



REAL-TIME ALGORITHMS FOR DIGITAL TWINS

Regression-based Modelling

UKACM Autumn School 2026
Prof. Dirk Hartmann



1

Regression-based Modelling – Introduction

2

Singular Value Decomposition & Proper Orthogonal Decomposition

3

Operator Inference

4

Neural Networks

Task [Regression]:

Regression aims to identify the relationship between independent variables $X = (x_1, \dots, x_n)$ and dependent variables $Y = (y_1, \dots, y_n)$ and some unknown parameters β , i.e.,

$$Y = f(X, \beta),$$

where the regression function $f(\cdot)$ is typically prescribed and the parameters β are found by optimizing the *goodness-of-fit* of this function to the data.

- **Goodness-of-fit:** Maximum Error l_∞ : $E_\infty(f) = \max_{1 \leq k \leq n} |f(x_k) - y_k|$
Mean Absolute Error l_1 : $E_1(f) = \frac{1}{n} \sum_{k=1}^n |f(x_k) - y_k|$
Least Squares Error l_2 : $E_2(f) = \left(\frac{1}{n} \sum_{k=1}^n |f(x_k) - y_k|^2 \right)^{1/2}$

- ▶ Curve fitting is the most basic regression technique usually resulting in solutions that com from solving a linear system

$$\mathbf{A} \cdot \mathbf{x} = \mathbf{b}$$

Example [Least Squares Polynomial Fitting]

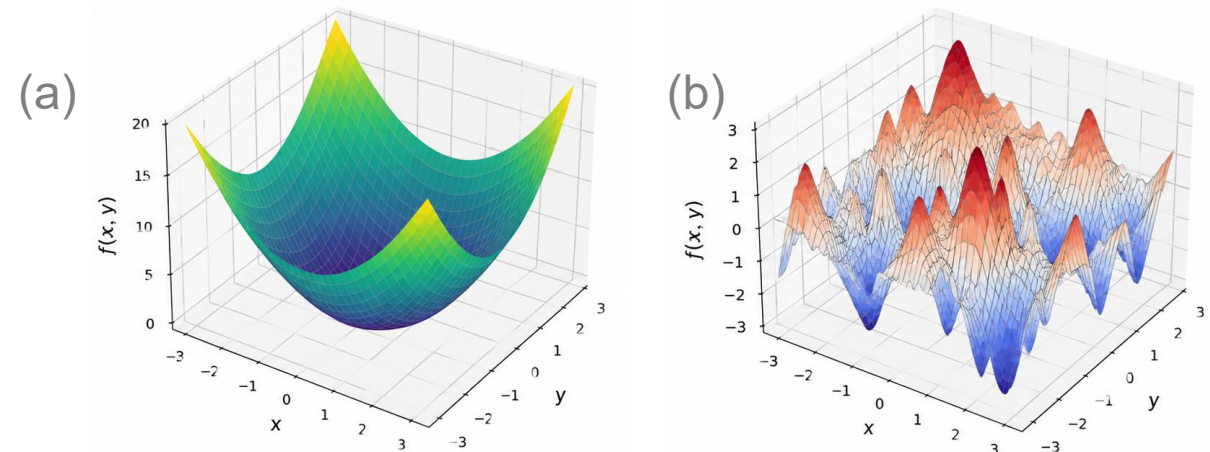
Given data (x_k, y_k) choose the parameter β to fit the curve $f(x) = \beta_0 + \beta_1 x + \dots + \beta_m x^m$ such that $E_2(f) = \sum_{k=1}^n |f(x_k) - y_k|^2$ is minimal.

- ▶ In the minimum it holds $\partial \beta_i E_2(f) = 0$ for $0 \leq i \leq m$, which leads to a linear system $\mathbf{A} \cdot \boldsymbol{\beta} = \mathbf{b}$, which can be solved effectively, i.e., do not need any optimization algorithm.
- ▶ Least-squares fitting has the critical advantage, that the optimization is often inexpensive.

Regression

Nonlinear Curve Fitting

- ▶ Choosing a non-linear function , e.g., $f(x) = \beta_1 \cos(\beta_2 x + \beta_3) + \beta_4$, the system to be solved becomes nonlinear and **appropriate optimization algorithms** are required.
- ▶ Nonlinear systems can have no solutions, several solutions, as well as an infinite number of solutions. Contrary to linear systems, the existence and potential uniqueness is more difficult to check.
- ▶ Convex functions are a class of functions where corresponding statements can be made. They also guarantee convergence, while non-convex functions can have many pitfalls



Two objective function landscapes representing (a) a convex function and (b) a non-convex function.

- In classical regression, we often work with underdetermined or overdetermined (linear) systems leading to the following optimization problems

$$\mathbf{x} = \underset{\mathbf{x}}{\operatorname{argmin}} (\|\mathbf{A} \cdot \mathbf{x} - \mathbf{b}\|_2 + \lambda g(\mathbf{x})) \quad [\text{Overdetermined System}]$$

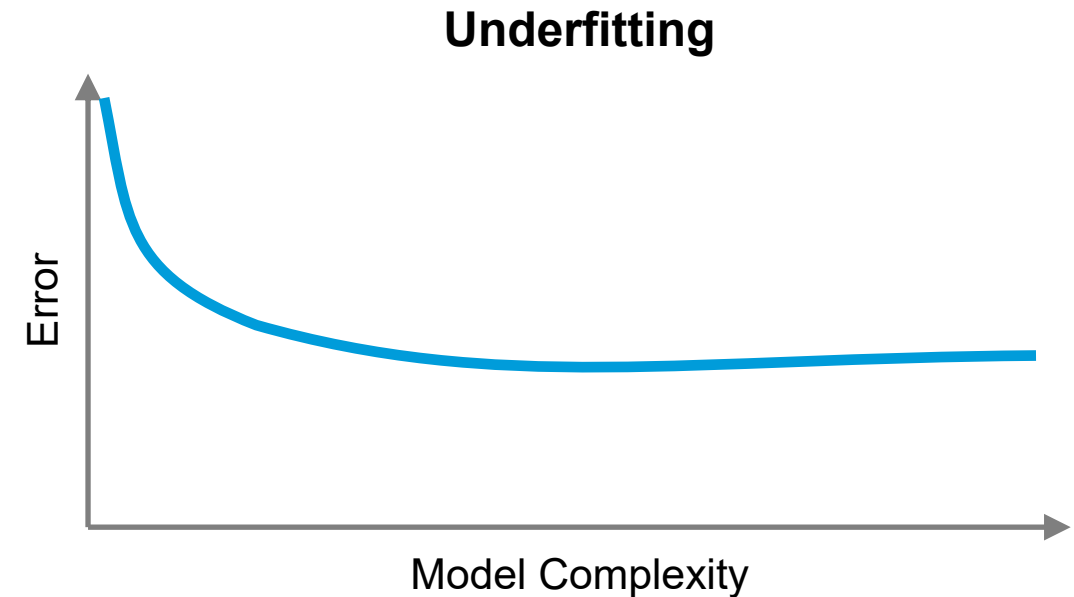
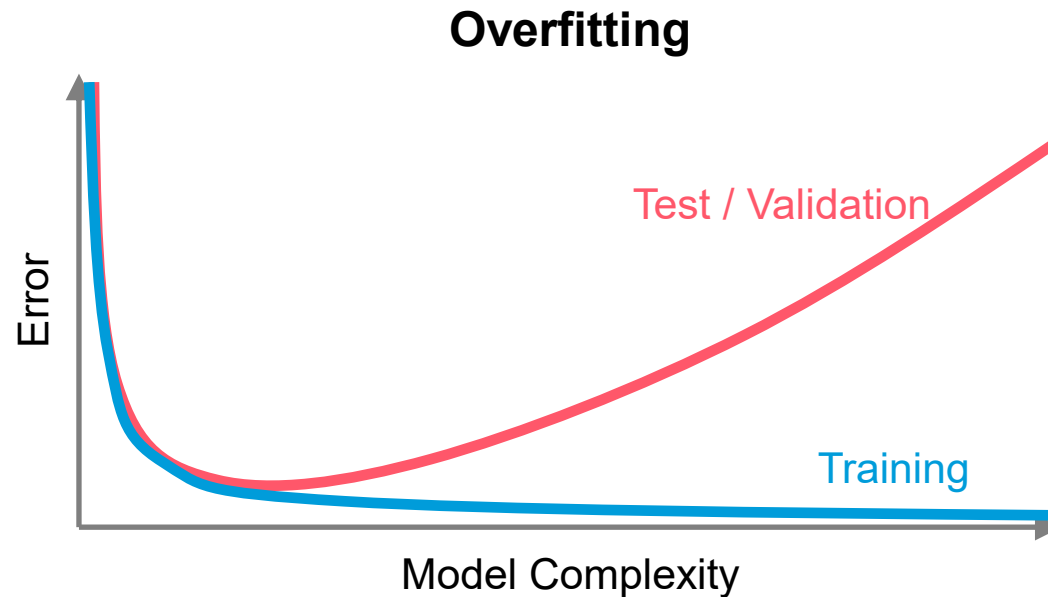
$$\text{or } \mathbf{x} = \underset{\mathbf{x}}{\operatorname{argmin}} g(\mathbf{x}) \quad \text{subject to } \|\mathbf{A} \cdot \mathbf{x} - \mathbf{b}\|_2 \leq \epsilon, \quad [\text{Underdetermined System}]$$

