

GenAI4Numerical/Math

Hanno Gottschalk (TU Berlin) joint work with C. Drygala, F. di Mare, E. Ross, M. de Campos, W. Krebs, L. Schwartz. . .

Institute of Mathematics, TU Berlin

ai4math WS U Potsdam | 5. August 2026

The Challenge of Generative Learning



From NVIDIA Style GAN, thispersondoesnotexist.com, Karras, Laine, Alia 2018

► **Challenge:** $X_j \in \mathbb{R}^d$ is high dimensional, $d \sim 10^6$.

Assumptions and Existence of a Generator

- ▶ Source measure ν and target measure μ have a density wrt the Lebesgue measure on Ω

$$d\nu(x) = f_\nu(x) d\lambda(x), \quad d\mu(x) = f_\mu(x) d\lambda(x)$$

- ▶ f_ν, f_μ is a k -times differentiable function
- ▶ $f(x) \geq \kappa > 0$ on Ω .

Theorem (Realizability):

- (i) There exists a k times differentiable function $\phi_0 : \Omega \rightarrow \Omega$ such that

$$\phi_{0*}\nu = \mu.$$

Flow maps

- Want to realize ϕ as a flow endpoint w.r.t a vector field $v(x, t)$,

$$\dot{\phi}_t(x) = v(\phi(x), t), \quad \phi = \phi_1$$

- ▶ Does a suitable $v(x, t)$ exist?
- ▶ Continuity equation

$$\partial_t f_t(x) = \nabla \cdot (v(x, t) f_t(x)), \quad f_0(x) = f_\nu(x), \quad f_1(x) = f_\mu(x) \quad (!)$$

Beckmann problem

- ▶ Model the influx $w(x) = v(x, t)f_t(x)$ time independent such that

$$\nabla \cdot w(x) = f_\mu(x) - f_\nu(x), \quad x \in \Omega, \quad n \cdot w = 0, \quad x \in \partial\Omega. \quad (1)$$

- ▶ Then from the continuity equation for $v(x, t) = \frac{w(x)}{f_t(x)}$

$$\partial_t f_t(x) = \nabla \cdot w(x) = f_\mu(x) - f_\nu(x) \quad \Rightarrow \quad f_t(x) = t f_\mu(x) - (1 - t) f_\nu(x)$$

- ▶ Beckmann's problem:

$$\frac{1}{2} \int_{\Omega} |w|^2 dx \rightarrow \min \quad \text{under constraints (1)}$$

Unrestricted Beckmann Problem

- Introduce Lagrangian multipliers $u : \Omega \rightarrow \mathbb{R}$, $\nabla u(x) \cdot n = 0$ on $\partial\Omega$, $v \in C(\partial\Omega)$, $f = f_\mu - f_\nu$

$$\mathcal{L}(w, u, v) = \frac{1}{2} \int_{\Omega} |w|^2 dx + \int_{\Omega} (\nabla \cdot w - f) u dx + \int_{\partial\Omega} n \cdot (wv) dA$$

- ## ► Optimality conditions

$$\frac{\delta}{\delta v} \mathcal{L}(w, u, v) = 0 \Rightarrow n \cdot w = 0 \text{ on } \partial\Omega, \quad \frac{\delta}{\delta v} \mathcal{L}(w, u, v) = 0 \Rightarrow \nabla \cdot w - f = 0 \text{ on } \Omega$$

$$\frac{\delta}{\delta v} \mathcal{L}(w, u, v) = 0 \quad \Rightarrow \quad -\nabla u = w \quad \text{on } \Omega$$

Solution of the Optimality Conditions

- $u : \Omega \rightarrow \mathbb{R}$ fulfills Poisson equation with zero Neumann boundary conditions

$$-\Delta u = f \quad \text{on } \Omega \quad \nabla u \cdot n = 0 \quad \text{on } \partial\Omega$$

- ▶ Weak formulation not coercive, but fulfills condition for Fredholm alternative with $g = 0$

$$\int_{\Omega} f \, dx - \int_{\partial\Omega} g \, dA = \int_{\Omega} f_{\mu} - f_{\nu} \, dx = 0.$$

- ▶ Classical solution exists (Nadri 2015).

Schauder Estimates

Theorem

- ▶ u exists and is $k + 2$ times differentiable (Schauder estimates, Agmon, Douglis, Nirenberg 59)
- ▶ $\Rightarrow w$ is $k + 1$ times differentiable
- ▶ $\Rightarrow v$ is k -times differentiable (Asatryan et al, 2023)
- ▶ $\Rightarrow \phi_t$ is k times differentiable
- ▶ Transport map is regular and can be well approximated by NN, (Yarotzky 2017).

