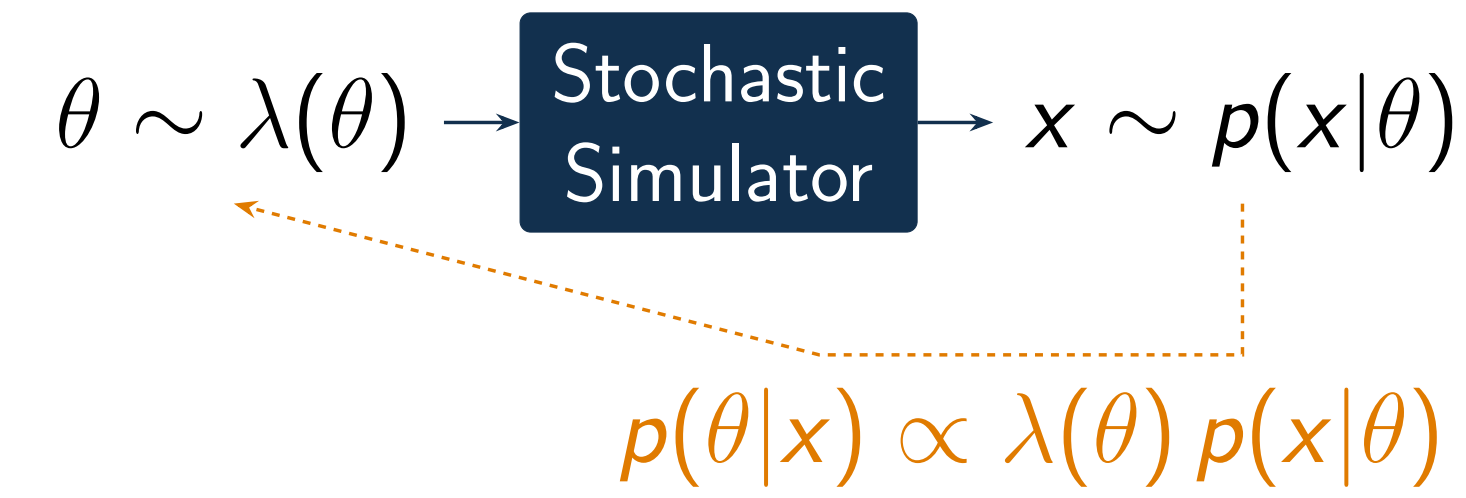


Background

Simulation-Based Inference. Scientific models are often defined as stochastic simulators:

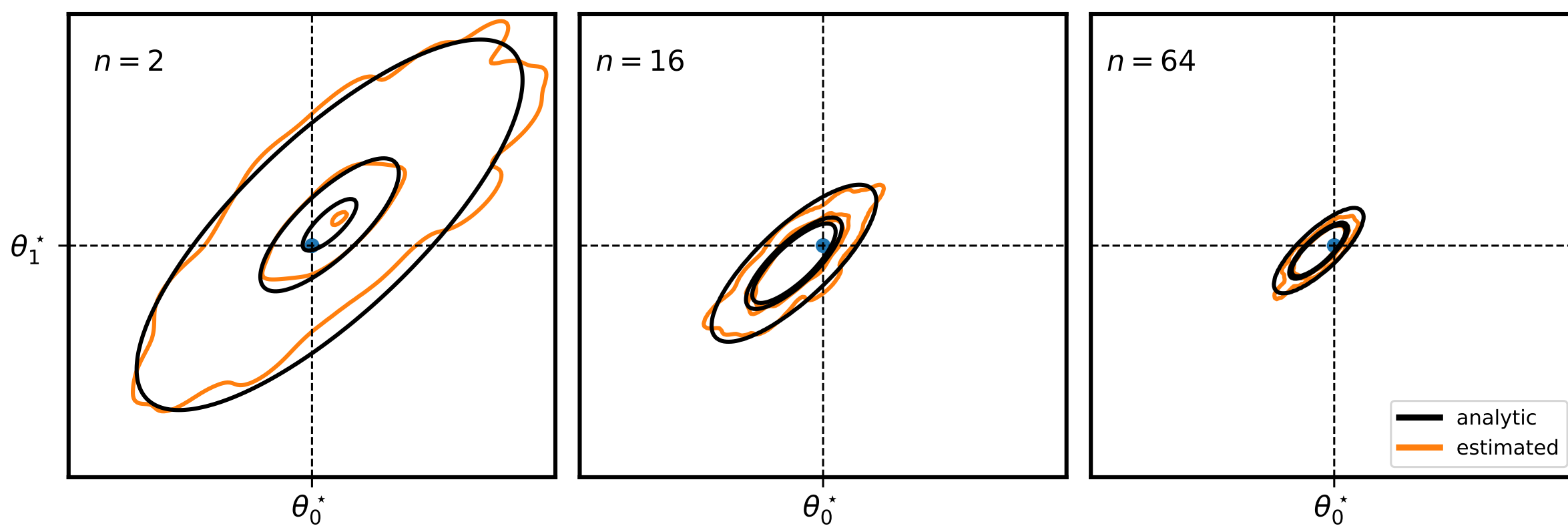


$p(x|\theta)$ is intractable, but we can simulate $(\theta_i, x_i) \sim p(\theta, x)$, and train a network *once* to approximate $p(\theta|x^*)$ for any x^* .

