



Stochastic Parameterization Using Compressed Sensing: Application to the Lorenz-96 Atmospheric Model

A. MUKHERJEE

Y. AYDOGDU

T. RAVICHANDRAN

N. SRI NAMACHCHIVAYA

**Author affiliations can be found in the back matter of this article*

ORIGINAL RESEARCH
PAPER



STOCKHOLM
UNIVERSITY PRESS

ABSTRACT

Growing set of optimization and regression techniques, based upon sparse representations of signals, to build models from data sets has received widespread attention recently with the advent of compressed sensing. This paper deals with the parameterization of the Lorenz-96 (Lorenz, 1995) model with two time-scales that mimics mid-latitude atmospheric dynamics with microscopic convective processes. Compressed sensing is used to build models (vector fields) to emulate the behavior of the fine-scale process, so that explicit simulations become an online benchmark for parameterization. We apply compressed sensing, where the sparse recovery is achieved by constructing a sensing/dictionary matrix from ergodic samples generated by the Lorenz-96 atmospheric model, to parameterize the unresolved variables in terms of resolved variables. Stochastic parameterization is achieved by auto-regressive modelling of noise. We utilize the ensemble Kalman filter for data assimilation, where observations (direct measurements) are assimilated in the low-dimensional stochastic parameterized model to provide predictions. Finally, we compare the predictions of compressed sensing and Wilks' polynomial regression to demonstrate the potential effectiveness of the proposed methodology.

CORRESPONDING AUTHOR:

N. Sri Namachchivaya

Department of Applied
Mathematics, University of
Waterloo, Waterloo, Ontario
N2L 3G1, Canada

navam@uwaterloo.ca

KEYWORDS:

compressed sensing; sparse
regression; ensemble Kalman
filter; auto-regression

TO CITE THIS ARTICLE:

Mukherjee, A, Aydogdu, Y,
Ravichandran, T and
Sri Namachchivaya, N. 2022.
Stochastic Parameterization
Using Compressed Sensing:
Application to the Lorenz-96
Atmospheric Model. *Tellus A:
Dynamic Meteorology and
Oceanography*, 74(2022):
300–317. DOI: [https://doi.
org/10.16993/tellusa.42](https://doi.org/10.16993/tellusa.42)

1 INTRODUCTION

Global weather and climate models are the best available tools for projecting likely climate changes. However, current climate model projections are still considerably uncertain, due to uncertainties and errors in mathematical models, “curse of dimensionality”, model initialization, as well as inadequate observations and computational errors. Multi-layer models are adequate to capture the dynamics of large-scale atmospheric phenomena. Solving for small-scale processes with help of large eddy simulation (LES) models, on a high-resolution grid embedded (with grid spacing of a few tens of meters and can explicitly simulate the clouds and turbulence represented by parameterization schemes) within the low-resolution grid is very useful, but its applicability remains limited to its high computational cost. In this case, learning methods are used to develop parameterization schemes that utilize the explicitly simulated clouds and turbulence. In practice, low-resolution models resolve the large-scale processes while the effects of the small-scale processes are often parameterized in terms of the large-scale variables.

In 1995, Ed Lorenz came up with a simple model that captured the essence of this problem with two stacked layers, which is referred to as the Lorenz-96 (Lorenz, 1995) model. Even though this simple model is still a long way from the Primitive Equations that we need to eventually study, it is hoped that the parameterization schemes and filtering methods presented here can be ultimately adapted for the realistic models that need to be computationally solved for weather and climate predictions. We use the Lorenz-96 model to demonstrate the sparse optimization methods and regression techniques to build models from data. This model, as discussed in Section 2, has nonlinearities resembling the advective nonlinearities of fluid dynamics and a multiscale coupling of slow and fast variables similar to what is seen in the two-layer atmospheric models. Resolving all the relevant length and time scales in simulations of the turbulent atmospheric circulations remains out of reach due to computational constraints. Hence, the small-scale processes must be represented indirectly, through parameterization schemes which estimate their net impact on the resolved atmospheric state.

The objective of this paper is to develop a data-driven method to approximate the effects of small-scale processes with simplified functions called parameterizations which are characterized by compressed sensing (CS). Such parameterization schemes are inherently approximate, and the development of schemes which produce realistic simulations is a central challenge of model development. Most parameterization schemes depend critically on various parameters whose values cannot be determined a priori but must instead be estimated

through data. Recent developments in optimization and sparse representations in high-dimensional spaces have provided powerful new techniques.

This paper deals with the parameterization of the Lorenz-96 model with two time-scale simplified ordinary differential equations describing advection, damping and forcing. We assume that the model that expresses the resolved atmospheric variables of the known physics or the resolved modes are well understood. Furthermore, Lorenz in introducing this multiscale model showed that the small-scale features of the atmosphere act like forcing. We present a CS and sparse recovery approach for constructing approximations to the Lorenz-96 model system that blends physics-informed candidate functions with functions that enforce sparse structure. The ultimate aim is to adapt the mathematical methods developed here for the realistic models that can be solved computationally to examine the turbulent atmospheric circulations. In the past few years, data-driven parameterization of the Lorenz-96 model using machine learning and deep learning techniques (Gagne et al., 2020) has shown promising results.

The aim of any parameterization is to develop enhanced models that capture the essence of dynamics and produce reliable predictions. Besides deterministic parameterization, stochastic parameterization was applied to the Lorenz-96 model with several kinds of noise and provided promising results (Arnold et al., 2013). In our problem, we extract the useful elements for parameterization by using CS that relies on the restricted isometry property (RIP) (Candès and Wakin, 2006). The RIP is satisfied due to the chaotic and ergodic nature of the Lorenz-96 model (numerically verified), showing that CS would be a great fit for stochastic parametrization. In this way, we are able to capture the key components that allow us to represent the unresolved variables in terms of resolved variables with only a few elements. After the parameterization of the deterministic part with compressed sensing, the stochastic part is achieved by auto-regressive modelling of noise. These parameterized models are then used for the predictions by using ensemble Kalman filter, where the observations are assimilated through the slow time scale.

The paper is organized as follows. In Section 2, we will introduce the 40-dimensional Lorenz-96 model, that was originally suggested by (Lorenz, 1995), and describe the characterization of complex dynamics including statistical properties, Lyapunov exponents, marginal probability density functions (PDFs) and auto-correlation functions (ACFs). Section 3 will present the details on data-driven reduction methods such as regression and compressed sensing with sparse optimization. In Section 4, stochastic parameterization using the compressed sensing approach will be illustrated on Lorenz-96 model. Section 5 will describe a general ensemble filtering algorithm, where the forecasting step is straightforward

because the Markov kernel is known, while the update step is much more sophisticated. This nonlinear filtering method will be applied to parameterized models. Finally, conclusions and future work discussion will follow in Section 6. To ensure that our results are reproducible, we have provided our code in (Mukherjee et al., 2021).

2 THE LORENZ-96 ATMOSPHERIC MODEL

Let us now turn to a simple atmospheric model, which nonetheless exhibits many of the difficulties arising in realistic models, to gain insight into predictability and data assimilation. The Lorenz-96 model was originally introduced in (Lorenz, 1995) to mimic multiscale mid-latitude atmospheric dynamics for an unspecified scalar meteorological quantity at K equidistant grid points along a latitude circle:

$$\dot{X}_k = -X_{k-1}(X_{k-2} - X_{k+1}) - X_k + F - \frac{hc}{b} \sum_{j=1}^J Z_{j,k} \quad (1)$$

$$\dot{Z}_{j,k} = -cbZ_{j+1,k}(Z_{j+2,k} - Z_{j-1,k}) - cZ_{j,k} + \frac{hc}{b} X_k \quad (2)$$

Equation (1) describes the dynamics of some atmospheric quantity X , and X_k can represent the value of this variable at time t in the k^{th} sector defined over a latitude circle in the mid-latitude region. (Note that we use subscripts k and j to conform with the typical spatial indexing notation used for the Lorenz-96 model. In sections that follow, subscripts k and j will be used as discrete time indices, not to be confused with the spatial indices of the Lorenz model). The latitude circle is divided into K sectors. For convenience, the vector field associated with the slow components or the resolved variables is assumed a priori known with representable physics, which is denoted by

$$\mathcal{G}_k(X) := -X_{k-1}(X_{k-2} - X_{k+1}) - X_k + F \quad (3)$$

Figure 1 illustrates the behavior of a generic slow state X_k (shown in the inner ring with $K = 8$ for illustration), the fast states in the 1st sector, that is, $Z_{1,1}, \dots, Z_{J,1}$ (shown in the outer ring with $J = 7$ for illustration), and the fast scale forcing that enters (1); for the X_k component the fast scale forcing is

$$U_k = -\frac{hc}{b} \sum_{j=1}^J Z_{j,k} \quad (4)$$

Each X_k is coupled to its neighbors X_{k+1} , X_{k-1} , and X_{k-2} by (1). X_k and $Z_{j,k}$ are periodic on the spatial domain, and (1) and (2) applies for all values of k and j by letting $X_{k+K} = X_k$, $Z_{j,k+K} = Z_{j,k}$, $Z_{j+J,k} = Z_{j,k+1}$ so, for example, in **Figure 1** we have $X_9 = X_1$, $Z_{1,9} = Z_{1,1}$, and $Z_{8,1} = Z_{1,2}$. Even though this model is not a simplification of any atmospheric system, it is designed to satisfy three basic properties: contains a

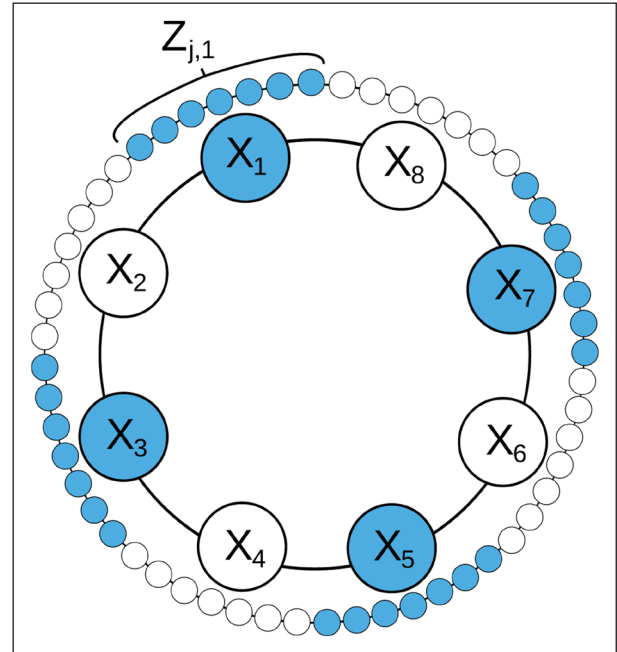


Figure 1 A recreation of a similar image by Wilks (2005), that illustrates the two timescale Lorenz-96 model.

quadratic discrete advection-like term that conserves the total energy (i.e. it does not contribute to dE/dt) just like many atmospheric models; it has linear dissipation (the $-X_k$ term) that decreases the energy defined as $E = \sum_{k=1}^K X_k$, and an external forcing term $F > 0$ that can increase or decrease the total energy.

