

Leveraging Self-Consistency for Data-Efficient Amortized Bayesian Inference



Marvin Schmitt
Stuttgart, GER



Desi Ivanova
Oxford, UK



Daniel Habermann
Dortmund, GER



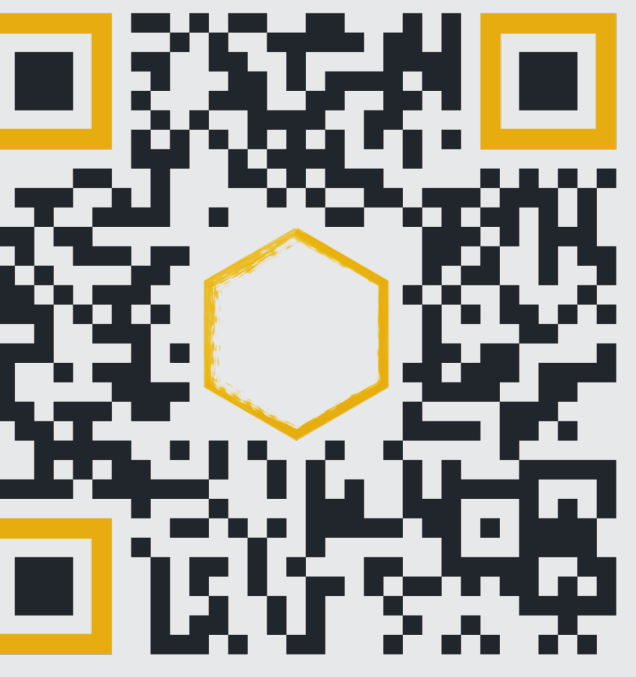
Ullrich Köthe
Heidelberg, GER



