



Rapid update cycling with delayed observations

By T. J. PAYNE^{1*}, ¹*Met Office, Exeter, UK*

(Manuscript received 10 August 2017; in final form 14 November 2017)

ABSTRACT

In this paper we examine the fundamental issues associated with the cycling of data assimilation and prediction in the case where observations are received after a delay, but we seek to assimilate them immediately on receipt, or within a short time of receipt. We obtain the optimal solution to this problem in the linear and non-linear cases, and explore its relation to simplified strategies which are adaptations of contemporary methods for large-scale data assimilation. We also discuss the challenges facing such cycling in large-scale numerical weather prediction.

Keywords: Data assimilation, rapid update cycling

1. Introduction

In the traditional cycling of forecasts and data assimilation (DA) for numerical weather prediction, the DA step for the global model has occurred every 6 or 12 hours. This was appropriate for an era when data was concentrated at the main synoptic times, and the limited area models (LAMs) for which the global model provides boundary conditions were cycled every 6 hours. However, in recent years data has become dominated by sources which are essentially continuous in time, and centres such as the Met Office will soon cycle their highest resolution LAM every hour.

By increasing the frequency of global analyses (e.g. to every hour) global forecasts can be based on more recent data, which is not only desirable in itself but provides timely lateral boundary conditions (LBCs) for high resolution LAMs. In one study of the Met Office's 1.5 km LAM covering the British Isles (Tang et al., 2013) it was found that replacing 3-hour and 6-hour old LBCs by 3-hour and fresh LBCs improved the UK index (a basket of scores measuring forecast skill) by 1.5% (Bruce Macpherson, *pers comm*). Furthermore, by having more frequent analyses the analysis increments will be smaller, which will improve the validity of the linear approximations in DA schemes. More frequent analyses may also improve the affordability of DA methods as the computational load is distributed more evenly in time.

We will see that, because of the delay in receiving some data, to ensure that all the data which are received are also assimilated, the assimilation windows will need to overlap. We obtain the optimal solution to this problem, which involves manipulating simultaneously all the states in the window and their joint errors.

We explore the relation between the optimal solution and simplified methods which are closer to current methods for large-scale DA.

We show that the current use of largely climatological prior error covariances may pose a challenge for high frequency cycling, and discuss how this may be overcome.

2. 'Traditional' vs. 'Rapid Update' cycling

An immediate issue is that observations are not received instantaneously. For example, by 09Z on 18 June 2015 the Met Office had received over 80 million observations valid between 09Z on 15 June and 03Z on 18 June 2015, including around 0.8 million surface, 2.1 million aircraft and sonde, 12.4 million satwind and 14 million ATOVS observations. The delay between validity time and receipt for these observation types is recorded in Fig. 1. We see that to receive 95% of aircraft and sonde, surface, ATOVS and satwind observation took respectively 0.6, 1.5, 3.4 and 4.1 hours.

This presents a quandary for traditional cycling which aims to produce an analysis every 6 (at some centres every 12) hours. For definiteness consider 4D-Var (e.g. Li and Navon, 2001) with a 6-hour window $[T-3, T+3]$, which generates an analysis at $T-3$ (in this discussion the units are hours).

One could perform the analysis at $T+3$ using all observations available by $T+3$, which would minimise the time delay to produce the analysis, but observations received after $T+3$ would not be assimilated. Alternatively, one could perform the analysis at $T+7$, by which time (in view of Fig. 1) almost all the observations valid in the window have been received, but the analysis is only

*Corresponding author. e-mail: tim.payne@metoffice.gov.uk

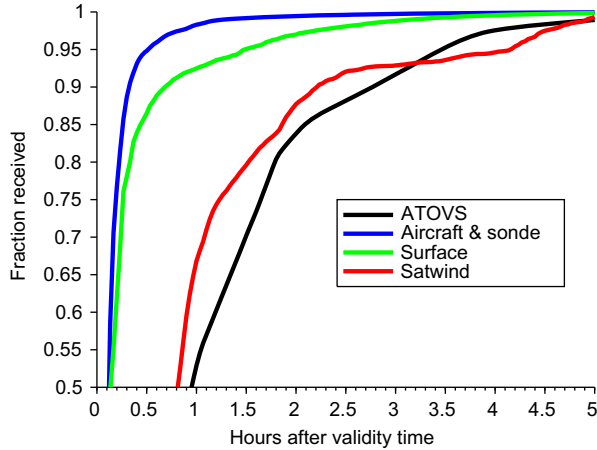


Fig. 1. Delay in receiving various observation types in the Met Office observation processing system, for observations valid between 9Z on 15 June 2015 and 3Z on 18 June 2015.

available 4 hours after the end of the window and 10 hours after its beginning. To generate an estimate of state at $T+7$ we could run a 10-hour forecast from the analysis, but compared with the estimate of state at $T-3$ this will be degraded by model error.

Centres such as the Met Office mitigate these issues by performing each analysis twice, a ‘late cut-off’ analysis currently at about $T + 6$, and an ‘early cut-off’ analysis at about $T+3$. This is illustrated somewhat schematically in Fig. 2, which shows two adjacent non-overlapping windows [3Z,9Z) and [9Z,15Z) (where $[T_1, T_2)$ denotes the time interval $T_1 \leq t < T_2$). For example, at 8Z we receive observations valid between 4Z and 8Z. Considering the window [3Z,9Z), for the ‘early cut-off’ run we perform the data assimilation at 9Z and use the observations in the dark blue region, and for the ‘late cut-off’ run perform it at about 12Z and use virtually all observations ever valid in the window (combined light and dark blue region).

In principle the late cut-off analysis could make use of the early cut-off one to reduce its work load, as is done in the ‘quasi-continuous’ approach of Järvinen et al. (1996) and Veerse and Thépaut (1998), but at the Met Office the late cut-off analyses start again from scratch, making no use of the work done for the early cut-off analysis.

Having both early and late cut-off analyses goes some way to mitigating the shortcomings of 6-hourly cycling. However, the analyses are still 6 hours apart, which makes them insufficiently timely for some purposes, notably (considering the comments in Section 1) the LBCs for hourly LAM analyses; the analysis increment is much larger than would be the case with an hourly update, so nonlinearity can be a significant problem, especially for the linear model in 4D-Var; and the approach is inefficient insofar as the early cut-off analyses are not used as part of a cycle.

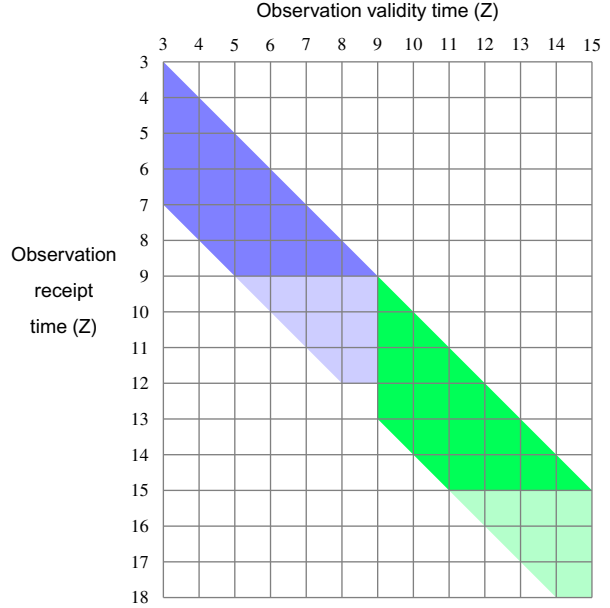


Fig. 2. Cycling with non-overlapping windows as used at the Met Office. Observations in the dark blue (dark green) region are assimilated for an ‘early cut-off’ run at 9Z (15Z). The extra observations in the lightly shaded regions are assimilated for ‘late cut-off’ runs at 12Z and 18Z.

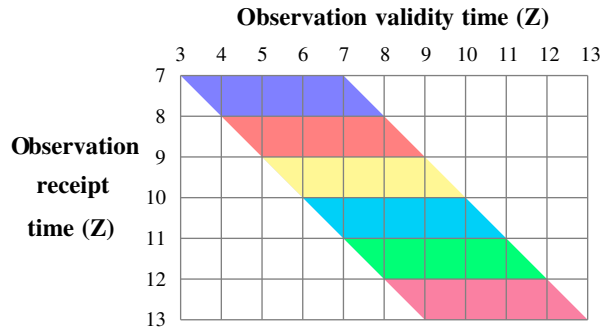


Fig. 3. Rapid update cycling with overlapping windows. Observations are assimilated within one hour of receipt.

In Fig. 3 we illustrate how we would like to deal with the same case: each hour we assimilate all observations received in the last hour, e.g. at 12Z we assimilate the observations received between 11Z and 12Z (green region); these are valid between 7Z and 12Z. In principle we do not re-assimilate the observations valid between 7Z and 12Z received at earlier times (blue, red yellow and cyan regions in Fig. 3) as the information from these observations has been transferred to previous analyses and thereby to the background for this cycle.

In the context of global NWP, which provides among other outputs LBCs for hourly cycling of LAMs, an hourly update cycle is natural, but in principle the updates could be as short as one model time step. For the purposes of this paper we will refer to any cycling where observations are assimilated as soon as they are received, or as in Fig. 3 within some short time of receipt, as *rapid update cycling* (RUC).

We should note that the term ‘Rapid Update Cycle’ has been employed in the past to denote specific rapidly cycled NWP systems, for example by the National Centres for Environmental Prediction in the USA. In that case it referred to an operational regional forecast-analysis system over North America, where data was assimilated by 3D-Var (originally optimal interpolation) using non-overlapping windows of length one hour (Benjamin et al., 2004b; Benjamin et al., 2004a).

3. Optimal RUC

To examine the rapid update cycling problem further we will idealise it slightly by supposing that observations are valid at exact multiples of a time increment δt (as opposed to continuously in time), and become available after delays of $0, \delta t, 2\delta t, \dots, N\delta t$. We will suppose that observations received at time $k\delta t$ are

$$\mathbf{y}_k^{(k)}, \mathbf{y}_{k-1}^{(k)}, \mathbf{y}_{k-2}^{(k)}, \dots, \mathbf{y}_{k-N}^{(k)} \quad (1)$$

Superscripts denote when the observations are received and subscripts their validity time, the longest delay being $N\delta t$.

The first problem is to develop an optimal method for assimilating observations as soon as they are available. At time $k\delta t$ we seek to estimate $\mathbf{x}_k, \dots, \mathbf{x}_{k-N}$, given observations (1) and our previous estimate of $\mathbf{x}_{k-1}, \dots, \mathbf{x}_{k-N-1}$.

We will suppose that for $i = k, k-1, \dots, k-N$ we have observation operators $\mathbf{h}_i^{(k)}$ such that

$$\mathbf{y}_i^{(k)} = \mathbf{h}_i^{(k)}(\mathbf{x}_i) + \mathbf{v}_i^{(k)} \quad (2)$$

and for each i a model $\mathbf{f}_i$

$$\mathbf{x}_{i+1} = \mathbf{f}_i(\mathbf{x}_i) + \boldsymbol{\omega}_i \quad (3)$$

where the distributions of the errors $\mathbf{v}_i^{(k)}, \boldsymbol{\omega}_i$ are supposed known.

3.1. Notation and problem formulation

The optimal method is obtained by formulating the problem in such a way that standard estimation theory can be applied.

