

The error of representation: basic understanding

By DANIEL HODYSS^{1*} and NANCY NICHOLS², ¹*Marine Meteorology Division, Naval Research Laboratory, Monterey, CA, USA;* ²*School of Mathematical and Physical Sciences, University of Reading, Reading, UK*

(Manuscript received 1 May 2014; in final form 10 December 2014)

ABSTRACT

Representation error arises from the inability of the forecast model to accurately simulate the climatology of the truth. We present a rigorous framework for understanding this kind of error of representation. This framework shows that the lack of an inverse in the relationship between the true climatology (true attractor) and the forecast climatology (forecast attractor) leads to the error of representation. A new gain matrix for the data assimilation problem is derived that illustrates the proper approaches one may take to perform Bayesian data assimilation when the observations are of states on one attractor but the forecast model resides on another. This new data assimilation algorithm is the optimal scheme for the situation where the distributions on the true attractor and the forecast attractors are separately Gaussian, and there exists a linear map between them. The results of this theory are illustrated in a simple Gaussian multivariate model.

Keywords: Representation error, data assimilation, correlated observations, model error, Bayesian

1. Introduction

Representation error in this manuscript will refer to the impact of the unavoidable misrepresentation of complex atmospheric flows by the inadequacies of the forecast model on the data assimilation. This misrepresentation of complex fluid flows arises for the most part from the inability of the forecast model, using the relatively coarse grids customarily employed in numerical weather prediction, to resolve small-scale properties of the boundary conditions as well as other small-scale properties of the turbulent flows in the interior of the fluid. This misrepresentation of the flow leads to an incompatibility between the observations of the true state, which see these small-scale processes, and the relatively coarser states achievable by the forecast model. The result of this incompatibility is that the forecast model can effectively consider the states implied by the observations to be inconsistent with its own attracting manifold, and therefore either ignore some or all of the information in the observations or react pathologically to them.

Recognition of this incompatibility between the observations of the true state and the states achievable by the forecast model goes back at least to Petersen and Middleton

(1963) with more thorough and modern treatments in Daley (1993), Mitchell and Daley (1997a, 1997b), Liu and Rabier (2002), Janjic and Cohn (2006) and Frehlich (2006). This body of work has correctly identified that because of the incompatibility mentioned above, performance gains in the quality of the analysis can be made by inflating the observation error variances and accounting for the implied correlations between observations owing to unresolved processes. In addition to these, more theoretical works there have recently been attempts at estimating the structure of representation error from observational data and numerical model output (e.g. Richman et al., 2005; Frehlich, 2008; Oke and Sakov, 2008; Waller et al., 2013). This work has shown that there are a number of ways to see this error of representation in data. For example, the standard observations of temperature and humidity as well as aircraft measurements of turbulence have been compared with forecasts and shown to contain a component consistent with errors in representation.

This paper intends to conjoin this past work under a single unifying theme. The basic idea is to extend the Kalman (1960) filter to explicitly account for the fact that the climate (attracting manifold) of the Earth's atmosphere is distinct from that of the forecast model. To understand how we will accomplish this, it will prove illustrative to review Kalman's original setup. In the work of Kalman,

*Corresponding author.
email: daniel.hodyss@nrlmry.navy.mil

the flaws in the model were accounted for using a white noise source, viz.

$$\mathbf{x}_f^{k+1} = \mathbf{M}\mathbf{x}_f^k + \mathbf{w}^k, \quad (1)$$

where $\mathbf{M}$ is a linear model, $\mathbf{x}_f^k$ is the forecast at time k and $\mathbf{w}^k$ is white noise drawn from $N(\mathbf{0}, \mathbf{Q})$. The idea behind eq. (1) is that the model $\mathbf{M}$ is not entirely adequate at producing the set of states that correctly describes all the potential true states (one of which is being observed by the observational instruments). In terms of the problem at hand, we may interpret eq. (1) as implying that the forecast model's climatology (attracting manifold) is not sufficient to encompass the complete set of potential true states. The inclusion of the white noise source is then there to enhance the spread of states in state space (with $\mathbf{Q}$ chosen large enough to more than fill this gap between the true and forecast attractors) such that this noisy forecast model does encompass the complete set of potential true states. In other words, the action of this noise is such that the forecast model's attracting manifold, which is distinct from that of the true attracting manifold, is in essence 'blurred' in phase-space until the states available to eq. (1) encompasses the states on the true attracting manifold and many others for that matter.

There are at least two downsides to this choice to render the model stochastic. The first is that while probability forecasts using this noisy model now correctly assign non-zero probability to states on the true attracting manifold, other forecasts from eq. (1) will assign non-zero probability to states off the true attracting manifold, which of course implies that implausible events are assigned a non-zero probability of occurrence. The second issue with this noisy model is that in fluids as complex as the general circulation of the atmosphere, it is well known that choosing the character of this noise is difficult and improper choices may lead to issues with the physical realism of the fluid evolution from the noisy forecast model (e.g. Hodyss et al., 2014).

While we believe that stochastic modelling can be useful, the tack taken here is to explore what happens when one does not add noise to the model but addresses the climatology (attractor) differences between the true physical system and that of the forecast model through the use of maps between distributions on each attractor. We will derive the best, linear unbiased estimate (minimum error variance) of the state *on the forecast model attractor* given observations of states on the true attractor. This new data assimilation method will be optimal when the distributions on the true attractor and the forecast attractor are Gaussian and the map between them is linear. By deriving the data assimilation method that produces the best state estimates *on the forecast model attractor*, we will see the

error of representation arise as a natural consequence of a specific form of attracting manifold difference.

Finally, we would like to point out here that one common viewpoint for the connection between the data assimilation process and the forecast process is the expectation that the most accurate forecast arises from the most accurate estimate of the true state. This desire for the data assimilation algorithm to produce an estimate that is close to the true state is conceptually easy to rationalise when the model is perfect. However, when the model is flawed, creating a data assimilation algorithm that produces an accurate estimate of the true state is likely to mean that this state is in some way incompatible (e.g. unbalanced) with the flawed forecast model. This notion that the true state might not be the best initial condition for the flawed forecast model has led us to choose to develop a data assimilation system that attempts to produce a state estimate on the forecast attractor. The estimation of the true state is then one relegated to post-processing of the resulting forecasts. More discussion of the ramifications of this choice will be made throughout the development and in the conclusions.

In Section 2, we derive the general theory for this new form of data assimilation algorithm. In Section 3, we apply the theory of Section 2 to representation error in the form of a smoothing operator that describes the differences between the true and forecast attractors. Section 4 applies the theory of Section 3 to a multivariate Gaussian model to illustrate the basic ideas in their simplest forms. Section 5 closes the manuscript with a summary of the results and a discussion of the conclusions.

2. The two attractor problem

In this section, we formulate a general theory for data assimilation in the situation where the observations are of states in one subset of state space but the forecast model resides in another. Our basic assumption throughout will be that our goal in such a situation is to develop a data assimilation method that will provide the best estimate of the state in the region of state space in which the forecast model resides. As we go we will illustrate the theory of this section using a very simple, example problem that we believe illustrates the basic ideas in their simplest form.

2.1. The true posterior

We imagine the true state, $\mathbf{x}_t$, to be an N -vector and that it is drawn from a climatological distribution whose probability density function (pdf) we label $\rho(\mathbf{x}_t)$. By 'climatological', we are referring to the pdf we would obtain if we ran the true model for a very long time, discretised state space, and counted the number of times the state entered each cell of our discretisation. In the limit as this true

model simulation becomes very long and the volume of each cell of our discretisation of state space becomes very small, we obtain this climatological pdf. We define from this climatological pdf the set $\mathbb{A}_t$ of states $\mathbf{x}_t$ with the property that $\rho(\mathbf{x}_t) > 0$ and will subsequently refer to the set $\mathbb{A}_t$ as the true ‘attracting manifold’ of the physical system under consideration. Therefore, the set $\mathbb{A}_t$ consists of all the states in which the true state at any point in the future may be found.

We simply use the phrase ‘attracting manifold’ throughout this manuscript as a convenient way to refer to the portion of state space in which the true physical system resides. We emphasise however that the following theory does not require the existence of a compactly supported region in state space with the properties normally associated with attracting manifolds as we will apply the theory to example problems which do not technically have this feature. More discussion of this set and its relationship to the ‘attracting manifold’ of the forecast model can be found in the next subsection.

We obtain a sequence of p -vector observations, $\mathbf{y}_j = \mathbf{H}\mathbf{x}_t + \boldsymbol{\varepsilon}_0$ that were taken at various times $j = 1, 2, \dots, J$. The object $\mathbf{H}$ is a vector-valued function and, for simplicity, will be assumed to be a linear matrix operator ($p \times N$) with the instrument errors drawn from $\boldsymbol{\varepsilon}_0 \sim \mathcal{N}(\mathbf{0}, \mathbf{R}_i)$. We emphasise however that the results presented below will not depend on a linear observation operator. For simplicity, we will assume that both the instrument error variance, $\mathbf{R}_i$, and the number of observations, p , per assimilation time, j , are fixed constants. Because we have observations at various times, j , this implies that the state must also be integrated through time. In the interest of simplicity of presentation, we do not attach a label to $\mathbf{x}_t$ that denotes its relevant time j . However, because we will refer to the ‘filtering’ data assimilation problem throughout, this should cause no confusion as $\mathbf{x}_t$ will always be considered to be at the time of the latest set of observations.

We begin by assimilating the $j=1$ set of observations using Bayes’ rule, viz.

$$\rho(\mathbf{x}_t | \mathbf{y}_1) = C_1 \rho(\mathbf{y}_1 | \mathbf{x}_t) \rho(\mathbf{x}_t), \quad (2)$$

where $C_1 = 1/\rho(\mathbf{y}_1)$. The density $\rho(\mathbf{y}_1 | \mathbf{x}_t)$ describes the conditional distribution of observations given a particular value of the state on the true attractor (often referred to as the observation likelihood). The interpretation of the act of employing Bayes’ rule in eq. (2) is simply as a ‘windowing’ function through the observation likelihood that acts to reduce the view of the climatological distribution to a portion of $\mathbb{A}_t$ in the vicinity of the observation. This is important because, under the assumption that the observation likelihood is Gaussian, which implies that $\rho(\mathbf{y}_1 | \mathbf{x}_t) > 0$ for all values of $(\mathbf{x}_t, \mathbf{y}_1)$, this means that the act of data

assimilation does not change the states in $\mathbb{A}_t$ but simply re-calculates their relative probability of being the true state. We will use this fact later in the next subsection.

By repeating the indicated operations in eq. (2), and using the Fokker-Planck equation to propagate the resulting distributions forward in time between each set of observations, we may repeat the process in eq. (2) up to the present time, $j=J$, for which we have observations $\mathbf{Y}_J$ such that:

$$\rho(\mathbf{x}_t | \mathbf{Y}_J) = C_J \rho(\mathbf{y}_J | \mathbf{x}_t) \rho(\mathbf{x}_t | \mathbf{Y}_{J-1}), \quad (3)$$

where the symbol $\mathbf{Y}_j$ denotes the set of all observations at all times up to and including the j th set and the $C_J = 1/\rho(\mathbf{y}_J | \mathbf{Y}_{J-1})$ is simply the normalisation. The density $\rho(\mathbf{x}_t | \mathbf{Y}_{J-1})$ will hereafter be referred to as the ‘prior’. The density $\rho(\mathbf{x}_t | \mathbf{Y}_J)$ describes the conditional distribution of all possible true states given all observations including the present set; this density will hereafter be referred to as the ‘posterior’.

