



Analog data assimilation for the selection of suitable general circulation models

Juan Ruiz, Pierre Ailliot, Thi Tuyet Trang Chau, Pierre Le Bras, Valérie Monbet, Florian Sévellec, Pierre Tandeo

► To cite this version:

Juan Ruiz, Pierre Ailliot, Thi Tuyet Trang Chau, Pierre Le Bras, Valérie Monbet, et al.. Analog data assimilation for the selection of suitable general circulation models. *Geoscientific Model Development*, 2022, 15, pp.7203-7220. 10.5194/gmd-15-7203-2022 . insu-03868833v2

HAL Id: insu-03868833

<https://insu.hal.science/insu-03868833v2>

Submitted on 24 Nov 2022

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



Distributed under a Creative Commons Attribution 4.0 International License



Analog data assimilation for the selection of suitable general circulation models

Juan Ruiz¹, Pierre Ailliot², Thi Tuyet Trang Chau³, Pierre Le Bras^{4,5}, Valérie Monbet⁶, Florian Sévellec^{4,7}, and Pierre Tandeo⁵

¹Centro de Investigaciones del Mar y la Atmósfera, Facultad de Ciencias Exactas y Naturales, Universidad de Buenos Aires, CONICET-UBA, Buenos Aires, Argentina, UMI-IFAECI (CNRS-CONICET-UBA), Buenos Aires, Argentina

²Univ Brest, CNRS UMR 6205, Laboratoire de Mathématiques de Bretagne Atlantique, France

³LSCE, IPSL-CEA Saclay, 91191 Gif-sur-Yvette cedex, France

⁴Laboratoire d'Océanographie Physique et Spatiale, IUEM, Univ. Brest, CNRS, IRD, Ifremer, Brest, France

⁵IMT Atlantique, Lab-STICC, UMR CNRS 6285, 29238, France

⁶INRIA & Univ Rennes, CNRS, IRMAR – UMR 6625, 35000 Rennes, France

⁷Ocean and Earth Science, University of Southampton, Southampton, UK

Correspondence: Juan Ruiz (jruiz@cima.fcen.uba.ar)

Received: 24 December 2021 – Discussion started: 10 February 2022

Revised: 16 July 2022 – Accepted: 18 July 2022 – Published: 26 September 2022

Abstract. Data assimilation is a relevant framework to merge a dynamical model with noisy observations. When various models are in competition, the question is to find the model that best matches the observations. This matching can be measured by using the model evidence, defined by the likelihood of the observations given the model. This study explores the performance of model selection based on model evidence computed using data-driven data assimilation, where dynamical models are emulated using machine learning methods. In this work, the methodology is tested with the three-variable Lorenz model and with an intermediate complexity atmospheric general circulation model (a.k.a. the SPEEDY model). Numerical experiments show that the data-driven implementation of the model selection algorithm performs as well as the one that uses the dynamical model. The technique is able to select the best model among a set of possible models and also to characterize the spatiotemporal variability of the model sensitivity. Moreover, the technique is able to detect differences among models in terms of local dynamics in both time and space which are not reflected in the first two moments of the climatological probability distribution. This suggests the implementation of this technique using available long-term observations and model simulations.

1 Introduction

Data assimilation (DA) methods aim to provide the best estimation of the state of a dynamical system based on a set of noisy and partial observations (see Carrassi et al., 2018; Reich, 2019; Van Leeuwen et al., 2019, and references therein). Current state of the art DA systems are based on robust mathematical grounds, allowing their use to be expanded beyond their original aim. One application that has recently received increasing attention is the use of DA methods for model optimization and model selection. The former is concerned with obtaining better estimates for model parameters and configuration with the ultimate goal of quantifying and reducing model error and dispersion in various applications (e.g., Schirber et al., 2013; Ruiz et al., 2013; Ruiz and Pulido, 2015; Lauvaux et al., 2019; Kotsuki et al., 2020). Model selection aims at identifying the model which best describes a set of observations in a finite set of possible models. In meteorology and oceanography, model selection may be useful for example to find the best physics for forecast applications or for detecting the forcing terms that better explain the evolution of a dynamical system for attribution purposes.

So far, model selection in the context of DA, seems to have received less attention than model optimization. In Hanhart et al. (2016), the authors showed that the fraction of at-

tributable risk in the context of climate change (Pearl, 2000) can be estimated as a by-product of two DA systems. One of these systems is run with a model which includes forcing consistent with anthropogenic emissions, and another is run without considering those emissions. This approach was proven to be more sensitive to differences between the two scenarios at a lower computational cost than other available attribution techniques. One usual approach to conduct model selection is based on the computation of model evidence, which is defined as the log-likelihood of the available observations for a given model configuration (Reich and Cotter, 2015). Carson et al. (2018) proposed using model evidence to select a model consistent with data records in paleoclimate science. This work succeeds in both fitting conceptual models and identifying the one with the most appropriate orbital forcing to represent the glacial-interglacial cycle. Estimating model evidence, however, is generally complex for practical applications as geophysical models are usually high-dimensional and nonlinear. In such circumstances, it is crucial to develop and to implement DA methods which allow accurate estimations of model evidence.

Carrassi et al. (2017) introduced the concept of contextual model evidence (CME), which can be roughly defined as the log-likelihood of a set of observations over a short time period for a given dynamical model. Based on this approach, the model selection can be obtained by running several DA systems (i.e., each one using a particular dynamical model) and then comparing their corresponding CME values. Metref et al. (2019) extended this idea to high-dimensional systems and studied the impact of domain localization upon the computation of the CME in the context of an ensemble Kalman filter (EnKF). A similar idea has been implemented by Otsuka and Miyoshi (2015) for the online optimization of a multimodel EnKF. They ran a particle filter that assigns weights to each model configuration based on the likelihood of the observations given different model configurations. The approach successfully identifies the most accurate model improving the performance of both the assimilation and the forecast.

The articles cited above perform model selection using classical DA methods where the different competing models that represent the dynamics of a particular system must be solved several times at each time step. Recently, Tandeo et al. (2015) and Lguensat et al. (2017) introduced the concept of analog DA (AnDA). This approach can be particularly beneficial in cases where the numerical model is not explicitly known or computationally extremely expensive (as is usually the case with state of the art numerical models of the climate system) or when such models are not available and only the observation dataset exists. In this case, AnDA can take advantage of existing long-term climate model simulations or observations and perform DA by emulating the dynamical model using nearest neighbor regression, also called analog forecasting in meteorology (Lorenz, 1969).

In this paper, we combine the AnDA method with the computation of the CME. The objective is to provide a proof of concept using numerical experiments on the low-dimensional modified Lorenz-63 toy model (Lorenz, 1963) and the intermediate complexity atmospheric general circulation model SPEEDY (Molteni, 2003). This proof of concept suggests the possible use of existing model simulations to efficiently compare different models and physics based on observations.

The paper is organized as follows. Section 2 provides a brief review of model evidence and introduces an algorithm to compute CME using AnDA. Section 3 presents the numerical results and Sect. 4 presents final remarks and perspectives for future research.

2 Methodology

Model evidence measures the ability of a dynamical model $\mathcal{M}$ to describe a sequence of multivariate, noisy, and partial observations $\mathbf{y}_{0:K} = (\mathbf{y}_0, \dots, \mathbf{y}_K)$ (from a sufficiently long time in the past and up to time K). It is a useful tool to identify the model best fitting a set of observations in a list of competing models. In this section, after defining model evidence, we discuss its computation by combining DA and analog forecasting.

2.1 Contextual model evidence

2.1.1 Definition

Model selection or comparison is usually performed using the climatological model evidence (ME), see Metref et al. (2019) and references therein. It corresponds to $\ln p(\mathbf{y}_{0:K} | \mathcal{M})$, the global log-likelihood of the observations $\mathbf{y}_{0:K}$ for the dynamical model $\mathcal{M}$. This ME metric roughly measures the adequacy between the observations available up to time K and the climatological distribution of the model, although this global metric is probably not descriptive enough for studying the model performance over a particular time period or the transition between different states of the system.

Alternatively, Carrassi et al. (2017) proposed to compute the contextual model evidence (CME) defined as the local log-likelihood on a short interval of time. More precisely, it is defined as

$$\text{CME}_{k:k+h}(\mathcal{M}) = \ln p(\mathbf{y}_{k:k+h} | \mathbf{y}_{0:k-1}; \mathcal{M}),$$

where $p(\mathbf{y}_{k:k+h} | \mathbf{y}_{0:k-1}; \mathcal{M})$ denotes the conditional (forecast) distribution of the observations $\mathbf{y}_{k:k+h}$ between times k and $k+h$ given the observations up to time $k-1$ for model $\mathcal{M}$ and h is the width of the evidencing window. As stated in Metref et al. (2019), a key difference between CME and ME is that the former takes into account the actual state of the system. This information is considered in the a priori estimation of the state of the system at the beginning of the

evidencing window, which is assumed to be approximately known. Therefore, CME computes the evidence taking the context into account, which can provide a more detailed local evaluation of the system dynamics.

Note that as shown by Carrassi et al. (2017), for a Markovian system and for independent observations

$$\begin{aligned} \text{CME}_{k:k+h}(\mathcal{M}) &= \sum_{i=k}^{k+h} \ln p(\mathbf{y}_i | \mathbf{y}_{0:i-1}; \mathcal{M}) \\ &= \sum_{i=k}^{k+h} \text{CME}_i(\mathcal{M}), \end{aligned} \quad (1)$$

where $\text{CME}_i(\mathcal{M}) = \ln p(\mathbf{y}_i | \mathbf{y}_{0:i-1}; \mathcal{M})$ is the CME in the particular case where the evidencing window is reduced to a single time. This last quantity corresponds to the logarithmic score (also known as the ignorance score, see e.g., Siegert et al., 2019 and references therein) of the forecast distribution $p(\mathbf{y}_i | \mathbf{y}_{0:i-1}; \mathcal{M})$ associated to model $\mathcal{M}$. Also when k is set to 0 in the CME (i.e., there is no observation prior to the evidencing window), we obtain the ME.

The CME cannot be evaluated directly because the observations $\mathbf{y}$ are often incomplete, intermittent and uncertain. To tackle these issues, a latent variable which represents the true state of the system is introduced leading to the following state-space model

$$\mathbf{x}_k = \mathcal{M}(\mathbf{x}_{k-1}) + \boldsymbol{\eta}_k, \quad (2)$$

$$\mathbf{y}_k = \mathcal{H}(\mathbf{x}_k) + \boldsymbol{\epsilon}_k, \quad (3)$$

where $\mathbf{x}_k$ denotes the latent state and $\boldsymbol{\eta}_k$ represents the model noise (i.e., the part of the system dynamics which is not represented by the numerical model) at time k . In Eq. (3), $\mathcal{H}$ is the observation operator, representing the link between the latent state $\mathbf{x}$ (i.e., what we want to estimate) and the observations $\mathbf{y}$. The additive term $\boldsymbol{\epsilon}_k$ represents observation errors.

For a state-space model $\text{CME}_k(\mathcal{M})$ can be expressed as follows:

$$\text{CME}_k(\mathcal{M}) = \ln \int p(\mathbf{y}_k | \mathbf{x}_k) p(\mathbf{x}_k | \mathbf{y}_{0:k-1}; \mathcal{M}) d\mathbf{x}_k. \quad (4)$$

The forecast distribution $p(\mathbf{x}_k | \mathbf{y}_{0:k-1}; \mathcal{M})$, which appears in the previous expression can again be decomposed into two terms

$$\begin{aligned} p(\mathbf{x}_k | \mathbf{y}_{0:k-1}; \mathcal{M}) &= \int p(\mathbf{x}_k | \mathbf{x}_{k-1}; \mathcal{M}) \\ &\quad p(\mathbf{x}_{k-1} | \mathbf{y}_{0:k-1}; \mathcal{M}) d\mathbf{x}_{k-1}, \end{aligned} \quad (5)$$

where the analysis distribution $p(\mathbf{x}_{k-1} | \mathbf{y}_{0:k-1}; \mathcal{M})$ represents the estimation of the state of the system at time $k-1$ given all the previous available observations (i.e., the context which is usually provided by a DA system).