Empirical and Theoretical Loss in Adv. Learning

Theorem (Goodfellow et al 2014):

(i) The expected loss

$$L(\phi, D, \{X_j\}) = \mathbb{E}_{X_j \sim \mu}[\hat{L}(\phi, D, \{X_j\})]$$

is maximized by $D_\phi = \frac{f(x)}{f(x) + f_\phi(x)}$, provided this is realizable in $\mathcal{H}_D$;

(ii)

$$L(\phi, D_\phi, \{X_j\}) = \frac{1}{2} \left[\mathfrak{d}_{KL} \left(\mu \left\| \frac{\phi_* \nu + \mu}{2} \right) + \mathfrak{d}_{KL} \left(\phi_* \nu \left\| \frac{\phi_* \nu + \mu}{2} \right) \right] + \log(2) = \mathfrak{d}_{JS}(\mu \| \phi_* \nu) + \log(2).$$

Error Decomposition for Generative Adversarial Learning

- Asatryan, G., Lippert, Rottmann (2020), Biau, Cadre, Sangnier, Tanielian (2020)

$$\begin{aligned}
\mathfrak{d}_{JS}(\mu || \hat{\phi}_* \nu) &= L(\hat{\phi}, D_{\hat{\phi}}) - \log(2) \leq \hat{L}(\hat{\phi}, D_{\hat{\phi}}) + \sup_{\phi, D} |L(\phi, D) - \hat{L}(\phi, D)| - \log(2) \\
&\leq \hat{L}(\hat{\phi}, \hat{D}) + \sup_{\phi, D} |L(\phi, D) - \hat{L}(\hat{\phi}, D_{\hat{\phi}})| - \log(2) \\
&\leq \hat{L}(\phi, \hat{D}) + \sup_{\phi, D} |L(\phi, D) - \hat{L}(\phi, D)| - \log(2) \\
&\leq L(\phi, \hat{D}) + 2 \sup_{\phi, D} |L(\phi, D) - \hat{L}(\phi, D)| - \log(2) \\
&\leq L(\phi, D_{\phi_0}) + 2 \sup_{\phi, D} |L(\phi, D) - \hat{L}(\phi, D)| - \log(2) \\
&= \mathfrak{d}_{JS}(\mu || \phi \nu) + 2 \sup_{\phi, D} |L(\phi, D) - \hat{L}(\phi, D)| \rightarrow 0 \quad (\text{ULLN})
\end{aligned}$$

Turbulence

Turbulence in Fluid Dynamics

- ▶ Turbulent flow is **chaotic** by nature
- ▶ Computational Fluid Dynamics (CFD): Simulation of turbulences by numerically solving Navier-Stokes equations:

$$\begin{aligned} \frac{\partial \rho}{\partial t} + \nabla \cdot (\rho u) &= 0 \\ \frac{\partial(\rho u)}{\partial t} + \nabla \cdot [\rho u \otimes u] &= -\nabla p + \nabla \cdot \tau + \rho f \\ \frac{\partial(\rho e)}{\partial t} + \nabla \cdot ((\rho e + p)u) &= \nabla \cdot (\tau \cdot u) + \rho f \cdot u + \nabla \cdot \dot{q} + r \end{aligned}$$

- But: Turbulent flow develops ever smaller and faster structures which are hard to resolve numerically

Ergodicity: What Does “Chaotic” Actually Mean?

- ▶ Vortices always look different but in the long run their statistical properties are determined (**ergodicity**):

$$\lim_{T \rightarrow \infty} \frac{1}{T} \int_0^T f \circ \varphi_t(x_0) dt = \int_{\Omega} f(x) d\mu(x) \quad \forall x_0 \in \Omega. \quad (2)$$

⇒ The time average of a dynamical system equates with the ensemble average of its invariant measure.

- ▶ Probability measure μ on the flow configurations x encodes the statistical properties of the chaotic flow/dynamics $\varphi_t(x_0)$
- ▶ Goal: Sampling from the **unknown** measure $\mu \Rightarrow$ Can we **learn μ from data?**

Changing the Data Model from i.i.d. to Ergodic Data

- ▶ **Theorem (Drygala, Winhardt, di Mare, G. 2021)** Consider the Min-Max two player game for ergodic data

$$\hat{\phi} \in \arg \min_{\phi \in \mathcal{H}} \sup_{D \in \mathcal{H}_D} \hat{L}_T(\phi, D, x_0, \{U_j\})$$

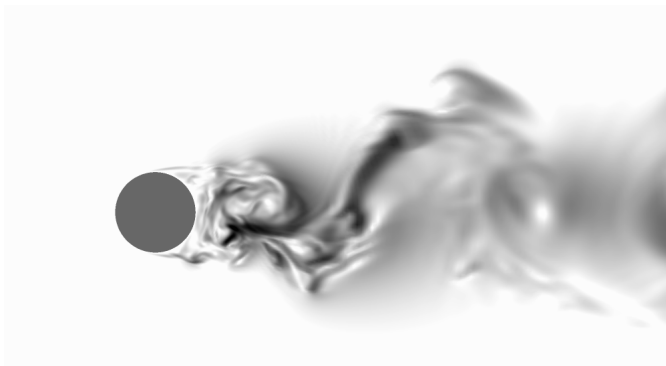
- with $x_0 \in \Omega$, $U_j \sim \nu$, the empirical loss **for ergodic flow** is defined as

$$\hat{L}_T(\phi, D, \{X_j\}) = \frac{1}{2T} \int_0^T \log(D(\varphi_t(x_0))) dt + \frac{1}{2T} \sum_{j=1}^{[T]} \log(1 - D(\phi(U_j))).$$

- **Proof:** Error decomposition + use uniform ergodicity instead of ULLN

Karman Vortex Street - Data

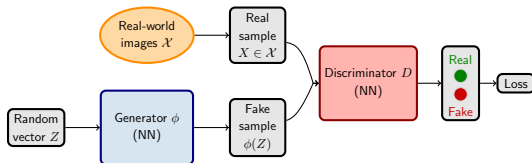
- ▶ Dataset: 5,000 images produced by LES
- ▶ Images: $w \times h = 1,000 \times 600$



