



MATHMOD

Short
Contribution
Volume

2025

EDITED BY
Andreas KÖRNER
Andreas KUGI
Wolfgang KEMMETMÜLLER
Andreas DEUTSCHMANN-OLEK
Andreas STEINBÖCK
Christian HARTL-NESIC
and Lukasz JADACHOWSKI

MATHMOD Short Contribution Volume

11th Vienna International Conference on Mathematical Modelling

February 19-21, 2025, TU Wien, Vienna, Austria.

<https://doi.org/10.34726/8999>



MATHMOD 2025

Vienna

Edited by

Andreas KÖRNER, Andreas KUGI, Wolfgang KEMMETMÜLLER, Andreas DEUTSCHMANN-OLEK, Andreas STEINBOECK, Christian HARTL-NESIC, and Lukasz JADACHOWSKI.

Preface

The 11th Vienna International Conference on Mathematical Modelling 2025 (MATHMOD2025) took place at TU Wien, 19-21 February 2025. Since 1994, the MATHMOD conference series has invited scientists, engineers, and experts to present their ideas, methods, applications, and results in mathematical modelling and share their experiences in different application domains. Like a mathematical model, MATHMOD has constants, parameters, and variables. One constant is the frequency - triennially in February. Only the last MATHMOD 2022 had to take place in July 2022 due to the coronavirus pandemic, so the 11th MATHMOD was scheduled again for February 2025. One variable in MATHMOD is the number of participants and the number of contributions. In order to maximise these two variables, MATHMOD offers three types of contributions: full contributions, short, and poster contributions. The short and poster contributions are reviewed based on a 2-page short paper, which, after acceptance, will be published in this MATHMOD Short Contribution Volume.

MATHMOD 2025 is organised by the Automation and Control Institute and the Institute of Analysis and Scientific Computing of TU Wien and co-sponsored by the International Federation of Automatic Control (IFAC). In addition to this main sponsor, we have MathWorks and the Vienna Convention Bureau as financial sponsors. Several co-sponsoring organisations supported the conference with advertising and announcements in their societies. These include the AIT Austrian Institute of Technology, OVE Austrian Electrotechnical Association, VDI/VDE Association of German Engineers, GAMM Association of Applied Mathematics and Mechanics, ARGESIM/ASIM German Simulation Society and EUROSIM Federation of European Simulation Societies.

The conference featured three plenary talks: *Advances in Interpretable Language Models* was given by Prof. Yulan He (King's College, London, United Kingdom), *Model-Based Control in Construction Robotics* by Prof. Oliver Sawodny (Institute for System Dynamics, University of Stuttgart, Germany) and *Modeling National Supply Chains with Data-Driven 1:1 Agent-Based Models – and why it is Important* given by Prof. Stefan Thurner (Medical University of Vienna and Complexity Science Hub, Vienna, Austria). Besides the scientific programme, MATHMOD 2025 also took care of the social programme, including a social plenary lecture about *Solving the Gravity-Quantum Dilemma in Experiments* given by Prof. Markus Aspelmeyer (University of Vienna and Institute for Quantum Optics and Quantum Information (IQOQI) Vienna, Vienna, Austria). *Café Simulation* opened the doors at the conference and the conference dinner was a *Viennese Heurigen Evening* at the "Zwölf Apostelkeller".

The MATHMOD Organiser would like to thank all the contributors of this MATHMOD Short Contribution Volume.

Table of Contents

A Trust Region RB-ML-ROM Approach for Parabolic PDE Constrained Optimization	1
<i>Benedikt Klein, Mario Ohlberger</i>	
Interpretable data-driven battery model based on tensor trains	3
<i>Emina Hadzialic, Alexander Ryzhov</i>	
Capturing Biocides Uptake: Model Development Under Uncontrolled Uncertainties	5
<i>E. Sangoi, F. Cattani, F. Galvanin</i>	
Generalizing the optimal interpolation points for IRKA	7
<i>Alessandro Borghi, Tobias Breiten</i>	
Model-Based Optimisation of Outbreak Detection via Wastewater Probing	9
<i>Martin Bicher, Fabian Amman, Gergely Ódor, Andreas Bergthaler, Niki Popper</i>	
Residual Data-Driven Variational Multiscale Reduced Order Models for Convection-Dominated Problems	11
<i>Birgul Koc, Samuele Rubino, Tomás Chacón, Traian Iliescu</i>	
Data based modeling and reduced order modeling for port-Hamiltonian descriptor systems	13
<i>C. A. Beattie, V. Mehrmann</i>	
Adaptive Model Hierarchies for Multi-Query Scenarios	15
<i>Hendrik Kleikamp, Mario Ohlberger</i>	
A Green Markup for the Assessment of Optimized Circulation Plans	17
<i>Matthias Röbler, Daniele Giannandrea, Niki Popper</i>	
Vector Field Construction for Football Game-Flow Evaluation	19
<i>Tenpei Morishita, Yuji Aruga, Masao Nakayama, Akifumi Kijima, Hiroyuki Shima</i>	

Learning non-intrusive ROMs from linear SDEs with additive noise	21
<i>M. A. Freitag, J. M. Nicolaus, M. Redmann</i>	
Failure Rates from Data of Field Returns	23
<i>Sven-Joachim Kimmerle, Karl Dvorsky, Hans-Dieter Liess</i>	
Data driven identification and model reduction for nonlinear dynamics	25
<i>Zlatko Drmač, Igor Mezić</i>	
Confronting knowledge-based and machine learning models in describing batch fermentation	27
<i>Núria Campo-Manzanares, Artai Rodríguez-Moimenta, Romain Minebois, Amparo Querol, Eva Balsa-Canto</i>	
Mathematical Modeling of Microbial Community Dynamics	29
<i>Ana Paredes, Eva Balsa-Canto, Julio R. Banga</i>	
Capacity Building Project in Higher Education: Leveraging Big Data and Engineering Tools to Transform Food Science Education in Indonesia	31
<i>Monika Polanska, Yoga Pratama, Setya Budi Abduh, Ahmad Ni'matullah Al-Baarri, Jan F.M. Van Impe</i>	
Sensitivity of inducible gene expression	33
<i>Satyajeet Bhonsale, Yadira Boada, Alejandro Vignoni, Jesus Pico, Jan Van Impe</i>	
Model order reduction for artificial neural networks generated from data driven state space models	35
<i>Wil Schilders</i>	

A Trust Region RB-ML-ROM Approach for Parabolic PDE Constrained Optimization^{*}

Benedikt Klein^{*} Mario Ohlberger^{*}

^{*} *Mathematics Münster, University of Münster, Münster, Germany,
(e-mail: {benedikt.klein, mario.ohlberger}@uni-muenster.de).*

1. INTRODUCTION & PROBLEM SETUP

Optimization problems constrained by parabolic PDEs are common in science and engineering, involving challenges like optimal control and inverse problems. Traditional methods require numerous iterations to solve discretized PDEs, often creating a computational bottleneck. Surrogate models, especially reduced order models (ROM) obtained from reduced basis (RB) methods, offer increased efficiency by approximating these high-fidelity solutions. This contribution explores how machine learning (ML) can enhance surrogate model construction in the framework of error aware trust region (TR) optimization methods.

Consider a bounded domain $\Omega \subset \mathbb{R}^d$ and let V be a Hilbert space, with $H_0^1(\Omega) \subset V \subset H^1(\Omega)$. Our goal is to find a parameter μ within a bounded parameter domain $\mathcal{P} \subset \mathbb{R}^P$, minimizing the least-squares objective functional, relative to a desired state $g_{\text{ref}} \in V^K$:

$$J(u(\mu); \mu) := \Delta t \sum_{k=1}^K \|u^k(\mu) - g_{\text{ref}}^k\|_V^2 + \lambda \mathcal{R}(\mu),$$

where $u(\mu) := (u^k(\mu))_{k \in \{0, \dots, K\}} \in V^{K+1}$ represents the solution trajectory to a time-discretized parametrized parabolic PDE (primal problem), i.e., $u(\mu)$ solves

$$\frac{(u^k(\mu) - u^{k-1}(\mu), v)_{L^2(\Omega)}}{\Delta t} + a(u^k(\mu), v; \mu) = b(t^k)f(v; \mu)$$

$$u^0(\mu) = 0$$

for all $v \in V$ and $k \in \{1, \dots, K\}$, where $a(\cdot, \cdot; \mu)$ and $f(\cdot; \mu)$ are parameter-dependent (bi)linear forms, and b a time-dependent forcing input. Regularization is provided by a smooth $\mathcal{R}(\mu)$ for $\lambda > 0$, cf. Qian et al. (2017).

The gradient of $\mathcal{J}(\mu) := J(u(\mu); \mu)$ will be computed via an adjoint approach, by $\nabla_{\mu} \mathcal{J}(\mu) = \nabla_{\mu} \mathcal{L}(u(\mu), p(\mu); \mu)$, with the Lagrangian $\mathcal{L}$ and the adjoint solution $p(\mu) \in V^{K+1}$ w.r.t. $u(\mu)$, cf. Qian et al. (2017) for details.

2. MODEL REDUCTION AND MACHINE LEARNING

To numerically solve the primal and adjoint problems, spatial discretization is employed, projecting the problems into a high-dimensional space $V_h \subset V$. RB methods, project these problems further into a space $V_{\text{RB}} \subset V_h$

with significant lower dimension, thus reducing the computational burdens. This approach yields RB approximations $u_{\text{RB}}(\mu)$ and $p_{\text{RB}}(\mu)$ to the high-fidelity solutions and thereby for the objective functional $\mathcal{J}_{\text{RB}}(\mu) := J(u_{\text{RB}}(\mu); \mu)$ and its gradient $\nabla_{\mu} \mathcal{J}_{\text{RB}}(\mu)$. Additionally, an a posteriori error estimator $\Delta_{\text{RB}}^J(\mu)$ measuring the error $|\mathcal{J}(\mu) - \mathcal{J}_{\text{RB}}(\mu)|$ is provided, cf. Qian et al. (2017); Keil et al. (2021).

A way to further reduce computational costs is to replace these RB-ROMs with a machine learning surrogate, cf. Fresca and Manzoni (2022); Haasdonk et al. (2023). These models allow, by design, faster evaluations, providing an approximation to $u_{\text{ML}}(\mu)$. Suitable ML models are, for example, kernel methods. Here, the Degree of Freedom vector for the primal RB trajectory $u_{\text{RB}}(\mu)$ is approximated by a linear combination of smooth kernel functions. The associated coefficients are learned using previously collected RB solutions as training data, $(\mu_1, u_{\text{RB}}(\mu_1)), \dots, (\mu_N, u_{\text{RB}}(\mu_N))$. The surrogate to the objective functional, utilizing an ML-ROM, is defined similarly to that of RB-ROMs. However, the gradient will be calculated directly by applying the chain rule.

3. TRUST REGION OPTIMIZATION

A trust region method iteratively replaces the global optimization problem by solving a sequence of sub-problems restricted to a local trust region $T^{(i)} \subset \mathcal{P}$, i.e.,

$$\mu^{(i+1)} := \underset{\mu \in \mathcal{P}}{\text{argmin}} J^{(i)}(\mu) \text{ s.t. } \mu \in T^{(i)}, \quad (1)$$

employing a local surrogate $J^{(i)}$ to $\mathcal{J}$, for $i \in \mathbb{I} := \{0, \dots, I\}$, with $I \in \mathbb{N}_0 \cup \{\infty\}$.

Problem (1) is addressed using a projected BFGS method with an Armijo-type backtracking line search, starting at $\mu^{(i)}$ and stopping if a suitable termination criterion is reached, generating a sequence $(\mu^{(i,l)})_{l \in \{0, \dots, L^{(i)}\}}$, cf. Keil et al. (2021). The guess $\mu^{(i, L^{(i)})}$ will be rejected and (1) solved again with a shrunken trust region, if a sufficient decay condition is not satisfied and accepted otherwise.

For each $i \in \mathbb{I}$, let $V_{\text{RB}}^{(i)} \subset V_h$ be fixed RB spaces, iteratively constructed by basis enrichment, starting with $V_{\text{RB}}^{(0)} := \langle \emptyset \rangle$. This enrichment is performed by computing high-fidelity solutions at $\mu^{(i)}$ and applying proper orthogonal decomposition (POD) to them. After orthogonalization, the so selected singular vectors are added to the basis of $V_{\text{RB}}^{(i)}$. This process uniquely defines for all $i \in \mathbb{I}$ the RB surrogate

^{*} Funded by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) under Germany's Excellence Strategy EXC 2044-390685587, Mathematics Münster: Dynamics–Geometry–Structure.

$J_{\text{RB}}^{(i)}(\mu)$, its gradient $\nabla_{\mu} J_{\text{RB}}^{(i)}(\mu)$, and the error estimator $\Delta_{\text{RB}}^{J, (i)}(\mu)$.

For the ML models, in contrast, it is essential to re-adapt (retrain) them to the RB solution manifold at a relatively high frequency compared to the RB space enrichment. This adaptation is necessary for ensuring that the error remains acceptably small.

Let $J_{\text{ML}}^{(i, m)}(\mu)$, $\nabla_{\mu} J_{\text{ML}}^{(i, m)}(\mu)$ for $m \in \mathbb{N}$, denote the approximate cost functional and its gradient given by trajectories in $V_{\text{RB}}^{(i)}$, yielded from the m -th ML-ROM. When the underlying ML model is (re)trained, the functionals are updated to $J_{\text{ML}}^{(i, m+1)}(\mu)$ and $\nabla_{\mu} J_{\text{ML}}^{(i, m+1)}(\mu)$.

In general, ML surrogate models lack an *efficient* a posteriori error estimator. For RB-ML-ROMs, however, it is possible to extend the error estimators of the RB-ROMs to the ML setting. Nevertheless, the evaluation of such estimators is usually very costly relative to the faster parameter inference of ML models Haasdonk et al. (2023). This could offset the reduction in computational time gained by the ML-ROM, when training and less accurate approximations of the ML model are counted in.

This raises the problem that a solely error-based trust region, such as

$$\tilde{T}^{(i)} := \left\{ \mu \in \mathcal{P} \mid \Delta_{\text{RB}}^{J, (i)}(\mu) / J_{\text{RB}}^{(i)}(\mu) \leq \epsilon^{(i)} \right\}$$

for some $\epsilon^{(i)} > 0$, as outlined in Qian et al. (2017), cannot be used due to the lack of efficient error estimations for the ML surrogate. To address this, we choose

$$T^{(i)} := \left(\tilde{T}^{(i)} \cup \bigcup_{\mu \in \partial \tilde{T}^{(i)}} \left\{ \tilde{\mu} \in \mathbb{R}^P \mid \|\tilde{\mu} - \mu\|_{\mathbb{R}^P} \leq \kappa^{(i)} \right\} \right) \cap \mathcal{P}$$

with $\kappa^{(i)} := \alpha_0^{(i)} l_{\text{check}}$, where $l_{\text{check}} \in \mathbb{N}$ and $\alpha_0^{(i)}$ is the step size from the Armijo-type line search. By verifying $\mu^{(i, l)} \in \tilde{T}^{(i)}$ every l_{check} -th step during the BFGS optimization, we ensure that all queried parameters are within $T^{(i)}$. If this condition fails, the last accepted parameter $\mu^{(i, l)}$ is returned as an approximate solution to (1). For shrinking the trust region, we set $\epsilon^{(i+1)} := \beta_1 \epsilon^{(i)}$ and $\alpha_0^{(i+1)} := \beta_2 \alpha_0^{(i)}$, where $\beta_1, \beta_2 \in (0, 1)$, if $\mu^{(i+1)}$ is rejected.

Utilizing ML-ROM poses another challenge: the line search may result in excessively small updates of the parameter, due to the generally larger error of the ML models compared to RB-ROMs. We therefore propose terminating the line search for ML surrogates if it does not succeeds within a reasonable number of steps. In this case, the line search is rerun using $J_{\text{RB}}^{(i)}(\mu)$ and $\nabla_{\mu} J_{\text{RB}}^{(i)}(\mu)$ and the so collected RB trajectories will be used to retrain the ML-ROM. If the second backtracking also fails, the last accepted step $\mu^{(i, l)}$ is returned. Initially, the first $l_{\text{warmup}} \in \mathbb{N}$ line searches use RB-ROMs only to gather training data. Details are given in Algorithm 1.

4. CONCLUSION

This contribution shows how machine learning can be integrated in a trust region method for parameter optimizations constrained by parabolic PDEs, by combining

reduced basis models and kernel-based ML surrogates. However, maintaining accuracy requires frequent retraining and error management due to inherent errors of the ML-ROM. This highlights the potential for efficiently solving high-dimensional problems, but emphasizing the need to carefully balance between efficiency and accuracy.

REFERENCES

- Fresca, S. and Manzoni, A. (2022). POD-DL-ROM: Enhancing deep learning-based reduced order models for nonlinear parametrized PDEs by proper orthogonal decomposition. *Comput. Methods Appl. Mech. Eng.*, 388, 114181.
- Haasdonk, B., Kleikamp, H., Ohlberger, M., Schindler, F., and Wenzel, T. (2023). A new certified hierarchical and adaptive RB-ML-ROM surrogate model for parametrized PDEs. *SIAM SISC*, 45(3), A1039–A1065.
- Keil, T., Mechelli, L., Ohlberger, M., Schindler, F., and Volkwein, S. (2021). A non-conforming dual approach for adaptive trust-region reduced basis approximation of PDE-constrained parameter optimization. *ESAIM: M2AN*, 55(3), 1239–1269. doi:10.1051/m2an/2021019.
- Qian, E., Grepl, M., Veroy, K., and Willcox, K. (2017). A certified trust region reduced basis approach to PDE-constrained optimization. *SIAM SISC*, 39(5), S434–S460.

Algorithm 1 Inner loop

```

1: Set  $l := 0$ ,  $m := 0$ ,  $\mu^{(i, 0)} := \mu^{(i)}$ , no_progress :=
   false, use_ML := true and choose  $l_{\text{warmup}}, k_{\text{max}} \in \mathbb{N}$ .
2: while a termination criteria is not met or  $l = 0$  do
3:   if  $l \bmod l_{\text{check}} = 0$  or no_progress then
4:     if  $\mu^{(i, l)} \notin \tilde{T}^{(i)}$  then Go to line 33.
5:   end if
6:   Set  $k := 0$  and no_progress  $\leftarrow$  false.
7:   while a line search stopping criteria is not met do
8:     Get  $\mu^{(i, l)}(k)$  by  $k$ -th line search iteration.
9:     if  $l \leq l_{\text{warmup}}$  or not use_ML then
10:      Get  $J_{\text{RB}}^{(i)}(\mu^{(i, l)}(k))$ ,  $\nabla_{\mu} J_{\text{RB}}^{(i)}(\mu^{(i, l)}(k))$ 
11:    else
12:      Get  $J_{\text{ML}}^{(i, m)}(\mu^{(i, l)}(k))$ ,  $\nabla_{\mu} J_{\text{ML}}^{(i, m)}(\mu^{(i, l)}(k))$ 
13:    end if
14:    if  $k = k_{\text{max}}$  and  $l > 0$  then
15:      no_progress  $\leftarrow$  true
16:      Go to line 20.
17:    end if
18:     $k \leftarrow k + 1$ 
19:  end while
20: if no_progress and not use_ML then
21:   Go to line 33.
22: else if no_progress then
23:   use_ML  $\leftarrow$  false
24:   Update the search direction, using RB-ROMs.
25: else
26:   use_ML  $\leftarrow$  true and no_progress  $\leftarrow$  false
27:    $\mu^{(i, l+1)} := \mu^{(i, l)}(k)$ 
28:   Update the search direction.
29:   Train ML-ROM.
30:    $l \leftarrow l + 1$  and  $m \leftarrow m + 1$ 
31: end if
32: end while
33: return  $\mu^{(i+1)} := \mu^{(i, l)}$ 

```

Interpretable data-driven battery model based on tensor trains

Emina Hadzialic*, Alexander Ryzhov*.

**Austrian Institute of Technology, Vienna, 1210
Austria (Tel: +43 664 88390026; e-mail: alexander.ryzhov@ait.ac.at).*

Abstract: The global energy transition increasingly relies on renewable energy sources and the use of batteries for electrical energy storage. Efficient battery utilization necessitates accurate state estimation algorithms and appropriate control mechanisms. This paper presents and evaluates a data-driven approach for estimating a battery's dynamic model using tensor trains, that efficiently reconstruct complex multidimensional systems with respect to time and memory, enabling the development of adaptive models capable of capturing real-time variations in system parameters. In this study, the proposed method is applied to reconstruct a dynamic battery model from operational data and is tested upon a solid-state lithium-ion battery cell. The method's explanatory capabilities are demonstrated through the extraction of key parameters such as open circuit voltage and impedance in the form of relaxation times distribution. The accuracy is further validated against the results of conventional battery characterization tests. Owing to its intrinsic scalability and low computational cost, this method holds potential for integration into artificial intelligence-driven battery management systems, enhancing battery longevity and safety while optimizing time-intensive battery characterization processes.

Keywords: mathematical models, machine learning, tensor trains, batteries.

1. INTRODUCTION

Batteries start to play a crucial role in the global energy system, enabling renewable energy integration and supporting electrification of transport. Grid-level storage demands scalability, durability, and cost-effectiveness, while electric vehicles prioritize safety, price, and energy density. Aviation requires the highest standards for safety and efficiency, which current battery technology struggles to meet (Bills et.al.).

Battery Management Systems (BMS) monitor battery performance, estimating key metrics like State of Charge (SOC) and State of Health (SOH) to improve safety and durability. These systems use mathematical models or machine learning techniques to optimize battery behaviour, and can be enhanced through real-time data and control. In aviation, rigorous certification standards are assumed to be partially addressed through solid-state batteries and advanced safety designs. Predictive BMS tools are also essential to meet the safety margins, requiring accuracy and adaptability.

This research introduces a real-time model estimation algorithm based on low-rank data decomposition based on tensor trains (TT), which efficiently estimates SOC as well as physically relevant parameters without time-consuming battery characterization tests for a preliminary tuning (Pattipati et.al.). Experiments with lithium-ion solid-state batteries are used to validate the method, demonstrating high accuracy and negligible sub-second training time. The model's ability to explain physical battery behaviour is vital for safety-critical applications, showing its potential in enhancing battery management across various sectors.

2. METHODS

TT are widely used in machine learning as allow to efficiently process sparse high-dimensional data (Oseledets). In the present work we use TT to build two models: battery dynamic mode, that includes an open circuit voltage (OCV) and impedance, depending on the relative state of charge (SoC), and SoC as a function of observables, that include in the present case voltage drop on the battery current collectors, load and the battery temperature. The first one can be used for predictive modelling and safety status assessments, while the later provides an adaptable battery state estimation tool for the BMS.

Temperature, voltage (for the state estimation) or SoC (for the dynamic model estimation) are represented using a series of Chebyshev polynomials of the order up to 8. The load, or a current, is recalculated into a number of "smoothed" functions to use a relaxation times distribution method (Heinzmann et.al.). In order to build a model, 3 rank-3 or 2 tensors are randomly generated on the initial phase, corresponding to temperature, load and voltage/SoC. Scalar product of inputs with the tensors, and further convolution of tensors with each other allows to calculate the output, that correspond either to SoC (state estimation) or to voltage (dynamical model). Each tensor is calculated using Ridge regression method (Marquardt et.al.) with quadratic regularization assuming the environment is frozen. Experiments demonstrate, that 2-5 sweeps over the tensors, each taking <1ms on a single-core 3GHz CPU, is enough to get a converged solution. Bonds dimensions of up to 3 is

enough to get the reconstruction accuracy within $\sim 1\%$ average and $\sim 5\%$ maximum error.

3. RESULTS

Experiments were conducted on 30 Ah lithium-ion solid-state cells with a voltage window from 2.75 to 4.2 V. The specific energy is 270 Wh/kg and 560 Wh/L. The batteries have a durability of 500 cycles with 80% capacity retention at 3C, and a maximum operation current of 7C. The cells use NMC cathodes, graphite anodes, and a solid-state electrolyte, though the exact formulation is undisclosed. Experiments are performed in climatic chambers with fixed temperatures.

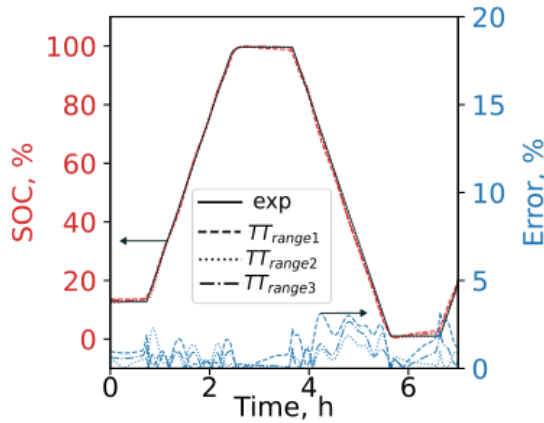


Fig. 1. Experimental SoC (left axis) for charge-discharge cycle (solid line) and reconstructed SoC using TT for 3 training data ranges, and corresponding error (right axis).

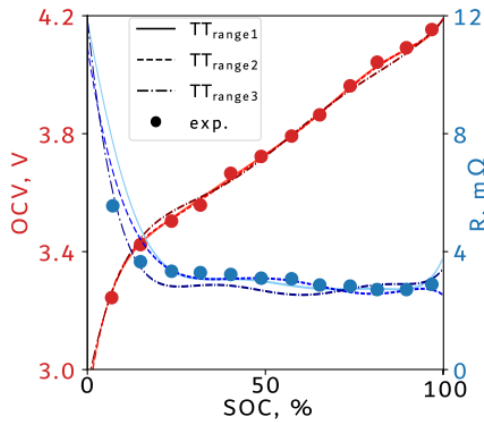


Fig. 2. OCV and resistance from characterization tests (dots) and extracted from TT for 3 training data ranges.

Standard battery characterization tests were conducted with different temperatures to obtain reference OCV and resistance values. Cycling tests, mimicking the battery operation in EV/HEA are performed to assess the method performance, with data collected every second.

