

Convolutional Autoencoder for the Spatiotemporal Latent Representation of Turbulence

Doan, Nguyen Anh Khoa; Racca, Alberto ; Magri, Luca

DOI

[10.1007/978-3-031-36027-5_24](https://doi.org/10.1007/978-3-031-36027-5_24)

Publication date

2023

Document Version

Final published version

Published in

Lecture Notes in Computer Science - ICCS 2023

Citation (APA)

Doan, N. A. K., Racca, A., & Magri, L. (2023). Convolutional Autoencoder for the Spatiotemporal Latent Representation of Turbulence. In J. Mikyška, C. de Mulatier, V. V. Krzhizhanovskaya, P. M. A. Soot, M. Paszynski, & J. J. Dongarra (Eds.), *Lecture Notes in Computer Science - ICCS 2023: Computational Science – ICCS 2023 23rd International Conference, Prague, Czech Republic, July 3–5, 2023, Proceedings, Part IV* (Vol. 10476, pp. 328-335). (Lecture Notes in Computer Science - ICCS2023; Vol. 10476). Springer. https://doi.org/10.1007/978-3-031-36027-5_24

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Green Open Access added to TU Delft Institutional Repository

'You share, we take care!' - Taverne project

<https://www.openaccess.nl/en/you-share-we-take-care>

Otherwise as indicated in the copyright section: the publisher is the copyright holder of this work and the author uses the Dutch legislation to make this work public.



Convolutional Autoencoder for the Spatiotemporal Latent Representation of Turbulence

Nguyen Anh Khoa Doan^{1(✉)}, Alberto Racca^{2,3}, and Luca Magri^{3,4}

¹ Delft University of Technology, 2629HS Delft, The Netherlands
`n.a.k.doan@tudelft.nl`

² University of Cambridge, Cambridge CB2 1PZ, UK

³ Imperial College London, London SW7 2AZ, UK

⁴ The Alan Turing Institute, London NW1 2DB, UK

Abstract. Turbulence is characterised by chaotic dynamics and a high-dimensional state space, which make this phenomenon challenging to predict. However, turbulent flows are often characterised by coherent spatiotemporal structures, such as vortices or large-scale modes, which can help obtain a latent description of turbulent flows. However, current approaches are often limited by either the need to use some form of thresholding on quantities defining the isosurfaces to which the flow structures are associated or the linearity of traditional modal flow decomposition approaches, such as those based on proper orthogonal decomposition. This problem is exacerbated in flows that exhibit extreme events, which are rare and sudden changes in a turbulent state. The goal of this paper is to obtain an efficient and accurate reduced-order latent representation of a turbulent flow that exhibits extreme events. Specifically, we employ a three-dimensional multiscale convolutional autoencoder (CAE) to obtain such latent representation. We apply it to a three-dimensional turbulent flow. We show that the Multiscale CAE is efficient, requiring less than 10% degrees of freedom than proper orthogonal decomposition for compressing the data and is able to accurately reconstruct flow states related to extreme events. The proposed deep learning architecture opens opportunities for nonlinear reduced-order modeling of turbulent flows from data.

Keywords: Chaotic System · Reduced Order Modelling · Convolutional Autoencoder

1 Introduction

Turbulence is a chaotic phenomenon that arises from the nonlinear interactions between spatiotemporal structures over a wide range of scales. Turbulent flows are typically high-dimensional systems, which may exhibit sudden and unpredictable bursts of energy/dissipation [2]. The combination of these dynamical properties makes the study of turbulent flows particularly challenging. Despite these complexities, advances have been made in the analysis of turbulent flows

through the identification of coherent (spatial) structures, such as vortices [15], and modal decomposition techniques [13]. These achievements showed that there exist energetic patterns within turbulence, which allow for the development of reduced-order models.