The original K -dimensional system was extended to study the influence of multiple spatio-temporal scales on the predictability of atmospheric flow by dividing each segment k into J subsectors ($J = 10$), and introducing a fast variable, $Z_{j,k}$ given by (2), associated with each subsector. Thus, each X_k represents a slowly-varying, large amplitude atmospheric quantity, with J fast-varying, low amplitude, similarly coupled quantities, $Z_{j,k}$, associated with it. In the context of climate modelling, the slow component is also known as the resolved climate modes while the rapidly varying component is known as the unresolved non-climate modes. The coupling terms between neighbors model advection between sectors and subsectors, while the coupling between each sector and its subsectors model damping within the model. The model is subjected to linear external forcing, F , on the slow timescale.

We consider 40-dimensional slow variable components and 400-dimensional fast variable components. With $K = 40$, one can think of $\mathbf{X}$ as representing an unspecified scalar meteorological quantity equally spaced along a mid-latitude belt of roughly 30,000 km, thereby giving the model the correct synodic scale to mimic Rossby waves. Our main research interest in this paper is to develop a dimensional reduction and parameterization method based on *compressed sensing* (Candès and Wakin, 2006; Donoho, 2006). We compare this compressed sensing scheme based on sparse regression with Wilks' *polynomial regression* (Wilks, 2005).

2.1 CHARACTERIZATION OF COMPLEX DYNAMICS

According to (Devaney, 1989), a chaotic dynamical system defined on a metric space X by a continuous map f is said to be chaotic on X if it exhibits the following three essential features; (a) f is transitive, (b) the periodic points of f are dense in X , and (c) f has sensitive dependence on initial conditions. These systems are completely deterministic in that if the current state of the system X_0 is known, all future states at time t , denoted by X_t , $t > 0$, are known. However, even small numerical errors in the computation of X_0 can lead to very large differences in the values of X_t , so while these systems are deterministic, they are not predictable. Hence, chaotic systems exhibit a large departure in the long-term evolution for very tiny changes in initial conditions, and hence to study the limiting behavior of these systems, one cannot study the global evolution “orbit by orbit”. Instead, one studies the *statistical properties* of these systems, as we describe in this section.

2.1.1 Statistical properties of chaotic systems

Although individual orbits of chaotic dynamical systems are by definition unpredictable, the average behavior of typical trajectories can often be given a precise statistical description. Anosov (1969), Pesin (2004) and Sinai (1977) showed that a wide class of “chaotic” maps exhibit natural ergodic invariant measures, the maps being “chaotic” in the sense that the orbits of two points, which may be arbitrarily close to each other with respect to a metric, may evolve very differently over time (positive Lyapunov exponents as discussed below). Loosely speaking, ergodicity is the “law of large numbers” for dynamical systems. While the ergodic theorems prove convergence to the mean, natural questions arise about the deviation of these processes from the mean and about other statistical properties such as rates of mixing. For a given invariant measure, and a class of observables, the correlation functions tell whether (and how fast) the system “mixes”, i.e. “forgets” its initial conditions.

Many statistical properties, such as the Central Limit Theorem and the Law of Iterated Logarithm, displayed by independent identically distributed stochastic processes, have been shown to hold for stochastic processes generated by chaotic dynamical systems. One of the better known techniques for establishing such results was proved by Young (1998), where it was shown that for such (chaotic) dynamical systems with certain “towers” property (a powerful tool for coding the dynamics of a system in such a way so as to enable the computation of its statistical properties), behave as Markov chains on spaces with infinitely many states. Should such dynamical systems mix sufficiently quickly, the process satisfies the *central limit theorem* and the completely deterministic process behaves, at least in the first two moments, as an independent stochastic process. Furthermore, the degree

of independence is characterized by the rate at which the process mixes.

2.1.2 Lyapunov Exponents (LEs)

Among the features listed above, sensitivity to initial conditions is widely known to be associated with chaotic dynamical systems, and a quantitative measure of this sensitivity is given by the *Lyapunov exponents* (LEs) which are the most commonly used quantities to characterize chaotic phenomena. The LEs are asymptotic measures characterizing the average rate of divergence (or convergence) of small perturbations to the solutions of a dynamical system. The value of the maximum LE is an indicator of the chaotic or regular nature of orbits.

Considering the phase space, which is assumed to be $\mathbb{R}^d$ or a Riemannian manifold, and denoting by M throughout, the Lebesgue or the Riemannian measure on M is denoted by m , and the dynamics are generated by iterating a self-map of M , written $f : M \rightarrow M$. Given a differentiable map $f : M \rightarrow M$, a point $x \in M$ and a tangent vector v at x , we define

$$\lambda(x, v) = \lim_{n \rightarrow \infty} \frac{1}{n} \log |Df_x^n(v)| \text{ if this limit exists} \quad (5)$$

where $Df_x^n(v)$ is the derivative of the n th iteration of f , (i.e., $f \circ \dots \circ f$), with respect to variable x at the vector v . And if the limit does not exist, then the limit is replaced by inf or sup and we write $\underline{\lambda}(x, v)$ and $\bar{\lambda}(x, v)$, respectively. Thus, $\lambda(x, v) > 0$ implies the exponential growth of $|Df_x^n(v)|$ and can be interpreted to mean exponential divergence of nearby orbits. The computational procedure for Lyapunov exponents using the standard method can be found in (Wolf et al., 1985; Skokos, 2010).

When the chaotic system is dissipative, solutions will eventually reside on an attractor. The dimension of this attractor, i.e., the number of “active” degrees of freedom, is also known as the Kaplan-Yorke dimension D_{KY} (Eckmann and Ruelle, 1985; Beeson and Namachchivaya, 2020) and is given as

$$D_{KY} = r + \frac{1}{|\tilde{\lambda}_{r+1}|} \sum_{i=1}^r \tilde{\lambda}_i \quad (6)$$

where $\tilde{\lambda}_1 \geq \tilde{\lambda}_2 \geq \dots \geq \tilde{\lambda}_d$ are the Lyapunov exponents counted with multiplicity ($d = \dim(M)$) and r is the largest integer such that $\sum_{i=1}^r \tilde{\lambda}_i > 0$.

The Kolmogorov-Sinai entropy H_{KS} is a measure of the diversity of the trajectories generated by the dynamical system. This entropy measure is determined through the Pesin formula, which provides an upper bound to H_{KS} (Eckmann and Ruelle, 1985) as

$$H_{KS} = \sum_{\tilde{\lambda}_i > 0} \tilde{\lambda}_i. \quad (7)$$

Complexity of the dynamics of the Lorenz-96 model varies considerably with different values of the constant

forcing term F . It is shown in (Majda et al., 2005), the model is in weakly chaotic regimes for $F = 5, 6$, strongly chaotic regime for the values of F from 8 to 10, and turbulent regimes for $F = 12, 16, 24$. We will present our results for a specific value of forcing $F = 10$.

In our numerical simulations, we begin from random initial conditions and use a fourth-order Runge-Kutta time integration with a time step of 0.01. Typically, we integrate forward for 500 time units before starting the calculation of Lyapunov exponents to ensure that all transients have decayed. At this point, we begin integrating the tangent space equations using the Gram-Schmidt procedure update every $\tau = 0.2$ time units. With $K = 40$, $F = 10$, $h = 0$ (just the slow variables), $c = 10$ and $b = 10$, the summary of Lyapunov Exponents obtained from the above simulation of the Lorenz-96 system is given in [Table 1](#). The time evolution of the Lyapunov exponents for the Lorenz-96 model

Largest Lyapunov exponent λ_1	2.3098
Error doubling time	0.3 time units
Number of strictly positive Lyapunov exponents	14
Number of neutral, $\lambda \in [-1E-2, 1E-2]$, exponents	1
Number of strictly negative Lyapunov exponents	26
Kaplan-Yorke dimension	29.4694
Kolmogorov-entropy	14.8409

Table 1 Summary of Lyapunov Exponents.

is shown in [Figure 2](#). As noted before, the calculation of the Lyapunov exponents begins at 500 time units and [Figure 2](#) depicts the convergence of the Lyapunov exponents around 1000 time units.

2.1.3 Marginal Probability Density Functions (PDFs) and Auto-correlation Functions (ACFs)

Typically, the PDFs are computed using bin-counting also known as histogram methods. For estimating the PDF of a continuous variable x , standard histogram methods simply partition x into distinct bins of width Δ_p , and then the number of observations n_i of x falling into bin i is converted into a normalized probability density for that bin.

More smoother estimates for PDFs can be obtained by using the Gaussian product kernel function which results in the estimation of the PDF known as the *kernel density estimation* (KDE) (Bishop, 2006):

$$p_x(x) = \frac{1}{N} \sum_{j=1}^N \frac{1}{(2\pi\lambda^2)^{1/2}} \exp\left\{-\frac{\|x_j - x\|^2}{2\lambda^2}\right\} \quad (8)$$

where λ represents the standard deviation of the Gaussian components and N is the total number of data points.

[Figure 3](#) depicts the marginal PDFs of some (randomly selected) slow variables X_1 , X_8 , X_{19} and X_{31} simulated as described above in Section 2. The convergence of the marginal PDFs shown in [Figure 3](#) demonstrates the existing symmetry among the slow variables.

