



Simulation-Driven Parameter Estimation with a Focus on Control and Privacy

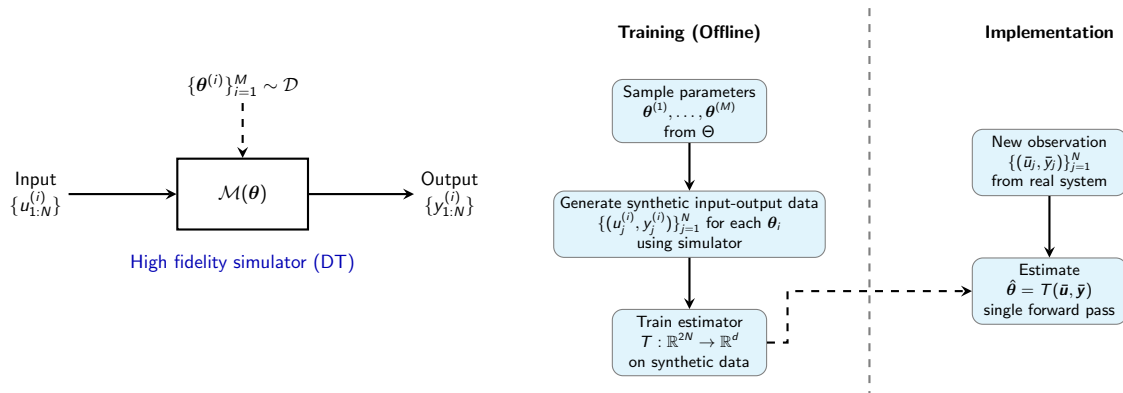
Braghadeesh Lakshminarayanan

Doctoral Thesis Defence
Department of Decision and Control Systems
KTH Royal Institute of Technology, Sweden

May 8, 2026

- ➊ **Simulation-Driven Estimation** — motivation and the Two-Stage (TS) estimator
- ➋ **Part 1: Analysis of the TS Estimator**
- ➌ **Part 2: Fine-Tuning a Simulation-Driven Estimator**
- ➍ **Part 3: Applications in Control**
- ➎ **Part 4: Privacy in Point Estimation**

Simulation-Driven Estimation



- S. Garatti, S. Bittanti. "A new paradigm for parameter estimation in system modeling". *Int. J. Adapt. Control Sig. Proc.*, 2013.
- T. Diskin et al. "Learning to estimate without bias". *Trans. Sig. Proc.*, 2022.
- A. Ghosh et al. "DeepBayes—An estimator for parameter estimation in stochastic nonlinear dynamical models". *Automatica*, 2024.

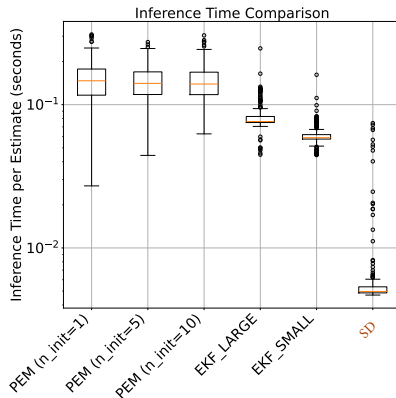
Why Simulation-Driven Estimation?

Traditional methods

- PEM, EKF, MLE, ...
- High inference time

Simulation-driven (TS)

- Train offline using synthetic data
- Inference = single forward pass
- **Low inference time**

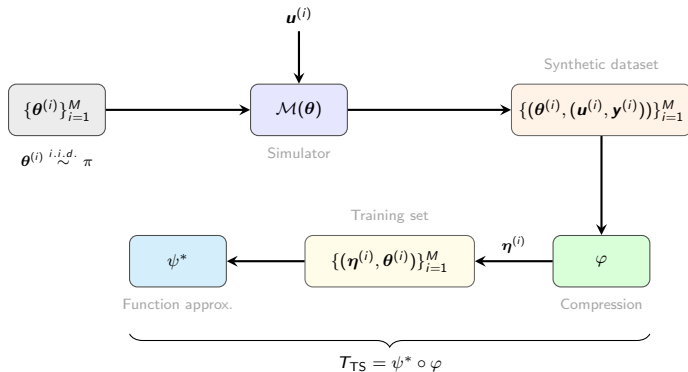


Part I

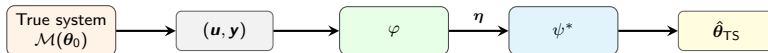
Analysis of a Simulation-Driven Estimator

The Two-Stage (TS) Estimator

Training phase:



Implementation (single forward pass):



S. Garatti, S. Bittanti. "A new paradigm for parameter estimation in system modeling" *Int. J. Adapt. Control Sig. Proc.*, 2013.