The developed method is first used as a state estimation tool. SoC is typically calculated as the cumulative sum of charge passed through the battery, normalized by its capacity, and

scaled so that SoC values range from 0 to 1, corresponding to the voltage window limits. For the model training, a charge conservation is used. This method calculates the change in capacity between two time stamps by integrating the current passed through the battery over that time period. The model allows to calculate observables into the latent space variable (SoC), with the error within 5% on the whole testing dataset (one cycle is plotted at Fig.1). Therefore, potentially the method can substitute such widely used methods as KF, that require preliminary characterization tests data, that are obtained on the beginning of the battery's lifetime and change during a long-term operation thus decreasing the state estimation accuracy.

Secondly, the method is applied to reconstruct a dynamic battery model. Here, SoC is assumed to be available along with load and temperature, while the voltage is the output value. The relaxation times distribution method allows to explicitly parse the trained TT to extract OCV and resistance functions vs. temperature and SoC (Fig.2). Note, that these functions correspond to the current battery "health" status and can be treated as synthetic characterization tests.

4. CONCLUSIONS

The safe and efficient use of batteries requires monitoring and control systems. Algorithms should be fast, adaptable, and representative. TT algorithm solves battery state estimation and dynamic model reconstruction problems. The state estimation algorithm is used in BMS to optimize battery loading, thermal management, and remaining capacity calculation. Dynamic model also allows extracting battery properties important for the health status calculation and early fault prediction. The proposed method does not require characterization data obtained under controlled conditions with predefined loads. The algorithm uses operational data on voltage, temperature and load to train in near-real time. The algorithm can be used in safety-critical applications, such as aviation, and in fully adaptive BMS. It can optimize or eliminate the need for time-consuming battery characterization tests.

REFERENCES

- Bills, A., Sripad, S., Fredericks, W. L., Singh, M. and Viswanathan, V. (2020). Performance metrics required of next-generation batteries to electrify commercial aircraft. *ACS Energy Letters*, (5), 663–668.
- Heinzmann, M., Weber, A., Ivers-Tiffée, E. (2018). Advanced impedance study of polymer electrolyte membrane single cells by means of distribution of relaxation times. *Journal of Power Sources*, (402), 24–33.
- Marquardt, D. W., Snee, R. D. (1975). Ridge regression in practice. *The American Statistician*, (29), 3–20.
- Oseledets, I.V. (2011). Tensor train decomposition. *SIAM Journal on Scientific Computing*, (33), 2295–2317.
- Pattipati, B., Balasingam, B., Avvari, G., Pattipati, K., Bar-Shalom, Y. (2014). Open circuit voltage characterization of lithium-ion batteries. *Journal of Power Sources*, (269), 317–333.

Capturing Biocides Uptake: Model Development Under Uncontrolled Uncertainties

E. Sangoi* F. Cattani** F. Galvanin*

* *Department of Chemical Engineering, University College London (UCL), Torrington Place, WC1E 7JE, London, United Kingdom (e-mail: enrico.sangoi.19@ucl.ac.uk, f.galvanin@ucl.ac.uk)*

** *Process Studies Group, Global Sourcing and Production, Syngenta, Jealott's Hill International Research Centre, Berkshire, Bracknell, RG42 6EY, United Kingdom (e-mail: federica.cattani@syngenta.com)*

1. INTRODUCTION

Crop protection science plays a role in responding to the challenge of food demand with growing population and climate change. Mathematical models able to predict the interactions between the biocides and the crops can be exploited to develop new products that are also safer for the environment, aligning with sustainable agricultural practices (Umetsu and Shirai (2020)), at a lower cost and time to the market.

This project focuses on modelling the foliar uptake of pesticides (Wang and Liu (2007)). The goal is to obtain a reliable model to describe the phenomena taking place when the formulated active ingredients (AI) are sprayed on the crops leaves and to predict the percentage of AI uptake by the crop. Since plants are biological systems, there is biological variability between different species, between plants of the same species, and also between leaves of the same plant. This variability is reflected in the large variance observed in the experimental data used to calibrate the model. Another source of uncertainty is the fact that the physico-chemical phenomena affecting it are strictly correlated and observing them independently is extremely challenging. Therefore, an objective is also to understand the trade-off between complexity and explainability of the model, to guarantee that eventually the uncertainty in the model predictions and the correlations between the model parameters are acceptable.

In the literature there has been several works towards the modelling of the foliar uptake of pesticides, ranging from simple empirical correlations (Forster et al. (2004)), to compartmental models (Bridges and Farrington (1974)), up until more detailed physics-based ones (Tredenick et al. (2019)). However, the effect of the parametric uncertainty in order to guarantee the applicability of the model in predictions has not been addressed in detail in the previous works.

The novelty in this research is that the predictive mathematical model is used also to optimize the experimental campaigns, allowing a better exploitation of time and resources, and to achieve this result is fundamental to systematically take in to account the uncertainty in predictions.

2. METHODOLOGY

The modelling procedure is adapted from Franceschini and Macchietto (2008) and the framework is shown in Fig. 1.

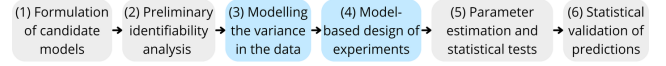


Fig. 1. Model building approach based on Franceschini and Macchietto (2008). The focus of this paper is on the steps 3 and 4, highlighted in blue.

The framework consists of 6 key steps:

- (1) Formulation of different candidate models to describe the system under study.
- (2) Preliminary analyses on the identifiability of the model parameters are conducted and any identifiability issue is addressed.
- (3) A model is fit to characterise the variability in the experimental data.
- (4) The application of Model-Based Design of Experiments (MBDoE) techniques for model discrimination and for parameter precision (Franceschini and Macchietto (2008)).
- (5) The model parameters are precisely estimated and validated statistically.
- (6) The model predictions are validated based on new experimental data and the statistics of model predictions.

Steps (1) and (2) of this procedure have been addressed by the authors in Sangoi et al. (2024a,b), where different formulations of dynamic models for foliar uptake mechanism are presented, while this paper focuses on the MBDoE (step 4 of the procedure).

2.1 Model-Based Design of Experiments

To describe the dynamics of AI leaf uptake we consider dynamic models generally formulated as in

$$\begin{aligned}\dot{\mathbf{x}}(t) &= \mathbf{f}(\mathbf{x}(t), \mathbf{u}(t), \boldsymbol{\theta}, t) \\ \hat{\mathbf{y}}(t) &= \mathbf{g}(\mathbf{x}(t), \mathbf{u}(t), \boldsymbol{\theta})\end{aligned}\quad (1)$$

where $\mathbf{x} \in \mathbb{R}^{N_x}$ is a vector of state variables, $\dot{\mathbf{x}} \in \mathbb{R}^{N_x}$ indicates the time derivatives of the states, $\hat{\mathbf{y}} \in \mathbb{R}^{N_y}$ the

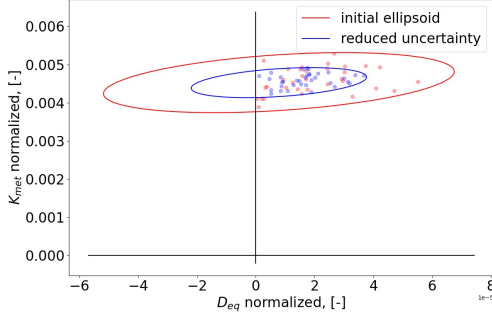


Fig. 2. Sampling of the parameter uncertainty region.

vector of predicted model outputs, $\mathbf{u} \in \mathbb{R}^{N_u}$ the vector of known system inputs, $\boldsymbol{\theta} \in \mathbb{R}^{N_\theta}$ the model parameters vector. The model considered in this paper is diffusion-based with no flux boundary conditions (Sangoi et al., 2024a), with three parameters: i) the partition between the droplet and the leaf K_{DL} , ii) the equivalent diffusion of AI through the leaf D_{eq} , iii) the metabolism rate constant (AI consumption) K_{met} . MBDoE techniques allow to exploit the mathematical formulation of the dynamic model to optimize the experimental campaigns, so that the data obtained are the most informative for the modelling task, e.g. to improve parameter precision (MBDoE-PP), and that model and experiments are coupled in a bi-directional way. To apply MBDoE-PP the information content of the experiments is generally quantified through a metric of the Fisher Information Matrix $\hat{\mathbf{H}}$ (2) or the covariance matrix of the parameter estimates $\hat{\mathbf{V}}_\theta$ (3).

$$\hat{\mathbf{H}}(\boldsymbol{\varphi}) = \nabla \hat{\mathbf{y}}(\boldsymbol{\varphi}, \hat{\boldsymbol{\theta}}) \boldsymbol{\Sigma}_y^{-1} \nabla \hat{\mathbf{y}}(\boldsymbol{\varphi}, \hat{\boldsymbol{\theta}}) \quad (2)$$

$$\hat{\mathbf{V}}_\theta(\boldsymbol{\varphi}) \simeq [\hat{\mathbf{H}}(\boldsymbol{\varphi})]^{-1} \quad (3)$$

In equation (2), the symbol $\boldsymbol{\varphi}$ stands for the experimental design vector, in this application defined by the sampling times in biokinetic experiments of uptake, and $\boldsymbol{\Sigma}_y$ represents the covariance matrix of the measurement error. Therefore in order to apply successfully MBDoE techniques it is crucial to characterise the uncertainty associated to the experiments (step 3 of the procedure in Fig. 1). MBDoE allows to optimise the experimental design vector $\boldsymbol{\varphi}$. This can be approached also by minimising the uncertainty in model predictions $\hat{\mathbf{V}}_y$ (Cenci et al. (2023)) instead of $\hat{\mathbf{V}}_\theta$, an option that can be useful in the foliar uptake case study to guarantee that the final model is reliable in its predictions.

3. RESULTS

As a preliminary study before the application of MBDoE, the results of the error propagation from the parameters to predictions are presented. Samplings are collected from the parameters uncertainty region, obtained from parameter estimation, with a Monte Carlo simulation (Fig. 2). An uncertainty reduction scenario is considered, assuming a 50% reduction in the standard deviation of parameters, and their effect is propagated to model predictions, i.e. the AI mass on the leaf surface and in the tissue (spatial

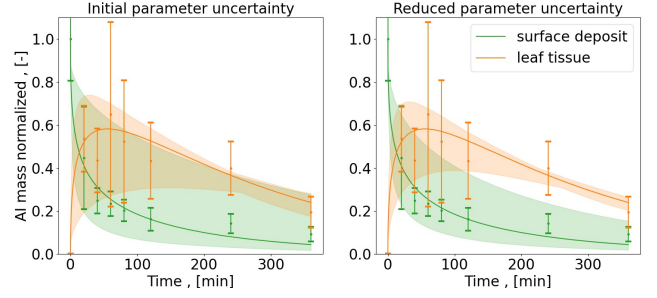


Fig. 3. Monte Carlo based uncertainty propagation from model parameters to the predictions.

integral of the discretized variable). Results on prediction uncertainty are shown in Fig. 3: the lower parameter uncertainty allows to sensibly reduce the prediction uncertainty for the AI on the leaf surface (green area), aligning with the observed experimental variability shown via the error bars. This analysis paves the way to the application of MBDoE techniques in the context of biological systems, in particular for the foliar application of biocides.

ACKNOWLEDGEMENTS

The authors gratefully acknowledge the support of the Department of Chemical Engineering - UCL and Syngenta.

REFERENCES

- Bridges, R. and Farrington, J. (1974). A compartmental computer model of foliar uptake of pesticides. *Pest. Sci.*, 5, 365—381.
- Cenci, F., Pankajakshan, A., Facco, P., and Galvanin, F. (2023). An exploratory model-based design of experiments approach to aid parameters identification and reduce model prediction uncertainty. *Comput. Chem. Eng.*, 177, 108353.
- Forster, W., Zabkiewicz, J., and Riederer, M. (2004). Mechanisms of cuticular uptake of xenobiotics into living plants: 1. influence of xenobiotic dose on the uptake of three model compounds applied in the absence and presence of surfactants into chenopodium album, hedera helix and stephanotis floribunda leaves. *Pest Manag. Sci.*, 60, 1105—1113.
- Franceschini, G. and Macchietto, S. (2008). Model-based design of experiments for parameter precision: State of the art. *Chem. Eng. Sci.*, 63, 4846—4872.
- Sangoi, E., Cattani, F., Padia, F., and Galvanin, F. (2024a). Foliar uptake models for biocides: Testing practical identifiability of diffusion-based models. *IFAC-PapersOnLine*, 58(23), 73—78.
- Sangoi, E., Cattani, F., Padia, F., and Galvanin, F. (2024b). Foliar uptake models for biocides: Testing structural and practical identifiability. *Comput. Aided Chem. Eng.*, 53, 37—42.
- Tredenick, E., Farrell, T., and Forster, W. (2019). Mathematical modelling of hydrophilic ionic fertiliser diffusion in plant cuticles: Lipophilic surfactant effects. *Plants*, 8, 202.
- Umetsu, N. and Shirai, Y. (2020). Development of novel pesticides in the 21st century. *J. Pest. Sci.*, 45, 54—74.
- Wang, C. and Liu, Z. (2007). Foliar uptake of pesticides—present status and future challenge. *Pestic. Biochem. Physiol.*, 87, 1—8.

Generalizing the optimal interpolation points for IRKA^{*}

Alessandro Borghi^{*} Tobias Breiten^{*}

^{*} *Mathematics Department, Technical University of Berlin, Straße des
17. Juni 136, 10623 Berlin, Germany (e-mail: borghi@tu-berlin.de,
tobias.breiten@tu-berlin.de).*

1. INTRODUCTION

The iterative rational Krylov algorithm (IRKA) was introduced in Gugercin et al. (2008) and has become a widely adopted method in the model reduction community for computing locally optimal solutions for the $\mathcal{H}_2$ model reduction problem. IRKA relies on the first-order optimality conditions for the solution of the $\mathcal{H}_2$ optimization problem derived in Meier and Luenberger (1967). One of the main assumptions needed for IRKA to work properly is for the error transfer function between the full and reduced order models to be in $\mathcal{H}_2$. Recently, in Borghi and Breiten (2024), this assumption was relaxed to account for non asymptotically stable systems, giving the possibility of designing an extended version of IRKA by developing a new framework and deriving its optimal interpolation conditions. Throughout this paper we refer to this algorithm as extended IRKA.

The contribution of this work is twofold: (1) We build upon the findings in Borghi and Breiten (2024) and show that the optimal interpolation points are related to the Schwarz function (see Davis (1974)); (2) We show numerically that we can use the adaptive Antoulas-Anderson (AAA) algorithm developed in Nakatsukasa et al. (2018) to approximately compute the interpolation points in extended IRKA for user-defined domains in the complex plane.

2. PRELIMINARIES

We consider the large-scale minimal single-input single-output (SISO) linear time invariant (LTI) dynamical system

$$\begin{cases} \dot{\mathbf{x}}(t) = \mathbf{A}\mathbf{x}(t) + \mathbf{b}u(t), \\ y(t) = \mathbf{c}\mathbf{x}(t), \quad \mathbf{x}(0) = 0, \end{cases} \quad G(s) = \mathbf{c}(s\mathbf{I} - \mathbf{A})^{-1}\mathbf{b}, \quad (1)$$

with $\mathbf{A} \in \mathbb{C}^{n \times n}$, $\mathbf{c}^\top \in \mathbb{C}^n$, and $\mathbf{b} \in \mathbb{C}^n$. In addition, for a fixed time t , $\mathbf{x}(t) \in \mathbb{C}^n$, $u(t) \in \mathbb{C}$, and $y(t) \in \mathbb{C}$, denote the state, input, and output of the system respectively. Here, G denotes the transfer function of the system in the frequency domain. We refer to (1) as the full-order model (FOM). We assume the eigenvalues $\{\lambda_j\}_{j=1}^n$ of $\mathbf{A}$ to be in a simply-connected open set $\mathbb{A}$ in the complex plane. Most importantly, $\mathbb{A}$ does not have to be in the left half-plane $\mathbb{C}_-$, which can result in (1) not being asymptotically stable. Our work deals with the computation of a reduced-order model (ROM) of a form analogous to (1) but with

system matrices $\mathbf{A}_r \in \mathbb{C}^{r \times r}$, $\mathbf{c}_r^\top \in \mathbb{C}^r$, $\mathbf{b}_r \in \mathbb{C}^r$, where $r \ll n$, and transfer function $G_r(s) = \mathbf{c}_r(s\mathbf{I} - \mathbf{A}_r)^{-1}\mathbf{b}_r$ with poles $\{\hat{\lambda}_j\}_{j=1}^r \in \mathbb{A}$. Similar to the $\mathcal{H}_2$ -optimal model reduction framework, we seek G_r such that

$$\min_{\deg(G_r)=r} \|G - G_r\|_*, \quad (2)$$

is solved, where $\| \cdot \|_*$ indicates a proper norm. In Borghi and Breiten (2024) an $\mathcal{H}_2(\bar{\mathbb{A}}^c)$ space for systems with poles in $\mathbb{A}$ and analytic in its exterior $\bar{\mathbb{A}}^c$ was introduced to include error transfer functions that are not in $\mathcal{H}_2$. For the development of the $\mathcal{H}_2(\bar{\mathbb{A}}^c)$ -optimal model reduction framework, the concept of conformal map (see Theorem 6.1.2 in Wegert (2012)) is pivotal. Under proper assumptions on the conformal maps and $G_r \in \mathcal{H}_2(\bar{\mathbb{A}}^c)$ being a local minimizer of (2), simplified $\mathcal{H}_2(\bar{\mathbb{A}}^c)$ optimality conditions were derived in (Borghi and Breiten, 2024, Corollary 3) resulting in

$$G_r(\varphi(\hat{\lambda}_j)) = G(\varphi(\hat{\lambda}_j)) \text{ and } G'_r(\varphi(\hat{\lambda}_j)) = G'(\varphi(\hat{\lambda}_j)), \quad (3)$$

for $j = 1, \dots, r$, with

$$\varphi(s) = \overline{\psi(-\psi^{-1}(s))} = \psi(-\overline{\psi^{-1}(s)}), \quad (4)$$

where $\overline{\psi}(s) = \overline{\psi(\bar{s})}$. The conditions in (3) then led to the development of extended IRKA. In the next sections, we leverage the connection between (4) and Schwarz functions to approximate φ given only points on the boundary of user defined domains. However, the approximation will not necessarily satisfy the assumptions made in Borghi and Breiten (2024).

3. OUR METHOD

Before introducing the result of this work, we give a brief summary of the concepts of Schwarz reflection and Schwarz function based on Davis (1974). For a more detailed description of this topic see Davis (1974) and Shapiro (1992). We consider the analytic arc Γ given by the parametrization $z = f(\theta)$ with $\theta \in [a, b]$, $a, b \in \mathbb{R}$. Given any point $\theta_0 \in [a, b]$, f is a bijective conformal map in the disk $D_{\theta_0} = \{\tau \in \mathbb{C} \mid |\tau - \theta_0| < \rho(\theta_0)\}$, with $\rho: [a, b] \rightarrow (0, \infty)$. Let $\tilde{z} = f(\bar{\tau})$ for $\tau \in D_{\theta_0}$. Then $\tilde{z}$ is called the *Schwarz reflection* of z with respect to the analytic arc Γ . The *Schwarz function* is defined as $S(z) = \tilde{z}$ and analytic in a neighborhood of Γ . For $z \in \Gamma$ we get $S(z) = \bar{z}$. The complex conjugate of the Schwarz function applied to a point is the Schwarz reflection of the point with respect to Γ (see Chapter 6 in Davis (1974)). For the ease of notation we use the term *anti-conformal reflection* $R(\cdot) = \bar{S}(\cdot)$ from Shapiro (1992) to indicate the

^{*} Funded by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) - 384950143 as part of GRK2433 DAEDALUS.

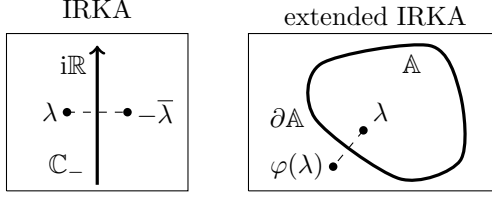


Fig. 1. Depiction of the interpolation points for IRKA (left) and extended IRKA (right).

Schwarz reflection. It is possible to use f to get an explicit formulation of the anti-conformal reflection for any point z in the neighborhood of Γ . The steps are the following

$$z \xrightarrow{f^{-1}} \tau \xrightarrow{\overline{(\cdot)}} \bar{\tau} \xrightarrow{f} \tilde{z} = R(z), \quad (5)$$

which lead to the definition

$$R(\cdot) = f(\overline{f^{-1}(\cdot)}), \quad (6)$$

and, consequently, $S(\cdot) = \bar{f}(f^{-1}(\cdot))$. Let us now set

$$f(\tau) = \psi(i\tau) \text{ and so } f^{-1}(z) = -i\psi^{-1}(z), \quad (7)$$

with ψ a conformal map satisfying the assumptions below.

Assumption 1. Let $\mathbb{X} \subseteq \mathbb{C}$ include $i\mathbb{R}$, and $\mathbb{Y} \subseteq \mathbb{C}$ include the analytic closed curve $\partial\mathbb{A}$. Furthermore, Let $\mathbb{X}_- = \mathbb{X} \setminus \mathbb{C}_+$, $\mathbb{X}_+ = \mathbb{X} \setminus \mathbb{C}_-$ such that $\{s \in \mathbb{C} \mid -\bar{s} \in \mathbb{X}_-\} \subseteq \mathbb{X}_+$, $\partial\mathbb{A}_- = \mathbb{Y} \setminus \mathbb{A}^c$, and $\partial\mathbb{A}_+ = \mathbb{Y} \setminus \mathbb{A}$. We assume $\psi: \mathbb{X} \rightarrow \mathbb{Y}$ to be a bijective conformal map such that (i) $\psi \circ i: \mathbb{R} \rightarrow \partial\mathbb{A}$, (ii) $\psi: \mathbb{X}_+ \rightarrow \partial\mathbb{A}_+$ and (iii) $\psi: \mathbb{X}_- \rightarrow \partial\mathbb{A}_-$.

With (7) we can connect the definition of φ in (4) and the anti-conformal reflection in (6). As a matter of fact, by substituting f in (6) with (7) we get

$$R(\cdot) = \psi(i(-i\psi^{-1}(\cdot))) = \psi(-\overline{\psi^{-1}(\cdot)}) = \varphi(\cdot).$$

The composition of φ is similar to (5) but instead of applying a complex conjugation we take the mirror image with respect to the imaginary axis. In more detail, we have

$$z \xrightarrow{\psi^{-1}} \tau \xrightarrow{-\overline{(\cdot)}} -\bar{\tau} \xrightarrow{\psi} \varphi(z).$$

While $\partial\mathbb{A}$ is defined by the user, the function φ is unknown a-priori. For this reason, now that the connection between the anti-conformal reflection and (4) has been established, we use the AAA algorithm to approximate φ given $\partial\mathbb{A}$. It is important to emphasize that the function S is solely determined by the chosen $\partial\mathbb{A}$. This gives us the possibility to approximate φ with AAA using only samples of $\partial\mathbb{A}$. To do so we take points $z \in \partial\mathbb{A}$, approximate $S(z) = \tilde{z}$ with AAA (see also Trefethen (2024)), and complex conjugate the resulting function. The approximated φ is then used for computing the interpolation points employed by extended IRKA. The main drawback of this approach is that, as S is defined in a neighborhood of $\partial\mathbb{A}$, the same applies for the approximation to φ .

4. NUMERICAL EXAMPLE

We test the extended IRKA with interpolation points computed through AAA on the controlled linear undamped wave equation from Borghi and Breiten (2024). After discretization by centered finite differences we get a FOM with $n = 400$ and poles on the imaginary axis. We are interested in the poles near the origin as they provide an approximately good description of the original system. For this reason, we use the ‘boomerang’ shape illustrated in

Fig. 2 on the left for $\partial\mathbb{A}$. We parametrized $\partial\mathbb{A}$ such that it has only one segment near the FOM poles and it is close to the origin (see Fig. 2 on the right). We do so in order for extended IRKA to identify these poles as dominant and place the reduced poles accordingly. In Fig. 2 we show the resulting ROM poles $\{\hat{\lambda}_j\}_{j=1}^r$ and interpolation points $\{\varphi(\hat{\lambda}_j)\}_{j=1}^r$ for $r = 18$. In addition, Fig. 3 shows that the impulse response of the resulting ROM well approximates the one of the FOM.

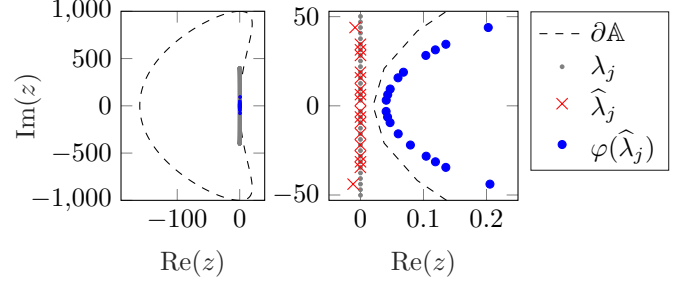


Fig. 2. Chosen $\partial\mathbb{A}$, poles of the FOM and ROM, and the computed interpolation points. The plot on the right is a magnification of the one the left near the origin.

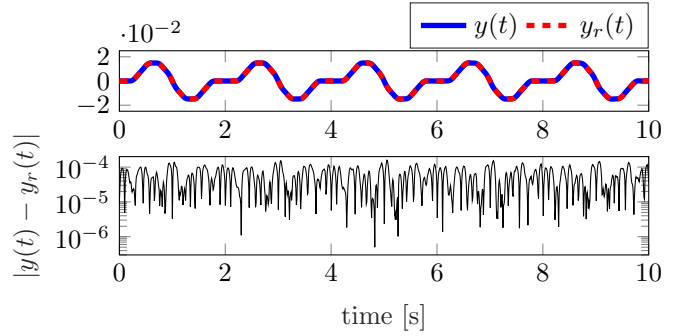


Fig. 3. (Top) output trajectories of the FOM (y) and ROM (y_r) with respect to time t . (Bottom) absolute error.

