

Identifiability-Guided Bound Tightening for Parameter Identification Under Plant-Model Mismatch

Terrance Wilms¹, Steffi Knorn¹

Abstract—Parameter identification (PI) for nonlinear bioprocess models is challenging due to wide parameter bounds, sparse and noise data sets, multimodal objectives, and plant–model mismatch. Hybrid global–local strategies combining stochastic search with gradient-based refinement are effective but can be sample-inefficient when many parameter sets yield infeasible simulations and different outputs inform different parameter subsets. We propose an identifiability-guided bound tightening strategy inspired by domain-reduction techniques from global optimization. Stage 1 performs Latin hypercube sampling in log-parameter space, filters feasible simulations, and ranks candidates by single-output objectives. For each output, elite parameter samples are pooled only for structurally identifiable parameters regarding the single output, yielding parameter-wise tightened bounds. Stage 2 solves the full multi-output PI problem via multi-start gradient-based optimization or hybrid stochastic-local refinement within these tightened bounds. In an in-silico fed-batch case study which takes plant–model mismatch into account, the identifiability-guided approach reduces prediction error by 20% and relative parameter error by 47% compared to standard multi-start baselines, achieving optimal performance with efficient screening budget.

I. INTRODUCTION

Parameter identification (PI) of nonlinear dynamic models is crucial for model-based monitoring and control in bioprocesses [1]–[3]. But PI in biotechnology is often ill-conditioned due to wide prior parameter bounds, sparse and noise data sets, multimodal objectives, and computationally expensive ODE simulations, causing standard local optimization to fail without careful initialization [4]–[6]. These challenges are particularly critical in online applications like adaptive experiment design, where parameters must be re-identified fast during ongoing cultivations [7]. Common practice thus relies on multi-start strategies, e.g. with Latin hypercube sampling (LHS) [8] in log-parameter space and hybrid global–local workflows combining stochastic metaheuristics with gradient-based (via forward sensitivity) refinement [4], [6], [9].

A key challenge in bioprocess PI is that many parameters are structurally or practically unidentifiable for a given set of measured outputs [10]. When adapting models to new strains, media, or operating regimes, prior parameter knowledge is weak and relevant parameter scales are often unknown *a priori*, motivating wide bounds spanning multiple orders of magnitude. These wide bounds increase both the risk of numerically ill-conditioned simulations (ODE failures) and the effective search space for nonlinear optimizers. Furthermore, unmodelled kinetics and physiological variability introduce

plant–model mismatch (PMM), in which the true process dynamics differ structurally from the candidate model [11].

In model-based PI, ill-conditioning and practical non-identifiability are commonly analyzed using sensitivity analysis, and are closely linked to ill-posedness of the least-squares problem [7], [12], [13]. Such sensitivity-based approaches reduce the effective identification problem either by selecting identifiable parameter subsets or by sequentially estimating parameter groups. However, local sensitivity-based strategies typically depend on a meaningful nominal parameter vector, which is often unavailable when dealing with new strains or operating conditions. Lie-derivative-based tools like STRIKE-GOLDD [14] enable output-wise structural identifiability analysis *a priori*, from model structure alone, determining which parameters are in principle uniquely identifiable from each measured output under noise-free conditions.

In global optimization, bound tightening (BT), also known as domain reduction, is a standard technique for focusing the search around optimal solutions in parameter space [15]. The BT strategy proposed here is conceptually related to bounded-error [16] and set-membership [17] identification, where feasible parameter sets are characterized from models, data, and error bounds. However, we do not compute guaranteed feasible-set enclosures, but instead use a heuristic, simulation-based BT, in which elite samples are selected according to objective values (screening) over wide parameter bounds and pooled only for structurally identifiable output–parameter pairs. By combining this screening approach with *a priori* structural identifiability (SI) information, we determine which outputs can identify which parameters, yielding an identifiability-guided (IG) BT stage before refinement.

Practitioners often address these difficulties heuristically, by first calibrating parameter subsets using single-output objectives (e.g., biomass or substrate alone) to localize plausible ranges, then re-running full multi-output PI within tightened bounds. Plausibility is assessed by numerical robustness (absence of ODE errors), realistic trajectories, and acceptable prediction errors. While effective, this trial-and-error workflow is rarely formalized and seldom exploits *a priori* identifiability information. We formalize this into a metaheuristic strategy by using structural identifiability to define output-wise screening subproblems, exploring these with space-filling sampling, and pooling elite candidates only for structural identifiable parameters to construct parameter-wise tightened bounds.

Since bioprocess models (especially simple unstructured

¹Chair of Control, Technische Universität Berlin, Germany. knorn@tu-berlin.de

kinetic models) capture only a fraction of true plant dynamics, ‘true’ parameter values may be undefined for the candidate model. Under PMM, PI converges to pseudo-true best-fit parameters of the misspecified model [18]. We therefore evaluate parameter accuracy relative to pseudo-true parameters in the PMM scenario, obtained by fitting the candidate model to noise-free plant data.