Chapter 3

Statistical Framework of the TS Estimator

B.Lakshminarayanan, C.R. Rojas. “A Statistical Decision-Theoretical Perspective on the Two-Stage Approach to Parameter Estimation” *IEEE CDC*, pages 5369–5374, 2022.

Statistical Frameworks of the TS Estimator

Bayes TS

Minimize the average risk:

$$\psi^* = \arg \min_{\psi \in \mathcal{G}} \frac{1}{M} \sum_{i=1}^M \ell(\boldsymbol{\theta}^{(i)}, \psi(\boldsymbol{\eta}^{(i)}))$$

- $\boldsymbol{\theta}^{(i)} \sim \pi$ (prior)
- Standard supervised learning

Minimax TS

Minimize the worst-case risk:

$$\psi^* = \arg \min_{\psi \in \mathcal{G}} \max_{i=1, \dots, M} \ell(\boldsymbol{\theta}^{(i)}, \psi(\boldsymbol{\eta}^{(i)}))$$

- $\boldsymbol{\theta}^{(i)} \sim s$ (proposal)
- Robust: no prior needed

Both share the same structure $T_{\text{TS}} = \psi^* \circ \varphi$; they differ only in how ψ^* is trained.

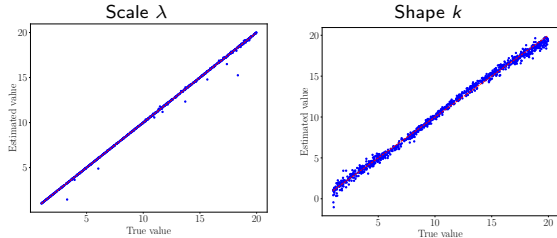
Minimax via epigraph formulation:

$$\min_{\boldsymbol{\beta}, t} \quad t \quad \text{s.t.} \quad \ell(\boldsymbol{\theta}^{(i)}, \boldsymbol{\beta}^\top J(\boldsymbol{\eta}^{(i)})) \leq t, \quad i = 1, \dots, M$$

Numerical Example

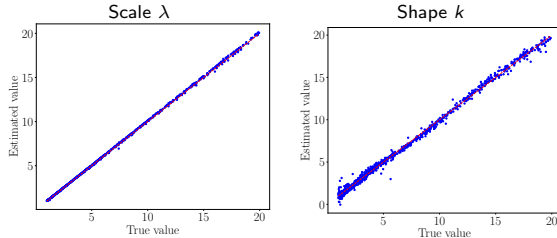
Model: $y_i \overset{i.i.d.}{\sim} \text{Weibull}(\lambda, k)$, λ : scale, k : shape

Bayes TS Estimator



performance depends on prior π .

Minimax TS Estimator



Robust (no prior); higher variance $\rightarrow$ motivates improved minimax (Ch. 4).

Chapter 4

Improved Minimax TS via Gradient Boosting

B. Lakshminarayanan and C. R. Rojas. “Minimax Two-Stage Gradient Boosting for Parameter Estimation”. *IEEE CDC*, pages 1189–1194, 2023.

Improved Minimax TS via Gradient Boosting

Minimax objective:

$$\min_{\psi} \max_{i=1, \dots, M} \ell(\boldsymbol{\theta}^{(i)}, \psi(\boldsymbol{\eta}^{(i)}))$$

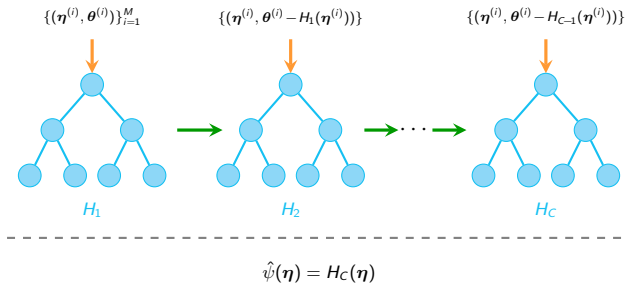
Problem: max is non-smooth.