2.1.2 Contextual model evidence and the ensemble Kalman filter

In an ensemble Kalman filter (EnKF), the forecast distribution $p(\mathbf{x}_k | \mathbf{y}_{0:k-1}; \mathcal{M})$ is approximated using a Monte Carlo approach performing multiple evaluations of the model $\mathcal{M}$, initialized from a sample of states drawn from the analysis distribution $p(\mathbf{x}_{k-1} | \mathbf{y}_{0:k-1}; \mathcal{M})$. More precisely, a sample of the forecast distribution is generated as follows

$$\mathbf{x}_{(j),k}^f = \mathcal{M}(\mathbf{x}_{(j),k-1}^a) + \boldsymbol{\eta}_{(j),k}, \quad (6)$$

where $\{\mathbf{x}_{(j),k-1}^a\}_{j=1:N}$ is a sample of N members from the analysis distribution at time $k-1$ and $\boldsymbol{\eta}_{(j),k}$ is a realization of a stochastic process representing model imperfections typically drawn from a Gaussian distribution with zero mean and covariance $\mathbf{Q}$ (see Tandeo et al., 2020). The different members $\{\mathbf{x}_{(j),k}^f\}_{j=1:N}$ of the forecast are used to approximate the first two moments of the forecast distribution through the sample mean

$$\bar{\mathbf{x}}_k^f = \frac{1}{N} \sum_{j=1}^N \mathbf{x}_{(j),k}^f, \quad (7)$$

and covariance

$$\boldsymbol{\Sigma}_k^f = \frac{1}{N-1} \sum_{j=1}^N (\mathbf{x}_{(j),k}^f - \bar{\mathbf{x}}_k^f)(\mathbf{x}_{(j),k}^f - \bar{\mathbf{x}}_k^f)^T. \quad (8)$$

When observations are available at time k , the state distribution can be updated based on the information provided by the observations. This update is performed based on the Bayes' theorem assuming that both the observation likelihood and the forecast distributions are Gaussian and that their moments are well approximated by the sample moments. The update can be conducted in different ways, one possible approach is the so-called stochastic EnKF update (Burgers et al., 1998) in which each member of the analysis sample is obtained as

$$\mathbf{x}_{(j),k}^a = \mathbf{x}_{(j),k}^f + \mathbf{K}_k (\mathbf{y}_k - \mathcal{H}(\mathbf{x}_{(j),k}^f) + \boldsymbol{\epsilon}_{(j),k}), \quad (9)$$

with $\mathbf{K}_k$ the Kalman filter gain defined as

$$\mathbf{K}_k = \boldsymbol{\Sigma}_k^f \mathbf{H}^T (\mathbf{H} \boldsymbol{\Sigma}_k^f \mathbf{H}^T + \mathbf{R})^{-1}, \quad (10)$$

where $\mathbf{H}$ denotes the tangent linear of the observation operator $\mathcal{H}$.

Equations (6)–(10) can be used sequentially to produce an estimate of the state conditioned on previous observations ($\mathbf{y}_{0:k-1}$) at each time k . Given an EnKF sequential cycle, based again on Gaussian assumptions $\text{CME}_k(\mathcal{M})$ can be approximated as (Carrassi et al., 2017):

$$\begin{aligned} \text{CME}_k(\mathcal{M}) &\approx -\frac{1}{2} (\mathbf{y}_k - \mathcal{H}(\bar{\mathbf{x}}_k^f)) (\mathbf{H} \boldsymbol{\Sigma}_k^f \mathbf{H}^T + \mathbf{R})^{-1} \\ &\quad (\mathbf{y}_k - \mathcal{H}(\bar{\mathbf{x}}_k^f))^T - \frac{1}{2} \ln |\mathbf{H} \boldsymbol{\Sigma}_k^f \mathbf{H}^T + \mathbf{R}| \\ &\quad - \frac{n}{2} \ln(2\pi), \end{aligned} \quad (11)$$

where n is the number of available observations at time k . Combining Eqs. (1) and (11) (also with Eqs. 6–10) we can compute $\text{CME}_{k:k+h}$ from an EnKF-based DA system for an arbitrary time window $[k, k+h]$. When both forecast mean and covariance are well estimated, Carrassi et al. (2017) showed that this CME approach is able to detect the model candidates which best describe a given set of observations.

The EnKF usually suffer from systematic underestimation of the forecast error variance. To compensate for this issue an ad hoc multiplicative inflation coefficient γ is usually applied such that

$$\Sigma_k^f \leftarrow \gamma \Sigma_k^f. \quad (12)$$

The inflation coefficient γ can be estimated online using different techniques, which are based on the work of Desroziers et al. (2005). In particular Li et al. (2009) proposed the following estimation approach, in which at each time k the inflation value is updated according to

$$\gamma_{k+1} = \rho \left(\frac{\text{Tr}(\mathbf{d}_{\text{o-f}}^T \mathbf{d}_{\text{o-f}} \circ \mathbf{R}^{-1}) - p}{\text{Tr}(\mathbf{H} \Sigma_k^f \mathbf{H}^T \circ \mathbf{R}^{-1})} \right) + (1 - \rho) \gamma_k, \quad (13)$$

with $0 \leq \rho \leq 1$ a time smoothing parameter, p the number of observations, $\mathbf{d}_{\text{o-f}}$ the difference between the forecast and the observations in the observation space and, $\circ$ the element-wise product.

2.2 Data-driven contextual model evidence

2.2.1 The analog EnKF

In a DA system the numerical model is required to propagate the information in time. Particularly, in the case of ensemble-based assimilation techniques, the numerical model has to be run several times (once for each of the ensemble members and each time step). An alternative to this computationally intensive approach is to use analog forecasting. This method, initially introduced by Lorenz (1969) in meteorology, is a data-driven procedure which uses a catalog of historical data, most of the time corresponding to simulation runs or analysis (i.e., outputs of DA procedures). The idea is to search in the catalog, using an appropriate distance, the nearest analogs of the initial condition from which we want to get a forecast. Then, the successors of these closest analogs are combined to get a probabilistic forecast. Analog forecasting is popular because of its simplicity and robustness. Several studies confirmed these advantages in the context of environmental sciences (Barnett and Preisendorfer, 1978; Bannayan and Hoogenboom, 2008; Yiou, 2014; Atencia and Zawadzki, 2015; Ayet and Tandeo, 2018; Sévellec and Drijfhout, 2018).

The combination of DA and analog forecasting has been proposed in Tandeo et al. (2015) and Lguensat et al. (2017) leading to the analog data assimilation (hereinafter AnDA) method. It is a flexible framework that can be adapted to a large set of problems. An interesting feature of AnDA is

that it can be applied locally without the need to approximate the full model. On the contrary, only a part of the state space (e.g., a particular region or physical variable) can be used to select the analogs and to emulate the dynamics of that particular part of the system. This is particularly advantageous when dealing with high-dimensional systems. Also, as for most DA systems AnDA can handle complex observations with irregular spatiotemporal distributions as long as an appropriate observation operator is available. More precisely, AnDA corresponds to running a DA algorithm using the state-space model (Eqs. 2–3) with (2) replaced by

$$\mathbf{x}_k = \widehat{\mathcal{M}}(\mathbf{x}_{k-1}) + \boldsymbol{\eta}_k, \quad (14)$$

where $\widehat{\mathcal{M}}$ denoting the analog-based approximation of the dynamical model $\mathcal{M}$.

An efficient statistical forecast operator used in Lguensat et al. (2017) and detailed in Platzer et al. (2021) is the local linear regression originally introduced in Cleveland and Devlin (1988). It consists in first searching in the catalog for the M -nearest neighbors (i.e., the analogs) of a given state $\mathbf{x}$ along with their corresponding successors in time. Then, a multiple linear regression is fitted between the M analogs and their successors. The coefficients of this regression are denoted $\beta(\mathbf{x})$ and $\alpha(\mathbf{x})$, where we stress that these coefficients depend on the state of the system. Note that this regression is able to emulate the model dynamics including nonlinearities as it is a linear approximation applied locally in state space and corresponds to a first-order expansion of the dynamical model (see Platzer et al., 2021).

The combination of analog forecasting, based on local linear regressions, and the EnKF leads to the AnEnKF algorithm where Eq. (6) is replaced by

$$\mathbf{x}_{(j),k}^f = \underbrace{\alpha_{(j)}(\mathbf{x}_{(j),k-1}^a) \mathbf{x}_{(j),k-1}^a + \beta_{(j)}(\mathbf{x}_{(j),k-1}^a)}_{\widehat{\mathcal{M}}_{(j)}(\mathbf{x}_{(j),k-1}^a)} + \boldsymbol{\eta}_{(j),k}, \quad (15)$$

where the model error $\boldsymbol{\eta}_{(j),k}$ is drawn from a multivariate Gaussian distribution $\mathcal{N}(0, \tilde{\Sigma}_{(j)}(\mathbf{x}_{(j),k-1}^a))$, where $\tilde{\Sigma}_{(j)}(\mathbf{x}_{(j),k-1}^a)$ denotes the sample covariance of the residuals of the multiple linear regression between the analogs and the successors (see Lguensat et al., 2017, their Sect. 3a for further details). Once the forecast is performed using the analog technique, the analysis update can be done as in the EnKF using Eq. (9). Multiplicative inflation (Eq. 12) could also be applied and estimated online using Eq. (13). The AnEnKF has been tested on toy dynamical models in Lguensat et al. (2017). Numerical results show that the performance of the classical EnKF (i.e., using the true model $\mathcal{M}$) and AnEnKF (i.e., using analog forecasts $\widehat{\mathcal{M}}$), is almost the same when the size of the catalog is large enough; however, efficiently applying analog forecasting for chaotic dynamics with strong nonlinearities and high-dimensional spaces is not always straightforward. As mentioned in Zhen et al. (2020) the analog space must be large enough to capture the dynamics of

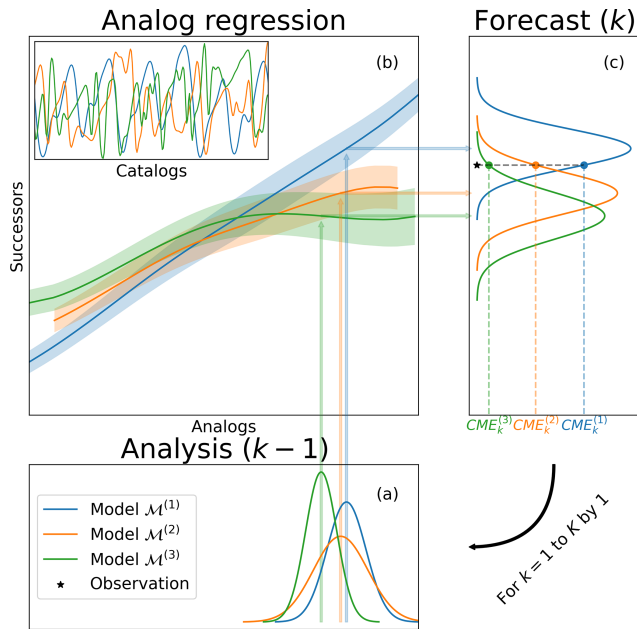


Figure 1. Schematic representation of the proposed methodology. The procedure is iterative, from an initial to a final time index. At time index $k-1$, the procedure starts with the results of different DA systems in (a), corresponding to different dynamical models. In (b), each analysis state is used to find the nearest analogs. Those analogs and corresponding successors, coming from catalogs of numerical simulations, are used to build analog regressions. The resulting probabilistic forecasts given in (c) are compared to the available observations at time index k . Then, the likelihoods (CMEs) are used to compute a weight for each dynamical model.

the system, but small enough to avoid the curse of dimensionality. This can be achieved using time delay embedding and projection in an appropriate subspace. In this work, this issue is avoided by running AnDA locally on a reduced set of the state-space variables (see Sect. 3.2).