As alluded to in the beginning of this section, we construct here a simple data assimilation example that we believe will illustrate the basic properties of the more complex concepts in the remainder of this section in the simplest way. To this end, we assume the true states to be characterised by a two-vector whose climatological distribution is illustrated in Fig. 1a. This distribution is Gaussian with a variance of 3 on each variable and a covariance between the two variables of 1. Note that the set $\mathbb{A}_t$ in this case is the entire plane as a Gaussian vanishes nowhere. An observation of only one of the variables is made and defines an observation likelihood with instrument error variance equal to 1 (Fig. 1b). Calculating eq. (2) from the climatological distribution and observation likelihood for an observation of $y_1 = 1$ obtains the posterior in Fig. 1c. For simplicity, we further assume that the true model that is propagating the states forward in time is simply the identity. This implies that the posterior from eq. (2) is simply the prior for the next assimilation step in eq. (3) and the result of this calculation with $y_2 = 3$ is plotted in Fig. 1d. One can see in Fig. 1c and 1d that the assimilation of the observations has reduced the variance greatly for the observed variable but less so for the unobserved variable. In the next subsection, we will return to this example and illustrate its relationship to the forecast posterior.

2.2. The forecast posterior

Because of our fundamental inability to construct an exact model of the evolution of the general circulation of the atmosphere, the posterior distribution described by eq. (3) can never be produced by a data assimilation algorithm. This is true for several reasons. First, even if we could give

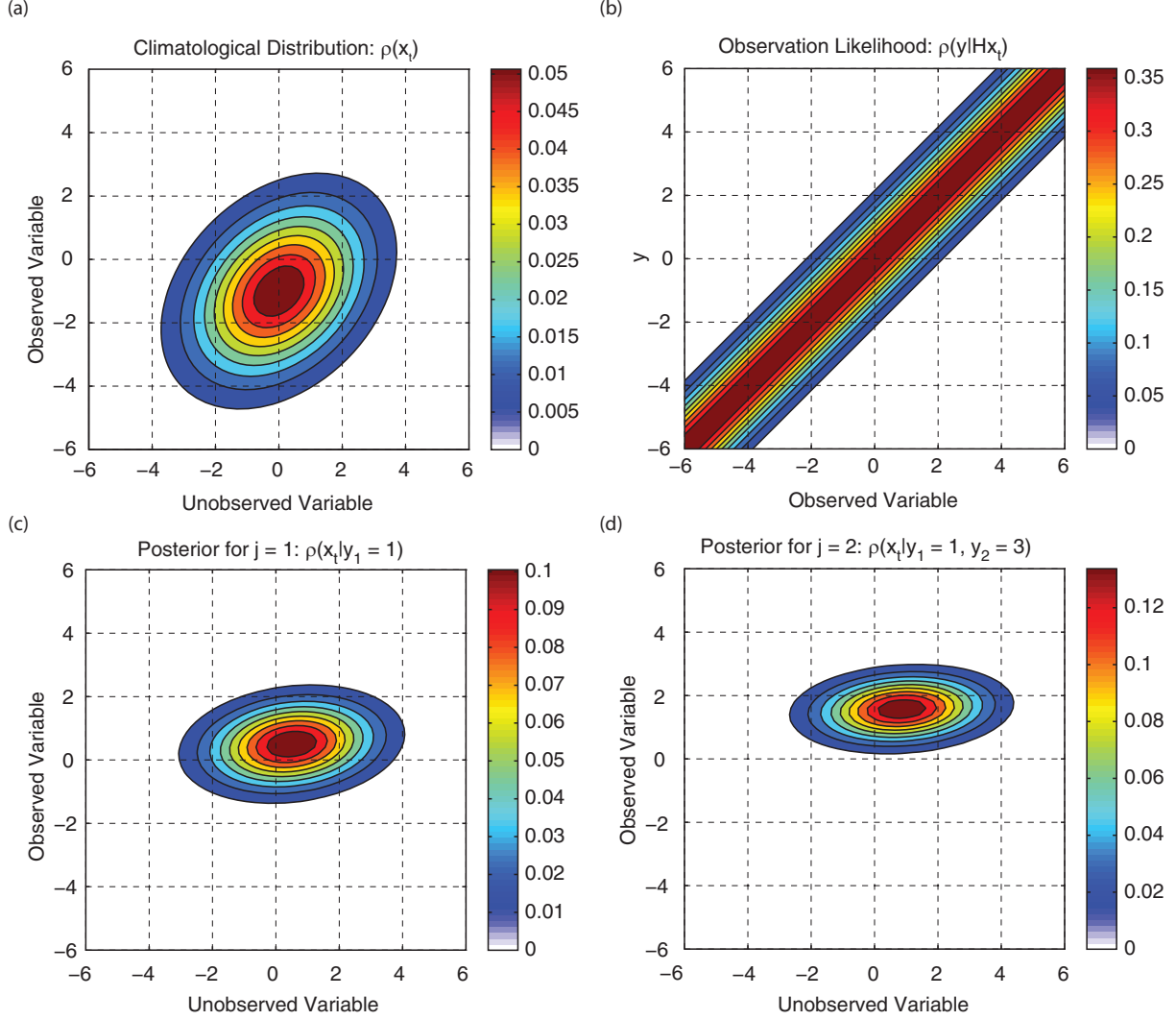


Fig. 1. Data assimilation on the high-resolution (true) attractor. (a) High-resolution climatological distribution, (b) high-resolution observation likelihood, (c) high-resolution posterior for $j=1$ and (d) high-resolution posterior for $j=2$.

the forecast model the true state, it would not produce the correct forecast evolution because of incorrectly specified parameters and altogether missing physics. Second, what exactly is the appropriate true state to give the forecast model as an initial condition is ambiguous given the fact that this model can only represent states that are ‘smoothed’, or said another way, truncated with respect to the truth (Frehlich, 2011). More on this notion will be presented in Sections 3 and 4.

Because the forecast model cannot represent the true attractor, we begin by defining the state on the forecast attractor, $\mathbf{x}_f$, as an M -vector and hypothesising that it too is drawn from a ‘climatological’ distribution whose pdf we label $\rho(\mathbf{x}_f)$. We emphasise that the ‘climatology’ here is different from that of the true distribution denoted above and results from running the forecast model for a very long

time, discretising state space, and counting the number of times the forecast state enters each cell of our discretisation. We again define from this forecast climatological distribution a new and distinct set of states $\mathbb{A}_f$ with the property that $\rho(\mathbf{x}_f) > 0$ and will subsequently refer to the set $\mathbb{A}_f$ as the ‘attracting manifold’ of the forecast model’s *representation* of the physical system under consideration.

Next, we hypothesise the existence of a map (function) over $\mathbb{A}_t \rightarrow \mathbb{A}_f$. We do this by relating the states on the forecast attractor and the states on the true attractor through a vector-valued function:

$$\mathbf{x}_f = \mathbf{F}(\mathbf{x}_t). \quad (4)$$

The function, $\mathbf{F}$, represents a mapping from the true attracting manifold to that of the forecast model. We emphasise that this is a mapping from one state space ($\mathbb{A}_t$)

to another ($\mathbb{A}_f$) and not a direct relationship between today's truth and today's forecast, which because of the noise in the observations could not possibly satisfy eq. (4). We will assume that this function, $\mathbf{F}$, has the property that for every $\mathbf{x}_t$ in $\mathbb{A}_t$ there is a corresponding element of $\mathbf{x}_f$ in $\mathbb{A}_f$. However, we will not in general assume that the converse is true and therefore we will discuss, $\mathbf{F}$, as lacking an inverse, but we will also examine specific situations where this function does have an inverse. Please see Fig. 2.

From a numerical weather prediction perspective, we view eq. (4) much like that of the algorithms in the ensemble post-processing and bias-correction literature (e.g. Glahn and Lowry, 1972; Raftery et al., 2005) in which climatological information is used to build relationships between forecasts and observations. Additionally, we view eq. (4) as having both a practical as well as a pedagogical application. From a practical perspective, one may wish to deduce the relationship [eq. (4)] from some climatological archive of forecast-truth pairs in order to implement the data assimilation algorithm discussed in this manuscript. By contrast, and from a pedagogic perspective, one may wish to simply assert a particular relationship in eq. (4) and subsequently use the framework presented below to understand its implications. This last tack will be the one illustrated in this manuscript as we will focus in Sections 3 and 4 on a linear map in eq. (4) that is to be interpreted as a smoothing operator as this was used in previous work in the representation error literature (e.g. Liu and Rabier, 2002; Waller et al., 2013). In a sequel to this manuscript, we will illustrate estimation techniques for $\mathbf{F}$ and the results of its application to different cycling data assimilation experiments.

Equation (4) allows for the definition of several new densities in this problem that will prove to be useful tools in the subsequent analysis. Equation (4) implies that the

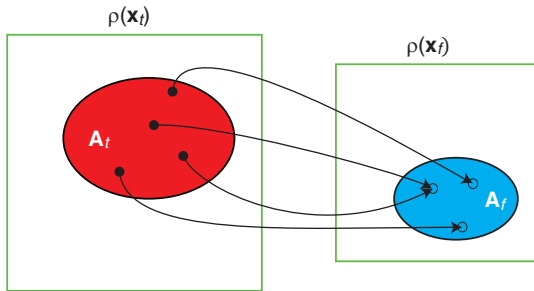


Fig. 2. The attractors and their map. The two green squares represent the domain of the true states (large square) and the forecast states (small square). The region for which $\rho(\mathbf{x}_t) > 0$ ($\rho(\mathbf{x}_f) > 0$) is denoted as the red (blue) shaded region. An example of the function in eq. (4) is denoted by the arrows, which travel from states in $\mathbb{A}_t$ denoted by filled circles to states in $\mathbb{A}_f$ denoted by open circles. Note that multiple states in $\mathbb{A}_t$ may map to the same point in $\mathbb{A}_f$. We will show that this results in an error of representation.

density that describes the distribution of states on the forecast attractor given a state on the true attractor, which we will refer to as a *conversion density*, is the Dirac delta function:

$$\rho(\mathbf{x}_f|\mathbf{x}_t) = \delta(\mathbf{x}_f - \mathbf{F}(\mathbf{x}_t)). \quad (5)$$

This is because given a particular true state $\mathbf{x}_t$ there is one and only one forecast state $\mathbf{x}_f$ that it is related to and hence there is no uncertainty in the position on the forecast attractor given a particular realisation of the state on the true attractor. It is important to realise that eq. (4) implies that while $\rho(\mathbf{x}_f|\mathbf{x}_t)$ has zero variance that in general the *converse* conversion density $\rho(\mathbf{x}_t|\mathbf{x}_f)$ has a non-zero variance because $\mathbf{F}$ does not necessarily have an inverse. Last, the assumption that the relationship between the two attractors is of the form [eq. (4)] allows for the subsequent simplification in eq. (5) that the conversion density is simply the Dirac delta function, but this assumption is technically not required by the following theory and is largely used to allow greater focus on the relationship between the lack of an inverse in eq. (4) and the representation error.

Bayes' rule tells us that the two conversion densities discussed above can be related to each other through their associated climatological distributions as

$$\rho(\mathbf{x}_f|\mathbf{x}_t)\rho(\mathbf{x}_t) = \rho(\mathbf{x}_t|\mathbf{x}_f)\rho(\mathbf{x}_f). \quad (6)$$

We may understand a little bit about the structure of the converse conversion density $\rho(\mathbf{x}_t|\mathbf{x}_f)$ by solving eq. (6):

$$\rho(\mathbf{x}_t|\mathbf{x}_f) = w(\mathbf{x}_t, \mathbf{x}_f)\delta(\mathbf{x}_t - \mathbf{F}(\mathbf{x}_f)), \quad (7)$$

where

$$w(\mathbf{x}_t, \mathbf{x}_f) = \frac{\rho(\mathbf{x}_t)}{\rho(\mathbf{x}_f)} \quad (8)$$

First, eq. (7) shows that when $\mathbf{F}$ does not have an inverse, the density $\rho(\mathbf{x}_t|\mathbf{x}_f)$ is a weighted collection of Dirac delta functions on the surface defined by fixing $\mathbf{x}_f$ and determining the collection of states $\mathbf{x}_t$ for which $\mathbf{F}(\mathbf{x}_t) = \mathbf{x}_f$. Note that because $\mathbf{F}$ only produces states on $\mathbb{A}_f$ we never need evaluate $w(\mathbf{x}_t, \mathbf{x}_f)$ for values of $\mathbf{x}_f$ for which $\rho(\mathbf{x}_f) = 0$. Second, as discussed in subsection 2.1, because the observations in the data assimilation cycles from times $j = 1, 2, \dots, J$ do not change the set $\mathbb{A}_t$, and the map [eq. (4)] is valid for all $\mathbb{A}_t$ and $\mathbb{A}_f$, we may apply the map [eq. (4)], and its corresponding conversion densities, to all data assimilation cycles, j .