↓ *soft-max approximation*

TSGBM objective:

$$\min_{\psi} \frac{1}{K} \log \left(\sum_{i=1}^M e^{K \ell(\boldsymbol{\theta}^{(i)}, \psi(\boldsymbol{\eta}^{(i)}))} \right)$$

GBM: iterative residual fitting

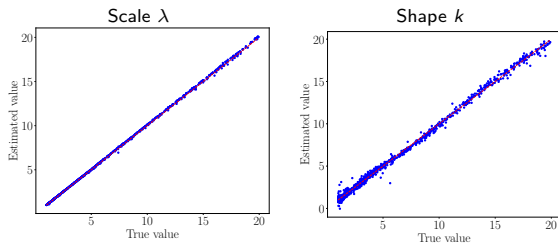


Smooth $\rightarrow$ usable as GBM loss | $K \rightarrow \infty$ recovers exact minimax

Numerical Example

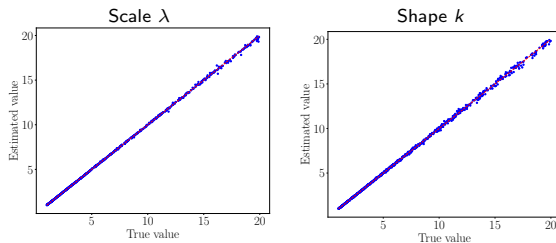
Model: $y_i \overset{i.i.d.}{\sim} \text{Weibull}(\lambda, k)$, λ : scale, k : shape

Minimax TS (Ch. 3)



High variance for k ; slow (epigraph + CVXPY per sample).

TSGBM (Ch. 4)



Low variance for k improved up to $10\times$; much faster via soft-max + LightGBM.

Chapter 5

Asymptotic Analysis of the TS Estimator

B. Lakshminarayanan and C. R. Rojas. "On Asymptotic Analysis of the Two-Stage Approach: Towards Data-Driven Parameter Estimation". Submitted to *Automatica*, 2024.

Asymptotic Analysis of the TS Estimator

Challenge: Unlike classical estimators (one limit: $N \rightarrow \infty$), the TS estimator involves *three* limits:

<i>Training:</i>	$M \rightarrow \infty$	number of parameter value draws
	$N \rightarrow \infty$	observation length per draw
<i>Implementation:</i>	$N \rightarrow \infty$	observation length at deployment

Result: $\hat{\theta}_{\text{TS}}: \mathbb{R}^N \times \mathbb{R}^M \rightarrow \mathcal{G} \circ h$ converges to $\bar{\theta}$ as $N, M \rightarrow \infty$

① **Consistency:** $\hat{\theta}_0 = \bar{\theta}\left((y_1^{(0)}, \dots, y_N^{(0)})^\top\right) \xrightarrow{\text{w.p.1}} \theta_0$ as $N \rightarrow \infty$

② **Asymptotic Normality:** $\sqrt{N}(\hat{\theta}_0 - \theta_0) \xrightarrow{d} \mathcal{N}(\mathbf{0}, \Sigma)$ as $N \rightarrow \infty$

Key insight: TS is asymptotically consistent and normal under simplified assumptions.

Chapter 6

Non-Asymptotic Analysis of TS

B. Lakshminarayanan and C. R. Rojas. “Finite Sample Analysis of a Data-Driven Estimator”. Submitted to *IEEE SPL*, 2026.

Finite-Sample Bounds for the TS Estimator

Goal: Non-asymptotic guarantee: $|\hat{\theta} - \theta_0| \leq \varepsilon$ with high probability.

Key idea: Use an IV regressor (two independent sequences per $\theta^{(i)}$) and decompose the error:

$$\hat{\beta}_{\text{IV}} - \beta^* = \mathbf{B}_{N,M}^{-1} \mathbf{S}_{N,M} \beta^* + \mathbf{B}_{N,M}^{-1} \mathbf{e}_{N,M}$$

where $\mathbf{B}_{N,M} \rightarrow \mathbf{B}_* := \mathbb{E}_\theta[g_*(\theta)g_*^\top(\theta)]$ as $M, N \rightarrow \infty$.

Result: For any $\varepsilon > 0$,

$$|\hat{\theta} - \theta_0| \leq \zeta + \varepsilon \|\beta^*\|_2 + \alpha \frac{\varepsilon \|\beta^*\|_2 + \alpha \zeta}{2 \sigma_{\min}(\mathbf{B}_*)}$$

except with probability at most

$$4p^2 \exp\left[-\frac{M\sigma_{\min}^2(\mathbf{B}_*)}{128p^2\alpha^4}\right] + 2p^2 \exp\left[-\frac{M\varepsilon^2}{32\alpha^4}\right] + \delta(N, \varepsilon)$$

Here ζ is the approximation bias: $\theta = \beta^{*\top} g_*(\theta) + r(\theta)$ with $|r(\theta)| \leq \zeta$.

Probability bound decays exponentially in M .

