
Are gradients worth the effort? Comparing automatic differentiation and simulation-based inference for agent-based models

Timothy James Hitge*
AIMS South Africa

Arnau Quera-Bofarull
Macrocosm

Elizaveta Semenova
Imperial College London

Abstract

Agent-based models (ABMs) are flexible tools for simulating complex systems, but their calibration is difficult because their likelihoods are intractable and simulations are expensive. Two modern approaches tackle this challenge: automatic differentiation (AD), which makes simulators differentiable to enable gradient-based optimisation, and simulation-based inference (SBI), which learns approximate posteriors from simulated data without changing the simulator. Despite their growing use for inferring belief distributions over model parameters, these methods have not been directly compared for ABMs. We present an empirical study comparing AD-based variational inference and SBI on a spatial SIRS model. We evaluate the methods based on the trade-offs they present between predictive accuracy, sample efficiency, and implementation complexity. While our results suggest that SBI is preferable for ABMs with low-dimensional parameter spaces, they also highlight the need for future research, which we outline in our discussion. Code for reproducibility is available at https://github.com/SteamedGit/ad_vs_sbi_workshop.

1 Introduction

Agent-based models (ABMs) simulate complex systems by specifying rules for individual agents and observing the emergent behaviour that arises from their interactions. This bottom-up modelling paradigm is widely used in epidemiology, ecology, and the social sciences because it naturally captures heterogeneity, spatial structure, and nonlinear feedbacks. However, calibrating ABMs to data remains challenging: their likelihood functions are typically intractable, and simulations are computationally expensive, making standard Bayesian or maximum-likelihood approaches impractical [1, 2].

Two modern approaches have emerged to address these problems. The first is to make simulators differentiable and perform gradient-based inference using automatic differentiation (AD). Differentiable ABMs [3, 4] enable gradient-based optimisation and variational inference that directly use gradient information, often by replacing discrete choices with continuous relaxations (e.g., Gumbel–Softmax [5]) or by applying reparameterisation tricks. The second approach is simulation-based inference (SBI) [6, 2], which treats the simulator as a black box and trains neural conditional density or ratio estimators from simulated parameter-data pairs (θ, x) to approximate the posterior distribution $p(\theta | x)$. SBI does not require modifying the simulator, allowing it to be applied to legacy code and to systems with discrete or otherwise non-differentiable dynamics.

These two paradigms trade different kinds of cost and benefit. AD-based inference can be sample-efficient and fast at test time because gradients guide optimisation, but it usually requires substantial engineering to make the simulator differentiable and can introduce bias from relaxations. SBI is easier to deploy on off-the-shelf simulators but can require large numbers of simulator runs and careful

*Correspondence: tim@aims.ac.za

representation learning or summary-statistic design, especially when outputs are high-dimensional and structured in space and time.

Despite their complementary strengths, AD and SBI have not been systematically compared on ABMs. In this paper, we provide such a comparison on a spatial SIRS-style ABM with local contact dynamics. Our contributions are:

- A controlled empirical comparison of AD-based and SBI-based variational inference on a spatiotemporal ABM.
- An open-source differentiable ABM implemented in JAX [7] with accompanying scripts to validate the correctness of the AD gradients against finite difference baselines to enable further work in this area.

2 Background

2.1 Simulation-based inference.

Simulation-based inference (SBI) provides Bayesian inference for complex simulators whose likelihood functions are intractable but from which data can be simulated. The goal is to infer parameters θ given observed data x_{obs} when one can only draw samples $x \sim p(x | \theta)$ from a simulator. Modern SBI algorithms train neural networks to approximate either the likelihood $p(x | \theta)$, the likelihood ratio $p(x | \theta)/p(x)$, or the posterior $p(\theta | x)$ directly from simulated pairs (θ, x) [16].