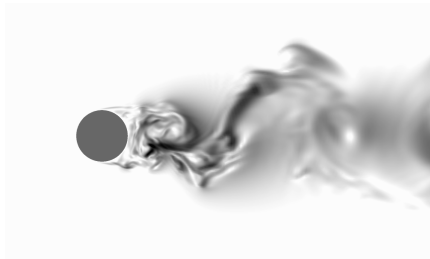
Karman Vortex Street - Training and Inference I

Training

- ▶ GAN frameworks:
 - Vanilla GAN
 - Wasserstein GAN
 - Deep convolutional GAN
- ▶ Epochs: 200
- ▶ Batch size: 20
- ▶ Input:
 - 5,000 LES images
 - Image size: $k \times k$,
 $k \in \{64, 128, 256, \mathbf{512}\}$
 - Noise vector $Z \in \mathbb{R}^{100}$, $Z \sim N(0, 1)$



Karman Vortex Street - Vanilla GAN

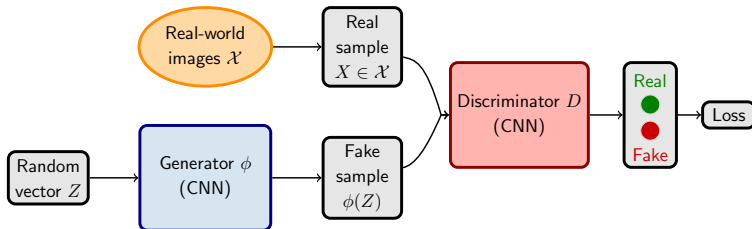


LES image

Fake image

- **Experiment settings:** Architecture of ϕ and D : Fully connected NN with 5 layers
- Optimizer: Adam, Learning rate: $2 \cdot 10^{-4}$

Deep Convolutional Generative Adversarial Network



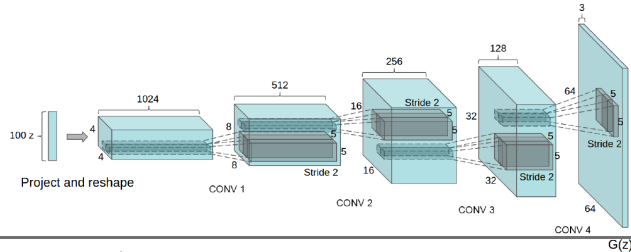
Architecture of DCGAN.

- ▶ Generator and discriminator are convolutional neural networks (CNNs)
- ▶ CNNs especially successful and applicable in field of image processing

Deep Convolutional Generative Adversarial Network

Guidelines to follow for **stable training** at **higher resolution** and **deeper architectures**:

- ▶ Stability: Apply batch normalization on the output layer of ϕ and the input layer of D
- ▶ Deeper architectures: Avoid fully-connected layers on top of convolutional features
- ▶ Higher resolution modeling: Leaky Rectified Linear Unit (ReLU) activation function for D
- ▶ ϕ and D learn own spatial up- or downsampling by replacing deterministic spatial pooling layers



Karman Vortex Street - Inference after 2,000 Epochs

LES



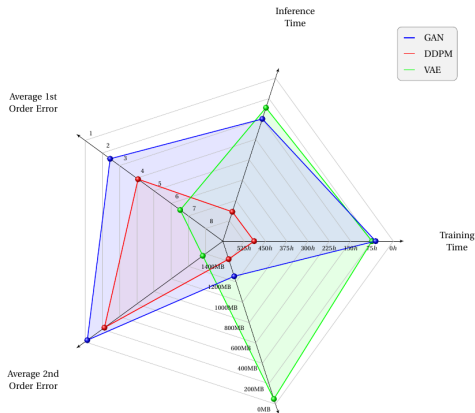
GAN



Comparison of LES (top) and fake images (bottom) produced by generator ϕ trained for 2 000 epochs.

Sampling Results for the Karman Street

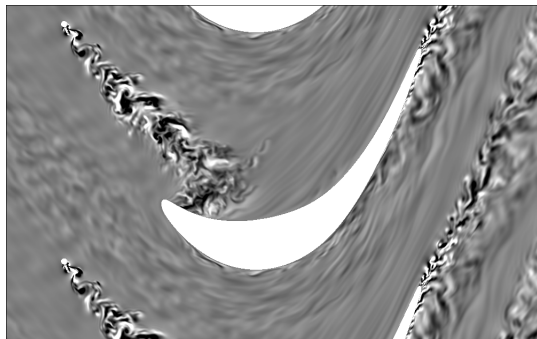
Why GAN?



