|} \mathbf{X}[i, J_k]$ 
        R = cholup(R, v,  $\alpha > 0$ )
      end if
    end for
     $\alpha = l_i^{jj} - b[j] \triangleright b[j]$  times block subtracted
    if  $(q = 1, \alpha > 0)$  or  $(q = 2, \alpha < 0)$  then
      v = 0;  $\mathbf{v}[J_j] = \sqrt{|\alpha|} \mathbf{X}[i, J_j]$ 
      R0 = cholup(R, v,  $\alpha > 0$ )
    end if
  end for
end for

```

The basic idea is that the rank one updates creating the off diagonal blocks of $\mathbf{H}$ also add unwanted contributions to the leading diagonal blocks. These extra contributions are tracked in the length K vector $\mathbf{b}$, and then corrected as part of each leading diagonal block's update. The two passes of the algorithm ensure that all guaranteed positive definite updates are made before any updates

that could potentially remove positive definiteness.

As with the simpler regression models, the algorithm can be modified to deal with indefinite updates. If a rank one update fails, we simply revert to the pre-attempt $\mathbf{R}$ and store the corresponding $\mathbf{v}$ as a column of $\mathbf{U}$ to employ with (4) or an equivalent iterative solver.

In summary, (3) is computed via two triangular solves using the cheaply computed Cholesky factor of $\mathbf{H}_{\lambda, \alpha(i)}$, or via (4) using the partially updated Cholesky factor and the skipped updates, $\mathbf{U}$, if indefiniteness is detected. The latter case warrants some sensitivity checking, perhaps checking the results with the offending points omitted entirely from fitting, or in some cases down-weighting them.

2.2 Derivatives of the NCV

Efficient and reliable optimization with respect to multiple smoothing/precision parameters is only possible if numerically exact derivatives of the NCV criterion with respect to the log smoothing parameters are available. The derivatives of V w.r.t. $\boldsymbol{\rho} = \log \boldsymbol{\lambda}$ follow from $d\hat{\boldsymbol{\beta}}/d\rho_j$ and

$$\begin{aligned} \frac{d\Delta^{-\alpha(i)}}{d\rho_j} &= -\mathbf{H}_{\lambda, \alpha(i)}^{-1} \frac{d\mathbf{H}_{\lambda, \alpha(i)}}{d\rho_j} \mathbf{H}_{\lambda, \alpha(i)}^{-1} \sum_{j \in \alpha(i)} \mathbf{g}_j \\ &\quad + \mathbf{H}_{\lambda, \alpha(i)}^{-1} \sum_{j \in \alpha(i)} \frac{d\mathbf{g}_j}{d\rho_j} \\ &= \mathbf{H}_{\lambda, \alpha(i)}^{-1} \left(\sum_{j \in \alpha(i)} \frac{d\mathbf{g}_j}{d\rho_j} - \frac{d\mathbf{H}_{\lambda, \alpha(i)}}{d\rho_j} \Delta^{-\alpha(i)} \right) \end{aligned}$$

all evaluated at $\hat{\boldsymbol{\beta}}$. Obviously $\partial \mathbf{H}_{\lambda, \alpha(i)} / \partial \rho_j = \partial \mathbf{H}_\lambda / \partial \rho_j - \partial \mathbf{H}_{\alpha(i), \alpha(i)} / \partial \rho_j$ with both terms on the right hand side depending in turn on $d\hat{\boldsymbol{\beta}}/d\rho_j$. The latter is obtained from (1), or its generalization with $\mathcal{D}(y_i, \boldsymbol{\theta}_i)$ as the loss, by implicit differentiation,

$$\frac{d\hat{\boldsymbol{\beta}}}{d\rho_j} = -\lambda_j \mathbf{H}_\lambda^{-1} \mathbf{S}_j \hat{\boldsymbol{\beta}}.$$

Notice how these computations add only $O(p^2)$ operations given the preceding methods for efficient computation with $\mathbf{H}_{\lambda, \alpha(i)}^{-1}$. Derivatives of V w.r.t. ρ_j then follow by routine application of the chain rule at $O(np^2)$ computational cost in the k -fold or leave one out cases. Appendix B provides expressions for $d\mathbf{g}/d\rho_j$ for some specific cases.

2.3 Low level computational considerations

The $O(np^2)$ floating point operation cost of NCV is similar to the cost for GCV or REML. However in real-

ity NCV is more costly in practice, since its leading order operations consist of matrix-vector operations, while GCV and REML computations can be structured so that the leading order operations are matrix-matrix. Matrix-matrix operations can typically be arranged in a way that is highly cache efficient and any optimized BLAS [e.g. 55] will exploit this. At time of writing such cache efficiency typically results in some 20-fold speed up for matrix cross products, for example. Matrix-vector operations can not be arranged in this cache efficient manner [see 14, for an introduction to the issues].

However, this disadvantage can be offset by relatively simple parallelization. The terms in the summations (2) can be trivially parallelized, as can the summations required for derivative computation. Here, OMP [34] in C was used. Realized scaling is then reasonable. This is in contrast to attempts to use simple parallelization approaches for matrix-matrix dominated GCV or REML computations. In such cases anything gained by parallelization is often immediately lost, as different execution threads compete for cache, destroying the cache efficiency on which the BLAS relies. For matrix-vector dominated computation this trade-off does not occur.

3 More robust versions of NCV

Given the tendency of prediction error criteria to be more sensitive to underfit than overfit, robust versions of GCV have a long history in the context of smooth mean regression [e.g 40, 46, 26, 27, 28]. The basic idea is to add, to the standard leave one out cross validation criterion, a penalty on the change in the estimate of $\mu_i = \mathbb{E}(y_i)$ on omission of y_i from the fit. This penalty is essentially a stability criterion, and the analyst has to choose how much weight to give the prediction error and stability terms in the robustified criterion. The choice is fundamental as neither the tendency for overfit, nor the effect of the stability penalty, vanish in the large sample limit. Often the comparison of robust and non robust cross validation results serves as a useful check of statistical stability.

Beyond mean regression a more general notion of model stability is required. For example

$$V_s = \sum_{k=1}^m \sum_{i \in \beta(k)} \mathcal{D}(y_i, \theta_i^{-\alpha(k)}) - \sum_{k=1}^m \sum_{i \in \beta(k)} \mathcal{D}(y_i, \hat{\theta}_i)$$

measures the sum of the difference in loss between fits with and without each $\alpha(i)$ omitted. This is a measure of how sensitive the model fit is to omission of data, which is naturally on the same scale as the NCV score. Letting $\gamma - 1$ denote the weight to give to V_s , the robust NCV

criterion then becomes

$$V_r = \gamma \sum_{k=1}^m \sum_{i \in \beta(k)} \mathcal{D}(y_i, \theta_i^{-\alpha(k)}) - (\gamma - 1) \sum_{k=1}^m \sum_{i \in \beta(k)} \mathcal{D}(y_i, \hat{\theta}_i).$$

The definition of γ such that $\gamma = 1$ corresponds to the usual usage for robustified GCV criteria. If the loss is a negative log likelihood and leave one out or k -fold cross validation is used, then to within an additive constant NCV is a direct estimate of the KL-divergence, and V_s might be thought of as estimating a sort of ‘KL-stability’.

A simpler alternative is as follows. Suppose that, rather than simply being dropped, the data for a neighbourhood are replaced by data having the opposite influence on the gradient of the loss w.r.t. linear predictor, η . This leads to a doubling of the change in gradient on perturbation of the neighbourhood, and using the same quadratic model as in the data omission case, a doubling of the step length. In short,

$$\eta^{[-\alpha(i)]} \leftarrow 2\eta^{[-\alpha(i)]} - \hat{\eta}.$$

More generally one might choose to move a different size of perturbation, leading to

$$\eta^{[-\alpha(i)]} \leftarrow \gamma \eta^{[-\alpha(i)]} - (\gamma - 1)\hat{\eta}.$$

again for $\gamma > 1$. The modified $\eta^{[-\alpha(i)]} \leftarrow 2\eta^{[-\alpha(i)]} - \hat{\eta}$ is then used in (2) to obtain an alternative V_r .

Computation of V_r requires only routine modifications of the computations for V .

3.1 Finite sample computational robustness

Sometimes models are specified in such a way that some β values are impossible. For example, using a regression model for Poisson data with a square root or identity link will predict negative expected values for some regression coefficient values. For model fitting this is not usually problematic, because a numerical optimizer can easily avoid regions of infinite negative log likelihood. However the proposed NCV requires taking Newton steps without checking that the loss function is finite at the step end. Indeed making such checks would break the differentiability of the NCV score on which efficient optimization relies. Further, while the high convergence rates on which the proposal relies suggest that this issue will be rare at large sample sizes, there are no guarantees at finite sample size, and in any case it only takes one problematic datum at some point in the NCV optimization to break the optimization.