REFERENCES

- Borghi, A. and Breiten, T. (2024). $\mathcal{H}_2$ optimal rational approximation on general domains. *Advances in Computational Mathematics*, 50, 28.
- Davis, P.J. (1974). *The Schwarz Function and its Applications*. Carus Mathematical Monographs. Mathematical Association of America.
- Gugercin, S., Antoulas, A.C., and Beattie, C. (2008). $\mathcal{H}_2$ model reduction for large-scale linear dynamical systems. *SIAM Journal on Matrix Analysis and Applications*, 30(2), 609–638.
- Meier, L. and Luenberger, D. (1967). Approximation of linear constant systems. *IEEE Transactions on Automatic Control*, 12, 585–588.
- Nakatsukasa, Y., Sète, O., and Trefethen, L.N. (2018). The AAA algorithm for rational approximation. *SIAM Journal on Scientific Computing*, 40(3), A1494–A1522.
- Shapiro, H. (1992). *The Schwarz Function and Its Generalization to Higher Dimensions*. The University of Arkansas Lecture Notes in the Mathematical Sciences. Wiley.
- Trefethen, L.N. (2024). Polynomial and rational convergence rates for Laplace problems on planar domains. *Proceedings of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 480(2295), 20240178.
- Wegert, E. (2012). *Visual Complex Functions*. Springer, Basel, Switzerland.

Model-Based Optimisation of Outbreak Detection via Wastewater Probing

Martin Bicher* Fabian Amman** Gergely Ódor**
Andreas Bergthaler** Niki Popper***

* *Institute of Information Systems Engineering, TU Wien,
Favoritienstraße 9-11, 1040 Vienna, Austria (e-mail:
martin.bicher@tuwien.ac.at)*

** *Center for Pathophysiology, Infectiology and Immunology, Medical
University of Vienna, Kinderspitalgasse 15, 1090 Vienna, Austria*

*** *dwh GmbH, Neustiftgasse 57-59, 1070 Vienna, Austria*

1. INTRODUCTION

At least since the COVID-19, crisis wastewater-based epidemiology (WBE) has been recognised internationally as a reliable surveillance and early warning system to track circulating pathogens (Singer et al. (2023)). By probing and sequencing wastewater, WBE allows to detect, quantify and characterise pathogens circulating in the connected catchment. Several health authorities have set up WBE programs in their surveillance efforts. For example in Austria, wastewater from 48 purification plants is probed on a weekly basis for SARS-CoV-2 and its variants (see Amman et al. (2022)), and in parts of Austria also for influenza and the respiratory syncytial virus. Expansion to other relevant pathogens is considered.

The main advantage of WBE over traditional case-based surveillance is its scalability, allowing one sample to cover thousands without requiring active participation, reducing costs and bias. While cost-effective, nationwide surveillance at reasonable granularity incurs significant taxpayer costs. Deciding which wastewater plants to sample and how often is crucial for public health but must be economically justified to convince policymakers and the public of WBE's value.

In this work, we present a model that simulates the regional spread of a new pathogen, its concentration at wastewater plants, and the limit of detection of pathogen specific assays as a function of wastewater catchment characteristics. The goal is to minimise detection time by optimising the selection of plants and sampling intervals. After introducing the model and showing preliminary results with manually varied strategies, we propose ideas for using a simheuristic to solve this optimisation problem.

2. METHODS

2.1 Simulation Model

We follow the network-based SIRS approach in Hethcote (1978) and specify a model with M nodes, representing households, and a corresponding vector $(N_i)_{i=1}^M$ of inhabitants. The nodes are connected by a weighted digraph

identified by adjacency matrix $A \in (\mathbb{R}^+)^{M \times M}$, whereas each entry $A_{i,j}$ corresponds to the average daily number of contacts between individuals in nodes i and j . With $S_i(0) + I_i(0) + R_i(0) = N_i$ the dynamics is defined by

$$\begin{aligned}\dot{S}_i(t) &= -S_i(t)\Theta(I, i) + \delta R_i(t), \\ \dot{I}_i(t) &= S_i(t)\Theta(I, i) - \gamma I_i(t), \\ \dot{R}_i(t) &= \gamma I_i(t) - \delta R_i(t),\end{aligned}\quad (1)$$

$$\Theta(I, i) = A_{i,i}\beta^{in}\frac{I_i(t)}{N_i} + \sum_{j \neq i} A_{i,j}\beta^{out}\frac{I_j(t)}{N_j}. \quad (2)$$

Hereby, β^{in} and β^{out} refer to the in-node and out-node transmission rate, γ to the recovery rate and δ to the immunity waning rate. Time-unit will be days.

Compartments $S_i(t), I_i(t), R_i(t)$ represent the expected number of susceptible, infectious and recovered persons in household i at time t . We argue, that I_i is proportional to the overall quantity of pathogen present in household i and to the pathogen load excreted into the wastewater.

Furthermore we introduce K purification plant nodes and include them into the model using adjacency matrix $B \in \{0, 1\}^{K \times M}$. Hereby $B_{k,i} = 1 \Leftrightarrow$ household j lies in the catchment area of plant k . We define

$$W_k(t) = \frac{\sum_{i=1}^M B_{k,i} I_i(t)}{\sum_{i=1}^M B_{k,i} N_i} \quad (3)$$

as the pathogen *ground truth* at plant k . It models the ratio of all pathogen excreted by households in the catchment area compared to a human control signal.

We finally define the *measured signal* at plant k via

$$\hat{W}_k(t) := \nu_k \begin{cases} W_k\left(\left\lfloor \frac{t}{\xi_k} \right\rfloor \xi_k\right), & \text{if } W_k\left(\left\lfloor \frac{t}{\xi_k} \right\rfloor \xi_k\right) > \mu, \\ 0, & \text{else.} \end{cases} \quad (4)$$

Hereby, μ defines a concentration threshold below which a given pathogen signal cannot reliably be detected by a probe. The parameter vectors ν and ξ will be regarded as *probing strategy*: $\nu_k \in \{0, 1\}$ defines if plant k is probed at all, $\xi_k \in \mathbb{N}$ defines the interval between taking two probes from the plant.

2.2 Optimisation Problem

For optimisation we focus on the *detection time*, i.e. the first time a signal is detected at any of the probed plants:

* This research was funded in part by the Austrian Science Fund (FWF) I 5908-G.

$$\tau := \min_{k \in \{1, \dots, K\}} \tau_k^i := \min_{k \in \{1, \dots, K\}} \min_{t > 0} (\hat{W}_k(t) > 0). \quad (5)$$

In addition to the probing strategy, the value of τ will also depend on the location of the outbreak. To simulate the latter we specify the initial conditions $R_i(0) := 0$,

$$I_i(0) := \delta_{i,i_0} I_\epsilon, \text{ and } S_i(0) := N_i(0) - I_i(0), \quad (6)$$

where node i_0 will be regarded as *index-household*. With this definition, (ν, ξ) and i_0 uniquely map to a detection time: $\tau = \tau^{i_0, \nu, \xi}$. Considering that the probing strategy should be able to quickly detect the pathogen independent of the choice of i_0 , we specify the optimisation target

$$F : \{0, 1\}^K \times \mathbb{N}^K \rightarrow \mathbb{R}^+ : F(\nu, \xi) \mapsto \frac{1}{M} \sum_{i_0=1}^M \tau^{i_0, \nu, \xi}. \quad (7)$$

Clearly, minimising F alone would be meaningless, since it does not incorporate any cost or sensitivity constraints. We define

$$C_1(\nu, \xi) = \sum_{i=1}^K \nu_i, \quad C_2(\nu, \xi) = 7.0 \cdot \sum_{i=1}^K \frac{\nu_i}{\xi}. \quad (8)$$

That means, $C_1(\nu, \xi) \leq c_1$ limits the total number of included plants to c_1 , to restrict the logistic efforts, and $C_2(\nu, \xi) \leq c_2$ limits the total number of probes taken per week to c_2 , to restrict the total costs.

In summary, we define the optimisation problem as follows:

$$(\nu, \xi)_{\text{opt}} = \operatorname{argmin}_{(\nu, \xi) \in \Omega_{c_1, c_2}} F(\nu, \xi), \text{ with } \Omega_{c_1, c_2} := \{(\nu, \xi) : C_1(\nu, \xi) \leq c_1, C_2(\nu, \xi) \leq c_2\}. \quad (9)$$

3. RESULTS

For the preliminary results shown in this paper, we extracted a synthetic contact network from an existing agent-based epidemics model. It was developed and applied in the course of the COVID-19 crisis, is specified in the SI of Bicher et al. (2021), and was used for export of synthetic data before (see Popper et al. (2021)). We counted, averaged and exported randomly sampled contacts between the roughly 4.5M model households over ten simulated days, leading to an integer vector $N \in \mathbb{N}^{4.5 \cdot 10^6}$ and a sparse matrix $A \in \mathbb{N}^{4.5 \cdot 10^6 \times 4.5 \cdot 10^6}$. Data about the catchment areas of the $K = 636$ largest purification plants in Austria was used to specify B . Finally, disease parameters $\beta^{\text{in}}, \beta^{\text{out}}, \gamma$ and δ were manually chosen to reflect R_0 and immunity waning behaviour of SARS-CoV-2. To make target function F computable, we ran the sum defined in (7) only over 5000 randomly drawn index households instead of all M possible ones.

In Figure 1 detection times $\tau^{i, \nu, \xi}$ are compared for two probing strategies. The *reference* strategy uses the 48 Austrian plants probed once per week, analogous to the currently implemented system in Austria. The *comparator* strategy uses 27 manually selected plants probed in intervals between 2 and 12 days. Both lie in $\Omega_{48, 48}$ and are therefore comparable with respect to costs.

4. DISCUSSION AND CONCLUSION

Preliminary results show that the currently implemented strategy for early pathogen detection can be improved. Manual tests reduced detection time by nearly three days,

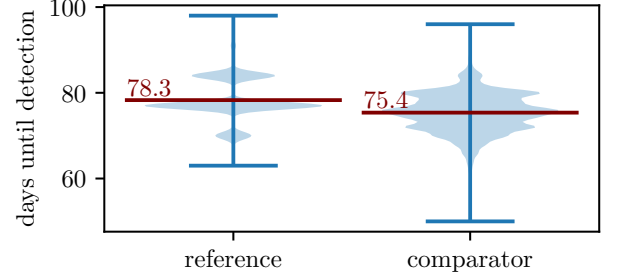


Fig. 1. Detection times for two probing strategies in $\Omega_{48, 48}$. The red line shows the target F computed as the mean of the detection-times for 5000 random index households.

offering policymakers a crucial time-advantage. Using well-defined simheuristics instead of manual methods would yield further improvements.

The optimisation challenge lies in the time-consuming simulations and vast search space. Despite efficient matrix multiplications and parallelisation using Numpy and Scipy, computing I, W , and $\hat{W}$ for one index household and probing strategy still takes about one minute on a well-equipped server.

Given the complexity of the search space Ω_{c_1, c_2} , traditional population-based metaheuristics like Genetic Algorithms (GAs) would require large populations to converge effectively. A more integrated approach is needed to reduce the number of simulations required. The plan is to exploit additional feedback from the simulation in addition to the detection time, such as particularly successful or unsuccessful plants, to guide more targeted crossovers and mutations in a GA setup.

REFERENCES

- Amman, F., Markt, R., Endler, L., Hupfaut, S., Agerer, B., Schedl, A., Richter, L., Zechmeister, M., Bicher, M., Heiler, G., et al. (2022). Viral variant-resolved wastewater surveillance of sars-cov-2 at national scale. *Nature biotechnology*, 40(12), 1814–1822.
- Bicher, M., Rippinger, C., Urach, C., Brunmeir, D., Siebert, U., and Popper, N. (2021). Evaluation of Contact-Tracing Policies against the Spread of SARS-CoV-2 in Austria: An Agent-Based Simulation. *Medical Decision Making*, 41(8), 1017–1032. doi: 10.1177/0272989X211013306.
- Hethcote, H.W. (1978). An immunization model for a heterogeneous population. *Theoretical population biology*, 14(3), 338–349.
- Popper, N., Zechmeister, M., Brunmeir, D., Rippinger, C., Weibrecht, N., Urach, C., Bicher, M., Schneckenreither, G., and Rauber, A. (2021). Synthetic reproduction and augmentation of covid-19 case reporting data by agent-based simulation. *Data Science Journal*, 20, 16. doi: 10.5334/dsj-2021-016.
- Singer, A.C., Thompson, J.R., Filho, C.R.M., Street, R., Li, X., Castiglioni, S., and Thomas, K.V. (2023). A world of wastewater-based epidemiology. *Nature Water*, 1(5), 408–415.

Residual Data-Driven Variational Multiscale Reduced Order Models for Convection-Dominated Problems ^{*}

Birgul Koc ^{*} Samuele Rubino ^{**} Tomás Chacón ^{***}
Traian Iliescu ^{****}

^{*} University of Seville, EDAN, Spain (e-mail: bkoc@us.es).
^{**} University of Seville, EDAN, Spain (e-mail: samuele@us.es)
^{***} University of Seville, EDAN, Spain (e-mail: chacon@us.es)
^{****} Virginia Tech, USA (e-mail: iliescu@vt.edu)

1. INTRODUCTION

As a mathematical model, we use Navier-Stokes equations (NSE) (1)-(2):

$$\frac{\partial \mathbf{u}}{\partial t} - Re^{-1} \Delta \mathbf{u} + \mathbf{u} \cdot \nabla \mathbf{u} + \nabla p = \mathbf{0}, \quad (1)$$

$$\nabla \cdot \mathbf{u} = 0, \quad (2)$$

where $\mathbf{u}$ is the velocity, p the pressure, t the continuous time instant, and Re the Reynolds number. Furthermore, we use homogeneous Dirichlet boundary conditions.

We use *proper orthogonal decomposition* (POD) to obtain the *reduced order model* (ROM) basis and operators for all ROMs. Thanks to the orthogonality of the ROM basis functions, we can decompose the ROM space into *large* and *small* spaces as follows: $\mathbf{X}^d = \mathbf{X}^L \oplus \mathbf{X}^S$, where $\mathbf{X}^d := \text{span}\{\boldsymbol{\varphi}_1, \dots, \boldsymbol{\varphi}_d\}$, $\mathbf{X}^L := \text{span}\{\boldsymbol{\varphi}_1, \dots, \boldsymbol{\varphi}_L\}$, and $\mathbf{X}^S := \text{span}\{\boldsymbol{\varphi}_{L+1}, \dots, \boldsymbol{\varphi}_d\}$.

When all the ROM modes are used, the ROM approximation $\mathbf{u}_d$, i.e.,

$$\mathbf{u}_d = \sum_{j=1}^d (\mathbf{a}_d)_j \boldsymbol{\varphi}_j \quad (3)$$

is the most accurate ROM approximation of the *full order model* (FOM) solution with the given data in the POD sense.

For laminar flows, a low-dimensional ROM solution $\mathbf{u}_L$, with small $L \ll d$, yields an accurate approximation of the FOM solution. In the resolved regime, the most straightforward model of ROMs, *Galerkin ROM* (G-ROM) can be used to obtain the ROM solution $\mathbf{u}_L$:

$$\dot{\mathbf{a}}_L = \mathbf{A}_{LL} \mathbf{a}_L + \mathbf{a}_L^\top \mathbf{B}_{LLL} \mathbf{a}_L, \quad (4)$$

where $(\mathbf{A}_{LL})_{ij} := -Re^{-1}(\nabla \boldsymbol{\varphi}_i, \nabla \boldsymbol{\varphi}_j)$ and $(\mathbf{B}_{LLL})_{ijk} := -(\boldsymbol{\varphi}_i, \boldsymbol{\varphi}_j \cdot \nabla \boldsymbol{\varphi}_k)$, respectively, $\forall i, j, k = 1, \dots, L$. The derivation of the G-ROM (4) is built by replacing $\mathbf{u}$ in (1)-(2) with $\mathbf{u}_L$ and projecting the resulting system onto the ROM space $\mathbf{X}^L$.

However, for turbulent flows, the low-dimensional ROM solution $\mathbf{a}_L$ of (4) is not an accurate approximation of the FOM solution. To increase the numerical accuracy of $\mathbf{a}_L$

without significantly increasing the computational cost, one needs to add a low-dimensional ROM *closure term* to the G-ROM (4).

2. ROM CLOSURE MODELS

The ROM closure modeling aims to model the closure term which is derived from a *variational multiscale* (VMS) setting (see Mou et al. (2021) and Ballarin et al. (2020)). To construct the ROM closure term, first, we need to define the large and sub-scale solutions of the most accurate ROM solution, $\mathbf{u}_d$, as follows:

$$\mathbf{u}_L := \sum_{j=1}^L (\mathbf{a}_L)_j \boldsymbol{\varphi}_j, \quad \mathbf{u}_S := \sum_{j=L+1}^d (\mathbf{a}_S)_j \boldsymbol{\varphi}_j. \quad (5)$$

Then, we obtain the large and sub-scale equations: (i) replace the $\mathbf{u}$ in (1)-(2) with $\mathbf{u}_d = \mathbf{u}_L + \mathbf{u}_S$ and project the resulting system onto the ROM spaces $\mathbf{X}^L$ and $\mathbf{X}^S$, respectively. Then, the large and sub-scale equations are:

$$\begin{aligned} \dot{\mathbf{a}}_L = & \mathbf{A}_{LL} \mathbf{a}_L + \mathbf{A}_{LS} \mathbf{a}_S + \mathbf{a}_L^\top \mathbf{B}_{LLL} \mathbf{a}_L \\ & + \mathbf{a}_L^\top \mathbf{B}_{LLS} \mathbf{a}_S + \mathbf{a}_S^\top \mathbf{B}_{LSL} \mathbf{a}_L + \mathbf{a}_S^\top \mathbf{B}_{LSS} \mathbf{a}_S, \end{aligned} \quad (6a)$$

$$\begin{aligned} \dot{\mathbf{a}}_S = & \mathbf{A}_{SS} \mathbf{a}_S + \mathbf{A}_{SL} \mathbf{a}_L + \mathbf{a}_S^\top \mathbf{B}_{SSS} \mathbf{a}_S \\ & + \mathbf{a}_S^\top \mathbf{B}_{SSL} \mathbf{a}_L + \mathbf{a}_L^\top \mathbf{B}_{SLS} \mathbf{a}_S + \mathbf{a}_L^\top \mathbf{B}_{SLL} \mathbf{a}_L. \end{aligned} \quad (6b)$$

In this work, we use two different ROM closure constructions, which yield two different ROM model: the *coefficient-based data-driven variational multiscale* ROM (C-D2-VMS-ROM) and the *residual-based data-driven variational multiscale* ROM (R-D2-VMS-ROM).

The C-D2-VMS-ROM (Mou et al. (2021)) is derived from the large-scale equation (6a) by defining the closure term as "**closure term** = $\mathbf{A}_{LS} \mathbf{a}_S + \mathbf{a}_L^\top \mathbf{B}_{LLS} \mathbf{a}_S + \mathbf{a}_S^\top \mathbf{B}_{LSL} \mathbf{a}_L + \mathbf{a}_S^\top \mathbf{B}_{LSS} \mathbf{a}_S$ ". Since the **closure term** is not in a closed form, to close it, we use a quadratic coefficient-based ansatz (Mou et al. (2021)), which depends on the large-scale solution $\mathbf{a}_L$: "**ansatz** = $\tilde{\mathbf{A}}_{LL} \mathbf{a}_L + \mathbf{a}_L^\top \tilde{\mathbf{B}}_{LLL} \mathbf{a}_L$ ".

In the new R-D2-VMS-ROM, we define the closure term and residual-based ansatz from the sub-scale equation (6b) as "**closure term** = $\mathbf{a}_S$ " and "**ansatz** = $\tilde{\mathbf{A}}_{SS} \text{ResS}(\mathbf{a}_L) + \text{ResS}(\mathbf{a}_L)^\top \tilde{\mathbf{B}}_{SSS} \text{ResS}(\mathbf{a}_L)$ ", where the residual is $\text{ResS}(\mathbf{a}_L) := \mathbf{A}_{SL} \mathbf{a}_L + \mathbf{a}_L^\top \mathbf{B}_{SLL} \mathbf{a}_L$.

To find the unknown operators $\tilde{\mathbf{A}}, \tilde{\mathbf{B}}$, we use a *data-driven* (D2) approach (Rebollo and Coronil (2024)). We obtain the D2 operators by solving the following minimization problem:

$$\min_{\tilde{\mathbf{A}}, \tilde{\mathbf{B}}} \sum_{k=1}^M \|\text{closure term}(\mathbf{a}_L^k, \mathbf{a}_S^k) - \text{ansatz}(\mathbf{a}_L^k)\|_{\mathcal{L}^2}^2, \quad (7)$$

where M represents the number of snapshots. By using the closure terms and ansatzes for the C-D2-VMS-ROM and R-D2-VMS-ROM, we solve (7) to obtain the corresponding D2 operators, i.e. $\tilde{\mathbf{A}}$ and $\tilde{\mathbf{B}}$. Then, by plugging the resulting ansatzes into (6a), C-D2-VMS-ROM and R-D2-VMS-ROM read as follows:

$$\dot{\mathbf{a}}_L = (\mathbf{A}_{LL} + \tilde{\mathbf{A}}_{LL})\mathbf{a}_L + \mathbf{a}_L^\top (\mathbf{B}_{LLL} + \tilde{\mathbf{B}}_{LLL})\mathbf{a}_L, \quad (8a)$$

$$\begin{aligned} \dot{\mathbf{a}}_L &= \mathbf{A}_{LL}\mathbf{a}_L + \mathbf{a}_L^\top \mathbf{B}_{LLL}\mathbf{a}_L + \mathbf{A}_{LS}\mathbf{a}_S^* + \mathbf{a}_L^\top \mathbf{B}_{LLS}\mathbf{a}_S^* \\ &\quad + (\mathbf{a}_S^*)^\top \mathbf{B}_{LSL}\mathbf{a}_L + (\mathbf{a}_S^*)^\top \mathbf{B}_{LSS}\mathbf{a}_S^*, \end{aligned} \quad (8b)$$

where approximated sub-scale coefficient $\mathbf{a}_S^*$ is computed as $\mathbf{a}_S^* := \tilde{\mathbf{A}}_{SS} \text{ResS}(\mathbf{a}_L) + \text{ResS}(\mathbf{a}_L)^\top \tilde{\mathbf{B}}_{SSS} \text{ResS}(\mathbf{a}_L)$.

3. NUMERICAL RESULTS

We investigate the numerical accuracy of G-ROM, C-D2-VMS-ROM, and new R-D2-VMS-ROM in the numerical simulation of a 2D channel flow past a circular cylinder at Reynolds numbers $\text{Re} = 1000$. We present the numerical accuracy of the ROM models for two different regimes: (i) **reconstructive regime**: we build the ROM basis and operators, and D2 operators by using the FOM snapshots from $t = 13$ to $t = 16$. Then, we test ROMs over the same time interval. (ii) **predictive regime**: we build the ROM basis and operators by using the FOM snapshots from $t = 13$ to $t = 16$, and D2 operators by using the FOM snapshots from $t = 13$ to $t = 13.134$. Then, we test ROMs over the longer time interval, $t = 16$ to $t = 23$.

Furthermore, in our numerical accuracy investigation of the ROMs, we use the average $\mathcal{L}^2$ projection error:

$$\mathcal{E}_{\text{avg} \mathcal{L}^2 \text{proj}} = \frac{1}{M} \sum_{k=1}^M \left\| \mathbf{u}_L(t_k) - \sum_{i=1}^L \left(\mathbf{u}^{\text{FOM}}(t_k), \boldsymbol{\varphi}_i \right) \boldsymbol{\varphi}_i \right\|_{\mathcal{L}^2}. \quad (9)$$

In Tables 1-2, we list the average $\mathcal{L}^2$ projection errors of G-ROM, C-D2-VMS-ROM, and R-D2-VMS-ROM for the reconstructive and predictive regimes, respectively. In Table 1, we observe that C-D2-VMS-ROM and R-D2-VMS-ROM yield much better accuracy (for some values, the improvement is more than 2 orders of magnitude) than G-ROM in the reconstructive regime. C-D2-VMS-ROM and R-D2-VMS-ROM have similar accuracy behavior. In Table 2, we still observe that C-D2-VMS-ROM and R-D2-VMS-ROM yield much better accuracy (for some values, the improvement is more than 1 order of magnitude) than G-ROM in the predictive regime. Furthermore, R-D2-VMS-ROM yields better accuracy than C-D2-VMS-ROM.

In Figures 1-2, we plot the kinetic energy of the FOM projection, G-ROM, C-D2-VMS-ROM, and R-D2-VMS-ROM for the reconstructive and predictive regimes, respectively. We fix the large-scale ROM dimension $L = 6$, to compare the kinetic energy behavior of C-D2-VMS-ROM and R-D2-VMS-ROM. We observe that R-D2-VMS-ROM is sig-

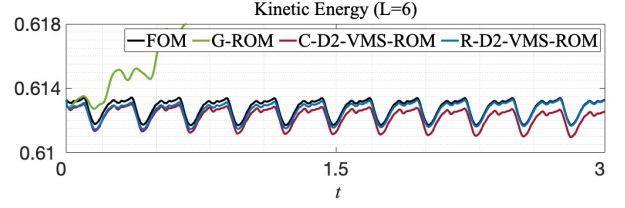


Fig. 1. Reconstructive regime; kinetic energy of ROMs.

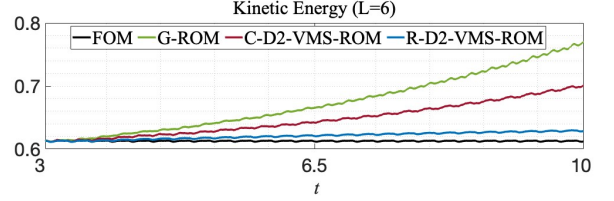


Fig. 2. Predictive regime; kinetic energy of ROMs.

nificantly more accurate than C-D2-VMS-ROM, especially in the predictive regime.