The strongly mixing character of the slow dynamics can be inferred from the correlation function, which means

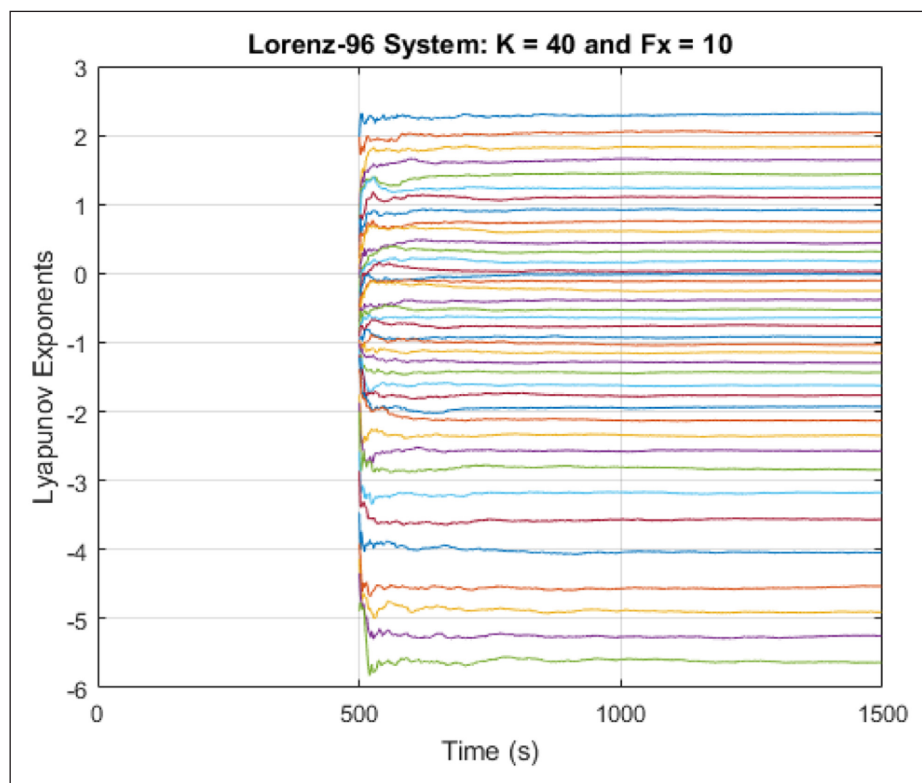


Figure 2 Time evolution of Lyapunov exponents for the Lorenz-96 system.

that the correlation function decays quickly to zero or near zero values. The discrete version of the auto-correlation functions (ACFs) for the slow variables is given by

$$C_{k,k}(m\Delta t) = \frac{1}{M-m} \sum_{h=1}^{M-m} (X_k(h\Delta t) - \bar{X})(X_k((h+m)\Delta t) - \bar{X}) \quad (9)$$

where

$$\bar{X} = \frac{1}{M} \sum_{m=1}^M X_k(m\Delta t) \quad (10)$$

and these formulae can be given similarly for the fast variables. **Figure 4** shows the ACF for the slow variable X_1 which decays very quickly and then oscillates around zero illustrating the strongly mixing character of the dynamics.

3 DATA-DRIVEN REDUCTION METHODS

The objective of this section is to compare model-based and data-driven tools for dimensional reduction and

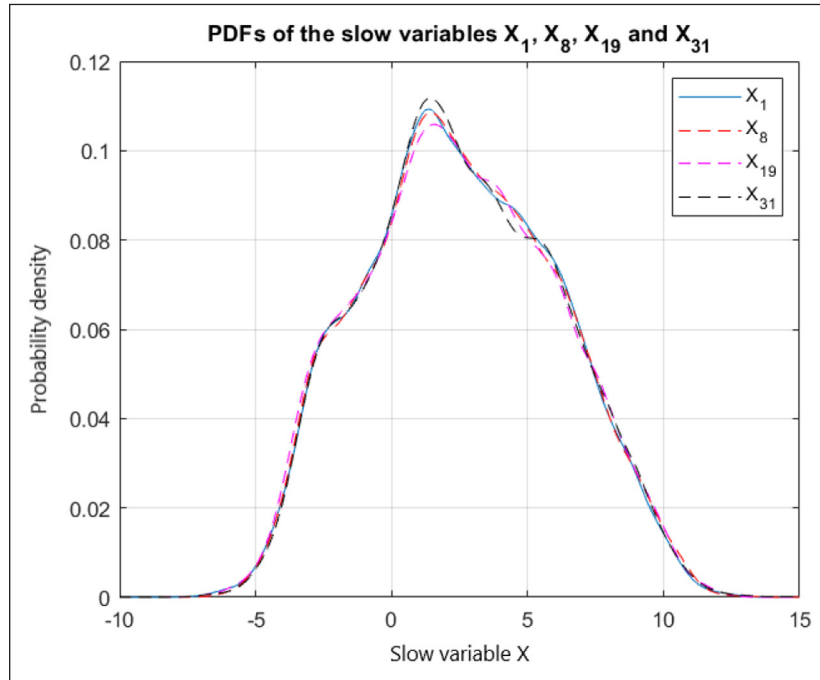


Figure 3 PDFs of the slow variables X_1 , X_8 , X_{19} and X_{31} .

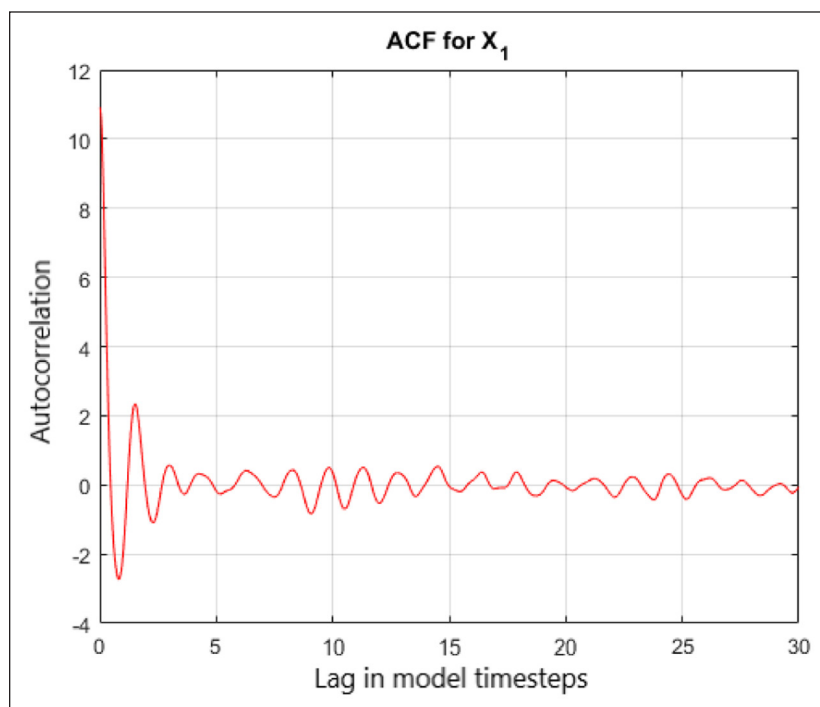


Figure 4 ACF of the slow variable X_1 .

reducing uncertainties in prediction. It is impossible to effectively “learn” from high-dimensional data unless there is some kind of implicit or explicit low-dimensional structure – this can be made mathematically precise in different ways. The methods presented below provide the low-dimensional structure either because of the presence of multi-scales in the model or due to the measurement matrix chosen to satisfy with high probability the so called “Restricted Isometry Property (RIP)” condition introduced in compressed sensing (Candès and Wakin, 2006).

Compressed sensing theory addresses the problem of recovering high-dimensional but sparse vectors from a limited number of measurements – the number of observed measurements m is significantly smaller than the dimension n of the original vector. We present a compressed sensing approach for sparse regression problem and using the ideas from concentration of measure, extend it to a case where the samples are generated from a chaotic time series, which are sufficiently mixing, thus ergodic, as shown in subsection 2.1. Here, the sparse recovery is achieved by constructing a matrix Θ in section 3.2 from ergodic samples generated by the Lorenz-96, such that it is suitable for reconstructing sparse parameter vectors via l_1 -minimization.

3.1 REGRESSION

This parameterization was proposed by Wilks (Wilks, 2005), and it uses a polynomial regression to parametrize the effects (4) of unresolved variables:

$$U_k(t) = P(X_k(t)) + e_k(t) \quad (11)$$

where P is a 4th degree polynomial in X_k and is given as:

$$P(X_k) = a_0 + a_1 X_k + a_2 X_k^2 + a_3 X_k^3 + a_4 X_k^4 \quad (12)$$

And e_k represents noise, which can be written as a first-order autoregressive model:

$$e_k(t_i) = \varphi e_k(t_{i-1}) + \sigma z_{k,i} \quad (13)$$

where $z_{k,i} \stackrel{i.i.d}{\sim} N(0,1)$ for all i, k . The method of calculating φ and σ is explained in Appendix A.

3.2 COMPRESSED SENSING AND SPARSE RECOVERY

Over the past two decades, researchers have focused on sparsity as one type of low-dimensional structure. Given the recent advances in both compressed sensing (Candès and Wakin, 2006; Donoho, 2006) and sparse regression (Tibshirani, 1996) it has become computationally feasible to extract system dynamics from large, multimodal datasets. The term sparse in signal processing context refers to the case where signals (or any type of data, in general) have merely few non-zero components with respect to the total number

of components. Sparsity plays a key role in optimization and data sciences. In the context of our work, these techniques rely heavily on the fact that many dynamical systems can be represented by governing equations that are sparse in the space of all possible functions of a given algebraic structure.

Compressed sensing (CS) is a technique for sampling and reconstructing signals as sparse signals, that is, the reconstructed signals have few non-zero components. More precisely, k -sparse signals are those that can be represented by $k \ll n$ significant coefficients over an n -dimensional basis. The central goal of CS is to capture attributes of a signal using very few measurements, which distinguishes CS from other dimensionality reduction techniques – a methodology for the recovery of sparse vectors from a small number of linear measurements. Hence, this allows for polynomial-time reconstruction of the sparse signal (Donoho, 2006).

In (Donoho, 2006) and (Candès and Tao, 2006), the original signal X is projected onto a lower-dimensional subspace via a random projection scheme, called the sensing matrix. This is used to reconstruct a signal $Y \in \mathbb{R}^m$. More precisely, this broader objective is exemplified by the important special case in which one is interested in finding a parameterization vector $S \in \mathbb{R}^n$ using the (noisy) observation or the measurement data

$$Y = \Theta S + \eta \quad (14)$$

where $\Theta (=:\Theta(X)) \in \mathbb{R}^{m \times n}$ with $k < m < n$ is the dictionary or sensing matrix and η is the measurement noise. In CS, we estimate a good sensing matrix Θ and a vector S that reconstructs the signal Y from X .

In general, this problem cannot be solved uniquely. However, if S is k -sparse i.e., if it has up to k nonzero entries, the theory of CS shows that it is possible to reconstruct Y uniquely, from m linear measurements even when $m \ll n$, by exploiting the sparsity of S . This can be achieved by finding the sparsest signal consistent with the vector of measurements (Donoho, 2006), i.e.

$$\arg \min_{S \in \mathbb{R}^n} \|S\|_0 \quad \text{subject to} \quad \|Y - \Theta S\|_2 \leq \varepsilon, \quad (15)$$

where $\|S\|_0$ denotes the l_0 norm for S (i.e., the number of non-zero entries of S), while ε denotes a parameter that depends on the level of measurement noise η . It can be shown that the l_0 minimization method can exactly reconstruct the original signal in the absence of noise using a properly chosen sensing matrix Θ whenever $m > 2k$. However, l_0 minimization problem (15) is non-convex and combinatorial, which is NP-hard.