To efficiently identify these patterns, recent works have used machine learning [3]. Specifically, Convolutional Neural Networks (CNNs) have been used to identify spatial features in flows [10] and perform nonlinear modal decomposition [7, 11]. These works showed the advantages of using CNNs over traditional methods based on Principal Component Analysis (PCA) (also called Proper Orthogonal Decomposition in the fluid mechanics community), providing lower reconstruction errors in two-dimensional flows. More recently, the use of such CNN-based architecture has also been extended to a small 3D turbulent channel flow at a moderate Reynolds number [12]. The works highlight the potential of deep learning for the analysis and reduced-order modelling of turbulent flows, but they were restricted to two-dimensional flows or weakly turbulent with non-extreme dynamics. Therefore, the applicability of deep-learning-based techniques to obtain an accurate reduced-order representation of 3D flows with extreme events remains unknown. Specifically, the presence of extreme events is particularly challenging because they appear rarely in the datasets. In this paper, we propose a 3D Multiscale Convolutional Autoencoder (CAE) to obtain such a reduced representation of the 3D Minimal Flow Unit (MFU), which is a flow that exhibit such extreme events in the form of a sudden and rare intermittent quasi-laminar flow state. We explore whether the Multiscale CAE is able to represent the flow in a latent space with a reduced number of degrees of freedom with higher accuracy than the traditional PCA, and whether it can also accurately reconstruct the flow state during the extreme events.

Section 2 describes the MFU and its extreme events. Section 3 presents in detail the Multiscale CAE framework used to obtain a reduced-order representation of the MFU. The accuracy of the Multiscale CAE in reconstructing the MFU state is discussed in Sect. 4. A summary of the main results and directions for future work are provided in Sect. 5.

2 Minimal Flow Unit

The flow under consideration is the 3D MFU [9]. The MFU is an example of prototypical near-wall turbulence, which consists of a turbulent channel flow whose dimensions are smaller than conventional channel flow simulations. The system is governed by the incompressible Navier-Stokes equations

$$\begin{aligned}\nabla \cdot \mathbf{u} &= 0, \\ \partial_t \mathbf{u} + \mathbf{u} \cdot \nabla \mathbf{u} &= \frac{1}{\rho} \mathbf{f}_0 - \frac{1}{\rho} \nabla p + \nu \Delta \mathbf{u},\end{aligned}\tag{1}$$

where $\mathbf{u} = (u, v, w)$ is the 3D velocity field and $\mathbf{f}_0 = (f_0, 0, 0)$ is the constant forcing in the streamwise direction, x ; ρ , p , and ν are the density, pressure, and kinematic viscosity, respectively. In the wall-normal direction, y , we impose a

no-slip boundary condition, $\mathbf{u}(x, \pm\delta, z, t) = 0$, where δ is half the channel width. In the streamwise, x , and spanwise, z , directions we have periodic boundary conditions. For this study, a channel with dimension $\Omega \equiv \pi\delta \times 2\delta \times 0.34\pi\delta$ is considered, as in [2] with $\delta = 1.0$. The Reynolds number of the flow, which is based on the bulk velocity and the half-channel width, is set to $Re = 3000$, which corresponds to a friction Reynolds number $Re_\tau \approx 140$. An in-house code similar to that of [1] is used to simulate the MFU, and generate the dataset on which the Multiscale CAE is trained and assessed.

The extreme events in the MFU are quasi-relaminarization events, which take place close to either wall. A typical evolution of the flow during a quasi-relaminarization event is shown in Fig. 1(a–g). Time is normalized by the eddy turnover time. During an extreme event, (i) the flow at either wall (the upper wall in Fig. 1) becomes laminar (Fig. 1(a–c)); (ii) the flow remains laminar for some time (Fig. 1(c–f)), which results in a larger axial velocity close to the centerline (and therefore an increase in kinetic energy); (iii) the greater velocity close to the centerline makes the effective Reynolds number of the flow larger, which in turn makes the flow prone to a turbulence burst on the quasi-laminar wall; (iv) the turbulence burst occurs on the quasi-laminar wall, which results in a large increase in the energy dissipation rate; and (v) the flow close to that quasi-laminar wall becomes turbulent again, which leads to a decrease in the kinetic energy (Fig. 1(g)).

