

The linearisation of maps in data assimilation

By TIMOTHY J. PAYNE*, *Met Office, Exeter, EX1 3PB, United Kingdom*

(Manuscript received 24 May 2012; in final form 11 March 2013)

ABSTRACT

For the purpose of linearising maps in data assimilation, the tangent-linear approximation is often used. We compare this with the use of the ‘best linear’ approximation, which is the linear map that minimises the mean square error. As a benchmark, we use minimum variance filters and smoothers which are non-linear generalisations of Kalman filters and smoothers. We show that use of the best linear approximation leads to a filter whose prior has first moment unapproximated compared with the benchmark, and second moment whose departure from the benchmark is bounded independently of the derivative of the map, with similar results for smoothers. This is particularly advantageous where the maps in question are strongly non-linear on the scale of the increments. Furthermore, the best linear approximation works equally well for maps which are non-differentiable. We illustrate the results with examples using low-dimensional chaotic maps.

Keywords: tangent-linear approximation, incremental 4D-Var, extended Kalman smoother

1. Introduction

The methods currently used for large-scale data assimilation in the earth sciences can be viewed as approximate implementations of the Kalman filter or Kalman smoother, and their variational equivalents. These are only optimal for linear forecast and observation operators, while the forecast maps encountered in the earth sciences are non-linear.

One important class of methodologies, including the extended Kalman filter and smoother and incremental four-dimensional variational assimilation (4D-Var), attempts to get round this by linearisation. In practice, this has nearly always meant approximating each map $\mathbf{f}(\mathbf{x} + \delta)$ by its first-order Taylor polynomial $\mathbf{f}(\mathbf{x}) + \mathbf{f}'(\mathbf{x})\delta$, the so-called ‘tangent-linear’ (TL) approximation. If the distribution of the perturbations δ is known, we may form instead the ‘best linear’ (BL) approximation $\mathbf{f}(\mathbf{x} + \delta) \approx \tilde{\mathbf{f}} + \tilde{\mathbf{T}}\delta$, in which $\tilde{\mathbf{f}}$, $\tilde{\mathbf{T}}$ are chosen to minimise the mean square error in this approximation. We will designate these approximations the TLA and BLA, respectively. In this article we consider the merits of the BLA in data assimilation.

We begin in Section 2 by giving some of the properties of the BLA. Versions of the same idea have been considered by others, such as by Gelb (1974) who call it ‘statistical linearisation’, and by Lorenc and Payne (2007) where it is termed ‘optimal regularisation’. The reason for this

latter terminology will be apparent in Section 3, where we explore the properties of the BLA through one application, the regularisation of a parametrisation of grid-box cloud fraction.

The main focus in this article is on the use of the BLA in the data assimilation-forecast cycle where the non-linearity is in the forecast map. The case where the observation operators are also non-linear is considered briefly in an appendix.

For optimal filtering, to form the prior pdf at the next time level the posterior pdf from the former time level is evolved forward in time by the Chapman–Kolmogorov equation [e.g. Jaswinski (1970)]. If the forecast map is linear and the posterior pdf is Gaussian, this is exactly accomplished by the predictive part of the Kalman filter. For non-linear maps, we should use the Chapman–Kolmogorov equation in its full form, but a popular approach in practice is to apply the Kalman filter to a linearised map, giving rise to the extended Kalman filter (EKF). Similarly, for the Kalman smoother (and its variational equivalent, 4D-Var) which is the optimal smoother for linear maps and Gaussian pdfs, a common approach is to linearise the non-linear map to form the (generally suboptimal) extended Kalman smoother (EKS) and incremental 4D-Var.

Normally for the EKF, EKS and incremental 4D-Var the TLA is used. In Lorenc (1997), the limitations of the TLA for 4D-Var were recognised and Lorenc coined the term ‘perturbation forecast’ (PF) model to describe a linear model appropriate for the evolution of finite perturbations. Errico and Raeder (1999) also recognised that it is the

*Correspondence.
email: tim.payne@metoffice.gov.uk

accuracy of the linear model in evolving finite amplitude perturbations which is important for applications, and assessed how well the tangent-linear to some meteorological models performed in this regard. Lawless et al. (2005) considered specifically the non-tangent-linear model arising from discretising a linearisation of the continuous non-linear equations (rather than linearising a discretisation, as in a TL model), from the point of view of showing that the use of such a model would not necessarily degrade the assimilation compared with a TL model.

It has long been accepted when linearising components of non-linear models for adjoint applications such as incremental 4D-Var, that physical parametrisations in particular may need to be ‘regularised’ rather than the TL taken directly. Lorenc and Payne (2007) used the BLA to construct an optimal regularisation, and suggested applying the principle to the whole model rather than merely individual components. The present work follows this up analytically, and thereby provides a justification for the ‘PF not TL’ philosophy in cycled data assimilation.

In order to develop and compare TLA and BLA based EKF, we begin in Section 4 by using the same approach we used to find the BLA to maps to derive a generalisation of the Kalman filter for non-linear propagating maps. Specifically, we find the analysis linear in the observations which minimises mean square analysis error. This is equivalent (see Appendix A.1) to specialising the optimal filter obtained from Bayes’ rule and the Chapman–Kolmogorov equation by approximating where necessary each pdf by a Gaussian one with the same first two moments. The resulting prediction step has been used before as the basis for non-linear generalisations of the Kalman filter, such as the ‘unscented Kalman filter’ of Julier and Uhlmann (1997).

In Section 5 the non-linear filter of Section 4 is approximated by the TL and BL approximations to form TLA and BLA based EKF. We show that the latter can be expected to outperform the former. In Section 6, we illustrate the methods and the results by comparing the TLA and BLA in filters assimilating noisy observations of states of the Duffing map, a chaotic non-linear map which discretises the Duffing oscillator.

In Sections 7–9 we treat smoothers in an analogous way to the treatment of filters in Sections 4–6. In Section 7, we develop a generalisation of the Kalman Smoother for non-linear propagating maps, and use it in Section 8 to develop TLA and BLA based EKS. These are of particular interest as they are equivalent to TLA and BLA based versions of incremental 4D-Var. Again we show the latter can be expected to outperform the former, and give examples in Section 9.

2. Best linear approximation

We begin by formulating and obtaining the best linear approximation, and comparing with the tangent-linear approximation. Suppose we have a function $\mathbf{f}$ for which we seek a local linear approximation

$$\mathbf{f}(\mathbf{x} + \boldsymbol{\delta}) \approx \tilde{\mathbf{f}} + \tilde{\mathbf{T}}\boldsymbol{\delta}$$

where $\tilde{\mathbf{f}}$ is a vector and $\tilde{\mathbf{T}}$ a matrix. If $\mathbf{f}$ is C^1 , we might choose $\tilde{\mathbf{f}} = \mathbf{f}(\mathbf{x})$, $\tilde{\mathbf{T}} = \mathbf{f}'(\mathbf{x})$, so the local linear approximation to $\mathbf{f}(\mathbf{x} + \boldsymbol{\delta})$ is the first-order Taylor polynomial $\mathbf{f}(\mathbf{x}) + \mathbf{f}'(\mathbf{x})\boldsymbol{\delta}$. This is the tangent-linear approximation or TLA, usually encountered in the adjoint applications literature in the form $\mathbf{f}(\mathbf{x} + \boldsymbol{\delta}) - \mathbf{f}(\mathbf{x}) \approx \mathbf{f}'(\mathbf{x})\boldsymbol{\delta}$.

However, if we know the pdf of $\boldsymbol{\delta}$ we may seek the ‘best local linear approximation’ in the sense of minimising

$$\mathcal{I} = E \left[\left(\mathbf{f}(\mathbf{x} + \boldsymbol{\delta}) - \tilde{\mathbf{f}} - \tilde{\mathbf{T}}\boldsymbol{\delta} \right)^T A \left(\mathbf{f}(\mathbf{x} + \boldsymbol{\delta}) - \tilde{\mathbf{f}} - \tilde{\mathbf{T}}\boldsymbol{\delta} \right) \right] \quad (1)$$

over $\tilde{\mathbf{f}}$ and $\tilde{\mathbf{T}}$ for a given symmetric positive-definite matrix A . This is termed ‘statistical linearisation’ in Gelb (1974) and, for reasons discussed in Section 3 below, ‘optimal regularisation’ in Lorenc and Payne (2007).

If $\mathcal{I}$ is minimised by $\tilde{\mathbf{f}} = \mathbf{f}^{(0)}(\mathbf{x})$ and $\tilde{\mathbf{T}} = \mathbf{f}^{(1)}(\mathbf{x})$, then differentiating $\mathcal{I}$ with respect to $\tilde{\mathbf{f}}$ and $\tilde{\mathbf{T}}$ and setting the derivatives to zero we obtain respectively the vector and matrix equations:

$$\begin{aligned} E[\mathbf{f}(\mathbf{x} + \boldsymbol{\delta}) - \mathbf{f}^{(0)}(\mathbf{x}) - \mathbf{f}^{(1)}(\mathbf{x})\boldsymbol{\delta}] &= \mathbf{0} \\ E[(\mathbf{f}(\mathbf{x} + \boldsymbol{\delta}) - \mathbf{f}^{(0)}(\mathbf{x}) - \mathbf{f}^{(1)}(\mathbf{x})\boldsymbol{\delta})\boldsymbol{\delta}^T] &= \mathbf{0}. \end{aligned} \quad (2)$$

Solving the simultaneous equations [eq. (2)] we see that $\mathbf{f}^{(0)}(\mathbf{x})$ and $\mathbf{f}^{(1)}(\mathbf{x})$ are independent of A and given by

$$\begin{aligned} \mathbf{f}^{(1)}(\mathbf{x}) &= (E[\mathbf{f}(\mathbf{x} + \boldsymbol{\delta})\boldsymbol{\delta}^T] - E[\mathbf{f}(\mathbf{x} + \boldsymbol{\delta})]E[\boldsymbol{\delta}]^T) \\ &\quad \times (E[\boldsymbol{\delta}\boldsymbol{\delta}^T] - E[\boldsymbol{\delta}]E[\boldsymbol{\delta}]^T)^{-1} \end{aligned} \quad (3)$$

$$\mathbf{f}^{(0)}(\mathbf{x}) = E[\mathbf{f}(\mathbf{x} + \boldsymbol{\delta}) - \mathbf{f}^{(1)}(\mathbf{x})\boldsymbol{\delta}]. \quad (4)$$

So the linear approximation to $\mathbf{f}(\mathbf{x} + \boldsymbol{\delta})$ which minimises mean square linearisation error, and which we call the best linear approximation or BLA, is $\mathbf{f}^{(0)}(\mathbf{x}) + \mathbf{f}^{(1)}(\mathbf{x})\boldsymbol{\delta}$. Noting that if this approximation is used then linearisation error itself is

$$\boldsymbol{\ell}^{BL} \equiv \mathbf{f}(\mathbf{x} + \boldsymbol{\delta}) - \mathbf{f}^{(0)}(\mathbf{x}) - \mathbf{f}^{(1)}(\mathbf{x})\boldsymbol{\delta} \quad (5)$$

it follows that eq. (2) may be written

$$E[\boldsymbol{\ell}^{BL}] = \mathbf{0}, \quad E[\boldsymbol{\ell}^{BL}\boldsymbol{\delta}^T] = \mathbf{0}. \quad (6)$$

The coefficients $\mathbf{f}^{(0)}(\mathbf{x})$ and $\mathbf{f}^{(1)}(\mathbf{x})$ in the BLA share with the coefficients in the TLA the property of being linear as a function of $\mathbf{f}$, i.e. if $\mathbf{h} = \alpha\mathbf{f} + \beta\mathbf{g}$ then for $j = 0, 1$

$$\mathbf{h}^{(j)}(\mathbf{x}) = \alpha\mathbf{f}^{(j)}(\mathbf{x}) + \beta\mathbf{g}^{(j)}(\mathbf{x})$$

which can be useful for constructing BLAs.