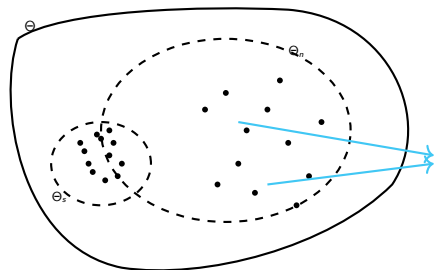
Chapter 7

Sampling from Jeffreys Prior

Y. Shi, **B. Lakshminarayanan**, and C. R. Rojas. “A Metropolis-Adjusted Langevin Algorithm for Sampling Jeffreys Prior”. *IEEE CDC*, pages 2359–2364, 2025.

Diverse Training via Jeffreys Prior Sampling

Question: How to sample $\theta^{(1)}, \dots, \theta^{(M)}$ to improve TS performance?



Θ_s : sensitive region (high FIM), needs *more* samples.

Θ_n : non-sensitive region (low FIM).

Uniform sampling ignores this geometry!

Proposed: Sample from Jeffreys prior

$$\pi_J(\theta) \propto \sqrt{\det \mathbf{I}_\theta}$$

- Concentrates tasks in Θ_s where data is most sensitive to θ

Sampling via Langevin SDE + MALA:

$$d\theta_t = -\nabla_\theta V(\theta_t) dt + \sqrt{2} d\mathbf{W}_t$$

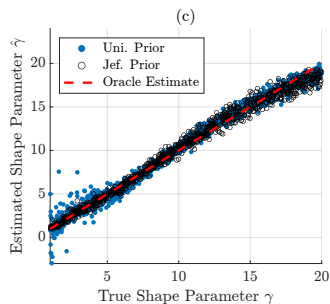
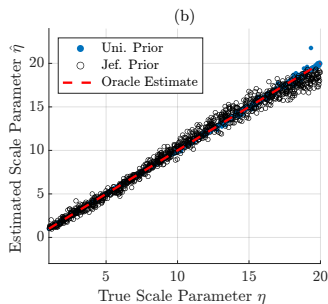
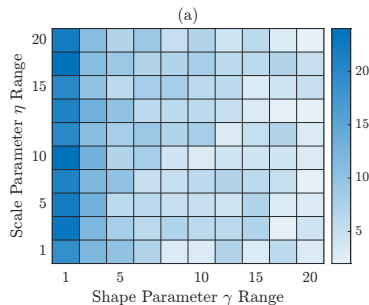
$$V(\theta) = -\frac{1}{2} \ln \det \mathbf{I}_\theta \Rightarrow \text{stationary dist.} = \pi_J(\theta)$$

When FIM is intractable (e.g. NLSS):

- Approximate $\mathbf{I}_\theta$ via particle filter
- MH accept/reject corrects errors

Numerical Example

Model: $y_i \stackrel{i.i.d.}{\sim} \text{Weibull}(\lambda, \kappa), \quad \lambda \in [1, 20], \quad \kappa \in [1, 20]$



(a) Jeffreys prior concentrates samples at low κ (high FIM); (b) Both priors comparable for λ ; (c) Jeffreys prior reduces variance and bias for $\kappa < 5$.

Part 2

Fine-Tuning a Simulation-Driven Estimator

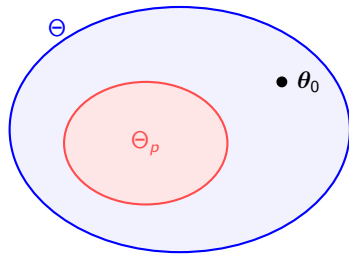
The Out-of-Distribution Problem

TS is pretrained on $\Theta_p \subseteq \Theta$.

If $\theta_0 \notin \Theta_p$:

- ψ has never seen features near θ_0
- Estimates would degrade significantly

Goal: Adapt the pretrained estimator using new observations.



Chapter 8

Supervised Fine-Tuning of a Simulation-Driven Estimator

B. Lakshminarayanan, M. A. Guerrero, and C. R. Rojas. “Fine-Tuning a Simulation-Driven Estimator”. *IEEE L-CSS*, vol. 9, pages 2975–2980, 2025.