Contributions. We (i) propose an IG, screening-based BT strategy that automates heuristic output-wise calibration workflows; (ii) benchmark six solver configurations (M1–M6) combining different BT strategies (none, joint, IG) with different refinement methods (gradient-based (GB) multi-start (MS), hybrid stochastic-local); and (iii) quantify the effect of screening budget on prediction and parameter accuracy under controlled PMM. Results from 100 Monte Carlo (MC) runs demonstrate that IG screening reduces prediction error by 20% and relative parameter error by 47% compared to standard baselines, with optimal screening budgets of 500–1000 simulations.

II. METHOD

A. Parameter identification problem

We consider a nonlinear ODE model

$$\dot{x} = f(x, u, \theta), \quad y = h(x), \quad (1)$$

with states $x \in \mathbb{R}^{n_x}$, inputs $u \in \mathbb{R}^{n_u}$, outputs $y \in \mathbb{R}^{n_y}$, and parameters $\theta \in \mathbb{R}^{n_\theta}$. Given measurements $y^M(t_k)$ at times t_k and diagonal measurement noise covariance matrix $\tilde{C}_k \in \mathbb{R}^{n_y \times n_y}$, we estimate parameters by minimizing the weighted maximum-likelihood estimation objective

$$\hat{\theta} = \arg \min_{\theta \in \Theta} \Phi_{\text{PI}}(\theta), \quad (2)$$

where

$$\Phi_{\text{PI}}(\theta) = \sum_{k=1}^N [y^M(t_k) - y(t_k; \theta)]^\top \times \tilde{C}_k^{-1} [y^M(t_k) - y(t_k; \theta)]. \quad (3)$$

The feasible parameter set is a box constraint

$$\Theta = \{\theta : \underline{\theta} \leq \theta \leq \bar{\theta}\}, \quad (4)$$

with wide but realistic initial bounds to reflect limited prior knowledge. All sampling and optimization is performed in $\log_{10}$ -space to handle high order-of-magnitude ranges, as suggested by [6].

For output-wise screening we use the single-output contribution

$$\Phi_{\text{PI},\ell}(\theta) = \sum_{k=1}^N [y_\ell^M(t_k) - y_\ell(t_k; \theta)]^\top \times \tilde{C}_{k,\ell}^{-1} [y_\ell^M(t_k) - y_\ell(t_k; \theta)]. \quad (5)$$

B. Gradient computation via forward sensitivities

For gradient-based optimization (M1–M6), we supply analytical gradients for $j = 1, \dots, n_\theta$ as

$$\nabla_{\theta_j} \Phi_{\text{PI}} = 2 \sum_{k=1}^N (y^M(t_k) - y(t_k))^\top \tilde{C}_k^{-1} \frac{\partial y}{\partial \theta_j}, \quad (6)$$

$$\frac{\partial y}{\partial \theta_j} = \frac{\partial h}{\partial x} \frac{\partial x}{\partial \theta_j} + \frac{\partial h}{\partial \theta_j}, \quad (7)$$

computed via forward sensitivity equations [19]:

$$\frac{d}{dt} \frac{\partial x}{\partial \theta_j} = \frac{\partial f}{\partial x} \frac{\partial x}{\partial \theta_j} + \frac{\partial f}{\partial \theta_j}, \quad (8)$$

integrated alongside the state equations. This avoids finite-difference approximations and accelerates convergence, especially in high-dimensional parameter spaces.

C. Stage 1: Bound Tightening

We first draw N_{sim} samples $\theta^{(i)}$ using LHS in $\log_{10}$ -space over the prior box Θ and discard infeasible samples (e.g. integration failures). For each feasible sample, we evaluate the joint PI objective $\Phi_{\text{PI}}(\theta^{(i)})$ via (3) and the single-output objectives $\Phi_{\text{PI},\ell}(\theta^{(i)})$ via (5) for all outputs $\ell = 1, \dots, n_y$.

Joint tightening (M3–M4). We form an elite set $\mathcal{E}_{\text{joint}}$ by keeping the n_{best} samples with the lowest joint objective value Φ_{PI} . For each parameter j , we define the pool $\mathcal{P}_j = \{\theta_j^{(i)} \mid i \in \mathcal{E}_{\text{joint}}\}$ and compute tightened bounds as

$$\underline{\theta}_j^{\text{init}} = \min(\mathcal{P}_j), \quad \bar{\theta}_j^{\text{init}} = \max(\mathcal{P}_j). \quad (9)$$