A simple fix for this problem is to replace the loss function in the NCV criterion with a quadratic approximation

about the full model fit. Consider $\mathcal{D}$ when several linear predictors are involved. For compactness denote derivatives by sub and superscripts so that $\mathcal{D}_{\hat{\eta}}^j = \partial \mathcal{D} / \partial \eta_j|_{\hat{\eta}_j}$, and use the convention that repeated indices in a product are to be summed over. Suppressing the indices relating to particular observations, we then have

$$\gamma \mathcal{D}(\eta_1, \eta_2, \dots) - (\gamma - 1) \mathcal{D}(\hat{\eta}_1, \hat{\eta}_2, \dots) \simeq \mathcal{D}(\hat{\eta}_1, \hat{\eta}_2, \dots) + \gamma \mathcal{D}_{\hat{\eta}}^j (\eta_j - \hat{\eta}_j) + \gamma \mathcal{D}_{\hat{\eta}\hat{\eta}}^{jk} (\eta_j - \hat{\eta}_j)(\eta_k - \hat{\eta}_k)/2.$$

This quadratic approximation on the right hand side is then used in place of $\mathcal{D}$ in (2). The resulting criterion, V_q say, is always finite. Its derivatives with respect to smoothing/precision parameters rely on differentiating the quadratic approximation to yield

$$\begin{aligned} \mathcal{D}_{\hat{\rho}}^i &\simeq \mathcal{D}_{\hat{\eta}}^j \frac{d\hat{\eta}_j}{d\rho_i} + \gamma \mathcal{D}_{\hat{\eta}\hat{\eta}}^{jk} \frac{d\hat{\eta}_k}{d\rho_i} (\eta_j - \hat{\eta}_j) \\ &\quad + \frac{\gamma}{2} \mathcal{D}_{\hat{\eta}\hat{\eta}\hat{\eta}}^{jkl} \frac{d\hat{\eta}_l}{d\rho_i} (\eta_j - \hat{\eta}_j)(\eta_k - \hat{\eta}_k) \\ &\quad + \gamma \mathcal{D}_{\hat{\eta}}^j \left(\frac{d\eta_j}{d\rho_i} - \frac{d\hat{\eta}_j}{d\rho_i} \right) + \frac{\gamma}{2} \mathcal{D}_{\hat{\eta}\hat{\eta}}^{jk} \left(\frac{d\eta_j}{d\rho_i} - \frac{d\hat{\eta}_j}{d\rho_i} \right) (\eta_k - \hat{\eta}_k) \\ &\quad + \frac{\gamma}{2} \mathcal{D}_{\hat{\eta}\hat{\eta}}^{jk} (\eta_j - \hat{\eta}_j) \left(\frac{d\eta_k}{d\rho_i} - \frac{d\hat{\eta}_k}{d\rho_i} \right) \end{aligned}$$

which, on collection of terms, is

$$\begin{aligned} \mathcal{D}_{\hat{\rho}}^i &\simeq \mathcal{D}_{\hat{\eta}}^j \left\{ (1 - \gamma) \frac{d\hat{\eta}_j}{d\rho_i} + \gamma \frac{d\eta_j}{d\rho_i} \right\} \\ &\quad + \frac{\gamma}{2} \mathcal{D}_{\hat{\eta}\hat{\eta}}^{jk} \left\{ \left(\frac{d\hat{\eta}_k}{d\rho_i} + \frac{d\eta_k}{d\rho_i} \right) (\eta_j - \hat{\eta}_j) \right. \\ &\quad \left. + \left(\frac{d\eta_j}{d\rho_i} - \frac{d\hat{\eta}_j}{d\rho_i} \right) (\eta_k - \hat{\eta}_k) \right\} \\ &\quad + \frac{\gamma}{2} \mathcal{D}_{\hat{\eta}\hat{\eta}\hat{\eta}}^{jkl} \frac{d\hat{\eta}_l}{d\rho_i} (\eta_j - \hat{\eta}_j)(\eta_k - \hat{\eta}_k) \end{aligned}$$

V_q is obviously never needed when the likelihood is finite for all finite β values, and even when this is not guaranteed seems to be needed only rarely. However there are cases where it is essential. For example for distributional regression with the generalized extreme value distribution it is very easy to find cases for which optimization of the original V fails.

4 NCV optimization

The preceding sections provide accurate, computationally efficient and robust, approximations for neighbourhood cross validation criteria and their first order derivatives with respect to smoothing parameters. However reliable optimization with respect to those smoothing parameters is still challenging. The main problem is that prediction

error based smoothness selection criteria are indefinite at very high or low values of the log smoothing parameters, because the model fit is then invariant to smoothing parameters changes. In practice problems caused by very low smoothing parameter values are extremely rare, but very high smoothing parameters occur often, corresponding to the case in which a model in the null space of a smoothing penalty is appropriate. For example, when using a cubic spline term, very large smoothing parameters are appropriate whenever a simple linear effect is supported by the data. Since such occurrences are routine, it is necessary to be able to deal with them efficiently.

The major problem with criteria indefiniteness at large smoothing parameters is that it can easily cause an optimizer to continually increase a smoothing parameter to the point at which numerical stability is lost as the smoothing penalty completely dominates the data. This issue is particularly acute when using Newton type optimizers, where low curvature on near flat sections of the criteria can lead to very large steps being taken. Simple steepest decent type optimizers, or conjugate gradient methods, avoid this large steps issue, but are extremely computationally inefficient for near indefinite problems.

Full Newton optimization requires the exact Hessian of the criterion to be optimized. If that is available then it can be used directly to detect indefiniteness, with optimization then proceeding only in the strictly positive definite subspace of the smoothing parameters. In the current context the full Hessian is both implementationally tedious and somewhat computationally expensive to obtain, making quasi-Newton optimization based on first derivatives a more appealing choice. However the approximate Hessian used by a quasi-Newton method is positive definite by construction, and therefore useless for detecting indefiniteness.

An alternative is to recognise that the indefiniteness arises because $\hat{\beta}$ ceases to depend meaningfully on one or more λ_j s. That is $d\hat{\beta}/d\rho_j \simeq 0$. Testing a vector condition is inconvenient and liable to involve arbitrary scaling choices. However, recognising that $\mathbf{H}_{\lambda}$ is the large sample approximate posterior precision matrix for β , then a suitable test for indefiniteness can be based on the scalar test of

$$\frac{d\hat{\beta}^T}{d\rho_j} \mathbf{H}_{\lambda} \frac{d\hat{\beta}}{d\rho_j} \simeq 0.$$

If this condition is met and $\partial V / \partial \rho_j \simeq 0$ then the j th element of the proposed quasi-Newton update step can be set to zero. It is readily shown that such a step is still a descent direction, and that as such its length can always be selected to satisfy the Wolfe conditions [see 31, section 3.1], the second of which is sufficient to ensure that the BFGS quasi-Newton update maintains positive definiteness of the Hessian approximation.

5 Cross validated uncertainty quantification and the jackknife

The link between cross validation and the jackknife is obvious, but is worth exploring for the NCV criterion based on predicting each single datum in turn (i.e. $\delta(k) = k$ for $k = 1, \dots, n$). Define $n \times p$ matrix $\mathbf{D}$ with i th row given by $\Delta^{-\alpha(i)}(n - |\alpha(i)|)^{1/2}/(n|\alpha(i)|)^{1/2}$, where $\Delta^{-\alpha(i)}$ is from (3). Then the Jackknife estimate of the covariance matrix of $\hat{\beta}$ is $\mathbf{V}_J = \mathbf{D}^\top \mathbf{D}$ [e.g. 41]. [7] (section 2.8.1) point out that for an estimator with convergence rate $O(n^{-c})$ the power of 1/2 in the definition of $\mathbf{D}$ should really be c : between 2/5 and 4/9 for cubic penalized regression splines, for example. However this correction turns out to be of negligible importance given what follows.

Unfortunately the jackknife estimator only gives well calibrated inference for independent response data. When there is residual auto-correlation then the estimate is better calibrated than standard results assuming independence, but performance is still far too poor for practical use. This behaviour is easy to understand, by expanding $\hat{\beta}(\mathbf{y})$ about $\boldsymbol{\mu} = E(\mathbf{y})$

$$\begin{aligned} \hat{\beta}(\mathbf{y}) &= \hat{\beta}(\boldsymbol{\mu}) + \sum_i \left. \frac{\partial \hat{\beta}}{\partial y_i} \right|_{\boldsymbol{\mu}} (y_i - \mu_i) \\ &\quad + \frac{1}{2} \sum_{ij} \left. \frac{\partial^2 \hat{\beta}}{\partial y_i \partial y_j} \right|_{\tilde{\mathbf{y}}} (y_i - \mu_i)(y_j - \mu_j) \end{aligned} \quad (5)$$

where $\tilde{\mathbf{y}}$ lies on a line between $\boldsymbol{\mu}$ and $\mathbf{y}$. In the case of a discrete y_i this expression should be understood as being based on estimation using quasi-likelihood. In that case estimates are identical whether the y_i are viewed as observations of discrete or continuous random variables, and the required derivatives exist. Typically $\partial \hat{\beta} / \partial y_i = O(n^{-1})$ and $\partial^2 \hat{\beta} / \partial y_i \partial y_j = O(n^{-2})$, as can readily be checked for single parameter exponential family distributions, or quasi-likelihoods, used with GLMs for example. Since $E(y_i - \mu_i) = 0$, and assuming bounded variance of $y_i - \mu_i$, the summations can be viewed as random walks. For independent y_i it then follows that $\sum_i \left. \partial \hat{\beta} / \partial y_i \right|_{\boldsymbol{\mu}} (y_i - \mu_i) = O_p(n^{-1/2})$, and $\sum_{ij} \left. \partial^2 \hat{\beta} / \partial y_i \partial y_j \right|_{\tilde{\mathbf{y}}} (y_i - \mu_i)(y_j - \mu_j) = O_p(n^{-1})$. For correlated y_i the rates may be lower, but the linear term still dominates the quadratic term asymptotically.

Concentrating on the dominant linear term, while writ-

ing $\hat{\beta}$ for $\hat{\beta}(\mathbf{y})$ and β for $\hat{\beta}(\boldsymbol{\mu})$ it follows that

$$\begin{aligned} \hat{\beta}^{-\alpha(i)} &\simeq \beta + \sum_{j \notin \alpha(i)} \frac{\partial \beta}{\partial y_j} (y_j - \mu_j) \\ \Rightarrow \Delta^{-\alpha(i)} &= \hat{\beta}^{-\alpha(i)} - \hat{\beta} \simeq - \sum_{j \in \alpha(i)} \frac{\partial \beta}{\partial y_j} (y_j - \mu_j). \end{aligned}$$

Hence while

$$\begin{aligned} \mathbf{V} &= \mathbb{E}\{(\hat{\beta} - \beta)(\hat{\beta} - \beta)^\top\} \\ &\simeq \mathbb{E}\left\{ \frac{\partial \beta}{\partial \mathbf{y}} (\mathbf{y} - \boldsymbol{\mu})(\mathbf{y} - \boldsymbol{\mu})^\top \frac{\partial \beta^\top}{\partial \mathbf{y}^\top} \right\} \end{aligned} \quad (6)$$

the components of the Jackknife estimator have expectation

$$\begin{aligned} &\mathbb{E}\{(\hat{\beta}^{-\alpha(i)} - \hat{\beta})(\hat{\beta}^{-\alpha(i)} - \hat{\beta})^\top\} \\ &\simeq \mathbb{E}\left\{ \frac{\partial \beta}{\partial \mathbf{y}_{\alpha(i)}} (\mathbf{y}_{\alpha(i)} - \boldsymbol{\mu}_{\alpha(i)})(\mathbf{y}_{\alpha(i)} - \boldsymbol{\mu}_{\alpha(i)})^\top \frac{\partial \beta^\top}{\partial \mathbf{y}_{\alpha(i)}^\top} \right\}. \end{aligned}$$

So the Jackknife estimator partly accounts for residual correlation in $\mathbf{y}$, but its component terms ignore the correlation between $\mathbf{y}_{\alpha(i)}$ and $\mathbf{y}_{-\alpha(i)}$, and are hence missing part of the correlation structure expected even if each y_i is only correlated with terms in neighbourhood $\alpha(i)$.