Fine-Tuning TS: Supervised Approach

Architecture:

Trunk $\hat{\mathbf{w}}_T$: frozen

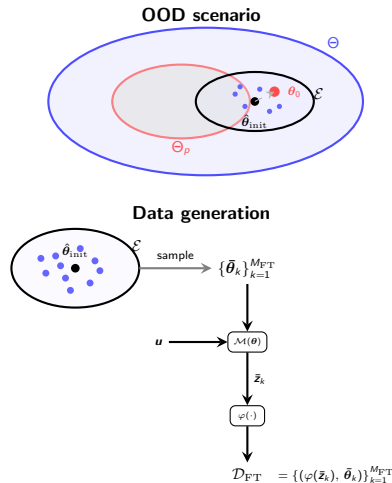
Heads $\hat{\mathbf{w}}_H$: retrained on $\mathcal{D}_{\text{FT}}$

Steps:

- 1 **Detect** OOD via whitened feature-space test
- 2 **Sample** cloud around $\hat{\theta}_{\text{init}}$, simulate DT $\rightarrow \mathcal{D}_{\text{FT}}$
- 3 **Retrain** heads with a weighted loss

$$T_{\text{TS}}^{\text{FT}} = \psi(\varphi(\cdot); \hat{\mathbf{w}}_T, \hat{\mathbf{w}}_H^{\text{FT}})$$

Only *forward simulation* needed.



Chapter 9

Unsupervised Fine-Tuning of a Simulation-Driven Estimator

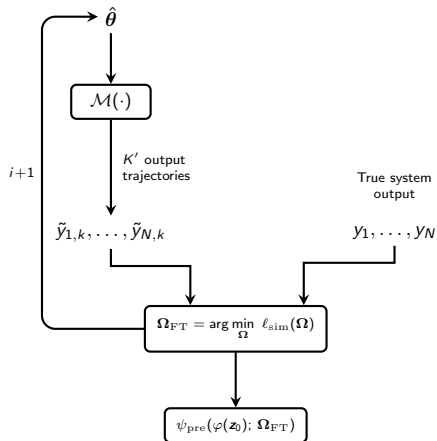
B. Lakshminarayanan and C. R. Rojas. “Trajectory-Level Self-Supervision for Simulation Driven Estimators”.
IFAC Journal of Systems and Control, vol. 35, page 100374, 2026.

Unsupervised Fine-Tuning via Trajectory Matching

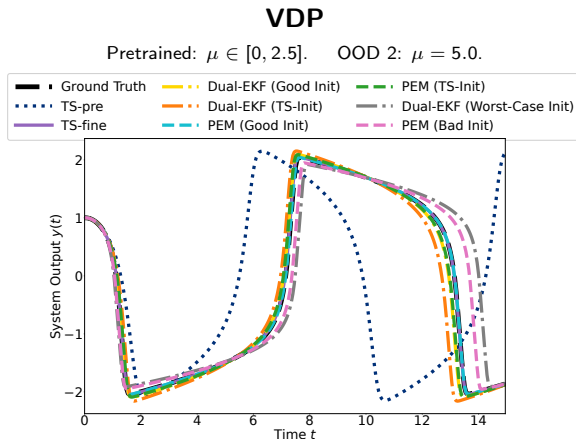
Idea: Directly minimize simulation error: no synthetic dataset needed for retraining.

$$\ell_{\text{sim}}(\Omega) = \frac{1}{K'N} \sum_{k=1}^{K'} \sum_{t=1}^N (\tilde{y}_{t,k} - y_t)^2$$

- Differentiable simulator needed
- Updates *all* weights Ω



Numerical example



- TS-FT closely tracks ground truth
- TS-pre and poorly initialized baselines diverge

TS-pre: < 2 ms. Fine-tuning overhead incurred *only* upon OOD detection.

Part 3

Applications in Control

Chapter 10: The Two-Stage Approach for Controller Tuning

F. Dettù, **B. Lakshminarayanan**, S. Formentin, and C. R. Rojas. “From Data to Control: A Two-Stage Simulation-Based Approach”. *ECC*, pages 3428–3433, 2024.

Chapter 11: Direct Learning of Controller Tuning Rules

B. Lakshminarayanan, F. Dettù, C. R. Rojas, and S. Formentin. “Inverse Supervised Learning of Controller Tuning Rules”. *Automatica*, vol. 178, p. 112356, 2025.