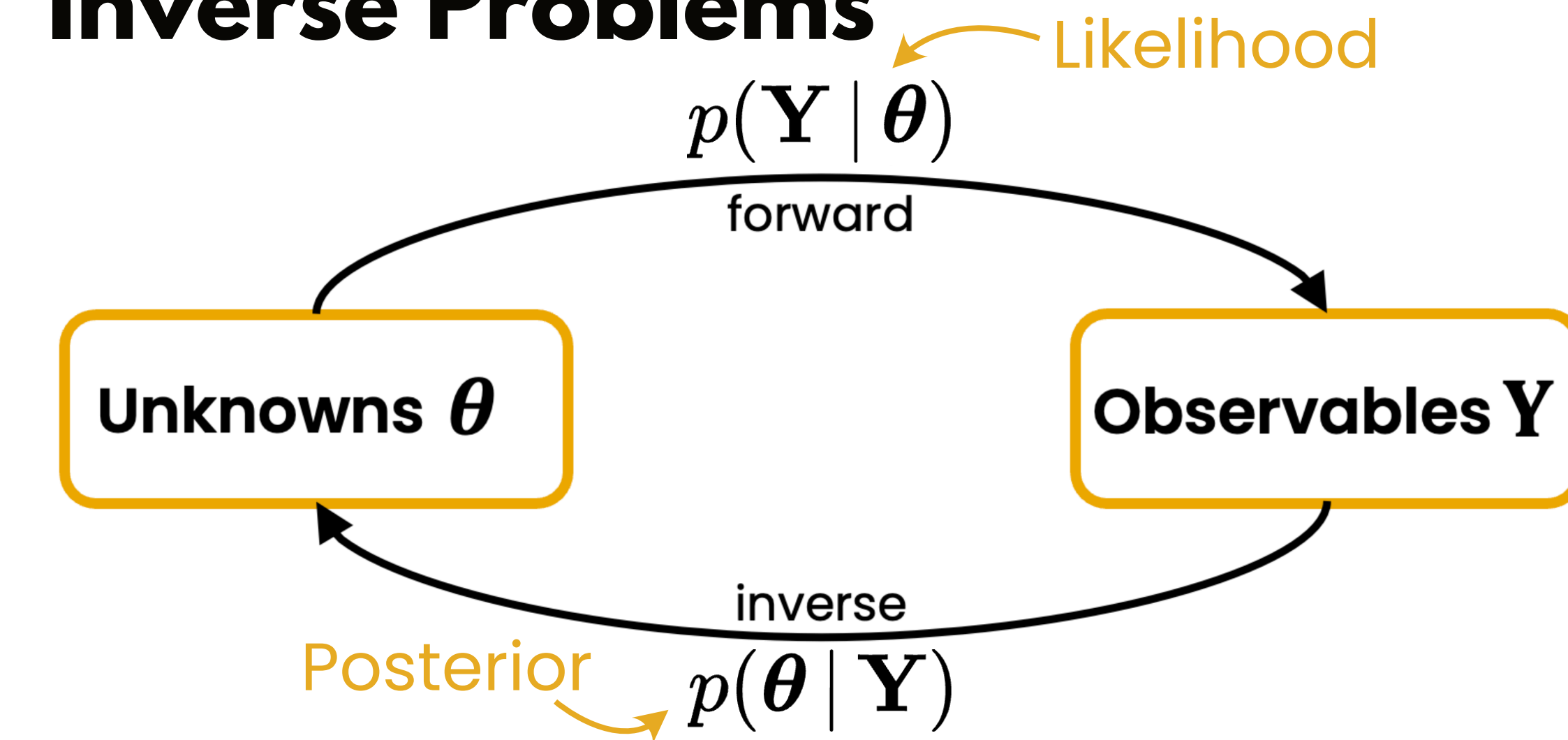
Paul Bürkner
Dortmund, GER



Stefan Radev
RPI, US



Inverse Problems

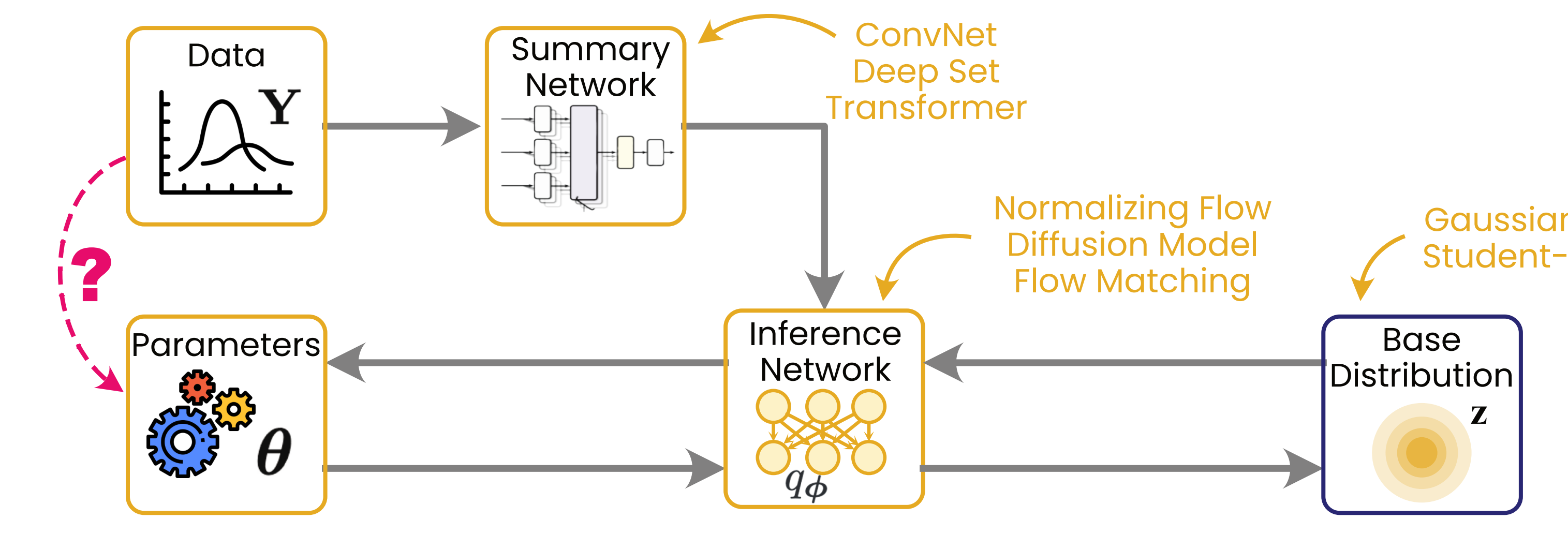


Examples: Gravitational Waves
Decision Making
Climate Modeling

Methods: MCMC, ABC, VI, Flows, ...