Hence, instead of problem (15) we consider its l_1 convex relaxation which may be stated as (Chen et al., 2003)

$$\arg \min_{S \in \mathbb{R}^n} \|S\|_1 \quad \text{subject to} \quad \|Y - \Theta S\|_2 \leq \varepsilon, \quad (16)$$

where $\|\cdot\|_1$ represents the l_1 norm, which is a convex function and problem (16) is a convex optimization problem which can accurately approximate the signal X (solution to problem (15)) in polynomial time if the sensing matrix Θ is chosen to satisfy the necessary RIP condition with high probability (Candès et al., 2005; Candès and Wakin, 2006). Loosely speaking, if Θ satisfies the RIP, then the measurement matrix approximately preserves the Euclidean length of every signal. Equivalently, all subsets of k columns taken from Θ are nearly orthogonal. One should note that the l_1 minimization in (16) is closely related to the Lasso problem (Tibshirani, 1996).

$$\arg \min_{S \in \mathbb{R}^n} \frac{1}{2} \|Y - \Theta S\|_2^2 + \lambda \|S\|_1 \quad (17)$$

where $\lambda \geq 0$ is a regularization parameter. If ε and λ in (16) and (17) satisfy some special conditions, the two problems are equivalent; however, characterizing the relationships between ε and λ is difficult, except for the special case of orthogonal sensing matrices Θ .

Even though there are several objectives in compressed sensing, we are mainly interested in estimating a good sensing matrix and efficiently reconstructing the k -sparse vector S from the measurements. Hence, the RIP plays an important role in our calculations. Unfortunately, direct verification of RIP for a given matrix is not feasible. Nevertheless, it has been shown that there are some matrices, particularly random matrices with independent and identically distributed (i.i.d) entries, satisfy RIP with high probability.

We now enforce this sparse structure within learning. One of the related problems in the context of sparsity is “dictionary learning”. We refer to the columns of the $m \times n$ matrix Θ , as a dictionary, and they are assumed to be a set of vectors capable of providing a highly succinct representation for the signal. One has the freedom to determine the particular dictionary elements; however, the choice of monomial terms provides a model for the interactions between each of the components of the time-series and is used for model identification of dynamical systems in the general setting. In dictionary learning, which we shall use to determine Θ , the aim is to find a dictionary that can sparsely represent a given set of training data for further analysis –“sparsifying matrix” from a given set of training data. Regularized optimization with polynomial dictionaries is used to approximate the generating function of some unknown dynamic process. When the dictionary is large enough so that the “true” function is in the span of the candidate space, the solutions produced by sparse optimization are guaranteed to be accurate. This means that any of the training data should be presentable using linear combinations of few columns from the dictionary. In Lorenz-96, the sensing matrix Θ is formed as the collection of matrix of all monomials (stored column-wise) of resolved modes and due to the ergodic nature

of the simulated time series the entries satisfy RIP with high probability.

In this paper, the effects of unresolved processes (2) on the resolved modes (1) in the Lorenz-96 model are denoted by $u_k = -\frac{hc}{b} \sum_{j=1}^J Z_{j,k}$. The CS method takes snapshot data $\mathbf{x}(t) \in \mathbb{R}^n$ and attempts to discover the structure of a potentially highly nonlinear mapping of the form

$$\mathbf{u}(t) \approx \mathbf{f}(\mathbf{x}(t)) \quad (18)$$

where the n -dimensional vector $\mathbf{x}(t) = [x_1(t) \ x_2(t) \ \dots \ x_n(t)]^T$ represents the state of the system at time t . The error $\mathbf{e}(t) = \mathbf{u}(t) - \mathbf{f}(\mathbf{x}(t))$ at time t will be used later for auto-regression. The basis of this construction lies on the fact that in the multi scale Lorenz-96 model, the unresolved dynamics depends on the value of the resolved variables $\mathbf{x}$.

In particular, this method searches for a sparse approximation of the function mapping $\mathbf{f}(\mathbf{x}) = [f_1(\mathbf{x}) \ f_2(\mathbf{x}) \ \dots \ f_n(\mathbf{x})]^T \in \mathbb{R}^n$ using a dictionary of basis functions yielding

$$f_k(\mathbf{x}) \approx \sum_{j=1}^p \theta_j(\mathbf{x}) s_{jk} = \Theta(\mathbf{x}) \mathbf{s}_k \quad (19)$$

where $\Theta(\mathbf{x}) = [\theta_1(\mathbf{x}), \theta_2(\mathbf{x}), \dots, \theta_p(\mathbf{x})] \in \mathbb{R}^{1 \times p}$ form the dictionary of basis functions, and $\mathbf{s}_k = [s_{1k} \ s_{2k} \ \dots \ s_{pk}]^T \in \mathbb{R}^p$ is the vector of coefficients. Sparsity is expressed “not” in terms of an orthonormal basis but in terms of an over-complete dictionary. The majority of the coefficients s_{jk} are zero while the remaining nonzero entries identify the active terms contributing to the sparse representation of $\mathbf{f}(\mathbf{x})$.

To learn the function mapping $\mathbf{f}$ from data, time-histories of vectors $\mathbf{x}(t)$ and $\mathbf{u}(t)$ are collected by performing simulations for the Lorenz-96 system. The input data matrix $\mathbf{X} \in \mathbb{R}^{m \times n}$ is made of m snapshots of the state vector sampled at several times $t_1, t_2, \dots, t_m$ where subscripts 1, 2, ..., n indicates the elements of the state vector. The output data matrix $\mathbf{U} \in \mathbb{R}^{m \times n}$ follows the same arrangement, and these two data matrices are given as follows:

$$\begin{aligned} \mathbf{X} &= \begin{bmatrix} \mathbf{x}^T(t_1) \\ \mathbf{x}^T(t_2) \\ \vdots \\ \mathbf{x}^T(t_m) \end{bmatrix} = \begin{bmatrix} x_1(t_1) & x_2(t_1) & \dots & x_n(t_1) \\ x_1(t_2) & x_2(t_2) & \dots & x_n(t_2) \\ \vdots & \vdots & \ddots & \vdots \\ x_1(t_m) & x_2(t_m) & \dots & x_n(t_m) \end{bmatrix}, \\ \mathbf{U} &= \begin{bmatrix} \mathbf{u}^T(t_1) \\ \mathbf{u}^T(t_2) \\ \vdots \\ \mathbf{u}^T(t_m) \end{bmatrix} = \begin{bmatrix} u_1(t_1) & u_2(t_1) & \dots & u_n(t_1) \\ u_1(t_2) & u_2(t_2) & \dots & u_n(t_2) \\ \vdots & \vdots & \ddots & \vdots \\ u_1(t_m) & u_2(t_m) & \dots & u_n(t_m) \end{bmatrix}, \\ \mathbf{X}^2 &= \begin{bmatrix} x_1^2(t_1) & x_1(t_1)x_2(t_1) & \dots & x_1^2(t_1) & \dots & x_n^2(t_1) \\ x_1^2(t_2) & x_1(t_2)x_2(t_2) & \dots & x_2^2(t_2) & \dots & x_n^2(t_2) \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots \\ x_1^2(t_m) & x_1(t_m)x_2(t_m) & \dots & x_2^2(t_m) & \dots & x_n^2(t_m) \end{bmatrix} \quad (21) \end{aligned}$$

3.2.1 Dictionary of Basis Functions

First, we construct an augmented dictionary matrix $\Theta(\mathbf{X}) \in \mathbb{R}^{m \times p}$ where each column represents a basis function that can be a potential candidate for the terms in $\mathbf{f}$ to be discovered from data. p is the maximal number of n -multivariate monomials of degree at most d . As stated in (Brunton et al., 2016), there is huge flexibility in choosing the basis functions to populate the dictionary matrix $\Theta(\mathbf{X})$. For example, the dictionary matrix $\Theta(\mathbf{X})$ may consist of constant, polynomial and trigonometric terms as shown in (20).

$$\Theta(\mathbf{X}) = \begin{bmatrix} | & | & | & \dots & | \\ \mathbf{1} & \mathbf{X} & \mathbf{X}^2 & \dots & \mathbf{X}^d \\ | & | & | & \dots & | \end{bmatrix}_{m \times p} \quad (20)$$

The dictionary matrix $\Theta(\mathbf{X})$ is constructed by stacking together, column-wise, candidate nonlinear functions of $\mathbf{X}$. Here, higher order polynomials are denoted as $\mathbf{X}^d$ where d is the order of the polynomial considered. For example, element 1 is a column vector of ones, element $\mathbf{X}$ is as defined above, element $\mathbf{X}^2$ is the matrix containing the set of all quadratic polynomial functions of the state vector $\mathbf{x}$, and is constructed as shown in (21).

3.2.2 Sparse Optimization

Sparse regression is performed to approximately solve

$$\mathbf{U} \approx \Theta(\mathbf{X})\mathbf{S} \quad (22)$$

to determine the vector coefficients $\mathbf{S} = [\mathbf{s}_1 \ \mathbf{s}_2 \ \dots \ \mathbf{s}_n] \in \mathbb{R}^{p \times n}$ that determines the active terms in $\mathbf{f}$. Since only a few of the candidate functions are expected to have an effect on the function mapping $\mathbf{f}$, all coefficient vectors $\mathbf{s}_i$ are expected to be sparse. In this work, we use the Lasso algorithm (Tibshirani, 1996) for solving the above sparse regression problem.

The goal of the sparse optimization to solve the regression problem (22) can be defined using the Lasso form (17) for $k = 1, \dots, n$

$$\mathbf{s}_k = \underset{\mathbf{s}'_k}{\operatorname{argmin}} \|\mathbf{u}_k - \Theta(\mathbf{X})\mathbf{s}'_k\|_2^2 + \lambda \|\mathbf{s}'_k\|_1 \quad (23)$$

which uses l_1 norm constraint on the coefficient vector $\mathbf{s}_k$. The term $\lambda \|\mathbf{s}'_k\|_1$ is the regularization term, which penalizes coefficients different from 0 in a linear fashion. It is the actual promoter of sparsity in the minimization problem.

In the Lasso method, given a collection of m time snapshot pairs $(\mathbf{x}(t_i), \mathbf{u}(t_i))_{i=1}^m$, we estimate the coefficients in (22) by solving the following optimization problem for $k = 1, \dots, n$

$$\underset{\mathbf{s}_k}{\operatorname{minimize}} \sum_{i=1}^m \left(u_k(t_i) - \sum_{j=1}^p \theta_j(\mathbf{x}(t_i)) s_{jk} \right)^2 + \lambda \sum_{j=1}^p |s_{jk}| \quad (24)$$

for some Lagrange multiplier $\lambda \geq 0$. In this paper, we use $\lambda = 10^{-3}$.