Inverse Supervised Learning for Controller Tuning



Parameter	Trajectory	Design rule
$\tilde{\theta}_1$	$\mathcal{D}_1^N = \{(u_k^{(1)}, y_k^{(1)})\}_{k=1}^N$	$C(\Psi_1, \mathcal{M}(\tilde{\theta}_1))$
$\tilde{\theta}_2$	$\mathcal{D}_2^N = \{(u_k^{(2)}, y_k^{(2)})\}_{k=1}^N$	$C(\Psi_2, \mathcal{M}(\tilde{\theta}_2))$
$\vdots$	$\vdots$	$\vdots$
$\tilde{\theta}_M$	$\mathcal{D}_M^N = \{(u_k^{(M)}, y_k^{(M)})\}_{k=1}^N$	$C(\Psi_M, \mathcal{M}(\tilde{\theta}_M))$

$$f^* = \arg \min_f \frac{1}{M} \sum_{i=1}^M \left\| \Psi_i - f\left(y_1^{(i)}, u_1^{(i)}, \dots, y_N^{(i)}, u_N^{(i)}\right) \right\|_2^2$$

TS-ISL: $f^* = \psi^* \circ \varphi$, φ : ARX compression, ψ^* : FFNN or GBM

Direct: f^* : CNN or WaveNet (end-to-end, no compression)

A Simple Case Study: Yaw-Rate Tracking

Objective: Achieve desired yaw-rate r under uncertain vehicle parameters.

Vehicle dynamics:

$$\dot{\beta} = -r - \frac{C_f \alpha_f}{M_{\text{veh}} v_x} - \frac{C_r \alpha_r}{M_{\text{veh}} v_x}$$
$$\dot{r} = -\frac{l_f C_f \alpha_f}{J_z} + \frac{l_r C_r \alpha_r}{J_z}$$

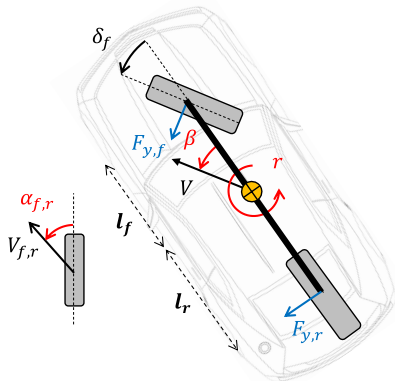
State: $\mathbf{x} = (\beta, r, \alpha_f, \alpha_r, z, \delta_f)^\top$

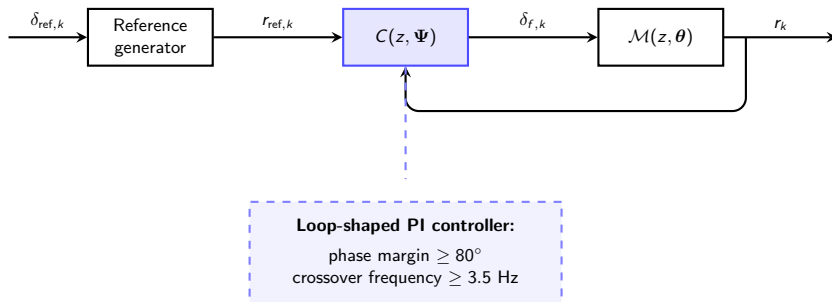
Input: $u = \delta_f^{\text{cmd}}$, Output: $y = r$

Uncertain parameters:

$$M_{\text{veh}} = M_0 + M_\delta, \quad M_\delta \in [0, 300] \text{ kg}$$

$$C_{f,r} = C_{f,r}^{\text{nom}} \mu_s, \quad \mu_s \in [0.3, 1.1]$$

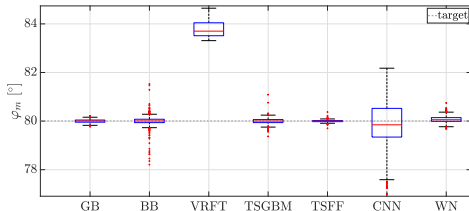




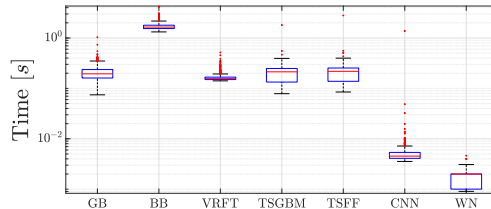
Controller parameters $\Psi = (K_p, K_i)$ are tuned to meet the loop-shaping specifications for each uncertain θ .

Results

Performance index



Computation (inference) time



GB: Grey Box; BB: Black Box; VRFT: Virtual Reference Feedback Tuning.

Our methods: TSFF, TSGBM, CNN, WN.

For the same control performance, re-tuning time is significantly lower for our approach.

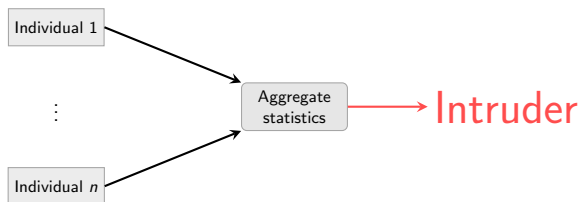
Part 4

Privacy in Point Estimation

Chapter 12: Differentially Private Point Estimation

B. Lakshminarayanan and C. R. Rojas. “A Unified Approach to Differentially Private Bayes Point Estimation”. *IFAC WC*, pp. 8375–8380, 2023.