Tall Data Setting. For n i.i.d. observations $x_{1:n}^* \sim p(x|\theta^*)$, the tall posterior factorizes as

$$p(\theta | x_{1:n}^*) \propto \lambda(\theta)^{1-n} \prod_{j=1}^n p(\theta | x_j^*) \quad (1)$$

and concentrates around θ^* as $n \rightarrow \infty$ (Bernstein-von Mises).



Related Work

Method	Direct post.	Sim.-effic.	MCMC-free
Aug. NPE + DeepSet (Radev et al., 2020)	✓	✗	✓
Factorized NLE/NRE (Hermans et al., 2020)	✗	✓	✗
F-NPSE (Geffner et al., 2023)	✓	✓	✗
Ours (JAC / GAUSS)	✓	✓	✓

- Simulation-efficient methods require **MCMC** at inference.
- MCMC-free methods rely on **augmented datasets**.



Can we compose individual NPSE scores to directly sample the tall posterior — without MCMC?

Method

Setup. Forward diffusion: $q_{t|0}(\theta|\theta_0) = \mathcal{N}(\sqrt{\alpha_t}\theta_0, v_t\mathbf{I})$. Train NPSE via DSM on $(\theta_i, x_i) \sim p(\theta, x)$: $s_{\varphi}(\theta, x, t) \approx \nabla_{\theta} \log p_t(\theta | x)$. At inference, $s_{\varphi}(\theta, x_j^*, t) =: s_{\varphi,j,t}(\theta)$ gives the score at each x_j^* .

Key result — Tractable approximate diffused tall posterior score

$$\nabla_{\theta} \log p_t(\theta | x_{1:n}^*) \approx \Lambda^{-1} \left[\sum_{j=1}^n \Sigma_{t,j}^{-1} \underbrace{s_{\varphi,j,t}(\theta)}_{\text{NPSE}} + (1-n) \Sigma_{t,\lambda}^{-1} \underbrace{\nabla_{\theta} \log p_t^{\lambda}(\theta)}_{\text{prior score}} \right] \quad (2)$$

with $\Lambda = \sum_j \Sigma_{t,j}^{-1} + (1-n) \Sigma_{t,\lambda}^{-1}$ positive definite and $\Sigma_{t,j/\lambda}$ the backward diffusion covariances.

- ✓ Explicit approx. of the diffusion process — enable DDIM, no MCMC.
- ✓ Single amortized NPSE — no augmented dataset.

Derivation in 3 steps.

- Exact decomposition into individual scores (via (1) & Fisher's identity):

$$\begin{aligned} \nabla_{\theta} \log p_t(\theta | x_{1:n}^*) &= \nabla_{\theta} \log \int q_{t|0}(\theta | \theta_0) p(\theta_0 | x_{1:n}^*) d\theta_0 \\ &= (1-n) \nabla_{\theta} \log p_t^{\lambda}(\theta) + \sum_{j=1}^n \nabla_{\theta} \log p_t(\theta | x_j^*) + \underbrace{\nabla_{\theta} \log \mathcal{L}_{\lambda}(\theta, x_{1:n}^*)}_{\text{intractable}} \end{aligned}$$

where $\mathcal{L}_{\lambda}$ is defined by the backward kernels needed to run DDIM.

- Second-order Tweedie approximation of backward kernels:

$$\begin{aligned} \hat{q}_{0|t}(\theta_0 | \theta, x_j^*) &= \mathcal{N}(\mu_{t,j}(\theta), \Sigma_{t,j}(\theta)) \\ \hat{q}_{0|t}^{\lambda}(\theta_0 | \theta) &= \mathcal{N}(\mu_{t,\lambda}(\theta), \Sigma_{t,\lambda}(\theta)) \end{aligned} \Rightarrow (2) + F$$

with $\mu_{t,j}(\theta), \Sigma_{t,j}(\theta)$ in closed form (Boys et al., 2023), making (2) tractable up to a residual term F .