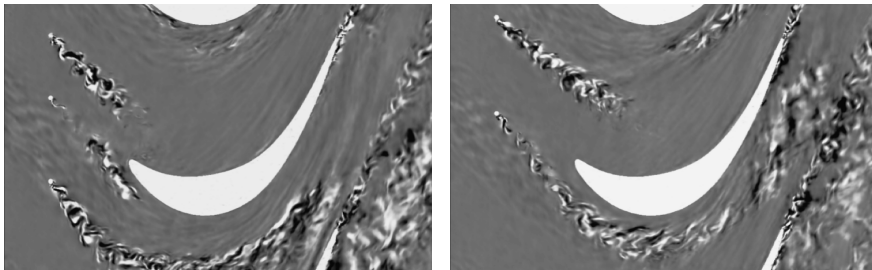
LPT Stator Under Periodic Wake Impact (LPT Stator) - Data

Technische
Universität
Berlin

- ▶ Dataset: 2,250 images produced by LES
- ▶ Images: $w \times h = 1,000 \times 625$

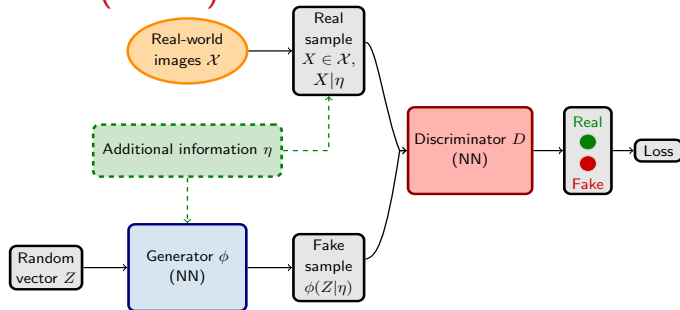


LPT Stator - DCGAN: Inference Results



Examples of images generated by ϕ trained for 2 000 epochs with 2 250 images.

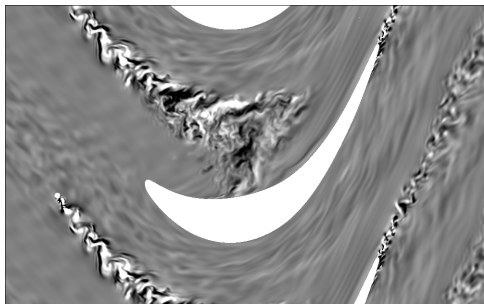
Conditional GAN (cGAN)



- ▶ Take control over the data production process by conditioning GAN framework
- ▶ Extension of loss function:

$$\mathcal{L}_{cond}(D, \phi) = \mathbb{E}_{\substack{X \sim \mu \\ \eta \sim \nu}} [\log(D(X|\eta))] + \mathbb{E}_{\substack{Z \sim \lambda \\ \eta \sim \nu}} [\log(1 - D(\phi(Z|\eta)))] \quad (3)$$

LPT Stator - Conditional Training



LES image



Binary segmentation mask

High-Resolution Image Synthesis with Conditional

- ▶ Special form of cGAN - pix2pixHD (NVIDIA)
- ▶ Generation of **high-resolution photo-realistic** images by **conditioning** the input of the adversarial network on the corresponding **semantic label maps**
- ▶ Based on its former version pix2pix whose optimization problem is given as

$$\min_{\phi} \max_D \mathcal{L}_{cond}(D, \phi) \quad (4)$$

pix2pixHD Innovations

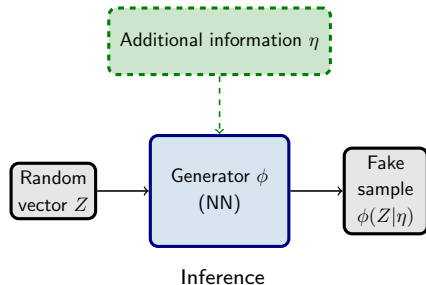
- Introduction of **three innovations** for **improvement** of **photorealism** and **resolution** of the synthesized images:
- 1 Coarse-to-fine generator: Decomposition of the generator into two sub-networks
 $\Rightarrow \phi = \{\phi_1, \phi_2\}$
 - 2 Multi-scale discriminators: Three discriminators D_1, D_2 and D_3 with same architecture but operating on different scales $\Rightarrow$ Modification of loss function:

$$\min_{\phi} \max_{D_1, D_2, D_3} \sum_{i=1}^3 \mathcal{L}_{cond}(\phi, D_i) + \mathcal{L}_{FM}(D_i, \phi) \quad (5)$$

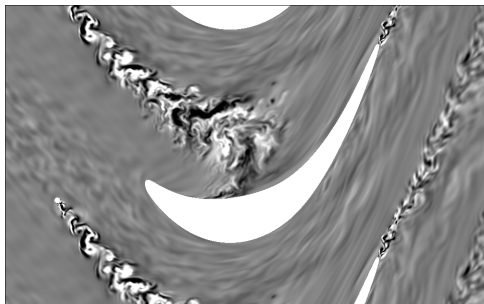
- 3 Feature matching loss $\mathcal{L}_{FM}$: Stabilize training

LPT-Stator - Experiment Settings

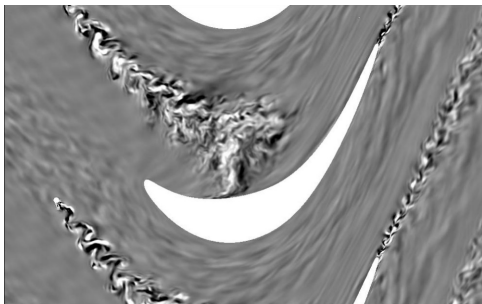
- ▶ Architecture: Adopted from original authors with small changes to avoid artifacts
- ▶ Epochs: 200
- ▶ Batch size: 10
- ▶ Input: 2,000 LES images, noise vector and masks of size $k \times k' = 992 \times 624$
- ▶ Optimizer: Adam
- ▶ Learning rate: $2 \cdot 10^{-4}$