To avoid the problem of neglected residual auto-correlation, a simpler estimator of the coefficient covariance matrix is based on the observed version of (6), reusing the assumption of negligible residual correlation between y_i and $\mathbf{y}^{-\alpha(i)}$ (and assuming that $j \in \alpha(i) \Leftrightarrow i \in \alpha(j)$). Computation is especially simple if we note that from application of the preceding expansion to the leave one out case we can write $\Delta^{-i} \simeq \left. \partial \hat{\beta} / \partial y_i \right|_{\hat{\boldsymbol{\mu}}_i} (y_i - \hat{\mu}_i)$, so that a sample version of (6) would be:

$$\bar{\mathbf{V}} = \sum_i \Delta^{-i} \sum_{j \in \alpha(i)} \Delta^{-j^\top}. \quad (7)$$

Unfortunately the estimator underestimates variances. This can be seen by considering the residual products involved in (7). Denote the true residual product by $a = (y_i - \mu_i)(\mathbf{y}_{\alpha(i)} - \boldsymbol{\mu}_{\alpha(i)})^\top$, and let $\hat{\mu}_i = \mu_i + \delta_i$, where δ_i is the error in $\hat{\mu}_i$.

$$\begin{aligned} (y_i - \hat{\mu}_i)(\mathbf{y}_{\alpha(i)} - \hat{\boldsymbol{\mu}}_{\alpha(i)})^\top &= a - (y_i - \mu_i)\hat{\boldsymbol{\delta}}_{\alpha(i)}^\top \\ &\quad - \hat{\delta}_i(\mathbf{y}_{\alpha(i)} - \boldsymbol{\mu}_{\alpha(i)})^\top + \hat{\delta}_i\hat{\boldsymbol{\delta}}_{\alpha(i)}^\top \approx a - c\hat{\delta}_i\hat{\boldsymbol{\delta}}_{\alpha(i)}^\top \end{aligned}$$

where $c \geq 1$ and the final approximation follows from the fact that the error in $\hat{\mu}_i$ is a shrunken version of the local average of $y_i - \mu_i$ (short range autocorrelation increases both the local deviation of the average from zero and the degree of shrinkage under NCV).

To avoid this underestimation of variance we could base variance estimation on the cross validated residuals $\tilde{\epsilon}_k = y_k - \tilde{\mu}_k$ where $\tilde{\mu}_i = \hat{\mu}_i^{-\alpha(i)}$. Then we can define

$$\tilde{\Delta}^{-k} = \Delta^{-k} \tilde{\epsilon}_k / \hat{\epsilon}_k \quad (8)$$

which can be substituted for Δ^{-k} ($k = i, j$) in (7) to give the cross-validated estimate of $\tilde{\mathbf{V}}$. By definition $\mathbf{y}_{\alpha(i)} - \boldsymbol{\mu}_{\alpha(i)}$ plays no part in the estimation of $\tilde{\mu}_i$ - there is no relationship between $\tilde{\delta}_i$, the error in $\tilde{\mu}_i$, and the local average of $\mathbf{y}_{\alpha(i)} - \boldsymbol{\mu}_{\alpha(i)}$. Hence

$$\mathbb{E}\{(y_i - \tilde{\mu}_i)(\mathbf{y}_{\alpha(i)} - \tilde{\boldsymbol{\mu}}_{\alpha(i)})^T\} = \mathbb{E}(a) + \mathbb{E}\{\tilde{\delta}_i \tilde{\boldsymbol{\delta}}_{\alpha(i)}^T\},$$

and this estimate overestimates variance, by roughly the same amount as the previous estimate underestimated in the case that $c \approx 1$. This suggests using the average of the cross validated and direct estimators

$$\hat{\mathbf{V}} = (\tilde{\mathbf{V}} + \bar{\mathbf{V}})/2 \quad (9)$$

which will underestimate variances less than $\bar{\mathbf{V}}$ and overestimate less than the cross validated equivalent $\tilde{\mathbf{V}}$. An interesting question is whether a reliable estimator of c could be used to improve on this, by re-weighting the average, for cases where c is not close to one. An immediate generalization replaces the simple residuals used above with generalized residuals based on the loss itself and constructed such that zeros for all residuals would imply $\hat{\boldsymbol{\beta}} = \boldsymbol{\beta}$, so that (5) can be replaced by an expansion (about 0) directly in terms of these residuals. Deviance residuals can be used, for example.

So we have three distinct cases

1. For independent data using a general loss, then the jackknife estimator of the covariance matrix can be used.
2. For data subject to unmodelled short range autocorrelation (and sufficiently large n) then (9) is used, or conservatively the analyst might choose to use $\tilde{\mathbf{V}}$.
3. For independent data using a negative log likelihood loss then standard results can still be used; in particular the large sample approximation $\boldsymbol{\beta}|\mathbf{y} \sim N(\hat{\boldsymbol{\beta}}, \mathbf{H}_\lambda^{-1})$.

1 and 2 result in frequentist covariance matrices, which do not account for smoothing bias, and are therefore likely to result in miscalibrated inference. In case 3 this problem is often addressed by using the Bayesian posterior covariance matrix of the model coefficients for inference. The Bayesian covariance matrix can be interpreted as the frequentist covariance matrix plus a Bayesian bias correction: the expectation, under the prior, of the squared bias

matrix [see in particular 32, 29]. Computation of an appropriate Bayesian bias correction for a general regular loss can be conducted under the coherent belief updating framework of [3]. To do so requires an estimate of the learning rate parameter, ν , which is the weight multiplying the loss in the belief updating framework. In the current context it is hence a multiplier on the unpenalized Hessian of the loss. Obviously NCV is insensitive to ν - any multiplicative change in ν simply results in the same change in the estimated $\boldsymbol{\lambda}$. However, the asymptotic Bayesian and frequentist covariance matrices, $\mathbf{V}_b = \mathbf{H}_\lambda^{-1}$ and $\mathbf{V}_f = \mathbf{H}_\lambda^{-1} \mathbf{H} \mathbf{H}_\lambda^{-1}$, scale directly as ν^{-1} . Hence ν can be estimated from the relative scaling of $\mathbf{V}_f$ and $\hat{\mathbf{V}}$ (e.g. from the ratio of their traces), and the resulting $\hat{\nu}$ used to compute the Bayesian bias correction, $(\mathbf{V}_b - \mathbf{V}_f) \hat{\nu}^{-1}$, to be added to $\hat{\mathbf{V}}$ (or, more conservatively to $\tilde{\mathbf{V}}$).

6 Example simulations

By way of illustration of the ability of NCV to cope with short range auto-correlation, x, y data were simulated from the function $f(x) = 2.5 \sin(4\pi x) \exp(-2x)$ evaluated at x values spaced equally over $[0, 1]$. Gaussian, Poisson and gamma simulations were run, with short range autocorrelation generated from Gaussian AR or MA processes. In the AR1 case the autocorrelation model was $e_{i+1} = 0.6e_i + \epsilon_i$, where ϵ_i are i.i.d $N(0, 1)$ deviates. In the MA case the model was $e_i = \sum_{j=i-2}^{i+2} \epsilon_j$, but linearly scaled so that $\text{var}(e_i) = 1$. Then three response models were used: $y_i = f(x_i) + e_i$, $y_i \sim \text{Poi}\{\exp(f(x_i) + e_i - \sigma_e^2/2)\}$ or $y_i \sim \text{Gamma}\{\exp(f(x_i) + e_i - \sigma_e^2/2), \phi\}$ where $\phi = 0.1$.