Some obvious but important points of difference or comparison between the BLA and TLA are:

(a) Whereas the coefficients of the TLA are a function only of $\mathbf{f}$, those of the BLA are a function of $\mathbf{f}$ and the pdf of the increments.

(b) In forming the BLA, unlike the TLA, $\mathbf{f}$ does not need to be differentiable. The BLA captures the local form of $\mathbf{f}$ on the scale of the perturbations δ , rather than the instantaneous derivative. This can be an important distinction where $\mathbf{f}$ contains more than one scale or contains sudden changes in derivative, as for example is often the case with physical parametrisations.

(c) In the limit that the perturbations become small compared with the scale of the non-linearity, for example if $\mathbf{f}$ is C^1 and the distribution of δ tends to a Dirac delta function, then $\mathbf{f}^{(0)}(\mathbf{x}) \rightarrow \mathbf{f}(\mathbf{x})$ and $\mathbf{f}^{(1)}(\mathbf{x}) \rightarrow \mathbf{f}'(\mathbf{x})$. Thus the TLA is appropriate for infinitesimal perturbations.

These points are illustrated by the example which follows.

3. Example of best linear approximation: regularisation of grid-box cloud fraction

As an example, consider the parametrisation of grid-box cloud fraction Φ as a function of total water q_t , liquid temperature T_L and pressure p given in Smith (1990):

$$\Phi = \begin{cases} 0 & Q_N \leq -1 \\ \frac{1}{2}(1 + Q_N)^2 & -1 < Q_N \leq 0 \\ 1 - \frac{1}{2}(1 - Q_N)^2 & 0 < Q_N < 1 \\ 1 & 1 \leq Q_N \end{cases} \quad (7)$$

where

$$Q_N = \frac{q_t - q_{sat}(T_L, p)}{(1 - RH_c)q_{sat}(T_L, p)} \quad (8)$$

where q_{sat} is the saturation specific humidity and RH_c is a constant. Suppose we require a linear approximation

$$\Phi(Q_N + \delta Q_N) \approx \Phi_0(Q_N) + \Phi_1(Q_N)\delta Q_N \quad (9)$$

to this function, e.g. for use in the linear model in incremental 4D-Var. Note that if $Q_N \leq -1$ or $Q_N \geq 1$ then $d\Phi/dQ_N = 0$, so the TLA $\Phi(Q_N + \delta Q_N) \approx \Phi(Q_N) + \Phi'(Q_N)\delta Q_N$ will simply return $\Phi(Q_N)$ for Q_N in that range. Traditionally such functions have been ‘regularised’, i.e. replaced by slightly smoothed versions whose first-order Taylor polynomial better approximates the non-linear function. In this case, a regularisation which has been used is to replace Φ by Φ_{reg}

$$\Phi \approx \Phi_{reg} = \frac{1}{2}[1 + \tanh(2Q_N)]. \quad (10)$$

Clearly this regularisation is ad hoc; there are any number of other smooth functions which could have been chosen to

approximate Φ . However, if we know the distribution of δQ_N then we see the BLA of Section 2 constitutes an ‘optimal’ regularisation, in the sense of minimising mean square linearisation error.

In the 4D-Var context δQ_N are, at convergence of the 4D-Var minimisation, analysis increments, and we can accumulate statistics on them from past analyses. The simplest way to then form $\Phi^{(0)}$ and $\Phi^{(1)}$ is by discretising eqs. (3) and (4) directly (i.e. replace the integrals defining the expectations by finite summations using the empirical data). However, to illustrate some features of the BLA we will approximate the empirical distribution by an analytic one. The *generalised normal distribution* is a family of distributions which for a one-dimensional independent variable x with mean x_0 have pdf

$$f_p(x) = \frac{p^{1-\frac{1}{p}}}{2\sigma\Gamma(\frac{1}{p})} \exp\left\{-\frac{1}{p} \frac{|x - x_0|^p}{\sigma^p}\right\} \quad (11)$$

where p, σ are parameters. In the Gaussian or standard normal distribution, $p = 2$. We find empirically that δQ_N is approximately distributed with a generalised normal distribution with $p = 1$, zero mean and $\sigma \approx 0.4$, independently of Q_N .

We consider four linearisations [eq. (9)] of the cloud function [eq. (7)] about Q_N , which will be denoted (i)–(iv). For (i), we use the TLA, so $\Phi_0(Q_N) = \Phi(Q_N)$ and $\Phi_1(Q_N) = \Phi'(Q_N)$. For (ii)–(iv), we use the BLA of $\Phi(Q_N)$ calculated for δQ_N distributed with a range of pdfs. In detail, we will in each case compute the BLA for δQ_N distributed with a generalised normal distribution [eq. (11)] with $p = 1$, and for (ii) $\sigma = 0.8$, for (iii) $\sigma = 0.4$ (the correct value), and for (iv) $\sigma = 0.1$. Fig. 1 shows

(a) The generalised normal distribution with $p = 1$ for $\sigma = 0.8$ (blue), $\sigma = 0.4$ (green) and $\sigma = 0.1$ (red);

(b) The constant term $\Phi_0(Q_N)$ in eq. (9) computed according to the methods described above. The black curve uses the TLA as in (i), so is simply $\Phi(Q_N)$. The other curves use the BLA with $\sigma = 0.8$ (blue), $\sigma = 0.4$ (green), $\sigma = 0.1$ (red).

(c) The coefficient $\Phi_1(Q_N)$ in eq. (9) computed according to the methods described above, same colour scheme as (b).

(d) For the last plot we suppose δQ_N has a generalised normal distribution [eq. (11)] with $p = 1$ and $\sigma = 0.4$, and compute the mean square linearisation error using the linearisations (i)–(iv). Note that in this scalar case mean square linearisation error is the function $\mathcal{I}$ of eq. (1) with $A = 1$.

3.1. Remarks

- (1) As noted above, as $\sigma \rightarrow 0$ we see [Fig. 1b] that $\Phi^{(0)}$ converges to Φ and [Fig. 1c] that $\Phi^{(1)}$ converges to Φ' .

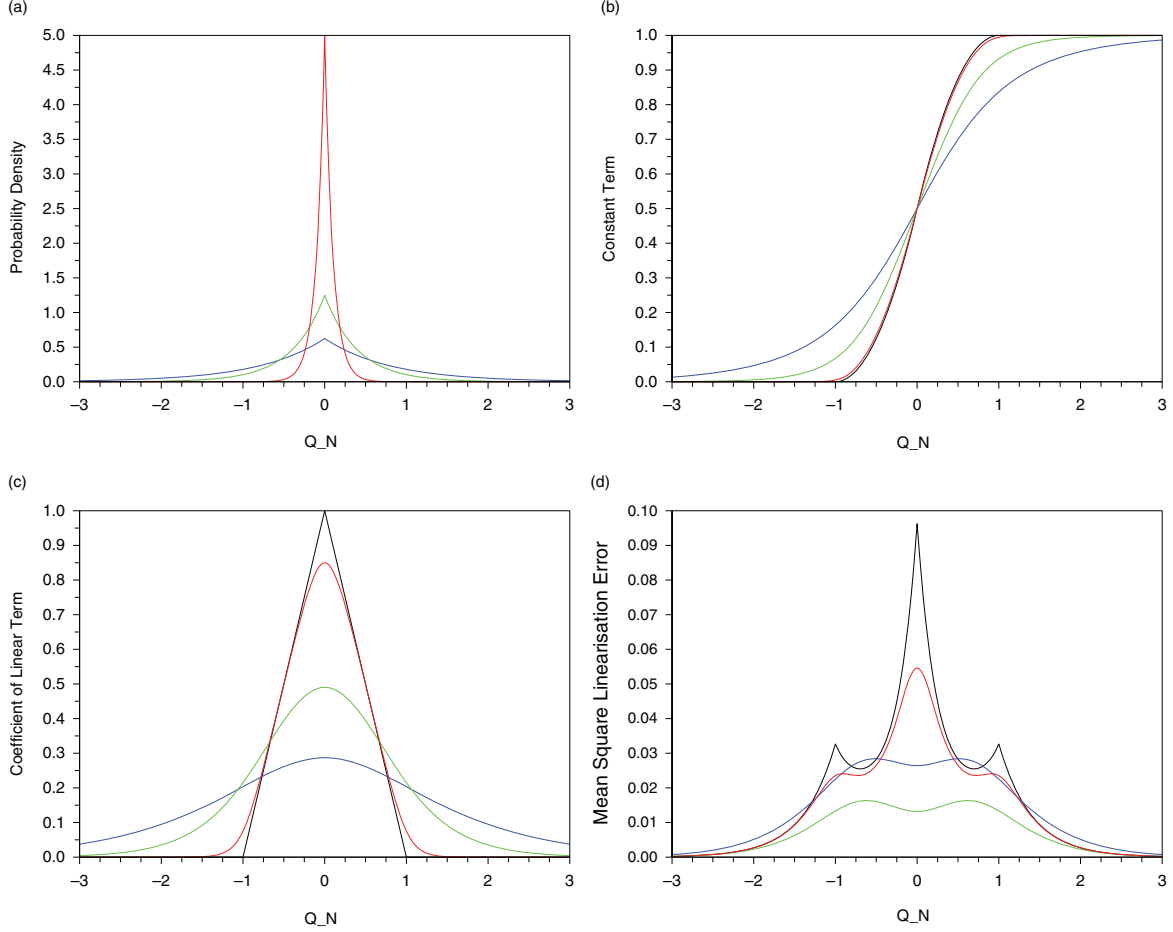


Fig. 1. (a) pdf of generalised normal distribution (11) with $p=1$ for $\sigma=0.8$ (blue), $\sigma=0.4$ (green) and $\sigma=0.1$ (red); (b) cloud fraction Φ (black) and $\Phi^{(0)}$ for various σ as in (a); (c) Φ' (black) and $\Phi^{(1)}$ for various σ as in (a); (d) linearisation error assuming true pdf has $\sigma=0.4$, using the TLA (black) and the BLA with various σ as in (a).

- (2) With conventional regularisation Φ is approximated by smooth Φ_{reg} and one would use $\Phi_0(Q_N) = \Phi_{reg}$ and $\Phi_1(Q_N) = \Phi_{reg}'$, so $\Phi_1(Q_N) = \Phi_0'(Q_N)$. However for the BLA, in general $\Phi^{(1)}$ is not the derivative of $\Phi^{(0)}$.
- (3) $\Phi^{(0)}$ is formed by taking a weighted average of nearby values of Φ , and as σ increases greater distances are given greater weights, so for fixed Q_N we have $\Phi^{(0)}(Q_N) \rightarrow \frac{1}{2}$ as $\sigma \rightarrow \infty$ [Fig. 1b]. Similarly, for fixed Q_N we have $\Phi^{(1)}(Q_N) \rightarrow 0$ as $\sigma \rightarrow \infty$ [Fig. 1c].
- (4) The mean square linearisation error using the BLA and the correct pdf (green curve in Fig. 1d) is always less and mostly a small fraction of that using the TLA [black curve in Fig. 1d]. The regularisation $\Phi \approx \frac{1}{2}[1 + \tanh(2Q_N)]$ is typical in that it smooths Φ but is everywhere close to Φ in the C^0 norm, and actually affords very little benefit compared with the TLA – its linearisation error is nearly as large as that of the tangent linear.