Acknowledgments: Research is partially funded by the Spanish Research Agency Juan de la Cierva 2022 with 2023/1061 and PID2021-123153OB-C21 - Feder Fund Grants.

REFERENCES

- Ballarin, F., Rebollo, T.C., Avila, E.D., Mármol, M.G., and Rozza, G. (2020). Certified reduced basis vms-smagorinsky model for natural convection flow in a cavity with variable height. *Computers & Mathematics with Applications*, 80(5), 973–989.
- Mou, C., Koc, B., San, O., Rebholz, L.G., and Iliescu, T. (2021). Data-driven variational multiscale reduced order models. *Computer Methods in Applied Mechanics and Engineering*, 373, 113470.
- Rebollo, T.C. and Coronil, D.F. (2024). Data-driven stabilized finite element solution of advection-dominated flow problems. *Mathematics and Computers in Simulation*, 226, 540–559.

Table 1. Reconstructive regime; average $\mathcal{L}^2$ projection error (9) for different L values.

L	G-ROM	C-D2-VMS-ROM	R-D2-VMS-ROM
2	4.94e-01	4.00e-03	5.06e-03
3	5.11e-01	3.09e-03	4.17e-03
4	5.98e-01	2.89e-03	1.45e-03
5	6.58e-01	6.07e-03	1.31e-03
6	1.50e-01	2.62e-03	9.83e-04
7	1.36e-01	2.76e-03	4.42e-03
8	7.08e-02	3.14e-03	1.32e-03

Table 2. Predictive regime; average $\mathcal{L}^2$ projection error (9) for different L values.

L	G-ROM	C-D2-VMS-ROM	R-D2-VMS-ROM
2	1.15e+00	4.11e-01	3.59e-01
3	9.22e-01	5.51e-01	6.16e-02
4	7.21e-01	1.98e-01	1.18e-01
5	7.28e-01	5.81e-01	3.11e-01
6	3.54e-01	1.48e-01	4.36e-02
7	3.02e-01	2.81e-01	5.29e-02
8	1.59e-01	9.44e-02	2.53e-02

Data based modeling and reduced order modeling for port-Hamiltonian descriptor systems

C. A. Beattie* V. Mehrmann**

* *Department of Mathematics, Virginia Tech, Blacksburg, VA 24061, USA. (e-mail: beattie@vt.edu).*

** *Institut für Mathematik MA 4-5, TU Berlin, Str. des 17. Juni 136, D-10623 Berlin, FRG. (e-mail: mehrmann@math.tu-berlin.de)*

1. INTRODUCTION

We present a method to construct low order linear time-invariant (LTI) port-Hamiltonian (pH) descriptor systems from observed response data in the frequency domain.

The simplest form of an LTI pH descriptor system is

$$\begin{aligned} \mathbf{E}\dot{\mathbf{x}} &= (\mathbf{J} - \mathbf{R})\mathbf{x} + \mathbf{B}\mathbf{u} \\ \mathbf{y} &= \mathbf{B}^T\mathbf{x} \end{aligned} \quad \text{and} \quad \mathbf{x}(0) = \mathbf{x}_0, \quad (1)$$

with $\mathbf{J} = -\mathbf{J}^T$, $\mathbf{E} = \mathbf{E}^T \geq 0$, $\mathbf{R} = \mathbf{R}^T \geq 0 \in \mathbb{R}^{n,n}$, $\mathbf{B} \in \mathbb{R}^{n,m}$, and the quadratic Hamiltonian (modeling the internal energy stored in the system) is $\mathcal{H} = \frac{1}{2}\mathbf{x}^T\mathbf{E}\mathbf{x}$, see e.g. Mehrmann and Unger (2023). PH descriptor systems are closely related to passive and positive-real systems, see Cherifi et al. (2023). We assume that the underlying physical system is passive and that it can be represented as a pH system of the form (1) with positive definite $\mathbf{E}$.

A classical approach to derive a low-order pH descriptor system from observed data is to first derive a realization

$$\begin{aligned} \dot{\mathbf{x}} &= \tilde{\mathbf{A}}\mathbf{x} + \tilde{\mathbf{B}}\mathbf{u}, \quad \mathbf{x}(0) = \mathbf{x}_0, \\ \mathbf{y} &= \tilde{\mathbf{C}}\mathbf{x}, \end{aligned} \quad (2)$$

with $\tilde{\mathbf{A}} \in \mathbb{R}^{n,n}$, $\tilde{\mathbf{B}}, \tilde{\mathbf{C}}^T \in \mathbb{R}^{n,m}$, e.g. via the Loewner approach Mayo and Antoulas (2007). The resulting system is typically close to a passive system. If the system is passive, see e.g. Cherifi et al. (2023), then there exists a matrix $\mathbf{E} = \mathbf{E}^T > 0$ such that

$$\begin{bmatrix} -\mathbf{E}\tilde{\mathbf{A}} - \tilde{\mathbf{A}}^T\mathbf{E} & \tilde{\mathbf{C}}^T - \mathbf{E}\tilde{\mathbf{B}} \\ \tilde{\mathbf{C}} - \mathbf{E}\tilde{\mathbf{B}}^T & \mathbf{0} \end{bmatrix} \geq \mathbf{0}. \quad (3)$$

Setting $\mathbf{B} := \mathbf{E}\tilde{\mathbf{B}}$ and $\mathbf{E}\tilde{\mathbf{A}} = \mathbf{J} - \mathbf{R}$, with $\mathbf{J} = -\mathbf{J}^T$ and $\mathbf{R} = \mathbf{R}^T \geq 0$, see e.g. Beattie et al. (2019), gives the pH descriptor system

$$\begin{aligned} \mathbf{E}\dot{\mathbf{x}} &= (\mathbf{J} - \mathbf{R})\mathbf{x} + \mathbf{B}\mathbf{u}, \\ \mathbf{y} &= \mathbf{B}^T\mathbf{x}. \end{aligned} \quad (4)$$

For a non-passive realization one typically applies small perturbations Alam et al. (2011); Brüll and Schröder (2013); Grivet-Talocia (2004); Freund et al. (2007) to obtain a nearby passive system for which the approach can be applied.

Unfortunately, proceeding in this way requires the solution of an optimization problem to enforce passivity and a solution of (3), both of which is prohibitive for large scale

systems. So typically it is necessary to first perform a model reduction, see e.g. Antoulas et al. (2020), which may, however, destroy the pH structure. For pH descriptor systems, such model reduction methods are well established, see e.g. Beattie et al. (2022b); Hauschild et al. (2019).

All these approaches, however, require an explicit realization of the system. An ideal procedure would directly generate a pH descriptor system from data. But so far only indirect methods are available Benner et al. (2020); Cherifi et al. (2019). We propose a direct approach.

2. A DATA BASED DIRECT APPROACH

Suppose that we are able to sample the system response of a passive system at (complex) frequencies $\{\sigma_1, \sigma_2, \dots, \sigma_r\} \subset \mathbb{C}_+$ with corresponding input profiles $\{b_1, b_2, \dots, b_r\}$. The input-output map of (2) in frequency domain is associated to the transfer function $\mathcal{G}(s) = \tilde{\mathbf{C}}(s\mathbf{I} - \tilde{\mathbf{A}})^{-1}\tilde{\mathbf{B}}$, so we have access to the data, $g_k = \mathcal{G}(\sigma_k)b_k$, $k = 1, \dots, r$ but not to a realization (2).

We may construct an interpolating subspace, $\mathfrak{W}_r = \text{Ran}(\mathbf{V}_r)$, where $\mathbf{V}_r = [\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r]$ and $\mathbf{v}_k = (\sigma_k\mathbf{I} - \tilde{\mathbf{A}})^{-1}\tilde{\mathbf{B}}b_k$, $k = 1, \dots, r$; or equivalently,

$$\mathbf{V}_r\boldsymbol{\Sigma}_r - \tilde{\mathbf{A}}\mathbf{V}_r = \tilde{\mathbf{B}}\mathbb{B}_r \quad (5)$$

with $\boldsymbol{\Sigma}_r = \text{diag}\{\sigma_1, \sigma_2, \dots, \sigma_r\}$ and $\mathbb{B}_r = [b_1, b_2, \dots, b_r]$. Then define $\mathbf{E}_r = \mathbf{V}_r^*\mathbf{E}\mathbf{V}_r$ and the auxiliary modeling space, $\mathfrak{W}_r = \text{Ran}(\mathbf{W}_r)$, with $\mathbf{W}_r = \mathbf{E}\mathbf{V}_r$. Using the ansatz, $\mathbf{x} \approx \mathbf{V}_r\mathbf{x}_r$, in (4) and applying a Petrov-Galerkin condition that forces residual orthogonality (in $\mathbb{C}^n$) to $\mathfrak{W}_r$, we obtain a reduced order model

$$\begin{aligned} \mathbf{E}_r\dot{\mathbf{x}}_r &= (\mathbf{J}_r - \mathbf{R}_r)\mathbf{x}_r + \mathbf{C}_r^T\mathbf{u}, \\ \mathbf{y} &= \mathbf{C}_r\mathbf{x}_r, \end{aligned} \quad (6)$$

where $\mathbf{J}_r = \mathbf{V}_r^*\mathbf{J}\mathbf{V}_r = -\mathbf{J}_r^*$, $\mathbf{R}_r = \mathbf{V}_r^*\mathbf{R}\mathbf{V}_r = \mathbf{R}_r^* \geq 0$, and $\mathbf{C}_r = \tilde{\mathbf{C}}\mathbf{V}_r = \tilde{\mathbf{B}}^T\mathbf{E}\mathbf{V}_r = \tilde{\mathbf{B}}^T\mathbf{W}_r$. Note that $g_k = \mathcal{G}(\sigma_k)b_k = \tilde{\mathbf{C}}(\sigma_k\mathbf{I} - \tilde{\mathbf{A}})^{-1}\tilde{\mathbf{B}}b_k = \tilde{\mathbf{C}}\mathbf{v}_k$, so with $\mathbb{G}_r = [g_1, g_2, \dots, g_r]$, we have $\mathbf{C}_r = \mathbb{G}_r$, and likewise, applying the same Petrov-Galerkin condition to (5) implies

$$\mathbf{V}_r^*\mathbf{E}\mathbf{V}_r\boldsymbol{\Sigma}_r - \mathbf{V}_r^*\mathbf{E}\tilde{\mathbf{A}}\mathbf{V}_r = \mathbf{V}_r^*\mathbf{E}\tilde{\mathbf{B}}\mathbb{B}_r = \mathbf{C}_r^*\mathbb{B}_r$$

which gives

$$\mathbf{E}_r\boldsymbol{\Sigma}_r - (\mathbf{J}_r - \mathbf{R}_r) = \mathbf{C}_r^*\mathbb{B}_r. \quad (7)$$

Summing with the conjugate transpose of (7) we have

$$\mathbf{E}_r\boldsymbol{\Sigma}_r + \overline{\boldsymbol{\Sigma}_r}\mathbf{E}_r + 2\mathbf{R}_r = \mathbb{G}_r^*\mathbb{B}_r + \mathbb{B}_r^*\mathbb{G}_r. \quad (8)$$

Denote by $\mathbb{S}^{r \times r}$ the set of Hermitian $r \times r$ matrices, by $\mathcal{P}_r$ the closed, convex cone of positive semidefinite (Hermitian) $r \times r$ matrices, and by $\mathcal{P}_r^\circ$ its interior, i.e., the open cone of $r \times r$ strictly positive-definite Hermitian matrices. Then define the linear operator $\mathcal{L} : \mathbb{S}^{r \times r} \rightarrow \mathbb{S}^{r \times r}$ as

$$\mathcal{L}(\mathbf{M}) = \mathbf{M} \Sigma_r + \overline{\Sigma_r} \mathbf{M}$$

and observe that since $\{\sigma_1, \sigma_2, \dots, \sigma_r\}$ is contained in the open right half-plane, $\mathcal{L}^{-1}$ is well-defined and cone-preserving: $\mathcal{L}^{-1} : \mathcal{P}_r \rightarrow \mathcal{P}_r$ and $\mathcal{L}^{-1} : \mathcal{P}_r^\circ \rightarrow \mathcal{P}_r^\circ$. Then (8) leads to

$$\mathbf{E}_r + 2\mathcal{L}^{-1}(\mathbf{R}_r) = \mathcal{L}^{-1}(\mathbb{G}_r^* \mathbb{B}_r + \mathbb{B}_r^* \mathbb{G}_r). \quad (9)$$

Note that in all nontrivial circumstances, $\mathbb{G}_r^* \mathbb{B}_r + \mathbb{B}_r^* \mathbb{G}_r$ is indefinite, however, since $\mathbf{E}_r > 0$ and $\mathcal{L}^{-1}(\mathbf{R}_r) \geq \mathbf{0}$, whenever the original system (2) is passive, then $\mathcal{L}^{-1}(\mathbb{G}_r^* \mathbb{B}_r + \mathbb{B}_r^* \mathbb{G}_r)$ itself must be positive definite if the original system (2) is passive. Note also that the condition $\mathcal{L}^{-1}(\mathbb{G}_r^* \mathbb{B}_r + \mathbb{B}_r^* \mathbb{G}_r) > 0$ involves only observed quantities, and it can be checked computationally using methods for the numerical solution of Sylvester equations Golub and Van Loan (1996).

Theorem 1. Given complex frequencies $\{\sigma_1, \sigma_2, \dots, \sigma_r\} \subset \mathbb{C}_+$ and corresponding input profiles $\{b_1, b_2, \dots, b_r\}$ together with the induced system responses $\{g_1, g_2, \dots, g_r\}$, where $g_k = \mathcal{G}(\sigma_k) b_k$, $k = 1, 2, \dots, r$. Let the matrices, $\mathbb{B}_r$ and $\mathbb{G}_r$, as well as the linear operator, $\mathcal{L}$, be defined as above. If the original system, (2), is passive then $\mathcal{L}^{-1}(\mathbb{G}_r^* \mathbb{B}_r + \mathbb{B}_r^* \mathbb{G}_r)$ is positive-definite.

Equivalently, if the data-based quantity, $\mathcal{L}^{-1}(\mathbb{G}_r^* \mathbb{B}_r + \mathbb{B}_r^* \mathbb{G}_r)$, fails to be positive-definite then the observed responses of the original system (2) are incompatible with passivity of the system; it cannot be expressed as a port-Hamiltonian system.

3. ANALYSIS OF THE PROCEDURE AND OPEN QUESTIONS

The described approach allows to produce a reduced order port-Hamiltonian from input-output data in frequency domain. However, there is freedom in the approach. since the representation of a standard system as port-Hamiltonian system is not unique. Any solution $\mathbf{E} > 0$ of (3) will lead to a port-Hamiltonian formulation. This freedom can be used to make the representation robust against perturbations, see Bankmann et al. (2020); Mehrmann and Dooren (2020). So one could first produce any reduced order model of the form (7) and then make it robust. We can also use a perturbation approach if the conditions of Theorem 1 are not met, see Beattie et al. (2022a).

Another open question is whether we can take a greedy approach by first taking a small number of sampling data and then increase the model order to achieve better approximations.

REFERENCES

Alam, R., Bora, S., Karow, M., Mehrmann, V., and Moro, J. (2011). Perturbation theory for Hamiltonian matrices and the distance to bounded-realness. *SIAM J. Matrix Anal. Appl.*, 32, 484–514.

Antoulas, A.C., Beattie, C.A., and Gugercin, S. (2020). *Interpolatory Methods for Model Reduction*. Computational Science & Engineering. Society for Industrial

and Applied Mathematics, Philadelphia, PA, USA. doi: 10.1137/1.9781611976083.

Bankmann, D., Mehrmann, V., Nesterov, Y., and Van Dooren, P. (2020). Computation of the analytic center of the solution set of the linear matrix inequality arising in continuous- and discrete-time passivity analysis. *Vietnam J. Math.*, 48, 633–660. doi:10.1007/s10013-020-00427-x.

Beattie, C., Mehrmann, V., and Van Dooren, P. (2019). Robust port-Hamiltonian representations of passive systems. *Automatica*, 100, 182–186. doi: 10.1016/j.automatica.2018.11.013.

Beattie, C., Mehrmann, V., and Xu, H. (2022a). Port-Hamiltonian realizations of linear time invariant systems. *ArXiv e-print 2201.05355*. URL <http://arxiv.org/abs/2201.05355>.

Beattie, C.A., Gugercin, S., and Mehrmann, V. (2022b). Structure-preserving interpolatory model reduction for port-Hamiltonian differential-algebraic systems. In C. Beattie, P. Benner, M. Embree, S. Gugercin, and S. Lefteriu (eds.), *Realization and Model Reduction of Dynamical Systems*, 235–254. Springer, Cham.

Benner, P., Goyal, P., and Van Dooren, P. (2020). Identification of port-Hamiltonian systems from frequency response data. *Systems Control Lett.*, 143, 104741. doi: 10.1016/j.sysconle.2020.104741.

Brüll, T. and Schröder, C. (2013). Dissipativity enforcement via perturbation of para-Hermitian pencils. *IEEE Trans. Circuits Syst. I. Regul. Pap.*, 60(1), 164–177.

Cherifi, K., Gernandt, H., and Hinsen, D. (2023). The difference between port-Hamiltonian, passive and positive real descriptor systems. *Math. Control Signals Systems*, 1–32. doi:10.1007/s00498-023-00373-2.

Cherifi, K., Mehrmann, V., and Hariche, K. (2019). Numerical methods to compute a minimal realization of a port-Hamiltonian system. *ArXiv e-print 1903.07042*. URL <http://arxiv.org/abs/1903.07042>.

Freund, R., Jarre, F., and Vogelbusch, C. (2007). Nonlinear semidefinite programming: Sensitivity, convergence and an application in passive reduced-order modeling. *Math. Programming*, 109, 581–611.

Golub, G.H. and Van Loan, C.F. (1996). *Matrix Computations*. The Johns Hopkins University Press, Baltimore, MD, 4th edition.

Grivet-Talocia, S. (2004). Passivity enforcement via perturbation of Hamiltonian matrices. *IEEE Trans. Circuits and Systems*, 51, 1755–1769.

Hauschild, S.A., Marheineke, N., and Mehrmann, V. (2019). Model reduction techniques for linear constant coefficient port-Hamiltonian differential-algebraic systems. *Control & Cybernetics*, 48, 125–152.

Mayo, A.J. and Antoulas, A.C. (2007). A framework for the solution of the generalized realization problem. *Linear Algebra Appl.*, 425(2-3), 634–662. doi: 10.1016/j.laa.2007.03.008.

Mehrmann, V. and Dooren, P.V. (2020). Optimal robustness of port-Hamiltonian systems. *SIAM J. Matrix Anal. Appl.*, 41(1), 134–151. doi:10.1137/19M1259092.

Mehrmann, V. and Unger, B. (2023). Control of port-Hamiltonian differential-algebraic systems and applications. *Acta Numerica*, 395–515. doi: 10.1017/S0962492922000083.

Adaptive Model Hierarchies for Multi-Query Scenarios^{*}

Hendrik Kleikamp^{*} Mario Ohlberger^{*}

^{*} *Institute for Analysis and Numerics, Mathematics Münster,
University of Münster, Einsteinstrasse 62, 48149 Münster, Germany
(e-mail: hendrik.kleikamp@uni-muenster.de).*

1. INTRODUCTION

In this contribution we present an abstract framework for adaptive model hierarchies together with several instances of hierarchies for specific applications. The hierarchy is particularly useful when integrated within an outer loop, for instance an optimization iteration or a Monte Carlo estimation where for a large set of requests answers fulfilling certain criteria are required. Within the hierarchy, multiple models are combined and interact with each other pursuing the overall goal to reduce the run time in a multi-query context. To this end, models with different accuracies and effort for evaluation are used in such a way that the cheapest (and typically least accurate) models are evaluated first when a request comes in. If the result fulfills a prescribed criterion, it can be returned to the outer loop. Otherwise, the model hierarchy falls back to more costly, but at the same time more accurate, models. The cheaper models are improved by means of training data gathered whenever the more accurate models are evaluated.

In the next section we provide an abstract and detailed description of the components of the hierarchy and their interaction. Subsequently, various applications are briefly discussed for which hierarchies with different numbers of stages were developed.

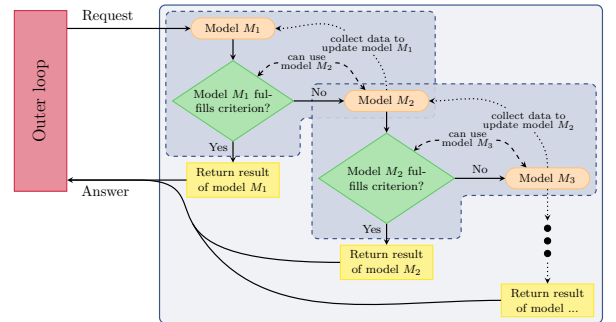
2. ABSTRACT DESCRIPTION

The idea of a model hierarchy in the context of parameterized partial differential equations (PDEs) was originally introduced in Haasdonk et al. (2023). Here we describe the concept in a general form that is applicable in a wide range of scenarios and for several types of models.

In our abstract description we consider a solution operator $S: \mathcal{P} \rightarrow \mathcal{V}$ that maps from an admissible input space $\mathcal{P}$ to a possibly infinite dimensional solution space $\mathcal{V}$, where usually we know that S exists, but it might not be accessible. A typical example would be the solution operator of a parameterized PDE where $\mathcal{P}$ corresponds to the parameter set. Furthermore, we assume that we are given a hierarchy of approximate models $M_1, M_2, \dots$ that approximate the map S , where two successive models M_l and M_{l+1} in the hierarchy satisfy the following multi-fidelity assumptions:

- $\mathcal{C}(M_l) < \mathcal{C}(M_{l+1})$, where $\mathcal{C}(M_l)$ denotes the model complexity as a measure for the runtime.
- $\mathcal{E}(M_l(\mu), \mu) \geq \mathcal{E}(M_{l+1}(\mu), \mu)$, where $\mathcal{E}(M_l(\mu), \mu)$ denotes an error measure w.r.t. $S(\mu)$ for $\mu \in \mathcal{P}$.
- Model M_l can be improved by means of information from model M_{l+1} .

Assume now that a request $\mu \in \mathcal{P}$ in an outer multi-query loop needs to be processed. The request is first passed to model M_1 which produces a result $M_1(\mu)$ for the request. This result is evaluated using the error measure, i.e. it is verified whether $\mathcal{E}(M_1(\mu), \mu) \leq \text{TOL}$ is satisfied. In order to check the criterion it might be necessary to also retrieve additional information from model M_2 . In general, if model M_l fulfills the criterion, the result of model M_l is returned to the outer loop. If the criterion is not met, the request is passed to model M_{l+1} which is assumed to be more accurate and is therefore more likely to fulfill the prescribed criterion. Model M_{l+1} now proceeds similar to model M_l , i.e. the request is processed resulting in an answer of model M_{l+1} . When evaluating model M_{l+1} , data is collected that can be used, according to the third assumption from above, to improve model M_l . Hence, model M_l is constructed and enhanced in an adaptive manner. The result of M_{l+1} might now be passed on, depending on the structure of the remaining parts of the hierarchy. Due to the involved check of the accuracy criterion for all results, the output of the model hierarchy is certified. The overall hierarchical structure of the multi-fidelity algorithm is shown in Fig. 1 when applied in an outer loop for a hierarchy consisting of multiple stages. For the algorithm performing the outer loop, the hierarchy behaves like a single model that returns a certain result of guaranteed accuracy. All the internal model selection and adaptation is invisible from the outside.



3. APPLICATIONS

In the following paragraphs we discuss applications where the concept of an adaptive model hierarchy was utilized successfully to speed up different computational tasks.

PDE-constrained optimization. In Keil et al. (2022) we introduced a two-stage hierarchy consisting of a full order model and a machine learning surrogate for a PDE-constrained optimization problem that occurs in enhanced oil recovery. The machine learning surrogate approximates the objective function and is based on training data gathered when evaluating the full order model (involving the costly simulation of a three-phase flow in a porous medium). From the point of view of the model hierarchy as shown in Fig. 1, the machine learning surrogate corresponds to model M_1 and during its evaluation an optimization problem for the approximate objective function is solved. The accuracy of the result of this inner optimization loop is then evaluated by computing an approximate gradient using the full model M_2 .

Parametrized parabolic PDEs. As a second example, we considered in Haasdonk et al. (2023) parametrized parabolic PDEs where the hierarchy consists of a full order model, a reduced basis reduced order model and a machine learning model. The latter model is based on the approach of learning the reduced coefficients with respect to a reduced basis as introduced in Hesthaven and Ubbiali (2018). The reduced basis is computed using evaluations of the full order model whereas the machine learning surrogate is trained based on solutions of the reduced basis model and therefore contains an additional layer of approximation. Accuracy of the reduced basis and the machine learning model is verified by means of an a posteriori error estimator for reduced models of parabolic problems. Since the machine learning surrogate uses the same reduced space as the reduced basis model, the a posteriori error estimator is applicable also to the machine learning approximation. Hence, a close connection between the two surrogate models facilitates their interaction in the hierarchy in this case. The full order model here serves as reference and is therefore assumed to be arbitrarily accurate. Hence, no accuracy check of the full order solution is performed.

Parametrized optimal control problems. A three-stage adaptive model hierarchy for linear-quadratic optimal control problems with parameter-dependent system components was developed in Kleikamp (2024). The general structure is comparable to the one for parabolic PDEs. In particular, the three involved models and their interaction are similar and an a posteriori error estimator is used to certify the results obtained by the reduced models. The special structure of the considered optimal control problems allows to identify solutions to the associated optimality system by the optimal adjoint at final time. The reduced basis model thus builds on an approximation of the set of optimal final time adjoints by linear subspaces. As before, the machine learning surrogate makes use of the same reduced space which allows to reuse the a posteriori

error estimate of the reduced basis model.