where $g(\mathbf{x})$ is a given penalization – also known as regularization – with penalty parameter λ for overdetermined systems, e.g., $g(\mathbf{x}) = \|\mathbf{x}\|_2$

- Often the exact family of curve is not known, e.g., the detailed polynomial degree, and optimization methods are used to choose the best model.
- A concern thus is to not **overfit** or **underfit** the data.

Model Selection

Overfitting and Underfitting



*Prototypical behavior of over- and underfitting data. **Overfitting:** Increasing the number of training iterations leads to an improved reduction on the training data while increasing the error on the validation / test data. **Underfitting:** the error performance is limited due to the restrictions on model complexity,*

Model Selection

Overfitting and Underfitting

► How to choose the right model?

- **Occam's razor / Lex Parsimoniae:** Among competing hypotheses, the one with the fewest assumptions should be selected, or when you have two competing theories that make exactly the same predictions, the simpler one is the more likely.

- In regression, we optimize

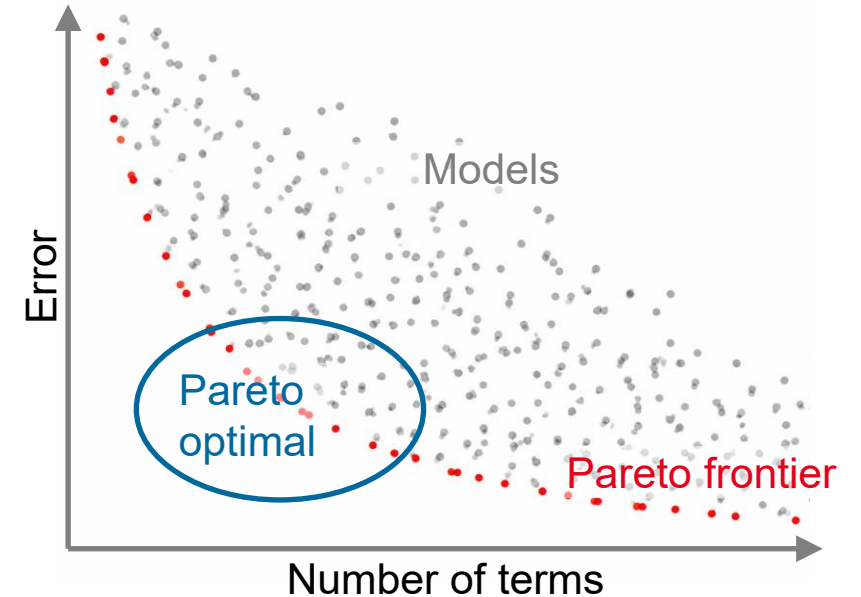
$$\underset{x}{x} = \operatorname{argmin}(\|A \cdot x - b\|_2),$$

which does not explicitly enforce any constraints on x .

- To have some control for avoiding under or over fitting one typically controls the “complexity” of x by appropriate regularization, i.e.,

$$\underset{x}{x} = \operatorname{argmin} \|A \cdot x - b\|_2 + \lambda_1 \|x\|_1 + \lambda_2 \|x\|_2$$

where the parameters λ_1 and λ_2 control the penalization of norms.





1

Regression-based Modelling – Introduction

2

Singular Value Decomposition & Proper Orthogonal Decomposition

3

Operator Inference

4

Neural Networks

Singular Value Decomposition

Introduction

- ▶ Singular Value Decomposition is among the most important matrix factorizations in the computational era. It is a cornerstone of science and engineering.¹ It is the basis for many techniques in dimension reduction.
- ▶ In the context of Data-Driven methods we are interested in analyzing large data sets $X \in \mathbb{R}^{n \times m}$:

$$X = \begin{bmatrix} | & & | \\ \vec{x}_1 & \dots & \vec{x}_m \\ | & & | \end{bmatrix}$$

where each data point $1 \leq i \leq m$ is an n -dimensional vector, i.e., $\vec{x}_i \in \mathbb{R}^n$. If not given naturally in this form, many data can be transferred to this.

- ▶ Often the individual data points will represent simulation snapshots – a time series of simulation data. That is, $\vec{x}_k = \vec{x}(k\Delta t)$ and $n \gg m$.

