

Learning Stochastic Multiscale Models

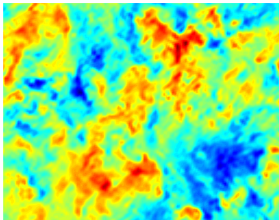
Andrew F. Ilersich
andrew.ilersich@mail.utoronto.ca

Prasanth B. Nair
prasanth.nair@utoronto.ca

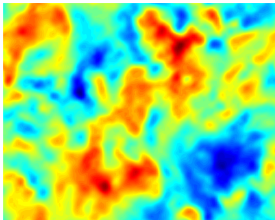
NeurIPS 2025

Background

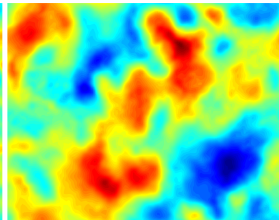
DNS



LES (fine)



LES (coarse)



1

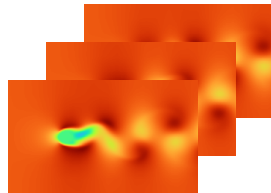
- ▶ Complex physical systems exhibit a vast range of scales
- ▶ Resolving the smallest scales is computationally infeasible, requiring approximation
- ▶ Closure models are a common approach that uses coarsened representations

¹Tom-Robin Teschner. The introduction to LES I wish I had. CFD University. 2025. 

Problem Setting

$u : \Omega \times [t_0, T] \rightarrow \mathbb{R}^{d_u}$
is a spatio-temporal field

Sample on fine spatial
grid, over n_t time steps

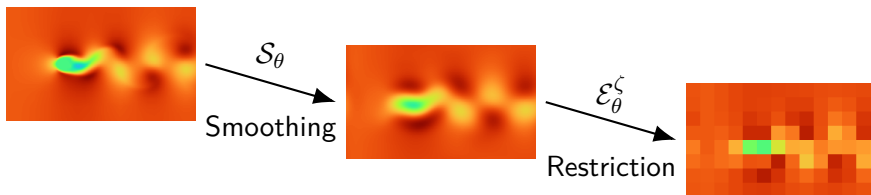


Training dataset $\mathcal{Y} = \{(t_i, y_i)\}_{i=1}^{n_t}$ where each snapshot $y_i \in \mathbb{R}^{n_y}$.

Goal

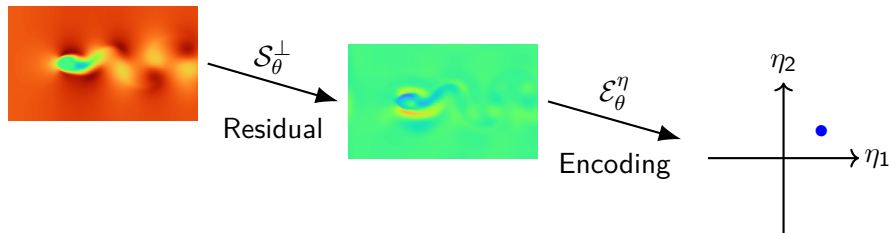
Given dataset $\mathcal{Y}$, learn a stochastic multiscale model that explicitly captures the coupled dynamics of **macroscale** and **microscale** features

Forming the Macroscale State with Probabilistic Encoder



Macroscale state $\zeta \sim p_\theta(\zeta | y)$, where $\mathbb{E}[\zeta | y] = \mathcal{E}_\theta^\zeta(\mathcal{S}_\theta(y))$

Forming the Microscale State with Probabilistic Encoder



Microscale state $\eta \sim p_\theta(\eta | y)$, where $\mathbb{E}[\eta | y] = \mathcal{E}_\theta^\eta(\mathcal{S}_\theta^\perp(y))$ and $\mathcal{S}_\theta^\perp(y) = y - \mathcal{S}_\theta(y)$

$$d \begin{bmatrix} \zeta \\ \eta \end{bmatrix} = \begin{bmatrix} f_{\theta}(\zeta, \phi(\eta)) \\ g_{\theta}(\eta, \psi(\zeta)) \end{bmatrix} dt + L_{\theta}(t) d\beta$$

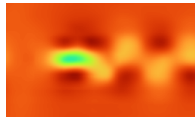
- ▶ Macroscale drift f_{θ} , microscale drift g_{θ}
- ▶ Dispersion matrix L_{θ}
- ▶ Wiener process β
- ▶ Inter-scale modulation terms ϕ and ψ