IG BT (M5–M6). For each output ℓ , we form an elite set $\mathcal{E}_\ell$ by keeping the n_{best} samples with the smallest $\Phi_{\text{PI},\ell}$. We then pool parameter values only from outputs that structurally identify parameter j

$$\mathcal{P}_j = \bigcup_{\ell: \mathbf{M}_{\text{id}}[\ell, j] = 1} \{\theta_j^{(i)} \mid i \in \mathcal{E}_\ell\}, \quad (10)$$

where $\mathbf{M}_{\text{id}} \in \{0, 1\}^{n_y \times n_\theta}$ encodes SI: $\mathbf{M}_{\text{id}}[\ell, j] = 1$ if parameter θ_j is structurally identifiable from output ℓ (computed once *a priori* using STRIKE-GOLDD [14]). Tightened bounds are computed as in M3–M4. For empty pools (non-identifiable parameters), we retain the original bounds $\underline{\theta}_j$ and $\bar{\theta}_j$. The resulting tightened box is denoted $\Theta^{\text{init}} = [\underline{\theta}^{\text{init}}, \bar{\theta}^{\text{init}}]$, which is used for the initialization of Stage 2.

D. Stage 2: Refinement Within Tightened Bounds

In Stage 2 we refine the best solution by running the chosen solver on Φ_{PI} constrained to Θ^{init} . Initial points (multi-start methods M1, M3, M5) or initial populations (hybrid methods M2, M4, M6) are sampled within Θ^{init} . All refinement is performed within the tightened bounds.

Algorithm 1 IG bound tightening and refinement (M5–M6)

Require: Bounds Θ , data $\{y^M(t_k)\}_{k=1}^N$, budgets N_{sim} , N_{ref} , elite size n_{best}

- 1: Compute output–parameter identifiability matrix $\mathbf{M}_{\text{id}}$
- 2: **Stage 1 (bound tightening):** Sample $\{\theta^{(i)}\}_{i=1}^{N_{\text{sim}}}$ by LHS in $\log_{10}$ -space over Θ
- 3: **for** $i = 1$ to N_{sim} **do**
- 4: Simulate candidate model; discard $\theta^{(i)}$ if infeasible
- 5: Compute single-output costs $\Phi_{\text{PI},\ell}(\theta^{(i)})$ for outputs ℓ
- 6: **end for**
- 7: **IG tightening:** For each output ℓ , keep elite set $\mathcal{E}_\ell$ of size n_{best} with smallest $\Phi_{\text{PI},\ell}$
- 8: **for** each parameter j **do**
- 9: Form pool $\mathcal{P}_j \leftarrow \bigcup_{\ell: \mathbf{M}_{\text{id}}[\ell,j]=1} \{\theta_j^{(i)} \mid i \in \mathcal{E}_\ell\}$
- 10: Compute tightened bounds $(\theta_j^{\text{init}}, \bar{\theta}_j^{\text{init}})$ from $\mathcal{P}_j$
- 11: **end for**
- 12: **Stage 2 (refinement):** Sample N_{ref} initial points in Θ^{init}
- 13: Run refinement solver on Φ_{PI} within Θ^{init} ; return best solution

E. Solver configurations (M1–M6)

We compare six configurations under a fixed total evaluation budget. Methods M1–M2 serve as baselines without BT. Methods M3–M4 implement joint BT (ranking all samples by the multi-output objective Φ_{PI}). Methods M5–M6 implement the proposed IG BT, which pools elite samples per output according to SI. Algorithm 1 details the IG approach.

- **M1: GB MS (baseline).** Draw N_{ref} ; LHS starts over full bounds Θ ; run gradient-based local optimization.
- **M2: Hybrid stochastic-local (baseline).** Run GA over full bounds Θ , followed by gradient-based refinement.
- **M3: Joint BT + GB MS.** Stage 1 ranks samples by joint objective Φ_{PI} , forms Θ^{init} from n_{best} elites; Stage 2 performs gradient-based multi-start within Θ^{init} .
- **M4: Joint BT + hybrid stochastic-local.** As M3, with GA-hybrid refinement.
- **M5: IG BT + GB MS.** Output-wise elite selection with IG pooling (Algorithm 1), refined by gradient-based MS.
- **M6: IG BT + hybrid stochastic-local.** As M5, with GA-hybrid refinement.

III. CASE STUDY AND EXPERIMENTAL SETUP

We evaluate **M1–M6** on an unstructured fed-batch model based on [20]:

$$\dot{x} = f(x, u_S, \theta), \quad x = [m_S, m_X, m_P, V]^\top \quad (11)$$

with biomass (m_X), substrate (m_S), product (m_P), and volume V , where

$$\dot{m}_S = (-Y_{S,P} r_P - Y_{S,X} r_X) m_X + u_S c_{S,\text{in}}, \quad (12)$$

$$\dot{m}_X = r_X m_X, \quad (13)$$

$$\dot{m}_P = r_P m_X, \quad (14)$$

$$\dot{V} = u_S. \quad (15)$$

TABLE I

MODEL PARAMETERS: PLANT/REFERENCE VALUES θ^{plant} , PSEUDO-TRUE PARAMETERS θ^{pseudo} AND PRIOR BOUNDS FOR PI.