To illustrate the properties of these conversion densities, we relate these densities to the two-vector example of subsection 2.1. Here, we define a forecast (low-resolution) state to be a scalar that arises from a smoothing operator that

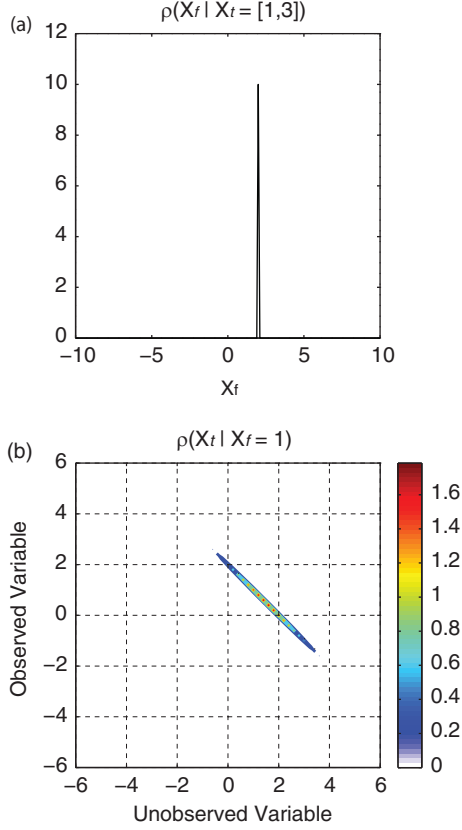


Fig. 3. The conversion densities. (a) The conversion density describing the distribution of forecast states given a state on the true attractor ($\mathbf{x}_t = [1 \ 3]^T$), and (b) the conversion density describing the distribution of true states given a state on the forecast attractor ($x_f = 1$).

is an arithmetic mean: $\mathbf{F} = [0.5 \ 0.5]$. We use this operator in eq. (4) to define the forecast (scalar) states that are obtained from the high-resolution true (two-vector) states. Equation (5) is plotted in Fig. 3a for an example high-resolution state of $\mathbf{x}_t = [1 \ 3]^T$, which implies a low-resolution state of $x_f = 2$ because $\mathbf{F}\mathbf{x}_t = 2$ in this case. Note that the delta function in Fig. 3a is placed at $x_f = 2$ and has amplitude 10. The amplitude of 10 arises because we have defined integration numerically here as the trapezoidal rule. For the trapezoidal rule, the Dirac delta function is inversely proportional to the grid resolution that we are using to represent these densities. Here we have used a grid resolution of 0.1, which implies that the Dirac delta has amplitude 10. As described in the previous section, while $\rho(x_f | \mathbf{x}_t = [1 \ 3]^T)$ is the Dirac delta function, the converse conversion density $\rho(\mathbf{x}_t | x_f)$ is not. As an example, the converse conversion density, $\rho(\mathbf{x}_t | x_f = 1)$, has non-zero probability on a line defined by $\mathbf{F}\mathbf{x}_t = 1$ weighted by the climatological distributions through w and is shown in Fig. 3b.

These conversion densities are useful because they allow one to convert between the true and forecast densities. For

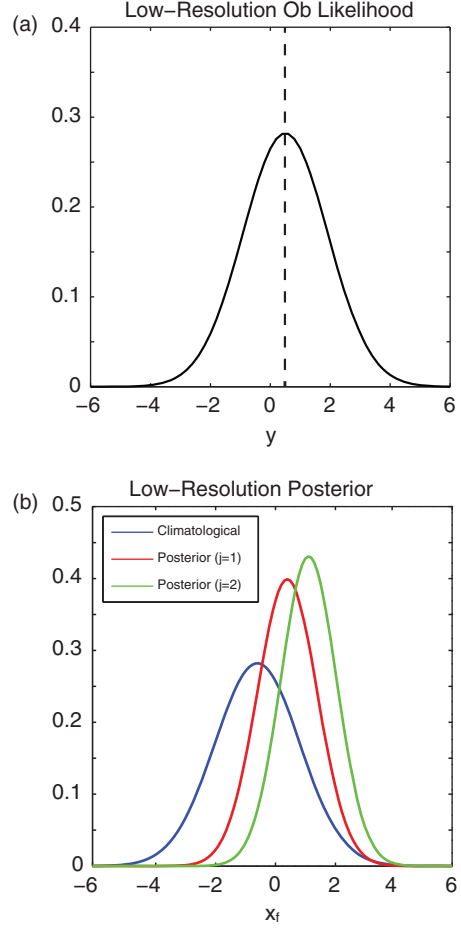


Fig. 4. Data assimilation on the low-resolution (forecast) attractor. The low-resolution (a) observation likelihood and (b) climatological and posterior distributions for $j = 1$ and 2.

example, the prior density on the attracting manifold of the forecast model is:

$$\rho(\mathbf{x}_f | \mathbf{Y}_{J-1}) = \int_{-\infty}^{\infty} \rho(\mathbf{x}_f | \mathbf{x}_t) \rho(\mathbf{x}_t | \mathbf{Y}_{J-1}) d\mathbf{x}_t. \quad (9)$$

Equation (9) describes the distribution of forecasts $\mathbf{x}_f$ that obtains from sampling $\rho(\mathbf{x}_t | \mathbf{Y}_{J-1})$ for $\mathbf{x}_t$ and using these samples of $\mathbf{x}_t$ in eq. (4) to obtain samples of $\mathbf{x}_f$. We may apply eq. (9) to our simple example problem to find the prior distribution of forecasts for the $j = 2$ cycle (Fig. 4b). Recall that the prior distribution for our simple example problem is the previous ($j = 1$) posterior, $\rho(\mathbf{x}_f | y_1 = 1)$. Note that the mode of this low-resolution prior is not the mode of either variable in the high-resolution prior. In fact, the mode of the low-resolution prior is the arithmetic mean of the mode of each high-resolution variable in the high-resolution prior.

The conversion densities are central to our development because through them we are able to build one of the most important components of this work. These conversion densities allow us to show that the forecasts have their own posterior density defined as

$$\rho(\mathbf{x}_f | \mathbf{Y}_J) = \int_{-\infty}^{\infty} \rho(\mathbf{x}_f | \mathbf{x}_t) \rho(\mathbf{x}_t | \mathbf{Y}_J) d\mathbf{x}_t, \quad (10)$$

which, upon using (3) and (5), implies

$$\rho(\mathbf{x}_f | \mathbf{Y}_J) = C_J \rho(\mathbf{y}_J | \mathbf{x}_f) \rho(\mathbf{x}_f | \mathbf{Y}_{J-1}), \quad (11)$$

where $\rho(\mathbf{y}_J | \mathbf{x}_f)$ is the conditional density for the observation conditioned on the forecast state (which we shall refer to as the *forecast* observation likelihood) and $\rho(\mathbf{x}_f | \mathbf{Y}_J)$ represents the density that describes the conditional distribution of forecast states given the entire observational record. This notion that the states on the forecast attractor have a posterior density that has been conditioned upon observations of the state on a different attractor is a unique aspect of this work and will allow us to explicitly show how the difference in the two attractors leads to an error of representation.

Equation (11) shows that the data assimilation procedure, starting from the climatological distribution and proceeding for J assimilation steps, applies to the forecast model in so far as one simply replaces the true states, $\mathbf{x}_t$, with the forecast states, $\mathbf{x}_f$, in eqs. (2) and (3) to define Bayesian data assimilation on the forecast attractor. It is however important to recognise that while eq. (11) looks superficially similar to eq. (3), it is in fact quite different. This is true for two reasons: (1) the data assimilation procedure for the forecast model begins with the forecast climatological distribution, which may be significantly different from the true climatological distribution, and (2) the forecast observation likelihood is actually very different from the true observation likelihood. This difference between the true observation likelihood and the forecast observation likelihood will be described next.

2.3. The forecast observation likelihood

Central to understanding the forecast posterior distribution, and the manifestation of representation error within it, is the forecast observation likelihood. The observation likelihood in eq. (11) obtains from application of the chain rule of probability through the use of the conversion density as

$$\rho(\mathbf{y}_J | \mathbf{x}_f) = \int_{-\infty}^{\infty} \rho(\mathbf{y}_J | \mathbf{x}_t) \rho(\mathbf{x}_t | \mathbf{x}_f) d\mathbf{x}_t. \quad (12)$$

An important difference between $\rho(\mathbf{y}_J | \mathbf{x}_f)$ and $\rho(\mathbf{y}_J | \mathbf{x}_t)$ is that while the mean of $\rho(\mathbf{y}_J | \mathbf{x}_t)$ is the true state ($\mathbf{H}\mathbf{x}_t$)

and its variance is the instrument error, $\mathbf{R}_i$, this is not true of $\rho(\mathbf{y}_J | \mathbf{x}_f)$. The mean of $\rho(\mathbf{y}_J | \mathbf{x}_f)$ is, after use of eq. (12):

$$\mathbf{y}_f(\mathbf{x}_f) = \int_{-\infty}^{\infty} \mathbf{y}_J \rho(\mathbf{y}_J | \mathbf{x}_f) d\mathbf{y}_J = \mathbf{H}\bar{\mathbf{x}}_c, \quad (13)$$

$$\bar{\mathbf{x}}_c(\mathbf{x}_f) = \int_{-\infty}^{\infty} \mathbf{x}_t \rho(\mathbf{x}_t | \mathbf{x}_f) d\mathbf{x}_t, \quad (14)$$

where we will explicitly write this as a state-dependent bias (from the perspective of the forecast model attractor and its associated prior) as

$$\mathbf{y}_f(\mathbf{x}_f) = \mathbf{H}_f \mathbf{x}_f + \mathbf{b}, \quad (15)$$

with this bias defined as

$$\mathbf{b} = \mathbf{H}\bar{\mathbf{x}}_c - \mathbf{H}_f \mathbf{x}_f, \quad (16)$$

and $\mathbf{H}_f$ ($p \times M$) is the observation operator in the space of the forecast model. At this point, we will leave the definition and the distinction between $\mathbf{H}$ and $\mathbf{H}_f$ undefined. We emphasise here however that the situation we will examine is one in which the difference between $\mathbf{H}$ and $\mathbf{H}_f$ is because $\mathbf{H}_f$ operates on a truncated state vector and is not because $\mathbf{H}_f$ has been approximated or created in error. Nevertheless, we will see subsequently that this difference between these two observation operators will turn out to be a central feature of the analysis and therefore we shall return to discuss their differences at several points below.