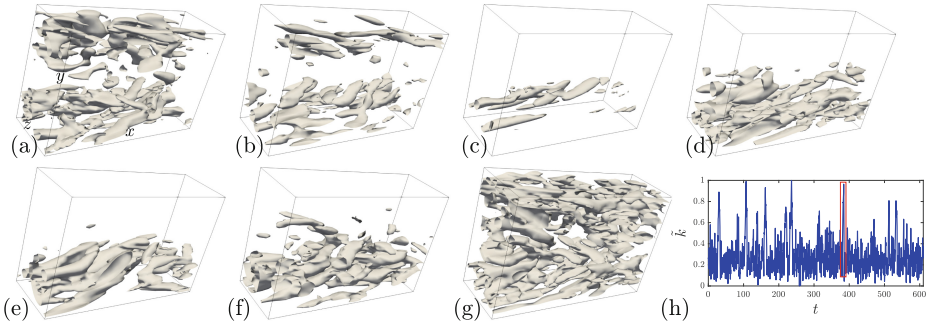


Fig. 1. (a–g) Snapshots of the Q -criterion isosurface (with value $Q = 0.1$) during an extreme event, where $Q = 0.5(\|\boldsymbol{\omega}\|^2 - \|\mathbf{S}\|^2)$, $\boldsymbol{\omega}$ is the vorticity vector, and $\mathbf{S}$ is the strain-rate tensor. (h) Evolution of kinetic energy, k , of the MFU. The red box indicates the event whose evolution is shown in (a–g). (Color figure online)

These quasi-relaminarisation are accompanied by bursts in the total kinetic energy, $k(t) = \int \int \int_{\Omega} \frac{1}{2} \mathbf{u} \cdot \mathbf{u} dx dy dz$, where Ω is the computational domain. This can be seen in Fig. 1h, where the normalized kinetic energy, $\tilde{k} = (k - \min(k))/(\max(k) - \min(k))$ is shown.

The dataset of the MFU contains 2000 eddy turnover times (i.e., 20000 snapshots) on a grid of $32 \times 256 \times 16$, which contains 50 extreme events. The first

200 eddy turnover times of the dataset (2000 snapshots) are employed for the training of the Multiscale CAE, which contains only 4 extreme events.

3 Multiscale Convolutional Autoencoder

We implement a 3D convolutional autoencoder. It should be noted that we only consider CNN-based autoencoder here and not other approaches such as transformer-based ones (like in [4]) as the latter have not yet shown to be widely applicable to dataset with a strong spatial based information (such as images or flow dataset). A schematic of the proposed architecture is shown in Fig. 2.

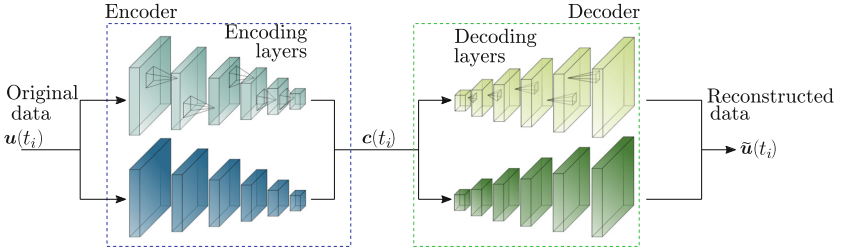


Fig. 2. Schematic of the 3D multiscale convolutional autoencoder.

The three dimensional convolution autoencoder (CAE) learns an efficient reduced-order representation of the original data, which consists of the flow state $\mathbf{u} \in \mathbb{R}^{N_x \times N_y \times N_z \times N_u}$, where N_x , N_y and N_z are the number of grid points, and $N_u = 3$ is the number of velocity components. On one hand, the encoder (blue box in Fig. 2) reduces the dimension of the data down to a latent state, $\mathbf{c}$, with small dimension $N_c \ll N_x N_y N_z N_u$. This operation can be symbolically expressed as $\mathbf{c} = \mathcal{E}(\mathbf{u}; \phi_E)$, where ϕ_E represents the weights of the encoder. On the other hand, the decoder (green box in Fig. 2) reconstructs the data from the latent state back to the original full flow state. This operation is expressed as $\tilde{\mathbf{u}} = \mathcal{D}(\mathbf{c}; \phi_D)$, where ϕ_D are the trainable weights of the decoder. We employ a multiscale autoencoder, which was originally developed for image-based super-resolution analysis [5]. We use here a multiscale autoencoder and not a standard one as previous works [6, 8] have demonstrated the ability of the multiscale version in leveraging the multiscale information in turbulent flows to better reconstruct the flow from the latent space. It relies on the use of convolutional kernels of different sizes to analyse the input and improve reconstruction in fluids [8, 14]. In this work, two kernels, $(3 \times 5 \times 3)$ and $(5 \times 7 \times 5)$, are employed (represented schematically by the two parallel streams of encoder/decoder in the blue and green boxes in Fig. 2). This choice ensures a trade-off between the size of the 3D multiscale autoencoder and the reconstruction accuracy (see Sect. 4). To reduce the dimension of the input, the convolution operation in each layer of the encoder is applied in a strided manner, which means that the convolutional