An additional speedup can be obtained by also reducing the primal and adjoint trajectories in an efficient manner as described in Kleikamp and Renelt (2024). The resulting fully reduced model is based on the reduced basis model for approximate final time adjoints. It is moreover possible to incorporate machine learning in the fully reduced model. We hence obtain a four-stage hierarchy consisting of the full order model (FOM), the reduced basis reduced order model (RB-ROM), the fully reduced model (F-ROM) and the machine learning fully reduced model (ML-F-ROM). In Tab. 1 we present the results in terms of number of evaluations and average run time of the individual models within the four-stage hierarchy, when querying the hierarchy for 10,000 randomly chosen parameters and a fixed error tolerance of 10^{-4} in the cookie baking example described in Kleikamp and Renelt (2024). As can be seen

Table 1. Results of the four-stage model hierarchy applied to the cookie baking test case

Model	Number of solves	Average time for error estimation and solving [s]
FOM	4	76.24
RB-ROM	12	19.55
F-ROM	437	1.03
ML-F-ROM	9,547	0.54

from the numerical results depicted in Tab. 1, the ML-F-ROM, which is the fastest of the four involved models, is sufficiently accurate in more than 95% of the calls to the hierarchy. In contrast, the full order model has to be solved only four times in order to meet the accuracy requirements.

4. CONCLUSION

The introduced concept of adaptive model hierarchies provides a possibility to combine different models of varying complexity within a joint hierarchy that can be evaluated efficiently. As discussed in the last section, model hierarchies are applicable in different contexts and make use of the advantages of all involved models.

REFERENCES

- Haasdonk, B., Kleikamp, H., Ohlberger, M., Schindler, F., and Wenzel, T. (2023). A new certified hierarchical and adaptive RB-ML-ROM surrogate model for parametrized PDEs. *SIAM J. Sci. Comput.*, 45(3), A1039–A1065.
- Hesthaven, J. and Ubbiali, S. (2018). Non-intrusive reduced order modeling of nonlinear problems using neural networks. *J. Comput. Phys.*, 363, 55–78.
- Keil, T., Kleikamp, H., Lorentzen, R.J., Oguntola, M.B., and Ohlberger, M. (2022). Adaptive machine learning-based surrogate modeling to accelerate PDE-constrained optimization in enhanced oil recovery. *Adv. Comput. Math.*, 48(6), 73.
- Kleikamp, H. (2024). Application of an adaptive model hierarchy to parametrized optimal control problems. *Proceedings of the Conference Algorithmy*, 66–75.
- Kleikamp, H. and Renelt, L. (2024). Two-stage model reduction approaches for the efficient and certified solution of parametrized optimal control problems. *arXiv preprint, DOI: 10.48550/ARXIV.2408.15900*.

A Green Markup for the Assessment of Optimized Circulation Plans^{*}

Matthias Rößler^{*} Daniele Giannandrea^{**} Niki Popper^{*,**}

^{*} DWH Simulation Services, Neustiftgasse 57-59, 1070 Vienna, Austria (e-mail: matthias.roessler@dwh.at).

^{**} TU Wien, Institute of Information Systems Engineering, Favoritenstraße 9-11, 1040 Vienna, Austria

1. INTRODUCTION

Due to the climate crisis, reducing energy consumption and greenhouse gas emissions has gained more and more relevance in railway traffic over the past years. While many studies propose optimization methods that consider energy consumption during circulation planning directly (Fernández et al., 2019), there is still an abundance of other methods (Piu and Speranza, 2014) that do not take energy consumption into account, especially for real-world applications. We propose a simulation-based approach to assess the quality of such circulation plans from an energy consumption and a robustness perspective, the *Green Markup*, which enables comparison of circulation plans and may be considered for optimization within a feedback loop.

2. DEFINITION

For a markup we want to compare the effects of a certain circulation plan (circulation scenario; *cs*) to an idealized base scenario (*bs*) where traction units are available if needed, in a realistic setting. For that, we use a simulation model that calculates the delay propagation within a time table based on injected primary delays (Rößler et al., 2020). In both scenarios, the same trains are simulated using the same primary delays. In the circulation scenario, additional empty runs are introduced through the circulation plan, for which no primary delays are added.

A green markup should indicate the performance of circulation plans both in terms of the corresponding energy consumption as well as their robustness against delays. For each of the two characteristics, robustness and energy consumption, we define a separate markup, in which we compare the two scenarios.

The markups can be calculated for different subsets within the time table. While in principle all possible subsets are feasible, the following are the most reasonable:

- *Global Markup*: sums up all relevant values for the whole time table.
- *Circulation Markup*: sums up all relevant values for the tasks used in the circulation plan.

^{*} This research was funded through the Green-TrAIIn-Plan project (FFG project number 892235) by the Federal Ministry for Climate Protection, Environment, Energy, Mobility, Innovation and Technology (BMK) as part of the IKT der Zukunft initiative AI for Green. We thank ÖBB and its digitalisation program ARP for providing the data and valuable insights used in this study.

- *Train Markup*: sums up all relevant values for the tasks driven by a certain train.

While Global Markup and Circulation Markup may be used to assess the overall quality of a circulation plan, the Train Markup can be used within a feedback loop with an optimization algorithm.

2.1 Delay Markup

The delay markup was already introduced in Rößler et al. (2020) and is defined as

$$m_d := \frac{\sum_{t \in T} SD_{t;cs}}{\sum_{t \in T} SD_{t;bs}}, \quad (1)$$

with $SD_{t;cs}$ and $SD_{t;bs}$ referring to the secondary delays of a task t for the circulation scenario and the base scenario, respectively. The set T depends on which type of markup should be calculated.

As delayed empty runs only impact the quality of a circulation plan if they affect other trains, their secondary delays are not directly considered in the markup calculation.

2.2 Energy Markup

Analogously, we define the energy markup:

$$m_e := \frac{\sum_{t \in T} EC_{t;cs} + \sum_{r \in ER} EC_{r;cs}}{\sum_{t \in T} EC_{t;bs} + \sum_{r \in ER} \tilde{EC}_r} \quad (2)$$

with EC_{cs} and EC_{bs} denoting the energy consumption (also including recuperation) of the circulation scenario and the base scenario, respectively. Again, the set T depends on the type of markup to be calculated. Additionally, for the energy markup the energy consumption of the empty runs (ER) in the circulation scenario ($EC_{r;cs}$) contributes to the markup. It is compared to an idealized energy consumption value, that the empty run would need with no interference from the rest of the time table ($\tilde{EC}_r$).

The energy consumption values are approximated based on historical energy data from the Austrian railway system. For the calculation, geographical information of the tracks, travel times, weight and length of the train, technical data of the locomotives, and also planned (i.e. given in the time table) and unplanned (i.e. made necessary during simulation) stops during a trip are taken into account.

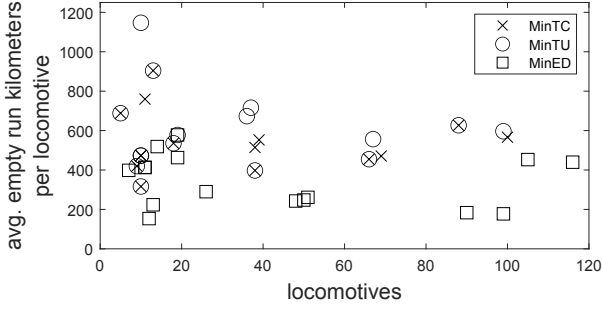


Fig. 1. Number of locomotives and average empty run kilometers per locomotive of different circulation plans

2.3 Green Markup

Finally, we define the green markup as the weighted sum of the delay and energy markups m_d and m_e , namely

$$m := \frac{w_d m_d + w_e m_e}{w_d + w_e} = \frac{m_d + \lambda m_e}{1 + \lambda} \quad (3)$$

where w_d and w_e are the corresponding weights which can also be consolidated in the single parameter $\lambda = w_e/w_d$. Evaluating several circulation plans in terms of their energy and delay markups will give us a sense of the magnitudes of the markups, which in turn will allow us to choose a suitable value λ for the green markup.

All three markups are defined such that they are equal to or greater than 1, taking the value 1 for a circulation plan that does not lead to additional energy consumption or delay compared to the base scenario. A meaningful comparison of markups is only possible for evaluations based on the same time tables and primary delay distributions.

3. EXPERIMENTS

The circulation plans are created using the optimization approach presented in Frisch et al. (2021). The objective function is a weighted sum of the number of locomotives (l) and the total amount of empty run kilometers (km)

$$\min \quad w_l \cdot l + w_{km} \cdot km, \quad (4)$$

with weights w_l and w_{km} . We vary the weights to achieve different results, that mimic different requirements.

- MinTC: Minimize total costs.
- MinTU: Minimize the number of used traction units.
- MinED: Minimize empty run distance.

The weights for MinTC are chosen to roughly reflect the cost difference between a locomotive and a driven empty run kilometer, for MinTU and MinED the weight for the not-prioritized component is multiplied by a small $\varepsilon > 0$ to limit the respective usage.

The circulation plans are based on real-world time tables for a reference week, that contain both passenger and freight traffic. For the analyzed instances we only use freight trains, as the amount and distribution of energy consumption for freight and passenger traffic is vastly different. We create instances for several traction unit classes using different additional filters to gather a variety of circulation plans with different sizes. Figure 1 shows an overview of the KPIs for the different circulation plans.

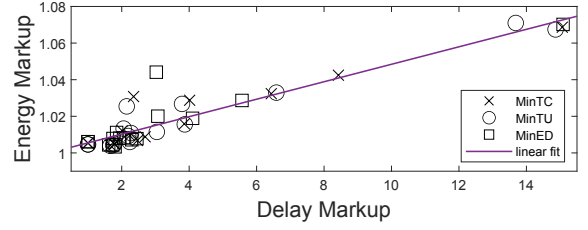


Fig. 2. Energy Markup plotted against Delay Markup

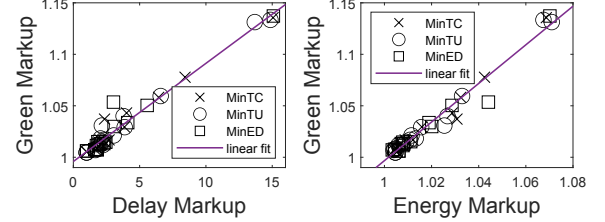


Fig. 3. Green Markup plotted against Delay Markup and Energy Markup

For the simulation, all trains present in the time tables are used to create a realistic model of the railway traffic. Each circulation plan instance is simulated on its own, and the delay and energy markup are calculated separately.

4. RESULTS

As Figure 2 shows, the delay markup for the chosen circulation plans lies in the range $[1, 16]$ whereas the energy markup lies within $[1, 1.08]$. The plot indicates a linear relation between the two values, and a group of slightly deviating values where the energy markups are greater (or delay markups smaller) than the linear fit would suggest.

The selection of λ must take into account the variations of the markups. Choosing

$$\lambda = \frac{\max(m_d) - \min(m_d)}{\max(m_e) - \min(m_e)} = 208 \quad (5)$$

yields the green markup presented in Figure 3. As can be seen, the relation between the delay and energy markups is preserved in the green markup, as the group of values deviating from the linear relation is visible in all plots.

REFERENCES

- Fernández, P.M., Sanchís, I.V., Yepes, V., and Franco, R.I. (2019). A review of modelling and optimisation methods applied to railways energy consumption. *Journal of Cleaner Production*, 222, 153–162.
- Frisch, S., Hungerländer, P., Jellen, A., Primas, B., Steininger, S., and Weinberger, D. (2021). Solving a real-world locomotive scheduling problem with maintenance constraints. *Transportation Research Part B: Methodological*, 150, 386–409.
- Piu, F. and Speranza, M.G. (2014). The locomotive assignment problem: a survey on optimization models. *International Transactions in Operational Research*, 21(3), 327–352.
- Rößler, M., Wastian, M., Jellen, A., Frisch, S., Weinberger, D., Hungerländer, P., Bicher, M., and Popper, N. (2020). Simulation and optimization of traction unit circulations. In *2020 Winter Simulation Conference (WSC)*, 90–101. doi:10.1109/WSC48552.2020.9383926.

Vector Field Construction for Football Game-Flow Evaluation

Tenpei Morishita*, Yuji Aruga**, Masao Nakayama***, Akifumi Kijima**, Hiroyuki Shima*

* *Department of Environmental Sciences, University of Yamanashi, Kofu, Yamanashi 400-8510, Japan
(e-mail: t.morishita@yamanashi.ac.jp) (e-mail: hshima@yamanashi.ac.jp)*

** *Faculty of Education, University of Yamanashi, Kofu, Yamanashi 400-8510, Japan*

*** *Faculty of Health and Sport Sciences, University of Tsukuba, Tsukuba, Ibaraki 305-8574, Japan*

Abstract: Traditional methods for analysing football gameplay involve observing players' movements in real time and using the data to explain tactics and player interactions. In contrast, this study introduces a more sophisticated mathematical approach to elucidate typical game flows and tactical features. Specifically, vector analysis was applied to the direction and length of the last pass observed in real football games, and potential fields were derived from the vector fields of the last pass. This approach enabled a visual distinction between unconscious pass flows that occur along potential gradients and conscious pass flows that do not follow gradients. The vector analysis also visualised the spontaneous formation of low-potential areas where the last passes are concentrated in front of the goal area, as well as the characteristics of cross passes near the penalty area. The results clearly show that the vector field-based approach provides useful insights into tactical analysis and strategy optimisation and offers a new perspective to sports science.

Keywords: Data-Driven Models; Modelling for Control and Real-Time Applications; Numerical Simulation.

1. INTRODUCTION

Understanding player behaviour and team dynamics in football has long been of interest to sports scientists and analysts. Traditional research has focused on observing player and ball movement and using statistical analysis to identify patterns and improve team strategy (Duarte 2012, Kijima 2014, Yokoyama 2018, Chacoma 2021, Welch 2021, Narizuka 2023). In these studies, spatio-temporal tracking methods have frequently been used to analyse formations and player interactions, providing insights into player behaviour and performance during matches.

In this study, a novel mathematical approach was applied to this problem. In order to elucidate the basic principles governing the collective behaviour of footballers immediately before a shot, a vector analysis method was devised to be applied to the measured data of the last pass. The reason for focusing on the last pass is that it is an important play that directly leads to a goal and influences the flow of the game and team tactics. As the last pass leads directly to the shooter, most passes occur in one half of the pitch, i.e. in situations where one team is attacking the opposing team's position. In this situation, the balance between offence and defence is greatly disrupted. Therefore, it is statistically expected that attacking action in this situation will have the consistent objective of shooting and scoring a goal. Focusing on the last pass, one step before the shot, simplifies the purpose of the action to be analysed and leads to a proper interpretation of the numerical analysis results.

Our theoretical approach, based on vector analysis, boasts three advantages in analysing the dynamics of football

(Morishita 2024). First, it introduces new methods to apply the concept of potential fields to football analysis and provides new perspectives for interpreting player and ball movement. Second, by exploring these new perspectives in football analysis, this study expands knowledge in sports science and provides practical insights to improve team performance. Finally, these statistical features are expected to visualise the universal typicality of collective movements in the game of football and provide new means to deepen our understanding of sports science and collective action studies.

2. GAME DATA ACQUISITION

Of the total of 64 matches in the 2022 World Cup in Qatar, 39 matches were included in the analysis. The data was visually confirmed by video footage taken from a position from which the entire pitch could be seen, and the coordinates of the position of the last pass by an attacking player and the coordinates of the position of another player who received the pass were all recorded in pairs.

3. VECTOR FIELD OF LAST PASSES

Using the above measured data, the vector field of the last pass was defined in the following way. The starting point of the vector corresponds to the position where the last pass occurs, and the direction and length of the vector correspond to the direction and distance of the pass. Since the sparsity of the last-pass vector field can cause noise in numerical analysis, smoothing was applied using a Gaussian filter as shown in Figure 1.

From the vector fields thus obtained, the following two types

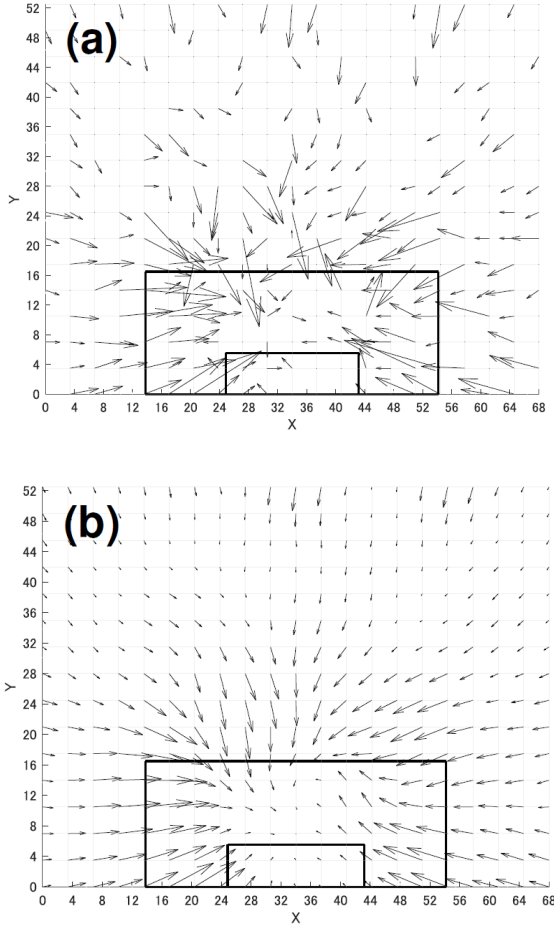


Fig. 1. (a) Last-pass vector fields constructed from measured football data. (b) Visualization of the last-pass vector field after smoothing.

of potential fields were derived: the scalar potential field ϕ and the vector potential field $\mathbf{A}$. Potential fields are mathematical concepts widely used in physics and engineering to analyse the properties of vector fields and their variations. It can be proved that any vector field can be decomposed into an orthogonal sum of the two components: the gradient one $\nabla\phi$ and rotational one $\nabla \times \mathbf{A}$; therefore, the two potential fields can be derived from each decomposed vector field. The scalar potential field ϕ represents the area of concentration of the last-pass and the natural convergence to that point, whereas the rotational component of the vector potential field, $\mathbf{A}$, is assumed to reflect tactics, including the intention to deviate from the gradient direction and the reaction of the defender to obstacles.

4. RESULTS AND DISCUSSIONS

Analysis of scalar potential fields showed that areas of low potential were concentrated in front of the goal area and that last passes tended to converge in this area. It was also found that the potential minima formed spontaneously on the slightly left side in front of the goal area. This fact means that there are more crossing passes released from the left side (right side from the offence's point of view), and the spatial distribution

of scalar potential field reflects the tactical behaviour of players aiming to cross passes to the advantageous position in front of the goal area.

Meanwhile, analysis of vector potential fields allowed the identification of areas with a significant rotational component. In particular, a strong rotational component was found near the centre of the penalty area and on both sides of it. This result may reflect the tactical intentions of the attackers. In other words, the large number of defensive players in the centre of the court means that it is difficult to pass from the front, and the attackers adopt a strategy of sending crosses in a large circle from the left and right sides. These features of the potential field distributions are based on the last-pass data of all teams in 39 matches, being not dependent on the individuality of a particular player or team but show a universal typology as football.

The following conclusions can be drawn from the results of the present analysis and the discussion on them. The potential field analysis method for the game of football provides a quantitative means of evaluating complex last-pass patterns by reducing them to simple mathematical models. This allows for the clear visualisation of trends in tactical behaviour that are universal to the game of football, without being dependent on players or teams. The analysis of scalar and vector potentials also provides an important perspective in the design of attacking and defensive tactics, contributing to the optimisation of match strategy.

ACKNOWLEDGEMENT

We thank Prof. Keiko Yokoyama in Nagoya University for fruitful discussions. This study was financially supported by JPSJ KAKENHI, Grant Numbers: 22K19727 and 23K20369.

REFERENCES

- Chacoma, A., Almeida, N., Perotti, J. I. and Billoni, O.V. (2021) Stochastic model for football's collective dynamics. *Physical Review E* 104, 024110.
- Duarte, R., Araújo, D., Correia, V. and Davids, K. (2012). Intra- and inter-group coordination patterns reveal collective behaviors of football players near the scoring zone. *Human Movement Science*, 31, 1639.
- Kijima, A., Yokoyama, K., Shima, H., and Yamamoto, Y. (2014). Emergence of self-similarity in football dynamics. *The European Physical Journal B*, 87, 1-6.
- Morishita, T., Aruga, Y., Nakayama, M., Kijima, A. and Shima, H. (2024). Tactical Analysis of Football Games by Vector Calculus of Last-Pass Performance. *SSRN*, 5015667 (<https://ssrn.com/abstract=5015667>).
- Narizuka, T., Takizawa, K. and Yamazaki, Y. (2023). Validation of a motion model for soccer players' sprint by means of tracking data. *Scientific Reports*, 13, 865.
- Welch, M., Schaerf, T. M. and Murphy, A. (2021) Collective states and their transitions in football. *PLoS ONE* 16, e0251970.
- Yokoyama, K., Shima, H., Fujii, K., Tabuchi, N., and Yamamoto, Y. (2018). Social forces for team coordination in ball possession game. *Physical Review E*, 97, 022410.

Learning non-intrusive ROMs from linear SDEs with additive noise

M. A. Freitag^{*} J. M. Nicolaus^{**} M. Redmann^{***}

^{*} University of Potsdam, Karl-Liebknecht-Str. 24-25, 14476 Potsdam

^{**} University of Potsdam, Karl-Liebknecht-Str. 24-25, 14476
Potsdam (e-mail: jan.martin.nicolaus@uni-potsdam.de)

^{***} Martin-Luther-University Halle-Wittenberg, Theodor-Lieser-Str. 5,
06120 Halle (Saale)

Keywords: black box modelling, data-driven modelling, non-intrusive model order reduction, stochastic linear systems

1. INTRODUCTION

Model Order Reduction (MOR) for stochastic linear systems is concerned with approximating the high-dimensional Full Order Model (FOM)

$$dx(t) = [Ax(t) + Bu(t)]dt + MdW(t), \quad x(0) = x_0, \quad (1)$$

by a surrogate model. Such arise, for example, from spatial discretisations of PDEs, with noisy boundary conditions. Here, the so-called drift of (1) is a function of the state $x(t) \in \mathbb{R}^n$ and the control $u(t) \in \mathbb{R}^m$ given by the respective multiplication with $A \in \mathbb{R}^{n \times n}$ and $B \in \mathbb{R}^{n \times m}$. The term $MdW(t)$ is called the diffusion term and describes the influence of the noise-generating process $W(t) \in \mathbb{R}^d$ on the state variable. In this case the process is a d -dimensional standard *Brownian motion*. The matrix $M \in \mathbb{R}^{n \times d}$ is called the diffusion coefficient. Due to the structure of the SDE (1) the FOM state variable $x(t)$ is Gaussian for each fixed time t if x_0 is Gaussian or constant. Hence, the distribution of x at the time $t \in [0, T]$ is completely determined by the expectation $E(t) \in \mathbb{R}^n$ and covariance $C(t) \in \mathbb{R}^{n \times n}$.

2. PROJECTION-BASED MOR

Projection-based MOR constructs such Reduced Order Models (ROMs) by approximating the FOM state variable in an r -dimensional subspace $\mathcal{V}_r$, that is, one assumes $x(t) \approx V_r V_r^T x(t) = V_r x_r(t)$, where orthogonal columns of $V_r = [v_1, \dots, v_r] \in \mathbb{R}^{n \times r}$ span the subspace $\mathcal{V}_r$. By the requiring a Galerkin condition on the residual of the FOM dynamics, one obtains

$$dx_r(t) = [A_r x_r(t) + B_r u(t)]dt + M_r dW(t), \quad (2)$$

with the *reduced coefficients*

$$A_r := V_r^T A V_r \in \mathbb{R}^{r \times r}, \quad B_r := V_r^T B \in \mathbb{R}^{r \times m},$$

$$M_r := V_r^T M \in \mathbb{R}^{r \times d}, \quad x_r(0) := V_r^T x_0 \in \mathbb{R}^r.$$

If $r \ll n$, then (2) is much cheaper to compute than the FOM (1). Since projections retain the linear structure

^{*} The research has been partially funded by the Deutsche Forschungsgemeinschaft (DFG) - Project-ID 318763901 - SFB1294 as well as by the DFG individual grant "Low-order approximations for large-scale problems arising in the context of high-dimensional PDEs and spatially discretized SPDEs" - project number 499366908.

of the original FOM equations, the ROM state $x_r(t)$ is Gaussian as well for each fixed $t \geq 0$. One can show that the expectation of the projected state $E_r(t) := E[x_r(t)]$ satisfies the ODE

$$\dot{E}_r(t) = A_r E_r(t) + B_r u(t), \quad E_r(0) = E[x_{r,0}]$$

and approximates the FOM expectation after lifting with V_r , since $V_r E_r(t) = V_r E[x_r(t)] = E[V_r V_r^T x(t)] \approx E[x(t)]$. Analogous results for the covariance matrix $C_r(t) = \text{cov}[x_r(t)]$ hold, see Freitag et al. (2024).

A popular data-driven method to construct $\mathcal{V}_r$ is *Proper Orthogonal Decomposition* (POD) method. In this method, the dominant subspace of observed *snapshots* $X^s = [x(t_1), \dots, x(t_s)]$ is chosen as $\mathcal{V}_r$. This is achieved by taking the r leading left-singular vectors of X^s as the columns of $V_r = [v_1, \dots, v_r]$. To construct a ROM in such a way, it is necessary to have access to the FOM matrices A, B , and M . Such methods are called *intrusive* and can be infeasible in the case of, for instance, black-box or legacy code.