1) GW Stewart (1993): [On the Early History of the Singular Value Decomposition](#). SIAM Review

Singular Value Decomposition

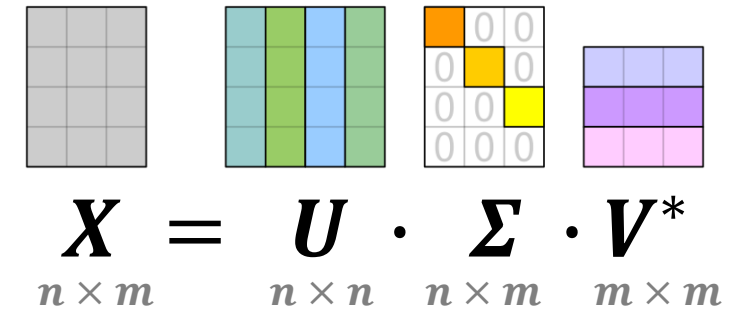
Definition

Singular Value Decomposition (SVD): The SVD is a unique matrix decomposition that exists for every complex-valued matrix

$X \in \mathbb{C}^{n \times m}$ such that

$$X = U \cdot \Sigma \cdot V^*.$$

Here, $U \in \mathbb{C}^{n \times n}$ and $V \in \mathbb{C}^{m \times m}$ are unitary matrices with orthonormal columns / rows and $\Sigma \in \mathbb{R}^{n \times m}$ is a real non-negative diagonal matrix.



$$\begin{matrix} X & = & U & \cdot & \Sigma & \cdot & V^* \\ n \times m & & n \times n & & n \times m & & m \times m \end{matrix}$$

- The columns of U are called the left singular vectors of X , the columns of V singular vectors of X , and entries in Σ singular values (usually arranged in decreasing order $\sigma_1 \geq \sigma_2 \geq \dots \sigma_m \geq 0$)

$$X = \sum_{k=1}^m \sigma_k \vec{u}_k \vec{v}_k.$$

- The SVD is essentially a separation of variables between the columns and rows of a data matrix, where U *encodes spatial patterns while V encodes temporal patterns.*

Singular Value Decomposition

Computation

- ▶ The SVD is numerically stable and albeit similarities to eigen decomposition is guaranteed to exist for any matrix.
- ▶ Efficient implementations for SVD are available in nearly all modern computer languages, thus we will take this for granted. (Albeit the specific implementations are important and also enlightening).

```
43
44  #Assemble the data matrix
45   $\varphi$ ,  $\varphi_t$ ,  $\varphi_v$  = assemble_data_matrices(sol);
46
47  #Perform an SVD
48   $U, \Sigma, V$  = svd( $\varphi_t$ );
49
50  #Analyze the SVD / Plot SVD information
51  plot_svd_spectrum("Fig\\svd_spectrum.png",  $\Sigma$ );
52  plot_pod_error("Fig\\pod_error.png",  $U$ ,  $\varphi_t$ ,  $\varphi_v$ )
53
```

Singular Value Decomposition

Optimal low-rank Matrix Approximations (i)

- SVD provides a hierarchy of **optimal low-rank approximations** $\tilde{X}$ to the matrix X , keeping the leading r values and vectors and disregarding the rest (often σ_k decrease rapidly):

$$\tilde{X} = \sum_{k=1}^r \sigma_k \vec{u}_k \vec{v}_k.$$

Theorem [Eckart-Young]: The optimal rank- r approximation to X is given by the rank- r SVD truncation $\tilde{X}$:

$$\underset{\tilde{X}, \text{ s.t. } \text{rank}(\tilde{X})=r}{\text{argmin}} \quad \|X - \tilde{X}\|_F = \tilde{U} \cdot \tilde{\Sigma} \cdot \tilde{V}^*,$$

with the Frobenius norm $\|X\|_F = \left(\sum_{i=1}^n \sum_{j=1}^m |x_{ij}|^2 \right)^{1/2}$.

- It is possible to exactly quantify the error, i.e., $\|X - \tilde{X}\|_F^2 = \sum_{k=r+1}^m \sigma_k^2$.
- The corresponding relative error, can be interpreted as the fraction of kinetic energy missing in the approximation $\tilde{X}$.

Singular Value Decomposition

Optimal low-rank Matrix Approximations (ii)

- ▶ The SVD also provides an optimal rank- r approximation in the matrix 2 norm / Spectral norm, i.e.,

$$\underset{\tilde{X}, \text{ s.t. } \text{rank}(\tilde{X})=r}{\text{argmin}} \quad \|\mathbf{X} - \tilde{\mathbf{X}}\|_2 = \tilde{\mathbf{U}} \cdot \tilde{\mathbf{\Sigma}} \cdot \tilde{\mathbf{V}}^*$$

$$\text{with } \|\mathbf{X}\|_2 = \max_{\vec{v} \neq 0} \frac{\|\mathbf{X} \cdot \vec{v}\|_2}{\|\vec{v}\|_2}$$

- ▶ The error expression is rather simple

$$\|\mathbf{X} - \tilde{\mathbf{X}}\|_2 = \sigma_{r+1}$$

which follows from the expansion of the SVD and orthogonality of vectors $\vec{v}_k$.

- ▶ There are many applications of SVD, e.g.,
 - **Linear Algebra:** Improve the condition number of matrices (effectively cutting off the small singular values)
 - **Data Analysis:** Data transformation to more interpretable representations (Principal Component Analysis)
 - **Reduced Order Modelling:** Identification of dominant dynamical modes in which the system can be represented
 - ...

...

Proper Orthogonal Decomposition (POD)

Proper Orthogonal Decomposition:

Given a dynamical system

$$\partial_t \vec{x}(t) = \mathbf{A} \cdot \vec{x}(t) + \vec{F}(\vec{x}(t)) + \mathbf{B}c(t)$$

and trajectory data $\mathbf{X} = [\vec{x}(t_0), \dots, \vec{x}(t_m)]$, we use SVD $\mathbf{X} = \mathbf{U} \cdot \mathbf{\Sigma} \cdot \mathbf{V}^T$ to obtain the dominant orthogonal components $\tilde{\mathbf{U}}$. We then model the system only in these dominant modes, i.e., $\vec{x}(t) = \sum_{k=1}^r a_k(t) \vec{u}_k$.

- ▶ Typically $\vec{x} \in \mathbb{R}^n$ with $n \gg 1$ obtained from the discretization of a PDE, i.e., the individual entries represent nodal values.
- ▶ Once the modal basis is decided the governing equations for the $a_k(t)$ can be determined by projection of the full system. The corresponding Reduced Order Model is then

$$\partial_t \tilde{\vec{x}}(t) = \tilde{\mathbf{A}} \cdot \tilde{\vec{x}}(t) + \tilde{\vec{F}}(\tilde{\vec{x}}(t)) + \tilde{\mathbf{B}}u(t)$$

with $\tilde{\vec{x}}(t) = \mathbf{U}^* \cdot \vec{x}(t)$, $\tilde{\mathbf{A}} = \mathbf{U}^* \cdot \mathbf{A} \cdot \mathbf{U}$, $\tilde{\mathbf{B}} = \mathbf{U}^*$, and $\tilde{\vec{F}}(\tilde{\vec{x}}(t)) = \mathbf{U}^* \cdot \vec{F}(\mathbf{U} \cdot \tilde{\vec{x}}(t))$.