4 STOCHASTIC PARAMETERIZATION: APPLICATION TO LORENZ-96

Having introduced the data-driven methods for model reduction in the presence of time-scale separation, let us now turn to a simple atmospheric model, which nonetheless exhibits many of the difficulties arising in realistic models, to gain insight into predictability and data assimilation. The dynamics of the Lorenz-96 equation are shown in (1) and (2). Let X_k and $Z_{j,k}$ be the numerical solutions of the Lorenz-96 equations, generated using the 4th order Runge-Kutta method, for $k = 1, \dots, K, j = 1, \dots, J$. We define U_k as in Equation (4):

$$U_k = -\frac{hc}{b} \sum_{j=1}^J Z_{j,k}$$

The resolved modes in (1) is now modified as:

$$\dot{X}_k = -X_{k-1}(X_{k-2} - X_{k+1}) - X_k + F + U_k = \mathcal{G}_k(X) + U_k \quad (25)$$

Our goal is to use the methods discussed in Section 3 to form parameterizations, in other words, express U_k as a function of $X_1, \dots, X_K$ and compare each of them. Based on the ergodic properties of the Lorenz-96 system shown in Section 2.1, the Lyapunov exponents stabilize at $t = 500$. The training set is set to $t_{\text{train}} = [500, 500 + h, 500 + 2h, \dots, 1000]$ where $h = 0.01$. The test set is set to $t_{\text{test}} = [1000, 1000 + h, 1000 + 2h, \dots, 2000]$.

Using polynomial regression to solve stochastic parameterizations in the Lorenz-96 system was proposed by Wilks (2005). The goal of this method is to find a degree-4 polynomial P such that U_k can be approximated by $P(X_k)$ for $k = 1, \dots, K$. This requires solving a least squares problem:

$$\sum_{k=1}^K \sum_{i \in t_{\text{train}}} (U_k(t_i) - P(X_k(t_i)))^2 \quad (26)$$

The polynomial obtained from this method is:

$$P(X_k) = -0.0003142X_k^4 + 0.004882X_k^3 + 0.005197X_k^2 - 0.47362X_k - 0.03779 \quad (27)$$

The goal of compressed sensing in this context is to find a sparse matrix S that solves for $\mathbf{Y} = \Theta\mathbf{S}$. Let $\mathbf{Y}$ be the matrix with columns $U_1, \dots, U_K$, each divided by its L_2 norm. Every column of Θ is a basis function of $X_1, \dots, X_K$ divided by its L_2 norm. The goal of compressed sensing is to find unique equations for every U_k for $k = 1, \dots, K$. Since regression gave us a polynomial of degree 4, our goal is to express U_k as a function of X values of order 4. However, with $K = 40$, this

means the Θ matrix will have 135750 columns, and thus, compressed sensing will be computationally expensive. Each of the parameterizations we get are shown in the “CompressedSensingEquations” folder in (Mukherjee et al., 2021).

A computationally less expensive approach will be to find a solution of order 2. The 40-dimensional CS model is used for comparison of distributions, but we present below the single equation with averaged coefficients that will be used later for filtering:

$$f_k(\mathbf{X}) = -0.376X_k + 0.0166X_k^2 \quad (28)$$

This parameterization shows that every f_k is expressed as a function of X_k and X_k^2 , thus showing that every basis function $X_i X_j$ where $i \neq j$ is not relevant, and the Θ matrix can be reduced to a matrix with X_k and X_k^2 as the columns.

We then explore solutions of order 3. Given $K = 40$, X^{P_3} has 11480 columns, thus finding a solution of order 3 is computationally expensive. However, creating multiple parameterizations, each with a subset of columns of X^{P_3} is computationally less expensive. Thus, in this section, we create 100 parameterizations, each of which the Θ matrix contains X_k^3 and 100 random columns from X^{P_3} . If we average the coefficients of the model, we get:

$$f_k(\mathbf{X}) = -0.0338X_k + 0.00452X_k^2 + 0.00142X_k^3 \quad (29)$$

This shows that every f_k is expressed as a function of X_k , X_k^2 and X_k^3 , thus showing that all remaining columns in X^{P_3} are not relevant. Based on the results of the previous models, we expect $X_1^4, \dots, X_k^4$ to be the only relevant columns from X^{P_4} . This information captures the underlying structure in the Lorenz-96 equation since, in equation (2), $\dot{z}_{j,k}$ only depends on X_k .

Thus, our Θ matrix contains the basis functions $1, X_k, X_k^2, X_k^3, X_k^4$ for $k = 1, \dots, K$. We will call this parameterization the raw degree-4 compressed sensing model. Since this algorithm will be used for testing purposes, the compressed sensing algorithm also calculates biases for each f_k . If we average the coefficients, we get:

$$f_k(\mathbf{X}) = -0.416X_k + 0.00277X_k^2 + 0.00509X_k^3 - 0.000306X_k^4 - 0.0243 \quad (30)$$

The biases are calculated by the compressed sensing algorithm as [2.173746, 1.393404, ..., -2.973427] and their mean is -0.0243. The problem with this parameterization is that there is too much variation in the biases.

We investigate other ways to modify the biases. In the following model, the biases are strictly set to zero while keeping other coefficients constant. If we average the coefficients, we get:

$$f_k(\mathbf{X}) = -0.474X_k + 0.00277X_k^2 + 0.00509X_k^3 - 0.000306X_k^4 \quad (31)$$

In the next model, the biases are strictly set to the average of the biases while keeping other coefficients constant. If we average the coefficients, we get the same equation as (30). Both of these models perform equally well in predicting the trajectory of the system and significantly better than the raw 4-dimensional model. The model where biases are strictly set to the average of the biases performs better in the Kullback-Leibler divergence measure, but does not outperform linear regression.

A small variability was introduced in the biases with the goal of outperforming linear regression in terms of the Kullback-Leibler divergence measure. The biases were selected by sampling from $N(\mu, \sigma^2)$, where μ is the average of the biases ($\mu = -0.0243$). We used $\sigma = 0.07$ as it gave the lowest prediction error on the training set. This leads to a new model, which outperforms all other models. The model is shown in Appendix B. If we average the coefficients, we get:

$$f_k(\mathbf{X}) = -0.474X_k + 0.00277X_k^2 + 0.00509X_k^3 - 0.000306X_k^4 - 0.0241 \quad (32)$$

We will use (32) as our compressed sensing model in the later sections.

The equations of the resolved variables associated with the vector field denoted by $\mathcal{G}_k(\mathbf{X})$ are equivariant with respect to a cyclic permutation of the variables. Therefore, the system with dimension K is completely determined by the equation for the k^{th} variable. Even though the unresolved dynamics depend on the value of the resolved variables $\mathbf{X}$, while computing the parameterization of U_k by CS, there were no restrictions placed on the form of the polynomial on the sparse regression. The consequences of $\Theta(\mathbf{X})$ satisfying RIP (hence, sparsity) are quite profound as shown in results provided in Appendix B. Compared to regression, compressed sensing also has the advantage that each unresolved variable U_k is given a unique equation that best represents its dynamics as a function of resolved variables.

4.1 COMPARISON OF TRAJECTORIES

Section 2.1.2 shows that the largest Lyapunov exponent is 2.3. This shows that the error-doubling time is $\frac{\ln(2)}{2+2.3} \approx 0.3$. To show the practicality of each method applied to Lorenz-96, we predict the trajectory of each X_k for a time interval of three times the error doubling time (0.9) after the end of the training set. The measure of inaccuracy used in this paper is the Mean Squared Prediction Error (MSPE) (Pindyck and Rubinfeld, 1991).

$$MSPE = \frac{1}{K} \sum_{k=1}^K \frac{1}{I} \sum_{i \in I} (\hat{X}_k(t_i) - X_k(t_i))^2 \quad (33)$$

where X_k is derived from (1) and (2) and $\hat{X}_k$ is the approximation to X_k derived from a parametrized model (25) through the 4th order Runge-Kutta method, I is the time interval where we compare the trajectories. In this case, $I = [1000, 1000 + h, 1000 + 2h, \dots, 1000 + 0.9]$ where $h = 0.01$. The initial conditions are set as: $\hat{X}_k(1000) = X_k(1000)$, $k = 1, \dots, K$. The MSPEs derived for each method are shown in [Table 2](#).

4.2 COMPARISON OF PROBABILITY DENSITY FUNCTIONS

This section will compare the probability density functions (PDFs) of $\hat{X}_k$ and X_k for $k = 1, \dots, K$ during the time interval t_{test} . The PDFs are computed using the kernel density estimation, given in (8). PDFs are computed for every X_k and $\hat{X}_k$ using their data in the time interval [500,2000].

The Kullback-Leibler divergence was introduced by Kullback and Leibler (1951) and measures the discrepancy between two PDFs. The Kullback-Leibler divergence is calculated as

$$D_{KL}(P||Q) = \sum_{x \in X} P(x) \log\left(\frac{P(x)}{Q(x)}\right) \quad (34)$$

where P is the PDF calculated from the data of X_k and Q is the PDF calculated from $\hat{X}_k$. We average the Kullback-Leibler divergence in the PDFs of $\hat{X}_k$ and X_k for all k to compare each model. The average Kullback-Leibler Divergence derived for each method is shown in [Table 2](#).

4.3 AUTO-REGRESSIVE MODELS IN RESIDUALS

So far, we have investigated the deterministic part in stochastic parameterization. The goal of this subsection is to investigate the stochastic part in stochastic parameterization to further improve the regression and compressed sensing models explored so far. This is done by obtaining φ , σ and σ_e from the first-order AR model in

(13). The method of calculating these values is explained in Appendix A. Throughout this subsection, we will refer to the compressed sensing model with noise added to the average bias as the compressed sensing model.

The values for φ , σ and σ_e that we obtain are shown in [Table 3](#). Given the values of φ and σ , we can run a numerical simulation of the reduced dynamical system with AR residuals. The equation is shown below:

$$\dot{X}_k = -X_{k-1}(X_{k-2} - X_{k+1}) - X_k + F - f_k(\mathbf{X}(t)) + e_k(t_i) \quad (35)$$

where $f_k(\mathbf{X}(t))$ represents either the Wilks' parameterization (27) or the CS parameterization (32) and $e_k(t_i)$ is generated using the first order AR model below:

$$\begin{aligned} e_k(t_i) &= \varphi e_k(t_{i-1}) + \sigma z_k(t_i) \\ &= \varphi e_k(t_{i-1}) + \sigma_e(1 - \varphi^2)^{1/2} z_k(t_i) \end{aligned} \quad (36)$$

$z_k(t_i)$ is sampled randomly from $N(0, 1)$.

Numerical simulations are done using the 4th order Runge-Kutta method. $e_k(t_0)$ is set to 0 for all k . Similar to section 4.2, we evaluate our model by assessing its trajectory as well as its PDF. The results are summarized in [Table 3](#).