LPT Stator - Inference Results I

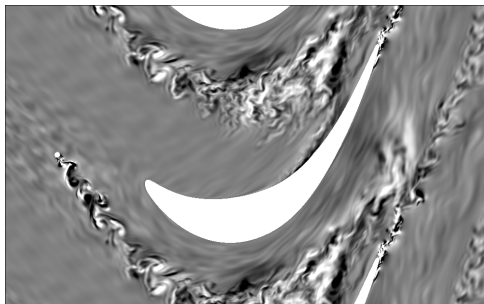


LES image

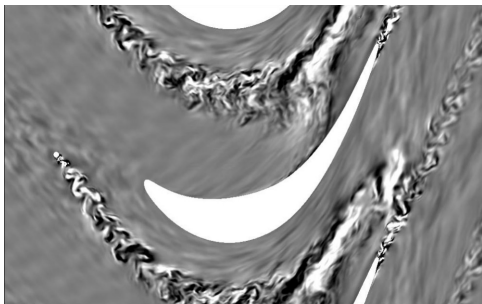


GAN image

LPT Stator - Results II

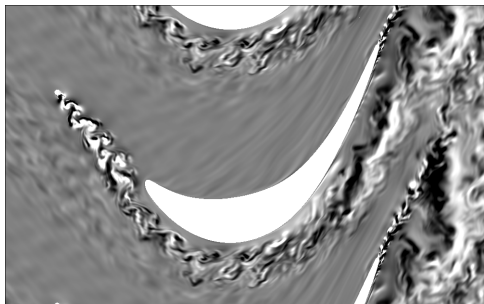


LES image

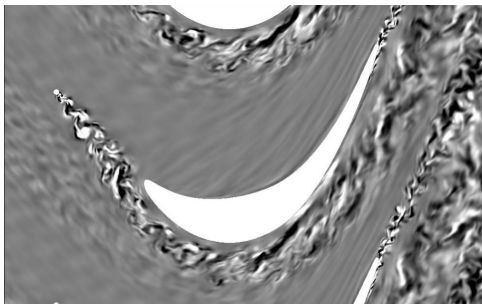


GAN image

LPT Stator - Results III

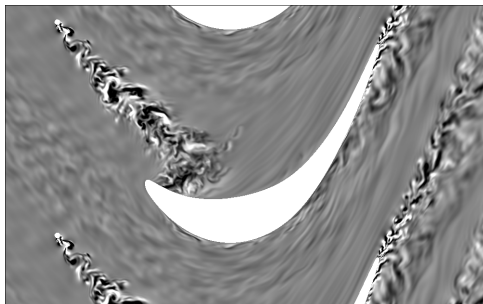


LES image

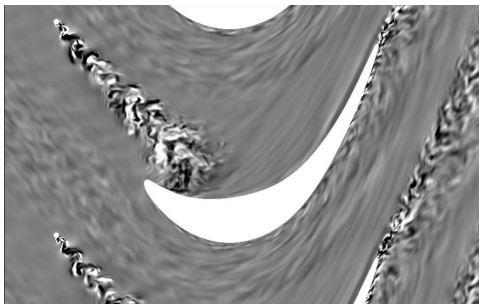


GAN image

LPT Stator - Results IV



LES image



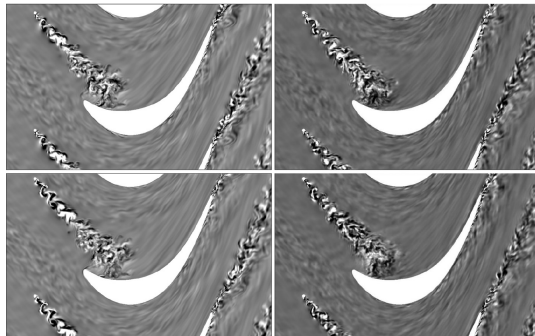
GAN image

Computational Costs

- ▶ **LES:** Performed on 560 CPU cores of the CPU type Intel Xeon S"SkyLake" Gold 6132
 - Karman vortex street (5 000 images): 72 core weeks $\hat{=}$ 21.6h
 - LPT Stator (2250 images): 640 core weeks $\hat{=}$ 192h
- ▶ **GAN-Training:** Performed on GPU of type Quadro RTX 8000 with 48 GB
 - Karman vortex street (DCGAN, 2 000 epochs): 1.5 min/epoch = 50h
 - LPT Stator (pix2pixHD, 200 epochs): 17 min/epoch = 56h
- ▶ **GAN-Inference:** Performed on GPU of type Quadro RTX 8000 with 48 GB
 - Karman vortex street (DCGAN): 0.001 sec/image $\Rightarrow$ 5 seconds
 - LPT Stator (pix2pixHD): 0.01 sec/image $\Rightarrow$ 22.5 seconds

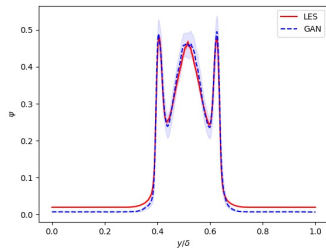
Sampling results for the LPT

Generalization for Unseen Geometry Configurations



- ▶ Repeat experiment with 5% of the images omitted around specific wake position
- ▶ During inference, position the wake at exactly the middle of that position

