

A Hybrid Data-Driven Approach for the Optimization of an Industrial Alkylation Unit

Rastislav Fáber^{a*}, Karol Ľubušký^b, and Radoslav Paulen^a

^a Faculty of Chemical and Food Technology, Slovak University of Technology in Bratislava, 812 37 Bratislava, Slovakia

^b Slovnaft, a.s., 824 12 Bratislava, Slovakia

* Corresponding Author: rastislav.faber@stuba.sk

ABSTRACT

We develop a multi-fidelity soft-sensing framework to reconcile online (low-fidelity) industrial measurements with sparse (high-fidelity) laboratory samples from an alkylation unit in a refinery. A first-principles model is used to generate an additional low-fidelity dataset and train a surrogate that predicts the output variable. We investigate whether incorporating sparse high-fidelity laboratory data with the low-fidelity data improves prediction accuracy. A multi-fidelity predictor forms a corrected output by learning the residual between the high-fidelity observations and the low-fidelity surrogate using Gaussian process regression. The simplest model structure performs best, reducing test prediction error (computed against laboratory samples) by 31.1% relative to the currently deployed industrial analyzer and outperforming a standard high-fidelity-only model trained on laboratory data. Overall, the simplified surrogate model captures the main industrial trends well enough to serve as a reliable low-fidelity input for the multi-fidelity soft sensor.

Keywords: Process Simulation, Deep Learning, Alkylation, Process Optimization, Energy Efficiency, Data-Driven Modeling

INTRODUCTION

Industrial processes operate under tight economic and safety constraints, and their units are expected to run close to optimal conditions for long periods [1]. Deviations from the desired operating point reduce yield, increase utility consumption, and can accelerate equipment wear [2]. At the same time, plants face feed variability, seasonal effects, catalyst aging, and operational interventions that change the process behavior over time [3]. This combination drives operators to rely on monitoring and decision support tools [4] that are reliable, and easy to maintain.

Modern plants are heavily instrumented. Flow, pressure, and temperature measurements are usually available at high sampling rates and are directly usable for monitoring and control. Other variables readings can be harder to obtain. Product quality indicators and stream compositions are often measured by online sensors or by laboratory analyses [5]. Online analyzers provide frequent updates but may drift, show bias, or suffer from maintenance related faults, whereas lab measurements are less frequent but are typically treated as a reference

signal. This results in a practical mismatch between what is needed for operation, and what is available as reliable ground truth [6]. Since quality measurements are not collected continuously, plants often operate with wider safety margins, which typically reduces profitability [7].

Data-driven models are a common way to bridge the gap between high-rate process measurements and sparse, reliable quality data, for example through soft sensors that can infer difficult-to-measure outputs from easy-to-measure inputs. In industrial operation, the implemented soft sensor must handle missing values, sensor faults, and operating mode changes. They also need to be interpretable to diagnose failures. Moreover, many industrial soft sensors are deliberately simple and use a small set of engineered regressors that aim to approximate the dominant material and energy balances [8, 9]. Even the best-performing offline model is not useful if it depends on variables that are unreliable, unavailable in real time, or unstable after maintenance [10]. Some recent deep learning soft sensors aim to improve prediction by extracting relevant hidden features while eliminating redundancy and maintaining structure simplicity [11, 12].

First-principles process models can also address

some of these issues because they encode mass and energy conservation laws, phase equilibria, and known reaction pathways. They can provide consistent data under systematic perturbations. In large-scale industrial settings, however, fully accurate mechanistic models are difficult to obtain. Parameter uncertainty, unknown disturbances, unmeasured states, and changing plant conditions often limit predictive accuracy. This makes a direct mechanistic model unrealistic when the goal is a reliable, daily decision-making support [13]. Broader reviews report that implementations of many first-principles models remain constrained by uncertainty and integration or maintenance burden [14].