3. NON-INTRUSIVE MOR

To address this issue, so-called *non-intrusive* methods have been developed. These methods do not require the availability of the FOM system coefficients, but instead rely on the availability of large amounts of data or the ability to query the FOM. One well-known method is the Operator Inference (OpInf) approach by Peherstorfer and Willcox (2016), which recently has been extended to the SDE setting by Freitag et al. (2024). We briefly illustrate this extension. In the standard OpInf approach for SDEs, one first collects L samples of s trajectory observations $x(t_1), \dots, x(t_s)$ of the FOM state, which are then used to compute approximations of the reduced expectation

$$E_{r,i}^L = V_r^T E_i^L, \quad E_i^L \approx E(t_i) := E[x_{t_i}], \quad i \in \{1, \dots, s\}.$$

of the ROM state variable x_r at the observation times. An approximation of the time derivative $\dot{E}_r(t)$ of the reduced expectation can be obtained by a finite difference approximation $E_{r,i}^{L,h}$ using $E_{r,i}^L := V_r^T E_i^L$, where h is given by the difference between the (equidistant) observation times t_i . Thus, to obtain approximations to A_r and B_r , one can solve the least-squares problem

$$[A_r^* \ B_r^*] = \underset{\tilde{O} \in \mathbb{R}^{r \times (r+m)}}{\text{argmin}} \|\tilde{O}^L - R^{L,h}\|_F, \quad (3)$$

where

$$D^L = \begin{bmatrix} E_{r,1}^L & \cdots & E_{r,s}^L \\ u(t_1) & \cdots & u(t_s) \end{bmatrix} \text{ and } R^{L,h} = [E_{r,1}^{L,h}, \dots, E_{r,s}^{L,h}].$$

Since D^L and $R^{L,h}$ are constructed from approximations, one can view (3) as a perturbed version of an unperturbed least-squares problem, where direct observations would be available. Freitag et al. (2024) show that if the data matrix $D^L \in \mathbb{R}^{(r+m) \times s}$ is of full rank, the unique solution is an almost-surely convergent estimator of the unperturbed least squares solution in the limit of $L \rightarrow \infty$ and $h \rightarrow 0$. Using an analogous approach, Freitag et al. (2024) obtain an estimation of the product $M_r M_r^T \in \mathbb{R}^{d \times d}$ for the diffusion coefficient by utilising the inferred A_r^* instead of A_r .

One drawback of the standard OpInf formulation, as Peherstorfer (2020) points out, is that the projected FOM trajectories can differ from the trajectories of the intrusive ROM (2). This so-called *closure error* arises due to the inability of ROMs of the form (2) to model the non-Markovian dynamics of the projected FOM state variable with respect to the subspace $\mathcal{V}_r$. Thus, the ROM obtained from standard OpInf can fail to approximate the reduced dynamics in $\mathcal{V}_r$. To address this issue, Peherstorfer (2020) proposes a modified sampling scheme called *re-projection*. We illustrate this method in the SDE setting of this paper. The core idea is to estimate the re-projection sampling, performed directly on the expectation, by computing the empirical mean of re-projected samples. To perform the re-projection scheme, access to the stepping function $f(x, u)$

$$\tilde{x}_{i+1} = f(\tilde{x}_i, u_i) = \tilde{A}\tilde{x}_i + \tilde{B}u_i + \tilde{M}z_i, z_i \sim \mathcal{N}(0_d, I_d),$$

of the time-discretised FOM is required. Here, 0_d is the d -dimensional zero vector and I_d the identity matrix of size $d \times d$. The sampling algorithm then computes trajectories $\{\hat{x}_i, i = 1, \dots, s\} \subset \mathbb{R}^r$ by projecting each query result onto $\mathcal{V}_r$, that is, one computes $\hat{x}_{i+1} = V_r^T f(V_r \hat{x}_i, u_i)$ and constructs the matrices

$$\hat{X}^s = \begin{bmatrix} \hat{x}_1^E & \cdots & \hat{x}_{s-1}^E \\ u(t_1) & \cdots & u(t_{s-1}) \end{bmatrix} \in \mathbb{R}^{(r+m) \times (s-1)} \text{ and } \hat{Y}^s = [\hat{x}_2^E \cdots \hat{x}_s^E] \in \mathbb{R}^{r \times (s-1)}$$

from the empirical estimation $\hat{x}_i^E$ of the expectation of $\hat{x}_i$. An approximation to the time-discrete reduced operators is obtained by solving the least-squares problem

$$[\hat{A}_r^* \hat{B}_r^*] = \underset{\hat{O} \in \mathbb{R}^{r \times (r+m)}}{\operatorname{argmin}} \|\hat{O}\hat{X}^s - \hat{Y}^s\|_F. \quad (4)$$

As in the standard OpInf method, the condition number of the data-matrix X^s can be improved by sampling linearly independent pairs of initial conditions and control. Lastly, the availability of f enables us to easily obtain an estimation of the projected time-discrete diffusion operator $\tilde{M}_r$. While one could proceed as in Freitag et al. (2024) by using the covariance matrices of the re-projected time-steps, it is much simpler to sample the projected time-stepping function f with a zero initial condition and control, since

$$V_r^T f(0_r, 0) = V_r^T \tilde{M}z, z \sim \mathcal{N}(0_d, I_d).$$

The covariance matrix $\tilde{C}_f$ of such samples is an estimation of $V_r^T \tilde{M} \tilde{M}^T V_r$. An approximation of $\tilde{M}_r$ is then obtained by, e.g., an eigenvalue decomposition of $\tilde{C}_f$. Note, that this approach approximates the reduced system coefficients of the time-discretised FOM, instead of the reduced system coefficients of the time-continuous FOM.

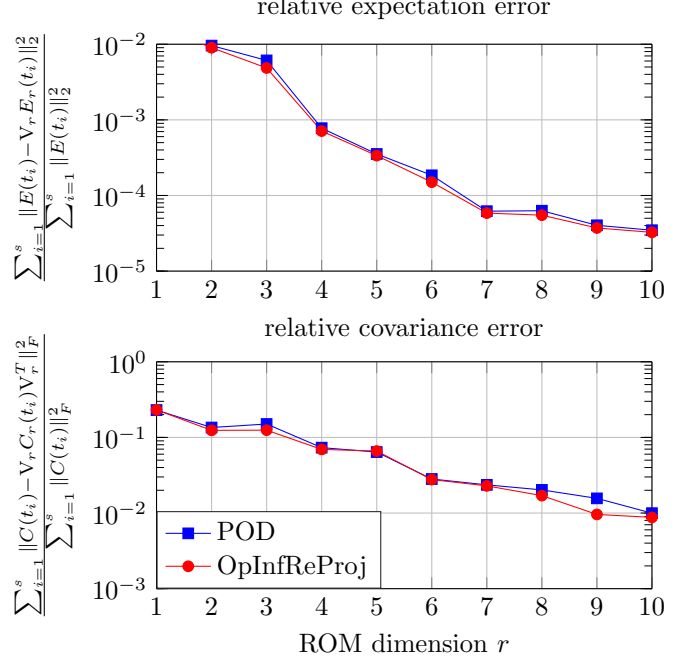


Fig. 1. Relative errors of expectation and covariance. FOM dimension $n = 1357$. Number of samples and time steps for inference and testing $L = 10^4$ and $s = 10^3$. Time-step size $h = 10^{-3}$.

4. NUMERICAL EXAMPLE

The FOM is obtained from the *Steel Profile* Benchmark from the Oberwolfach Benchmark Collection (2005). The control function models the temperature controls, which can be applied on $m = 7$ sections of the boundary of the profile. One can model a noisy control u , which is perturbed by white noise, by choosing the diffusion coefficient to be $M = \frac{B}{\|B\|}$. The step function f is given by a semi-implicit Euler-Maruyama time-discretisation of the corresponding SDE. Figure 1 reports the relative summed errors in the expectation and covariance between the FOM and the POD and the FOM and Operator Inference with re-projection ROMs. The code used to perform the experiment displayed in Figure 1 is available at https://github.com/JMNicolaus/SDE_OpInfRP

REFERENCES

- Freitag, M.A., Nicolaus, J.M., and Redmann, M. (2024). Learning stochastic reduced models from data: A non-intrusive approach. URL <https://arxiv.org/abs/2407.05724>.
- Oberwolfach Benchmark Collection (2005). Steel profile. hosted at MORwiki – Model Order Reduction Wiki. URL http://modelreduction.org/index.php/Steel_Profile.
- Peherstorfer, B. (2020). Sampling low-dimensional markovian dynamics for preasymptotically recovering reduced models from data with operator inference. *SIAM Journal on Scientific Computing*, 42(5), A3489–A3515.
- Peherstorfer, B. and Willcox, K. (2016). Data-driven operator inference for nonintrusive projection-based model reduction. *Computer Methods in Applied Mechanics and Engineering*, 306, 196–215.

Failure Rates from Data of Field Returns

Sven-Joachim Kimmerle^{*,***}, Karl Dvorsky^{***},
Hans-Dieter Liess^{*,****}

^{*} PSS GmbH, Pfaffenkammer Str. 5, 82541 Münsing, Germany
^{**} (e-mail: k.dvorsky@physsolutions.com)

^{***} TH Rosenheim, Hochschulstr. 1, 83024 Rosenheim, Germany
(e-mail: sven-joachim.kimmerle@th-rosenheim.de)

^{****} Universität der Bundeswehr München, Werner-Heisenberg-Weg
39, 85577 Neubiberg, Germany (e-mail: hдлиess@unibw.de)

Abstract: To determine failure rates is a challenge, if there are only a few failures and typical failure rates are low. As an application example we are interested in failure rates of electrical automotive components for automated/autonomous driving. As method we focus here on the exploitation of field data. Our contribution classifies different approaches from statistics and shows how this can be applied to real-world production figures as available in industry.

Keywords: modelling uncertainties and stochastic systems, automotive vehicle electrical systems, autonomous driving, reliability, failure rates, confidence intervals, field data

1. APPROACHES TO FAILURE RATES

For applications like autonomous/automated driving a high reliability of the used components is required for functional safety. Therefore long survival times with a high statistical confidence should be known for the components before the design of the vehicle electrical system and its production. So-called *FIT rates*, i.e. failure rates in the order of 1 FIT, i.e. 1 failure in 1 billion ($=10^9$) hours of lifetime seem to be acceptable for the envisaged application. Since components are under certain stresses not only when in use and ageing happens all the time, failure rates are calculated w.r.t. hours of lifetime, not hours of operation.

Moreover, it is necessary to define a failure precisely by a so-called *failure criterion*. This might depend on the type of the component, e.g. the doubling of the Ohmic resistance for a welded splice. Furthermore, we could differentiate between failure modes like “open” or “short-circuit”.

There are three important approaches to determine failure rates: (i) exploitation of data from field returns (so-called field data), (ii) standardized handbooks with failure rates, and (iii) laboratory exposure tests. In approach (i) the field data, i.e. failures and survivals from components used in the field, is evaluated statistically. The number of legitimate failures (e.g. a wrong coloring of the component might not enhance its functionality) compared with the number of components and hours of lifetime allow to estimate a failure rate (not necessarily constant). This approach is also known as REX (return of experience). The handbooks in (ii) rely on expert opinions and previous records (including statistics and partially field data). There are several international standards for failure rates, FIDES (2022) being the most recent. For approach (iii) long term exposure tests are designed, that try to trigger each a specific physical failure mechanism. In order to observe failures in reasonable time and for a sufficient number of specimens, exposure tests that can speed-up time due

to overstresses are crucial. is then analyzed statistically. Here the focus is on the approach (i) using field data to determine low failure rates, in particular in the case of zero or only a few observed failures.

2. PROBABILISTIC AND STATISTICAL MODELS

We examine a family of identical, independent, newly manufactured specimens over a given time period where it is possible to track the specimen for failures. As collective we consider the number of specimen N times the time period under consideration T (in h). We consider X as the random variable with values in $\mathbb{N}_0$, x denotes the observed realizations. A point estimate of the failure rate in our collective serving as a sample is

$$\hat{\Theta} = \frac{X}{n} = \frac{X}{NT} \quad (\text{FIT} = \# \text{ failures}/(10^9 \text{h})). \quad (1)$$

However, a point estimate lacks a statement about the statistical confidence of the result. If, let's say we observe 0 failure among a sample of length $N = 1000$ or $N = 10^9$ (for the same T) should make a difference, but in both cases we estimate $\hat{\Theta} = 0$ FIT for the failure rate.

We assume that a required confidence level ν , e.g. 90%, is prescribed (by rules or economically) for the application. If we consider a suitable confidence interval $[\underline{\Theta}, \bar{\Theta}]$ for the unknown parameter Θ for the confidence level ν , than the upper bound $\bar{\Theta}$ of the confidence interval may be considered as a conservative estimate for the failure rate λ_{total} that here incorporates all influences on the failure rate. However, the construction of a confidence interval depends on the underlying model, e.g. whether a non-parametric or parametric estimate is appropriate. We will discuss shortly the two models in the following.

2.1 Binomial Distribution as Model

We suppose that Θ is the probability that an event occurs among N components in a given time interval of length

T . Let X denote the random variable of the sum of events E within $n = NT$ trials. Accordingly the probability for observing events per time is given by (1).

We follow Stange (1970, p. 436) for the construction of the corresponding (typically non-unique) confidence interval. If a required confidence level ν or vice versa the probability of error $\alpha_1 + \alpha_2 = 1 - \nu$ is given¹, then there holds

$$\mathbb{P}(\underline{x} \leq X \leq \bar{x}) = \sum_{i=\underline{x}}^{\bar{x}} \binom{n}{i} \Theta^i (1-\Theta)^{n-i} = 1 - \alpha_1 - \alpha_2, \quad (2)$$

where $\underline{x}$ or $\bar{x}$ is the smallest / largest natural number greater / less than or equal to $n\Theta \mp q_{1/2} \sqrt{n\Theta(1-\Theta)}$, where $q_i = q_{1-\alpha_i}^{\mathcal{B}(\Theta, n)}$ is the quantile of the binomial distribution with parameters Θ and n w.r.t. $1 - \alpha_i$, $i = 1, 2$. The crucial estimate $\bar{\Theta}$ for $X = x$ is determined by

$$\sum_{i=0}^x \binom{n}{i} \bar{\Theta}^i (1 - \bar{\Theta})^{n-i} = \alpha_1. \quad (3)$$

In particular, in the case $x = 0$, that is important for the applications, (3) simplifies to $\bar{\Theta} = 1 - \sqrt[n]{\alpha_1}$, where we have naturally $\alpha_1 = \alpha$ and $\alpha_2 = 0$. Thus we have the confidence interval $[0, \bar{\Theta}]$ for zero failures.

If we have zero failures the question whether failed components are (instantaneously) replaced is obsolete. Moreover, for a few failures, i.e. $x \ll n$ the question of replacing components has no numerical influence on the results.

We see that the construction of a confidence interval as for the binomial distribution may yield technical challenges. By using Fisher's F -distribution this might be overcome. Further details yielding the so-called Pearson-Clopper values, see Stange (1970, p. 433 ff.), will be presented on site.

2.2 Maximum Likelihood

In addition, we have modelling challenges due to the *censoring* of the data. In statistics censored data means that random variables, as the survival times here might not be observed/measured over the whole time. If we cannot track each sample until a failure occurs, then this is a typical example for *right-censoring*, whereas if the starting time of the observation/measurement cannot be traced back yields a so-called *left-censoring*.

Considering a constant (random) failure rate, we illustrate here the case of the exponential distribution. For $X \geq 1$, the maximum likelihood estimate is in the uncensored case identical to (1), in the general case of censored data (no replacement of the component)

$$\hat{\Theta} = \frac{X}{\sum_{i=1}^N t_i} = \frac{X}{\sum_{i=1}^X t_i + (N - X)T}, \quad (4)$$

where t_i denotes the random survival time of component i , being T , if component i does not fail in the observed time period. W.l.o.g. the components with failures get the lowest indices. For the observed times with components in function, appearing in the denominator, we abbreviate $T_{obs} = \sum_{i=1}^X t_i + (N - X)T$. We obtain (Sundberg (2001)) the confidence interval

$$\mathbb{P}\left(\frac{q_{\alpha_1}^{\chi^2(2X)}}{T_{obs}(X)} \leq \hat{\Theta} \leq \bar{\Theta} := \frac{q_{1-\alpha_2}^{\chi^2(2X)}}{2T_{obs}(X)}\right) = 1 - \alpha_1 - \alpha_2, \quad (5)$$

¹ α_1 is the probability of an error above the sample mean plus a margin of error and α_2 is the probability of the error below.

where $q_{\tilde{\nu}}^{\chi^2(2X)}$ is the quantile of the χ^2 -distribution with $2X$ degrees of freedom w.r.t. the confidence level $\tilde{\nu}$. We see that the results for the sample mean for the binomial as for the exponential distribution coincide, whereas the relevant estimate in (5) might be different already in the uncensored case.

3. DATA OF FIELD RETURNS AND OUTLOOK

On site we present data of field returns for welded automotive splice and follow approach (i). Moreover, we include here so-called dark figures for N and X , these percentages model that not all components in the collective may be tracked and that not all failures might be reported in the real world. Finally, we will discuss phenomena due to the size of the collective and to the split of components into smaller groups.

It turns out in this example that the binomial model and the maximum likelihood method yield the same estimates for the failure rate λ_{total} . Following approach (ii) we obtain a higher value for the failure rate. However, both values were undercut by laboratory long exposure tests following approach (iii) for this component (ZVEI-BI, 2021, Section 7.5). Note that the three approaches yield different FIT rates not only here and to be on the safe side, the worst (highest) FIT rate is used as prescribed in the FIDES. The reason for this is that a chain is only as strong as its weakest link. However, it is recommendable to use several approaches as pillars for the FIT rates.

Finally, we will discuss the stated results and close with an outlook. It should also be mentioned that we consider here only constant failure rates, modelling random errors. Systematic errors are assumed to be avoided by a strong quality management. This approach has been applied by the authors together with industrial partners for several electric components as splice, power and data cables, fuses, and mass connections in automotive cable harnesses, see, e.g., ZVEI-BI (2021); Kimmerle & Liess (2019).

ACKNOWLEDGEMENTS

We acknowledge support from several industrial partners by data and within projects.

In particular, we would like to thank very much Rudolf Avenhaus at UniBw München for stimulating discussions.

REFERENCES

- IMdR - Institut pour la Maîtrise des Risques (Ed.). *FIDES Guide 2022 Edition A, Reliability Methodology for Electrical Systems* (revised version, English translation). Standardized as UTE C80-811. IMdR, Gentilly, 2023.
- S.-J. Kimmerle and H.-D. Liess. Zuverlässigkeit elektrischer Komponenten im Bordnetz: Jedes Bauteil zählt. *Elektronik automotive, Sonderheft 4*, 42–45, 2019.
- K. Stange. *Angewandte Statistik. Teil 1: Eindimensionale Probleme*. Springer, Heidelberg, 1970.
- R. Sundberg. Comparison of Confidence Procedures for Type I Censored Exponential Lifetimes. *Lifetime Data Analysis* 7(4), 393–413, 2001.
- ZVEI & Bayern Innovativ, Cluster Automotive (Eds.). *Technischer Leitfaden: Ausfallraten für Bordnetz-Komponenten - Erwartungswerte und Bedingungen (1st ed.)*. ZVEI e.V., Köln, 2021.

Data driven identification and model reduction for nonlinear dynamics^{*}

Zlatko Drmač^{*} Igor Mezić^{**}

^{*} *Department of Mathematics, Faculty of Science, University of Zagreb, Croatia (e-mail: drmac@math.hr).*

^{**} *Department of Mechanical Engineering, University of California, Santa Barbara, (e-mail: mezici@aimdyn.com)*

1. INTRODUCTION

The Dynamic Mode Decomposition (DMD) Schmid and Sesterhenn (2008), Schmid (2010) is a versatile computational tool for data driven analysis of nonlinear dynamical systems, with applications in e.g. computational fluid dynamics, aeroacoustics, robotics. It can be used for model order reduction, analysis of latent structures in the dynamics, and e.g. for data driven identification, forecasting and control. The theoretical underpinning of the DMD is in the framework of the Koopman (composition) operator theory Rowley et al. (2009), Mezić (2013).

The Koopman (composition) operator provides an infinitely dimensional linearization of nonlinear dynamical systems, and it is a tool of the trade for computational data driven analysis (identification, prediction, control) of nonlinear dynamics. For instance, consider the discrete dynamical system

$$\mathbf{x}_{i+1} = \mathbf{T}(\mathbf{x}_i), \quad (1)$$

where $\mathbf{T} : \mathcal{X} \rightarrow \mathcal{X}$ is a map on a state space $\mathcal{X} \subseteq \mathbb{R}^n$ and $i \in \mathbb{Z}$. The $\mathbf{x}_i$'s are e.g. obtained by numerical simulations of a continuous system (i.e. system of differential equations), or by measuring experimental data. Further, the mapping $\mathbf{T}$ might be unknown, but an abundance of data snapshots $\mathbf{x}_i$ is available. The Koopman operator $\mathcal{K} \equiv \mathcal{K}_{\mathbf{T}}$ for the discrete system (1) is defined on a suitable (Hilbert) space of observables $\mathcal{F}$ by

$$\mathcal{K}f = f \circ \mathbf{T}, \quad f \in \mathcal{F}. \quad (2)$$

The key observation is that for a vector valued observable $\mathbf{f} = (f_1, \dots, f_n)^T$ of interest, its value along the trajectory of (1) can be represented as $\mathbf{f}(\mathbf{x}_1) = (\mathcal{K}^0 \mathbf{f})(\mathbf{x}_1)$, $\mathbf{f}(\mathbf{x}_2) = (\mathcal{K} \mathbf{f})(\mathbf{x}_1)$, $\mathbf{f}(\mathbf{x}_3) = (\mathcal{K}^2 \mathbf{f})(\mathbf{x}_1)$, ..., $\mathbf{f}(\mathbf{x}_{m+1}) = (\mathcal{K}^m \mathbf{f})(\mathbf{x}_1)$, where the action of $\mathcal{K}$ defined component-wise. Hence, to reveal the latent structure of (1) and to develop forecasting skills, or to identify $\mathbf{T}$, it is plausible to try to identify $\mathcal{K}$ (based on the data only) and compute its approximate eigenvalues and eigenvectors (using a data driven compression of $\mathcal{K}$ and the well known procedures from numerical linear algebra, but adapted to the data driven scenario).

The available data are stored in the snapshot matrix F with columns $\mathbf{f}(\mathbf{x}_1), \mathbf{f}(\mathbf{x}_{k+1}) = (\mathcal{K} \mathbf{f})(\mathbf{x}_k)$, $\mathbf{x}_{k+1} = \mathbf{T}(\mathbf{x}_k)$:

$$F = (\mathbf{f}(\mathbf{x}_1) \dots \mathbf{f}(\mathbf{x}_m) \mathbf{f}(\mathbf{x}_{m+1})) = \begin{pmatrix} f_1(\mathbf{x}_1) & \dots & f_1(\mathbf{x}_m) & f_1(\mathbf{x}_{m+1}) \\ f_2(\mathbf{x}_1) & \dots & f_2(\mathbf{x}_m) & f_2(\mathbf{x}_{m+1}) \\ \vdots & \vdots & \vdots & \vdots \\ f_d(\mathbf{x}_1) & \dots & f_d(\mathbf{x}_m) & f_d(\mathbf{x}_{m+1}) \end{pmatrix}.$$

- (i) The snapshots are generated by a nonlinear system.
- (ii) The snapshots are a Krylov sequence $\mathbf{f}, \mathcal{K}\mathbf{f}, \mathcal{K}^2\mathbf{f}, \dots$, driven by the linear operator $\mathcal{K}$ and evaluated along a trajectory initialized at $\mathbf{x}_1$.

It makes sense to find a matrix $\mathbf{A}$ that reproduces the Krylov sequence (over available data), i.e. such that

$$\mathbf{A}\mathbf{f}(\mathbf{x}_k) = (\mathcal{K}\mathbf{f})(\mathbf{x}_k) = \begin{pmatrix} (\mathcal{K}f_1)(\mathbf{x}_k) \\ \vdots \\ (\mathcal{K}f_n)(\mathbf{x}_k) \end{pmatrix} = \mathbf{f}(\mathbf{T}(\mathbf{x}_k)), \quad k = 1, \dots, m. \quad (3)$$

The Koopman Mode Decomposition (KMD) represents the scalar observables f_i in terms of the eigenfunctions of $\mathcal{K}$, so that for an $\mathbf{x}$

$$(\mathcal{K}^k \mathbf{f})(\mathbf{x}) = \begin{pmatrix} (\mathcal{K}^k f_1)(\mathbf{x}) \\ \vdots \\ (\mathcal{K}^k f_n)(\mathbf{x}) \end{pmatrix} \approx \sum_{i=1}^m \mathbf{z}_i \phi_i(\mathbf{x}) \lambda_i^k, \quad k = 0, 1, \dots \quad (4)$$

where $(\mathcal{K}\phi_i)(\mathbf{x}) \approx \lambda_i \phi_i(\mathbf{x})$. It can be shown that $(\mathbf{z}_i, \lambda_i)$'s are approximate eigenpairs of $\mathbf{A}$ ($\mathbf{A}\mathbf{z}_i \approx \lambda_i \mathbf{z}_i$). This requires solving the eigenvalue problem for the matrix $\mathbf{A}$ defined in (3).

2. THE DMD AND THE KMD: NUMERICAL ALGORITHMS

The application of the KMD introduced in §1 is based on a supplied sequence of snapshots $\mathbf{f}_i \in \mathbb{C}^n$ of an underlying dynamics, that are driven by an inaccessible *black box* linear operator $\mathbf{A}$;

$$\mathbf{f}_{i+1} \approx \mathbf{A}\mathbf{f}_i, \quad i = 1, \dots, m, \quad m < n, \quad (5)$$

with some initial $\mathbf{f}_1$. No other information is available.

The two basic tasks are then:

- (1) Identify approximate eigenpairs $(\lambda_j, \mathbf{z}_j)$ such that $\mathbf{A}\mathbf{z}_j \approx \lambda_j \mathbf{z}_j$, $j = 1, \dots, k$; $k \leq m$, (6)
This is solved by a data driven Rayleigh–Ritz procedure introduced by Schmid (2010).
- (2) Derive a spectral spatio-temporal representation of the snapshots $\mathbf{f}_i$ (KMD):

$$\mathbf{f}_i \approx \sum_{j=1}^k \mathbf{z}_{\zeta_j} \alpha_j \lambda_{\zeta_j}^{i-1}, \quad i = 1, \dots, m, \quad (7)$$

^{*} Supported by the DARPA Small Business Innovation Research Program (SBIR) Program Office under Contract No. W31P4Q-21-C-0007 to AIMdyn, Inc. The second author is also supported by the AFOSR Award FA9550-22-1-0531 and the ONR Award N000142112384.