Proper Orthogonal Decomposition (POD)

- ▶ Reduced Order Model:

$$\partial_t \vec{\tilde{x}}(t) = \boxed{\tilde{\mathbf{A}} \cdot \vec{\tilde{x}}(t)} + \boxed{\vec{\tilde{F}}(\vec{\tilde{x}}(t))} + \boxed{\tilde{\mathbf{B}} u(t)}$$

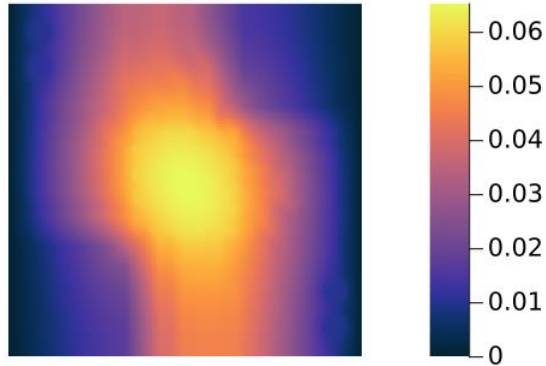
with $\vec{\tilde{x}}(t) = \mathbf{U}^* \cdot \vec{x}(t)$, $\tilde{\mathbf{A}} = \mathbf{U}^* \cdot \mathbf{A} \cdot \mathbf{U}$, $\tilde{\mathbf{B}} = \tilde{\mathbf{U}}^*$, and $\vec{\tilde{F}}(\vec{\tilde{x}}(t)) = \mathbf{U}^* \cdot \vec{F}(\tilde{\mathbf{U}} \cdot \vec{\tilde{x}}(t))$.

- ▶ If we have access to the full system / simulation code $\tilde{\mathbf{A}} = \mathbf{U}^* \cdot \mathbf{A} \cdot \mathbf{U}$ can be determined by simple matrix multiplication. However, the reduction of $\vec{\tilde{F}}(\vec{\tilde{x}}(t)) = \mathbf{U}^* \cdot \vec{F}(\tilde{\mathbf{U}} \cdot \vec{\tilde{x}}(t))$ might be very costly.
- ▶ Under certain circumstances efficient methods for reduction / approximation of $\vec{\tilde{F}}(\vec{\tilde{x}}(t)) = \mathbf{U}^* \cdot \vec{F}(\tilde{\mathbf{U}} \cdot \vec{\tilde{x}}(t))$ are available, these are referred to hyper reduction methods, one of these being *Discrete Empirical Interpolation (DEIM)* leveraging concepts from sparse sensing.

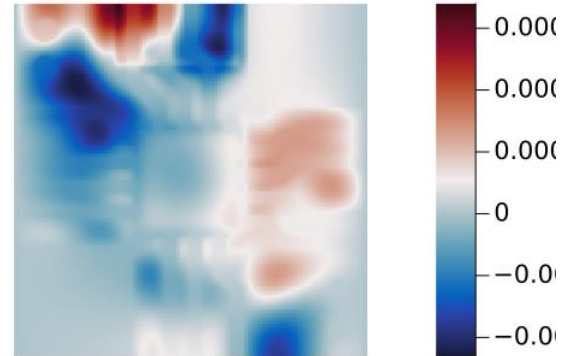
Reduced Order Modelling

Proper Orthogonal Decomposition + Galerkin Projection

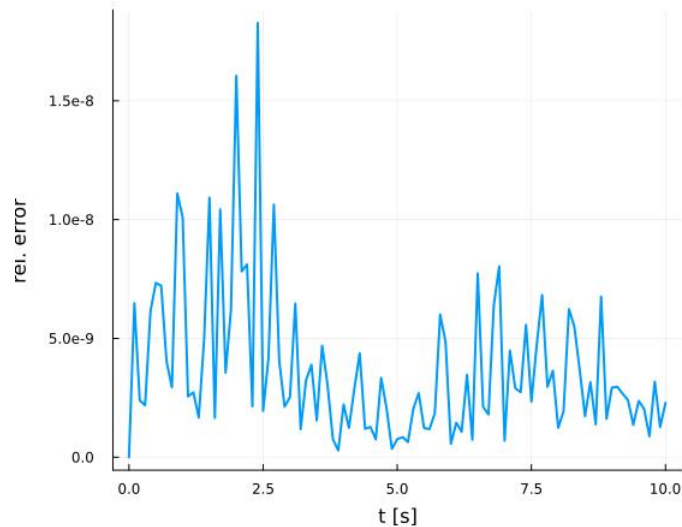
T POD ROM ($t = 10.0$ s, $r = 15$)



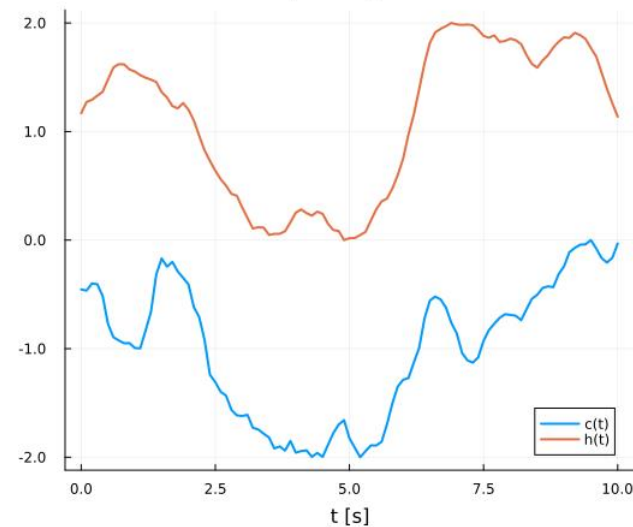
difference to full solution



ROM error over time



input signals



[https://github.com/Dirk-Hartmann/
real-time-digital-twins/blob/main/
scripts/rom/PCB_PODGalerkinROM.jl](https://github.com/Dirk-Hartmann/real-time-digital-twins/blob/main/scripts/rom/PCB_PODGalerkinROM.jl)



1

Regression-based Modelling – Introduction

2

Singular Value Decomposition & Proper Orthogonal Decomposition

3

Operator Inference

4

Neural Networks

Proper Orthogonal Decomposition:

Given a dynamical system $\partial_t \vec{x}(t) = \mathbf{A} \cdot \vec{x}(t) + \vec{F}(\vec{x}(t)) + \mathbf{B}c(t)$ and trajectory data $\mathbf{X} = [\vec{x}(t_0), \dots, \vec{x}(t_m)]$, we use SVD $\mathbf{X} = \mathbf{U} \cdot \mathbf{\Sigma} \cdot \mathbf{V}^T$ to obtain the dominant orthogonal components $\tilde{\mathbf{U}}$. We then model the system only in these dominant modes, i.e., $\vec{x}(t) = \sum_{k=1}^r a_k(t) \vec{u}_k$.

Operator Inference (OI):

Learn a quadratic model, i.e.,

$$\partial_t \tilde{\vec{x}}(t) = \tilde{\mathbf{A}} \cdot \tilde{\vec{x}}(t) + \tilde{\mathbf{H}} \cdot (\tilde{\vec{x}}(t) \otimes \tilde{\vec{x}}(t)) + \tilde{\mathbf{B}}u(t)$$

with $\tilde{\vec{x}}(t) = \mathbf{U}^* \cdot \vec{x}(t)$, etc as usual and $P \otimes Q \equiv [p_{ij}Q]_{ij}$, via a regression problem

$$(A, H, B) = \underset{A, H, B}{\operatorname{argmin}} \left\| \partial_t X - [A \ H \ B] \cdot \begin{bmatrix} X \\ X * X \\ U \end{bmatrix} \right\|_F^2 + \mathcal{R}(A, H, B)$$

with the Khatri-Rao product $X * X := [x_0 \otimes x_0, \dots, x_K \otimes x_K]$ and an appropriate regularization $\mathcal{R}(A, H, B)$.

- ▶ OI yields an optimal approximation in the L_2 obtainable from the available data

Theorem [Willcox and Peherstorfer 2016]

In the limit of an infinite amount of stat, the operator inference model of order r recovers the POD model of order r in the case of a quadratic model.