The data were then fitted using appropriate spline models with the autocorrelation handled in one of four ways. The first option used penalized quasi-likelihood based fitting, with an AR1 correlation model on the working model scale. This is the correct model only for the Gaussian AR1 simulation, but is typical of what is often done in applications to somewhat account for otherwise un-modelled autocorrelation. The second option used NCV in which $\delta(i) = i$ and $\alpha(i)$ consisted of i and its 8 nearest x -neighbours. The final options simply ignored the autocorrelation and estimated smoothing parameters using REML or GCV. Figure 1 shows reasonably typical reconstructions (the second replicate of each case). Table 1 gives coverage probability and MSE results. NCV generally outperforms the alternatives by some margin, with coverage probabilities close to nominal at larger sample sizes for the moving average correlation process for which it is ideally suited. It is of course not a panacea: when there is some remaining auto-correlation between each

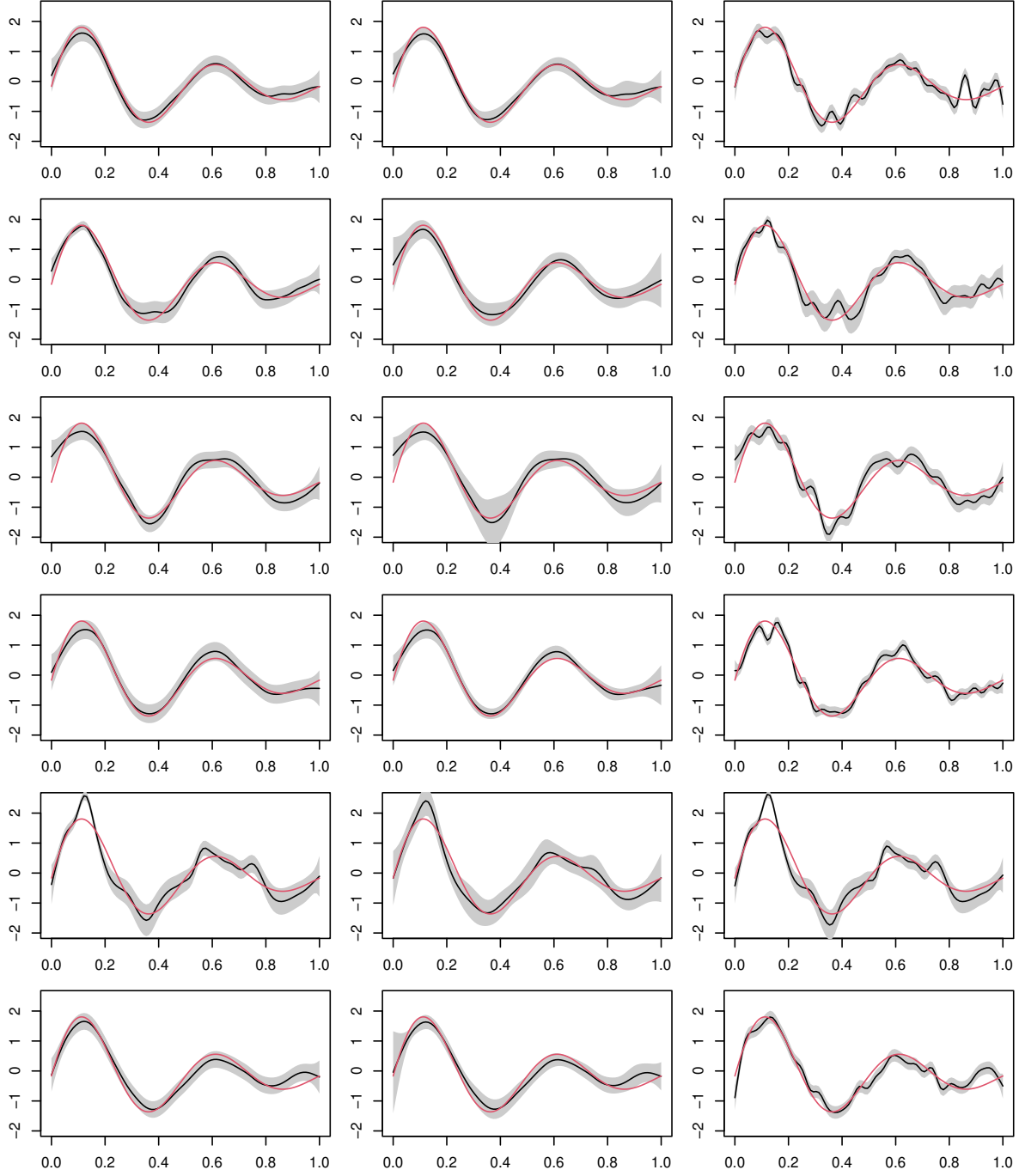


Figure 1: Reasonably typical example smooth estimates using PQL with working AR1 model (left), NCV (middle) and REML (right). In each panel nominal 95% confidence bands are shown in grey, truth is in red. For the top three rows autocorrelation was generated using an AR1 process, and in the bottom three an MA process. The top three rows are for Gaussian, Poisson and gamma responses respectively and the same ordering is repeated for the bottom 3 rows.

point i and the points outside $\alpha(i)$, then coverage probabilities are somewhat degraded.

Spatial example. Data were also simulated from a true function, $f(x, z)$, that was the sum of two scaled Gaussian p.d.f.s over a square x, z grid, with equal grid spacing in x and z . For each grid point independent $N(0, 1)$ deviates ϵ' were simulated. Dependent deviates ϵ were then produced for each grid point by averaging ϵ' values for the grid point and its 8 nearest neighbours. The average was weighted: 1 for the point itself, 0.5 for each of the 4 nearest neighbours and 0.3 for the 4 next-nearest neighbours. Edges were dealt with by simply extending the grid for the ϵ' , so that the required components of ϵ were always available. 3 models were simulated from: the Gaussian model $y_i = f(x_i, z_i) + \epsilon_i$, the Poisson model $y_i \sim \text{Poi}\{\exp(f(x_i, z_i) + \epsilon_i - \sigma_\epsilon^2/2)\}$ and the gamma model $y_i \sim \text{Gamma}\{\exp(f(x_i, z_i) + \epsilon_i - \sigma_\epsilon^2/2), \phi\}$ where $\phi = 0.1$. Grids of size 18×18 and 36×36 were used.

Models were then estimated in which the smooth function of x and z were represented using penalized thin plate regression splines, with basis dimensions of 100 for the smaller grid and 150 for the larger. Smoothing parameters were estimated by GCV, REML and by NCV with the omitted neighbourhoods for each point consisting of the square of 25 points centred on the point. Typical reconstructions using REML and NCV, for the larger grid, are shown in Figure 2. The tendency of REML to severely over fit when there is un-modelled short range residual auto-correlation is shown even more strongly by GCV [see 24, for an explanation]. Performance of the model fits over 500 replicates is shown in table 2.

Quantile regression example. Smooth quantile regression models provide an example where attempts to use generalized cross validation perform exceptionally poorly. Moving away from the 50% quantile increasingly severe overfit is observed, as discussed in [37]. The difficulty appears to relate to the averaging of leverage used in generalized cross validation, which is inappropriate for a highly asymmetric loss. The NCV approach avoids this problem and gives well calibrated inference. For example, the model $y \sim N(1 + x + x^2, \{1.2 + \sin(2x)\}^2)$ was used to simulate x, y data for 2000 x data evenly spaced over $[-4, 3]$: see figure 3. Smooth quantile regression was used to estimate the 95% quantile, with the quantile represented using a rank 50 penalized cubic spline. The ELF loss of [12] was used as detailed in appendix A. The smoothing parameter was selected by leave-one-out NCV, and the coefficient covariance matrix computed using the method in section 5. 200 replicates were run, with a typical reconstruction and 95% confidence band shown in figure 3. Over 200 replicates the average across the

function coverage probability of the nominally 95% credible bands was 0.947, with on average 4.82% of data lying above the estimated quantile. In this case the NCV criterion corresponds to the direct cross validation used by [33] (to within $O(p^3 n^{-2})$ error), but they did not consider interval estimation and used direct calculation and grid search for the smoothing parameter optimization, which would be infeasibly expensive for multiple smoothing parameters, unlike NCV.

7 Examples

7.1 UK electricity load

Figure 4 shows daily data on load on the UK national electricity grid between 6AM and 6.30AM¹. Operational prediction of the load one day ahead is often based on generalized additive models, either alone or as the major component of a mixture of experts approach. There is typically a degree of un-modelled autocorrelation in the residuals from such models. This can cause over fit, increasing prediction error and degrading estimates of predictive uncertainty. A typical model structure for the data shown is

$$\begin{aligned} \mathbb{E}(\text{load}_i) = & \alpha_{\text{dow}(i)} + f_{\text{dt}(i)}(\text{load48}_i) + f_1(\text{toy}_i) \\ & + f_2(\text{timeCount}_i) + f_3(\text{temp}_i) + f_4(\text{temp95}_i) \end{aligned}$$

where $\text{dow}(i)$ is the day of the week of observation i and $\alpha_{\text{dow}(i)}$ a parameter. The f_j are all unknown smooth functions. $\text{dt}(i)$ is a label taking one of 4 values: Mon, Sat, Sun or 'ww' corresponding to any working day that follows a working day. load48_i is the load 48 half hours before point load_i occurred. So the idea is that load is predicted by the load at the same time the previous day, but that the relationship varies from day to day. toy_i is time of year, timeCount_i is scaled total elapsed time, temp_i is average current temperature and temp95_i is exponentially smoothed lagged temperature.

Fitting the model to pre-2016 data, assuming load is normally distributed with constant variance, and estimating smoothing parameters by standard marginal likelihood maximization gives significant residual autocorrelation to lag 9. Refitting the model by penalized least squares with NCV using the 9 neighbouring points on either side of point i as $\alpha(i)$, results in smoother estimates for some model terms, as shown in figure 5. It also yields improved prediction of load over 2016. Both models have mean absolute percentage prediction error

¹obtained from the predecessor of <https://data.nationalgrideso.com/data-groups/demand> and https://demandforecast.nationalgrid.com/efs_demand_forecast/faces/DataExplorer

	AR1 correlation						Moving average correlation					
	$n = 250$			$n = 1000$			$n = 250$			$n = 1000$		
	Gau	Poi	Γ	Gau	Poi	Γ	Gau	Poi	Γ	Gau	Poi	Γ
S/N	1.14	0.76	0.87	1.14	0.73	0.79	1.44	0.89	1.07	1.42	0.87	0.97
LR	1.13	1.11	1.14	1.04	1.04	1.05	1.12	1.10	1.13	1.04	1.03	1.05
NCV CP	.919	.897	.904	.940	.933	.921	.934	.919	.925	.954	.945	.947
NCV ⁺ CP	.943	.918	.927	.947	.941	.931	.954	.939	.945	.960	.952	.954
PQL-AR CP	.744	.772	.831	.965	.779	.930	.610	.820	.728	.994	.840	.959
REML CP	.643	.741	.676	.695	.709	.733	.470	.773	.562	.617	.752	.679
GCV CP	.610	.748	.635	.666	.693	.706	.453	.774	.533	.593	.734	.649
NCV MSE	.085	.130	.103	.024	.036	.030	.068	.105	.080	.019	.030	.022
PQL-AR MSE	.297	.197	.166	.022	.059	.029	.611	.142	.322	.018	.042	.022
REML MSE	.159	.237	.174	.046	.077	.058	.194	.191	.190	.046	.060	.051
GCV MSE	.241	.629	.268	.072	.124	.085	.223	.599	.248	.061	.098	.070

Table 1: Nominal 95% confidence interval coverage and mean squared error for penalized regression splines with smoothness selection by GCV, REML and NCV with $\alpha(i)$ consisting of the 4 neighbours of i on either side plus i itself. The same spline was also fitted using PQL with an AR1 residual correlation model on the working model scale. Notice that NCV achieves close to nominal coverage for the larger sample size and MA process, for which it is eliminates all residual autocorrelation effects in smoothness selection. NCV⁺ is the case in which the conservative covariance matrix based purely on cross-validated residuals is used. LR is the ratio of conservative to basic interval widths. Cases in which a non-NCV method outperformed the basic NCV method are highlighted in bold. For the most part NCV achieves better coverage probabilities and lower MSE than the PQL approach, even for the Gaussian - AR1 simulations for which the latter model is correct.