- Constant covariance assumption:

$$\nabla_{\theta} \Sigma_{t,j}(\theta) = \nabla_{\theta} \Sigma_{t,\lambda}(\theta) = 0 \Rightarrow F = 0 \Rightarrow \text{Key result (2)}$$

Algorithm

Input: NPSE s_{φ} , observations $x_{1:n}^*$

Sample $\theta_T \sim \mathcal{N}(0, \mathbf{I})$

for $t = T, \dots, 1$:

- Get scores $s_{\varphi,j,t}(\theta_t)$; prior score
- Estimate $\hat{\Sigma}_{t,j}$: **JAC** or **GAUSS**
- Compute s^* via (2)
- DDIM: $\theta_{t-1} \leftarrow \text{DDIM}(\theta_t, s^*)$

return $\theta_0 \sim p(\theta | x_{1:n}^*)$

Strategies for $\hat{\Sigma}_{t,j}$

JAC — theoretically exact

$$\hat{\Sigma}_{t,j}^{-1} = \frac{\alpha_t}{v_t} (\mathbf{I} + v_t \nabla_{\theta} s_{\varphi,j,t})^{-1}$$

via Jacobian of NN — expensive, unstable · constant via stop-grad

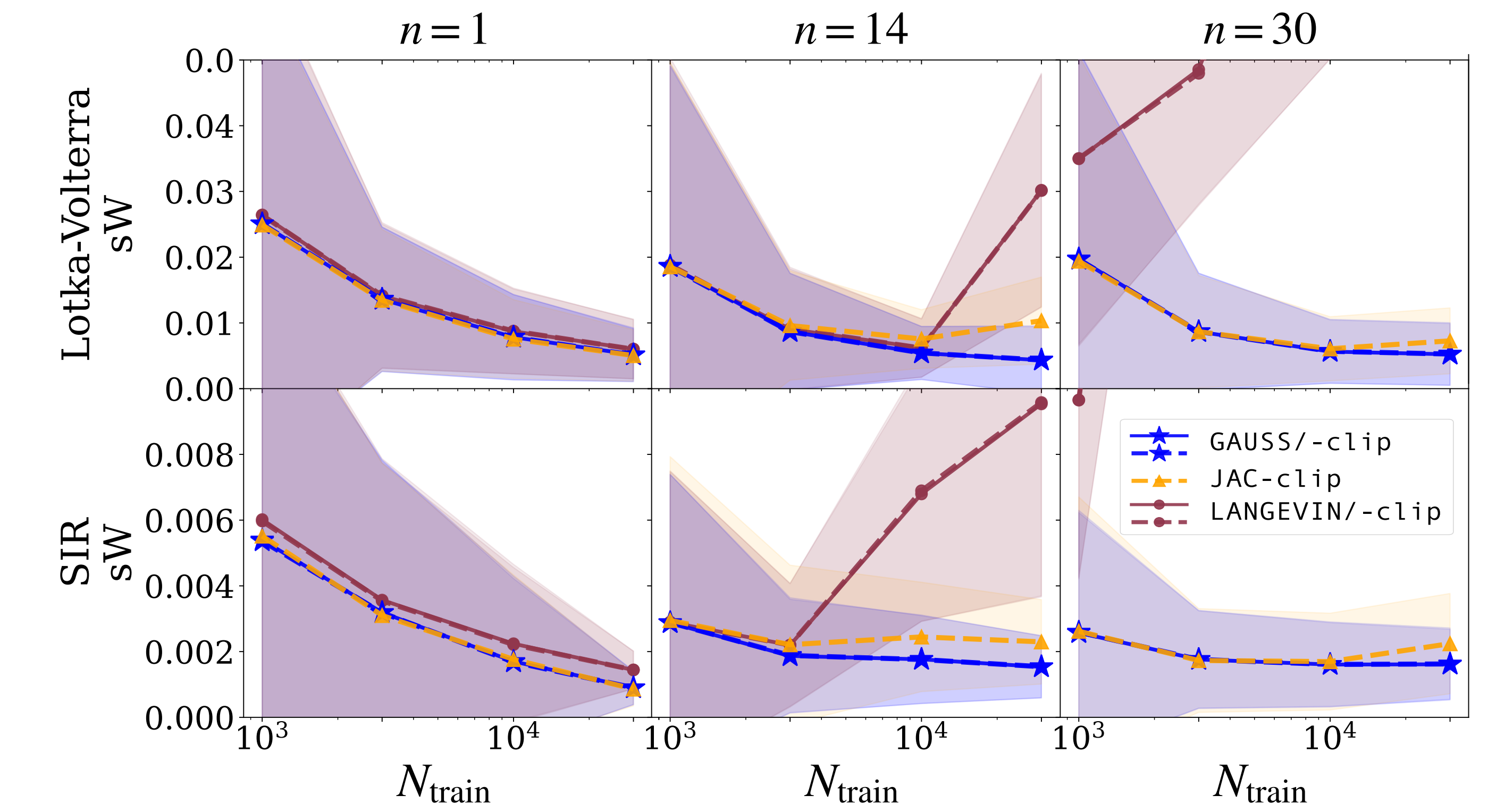
GAUSS — closed-form approximation

$$\hat{\Sigma}_{t,j}^{-1} = \hat{\Sigma}_{0,j}^{-1} + \frac{\alpha_t}{v_t} \mathbf{I}$$