Equations (13) and (14) show that the observation conditioned on the forecast is biased with respect to the truth (because the observation is of an object on the true attracting manifold and not on the forecast attracting manifold) and that bias is described by the mean of the conversion density. In addition, the variance about the mean state, $\mathbf{y}_f(\mathbf{x}_f)$, is

$$\begin{aligned} \mathbf{R}_f(\mathbf{x}_f) &= \int_{-\infty}^{\infty} (\mathbf{y}_J - \mathbf{y}_f)(\mathbf{y}_J - \mathbf{y}_f)^T \rho(\mathbf{y}_J | \mathbf{x}_f) d\mathbf{y}_J, \\ &= \mathbf{R}_i + \mathbf{H}\mathbf{P}_c\mathbf{H}^T \end{aligned} \quad (17)$$

where

$$\mathbf{P}_c(\mathbf{x}_f) = \int_{-\infty}^{\infty} (\mathbf{x}_t - \bar{\mathbf{x}}_c)(\mathbf{x}_t - \bar{\mathbf{x}}_c)^T \rho(\mathbf{x}_t | \mathbf{x}_f) d\mathbf{x}_t, \quad (18)$$

is the covariance matrix of the conversion density. Equation (18) carries the information that relates the forecast model states to the true attracting manifold. The term $\mathbf{H}\mathbf{P}_c\mathbf{H}^T$ is the manifestation of the representation error in the observation covariance matrix and shows that representation error occurs when the function $\mathbf{F}$ does not have

an inverse. To see this note that when the function $\mathbf{F}$ has an inverse, then $w = 1$ and

$$\rho(\mathbf{x}_t | \mathbf{x}_f) = \delta(\mathbf{x}_f - \mathbf{F}(\mathbf{x}_t)) = \delta(\mathbf{x}_t - \mathbf{F}^{-1}(\mathbf{x}_f)), \quad (19)$$

which clearly has a variance of zero and hence eq. (18) would vanish. Furthermore, note that if we use eq. (19) in eqs. (13) and (14), then we obtain $\mathbf{y}_f = \mathbf{H}\mathbf{F}^{-1}(\mathbf{x}_f)$, which is biased from the perspective of the forecast attractor, where that bias is $\mathbf{b} = \mathbf{H}\mathbf{F}^{-1}(\mathbf{x}_f) - \mathbf{H}_f\mathbf{x}_f$. We may however remove this bias by defining the observation operator in the space of the forecast model as $\mathbf{H}_f = \mathbf{H}\mathbf{F}^{-1}$. This definition of the forecast observation operator is novel as it implies that we extend our view of what an observation operator does. The observation operator, $\mathbf{H}_f = \mathbf{H}\mathbf{F}^{-1}$, implies that the forecast observation operator should include a ‘bias’ correction for the particular forecast model that we are using. In Section 4, this definition of the forecast observation operator will be generalised to linear equations [eq. (4)] that do not contain an explicit inverse and subsequently shown to be the proper way to account for representation error. In the meantime, however, we will maintain the general form for $\mathbf{H}_f$ that we have been using thus far.

In order to understand the forecast observation likelihood better, we apply our example problem to eq. (12) to find the low-resolution observation likelihood. This low-resolution observation likelihood is plotted in Fig. 4a as a function of observation and conditioned on the low-resolution state of $x_f = 1$. Because the observation is actually of the high-resolution state, the low-resolution observation likelihood is biased (with respect to the low-resolution states), which can be seen by the mode of the distribution not being centred on the low-resolution state of $x_f = 1$. It is actually centred at $x_f = 1/2$, where this bias may be seen as coming from the bias in the forecast climatological distribution whose mean is at $\bar{x}_f = -1/2$. The forecast observation likelihood has a variance of 2, of which $1/2$ of this variance (recall that the instrument error is 1) can be attributed to representation error and eq. (18). Finally, the low-resolution posterior [eqs. (10) and (11)] for the observation $j = 1$ and 2 cycles is plotted in Fig. 4b. The posterior for the $j = 1$ cycle has its mode at $1/2$, and the mode of the posterior for the $j = 2$ cycle is at 1.2 . Again, the mode at each cycle is located at the arithmetic mean of the location of the modes in the true (high-resolution) posteriors.

2.4. Data assimilation

Because we are interested in doing data assimilation with these concepts, we will now derive the minimum error variance estimate of a linear estimator on the forecast models attractor, which, in this context, is the Kalman filter

(Kalman, 1960) for the states on the forecast attractor. This formula is derived by minimising the trace of the *expected* posterior forecast covariance matrix about a linear estimator. The expected posterior forecast covariance matrix may be written as:

$$\bar{\mathbf{P}}_a = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (\mathbf{x}_f - \bar{\mathbf{x}}_a)(\mathbf{x}_f - \bar{\mathbf{x}}_a)^T \rho(\mathbf{x}_f | \mathbf{Y}_J) d\mathbf{x}_f \rho(\mathbf{y}_J | \mathbf{Y}_{J-1}) d\mathbf{y}_J, \quad (20)$$

where the analysis update equation, whose ‘error’ variance is being measured, takes the form of a linear estimation equation

$$\bar{\mathbf{x}}_a = \bar{\mathbf{x}}_f + \mathbf{G}[\mathbf{v} - \langle \mathbf{v} \rangle], \quad (21)$$

where

$$\bar{\mathbf{x}}_f = \int_{-\infty}^{\infty} \mathbf{x}_f \rho(\mathbf{x}_f | \mathbf{Y}_{J-1}) d\mathbf{x}_f = \int_{-\infty}^{\infty} \mathbf{F}(\mathbf{x}_t) \rho(\mathbf{x}_t | \mathbf{Y}_{J-1}) d\mathbf{x}_t. \quad (22)$$

The overbar on $\bar{\mathbf{P}}_a$ in eq. (20) is there to make clear that this is the expression for the expected posterior covariance matrix as it has been averaged over all the observations [see Hodyss and Campbell (2013) for further discussion]. The innovation in eq. (21) is $\mathbf{v} = \mathbf{y}_J - \mathbf{H}_f \bar{\mathbf{x}}_f$, which we emphasise uses the observation operator in the space of the forecast model. The expected innovation, $\langle \mathbf{v} \rangle$, in eq. (21) is there because the observations with respect to the forecasts are biased because the observations are taken on the true attracting manifold. This de-biasing of the innovation in eq. (21) is critical to get the data assimilation to put the analysis at the centre of $\rho(\mathbf{x}_f | \mathbf{Y}_J)$ and therefore the posterior distribution on the desired forecast attractor.

It is important to realise that while the quantities derived in eqs. (13) and (14) through eq. (17) are interesting descriptions of the corresponding densities, they are in fact not the ones that would be used in a data assimilation algorithm. This is because as shown in eq. (20) those expressions need to be averaged over the forecasts in the derivation of the update equations. This averaging for eqs. (13) and (14) would take the form:

$$\bar{\mathbf{y}}_f = \int_{-\infty}^{\infty} \mathbf{y}_f(\mathbf{x}_f) \rho(\mathbf{x}_f | \mathbf{Y}_{J-1}) d\mathbf{x}_f = \mathbf{H}_f \bar{\mathbf{x}}_f + \bar{\mathbf{b}}, \quad (23)$$

where $\bar{\mathbf{b}} = \langle \mathbf{v} \rangle = \mathbf{H} \bar{\mathbf{x}}_t - \mathbf{H}_f \bar{\mathbf{x}}_f$ and

$$\bar{\mathbf{x}}_t = \int_{-\infty}^{\infty} \mathbf{x}_t \rho(\mathbf{x}_t | \mathbf{Y}_{J-1}) d\mathbf{x}_t. \quad (24)$$

In addition, the observation error variance that would be used in a data assimilation algorithm is obtained from

$$\bar{\mathbf{R}}_f = \int_{-\infty}^{\infty} \mathbf{R}_f(\mathbf{x}_f) \rho(\mathbf{x}_f | \mathbf{Y}_{J-1}) d\mathbf{x}_f. \quad (25)$$

Note that the result of the integration in eq. (25) is an observation error variance that is not state-dependent. The derivation of the Kalman gain requires the use of an observation error covariance matrix that is the result of a weighted average over all possible observation error covariance matrices and therefore cannot depend on the state.

Equation (25) is interesting because we may use eq. (17) in eq. (25) to obtain:

$$\bar{\mathbf{R}}_f = \mathbf{R}_i + \int_{-\infty}^{\infty} \mathbf{H} \mathbf{P}_c \mathbf{H}^T \rho(\mathbf{x}_f | \mathbf{Y}_{J-1}) d\mathbf{x}_f, \quad (26)$$

where

$$\begin{aligned} \int_{-\infty}^{\infty} \mathbf{H} \mathbf{P}_c \mathbf{H}^T \rho(\mathbf{x}_f | \mathbf{Y}_{J-1}) d\mathbf{x}_f &= \mathbf{H} \mathbf{P}_t \mathbf{H}^T - \mathbf{H}_f \mathbf{P}_f \mathbf{H}_f^T - \mathbf{H}_f \mathbf{P}_{fb} \\ &\quad - \mathbf{P}_{fb}^T \mathbf{H}_f^T - \mathbf{P}_{bb}, \end{aligned} \quad (27)$$

$$\mathbf{P}_t = \int_{-\infty}^{\infty} (\mathbf{x}_t - \bar{\mathbf{x}}_t)(\mathbf{x}_t - \bar{\mathbf{x}}_t)^T \rho(\mathbf{x}_t | \mathbf{Y}_{J-1}) d\mathbf{x}_t, \quad (28)$$

$$\mathbf{P}_f = \int_{-\infty}^{\infty} (\mathbf{x}_f - \bar{\mathbf{x}}_f)(\mathbf{x}_f - \bar{\mathbf{x}}_f)^T \rho(\mathbf{x}_f | \mathbf{Y}_{J-1}) d\mathbf{x}_f, \quad (29)$$

$$\mathbf{P}_{fb} = \int_{-\infty}^{\infty} (\mathbf{x}_f - \bar{\mathbf{x}}_f)(\mathbf{b} - \bar{\mathbf{b}})^T \rho(\mathbf{x}_f | \mathbf{Y}_{J-1}) d\mathbf{x}_f, \quad (30)$$

$$\mathbf{P}_{bb} = \int_{-\infty}^{\infty} (\mathbf{b} - \bar{\mathbf{b}})(\mathbf{b} - \bar{\mathbf{b}})^T \rho(\mathbf{x}_f | \mathbf{Y}_{J-1}) d\mathbf{x}_f. \quad (31)$$

Equation (27) shows that representation error is the difference between the true covariance, the covariance of the forecasts and the bias, and the covariance matrix of the state-dependent bias all mapped to observation space.

The steps required to perform the minimisation of the trace of eq. (20) are well known, can be found in Hodyss (2011) and will not be repeated here. The result of this minimisation is that the *gain matrix*, $\mathbf{G}$, can be written in two equivalent ways:

$$\mathbf{G}_1 = \mathbf{P}_f \mathbf{H}^T [\mathbf{H} \mathbf{P}_t \mathbf{H}^T + \mathbf{R}_i]^{-1}, \quad (32)$$

$$\begin{aligned} \mathbf{G}_2 &= [\mathbf{P}_f \mathbf{H}_f^T + \mathbf{P}_{fb}] \\ &\quad \times [\mathbf{H}_f \mathbf{P}_f \mathbf{H}_f^T + \mathbf{H}_f \mathbf{P}_{fb} + \mathbf{P}_{fb}^T \mathbf{H}_f^T + \mathbf{P}_{bb} + \bar{\mathbf{R}}_f]^{-1}, \end{aligned} \quad (33)$$

where

$$\mathbf{P}_{fi} = \int_{-\infty}^{\infty} (\mathbf{F}(\mathbf{x}_t) - \bar{\mathbf{x}}_f)(\mathbf{x}_t - \bar{\mathbf{x}}_t)^T \rho(\mathbf{x}_t | \mathbf{Y}_{J-1}) d\mathbf{x}_t, \quad (34)$$

$$\mathbf{P}_{fb} \mathbf{H}^T = \mathbf{P}_f \mathbf{H}_f^T + \mathbf{P}_{fb}. \quad (35)$$

The first gain matrix, $\mathbf{G}_1$, uses information from the true attractor that is subsequently mapped back to the forecast attractor through the covariance between the states on the two attractors, $\mathbf{P}_{fi}$. This version of the gain is most similar to the traditional Kalman gain whereby the numerator relates the (true attractor's) observation space to the state space to be updated (forecast attractor) using the covariance [eq. (34)]. This version of the gain matrix is of course impossible to implement in practice as it requires use of the true prior distribution.