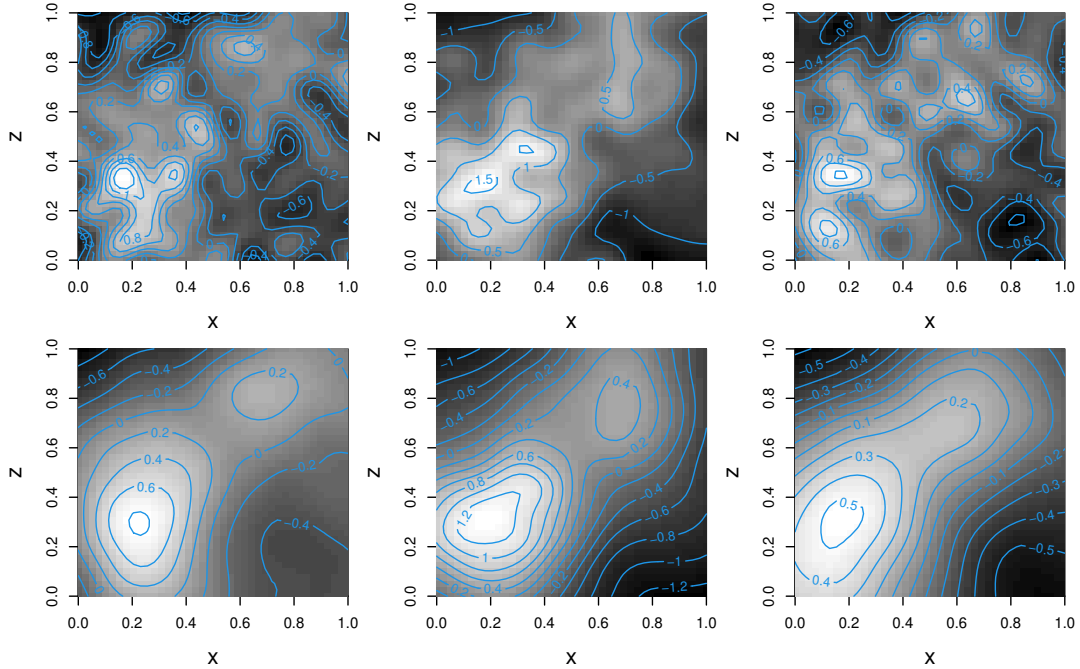


Figure 2: Typical example smooth estimates using REML (top) and NCV (bottom) smoothing selection, when spatial data are auto-correlated, and the mean field is a smooth sum of two Gaussians. Data are simulated on a grid, with errors that include a moving average process, where the average is over each grid point and its 8 nearest neighbours. NCV omits each point and its (24) auto-correlated neighbours in turn. Left is for Gaussian response, middle Poisson and right gamma.

	18 × 18			36 × 36		
	Gaussian	Poisson	gamma	Gaussian	Poisson	gamma
S/N	0.95	1.11	0.68	0.94	1.08	0.66
LR	1.16	1.16	1.15	1.07	1.09	1.07
NCV CP	.940	.966	.928	.948	.948	.948
NCV ⁺ CP	.960	.978	.949	.960	.962	.957
REML CP	.321	.827	.671	.467	.763	.672
GCV CP	.319	.826	.653	.459	.749	.651
NCV MSE	.040	.064	.037	.012	.021	.013
REML MSE	.111	.093	.091	.063	.061	.056
GCV MSE	.114	.116	.110	.068	.082	.069

Table 2: Spatial moving average simulation example results over 500 replicates, for two grid sizes and 3 response distributions, reporting signal to noise ratio (ratio of standard deviation of f to standard deviation of residual), coverage probabilities of nominally 95% CIs and mean square error, for NCV, REML and GCV smoothing parameter estimation. NCV⁺ gives coverage for nominally 95% intervals based on conservative use of the cross-validated residual based covariance. LR is the ratio of the average width of the conservative intervals to the standard NCV intervals.

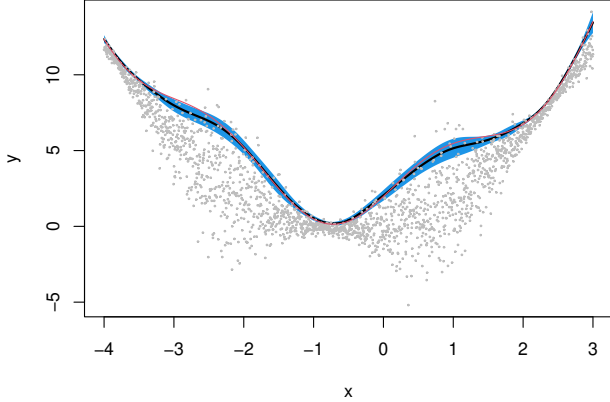


Figure 3: Typical example NCV based smooth estimate of 95% quantile for the distribution of the data shown in grey as a smooth function of x . The smooth is represented using a rank 50 penalized spline and the ELF loss is used. The fit is shown in black, with the 95% credible band in blue and the truth in red.

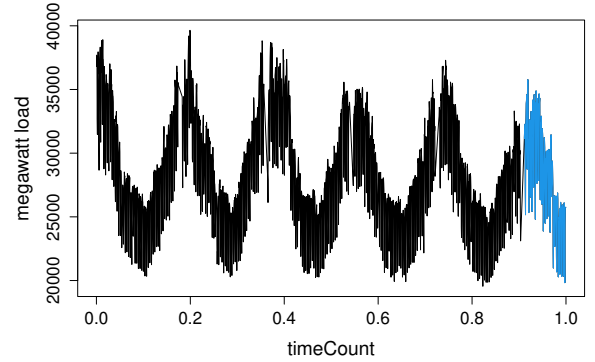


Figure 4: Megawatt load on the UK electricity grid from 6AM until 6.30AM each day from 2011 to mid 2016, plotted against scaled time (`timeCount`). The 2016 data shown in blue is to be predicted 24 hours ahead, from a model fitted to the data shown in black.

(MAPE) of 1.1% on the fit data. For prediction of 2016 this increases to 3.0% for the standard approach, but only 1.6% for the model estimated using NCV. Nominally 95% prediction intervals have realized coverage of 65% for the standard approach, and 91% for the NCV based model.

In this case it is clear that the improved prediction performance relates to the smoother `timeCount` and `toy` effects. Careful model construction could have imposed this in an ad hoc manner, and achieved MAPE performance similar to the NCV fit, but confidence bands would still be neglecting the residual autocorrelation. In addition operational use of these models typically involves substantial automation. For example, a full day's prediction would involve 48 models of the type shown here,

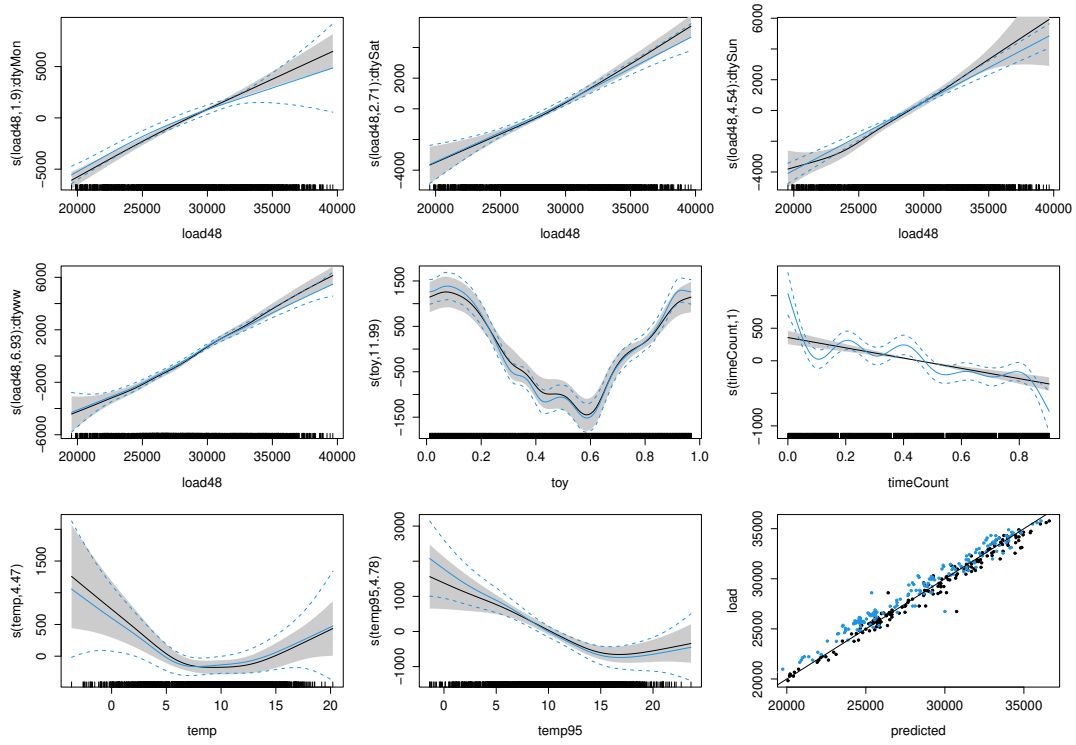


Figure 5: Estimated smooth effects for the grid load prediction models. For the first 8 plots from top left, the black curves and grey bands show estimates and 95% confidence bands from the NCV based model fit, while the superimposed blue curves show the equivalent for the marginal likelihood based fit, ignoring auto-correlation. The numbers in the y-axis labels are effective degrees of freedom. The plot at bottom right compares the models' predictive performance for the 2016 data. The black points show load versus predicted load for the NCV estimated model, and the blue points are for the marginal likelihood estimated model. The line illustrates where perfect prediction would lie.

which may in turn only be components of mixtures of experts, and which need to be regularly re-estimated as new training data accumulate. More local forecasting may involve thousands of such models. Hence the scope for careful ad hoc adjustment to deal with auto-correlation problems is limited, while the computational and implementational burden of fully specifying models for the auto-correlation appears prohibitive.