We will use the convention that the underlined vector $\underline{\mathbf{x}}_k$ denotes the concatenation of the $N+1$ vectors $\{\mathbf{x}_i, k-N \leq i \leq k\}$:

$$\underline{\mathbf{x}}_k = \begin{pmatrix} \mathbf{x}_{k-N} \\ \vdots \\ \mathbf{x}_k \end{pmatrix} \quad (4)$$

and similarly the underlined matrix $\underline{A}_k$ is formed from matrices $\{A_{i,j}, k-N \leq i, j \leq k\}$ of compatible size:

$$\underline{A}_k = \begin{pmatrix} A_{k-N,k-N} & \dots & A_{k-N,k} \\ \vdots & \ddots & \vdots \\ A_{k,k-N} & \dots & A_{k,k} \end{pmatrix}$$

For example, if $\{\mathbf{x}_i, k-N \leq i \leq k\}$ are vectors of length n and $\{A_{i,j}, k-N \leq i, j \leq k\}$ are matrices of size $n \times n$, then $\underline{\mathbf{x}}_k, \underline{A}_k$ are of size respectively $n(N+1) \times 1$ and $n(N+1) \times n(N+1)$.

Define $\underline{\mathbf{y}}_k$ to be the observations received at time $k\delta t$, so

$$\underline{\mathbf{y}}_k = \begin{pmatrix} \mathbf{y}_{k-N}^{(k)} \\ \vdots \\ \mathbf{y}_k^{(k)} \end{pmatrix} \quad (5)$$

We seek the conditional expectations of $\underline{\mathbf{x}}_k$ and $\underline{\mathbf{x}}_{k+1}$ given observations received up to time $k\delta t$

$$E[\underline{\mathbf{x}}_k | \underline{\mathbf{y}}_0, \underline{\mathbf{y}}_1, \dots, \underline{\mathbf{y}}_k], \quad E[\underline{\mathbf{x}}_{k+1} | \underline{\mathbf{y}}_0, \underline{\mathbf{y}}_1, \dots, \underline{\mathbf{y}}_k] \quad (6)$$

Note that, given $\underline{\mathbf{x}}_k$, and assuming $\boldsymbol{\omega}_k$ has zero mean, the best estimate of $\underline{\mathbf{x}}_{k+1}$ before observations received at $k\delta t$ are assimilated is

$$\underline{\mathbf{f}}_k(\underline{\mathbf{x}}_k) \equiv \begin{pmatrix} \mathbf{x}_{k-N+1} \\ \vdots \\ \mathbf{x}_k \\ \mathbf{f}_k(\mathbf{x}_k) \end{pmatrix} \quad (7)$$

If we now define

$$\underline{\boldsymbol{\omega}}_k = \begin{pmatrix} \mathbf{0} \\ \vdots \\ \mathbf{0} \\ \boldsymbol{\omega}_k \end{pmatrix}, \quad \underline{\mathbf{v}}_k = \begin{pmatrix} \mathbf{v}_{k-N}^{(k)} \\ \vdots \\ \mathbf{v}_{k-1}^{(k)} \\ \mathbf{v}_k^{(k)} \end{pmatrix}, \quad \underline{\mathbf{h}}_k(\underline{\mathbf{x}}_k) = \begin{pmatrix} \mathbf{h}_{k-N}^{(k)}(\mathbf{x}_{k-N}) \\ \vdots \\ \mathbf{h}_{k-1}^{(k)}(\mathbf{x}_{k-1}) \\ \mathbf{h}_k^{(k)}(\mathbf{x}_k) \end{pmatrix} \quad (8)$$

Then we may write (2) and (3) as respectively

$$\underline{\mathbf{y}}_k = \underline{\mathbf{h}}_k(\underline{\mathbf{x}}_k) + \underline{\mathbf{v}}_k \quad (9)$$

$$\underline{\mathbf{x}}_{k+1} = \underline{\mathbf{f}}_k(\underline{\mathbf{x}}_k) + \underline{\boldsymbol{\omega}}_k \quad (10)$$

which are in the standard form for observation and signal map equations in estimation theory (e.g. Jazwinski, 1970).

3.2. Linear Gaussian case

It is illuminating to first work out the details in the simplest case, where in (2) and (3) the observation operators $\mathbf{h}_i^{(k)}$ and model $\mathbf{f}_i$ are linear and the errors $\mathbf{v}_i^{(k)}$, $\boldsymbol{\omega}_i$ are zero-mean, Gaussian and uncorrelated.

In this case

$$\begin{aligned}\mathbf{y}_i^{(k)} &= H_i^{(k)} \mathbf{x}_i + \mathbf{v}_i^{(k)} \\ \mathbf{x}_{i+1} &= M_i \mathbf{x}_i + \boldsymbol{\omega}_i\end{aligned}\quad (11)$$

where

$$\mathbf{v}_i^{(k)} \sim N(\mathbf{0}, R_i^{(k)}), \quad \boldsymbol{\omega}_i \sim N(\mathbf{0}, Q_i)$$

for some matrices $R_i^{(k)}$, Q_i (where $\sim N(\boldsymbol{\mu}, \Sigma)$ denotes normally distributed with mean $\boldsymbol{\mu}$ and variance Σ), and setting

$$\begin{aligned}\underline{H}_k &= \begin{pmatrix} H_{k-N}^{(k)} & & & \\ & H_{k-N+1}^{(k)} & & \\ & & \ddots & \\ & & & H_k^{(k)} \end{pmatrix}, \\ \underline{M}_k &= \begin{pmatrix} 0 & I & & \\ & \ddots & \ddots & \\ & & 0 & I \\ & & & M_k \end{pmatrix}\end{aligned}\quad (12)$$

and

$$\underline{R}_k = \begin{pmatrix} R_{k-N}^{(k)} & & & \\ & R_{k-N+1}^{(k)} & & \\ & & \ddots & \\ & & & R_k^{(k)} \end{pmatrix}, \quad \underline{Q}_k = \begin{pmatrix} 0 & & & \\ & \ddots & & \\ & & 0 & \\ & & & Q_k \end{pmatrix}\quad (13)$$

(9,10) become

$$\begin{aligned}\underline{\mathbf{y}}_k &= \underline{H}_k \underline{\mathbf{x}}_k + \underline{\mathbf{v}}_k \\ \underline{\mathbf{x}}_{k+1} &= \underline{M}_k \underline{\mathbf{x}}_k + \underline{\boldsymbol{\omega}}_k\end{aligned}\quad (14)$$

with

$$\underline{\mathbf{v}}_k \sim N(\mathbf{0}, \underline{R}_k), \quad \underline{\boldsymbol{\omega}}_k \sim N(\mathbf{0}, \underline{Q}_k)$$

and the problem of finding the conditional expectations (6) of $\underline{\mathbf{x}}_k$ and $\underline{\mathbf{x}}_{k+1}$ given observations received up to time k , which may be denoted $\hat{\underline{\mathbf{x}}}_k|k$, $\hat{\underline{\mathbf{x}}}_{k+1}|k$, is solved by a standard Kalman Filter, as in Table 1.

The basic objects manipulated are whole windows of states and observations and the covariances of the errors in these objects. The symbols in Table 1 are whole-window analogues of

Table 1. Optimal filtering with overlapping windows in linear-Gaussian case.

Given first guesses $\hat{\underline{\mathbf{x}}}_{N|N-1}$ with error covariance $\underline{P}_{N|N-1}$

For $k = N, N+1, N+2, \dots$

Assimilate:

$$\underline{K} = \underline{P}_{k|k-1} \underline{H}_k^T (\underline{R}_k + \underline{H}_k \underline{P}_{k|k-1} \underline{H}_k^T)^{-1} \quad (15)$$

$$\hat{\underline{\mathbf{x}}}_k|k = \hat{\underline{\mathbf{x}}}_k|k-1 + \underline{K} (\underline{\mathbf{y}}_k - \underline{H}_k \hat{\underline{\mathbf{x}}}_k|k-1) \quad (16)$$

$$\underline{P}_{k|k} = (\underline{I} - \underline{K} \underline{H}_k) \underline{P}_{k|k-1} \quad (17)$$

Predict:

$$\hat{\underline{\mathbf{x}}}_{k+1}|k = \underline{M}_k \hat{\underline{\mathbf{x}}}_k|k \quad (18)$$

$$\underline{P}_{k+1|k} = \underline{M}_k \underline{P}_{k|k} \underline{M}_k^T + \underline{Q}_k \quad (19)$$

End for $k = N, N+1, N+2, \dots$

their usual values, e.g. $\underline{P}_{k|k-1}$ and $\underline{P}_{k|k}$ are $n(N+1) \times n(N+1)$ prior and posterior error covariance matrices of the estimated $\underline{\mathbf{x}}_k$. In special circumstances simplification is possible. For example, if new observations only occur in the final time slot, i.e. if the only observations which become available at k are $\underline{\mathbf{y}}_k^{(k)}$ and there are no observations $\underline{\mathbf{y}}_{k-1}^{(k)}, \dots, \underline{\mathbf{y}}_{k-N}^{(k)}$, it is straightforward to show that (15)–(19) simplifies to a conventional Kalman smoother.

For linear $M_i, H_i^{(k)}$ the algorithm (15)–(19) finds $E[\underline{\mathbf{x}}_k | \underline{\mathbf{y}}_0, \underline{\mathbf{y}}_1, \dots, \underline{\mathbf{y}}_k]$, $E[\underline{\mathbf{x}}_{k+1} | \underline{\mathbf{y}}_0, \underline{\mathbf{y}}_1, \dots, \underline{\mathbf{y}}_k]$ if the errors $\mathbf{v}_i^{(k)}, \boldsymbol{\omega}_i$ are zero-mean, Gaussian and uncorrelated. As noted in Anderson and Moore (1979), if we restrict attention to analysis-prediction equations of form (16) and (18) then we may drop the Gaussian assumption on $\mathbf{v}_i^{(k)}, \boldsymbol{\omega}_i$, merely requiring them to be zero-mean and uncorrelated, and (15)–(19) still minimises the expected error variance.

3.3. Variational equivalent of analysis step in linear Gaussian case

For large-scale data assimilation variational methods are almost universally used, so it is of interest to cast the optimal RUC analysis step (16) in variational form. Define

$$\underline{\boldsymbol{\delta}} \equiv \begin{pmatrix} \boldsymbol{\delta}_{k-N} \\ \boldsymbol{\delta}_{k-N+1} \\ \vdots \\ \boldsymbol{\delta}_k \end{pmatrix}$$

$$J_b(\underline{\boldsymbol{\delta}}) = \frac{1}{2} \begin{pmatrix} \boldsymbol{\delta}_{k-N} \\ \vdots \\ \boldsymbol{\delta}_{k-1} \end{pmatrix}^T \hat{A}^{-1} \begin{pmatrix} \boldsymbol{\delta}_{k-N} \\ \vdots \\ \boldsymbol{\delta}_{k-1} \end{pmatrix} \quad (20)$$

$$J_o(\underline{\delta}) = \frac{1}{2}(\underline{\mathbf{y}}_k - \underline{H}_k(\underline{\hat{\mathbf{x}}}_{k|k-1} + \underline{\delta}))^T \underline{R}_k^{-1}(\underline{\mathbf{y}}_k - \underline{H}_k(\underline{\hat{\mathbf{x}}}_{k|k-1} + \underline{\delta})) \quad (21)$$

$$J_q(\underline{\delta}) = \frac{1}{2}(\underline{\delta}_k - \underline{M}_{k-1}\underline{\delta}_{k-1})^T \underline{Q}_{k-1}^{-1}(\underline{\delta}_k - \underline{M}_{k-1}\underline{\delta}_{k-1}) \quad (22)$$

where $\hat{A}$ is the bottom right $Nn \times Nn$ submatrix of $\underline{P}_{k-1|k-1}$. Then the analysis step (16) is equivalent to

$$\hat{\underline{\mathbf{x}}}_{k|k} = \hat{\underline{\mathbf{x}}}_{k|k-1} + \underline{\delta} \quad (23)$$

where $\underline{\delta}$ minimises

$$J(\underline{\delta}) = J_b(\underline{\delta}) + J_o(\underline{\delta}) + J_q(\underline{\delta}) \quad (24)$$

This is proved in Appendix A. The J_b term (20) constrains the N states in the *intersection* between the old and new windows by the inverse of the ‘background error covariance matrix’ $\hat{A}$. This ‘big B’ of size $Nn \times Nn$ is formed by taking the analysis error covariance from the previous stage and shearing off the oldest row and column.