Amortized Bayesian Inference

Neural Network Training → Instant Posterior Sampling



Neural Network Training with Maximum Likelihood

$$\begin{aligned}\mathcal{L}_{\text{NPE}}(\phi) &= \mathbb{E}_{p(\mathbf{Y})} [\text{KL}(p(\theta | \mathbf{Y}) \| q_\phi(\theta | \mathbf{Y}))] \\ &= \mathbb{E}_{p(\theta, \mathbf{Y})} [-\log q_\phi(\theta | \mathbf{Y})] + \text{const}\end{aligned}$$

Bayes' Theorem

$$p(\theta | \mathbf{Y}) = \frac{p(\theta) p(\mathbf{Y} | \theta)}{p(\mathbf{Y})}$$

re-arrange

$$p(\mathbf{Y}) = \frac{p(\theta) p(\mathbf{Y} | \theta)}{p(\theta | \mathbf{Y})} = \text{const} \forall \theta \in \Theta$$

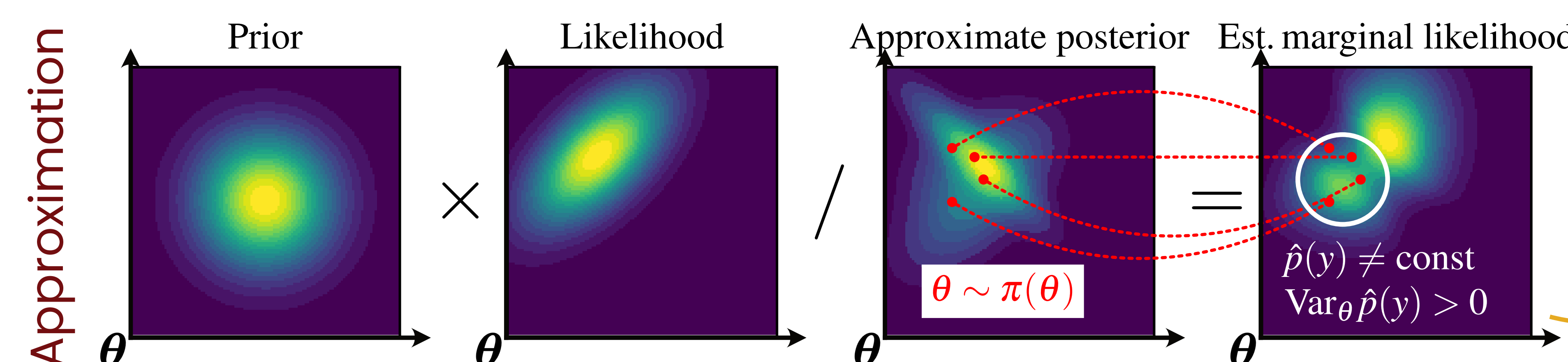
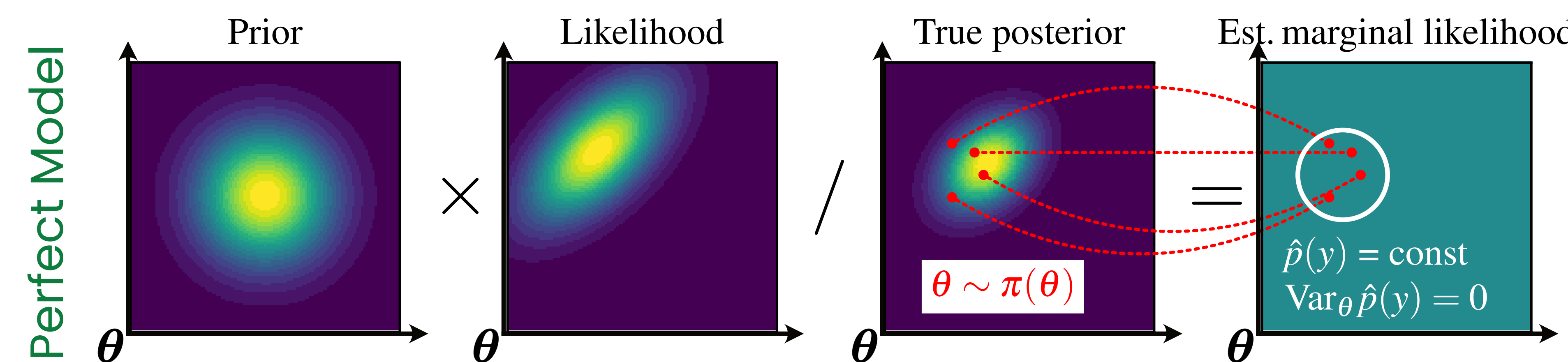
The marginal likelihood is a constant function in θ .