Parameter (unit)	Plant/ref.	Pseudo-true	LB	UB
$Y_{S,X}$ (g _X /g _S)	5.00	4.98	10^{-2}	10^1
$\mu_{\text{max},X}$ (h ⁻¹)	0.20	0.094	10^{-3}	1
$K_{S,X}$ (g/L)	6.00	1.37	10^{-3}	10^2
$K_{P,X}$ (g/L)	0.10	0.065	10^{-3}	10^2
$Y_{S,P}$ (g _P /g _S)	1.00	1.63	10^{-2}	10^1
$\mu_{\text{max},P}$ (h ⁻¹)	0.01	0.014	10^{-4}	0.5
k_m (g/L)	1.00	1.75	10^{-3}	10^2
k_i (g/L)	0.50	0.90	10^{-3}	10^2

The measured outputs are concentrations $y = [c_S, c_X, c_P]^\top$ with $c_i = m_i/V$. The feed flow rate u_S with inlet substrate concentration $c_{S,\text{in}} = 200$ g/L is the known input. Kinetics follow a Monod term with product inhibition for growth and a substrate-saturated product term based on [20]:

$$r_X = \mu_{\text{max},X} \frac{c_S}{c_S + K_{S,X}} \frac{K_{P,X}}{c_P + K_{P,X}}, \quad (16)$$

$$r_P = \mu_{\text{max},P} \frac{c_S}{k_m + c_S + c_S^2/k_i}. \quad (17)$$

The parameter vector of the candidate model is

$$\theta = [Y_{S,X}, \mu_{\text{max},X}, K_{S,X}, K_{P,X}, Y_{S,P}, \mu_{\text{max},P}, k_m, k_i]^\top.$$

Note that volume V is not included in the PI, as it contains no information about the parameters, merely represents the sum of known input variables and is therefore known.

Plant–model mismatch (PMM) and pseudo-true parameters. Synthetic data are generated with a *plant model* that includes a Haldane-type substrate inhibition term:

$$r_X^{\text{plant}} = \mu_{\text{max},X} \frac{c_S}{K_{S,X} + c_S + \frac{c_S^2}{K_{I,S}}} \frac{K_{P,X}}{c_P + K_{P,X}}, \quad (18)$$

using a fixed substrate inhibition constant $K_{I,S} = 10$ g/L. In this setting PI is performed with the simplified candidate model ((16)-(17)) and therefore, the ‘true’ plant parameters are generally not recoverable by the candidate model. Instead, identification converges to *pseudo-true* (θ^{pseudo}) best-fit parameters. We thus define

$$\theta^{\text{pseudo}} = \arg \min_{\theta \in \Theta} \Phi_{\text{PI}}(\theta; y^{\text{plant, clean}}), \quad (19)$$

where $y^{\text{plant, clean}}$ are noise-free outputs generated by the plant model (18), and evaluate parameter accuracy relative to θ^{pseudo} . These pseudo-true parameters were obtained via extensive optimization ($N_{\text{LHS}} = 500$, population size $N_{\text{pop}} = 1000$, $N_{\text{gen}} = 500$ generations) on noise-free outputs, representing the best achievable fit under structural PMM.

Table I lists the plant/reference parameter values θ^{plant} used to generate the PMM trajectories, the pseudo-true best-fit parameters θ^{pseudo} of the candidate model under mismatch, and the prior bounds defining Θ for PI as lower bounds (LB) and upper bounds (UB).

To obtain realistic synthetic data, we add signal-dependent noise [21]:

$$\sigma_i(t) = a_i y_i(t) + b_i, \quad (20)$$

TABLE II
NOISE MODEL COEFFICIENTS IN (20) FROM [21].

Variable	a_i [-]	b_i [g L ⁻¹]
c_S	0.06	0.25
c_X	0.04	0.03
c_P	0.05	0.25

using values from Table II and sample $y_{k,i}^M = y_i(t_k) + \varepsilon_{k,i}$ with $\varepsilon_{k,i} \sim \mathcal{N}(0, \sigma_i(t_k)^2)$, where a_i is dimensionless and b_i has the same unit as y_i (here g L⁻¹). Such signal-dependent noise captures the empirically observed increase of measurement variance with concentration in typical bioprocess analytics and PAT signals.

A. Performance metrics and evaluation protocol

Computational budget. To ensure fair comparison, we control computational effort via the total number of objective evaluations N_{eval} , directly reported by each solver. For multi-start methods (M1, M3, M5):

$$N_{\text{eval}} = N_{\text{sim}} + \sum_{i=1}^{N_{\text{ref}}} N_{\text{fmincon}}^{(i)}, \quad (21)$$

where N_{sim} is the screening budget and $N_{\text{fmincon}}^{(i)}$ the evaluations in the i -th local run (reported by solver). For GA-hybrid methods (M2, M4, M6):

$$N_{\text{eval}} = N_{\text{sim}} + N_{\text{GA}}, \quad (22)$$

where N_{GA} is the total GA evaluations (reported by solver). For baselines without screening (M1-M2), $N_{\text{sim}} = 0$.