3.4. General (nonlinear, non-Gaussian) case

We saw in Section 3.1 how the problem of assimilating data immediately it becomes available can be cast into a standard signal model/observation model form (9) and (10), to which we can then apply well-established theory, e.g. Jazwinski (1970). In particular, given (9) and (10) we can compute (sequentially in k) the conditional pdfs $p(\underline{\mathbf{x}}_k|\underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_k)$ and $p(\underline{\mathbf{x}}_{k+1}|\underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_k)$. The novel feature for us is that $\underline{\mathbf{f}}_k$ and $\underline{\mathbf{h}}_k$ in (9) and (10) have a special structure which leads to significant simplifications, in particular enabling us to express these pdfs in terms of the original (as opposed to block) variables.

In general, given the prior pdf $p(\underline{\mathbf{x}}_k|\underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_{k-1})$ and the conditional pdf of the observations given the state $p(\underline{\mathbf{y}}_k|\underline{\mathbf{x}}_k)$, Bayes’ theorem tells us that the posterior pdf $p(\underline{\mathbf{x}}_k|\underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_k)$ is

$$p(\underline{\mathbf{x}}_k|\underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_k) = \frac{p(\underline{\mathbf{y}}_k|\underline{\mathbf{x}}_k)p(\underline{\mathbf{x}}_k|\underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_{k-1})}{\mathcal{N}} \quad (25)$$

where the normalisation $\mathcal{N}$ is

$$\mathcal{N} = \int_{\underline{\mathbf{x}}_k \in \mathcal{U}} \left[p(\underline{\mathbf{y}}_k|\underline{\mathbf{x}}_k) p(\underline{\mathbf{x}}_k|\underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_{k-1}) \right] d\underline{\mathbf{x}}_k \quad (26)$$

and the domain of integration $\mathcal{U} = R^{(N+1)n}$. We will suppose the basic process satisfies the Markov property

$p(\underline{\mathbf{x}}_k|\underline{\mathbf{x}}_{k-1}, \underline{\mathbf{x}}_{k-2}, \dots) = p(\underline{\mathbf{x}}_k|\underline{\mathbf{x}}_{k-1})$ which implies the same for underlined states $p(\underline{\mathbf{x}}_k|\underline{\mathbf{x}}_{k-1}, \underline{\mathbf{x}}_{k-2}, \dots) = p(\underline{\mathbf{x}}_k|\underline{\mathbf{x}}_{k-1})$. Given $p(\underline{\mathbf{x}}_{k-1}|\underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_{k-1})$ (the posterior pdf at $k-1$) and $p(\underline{\mathbf{x}}_k|\underline{\mathbf{x}}_{k-1})$, the Chapman-Kolmogorov equation then gives us for the prior pdf $p(\underline{\mathbf{x}}_k|\underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_{k-1})$

$$\begin{aligned} & p(\underline{\mathbf{x}}_k|\underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_{k-1}) \\ &= \int_{\underline{\tilde{\mathbf{x}}}_{k-1} \in \mathcal{U}} p(\underline{\mathbf{x}}_k|\underline{\tilde{\mathbf{x}}}_{k-1}) p(\underline{\tilde{\mathbf{x}}}_{k-1}|\underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_{k-1}) d\underline{\tilde{\mathbf{x}}}_{k-1} \end{aligned} \quad (27)$$

We could cycle (27) and (25) to obtain the posterior pdf for every k . However, as mentioned there are simplifications arising in this case.

By virtue of the fact that in (8) the i th sub-vector of $\underline{\mathbf{h}}_k$ depends only on $\underline{\mathbf{x}}_{k-N+i-1}$, $p(\underline{\mathbf{y}}_k|\underline{\mathbf{x}}_k)$, the conditional pdf of the observations given the state, factors into

$$p(\underline{\mathbf{y}}_k|\underline{\mathbf{x}}_k) = p(\underline{\mathbf{y}}_k^{(k)}|\underline{\mathbf{x}}_k) \times \dots \times p(\underline{\mathbf{y}}_{k-N}^{(k)}|\underline{\mathbf{x}}_{k-N}) \quad (28)$$

Additionally, the transition pdf $p(\underline{\mathbf{x}}_k|\underline{\tilde{\mathbf{x}}}_{k-1})$ may be written

$$\begin{aligned} & p(\underline{\mathbf{x}}_k|\underline{\tilde{\mathbf{x}}}_{k-1}) \\ &= p(\underline{\mathbf{x}}_k|\underline{\tilde{\mathbf{x}}}_{k-1}) \delta(\underline{\mathbf{x}}_{k-1} - \underline{\tilde{\mathbf{x}}}_{k-1}) \dots \delta(\underline{\mathbf{x}}_{k-N} - \underline{\tilde{\mathbf{x}}}_{k-N}) \end{aligned} \quad (29)$$

Combined with (27) this gives

$$\begin{aligned} & p(\underline{\mathbf{x}}_k|\underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_{k-1}) \\ &= \int_{\underline{\tilde{\mathbf{x}}}_{k-1} \in \mathcal{U}} p(\underline{\mathbf{x}}_k|\underline{\tilde{\mathbf{x}}}_{k-1}) p \left(\begin{array}{c} \underline{\tilde{\mathbf{x}}}_{k-N-1} \\ \underline{\tilde{\mathbf{x}}}_{k-N} \\ \vdots \\ \underline{\tilde{\mathbf{x}}}_{k-1} \end{array} \middle| \underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_{k-1} \right) d\underline{\tilde{\mathbf{x}}}_{k-1} \\ &= \int_{\underline{\tilde{\mathbf{x}}}_{k-N-1} \in R^n} p(\underline{\mathbf{x}}_k|\underline{\mathbf{x}}_{k-1}) p \left(\begin{array}{c} \underline{\tilde{\mathbf{x}}}_{k-N-1} \\ \underline{\mathbf{x}}_{k-N} \\ \vdots \\ \underline{\mathbf{x}}_{k-1} \end{array} \middle| \underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_{k-1} \right) d\underline{\tilde{\mathbf{x}}}_{k-N-1} \\ &= p(\underline{\mathbf{x}}_k|\underline{\mathbf{x}}_{k-1}) \\ &\quad \times \int_{\underline{\tilde{\mathbf{x}}}_{k-N-1} \in R^n} p \left(\begin{array}{c} \underline{\tilde{\mathbf{x}}}_{k-N-1} \\ \underline{\mathbf{x}}_{k-N} \\ \vdots \\ \underline{\mathbf{x}}_{k-1} \end{array} \middle| \underline{\mathbf{y}}_0, \dots, \underline{\mathbf{y}}_{k-1} \right) d\underline{\tilde{\mathbf{x}}}_{k-N-1} \end{aligned} \quad (30)$$

We may cycle (30) and (25), (28) to obtain the posterior pdf for every k . For example, if the observation operators $\underline{\mathbf{h}}_i^{(k)}$ and

model $\mathbf{f}_k$ are linear, and the errors $\mathbf{v}_i^{(k)}$, ω_i are Gaussian, then one may check that (30), (25) and (28) imply that the posterior pdf

$$p(\mathbf{x}_k | \mathbf{y}_0, \dots, \mathbf{y}_k) \propto e^{-(J_b + J_q + J_o)}$$

where J_b , J_q , J_o are given by (20)–(22).

4. Example

We illustrate some of the foregoing with a small example, in which the model is the 40-dimensional chaotic model proposed by Lorenz (1996):

$$\frac{dx_i}{dt} = -x[i-2]x[i-1] + x[i-1]x[i+1] - x_i + F, \quad i = 0, \dots, n-1 \quad (31)$$

with $F = 8$, $n = 40$, where $[i]$ denotes $i, \text{mod } n$. This system is integrated using fourth order Runge–Kutta with a time step of 0.05/6 during which time errors grow at a rate corresponding to order one hour in an atmospheric system. We therefore refer to time step k as time k . The truth is obtained by integrating (31) and to each component of x adding Gaussian model error every time step with variance σ_q^2 .

We suppose that at time k eight observations have just become available, at:

- points $x[k+1]$, $x[k+11]$, $x[k+21]$ and $x[k+31]$, valid at time $k-1$
- point $x[k+6]$, valid at time $k-2$
- point $x[k+16]$, valid at time $k-3$
- point $x[k+26]$, valid at time $k-4$
- point $x[k+36]$, valid at time $k-5$

In this example we suppose there are no ‘instantaneously available’ observations (i.e. available at time k and also valid at k). Each observation has Gaussian error with variance σ_o^2 where $\sigma_o = 0.546$. Every grid point is observed every 5 time steps and the observation network repeats itself exactly every 10 time steps. The system is well-observed and for the values of σ_o , σ_q used here the departure from linearity small enough that the foregoing linear theory can be well applied to the linearised model.

4.1. Non-overlapping windows with different lags

Consider first traditional non-overlapping window strategies. Suppose we have a 6-hour cycle with 6-hour windows. For the window $[k-6, k)$ we wish to assimilate observations valid in $k-6 \leq t < k$. For non-overlapping windows we use an optimal¹ smoother, i.e. 4D-Var (Li and Navon, 2001) with model error correctly accounted for and correct cycling of background error covariances. For the window $[k-6, k)$ this produces analyses $\{\mathbf{x}_{k-6}^a, \dots, \mathbf{x}_{k-1}^a\}$.

As discussed in Section 2, the number of observations which are valid in the window and available for the analysis increases the longer the interval of time between the end of the window

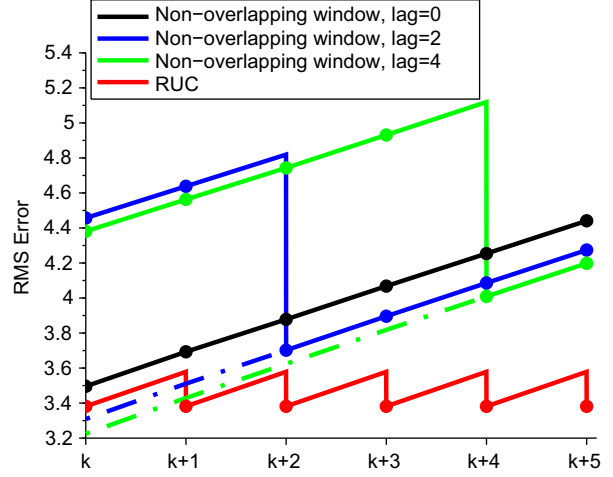


Fig. 4. RMS forecast error from most recent available analysis at times 0–5 hours into the next window following an assimilation window $[k-6, k)$, for lags 0 (black), 2 (blue) and 4 (green). Dashed lines denote RMS error at times *before* the analysis is performed. Also shown in red is the RMS error in the optimal RUC analysis using the data available at each time. In this example $\sigma_q = 0.182$.

and when the analysis is performed, which we term the *lag*. For the present example the number of observations available at each time in the window if the analysis is performed at k , $k+2$, $k+4$ is shown in Table 2.

Taking for example the lag=2 case, at times k and $k+1$ the most recently available analyses are those in the window $[k-12, k-6)$ with last analysis at $k-7$, while at $k+2, \dots, k+5$ the most recently available analysis is at $k-1$. Table 3 shows the validity time of the most recent analysis available at times $k, \dots, k+5$ for lags of 0, 2 and 4 hours.

Fig. 4 shows the RMS forecast error (using $\sigma_q = 0.182$ and averaged over 2000 cycles) in forecasts valid at times $t = k, k+1, \dots, k+5$ taken from the latest available analysis using lag=0 (in black), lag=2 (in blue) and lag=4 (in green). Fig. 4 illustrates a point about non-overlapping windows made in Section 2 above, that we choose between a short lag between observations and when the analysis is performed, giving timely analyses but not using all the observations, or a longer lag using more observations, but which at any given time requires longer forecasts which will be more degraded by model error.