Algorithm 1. $(Z_k, \Lambda_k, r_k, [B_k], [Z_k^{(ex)}]) = \text{xGEDMD}(X, Y; \text{tol})$

Input: $X = (x_1, \dots, x_m), Y = (y_1, \dots, y_m) \in \mathbb{C}^{N \times m}$ that define a sequence of snapshot pairs (x_i, y_i) . Tolerance tol for the numerical rank of X .

- 1: $D = (\text{diag}(\|X(:, i)\|_2)_{i=1}^m)^\dagger$; $X_c = XD$; $Y_c = YD$.
- 2: $[U, \Sigma, V] = \text{svd}(X_c)$; *{The thin SVD: } $X_c = U\Sigma V^*$, $U \in \mathbb{C}^{n \times m}$, $\Sigma = \text{diag}(\sigma_i)_{i=1}^m$.*
- 3: Determine numerical rank k , using the threshold tol .
- 4: Set $U_k = U(:, 1:k)$, $V_k = V(:, 1:k)$, $\Sigma_k = \Sigma(1:k, 1:k)$.
- 5: $B_k = Y_c(V_k \Sigma_k^{-1})$; *{Schmid's data driven formula for } AU_k [optional output].}*
- 6: $S_k = U_k^* B_k$ *{ } $S_k = U_k^* A U_k$ is the Rayleigh quotient.}*
- 7: $[W_k, \Lambda_k] = \text{eig}(S_k)$ *{ } $\Lambda_k = \text{diag}(\lambda_i)_{i=1}^k$; } $S_k W_k(:, i) = \lambda_i W_k(:, i)$; $\|W_k(:, i)\|_2 = 1$.*
- 8: $Z_k = U_k W_k$ *{The Ritz vectors.}*
- 9: $Z_k^{(ex)} = B_k W_k$ *{The (unscaled) Exact DMD vectors [optional output].}*
- 10: $r_k(i) = \|B_k W_k(:, i) - \lambda_i Z_k(:, i)\|_2$, $i = 1, \dots, k$. *{The residuals.}*

Output: $Z_k, \Lambda_k, r_k, [B_k], [Z_k^{(ex)}]$.

for some suitable selection of the modes $\mathbf{z}_{\zeta_j}$. The coefficients are computed by using a sparsity promoting optimization Jovanović et al. (2014), or by solving a Khatri–Rao structured least squares problem Drmač et al. (2020).

2.1 An improved DMD/KMD

The original method Schmid (2010) is considerably improved in Drmač et al. (2018), and a robust software implementation is available in the LAPACK library Drmač (2024a). One of the key features of the modified DMD is that it provides computable residuals $(r_k(\zeta_j) = \|\mathbf{A}\mathbf{z}_{\zeta_j} - \lambda_{\zeta_j} \mathbf{z}_{\zeta_j}\|_2)$, that can be used to select physically meaningful eigenvalues and modes, and to guide sparse representation of the snapshot in the KMD (7).

The improved version of the DMD is summarized in Algorithm 1. In the case of physics-informed DMD, where it is known that the underlying operator is Hermitian, a Hermitian version of the DMD requires careful implementation as in Baddoo et al. (2021), Drmač (2024b).

2.2 An example

An example of DMD/KMD is illustrated in Figure 1. The data are collected by solving the two-dimensional Navier–Stokes equation for 150 discrete time steps. The grid data are reshaped into columns and arranged column-wise in the 89351×151 matrix F . The input to DMD is $X = F(:, 1 : 150)$, $Y = F(:, 2 : 151)$. Only nine modes (eigenvectors of A) are enough to represent the entire simulation with high fidelity, and to provide very good forecasting skill.

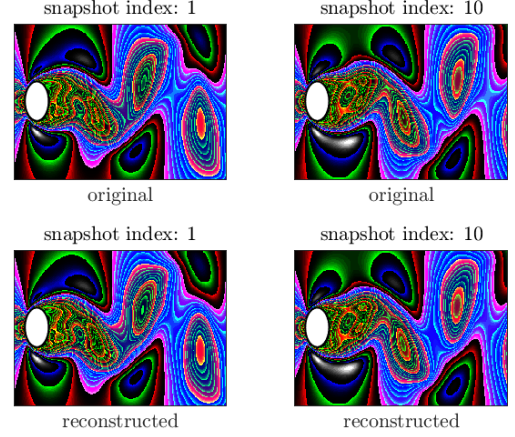


Fig. 1. Flow around a circular cylinder. The observable is the norm of the velocity. The snapshots obtained in the simulation are represented using only 9 Koopman modes. The accuracy of such representation is illustrated in the first plot for snapshots indexed as $i = 1$ and $i = 10$.

- Drmač, Z. (2024a). A LAPACK implementation of the Dynamic Mode Decomposition. *ACM Trans. Math. Softw.* doi:10.1145/3640012.
- Drmač, Z. (2024b). Hermitian Dynamic Mode Decomposition – numerical analysis and software solution. *ACM Trans. Math. Softw.* doi:10.1145/3641884.
- Drmač, Z., Mezić, I., and Mohr, R. (2018). Data driven modal decompositions: Analysis and enhancements. *SIAM Journal on Scientific Computing*, 40(4), A2253–A2285. doi:10.1137/17M1144155.
- Drmač, Z., Mezić, I., and Mohr, R. (2020). On least squares problems with certain Vandermonde–Khatri–Rao structure with applications to dmd. *SIAM Journal on Scientific Computing*, 42(5), A3250–A3284. doi: 10.1137/19M1288474.
- Jovanović, M.R., Schmid, P.J., and Nichols, J.W. (2014). Sparsity-promoting dynamic mode decomposition. *Physics of Fluids*, 26(2), 024103. doi:10.1063/1.4863670. URL <https://doi.org/10.1063/1.4863670>.
- Mezić, I. (2013). Analysis of Fluid Flows via Spectral Properties of the Koopman Operator. *Annual Review of Fluid Mechanics*, 45(1).
- Rowley, C.W., Mezić, I., Bagheri, S., Schlatter, P., and Henningson, D.S. (2009). Spectral analysis of nonlinear flows. *Journal of Fluid Mechanics*, 641, 115–127. doi: 10.1017/S0022112009992059.
- Schmid, P. and Sesterhenn, J. (2008). Dynamic Mode Decomposition of numerical and experimental data. *Bull. Amer. Phys. Soc., 61st APS meeting, San Antonio.*, 208.
- Schmid, P.J. (2010). Dynamic mode decomposition of numerical and experimental data. *Journal of Fluid Mechanics*, 656, 5–28. doi:10.1017/S0022112010001217.

REFERENCES

- Baddoo, P.J., Herrmann, B., McKeon, B.J., Kutz, J.N., and Brunton, S.L. (2021). Physics-informed dynamic mode decomposition (piDMD). *arXiv e-prints*, arXiv:2112.04307.

Confronting knowledge-based and machine learning models in describing batch fermentation.

Núria Campo-Manzanares^{*,**} Artai Rodríguez-Moimenta^{*,**}
Romain Minebois^{***} Amparo Querol^{***} Eva Balsa-Canto^{*,†}

^{*} *Bioprocess and Biosystems Engineering, IIM-CSIC, Vigo, Spain.*

^{**} *Applied Mathematics Dept. II, University of Vigo, Spain.*

^{***} *YeastOmics, IATA-CSIC, Valencia, Spain*

[†] Corresponding author: ebalsa@iim.csic.es

1. INTRODUCTION

Batch fermentation processes are widely used in industry to produce antibiotics, enzymes, biofuels, and fermented foods and beverages such as wine, yogurt, bread, and beer. The underlying concept is to introduce specific microorganisms into a medium with the nutrients necessary for cells to grow. During growth, microorganisms transform nutrients into biomass, releasing the desired products.

Optimizing fermentation processes, including species and conditions, is essential for improving yield and productivity. Knowledge-based models facilitate decision making while minimizing experiments (Lopatkin and Collins, 2020; Wang et al., 2023). Although they offer many advantages, the formulation of such models requires data, time, and insight.

In recent years, machine learning (ML) algorithms have gained significant attention due to their ability to analyze vast amounts of data and uncover patterns that traditional methods might miss. ML has the potential to optimize workflows, reduce costs, and improve product quality. In the context of fermentation, ML has been used to predict production when historical data are available (Shah et al., 2022) or as a surrogate model for scaling up (del Rio-Chanona et al., 2019). However, its effectiveness relies on the availability and quality of the data.

Fermentation knowledge-based models are often formulated using time-series data for biomass, substrate, and product dynamics, with fewer than ten sampling points. Experiments may vary temperature or pH to explain their impact. The laws of mass and energy conservation compensate for the limited data. Can ML be applied under these conditions?

This work addresses this question by considering a case study related to yeast fermentation. We first built a knowledge-based model to describe the process under temperature-varying conditions and then used the same data to formulate an ML model of the process.

Our results show that: i) building a knowledge-based model is an iterative, time-consuming process; ii) ML model design is easier, needing no specific process knowledge, but requires testing multiple architectures; iii) ML models simulation is faster; but iv) ML is not competitive for the same amount of data, offering worse predictive capabilities.

2. RESULTS

2.1 Kinetic model

We have generalized the model by Moimenta et al. (2023) to account for the effect of temperature. The model consists of a set of ordinary differential equations (ODEs) describing biomass growth phases, uptake of sugars and yeast assimilable nitrogen, and relevant products.

The model was built using a multi-experiment identification approach. Six experiments, performed at three different constant temperatures with two different levels of sugars, were used for model formulation and calibration; and two additional experiments, performed at a time-dependent temperature profile were used for validation. We considered five different sets of fermentations led by five industrial yeast species to test the generalizability of the model. Model identification was implemented in the AMIGO2 toolbox (Balsa-Canto et al., 2016).

The model successfully explained the data, with a normalized mean square error in the prediction of less than 10% for all species.

2.2 Machine learning models

We used regression models based on artificial neural networks (ANN), specifically the multilayer perceptron (MLP). We first consider the case of a particular yeast strain (*S. cerevisiae* GALA) and try to predict the production of ethanol, glycerol, acetate, and succinate. We compared three different scenarios:

- (1) **Model ML-KIn**: using the same inputs as the kinetic model, including time, initial conditions of temperature and sugar, and yeast assimilable nitrogen (YAN). The model architecture includes 4 input variables, a hidden layer with 2 neurons, and 4 output variables.

^{*} This work was funded by MCIN/AEI/10.13039/501100011033 and EU NextGenerationEU/PRTR grant PLEC2021-007827, and Xunta de Galicia (IN607B 2023/04).

- (2) **Model ML-12In**: using twelve inputs, including time, initial conditions, YAN, and the following amino acids: cysteine, glycine, histidine, methionine, aspartate, phenylalanine, isoleucine, and leucine. Note that selected amino acids present distinct dynamic profiles. The model architecture includes 12 input variables, a hidden layer with 3 neurons, and 4 output variables.
- (3) **Model ML-22In**: using 22 inputs including time, initial conditions, and all amino acids measured experimentally, except proline, ornithine. The model architecture includes 22 input variables, a hidden layer with 5 neurons, and 4 output variables.

The selected network structures demonstrated minimal overfitting after testing various combinations with the dataset. We also explored combining data from five industrial yeast *S. cerevisiae* strains, which, despite similar ethanol yields, produced varying amounts of glycerol and succinate, enriching the data. The resulting models, ML-KI-AllSp, ML-12I-AllSp, and ML-22I-AllSp, shared the architecture with those obtained with a single species dataset.

For the training of ML models, the input and output data were normalized; missing data was imputed, and outliers were removed from the dataset. The six models were trained using the mean square error (MSE) metric as the loss function, a learning rate of 0.1, a sigmoid activation function and the Stochastic Gradient Descent optimizer. The modeling workflow was implemented in Python using the Keras toolbox (Chollet et al., 2015).

2.3 Comparative analysis

Our results show that when used under the same conditions, ML offers poor performance. Only when data for multiple species were combined, the ML became more accurate (see Figure 1). **Model ML-KI-AllSp** shows an overall normalized mean square error of 18%, while **Model ML-12I-AllSp** shows a 13%, attributed to the increase in input data from 4 to 12 inputs. The addition of data in **Model ML-22I-AllSp** did not result in further improvements. Even with five times more data, the top ML model underperformed compared to the kinetic model.

3. CONCLUSION

The widespread enthusiasm for machine learning (ML) has led to its use in numerous fields. Although ML provides powerful tools for modeling, balancing this excitement with a clear understanding of its limitations and the contexts in which it can truly add value is essential.

In this work, we have confronted knowledge-based kinetic models with ML models in the prediction of yeast batch fermentation. Our results showed that the kinetic model outperformed the ML models, despite the latter being trained on a larger dataset. This is attributed to the fact that ML models rely solely on experimental data and lack prior knowledge, making them susceptible to errors and bias. This highlights the need for a hybrid approach that combines ML and knowledge-based models to exploit their individual advantages and compensate for their individual limitations (Procopio et al., 2023).

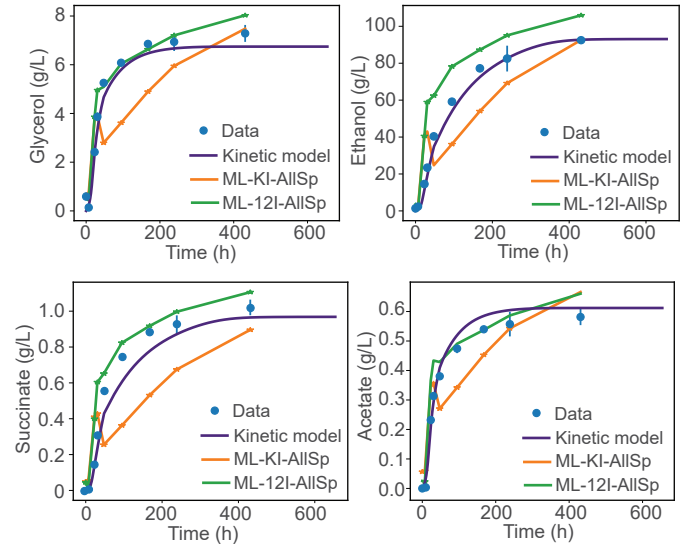


Fig. 1. Kinetic versus ML models for selected examples.

ACKNOWLEDGEMENTS

Funded by MCIN/AEI/10.13039/501100011033 and the EU NextGenerationEU/PRTR grant PLEC2021-007827 and Xunta de Galicia (IN607B 2023/04).

REFERENCES

- Balsa-Canto, E., Henriques, D., Gabor, A., and Banga, J. (2016). AMIGO2, a toolbox for dynamic modeling, optimization and control in systems biology. *Bioinformatics*, 32(21), 3357–3359.
- Chollet, F. et al. (2015). Keras. <https://keras.io>.
- del Rio-Chanona, E.A., Wagner, J.L., Ali, H., Fiorelli, F., Zhang, D., and Hellgardt, K. (2019). Deep learning-based surrogate modeling and optimization for microalgal biofuel production and photobioreactor design. *AIChE J.*, 65(3), 915–923.
- Lopatkin, A.J. and Collins, J.J. (2020). Predictive biology: modelling, understanding and harnessing microbial complexity. *Nat. Rev. Microbiol.*, 18(9), 507–520.
- Moimenta, A.R., Henriques, D., Minebois, R., Querol, A., and Balsa-Canto, E. (2023). Modelling the physiological status of yeast during wine fermentation enables the prediction of secondary metabolism. *Microb. Biotechnol.*, 16(4), 847 – 861.
- Procopio, A., Cesarelli, G., Donisi, L., Merola, A., Amato, F., and Cosentino, C. (2023). Combined mechanistic modeling and machine-learning approaches in systems biology—a systematic literature review. *Comp M Prog Biomed*, 240, 107681.
- Shah, P., Sheriff, M.Z., Bangi, M.S.F., Kravaris, C., Kwon, J.S.I., Botre, C., and Hirota, J. (2022). Deep neural network-based hybrid modeling and experimental validation for an industry-scale fermentation process: Identification of time-varying dependencies among parameters. *Chem Eng J*, 441, 135643.
- Wang, X., Mohsin, A., Sun, Y., Li, C., Zhuang, Y., and Wang, G. (2023). From Spatial-Temporal Multiscale Modeling to Application: Bridging the Valley of Death in Industrial Biotechnology. *Bioeng.*, 10(6), 744.

Mathematical Modeling of Microbial Community Dynamics

Ana Paredes^{*,***} Eva Balsa-Canto^{**} Julio R. Banga^{*}

^{*} *Computational Biology Lab, MBG-CSIC (Spanish National Research Council), Pontevedra, Spain (j.r.banga@csic.es)*

^{**} *Biosystems and Bioprocess Engineering, IIM-CSIC (Spanish National Research Council), Vigo, Spain (ebalsa@iim.csic.es)*

^{***} *Applied Mathematics Dept., Universidade Santiago de Compostela, Santiago de Compostela, Spain*

1. INTRODUCTION

Microorganisms rarely exist in isolation. Instead, they form complex networks of ecological interactions, known as microbial communities. These communities, composed of bacteria, fungi, viruses, and other microorganisms, are ubiquitous across diverse environments, from soil and water to extreme habitats such as hot springs and acidic mines. Microbial communities also thrive within and on plants and animals, including humans, where they play essential roles, particularly in the gut microbiome. Systems biology provides a quantitative framework to study these communities through mathematical models, enabling a structured and nuanced understanding of their dynamics.

2. MODEL CALIBRATION

Our research focuses on the temporal dynamics of microbial communities, for which Ordinary Differential Equations (*ODEs*) provide an appropriate modeling framework. A critical component of *ODE*-based modeling is the calibration process, also known as parameter estimation. Calibration involves identifying unknown or non-measurable parameters by adjusting the model to fit experimental data. Typically, this is an iterative process that encompasses several steps (Balsa-Canto et al., 2010; Villaverde et al., 2021).

However, this process is full of possible pitfalls and challenges due to the potential non-uniqueness, ill-conditioning and non-convexity of the estimation problem. Here, we focus on issues related to (i) lack of identifiability and (ii) convergence difficulties during the parameter estimation (under- and over-fitting).

We investigate a set of canonical models with increasing complexity that represent the most common frameworks in microbial ecology, from the most classical and simple ecological models (such as *Generalized Lotka-Volterra (GLV)* models), to more complex models accounting for nutrients dynamics (such as food web models), and coarse-grained

models incorporating different regulatory mechanisms and nutrients' dynamics.

2.1 Structural Identifiability Analysis

Structural identifiability (SI) in the context of *ODE*-based dynamic models refers to the theoretical possibility of uniquely determining parameter values from ideal model outputs. This assumes perfect, noise-free, and continuous measurements, allowing for an assessment of whether the model structure itself permits unique parameter estimation, independent of data quality or experimental conditions. This concept is crucial because if a model is not structurally identifiable, it means that there could be multiple sets of parameter values that produce the same output, making it impossible to accurately estimate those parameters. Although crucial, structural identifiability analysis (SIA) has been the focus of only a few studies (Balsa-Canto et al., 2020; Remien et al., 2021; Díaz-Seoane et al., 2023). This concept is extremely important because if a model is not structurally identifiable, it means that there could be multiple sets of parameter values that produce the same output, making it impossible to accurately estimate them.

SIA classifies unknown parameters into three groups: globally identifiable, locally identifiable, and non-identifiable. If, after performing the analysis, some parameters are classified as non-identifiable, possible solutions are a reformulation of the model, fixing the non-identifiable parameters to realistic values, or planning additional experiments (if possible). These new experiments could include new observables, experimental conditions, or initial conditions.

The general question of *SIA* for arbitrary non-linear dynamic models described by *ODEs* remains an open and unresolved matter. Nevertheless, significant progress has been made over the past two decades, leading to the development of several promising software tools (Rey Barreiro and Villaverde, 2023). In this study, we focus on three state-of-the-art tools: **Structural Identifiability**, **GenSSI2** and **SIAN**. After testing these software with the selected case studies, our analysis indicates that the most efficient and robust tool is **Structural Identifiability**, while both **SIAN** and **GenSSI2** are still reasonable options. When dealing with the most complex models, these tools encountered several difficulties.

^{*} This work was supported by grant PID2020-117271RB-C22 (BIO-DYNAMICS) funded by MCIN/AEI/10.13039/501100011033, grant PID2023-146275NB-C22 (DYNAMO-bio) funded by MICIU/AEI/10.13039/501100011033 and ERDF/EU, grant IN607B 2023/04 funded by Xunta de Galicia and grant CSIC PIE 202470E108 (LARGO).

Specifically, GenSSI2 (a fully symbolic tool) experienced a significant surge in memory consumption and computation time, and sometimes failed to guarantee the uniqueness of the solution.

2.2 Practical Identifiability Analysis

Model calibration involves finding the optimal parameter values that best align the model outputs with experimental data. This process is formulated as a non-linear optimization problem. The objective is to estimate the parameters which minimize a cost function that quantifies the discrepancy between model predictions and observed data. This optimization is conducted subject to the constraints imposed by the system of ordinary differential equations (ODEs) that define the model, as well as any additional algebraic constraints that may apply.

Even with structural identifiability, practical identifiability can be compromised by insufficient data or noise, affecting parameter uniqueness and model reliability. There are several methods to assess Practical Identifiability, including the use of the Fisher Information Matrix (FIM), Profile Likelihoods or Bayesian sampling based procedures, for example.

To perform a Practical Identifiability Analysis (PIA), here we employ the **AMIGO2** (*Advanced Model Identification using Global Optimization*) toolbox for **Matlab**, which facilitates parameter estimation using global optimization, followed by sensitivity analyses and FIM-based PIA (Balsa-Canto et al., 2016).

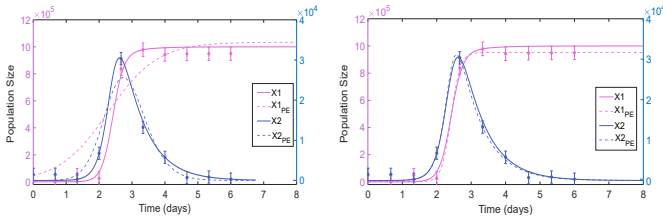


Fig. 1. Two solutions from the calibration of a *GLV* model considering two-species: local optimum (left), and global optimum (right).

Underfitting (Figure 1) occurs when the estimation algorithm converges to a local optimum. As a consequence, the calibrated model fails to capture the underlying dynamics of the data, leading to inaccurate parameter estimates. Utilizing global optimizers in **AMIGO2** allowed us to sidestep these local solutions.

Due to the flexibility and oscillatory nature of several of the models considered, we also observed that their calibration can result in overfitting, i.e. fitting the noise instead of the signal (see the example in Figure 2). In other words, the fit is very good, but the predictive power of the calibrated model is very poor. To surmount this common pitfall, at least two strategies are possible: (i) simplifying the model (sensitivity analyses can help to select the parameters to be fixed or removed); (ii) use regularization techniques (to reduce the ill-conditioning of the problem).

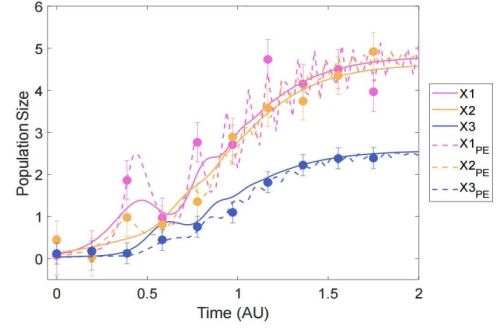


Fig. 2. Overfitting in a *GLV* model for a three species system, showing spurious oscillations in the dynamics.

3. CONCLUSION

In this study, we have addressed several key issues involved in the mathematical modeling of the dynamics of microbial communities. In particular, we considered the calibration of dynamic models composed of deterministic non-linear ordinary differential equations. First, we illustrated why Structural Identifiability Analysis (*SIA*) is a critical step in model calibration. After testing the latest available software tools, our results indicate that **Structural Identifiability** is the most robust and efficient. Second, we also illustrated two other potential pitfalls during parameter estimation, underfitting and overfitting, which can compromise the calibrated model accuracy. Addressing these challenges strengthens the predictive power of the model, facilitating more effective applications in microbial ecosystem management.

REFERENCES

- Balsa-Canto, E., Alonso, A.A., and Banga, J.R. (2010). An iterative identification procedure for dynamic modeling of biochemical networks. *BMC systems biology*, 4, 1–18.
- Balsa-Canto, E., Alonso-Del-Real, J., and Querol, A. (2020). Mixed growth curve data do not suffice to fully characterize the dynamics of mixed cultures. *Proc. Natl. Acad. Sci. U. S. A.*, 117(2), 811–813.
- Balsa-Canto, E., Henriques, D., Gábor, A., and Banga, J.R. (2016). **AMIGO2**, a toolbox for dynamic modeling, optimization and control in systems biology. *Bioinformatics*, 32(21), 3357–3359.
- Díaz-Seoane, S., Sellán, E., and Villaverde, A.F. (2023). Structural identifiability and observability of microbial community models. *Bioengineering (Basel)*, 10(4).
- Remien, C.H., Eckwright, M.J., and Ridenhour, B.J. (2021). Structural identifiability of the generalized lotka–volterra model for microbiome studies. *Royal Society Open Science*, 8(7), 201378.
- Rey Barreiro, X. and Villaverde, A.F. (2023). Benchmarking tools for a priori identifiability analysis. *Bioinformatics*, 39(2), btad065.
- Villaverde, A.F., Pathirana, D., Fröhlich, F., Hasenauer, J., and Banga, J.R. (2021). A protocol for dynamic model calibration. *Briefings in Bioinformatics*, 23(1).