A comparison of [Table 3](#) with [Table 2](#) shows that using an AR model in compressed sensing and regression has significantly improved the prediction of the trajectory of resolved variables as the MSPE is lower. Compressed sensing shows the lowest MSPE and Kullback-Leibler divergence.

The goal of this section was to create a deterministic and AR model to improve the prediction of resolved variables. Based on the results shown, compressed sensing improves the prediction significantly compared to regression. The coefficients of the deterministic model as well as the values of φ , σ and σ_e will be useful in the following section where filtering methods will be used to further improve the prediction of resolved variables.

METHOD	MSPE	AVERAGE K-L DIVERGENCE
Regression	0.03606	0.06729
Compressed sensing (Raw)	5.9993	1.9891
Compressed sensing [1] (Setting biases to zero)	0.03398	0.10099
Compressed sensing [2] (Setting biases to average bias)	0.03440	0.07351
Compressed sensing [3] (Adding noise to average bias)	0.03387	0.05919

Table 2 Evaluation of each method.

METHOD	φ	σ	σ_e	MSPE	AVE. K-L DIVERGENCE
Regression	0.9453	0.2265	0.6945	0.02624	0.07692
Compressed sensing	0.9981	0.1710	2.775	0.02066	0.07142

Table 3 AR evaluation of residuals.

5 NONLINEAR FILTERING APPLIED TO PARAMETERIZED MODELS

The second component of this paper deals with nonlinear filtering or data assimilation (DA). The main focus of filtering is to combine computational models with sensor data (observations) to predict the dynamics of large-scale evolving systems. Filtering provides a recursive algorithm for estimating a signal or state of a random dynamical system based on noisy measurements. The signal that is represented by a Markov process cannot be accessed or observed directly and is to be “filtered” from the trajectory of the observation process which is statistically related to the signal. Suppose we make a forecast about the behavior at a future time of a complex system with some uncertainties (randomness) and there is near-continuous data available from remote sensing instrumentation networks. Then, as new information becomes available through observations, it is natural to ask how to best incorporate this new information into the prediction. Data assimilation can be simply defined as the combination of observations with the model output to produce an optimal estimate of a dynamical model.

The standard notation for the discrete DA problems from the time t_k to t_{k+1} can be formulated as follows:

$$x^f(t_{k+1}) = M[x^f(t_k)] + w_k \quad (37)$$

where x is the model state, M is the model dynamics and w is the model error. The operator M is represented with matrix in discrete case and differential operator in continuous case. In this section, for notational simplicity, x represents the model state $\mathbf{X}$ as described in previous sections. Observations at time t_k can be defined by:

$$y_k = H[x^f(t_k)] + v_k \quad (38)$$

where y is the observation and H is the observation operator, which is identity $\mathbb{I}$ in our problem. v is the white noise with associated covariance matrix describing the observation error.

Data assimilation methods consist of *analysis* (or correction/filtering) and *forecast* (or prediction) steps. At time t_k , we have the output state of the background forecast x_k^f , and observation state y_k . x^t is the truth generated by equation 1 and 2. The (linear or nonlinear) analysis step is based on the outputs of x_k^f and y_k and produces analysis x_k^a . Then, applying the model dynamics on the analysis step generates the next forecast output as x_{k+1}^f .

5.1 ENSEMBLE KALMAN FILTER

It is well documented that particle filters (Gordon et al., 1993) suffer from the well-known problem of sample degeneracy (Bengtsson et al., 2008; Snyder et al., 2008). Recently, (Yeong et al., 2020) presented reduced-order nonlinear filtering schemes based on a theoretical

framework that combined stochastic parametrization based on homogenization and nonlinear filtering for the implementation of particle filters. In contrast, the ensemble Kalman filter (EnKF) (Evensen, 1994; Evensen, 2003) can handle some of these difficulties, where the dimensions of states and observations are large and the number of replicates is small. However, EnKF incorrectly treats the non-Gaussian features of the forecast distribution that arise in certain nonlinear systems. Nevertheless, we apply the ensemble Kalman filter for data assimilation, because they are easy to implement in large systems.

For the Kalman filter (Kalman, 1960), the modelling and observation errors are assumed to be independent with Gaussian/normal probability distributions:

$$w_k \sim N(0, Q_k) \text{ and } v_k \sim N(0, R_k) \quad (39)$$

where Q and R are the covariance matrices of the model error and observation error, respectively. The Kalman filter is of the form:

$$x_k^a = x_k^f + K_k(y_k - H_k x_k^f) \quad (40)$$

where K_k is known as the *Kalman gain* and $(y_k - H_k x_k^f)$ is called the *innovation*. After some manipulations, the Kalman gain can be obtained as:

$$K_k = P_k^f H^T (H P_k^f H^T + R)^{-1} \quad (41)$$

where P_k^f is the forecast error covariance matrix.

The idea behind EnKF (Evensen, 2003) is the propagation and analysis of ensemble members instead of full-rank covariance matrices. Unlike Kalman filter, the EnKF is applicable for nonlinear dynamical models and computationally efficient for high-dimensional systems (Asch et al., 2016; Law et al., 2015). The forecast step of EnKF can be formulated as follows:

$$\begin{aligned} \hat{x}_{j+1}^n &= \Psi(x_j^n) + e_j^n \\ \hat{m}_{j+1} &= \frac{1}{N} \sum_{n=1}^N \hat{x}_{j+1}^n \\ \hat{C}_{j+1} &= \frac{1}{N-1} \sum_{n=1}^N (\hat{x}_{j+1}^{(n)} - \hat{m}_{j+1})(\hat{x}_{j+1}^{(n)} - \hat{m}_{j+1})^T \end{aligned} \quad (42)$$

where x is the model state that is propagated by Ψ and e_j is the auto-regressive modeling of the noise as described before. m is the model state mean, C is the model covariance matrix and $n = 1, \dots, N$ is the number of ensembles. Then, the analysis step can be performed as follows:

$$\begin{aligned} S_{j+1} &= H \hat{C}_{j+1} H^T + \Gamma \\ K_{j+1} &= \hat{C}_{j+1} H^T S_{j+1}^{-1} \\ x_{j+1}^n &= (I - K_{j+1} H) \hat{x}_{j+1}^{(n)} + K_{j+1} y_{j+1}^{(n)} \\ y_{j+1}^{(n)} &= y_{j+1} + \eta_{j+1}^n, \end{aligned} \quad (43)$$

where Γ and η are i.i.d's draws from normal distribution and y_t 's are the observation states. In our problem, the truth is generated by the slow variables of non-parametric Lorenz-96 system, which is given by (1) and (2). Then, observations (direct measurements) are produced by adding small sensor noise to the slow variables X . And, model used for the propagation of ensemble Kalman filter is either given by compressed sensing or polynomial regression.

5.2 APPLICATION TO LORENZ-96

In this section, we compare the filtering results of Wilks' (polynomial) parameterization and compressed sensing parameterization. For each filtering implementation, 40 ensembles are used and observations are included in every 20 time steps. We observe that compressed sensing dynamics provides better prediction results compared to other parameterization methods. There are several software that can be used for filtering problems. Here, we benefit from the open source provided in (Pawar and San, 2020) regarding our parameterized Lorenz-96 models. Observations are generated by adding normal noise to the non-parameterized Lorenz-96 model (truth) and filtering trajectories are obtained by using the set of equations given in (42) and (43) for all cases.

5.2.1 Autoregressive Wilks' parameterization

Here, we analyze filtering results of Wilks' parameterization with auto-regressive models. Instead of using random noise for the prediction, we use the auto-regressive modelling of noise as obtained in Section 4. While keeping the parametrized model same, we include modelled noise to the ensemble dynamics as the following:

$$\begin{aligned} \dot{X}_k = & \mathcal{G}_k(X) - 0.03779 - 0.47362X_k + 0.005197X_k^2 \\ & + 0.004882X_k^3 - 0.0003142X_k^4 + e_k(t) \end{aligned} \quad (44)$$

where $\mathcal{G}_k(X)$ is defined in (3) and

$$e_k(t_i) = 0.9453e_k(t_{i-1}) + 0.2265z_k(t_i) \quad (45)$$

$z_k(t_i)$ is sampled randomly from $N(0, 1)$.

Figure 5 shows the trajectories of truth and observations generated from the truth for the first 3 components of Lorenz-96 model. **Figure 7** depicts the truth and prediction (filtering) results of 40 components and MSPE is 0.00940 which is slightly better than deterministic parameterization.

5.2.2 Autoregressive Compressed Sensing

In addition to compressed sensing method in Section 4, we consider auto-regressive modelling of noise. While keeping the parametrized model same, we include modelled noise to the ensemble dynamics as the following:

$$\begin{aligned} \dot{X}_k = & \mathcal{G}_k(X) - 0.0241 - 0.474X_k + 0.00277X_k^2 \\ & + 0.00509X_k^3 - 0.000306X_k^4 + e_k(t) \end{aligned} \quad (46)$$

where $\mathcal{G}_k(X)$ is defined in (3) and

$$e_k(t_i) = 0.9981e_k(t_{i-1}) + 0.1710z_k(t_i) \quad (47)$$

$z_k(t_i)$ is sampled randomly from $N(0, 1)$. **Figure 6** shows the trajectories of truth and observations generated from the truth for the first 3 components of Lorenz-96 model. **Figure 8** depicts the truth and prediction (filtering) results of 40 components and MSPE is 0.00940 which is slightly better than deterministic parameterization.

In **Figures 7** to **8**, we see the depictions of filtering results for 40 slow components of Lorenz96 model. The first plot of each figure represent the true trajectories of components where the observations are generated. The second plot of each figure show the filtering (prediction) results for the system. Finally, the third plot of each figure

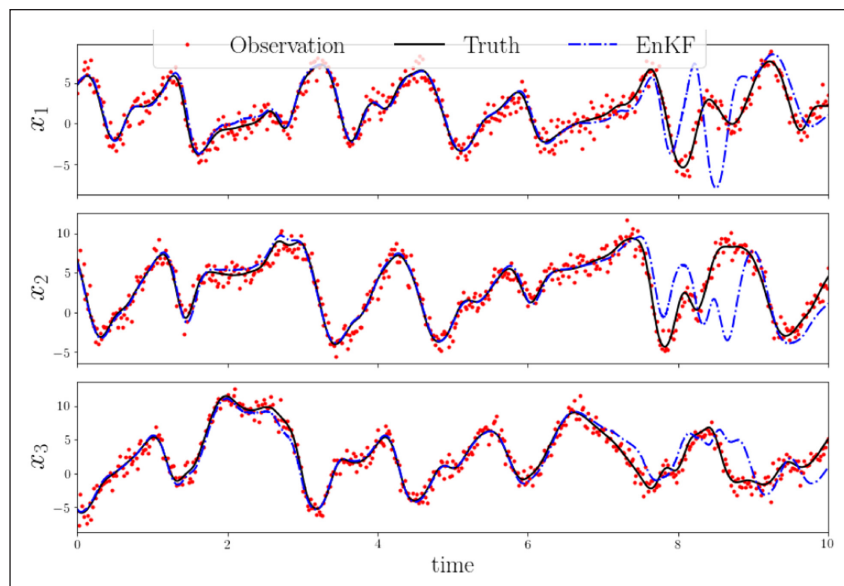


Figure 5 True trajectories, observations and EnKF with Wilks' parametrized model.