7.2 Swiss extreme rainfall

Now considering a smooth location-scale model for the most extreme 12 hour total rainfall each year from 65 Swiss weather stations from 1981 to 2015². After some model selection [52, §7.9.1], the data were modelled using a generalized extreme value distribution with location parameter model $\mu = f_1(E, N) + f_2(\text{nao}) + f_3(\text{elevation}) + \gamma_{\text{c.r.}}$; log scale parameter model $\rho = \log \sigma = f_4(E, N) + f_5(\text{year}) + f_6(\text{elevation}) + \beta_{\text{c.r.}}$ and modified logit shape parameter model $g(\xi) = \alpha_{\text{c.r.}}$. Function g is a logit link modified to constrain ξ between -1 and .5 (so variance remains finite and maximum likelihood estimation is consistent). The f_j are smooth functions, represented using thin plate regression splines. There is one parameter α_j , β_j and γ_j for each level of the 12 level factor **c.r.**, the climate region. Variable **nao** is the north Atlantic oscillation index. **N** and **E** are northing and easting in degrees. The station locations are shown as black points in the lower row of figure 6.

The model can be estimated using the methods of [54], with smoothing parameters estimated by Laplace approximate marginal likelihood. The resulting fit is shown in the left 2 columns of figure 6. However, estimation of GEV models is somewhat statistically and numerically taxing, and it is hence useful to be able to check results by re-estimating with an alternative smoothing parameter estimation criterion. Leave one out NCV offers this possibility, using the variant of section 3.1, and in this case give results visually almost identical to those obtained via marginal likelihood. An additional concern for these data is the possibility of localised short range spatial autocorrelation in the extremes, with the accompanying danger of over-fit. To investigate this NCV was used with neighbourhoods defined by stations being within 0.3 degrees of each other and year being the same (giving neighbourhood sizes from 1 to 6). The right 2 columns of figure 6 show the results, which are little changed, suggesting little effect of short range spatial correlation here. So in this case the new methods allow considerably more confidence about the practical reliability of the results than was previously possible.

²data are available as dataset **swr** in R package **gamair**.

7.3 Spatio-temporal models of forest health

The final example concerns spatio-temporal modelling of forest health data for Norway spruce trees in Baden-Württemberg, Germany. The data ($n = 5876$) were gathered annually from 1991-2020 on a semi-regular 4km $\times$ 4km spatial grid, but with coarser (8km $\times$ 8km or 16km $\times$ 16km) sub-grids used in some years. The nodes of the grid are tree stands, and the response variable of interest is average defoliation of trees at the stand, as a proportion between 0 and 1 (average number of spruce trees per stand is 17). The survey is in alignment with the International Co-operative Programme on Assessment and Monitoring of Air Pollution Effects on Forests and thus uses the same survey protocol [9]. There are reasons to expect some short range temporal autocorrelation at the stand level as spruce do not shed all their needles each year, and within year local spatial correlation might also be present. Simply neglecting this and treating the data as independent risks over fit and over-selection of model covariates. Neighbourhoods $\alpha(i)$ were hence constructed for each data point, consisting of the data from the same year and from a stand 10km or less away (on average 4 such neighbours per point), and from data less than 7 years apart at the same stand (average 5.6 such neighbours per point). The corresponding neighbourhood sizes range from 1 to 31 with a mean of 11.2.

Alongside spatial location and stand age (these are managed forests, so surveyed trees within a stand have very similar ages), a number of other covariates were measured, relating to the physical nature of the site (slope, elevation, aspect, wetness, soil composition, underlying geology, etc) and the tree properties (crown closure, level of fructification, forest diversity etc). Preliminary screening suggests some 16 covariates potentially of real interest (in addition to year and location). Generically the model for these data has the form

$$\mathbb{E}(\text{defoliation}_i) = f_0(\text{km.n}_i, \text{km.e}_i, \text{year}_i) + \sum_{j=1}^k f_j(x_{ji})$$

where x_j is a covariate and the f_j are smooth functions. f_0 is a smooth space-time interaction, represented by a tensor product of a thin plate spline of space with a cubic spline of year [2, 10]. The model can be estimated by penalized least squares with NCV smoothing parameter estimation, and a basic backward selection strategy used, which sequentially removes the least significant terms [using the approximate p-values of 51] until all are significant at the 10% level. This approach retains 10 covariates: tree age, previous year's mean temperature, precipitation evaporation index in May, topographic wetness index, topographic position index, exchangeably bound

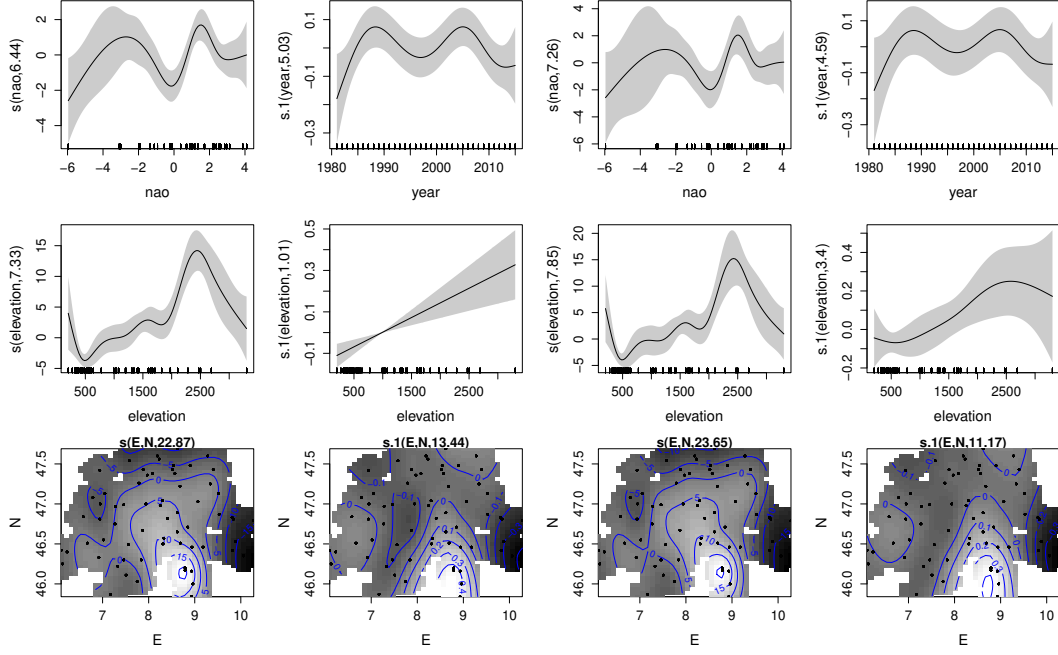


Figure 6: Swiss extreme rainfall location scale model. The left two columns are for Laplace approximate marginal likelihood based smoothing parameter estimation, the right two for NCV with neighbourhoods based on stations closer than 0.3 degrees. In each column pair the left column shows smooth effect estimates for the location parameter, and the right for the log scale parameter. In this case the estimates are very similar between the methods, with little evidence for un-modelled short range spatial correlation being important here. A fit using leave one out NCV is visually almost indistinguishable from the marginal likelihood estimates.

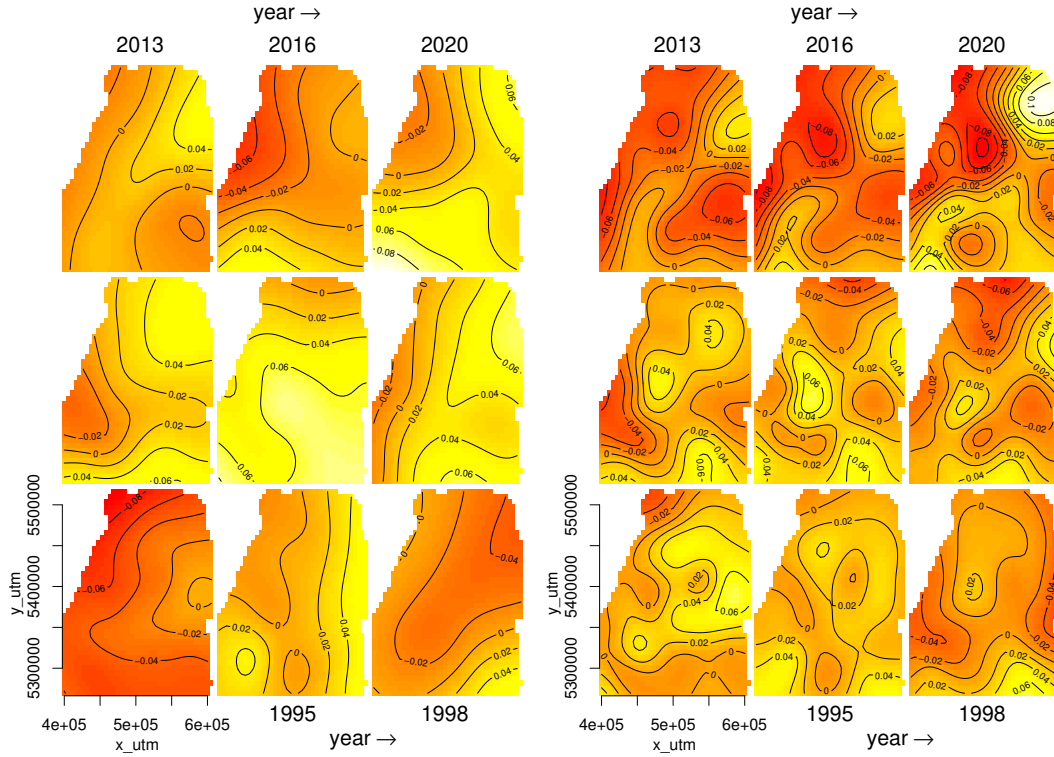


Figure 7: Forest health model space time interaction effect estimates. Left uses NCV accounting for the possibility of spatial and temporal autocorrelation as described in the text, and right uses REML assuming independence. Broadly, the NCV estimate is smoother in space, but less so in time.