neural network (CNN) kernel is progressively applied to the input by moving the CNN kernel by $(s_x, s_y, s_z) = (2, 4, 2)$ grid points. This results in an output of a smaller dimension than the input. After each convolution layer, to fulfill the boundary conditions of the MFU, periodic padding is applied in the x and z directions, while zero padding is applied in the y direction. Three successive layers of CNN/padding operations are applied to decrease the dimension of the original field from $(32, 256, 16, 3)$ to $(2, 4, 2, N_f)$, where N_f is the specified number of filters in the last encoding layer. As a result, the dimension of the latent space is $N_c = 16 \times N_f$. The decoder mirrors the architecture of the encoder, where transpose CNN layers [16] are used, which increase the dimension of the latent space up to the original flow dimension. The end-to-end autoencoder is trained by minimizing the mean squared error (MSE) between the reconstructed velocity field, $\hat{\mathbf{u}}$, and the original field, $\mathbf{u}$ using the ADAM optimizer.

4 Reconstruction Error

We analyze the ability of the CAE to learn a latent space that encodes the flow state accurately. To do so, we train three CAEs with latent space dimensions of $N_c = 384, 768$, and 1536 . The training of each CAE took between 8 and 16 h using 4 Nvidia V100S (shorter training time for the model with the smaller latent space). After training, the CAE can process 100 samples in 1.8 s to 5 s depending on the dimension of the latent space (the faster processing time corresponding to the CAE with the smaller latent space). We compute their reconstruction errors on the test set based on the MSE. A typical comparison between a reconstructed velocity field obtained from the CAE with $N_c = 1536$ is shown in Fig. 3. The CAE is able to reconstruct accurately the features of the velocity field.

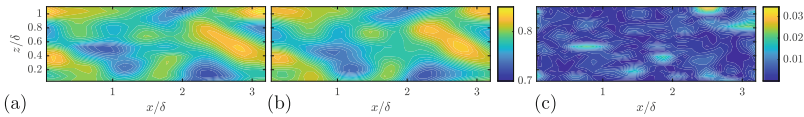


Fig. 3. Comparison of (a) the actual velocity magnitude (ground truth), (b) the CAE-reconstructed velocity magnitude, (c) the root-squared difference between (a) and (b) in the mid- y plane for a typical snapshot in the test set.

To provide a comparison with the CAE, we also compute the reconstruction error obtained from PCA, whose principal directions are obtained with the method of snapshots [13] on the same dataset. Figure 4 shows the reconstruction error. PCA decomposition requires more than 15000 PCA components to reach the same level of accuracy as the CAE with a latent space of dimension 1536. This highlights the advantage of learning a latent representation with nonlinear operations, as in the autoencoder, compared with relying on a linear combination of components, as in PCA.