For the lag=2 and lag=4 cases we also show (dashed lines) the RMS forecast error for times between the end of the analysis window and the time the analysis is performed.

4.2. Optimal RUC

Since in our example the longest delay in receipt of observations is 5 h, for the optimal RUC method of Section 3 we have $N = 5$. At time j this produces analyses $\{\mathbf{x}_{j-5}^a, \dots, \mathbf{x}_j^a\}$.

Table 2. Number of observations valid at times $k - 6, \dots, k - 1$, available for analyses performed at $k, k + 2, k + 4$.

Analysis performed at	Lag (hours)	No of obs available with validity time:						Total no. of obs available
		k-6	k-5	k-4	k-3	k-2	k-1	
k	0	8	8	7	6	5	4	38
k+2	2	8	8	8	8	7	6	45
k+4	4	8	8	8	8	8	8	48

Table 3. Validity time of the most recent analysis available at times $k, \dots, k + 5$ for lags of 0, 2 and 4 hours.

Lag (hours)	Validity time of most recent analysis available at:					
	k	k+1	k+2	k+3	k+4	k+5
0	k-1	k-1	k-1	k-1	k-1	k-1
2	k-7	k-7	k-1	k-1	k-1	k-1
4	k-7	k-7	k-7	k-7	k-1	k-1

A comparison of the observation usage of non-overlapping windows and optimal RUC was illustrated in Figs. 2 and 3. In optimal RUC all observations are used, as soon as they are received.

The red curve in Fig. 4 shows the RMS error in the RUC analysis at $\mathbf{x}_j^a$ for $j = k, \dots, k + 5$. From the foregoing we know this will always be less than the RMS error at j using any available analysis using a non-overlapping window with any lag. Note however that this error can be greater than that from the lagged analyses run at a later time, eg, in this example for the lag-2 and lag-4 analyses at time k . This is because the lagged analyses are using observations not available at time k .

5. Suboptimal methods for RUC

In Section 3 above we derived the optimal solution to the problem of assimilating data as soon as it becomes available, and saw that if the maximum delay is N and the state is described by n variables that this involved manipulating vectors of size $n(N+1)$ and their error covariances of size $n(N+1) \times n(N+1)$.

If observation and model errors are uncorrelated in time then optimal data assimilation methods for non-overlapping windows only involve vectors and matrices of size n and $n \times n$.

For large-scale systems manipulating vectors of size $n(N+1)$ and more particularly matrices of size $n(N+1) \times n(N+1)$ may not be manageable. Furthermore, NWP centres already have methods implemented for non-overlapping windows (we will refer to these as ‘traditional methods’) and will naturally seek ways of adapting these to the RUC problem. Hence a topic of practical importance is the relationship between the optimal solution to RUC and traditional methods applied to RUC.

For simplicity we restrict attention to the case of linear forecast and observation operators whereas in Section 3.2 the errors are Gaussian and uncorrelated. We will designate the optimal solution for RUC in this case (i.e. Table 1) as Method 0. We will develop suboptimal methods for RUC based on traditional methods for non-overlapping windows, with Method 3 a ‘naive’ application of such a method to RUC, and Methods 2 and 1 adaptations of this which are progressively closer to the optimal solution. We will then examine the relation between the four methods.

5.1. ‘Traditional’ methods as suboptimal methods for RUC

Suppose that at time k we have prior estimates

$$\{\mathbf{x}_j^b, j = k - N, \dots, k\} \quad (32)$$

and we have just received observations

$$\{\mathbf{y}_j^{(k)}, j = k - N, \dots, k\}$$

A natural extension of the 4D-Var method (e.g. Li and Navon (2001)) as used for non-overlapping windows is to form analyses at $j = k - N, \dots, k$

$$\mathbf{x}_j^a = \mathbf{x}_j^b + \delta_j, j = k - N, \dots, k \quad (33)$$

where

$$\underline{\delta} = \begin{pmatrix} \delta_{k-N} \\ \vdots \\ \delta_k \end{pmatrix}$$

minimises

$$\begin{aligned} J(\underline{\delta}) &= \frac{1}{2} (\delta_{k-N})^T B^{-1} \delta_{k-N} \\ &+ \frac{1}{2} \sum_{j=k-N}^{j=k-1} (\delta_{j+1} - M_j \delta_j)^T Q_j^{-1} (\delta_{j+1} - M_j \delta_j) \\ &+ \frac{1}{2} \sum_{j=k-N}^{j=k} (\mathbf{y}_j^{(k)} - H_j^{(k)} (\mathbf{x}_j^b + \delta_j))^T \\ &\times R_j^{-1} (\mathbf{y}_j^{(k)} - H_j^{(k)} (\mathbf{x}_j^b + \delta_j)) \end{aligned} \quad (34)$$

B in (34) is the error covariance of $\mathbf{x}_{k-N}^b$, if this is known, or some approximation otherwise; we return to this point below. All our suboptimal methods 1–3 use (33) and (34) for the analysis. There are many different ways of forming new priors $\{\mathbf{x}_j^b, j = k - N + 1, \dots, k + 1\}$. A non-exhaustive selection of possibilities is:

Method 3: The most similar to traditional 4D-Var, in which the only state saved from the above analyses is the first one $\mathbf{x}_{k-N}^a$ (i.e. at the beginning of the window). Denoting the model evolution from $k - N$ to j by $M_{k-N}^j = M_{j-1} \dots M_{k-N+1} M_{k-N}$ this gives us priors

$$\mathbf{x}_j^b = M_{k-N}^j \mathbf{x}_{k-N}^a, \quad j = k - N + 1, \dots, k + 1 \quad (35)$$

Method 2: Slightly better is to save and use the second analysed state, i.e. at time $k - N + 1$, which will be at the beginning of the next window, giving us priors

$$\mathbf{x}_j^b = M_{k-N+1}^j \mathbf{x}_{k-N+1}^a, \quad j = k - N + 1, \dots, k + 1 \quad (36)$$

Method 1: Finally, we could follow the optimal solution and save and use all the analysed states, giving us priors

$$\mathbf{x}_j^b = \begin{cases} \mathbf{x}_j^a & j = k - N + 1, \dots, k \\ M_k^{k+1} \mathbf{x}_k^a & j = k + 1 \end{cases} \quad (37)$$

The number of prior states which are simply analysis states from the previous cycle is therefore 0, 1 and N for Methods 3, 2 and 1, respectively. The four RUC methods (with the optimal one of Section 3.2 labelled ‘Method 0’) are summarised in Table 4. The formation of backgrounds is shown in more detail for $N = 2$ in Table 5.

Table 4. Summary of methods 0,1,2,3.

Method	0	1	2	3
Number of prior states which are analysis states from previous cycle	N	N	1	0
Formulation of $\mathbf{x}_k^b$	(37)	(37)	(36)	(35)
Formulation of $\mathbf{x}_k^a$	(16)	(33), (34)	(33), (34)	(33), (34)

Table 5. The background states for the analysis at $t = 3$ given the analysed states available from $t = 2$, illustrated with $N = 2$ for Methods 0,1,2,3. After the analysis at $t = 2$ we have analyses $\mathbf{x}_0^a, \mathbf{x}_1^a$ and $\mathbf{x}_2^a$. The backgrounds $\mathbf{x}_1^b, \mathbf{x}_2^b$ and $\mathbf{x}_3^b$ at $t = 3$ are formed from these analyses as shown in the table.

Time	0	1	2	3
	Analyses at $t = 2$			
	$\mathbf{x}_0^a$	$\mathbf{x}_1^a$	$\mathbf{x}_2^a$	
	Backgrounds at $t = 3$			
		$\mathbf{x}_1^b$	$\mathbf{x}_2^b$	$\mathbf{x}_3^b$
Methods 0 & 1		$\mathbf{x}_1^a$	$\mathbf{x}_2^a$	$M_2^3 \mathbf{x}_2^a$
Method 2		$\mathbf{x}_1^a$	$M_1^2 \mathbf{x}_1^a$	$M_1^3 \mathbf{x}_1^a$
Method 3	$M_0^1 \mathbf{x}_0^a$	$M_0^2 \mathbf{x}_0^a$	$M_0^3 \mathbf{x}_0^a$	

5.2. Covariance of analysis and background errors in methods 1–3

To compare the various methods we will need the covariance of their errors. We may write (33) and (34) as

$$\mathbf{x}_k^a = \mathbf{x}_k^b + \underline{K}_k (\mathbf{y}_k - \underline{H}_k \mathbf{x}_k^b) \quad (38)$$

where

$$\begin{aligned} \underline{H}_k &= \text{diag}(H_{k-N}^{(k)}, \dots, H_k^{(k)}) \\ (\underline{U}_k)_{i,j} &= \begin{cases} M_{k-N+j-1}^{k-N+i-1}, & \text{for } 1 \leq j \leq N+1 \\ 0 & \text{for } j > i \end{cases} \end{aligned} \quad (39)$$

$$\underline{B}_k^v = \underline{U}_k \text{diag}(B, Q_{k-N}, \dots, Q_{k-1}) \underline{U}_k^T \quad (40)$$

$$\underline{K}_k = \underline{B}_k^v \underline{H}_k^T (\underline{H}_k \underline{B}_k^v \underline{H}_k^T + \underline{R})^{-1} \quad (41)$$

where we use the notation $(\underline{U}_k)_{i,j}$ to denote the i, j $n \times n$ submatrix of $\underline{U}_k$. Denoting the truth by $\mathbf{x}_k^t$ and

$$\underline{\epsilon}_k^b = \mathbf{x}_k^b - \mathbf{x}_k^t, \quad \underline{B}_k = E[\underline{\epsilon}_k^b (\underline{\epsilon}_k^b)^T], \quad \underline{\epsilon}_k^a = \mathbf{x}_k^a - \mathbf{x}_k^t$$

Table 6. Values of $\underline{M}_k^{(\ell)}$, for $N = 2$, following Table 5.

	Method $\ell = 1$	Method $\ell = 2$	Method $\ell = 3$
$\underline{M}_k^{(\ell)}$	$\begin{pmatrix} 0 & I & 0 \\ 0 & 0 & I \\ 0 & 0 & M_k^{k+1} \end{pmatrix}$	$\begin{pmatrix} 0 & I & 0 \\ 0 & M_{k-1}^k & 0 \\ 0 & M_{k-1}^{k+1} & 0 \end{pmatrix}$	$\begin{pmatrix} M_{k-2}^{k-1} & 0 & 0 \\ M_{k-2}^k & 0 & 0 \\ M_{k-2}^{k+1} & 0 & 0 \end{pmatrix}$

it follows that for all Methods 1–3 the **analysis error covariance** $\underline{A}_k$ is, from (38)

$$\underline{A}_k \equiv E[\underline{\epsilon}_k^a (\underline{\epsilon}_k^a)^T] = (I - \underline{K}_k \underline{H}_k) \underline{B}_k (I - \underline{K}_k \underline{H}_k)^T + \underline{K}_k \underline{R} \underline{K}_k^T \quad (42)$$

We note this depends both on $\underline{B}_k$ and (via $\underline{K}_k$ and $\underline{B}_k^v$) the prescribed $\underline{B}$.