2.2.2 Contextual model evidence and the analog EnKF

It is straightforward to combine the CME and AnEnKF procedures in the same DA scheme. The procedure is summarized in Fig. 1 and detailed in algorithm 1.

In the next section, several models $\{\mathcal{M}^{(i)}\}_{i=1:L}$ are in competition and CME is used to identify those which best describe the dynamics of the observations. Note that the time series $\text{CME}_k(\mathcal{M}^{(i)})$ are computed independently for each model. The combination of CME and AnEnKF has several advantages in this context. Firstly, it is a fast procedure because it avoids running the different models at each time step and for each ensemble member. Secondly, uncertainty in the local linear regression is considered as a robust estimation of model errors and contributes to increase the forecast ensemble spread (Lguensat et al., 2017; Platzer et al., 2021). Thirdly, it allows the CME to be computed in specific re-

Algorithm 1 Contextual model evidence using EnKF and local linear regression forecasting.

- Initialization: sample the first ensemble, $\{\mathbf{x}_{(j),0}^a\}_{j=1:N} \sim p(\mathbf{x}_0)$.
- For $k = 1 : K$,
 - + **Forecast step:**
 - For each member $j = 1, \dots, N$ propagate the previous analysis member using the analog forecasting operator Eq. (15).
 - Compute the empirical mean forecast $\bar{\mathbf{x}}_k^f$ (7) and covariance forecast Σ_k^f (8) from the ensemble members.
 - Apply the multiplicative inflation factor Eq. (12).
 - Compute the contextual model evidence Eq. (11).
 - + **Analysis step:**
 - Update the state distribution Eqs. (9)–(10).
 - Update the inflation factor Eq. (13).

gions of the state space (e.g., in specific areas or integrated variables), only where observations are available. This last point is discussed in Sect. 3.2 using the SPEEDY model.

3 Results

This section presents numerical results which show that the proposed methodology is able to identify the model which is the most compatible with a given set of observations. A first set of experiments is performed using a modified version of the Lorenz-63 model. Then we focus on the more challenging intermediate complexity atmospheric general circulation model SPEEDY.

3.1 Modified Lorenz-63 model

In this section the method is tested on a modified Lorenz-63 model $\mathcal{M}^{(\lambda)}$ originally introduced in Palmer (1999). It is defined by the following system of differential equations:

$$\begin{aligned} \frac{dx_1}{dt} &= 10(x_2 - x_1) + \lambda \cos\left(\frac{7}{9}\pi\right), \\ \frac{dx_2}{dt} &= x_1(28 - x_3) - x_2 + \lambda \sin\left(\frac{7}{9}\pi\right), \\ \frac{dx_3}{dt} &= x_1x_2 - \frac{8}{3}x_3, \end{aligned} \quad (16)$$

where λ controls the magnitude of the external forcing term which is added onto the first two components. The particular case $\lambda = 0$ corresponds to the classical Lorenz-63 model (Lorenz, 1963). Figure 2 shows two trajectories simulated from Eq. (16) with $\lambda = 0$ and $\lambda = 8$. When $\lambda > 0$, the additional forcing term “pushes” the trajectories towards the right wing and the left wing is less often visited than the right one, the opposite holding true when $\lambda < 0$. This behavior can be clearly seen on the right hand plots in Fig. 2, which show

the number of times the system visits different regions of the state space.

A set of observations is generated using the classical Lorenz-63 model $\mathcal{M}^{(0)}$, referred to as the correct model hereafter. We conduct several experiments to check if the proposed methodology is able to identify that the observations were actually generated using the correct model $\mathcal{M}^{(0)}$ and not from another model in the list of competing models which corresponds to the modified Lorenz models $\mathcal{M}^{(\lambda)}$ with $\lambda = -8, \dots, 8$. In practice, the Lorenz model is integrated using a Runge-Kutta 4 scheme with a time step $dt = 0.1$, and the 3 components are observed and assimilated at every model time step. The observation operator is $\mathcal{H} = \mathbf{I}_3$ (the 3×3 identity matrix) and an additive Gaussian error with covariance $\mathbf{R} = 2\mathbf{I}_3$ is used. Hereafter each experiment is based on $K = 10^4$ DA cycles. Unless stated otherwise, the AnEnKF is run using catalogs of size $T = 10^4$ and the evidencing window size is $h = 1$. For the EnKF, experiments with $\mathbf{Q} = \mathbf{0}$ are used independently of the value of λ , while in the AnEnKF the approach described in the previous section is used, which in practice corresponds to a robust estimate of the model noise.

Figure 3 shows the mean CME for different values of λ , for different configurations and for the EnKF and the AnEnKF. Comparing the left panels of Fig. 3 shows that the EnKF is very sensitive to the ensemble member size N , especially when the strength of the forcing in the modified Lorenz-63 model (i.e., the value $|\lambda|$) increases. Applying a multiplicative inflation factor (see last row of Fig. 3) permits this issue to be solved. Hereafter, we discuss the results obtained when both EnKF and AnEnKF algorithms are run with $N = 20$ ensemble members and a multiplicative inflation factor is applied. This allows us to obtain robust results for the EnKF at a reasonable computational cost. Note that the AnEnKF is less sensitive to the ensemble size and that adding inflation does not seem necessary here (see right panel of Fig. 3). This may be explained by the adaptive data-driven procedure which is used to estimate the forecast error in the AnEnKF procedure. As expected, the CME has a larger mean value when the correct model $\mathcal{M}^{(0)}$ is used in the DA procedures and decreases when the value of $|\lambda|$ increases. Both EnKF and AnEnKF give comparable results when multiplicative inflation (Eq. 13) is applied (see bottom panels of Fig. 3).

As an example, Fig. 4 (top panel) shows the analysis for the second component x_2 obtained when running the AnEnKF with a catalog of the correct model $\mathcal{M}^{(0)}$. It suggests that the AnEnKF is generally able to reconstruct the true state of the system. This can be assessed more precisely by computing the RMSE between the true state $\mathbf{x}$ and the analyzed state $\mathbf{x}^a$ (0.50) and the mean coverage probability of the 95 % prediction interval (84.26 %). Then the AnEnKF was run with a catalog of the incorrect model $\mathcal{M}^{(8)}$. As expected, the RMSE is larger (0.72) and the mean coverage probability is also degraded (63.86 %), and thus less accurate estimates of the mean and the variance of the true state are

obtained. Note that these two quantities cannot be computed in practical applications since the true state is not known. The bottom panel of Fig. 4 shows that the CME obtained with the AnEnKF run with a catalog of the incorrect model are generally smaller than the ones obtained with a catalog of the correct model. It illustrates again that the CME computed with AnEnKF can be used to identify the correct model using only a sequence of noisy observations and catalogs of competing models.

This is confirmed by the results given in Table 1. The CME associated with the correct model is larger than the one associated with the incorrect model for 68 % of the assimilation cycles. Note that this percentage of correct identification increases with the catalog size T and seems to converge to the percentage of correct identification obtained with the EnKF (68 %) when T becomes large. This is not surprising since using a larger catalog provides a better approximation of the dynamical model and thus similar results when using EnKF and AnEnKF (see the discussion in Lguensat et al., 2017). Table 1 also shows that CME is more precise in identifying the correct model compared to the root mean squared error of the forecast in the observation space defined as $\text{RMSE}_k^f = \|\mathbf{y}_k - \mathcal{H}(\bar{\mathbf{x}}_k^f)\|$. Note that according to Eq. (13), CME depends not only on the forecast error but also on the variance of the forecast error and this may help to identify the correct model. Also, to highlight the advantage of using DA to identify the correct model we perform an experiment using only the pure analog forecasting (AnF), without DA. For each time k , we select analogs based only on the available noisy observations and we propagate the state using the local linear regression approach. The RMSE of each forecast at time $k + 1$, denoted $\text{RMSE}_k^{\text{AnF}}$, is reported in Table 1. The maximum selection probability obtained with the AnF is 57 %, which is smaller compared to the experiments using AnDA (64 % using RMSE of the forecast error and 68 % using CME). This stresses the importance of having an accurate initial conditions in order to be able to compare the forecasts obtained with two competing models.

3.2 SPEEDY model

In this section, we discuss the implementation of CME with AnDA for the Simplified Parameterizations, primitiveEquation DYnamics (SPEEDY; Molteni, 2003) which is an intermediate complexity, atmospheric general circulation model. The grid in SPEEDY consists of 48 points in the south-north direction and 96 points in the west-east direction and of 7 vertical σ levels. SPEEDY has a set of simplified parameterizations to represent unresolved scale processes including radiation, large scale condensation, soil-sea-ice-atmosphere energy fluxes, boundary layer, and moist convection. A brief description of these schemes can be found in the appendix of Molteni (2003).

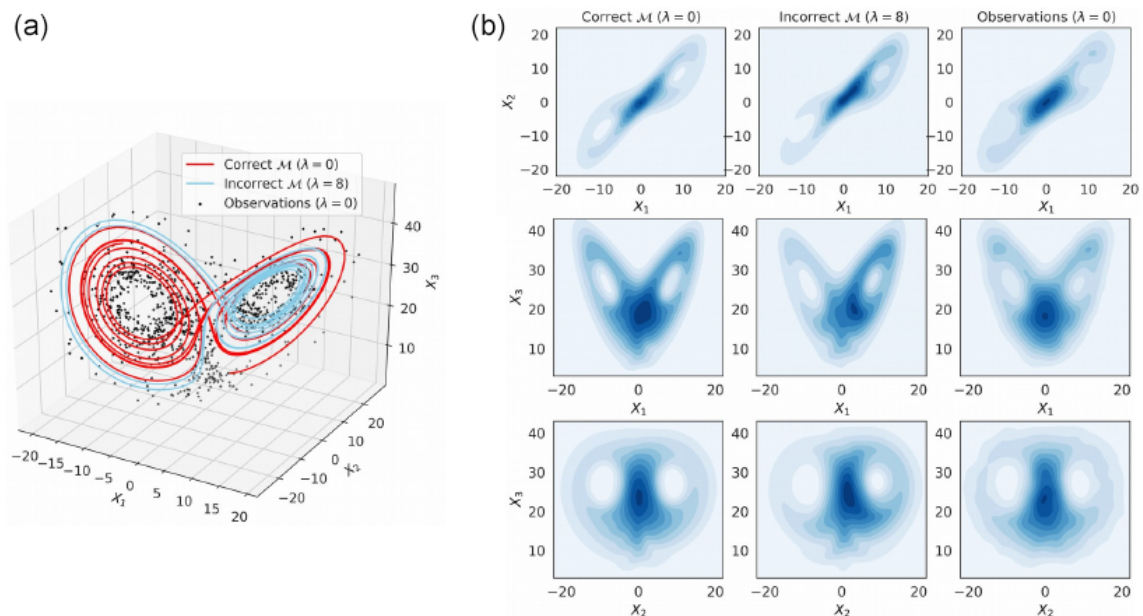


Figure 2. Simulated trajectories (a) and bivariate distributions (b) from the modified Lorenz-63 model Eq. (16) with $\lambda = 0$ and $\lambda = 8$. The observations (dots) are generated with the correct model with $\lambda = 0$ and an additive Gaussian noise with mean 0 and variance $2I_3$. Figure obtained using a time step $dt = 0.01$.