These constraints motivate hybrid workflows that combine mechanistic structure with data-driven learning. In such workflows, mechanistic models are used to provide physically consistent trends and to expand coverage of the operating space, while plant data compensate unmeasured disturbances [15]. Surrogate model frameworks (e.g., deep neural networks trained on simulation data) are commonly used to retain process trends while reducing computational cost for downstream tasks [16]. In this work, a deliberately simplified first-principles simulator is used to generate an additional fidelity level that supplies trend information beyond routine operation and supports surrogate training. Each of these data sources can be treated as a separate information stream within an MF framework. High-frequency measurements act as a low-fidelity (LF) data that capture dynamics and disturbances. Sparse laboratory samples provide high-fidelity (HF) reference to correct bias. The objective is to explore whether a simple surrogate trained on simulated inputs can explain the measured industrial output.

In this work, we build a workflow that (i) uses a mechanistic process model to generate a large synthetic dataset under controlled variations of key control parameters, (ii) trains a surrogate model that reproduces the simulator trends at low computational cost, and (iii) integrates the predictions with industrial measurements in an MF framework. The main focus is to improve monitoring and a deployable structure with limited number of input variables, and a predictor that can be retrained when an operation point changes.

PROBLEM STATEMENT

Many industrial regression problems involve multiple information sources with different fidelity. Let $x \in \mathbb{R}^p$ denote the regressor vector (measured variables) and $y \in \mathbb{R}$ the target output variable. An industrial LF dataset provides T online samples $\{(x_t^{\text{LF}}, y_t^{\text{LF}})\}_{t=1}^T$, while a HF reference provides only K trusted laboratory samples $\{(x_k^{\text{HF}}, y_k^{\text{HF}})\}_{k=1}^K$, with $K \ll T$. In addition, an auxiliary LF' dataset may be available from an alternative source, or model

simulation $\{(x_n^{\text{LF}'}, y_n^{\text{LF}'})\}_{n=1}^N$, with $K \ll T \ll N$. The objective is to use the sparse HF reference to correct predictions from the frequent LF data within an MF framework.

We consider three information sources. Simulation dataset (LF' hereafter referenced as X^{SIM}) generated from variations of key control variables within defined limits. Here, p denotes the number of monitored variables:

$$X^{\text{SIM}} = \{(x_n^{\text{SIM}}, y_n^{\text{SIM}})\}_{n=1}^N, x_n^{\text{SIM}} \in \mathbb{R}^p, y_n^{\text{SIM}} \in \mathbb{R}. \quad (1)$$

Industrial LF variables $X^{\text{IND}} = \{(x_t^{\text{LF}}, y_t^{\text{LF}})\}_{t=1}^T$ consist of a time series of online measurements with regressors x_t^{IND} , and the output approximation y^{LF} (e.g., online analyzer) that may be biased or faulty available in real time. Industrial HF samples $X^{\text{LAB}} = \{(x_k^{\text{HF}}, y_k^{\text{HF}})\}_{k=1}^K$ are available at sparse time instants, where y^{HF} is obtained from laboratory analyses.

Let θ^{LF} denote the parameters of the surrogate $f^{\text{SIM}}(\cdot; \theta^{\text{LF}})$ learned from the simulation dataset X^{SIM} and then evaluated on the industrial data X^{IND} :

$$\hat{y}_t^{\text{LF}} = f^{\text{SIM}}(x_t^{\text{LF}}; \theta^{\text{LF}}). \quad (4)$$

Two model classes, linear and nonlinear, are considered. The role of this step is to provide a dense, physics-consistent prediction over all industrial time points $t = 1, \dots, T$.