- ▶ OI can be combined with hyper-reduction, specifically with Discrete Empirical Interpolation Method (DEIM)¹.
- ▶ Using the concept of lifting, one can cover any polynomial model, not only quadratic ones².
- ▶ Specific constraints on the matrices can be imposed to enforce stability^{2,3}, e.g., ensuring that eigenvalues of A are nonnegative.

Sources: 1) P Benner, P Goyal, B Kramer, B Peherstorfer, K Willcox (2020): [Operator inference for non-intrusive model reduction of systems with non-polynomial nonlinear terms](#). CMAME
2) E Qian, B Kramer, B Peherstorfer, K Willcox (2020): [Lift & Learn: Physics-informed machine learning for large-scale nonlinear dynamical systems](#). Physica D: Nonlinear Phenomena
3) P Goyal, I Pontes Duff, P Benner (2023): [Inference of Continuous Linear Systems from Data with Guaranteed Stability](#). arXiv
4) P Goyal, I Pontes Duff, P Benner (2023): [Guaranteed Stable Quadratic Models and their applications in SINDy and Operator Inference](#). arXiv

Challenges:

in addition to the challenges associated with the SVD

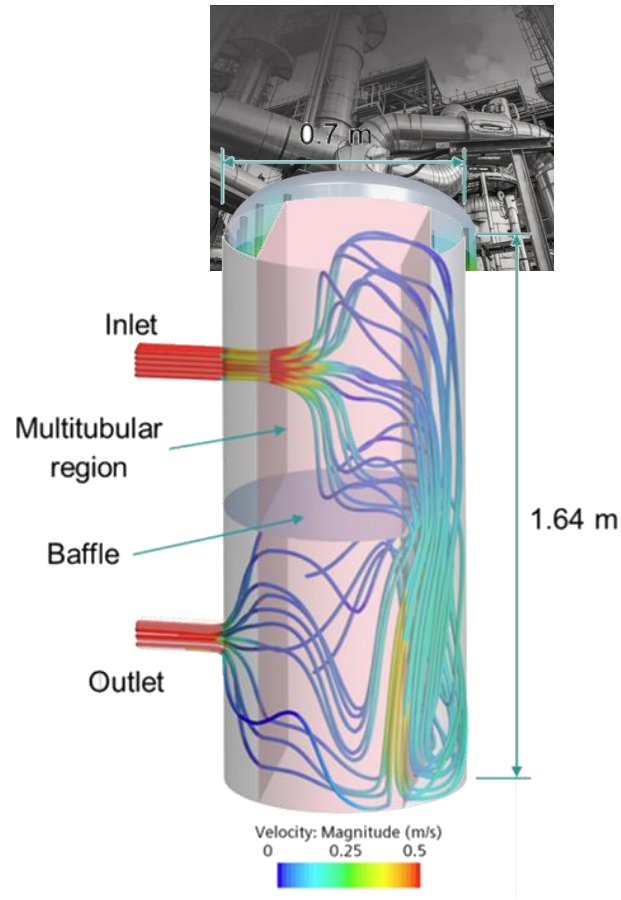
- ▶ Good estimates of time derivatives $\partial_t X$ in the data are crucial for success, since $[A, H, B] = \underset{A, H, B}{\operatorname{argmin}} \|\partial_t X - \dots\|_F^2$. In some circumstances very high approximation orders might be needed.
- ▶ The choice of the right regularization $\mathcal{R}(A, H, B)$ is challenging. Typically, this is done by automated hyper-parameter studies.
- ▶ Albeit the identification of the right operator structure should be straight forward, practice has shown that this is often a source of error due to limited knowledge of the underlying physics / systems.

Software:

- ▶ Availability of several open-source packages, e.g., the Operator Inference Libraries (for MATLAB and Python) are provided by Karen Willcox Lab, c.f., <https://github.com/Willcox-Research-Group>.

Operator Inference

Example (i)



Cooling flow velocities and pressure are modelled by Navier Stokes plus Brinkman law for the porous part

$$\rho_c \left(\frac{\partial \mathbf{v}}{\partial t} + (\mathbf{v} \cdot \nabla) \mathbf{v} \right) = -\nabla p + \mu \Delta \mathbf{v} + \boldsymbol{\alpha}(\mathbf{x}) \mathbf{v} \quad \text{in } \Omega$$

$$\nabla \cdot \mathbf{v} = 0 \quad \text{in } \Omega$$

Cooling fluid temperature and solid temperature are modelled as isotropic mediums

$$\rho_c c_{p,c} \left(\frac{\partial T_c}{\partial t} + (\mathbf{v} \cdot \nabla) T_c \right) = \nabla \cdot (\mathcal{K}_c \nabla T_c) + \chi_{\Omega_s} q(T_s - T_c) \quad \text{in } \Omega$$

$$\rho_s c_{p,s} \frac{\partial T_s}{\partial t} = \nabla \cdot (\mathcal{K}_s \nabla T_s) - q(T_s - T_c) + \mathcal{P}(t, T_s) \quad \text{in } \Omega_s$$

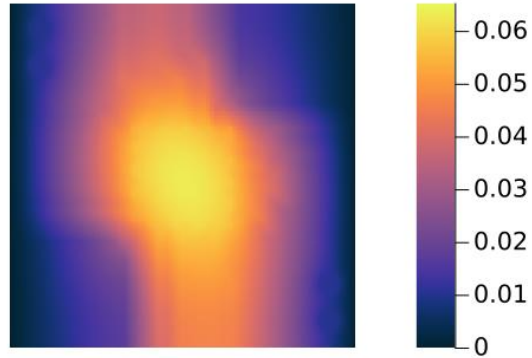
Data-based Hyper-Reduction Alternative Methods

- ▶ The field of Operator-Inference is a very active field of research with many more new developments, e.g., non-linear POD extensions, Exact Operator Inference.
- ▶ Operator-Inference is just one of many other options such as Dynamic Mode Decomposition or **Sparse** Identification of Nonlinear Dynamics (SINDy).

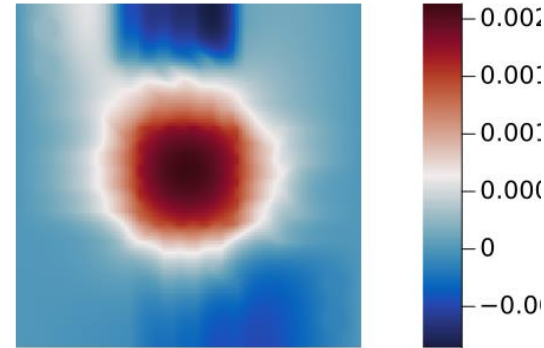
Reduced Order Modelling

Proper Orthogonal Decomposition + Operator Inference

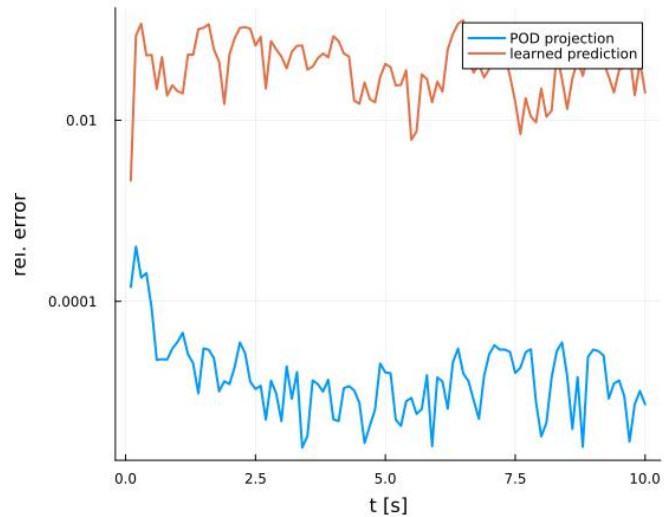
T operator inference ($t = 10.0$ s, $r = 8$)