Statistical Similarity vs Physics



- Measure strength of turbulence (variation of $c = |u|$ as a function of ξ)

$$\text{Var}[c(\xi)] = \lim_{T \rightarrow \infty} \frac{1}{T} \int_0^T |c(\xi, t) - \bar{c}(\xi)|^2 dt$$

Estimation of Physical Quantities

- ▶ All statistical physical quantities of the turbulent flow can be expressed through evaluation of functionals $\Psi(x)$ and averaging

$$\bar{\Psi} = \mathbb{E}_{X \sim \mu}[\Psi(X)]$$

- ▶ We can assume that $\Psi(x)$ is bounded on Ω . Then estimate $\bar{\Psi}$ by $\hat{\Psi}$

$$\hat{\Psi} = \mathbb{E}_{Z \sim \nu}[\Psi(\hat{\phi}(Z))] = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{j=1}^N \Psi(\hat{\phi}(U_j)) \quad \text{a.s.} \Rightarrow$$

$$|\bar{\Psi} - \hat{\Psi}| \leq C \sup_{\|\Psi'\| \leq 1} \left| \int \Psi' d\mu - \int \Psi' d\hat{\phi}_* \nu \right| = C \|\mu - \hat{\phi}_* \nu\|_{TV}.$$

Convergence of Physical Quantities

► **Lemma (Drygala, Winhard, di Mare, G. 2021):**

$$|\bar{\Psi} - \hat{\Psi}| \rightarrow 0 \quad \text{as} \quad n \rightarrow \infty$$

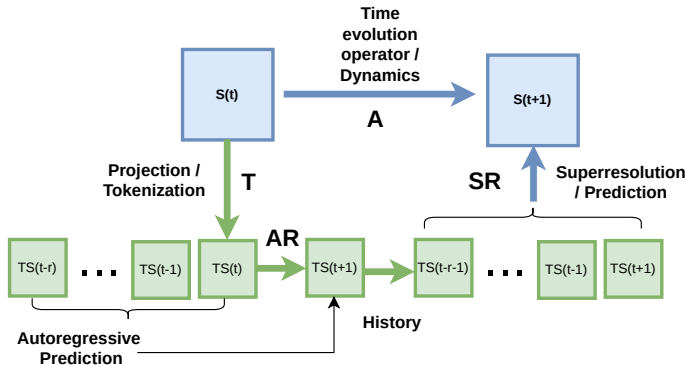
► **Proof:** Enough to show: $\|\mu - \hat{\phi}_* \nu\|_{TV} \rightarrow 0$. $\hat{\mu} = \hat{\phi}_* \nu$.

$$\begin{aligned} \|\hat{\mu} - \mu\|_{TV} &\leq \left\| \hat{\phi}_* \nu - \frac{\mu + \hat{\mu}}{2} \right\|_{TV} + \left\| \mu - \frac{\mu + \hat{\mu}}{2} \right\|_{TV} \\ &\leq \sqrt{2\mathfrak{d}_{KL} \left(\hat{\mu} \left\| \frac{\mu + \hat{\mu}}{2} \right\| \right)} + \sqrt{2\mathfrak{d}_{KL} \left(\mu \left\| \frac{\mu + \hat{\mu}}{2} \right\| \right)} \\ &\leq 2\sqrt{\mathfrak{d}_{KL} \left(\hat{\mu} \left\| \frac{\mu + \hat{\mu}}{2} \right\| \right) + \mathfrak{d}_{KL} \left(\mu \left\| \frac{\mu + \hat{\mu}}{2} \right\| \right)} = 2\sqrt{\mathfrak{d}_{JS}(\hat{\mu} \parallel \mu)} \rightarrow 0. \end{aligned}$$

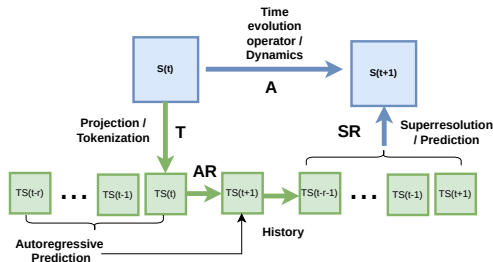
Learning Dynamics

So far so good - but this is no flow

- ▶ Try to use **Word Models** for flow generation
- ▶ Can we find a mathematical justification?



Observability of Dynamical Systems



- ▶ The autoregressive dynamic time step **AR** exists, if and only **SR** exists.
- ▶ In system theory, the existence of **SR** is known as **observability**!

$$\mathbf{SR}(TS, TA^{-1}S, TA^{-2}S, \dots, TA^{-r}S) = S$$

When is a Linear Dynamical System Observable wrt a 'Tokenization'?

- If for every ξ be an eigen vector of the linear dynamic map A

$$T\xi \neq 0$$

then the dynamical system is observable.

- **Example:** The discretized heat equation $A = e^{-t\Delta_\Gamma}$ is **not** observable wrt average pooling = projection to constant functions on pixel squares as 'tokenization'.