4. A non-linear generalisation of the Kalman Filter

The main purpose of this article is to examine and improve the use of linearised maps in data assimilation. To do this we need a non-linear generalisation of the predictive step of the Kalman filter, which we can approximate using the TLA and BLA, and then use as a benchmark against which to compare the impact of the two approximations.

The generalisation we use is obtained by applying the BLA approach to the state estimation problem in the forecast-assimilation cycle. This is equivalent to the familiar linear minimum variance estimator of, for example, Sage and Melsa (1971: Section 6.6). To keep the article self-contained and highlight the connection with the BLA we derive this filter in the same way as we did the BLA. The resulting generalisation of the prediction stage of the Kalman filter [eq. (20) below] has been used previously, for filtering polynomial systems in Luca et al. (2006), and

notably in Julier and Uhlmann (1997) as the basis for the unscented Kalman filter.

Suppose we have states $\mathbf{x}_i$, noisy observations $\mathbf{y}_i$ of these states and a noisy model $\mathbf{f}_i$ which evolves the state from time level i to $i+1$. Denote our best estimate or analysis of $\mathbf{x}_i$ given observations $\mathbf{y}_1, \dots, \mathbf{y}_i$ by $\hat{\mathbf{x}}_{i|i}$. We suppose we have an analysis $\hat{\mathbf{x}}_{i|i}$ at time i , and we seek to use an observation at time $i+1$ to form an analysis at $i+1$. In summary, we have

$$\begin{aligned}\mathbf{x}_i &= \hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i \\ \mathbf{x}_{i+1} &= \mathbf{f}_i(\mathbf{x}_i) + \mathbf{w}_i \\ \mathbf{y}_{i+1} &= \mathbf{x}_{i+1} + \mathbf{v}_{i+1}\end{aligned}\quad (12)$$

where $\boldsymbol{\epsilon}_i$, $\mathbf{w}_i$, $\mathbf{v}_{i+1}$ are random variables, the pdfs of which are supposed known. At this stage we do not require these pdfs to be Gaussian, though for simplicity we shall suppose that $\boldsymbol{\epsilon}_i$, $\mathbf{w}_i$, $\mathbf{v}_{i+1}$ have zero mean and are uncorrelated with each other and with all other quantities, though weaker hypotheses are possible. The covariance of the analysis error $\boldsymbol{\epsilon}_i$ is denoted $P_{i|i}$.

In the same spirit as the BLA for maps we may seek the best linear analysis $\mathbf{x}_{i+1}^*$, that is, we seek K_{i+1} , $\mathbf{c}_{i+1}$ depending on $\hat{\mathbf{x}}_{i|i}$ and the pdfs of $\boldsymbol{\epsilon}_i$, $\mathbf{w}_i$, $\mathbf{v}_{i+1}$ such that if $\mathbf{x}_{i+1}^*$ is of the form

$$\mathbf{x}_{i+1}^* = K_{i+1}\mathbf{y}_{i+1} + \mathbf{c}_{i+1} \quad (13)$$

then

$$\mathcal{E} = E[(\mathbf{x}_{i+1}^* - \mathbf{x}_{i+1})^T A(\mathbf{x}_{i+1}^* - \mathbf{x}_{i+1})] \quad (14)$$

is minimised, where the expectation is over $\boldsymbol{\epsilon}_i$, $\mathbf{w}_i$, $\mathbf{v}_{i+1}$. As before, the solution is independent of A . Set

$$E[\mathbf{v}_{i+1}\mathbf{v}_{i+1}^T] = R_{i+1}, \quad E[\mathbf{w}_i\mathbf{w}_i^T] = Q_i. \quad (15)$$

Setting the derivatives of eq. (14) with respect to $\mathbf{c}_{i+1}$, K_{i+1} to zero, solving the resulting simultaneous equations and using eq. (12) to express all expectations in terms of $\boldsymbol{\epsilon}_i$, $\mathbf{w}_i$, $\mathbf{v}_{i+1}$, we obtain after some manipulation

$$K_{i+1} = P_{i+1|i}(P_{i+1|i} + R_{i+1})^{-1} \quad (16)$$

$$\mathbf{c}_{i+1} = (I - K_{i+1})\hat{\mathbf{x}}_{i+1|i} \quad (17)$$

where

$$\begin{aligned}\hat{\mathbf{x}}_{i+1|i} &= E[\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i)] = E[\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) + \mathbf{w}_i] \\ P_{i+1|i} &= E\left[\left(\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) - E[\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i)]\right) \right. \\ &\quad \left. \times \left(\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) - E[\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i)]\right)^T\right] + Q_i\end{aligned}\quad (18)$$

the expectations in eq. (18) being over $\boldsymbol{\epsilon}_i$. The $\hat{\mathbf{x}}_{i+1|i}$, $P_{i+1|i}$ notation is consistent with standard filtering theory (see Appendix A.1).

None of this requires Gaussianity of $\boldsymbol{\epsilon}_i$, $\mathbf{w}_i$, $\mathbf{v}_{i+1}$. However, while the pdfs of $\mathbf{w}_i$, $\mathbf{v}_{i+1}$ are externally given, the pdf of $\boldsymbol{\epsilon}_i$ (which we need to compute the expectations) is the analysis error from the previous time level. By direct calculation, we find the covariance of $\boldsymbol{\epsilon}_{i+1} = \mathbf{x}_{i+1} - \mathbf{x}_{i+1}^*$ to be

$$(I - K_{i+1})P_{i+1|i}. \quad (19)$$

To close the system we will approximate the pdf of $\boldsymbol{\epsilon}_{i+1}$ as a Gaussian with zero mean and covariance [eq. (19)], and use this as the pdf at the next time level.

Putting this all together we obtain what we will term the *minimum variance filter* given by eqs. (20) and (21) below, which generalises the Kalman filter in the case where the propagating map $\mathbf{f}_i$ is non-linear. The notation $\boldsymbol{\epsilon}_i \sim N(\mathbf{0}, P_{i|i})$ signifies that $\boldsymbol{\epsilon}_i$ is normally distributed with zero mean and covariance $P_{i|i}$. We show in Appendix A.1 how for Gaussian $\mathbf{v}_{i+1}$ and $\mathbf{w}_i$ we obtain the same algorithm from Bayes' rule and the Chapman–Kolmogorov equation if we approximate the prior pdf by a Gaussian with the same first two moments.

For $i=0,1,\dots$

Predict

$$\begin{aligned}\hat{\mathbf{x}}_{i+1|i} &= E[\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i)] \\ P_{i+1|i} &= E\left[\left(\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) - \hat{\mathbf{x}}_{i+1|i}\right) \right. \\ &\quad \left. \times \left(\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) - \hat{\mathbf{x}}_{i+1|i}\right)^T\right] + Q_i\end{aligned}\quad (20)$$

where $\boldsymbol{\epsilon}_i \sim N(\mathbf{0}, P_{i|i})$

Update:

$$\begin{aligned}K_{i+1} &= P_{i+1|i}(P_{i+1|i} + R_{i+1})^{-1} \\ \hat{\mathbf{x}}_{i+1|i+1} &= \hat{\mathbf{x}}_{i+1|i} + K_{i+1}(\mathbf{y}_{i+1} - \hat{\mathbf{x}}_{i+1|i}) \\ P_{i+1|i+1} &= (I - K_{i+1})P_{i+1|i}.\end{aligned}\quad (21)$$

Because our observation operator is linear (in fact, here the identity) the update step eq. (21) is unchanged from the standard Kalman filter. As noted above, the predict stage eq. (20) has been used before as the basis for non-linear generalisations of the Kalman filter, such as in Julier and Uhlmann (1997) as the basis for the unscented Kalman filter.

If $\mathbf{f}_i$ is linear then eq. (20) simplifies to the usual Kalman filter, and if $\mathbf{f}_i$ is non-linear and we replace it by its TL approximation we obtain the usual form of the EKF, a point expanded upon in Section 5 below.

Readers familiar with ensemble filters may be interested in the relationship between the minimum variance filter, eqs. (20) and (21), and an ensemble Kalman Filter (EnKF) with infinite ensemble size. Assuming the EnKF is cycled in the usual way it differs in the following respect. We may choose the perturbed analyses used in the initial cycle of the EnKF to be distributed as $N(\hat{\mathbf{x}}_{0|0}, P_{0|0})$, but on cycles after

the first the perturbed analyses in the EnKF will only have a Gaussian distribution if the forecast map is linear. This distinguishes it from the minimum variance filter [eq. (20)] which always uses Gaussian analysis errors. The two filters would coincide (in the limit of an infinite ensemble size for the EnKF) if the EnKF was modified so that on every cycle its analysis ensemble was formed as a random sample drawn from a Gaussian distribution with covariance equal to the usual EnKF analysis covariance.

5. Linearisation stratagems for the prediction part of the update-prediction cycle

For small systems one might be able to compute the covariance in eq. (20) exactly, and for large systems it could be approximated by a small ensemble. In the standard form of the EKF, $\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i)$ is replaced [in both parts of eq. (20)] by its TL approximation. Here we shall compare this approximation with the BL approximation. By examining the impact of these approximations on eq. (20), we can compare the impact of making the TLA and BLA on the quality of the data assimilation system.

If $\mathbf{f}_i$ is C^1 and we apply Taylor's theorem to expand it about the analysis $\hat{\mathbf{x}}_{i|i}$, we have

$$\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) = \mathbf{f}_i(\hat{\mathbf{x}}_{i|i}) + \mathbf{f}_i'(\hat{\mathbf{x}}_{i|i})\boldsymbol{\epsilon}_i + \boldsymbol{\ell}_i^{TL} \quad (22)$$

where $\boldsymbol{\ell}_i^{TL}$ is the error in the TL approximation. Substituting eq. (22) into eq. (20), we get

$$\begin{aligned} \hat{\mathbf{x}}_{i+1|i} &= \mathbf{f}_i(\hat{\mathbf{x}}_{i|i}) + E[\boldsymbol{\ell}_i^{TL}] \\ P_{i+1|i} &= \mathbf{f}_i'(\hat{\mathbf{x}}_{i|i})P_{i|i}\mathbf{f}_i'(\hat{\mathbf{x}}_{i|i})^T + Q_i^{TL} + Q_i \end{aligned} \quad (23)$$

where Q_i^{TL} is the error in $P_{i+1|i}$ arising through the use of the TLA:

$$Q_i^{TL} = E[\mathbf{r}_i\mathbf{r}_i^T] + \mathbf{f}_i'(\hat{\mathbf{x}}_{i|i})E[\boldsymbol{\epsilon}_i\mathbf{r}_i^T] + E[\mathbf{r}_i\boldsymbol{\epsilon}_i^T]\mathbf{f}_i'(\hat{\mathbf{x}}_{i|i})^T \quad (24)$$

where $\mathbf{r}_i = \boldsymbol{\ell}_i^{TL} - E[\boldsymbol{\ell}_i^{TL}]$.