Table 1. Sensitivity of the percentage of correct model identification based on the CME and on the RMSE (values within parentheses) with respect to the length of the catalog (T) used in the AnEnKF and in the AnF. The mean and 95 % confidence interval (CI) of the percentage of correct model identification are computed for each of the algorithms using 10 repetitions of the assimilation experiment.

Approaches		$T = 10^2$	$T = 10^3$	$T = 10^4$
EnKF	Mean (%)	68.64 (64.37)		
	CI (%)	[67.59, 69.70] ([63.59, 65.25])		
AnEnKF	Mean (%)	59.85 (58.99)	67.02 (63.33)	68.44 (64.37)
	CI (%)	[58.74, 60.96] ([58.00, 59.99])	[65.43, 68.60] ([61.92, 64.73])	[67.09, 69.79] ([63.59, 65.25])
AnF	Mean (%)	– (53.38)	– (57.46)	– (57.71)
	CI (%)	[–, –] ([52.11, 54.60])	[–, –] ([56.61, 58.30])	[–, –] ([56.79, 58.63])

3.2.1 Data-driven model selection with SPEEDY

In this work, we conduct observation simulation experiments using this model. A 30-year run is performed using the default configuration of the model which has been shown to produce a good representation of the main features of the current climate (Molteni, 2003). The model is integrated with a time step of 40 min and model states are archived at 6-h intervals. This simulation is hereafter referred to as the TRUE simulation.

To simulate imperfections in the model formulation, we modify the value of the RH_{cnv} parameter, which is related to the deep convection parameterization in SPEEDY (Molteni, 2003). This parameter controls the activation of the convective parameterization and the intensity of the convective overturning. Lower values of the parameter lead to more frequent

convection activation and stronger vertical mass fluxes. The TRUE simulation uses $RH_{cnv} = 0.9$. Two additional simulations are performed using $RH_{cnv} = 0.8$ and $RH_{cnv} = 0.7$ which are referred to as RH08 and RH07, respectively. Figure 5 shows that reducing RH_{cnv} leads to an increase in mid-level mean temperature within in mid-latitudes and in the tropics (with some exceptions on the western Pacific and northern Africa). This increase is related to stronger latent heat release within the intertropical convergence zone (ITCZ, Fig. 5e and f) and to enhanced subsidence in the Polar side of the Hadley circulation. Precipitation produced by the convective scheme is enhanced over tropical regions which can also contribute to the mid-level warming in this region. RH07 and RH08 produce 12 % and 6 % more convective precipitation than the TRUE experiment, respectively. However, the sensitivity in total precipitation is much smaller due to a decrease

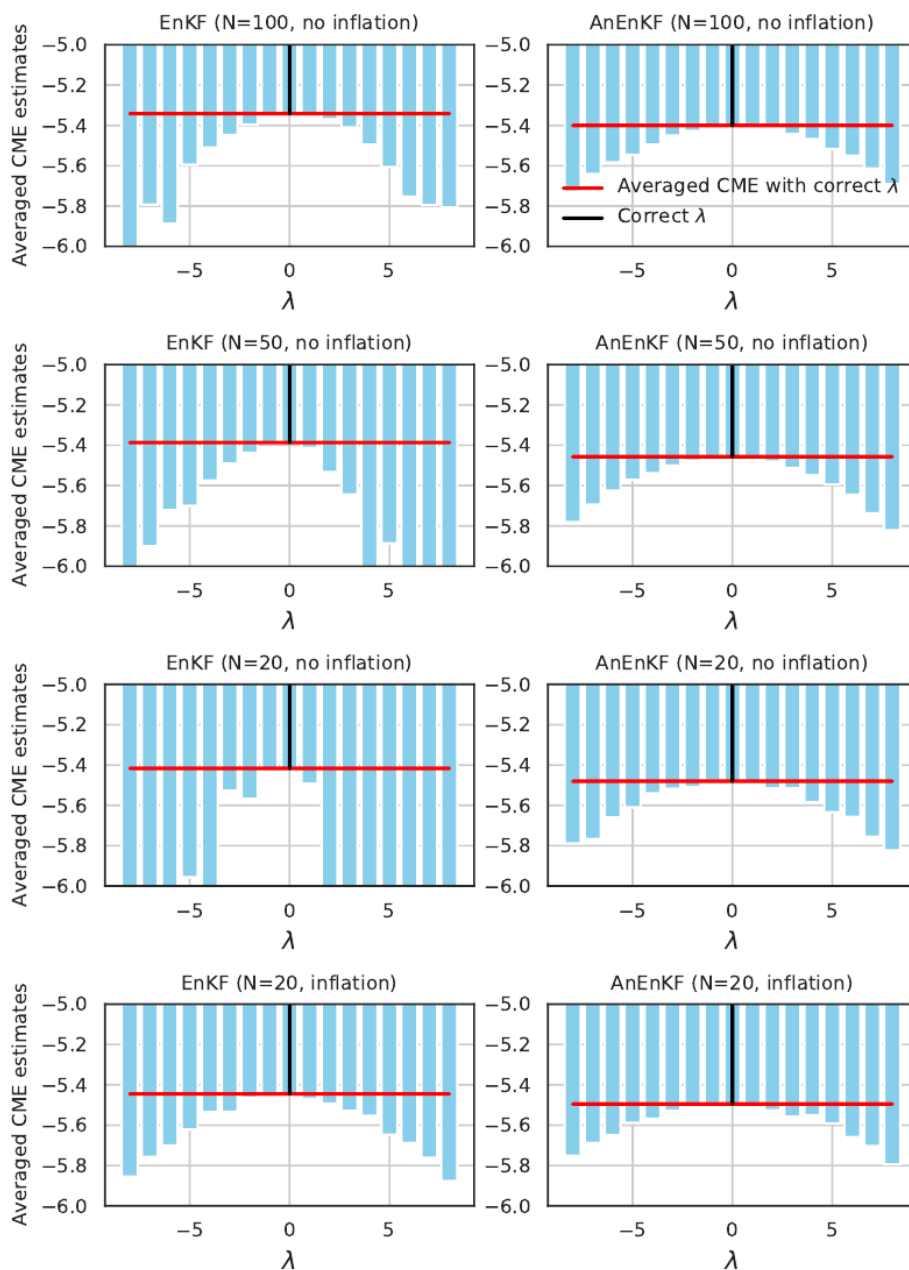


Figure 3. Time-averaged CME as a function of the parameter λ used to generate the incorrect model for the EnKF (left) and AnEnKF (right) approaches using different number of members N (rows). Adaptive inflation is used only for the results in the bottom row. The red line indicates the time-averaged CME for the correct model ($\lambda = 0$).

in the large scale precipitation as the convective precipitation increases. The impact of this parameter upon the mean distribution of temperature is somehow linear, since both RH08 and RH07 produce anomalies with a similar spatial pattern but a larger amplitude for the latter.

Temperature observations are generated from the TRUE run at each model grid point adding uncorrelated Gaussian random errors with a standard deviation of 0.7 K. The assumed error standard deviation is similar to the assumed error of several real temperature observations provided by

radiosondes and satellite retrievals. Temperature is selected since its horizontal distribution and its vertical gradient are directly affected by convective processes, particularly in the tropics, as is shown in Fig. 5.

The AnDA approach in SPEEDY is implemented over local domains centered at a given model grid point. In this work and for each grid point in the horizontal domain, analogs are defined based on the value of the temperature on a 3-dimensional box of 3×3 horizontal grid points and 3 vertical σ levels ($\sigma = 0.77, 0.6$, and 0.42). It is important to note here

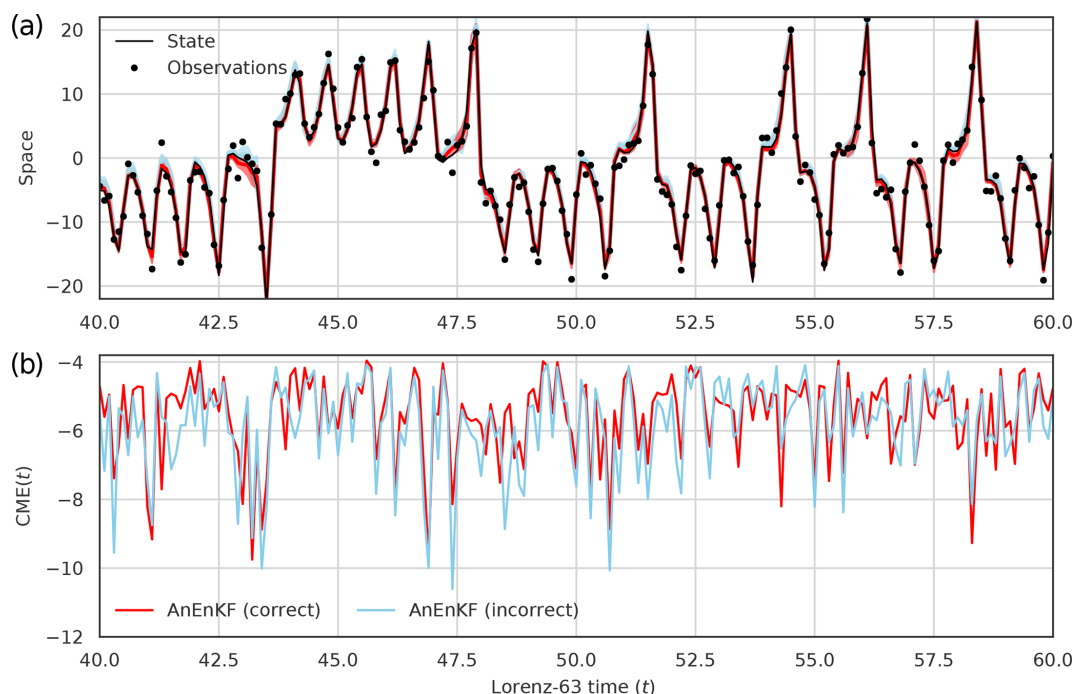


Figure 4. (a) Time series of the second component of the Lorenz-63 model with mean analysis and 95 % prediction interval obtained using AnEnKF with the correct $\mathcal{M}^{(0)}$ model (red) or using AnEnKF with the incorrect $\mathcal{M}^{(8)}$ model (blue). (b) CME time series for the correct $\mathcal{M}^{(0)}$ model and the incorrect $\mathcal{M}^{(8)}$ model. $N = 20$ members are used and an inflation factor is applied in AnEnKF.