SBI methods can be either amortised or sequential. In amortised SBI, the estimator is trained once using parameters sampled from the prior and can then infer posteriors for arbitrary observations. Sequential variants, such as Sequential Neural Posterior Estimation (SNPE), Sequential Neural Likelihood Estimation (SNLE), and Sequential Neural Ratio Estimation (SNRE), refine the proposal distribution over several rounds so that subsequent simulations concentrate in regions of high posterior probability [2]. Sequential Neural Variational Inference (SNVI) [8] unifies these approaches under a single variational framework. It combines SNLE (or SNRE) with variational inference; after training an approximate likelihood $\ell_\psi(x|\theta)$, the divergence D between a neural posterior estimator q_ϕ and an approximate posterior with normalizing constant Z is minimised,

$$\operatorname{argmin}_{\phi} D(q_\phi(\theta) \parallel \ell_\psi(x_{\text{obs}}|\theta)p(\theta)/Z).$$

This objective reveals that the main SBI families—posterior, likelihood, and ratio estimation—can all be interpreted within a common variational framework. SNVI provides a unifying view of SBI objectives and training procedures, clarifying how existing methods relate under a shared variational principle. Like other SBI approaches, SNVI operates on black-box stochastic simulators and does not require access to simulator gradients.

2.2 Generalized Variational Inference.

Generalized Variational Inference (GVI) [9] extends standard variational inference by replacing both the log-likelihood and the Kullback–Leibler (KL) divergence in the evidence lower bound with more general loss and divergence functions. In standard variational inference, the goal is to approximate the posterior $p(\theta | x_{\text{obs}})$ with a tractable distribution $q_\phi(\theta)$ by maximising the ELBO,

$$\mathcal{L}_{\text{VI}}(q_\phi) = \mathbb{E}_{q_\phi(\theta)} [\log p(x_{\text{obs}} | \theta)] - \text{KL}(q_\phi(\theta) \parallel p(\theta)).$$

This objective corresponds to minimising a reverse-KL divergence between $q_\phi(\theta)$ and the true posterior and is optimal when the model is well specified. However, in practice, simulators and generative models are often misspecified or contain unmodelled stochasticity, in which case strict likelihood-based objectives can be overly restrictive.

GVI introduces flexibility by defining a broader objective

$$\mathcal{L}_{\text{GVI}}(q_\phi) = \mathbb{E}_{q_\phi(\theta)} [d(x_{\text{obs}}, \theta)] + D(q_\phi(\theta) \parallel p(\theta)),$$

where $d(x_{\text{obs}}, \theta)$ is any user-defined loss measuring data–model fit (“pseudo-likelihood”) and D is any divergence or regulariser on the variational distribution. Different choices of (d, D) recover

many known methods: the standard ELBO (log-loss with KL), β -VAEs (scaled KL), power posteriors (α -divergences), and robust generalised Bayes updates. The GVI framework therefore encompasses both classical and robust Bayesian updating while preserving desirable theoretical properties such as coherence and posterior consistency under mild assumptions.

For differentiable simulators, the GVI objective can be optimised directly using AD through the simulator, enabling gradient-based updates of variational parameters. For black-box simulators, GVI principles can be combined with neural density estimators to approximate the expectations in $\mathcal{L}_{\text{GVI}}$, as in recent Generalized Bayesian Inference (GBI) methods [10]. This connection places both AD-based variational approaches and SBI methods such as SNVI within a shared variational family: they differ mainly in how the expected loss is estimated: via simulator gradients in AD-based GVI, or via amortised neural estimators in SBI.

3 Methods

We compare two inference strategies on the same spatial SIRS agent-based model: (i) simulation-based inference using Sequential Neural Variational Inference (SNVI) and (ii) automatic-differentiation-based Generalized Variational Inference (GVI).

SIRS Model. We use a grid-based spatial Susceptible-Infected-Recovered-Susceptible (SIRS) model on a $N \times N$ grid with M agents. Infection occurs with probability p_{infect} when susceptible agents have an infected neighbour. Infected agents recover with probability p_{recover} , and recovered agents lose immunity with probability p_{wane} . We make the model differentiable using Gumbel-Softmax sampling for all discrete transitions [5] (Figure 1b). Although the underlying ABM operates on a grid, the observations with which we calibrate are ABM state-count timeseries (Figure 1a).