Performance metrics. For each solver configuration and budget, we perform 100 MC runs with different random seeds. Random seeds refer to the pseudo-random number generator states used for sampling (LHS) and stochastic optimization components (GA), ensuring independent MC trials. We report:

- **Prediction error (NRMSE).** Normalized root mean squared error per output ℓ and experiment e :

$$\text{NRMSE}_{e,\ell} = \frac{\|y_{e,\ell}^{\text{model}} - y_{e,\ell}^M\|_2}{\|y_{e,\ell}^M - \bar{y}_{e,\ell}^M \mathbf{1}\|_2}, \quad (23)$$

where $\bar{y}_{e,\ell}^M$ is the sample mean and $\mathbf{1}$ is a vector of ones of matching length. We report the total NRMSE

$$\text{NRMSE}_{\text{total}} = \sum_{e \in \{\text{AV1}, \text{AV2}, \text{CV}\}} \sum_{\ell=1}^{n_y} \text{NRMSE}_{e,\ell} \quad (24)$$

with median and interquartile range (IQR) over runs.

- **Parameter error.** Relative error w.r.t. pseudo-true values:

$$\varepsilon_j = \frac{|\hat{\theta}_j - \theta_j^{\text{pseudo}}|}{\theta_j^{\text{pseudo}}}, \quad j = 1, \dots, n_{\theta}. \quad (25)$$

We report the mean absolute relative parameter error

$$\bar{\varepsilon} = \frac{1}{n_{\theta}} \sum_{j=1}^{n_{\theta}} \varepsilon_j \quad (26)$$

with median and IQR over runs.

TABLE III
 $\mathbf{M}_{\text{id}}$ (X = IDENTIFIABLE FROM OUTPUT).

Parameter	c_S (Output 1)	c_X (Output 2)	c_P (Output 3)
$Y_{S,X}$		X	
$\mu_{\text{max},X}$	X	X	
$K_{S,X}$	X	X	
$K_{P,X}$			X
$Y_{S,P}$			X
$\mu_{\text{max},P}$			X
k_m	X	X	X
k_i	X	X	X

In-Silico Experimental setup. Three in-silico experiments are generated with the PMM model: one batch and one fed-batch experiment (Exp. 1-2) for PI (auto-validation, AV1 and AV2), and one fed-batch (Exp. 3) for independent cross-validation (CV). All trajectories include signal-dependent measurement noise (Table II).

B. Solver parameter settings

To ensure fair comparison across methods, we control the total computational budget via objective evaluations N_{eval} as defined in Section III. The specific parameter settings are

Baseline methods (M1–M2):

- **M1:** $N_{\text{ref}} \in \{5, 10, 20\}$ LHS optimizes with initial values over full bounds Θ via GB local optimization (MATLAB `fmincon` with forward sensitivities).
- **M2:** GA with population sizes $N_{\text{pop}} \in \{20, 50, 100\}$ and corresponding generations $N_{\text{gen}} \in \{10, 20, 50\}$, followed by GB local refinement of the best individual. Under a fixed evaluation budget, we emphasize larger populations over many generations to maintain diversity and reduce premature convergence, consistent with the trade-off discussed in [22].

Bound tightening methods (M3–M6): N_{sim} and elite set size are set to $N_{\text{sim}} \in \{100, 300, 500, 1000, 2000\}$ and $n_{\text{best}} = 10$.

- **M3, M5:** $N_{\text{ref}} = 10$ LHS start points within Θ^{init} (settings as in M1).
- **M4, M6:** GA refinement within Θ^{init} using $N_{\text{pop}} = 20$, $N_{\text{gen}} = 10$ ($N_{\text{GA}} = 200$) (settings as in M2).

For methods M3–M6, N_{sim} is varied to assess its impact on performance (Fig. 3), while refinement budgets are held constant to isolate the effect of BT quality. All gradient-based optimizations use analytical gradients computed via forward sensitivity equations (8). The implementation was validated using `fmincon`'s built-in gradient checking (`CheckGradients`).

IV. RESULTS

Table III reports the structural identifiability matrix $\mathbf{M}_{\text{id}}$ obtained with STRIKE-GOLDD. Growth parameters ($Y_{S,X}$, $\mu_{\text{max},X}$) are informed primarily by biomass (c_X), production parameters ($Y_{S,P}$, $\mu_{\text{max},P}$) by product (c_P), while substrate kinetics ($K_{S,X}$, k_m , k_i) contribute to multiple outputs.

Parameter localization under Stage 1. Fig. 1 visualizes normalized parameter distributions from Stage 1 screening

TABLE IV
PMM RESULTS OVER 100 RUNS (MEDIAN [IQR]).