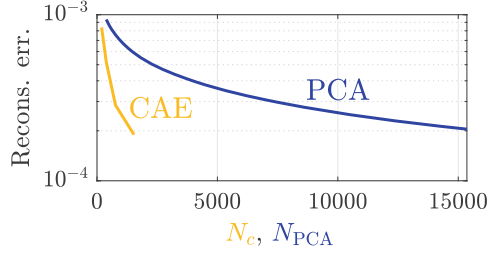


Fig. 4. Reconstruction error with a PCA-based method (blue) and the autoencoder (yellow) for different dimensions of the latent space, N_c , or number of retained PCA components, N_{PCA} . The reconstruction error is computed as the mean squared error between the reconstructed velocity field and the exact field, averaged over the test set. (Color figure online)

The better performance of the CAE with respect to PCA is evident when we consider the reconstruction accuracy for flow states that correspond to extreme events, which we extract from the test set. Here, we define an extreme event as a flow state with normalized kinetic energy above a user-selected threshold value of 0.7. Hence, for the selected snapshots in the test set, the mean squared error (MSE) between the reconstructed velocity and the truth is computed using the CAE and PCA as a function of the latent space size. The resulting MSE is shown in Fig. 5, where the CAE exhibits an accuracy for the extreme dynamics similar to the reference case of the entire test set (see Fig. 4). On the other hand, the accuracy of the PCA is lower than the reference case of the entire dataset. This lack of accuracy is further analysed in Fig. 6, where typical velocity magnitudes, i.e. the norm of the velocity field $\mathbf{u}$, are shown for the mid- z plane during a representative extreme event. The velocity reconstructed with the CAE (Fig. 6b) is almost identical to the true velocity field (Fig. 6a). The CAE captures the smooth variation at the lower wall indicating a quasi-laminar flow state in that region. In contrast, the velocity field reconstructed using PCA (Fig. 6c) completely fails at reproducing those features. This is because extreme states are rare in the training set used to construct the PCA components (and the CAE). Because of this, the extreme states only have a small contribution to the PCA components and are accounted for only in higher order PCA components, which are neglected in the 1536-dimensional latent space.

This result indicates that only the CAE and its latent space can be used for further applications that requires a reduced-order representation of the flow, such as when trying to develop a reduced-order model. This is also supported by findings in [14] where it was shown that a similar CAE could be used in combination with a reservoir computer to accurately forecast the evolution of a two dimensional flow.

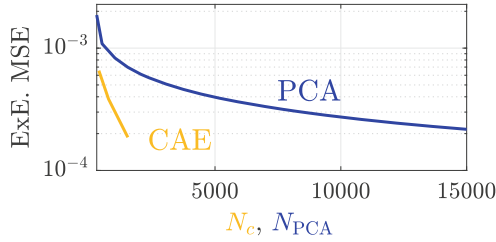


Fig. 5. Reconstruction error on the extreme events flow states with PCA-based method (blue) and autoencoder (yellow) for different dimensions of the latent space, N_c , or number of retained PCA components, N_{PCA} . The reconstruction error is computed as in Fig. 4 but only on the flow state corresponding to extreme events. (Color figure online)

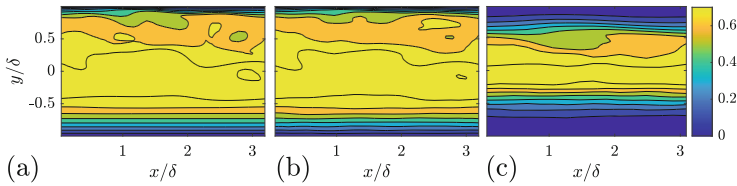


Fig. 6. Comparison of (a) the velocity magnitude (ground truth), (b) the CAE-reconstructed velocity magnitude, (c) the PCA-reconstructed velocity magnitude in the mid- z plane for a typical extreme event snapshot in the test set.

5 Conclusion

In this work, we develop a nonlinear autoencoder to obtain an accurate latent representation of a turbulent flow that exhibits extreme events. We propose the 3D Multiscale CAE to learn the spatial features of the MFU, which exhibits extreme events in the form of near-wall quasi-relaminarization events. The model consists of a convolutional autoencoder with multiple channels, which learn an efficient reduced latent representation of the flow state. We apply the framework to a three-dimensional turbulent flow with extreme events (MFU). We show that the Multiscale CAE is able to compress the flow state to a lower-dimensional latent space by three orders of magnitude to accurately reconstruct the flow state from this latent space. This constitutes a key improvement over principal component analysis (PCA), which requires at least one order of magnitude more PCA components to achieve an accuracy similar to the CAE. This improvement in reconstruction accuracy is crucial for the reconstruction of the flow state during the extreme events of the MFU. This is because extreme states are rare and, thus, require a large number of PCA components to be accurately reconstructed.