Probabilistic Symmetries

From here on, fix $\mathbf{Y}$. The marginal likelihood can be computed as a function of θ .

$$p(\mathbf{Y}) = \frac{p(\theta_1) p(\mathbf{Y} | \theta_1)}{p(\theta_1 | \mathbf{Y})} = \dots = \frac{p(\theta_K) p(\mathbf{Y} | \theta_K)}{p(\theta_K | \mathbf{Y})} \text{ for } \theta_1, \dots, \theta_K \in \Theta$$

It is **constant** in a **perfect joint model**, but not under **imperfect approximations**!



Frame as a loss!

Self-Consistency Loss

Neural Posterior Estimation

$$\begin{aligned}\mathcal{L}_{\text{SC-NPE}}(\phi) &= \mathbb{E}_{p(\mathbf{Y})} \left[\mathbb{E}_{p(\theta | \mathbf{Y})} [-\log q_\phi(\theta | \mathbf{Y})] \right. \\ &\quad \left. + \lambda \text{Var} \left(\log \frac{p(\theta) p(\mathbf{Y} | \theta)}{q_\phi(\theta | \mathbf{Y})} \right) \right]\end{aligned}$$

Self-Consistency Loss with Explicit Likelihood

Neural Posterior and Likelihood Estimation

$$\begin{aligned}\mathcal{L}_{\text{SC-NPLE}}(\phi) &= \mathbb{E}_{p(\mathbf{Y})} \left[\mathbb{E}_{p(\theta | \mathbf{Y})} \left[-\log q_\phi(\theta | \mathbf{Y}) \right. \right. \\ &\quad \left. \left. - \log q_\eta(\mathbf{Y} | \theta) \right] \right. \\ &\quad \left. + \lambda \text{Var} \left(\log \frac{p(\theta) q_\eta(\mathbf{Y} | \theta)}{q_\phi(\theta | \mathbf{Y})} \right) \right]\end{aligned}$$