Method	NRMSE _{total} [-]	Rel. param. error $\bar{\varepsilon}$ [-]
M1	8.9 [0.1]	1.5 [0.1]
M2	8.9 [0.1]	1.5 [0.1]
M3	6.6 [0.3]	1 [0.2]
M4	7.6 [0.3]	1.2 [0.3]
M5	7.5 [0.2]	1.1 [0.2]
M6	7.1 [0.2]	0.8 [0.2]

($N_{\text{sim}} = 300$, $n_{\text{best}} = 10$) on noisy in-silico data. Gray shows the initial LHS samples, yellow the joint elites (M3–M4), and blue the IG elites (M5–M6). Across parameters, IG pooling yields tighter and more reliable localization around the true values, while joint screening can mis-localize weakly informed parameters by mixing outputs that do not identify them.

Detailed performance analysis (Table IV, Fig. 2). We evaluate M1–M6 on the PMM scenario with 100 MC runs (random seeds). Performance is reported as NRMSE_{total} and $\bar{\varepsilon}$ (w.r.t. pseudo-true parameters), summarized by median and IQR in Table IV. Baselines without bound tightening (M1–M2) plateau at NRMSE_{total} $\approx$ 8.9 and $\bar{\varepsilon} \approx$ 1.5, indicating that wide bounds and infeasible regions lead to inefficient search. Joint tightening improves prediction accuracy (M3 achieves the lowest NRMSE_{total}), but parameter recovery remains less robust (IQR, Table IV). In contrast, IG BT (M5–M6) yields the lowest $\bar{\varepsilon}$ (M6 reaches the best $\bar{\varepsilon}$), while maintaining competitive NRMSE_{total} and low IQR, demonstrating improved pseudo-true parameter localization under mismatch. Fig. 2 further shows that these improvements are achieved at comparable evaluation counts. Note that each point summarizes stochastic (LHS and GA-based) configurations for 100 MC runs. Hence, larger N_{eval} does not necessarily yield monotonic NRMSE_{total} improvements.

Effect of N_{sim} (Fig. 3). Fig. 3 shows performance as a function of N_{sim} for the BT methods M3–M6. For IG methods (M5–M6), best performance is typically reached around $N_{\text{sim}} \approx 500$ –1000, with diminishing returns beyond. Overall, the IG variants achieve comparable or better accuracy at lower N_{sim} , indicating improved sample-efficiency.

Trajectory-level comparison (Fig. 4). Fig. 4 compares biomass trajectories $c_X(t)$ for the three PMM in-silico experiments (AV1, AV2, CV). The IG BT configuration (M6) yields a closer match to the measured c_X profiles in all experiments compared to the baseline (M2), consistent with improved parameter recovery and generalization under mismatch. For readability in the two-column format, only c_X is shown, because substrate and product show similar improvements. Furthermore, Fig. 4 shows M6 as the representative IG configuration because it achieved the lowest $\bar{\varepsilon}$ among all tested methods.

V. CONCLUSION

We propose an IG screening-based bound tightening strategy that automates heuristic output-wise calibration workflows. By combining simulation-based domain reduction

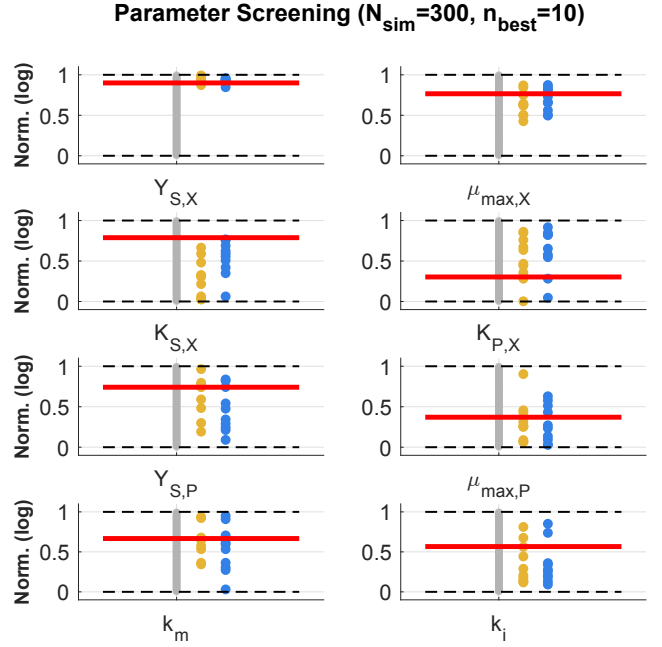


Fig. 1. Normalized parameter distributions from Stage 1 screening ($n_{\text{sim}} = 300$, $n_{\text{best}} = 10$) on in-silico data with measurement noise. Gray: initial LHS samples; yellow: joint screening elites (M3–M4); blue: IG elites (M5–M6); red line: true parameter value.