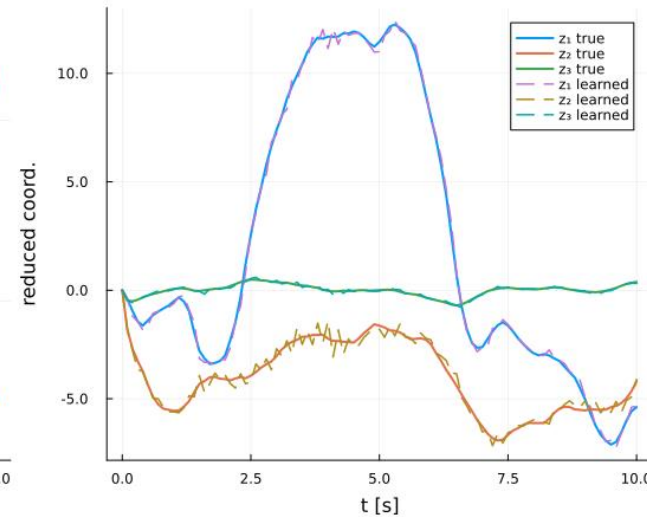
difference to full solution



ROM error over time



POD coordinates: true vs learned



[https://github.com/Dirk-Hartmann/
real-time-digital-twins/blob/main/scripts/
rom/PCB_OperatorInference.jl](https://github.com/Dirk-Hartmann/real-time-digital-twins/blob/main/scripts/rom/PCB_OperatorInference.jl)



1

Regression-based Modelling – Introduction

2

Singular Value Decomposition & Proper Orthogonal Decomposition

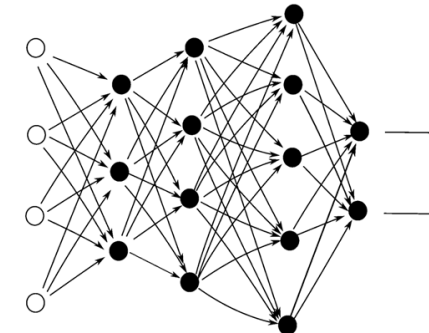
3

Operator Inference

4

Neural Networks

- ▶ The theoretical inception of Neural Networks has a more than 4 decades of history. However, only the emergence of effectively programable massively parallel compute units, i.e., CUDA in 2004, and the availability of abundant data, e.g., Image Net in 2009, made Neural Networks popular / scalable.
- ▶ Neural Networks were not even listed as one of the top 10 algorithms in data mining in 2008¹.
- ▶ Deep Learning is a loosely defined term and refers to Neural Networks with “many” layers.
- ▶ A **Neural Network** is a specific parameterized function $F(x, p)$, where $p \in \mathbb{R}$ indicates the parameters.
- ▶ The parameters p are then chosen such that the function regresses or classifies data optimally

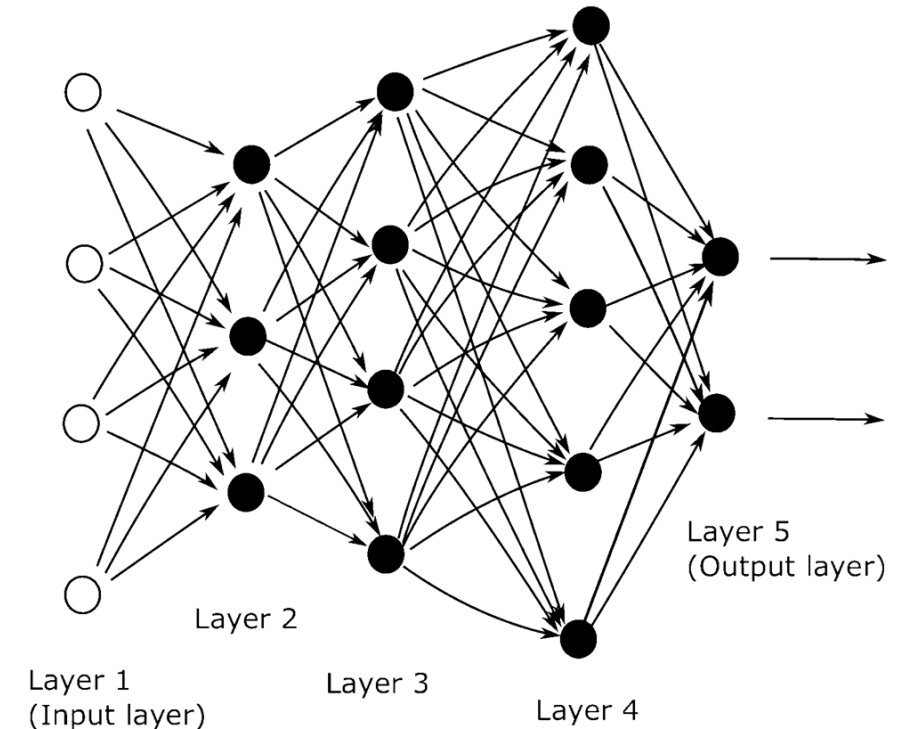


Sources: 1) X Wu, V Kumar, JR Quinlan, et al (2008): Top 10 algorithms in data mining. Knowledge and Information Systems
2) J Deng, W Dong, R Socher, LJ Li, K Li, L Fei-Fei (2009): ImageNet: A Large-Scale Hierarchical Image Database. IEEE Computer Vision and Pattern Recognition (CVPR)

Neural Networks

A simple example (i)

- ▶ A Neural Network has at least two but usually multiple layers L . Each layer l , for $l = 1, 2, \dots, L$, contains n_l neurons.
- ▶ The first layer $l = 1$ is called **Input Layer** and n_1 is the dimension of the input data. The layer $l = L$ is called the **Output Layer** and n_L is the dimension of the output data. All other layers are referred to as **Hidden Layers**. Overall, the network maps from $\mathbb{R}^{n_1}$ to $\mathbb{R}^{n_L}$.
- ▶ In each layer l , every neuron j outputs a single real number, which in a **densely connected Neural Network** is passed to every neuron in the next layer.
- ▶ This is a very small Neural Network, e.g., the Alex Net architecture which achieved groundbreaking image classification results in 2012 used 650,000 neurons and today transformer networks have billions of neurons.