emp. cov. of 100-step DDIM samples per obs. · constant by definition

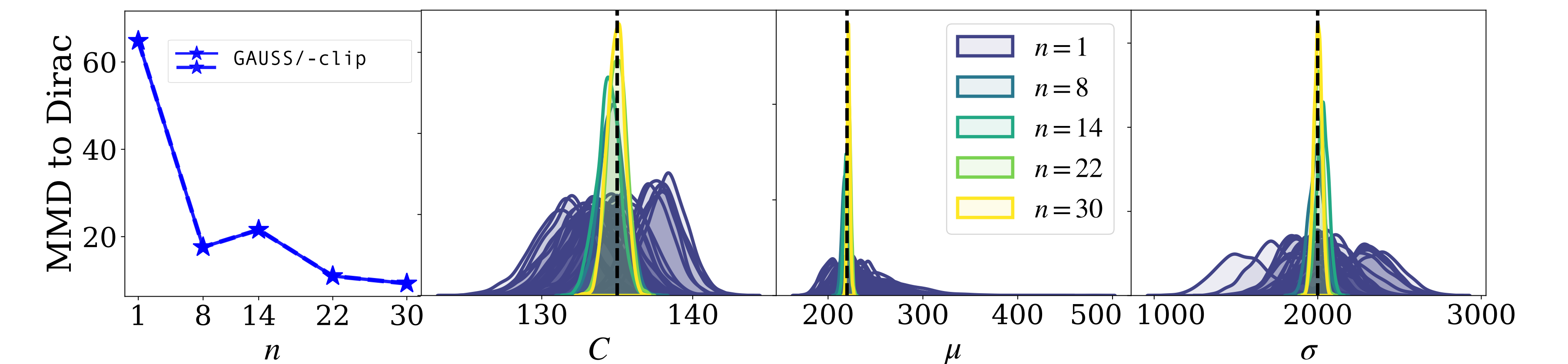
Results

SBI Benchmarks. SIR and Lotka-Volterra. Metric: Sliced Wasserstein vs. MCMC reference as a function of N_{train} and for increasing $n \in \{1, 14, 30\}$.



✗ **LANGEVIN** diverges for $n \geq 14$ ✓ **GAUSS** converges for all n ✓ **JAC**-clip competitive

Real-world application: Jansen-Rit Neural Mass Model. Stochastic simulator of EEG-like signals, with parameters $\theta = (C, \mu, \sigma, g)$ and fixed $g = 0$, $N_{\text{train}} = 50,000$.



✓ MMD decreases with n ✓ Marginals concentrate around true θ^*

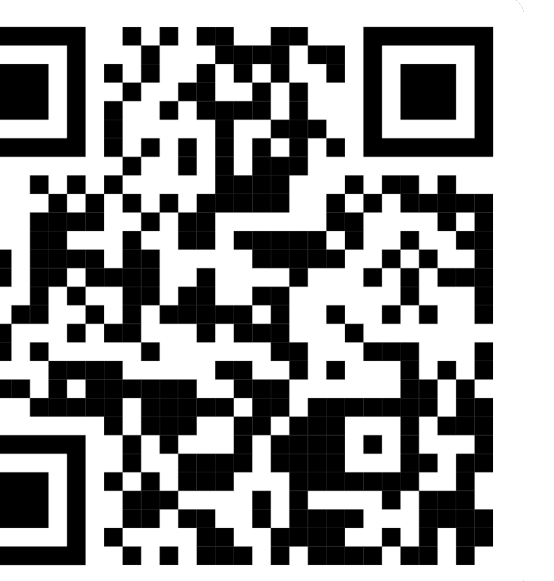
💡 indeterminacy for (μ, σ) in 4D $\Rightarrow$ coupling with $g \Rightarrow$ hierarchical model (Rodrigues et al., 2021)

Conclusion

- ✓ Tractable tall posterior score — one network, no MCMC, no augmented dataset
- ✓ **GAUSS**: more robust, accurate and $\sim 2.5\times$ faster than F-NPSE
- ✓ Validated on synthetic benchmarks and a real-world neuroscience simulator

Extensions & Perspectives.

- Partial factorization to mitigate error acc. for large n (Appendix)
- JAC** stabilization strategies (Gloeckler et al., 2025)
- Extension to large-scale hierarchical models (Arruda et al., 2025)
- Accurate and efficient model criticism...? (ongoing work)



Link to paper!