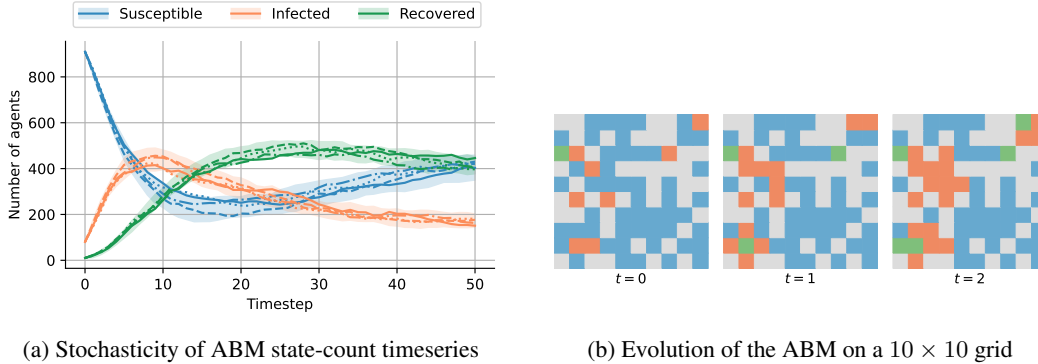


Figure 1: (a) Several realisations drawn from the ABM with the shaded region indicating the 95th percentile interval. (b) The underlying spatiotemporal structure. Grey regions are unoccupied.

SNVI. SNVI [8] performs Bayesian inference by optimising a variational lower bound on the marginal likelihood. The variational posterior q_ϕ is parametrised by a normalising flow and a neural likelihood (or neural ratio) estimator ℓ_ψ is substituted for the likelihood. We adopt the formulation based on a variational bound of the Rényi α -divergence² [11, 12, 8], which generalises the standard ELBO objective:

$$\mathcal{L}_{\text{SNVI},\alpha}(q_\phi) = \frac{1}{1-\alpha} \log \left(\mathbb{E}_{q_\phi(\theta)} \left[\left(\frac{\ell_\psi(x|\theta)p(\theta)}{q_\phi(\theta)} \right)^{1-\alpha} \right] \right).$$

Lower α values ($\alpha < 1$) emphasise coverage, while higher values focus on mode-seeking accuracy. Training proceeds for T sequential rounds: at each round we sample $\{\theta_i\}_{i=1}^N$ from the current proposal, generate simulations, update the neural likelihood (or neural ratio) estimator, update the neural posterior estimator by gradient ascent on $\mathcal{L}_{\text{SNVI},\alpha}$, and set the next-round proposal to the current posterior estimate. This focuses simulation effort on high-posterior regions without requiring

²Our experiments with the default forward-KL variant proved to be too unstable.

simulator gradients. In this work, we use SNVI with a neural likelihood estimator and henceforth refer to the method as SNVLI.

GVI. For the differentiable simulator we perform Generalized Variational Inference [9, 4]. We approximate the posterior by a normalizing-flow variational family $q_\phi(\theta)$ and minimise the objective proposed in [13]:

$$\mathcal{L}_{\text{GVI}} = w \mathbb{E}_{q_\phi(\theta)}[d(x_{\text{obs}}, \theta)] + \text{KL}(q_\phi(\theta) \parallel p(\theta)),$$

where $d(x_{\text{obs}}, \theta)$ is a differentiable discrepancy loss (e.g., squared error between simulated and observed summary statistics) and w is a scalar weighting that controls the influence of the data-fit term relative to the prior. Gradients of $\mathcal{L}_{\text{GVI}}$ with respect to ϕ are obtained via automatic differentiation through the differentiable ABM using reparameterised samples from $q_\phi(\theta)$.

Both methods therefore optimise variational objectives but differ in how they estimate expectations: SNVLI uses neural estimators on black-box simulators, while GVI uses pathwise gradients through a differentiable simulator.

4 Experiments

Since we do not have access to the classical posterior and GVI targets a generalised posterior, we opted to compare the methods in terms of the quality of their predictions using the Negative Log Predictive Density (NLPD). These comparisons were also made across different ABM sampling budgets in order to characterise the sample efficiency of the calibration techniques. Sample efficiency is important, because in practice, a major bottleneck is the computational cost of sampling from the ABM. The learning rates of SNVLI and AD GVI were tuned for each of the sampling budgets and since it was unclear *a priori* what the appropriate setting of w in $\mathcal{L}_{\text{GVI}}$ should be, we tuned two different variants. In Figure 2, SNVLI outperforms both variants of AD GVI at each sampling budget, is more sample efficient than the $w = 1e3$ variant, and unlike AD GVI, appears to monotonically improve as the sample budget increases. We provide more detail on the experimental setup in the Appendices.