We model the HF-to-LF discrepancy at HF time points k as:

$$\Delta_k = y_k^{\text{HF}} - \hat{y}_k^{\text{LF}}. \quad (5)$$

We train a correction model $g^{\text{IND}}(\cdot)$ on HF residuals, and predict all differences as:

$$\hat{\Delta}_t = g^{\text{IND}}(z_t^{\text{LF}}; \phi_t^{\text{LF}}), \quad (6)$$

where z_t^{LF} is defined as a representation of inputs with added information from a best performing surrogate model $f^{\text{SIM}}(\cdot)$:

$$z_t^{\text{LF}} = [x_t^{\text{LF}^\top}, \hat{y}_t^{\text{LF}^\top}]^\top. \quad (7)$$

The final MF predictor is then defined as:

$$\hat{y}_t^{\text{MF}} = \hat{y}_t^{\text{LF}} + \hat{\Delta}_t. \quad (8)$$

If we assume a perfect prediction that satisfies $\hat{\Delta}_t \approx \Delta_t$, then, based on (5) and (8), we can state:

$$\hat{y}_t^{\text{MF}} = \hat{y}_t^{\text{LF}} + \hat{\Delta}_t \approx \hat{y}_t^{\text{LF}} + (y_t^{\text{HF}} - \hat{y}_t^{\text{LF}}) = y_t^{\text{HF}}, \quad (9)$$

where y_t^{HF} denotes the theoretical HF data. Therefore, the MF prediction can be interpreted as a real-time estimate of the HF output $\hat{y}_t^{\text{MF}} \approx y_t^{\text{HF}}$.

METHODOLOGY

Following our previous work [17], all models use the

same feature-engineered regressor vector $x \in \mathbb{R}^p$, consisting only of variables that are both statistically relevant and operationally trusted.

Rows containing unusable (i.e. NaN, or out-of-range outliers) values in any regressor or in the corresponding output are removed prior to model training and evaluation. A standardization (zero mean, unit variance) is applied to scale and anonymise the data.

Principal Component Analysis (PCA) is a linear dimensionality-reduction method applied to z_t^{LF} as a pre-processing step to obtain a smaller set of principal components (PCs). By keeping only the most relevant q components ($q \ll p$). PCA reduces noise and collinearity and yields a compact representation that is easier to model.

We train the surrogate model $f^{SIM}(\cdot)$ model using two separate approaches:

1. **Linear Ordinary Least Squares (OLS) regression** is trained and validated on the SIM inputs x_n^{SIM} and output y_n^{SIM} .
2. **Nonlinear Neural network (NN)** uses a feed-forward network trained on the same SIM pairs (x_n^{SIM}, y_n^{SIM}) to provide an alternative $\hat{y}_t^{LF}$ prediction that, unlike linear models, can capture complex, non-linear relationships and feature interactions within the dataset.

Gaussian Process regression (GPR) is a nonlinear regression method that models an unknown function as a distribution over functions. The model is fully specified by mean $m(x)$, and a covariance function $k(x, x')$ (kernel).

$$g(x) = \mathcal{GP}(m(x), k(x, x')) \quad (10)$$

The kernel defines how strongly two inputs are correlated based on their distance, and tuned hyperparameters (e.g., length scale ℓ and signal variance δ^2). A commonly used RBF kernel can be defined as:

$$k(x, x') = \delta^2 \exp(-||x - x'||^2 / 2\ell^2). \quad (11)$$

In this work, the GP is used to learn the correction term Δ_k .

Performance is evaluated only at time indices $t(k)$, where HF measurements are available, using Root Mean Squared Error (RMSE) regression metric:

$$\text{RMSE} = \sqrt{\frac{1}{K} \sum_{vk} (y_k^{\text{HF}} - \hat{y}_{t(k)})^2}. \quad (12)$$

CASE STUDY

We study an industrial alkylation unit of our industrial partner, Slovnaft, a.s., (Bratislava, Slovakia), that produces high-octane alkylate by reacting C_3 - C_4 olefins with i - C_4 isobutane in an H_2SO_4 acid-catalyzed process. Fig. 1 shows the simplified layout, where the reactant mixture is cooled and fed to the reactor section, and the effluent is separated to recover and recycle isobutane while the

alkylate product is withdrawn.