The second gain matrix, $\mathbf{G}_2$, is novel and only uses the information on the forecast model attractor along with the function $\mathbf{F}$ to infer the relationship between the true and the forecast attractors. This second gain matrix has two terms in its numerator. The first term is the standard term that maps the observation space to the state space to be updated, but totally from within the forecast attractor. The second term is new and accounts for the possibility that on the forecast attractor there is no covariance between the observation space and state space to be updated but, because the observation is actually of a state on the true attractor and there may be a covariance between the true attractor and the forecast state space to be updated, there should in fact be a correction at that location. This new term is shown in eq. (35) to be precisely the difference between the forecast state-space/true observation-space covariance ($\mathbf{P}_{fi} \mathbf{H}^T$) and the forecast state-space/forecast observation-space covariance ($\mathbf{P}_f \mathbf{H}_f^T$). Equation (35) shows that the numerators of eqs. (32) and (33) would be identical if the forecast-bias covariance matrix, $\mathbf{P}_{fb}$, would vanish. In Section 4, we show how to choose the forecast observation operator to eliminate this forecast-bias covariance matrix.

Last, it is important to realise that the gain matrices [eqs. (32) and (33)] when used in eq. (21) do not in fact provide an estimate of the true state on the true attractor. The gain matrices [eqs. (32) and (33)] when used in eq. (21) find the state estimate that is closest (in the sense of the function $\mathbf{F}$ and in mean-square) to the *forecast* that obtains from mapping the truth through eq. (4). Hence, the gain matrices [eqs. (32) and (33)] push the state estimate from the true attracting manifold onto the forecast attracting manifold and is actually likely to push the state estimate further from the observations than it would have been if we had not altered the numerator and denominator of the gain matrix to account for this error in representation. The benefit from using these new gain matrices is therefore

not their distance from the true attractor but actually the fact that they produce a state estimate near to the forecast attractor, which is likely to produce a better, and more balanced, forecast.

3. Smoothing operators

Previous work in the representation error literature has examined the situation where the relationship between the states being observed by the observational instruments and that produced by the forecast model are related by a smoothing operator (e.g. Mitchell and Daley, 1997a, 1997b; Liu and Rabier, 2002; Waller et al., 2013). Following this work, we will assume that the forecast model resolution is coarse as compared to the true model resolution, that is, $M < N$. In addition, we will assume the function $\mathbf{F}$ is linear, but we make no assumptions on the shape of the prior distributions on either attractor.

The function $\mathbf{F}$ will be assumed to be a linear matrix operator that acts to ‘smooth’ the true state to the resolution of the forecast model, viz.

$$\mathbf{x}_f = \mathbf{S}\mathbf{x}_t, \quad (36)$$

where $\mathbf{S}$ is an $M \times N$ matrix whose singular value decomposition results in

$$\mathbf{S} = \mathbf{U}[\Lambda^{1/2} \quad \mathbf{0}]\mathbf{V}^T, \quad (37)$$

with the left singular vectors contained in $\mathbf{U}$ ($M \times M$), right singular vectors in $\mathbf{V}$ ($N \times N$), Λ ($M \times M$) the diagonal matrix of singular values and $\mathbf{0}$ [$M \times (N - M)$] the zero matrix. The bracket in eq. (37) is a truncation operator, and along with Λ , which we assume to be either constant (white) or steeply sloped (red), we view as a ‘smoothing’ operator. We attach to this notion of ‘smoothing’ the philosophy prevalent in numerical modelling (e.g. Lilly, 2005) in which forecast model output is assumed to represent grid-cell averages of the true state around each nodal point of the model’s grid representation of the spatial domain. Note that the statement [eq. (36)] states that the only difference between the forecast and the truth is through the smoothing of that state, which ignores the fact that a truncated model will differ in its cascades of energy and enstrophy (i.e. the nonlinear interaction between scales and their subsequent evolution).

Because $\mathbf{S}$ is $M \times N$ and $M < N$, the matrix $\mathbf{S}$ will not generally have an inverse. This has important implications for the structure of $\rho(\mathbf{x}_t|\mathbf{x}_f)$ as it implies that there exists a hyperplane defined by all of the states, $\mathbf{x}_t$, that satisfy eq. (36) for a given forecast state $\mathbf{x}_f$. Along this plane, $\rho(\mathbf{x}_t|\mathbf{x}_f)$ has non-zero probability, and this results in a non-zero variance of the converse conversion density, which as shown in eq. (18) leads to the error of representation.

An explicit example of $\rho(\mathbf{x}_t|\mathbf{x}_f)$ that had non-zero variance was presented in Fig. 3, and in this case one could see in that figure that the hyperplane alluded to above was reduced to a line in the two-dimensional plane of $\mathbf{x}_t$.

The smoothing operator in eq. (37) will lead to a bias in the observation mean [eq. (23)]. This bias leads to the expected innovation being

$$\langle \mathbf{v} \rangle = \bar{\mathbf{b}} = (\mathbf{H} - \mathbf{H}_f\mathbf{S})\bar{\mathbf{x}}_t. \quad (38)$$

The bias results from the difference between the true prior mean and the ‘smoothed’ one obtained after use of eq. (36).

Similarly, the use of eq. (36) in the *gain matrix*, $\mathbf{G}$ [eqs. (32) and (33)], obtains:

$$\mathbf{G}_1 = \mathbf{S}\mathbf{P}_t\mathbf{H}^T[\mathbf{H}\mathbf{P}_t\mathbf{H}^T + \mathbf{R}_t]^{-1} = \mathbf{S}\mathbf{G}_t, \quad (39)$$

$$\mathbf{G}_2 = [\mathbf{P}_f\mathbf{H}_f^T + \mathbf{P}_{fb}][\mathbf{H}_f\mathbf{P}_f\mathbf{H}_f^T + \bar{\mathbf{R}}_f]^{-1}, \quad (40)$$

where $\mathbf{P}_t$ is the covariance matrix of the true prior,

$$\mathbf{P}_f = \mathbf{S}\mathbf{P}_t\mathbf{S}^T, \quad (41)$$

$$\mathbf{P}_{fb} = \mathbf{S}\mathbf{P}_t(\mathbf{H} - \mathbf{H}_f\mathbf{S})^T, \quad (42)$$

$$\bar{\mathbf{R}}_f = \mathbf{R}_t + (\mathbf{H} - \mathbf{H}_f\mathbf{S})\mathbf{P}_t(\mathbf{H} - \mathbf{H}_f\mathbf{S})^T - \mathbf{P}_{bb}, \quad (43)$$

$$\bar{\mathbf{R}}_f^* = \mathbf{H}_f\mathbf{P}_{fb} + \mathbf{P}_{fb}^T\mathbf{H}_f^T + \mathbf{P}_{bb} + \bar{\mathbf{R}}_f = \mathbf{R}_t + \mathbf{H}\mathbf{P}_t\mathbf{H}^T - \mathbf{H}_f\mathbf{P}_f\mathbf{H}_f^T. \quad (44)$$

In eq. (39), the gain matrix $\mathbf{G}_t$, which is the true gain matrix for the true attractor, is mapped into the forecast space through application of the smoothing matrix $\mathbf{S}$. In eq. (40), the gain matrix is calculated using only quantities known from the forecast distributions and the matrix $\mathbf{S}$. The quantity $\bar{\mathbf{R}}_f^*$ is an abstraction of the observation error covariance matrix to include all the terms except the forecast covariance matrix in observation space in the denominator of eq. (40). We will refer to this quantity as the *effective* observation error covariance matrix. The effective observation error covariance matrix reduces to the sum of the instrument error and the difference between the true and forecast covariance matrices in observation space. Note that the difference between the effective observation covariance matrix, $\bar{\mathbf{R}}_f^*$, and the actual observation covariance matrix, $\bar{\mathbf{R}}_f$, is entirely a result of the matrices $\mathbf{P}_{fb}$ and $\mathbf{P}_{bb}$. Given eqs. (36) and (37), it may be shown that the *trace* ($\mathbf{P}_t/N \geq \text{trace}(\mathbf{P}_f/M)$), and therefore the diagonal of $\bar{\mathbf{R}}_f^*$ is equal to or greater than the instrument error, $\mathbf{R}_t$. This fact about smoothing matrices of the form [eq. (37)] implies that high-resolution models should as a general rule have greater variance than low-resolution models.

This matrix $\bar{\mathbf{R}}_f^*$ is important because it makes the connection between the Kalman gains $\mathbf{G}_1$ and $\mathbf{G}_2$. The gain matrix in eq. (39) operates on the same innovation as the

gain matrix in eq. (40). This implies that the innovation variance in the denominator, as calculated both ways, must be equal, viz.

$$\mathbf{H}_f \mathbf{P}_f \mathbf{H}_f^T + \bar{\mathbf{R}}_f^* = \mathbf{H}_f \mathbf{P}_f \mathbf{H}_f^T + \mathbf{R}_f. \quad (45)$$

This relationship may be proven by using eq. (44) in eq. (45). This shows that if we define the error of representation as a property of the covariance matrix of the forecast observation likelihood, then one cannot actually deduce it straightforwardly from the innovation variance (e.g. Hollingsworth and Lönnberg, 1986; Desroziers et al., 2006). The object that can be deduced directly from the innovation variance is $\bar{\mathbf{R}}_f^*$ and not $\bar{\mathbf{R}}_f$. More discussion of the differences between $\bar{\mathbf{R}}_f^*$ and $\bar{\mathbf{R}}_f$ will be presented in Section 4.

4. Gaussian statistics

In this section, we add to the development of Section 3 the assumption of Gaussianity to the climatological distributions and linearity in the true and forecast model evolution equations. This implies that the prior distributions of both the forecast and true systems must also be Gaussian. It is important to point out that the assumption of Gaussian error statistics in the climatological distributions and linear error evolution implies that the sets $\mathbb{A}_t$ and $\mathbb{A}_f$ are no longer bounded in state space as is implied by Fig. 2 but now extend to include all possible states in their domain.

4.1. Theory

A simple way to construct a Gaussian problem that is amenable to analysis is through the use of a discrete Fourier series representation. To this end, we assert a Gaussian covariance model for the true (high-resolution) states of the form

$$\mathbf{x}_t = \bar{\mathbf{x}}_t + \mathbf{Z}\boldsymbol{\eta}, \quad (46)$$

where $\bar{\mathbf{x}}_t$ is an N -vector, $\mathbf{Z}$ is the square-root of the true covariance matrix,

$$\mathbf{P}_t = \mathbf{Z}\mathbf{Z}^T, \quad (47)$$

and $\boldsymbol{\eta}$ is an N -vector of random numbers drawn from $N(\mathbf{0}, \mathbf{I})$. We construct eq. (47) using a sinusoidal basis in which the columns of $\mathbf{E}$ ($N \times N$) contain the sinusoids such that

$$\mathbf{P}_t = \mathbf{E}\boldsymbol{\Gamma}\mathbf{E}^T, \quad (48)$$

$\boldsymbol{\Gamma}$ is a diagonal matrix whose i th element of the diagonal is defined from $\Gamma_i = ae^{-\alpha^2 k_i^2}$, k_i is the wavenumber associated with the i th basis function of $\mathbf{E}$, and $a = N \sum_{i=1}^N e^{-\alpha^2 k_i^2}$. The parameter α determines the slope of the true spectrum,

where large α is associated with a red spectrum and small α is associated with a white spectrum.

We connect the true (high-resolution) states to the forecast (low-resolution) states through a smoother that operates as:

$$\mathbf{S} = \mathbf{E}_L [\mathbf{D}^{1/2} \mathbf{T} \quad \mathbf{0}] \mathbf{E}^T, \quad (49)$$

where $\mathbf{D}$ ($M \times M$) is a diagonal matrix whose diagonal elements are $d_i = e^{-\beta^2 k_i^2}$, $\mathbf{E}_L$ ($M \times M$) is the low-resolution basis whose columns are also the sinusoids and $\mathbf{T}$ ($M \times M$) is a diagonal matrix with the value $\sqrt{M/N}$ along the diagonal. The matrix $\mathbf{D}$ represents the climatological ‘model error’ on the resolved scales and would be equal to the identity matrix if the forecast model’s climate at the resolved scales was identical to the true model’s climate at those same scales. The matrix implied by the bracket in eq. (49) performs a truncation of the high-resolution basis to the M -dimensional subspace, while the matrix $\mathbf{T}$ assures that the Fourier coefficients calculated from the high-resolution basis are reweighted consistently with respect to the low-resolution basis.