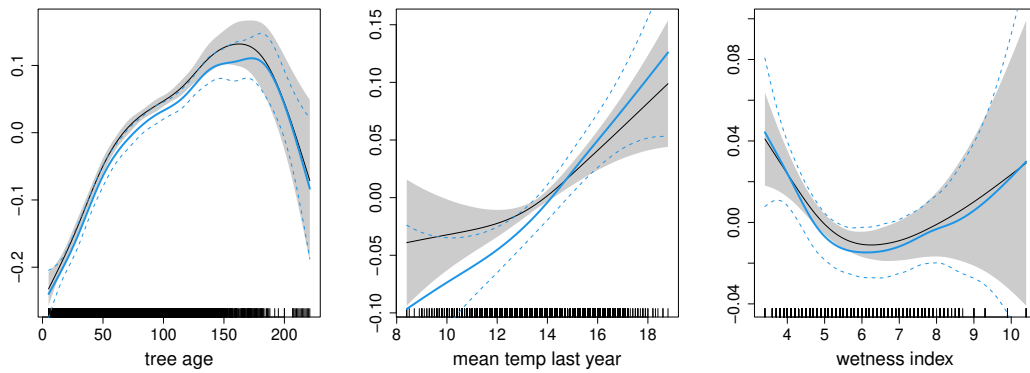


Figure 8: The three most interesting smooth covariate effects for forest defoliation estimated by NCV (black and grey band) compared to REML assuming independence (blue).

soil cations, days with max temperature $> 25^{\circ}\text{C}$ and 3 others that barely reach the inclusion threshold. By contrast, repeating the same exercise, but neglecting the possibility of auto-correlation and using REML smoothness selection, retains 13 covariates, with the May precipitation/evaporation index replaced by the same for September, and a measure of forest diversity, net evapotranspiration the preceding year, and summed nitrate over critical load.

Figure 7 compares the estimated spatio-temporal effects using NCV to those when the possibility of residual auto-correlation is ignored. The NCV estimates are smoother in space, but slightly less so in time. Figure 8 picks out the three most interesting smooth covariate effect estimates and compares them when estimation is by NCV and REML assuming independence. There is slight attenuation of the time related effects, whereas the topographic wetness index effect appears stable. The narrowing of credible intervals probably relates to the removal of other correlated covariates from the model. The NCV model has about 20 fewer effective degrees of freedom, and an r^2 about 1.4% lower than the alternative at just over 0.59. Clearly in this case there is evidence for short range residual auto-correlation in the data, as would be expected on a priori grounds, and this has some effect both on model estimates and the selection of model covariates.

8 Conclusions

NCV is both simple and general. It offers a principled method for choosing smoothing/precision parameters for models estimated with a variety of loss functions, in addition to negative log likelihoods and pseudo-likelihoods, and is also a model selection criterion [1]. In the case in which the the loss is a sum of negative log likelihood components for each datum it is applicable in all situations for which Laplace approximate marginal likelihood (LAML) might otherwise be used [see e.g. 54], but with the advantage over LAML of being directly useful for comparing models with differing fixed effects (unpenalized) structures. In any case, the availability of a second criterion for smoothing parameter selection in these cases offers a useful check of statistical stability in applications. The link to Jackknife estimation also leads to useful variance estimates.

One immediate application is to the estimation of smooth models in the presence of nuisance auto-correlation. Addressing short range auto-correlation without requiring formulation and estimation of a high rank correlation model can offer some application advantages, potentially allowing fuller exploration of the model elements of most direct scientific interest. NCV always of-

fers a way to check the sensitivity of model conclusions to the assumption of no unmodelled auto-correlation. However when such auto-correlation is present, it only offers mitigation when there is reason to expect that there is a separation of scales between the nuisance auto-correlation and the model components of interest. Otherwise there is no alternative but to fully model the correlation process. When the separation of scales assumption is reasonable, the question of whether better calibrated uncertainty quantification is possible for the NCV approach is interesting and open. A further open question is the feasibility of a general treatment of smoothing bias without requiring a Bayesian bias estimate via coherent belief updating.

The second immediate application is to non-likelihood based smooth loss functions. For example, in the smooth quantile regression case the approach offers a frequentist alternative to [12] without the need to subscribe to the Bayesian belief updating framework of [3], except for the limited purposes of accounting for smoothing bias in further inference. This less Bayesian approach aligns more readily with quantile regression’s avoidance of a full probabilistic model.

While the computational cost of NCV is the $O(np^2)$ of marginal likelihood or GCV, the constant of proportionality is higher and a disadvantage is that it is not obvious how to achieve the sort of scalability achieved for marginal likelihood by [53] and [25] in big data settings.

NCV is available in R package `mgcv`. Replication code is supplied as supplementary material.

Acknowledgements

I thank Heike Puhlmann and Simon Trust at the Forest Research Institute Baden-Württemberg (Germany) for making the Terrestrial Crown Condition Inventory (TCCI) forest health monitoring survey data available.

A Quantile regression

Following [12] the ELF loss targeting the τ quantile, q , is

$$\rho(y - \mu) = (\tau - 1) \frac{y - q}{\sigma} + \lambda \log \left(1 + e^{\frac{y - q}{\lambda \sigma}} \right).$$

This can be viewed as a smoothed version of the usual quantile regression pinball loss [23]. As [20] show, smoothing the loss generally results in lower MSE than using the unsmoothed loss, and the degree of smoothing can be chosen to optimize MSE, with the optimum depending on the number of data per model degree of freedom. A simple approach is to fit a pilot location scale model to the data, to obtain standardized residuals $z_i = (y_i - \hat{\mu}_i)/\hat{\sigma}_i$: let k be the estimated degrees

of freedom of the pilot mean model. Taking size $\lceil n/k \rceil$ bootstrap samples, the MSE of $\hat{q}$ based on optimization of the ELF loss for each sample can be computed, where the τ quantile of the z_i is the target truth. The λ minimizing this MSE is the estimate of the optimal degree of loss smoothing, and is used for fitting the full model, along with the $\hat{\sigma}_i$ from the pilot fit.

The approach of a pilot location-scale model fit and then computing the optimal amount of smoothing from the pooled pilot residuals follows [12], but the bootstrap step considerably simplifies the method for optimizing the smoothing of the loss, which is otherwise reliant on asymptotic results. Computationally, the same n_b bootstrap resamples are used with each trial λ . This ensures that the MSE is smooth in λ . The optimizing quantiles for each replicate can be found most efficiently by Newton optimization, starting from the true target quantile. The pilot fit consisted of a smooth model for the mean, using NCV, and then fitting a second smooth model to the resulting squared residuals to estimate $\hat{\sigma}_i^2$, again using NCV. A location scale model such as the `gauss` model in R package `mgcv` would be an alternative. When using this smoothed loss, an obvious check is that $\sum_i \mathbb{I}(y_i < q_i) \approx \tau$.

B Example likelihood derivatives

This appendix supplies examples of the loss specific derivatives required in practice to compute the NCV derivatives as described in section 2.2. First consider single parameter exponential family regression where $\mu_i = \mathbb{E}(y_i)$, $g(\mu_i) = \eta_i$, $\boldsymbol{\eta} = \mathbf{X}\boldsymbol{\beta}$ and $\text{var}(y_i) = V(\mu_i)\phi$, for link function g , variance function V and scale parameter ϕ (the loss is the negative log likelihood or the deviance). Then for, l_i , the contribution to the log likelihood from y_i we have

$$\frac{\partial l_i}{\partial \beta_j} = \frac{y_i - \mu_i}{\phi g'(\mu_i) V(\mu_i)} X_{ij}$$

and

$$\frac{\partial^2 l_i}{\partial \beta_j \partial \rho_q} = -\frac{X_{ij} \alpha(\mu_i)}{g'(\mu_i)^2 V(\mu_i)} \sum_k X_{ik} \frac{d\hat{\beta}_k}{d\rho_q}.$$

The latter term provides $d\mathbf{g}/d\rho_j$. More generally for other location parameter regression cases,

$$\frac{\partial l_i}{\partial \beta_j} = \frac{\partial l_i}{\partial \mu_i} \frac{d\mu_i}{d\eta_i} X_{ij}$$

and

$$\frac{\partial^2 l_i}{\partial \beta_j \partial \rho_q} = \left\{ \frac{\partial^2 l_i}{\partial \mu_i^2} \left(\frac{d\mu_i}{d\eta_i} \right)^2 + \frac{\partial l_i}{\partial \mu_i} \frac{d^2 \mu_i}{d\eta_i^2} \right\} X_{ij} \sum_k X_{ik} \frac{d\hat{\beta}_k}{d\rho_q}.$$

In both these cases the NCV criterion has the form $V = \sum_{k=1}^m \sum_{i \in \beta(k)} \mathcal{D}(y_i, \eta_i^{-\alpha(k)})$ so that

$$\frac{\partial V}{\partial \rho_j} = \sum_{k=1}^m \sum_{i \in \beta(k)} \frac{\partial \mathcal{D}}{\partial \eta_i} \Big|_{\eta_i^{-\alpha(i)}} \frac{\partial \eta_i^{-\alpha(i)}}{\partial \rho_j}$$

Some location parameter regressions may also depend on extra parameters, $\boldsymbol{\theta}$ (not varying with i). Then $V = \sum_{k=1}^m \sum_{i \in \beta(k)} \mathcal{D}(y_i, \eta_i^{-\alpha(i)}, \boldsymbol{\theta})$ so that

$$\frac{\partial V}{\partial \theta_j} = \sum_{k=1}^m \sum_{i \in \beta(k)} \frac{\partial \mathcal{D}}{\partial \eta_i} \Big|_{\eta_i^{-\alpha(i)}} \frac{\partial \eta_i^{-\alpha(i)}}{\partial \theta_j} + \frac{\partial \mathcal{D}}{\partial \theta_j} \Big|_{\eta_i^{-\alpha(i)}}$$

and

$$\begin{aligned} \frac{\partial^2 l_i}{\partial \beta_j \partial \theta_q} = & \left\{ \frac{\partial^2 l_i}{\partial \mu_i^2} \left(\frac{d\mu_i}{d\eta_i} \right)^2 + \frac{\partial l_i}{\partial \mu_i} \frac{d^2 \mu_i}{d\eta_i^2} \right\} X_{ij} \sum_k X_{ik} \frac{d\hat{\beta}_k}{d\theta_q} \\ & + \frac{\partial^2 l_i}{\partial \mu_i \partial \theta_q} \frac{d\mu_i}{d\eta_i} X_{ij}. \end{aligned}$$

The Hessian and its derivative w.r.t. ρ_j follow in a similar manner.