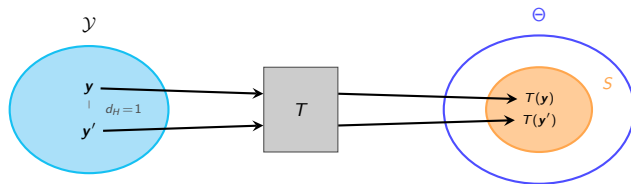
Need for Privacy in Point Estimates



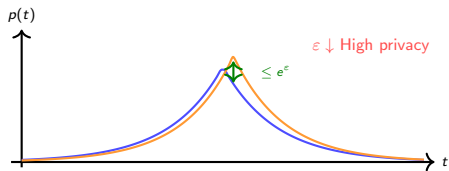
- Aggregate statistics: sample mean, sample covariance
- **Possible to infer an individual**¹
- Differential privacy (DP) overcomes this challenge
 - Still “optimality” in the sense of balancing accuracy-privacy trade-off is challenging!

¹ Homer, N. et al. “Resolving individuals contributing trace amounts of DNA to highly complex mixtures using high-density SNP genotyping microarrays”. *PLOS Genetics*, 2008.

Differential Privacy (DP)



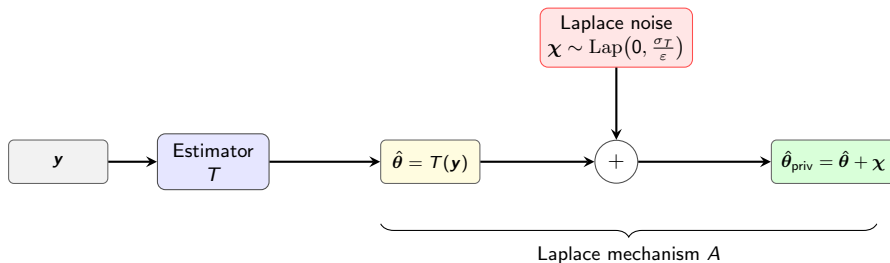
$$\mathbb{P}[T(y) \in S] \leq e^\epsilon \mathbb{P}[T(y') \in S]$$



Smaller $\epsilon \rightarrow$ output distributions become indistinguishable $\rightarrow$ stronger privacy.

Laplace Mechanism

Standard approach to achieve ϵ -DP: Add calibrated noise to the estimate.



$$\sigma_T = \sup_{d_H(y, y')=1} \|T(y) - T(y')\|_1 \quad (\ell_1\text{-sensitivity})$$

Two questions:

- 1 Is the Laplace noise addition *optimal* for Bayes estimation under DP?
- 2 What if T has no closed-form expression to compute σ_T ?

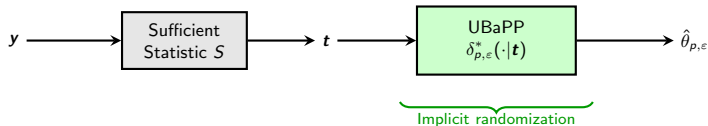
Unified Bayes Private Point (UBaPP) Estimator

Key idea: Integrate DP constraint directly into Bayes risk minimization.

$$\min_{\delta_{p,\varepsilon}} \int_{\Theta} \int_{\mathcal{T}} \int_{\Theta} \ell(\theta, \tilde{\theta}) \delta_{p,\varepsilon}(\tilde{\theta}|\mathbf{t}) \nu_{\theta}(\mathbf{t}) \pi(\theta) d\theta d\mathbf{t} d\tilde{\theta}$$

$$\text{s.t. } \delta_{p,\varepsilon}(\tilde{\theta}|S(\mathbf{y})) \leq e^{\varepsilon} \delta_{p,\varepsilon}(\tilde{\theta}|S(\mathbf{y}')), \quad \forall \tilde{\theta} \in \Theta, \quad \mathbf{y}, \mathbf{y}' \text{ s.t. } d_H(\mathbf{y}, \mathbf{y}') = 1 \quad \leftarrow \text{DP}$$

$$\int_{\Theta} \delta_{p,\varepsilon}(\tilde{\theta}|\mathbf{t}) d\tilde{\theta} = 1, \quad \delta_{p,\varepsilon}(\tilde{\theta}|\mathbf{t}) \geq 0$$



UBaPP is optimal by construction!