First Results on AR Operator Learning - Heat Equation

tokenized heat equation with random conductivity

$$\partial_t u(t, x) - \nabla \cdot a(x, \omega) \nabla u(t, x) = 0$$

First Results on AR Operator Learning - Heat Equation

reconstructed heat equation with random conductivity

$$\partial_t u(t, x) - \nabla \cdot a(x, \omega) \nabla u(t, x) = 0$$

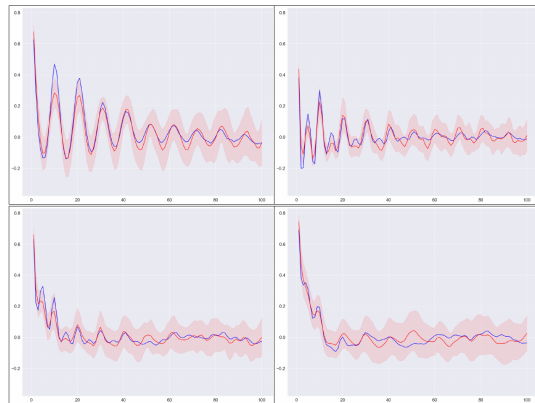
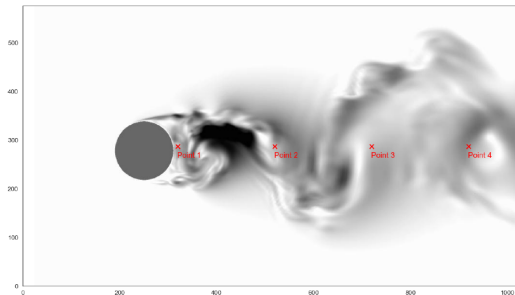
First Results on AR Operator Learning - Wave

tokenized wave equation with random Laplacian

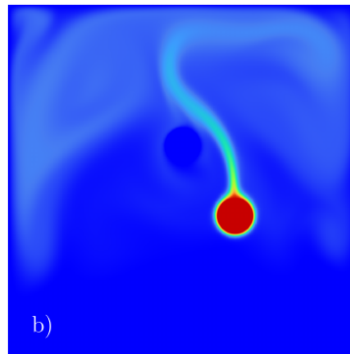
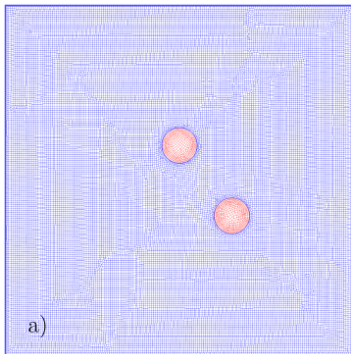
$$\partial_t^2 u(t, x) - \nabla \cdot a(x, \omega) \nabla u(t, x) = 0$$

VideoGAN Results with NVIDIA LongVideoGAN

Physical Evaluation of Turbulence Autocorrelations



Parametric Operator Learning

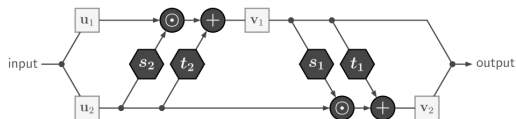


Numerical Results

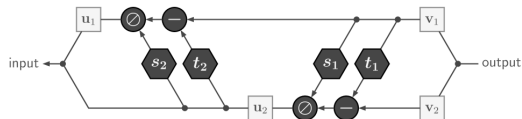
Invertible Neural Networks - INN

Forward Direction ():

Ardizzone et al., 2018

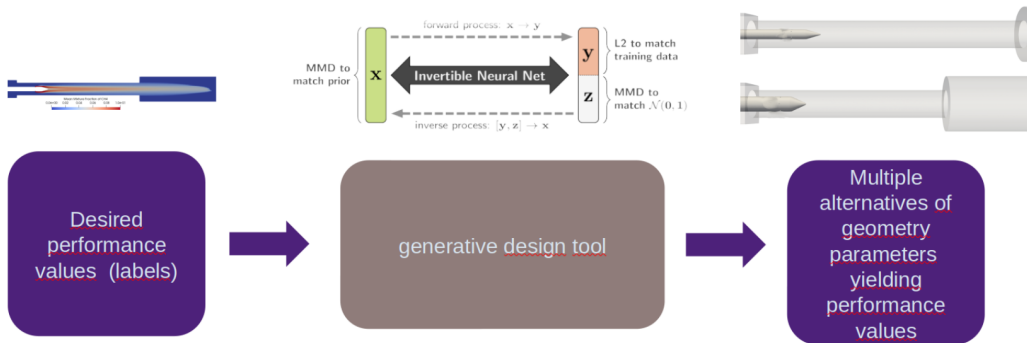


Backward Direction ():

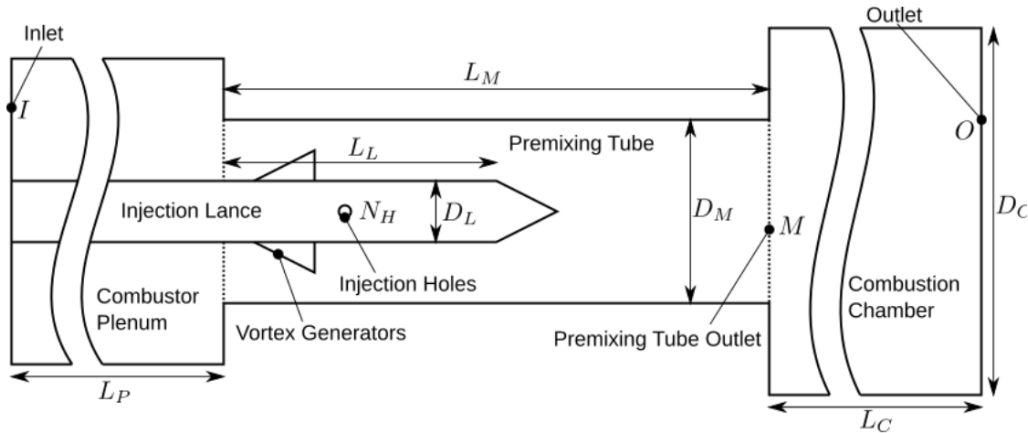


- ▶ Generative learning relates to inverse design (*Ardizzone et al. 2019*)