Capacity Building Project in Higher Education: Leveraging Big Data and Engineering Tools to Transform Food Science Education in Indonesia

Monika Polanska*, Yoga Pratama**, Setya Budi Abduh**, Ahmad Ni'matullah Al-Baarri**, Jan F.M. Van Impe*

* *BioTeC+, Chemical & Biochemical Process Technology & Control,
KU Leuven, Belgium (e-mail: jan.vanimpe@kuleuven.be)*

** *Department of Food Technology, Diponegoro University, Indonesia*

1. INTRODUCTION

The recently launched Erasmus+ project “Enhancing Higher Education Capacity for Sustainable Data Driven Food Systems in INDonesia” – FIND4S (“*FIND force*”) addresses EU overarching priorities including Green Deal and Digital Transformation to be applied within an Indonesian landscape.

Indonesia has poor performance in sustainable agriculture. A Barilla Foundation and Economist Impact report has ranked Indonesia 71 out of 78 countries assessed (The Economist Newspaper Limited, 2024). The low sustainability score in food production is largely influenced by the increasing deforestation due to massive plantation agriculture. For instance, oil palm plantations have been the biggest driver of Indonesian deforestation (Austin *et al.*, 2019). At the same time, it is a nationwide dilemma because palm oil, on one side, is an important export commodity and contributes to the country’s economy. Other significant problems undermining the sustainability of the food system also arise from overfishing, inadequate water management and high food loss/waste in the supply chain (Nurhasan *et al.*, 2021).

FIND4S approaches the sustainability problem in the Indonesian food system by providing suitable knowledge and skills. The courses in sustainable food systems have yet to become an integral part of Indonesian Higher Education Institutions (HEI) curricula. Likewise, Indonesian graduates need to be better equipped in data science/big data processing. Data mining and big data processing prove useful in closing the gap due to isolated studies. To illustrate, the sustainability of rice production at a national level can be simulated/extrapolated using machine learning from many region-specific studies. The result can then be used to formulate strategies for country-level sustainable rice production. Whilst Indonesia is diverse, and its suitable food system might be area-specific, similarities can be found in some respects. Similarities and dissimilarities in the archipelago need to be comprehensively understood to enable appropriate actions to deal with the issues in the food system. Hence, skills in data science are essential.

2. OBJECTIVES AND METHODOLOGY

2.1 Objectives

The main FIND4S objective is to increase the capacity of seven HEIs in Java by strengthening their institutional and administrative facilities. Sharing best practices of a consortium of four European universities, producing context-specific knowledge, and delivering and disseminating outcomes will enhance curricula relevance for the local labor market and impact society at large. The capacity, knowledge and skills developed at the regionally targeted HEIs will eventually be transferred throughout the country.

This capacity-building initiative aims to transform food systems education BSc/MSc in Indonesia by integrating cutting-edge technologies such as big data, quantitative modeling, and engineering tools into the core of the educational framework. By designing new curricula and upgrading existing programs, this project seeks to equip students and academic staff with the skills needed to harness these technologies, fostering a deeper understanding of food systems and their sustainable transformation.

2.2 Consortium composition

The required expertise is readily available at the level of a long lasting cooperation among four European partners KU Leuven (Belgium), UCD University College Dublin (Ireland), UCP Universidade Católica Portuguesa (Portugal), and Anhalt University of Applied Sciences (Germany). These European partners already jointly offer a European Master of Science in Sustainable Food Systems Engineering, Technology and Business (FOOD4S “*food force*”). FOOD4S adopts a transversal and multidisciplinary approach to a broad range of topics related to the 4S pillars Science (Food Science & Engineering Technology), Sustainability (Sustainable Food Product & Food Process Design), Safety (Food Safety & Quality), and Simulation (Computational Food Science & Technology) (www.food4s.eu). Seven Indonesian universities participate in FIND4S. Diponegoro University, a top ten university in Indonesia, will be the central hub of Indonesian HEIs forming a local cluster (Fig.1). Diponegoro offers both BSc and MSc, the other members offer BSc.

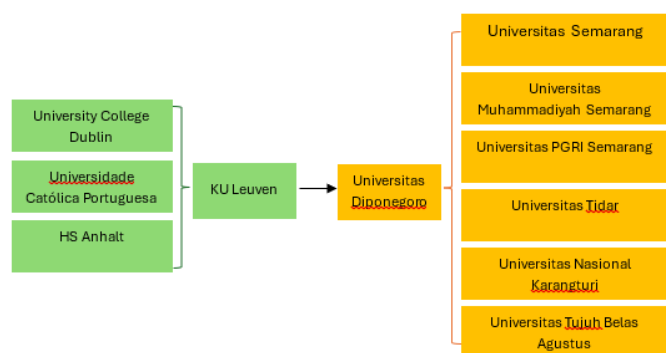


Fig. 1. FIND4S Consortium composition.

2.3 Methods and expected results

Higher Education plays a critical role in supporting the Green Deal by fostering knowledge, skills, and values that drive sustainability. The modernization of competitive and innovative curricula will promote the creation of green jobs and support the transition to sustainable food systems, with a focus on minimizing environmental impact. The project addresses significant environmental challenges, such as food safety and quality, water management, biodiversity loss, and the sustainable use of natural resources, while strengthening agri-food value chains at both national and regional levels.

The integration of risk assessment, predictive modeling, and computational optimization with sustainability principles in food production and processing is a core element of the new BSc/MSc curricula. These courses will encompass energy and food chain concepts, including Life Cycle Assessment, within a cohesive framework. By expanding the theoretical, research and policy discussions around sustainable agriculture and food production, the program aims to deepen understanding of ecological and food system dynamics. It will also explore strategies for regenerating natural systems through the use of big data and predictive tools for the food industry. These modeling tools will enable stakeholders, including industry players, to assess the impact of climate change on food safety and manage emerging threats.

At the heart of the project is the utilization of big data analytics, which will empower both educators and students to collect, analyze, and interpret vast amounts of data relevant to food science. By implementing quantitative modeling techniques, students will learn to predict and optimize processes in food production, distribution, and consumption, helping them solve complex problems faced by the food industry in real-time. Engineering tools will be incorporated into lab work and research activities, enabling the design and testing of innovative solutions to food system challenges.

The project also emphasizes the training of academic staff in these advanced technologies, ensuring they can effectively integrate them into their teaching methodologies. This will be further supported by establishing a dedicated research center and upgrading laboratory facilities to include the latest technological tools for data analysis and engineering simulations. Such infrastructure will allow students and

researchers to engage in hands-on learning, preparing them to apply these skills in real-world scenarios.

In collaboration with the aforementioned European HEIs, the initiative will foster an exchange of expertise, allowing Indonesian institutions to benefit from best practices in data-driven research and food system innovation. By building this international network, the project will ensure that Indonesian higher education stays at the forefront of global developments in food science.

A central component of the initiative is the development of a comprehensive MSc program at the central hub Diponegoro that embeds big data, quantitative modeling, and engineering tools throughout its curriculum. This advanced program will meet the growing demand for professionals equipped with modern technological and analytical skills, addressing critical issues in food security, sustainability, and innovation. Graduates will not only be able to analyze complex data sets but will also contribute to the design and implementation of sustainable food systems that are socially, economically, and environmentally responsible.

The strategic application of big data, computational methods and engineering tools will engage a broad range of stakeholders including industry partners, to ensure the program aligns with current and future market needs. These partnerships will enable the practical application of academic research, translating classroom knowledge into real-world solutions that promote a greener, more sustainable economy.

3. CONCLUSIONS

This capacity building project will serve as a transformative force in Indonesian higher education, equipping students and faculty with the tools needed to drive meaningful change. By integrating big data, quantitative modeling, and engineering tools, the initiative will support Indonesia's transition to sustainable food systems, fostering innovation and resilience in the country's food economy.

REFERENCES

- Austin, K.G., Schwantes, A., Gu, Y. and Kasibhatla, P.S. (2019) What causes deforestation in Indonesia? *Environ. Res. Lett.* 14, 024007. <https://doi.org/10.1088/1748-9326/aaf6db>
- Nurhasan, M., Samsudin, Y.B., McCarthy, J.F. (2021) Linking food, nutrition and the environment in Indonesia. *CIFOR* <https://doi.org/10.17528/cifor/008070>
- The Economist Newspaper Limited 2024. <https://impact.economist.com/projects/foodsustainability/interactive-world-map/>

ACKNOWLEDGEMENTS

This work was supported by the European Union within the framework of the Erasmus+ FOOD4S Programme (Erasmus Mundus Joint Master Degree in Food Systems Engineering, Technology and Business 619864-EPP-1-2020-1-BE-EPPKA1-JMD-MOB) and the Erasmus+ FIND4S Programme (Enhancing Higher Education Capacity for Sustainable Data Driven Food Systems in Indonesia Project 101179822 — ERASMUS-EDU-2024-CBHE). Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or the European Education and Culture Executive Agency (EACEA). Neither the European Union nor EACEA can be held responsible for them.

Sensitivity of inducible gene expression

Satyajeet Bhonsale*, Yadira Boada**, Alejandro Vignoni**,
Jesus Pico**, Jan Van Impe*

*BioTeC+, Chemical Engineering Department, KU Leuven, Technology Campus Ghent, Belgium

** SB2CL, Instituto de Automática e Informática Industrial, Universitat Politècnica de Valencia, Valencia, Spain

1. INTRODUCTION

Synthetic Biology is the engineering study of biology to enable (re)construction of cells to influence and control cellular behaviour. To avoid laborious trial and error experiments, synthetic biologists rely on mathematical models to design gene circuits. This belongs to the Design phase of the Design-Build-Test-Learn (DBTL) cycle. The use of modular modelling approaches allows synthetic biologists to compile various components in a variety of combinations and study cellular behaviour in silico. However, such use of mathematical models requires the bioparts to be appropriately characterized, i.e., the parameters which represent a particular biopart in the model must be uniquely and accurately identified.

Given the stochasticity of gene expression and measurement errors, these parameters are highly uncertain. Thus, in the design stage, it is imperative to understand the influence of this uncertainty. Moreover, it is important to know uncertainty in which of the bioparts has the largest influence.

In this paper, we demonstrate the utility of Sobol Global Sensitivity Analysis (GSA) in understanding how bioparts (i.e., model parameters) affect the engineered gene expression. As a case study, a synthetic gene circuit in which the expression of green fluorescence protein (GFP) by a transcription factor is considered. This circuit is illustrated in Figure 1. In the next sections, the gene circuit and the model used is described. The approach used to perform the GSA is then described briefly. Finally, the results of the sensitivity analysis are presented.

2. CASE STUDY

The gene circuit described here produces two proteins: LuxR (x_3) and GFP (x_8). LuxR is a constitutive protein produced in the cell (*Escherichia coli*). GFP on the other hand is an induced protein. The expression of GFP is triggered by addition of N-Acyl homoserine lactone (AHL) in the liquid medium (x_{10}). AHL diffuses through the cell membrane (x_9) and forms an AHL-LuxR complex is the transcription factor necessary to start GFP expression. Based on all the biochemical reactions

involved in the above circuit, a detailed mathematical model is obtained. Then a reduced order model is obtained using quasi-steady state assumptions (Pushkareva et al. 2023). For the cellular growth, the standard Baranyi-Roberts growth model is used. The model however contains a variety of parameters. Some of these parameters are listed in Table 1.

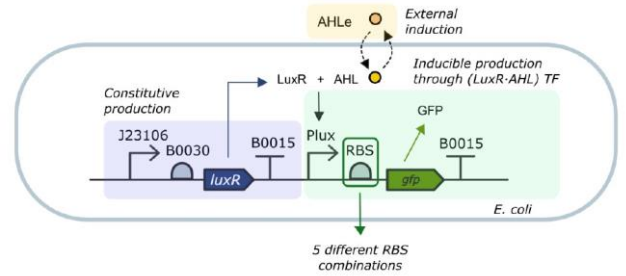


Figure 1. Gene circuit for expression of LuxR and GFP

$$\begin{aligned} \dot{x}_3 &= C_N \cdot \frac{k_1 \cdot k_2}{d_1 + \mu} - (d_2 + \mu) \cdot x_3 \\ \dot{x}_8 &= C_N \cdot \frac{k_3 \cdot k_4}{d_3 + \mu} \cdot (\alpha + (1 - \alpha) \cdot \frac{(x_9)^n}{(\frac{k_{dLux} \cdot k_1 \cdot C_N}{x_3})^n + (x_9)^n}) - (d_4 + \mu) \cdot x_8 \\ \dot{x}_9 &= D \cdot \frac{V_{cell}}{V_{ext}} \cdot x_{10} - (D + d_A + \mu) \cdot x_9 \\ \dot{x}_{10} &= -D \cdot x_{11} \cdot \frac{V_{cell}}{V_{ext}} \cdot x_{10} + D \cdot x_9 \end{aligned}$$

Figure 2. Reduced model for expression system

3. GLOBAL SENSITIVITY ANALYSIS

The Sobol indices based GSA is used in this study. These indices are determined via a variance decomposition where in the total variance in the model response can be decomposed into variations due to individual parameters, and their higher order interactions as

$$V[y] = \sum_i^n V_i + \sum_{i < j}^n V_{ij} + \dots$$

Where V_i is the variance in the output due to parameter i , V_{ij} is the variance due to parameters i and j . With this

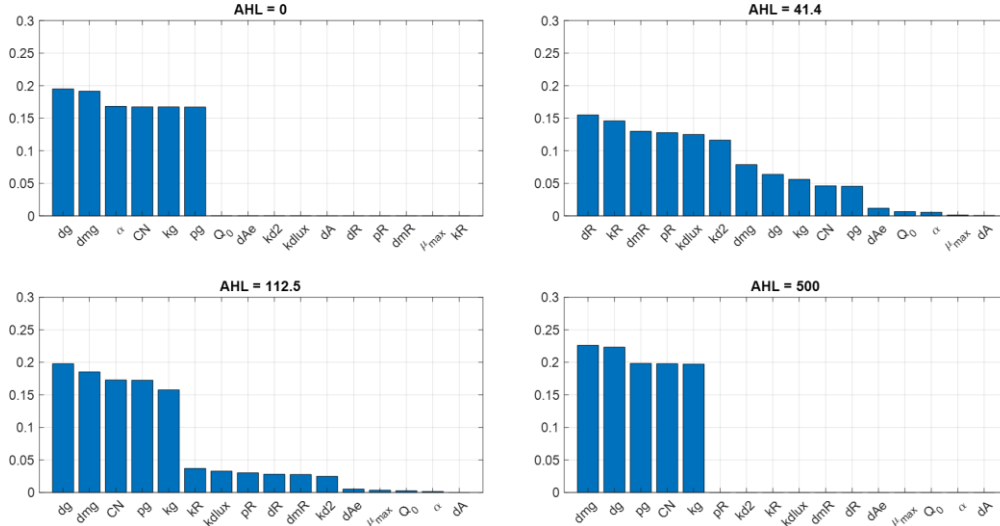


Figure 3. Sensitivity of GFP concentration at 48 h to model parameters

decomposition, the first order sensitivity index is defined as

$$S_i = \frac{V_i}{V[y]}$$

This index accounts for the variance in the model response *only* due to parameter i . S_i answers to the question “which parameter should be fixed first to reduce the variance of the output?”. In other words, a better estimation of this parameter (low parametric uncertainty) will reduce the output uncertainty. Higher order sensitivity indices can also be defined using the decomposed variance. In this study, the total and decomposed variances are computed using the polynomial chaos expansion approach (Bhonsale et al. 2019).

Table 1. Model Parameters

Parameter	Description	Parameter	Description
dg	GFP degradation	dmG	mGFP degradation
α	Basal GFP production	CN	Copy Number
kg	mGFP transcription	pg	GFP translation
dR	LuxR degradation	kR	mLuxR transcription
dmR	mLuxR degradation	pR	LuxR translation
KdLux	Half-life	μ	Growth rate

4. RESULTS

Figure 2 depicts the sensitivity of GFP concentration to all the parameters at end of 48 h. The sensitivity is reported for cases with induction by different 4 AHL concentrations. It can be observed that at AHL induction concentrations of 0 nM and 500 nM, the GFP concentration is sensitivity to parameters related to GFP production, but insensitive to LuxR production. At AHL inductions of 41.4 nM and 112.5 nM, the GFP concentration is sensitive to parameters related to both LuxR and GFP production. In all cases, the GFP concentration is insensitive to the growth rate parameters.

These results highlight the nature of the model. The relationship between AHL induction and synthesis rate is captured by a Hill function. At no induction or very small induction (i.e. the Hill function is close to 0), the basal production rate (α) is important and thus parameters related to GFP production (translation, transcription, base production rate of GFP as well as degradation rates of GFP and GFP associated mRNA) are highlighted as the important. When the AHL concentration is high, the Hill function saturates, and the maximum synthesis rate is achieved. Here, the influential parameters are related to GFP production but not the basal production rate. For intermediate concentrations, the parameters involved in LuxR production (i.e., transcription and translation rates for LuxR as well as degradation rates for LuxR and associated mRNA) become important.

4. CONCLUSION

The sensitivity indices show certain bioparts become important under certain environmental conditions (in this case AHL). Stochasticity in insensitive parts will not affect the final output. For example, if the circuit must produce GFP under high AHL concentrations, LuxR parts don't quantitatively influence the production (they are still necessary). Such information can be used to design robust gene circuits.

REFERENCES

- Bhonsale S., Munoz Lopez C. A., Van Impe J. (2019). Global Sensitivity Analysis of a Spray Drying Process. *Processes*, 7, 562.
- Pushkareva A., Beltran J., Diaz-Iza H., Arboleda-Garcia A., Boada Y., Vignoni A., Picó J. (2023). Towards automation of the Design-Build-Test-Learn (DBTL) bioengineering cycle: Application to the testing and characterization of standard bioparts. XLIV Jornadas de Automática 2023
- Saltelli A., Ratto M., Andres T., Campolongo F., Cariboni, J., Gatelli D., Saisana M., Tarantola S. (2008). *Global Sensitivity Analysis: The Primer*; John Wiley & Sons Ltd

Model order reduction for artificial neural networks generated from data driven state space models

Wil Schilders *

** Eindhoven University of Technology, Eindhoven, The Netherlands
(e-mail: w.h.a.schilders@tue.nl) & TU München – Institute for
Advanced Study, Munich, Germany*

1. INTRODUCTION

Model Order Reduction (MOR) for Artificial Neural Networks (ANNs) is an increasingly important field that aims to reduce the complexity and computational cost of ANNs while maintaining their predictive accuracy. This is particularly relevant for scientific machine learning, where neural networks are used in complex, high-dimensional tasks such as solving partial differential equations (PDEs), modelling physical processes, or real-time simulation and control in engineering systems. In this contribution, we provide an overview of the state of the art for MOR applied to ANNs, as well as some ideas we are pursuing for our novel ANNs generated from data-informed state space systems.

2. TECHNIQUES FOR MODEL ORDER REDUCTION IN ANNS

Various techniques have been developed to apply MOR in ANNs, which can be categorized into parameter reduction, layer-wise reduction, and structural simplifications.

2.1 Pruning Methods

Pruning removes unnecessary weights, neurons, or layers from neural networks while maintaining accuracy. Key methods include: (i) Magnitude-based pruning, which sets small-magnitude weights to zero; (ii) L1/L2 regularization, promoting sparsity by penalizing weight size; (iii) Structured pruning, targeting entire neurons, channels, or layers for more efficient models; and (iv) the Lottery Ticket Hypothesis, identifying small subnetworks that achieve full model performance when trained separately.

2.2 Low-rank Factorization

Low-rank factorization reduces large weight matrices into products of smaller ones, cutting parameters. Techniques include singular value decomposition (SVD) and advanced tensor factorization methods like Tucker decomposition and tensor train. SVD approximates weight matrices with low-rank representations, while tensor methods decompose higher-order tensors in convolutional layers. These approaches, applied during or after training, are effective for compressing convolutional layers in deep convolutional neural networks (CNNs) for image processing.

2.3 Quantization

Quantization reduces the precision of weights and activations, lowering memory use and computation complexity. Techniques include post-training quantization, where trained weights are mapped to lower precision, like 8-bit integers, without retraining, and quantization-aware training, where the network is trained with quantization constraints to maintain performance. It is commonly used for deploying ANNs on resource-limited devices like mobile phones and edge devices.

2.4 Knowledge Distillation

Knowledge distillation involves training a smaller network (student) to mimic the behavior of a larger network (teacher). The smaller network is trained to replicate the outputs (or feature maps) of the larger network, allowing for substantial reductions in model size while preserving accuracy. This technique is especially popular in reducing the size of very large models (such as BERT or GPT) for practical deployment.

2.5 Neural Network Compression via SVD and PCA

Principal Component Analysis (PCA) and Singular Value Decomposition (SVD) can be applied to the weight matrices of ANNs to reduce their dimensionality. These techniques work by identifying directions of variance in the data (or features) and projecting weights onto a lower-dimensional space, effectively reducing the number of parameters and improving computational efficiency.

2.6 Approximation via Surrogate Modeling

In scientific computing and physical simulations, reduced-order models serve as surrogates for neural networks. They approximate the network's behavior, especially in larger systems like physical process simulations or control applications. Surrogates, such as Polynomial Chaos Expansions or Gaussian Processes, are also used with ANNs to simplify input-output mappings, particularly for real-time settings.

3. MODEL ORDER REDUCTION IN THE CONTEXT OF PHYSICS-INFORMED NEURAL NETWORKS (PINNS)

PINNs represent a growing area where MOR is critically needed. PINNs embed physical laws (governed by PDEs)

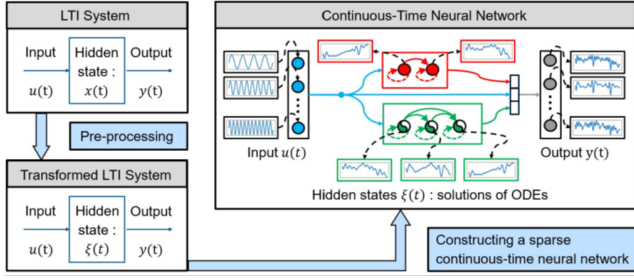


Fig. 1. Overview of our proposed workflow illustrating the systematic construction of continuous-time neural networks from Linear Time-Invariant (LTI) systems.

directly into the neural network loss function, making them suitable for solving complex physical problems like fluid dynamics, structural analysis, and electromagnetism. Key MOR approaches for PINNs include reduced basis methods and Galerkin projection. The former case involves identifying a low-dimensional subspace of the solution space, where solutions to the governing equations can be projected, thus reducing the computational load while maintaining accuracy. In the latter case, an MOR technique is used to project the dynamics of high-dimensional systems into a low-dimensional space by utilizing trial functions that satisfy the governing equations.

4. ANNS CONSTRUCTED FROM DATA-INFORMED STATE SPACE SYSTEMS

In Datar et al. (2024), we developed a systematic approach of constructing continuous-time ANNs for linear dynamical systems, based on original ideas in Meijer (1996). The idea is to first create a state-space system based on the available data, in our case using the so-called MOESP algorithm Verhaegen et al. (1992). Using a sequence of numerical methods, including QR decomposition and the Bartels-Stewart algorithm, the state-space system is transformed into an artificial neural network. The procedure is graphically illustrated in Figure 1. Special about the methodology is that horizontal layers are being formed instead of vertical layers, and that the networks are truly dynamic, i.e. non-recurrent: in the neurons, a first or second order scalar ODE needs to be solved.

This 1-1 relationship between state-space models and artificial neural networks will enable us to translate MOR methods for state-space models into MOR methods for artificial neural networks:

- First we transform the original state-space model into an equivalent ANN
- Next, we apply an arbitrary MOR method to the state-space model
- Then we translate the resulting lower-dimensional state-space model into a smaller ANN
- We then analyse how the smaller ANN can be obtained from the larger ANN, and how to formulate the corresponding MOR method for artificial neural networks.

In the talk, examples will be given of this procedure. It should be noted that the method described in Datar et al. (2024) in principle advocates the use of ANN with neuron activation functions that are special for the underlying

problem. In the case described in Meijer (1996), the first and second order ODEs are, in fact, similar to so-called low and high pass filters in electronics, and all applications were also in electronics. Recently, we have also worked on n-body dynamics to predict trajectories and masses of planets, and in that case we use neurons where the 2-body system is solved (Kepler system). This approach is actually also advocated in Ferrari et al. (2014), albeit without mentioning the use of artificial neural networks.

5. CHALLENGES AND FUTURE DIRECTIONS

Despite the progress in MOR for ANNs, several challenges remain:

- **Balancing accuracy and reduction:** Maintaining the accuracy of ANNs while reducing their order is a central challenge. Often, aggressive reductions can lead to significant degradation in performance.
- **Automated MOR techniques:** There is a need for automated techniques that can determine the optimal level of reduction for a given task, without requiring extensive hyperparameter tuning.
- **Generalization of reduced models:** Reduced models often perform well on training data but may generalize poorly to unseen data. Ensuring robust generalization is critical, particularly in safety-critical applications.
- **Integration with scientific machine learning:** As scientific ML grows, integrating MOR techniques seamlessly into hybrid methods (e.g., physics-informed ML, data-driven models) will be essential.

6. CONCLUSION

Model Order Reduction is a rapidly advancing field with significant applications for reducing the computational cost and complexity of ANNs. Techniques such as pruning, quantization, low-rank factorization, and knowledge distillation have made it possible to deploy smaller and more efficient neural networks in real-time and resource-constrained environments. These developments are especially crucial for applications in scientific machine learning, where high-fidelity simulations and real-time control require efficient approximations of large models.

REFERENCES

- Datar, C., Datar, A., Dietrich, F., and Schilders, W (2024). Systematic construction of continuous-time neural networks for linear dynamical systems. submitted to *SIAM Journal of Scientific Computing*.
- Meijer, P.B.L. (1996). Neural Network Applications in Device and Subcircuit Modelling for Circuit Simulation. PhD thesis, TU Eindhoven.
- Verhaegen, M., and Dewilde, P. (1992). Subspace model identification, part 1: The output error state space model identification class of algorithms. *International Journal of Control*, vol. 56, pp. 1187-1210.
- Goncalves Ferrari, G., Boekholt, T., and Zwart, S.F.P. (2014). A Keplerian-based Hamiltonian Splitting for Gravitational N-body Simulations. *Mon. Not. R. Astron. Soc.*, 000, pp. 1-13.