If we make an expansion based on the BLA of $\mathbf{f}_i$ about the analysis $\hat{\mathbf{x}}_{i|i}$

$$\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) = \mathbf{f}_i^{(0)}(\hat{\mathbf{x}}_{i|i}) + \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})\boldsymbol{\epsilon}_i + \boldsymbol{\ell}_i^{BL} \quad (25)$$

where $\boldsymbol{\ell}_i^{BL}$ is the error in the BL approximation, and substitute eq. (25) into eq. (20), we get

$$\begin{aligned} \hat{\mathbf{x}}_{i+1|i} &= \mathbf{f}_i^{(0)}(\hat{\mathbf{x}}_{i|i}) \\ P_{i+1|i} &= \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})P_{i|i}\mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})^T + Q_i^{BL} + Q_i. \end{aligned} \quad (26)$$

From eq. (6) we know that $E[\boldsymbol{\ell}_i^{BL}] = \mathbf{0}$, so does not appear in eq. (26). The term Q_i^{BL} representing the error in $P_{i+1|i}$ if

the BLA is made has some very interesting properties. We have

$$\begin{aligned} Q_i^{BL} &= E\left[\mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})\boldsymbol{\epsilon}_i(\boldsymbol{\ell}_i^{BL} - E[\boldsymbol{\ell}_i^{BL}])^T \right. \\ &\quad \left. + (\boldsymbol{\ell}_i^{BL} - E[\boldsymbol{\ell}_i^{BL}])\boldsymbol{\epsilon}_i^T\mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})^T \right. \\ &\quad \left. + (\boldsymbol{\ell}_i^{BL} - E[\boldsymbol{\ell}_i^{BL}])(\boldsymbol{\ell}_i^{BL} - E[\boldsymbol{\ell}_i^{BL}])^T\right] \end{aligned} \quad (27)$$

but by virtue of eq. (6) and the fact $\boldsymbol{\epsilon}_i$ has zero mean

$$E[\boldsymbol{\ell}_i^{BL}] = \mathbf{0}, E[\mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})\boldsymbol{\epsilon}_i(\boldsymbol{\ell}_i^{BL} - E[\boldsymbol{\ell}_i^{BL}])^T] = \mathbf{0} \quad (28)$$

so

$$Q_i^{BL} = E[\boldsymbol{\ell}_i^{BL}(\boldsymbol{\ell}_i^{BL})^T]. \quad (29)$$

In the light of the foregoing several key observations may be made:

(a) While Q_i^{TL} in eq. (24) is merely an error, and may be positive-definite, negative-definite or indefinite, Q_i^{BL} is a covariance; specifically, it is the covariance of linearisation error.

(b) As such Q_i^{BL} is positive semi-definite. Therefore if it is entirely neglected and for the prior covariance $P_{i+1|i}$ we use $\mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})P_{i|i}\mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})^T + Q_i$ then we *always underestimate the true prior covariance*.

(c) As noted above, if the BLA is used there is no approximation in the prior $\hat{\mathbf{x}}_{i+1|i}$ [first line of eq. (26)] compared with the minimum variance filter of Section 4.

In view of its formulation we also expect the BLA to lead to an improved prior covariance $P_{i+1|i}$ than the TLA, particularly when on the scale of the increments $\mathbf{f}_i$ is strongly non-linear. Specifically we have the following:

(d) While Q_i^{TL} in eq. (24) is a function of the pointwise derivative of $\mathbf{f}_i$ at $\hat{\mathbf{x}}_{i|i}$, which may be very large and erratic, Q_i^{BL} is bounded independently of $\mathbf{f}_i'$ as follows. Q_i^{BL} is positive semi-definite, so all its eigenvalues are non-negative, but the sum of its eigenvalues equals the trace of Q_i^{BL} . Using the definitions of BLA and $\boldsymbol{\ell}_i^{BL}$, and the fact that for any matrices A, B $\text{trace}(AB) = \text{trace}(BA)$, we have

$$\begin{aligned} \text{trace}(Q_i^{BL}) &= \text{trace}(E[\boldsymbol{\ell}_i^{BL}(\boldsymbol{\ell}_i^{BL})^T]) = \\ &= E[(\boldsymbol{\ell}_i^{BL})^T\boldsymbol{\ell}_i^{BL}] = \\ E\left[\left\|\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) - \mathbf{f}_i^{(0)}(\hat{\mathbf{x}}_{i|i}) - \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})\boldsymbol{\epsilon}_i\right\|^2\right] &= \\ \min_{\mathbf{f}, \mathbf{T}} E\left[\left\|\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) - \mathbf{f} - \mathbf{T}\boldsymbol{\epsilon}_i\right\|^2\right] \end{aligned} \quad (30)$$

where the expectation is over $\boldsymbol{\epsilon}_i$. Thus, the eigenvalues of Q_i^{BL} are all non-negative, but their sum is bounded by the *minimum possible mean square linearisation error*.

5.1. Filters using the TLA and BLA

At one extreme, if we compute the linearisation error matrices $\tilde{Q}_i^{TL}$ and $\tilde{Q}_i^{BL}$ exactly as defined by eqs. (24) and (29) we reproduce the minimum variance filter of Section 4 above. At the other extreme, if we set them to zero we obtain in the TLA case the standard EKF. However, as noted in Fisher et al. (2005), in this case analysis errors are very large. Even though they were working in a perfect model environment with $Q_i = 0$, they included a diagonal Q_i to allow for linearisation error, see especially their Fig. 1. Effectively they relabelled $\tilde{Q}_i^{TL}$ as Q_i and then parametrised it by one scalar.

We will follow a similar strategy and allow for simple parametrisations of $\tilde{Q}_i^{TL}$ and $\tilde{Q}_i^{BL}$. Fisher et al. (2005) showed this was important in the TL case; we anticipate it will be at least as important in the BL case, since as we noted in (b) above, if we neglect $\tilde{Q}_i^{BL}$ then the prior covariance obtained using the BLA is always underestimated.

With these considerations in view, the predict part of the filter using the TLA is:

$$\begin{aligned}\hat{\mathbf{x}}_{i+1|i} &= \mathbf{f}_i(\hat{\mathbf{x}}_{i|i}) \\ P_{i+1|i} &= \mathbf{f}_i'(\hat{\mathbf{x}}_{i|i}) P_{i|i} \mathbf{f}_i'^T(\hat{\mathbf{x}}_{i|i}) + \tilde{Q}_i^{TL} + Q_i\end{aligned}\quad (31)$$

and using the BLA is

$$\begin{aligned}\hat{\mathbf{x}}_{i+1|i} &= \mathbf{f}_i^{(0)}(\hat{\mathbf{x}}_{i|i}) \\ P_{i+1|i} &= \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i}) P_{i|i} \mathbf{f}_i^{(1)T}(\hat{\mathbf{x}}_{i|i}) + \tilde{Q}_i^{BL} + Q_i\end{aligned}\quad (32)$$

where $\tilde{Q}_i^{TL}$, $\tilde{Q}_i^{BL}$ are simple parametrisations of linearisation error. The update part of both filters is unchanged from the minimum variance filter listed in eq. (21) above.

The advantages of eq. (32) over eq. (31) were noted in (c) and (d) above. The relative importance of (c) and (d) is discussed in Section 10.

We make some additional comments on the BLA filter in Appendix A.2. Its existence is guaranteed under very mild conditions on $\mathbf{f}_i$ and the other quantities involved, as shown in Appendix A.2.1. Another filter, which like the BLA filter can be derived from the minimum variance filter [eqs. (20) and (21)] is the unscented Kalman filter (UKF) of Julier and Uhlmann (1997); we compare the BLA filter with the UKF in Appendix A.2.2. Finally, while for simplicity throughout the body of this article we only consider the case where the observation operator is linear (in fact the identity), we show in Appendix A.2.3 how to extend the BLA filter to cater for non-linear observation operators.

6. Example of filters using the TL and BL approximation

To illustrate the impact of different linearisation strategies for the predictive stage of cycled data assimilation we use the Duffing map [e.g. Saha and Tehri (2010)], which is a discretisation of the Duffing oscillator, and has a cubic non-linearity:

$$\mathbf{g}\begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} x_2 \\ -bx_1 + ax_2 - x_2^3 \end{pmatrix}. \quad (33)$$

The system has up to three fixed points along the line $x_1 = x_2$: at $x_1 = 0$ and $x_1^2 + 1 + b - a = 0$. For a range of choices of a , b , the system is chaotic, e.g. for $a = 2.75$, $b = 0.15$ [eq. (33)] has a leading Lyapunov exponent of 0.615. The first 100 000 iterates of the system with these parameters (started from $(\frac{1}{2}, \frac{1}{2})$) are shown in Fig. 2. Closer scrutiny shows the stable manifold to have an intricate structure which Fig. 2 can only partially reveal.

Set D to be the linear part of the Duffing map, i.e.

$$D = \begin{pmatrix} 0 & 1 \\ -b & a \end{pmatrix}$$

so $\mathbf{g}(\mathbf{x}) = D\mathbf{x} - x_2^3 \mathbf{e}_2$ where $\mathbf{e}_2$ is the unit vector $(0, 1)^T$. If δ is distributed normally with mean zero and variance C , then

$$\begin{aligned}\mathbf{g}(\mathbf{x} + \delta) &= D(\mathbf{x} + \delta) - (x_2^3 + 3x_2^2\delta_2 + 3x_2\delta_2^2 + \delta_2^3)\mathbf{e}_2 \\ \mathbf{g}'(\mathbf{x}) &= D - 3x_2^2\mathbf{e}_2\mathbf{e}_2^T \\ \mathbf{g}^{(0)}(\mathbf{x}) &= D\mathbf{x} - (x_2^3 + 3C_{22}x_2)\mathbf{e}_2 \\ \mathbf{g}^{(1)}(\mathbf{x}) &= D - (3x_2^2 + 3C_{22})\mathbf{e}_2\mathbf{e}_2^T.\end{aligned}\quad (34)$$

Using these relations we can write down the predictive parts of the TLA filter [eq. (31)] and of the BLA filter [eq. (32)], and with more work also compute exactly the

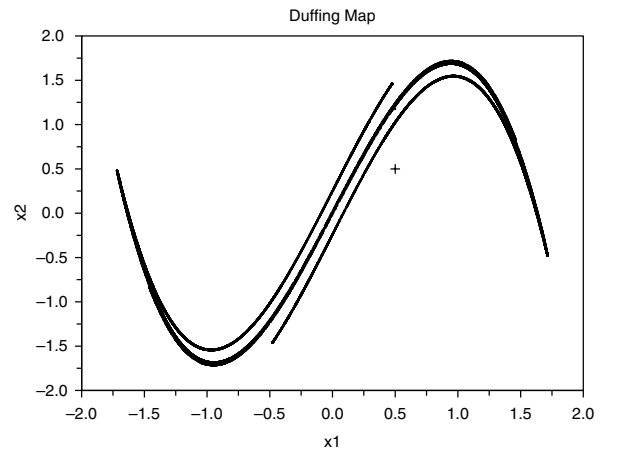


Fig. 2. 100 000 iterates of the Duffing map (33) with $a = 2.75$, $b = 0.15$ starting from $(\frac{1}{2}, \frac{1}{2})$ (plus sign).

predictive part of the minimum variance filter [eq. (20)]. We find that

$$Q_i^{BL} = \begin{pmatrix} 0 & 0 \\ 0 & 6(P_{i|i})_{22}^2(3(\hat{\mathbf{x}}_{i|i})_2^2 + (P_{i|i})_{22}) \end{pmatrix} \quad (35)$$

where since $P_{i|i}$ is positive semi-definite we have $(P_{i|i})_{22} \geq 0$, and therefore as expected Q_i^{BL} is positive semi-definite. Q_i^{TL} is a more complicated symmetric matrix with only the (1, 1) element zero. In view of this, we will set

$$\tilde{Q}_i^{BL} = \begin{pmatrix} 0 & 0 \\ 0 & \alpha \end{pmatrix}, \quad \tilde{Q}_i^{TL} = \begin{pmatrix} 0 & \beta \\ \beta & \alpha \end{pmatrix}$$

for real constants β, α .