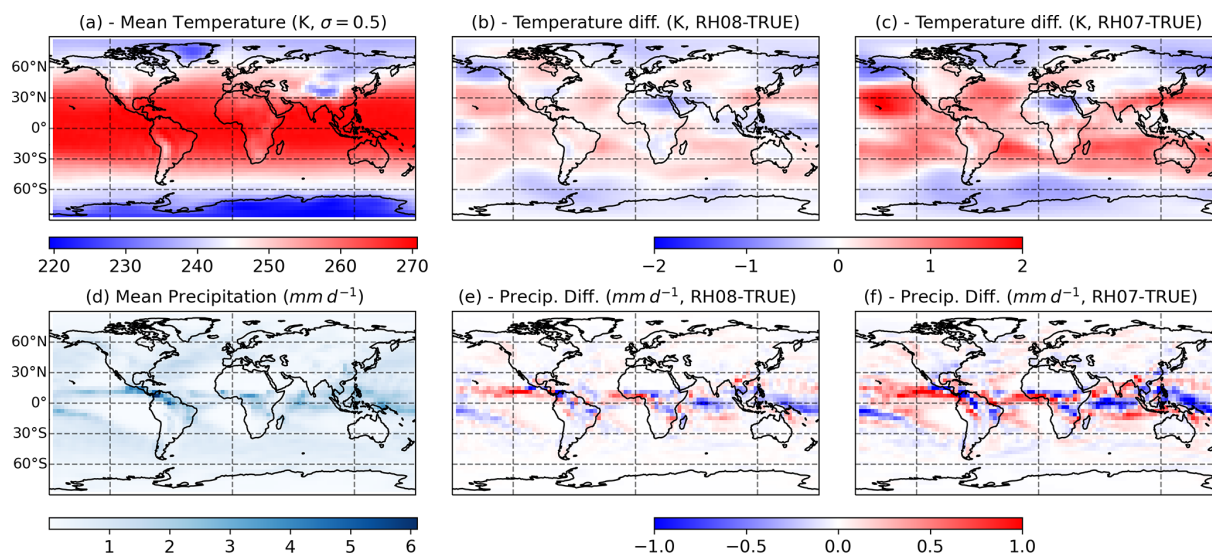


Figure 5. The 30-year mean for the TRUE model simulation (a, d) and differences between the TRUE model simulation and the RH08 and RH07 mode simulation (b, e, c, f, respectively), for temperature at the vertical level $\sigma = 0.5$ (a, b, c, K) and total precipitation (d, e, f, mm d^{-1}).

that AnDA is implemented at each local domain ($3 \times 3 \times 3$ grid points boxes) independently from other local domains, meaning that in fact AnDA can be applied at only one local domain or over all local domains over a particular region without the need to compute it over the entire global domain. Also, at each local domain, only the observations within the

local domain are assimilated (i.e., in this case 27 temperature observations per analysis cycle). By implementing AnDA locally we significantly increase the probability of finding relevant analogs reducing the size of the catalog required for an accurate approximation of the dynamical model. The local implementation allows us to avoid the global integration

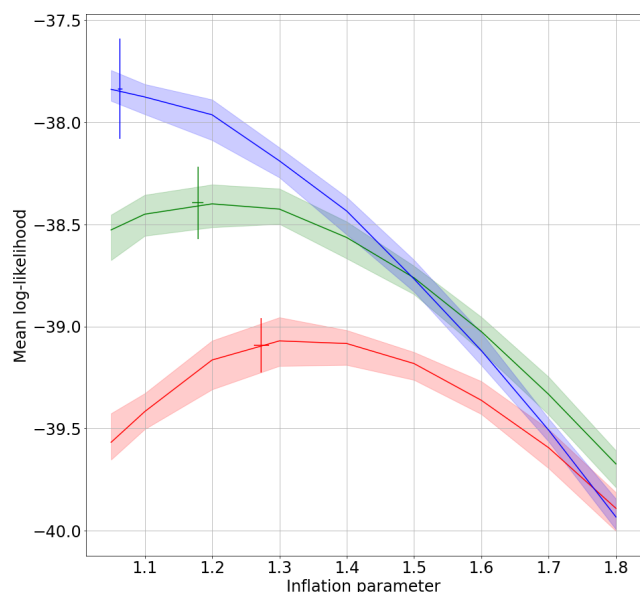


Figure 6. Mean log likelihood for grid points located over northern South America as a function of the inflation coefficient. The shade shows the range between the maximum and minimum values over 10 repetitions of the experiment. The vertical error bars represents the range between the maximum and minimum values of the mean log-likelihood for the experiments with adaptive multiplicative inflation over 10 repetitions of the experiment. The horizontal error bars indicate the range between the maximum and minimum values of the estimated inflation over 10 repetitions of the experiment. Results are shown for the TRUE (blue line), RH08 (green line), and RH07 (red line) model experiments.

of the model resulting in a substantial reduction of the computational cost and providing unprecedented flexibility to the computation of DA-based metrics.

This local implementation significantly differs from the localized CME implementation presented in Metref et al. (2019). In that paper the authors introduced a local computation of CME based on the localization of a global DA system. However, in their approach the DA is performed globally and using observations of several physical variables such as temperature and wind, while in the present work AnDA is implemented locally using only temperature observations within the local domain.

As in the Lorenz-63 experiments, the adaptive multiplicative inflation method indicated in algorithm 1 is used to find the optimal inflation value corresponding to each experiment. Figure 6 shows that CME is quite sensitive to the multiplicative inflation used in AnDA. As expected, the optimal inflation value (the one which produces the maximum CME) is larger for the catalogs corresponding to the imperfect models. Moreover, if a large inflation is assumed (e.g., larger than 1.5), CME associated with RH08 becomes larger than the one obtained with the perfect model catalog. Also, the adaptive inflation produces results which are close to the ones ob-

tained with the optimal fixed multiplicative inflation value thus avoiding the need to manually tune the inflation parameter. This is important considering the results of Miyoshi (2011), which show that for atmospheric general circulation models, the optimal inflation parameter depends on time and location.

In ensemble-based DA methods, the estimation of the forecast error covariance from a limited size ensemble usually leads to sampling noise. This is usually ameliorated by the use of localization schemes that reduce the amplitude of the covariance between distant variables. In this paper AnDA is implemented without localization given the relatively small size of the local domain in which DA is performed.

The AnDA experiments are conducted assimilating the observations generated from the last 3 years of the TRUE simulation. The catalogs for the analog forecasting are constructed from the first 25 years of the RH08, RH07, and TRUE model runs and 250 analogs are used for the forecast. In the SPEEDY experiments, the catalog contains over 36 000 samples (which is almost 4 times the size of the largest catalog which we tried with the Lorenz model). Although the local state space dimension that we used in SPEEDY is much larger (27 grid points), we argue that since there are substantial correlations among the state variables, the effective dimension can be significantly smaller.

The number of ensemble members is 30. To increase the evidence associated with the local dynamics of the models the assimilation frequency is set to 24 h. To take advantage of 6 h data, at each local domain, we perform four DA experiments which are run independently from each other starting at 00:00, 06:00, 12:00 and 18:00 UTC on the first day. These 4 DA cycles are performed over the same 3-year period. These configuration settings have been chosen based on preliminary experiments performed over a limited number of local domains in which the sensitivity of the results to these parameters has been explored. The analyses obtained from these experiments are merged to obtain a total of 4380 analysis cycles over the 3-year assimilation period (4 DA experiments $\times$ 1095 cycles each).

It is important to note that the generation of the catalog brings a significant computational cost in this approach since it requires running the global numerical model once over a long period of time. However, we argue that for the implementation of this technique in real data applications, available long model simulations like those produced by the Coupled Model Intercomparison Project (Eyring et al., 2016) can be used. Moreover, the lengths of these catalogs are of the same order of magnitude as the ones used in the idealized experiments with the SPEEDY model.

Preliminary experiments performed to optimize the configuration show that results are particularly sensitive to the assimilation frequency and to the size of the local domain used to identify analogs. In particular, using 3 different vertical levels and less frequent assimilation, results in a much stronger sensitivity to RH_{cnv} (i.e., larger difference in CME

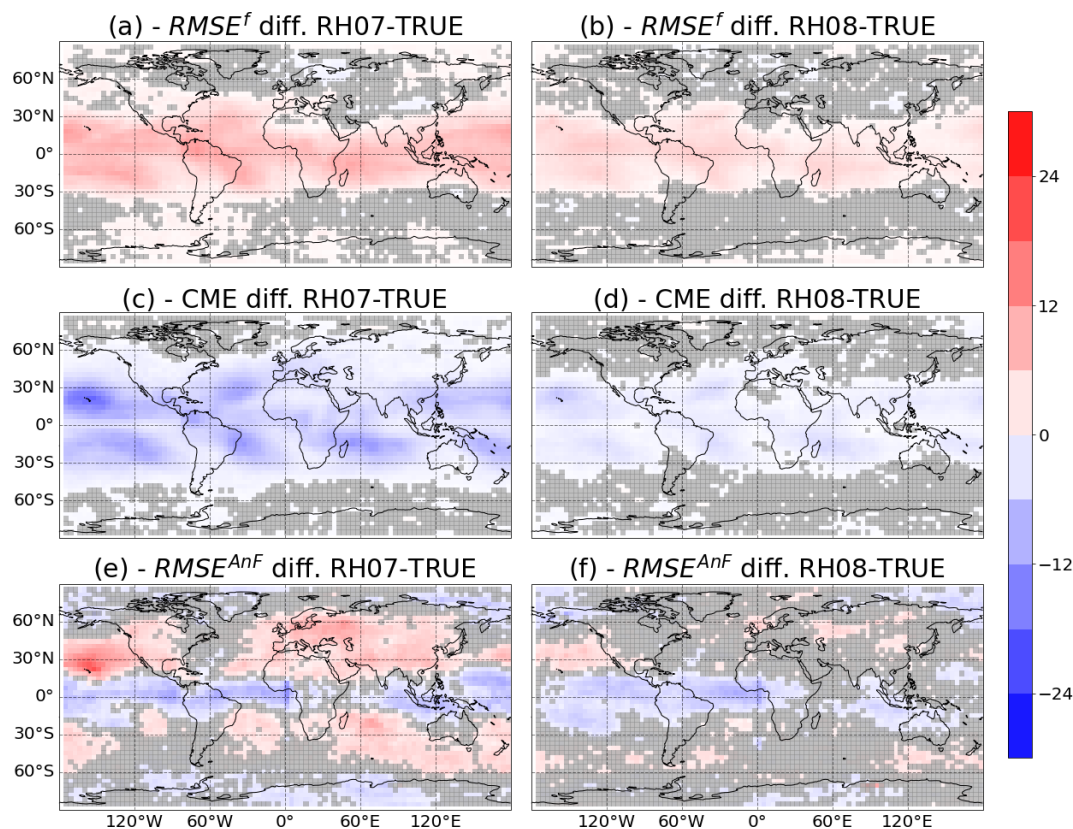


Figure 7. Percentage difference in $RMSE^f$ (a, b) and CME (c, d) for the AnDA experiments performed using the RH07 and RH08 catalogs (a, c, e and b, d, f, respectively) versus the experiment performed with the TRUE catalog. The images (e) and (f) correspond to the difference in $RMSE^{AnF}$ from the experiment based only on analog forecasting (AnF). In all panels, the gray shade indicates values which are below the 95 % statistical significance level.

associated with different catalogs). The number of analogs affects the performance of AnDA but has a lower impact on the relative performance of the different models.

First, it was checked that the proposed methodology is able to identify that the temperature observations were generated using the TRUE simulation and not the RH08 or RH07 runs. To evaluate the statistical robustness of the results, all the experiments were repeated 10 times using different realizations for the observation noise. The standard deviation of the values obtained from different experiments is used to estimate the 95 % confidence interval for the different metrics discussed in this section.

As described by Metref et al. (2019), model comparison in the context of DA can be conducted using different metrics. We compared the performance of the different models using RMSE of the forecast error in the observation space ($RMSE^f$) and CME. Figure 7a and b show the percentage difference in $RMSE^f$ and CME for the experiments with RH08 and RH07 with respect to the one using the TRUE catalog. The percent-

age difference is computed as

$$\%Diff_{RMSE} = (RMSE^f - RMSE^f_{TRUE}) / RMSE^f_{TRUE}, \quad (17)$$

$$\%Diff_{CME} = (CME - CME_{TRUE}) / CME_{TRUE}, \quad (18)$$

where $RMSE^f$ is the RMSE obtained with either the RH07 or RH08 catalogs and $RMSE^f_{TRUE}$ is the RMSE obtained with the TRUE catalog and the same naming convention is applied to the CME.