Neural Networks

A simple example (ii)

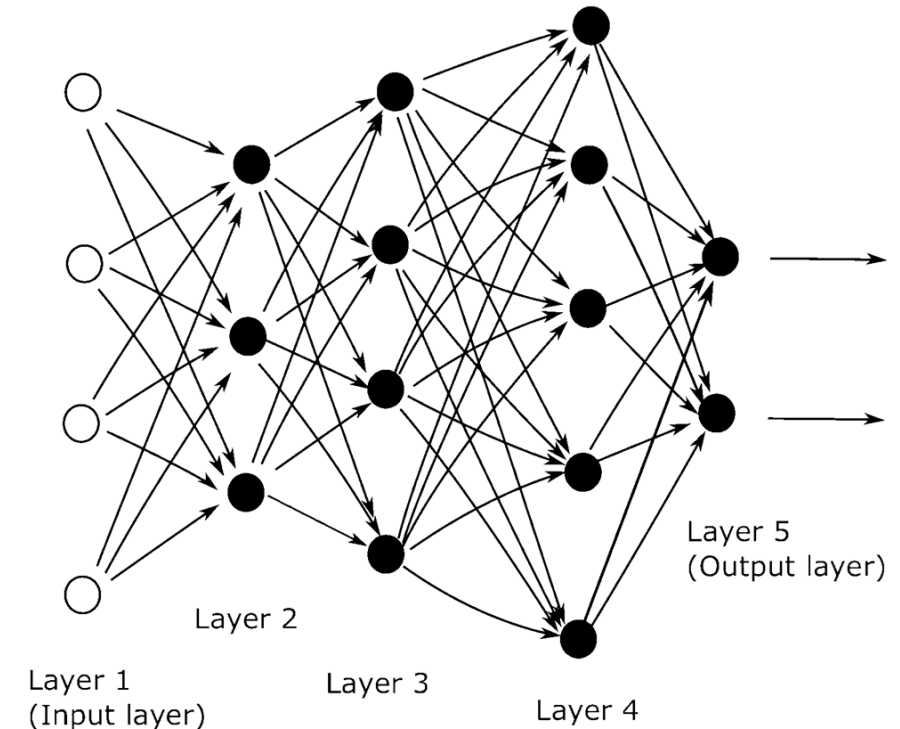
- ▶ Each neuron k in each layer l forms its own weighted linear combination of the inputs (**Weighted Input**) coming from the previous layer $l - 1$, i.e.,

$$z_j^{[l]} = \sum_k w_{jk} a_k + b_j$$

where $w_{jk}^{[l]}$ is the weight that neuron j at layer l applied to the output $a_k^{[l-1]}$ from neuron k at layer $l - 1$ and $b_j^{[l]}$ is an added bias.

- ▶ $z_j^{[l]}$ is then passed through the **Activation Function** σ providing for each neuron j of the layer l the corresponding output or activation

$$a_j^{[l]} = \sigma(z_j^{[l]})$$



Neural Networks

A simple example (iv)

Example:

Neural Network with 5 layers:

$$n_1 = 4, n_2 = 3, n_3 = 4, n_4 = 5, \text{ and } n_5 = 2$$

Thus,

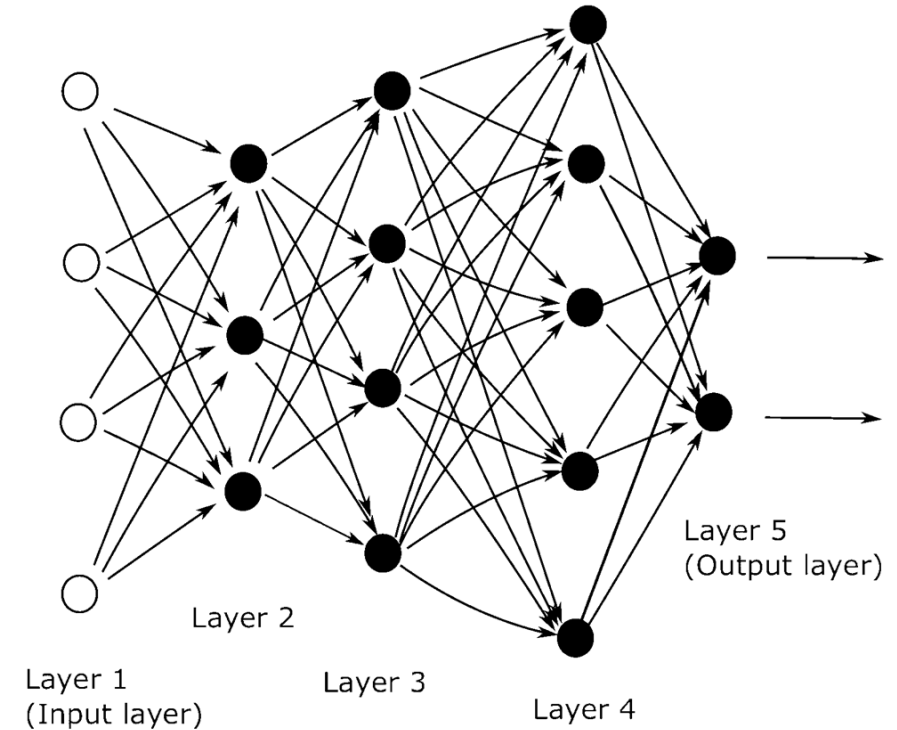
► $W^{[2]} \in \mathbb{R}^{3 \times 4}, W^{[3]} \in \mathbb{R}^{4 \times 3}, W^{[4]} \in \mathbb{R}^{5 \times 4}, W^{[5]} \in \mathbb{R}^{2 \times 5}$

► $b^{[2]} \in \mathbb{R}^3, b^{[3]} \in \mathbb{R}^4, b^{[4]} \in \mathbb{R}^5, b^{[5]} \in \mathbb{R}^2$

For $x \in \mathbb{R}^4$

$$F(x) = \sigma(W^{[5]}\sigma(W^{[4]}\sigma(W^{[3]}\sigma(W^{[2]}x + b^{[2]}) + b^{[3]}) + b^{[4]}) + b^{[5]}) \in \mathbb{R}^2$$

is a function of $3 \cdot (4 + 1) + 4 \cdot (3 + 1) + 5 \cdot (4 + 1) + 2 \cdot (5 + 1) = 74$ parameters



Neural Networks

An overview of activation functions



Typical activation functions $\sigma(x)$ used in Neural Networks

► **Linear:**

$$\sigma(x) = x$$

► **Sigmoid Function:**

$$\sigma(x) = \frac{1}{1 + \exp -x}$$

► **Tanh:**

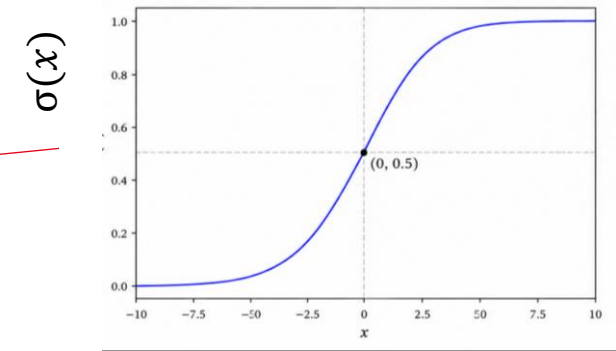
$$\sigma(x) = \tanh(x)$$

► **Rectified Linear Unit (ReLU):**

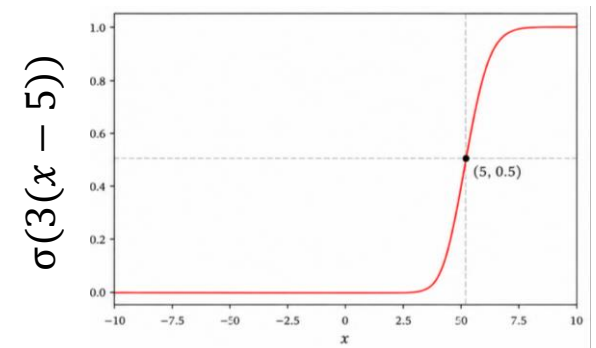
$$\sigma(x) = \begin{cases} 0 & \text{for } x \leq 0 \\ x & \text{for } x > 0 \end{cases}$$