We focus on the isobutane-to-olefin ratio control as it strongly influences product quality and plant economy. Excess of isobutane is needed to suppress undesired side reactions, however, a key online analyzer (y_t^{LF}) on a recycle stream provides unreliable measurements of i - C_4 affecting the amount of recycled material. Simply increasing the recycle flows would put unnecessary load on the downstream separation, deteriorating plant economics. A better monitoring strategy is needed.

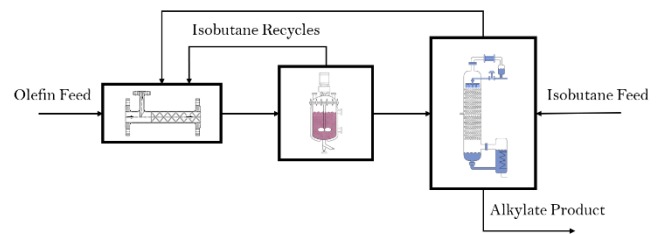


Figure 1. Overview of the simplified industrial process.

IMPLEMENTATION

A simplified plant flow-sheet model is built in the gPROMS ProcessBuilder [18]. A simulation dataset X^{SIM} is then generated by systematically varying operating parameters in the range of plant operating conditions to obtain the simulation dataset $x_n^{SIM} \in \mathbb{R}^p$ ($p = 6$) and y_n^{SIM} . Although more steady states could be generated, we limited the simulation set to $N = 5,000$ because the simplified first-principles model is used only to supply

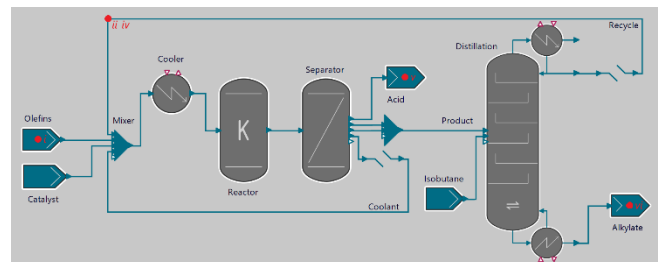


Figure 2. Overview of the simplified alkylation unit model developed in the gPROMS environment.

mass-balance-consistent trends for training, while the industrial dataset already provides dense coverage in time ($T > N$). The inputs (red dots in Fig. 2) were iteratively selected in previous work as most relevant to explain the output variable y_t^{IND} . They include four concentrations measured by online analyzers (olefin feed propylene (i), deisobutanizer recycle propane (ii), isobutane (iii), and n-butane (iv)) and two flow rates (recycled fresh acid (v), and alkylate to storage (vi)).

We compare the simulated and industrial data distributions. The simplified gPROMS model is able to

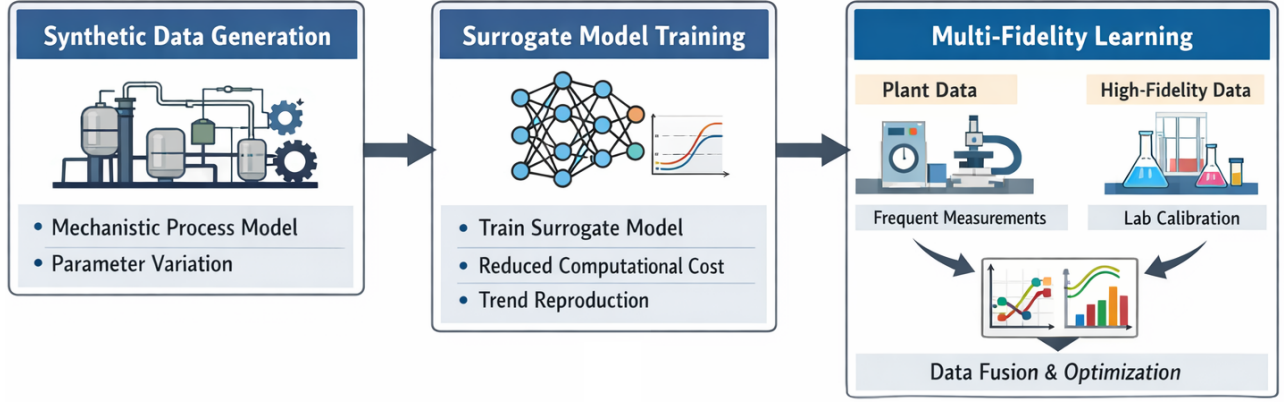


Figure 3. Overview of the proposed hybrid multi-fidelity workflow combining multiple fidelity sources.

reproduce the dominant trends of the plant, although some differences in variability and distribution remain.