The **background error covariance** $\underline{B}_{k+1}$ depends on which method is used. For method ℓ we may write

$$\underline{x}_{k+1}^b = \underline{M}_k^{(\ell)} \underline{x}_k^a \quad (43)$$

where $\underline{M}_k^{(1)} = \underline{M}_k^{(0)} = \underline{M}_k$ as specified in (12) and $\underline{M}_k^{(2)}, \underline{M}_k^{(3)}$ are illustrated for $N = 2$ in Table 6. Since the error in $\underline{x}_{k+1}^b$ using method ℓ is

$$\begin{aligned} \underline{\epsilon}_{k+1}^b &\equiv \underline{x}_{k+1}^b - \underline{x}_{k+1}^t = \underline{M}_k^{(\ell)} (\underline{x}_k^a - \underline{x}_k^t) + \underline{M}_k^{(\ell)} \underline{x}_k^t - \underline{x}_{k+1}^t \\ &= \underline{M}_k^{(\ell)} \underline{\epsilon}_k^a + \underline{\epsilon}_k^q \end{aligned} \quad (44)$$

where $\underline{\epsilon}_k^q = \underline{M}_k^{(\ell)} \underline{x}_k^t - \underline{x}_{k+1}^t$ we can express $\underline{B}_{k+1}$ in terms of $\underline{M}_k^{(\ell)}, \underline{A}_k$, model error covariance and the cross-covariance of analysis error $\underline{\epsilon}_k^a$ and model error $\underline{\epsilon}_k^q$:

$$\begin{aligned} \underline{B}_{k+1} &= E \left[\underline{\epsilon}_{k+1}^b (\underline{\epsilon}_{k+1}^b)^T \right] = \underline{M}_k^{(\ell)} E \left[\underline{\epsilon}_k^a (\underline{\epsilon}_k^a)^T \right] (\underline{M}_k^{(\ell)})^T \\ &\quad + E \left[\underline{\epsilon}_k^q (\underline{\epsilon}_k^q)^T \right] + \underline{M}_k^{(\ell)} E \left[\underline{\epsilon}_k^a (\underline{\epsilon}_k^q)^T \right] \\ &\quad + E \left[\underline{\epsilon}_k^q (\underline{\epsilon}_k^a)^T \right] (\underline{M}_k^{(\ell)})^T \end{aligned} \quad (45)$$

In order to cycle Methods 1–3 we need to specify $\underline{B}$ in (34). For the rest of this section, for the purposes of comparing the four methods, we will suppose that in Methods 1–3 we use $\underline{B} = (\underline{B}_k)_{1,1}$, which can be obtained from (45) (as above, underlined subscripts refer to submatrices, so $(\underline{B}_k)_{1,1}$ is the top left $n \times n$ submatrix of $\underline{B}_k$). It can be shown that for methods 1 and 2 that $(\underline{B}_k)_{1,1} = (\underline{A}_{k-1})_{2,2}$.

5.3. Relation between methods 0 and 1

An important comparison is between optimal Method 0 and suboptimal Method 1. They share the same background step (43) with the same $\underline{M}_k$. In both cases the analysis step may be written in the form (38), though whereas in Method 0 the gain $\underline{K}_k$

$$\underline{K}_k = \underline{B}_k \underline{H}_k^T (\underline{H}_k \underline{B}_k \underline{H}_k^T + \underline{R}_k)^{-1}$$

uses true $\underline{B}_k$, i.e. the covariance of the error in $\underline{x}_k^b$, for Method 1 the gain $\underline{K}_k$ is

$$\underline{K}_k = \underline{B}_k^v \underline{H}_k^T (\underline{H}_k \underline{B}_k^v \underline{H}_k^T + \underline{R}_k)^{-1}$$

where $\underline{B}_k^v$ is defined in (40). Because of the similarity in the structures of Methods 0 and 1 there is a simple and strong relation between their errors. If the sequence of background error covariances using Methods 0 and 1 are designated respectively $\underline{B}_k$ and $\tilde{\underline{B}}_k$, and we start from the same prior error covariance $\tilde{\underline{B}}_N = \underline{B}_N$, then for all $k \geq N$ the difference $\tilde{\underline{B}}_k - \underline{B}_k$ is positive semi-definite, usually written

$$\underline{B}_k \preceq \tilde{\underline{B}}_k$$

This is proved in Appendix B.

5.4. Relation between methods in limit $Q_k \rightarrow \infty$

We have four RUC methods, the optimal and three suboptimal ones. We can cycle each as described above, for the suboptimal methods using $\underline{B} = (\underline{B}_k)_{1,1}$ (Section 5.2).

A limiting case which exhibits some of the differences between them, in particular how information is saved from previous cycles, is obtained by letting model error covariance $Q_k \rightarrow \infty$ for all k .

For simplicity suppose $\underline{H} = I$ and $\underline{R}_k$ and Q_k are independent of k . If $Q = \infty$ then after N cycles for method 0, one cycle for Methods 1 and 2, and immediately for Method 3, all knowledge of the initial background state $\underline{x}_0^b$ and its error covariance is lost. In Table 7 we show the analysed state $\underline{x}_{j+N}^a$ produced by the four methods for any $j \geq N$ if model error is infinite, and the corresponding background error and analysis error covariances.

The optimal Method 0 retains all the observation information ever received; at time $j + N$ the estimate of state at any time between j and $j + N$ is simply the average of all the observations ever received valid at that time. At the other extreme, Method 3 ‘forgets’ all the observation information from previous cycles: at time $j + N$ the estimate of state at any time between j and $j + N$ is just the value of the observation valid at that time and received at time $j + N$. Methods 1 and 2 retain observation information from the previous cycle only at the initial time. These different

Table 7. $\underline{\mathbf{x}}_{j+N}^a$, $\underline{B}_{N+j}^{-1}$ and $\underline{A}_{N+j}^{-1}$ for different methods in limit $Q \rightarrow \infty$, for $j \geq N$.

	Method 0	Methods 1 and 2	Method 3
$\underline{\mathbf{x}}_{j+N}^a$	$\begin{pmatrix} \frac{1}{N+1} (\mathbf{y}_j^{(j)} + \mathbf{y}_j^{(j+1)} + \dots + \mathbf{y}_j^{(j+N)}) \\ \frac{1}{N} (\mathbf{y}_{j+1}^{(j+1)} + \dots + \mathbf{y}_{j+1}^{(j+N)}) \\ \vdots \\ \frac{1}{2} (\mathbf{y}_{j+N-1}^{(j+N-1)} + \mathbf{y}_{j+N-1}^{(j+N)}) \\ \mathbf{y}_{j+N}^{(j+N)} \end{pmatrix}$	$\begin{pmatrix} \frac{1}{2} (\mathbf{y}_j^{(j+N-1)} + \mathbf{y}_j^{(j+N)}) \\ \mathbf{y}_{j+1}^{(j+N)} \\ \vdots \\ \mathbf{y}_{j+N-1}^{(j+N)} \\ \mathbf{y}_{j+N}^{(j+N)} \end{pmatrix}$	$\begin{pmatrix} \mathbf{y}_j^{(j+N)} \\ \mathbf{y}_{j+1}^{(j+N)} \\ \vdots \\ \mathbf{y}_{j+N-1}^{(j+N)} \\ \mathbf{y}_{j+N}^{(j+N)} \end{pmatrix}$
$\underline{B}_{N+j}^{-1}$	$\text{diag} \left(NR^{-1}, (N-1)R^{-1}, \dots, R^{-1}, 0 \right)$	$\text{diag} \left(\underbrace{R^{-1}, \dots, R^{-1}}_{N \text{ times}}, 0 \right)$ and $\text{diag} \left(R^{-1}, \underbrace{0, \dots, 0}_{N \text{ times}} \right)$	$\text{diag} \left(\underbrace{0, \dots, 0}_{N+1 \text{ times}} \right)$
$\underline{A}_{N+j}^{-1}$	$\text{diag} \left((N+1)R^{-1}, NR^{-1}, \dots, 2R^{-1}, R^{-1} \right)$	$\text{diag} \left(2R^{-1}, \underbrace{R^{-1}, \dots, R^{-1}}_{N \text{ times}} \right)$	$\text{diag} \left(\underbrace{R^{-1}, \dots, R^{-1}}_{N+1 \text{ times}} \right)$

behaviours are reflected in the analysis error covariances shown in Table 7.

Comparing Methods 1 and 2, while in the limit $Q \rightarrow \infty$ Method 1 analyses are no better than those of Method 2, we note in Table 7 that it has better backgrounds than Method 2.

5.5. Relation between methods in the limit $Q_k \rightarrow 0$

Suppose that $Q_k = 0$ for all k , and we are given an estimate $\mathbf{x}_0^b$ of $\mathbf{x}_0$ with error covariance B_0 . Estimate the remaining states in the window $\{0, \dots, N\}$ by

$$\{\mathbf{x}_j^b = M_0^j \mathbf{x}_0^b, j = 1, \dots, N\}$$

If $Q_k = 0$ for all k then Methods 1–3 simplify to their ‘strong constraint’ forms, in which (34) simplifies to

$$\begin{aligned} J(\delta_{k-N}) &= \frac{1}{2} (\delta_{k-N})^T B_{k-N}^{-1} \delta_{k-N} \\ &+ \frac{1}{2} \sum_{j=k-N}^{j=k} \left(\mathbf{y}_j^{(k)} - H_j^{(k)} M_{k-N}^j \left(\mathbf{x}_{k-N}^b + \delta_{k-N} \right) \right)^T \\ &\times R_j^{-1} \left(\mathbf{y}_j^{(k)} - H_j^{(k)} M_{k-N}^j \left(\mathbf{x}_{k-N}^b + \delta_{k-N} \right) \right) \end{aligned} \quad (46)$$

Crucially, in the limit $Q \rightarrow 0$ Methods 0–3 coincide, so in particular Methods 1–3 are now optimal. This is proved in Appendix C. In the absence of model error the ‘suboptimal’ methods all coincide with each other, and in fact are in this case optimal.

5.6. Comparison of methods 0–3 for the example of Section 4

We may apply *linearised* versions of Methods 0–3 to our non-linear chaotic example of Section 4. In an attempt to mitigate the effects of linearisation error one can formulate outer loop-style iterations for these strategies, which may be worth implementing in more non-linear systems (for the examples here they made negligible difference). Alternatively one could use the ‘best linear approximation’ (Payne, 2013).

In our example of Section 4, at time k observations have just become available which are valid at $k-1, \dots, k-5$ so $N = 5$. The optimal Method 0 and suboptimal Methods 1–3 all provide analyses $\mathbf{x}_k^a, \mathbf{x}_{k-1}^a, \dots, \mathbf{x}_{k-5}^a$. (Since in our example no observations are instantaneously available, $\mathbf{x}_k^a$ is here a forecast from $\mathbf{x}_{k-1}^a$).

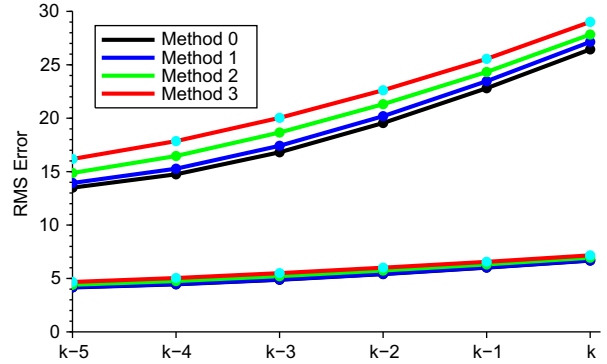


Fig. 5. RMS error in analysis at $k-5, \dots, k$, various RUC strategies using all observations available at time k , $\sigma_q = 1.82$ (upper set) and $\sigma_q = 0.455$ (lower set).

Each strategy is cycled 10,000 times and the first hundred cycles disregarded. Figure 5 shows the RMS error in the analyses at $k-5, \dots, k$, for the optimal Method 0 (black), and suboptimal Methods 1 (blue), 2 (green) and 3 (red), for $\sigma_q = 1.82$ and 0.455 (upper and lower set respectively).

As expected from the foregoing, the errors $\mathcal{E}$ are ordered

$$\mathcal{E}(\text{Method 0}) < \mathcal{E}(\text{Method 1}) < \mathcal{E}(\text{Method 2}) < \mathcal{E}(\text{Method 3})$$

Furthermore, the analyses and therefore their errors converge as $\sigma_q \rightarrow 0$.