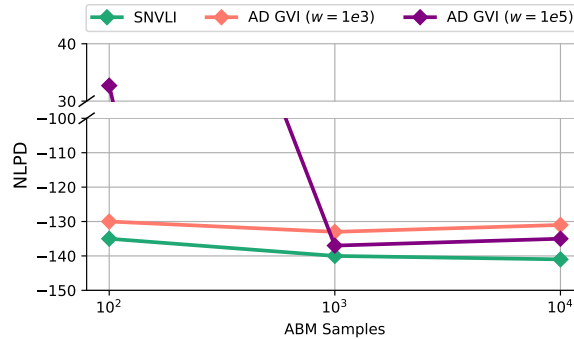


Figure 2: NLPD of the calibration techniques as the ABM sample budget increases

5 Discussion and future work

Our results suggest that for ABMs with low-dimensional parameter spaces, the effort required to make them differentiable may not be worthwhile. However, our experimentation is not comprehensive; we compared these methods on a single calibration task and only considered a single way of scaling the ABM sample budget. For SNVLI, we increased the number of simulations per round and for AD GVI we scaled the number of optimisation steps; however, we could have instead scaled the number of samples per optimisation step. It is also important that principled methods for determining the GVI objective’s data-fit term are developed as it directly impacts the concentration of the generalised posterior. Finally, future work would do well to include GVI with unbiased score-based gradient estimators in the comparison, as previous work [4] has shown it to work well for low-dimensional parameter spaces.

Acknowledgments and Disclosure of Funding

T.H. is a Google DeepMind scholar and acknowledges the African Institute for Mathematical Sciences (AIMS), South Africa, for the opportunity to conduct this research as a part of his MSc. T.H. acknowledges travel support from OJ Watson and the Academy of Medical Sciences/the Wellcome Trust/ the Government Department of Science Innovation and Technology/the British Heart Foundation/Diabetes UK Springboard Award SBF0010\1221. E.S. acknowledges AIMS South Africa for an opportunity to teach and meet students. E.S. acknowledges support in part by the AI2050 program at Schmidt Sciences (Grant [G-22-64476]).

References

- [1] Andreas Christ Sølvsten Jørgensen, Atiyo Ghosh, Marc Sturrock, and Vahid Shahrezaei. Efficient Bayesian inference for stochastic agent-based models. *PLoS Computational Biology*, 18(10):e1009508, October 2022.
- [2] Joel Dyer, Patrick Cannon, J. Doyne Farmer, and Sebastian M. Schmon. Black-box Bayesian inference for agent-based models. *Journal of Economic Dynamics and Control*, 161:104827, April 2024.
- [3] Ayush Chopra, Alexander Rodríguez, Jayakumar Subramanian, Arnau Quera-Bofarull, Balaji Krishnamurthy, B. Aditya Prakash, and Ramesh Raskar. Differentiable Agent-based Epidemiology, May 2023. arXiv:2207.09714 [cs].
- [4] Arnau Quera-Bofarull, Nicholas Bishop, Joel Dyer, Daniel Jarne Ornia, Anisoara Calinescu, Doyne Farmer, and Michael Wooldridge. Automatic differentiation of agent-based models. *arXiv preprint arXiv:2509.03303*, 2025.
- [5] Eric Jang, Shixiang Gu, and Ben Poole. Categorical reparameterization with gumbel-softmax. In *International Conference on Learning Representations*, 2017.
- [6] Kyle Cranmer, Johann Brehmer, and Gilles Louppe. The frontier of simulation-based inference. *Proceedings of the National Academy of Sciences*, 117(48):30055–30062, 2020.
- [7] James Bradbury, Roy Frostig, Peter Hawkins, Matthew James Johnson, Chris Leary, Dougal Maclaurin, George Necula, Adam Paszke, Jake VanderPlas, Skye Wanderman-Milne, and Qiao Zhang. JAX: composable transformations of Python+NumPy programs, 2018.
- [8] Manuel Glöckler, Michael Deistler, and Jakob H. Macke. Variational methods for simulation-based inference. In *International Conference on Learning Representations*, 2022.
- [9] Jeremias Knoblauch, Jack Jewson, and Theodoros Damoulas. Generalized variational inference: Three arguments for deriving new posteriors. *arXiv preprint arXiv:1904.02063*, 2019.
- [10] Richard Gao, Michael Deistler, and Jakob H Macke. Generalized bayesian inference for scientific simulators via amortized cost estimation. *Advances in Neural Information Processing Systems*, 36:80191–80219, 2023.
- [11] Alfréd Rényi. On Measures of Entropy and Information. In *Proceedings of the Fourth Berkeley Symposium on Mathematical Statistics and Probability, Volume 1: Contributions to the Theory of Statistics*, volume 4.1, pages 547–562. University of California Press, January 1961.
- [12] Yingzhen Li and Richard E Turner. Rényi Divergence Variational Inference. In D. Lee, M. Sugiyama, U. Luxburg, I. Guyon, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 29. Curran Associates, Inc., 2016.
- [13] Arnau Quera-Bofarull, Ayush Chopra, Anisoara Calinescu, Michael Wooldridge, and Joel Dyer. Bayesian calibration of differentiable agent-based models, May 2023. arXiv:2305.15340 [cs].
- [14] T. Toffoli and N. Margolus. *Cellular Automata Machines: A New Environment for Modeling*. MIT Press series in scientific computation. Cambridge, 1987.
- [15] Jan Boelts, Michael Deistler, Manuel Gloeckler, Álvaro Tejero-Cantero, Jan-Matthis Lueckmann, Guy Moss, Peter Steinbach, Thomas Moreau, Fabio Muratore, Julia Linhart, Conor Durkan, Julius Vetter, Benjamin Kurt Miller, Maternus Herold, Abolfazl Ziaemehr, Matthijs Pals, Theo Gruner, Sebastian Bischoff, Nastya Krouglova, Richard Gao, Janne K. Lappalainen, Bálint Mucsányi, Felix Pei, Auguste Schulz, Zinovia Stefanidi, Pedro Rodrigues, Cornelius Schröder, Faried Abu Zaid, Jonas Beck, Jaivardhan Kapoor, David S. Greenberg, Pedro J. Gonçalves, and Jakob H. Macke. sbi reloaded: a toolkit for simulation-based inference workflows. *Journal of Open Source Software*, 10(108):7754, 2025.
- [16] Durk P Kingma, Tim Salimans, Rafal Jozefowicz, Xi Chen, Ilya Sutskever, and Max Welling. Improved variational inference with inverse autoregressive flow. In D. Lee, M. Sugiyama, U. Luxburg, I. Guyon, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 29. Curran Associates, Inc., 2016.