We partition the simulation data X^{SIM} (80/20%) into training (TR) and validation subsets to learn the initial surrogate model. Similarly, the industrial dataset X^{IND} provides $T = 13,000$ online samples, while the laboratory dataset X^{LAB} provides $K = 18$ reference samples. Data is split (60/40%) based on balanced HF indices into training and testing (TS) sets by alternating the HF measurements. Following the flowchart shown in Fig. 3, we train and validate the surrogate models $f^{\text{SIM}}(\cdot)$ in Python using a nonlinear feed-forward NN implemented by Keras [19] with four fully connected layers of sizes 6-3-2-1 and activation functions SELU, tanh, SELU, and linear, respectively. The linear OLS regression, i.e., minimizing the sum of squares, is implemented by scikit-learn [20].

Trained prediction models are evaluated on the LF industrial dataset X^{IND} to predict the output $\hat{y}_t^{\text{LF}}$; see (4). The most accurate prediction is subsequently used as an additional input in the MF correction model (8).

To reduce dimensionality, we use PCA on this extended input vector $z_t^{\text{LF}} \in \mathbb{R}^7$. The resulting 4 PCs are used to train a GPR on the HF training points to predict the HF-to-LF differences; as defined in (6).

The MF estimate (8), is obtained by correcting the $\hat{y}_t^{\text{LF}}$ prediction, and performance is reported using RMSE and parity plots. A simple HF model (GPR trained only on HF data) is also trained directly on the HF data X^{LAB} to replicate standard practice and serve as a reference. The final soft-sensor predictions are evaluated against laboratory measurements y_k^{HF} .

RESULTS

The presented workflow encompasses two evaluation stages. First, we verify surrogate models on the SIM dataset, where OLS model achieves RMSE (TR) = 0.1772 and RMSE (TS) = 0.1879, while the NN model achieves

much lower RMSE (TR) = 0.0136 and RMSE (TS) = 0.0127.

Secondly, we evaluate the trained predictors on the IND dataset against laboratory measurements y_k^{HF} using RMSE on the TR and TS splits. This is visualized in Figure 4 and quantitatively summarized in Table 1. The OLS-based MF model (black points) and NN-based MF model (red points) are presented in the left and right columns, respectively. The top row is TR and the bottom row is TS. The blue diagonal represents a perfect fit.

As a reference, the online analyser y_t^{LF} evaluated at the HF indices k yields RMSE (TR) = 0.6759 and RMSE (TS) = 0.5557. This confirms a noticeable mismatch between the LF and the HF data.

Table 1: Predictors compared to laboratory data y_k^{HF} .

	RMSE (TR)	RMSE (TS)	reduction % (TS)
$y_{t(k)}^{\text{LF}}$	0.6759	0.5557	-
$\hat{y}_{t(k)}^{\text{HF}}$	0.4757	0.4735	↓ 14.8%
OLS			
$\hat{y}_{t(k)}^{\text{LF}}$	0.5373	0.4741	↓ 14.7%
$\hat{y}_{t(k)}^{\text{MF}}$	0.4605	0.3831	↓ 31.1%
NN			
$\hat{y}_{t(k)}^{\text{LF}}$	1.1921	1.3297	↑ 139.3%
$\hat{y}_{t(k)}^{\text{MF}}$	0.5566	0.6127	↑ 10.2%

The simple OLS predictor $\hat{y}_t^{\text{LF}}$ improves the output monitoring, achieving RMSE (TR) = 0.5373 and RMSE (TS) = 0.4741, i.e., a 14.7% reduction on the TS split relative to the analyser. This suggests that, even with a simple linear mapping, the simulator trends and engineered regressors carry useful information that partially explains the HF output.