6.1. Numerical results

We consider the Duffing system eq. (33) with $a=2.75$, $b=0.15$ as the truth. The observations are normally distributed about the truth with covariance $R = \sigma_o^2 I$ where $\sigma_o=0.3$. There is no model error so $Q_i=0$. For each assimilation method, the square analysis error $|\hat{\mathbf{x}}_{i|i} - \mathbf{x}_i|^2$ is computed for each i and averaged over 100 000 consecutive cycles.

If the minimum variance filter [eqs. (20) and (21)] is used, the mean square analysis error is 0.082. The mean square analysis error for various α (and $\beta=0$) using the TLA filter [eqs. (31) and (21)] is shown in solid black and using the BLA filter [eqs. (32) and (21)] in blue in Fig. 3. We have noted that Q_i^{TL} has non-zero off-diagonal elements, but in fact the dependence of the TLA filter on β is very slight – if for each α we find the mean square analysis error for the

TLA filter *minimised over all real β* we obtain the black dotted curve in Fig. 3.

6.2. Remarks

As expected, for a wide range of α the BLA filter has lower analysis errors than the TLA one. When α is large $P_{i+1|i}$ in eqs. (31) and (32) is dominated by the large $\tilde{Q}_i^{TL}$ or $\tilde{Q}_i^{BL}$, the analysis errors become large and increasingly independent of whether the TLA or BLA is used to evolve the posterior covariance from the former time level. In fact, denoting the analyses obtained using eqs. (31) and (32) as, respectively, $\hat{\mathbf{x}}_{i|i}^{TL}$, $\hat{\mathbf{x}}_{i|i}^{BL}$, one can show analytically that starting from the same data at $i=0$ then as $\alpha \rightarrow \infty$ that $\hat{\mathbf{x}}_{i|i}^{BL} - \hat{\mathbf{x}}_{i|i}^{TL} \rightarrow \mathbf{0}$ for all i .

For small α the analysis errors also become large, as noted in Fisher et al. (2005) and above. Analysis errors are smallest for intermediate values of α , and in fact for $\alpha \approx 0.07$ the BLA filter has an analysis error very close to that achieved by the minimum variance filter [eqs. (20) and (21)].

Note the log scale for α . As $\alpha \rightarrow 0$ analysis errors using both the TLA and BLA become very large, and in the limit there is no guarantee that the BLA filter will have smaller errors than the TLA one, and in fact does not for the parameters $a=2.75$, $b=0.15$, $\sigma_o=0.3$ as used in Fig. 3.

7. A non-linear generalisation of the Kalman smoother

We will now extend the theory of Sections 4 and 5 from filters to smoothers. As is well known, the EKS is equivalent to a form of 4D-Var, a method of particular interest as it is much the most popular for data assimilation in weather forecasting. For the links between Kalman smoothing and 4D-Var, see Bell (1994).

4D-Var in its simplest form has a fixed prior covariance, often designated B . If one runs 4D-Var this way the choice of B tends to be a large determinant of the quality of the assimilation, rendering comparison of other aspects of the system problematic. To compare the impact of the TL and BL approximations on incremental 4D-Var we therefore prefer to do so in the context of (the variational equivalent to) the EKS.

As with the filters, we first develop a minimum variance smoother from which TLA and BLA smoothers will (in Section 8 below) be derived by making the corresponding linearisations, and which will serve as a benchmark with which to assess the impact of the approximations.

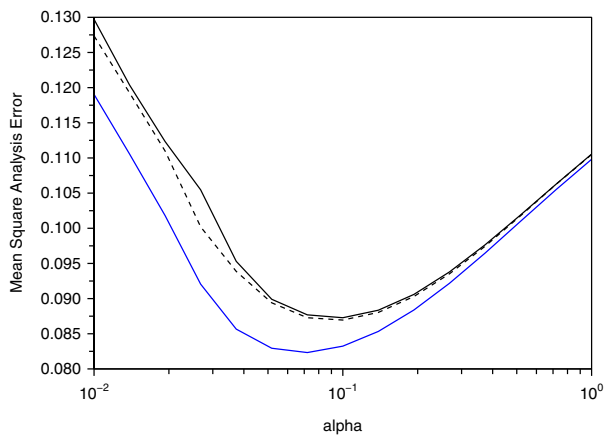


Fig. 3. Mean square analysis error for the Duffing map example, using TLA filter (black) and BLA filter (blue) for various values of the linearisation error parameter α . The dotted curve is the lowest mean square analysis error achievable using the TLA and any value of β .

7.1. Minimum variance smoother

The framework is identical to that in Section 4, but now instead of using the observation at time $i+1$ to find an analysis at time $i+1$, we are using it to improve the existing analysis $\hat{\mathbf{x}}_{i|i}$ at time i . Using exactly the same notation as Section 4, we now seek $L_i, \mathbf{d}_i$ depending on $\hat{\mathbf{x}}_{i|i}$ and the pdfs of $\epsilon_i, \mathbf{w}_i, \mathbf{v}_{i+1}$ such that if the smoothed analysis $\hat{\mathbf{x}}_{i|i+1}$ is of the form

$$\hat{\mathbf{x}}_{i|i+1} = L_i \mathbf{y}_{i+1} + \mathbf{d}_i \quad (36)$$

then

$$\mathcal{E} = E[(\hat{\mathbf{x}}_{i|i+1} - \mathbf{x}_i)^T A (\hat{\mathbf{x}}_{i|i+1} - \mathbf{x}_i)] \quad (37)$$

is minimised, where the expectation is over $\epsilon_i, \mathbf{w}_i, \mathbf{v}_{i+1}$. As usual, the solution will be independent of A .

The derivation follows the same pattern as Section 4. We suppose the same error structure and definitions as eq. (12), where as before the errors are zero mean and uncorrelated with each other, and again define R_{i+1}, Q_i by eq. (15). Setting the derivatives of eq. (37) with respect to $\mathbf{d}_i, L_i$ to zero, solving the resulting simultaneous equations and using eq. (12) to express all expectations in terms of $\epsilon_i, \mathbf{w}_i, \mathbf{v}_{i+1}$, we obtain after some manipulation

$$L_i = U_i (P_{i+1|i} + R_{i+1})^{-1} \quad (38)$$

$$\mathbf{d}_i = \hat{\mathbf{x}}_{i|i} - L_i E[\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \epsilon_i)] \quad (39)$$

where $\hat{\mathbf{x}}_{i+1|i}$ and $P_{i+1|i}$ have the same definitions as in eq. (18), and

$$U_i = E \left[\epsilon_i \left(\mathbf{f}_i'(\hat{\mathbf{x}}_{i|i} + \epsilon_i) \right)^T \right]. \quad (40)$$

We can easily compute the covariance of the error in the *smoothed* analysis $\hat{\mathbf{x}}_{i|i+1}$ to be

$$P_{i|i+1} = P_{i|i} - U_i (P_{i+1|i} + R_{i+1})^{-1} U_i^T \quad (41)$$

which evidently is less than the old one $P_{i|i}$.

Inserting the smoothing stage we have just derived between the prediction and update steps of the minimum variance filter [eqs. (20) and (21)], we obtain the minimum variance smoother of Table 1, where all the expectations are over $\epsilon_i \sim N(\mathbf{0}, P_{i|i})$. This algorithm generalises the Kalman smoother in the case the propagating map $\mathbf{f}_i$ is non-linear. If $\mathbf{f}_i$ is linear, then it simplifies to the usual Kalman smoother, and if $\mathbf{f}_i$ is non-linear and we replace $\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \epsilon_i)$ by its TL approximation we obtain the usual form of the EKS, a point we allude to in Section 8 below.

Table 1. Minimum variance smoother

Given first guesses $\hat{\mathbf{x}}_{0|0}, P_{0|0}$:

For $i=0, 1, \dots$

$$\begin{aligned} \text{Predict:} \quad & \hat{\mathbf{x}}_{i+1|i} = E \left[\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \epsilon_i) \right] \\ & P_{i+1|i} = E \left[\left(\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \epsilon_i) - \hat{\mathbf{x}}_{i+1|i} \right) \right. \\ & \quad \left. \times \left(\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \epsilon_i) - \hat{\mathbf{x}}_{i+1|i} \right)^T \right] + Q_i \\ \text{Smooth:} \quad & U_i = E \left[\epsilon_i \left(\mathbf{f}_i'(\hat{\mathbf{x}}_{i|i} + \epsilon_i) \right)^T \right] \\ & L_i = U_i (P_{i+1|i} + R_{i+1})^{-1} \\ & \hat{\mathbf{x}}_{i|i+1} = \hat{\mathbf{x}}_{i|i} + L_i (\mathbf{y}_{i+1} - \hat{\mathbf{x}}_{i+1|i}) \\ & P_{i|i+1} = P_{i|i} - U_i (P_{i+1|i} + R_{i+1})^{-1} U_i^T \\ \text{Update:} \quad & K_{i+1} = P_{i+1|i} (P_{i+1|i} + R_{i+1})^{-1} \\ & \hat{\mathbf{x}}_{i+1|i+1} = \hat{\mathbf{x}}_{i+1|i} + K_{i+1} (\mathbf{y}_{i+1} - \hat{\mathbf{x}}_{i+1|i}) \\ & P_{i+1|i+1} = (I - K_{i+1}) P_{i+1|i} \end{aligned}$$

End For i

8. Form of smoother where the forecast map is linearised using the TLA or BLA and variational equivalents

In the same way that in Section 5 we derived TLA and BLA filters from the minimum variance filter, we will now derive TLA and BLA smoothers from the minimum variance smoother of Table 1, by replacing $\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \epsilon_i)$ wherever it appears in Table 1 by expansions based on the TL and BL approximations. By examining the impact of these approximations on $\hat{\mathbf{x}}_{i+1|i}, P_{i+1|i}$ and U_i in the minimum variance smoother, we can compare the impact of making the TLA and BLA on the quality of the data assimilation system.

If in Table 1 we replace $\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \epsilon_i)$ by its first-order Taylor expansion [eq. (22)], then $\hat{\mathbf{x}}_{i+1|i}$ and $P_{i+1|i}$ are as given by eq. (23). The only new equation is

$$\begin{aligned} U_i &= E \left[\epsilon_i \left(\mathbf{f}_i'(\hat{\mathbf{x}}_{i|i}) \epsilon_i + \ell_i^{TL} \right)^T \right] \\ &= P_{i|i} \mathbf{f}_i'(\hat{\mathbf{x}}_{i|i})^T + E \left[\epsilon_i (\ell_i^{TL})^T \right]. \end{aligned} \quad (42)$$

Similarly, an expansion based on the BLA [eq. (25)] yields $\hat{\mathbf{x}}_{i+1|i}$ and $P_{i+1|i}$ as given by eq. (26), the only new equation is

$$U_i = E \left[\epsilon_i \left(\mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i}) \epsilon_i + \ell_i^{BL} \right)^T \right] = P_{i|i} \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})^T. \quad (43)$$

Interestingly, there is no term $E[\epsilon_i (\ell_i^{BL})^T]$ on the right hand side of eq. (43), as by virtue of eq. (6) it vanishes. Thus, in the smoother case use of the BLA has not only the advantages over the TLA listed in (c) and (d) of Section 5, but removes all approximation in U_i .