6. Impact of climatological background error covariances

A significant difference between the methods of the preceding sections and those used for large scale systems is that in the latter the prior error covariance (usually denoted B) is not cycled, but is either constant or is a convex combination of a constant and an estimate of cycled B (see (47) below).

It is important to note that insofar as B is fixed it is advantageous to assimilate data as far as possible simultaneously in larger units rather than split it up into smaller volumes and assimilate it in smaller units. Intuitively, by assimilating many observations simultaneously the deficiencies of the fixed B are reduced.

Illustrating this point is complicated by the fact that increasing observation batch size tends to involve making other changes which themselves have an impact. If we compare cycling using non-overlapping windows with RUC then at no instant in time are the two methodologies assimilating the same observations (see Section 2). If we use 4D-Var to compare cycling with windows of length 1 (assimilating one observation every time step) with windows of length 2 (two observations every two time steps) then the latter has the advantage of covariances evolved through the window, which is a different point to the one being made.

If we compare assimilating two observations simultaneously every time step with assimilating one after the other, in the latter case we have to decide what B to use for the second observation. In Appendix D we show that if in a scalar system we assimilate two observations every cycle, and have a choice between assimilating them

- (a) simultaneously using a fixed background error covariance B , or
- (b) separately, the first with fixed background error covariance B_1 and the second with fixed background error covariance B_2 ,

then it is possible to choose B so that no matter how well B_1 and B_2 are chosen strategy (a) will always outperform strategy (b).

We may contrast this result concerning the use of *fixed* background covariances with the fact that, if B is chosen optimally

every cycle, and B_1 and B_2 are chosen optimally every cycle, the two strategies will produce identical (and optimal) results.

This means that if B is fixed then, in this respect, RUC is at a disadvantage compared with conventional cycling as now there are more cycles with fewer observations used every cycle. In practice this effect may be dwarfed by the advantages of RUC as discussed in Section 2. If not, the obvious remedy is to improve the cycling of the background error covariances.

As noted above, the effect is due to using a fixed B , and is removed if B is cycled properly. This is unattainable in current large-scale NWP, but centres such as the Met Office already employ a ‘hybrid’ B (Clayton et al., 2013)

$$B = \beta_c^2 B_c + \beta_e^2 B_e \quad (47)$$

where B_c is fixed but B_e is an estimate (from an ensemble) of the true prior error covariance, with $\beta_c^2 + \beta_e^2 = 1$. For best performance $\beta_e \rightarrow 1$ as ensemble size increases, with the fixed part having no weight in the limit of an infinitely large ensemble.

There are other possible ways of introducing adaptivity into B , such as the ‘ensemble-variation integrated localised’ method (Auligné et al., 2016) and the so-called variational Kalman filters (Auvinen et al., 2010). In the latter the limited-memory quasi-Newton method is used to build a low-storage approximation to the Hessian of the analysis cost function, which therefore approximates the inverse of the analysis error covariance matrix, and can also be used to evolve the covariance forward to approximate B at the next analysis time.

7. Concluding remarks

Rapid update cycling (RUC) is the process by which we assimilate observations into a model as soon as they become available, or more practically, within some (short) time interval δt . We have seen that if the greatest delay in receiving observations is $N\delta t$ then the optimal solution to RUC at time k involves manipulating the $(N+1)n$ vectors $\underline{x}$ formed from $\mathbf{x}_{k-N}, \dots, \mathbf{x}_k$ and the moments of the errors in $\underline{x}$, such as their $(N+1)n \times (N+1)n$ error covariances.

Compared with ‘traditional’ cycling RUC makes more timely use of observations, which is particularly important for the provision of LBCs for LAMs. Another advantage of RUC is that the increments are smaller and hence linearisation error is reduced.

We have purposely concentrated on fundamental topics and avoided such practically important matters as efficiency and cost. The fact that for each analysis observation volumes and increments are smaller, and we always have a recent one available, suggest that it should be possible to reduce the cost per analysis.² On contemporary HPCs where the increased power comes through higher numbers of processors rather than increased clock speeds this is more important than the total cost per day.³

We may adapt ‘traditional’ methods designed for non-overlapping windows to RUC. These methods are suboptimal, but in all cases considered in this paper (Methods 1,2 and 3 in Section 4) they coincide with the optimal solution in the limit where model error vanishes.

Assimilating observations in smaller batches can be disadvantageous if climatological background error variances are used. This potentially poses a challenge for RUC, which could perhaps be met by improved cycling of error covariances.

Acknowledgements

The author thanks Mike Cullen and Andrew Lorenc for useful discussions on this topic, and the referees appointed by Tellus for their comments.

Disclosure statement

No potential conflict of interest was reported by the author.

Notes

1. Optimal except that we ignore linearisation error, which is small for our example.
2. Note also that there is scope for preconditioning using the work already done for recent analyses. This preconditioning could be based on Hessian eigenvectors (Fisher and Courtier, 1995), or on the vectors approximating the Hessian in the limited memory quasi-Newton method (Courtier et al., 1998).
3. We have also noted that the window length of RUC is determined by the longest delay in receiving observations, which in the operational example in Section 1 implies a window of 4 hours compared with the current 6 hours.

References

- Anderson, B. and Moore, J. 1979. *Optimal Filtering* Prentice-Hall, Englewood Cliffs, NJ. 357 pp.
- Auligné, T., Ménétrier, B., Lorenc, A. C. and Buehner, M. 2016. Ensemble-variational integrated localized data assimilation. *Mon. Weather Rev.* **144**, 3677–3696. DOI:10.1175/mwr-d-15-0252.1.
- Auvinen, H., Bardsley, J. M., Haario, H. and Kauranne, T. 2010. The Variational Kalman filter and an efficient implementation using limited memory BFGS. *Int. J. Numer. Methods Fluids* **64**, 314–335. DOI:10.1002/fld.2153.
- Benjamin, S. G., Dévényi, D., Weygandt, S. S., Brundage, K. J., Brown, J. M., and co-authors. 2004a. An hourly assimilation-forecast cycle: the RUC. *Mon. Weather Rev.* **132**, 495–518. DOI:10.1175/1520-0493.
- Benjamin, S. G., Grell, G. A., Brown, J. M., Smirnova, T. G. and Bleck, R. 2004b. Mesoscale weather prediction with the RUC hybrid isentropic terrain-following coordinate model. *Mon. Weather Rev.* **132**, 473–494. DOI:10.1175/1520-0493.
- Boyd, S. and Vandenberghe, L. 2004. *Convex Optimization* Cambridge University Press, New York, NY. 716 pp.
- Clayton, A. M., Lorenc, A. C. and Barker, D. M. 2013. Operational implementation of a hybrid ensemble/4D-Var global data assimilation system at the Met Office. *Q. J. R. Meteorol. Soc.* **139**, 1445–1461. DOI:10.1002/qj.2054.
- Courtier, P., Andersson, E., Heckley, W., Vasiljevic, D., Hamrud, M., and co-authors. 1998. The ECMWF implementation of three-dimensional variational assimilation (3D-Var). I: formulation. *Q. J. R. Meteorol. Soc.* **124**, 1783–1807. DOI:10.1002/qj.49712455002.
- Fisher, M. and Courtier, P. 1995. Estimating the covariance matrices of analysis and forecast error in variational data assimilation. *Tech. Memo.* **220**, 28. ECMWF, Shinfield Park, Reading, UK.
- Gallier, J. 2010. The Schur complement and symmetric positive semidefinite (and definite) matrices. *Penn Eng.* Online at: www.cis.upenn.edu/~jean/schur-comp.pdf
- Järvinen, H., Thépaut, J.-N. and Courtier, P. 1996. Quasi-continuous variational data assimilation. *Q. J. R. Meteorol. Soc.* **122**, 515–534. DOI:10.1002/qj.49712253011.
- Jazwinski, A. H. 1970. *Stochastic Processes and Filtering Theory* Academic Press, New York. 376 pp.
- Li, Z. and Navon, I. M. 2001. Optimality of variational data assimilation and its relationship with the Kalman filter and smoother. *Q. J. R. Meteorol. Soc.* **127**, 661–683. DOI:10.1002/qj.49712757220.
- Lorenz, E. 1996. Predictability – a problem partly solved. *Proceedings, Seminar on Predictability* Vol. 1, ECMWF, Reading, UK, pp. 1–18.
- Payne, T. J. 2013. The linearisation of maps in data assimilation. *Tellus A: Dyn. Meteorol. Oceanogr.* **65**, DOI:10.3402/tellusa.v65i0.18840.
- Tang, Y., Lean, H. W. and Bornemann, J. 2013. The benefits of the Met Office variable resolution NWP model for forecasting convection. *Met. Apps* **20**, 417–426. DOI:10.1002/met.1300.
- Veerse, F. and Thépaut, J.-N. 1998. Multiple-truncation incremental approach for four-dimensional variational data assimilation. *Q. J. R. Meteorol. Soc.* **124**, 1889–1908. DOI:10.1002/qj.49712455006.

Appendix A. Proof of variational form of optimal solution given in Section 3.3

Continuing with the notation of Section 3.3 we readily obtain identities (A1)–(A3) below:

(i)

$$J_b(\delta) + J_q(\delta) = \frac{1}{2} \delta^T \left[\begin{pmatrix} \hat{A}^{-1} & 0 \\ 0 & 0 \end{pmatrix} + \begin{pmatrix} -f^T \\ I \end{pmatrix} Q_{k-1}^{-1} \begin{pmatrix} -f & I \end{pmatrix} \right] \delta \quad (\text{A1})$$

where f is the $n \times nN$ matrix

$$f = \underbrace{(0 \ \cdots \ 0 \ M_{k-1})}_{N-1 \text{ times}}$$

(ii) By elementary algebra

$$\begin{aligned} & \left[\begin{pmatrix} I \\ f \end{pmatrix} \hat{A} \begin{pmatrix} I & f^T \end{pmatrix} + \begin{pmatrix} 0 & 0 \\ 0 & Q \end{pmatrix} \right] \\ & \times \left[\begin{pmatrix} -f^T \\ I \end{pmatrix} Q^{-1} \begin{pmatrix} -f & I \end{pmatrix} + \begin{pmatrix} \hat{A}^{-1} & 0 \\ 0 & 0 \end{pmatrix} \right] \\ & = \begin{pmatrix} I & 0 \\ 0 & I \end{pmatrix} \end{aligned} \quad (\text{A2})$$

(iii) From the definition of $\underline{M}_{k-1}$ in (12) we have

$$\underline{M}_{k-1} = \begin{pmatrix} 0 & I \\ 0 & f \end{pmatrix}$$

so by construction of $\hat{A}$

$$\begin{pmatrix} I \\ f \end{pmatrix} \hat{A} \begin{pmatrix} I & f^T \end{pmatrix} = \underline{M}_{k-1} \underline{P}_{k-1|k-1} \underline{M}_{k-1}^T \quad (\text{A3})$$

It follows from (A1)–(A3) and (13) that

$$\begin{aligned} & J_b(\underline{\delta}) + J_q(\underline{\delta}) \\ & = \frac{1}{2} \underline{\delta}^T \left[\left(\underline{M}_{k-1} \underline{P}_{k-1|k-1} \underline{M}_{k-1}^T + \underline{Q}_{k-1} \right)^{-1} \right] \underline{\delta} \end{aligned} \quad (\text{A4})$$

At a minimum of (24) we have $J'(\underline{\delta}) = \mathbf{0}$ so

$$\begin{aligned} & \left[\left(\underline{M}_{k-1} \underline{P}_{k-1|k-1} \underline{M}_{k-1}^T + \underline{Q}_{k-1} \right)^{-1} + \underline{H}_k^T \underline{R}_k^{-1} \underline{H}_k \right] \\ & \times \underline{\delta} = \underline{H}_k^T \underline{R}_k^{-1} (\mathbf{y}_k - \underline{H}_k \hat{\mathbf{x}}_{k|k-1}) \end{aligned} \quad (\text{A5})$$

and using this value of $\underline{\delta}$ in (23) is equivalent to (16).