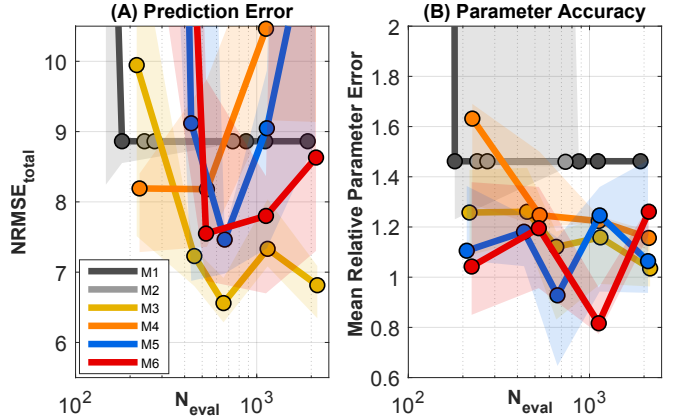


Fig. 2. PMM convergence over objective evaluations (100 runs; median and IQR band). M1–M2 are baselines without BT, M3–M4 use joint tightening, and M5–M6 use IG tightening.

with structural identifiability analysis, the method focuses computational effort on parameter subsets genuinely informed by measurements.

The parameter distribution analysis reveals the mechanism underlying these improvements: IG screening avoids the output-averaging bias of joint methods, ensuring that elite samples preserve structural information about parameter subsets. This leads to higher convergence rates and lower $\bar{\varepsilon}$ in the refinement phase. At trajectory level, we show that the improved parameter accuracy translates into visibly better biomass predictions in AV01, AV02 and CV, supporting the method’s robustness under structural mismatch.

In a controlled in-silico study with PMM, IG screening reduces prediction error by 20% and parameter error by 47% compared to standard multi-start baselines, achieving

Effect of Screening Budget on Parameter Accuracy

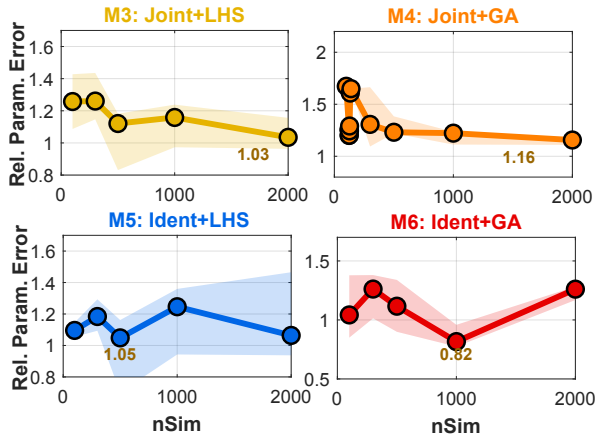


Fig. 3. Effect of N_{sim} (median and IQR) for M3-M6. IG screening reaches optimal performance at $N_{\text{sim}} \approx 500$ –1000.

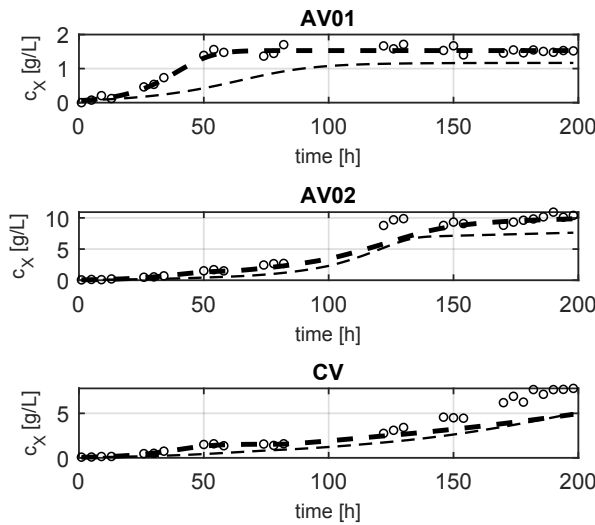


Fig. 4. Biomass concentration c_X for AV01, AV02 and CV. Markers: noisy measurements. Lines: simulated trajectories using parameter estimates from the median run of the baseline (M2, thin dashed lines) and the IG BT (M6, thick dashed lines), each selected from the corresponding MC batch.

optimal performance with 500–1000 screening simulations. The best configuration reaches $\bar{\varepsilon}$ of 0.82 with low IQR, demonstrating robust parameter recovery despite structural mismatch. In practice, the presented method is most useful for early stage bioprocess model calibration of moderate-size nonlinear models with large prior parameter bounds, sparse data, and a need for repeated and rapid PI (e.g. for online PI). The dominant computational cost is the screening (repeated model simulations) and refinement. Sorting elite sets and applying the *a-priori* identifiability matrix are negligible compared with simulation and optimization.

Future work will evaluate the approach on: (i) higher-dimensional industrial models with stiff dynamics (e.g. incorporation of dissolved oxygen tension) and multi-rate measurements; (ii) unmeasurable disturbances in addition to PMM; (iii) scenarios with clustered outputs exhibiting correlated identifiability patterns (e.g. biomass and substrates); (iv) practical identifiability analysis under limited data and

high noise (e.g. FIM analysis); and (v) validation and application on real experimental bioprocess data. Integration with adaptive online experimental design and online state estimation (e.g., Extended Kalman Filter) is a promising direction for closed-loop process optimization.