In summary, the smoother using the BLA is as given by Table 2. The smoother using the TLA is the same with

$\mathbf{f}_i^{(0)}(\hat{\mathbf{x}}_{i|i})$, $\mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})$, $\tilde{Q}_i^{BL}$ replaced by respectively $\mathbf{f}_i(\hat{\mathbf{x}}_{i|i})$, $\mathbf{f}'_i(\hat{\mathbf{x}}_{i|i})$, $\tilde{Q}_i^{TL}$.

We can identify three advantages of the BLA over the TLA for the smoother. Like the BLA filter in Section 5, use of the BLA leads to no approximation in the new prior $\hat{\mathbf{x}}_{i+1|i}$ compared with the minimum variance smoother, and the approximation to the prior covariance $P_{i+1|i}$ is not dependent on the instantaneous derivative. Furthermore, use of the BLA has the additional advantage that U_i is unapproximated compared with the minimum variance smoother.

8.1. Variational equivalents of the smoother

As discussed at the beginning of Section 7, a key motivation behind extending our study from filters to smoothers is the latter's well known equivalence to forms of 4D-Var (see e.g. Bell (1994)). The non-linear smoother of Section 7 has no immediate variational equivalent, but its TLA and BLA linearisations derived above do, at least so long as $Q_i + \tilde{Q}_i^{TL}$, $Q_i + \tilde{Q}_i^{BL}$ are invertible. Set

$$\begin{aligned} J(\delta_i, \delta_{i+1}) = & \frac{1}{2} \delta_i^T P_{i|i}^{-1} \delta_i \\ & + \frac{1}{2} (\mathbf{y}_{i+1} - \mathbf{f}_i^{(0)}(\hat{\mathbf{x}}_{i|i}) - \delta_{i+1})^T R_{i+1}^{-1} (\mathbf{y}_{i+1} - \mathbf{f}_i^{(0)}(\hat{\mathbf{x}}_{i|i}) - \delta_{i+1}) \\ & + \frac{1}{2} (\delta_{i+1} - \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i}) \delta_i)^T (Q_i + \tilde{Q}_i^{BL})^{-1} (\delta_{i+1} - \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i}) \delta_i). \end{aligned} \quad (44)$$

If the state vector is of length n the Hessian (i.e. second derivative) of J in eq. (44) is a matrix of size $2n \times 2n$. By computing its inverse in block form, one may show that the BLA smoother of Table 2 above is identical to the algorithm shown in Table 3, where $[\cdot]_{11}$ and $[\cdot]_{22}$ refer respectively to the upper left and lower right $n \times n$ submatrices. Note that in this version of incremental 4D-Var the background error covariance is cycled, instead of the static background error covariance used in early implementations.

Clearly the variational equivalent of the TLA smoother is as given by eq. (44) and Table 3 with $\mathbf{f}_i^{(0)}(\hat{\mathbf{x}}_{i|i})$, $\mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})$, $\tilde{Q}_i^{BL}$ replaced by respectively $\mathbf{f}_i(\hat{\mathbf{x}}_{i|i})$, $\mathbf{f}'_i(\hat{\mathbf{x}}_{i|i})$, $\tilde{Q}_i^{TL}$.

The matrix $Q_i + \tilde{Q}_i^{BL}$ in the last line of eq. (44) has appeared frequently in this article, as is to be expected given that from eqs. (12) and (25) we have $\mathbf{x}_{i+1} = \mathbf{f}_i^{(0)}(\hat{\mathbf{x}}_{i|i}) + \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})(\mathbf{x}_i - \hat{\mathbf{x}}_{i|i}) + \ell_i^{BL} + \mathbf{w}_i$. As linearisation error ℓ_i^{BL} and model error $\mathbf{w}_i$ are supposed uncorrelated the covariance of their sum is $Q_i + Q_i^{BL}$, which is parametrised as $Q_i + \tilde{Q}_i^{BL}$. One notable difference between the two errors making up this sum is that model error is a property solely of the model, while linearisation error is a property of the whole forecast-analysis system.

Table 2. Smoother using the BLA

Given first guesses $\hat{\mathbf{x}}_{0|0}$, $P_{0|0}$:

For $i = 0, 1, \dots$

$$\begin{aligned} \text{Predict:} \quad & \hat{\mathbf{x}}_{i+1|i} = \mathbf{f}_i^{(0)}(\hat{\mathbf{x}}_{i|i}) \\ & P_{i+1|i} = \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i}) P_{i|i} \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})^T + Q_i + \tilde{Q}_i^{BL} \\ \text{Smooth:} \quad & U_i = P_{i|i} \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i})^T \\ & L_i = U_i (P_{i+1|i} + R_{i+1})^{-1} \\ & \hat{\mathbf{x}}_{i|i+1} = \hat{\mathbf{x}}_{i|i} + L_i (\mathbf{y}_{i+1} - \hat{\mathbf{x}}_{i+1|i}) \\ & P_{i|i+1} = P_{i|i} - U_i (P_{i+1|i} + R_{i+1})^{-1} U_i^T \\ \text{Update:} \quad & K_{i+1} = P_{i+1|i} (P_{i+1|i} + R_{i+1})^{-1} \\ & \hat{\mathbf{x}}_{i+1|i+1} = \hat{\mathbf{x}}_{i+1|i} + K_{i+1} (\mathbf{y}_{i+1} - \hat{\mathbf{x}}_{i+1|i}) \\ & P_{i+1|i+1} = (I - K_{i+1}) P_{i+1|i} \end{aligned}$$

End For i

To obtain the strong constraint form of the variational smoother we will need both Q_i and $\tilde{Q}_i^{BL}$ to tend to zero. In the limit as $Q_i + \tilde{Q}_i^{BL} \rightarrow 0$ finding δ_i , δ_{i+1} which minimise eq. (44) is equivalent to finding δ_i which solves the familiar strong constraint 4D-Var problem:

$$\begin{aligned} \min_{\delta_i} \left\{ \frac{1}{2} \delta_i^T P_{i|i}^{-1} \delta_i + \frac{1}{2} (\mathbf{y}_{i+1} - \mathbf{f}_i^{(0)}(\hat{\mathbf{x}}_{i|i}) - \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i}) \delta_i)^T \right. \\ \left. \times R_{i+1}^{-1} (\mathbf{y}_{i+1} - \mathbf{f}_i^{(0)}(\hat{\mathbf{x}}_{i|i}) - \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i}) \delta_i) \right\} \end{aligned} \quad (45)$$

followed by

$$\delta_{i+1} = \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i}) \delta_i.$$

9. Example of smoothers using the TL and BL approximation

We use the BLA and TLA smoothers of Section 8, which we have seen are equivalent to a form of incremental 4D-Var, to assimilate data into the Duffing system of Section 6 (Example 1) and a roughened version (Example 2).

9.1. Example 1: standard version of Duffing map

We assimilate noisy observations into the Duffing map of Section 6, with the same parameters as there: $a = 2.75$, $b = 0.15$, $\sigma_o = 0.3$ and with no model error, and for both

Table 3. Variational form of smoother using the BLA

Given first guesses $\hat{\mathbf{x}}_{0|0}$, $P_{0|0}$:

For $i = 0, 1, \dots$

Find δ_i , δ_{i+1} which minimise (44)

$$\begin{aligned} \hat{\mathbf{x}}_{i|i+1} &= \hat{\mathbf{x}}_{i|i} + \delta_i \\ P_{i|i+1} &= [(J''(\delta_i, \delta_{i+1}))^{-1}]_{11} \\ \hat{\mathbf{x}}_{i+1|i+1} &= \mathbf{f}_i^{(0)}(\hat{\mathbf{x}}_{i|i}) + \delta_{i+1} \\ P_{i+1|i+1} &= [(J''(\delta_i, \delta_{i+1}))^{-1}]_{22} \end{aligned}$$

End For i

TLA and BLA filters/smoother a parametrisation of linearisation error of the form

$$\tilde{Q}^{TL}, \tilde{Q}^{BL} = \begin{pmatrix} 0 & 0 \\ 0 & \alpha \end{pmatrix}.$$

As for the filters in Section 6 we can easily compute exact analytic expressions for all the quantities required for the TLA and BLA smoothers of Section 8. The mean square analysis error obtained for 100 000 consecutive cycles is shown in Fig. 4, using the TLA and BLA filters of Section 5 (black), and the TLA and BLA smoothers of Section 8 (blue). Dotted curves use the TLA and solid curves use the BLA. As we have noted, the smoothers are equivalent to a form of 4D-Var.

As expected the smoothers outperform the filters. We also see that use of the BLA outperforms use of the TLA, even in this smooth slowly varying map for which the TLA is well-suited. However, the context for which the BLA was most intended is in systems where the propagating map varies rapidly on the scale of the increments, as in the next example.

9.2. Example 2: ‘roughened’ version of Duffing map

To illustrate the use of the BLA in cases where the derivative is unrepresentative of variation in the propagating map $\mathbf{f}_i$ on the scale of the increments, e.g. because of the presence of small scales, we consider a roughened version of the Duffing map.

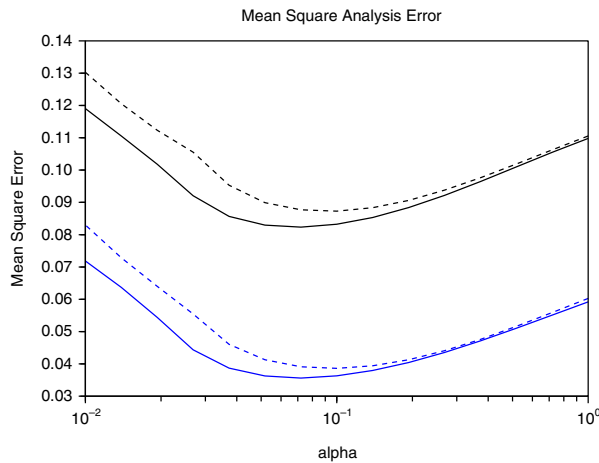


Fig. 4. Mean square analysis error for the Duffing map example for various values of the linearisation error parameter α . Filters are in black and smoothers in blue. Dotted curves use the TLA and solid curves use the BLA.

If the greatest integer less than or equal to a real scalar x is denoted $\text{floor}(x)$, for positive integer p and any real x set

$$[x]^p = \frac{\text{floor}(px)}{p}.$$

Thus the graph of $x \rightarrow [x]^p$ is a ‘staircase’ of square steps of step-size $1/p$. Note that $|x - [x]^p| \leq 1/p$ for any x , and so as $p \rightarrow \infty$ we have $|x - [x]^p| \rightarrow 0$ uniformly in x .

We now define a ‘roughened’ version of the Duffing map as

$$\mathbf{g}^p \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} [x_2]^p \\ -b[x_1]^p + a[x_2]^p - x_2^3 \end{pmatrix}. \quad (46)$$

This then is a roughened or pixelated version of the Duffing map which converges in the C^0 norm to the Duffing map as $p \rightarrow \infty$. It is still chaotic for $a = 2.75$, $b = 0.15$, at least if p is larger than about 50. Note that for any p the derivative of $\mathbf{g}^p$ is

$$\mathbf{g}^{p'}(\mathbf{x}) = -3x_2^2 \mathbf{e}_2 \mathbf{e}_2^T \quad (47)$$

so is zero except in the (2,2) component.

As before we can compute analytic expressions for all the quantities required for the TLA and BLA smoothers of Section 8. The mean square analysis error (averaged over 30 000 consecutive assimilations) for the roughened Duffing map with $a = 2.75$, $b = 0.15$, $p = 100$ is shown in Fig. 5 using the TLA and BLA filters (black) and smoothers (blue), where dotted curves use the TLA and solid curves use the BLA.