Appendix B. Proof of theorem of Section 5.3

If the sequence of background error covariances using Methods 0 and 1 are designated, respectively, $\underline{B}_k$ and $\tilde{\underline{B}}_k$, and we start from the same prior error covariance $\underline{B}_N = \tilde{\underline{B}}_N$, then for all $k \geq N$

$$\underline{B}_k \leq \tilde{\underline{B}}_k$$

Proof. For simplicity we shall suppose that both $\underline{R}$ and $\underline{B}_k^v$ are non-singular and that $\underline{H} = I$, so denoting Method 1 quantities with tildes we have $\underline{K} = (\underline{B}^{-1} + \underline{R}^{-1})^{-1} \underline{R}^{-1}$ and $\tilde{\underline{K}} = ((\underline{B}^v)^{-1} + \underline{R}^{-1})^{-1} \underline{R}^{-1}$. Hence from (17) and (42) we have

$$\underline{A} = (\underline{B}^{-1} + \underline{R}^{-1})^{-1}$$

$$\begin{aligned} \tilde{\underline{A}} &= ((\underline{B}^v)^{-1} + \underline{R}^{-1})^{-1} \left[(\underline{B}^v)^{-1} \tilde{\underline{B}} (\underline{B}^v)^{-1} + \underline{R}^{-1} \right] \\ &\quad \times ((\underline{B}^v)^{-1} + \underline{R}^{-1})^{-1} \end{aligned} \quad (\text{B1})$$

and therefore

$$\begin{aligned} \underline{A}^{-1} - \tilde{\underline{A}}^{-1} &= (\underline{B}^{-1} - \tilde{\underline{B}}^{-1}) \\ &+ \left[\tilde{\underline{B}}^{-1} + \underline{R}^{-1} - ((\underline{B}^v)^{-1} + \underline{R}^{-1}) \right] \\ &\times \left((\underline{B}^v)^{-1} \tilde{\underline{B}} (\underline{B}^v)^{-1} + \underline{R}^{-1} \right)^{-1} \left((\underline{B}^v)^{-1} + \underline{R}^{-1} \right) \end{aligned} \quad (\text{B2})$$

Now consider the matrix

$$\begin{aligned} \underline{X} &= \begin{pmatrix} \tilde{\underline{B}}^{-1} + \underline{R}^{-1} & (\underline{B}^v)^{-1} + \underline{R}^{-1} \\ (\underline{B}^v)^{-1} + \underline{R}^{-1} & (\underline{B}^v)^{-1} \tilde{\underline{B}} (\underline{B}^v)^{-1} + \underline{R}^{-1} \end{pmatrix} \\ &= \begin{pmatrix} \tilde{\underline{B}}^{-1} \\ (\underline{B}^v)^{-1} \end{pmatrix} \tilde{\underline{B}} \begin{pmatrix} \tilde{\underline{B}}^{-1} & (\underline{B}^v)^{-1} \end{pmatrix} + \begin{pmatrix} I \\ I \end{pmatrix} \underline{R}^{-1} \begin{pmatrix} I & I \end{pmatrix} \end{aligned} \quad (\text{B3})$$

Since $\underline{X}$ is the sum of two positive semi-definite matrices it is positive semi-definite. Note also that since $\underline{R}$ is positive definite that $(\underline{B}^v)^{-1} \tilde{\underline{B}} (\underline{B}^v)^{-1} + \underline{R}^{-1}$ is invertible, so its Moore–Penrose inverse is its usual matrix inverse.

The term in square brackets in (B2) is the Schur complement of $(\underline{B}^v)^{-1} \tilde{\underline{B}} (\underline{B}^v)^{-1} + \underline{R}^{-1}$ in $\underline{X}$ so by Theorem 4.3 of Gallier (2010) is positive semi-definite. Therefore if $\underline{B}_k^{-1} \succeq \tilde{\underline{B}}_k^{-1}$ then $\underline{A}_k^{-1} - \tilde{\underline{A}}_k^{-1}$ is the sum of two positive semi-definite matrices, so is positive semi-definite. We conclude that if $\underline{B}_k \leq \tilde{\underline{B}}_k$ then $\underline{A}_k \leq \tilde{\underline{A}}_k$. Since both methods satisfy (43) with the same $\underline{M}_k$, both satisfy the same relation

$$\underline{B}_{k+1} = \underline{M}_k \underline{A}_k \underline{M}_k^T + \underline{Q}, \quad \tilde{\underline{B}}_{k+1} = \underline{M}_k \tilde{\underline{A}}_k \underline{M}_k^T + \underline{Q}$$

Therefore $\underline{A}_k \leq \tilde{\underline{A}}_k$ implies $\underline{B}_{k+1} \leq \tilde{\underline{B}}_{k+1}$ and the claim follows. $\square$

Appendix C. Proof of equivalence of methods 0–3 if $\underline{Q} = \mathbf{0}$

Set $\underline{V}_k$ to be the first n columns of the $n(N+1) \times n(N+1)$ matrix $\underline{U}_k$, i.e.

$$\underline{V}_k = \begin{pmatrix} \underline{M}_{k-N}^{k-N} \\ \underline{M}_{k-N}^{k-(N-1)} \\ \vdots \\ \underline{M}_{k-N}^k \end{pmatrix}$$

Because $Q = 0$ it follows from (40) that

$$\underline{B}_k^v = \underline{V}_k \underline{B} \underline{V}_k^T \quad (C1)$$

We suppose inductively that for some $k \geq N$ Method 0 and Methods 1–3 have the same $\underline{x}_k^b, \underline{B}_k$ and that

$$\begin{aligned} \underline{x}_k^b &= \underline{V}_k \underline{x}_{k-N}^b \\ \underline{B}_k &= \underline{B}_k^v = \underline{V}_k \underline{B}_{k-N} \underline{V}_k^T \end{aligned} \quad (C2)$$

This is true for $k = N$ by construction. For all methods we therefore have

$$\underline{K}_k = \underline{V}_k \underline{B}_{k-N} \underline{V}_k^T \underline{H}_k^T (\underline{H}_k \underline{V}_k \underline{B}_{k-N} \underline{V}_k^T \underline{H}_k^T + \underline{R}_k)^{-1}$$

Using the ‘Kalman identity’

$$\begin{aligned} \underline{B} \underline{V}^T \underline{H}^T (\underline{H} \underline{V} \underline{B} \underline{V}^T \underline{H}^T + \underline{R})^{-1} &= \\ (\underline{B}^{-1} + \underline{V}^T \underline{H}^T \underline{R}^{-1} \underline{H} \underline{V})^{-1} \underline{V}^T \underline{H}^T \underline{R}^{-1} & \end{aligned}$$

it follows from (17), (42) that for all methods

$$\underline{A}_k = \underline{V}_k \left[(\underline{B}_{k-N}^{-1} + \underline{V}_k^T \underline{H}_k^T \underline{R}^{-1} \underline{H}_k \underline{V}_k)^{-1} \right] \underline{V}_k^T \quad (C3)$$

$$\underline{K}_k = \underline{A}_k \underline{H}_k^T \underline{R}^{-1} \quad (C4)$$

and therefore from (16), (38) that in all cases

$$\begin{aligned} \underline{x}_k^a &= \underline{V}_k \left[\underline{x}_{k-N}^b + (\underline{B}_{k-N}^{-1} + \underline{V}_k^T \underline{H}_k^T \underline{R}^{-1} \underline{H}_k \underline{V}_k)^{-1} \right. \\ &\quad \times \left. \underline{V}_k^T \underline{H}_k^T \underline{R}^{-1} (\underline{y}_k - \underline{H}_k \underline{x}_k^b) \right] \end{aligned} \quad (C5)$$

Recall that the superscript ℓ in $\underline{M}^{(\ell)}$ in (43) denotes the method used, with $\underline{M}_k^{(1)} = \underline{M}_k^{(0)} = \underline{M}_k$ as defined for the optimal solution in (12). For all $\ell = 0, 1, 2, 3$

$$\underline{M}_k^{(\ell)} \underline{V}_k = \begin{pmatrix} \underline{M}_{k-N}^{k+1-N} \\ \underline{M}_{k-N}^{k+1-(N-1)} \\ \vdots \\ \underline{M}_{k-N}^{k+1} \end{pmatrix} = \underline{V}_{k+1} \underline{M}_{k-N}^{k+1-N}$$

Let $\tilde{\underline{x}}_k^a$ denote the vector in square brackets in (C5) and $\tilde{\underline{A}}_k$ the matrix in square brackets in (C3). It follows from (18), (19) and (43), (45) that both for Method 0 and for Methods $\ell = 1, 2, 3$ we have

$$\begin{aligned} \underline{x}_{k+1}^b &= \underline{M}_k^{(\ell)} \underline{x}_k^a = \underline{M}_k^{(\ell)} \underline{V}_k \tilde{\underline{x}}_k^a = \underline{V}_{k+1} \underline{M}_{k-N}^{k+1-N} \tilde{\underline{x}}_k^a \\ \underline{x}_{k+1-N}^b &= (\underline{x}_{k+1}^b)_\perp = \underline{M}_{k-N}^{k+1-N} \tilde{\underline{x}}_k^a \\ \underline{B}_{k+1} &= \underline{M}_k^{(\ell)} \underline{A}_k (\underline{M}_k^{(\ell)})^T \\ &= \underline{V}_{k+1} \underline{M}_{k-N}^{k+1-N} \tilde{\underline{A}}_k (\underline{M}_{k-N}^{k+1-N})^T \underline{V}_{k+1}^T \\ \underline{B}_{k+1-N} &= (\underline{B}_{k+1})_{\perp, \perp} = \underline{M}_{k-N}^{k+1-N} \tilde{\underline{A}}_k (\underline{M}_{k-N}^{k+1-N})^T \end{aligned} \quad (C6)$$

(where $(\underline{x}_{k+1}^b)_\perp$ is the first $n \times 1$ sub-vector of $\underline{x}_{k+1}^b$). Therefore the inductive hypothesis holds for $k+1$, and therefore for all k .

Appendix D. 4D-Var with a fixed background error covariance: impact of observation batch size

Suppose we have a linear system, observations in some time interval $[0, T]$ and all errors are Gaussian. If we assimilate the observations using an optimal method, such as 4D-Var with correctly cycled prior and posterior error covariances, using m assimilation windows

$$[0, t_1], [t_1, t_2], \dots, [t_{m-1}, T]$$

then the estimate of state at time T is independent of m and of how we choose

$$t_1 < t_2 < \dots < t_{m-1}$$

However, if instead of properly cycling the error covariances the background error covariances are fixed, it is often advantageous to assimilate data in larger batches.

We illustrate this by considering a case where x is a scalar quantity, which evolves in time according to

$$x(i+1) = \mu x(i)$$

for some constant μ (which is supposed known, so there is no model error) with $|\mu| > 1$. At each time i we wish to assimilate an observation $y_1(i)$ with error variance r_i and an observation $y_2(i)$ with error variance s_i . To make the problem analytically tractable we will suppose that $\{(r_i, s_i), i = 1, \dots, \infty\}$ are drawn

Table D1. Parameters for Equation (D1).

Parameter	Simultaneous	Sequential
β_i	$\frac{1/b}{1/b + 1/r_i + 1/s_i}$	$\frac{r_i s_i}{(b_1 + r_i)(b_2 + s_i)}$
ρ_i	$\frac{1/r_i}{1/b + 1/r_i + 1/s_i}$	$\frac{s_i b_1}{(b_1 + r_i)(b_2 + s_i)}$
σ_i	$\frac{1/s_i}{1/b + 1/r_i + 1/s_i}$	$\frac{b_2}{b_2 + s_i}$

(independently of i) from $\{(R_j, S_j), j = 1 \dots, k\}$, where for any i $\{r_i = R_j, s_i = S_j\}$ with probability $p_j, j = 1, \dots, k$.