► **Leaky ReLU:**

$$\sigma(x) = \begin{cases} 0.001x & \text{for } x \leq 0 \\ x & \text{for } x > 0 \end{cases}$$



Sigmoid Function

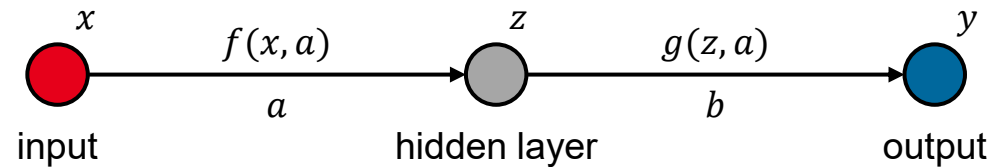


Sigmoid Function with sifted and scaled input

Task [Backpropagation]:

Calculate the partial derivatives of the cost function with respect to each $w_{jk}^{[L]}$ and $b_j^{[L]}$. For this derivation, it is sufficient to consider only one training point, i.e., we consider the cost function $C = \frac{1}{2} \|y - a^{[L]}\|_2^2$

Example [Most Simple Neural Network]:



$$y = g(z, b) = g(f(x, a), b)$$

With the cost function $C = \frac{1}{2} (y_0 - y)^2$. It thus holds

$$\frac{\partial C}{\partial b} = -(y_0 - y) \frac{\partial y}{\partial b}$$
$$\frac{\partial C}{\partial a} = -(y_0 - y) \frac{\partial y}{\partial z} \frac{\partial z}{\partial a}$$

Algorithm [Stochastic Gradient “Descent”]:

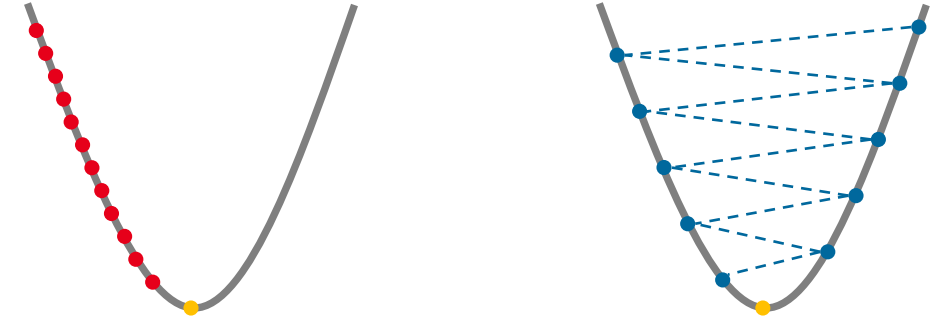
```
1. For counter=1 up to Niter
2.     Choose an integer  $k$  uniformly at random from  $\{1,2,3, \dots, N\}$ 
3.      $a^{[1]} = x^{\{k\}}$  with current training point  $x^{\{k\}}$ 
4.     for  $l = 2$  up to  $L$ 
5.          $z^{[l]} = W^{[l]}a^{[l-1]} + b^{[l]}$ 
6.          $a^{[l]} = \sigma(z^{[l]})$ 
7.          $\Sigma^{[l]} = \text{diag}(\sigma'(z^{[l]}))$ 
8.     end
9.      $\delta^{[L]} = \Sigma^{[L]}(a^{[L]} - y(x^{\{k\}}))$ 
10.    for  $l = L - 1$  downto 2
11.         $\delta^{[l]} = \Sigma^{[l]}(W^{[l+1]})^T \delta^{[l+1]}$ 
12.    end
13.    for  $l = L - 1$  downto 2
14.         $W^{[l]} \rightarrow W^{[l]} - \eta \delta^{[l]} a^{[l-1]T}$ 
15.         $\delta^{[l]} \rightarrow \delta^{[l]} - \eta \delta^{[l]}$ 
16.    end
17.end
```

► For simplicity we consider the basic version where single samples are chosen with replacement

Machine Learning

Choice of Learning Rates

- ▶ Choosing an appropriate learning rate is crucial for an effective optimization / machine learning model training.
- ▶ The convergence rate of decent methods and consequently the model's performance are strongly sensitive to the learning rate.
- ▶ While classical optimization methods typically adapt learning rates by means of line-search methods to achieve optimal convergence, the very large optimization problems involved in training neural networks are prohibitively large for corresponding approaches.
- ▶ The most popular and simple setting for the learning rate in SGD consists of taking constant values, i.e., $\eta_k = \eta \forall k$.

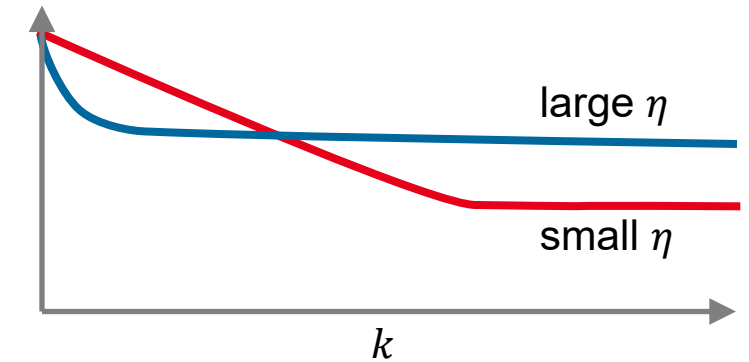


- ▶ To avoid slow convergence, another strategy is to start with a very large η_0 and reduce it progressively when the loss function features a plateau or reducing it with an inverse decay or an exponential rule

$$\eta_k = \frac{\beta}{\gamma + k} \text{ (inverse decay)}$$

$$\eta_k = \eta_0 \exp(-\gamma k) \text{ (exponential decay),}$$

where γ is a positive hyper parameter.



- ▶ Instead of choosing the same learning rate for the complete update vector, another alternative are component wise step directions. I.e., take small learning rates for components with high gradients and vice versa – though at the price of introducing additional hyper parameters.

Examples are the Root Mean Square Propagation method or the Adams method

Dirk Hartmann

Contact



Dr. Dirk Hartmann

S2|17

Schlossgartenstraße 8
64289 Darmstadt

E-mail dirk.hartmann@tu-darmstadt.de

Homepage: <https://dirk-hartmann.github.io>

TU Darmstadt: <https://www.temf.tu-darmstadt.de/digitaltwin>

LinkedIn: <https://www.linkedin.com/in/dirkhartmann>

The Art of the Possible – Blog: <https://blogs.sw.siemens.com/art-of-the-possible>