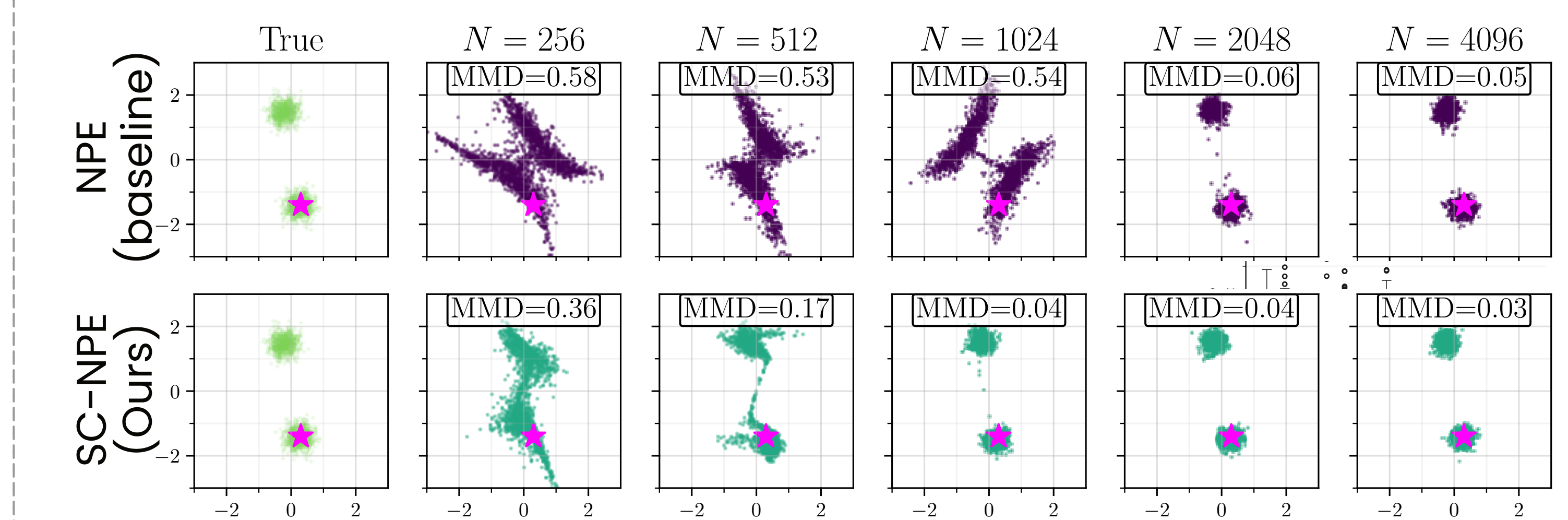
Self-Consistency Loss with Learned Likelihood