Adding the OLS $\hat{y}_t^{\text{LF}}$ to the MF correction provides the best overall performance. The MF estimate reaches RMSE (TR) = 0.4605 and RMSE (TS) = 0.3831, corresponding to a 31.1% TS improvement. The improvement over the HF-only model (0.3831 vs 0.4735 on TS)

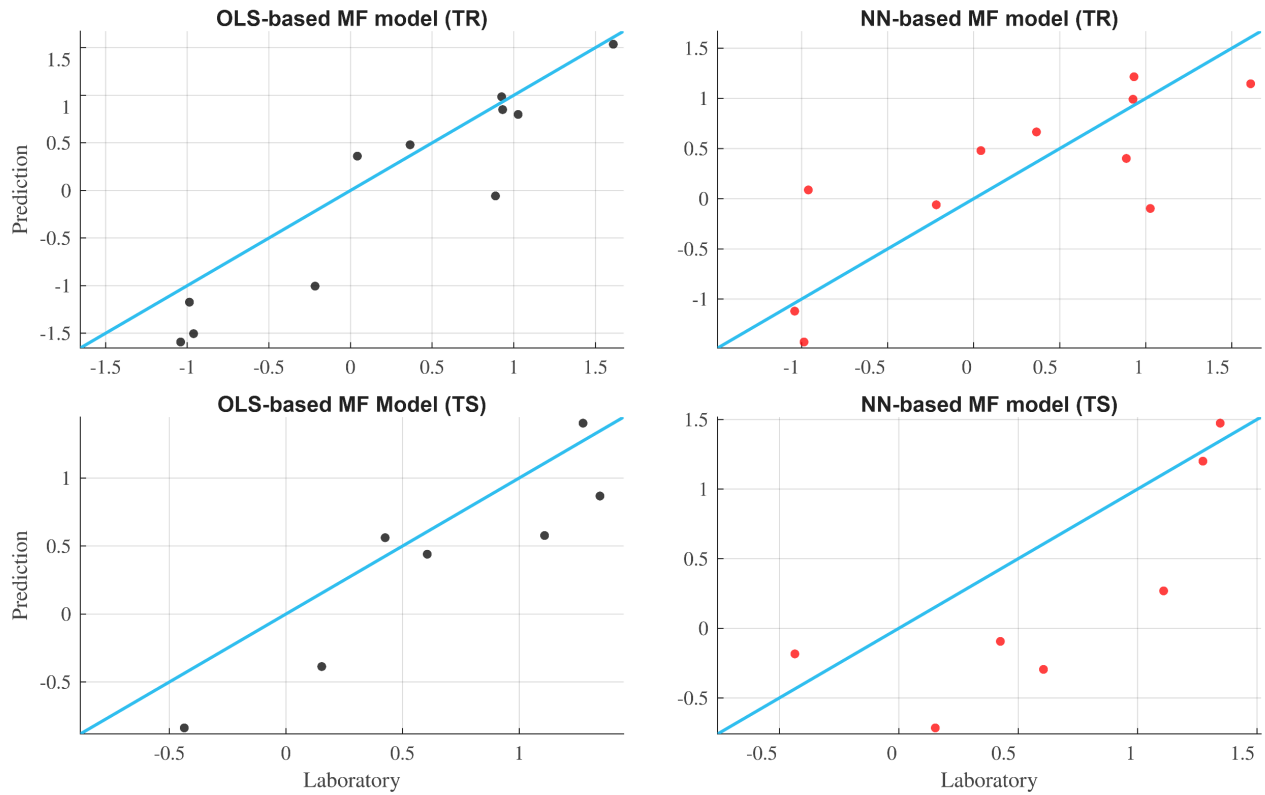


Figure 4. Parity plots of MF predictions $\hat{y}_{t(k)}^{MF}$ compared to the laboratory data y_k^{HF} .

indicates that the correction benefits from the dense LF data, rather than relying on sparse HF data alone.