We compare two assimilation strategies using 4D-Var with a non-cycled background:

Simultaneous (batch size of 2): at each time i we assimilate $y_1(i)$ and $y_2(i)$ simultaneously using fixed background error covariance b , i.e. $x_a = x_b + \delta$ where δ minimises

$$J(\delta) = \frac{1}{2} \frac{\delta^2}{b} + \frac{1}{2} \frac{(y_1 - x_b - \delta)^2}{r_i} + \frac{1}{2} \frac{(y_2 - x_b - \delta)^2}{s_i}$$

Sequential (batch size of 1): at each time i we assimilate $y_1(i)$ using fixed background error covariance b_1 to give intermediate analysis $x_{a1}(i)$, then assimilate observation $y_2(i)$ using fixed background error variance b_2 to give final analysis $x_a(i)$, i.e. $x_{a1} = x_b + \delta_1, x_a = x_{a1} + \delta_2$, where δ_1, δ_2 minimise

$$J(\delta_1) = \frac{1}{2} \frac{\delta_1^2}{b_1} + \frac{1}{2} \frac{(y_1 - x_b - \delta_1)^2}{r_i}$$

$$J(\delta_2) = \frac{1}{2} \frac{\delta_2^2}{b_2} + \frac{1}{2} \frac{(y_2 - x_{a1} - \delta_2)^2}{s_i}$$

We will show the following: we can choose b so that, however b_1 and b_2 are chosen, the mean square error using the simultaneous method is lower than that using the sequential method (and strictly lower if $k \geq 2$ and $R_1 \neq R_2$).

D.1. Proof (summary)

- (1) Denoting $E[(x_i^a - x_i^f)^2]$ by A_i it is readily shown that for both the simultaneous and sequential methods

$$A_i = \beta_i^2 \mu^2 A_{i-1} + \rho_i^2 r_i + \sigma_i^2 s_i \quad (\text{D1})$$

where the parameters $\beta_i, \rho_i, \sigma_i$ are as listed in Table D1.

- (2) It follows from (D1) that for any $i > 1$

$$A_i = +(\beta_i \beta_{i-1} \dots \beta_2 \mu^{i-1})^2 A_1$$

$$+ (\beta_i \beta_{i-1} \dots \beta_3 \mu^{i-2})^2 (\rho_2^2 r_2 + \sigma_2^2 s_2)$$

$$+ (\beta_i \beta_{i-1} \dots \beta_4 \mu^{i-3})^2 (\rho_3^2 r_3 + \sigma_3^2 s_3)$$

$$+ \dots$$

$$+ (\beta_i \mu)^2 (\rho_{i-1}^2 r_{i-1} + \sigma_{i-1}^2 s_{i-1})$$

$$+ \rho_i^2 r_i + \sigma_i^2 s_i \quad (\text{D2})$$

So A_i is a function of $(r_1, s_1), (r_2, s_2), \dots, (r_i, s_i)$ which we are considering as independent random variables, so

$$E \left[(\beta_i \beta_{i-1} \dots \beta_{j+1})^2 (\rho_j^2 r_j + \sigma_j^2 s_j) \right]$$

$$= E \left[\beta_i^2 \right] E \left[\beta_{i-1}^2 \right] \dots E \left[\beta_{j+1}^2 \right] E \left[\rho_j^2 r_j + \sigma_j^2 s_j \right]$$

Denote expectations over $(r_1, s_1), (r_2, s_2), \dots$ by $E_{\{r_j, s_j\}}$. If $|E_{\{r_j, s_j\}}[(\beta_j \mu)^2]| < 1$, and noting $\beta_i = 1 - \sigma_i - \rho_i$, we have

$$\lim_{i \rightarrow \infty} E_{\{r_j, s_j\}}[A_i] = \frac{E_{\{r_j, s_j\}}[\rho_j^2 r_j + \sigma_j^2 s_j]}{1 - E_{\{r_j, s_j\}}[(\beta_j \mu)^2]}$$

$$= \frac{\sum_{j=1}^k p_j [\rho(R_j, S_j)^2 R_j + \sigma(R_j, S_j)^2 S_j]}{1 - \mu^2 \sum_{j=1}^k p_j (1 - \rho(R_j, S_j) - \sigma(R_j, S_j))^2} \quad (\text{D3})$$

where in the simultaneous case ρ, σ are the functions

$$\rho_{sim}(R, S, b) = \frac{1/R}{1/b + 1/R + 1/S},$$

$$\sigma_{sim}(R, S, b) = \frac{1/S}{1/b + 1/R + 1/S} \quad (\text{D4})$$

and in the sequential case ρ, σ are the functions

$$\rho_{seq}(R, S, b_1, b_2) = \frac{S b_1}{(b_1 + R)(b_2 + S)},$$

$$\sigma_{seq}(R, S, b_1, b_2) = \frac{b_2}{b_2 + S} \quad (\text{D5})$$

- (3) Lemma. If

$$R_1, R_2, \dots, R_k > 0$$

$$S_1, S_2, \dots, S_k > 0$$

$$p_1, p_2, \dots, p_k > 0 \quad (\text{D6})$$

with $\sum p_i = 1$ and $\mu > 1$, then

$$\min_{b>0} \left\{ \frac{\sum p_j [\rho_{sim}(R_j, S_j, b)^2 R_j + \sigma_{sim}(R_j, S_j, b)^2 S_j]}{1 - \mu^2 \sum p_j (1 - \rho_{sim}(R_j, S_j, b) - \sigma_{sim}(R_j, S_j, b))^2} \right\}$$

subject to

$$\mu^2 \sum p_j (1 - \rho_{sim}(R_j, S_j, b) - \sigma_{sim}(R_j, S_j, b))^2 < 1 \quad (D7)$$

is less than or equal to

$$\min_{\substack{b_1>0 \\ b_2>0}} \left\{ \frac{\sum p_j [\rho_{seq}(R_j, S_j, b_1, b_2)^2 R_j + \sigma_{seq}(R_j, S_j, b_1, b_2)^2 S_j]}{1 - \mu^2 \sum p_j (1 - \rho_{seq}(R_j, S_j, b_1, b_2) - \sigma_{seq}(R_j, S_j, b_1, b_2))^2} \right\}$$

subject to

$$\mu^2 \sum p_j (1 - \rho_{seq}(R_j, S_j, b_1, b_2) - \sigma_{seq}(R_j, S_j, b_1, b_2))^2 < 1 \quad (D8)$$

where the inequality is strict if $k \geq 2$ and $R_1 \neq R_2$.

This Lemma proves the claim, and is itself proved in (3a)–(3d).

(3a) Given constants

$$\{R_j > 0, S_j > 0, p_j > 0; j = 1, \dots, k\}$$

$$\mu > 0 \quad (D9)$$

Define the function

$$\Gamma(\rho_1, \sigma_1, \dots, \rho_k, \sigma_k) = \frac{\sum_{j=1}^k p_j [\rho_j^2 R_j + \sigma_j^2 S_j]}{1 - \mu^2 \sum_{j=1}^k p_j (1 - \rho_j - \sigma_j)^2} \quad (D10)$$

on

$$\mathcal{D} = \{\rho_1, \sigma_1, \dots, \rho_k, \sigma_k : \rho_j > 0, \sigma_j > 0, \mu^2 \sum p_j (1 - \rho_j - \sigma_j)^2 < 1\}$$

We may readily check that $\mathcal{D}$ is convex and that Γ is convex on $\mathcal{D}$. At any local extremum of Γ the $2k$ conditions hold

$$\left(\frac{\partial}{\partial \rho_i}, \frac{\partial}{\partial \sigma_i} \right) \Gamma = 0, \quad i = 1, \dots, k \quad (D11)$$

(3b) Consider the 1-parameter family of functions

$$\rho_1(b), \sigma_1(b), \dots, \rho_k(b), \sigma_k(b)$$

given by

$$\rho_i(b) = \rho_{sim}(R_i, S_i, b) = \frac{1/R_i}{1/b + 1/R_i + 1/S_i}$$

$$\sigma_i(b) = \sigma_{sim}(R_i, S_i, b) = \frac{1/S_i}{1/b + 1/R_i + 1/S_i} \quad (D12)$$

Define

$$g(b) = 1 - \mu^2 \sum_{j=1}^k p_j (1 - \rho_j(b) - \sigma_j(b))^2 \quad (D13)$$

$$h(b) = \sum_{j=1}^k p_j [\rho_j(b)^2 R_j + \sigma_j(b)^2 S_j] \quad (D14)$$

Then we may check that if $\{\rho_j(b), \sigma_j(b)\}$ are of the form (D12), if $g(b) \neq 0$ and b satisfies the single condition

$$f(b) \equiv bg(b) - \mu^2 h(b) = 0 \quad (D15)$$

then the $2k$ conditions (D11) hold.

(3c) We may also check that we have

$$f'(b) = g(b) \text{ for all } b$$

$$g'(b) > 0 \text{ if } b > 0 \quad (D16)$$

Noting that as $b \rightarrow 0$ we have

$$f(b) \rightarrow 0$$

$$f'(b) = g(b) \rightarrow 1 - \mu^2 \quad (D17)$$

it follows by elementary arguments that so long as $|\mu| > 1$ there must exist $b_{**} > b_* > 0$ such that

$$f'(b_*) = g(b_*) = 0$$

$$f(b_{**}) = 0 \quad (D18)$$

and since for all $b > 0$ we have $g'(b) > 0$ we must have

$$f'(b) = g(b) > 0, \text{ for all } b > b_*$$

Furthermore, since $g(b) > 0$ if and only if

$$\{\rho_j(b), \sigma_j(b)\}_{j=1, \dots, k} \in \mathcal{D}$$

it follows that

$$\{\rho_j(b_{**}), \sigma_j(b_{**})\}_{j=1,\dots,k} \in \mathcal{D}$$

(3d) We have found a value b_{**} such that

$$\{\rho_{sim}(R_j, S_j, b_{**}), \sigma_{sim}(R_j, S_j, b_{**})\}, j = 1, \dots, k \quad (\text{D19})$$

is in $\mathcal{D}$ and such that at this point $\nabla \Gamma = \mathbf{0}$. Because Γ is convex in the convex set $\mathcal{D}$ any point where $\nabla \Gamma = \mathbf{0}$ is the global minimum of Γ on $\mathcal{D}$ (Boyd and Vandenberghe (2004)). Therefore the global minimum of Γ over $\mathcal{D}$ is achieved by (D19).

It remains to show that this minimum is not achievable by ρ_{seq}, σ_{seq} , for any $b_1 > 0, b_2 > 0$, if $R_1 \neq R_2$.

At any extremum of Γ Equation (D11) must hold, and in particular for every $i = 1, \dots, k$

$$\frac{\partial \Gamma}{\partial \rho_i} - \frac{\partial \Gamma}{\partial \sigma_i} = \frac{2p_i(\rho_i R_i - \sigma_i S_i)}{1 - \mu^2 \sum_{j=1}^k p_j(1 - \rho_j - \sigma_j)^2} = 0$$

The denominator is positive for all points in $\mathcal{D}$, therefore at any extremum in $\mathcal{D}$

$$\rho_j R_j = \sigma_j S_j, \quad j = 1, \dots, k$$

Inserting this requirement with $j = 1, 2$ into the expressions for ρ_{seq}, σ_{seq} in (D5) it follows we would need

$$R_1 b_1^2 = R_2 b_1^2 \quad (\text{D20})$$

which for $b_1 > 0$ is only possible if $R_1 = R_2$.