Equation (49) allows for the creation of the low-resolution forecast states from the true states in eq. (46) using eq. (36). This implies that the forecast error covariance matrix may be written as

$$\mathbf{P}_f = \mathbf{S} \mathbf{P}_t \mathbf{S}^T = \mathbf{E}_L [\mathbf{D}^{1/2} \mathbf{T} \quad \mathbf{0}] \boldsymbol{\Gamma} [\mathbf{D}^{1/2} \mathbf{T} \quad \mathbf{0}]^T \mathbf{E}_L^T. \quad (50)$$

Because $\boldsymbol{\Gamma}$ are the true (high-resolution) eigenvalues, eq. (50) shows that the forecast covariance matrix would be correct up to its M eigenvalues if the climatological model error $\mathbf{D}$ could be removed. We show next how to remove this climatological model error by accounting for the error of representation.

Because the true states and the forecast states are Gaussian and their relationship is described by a linear operator, we know that the mean of the converse conversion density $[\rho(\mathbf{x}_t|\mathbf{x}_f)]$ is a linear function of the vector it is conditioned upon. This presents us with a direct method to calculate the prediction of the observation by the forecast state:

$$\mathbf{y}_f(\mathbf{x}_f) = \mathbf{H}\bar{\mathbf{x}}_t = \mathbf{H}\bar{\mathbf{x}}_t + \mathbf{H}\mathbf{G}_p [\mathbf{x}_f - \bar{\mathbf{x}}_f], \quad (51)$$

where

$$\mathbf{G}_p = \mathbf{Z}[\mathbf{S}\mathbf{Z}]^\dagger = \mathbf{S}^\dagger. \quad (52)$$

The superscript $\dagger$ in eq. (52) denotes the Moore-Penrose pseudo-inverse (Golub and Van Loan, 1996). The pseudo-inverse is constructed by finding the singular value decomposition, viz.

$$\mathbf{S}\mathbf{Z} = \mathbf{E}_L [\mathbf{D}^{1/2} \mathbf{T} \quad \mathbf{0}] \boldsymbol{\Gamma}^{1/2} \mathbf{E}^T, \quad (53)$$

and then noting that its pseudo-inverse is therefore

$$[\mathbf{SZ}]^\dagger = \mathbf{E}\mathbf{\Gamma}^{-1/2}[\mathbf{D}^{-1/2}\mathbf{T}^{-1} \quad \mathbf{0}]^T \mathbf{E}_L^T. \quad (54)$$

Equation (54) is the *pseudo*-inverse of eq. (53) in the sense that $\mathbf{SZ}[\mathbf{SZ}]^\dagger \mathbf{r} = \mathbf{r}$ but $[\mathbf{SZ}]^\dagger \mathbf{SZ} \mathbf{q} \neq \mathbf{q}$ for arbitrary vectors $\mathbf{r}$ and $\mathbf{q}$. Please see Appendix A for a brief derivation of eq. (51). It is important to note in eq. (52) that the matrix $\mathbf{Z}$ is cancelled by the pseudo-inverse operation such that the only information required to determine $\mathbf{G}_p$ is the smoothing matrix [eq. (49)]. In practice, one would build eq. (49) by estimating its components using an archive of truth-forecast pairs as is done in bias-correction algorithms (Glahn and Lowry, 1972). Moreover, this implies that we may view eq. (51) as simply a bias-correction algorithm that we build into the data assimilation system to account for the fact that the forecast model's estimate of the observation is, in effect, biased.

Equation (51) is remarkable in so far as this construction allows us to calculate the important quantities from Sections 2 and 3 without the need to explicitly develop the conversion densities. For example, the representation error is therefore:

$$\bar{\mathbf{R}}_f - \mathbf{R}_i = \mathbf{H}\mathbf{E}\mathbf{\Theta}\mathbf{E}^T \mathbf{H}^T, \quad (55)$$

where $\mathbf{\Theta}$ is a diagonal matrix with the value of the diagonal of $\mathbf{\Theta}$ satisfying:

$$\Theta_i = \begin{cases} 0, & i = 1, \dots, M \\ \Gamma_i, & i = M + 1, \dots, N \end{cases}. \quad (56)$$

The term on the right-hand side of eq. (55) arises from eq. (43). This term clearly shows that the representation error [eq. (55)] is simply the portion of the high-resolution spectrum that is missing from the forecast states. Note that for $M = N$, and therefore no truncation, the elements of $\mathbf{\Theta}$ vanish and there is no representation error. This is again a result of there being, in that case, an inverse available when $M = N$ and, as discussed in Section 2, this implies that the representation error must vanish. This point is important because it shows that the climatological model error on the resolved scales ($\mathbf{D}$) is irrelevant to both the existence of representation error and to the structure of the representation error. By contrast, eq. (44) for which we are interested in $\bar{\mathbf{R}}_f^* - \mathbf{R}_i$ does depend on the climatological model error on the resolved scales ($\mathbf{D}$) through eq. (50).

An important quantity in our development of Section 2 was the state-dependent bias of the forecast model's estimate of the observation [eq. (16)], $\mathbf{b}$, which given eq. (51), implies:

$$\mathbf{b} = [\mathbf{HS}^\dagger - \mathbf{H}_f] \mathbf{x}_f + \mathbf{H}[\bar{\mathbf{x}}_t - \mathbf{S}^\dagger \bar{\mathbf{x}}_f]. \quad (57)$$

The first term in eq. (57) is interesting because it corresponds to the mismatch between the forecast model's estimate of the true observation operator $\mathbf{HS}^\dagger$ in eq. (51) and the observation operator we are actually using $\mathbf{H}_f$. We emphasise here that this mismatch is *not* one in which we are implying that $\mathbf{H}_f$ is incorrect in the sense that if, for example, we had a point measurement that there would be some form of inaccuracy in the interpolation to the observation location. Indeed, even in this case of a point measurement, in which we are assuming we have a perfect interpolation, the mismatch inferred by eq. (57) is between the statistically derived observation operator $\mathbf{HS}^\dagger$, which now corresponds to more than just an interpolation, and the operator, $\mathbf{H}_f$.

Equation (57) suggests that if we define $\mathbf{H}_f$ such that,

$$\mathbf{H}_f \equiv \mathbf{HS}^\dagger, \quad (58)$$

we may remove this state-dependent bias term, which subsequently renders $\mathbf{P}_{fb} = \mathbf{0}$ and $\mathbf{P}_{bb} = \mathbf{0}$. This implies that in this data assimilation system the observation operator does *not* simply map the truth to the observation, but rather it maps the forecast to the observation and, because the forecast is in a different portion of state space than the truth, this requires the matrix operator, $\mathbf{HS}^\dagger$ rather than just $\mathbf{H}$.

One of the most important results of choosing eq. (58), and subsequently rendering $\mathbf{P}_{fb} = \mathbf{0}$ and $\mathbf{P}_{bb} = \mathbf{0}$, is that this results in $\bar{\mathbf{R}}_f^* = \bar{\mathbf{R}}_f$, and therefore the choice [eq. (58)] has removed the impact of the climatological model error on the resolved scales, $\mathbf{D}$, in the forecast covariance matrix and therefore also in $\bar{\mathbf{R}}_f^*$. This has two important consequences: (1) innovation-based techniques (e.g. Hollingsworth and Lönnerberg, 1986; Desroziers et al., 2006) for the estimation of $\bar{\mathbf{R}}_f^*$ are therefore self-consistent estimators of the representation error *only when the choice (58) has been implemented* and (2) as we show next this provides a direct way to remove the deleterious impact of the climatological model error on the state estimate from the data assimilation algorithm.

Subsequently, by employing eq. (58), it can be shown that the gain matrix written in terms of the forecast quantities [eq. (40)] is now:

$$\mathbf{G}_2 = \mathbf{P}_f \mathbf{S}^{\dagger T} \mathbf{H}^T [\mathbf{HS}^\dagger \mathbf{P}_f \mathbf{S}^{\dagger T} \mathbf{H}^T + \bar{\mathbf{R}}_f^*]^{-1}. \quad (59)$$

This new gain matrix is the most important result of this manuscript. The calculation of the operator $\mathbf{G}_p$ has allowed for the creation of a new observation operator $\mathbf{HS}^\dagger$ that represents the forecast model's estimate of the observation of states on the true attractor. Subsequently, this has allowed the calculation of the correct *numerator* in eq. (59) that represents the covariance between the states on the

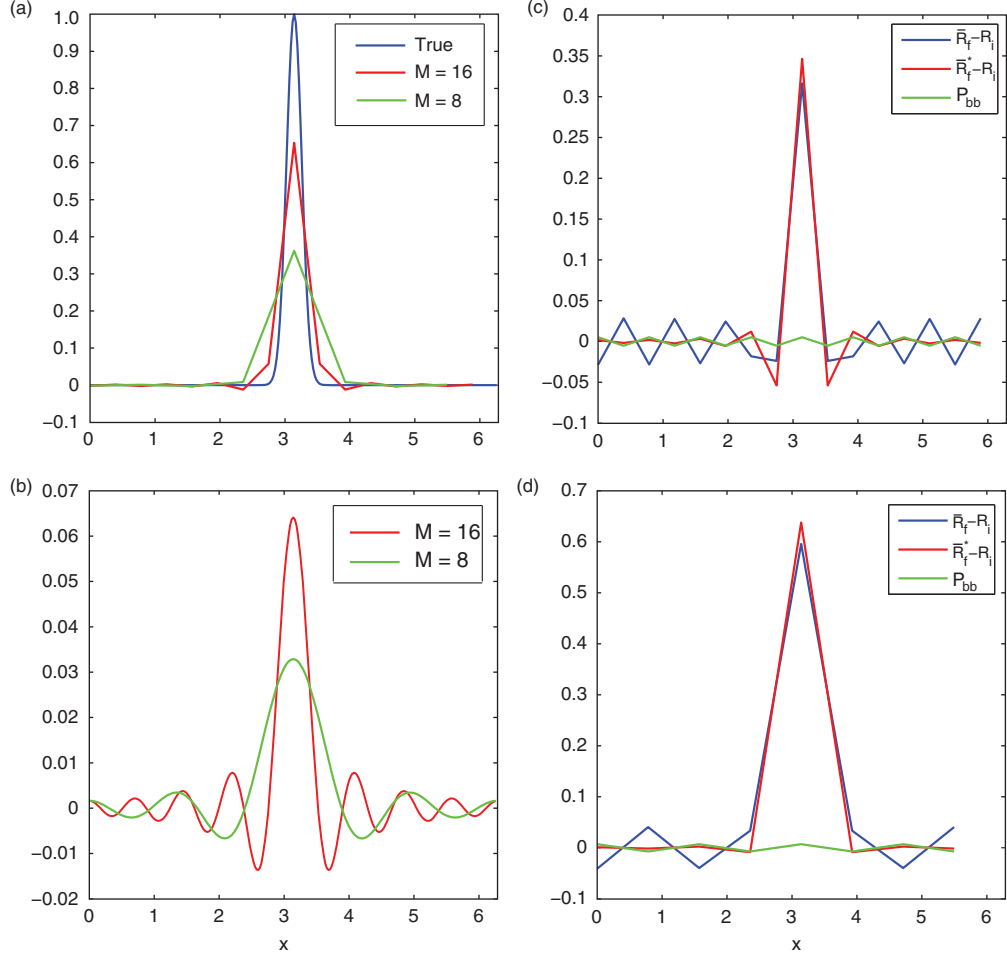


Fig. 5. Pure spectral truncation. (a) Three one-point prior covariance functions: blue is the high-resolution (true) covariance model, red is the $M=16$ low-resolution (forecast) covariance model and green is the $M=8$ low-resolution (forecast) covariance model. (b) One column of the smoother matrix [eq. (62)] for $M=16$ (red) and $M=8$ (green). (c, d) The main components of the theory for $M=16$ and $M=8$, respectively: blue is the representation error covariance ($\bar{\mathbf{R}}_f - \mathbf{R}_i$), red is the *effective* representation error covariance ($\bar{\mathbf{R}}_f^* - \mathbf{R}_i$) and green is the bias covariance matrix ($\mathbf{P}_{bb}$).

forecast attractor and the observations of the true attractor. Equation (58) results in

$$\mathbf{H}_f \mathbf{P}_f \mathbf{H}_f^T = \mathbf{H} \mathbf{S}^\dagger \mathbf{P}_f \mathbf{S}^{\dagger T} \mathbf{H}^T = \mathbf{H} \mathbf{E} \begin{bmatrix} \Gamma_M & \mathbf{0} \\ \mathbf{0} & \mathbf{0} \end{bmatrix} \mathbf{E}^T \mathbf{H}^T, \quad (60)$$

$$\langle \mathbf{v} \rangle = \mathbf{H} \mathbf{E} \begin{bmatrix} \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \mathbf{I} \end{bmatrix} \mathbf{E}^T \bar{\mathbf{x}}_i, \quad (61)$$

where Γ_M denotes the first M eigenvalues of Γ , and $\mathbf{0}$ is the zero matrix. Equation (60) shows that the choice [eq. (58)] results in $\mathbf{H}_f \mathbf{P}_f \mathbf{H}_f^T$ being correct up to the resolution of the forecast model in the sense that the climatological model error denoted by $\mathbf{D}$ has been removed. Equation (51), which describes the bias in the innovation owing to the unresolved scales of motion, is a result of the scales in the true prior mean that are missing in the forecast prior mean.