UBaPP: Finite Case

Finite case: When $|\Theta|, |\mathcal{T}|$ are finite, $\delta_{p,\varepsilon}$ is described by a matrix $\mathbf{P} \in \mathbb{R}^{|\Theta| \times |\mathcal{T}|}$.

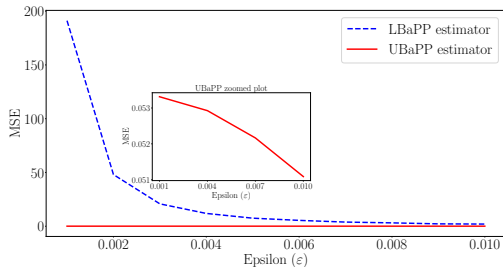
UBaPP reduces to a **linear program**:

$$\begin{aligned} \min_{\mathbf{P}} \quad & \text{tr}(\mathbf{V} \text{diag}(\boldsymbol{\pi}) \mathbf{L} \mathbf{P}) \\ \text{s.t.} \quad & P_{k,i} \leq e^\varepsilon P_{k,i'}, \quad \forall k, i, i' \text{ s.t. } d_H(\mathbf{y}_i, \mathbf{y}_{i'}) = 1 \\ & \mathbf{1}^\top \mathbf{P} = \mathbf{1}^\top, \quad \mathbf{P} \geq 0 \end{aligned}$$

$\mathbf{L}$: loss matrix, $\mathbf{V}$: likelihood matrix, $\boldsymbol{\pi}$: prior. Solvable via CVXPY.

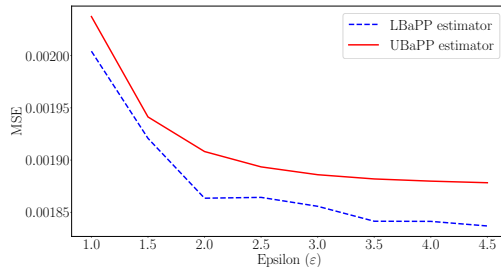
Privacy Results: MSE vs. ε

High privacy regime



LBaPP MSE $\rightarrow \infty$ as $\varepsilon \rightarrow 0$.
UBaPP MSE stays bounded.

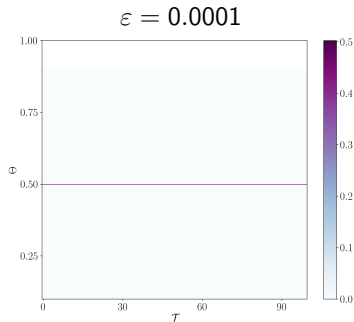
Low privacy regime



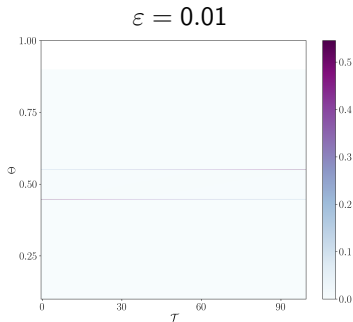
Both estimators perform similarly.
Small gap due to discretization of Θ .

High accuracy is achieved by our approach in the high privacy regime!

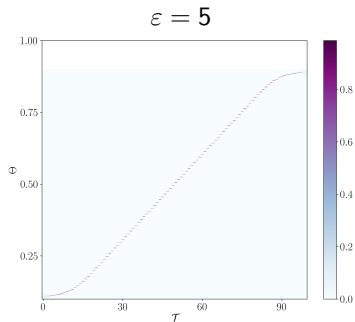
Privacy Results: Heat Maps of $\delta_{p,\epsilon}^*$



Deterministic at prior mean $\theta = 0.5$
for all t .



Implicit randomization: probability
spread over Θ .



Converges to non-private Bayes
(deterministic, data-dependent).

$\epsilon \rightarrow 0$: data-independent $\longrightarrow$ $\epsilon \rightarrow \infty$: data-dependent (non-private)

① Theoretical analysis of TS estimator

Statistical frameworks, consistency, asymptotic normality

② Fine-tuning for OOD adaptation

Supervised (OOD detection + retraining) and unsupervised (trajectory matching via differentiable simulator)

③ Data-driven controller tuning

Meta-learning design rules via TS-ISL, CNN, and WaveNet

④ Differentially private Bayes estimation

UBaPP: unified risk–privacy optimization, tractable LP for finite case

Thank You for Your Attention!

Braghadeesh Lakshminarayanan

`blak@kth.se`