Given the form of eq. (47) it is no surprise that the assimilation methods which use the TLA struggle. Errors

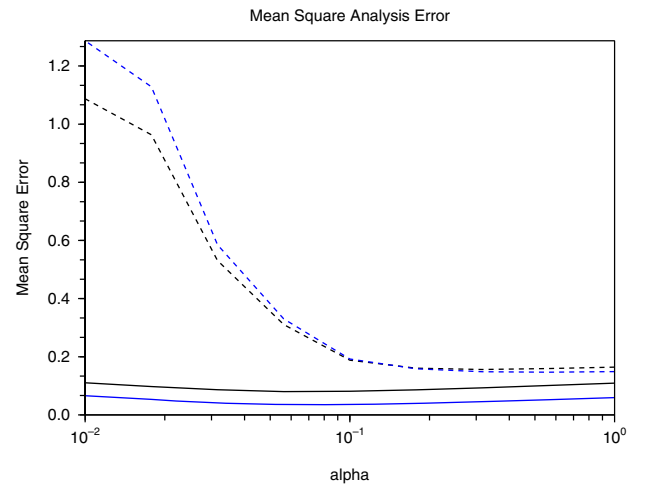


Fig. 5. Mean square analysis error in roughened Duffing map (with $p = 100$) for various values of the linearisation error parameter α . Filters are in black and smoothers in blue. Dotted curves use the TLA and solid curves use the BLA.

for the TLA filter are very large and the smoother is no better than the filter. Indeed, in view of eq. (47) we cannot expect that iterating the predict and smooth stages of the TLA smoother for the same i (equivalent to iterating the outer loop in incremental 4D-Var) will give any benefit, and this is confirmed by experiment. On the other hand when the BLA is used the errors are much lower, and (comparing the solid curves in Figs. 4 and 5) are in fact almost unaffected by the transition from smooth to roughened Duffing map.

It is worth noting that this example does not depend critically on the discontinuities in the construction of $\mathbf{g}^p$. For $\epsilon \in (0,1)$, let $\psi_\epsilon : \mathbb{R} \rightarrow [0,1]$ be a C^∞ ‘cut-off’ function: $\psi_\epsilon(x) = 1$ for $x \leq 0$, $\psi_\epsilon(x) = 0$ for $x \geq \epsilon$, ψ_ϵ is monotonically decreasing for $0 < x < \epsilon$, and all derivatives of $\psi_\epsilon(x)$ vanish at $x=0, \epsilon$. If one then sets

$$\text{floor}_\epsilon(x) = \begin{cases} \text{floor}(x) - \psi_\epsilon(x - \text{floor}(x)) & \text{for } x - \text{floor}(x) \leq \epsilon \\ \text{floor}(x) & \text{for } x - \text{floor}(x) \geq \epsilon \end{cases}$$

then floor_ϵ is a smooth function which ‘rounds off’ the steps of the function floor and converges to floor as $\epsilon \rightarrow 0$. Now define $[x]_\epsilon^p$ in terms of floor_ϵ and $\mathbf{g}_\epsilon^p$ in terms of $[x]_\epsilon^p$. Then we could repeat the example using smooth $\mathbf{g}_\epsilon^p$, and for sufficiently small ϵ the results will not differ appreciably from those we found using $\mathbf{g}^p$.

10. Summary and remarks

To form a local linear approximation for *finite-sized* perturbations of a non-linear function $\mathbf{f}$ the tangent-linear approximation may be inappropriate. The best linear approximation is a readily calculable function of $\mathbf{f}$ and the pdf of the increments. It does not require differentiability of $\mathbf{f}$. The fact that the best linear approximation is independent of the smoothness of $\mathbf{f}$ makes it particularly attractive for the linearisation of physical parametrisations, and as we saw in the cloud example the reduction in linearisation error in such cases can be considerable.

The main focus of the article is on applying the BLA in the data assimilation cycle. As a benchmark we use minimum variance filters and smoothers which are non-linear generalisations of Kalman filters and smoothers. In our context, the predictive part of the minimum variance filter is equivalent to evolving the posterior pdf by the Chapman–Kolmogorov equation, and approximating the new prior pdf by a Gaussian one with the same first two moments.

If into the minimum variance filter we replace the forecast map by its TL approximation we obtain the usual EKF. If instead we use the BL approximation, we noted two advantages, the first being that the prior is unapproximated compared with the minimum variance filter. This

occurs because the expectation of $\mathbf{f}_i$ is also the expectation of the BLA of $\mathbf{f}_i$. The covariance of the prior in the minimum variance filter involves the expectation of the square of $\mathbf{f}_i$, so using the BLA of $\mathbf{f}_i$ leads to an approximation in this. However, the second advantage of the BLA is that, by contrast with the TLA, the error in this approximation is bounded independently of the vagaries of the instantaneous derivative.

The first advantage removes an approximation while the second greatly reduces the damage a bad approximation can do. Their relative importance tends to depend on how poor the TL approximation is. If it is quite good, the main advantage is in the improved prior itself. On the other hand, if the TL approximation is poor the improved prior covariance of the BL filter becomes very important.

Interestingly, while the BL approximation yields a prior covariance which is normally closer to the benchmark than that from the TLA, it always underestimates it, by the covariance of the linearisation error. This means that allowing for linearisation error, which is already known to be important for the EKF using the TLA, is at least as important if the BLA is used. In the latter case, this involves adding an estimate of the covariance of linearisation error to the prediction of the covariance of state error. This is especially important in the perfect model context. With an imperfect model the approximations made in parametrising the true model error covariance may in some cases render separate allowance for linearisation error unnecessary.

We also developed a minimum variance smoother. Inserting the BL and TL approximations into this, we see that BLA smoothers have the same advantages over TLA smoothers that the filters do, and additionally that the smoothing step in the BLA smoother is unapproximated compared with the minimum variance smoother.

We have concentrated on non-linearity in the forecast map, but there is no difficulty (the details are in an appendix) applying the same principles to non-linearity in the observation operator.

In some large-scale applications the computation of $\mathbf{f}^{(0)}$ and $\mathbf{f}^{(1)}$ will need to be done approximately, for example by use of an ensemble of forecasts from perturbed analyses which may already exist for ensemble forecasting. Alternatively, in some instances one may be able to get a useful approximation to $\mathbf{f}^{(1)}$ by using a climatological pdf of perturbations, and so $\mathbf{f}^{(1)}$ can be computed entirely off-line. One can interpret the ‘perturbation forecast’ model approach of Lorenc (1997) and adopted at the Met Office as being in this spirit.

The analysis in this article has been for filters and smoothers in their discrete form, but a similar approach can be followed in the continuous case, by using the BLA in place of the TLA when linearising the non-linear terms in the PDEs describing the forecast model.

11. Acknowledgments

The author thanks Andrew Lorenc for useful discussions, and three referees appointed by *Tellus* for their careful reviews and helpful suggestions.

Appendix

A.1. Equivalence of the minimum variance algorithm of Section 4 and optimal Bayesian dynamic state estimation with Gaussianised pdfs

We use the framework described at the beginning of Section 4 and in particular in eq. (12), and suppose also that observation and model errors are Gaussian, i.e.

$$\mathbf{v}_{i+1} \sim N(\mathbf{0}, R_{i+1}), \mathbf{w}_i \sim N(\mathbf{0}, Q_i). \quad (\text{A1})$$

As usual we denote the conditional pdf of $\mathbf{x}_i$ given $\mathbf{y}_1, \dots, \mathbf{y}_i$ by $p(\mathbf{x}_i | \mathbf{y}_{1:i})$ or $p(\mathbf{x}_{i|i})$. We show that the minimum variance filter of Section 4 is equivalent to optimal Bayesian dynamic state estimation, if in the latter case each prior pdf is approximated by a Gaussian with the same first two moments.

In Bayesian dynamic state estimation (e.g. Arulampalam et al. (2002)), Bayes' rule

$$p(\mathbf{x}_{i|i}) = \frac{p(\mathbf{y}_i | \mathbf{x}_i) p(\mathbf{x}_{i|i-1})}{\int p(\mathbf{y}_i | \mathbf{x}_i) p(\mathbf{x}_{i|i-1}) d\mathbf{x}_i} \quad (\text{A2})$$

is used to form the posterior pdf $p(\mathbf{x}_{i|i})$ from the prior pdf $p(\mathbf{x}_{i|i-1})$. With our assumption [eq. (A1)] this posterior will be Gaussian if $p(\mathbf{x}_{i|i-1})$ is. The pdf $p(\mathbf{x}_{i|i})$ is then evolved forward to form the prior pdf $p(\mathbf{x}_{i+1|i})$ at the next time level by

$$p(\mathbf{x}_{i+1|i}) = \int p(\mathbf{x}_{i+1} | \mathbf{x}_i) p(\mathbf{x}_{i|i}) d\mathbf{x}_i \quad (\text{A3})$$

Equation (A3) can be viewed as a form of the Chapman–Kolmogorov equation, and we follow Arulampalam et al. (2002) by referring to it as such. Observe that by virtue of eqs. (12) and (A1) the transition pdf $p(\mathbf{x}_{i+1} | \mathbf{x}_i)$ used in (A3) is Gaussian with mean $\mathbf{f}_i(\mathbf{x}_i)$ and covariance Q_i , so if $\mathbf{f}_i$ is non-linear then $p(\mathbf{x}_{i+1|i})$ is non-Gaussian even if $p(\mathbf{x}_{i|i})$ was Gaussian.

Suppose that the prior pdf $p(\mathbf{x}_{1|0})$ is Gaussian with mean $\hat{\mathbf{x}}_{1|0}$ and covariance $P_{1|0}$. Applying Bayes Rule [eq. (A2)], we find $p(\mathbf{x}_{1|1})$ is Gaussian with mean $\hat{\mathbf{x}}_{1|1}$ and covariance $P_{1|1}$ given by eq. (21) with $i=0$, so the optimal Bayesian approach and the minimum variance filter of Section 4 coincide at the outset.

Now suppose inductively that at the i th stage the posterior pdf $p(\mathbf{x}_{i|i})$ is Gaussian with mean $\hat{\mathbf{x}}_{i|i}$ and covariance $P_{i|i}$ as given by eq. (21). Then using eq. (12),

the definition of conditional mean and the inductive hypothesis it follows that the conditional mean of $\mathbf{x}_{i+1}$ given $\mathbf{y}_1, \dots, \mathbf{y}_i$

$$E[\mathbf{x}_{i+1} | \mathbf{y}_{1:i}] = E[\mathbf{f}_i(\tilde{\mathbf{x}}_i)] \quad (\text{A4})$$

where the expectation on the right hand side is over $\tilde{\mathbf{x}}_i \sim N(\hat{\mathbf{x}}_{i|i}, P_{i|i})$. Denoting $E[\mathbf{x}_{i+1} | \mathbf{y}_{1:i}]$ by $\hat{\mathbf{x}}_{i+1|i}$, we find by similar means that the conditional covariance of $\mathbf{x}_{i+1}$ given $\mathbf{y}_1, \dots, \mathbf{y}_i$

$$\begin{aligned} E[(\mathbf{x}_{i+1} - \hat{\mathbf{x}}_{i+1|i})(\mathbf{x}_{i+1} - \hat{\mathbf{x}}_{i+1|i})^T | \mathbf{y}_{1:i}] = \\ E[(\mathbf{f}_i(\tilde{\mathbf{x}}_i) - \hat{\mathbf{x}}_{i+1|i})(\mathbf{f}_i(\tilde{\mathbf{x}}_i) - \hat{\mathbf{x}}_{i+1|i})^T] + Q_i \end{aligned} \quad (\text{A5})$$

where again the expectation on the right hand side is over $\tilde{\mathbf{x}}_i \sim N(\hat{\mathbf{x}}_{i|i}, P_{i|i})$. Evidently the right hand sides of eqs. (A4) and (A5) are identical to the expressions for $\hat{\mathbf{x}}_{i+1|i}$ and $P_{i+1|i}$ in eq. (20). If we now replace $p(\mathbf{x}_{i+1|i})$ by a Gaussian pdf with mean [eq. (A4)] and covariance [eq. (A5)] and use this modified pdf in Bayes' rule [eq. (A2)], we find $p(\mathbf{x}_{i+1|i+1})$ is Gaussian with mean $\hat{\mathbf{x}}_{i+1|i+1}$ and covariance $P_{i+1|i+1}$ given by eq. (21). By induction, it follows that if at every stage we approximate the prior pdf $p(\mathbf{x}_{i+1|i})$ given by the Chapman–Kolmogorov equation [eq. (A3)] by a Gaussian one with the same mean and covariance, then the resulting prediction-update cycle coincides with the minimum variance filter of Section 4.