4.2. Spatially extended example

We now explore the theory discussed above for the spatially extended problem described in this section by employing some simple example problems. We will not redefine the observation operator in order to clearly show the differences between $\bar{\mathbf{R}}_f - \mathbf{R}_i$ and $\bar{\mathbf{R}}_f^* - \mathbf{R}_i$, which will underscore the importance of applying eq. (58), as it was already proven above to remove this difference. We will define $\mathbf{H}$ to be the operator that observes each point of the low-resolution (forecast) state. Hence, in these experiments, $\mathbf{H}_f$ will simply be the identity operator for the point measurements we have available. We emphasise again that employing eq. (58) would lead to an observation operator for these point measurements that is not the identity operator even though these are point measurements, which as we have proven above corrects the problems to be

illustrated next. In order to maintain a connection with previous work (e.g. Mitchell and Daley, 1997a, 1997b; Liu and Rabier, 2002; Waller et al., 2013), we begin by defining the low-resolution (forecast) states as obtained from a smoother that operates as a pure spectral truncation such that we set $\beta = 0$ in eq. (49), which obtains $\mathbf{D} = \mathbf{I}$ and

$$\mathbf{S} = \mathbf{E}_L[\mathbf{T} \quad \mathbf{0}]\mathbf{E}^T. \quad (62)$$

We set $\alpha = 1/12$ and plot the central column of $\mathbf{P}_t$ in Fig. 5a, which represents the one-point covariance function for the point $x = \pi$. The length of the state vector for the high-resolution (true) states will be $N = 256$. In Fig. 5a, we also plot the central column of $\mathbf{P}_f$ for two different truncation matrices $\mathbf{S}$: one is for $M = 16$ and the other is for $M = 8$. Fig. 5b shows the corresponding central column of $\mathbf{S}$, which maps the high-resolution (true) states to the low-resolution (forecast) states. The smoothing functions in Fig. 5b show that truncating the true spectrum results in smoothing kernels that average the true state globally to produce the local low-resolution (forecast) state. In addition, note that the smoothing kernel weights the true state both positively and negatively, which we will see shortly results in long distance negative correlations. Figure 5c and 5d shows the one-point covariance function (again for the central column) for $\bar{\mathbf{R}}_f - \mathbf{R}_i$ and $\bar{\mathbf{R}}_f^* - \mathbf{R}_i$ for the $M = 16$ and $M = 8$ cases. In both cases, the representation error $\bar{\mathbf{R}}_f - \mathbf{R}_i$ and the effective representation error $\bar{\mathbf{R}}_f^* - \mathbf{R}_i$ are very nearly equal in the centre of the domain but differ in the far-field. We will show by contrast below that this equality between these two matrices is due to the pure truncation in this case. When we invoke a Gaussian smoother (i.e. $\mathbf{D} \neq \mathbf{I}$), which as noted above corresponds to a climatological model error on the resolved scales, $\bar{\mathbf{R}}_f - \mathbf{R}_i$ and $\bar{\mathbf{R}}_f^* - \mathbf{R}_i$ will become substantially different. Also, note that there exists a strong far-field positive-negative wave pattern in $\bar{\mathbf{R}}_f - \mathbf{R}_i$ and less so for $\bar{\mathbf{R}}_f^* - \mathbf{R}_i$. This is due to the aforementioned positive-negative weightings in the smoothing kernels of Fig. 5b. An example of a randomly constructed high-resolution (true) state [eq. (46)] is shown in Fig. 6. Using eq. (62), we may produce the corresponding low-resolution (forecast) state from the $M = 8$ forecast state shown in Fig. 6. Also, shown in Fig. 6 is the estimate of the true state, $\bar{\mathbf{x}}_c$, given the low-resolution forecast state from eq. (51) using the $M = 8$ forecast state shown in that figure. This ‘best estimate’ from eq. (51) makes use of the linear regression in eq. (51) but does not make use of the observation operator, $\mathbf{H}$. In this sense, this ‘best estimate’ is simply a state-dependent bias correction of the forecast.

The next set of experiments will make use of the matrix $\mathbf{D}$. Here, we study eq. (49) for $\beta = 1/6$, which implies a Gaussian smoother on the degrees of freedom that are retained *after* truncation. The same true covariance model

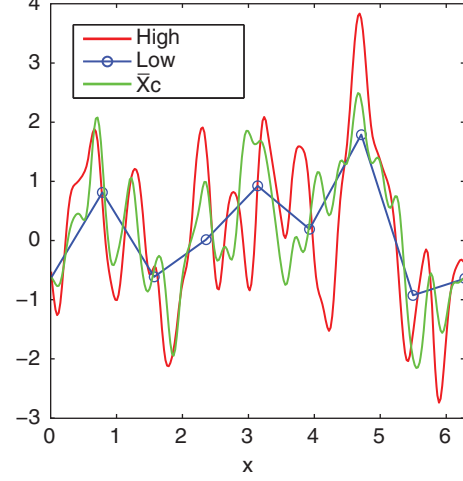


Fig. 6. Example high-resolution state, low-resolution state ($M = 8$) and mean of the conversion density, which is labelled above as the ‘best estimate’. The low-resolution state is defined only at the grid-points denoted by the open circles.

is used in this experiment, $\mathbf{P}_t$, for which the one-point covariance function is repeated in Fig. 7a. Again, in Fig. 7a, we also plot the central column of $\mathbf{P}_f$ for two different truncation matrices $\mathbf{T}$: one is for $M = 256$ and the other is for $M = 16$. Note that the $M = 256$ case corresponds to no truncation at all as the forecast state vector and the true state vector are equal ($M = N = 256$). The $M = 256$ case provides an example of a forecast model error that does not arise from a reduction (truncation) of the number of degrees of freedom in the forecast model. The smoothing kernels for these two cases are shown in Fig. 7b. The truncated smoothing kernel ($M = 16$) shows the positive-negative oscillations in the far-field that we saw previously. Contrast this with the non-truncated smoothing kernel ($M = 256$) that is completely local.

The application of these smoothing kernels produces distinctly different responses in $\bar{\mathbf{R}}_f - \mathbf{R}_i$ and $\bar{\mathbf{R}}_f^* - \mathbf{R}_i$. For the $M = 256$ case, the $\mathbf{S}$ is full-rank, has an inverse and results in $\bar{\mathbf{R}}_f - \mathbf{R}_i = \mathbf{0}$ as shown in Fig. 7c. However, because the smoothing matrix $\mathbf{S}$ still produces a difference between $\mathbf{P}_t$ and $\mathbf{P}_f$, which implies that $\bar{\mathbf{R}}_f^* - \mathbf{R}_i$ is non-zero as shown in Fig. 7c. Hence, when the forecast model is full-rank, but fails to reproduce the climatology of the true attractor, the representation error nonetheless vanishes. This type of climatological error in the model must still be accounted for using $\bar{\mathbf{R}}_f^* - \mathbf{R}_i$ and $\mathbf{G}_p$. In the case where $M = 16$, which implies both a Gaussian smoothing and a truncation, one can see in Fig. 7d that both representation error and $\bar{\mathbf{R}}_f^* - \mathbf{R}_i$ are non-zero and quite different. This difference between $\bar{\mathbf{R}}_f - \mathbf{R}_i$ and $\bar{\mathbf{R}}_f^* - \mathbf{R}_i$ illustrates that one can only estimate the representation error from

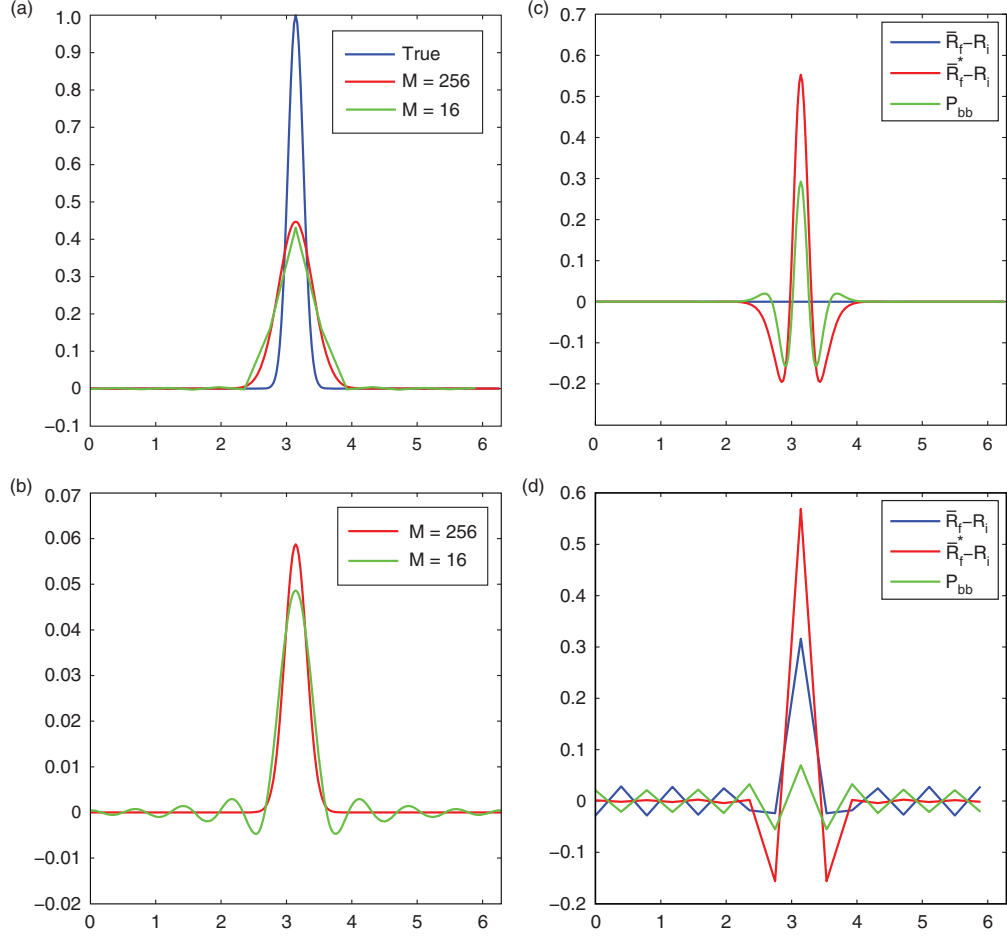


Fig. 7. Gaussian Smoothing. (a) Three one-point prior covariance functions: blue is the high-resolution (true) covariance model, red is the $M=256$ low-resolution (forecast) covariance model and green is the $M=16$ low-resolution (forecast) covariance model. (b) One column of the smoother matrix for $M=256$ (red) and $M=16$ (green). (c, d) The main components of the theory for $M=256$ and $M=16$, respectively: blue is the representation error covariance ($\bar{\mathbf{R}}_f - \mathbf{R}_i$), red is the *effective* representation error covariance ($\bar{\mathbf{R}}_f^* - \mathbf{R}_i$) and green is the bias covariance matrix ($\mathbf{P}_{bb}$).

innovation statistics when eq. (58) is implemented because this is the only way to render $\mathbf{P}_{fb} = \mathbf{0}$ and $\mathbf{P}_{bb} = \mathbf{0}$.