Reconstructing the Full State with Probabilistic Decoder

Prolong macroscale state mean



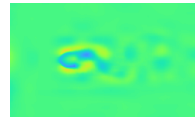
$\mathcal{D}_{\theta}^{\zeta}$



Decode microscale state mean

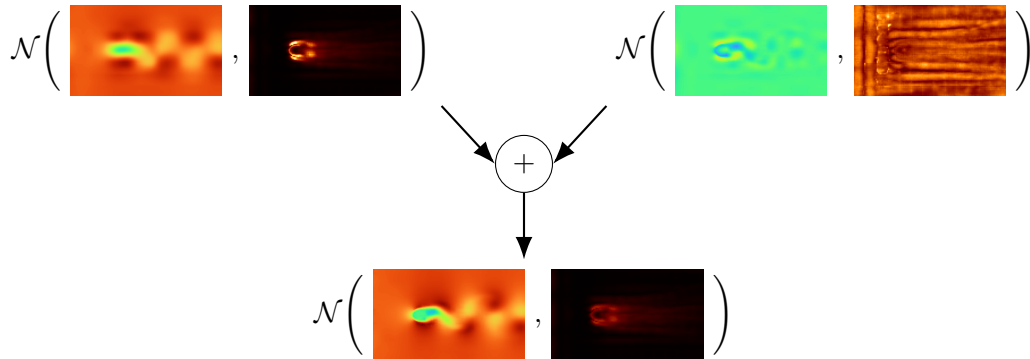
$\{\eta_1, \eta_2, \dots\}$

$\mathcal{D}_{\theta}^{\eta}$



Reconstructing the Full State with Probabilistic Decoder

Combine large-scale and small-scale features



Multiscale Likelihood

Goal

Learn a multiscale representation with explicit scale separation

Multiscale Likelihood

Goal

Learn a multiscale representation with explicit scale separation

Standard likelihoods do not ensure scale separation

$$p_{\theta}(y \mid \zeta, \eta) = \mathcal{N}(y \mid \mathcal{D}_{\theta}^{\zeta}(\zeta) + \mathcal{D}_{\theta}^{\eta}(\eta), \Sigma_y) \quad \text{X}$$

Multiscale Likelihood

Goal

Learn a multiscale representation with explicit scale separation

Standard likelihoods do not ensure scale separation

$$p_{\theta}(y \mid \zeta, \eta) = \mathcal{N}(y \mid \mathcal{D}_{\theta}^{\zeta}(\zeta) + \mathcal{D}_{\theta}^{\eta}(\eta), \Sigma_y) \quad \times$$

We formulate a multiscale likelihood based on Product of Experts (PoE) approach:

$$p_{\theta}(y \mid \zeta, \eta) = \frac{1}{Z(\zeta, \eta)} \times \underbrace{\mathcal{N}(y \mid \mathcal{D}_{\theta}^{\zeta}(\zeta), \Sigma_{\zeta})}_{\text{Macroscale expert}} \times \underbrace{\mathcal{N}(y - \mathcal{D}_{\theta}^{\zeta}(\zeta) \mid \mathcal{D}_{\theta}^{\eta}(\eta), \Sigma_{\eta})}_{\text{Microscale expert}}$$

The microscale state only captures what the macroscale cannot

Multiscale Likelihood

Goal

Learn a multiscale representation with explicit scale separation

$$\begin{aligned} \log p_{\theta}(y \mid \zeta, \eta) = & \underbrace{-\frac{1}{2} \|y - \mathcal{D}_{\theta}^{\zeta}(\zeta)\|_{(\Sigma_{\zeta}^{-1} + \Sigma_{\eta}^{-1})}^2}_{\text{Macroscale residual}} + \underbrace{\left(y - \mathcal{D}_{\theta}^{\zeta}(\zeta), \mathcal{D}_{\theta}^{\eta}(\eta)\right)_{\Sigma_{\eta}^{-1}}}_{\text{Microscale correlation}} \\ & \underbrace{-\frac{1}{2} \|\mathcal{D}_{\theta}^{\eta}(\eta)\|_{\Lambda}^2}_{\text{Microscale norm}} + \underbrace{\frac{1}{2} \log |\Sigma_{\zeta}^{-1} + \Sigma_{\eta}^{-1}|}_{\text{Cov. determinant}} - \frac{n_y}{2} \log(2\pi) \end{aligned}$$

where $\|v\|_A^2 = v^T A v$, $(u, v)_A = u^T A v$, and $\Lambda = \Sigma_{\eta}^{-1}(\Sigma_{\zeta}^{-1} + \Sigma_{\eta}^{-1})^{-1}\Sigma_{\eta}^{-1}$.

Numerical Results

Comparison of error statistics for each model on each test case, OA = our approach

Model type	Error mean $\pm$ std. dev. on test set			
	Wave 1D	KdV 1D	Cylinder 2D	SWE 2D
Coarse DNS	1.381 ± 0.382	0.523 ± 0.196	0.130 ± 0.040	0.861 ± 0.190
DMD	0.492 ± 0.066	0.163 ± 0.046	0.016 ± 0.006	0.531 ± 0.227
POD-SINDy	0.469 ± 0.054	0.057 ± 0.028	0.051 ± 0.024	0.535 ± 0.210
Macro-only	0.565 ± 0.116	0.324 ± 0.078	0.063 ± 0.005	0.426 ± 0.099
OA, $n_\eta = 5$	0.079 ± 0.031	0.042 ± 0.017	0.009 ± 0.001	0.152 ± 0.047

Explicitly modelling the microscale state improves error across all test cases

Visualization of Scale Separation

Learning Stochastic Multiscale Models

Code and data available:

`github.com/ailersic/multiscale-visde`