Experiments

Gaussian Mixture Model

Posterior estimation (NPE) using explicit likelihood

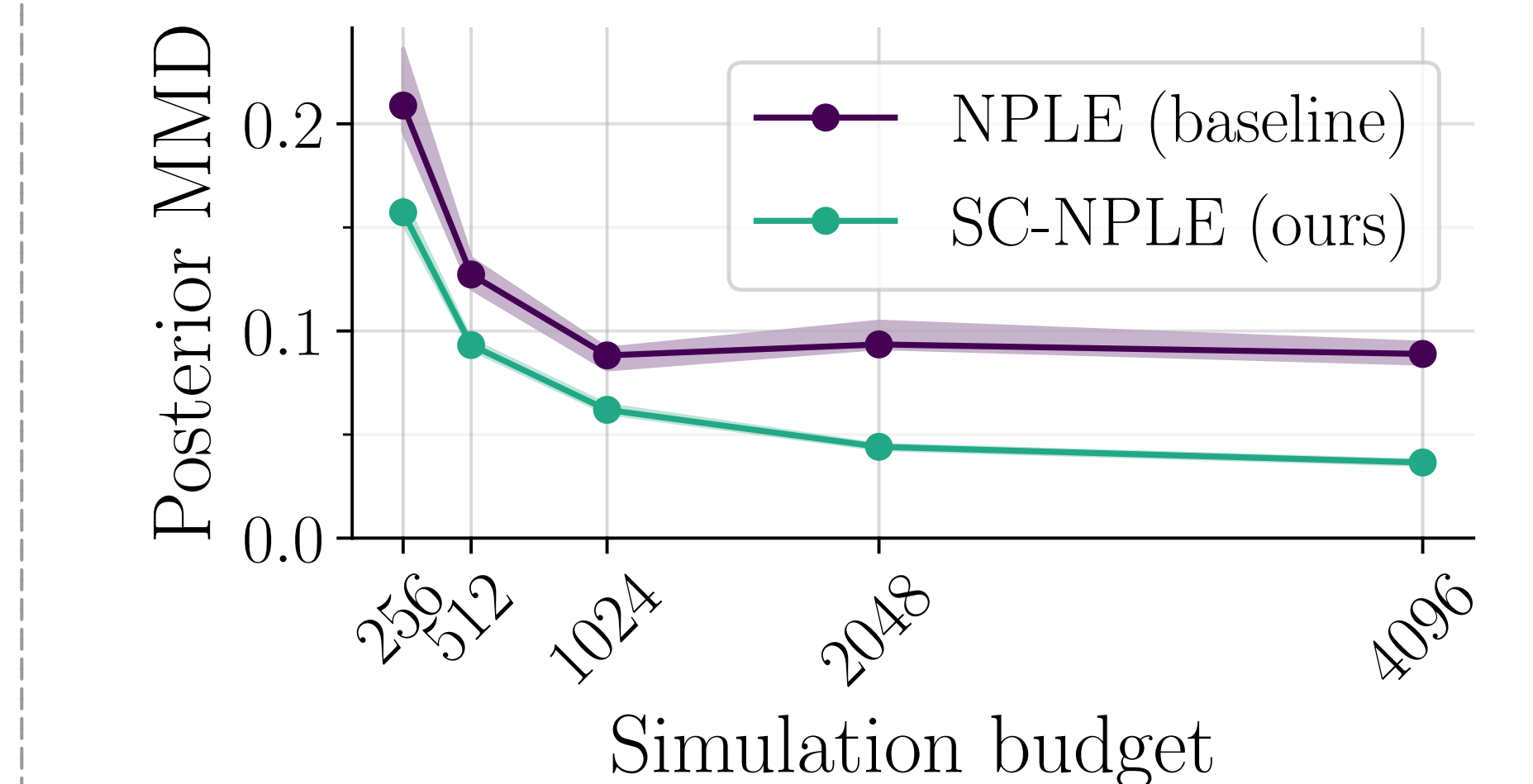
Better posterior samples



Two Moons

Posterior and likelihood estimation (NPLE)

Closer to true posterior (lower MMD)