A.2. The BLA filter – existence, comparison with UKF, and extension to non-linear observation operators

A.2.1. Conditions for the existence of the filter using the BLA. We compare the TLA and BLA filters [eqs. (31) and (32)] in terms of what restrictions must be placed on $\mathbf{f}_i$ and the other quantities involved for the filters to exist. Clearly, if the TLA is used as in eq. (31) then we need $\mathbf{f}_i$ to be differentiable. For the BLA filter [eq. (32)], $\mathbf{f}_i$ does not even need to be continuous, which can be a significant advantage. We merely require the existence of

$$\begin{aligned} \mathbf{f}_i^{(0)}(\hat{\mathbf{x}}_{i|i}) &= E[\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i)] \\ \mathbf{f}_i^{(1)}(\hat{\mathbf{x}}_{i|i}) &= E[\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i)\boldsymbol{\epsilon}_i^T] \{E[\boldsymbol{\epsilon}_i\boldsymbol{\epsilon}_i^T]\}^{-1} \end{aligned}$$

where the expectations are over $\boldsymbol{\epsilon}_i \sim N(\mathbf{0}, P_{i|i})$. For this, we require (a) that for every i the expectations $E[\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i)]$ and $E[\mathbf{f}_i(\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i)\boldsymbol{\epsilon}_i^T]$ exist, which is ensured under very mild restrictions on $\mathbf{f}_i$ by virtue of the rapidity with which a Gaussian pdf $p(\mathbf{x})$ tends to zero as $|\mathbf{x}| \rightarrow \infty$, and (b) the non-singularity of $E[\boldsymbol{\epsilon}_i\boldsymbol{\epsilon}_i^T] = P_{i|i}$. Suppose $P_{i-1|i-1}$ is

non-singular. From eq. (32) we see that $P_{i|i-1}$ is a sum of three positive semi-definite matrices,

$$P_{i|i-1} = E \left[\mathbf{f}_{i-1} (\hat{\mathbf{x}}_{i-1|i-1} + \boldsymbol{\epsilon}_{i-1}) \boldsymbol{\epsilon}_{i-1}^T \right] (P_{i-1|i-1})^{-1} \\ \times E \left[\mathbf{f}_{i-1} (\hat{\mathbf{x}}_{i-1|i-1} + \boldsymbol{\epsilon}_{i-1}) \boldsymbol{\epsilon}_{i-1}^T \right]^T + \tilde{Q}_{i-1}^{BL} + Q_{i-1}$$

and so if $P_{i-1|i-1}$ is non-singular then so is $P_{i|i-1}$ so long as at least one of $E \left[\mathbf{f}_{i-1} (\hat{\mathbf{x}}_{i-1|i-1} + \boldsymbol{\epsilon}_{i-1}) \boldsymbol{\epsilon}_{i-1}^T \right]$, $\tilde{Q}_{i-1}^{BL}$ or Q_{i-1} is. From eq. (21) we have $P_{i|i} = R_i (P_{i|i-1} + R_i)^{-1} P_{i|i-1}$. Since we may assume R_i is non-singular it follows that if $P_{i|i-1}$ is non-singular then so is $P_{i|i}$. Inductively, $P_{i|i}$ is non-singular for all i so long as for every i at least one of $E \left[\mathbf{f}_i (\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) \boldsymbol{\epsilon}_i^T \right]$, $\tilde{Q}_i^{BL}$ or Q_i is non-singular.

A.2.2. Remarks on the unscented Kalman filter and comparison with the BLA filter. Another filter, which like eqs. (31) and (32) is based on eq. (20) for its predictive part, is the unscented Kalman filter (UKF) of Julier and Uhlmann (1997). If the state space has dimension n , this estimates the expectations in eq. (20) by evaluating $\mathbf{f}_i$ at $2n+1$ carefully chosen ‘sigma-points’. Julier et al. (2000) compute Taylor series expansions of the mean and covariance as predicted by eq. (20) and compare with similar expansions for the standard EKF [i.e. using the TLA as in eq. (31)] and the UKF, where in all cases, as in eq. (20), the analysis errors have a Gaussian distribution. They find that using the TLA the mean is correct to first order and the covariance to third order, while using the UKF the mean is correct to second order and the covariance again to third order. For the BLA filter, we have already noted that $\hat{\mathbf{x}}_{i+1|i}$ in eq. (32) is identical to the value given by eq. (20), i.e. the mean is correct to every order, and by a similar analysis to that in Julier et al. (2000) we find that the BLA filter also gives third-order accuracy for the covariance. More importantly, we have seen that the departure of the covariance of the BLA filter from the benchmark [eq. (20)] is bounded by [eq. (30)].

A.2.3. Note on extension of the BLA filter to non-linear observation operators. For simplicity we have taken the observation operator to be the identity. More generally the last line of eq. (12) takes the form

$$\mathbf{y}_{i+1} = \mathbf{h}_{i+1}(\mathbf{x}_{i+1}) + \mathbf{v}_{i+1}$$

where the observation operator $\mathbf{h}_{i+1}$ may be non-linear. In this case, one can repeat the analysis of Section 4 to generalise our minimum variance filter [eqs. (20) and (21)].

One finds that eq. (20) is unchanged but the update equations in eq. (21) become

$$K_{i+1} = C_{i+1}^{xy} (C_{i+1}^{yy})^{-1} \\ \hat{\mathbf{x}}_{i+1|i+1} = \hat{\mathbf{x}}_{i+1|i} + K_{i+1} (\mathbf{y}_{i+1} - \hat{\mathbf{y}}_{i+1|i}) \\ P_{i+1|i+1} = P_{i+1|i} - K_{i+1} (C_{i+1}^{xy})^T \quad (\text{A6})$$

where denoting $E[(\mathbf{x} - E[\mathbf{x}])(\mathbf{y} - E[\mathbf{y}])^T]$ by $\text{cov}\{\mathbf{x}, \mathbf{y}\}$ we have

$$\hat{\mathbf{y}}_{i+1|i} = E \left[\mathbf{h}_{i+1} \left(\mathbf{f}_i (\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) + \mathbf{w}_i \right) \right] \quad (\text{A7})$$

$$C_{i+1}^{xy} = \text{cov} \left\{ \mathbf{f}_i (\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) + \mathbf{w}_i, \mathbf{h}_{i+1} \left(\mathbf{f}_i (\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) + \mathbf{w}_i \right) \right\} \quad (\text{A8})$$

$$C_{i+1}^{yy} = \text{cov} \left\{ \mathbf{h}_{i+1} \left(\mathbf{f}_i (\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) + \mathbf{w}_i \right), \right. \\ \left. \mathbf{h}_{i+1} \left(\mathbf{f}_i (\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) + \mathbf{w}_i \right) \right\} + R_{i+1}. \quad (\text{A9})$$

The BLA of $\mathbf{h}_{i+1} \left(\mathbf{f}_i (\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i) + \mathbf{w}_i \right)$ has the form $\mathbf{a}_{hf} + A_{hf}^e \boldsymbol{\epsilon}_i + A_{hf}^w \mathbf{w}_i$ and can either be computed directly, or approximated by first computing the BLA of $\mathbf{f}_i (\hat{\mathbf{x}}_{i|i} + \boldsymbol{\epsilon}_i)$ by eq. (25), and then computing the BLA of $\mathbf{h}_{i+1}$ applied to this. Either way we can insert this linear approximation and eq. (25) into eqs. (A7)–(A9), and thereby generalise the BLA filter to allow for non-linear observation operators.

References

- Arulampalam, M. S., Maskell, S., Gordon, N. and Clapp, T. 2002. A tutorial on particle filters for online nonlinear/non-Gaussian Bayesian tracking. *IEEE Trans. Signal Process.* **50**, 174–188.
- Bell, B. M. 1994. The iterated Kalman smoother as a Gauss–Newton method. *SIAM J. Optim.* **4**, 626–636.
- Errico, R. M. and Raeder, K. D. 1999. An examination of the accuracy of the linearization of a mesoscale model with moist physics. *Q. J. Roy. Meteorol. Soc.* **125**, 169–195.
- Fisher, M., Leutbecher, M. and Kelly, G. A. 2005. On the equivalence between Kalman smoothing and weak-constraint four-dimensional variational data assimilation. *Q. J. Roy. Meteorol. Soc.* **131**, 3235–3246.
- Gelb, A. 1974. *Applied Optimal Estimation*. MIT Press, Cambridge, MA.
- Jaswinski, A. H. 1970. *Stochastic Processes and Filtering Theory*. Academic Press, New York.
- Julier, S. J. and Uhlmann, J. K. 1997. A new extension of the Kalman filter to nonlinear systems. In: *Proc. SPIE Vol. 3068, Signal Processing, Sensor Fusion, and Target Recognition VI*, (ed. I. Kadar). SPIE, Orlando, FL, pp. 182–193.
- Julier, S., Uhlmann, J. and Durrant-Whyte, H. F. 2000. A new method for the nonlinear transformation of means and covariances in filters and estimators. *IEEE Trans. Automat. Contr.* **45**, 477–482.
- Lawless, A. C., Gratton, S. and Nichols, N. K. 2005. An investigation of incremental 4D-Var using non-tangent linear models. *Q. J. Roy. Meteorol. Soc.* **131**, 459–476.

- Lorenc, A. C. 1997. Development of an operational variational assimilation scheme. *J. Met. Soc. Japan* **75**, 339–346.
- Lorenc, A. C. and Payne, T. J. 2007. 4D-Var and the butterfly effect. *Q. J. Roy. Meteorol. Soc.* **133**, 607–614.
- Luca, M. B., Azou, S., Burel, G. and Serbanescu, A. 2006. On exact Kalman filtering of polynomial systems. *IEEE Trans. Circ. Syst.* **53**, 1329–1340.
- Sage, A. P. and Melsa, J. L. 1971. *Estimation Theory with Applications to Communications and Control*. McGraw Hill, New York.
- Saha, L. M. and Tehri, R. 2010. Applications of recent indicators of regularity and chaos in discrete maps. *Int. J. Appl. Math. Mech.* **6**, 86–93.
- Smith, R. N. B. 1990. A scheme for predicting layer clouds and their water content in a general circulation model. *Q. J. Roy. Meteorol. Soc.* **116**, 435–460.