REFERENCES

- [1] D. Dochain and M. Perrier, “Bioprocess control,” in *Bioreaction engineering: modeling and control*. Springer, 2000.
- [2] R. King, “Control of biotechnological processes,” in *Encyclopedia of Systems and Control*. Springer, 2021.
- [3] S. Hellmann, J. Frontzek, D. M. Zarate, T. Wilms, K. Koch, S. Knorn, S. Streif, and S. Weinrich, “Multi-stage model predictive control of agricultural anaerobic digestion plant with uncertain substrate characterization,” *Bioresource Technology*, 2025.
- [4] M. F. de Dios, Á. M. González-Rueda, J. R. Banga, J. González-Díaz, and D. R. Penas, “Parameter estimation in odes: assessing the potential of local and global solvers: Mf de dios et al.” *Optimization and Engineering*, 2025.
- [5] C. G. Moles, P. Mendes, and J. R. Banga, “Parameter estimation in biochemical pathways: a comparison of global optimization methods,” *Genome research*, vol. 13, no. 11, 2003.
- [6] A. F. Villaverde, F. Fröhlich, D. Weindl, J. Hasenauer, and J. R. Banga, “Benchmarking optimization methods for parameter estimation in large kinetic models,” *Bioinformatics*, vol. 35, no. 5, 2019.
- [7] T. Barz, A. Sommer, T. Wilms, P. Neubauer, and M. N. C. Bournazou, “Adaptive optimal operation of a parallel robotic liquid handling station,” *IFAC-PapersOnLine*, vol. 51, no. 2, 2018.
- [8] M. McKay and W. Conover, “Rj beckman a comparison of three methods for selecting values of input variables in the analysis of output from a computer code,” *Technometrics*, vol. 21, no. 2, 1979.
- [9] T. A. El-Mihoub, A. A. Hopgood, L. Nolle, and A. Battersby, “Hybrid genetic algorithms: A review,” *Eng. Lett.*, vol. 13, no. 2, 2006.
- [10] A. F. Villaverde, A. Barreiro, and A. Papachristodoulou, “Structural identifiability of dynamic systems biology models,” *PLOS Computational Biology*, vol. 12, no. 10, 2016.
- [11] A. Tuveri, C. S. M. Nakama, J. Matias, H. E. Holck, J. Jaeschke, L. Imsland, and N. Bar, “A regularized moving horizon estimator for combined state and parameter estimation in a bioprocess experimental application,” *Computers & Chemical Engineering*, vol. 172, 2023.
- [12] M. Burth, G. C. Verghese, and M. Velez-Reyes, “Subset selection for improved parameter estimation in on-line identification of a synchronous generator,” *IEEE Transactions on Power Systems*, vol. 14, no. 1, 1999.
- [13] T. Barz, S. Körkel, G. Wozny, et al., “Nonlinear ill-posed problem analysis in model-based parameter estimation and experimental design,” *Computers & Chemical Engineering*, vol. 77, 2015.
- [14] S. Díaz-Seoane, X. Rey-Barreiro, and A. F. Villaverde, “STRIKE-GOLDD 4.0: user-friendly, efficient analysis of structural identifiability and observability,” *Bioinformatics*, vol. 39, no. 1, 2023.
- [15] A. M. Gleixner, T. Berthold, B. Müller, and S. Weltge, “Three enhancements for optimization-based bound tightening,” *Journal of Global Optimization*, vol. 67, no. 4, 2017.
- [16] E. Walter and H. Piet-Lahanier, “Estimation of parameter bounds from bounded-error data: a survey,” *Mathematics and Computers in simulation*, vol. 32, no. 5-6, pp. 449–468, 1990.
- [17] M. Milanese and A. Vicino, “Optimal estimation theory for dynamic systems with set membership uncertainty: An overview,” *Automatica*, vol. 27, no. 6, pp. 997–1009, 1991.
- [18] H. White, “Maximum likelihood estimation of misspecified models,” *Econometrica: Journal of the econometric society*, 1982.
- [19] E. Walter, L. Pronzato, and J. Norton, *Identification of parametric models from experimental data*. Great Britain, UK: Springer, 1997.
- [20] M. Kawohl, T. Heine, and R. King, “A new approach for robust model predictive control of biological production processes,” *Chemical Engineering Science*, vol. 62, no. 18-20, 2007.
- [21] D. Krämer and R. King, “On-line monitoring of substrates and biomass using near-infrared spectroscopy and model-based state estimation for enzyme production by *s. cerevisiae*,” *IFAC-PapersOnLine*, vol. 49, no. 7, 2016.
- [22] D. Vrajitoru, “Large population or many generations for genetic algorithms? implications in information retrieval,” in *Soft computing in information retrieval: Techniques and applications*. Springer, 2000.