The simple HF model yields RMSE (TR) = 0.4757 and RMSE (TS) = 0.4735 (14.8% improvement on TS) is comparable with the OLS predictor, but it is worse than MF on the TS split.

Although the NN fits the simulation data accurately (near-zero simulation RMSE), its prediction on the industrial data is poor; RMSE (TR) = 1.1921 and RMSE (TS) = 1.3297 (worse by 139.3%). This may be caused by the domain-shift, where the NN learns simulator-specific nonlinearities that do not map clearly to the real plant (different disturbances, unmodelled constraints).

The MF correction with the NN $\hat{y}_t^{LF}$ partially improves this mismatch (RMSE = 0.6127 on TS), but it still does not surpass the $\hat{y}_{t(k)}^{HF}$, nor the $y_{t(k)}^{LF}$.

Overall, the results show that MF correction is most effective when the predictor provides a reasonable trend-consistent estimate.

CONCLUSION

A multi-fidelity correction framework was applied to reconcile dense industrial measurements with sparse laboratory references using a simulator-trained surrogate and a Gaussian process corrector. The ordinary least

squares-based multi-fidelity model achieves the best performance, reducing test prediction error against laboratory measurements by 31% relative to the industrial analyzer at high-fidelity time points, and outperforming a standard high-fidelity model (soft sensor). These results indicate that, even with a deliberately simplified simulator, the surrogate is able to capture the underlying trend, enabling the multi-fidelity correction to further improve predictions toward laboratory accuracy. The neural network surrogate, despite excellent fit to simulation data, did not transfer robustly to the industrial dataset; limiting the multi-fidelity correction.

Future work should focus on increasing physical realism and coverage of the simulation model to provide more information in the training phase.

ACKNOWLEDGEMENTS

This work is funded by the Slovak Research and Development Agency under the project APVV-21-0019, by the Scientific Grant Agency of the Slovak Republic under the grant 1/0263/25. We acknowledge the financial support provided by the European Union through the NextGenerationEU programme under the Recovery and Resilience Plan for Slovakia, projects no. 09I01-03-V05-00002, no. 09I01-03-V04-00024, and the Early Stage

AUTHOR IDENTIFIERS

Author ORCID:

Fáber R: 0000-0003-0088-3467

Paulen R: 0000-0002-1599-2634

REFERENCES

1. Engell S. Feedback control for optimal process operation. *Journal of Process Control* 17:203-219 (2007).
<https://doi.org/10.1016/j.jprocont.2006.10.011>
2. Qin SJ. Survey on data-driven industrial process monitoring and diagnosis. *Annual Reviews in Control* 36:220-234 (2012).
<https://doi.org/10.1016/j.arcontrol.2012.09.004>
3. Adams JN, van Zelst SJ, Rose T, van der Aalst WMP. Explainable concept drift in process mining. *Information Systems* 114:102177 (2023).
<https://doi.org/10.1016/j.is.2023.102177>
4. Thalmann S, Mangler J, Schreck T, Huemer C, Streit M, Pauker F, Weichhart G, Schulte S, Kittl C, Pollak C, Vukovic M, Gashi M, Rinderle-Ma S, Suschnigg J, Jekic N, Lindstaedt S. Data Analytics for Industrial Process Improvement: A Vision Paper. *IEEE 20th Conference on Business Informatics (CBI)* 92-96
5. Kadlec P, Grbić R, Gabrys B. Review of adaptation mechanisms for data-driven soft sensors. *Computers & Chemical Engineering* 35:1-24 (2011).
<https://doi.org/10.1016/j.compchemeng.2010.07.034>
6. Farsang B, Gomori Z, Horvath G, Nagy G, Németh S, Abonyi J. Simultaneous Validation of Online Analyzers and Process Simulators by Process Data Reconciliation. *Chemical Engineering Transactions* 32, 1303-1308 (2013)
7. Colling E, Farenzena M, Trierweiler JO. Data-driven prediction of intermediate product properties in petroleum refining. *Ind. Eng. Chem. Res.* 64:23510-23521 (2025).
<https://doi.org/10.1021/acs.iecr.5c02838>
8. Syauqi A, Nagulapati VM, Chaniago YD, Ningtyas JA, Andika R, Lim H. Advancement in power-to-methanol integration with steel industry waste gas utilization through solid oxide electrolyzer cells: surrogate model-based approach for optimization. *Sustainable Energy Technologies and Assessments* 73:104160 (2025).
<https://doi.org/10.1016/j.seta.2024.104160>
9. Fáber R, Vaccari M, Bacci di Capaci R, Ůubuřký K, Pannocchia G, Paulen R. Data-based multi-fidelity modeling for online sensors correction. *Computers & Chemical Engineering* 211:109665 (2026).
<https://doi.org/10.1016/j.compchemeng.2026.109665>
10. Kadlec P, Gabrys B, Strandt S. Data-driven soft sensors in the process industry. *Computers & Chemical Engineering* 33:795-814 (2009).
<https://doi.org/10.1016/j.compchemeng.2008.12.012>
11. Jin H, Dong X, Qian B, Wang B, Yang B, Chen X. Soft sensor modeling using deep learning with maximum relevance and minimum redundancy for quality prediction of industrial processes. *ISA Transactions* 159:293-311 (2025).
<https://doi.org/10.1016/j.isatra.2025.02.010>
12. Zhang L, Ren G, Li S, Du J, Xu D, Li Y. A novel soft sensor approach for industrial quality prediction based TCN with spatial and temporal attention. *Chemometrics and Intelligent Laboratory Systems* 257:105272 (2025).
<https://doi.org/10.1016/j.chemolab.2024.105272>
13. Peterson L, Gosea IV, Benner P, Sundmacher K. Digital twins in process engineering: an overview on computational and numerical methods. *Computers & Chemical Engineering* 193:108917 (2025).
<https://doi.org/10.1016/j.compchemeng.2024.108917>
14. Iranshahi K, Brun J, Arnold T, Sergi T, Müller UC. Digital twins: recent advances and future directions in engineering fields. *Intelligent Systems with Applications* 26:200516 (2025).
<https://doi.org/10.1016/j.iswa.2025.200516>
15. Sun B, Yang C, Wang Y, Gui W, Craig I, Olivier L. A comprehensive hybrid first principles/machine learning modeling framework for complex industrial processes. *Journal of Process Control* 86:30-43 (2020).
<https://doi.org/10.1016/j.jprocont.2019.11.012>
16. Nargesian F, Samulowitz H, Khurana U, Khalil EB, Turaga D. Learning feature engineering for classification. *Proceedings of the Twenty-Sixth International Joint Conference on Artificial Intelligence* :2529-2535 (2017).
<https://doi.org/10.24963/ijcai.2017/352>
17. F?ber R, Vaccari M, Capaci RB, Lubu?k? K, Pannocchia G, Paulen R. Soft-sensor-enhanced monitoring of an alkylation unit via multi-fidelity model correction. *Systems and Control Transactions* 4:1389-1395 (2025).
<https://doi.org/10.69997/sct.125527>
18. Siemens, gPROMS Digital Process Design and Operations.
19. Abadi M. Tensorflow: learning functions at scale. *SIGPLAN Not.* 51:1-1 (2016).
<https://doi.org/10.1145/3022670.2976746>

20. Pedregosa F, *et al.*, Scikit-learn: Machine Learning in Python, *JMLR* 12 pp. 2825-2830 (2011)

© 2026 by the authors. Licensed to PSEcommunity.org and PSE Press. This is an open access article under the creative commons CC-BY-SA licensing terms. Credit must be given to creator and adaptations must be shared under the same terms. See <https://creativecommons.org/licenses/by-sa/4.0/>