The performances of the experiments using catalogs RH08 and RH07 are consistently worse than the one of the experiments using the TRUE catalog (positive $RMSE^f$ percentage differences). Larger differences are found in the tropics where convection is more frequent and stronger. This suggests that AnDA provides a valuable hint about the source of model imperfection.

Figure 7c and d show the percentage difference for CME. In the case of CME, negative percentage differences indicate a worse performance with respect to the experiment that uses the TRUE catalog. The CME has a similar spatial pattern as the difference in RMSE. However, for the CME, the area where results are statistically significant is larger (see for example Asia, Europe and Hawaii for the experiment

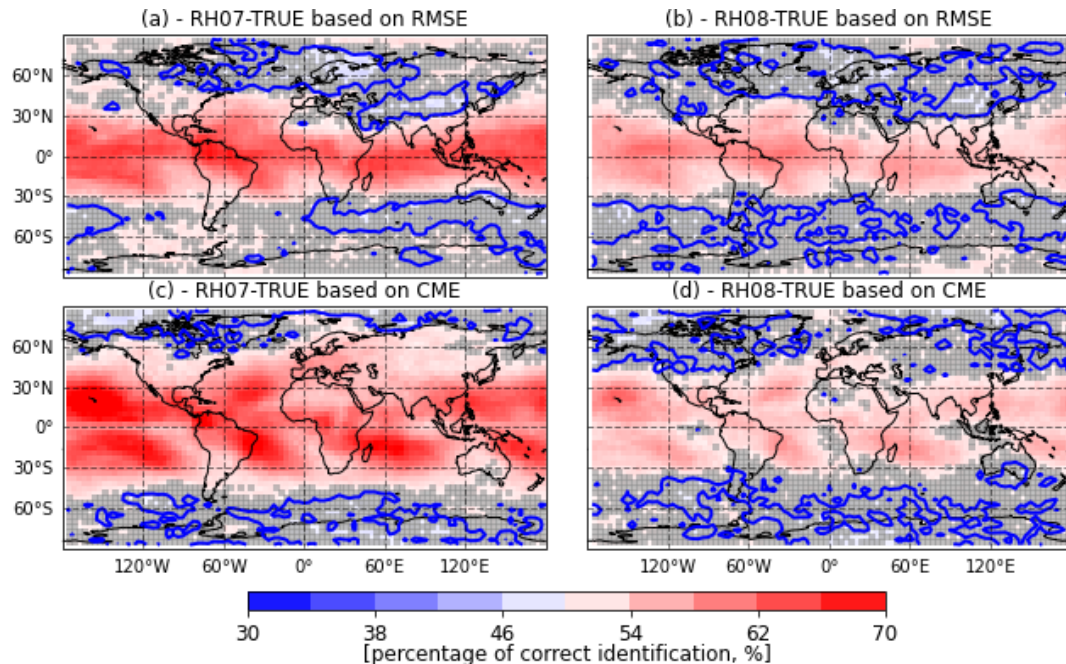


Figure 8. Percentage of correct identification between the RH07 and TRUE catalog (a, c) and between the RH08 and TRUE catalogs (b, d) using AnDA. The percentage of correct identification is computed using either RMSE^f or CME (a and b, c and d respectively). The gray shade indicate values which are below the 95 % statistical significance level. The blue line corresponds to a 50 % of correct identifications.

RH07). Figure 8 shows that the spatial pattern of the percentage of correct identification computed from RMSE or CME is similar. However, the CME generally identifies the right model more frequently and the number of grid points at which the percentage of correct identifications is significantly over 50 % is larger. This is mainly because CME incorporates the uncertainty in the quantification of the fit of the forecast to the observations. In these experiments, an accurate estimation of this uncertainty is achieved by combining the ensemble of analog forecasts and the adaptive multiplicative inflation. A mismatch between the forecast and the observations will contribute less to the CME when the forecast uncertainty is correctly specified. We achieved this by the implementation of AnDA with adaptive multiplicative inflation. This takes into account the effect of stochastic errors in the initial conditions and its amplification due to the chaotic nature of the system which is not explicitly considered in the RMSE metric. Although CME usually performs better than the RMSE at identifying the correct model, this is not always the case (see for example in Fig. 8 how the probability of correct identification is larger for the RMSE than for CME near the Equator). This result may be due to an overestimation of the forecast error covariance Σ^f , computed within the analog procedure. Indeed, as explained in Eq. (11), an augmentation of this error matrix implies a diminution of the CME, and thus a decrease of performance of this metric.

To quantify the advantage of using DA in identifying the catalog that best fits the observations, we compare AnDA

results with an experiment using the AnF approach (i.e., in which the analog forecast is initialized directly from the noisy observations and evaluated using the RMSE^{AnF}). During the same period corresponding to the DA experiment, 24 h analog forecasts are initialized every 6 h using the same 3 catalogs as in the AnDA experiments. The RMSE^{AnF} is used to compare the performance of the different catalogs. Figure 7e and f show that differences in performance are mixed when no DA is used. Larger RMSEs are associated with the catalogs generated with imperfect models at mid-latitudes. However, over the tropics the opposite is observed with the imperfect catalogs producing better forecasts. Moreover, the area in which differences between the imperfect model catalogs and the TRUE is larger in the experiments that use AnDA. These results highlight again the importance of using DA to quantify the performance of a dynamical model based on noisy observations. When DA is not used, the forecast associated to a model may be corrupted by a large initialization error and the resulting RMSE^{AnF} be sensitive not only to the model error but also to its sensitivity to initial conditions.

3.2.2 Standardized catalog experiments

To evaluate if the CME framework can detect differences in the system dynamics that go beyond the change in the probability distribution of the state variables (i.e., their climatological mean and standard deviation) we evaluate the ability of AnDA to detect differences in the model performance

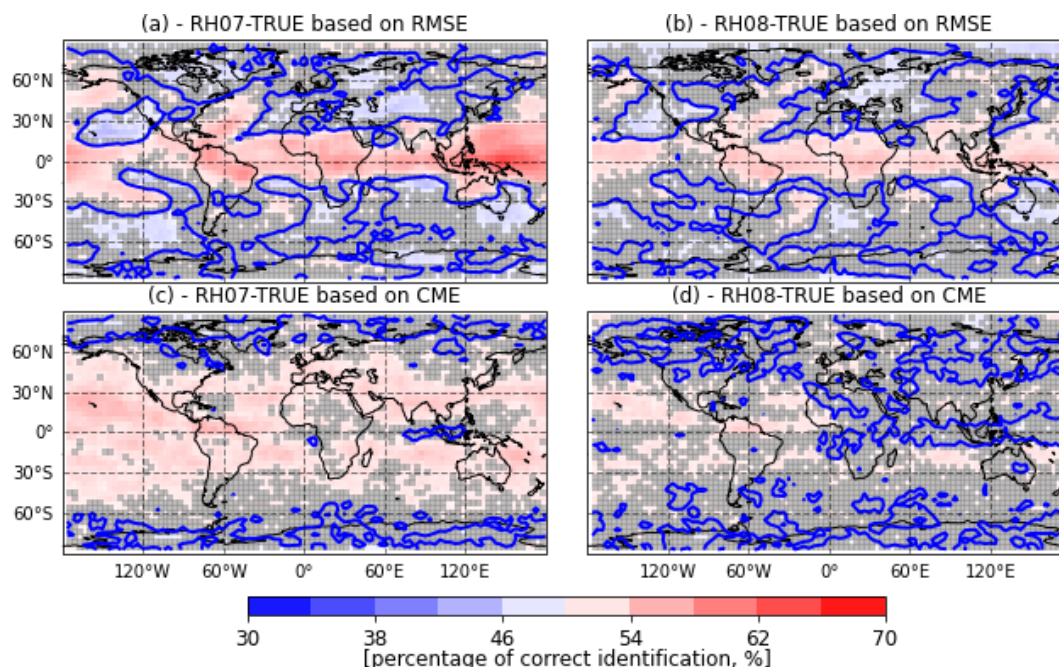


Figure 9. As in Fig. 8 but for the experiments in which the catalogs and observations have been standardized with respect of their respective climatological variability (see the text for details).

when the differences in the mean and in the standard deviation among the catalogs is removed. To achieve this we perform a standardization of the state variables and the true state prior to the generation of the observations. The seasonal cycle is taken into account in the standardization using a 60-d centered moving window to compute the climatological mean and standard deviation corresponding to different days of the year. In the AnDA experiments performed using standardized catalogs and observations, the observation error has been scaled accordingly. Results of the percentage of correct identification for both RMSE^f and CME are shown in Fig. 9. Differences in RMSE^f and CME among different catalogs are lower than in previous experiments. This indicates that the difference in the mean and standard deviation among the catalogs contributes to CME. However, when only standardized data are used, AnDA is still able to identify the correct model with the higher probabilities obtained by using the CME. RMSE^f shows correct model selection probabilities over the tropical regions where the signal is stronger; however, there are some areas over mid-latitudes where the percentage of correct identification is consistently below 50 % resulting in mixed results for this metric. For some of the areas where the RMSE provides the wrong answer, the CME is still able to provide percentages over 50 % (see for example the northern and southern Pacific).

Figure 10a and b show the percentage of correct identification based on CME using an evidencing window of $h = 7$ d. This is higher than the one obtained using a single day ($h = 1$) window (as in Fig. 8c and d). For the 7 d win-

dow, the percentage of correct identification is well above 70 % over a large part of the world suggesting that the impact of model imperfection on climatological events occurring at time scales in the order of weeks can be correctly detected most of the time. Also, the signal obtained when the standardized anomalies are assimilated is more clear when a 7 d window is used (compare Fig. 10c and d with Fig. 9c and d).

3.2.3 Seasonally dependent model errors

To evaluate if the CME computed using AnDA can capture the temporal variability of model errors, the percentage of correct identification is computed over two subperiods, one corresponding to the Northern Hemisphere summer (June, July, and August) and the other one corresponding to the Northern Hemisphere winter (December, January, and February). Figure 11 shows that the sensitivity in the mid-level mean temperature and mean precipitation in these two periods is quite different, with some areas showing the opposite sensitivity (e.g., central South America, southern North America, southern Africa and Australia).