5. Summary and conclusions

A framework has been presented to understand the origin of the representation error as well as to properly frame attempts at estimating and accounting for its effects. We have extended the work of Kalman (1960), whose data assimilation algorithm is optimal for Gaussian systems for which the flaws in the model were accounted for using a white noise forcing, to the case where the observations are of states on a true ‘attractor’ and the model evolution produces states on a forecast ‘attractor’ with both states Gaussian distributed and a linear map existing between them, but with no applied white noise forcing. In practical terms, when the distributions are not Gaussian and the

mapping between the two attractors is not linear, we have derived the best linear, unbiased estimation technique. For this case, we have shown that in this data assimilation algorithm the observation operator does *not* simply map the truth to the observation, but rather it maps the forecast to the observation, and because the forecast is in a different portion of state space than the truth, this requires a modified observation operator. We emphasise that the operation of this new data assimilation framework only requires the prior distribution on the forecast model attractor and the function $\mathbf{F}$, and does not need the prior distribution on the true attractor, to correctly assimilate observations of states on the true attractor. We view this map in eq. (4) much like that of the algorithms in the ensemble post-processing and bias-correction literature (e.g. Glahn and Lowry, 1972; Raftery et al., 2005) in which climatological information is used to build relationships

between forecasts and observations. As such, this new framework has shown how to properly include the information normally obtained from ‘bias-correction’ algorithms within the data assimilation system.

As discussed in the introduction, the notion that the true state might not be the best initial condition for the flawed forecast model has led us to choose to develop a data assimilation system that attempts to produce a state estimate on the forecast attractor. This obviously produces a state estimate that may not be as close to the truth as is required in some applications. Note however that we may use the state estimation procedure described in Appendix A to map our forecast back to the true attractor in order to produce a state estimate on the true attractor. This of course is nothing more than the well-known post-processing of a forecast but using the infrastructure that we have already developed for the data assimilation algorithm.

In any event, this new framework shows that the representation error arises from the lack of an inverse in the relationship between the true attracting manifold and that of the forecast models. This lack of an inverse in their relationship results when the forecast model is a truncated representation of the true states. In other words, representation error does not occur when the forecast model is simply in error in its representation of climatology. The error that the model must make for representation error to exist is one in which the forecast model has been truncated to represent fewer degrees of freedom than the number of degrees of freedom describing the true states. In this specific case, the forecast observation likelihood $\left[\rho(\mathbf{y}_f|\mathbf{x}_f)\right]$ will have a variance greater than that of the error of the instrument and is also likely to have a correlation between observations. In the Gaussian case, it was shown explicitly that this inflation and correlation structure arises as the covariance matrix of the portion of the true spectrum that goes missing from the truncated forecast model. Lastly, it was shown that innovation-based techniques (e.g. Hollingsworth and Lönnerberg, 1986; Desroziers et al., 2006) for the estimation of representation error covariance matrices are self-consistent estimators of the representation error *only when the observation operator has been modified to account for the attractor differences*.

Applying this new framework to specific problems in the geosciences will require estimation of the map between the true attractor and that of the forecast model [eq. (4)]. We imagine a proxy for such a model could be developed from the difference between high-resolution and low-resolution simulations. After development of the map [eq. (4)], one must develop the observation gain ($\mathbf{G}_p$) for each observation in which one is interested in accounting for errors in representation. We suggest performing this step using an observation-by-observation approach as this will likely lead to a reduction in the size of the matrices in the

calculation [eq. (52)]. Research into performing this estimation of $\mathbf{S}$ is underway and will be reported in a sequel.

6. Acknowledgements

We gratefully acknowledge support from the Chief of Naval Research PE-0601153N and from the NERC National Centre for Earth Observation in the UK and the European Space Agency. We thank Jeff Anderson for a thorough reading and insightful comments on this work.

7. Appendix

The Moore-Penrose Inverse and the Minimum Variance Estimate

We begin with eq. (36) and decompose it into prior mean and perturbation:

$$\bar{\mathbf{x}}_f = \mathbf{S}\bar{\mathbf{x}}_t, \mathbf{\epsilon}_f = \mathbf{S}\mathbf{\epsilon}_t \quad (\text{A1a,b})$$

where $\mathbf{\epsilon}_f = \mathbf{x}_f - \bar{\mathbf{x}}_f$ and $\mathbf{\epsilon}_t = \mathbf{x}_t - \bar{\mathbf{x}}_t$. Equation (46) tells us that

$$\mathbf{\epsilon}_t = \mathbf{Z}\boldsymbol{\eta} \quad (\text{A2})$$

where $\boldsymbol{\eta}$ is a N -vector whose elements are drawn from $N(0,1)$. Therefore, we attempt to minimise the cost function

$$J(\boldsymbol{\eta}) = \frac{1}{2} [\mathbf{\epsilon}_f - \mathbf{SZ}\boldsymbol{\eta}]^T [\mathbf{\epsilon}_f - \mathbf{SZ}\boldsymbol{\eta}] \quad (\text{A3})$$

to find that the minimum of eq. (A3) is defined by an infinite number of solutions of the form:

$$\hat{\boldsymbol{\eta}} = [\mathbf{SZ}]^\dagger \mathbf{\epsilon}_f + \mathbf{E}_2 \boldsymbol{\xi} \quad (\text{A4})$$

where

$$[\mathbf{SZ}]^\dagger = \mathbf{E}\boldsymbol{\Gamma}^{-1/2} [\mathbf{D}^{-1/2}\mathbf{T}^{-1} \quad \mathbf{0}] \mathbf{E}_L^T \quad (\text{A5})$$

$$\mathbf{E} = [\mathbf{E}_1 \quad \mathbf{E}_2] \quad (\text{A6})$$

and $\boldsymbol{\xi}$ is a vector of length $(N-M)$ that is composed of random noise with the property that it is a random draw from $N(\mathbf{0}, \boldsymbol{\Gamma}_{N-M})$ with $\boldsymbol{\Gamma}_{N-M}$ denoting the last $(N-M)$ elements of the diagonal of $\boldsymbol{\Gamma}$. In eq. (A6), $\mathbf{E}$ has been subdivided such that we define $\mathbf{E}_1$ as the first M columns of $\mathbf{E}$ and then place the remaining columns into $\mathbf{E}_2$. We choose the best solution from the set [eq. (A4)] by requiring the solution at the minimum of eq. (A3) to also minimise $\mathbf{\epsilon}_t^T \mathbf{P}_t^{-1} \mathbf{\epsilon}_t$, which may be shown to imply that $\boldsymbol{\eta}^T \boldsymbol{\eta}$ is also a minimum. The addition of this constraint chooses $\boldsymbol{\xi} = \mathbf{0}$, which then defines the solution as

$$\hat{\boldsymbol{\eta}} = \mathbf{Z}[\mathbf{SZ}]^\dagger \mathbf{\epsilon}_f \quad (\text{A7})$$

Equation (A7) defines the minimum variance estimate $\mathbf{\epsilon}_t$ under the constraint that $\mathbf{\epsilon}_t^T \mathbf{P}_t^{-1} \mathbf{\epsilon}_t$ is also a minimum. Hence,

eq. (A7) finds the minimum variance estimate of eq. (A1a,b) for the state ε_i given ε_f . The minimum variance estimate of a linear estimator will be equal to the mean of the conversion density when that density is Gaussian. When that density is not Gaussian, it then reduces to the best *linear*, unbiased estimate of the mean of the conversion density.

References

- Daley, R. 1993. Estimating observation error statistics for atmospheric data assimilation. *Ann. Geophys.* **11**, 634–647.
- Desroziers, G., Berre, L., Chapnik, B. and Poli, P. 2006. Diagnosis of observation, background and analysis-error statistics in observation space. *Q. J. Roy. Meteorol. Soc.* **132**, 3385–3396.
- Frehlich, R. 2006. Adaptive data assimilation including the effect of spatial variations in observation error. *Q. J. Roy. Meteorol. Soc.* **132**, 1225–1257.
- Frehlich, R. 2008. Atmospheric turbulence component of the innovation covariance. *Q. J. Roy. Meteorol. Soc.* **134**, 931–940.
- Frehlich, R. 2011. The definition of ‘truth’ for numerical weather prediction error statistics. *Q. J. Roy. Meteorol. Soc.* **137**, 84–98.
- Glahn, H. R. and Lowry, D. A. 1972. The use of model output statistics (MOS) in objective weather forecasting. *J. Appl. Meteorol.* **11**, 1203–1211.
- Golub, G. and Van Loan, C. F. 1996. *Matrix Computations*. The Johns Hopkins University Press, Baltimore, MD, USA, 687 pp.
- Hodyss, D. 2011. Ensemble state estimation for nonlinear systems using polynomial expansions in the innovation. *Mon. Weather Rev.* **139**, 3571–3588.
- Hodyss, D. and Campbell, W. F. 2013. Square root and perturbed observation ensemble generation techniques in Kalman and quadratic ensemble filtering algorithms. *Mon. Weather Rev.* **141**, 2561–2573.
- Hodyss, D., McLay, J. G., Moskaitis, J. and Serra, E. A. 2014. Inducing tropical cyclones to undergo Brownian motion: a comparison between Itô and Stratonovich in a numerical weather prediction model. *Mon. Weather Rev.* **142**, 1982–1996.
- Hollingsworth, A. and Lönnberg, P. 1986. The statistical structure of short-range forecast errors as determined from radiosonde data. Part I: The wind field. *Tellus A.* **38**, 111–136.
- Janjic, T. and Cohn, S. E. 2006. Treatment of observation error due to unresolved scales in atmospheric data assimilation. *Mon. Weather Rev.* **134**, 2900–2915.
- Kalman, R. E. 1960. A new approach to linear filtering and prediction problems. *J. Basic. Eng.* **82**, 35–45.
- Lilly, D. K. 1962. On the simulation of buoyant convection. *Tellus* **14**, 148–172.
- Liu, Z.-Q. and Rabier, F. 2002. The interaction between model resolution, observation resolution and observation density in data assimilation: a one-dimensional study. *Q. J. Roy. Meteorol. Soc.* **128**, 1367–1386.
- Mitchell, H. L. and Daley, R. 1997a. Discretization error and signal/error correlation in atmospheric data assimilation (I). All scales resolved. *Tellus A.* **49**, 32–53.
- Mitchell, H. L. and Daley, R. 1997b. Discretization error and signal/error correlation in atmospheric data assimilation (II). The effect of unresolved scales. *Tellus A.* **49**, 54–73.
- Oke, P. R. and Sakov, P. 2008. Representation error of oceanic observations for data assimilation. *J. Atmos. Oceanic Technol.* **25**, 1004–1017.
- Petersen, D. P. and Middleton, D. 1963. On representative observations. *Tellus A.* **15**, 387–405.
- Raftery, A. E., Gneiting, T., Balabdaoui, F. and Polakowski, M. 2005. Using Bayesian model averaging to calibrate forecast ensembles. *Mon. Wea. Rev.* **133**, 1155–1174.
- Richman, J. G., Miller, R. N. and Spitz, Y. H. 2005. Error estimates for assimilation of satellite sea surface temperature data I ocean climate models. *J. Geophys. Res.* **32**, L18608. DOI: 10.1029/2005 GL023591.
- Waller, J. A., Dance, S. L., Lawless, A. S., Nichols, N. K. and Eyre, J. R. 2013. Representivity error for temperature and humidity using the Met Office high-resolution model. *Q. J. Roy. Meteorol. Soc.* **140**, 1189–1197.