Case study for a Multi Fuel Burner

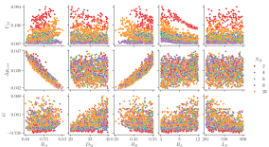


Design Parameters



Workflow Overview

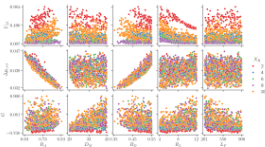
1) Create Combustor Dataset



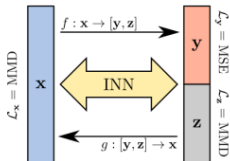
- Define combustor geometries by independent parameter vectors x
- Define performance label vectors y
- Sample designs and generate Labels via CFD and Acoustic Network Simulation

Workflow Overview

1) Create Combustor Dataset



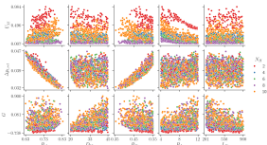
2) Train INN on Combustor Dataset



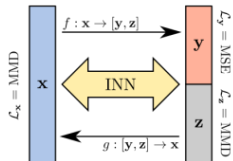
- INN Motivation & Concept
- Surrogate models for quick validation & data augmentation
- Ablation studies on surrogate model training and data augmentation size

Workflow Overview

1) Create Combustor Dataset

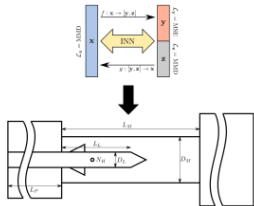


2) Train INN on Combustor Dataset



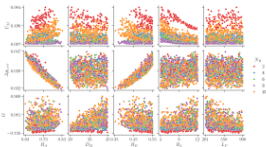
- Define targets for data generation
- Generate data
- Inspect distributions of generated data

3) Generate Designs

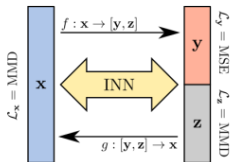


Workflow Overview

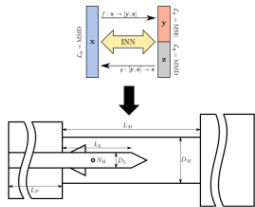
1) Create Combustor Dataset



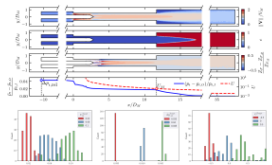
2) Train INN on Combustor Dataset



3) Generate Designs

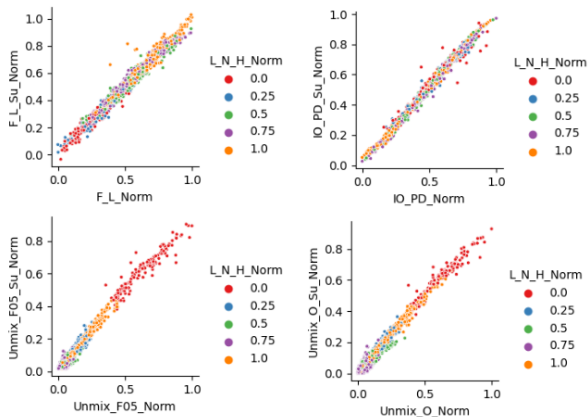


4) Validate Generated Designs

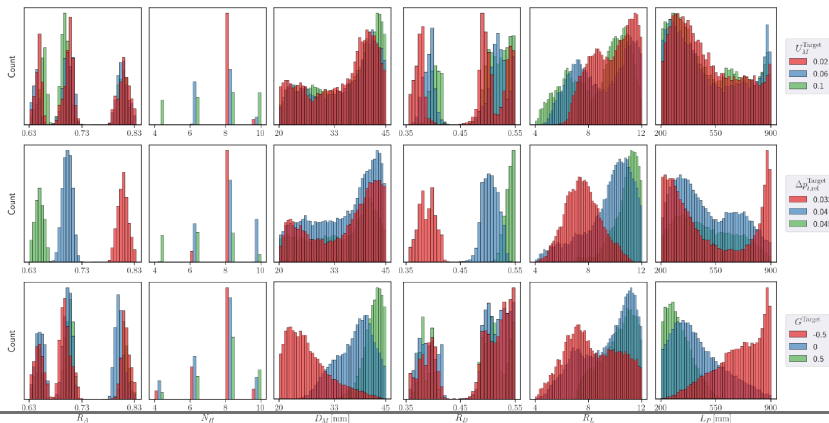


- Obtain true label values for generated designs
- Compare targets with obtained labels
- Outlook on label discretization

Trained INN - Forward Mode



Trained INN - Design Parameter Diversity in Backward Mode



The future I

The future II

Generative AI for numerical math – résumé

- ▶ Generative learning (small models) can be used for an **interactive** design process
- ▶ Generatives learning (large models) make simulations **real time capable**
- ▶ Generative AI is **no black box**, but one can and should understand the algorithms behind (and their properties)