Figure 12 shows that the CME for both seasons is quite different and that the evidence is generally stronger in that hemisphere where there is summer at that time, i.e., when convection is more frequent. Note in particular that over the North Atlantic, Europe and Asia, the CME is stronger during the summer. According to the SPEEDY model climatology convective rain is larger during this time of the year (not shown) and so is the sensitivity of this precipitation to the RH_{conv} parameter (see Fig. 11c and d). The seasonal cy-

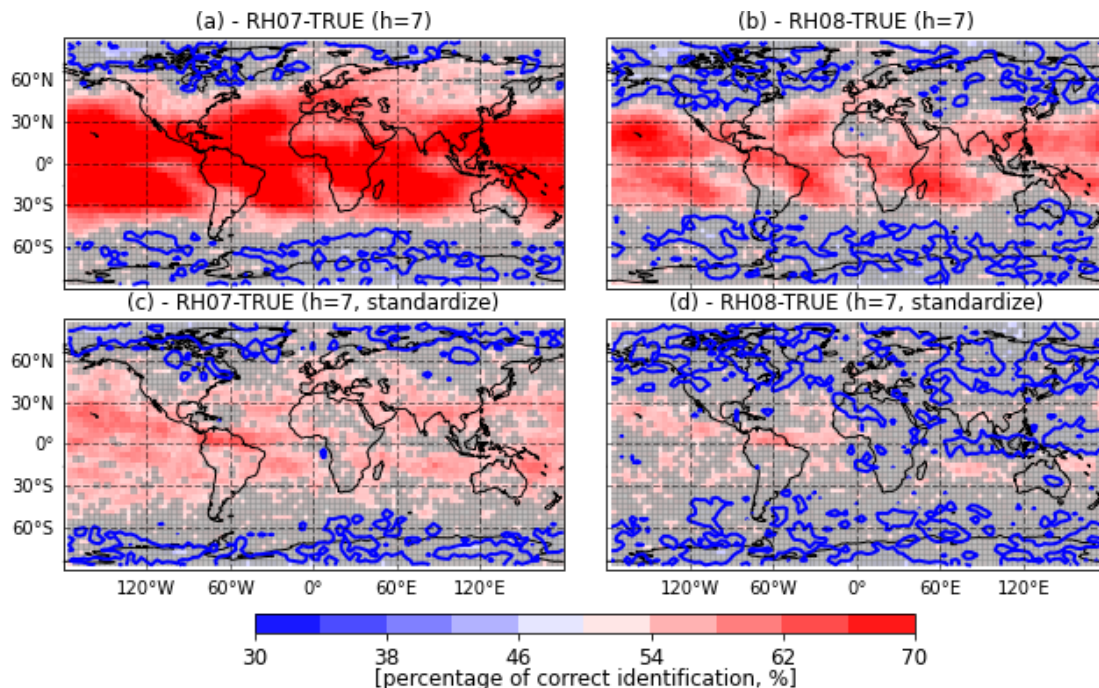


Figure 10. Percentage of correct identification based on the CME for a 7 d evidence window between the RH07 and TRUE catalog (a, c) and between the RH08 and TRUE catalogs (b, d) using AnDA. The percentage of correct identification is computed assimilating temperature observations or their standardized values (a and b, c and d, respectively). The gray shade indicates values which are below the 95 % statistical significance level. The blue lines correspond to a 50 % of correct identification.

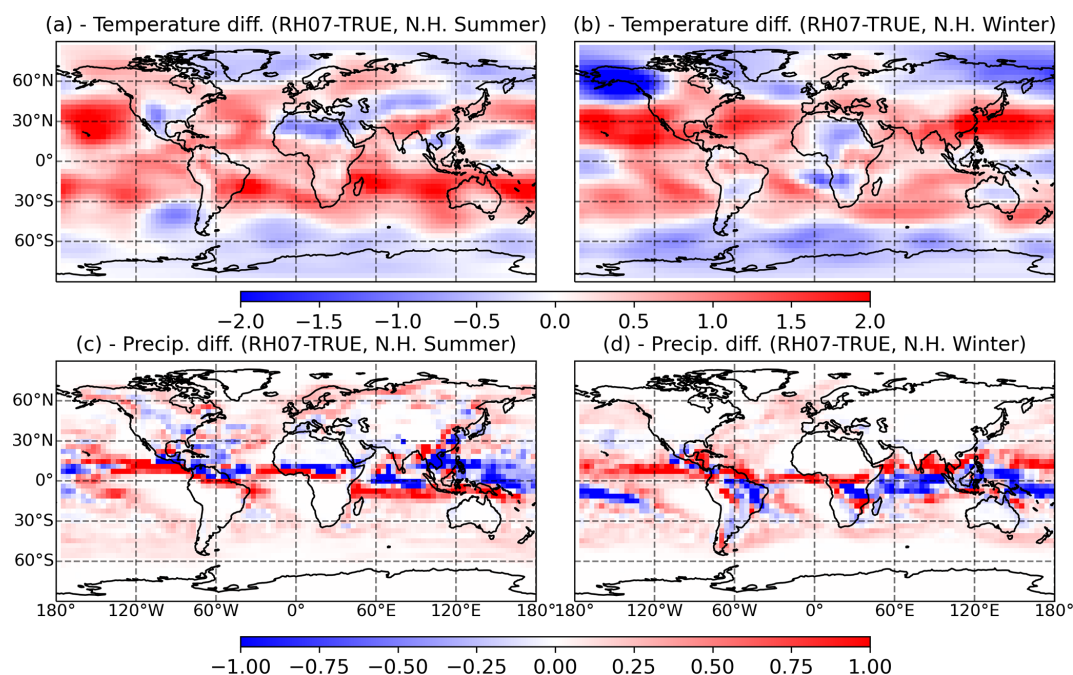


Figure 11. The 30-year mean difference between the TRUE model simulation and the RH07 for the Northern Hemisphere summer (June, July, and August, a, c) and the Northern Hemisphere winter (December, January, and February, b, d) for temperature at $\sigma = 0.5$ (a, c, K) and for convective precipitation (c, d, mm d^{-1}).

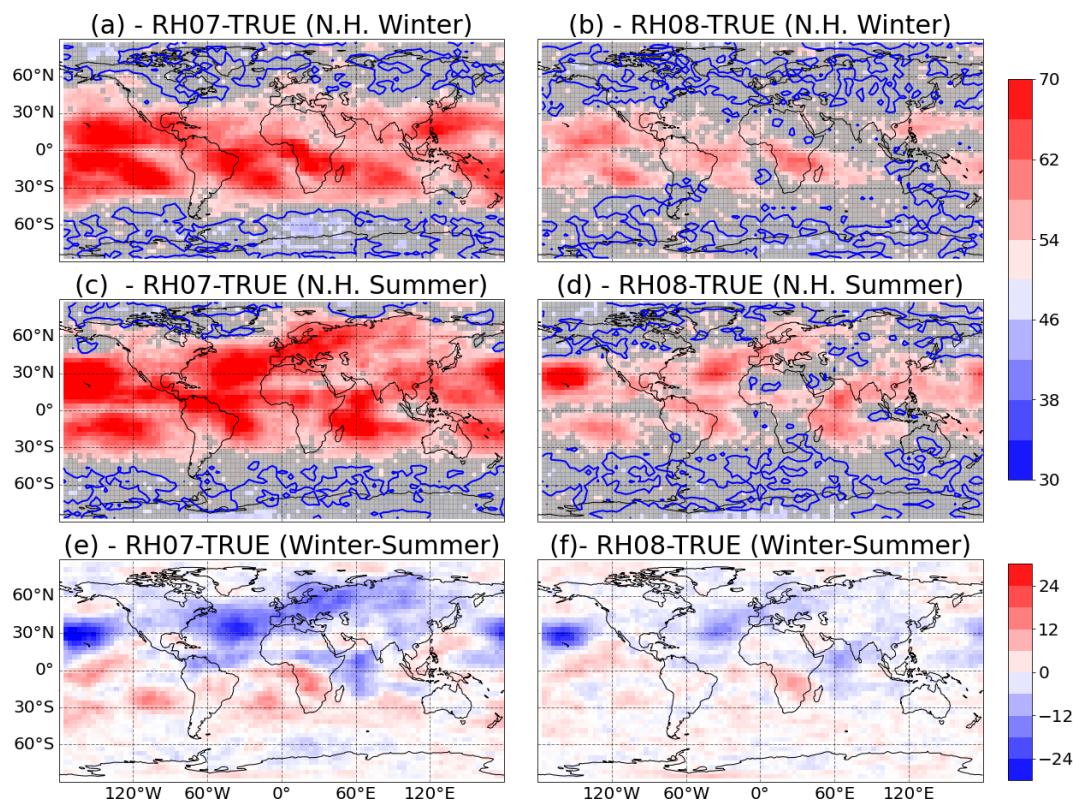


Figure 12. Percentage of correct identification between RH07 and TRUE catalog (**a**, **c**, **e**) and between the RH08 and TRUE catalogs (**b**, **d**, **f**) based on CME for the Northern Hemisphere winter (December, January, and February, **a**, **b**), the Northern Hemisphere summer (June, July, and August, **c**, **d**). The maps (**e**) and (**f**) show the difference between Northern Hemisphere winter minus summer in the percentage of correct identification. The gray shade indicates values which are below the 95 % statistical significance level. The blue lines correspond to 50 % of correct identification.

cle is acting as a state-dependent source of model error and the CME successfully recovered this characteristic providing useful information about the source of the difference among the models used to generate the different catalogs.

4 Discussion and conclusions

Model selection using data assimilation has been introduced in Carrassi et al. (2017), applied in the context of climate change detection and attribution in Hannart et al. (2016), and to complex or large state model evaluation in Metref et al. (2019). It consists of putting two or more dynamical models in competition and uses observations to compute a likelihood, also called CME, for each model, attributing the highest probability to the model which provides forecasts that better match the observations over a period of time. The CME compares the observations with the mean of the forecast distribution computed in the assimilation cycles, but also takes into account the uncertainties described by the variance of the forecast distribution.

The main issue related to model selection using data assimilation is the computational cost to perform several model

runs. Recently, a data-driven data assimilation method based on analog forecasting has been proposed in Tandeo et al. (2015) and Lguensat et al. (2017). It consists in replacing the physical model equations by a catalog of past observations or numerical simulations to statistically emulate the dynamics of the system. This current paper explores the use of such a model-free strategy in data assimilation to compute the CME of several dynamical models, represented by a set of numerical simulations.

The proposed methodology was assessed using numerical experiments. A first set of experiments was performed using a modified 3-variable Lorenz-63 model. It was found in particular that using analog data assimilation gives similar results to classical data assimilation, where the dynamical model is run at each time step. It indicates that the method is able to provide an accurate approximation of the CME metric of a dynamical model given a catalog and set of noisy observations. A second set of experiments was done using the intermediate complexity atmospheric general circulation model SPEEDY. They indicate that the proposed methodology is efficient in identifying the correct model using only observations of a small part of the state. The numerical results also highlight the importance of using data assimilation

compared to more direct approaches which rely only on the forecast sensitivity without proper state initialization. As a summary, the numerical results indicate that this technique is efficient in selecting the best model to describe a sequence of noisy observations, and it has various advantages: (i) it is a low computational cost method, (ii) it can be applied locally on a subpart of the system or using complex observation operators such as integrated parameters and (iii) it uses already existing model outputs.

This work is the first step of a more challenging project. It shows that selecting and weighting dynamical models can be performed inside a data assimilation framework, using analog forecasts, and thus avoiding the need to run numerical models to get predictions. This result opens up new perspectives for the use, for instance, of model simulations such as the Coupled Model Intercomparison Project (CMIP), see Eyring et al. (2016) for more details. The goal will be to propose weighted projections of climate indices, where the weights will be based on the ability of different climate models to represent the local dynamics and the current observations. The resulting weighted projections will improve the estimation of the mean and standard deviation of climate indices. Those results will be compared to model democracy strategy, for instance, which gives equal weights to all climate models, but is also somehow controversial (Knutti et al., 2019).

Implementing the combination of CME and AnDA in real-data cases brings additional challenges. For instance, in this work the application of the analog regression technique to a high-dimensional problem is achieved by using local domains. However, this approach does not take advantage of the covariance structure of the model output. This structure could be retrieved through a principal component analysis which may allow the implementation of the analog regression in a low-dimensional space while keeping the main aspects of large scale circulation patterns.

Another possible application of the AnDA-CME framework is in the context-weighted supermodels (Schevenhoven et al., 2019; Schevenhoven and Carrassi, 2022), which provides a way to combine the time derivatives of different models resulting in improved short-range and long-range predictions.