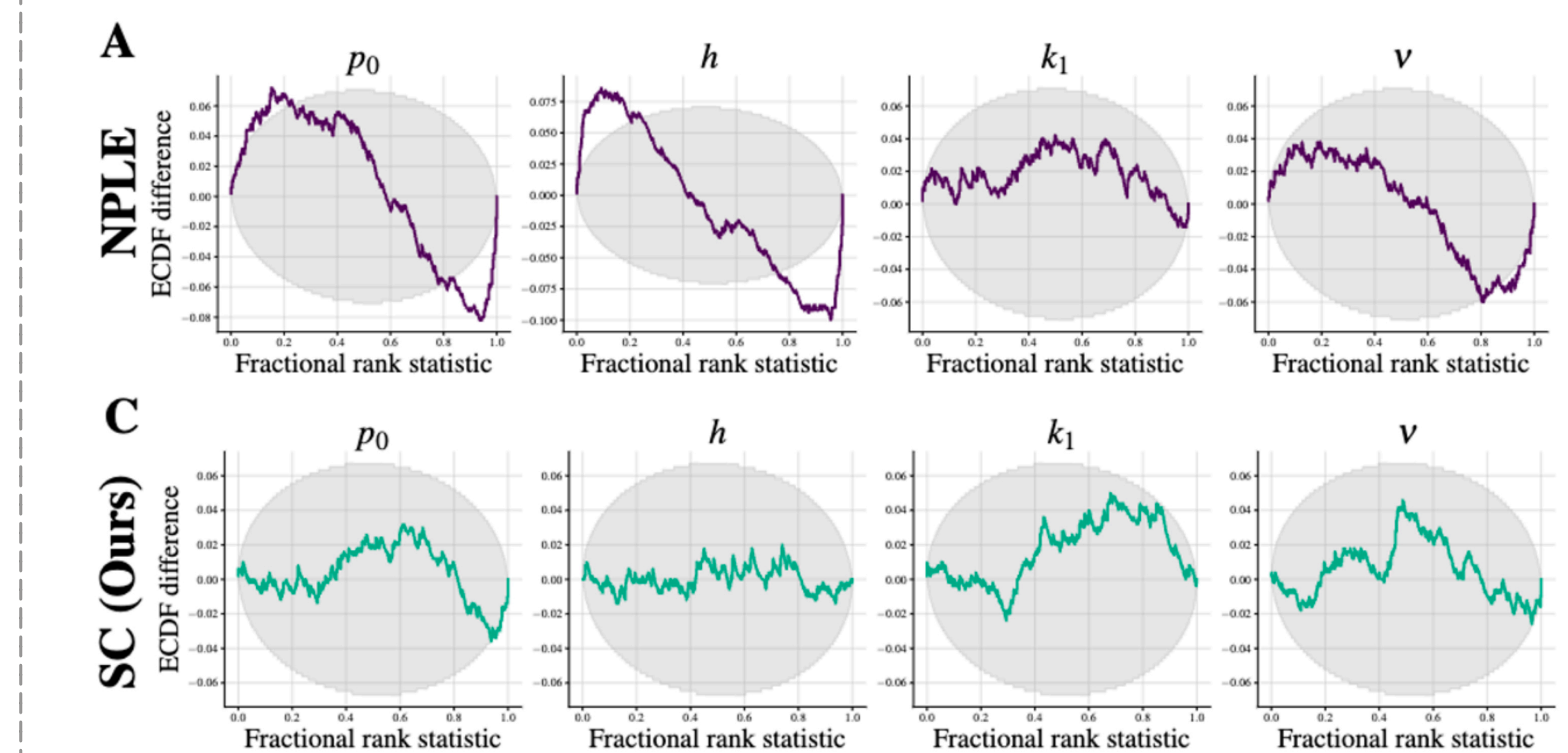
Other results

- Higher likelihood of ground-truths
- Sharper marginal likelihood estimates

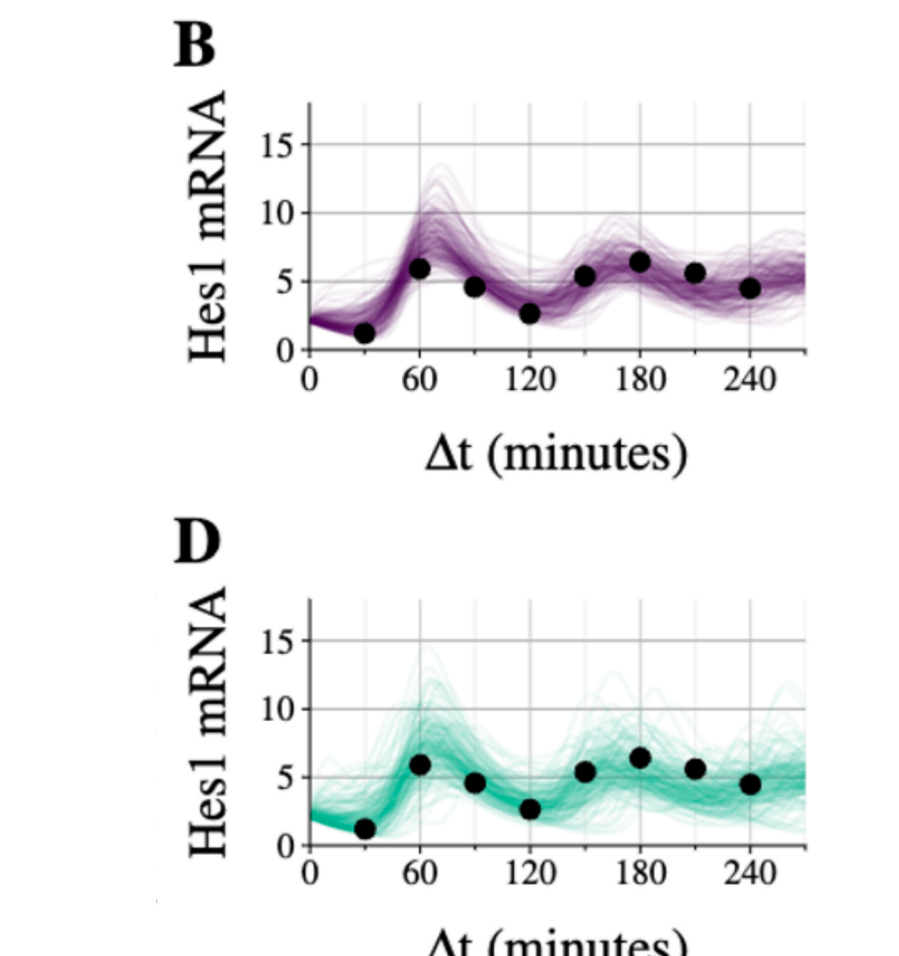
Hes1 Expression

Posterior and likelihood estimation (NPLE)

Better simulation-based calibration



Posterior predictive



— More experiments in the paper —