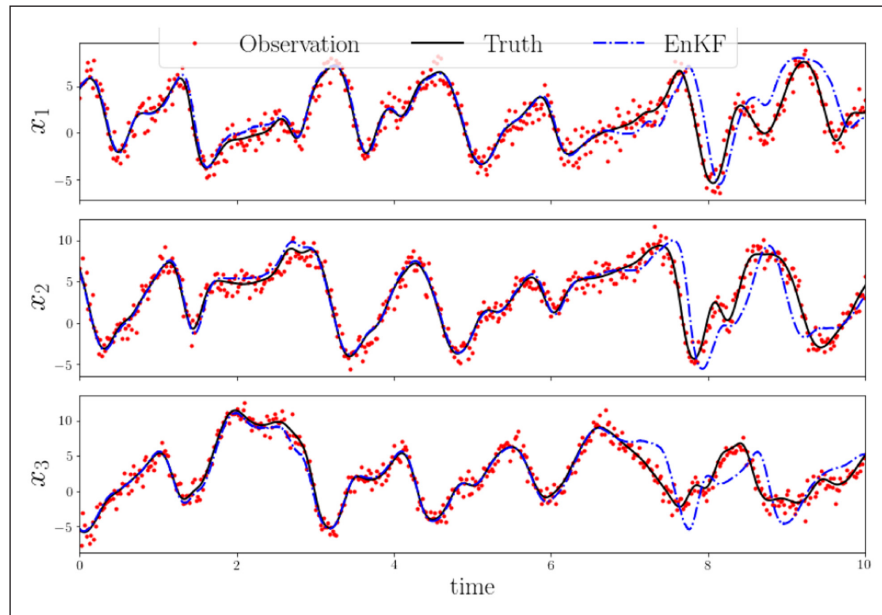


Figure 6 True trajectories, observations and EnKF with compressed sensing model.

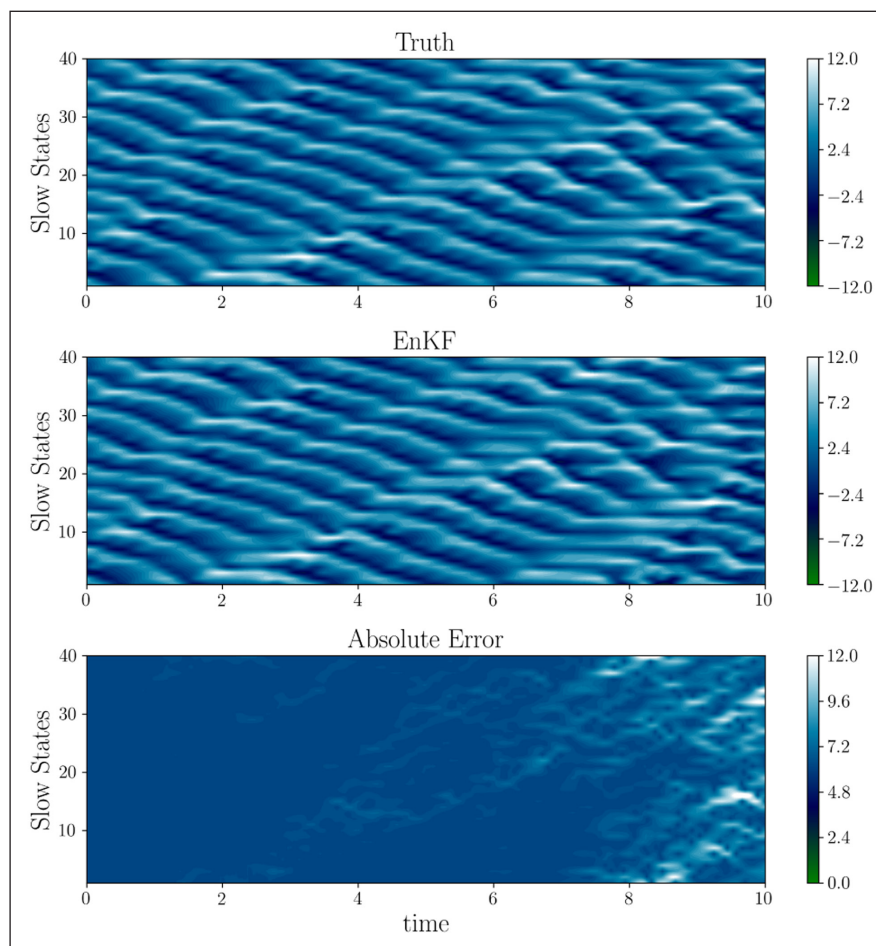


Figure 7 Depiction of Lorenz-96 true and EnKF prediction trajectories for 40 components (Auto-regressive Wilks' parameterization).

depict the L1 errors between the truth and prediction results. All the results about the figures shown in this section are discussed in the previous sections related to the individual parameterization methods.

Table 4 shows the parameterization methods with the MSPE errors. As a conclusion, stochastic parameterization methods provide slightly better prediction results compared to deterministic parameterizations.

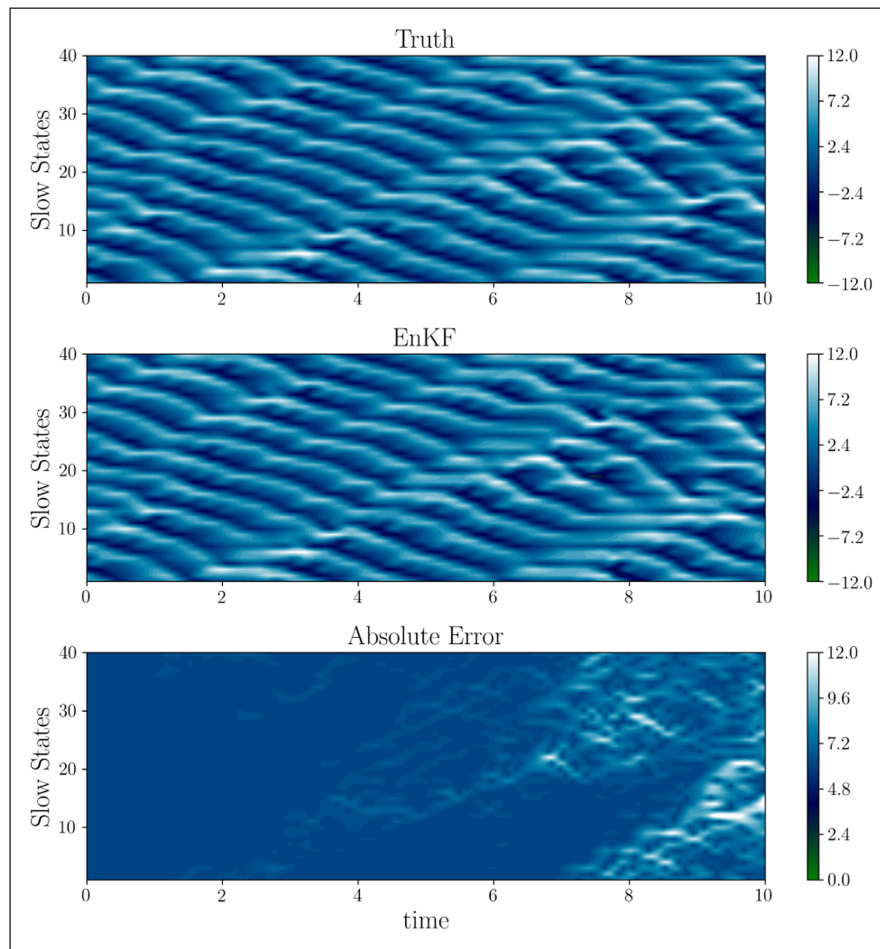


Figure 8 Depiction of Lorenz-96 true and EnKF prediction trajectories for 40 components (Auto-regressive compressed sensing).

METHOD	MSPE
Wilks' parameterization	0.01006
Compressed sensing	0.01005
Autoregressive Wilks' parameterization	0.00940
Autoregressive compressed sensing	0.00940

Table 4 Parameterization methods and MSPE results including random noise and modelled noise with autoregression.

6 CONCLUSION

The sheer number of calculations required in directly solving large-scale global weather and climate models becomes computationally overwhelming. Hence, in this paper, using the multi-scale Lorenz-96 system as a test problem, we introduced a data-driven stochastic-parameterization framework in which the equations of small-scales are integrated and incorporated using compressed sensing to reduce the cost of data assimilation. We presented a data-driven, SPR technique, where we learned to describe the corresponding unknown or unresolved physics from the data due to the fact the governing equations of the unresolved physics are too computationally expensive to solve. Since the

atmospheric model considered in this paper is inherently chaotic, it has several positive Lyapunov exponents and any small error in the initial conditions grow in finite time.

Random matrices with independent and identically distributed (i.i.d) entries, satisfy RIP with high probability, and the RIP plays an important role in our calculations. We provided a characterization of the dynamical system based on the Lyapunov exponents for the $F = 10$ case. In this case, the model is strongly chaotic and decorrelation times also suggested the model to be ergodic. Since the dictionary Θ was formed as the collection of matrices of all monomials (stored column-wise) of resolved modes and due to the ergodic nature of the simulated time series, the entries satisfied the RIP with high probability. The construction of the optimization problem and computational method were detailed in section 4.

For parameterization of U_k , while previous works mostly focused on suggesting pre-defined polynomial or other model form (or structure) and then estimating parameters, this work is the first attempt to use CS for *automatically* and simultaneously discovering polynomial or model form (or structure) and estimating parameters for U_k as illustrated by the results in section 4 and Appendix B.

Specifically, a reduced-order ensemble Kalman filter (EnKF) was applied to the stochastically parametrized model. For the EnKF, particle locations were corrected

based on their respective distances from the observation (innovation). Magnitude of the correction was proportional to the error covariance and inverse to the observation noise covariance (Kalman gain). In between observations, particles in the EnKF sample evolve according to the original signal dynamics, so particle trajectories are expected to deviate from the truth as time moves further away from the last observation correction.

Our results show that the compressed sensing model produces the best deterministic parameterization results. After incorporating auto-regression to model residuals, the stochastic parameterizations outperform deterministic parameterizations in terms of prediction accuracy. Compressed sensing with auto-regression produces the best stochastic parameterization results.