Code and data availability. Codes can be obtained from the GitHub open repository: https://github.com/gustfrontar/AnDA_SPEEDY.git (last access: 21 September 2022), <https://doi.org/10.5281/zenodo.5803356> (Ruiz and Tandeo, 2021). They include the codes for approximating the SPEEDY model using the analog forecasting technique, some simulated data and the specific version of AnDA used for the SPEEDY experiments. AnDA is an open-code system and can be obtained from the GitHub open repository: <https://github.com/ptandeo/AnDA> (last access: 21 September 2022), <https://doi.org/10.5281/zenodo.5795943> (Tandeo and Navaro, 2021).

Author contributions. All authors contributed in the model development, verification, discussions about the results, and preparation of the manuscript. TTTC and JR were responsible for the numerical experiments and prepared the figures.

Competing interests. The contact author has declared that none of the authors has any competing interests.

Disclaimer. Publisher's note: Copernicus Publications remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Acknowledgements. We thank the ECOS-Sud Program for its financial support through the project A17A08. We also thank financial support for the AMIGOS project from the Région Bretagne (ARED fellowship for PLB) and from the ISblue project (Interdisciplinary graduate school for the blue planet, ANR-17-EURE-0015) co-funded by a grant from the French government under the program Investissements d'Avenir. This research has also been supported by the National Agency for the Promotion of Science and Technology of Argentina (grant no. PICT 2233-2017), the University of Buenos Aires (grant no. UBACyT-20020170100504).

Financial support. This research has been supported by the Centre National de la Recherche Scientifique (grant no. ECOSud A17A08), the Région Bretagne (grant no. ARED fellowship), the Agencia Nacional de Promoción Científica y Tecnológica (grant no. PICT 2233-2017), the Universidad de Buenos Aires (grant no. UBACyT-20020170100504), and the Agence Nationale de la Recherche (grant no. ANR-17-EURE-0015).

Review statement. This paper was edited by Travis O'Brien and reviewed by two anonymous referees.

References

- Atencia, A. and Zawadzki, I.: A comparison of two techniques for generating nowcasting ensembles. Part II: Analogs selection and comparison of techniques, *Mon. Weather Rev.*, 143, 2890–2908, 2015.
- Ayet, A. and Tandeo, P.: Nowcasting solar irradiance using an analog method and geostationary satellite images, *Sol. Energy*, 164, 301–315, 2018.
- Bannayan, M. and Hoogenboom, G.: Predicting realizations of daily weather data for climate forecasts using the non-parametric nearest-neighbour re-sampling technique, *Int. J. Climatol.*, 28, 1357–1368, 2008.
- Barnett, T. and Preisendorfer, R.: Multifield analog prediction of short-term climate fluctuations using a climate state vector, *J. Atmos. Sci.*, 35, 1771–1787, 1978.

- Burgers, G., Jan van Leeuwen, P., and Evensen, G.: Analysis scheme in the ensemble Kalman filter, *Mon. Weather Rev.*, 126, 1719–1724, 1998.
- Carrassi, A., Bocquet, M., Hannart, A., and Ghil, M.: Estimating model evidence using data assimilation, *Q. J. Roy. Meteor. Soc.*, 143, 866–880, 2017.
- Carrassi, A., Bocquet, M., Bertino, L., and Evensen, G.: Data assimilation in the geosciences: An overview of methods, issues, and perspectives, *Wiley Interdisciplinary Reviews: Climate Change*, 0, e535, <https://doi.org/10.1002/wcc.535>, 2018.
- Carson, J., Crucifix, M., Preston, S., and Wilkinson, R. D.: Bayesian model selection for the glacial–interglacial cycle, *J. Roy. Stat. Soc. C*, 67, 25–54, 2018.
- Cleveland, W. S. and Devlin, S. J.: Locally weighted regression: an approach to regression analysis by local fitting, *J. Am. Stat. Assoc.*, 83, 596–610, 1988.
- Desroziers, G., Berre, L., Chapnik, B., and Poli, P.: Diagnosis of observation, background and analysis-error statistics in observation space, *Q. J. Roy. Meteor. Soc.*, 131, 3385–3396, 2005.
- Eyring, V., Bony, S., Meehl, G. A., Senior, C. A., Stevens, B., Stouffer, R. J., and Taylor, K. E.: Overview of the Coupled Model Intercomparison Project Phase 6 (CMIP6) experimental design and organization, *Geosci. Model Dev.*, 9, 1937–1958, <https://doi.org/10.5194/gmd-9-1937-2016>, 2016.
- Hannart, A., Carrassi, A., Bocquet, M., Ghil, M., Naveau, P., Pulido, M., Ruiz, J., and Tandeo, P.: DADA: data assimilation for the detection and attribution of weather and climate-related events, *Climatic Change*, 136, 155–174, 2016.
- Knutti, R., Baumberger, C., and Hadorn, G. H.: Uncertainty quantification using multiple models—Prospects and challenges, in: *Computer Simulation Validation*, 835–855, Springer, https://doi.org/10.1007/978-3-319-70766-2_34, 2019.
- Kotsuki, S., Sato, Y., and Miyoshi, T.: Data Assimilation for Climate Research: Model Parameter Estimation of Large-Scale Condensation Scheme, *J. Geophys. Res.-Atmos.*, 125, e2019JD031304, <https://doi.org/10.1029/2019JD031304>, 2020.
- Lauvaux, T., Díaz-Isaac, L. I., Bocquet, M., and Boussez, N.: Diagnosing spatial error structures in CO₂ mole fractions and XCO₂ column mole fractions from atmospheric transport, *Atmos. Chem. Phys.*, 19, 12007–12024, <https://doi.org/10.5194/acp-19-12007-2019>, 2019.
- Lguensat, R., Tandeo, P., Ailliot, P., Pulido, M., and Fablet, R.: The analog data assimilation, *Mon. Weather Rev.*, 145, 4093–4107, 2017.
- Li, H., Kalnay, E., and Miyoshi, T.: Simultaneous estimation of covariance inflation and observation errors within an ensemble Kalman filter, *Q. J. Roy. Meteor. Soc.*, 135, 523–533, <https://doi.org/10.1002/qj.371>, 2009.
- Lorenz, E. N.: Deterministic nonperiodic flow, *J. Atmos. Sci.*, 20, 130–141, 1963.
- Lorenz, E. N.: Atmospheric predictability as revealed by naturally occurring analogues, *J. Atmos. Sci.*, 26, 636–646, 1969.
- Metref, S., Hannart, A., Ruiz, J., Bocquet, M., Carrassi, A., and Ghil, M.: Estimating model evidence using ensemble-based data assimilation with localization – The model selection problem, *Q. J. Roy. Meteor. Soc.*, 145, 1571–1588, <https://doi.org/10.1002/qj.3513>, 2019.
- Miyoshi, T.: The Gaussian approach to adaptive covariance inflation and its implementation with the local ensemble transform Kalman filter, *Mon. Weather Rev.*, 139, 1519–1535, <https://doi.org/10.1175/2010MWR3570.1>, 2011.
- Molteni, F.: Atmospheric simulations using a GCM with simplified physical parametrizations. Part I: model climatology and variability in multi-decadal experiments, *Clim. Dynam.*, 20, 175–191, <https://doi.org/10.1007/s00382-002-0268-2>, 2003.
- Otsuka, S. and Miyoshi, T.: A bayesian optimization approach to multimodel ensemble kalman filter with a low-order model, *Mon. Weather Rev.*, 143, 2001–2012, 2015.
- Palmer, T. N.: A nonlinear dynamical perspective on climate prediction, *J. Climate*, 12, 575–591, 1999.
- Pearl, J.: *Causality: models, reasoning, and inference*, Cambridge University Press, ISBN-13 978-0521895606, 521, 8, 2000.
- Platzer, P., Yiou, P., Naveau, P., Tandeo, P., Zhen, Y., Ailliot, P., and Filipot, J.-F.: Using local dynamics to explain analog forecasting of chaotic systems, *J. Atmos. Sci.*, 2117–2133, <https://doi.org/10.1175/JAS-D-20-0204.1>, 2021.
- Reich, S.: Data assimilation: the schrödinger perspective, *Acta Numerica*, 28, 635–711, 2019.
- Reich, S. and Cotter, C.: Probabilistic forecasting and Bayesian data assimilation, Cambridge University Press, <https://doi.org/10.1017/CBO9781107706804>, 2015.
- Ruiz, J. and Pulido, M.: Parameter estimation using ensemble-based data assimilation in the presence of model error, *Mon. Weather Rev.*, 143, 1568–1582, 2015.
- Ruiz, J. and Tandeo, P.: AnDA-SPEEDY code version 1.0.0, Zenodo [code], <https://doi.org/10.5281/zenodo.5803356>, 2021.
- Ruiz, J., Pulido, M., and Miyoshi, T.: Estimating model parameters with ensemble-based data assimilation: A review, *J. Meteorol. Soc. Jpn.*, 91, 79–99, 2013.
- Schevenhoven, F. and Carrassi, A.: Training a supermodel with noisy and sparse observations: a case study with CPT and the synch rule on SPEEDO – v.1, *Geosci. Model Dev.*, 15, 3831–3844, <https://doi.org/10.5194/gmd-15-3831-2022>, 2022.
- Schevenhoven, F., Seltén, F., Carrassi, A., and Keenlyside, N.: Improving weather and climate predictions by training of supermodels, *Earth Syst. Dynam.*, 10, 789–807, <https://doi.org/10.5194/esd-10-789-2019>, 2019.
- Schirber, S., Klocke, D., Pincus, R., Quaas, J., and Anderson, J. L.: Parameter estimation using data assimilation in an atmospheric general circulation model: From a perfect toward the real world, *J. Adv. Model. Earth Sy.*, 5, 58–70, 2013.
- Sévellec, F. and Drijfhout, S. S.: A novel probabilistic forecast system predicting anomalously warm 2018–2022 reinforcing the long-term global warming trend, *Nat. Commun.*, 9, 1–12, 2018.
- Siebert, S., Ferro, C. A., Stephenson, D. B., and Leutbecher, M.: The ensemble-adjusted Ignorance Score for forecasts issued as normal distributions, *Q. J. Roy. Meteor. Soc.*, 145, 129–139, 2019.
- Tandeo, P. and Navaro, P.: AnDA code version 1.0.0, Zenodo [code], <https://doi.org/10.5281/zenodo.5795943>, 2021.
- Tandeo, P., Ailliot, P., Ruiz, J., Hannart, A., Chapron, B., Cu-zol, A., Monbet, V., Easton, R., and Fablet, R.: Combining analog method and ensemble data assimilation: application to the Lorenz-63 chaotic system, in: *Machine Learning and Data Mining Approaches to Climate Science*, 3–12, Springer, https://doi.org/10.1007/978-3-319-17220-0_1, 2015.
- Tandeo, P., Ailliot, P., Bocquet, M., Carrassi, A., Miyoshi, T., Pulido, M., and Zhen, Y.: A review of innovation-based meth-

- ods to jointly estimate model and observation error covariance matrices in ensemble data assimilation, *Mon. Weather Rev.*, 148, 3973–3994, 2020.
- Van Leeuwen, P. J., Künsch, H. R., Nerger, L., Potthast, R., and Reich, S.: Particle filters for high-dimensional geoscience applications: A review, *Q. J. Roy. Meteor. Soc.*, 145, 2335–2365, 2019.
- Yiou, P.: AnaWEGE: a weather generator based on analogues of atmospheric circulation, *Geosci. Model Dev.*, 7, 531–543, <https://doi.org/10.5194/gmd-7-531-2014>, 2014.
- Zhen, Y., Tandeo, P., Leroux, S., Metref, S., Penduff, T., and Le Sommer, J.: An adaptive optimal interpolation based on analog forecasting: application to SSH in the Gulf of Mexico, *J. Atmos. Ocean. Tech.*, 37, 1697–1711, 2020.