The proposed method and results open up possibilities for using deep learning to obtain an accurate latent reduced representation of 3D turbulent flows. Future work will be devoted to physically interpreting the latent space discovered by the Multiscale CAE, and learning the dynamics in this latent space.

Acknowledgements. The authors thank Dr. Modesti for providing the flow solver. N.A.K.D and L.M acknowledge that part of this work was performed during the 2022 Stanford University CTR Summer Program. L.M. acknowledges the financial support from the ERC Starting Grant PhyCo 949388. A.R. is supported by the Eric and Wendy Schmidt AI in Science Postdoctoral Fellowship, a Schmidt Futures program.

References

1. Bernardini, M., Pirozzoli, S., Orlandi, P.: Velocity statistics in turbulent channel flow up to $Re_\tau = 4000$. *J. Fluid Mech.* **742**, 171–191 (2014)
2. Blonigan, P.J., Farazmand, M., Sapsis, T.P.: Are extreme dissipation events predictable in turbulent fluid flows? *Phys. Rev. Fluids* **4**, 044606 (2019)
3. Brunton, S.L., Noack, B.R., Koumoutsakos, P.: Machine learning for fluid mechanics. *Annu. Rev. Fluid Mech.* **52**(1), 477–508 (2020)
4. Dosovitskiy, A., et al.: An image is worth 16x16 words: transformers for image recognition at scale. In: *ICLR2021* (2021)
5. Du, X., Qu, X., He, Y., Guo, D.: Single image super-resolution based on multi-scale competitive convolutional neural network. *Sensors* **18**(3), 1–17 (2018)
6. Fukami, K., Nabae, Y., Kawai, K., Fukagata, K.: Synthetic turbulent inflow generator using machine learning. *Phys. Rev. Fluids* **4**(6), 1–18 (2019)
7. Fukami, K., Nakamura, T., Fukagata, K.: Convolutional neural network based hierarchical autoencoder for nonlinear mode decomposition of fluid field data. *Phys. Fluids* **32**(9), 1–12 (2020)
8. Hasegawa, K., Fukami, K., Murata, T., Fukagata, K.: Machine-learning-based reduced-order modeling for unsteady flows around bluff bodies of various shapes. *Theor. Comput. Fluid Dyn.* **34**(4), 367–383 (2020). <https://doi.org/10.1007/s00162-020-00528-w>
9. Jiménez, J., Moin, P.: The minimal flow unit in near-wall turbulence. *J. Fluid Mech.* **225**, 213–240 (1991)
10. Morimoto, M., Fukami, K., Zhang, K., Nair, A.G., Fukagata, K.: Convolutional neural networks for fluid flow analysis: toward effective metamodeling and low-dimensionalization. *Theor. Comput. Fluid Dyn.* **35**, 633–658 (2021)
11. Murata, T., Fukami, K., Fukagata, K.: Nonlinear mode decomposition with machine learning for fluid dynamics. *J. Fluid Mech.* **882**, A13 (2020)
12. Nakamura, T., Fukami, K., Hasegawa, K., Nabae, Y., Fukagata, K.: Convolutional neural network and long short-term memory based reduced order surrogate for minimal turbulent channel flow. *Phys. Fluids* **33**, 025116 (2021)
13. Berkooz, G., Holmes, P., Lumley, J.L.: The proper orthogonal, decomposition in the analysis of turbulent flows. *Annu. Rev. Fluid Mech.* **25**, 539–575 (1993)
14. Racca, A., Doan, N.A.K., Magri, L.: Modelling spatiotemporal turbulent dynamics with the convolutional autoencoder echo state network. *arXiv* (2022)
15. Yao, J., Hussain, F.: A physical model of turbulence cascade via vortex reconnection sequence and avalanche. *J. Fluid Mech.* **883**, A51 (2020)
16. Zeiler, M.D., Krishnan, D., Taylor, G.W., Fergus, R.: Deconvolutional networks. In: *Proceedings of 2010 IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pp. 2528–2535 (2010)