Solving for small-scale processes with help of LES models, on a high-resolution grid embedded (with grid spacing of a few tens of meters and can explicitly simulate the clouds and turbulence represented by parameterization schemes) within the low-resolution grid is very useful, but its applicability remains limited to its high computational cost. A long term goal is to create a data-driven scheme, in which LES models embedded in selected grid columns of a global model explicitly simulate subgrid-scale processes which are represented by parameterization schemes in the other columns. The result is a seamless algorithm which only uses a small scale grid on a few selected columns that simulates interactively within a running global simulation, and is far cheaper than eddy-resolving simulation throughout the global simulation.

ADDITIONAL FILE

The additional file for this article can be found as follows:

- **Appendix.** Appendix A and B. DOI: <https://doi.org/10.16993/tellusa.42.s1>

ACKNOWLEDGEMENTS

The authors acknowledge partial support for this work from Natural Sciences and Engineering Research Council (NSERC) Discovery grant 50503-10802, TECSIS/Fields-CQAM Laboratory for Inference and Prediction and NSERC-CRD grant 543433-19. The authors would like to thank Professor Peter Baxendale, Department of Mathematics, University of Southern California, for stimulating discussions, contributions and suggestions on many topics of this paper.

COMPETING INTERESTS

The authors have no competing interests to declare.

AUTHOR AFFILIATIONS

A. Mukherjee  orcid.org/0000-0001-9962-8110
Department of Applied Mathematics, University of Waterloo,
Waterloo, Ontario N2L 3G1, Canada

Y. Aydogdu  orcid.org/0000-0002-2836-141X
Department of Applied Mathematics, University of Waterloo,
Waterloo, Ontario N2L 3G1, Canada

T. Ravichandran  orcid.org/0000-0002-3579-2832
Department of Applied Mathematics, University of Waterloo,
Waterloo, Ontario N2L 3G1, Canada

N. Sri Namachchivaya  orcid.org/0000-0002-4444-7344
Department of Applied Mathematics, University of Waterloo,
Waterloo, Ontario N2L 3G1, Canada

REFERENCES

- Anosov, DV.** 1969. Geodesic flows on closed Riemann manifold with negative curvature. *Volume 90 of Trudy ordena Lenina Matematicheskogo Instituta imeni V.A. Steklova.* American Mathematical Society.
- Arnold, HM, Moroz, IM and Palmer, TN.** 2013. Stochastic parametrizations and model uncertainty in the Lorenz '96 system *Philosophic Transactions of the Royal Society A*, 371: 1991. DOI: <https://doi.org/10.1098/rsta.2011.0479>
- Asch, M, Bocquet, M and Nodet, M.** 2016. Data assimilation: methods, algorithms, and applications. *Society for Industrial and Applied Mathematics.* DOI: <https://doi.org/10.1137/1.9781611974546>
- Beeson, R and Sri Namachchivaya, N.** 2020. Particle filtering for chaotic dynamical systems using future right-singular vectors. *Nonlinear Dyn*, 102: 679–696. DOI: <https://doi.org/10.1007/s11071-020-05727-y>
- Bengtsson, T, Bickel, P and Li, B.** 2008. Curse-of-dimensionality revisited: Collapse of the particle filter in very large scale systems. *IMS Collections: Probability and Statistics Essays in Honor of David A. Freedman*, 2: 316–334. DOI: <https://doi.org/10.1214/193940307000000518>
- Bishop, CM.** 2006. *Pattern Recognition and Machine Learning.* Springer.
- Brunton, SL, Proctor, JL and Kutz, JN.** 2016. Discovering governing equations from data by sparse identification of nonlinear dynamical systems. *Proceedings of the National Academy of Sciences*, 113(15): 3932–3937. DOI: <https://doi.org/10.1073/pnas.1517384113>
- Candès, EJ, Romberg, JK and Tao, T.** 2005. Decoding by linear programming. *IEEE Transactions on information theory*, 51(12): 4203–4215. DOI: <https://doi.org/10.1109/TIT.2005.858979>
- Candès, EJ and Tao, T.** 2006. Near-optimal signal recovery from random projections: Universal encoding strategies. *IEEE Transactions on Information Theory*, 52(12): 5406–5425. DOI: <https://doi.org/10.1109/TIT.2006.885507>
- Candès, EJ and Wakin, MB.** 2006. An introduction to compressed sampling. *IEEE signal processing magazine*, 25(2): 21–30. DOI: <https://doi.org/10.1109/MSP.2007.914731>

- Chen, SS, Donoho, DL and Saunders, MA.** 2003. Atomic decomposition by basis pursuit. *SIAM Journal on Scientific Computing*, 20(1): 33–61. DOI: <https://doi.org/10.1137/S1064827596304010>
- Devaney, RL.** 1989. *An Introduction to Chaotic Dynamical Systems*. Addison-Wesley.
- Donoho, DL.** 2006. Compressed sensing. *IEEE Transactions on information theory*, 52(4): 1289–1306. DOI: <https://doi.org/10.1109/TIT.2006.871582>
- Eckmann, P and Ruelle, D.** 1985. Ergodic theory of chaos and strange attractors. *Rev. Mod. Phys.*, 57(3): 617–656. DOI: <https://doi.org/10.1103/RevModPhys.57.617>
- Evensen, G.** 1994. Sequential data assimilation with a nonlinear quasi-geostrophic model using Monte Carlo methods to forecast error statistics. *Journal of Geophysical Research*, 99: 10143–10162. DOI: <https://doi.org/10.1029/94JC00572>
- Evensen, G.** 2003. The ensemble Kalman filter: Theoretical formulation and practical implementation. *Ocean dynamics*, 53(4): 343–367. DOI: <https://doi.org/10.1007/s10236-003-0036-9>
- Gagne, DJ, Christensen, H, Subramanian, A and Monahan, AH.** 2020. Machine learning for stochastic parameterization: Generative adversarial networks in the Lorenz '96 model. *Journal of Advances in Modeling Earth Systems*, 12(3). DOI: <https://doi.org/10.1029/2019MS001896>
- Gordon, NJ, Salmond, DJ and Smith, AFM.** 1993. Novel approach to nonlinear/non-gaussian Bayesian state estimation. *IEEE Proceedings*, F(2): 107–113. DOI: <https://doi.org/10.1049/ip-f-2.1993.0015>
- Kalman, RE.** 1960. A new approach to linear filtering and prediction problems. DOI: <https://doi.org/10.1115/1.3662552>
- Kullback, S and Leibler, RA.** 1951. On Information and Sufficiency. *The Annals of Mathematical Statistics*, 22: 79–86. DOI: <https://doi.org/10.1214/aoms/1177729694>
- Law, K, Stuart, A and Zygalakis, K.** 2015. *Data assimilation*. 214. Switzerland, Cham: Springer. DOI: <https://doi.org/10.1007/978-3-319-20325-6>
- Lorenz, EN.** 1995. Predictability: A problem partly solved. *ECMWF Seminar Proceedings on Predictability*, ECMWF.
- Majda, A, Abramov, R and Grote, M.** 2005. *Information Theory and Stochastic for Multiscale Nonlinear Systems* (CRM Monograph Series), 25. New York: American Mathematical Society. DOI: <https://doi.org/10.1090/crmm/025>
- Mukherjee, A, Aydogdu, Y and Ravichandran, T.** 2021. Stochastic Parameterization using Compressed Sensing: Application to the Lorenz-96 Atmospheric Model. *GitHub Repository*. <https://github.com/amartyamukherjee/StochasticParameterization-ApplicationToLorenz96>.
- Pawar, ASE and San, O.** 2020. A Hands-On Introduction to Dynamical Data Assimilation with Python. *Fluids*, 5(4): 225. DOI: <https://doi.org/10.3390/fluids5040225>
- Pesin, YB.** 2004. Lectures on partial hyperbolicity and stable ergodicity. *European Mathematical Society*. DOI: <https://doi.org/10.4171/003>
- Pindyck, RS and Rubinfeld, DL.** 1991. Forecasting with Time-Series Models. *Econometric Models & Economic Forecasts* 3rd ed., 516–535.
- Sinai, YG.** 1977. Introduction to ergodic theory. *Volume 18 of Mathematical Notes*. Princeton University Press.
- Skokos, Ch.** 2010. The Lyapunov Characteristic Exponents and Their Computation. *Lect. Notes Phys.*, 790: 63–135. DOI: https://doi.org/10.1007/978-3-642-04458-8_2
- Snyder, C, Bengtsson, T, Bickel, P and Anderson, J.** 2008. Obstacles to high-dimensional particle filtering. *Monthly Weather Review*, 136(12): 4629–4640. DOI: <https://doi.org/10.1175/2008MWR2529.1>
- Tibshirani, R.** 1996. Regression Shrinkage and Selection via the Lasso. *Journal of the Royal Statistical Society. Series B (Methodological)*, 58(1): 267–288. DOI: <https://doi.org/10.1111/j.2517-6161.1996.tb02080.x>
- Wilks, DS.** 2005. Effects of stochastic parameterizations in the Lorenz '96 system. *Royal Meteorological Society*, 131: 389–407. DOI: <https://doi.org/10.1256/qj.04.03>
- Wolf, AS, Swinney, JB and Vastano, HL.** 1985. Determining Lyapunov exponents from a time series. *Physica*, 16D: 285–317. DOI: [https://doi.org/10.1016/0167-2789\(85\)90011-9](https://doi.org/10.1016/0167-2789(85)90011-9)
- Yeong, HC, Beeson, R, Sri Namachchivaya, N and Perkowski, N.** 2020. Particle Filters with Nudging in Multiscale Chaotic Systems: with Application to the Lorenz-96 Atmospheric Model. *Journal of Nonlinear Science*, 30: 1519–1552. DOI: <https://doi.org/10.1007/s00332-020-09616-x>
- Young, LS.** 1998. Statistical properties of dynamical systems with some hyperbolicity. *The Annals of Mathematics*, 147(3): 585–650. DOI: <https://doi.org/10.2307/120960>

TO CITE THIS ARTICLE:

Mukherjee, A, Aydogdu, Y, Ravichandran, T and Sri Namachchivaya, N. 2022. Stochastic Parameterization Using Compressed Sensing: Application to the Lorenz-96 Atmospheric Model. *Tellus A: Dynamic Meteorology and Oceanography*, 74(2022): 300–317. DOI: <https://doi.org/10.16993/tellusa.42>

Submitted: 21 February 2022 Accepted: 21 February 2022 Published: 26 April 2022

COPYRIGHT:

© 2022 The Author(s). This is an open-access article distributed under the terms of the Creative Commons Attribution 4.0 International License (CC-BY 4.0), which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited. See <http://creativecommons.org/licenses/by/4.0/>.

Tellus A: Dynamic Meteorology and Oceanography is a peer-reviewed open access journal published by Stockholm University Press.