References

- [1] Arlot, S. and A. Celisse (2010). A survey of cross-validation procedures for model selection. *Statistics Surveys* 4, 40–79.
- [2] Augustin, N. H., M. Musio, K. von Wilpert, E. Kublin, S. N. Wood, and M. Schumacher (2009). Modeling spatiotemporal forest health monitoring data. *Journal of the American Statistical Association* 104(487), 899–911.
- [3] Bissiri, P. G., C. Holmes, and S. G. Walker (2016). A general framework for updating belief distributions. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 78(5), 1103–1130.
- [4] Chu, C.-K. and J. S. Marron (1991). Comparison of two bandwidth selectors with dependent errors. *The Annals of Statistics* 19(4), 1906–1918.
- [5] Claeskens, G., T. Krivobokova, and J. D. Opsomer (2009). Asymptotic properties of penalized spline estimators. *Biometrika* 96(3), 529–544.
- [6] Craven, P. and G. Wahba (1979). Smoothing noisy data with spline functions. *Numerische Mathematik* 31, 377–403.
- [7] Davison, A. C. and D. V. Hinkley (1997). *Bootstrap methods and their application*. Cambridge university press.

- [8] deHoog, F. and M. Hutchinson (1987). An efficient method for calculating smoothing splines using orthogonal transformations. *Numerische Mathematik* 50, 311–319.
- [9] Eichhorn, J., P. Roskams, N. Potočić, V. Timmermann, Ferretti, V. Mues, A. Szepesi, D. Durrant, I. Seletković, H-W.Schröck, S. Nevalainen, F. Bussotti, P. Garcia, and S. Wulff (Eds.) (2017). *ICP Forests manual on methods and criteria for harmonized sampling, assessment, monitoring and analysis of the effects of air pollution on forests*. Thünen Institute of Forest Ecosystems, Eberswalde, Germany.
- [10] Eickenscheidt, N., N. H. Augustin, and N. Wellbrock (2019). Spatio-temporal modelling of forest monitoring data: Modelling German tree defoliation data collected between 1989 and 2015 for trend estimation and survey grid examination using GAMMs. *iForest Biogeosciences and Forestry* 12, 338–348.
- [11] Eldén, L. (1984). A note on the computation of the generalized cross-validation function for ill-conditioned least squares problems. *BIT Numerical Mathematics* 24(4), 467–472.
- [12] Fasiolo, M., S. N. Wood, M. Zaffran, R. Nedellec, and Y. Goude (2021). Fast calibrated additive quantile regression. *Journal of the American Statistical Association* 116(535), 1402–1412.
- [13] Golub, G. H., M. Heath, and G. Wahba (1979). Generalized cross validation as a method for choosing a good ridge parameter. *Technometrics* 21(2), 215–223.
- [14] Golub, G. H. and C. F. van Loan (2013). *Matrix Computations* (4th ed.). Baltimore: Johns Hopkins University Press.
- [15] Gu, C. and D. Xiang (2001). Cross-validating non-gaussian data: Generalized approximate cross-validation revisited. *Journal of Computational and Graphical Statistics* 10(3), 581–591.
- [16] Harville, D. A. (1997). *Matrix Algebra from a Statistician’s Perspective*. New York: Springer.
- [17] Hastie, T. and R. Tibshirani (1990). *Generalized Additive Models*. Chapman & Hall.
- [18] Hastie, T., R. Tibshirani, and J. Friedman (2009). *The Elements of Statistical Learning*. Springer.
- [19] Hutchinson, M. F. and F. De Hoog (1985). Smoothing noisy data with spline functions. *Numerische Mathematik* 47(1), 99–106.
- [20] Kaplan, D. M. and Y. Sun (2017). Smoothed estimating equations for instrumental variables quantile regression. *Econometric Theory* 33(1), 105–157.
- [21] Klein, N., T. Kneib, S. Klasen, and S. Lang (2014). Bayesian structured additive distributional regression for multivariate responses. *Journal of the Royal Statistical Society: Series C (Applied Statistics)* 64, 569–591.
- [22] Klein, N., T. Kneib, S. Lang, and A. Sohn (2015). Bayesian structured additive distributional regression with an application to regional income inequality in Germany. *Annals of Applied Statistics* 9, 1024–1052.
- [23] Koenker, R. (2005). *Quantile Regression*. Cambridge University Press.
- [24] Krivobokova, T. and G. Kauermann (2007). A note on penalized spline smoothing with correlated errors. *Journal of the American Statistical Association* 102(480), 1328–1337.
- [25] Li, Z. and S. N. Wood (2020). Faster model matrix crossproducts for large generalized linear models with discretized covariates. *statistics and computing. Statistics and Computing*. 30(1), 19–25.
- [26] Lukas, M. A. (2006). Robust generalized cross-validation for choosing the regularization parameter. *Inverse Problems* 22(5), 1883.
- [27] Lukas, M. A. (2010). Robust GCV choice of the regularization parameter for correlated data. *The Journal of integral equations and applications*, 519–547.
- [28] Lukas, M. A., F. R. de Hoog, and R. S. Andersen (2016). Practical use of robust GCV and modified GCV for spline smoothing. *Computational Statistics* 31(1), 269–289.
- [29] Marra, G. and S. N. Wood (2012). Coverage properties of confidence intervals for generalized additive model components. *Scandinavian Journal of Statistics* 39(1), 53–74.
- [30] Mayr, A., N. Fenske, B. Hofner, T. Kneib, and M. Schmid (2012). Generalized additive models for location, scale and shape for high dimensional data — a flexible approach based on boosting. *Journal of the Royal Statistical Society: Series C (Applied Statistics)* 61(3), 403–427.
- [31] Nocedal, J. and S. Wright (2006). *Numerical Optimization* (2nd ed.). New York: Springer Verlag.
- [32] Nychka, D. (1988). Bayesian confidence intervals for smoothing splines. *Journal of the American Statistical Association* 83(404), 1134–1143.

- [33] Oh, H.-S., T. C. Lee, and D. W. Nychka (2011). Fast nonparametric quantile regression with arbitrary smoothing methods. *Journal of Computational and Graphical Statistics* 20(2), 510–526.
- [34] OpenMP Architecture Review Board (2008, May). OpenMP application program interface version 3.0.
- [35] Opsomer, J., Y. Wang, and Y. Yang (2001). Nonparametric regression with correlated errors. *Statistical Science*, 134–153.
- [36] Paige, C. C. and M. A. Saunders (1975). Solution of sparse indefinite systems of linear equations. *SIAM journal on numerical analysis* 12(4), 617–629.
- [37] Reiss, P. T. and L. Huang (2012). Smoothness selection for penalized quantile regression splines. *The international journal of biostatistics* 8(1).
- [38] Rigby, R. and D. M. Stasinopoulos (2005). Generalized additive models for location, scale and shape. *Journal of the Royal Statistical Society: Series C (Applied Statistics)* 54(3), 507–554.
- [39] Roberts, D. R., V. Bahn, S. Ciuti, M. S. Boyce, J. Elith, G. Guillera-Arroita, S. Hauenstein, J. J. Lahoz-Monfort, B. Schröder, W. Thuiller, et al. (2017). Cross-validation strategies for data with temporal, spatial, hierarchical, or phylogenetic structure. *Ecography* 40(8), 913–929.
- [40] Robinson, T. and R. Moyeed (1989). Making robust the cross-validators choice of smoothing parameter in spline smoothing regression. *Communications in Statistics-Theory and Methods* 18(2), 523–539.
- [41] Shi, X. (1988). A note on the delete-d jackknife variance estimators. *Statistics & probability letters* 6(5), 341–347.
- [42] Stasinopoulos, M. D., R. A. Rigby, G. Z. Heller, V. Voudouris, and F. De Bastiani (2017). *Flexible Regression and Smoothing: using GAMLSS in R*. Chapman and Hall/CRC.
- [43] Stephenson, W. and T. Broderick (2020). Approximate cross-validation in high dimensions with guarantees. In *International Conference on Artificial Intelligence and Statistics*, pp. 2424–2434. PMLR.
- [44] Stone, M. (1974). Cross-validators choice and assessment of statistical predictions (with discussion). *Journal of the Royal Statistical Society, Series B* 36, 111–147.
- [45] Stone, M. (1977). An asymptotic equivalence of choice of model by cross-validation and Akaike’s criterion. *Journal of the Royal Statistical Society, Series B* 39, 44–47.
- [46] van der Linde, A. (2000). *Smoothing and Regression*, Chapter Variance Estimation and Smoothing-Parameter Selection for Spline Regression, pp. 19–42. Wiley.
- [47] van der Vorst, H. A. (2003). *Iterative Krylov Methods for Large Linear Systems*. Cambridge.
- [48] Wahba, G. (1985). A comparison of GCV and GML for choosing the smoothing parameter in the generalized spline smoothing problem. *Annals of Statistics* 13(4), 1378–1402.
- [49] Wood, S. N. (2008). Fast stable direct fitting and smoothness selection for generalized additive models. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 70(3), 495–518.
- [50] Wood, S. N. (2011). Fast stable restricted maximum likelihood and marginal likelihood estimation of semi-parametric generalized linear models. *Journal of the Royal Statistical Society, B* 73(1), 3–36.
- [51] Wood, S. N. (2013). On p-values for smooth components of an extended generalized additive model. *Biometrika* 100(1), 221–228.
- [52] Wood, S. N. (2017). *Generalized Additive Models: An Introduction with R* (2 ed.). Boca Raton, FL: CRC press.
- [53] Wood, S. N., Z. Li, G. Shaddick, and N. H. Augustin (2017). Generalized additive models for gigadata: modelling the UK black smoke network daily data. *Journal of the American Statistical Association* 112(519), 1199–1210.
- [54] Wood, S. N., N. Pya, and B. Säfken (2016). Smoothing parameter and model selection for general smooth models (with discussion). *Journal of the American Statistical Association* 111, 1548–1575.
- [55] Xianyi, Z., W. Qian, and Z. Chothia (2014). OpenBLAS. URL: <http://xianyi.github.io/OpenBLAS>.
- [56] Yee, T. W. (2015). *Vector generalized linear and additive models: with an implementation in R*. Springer.
- [57] Yee, T. W. and C. Wild (1996). Vector generalized additive models. *Journal of the Royal Statistical Society. Series B (Methodological)* 58(3), 481–493.