- [17] George Papamakarios, Theo Pavlakou, and Iain Murray. Masked autoregressive flow for density estimation. In I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017.
- [18] Daniel Ward. Flowjax: Distributions and normalizing flows in jax, 2025.
- [19] Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. In *International Conference on Learning Representations*, 2019.
- [20] Geoffrey Roeder, Yuhuai Wu, and David Duvenaud. Sticking the landing: simple, lower-variance gradient estimators for variational inference. In *Proceedings of the 31st International Conference on Neural Information Processing Systems*, NIPS’17, page 6928–6937, Red Hook, NY, USA, 2017. Curran Associates Inc.

A Differentiable SIRS ABM

The ABM on which we test the calibration approaches is a grid-based spatial SIRS ABM implemented in JAX [7]. The state of the ABM with a $N \times N$ grid is represented by a $3 \times N \times N$ array where the agents have a one-hot encoding in the three channels. A susceptible agent can only be infected by infected agents within its von Neumann neighbourhood [14]. Furthermore, since the cell is infected by *each* neighbour with probability p_{infect} , the effective probability of infection is given by

$$p_{\text{effective}} = 1 - (1 - p_{\text{infect}})^{n_{\text{adj}}},$$

where n_{adj} is the number of infected neighbours within the von Neumann neighbourhood. This number can be obtained for all the susceptible agents by performing a 2D convolution between the infected channel of the grid and the neighbourhood matrix,

$$V_N = \begin{pmatrix} 0 & 1 & 0 \\ 1 & 0 & 1 \\ 0 & 1 & 0 \end{pmatrix}.$$

The spread of infection can then be computed by a single Bernoulli sampling operation across the grid followed by the masking out of entries that do not correspond to susceptible agents. Similarly, the recovery of agents and the waning of immunity can also be computed by Bernoulli sampling followed by masking. In order to make the ABM differentiable we replaced all the discrete sampling operations with an implementation of Gumbel-Softmax sampling [5] that allows for masking out parts of the sample by multiplying by a floating-point representation of a Boolean mask.

Table 1: True SIRS ABM parameters for synthetic data generation

Parameter	Value
<i>General Parameters</i>	
Grid Size	40×40
Number of Simulation Steps	20
Gumbel Softmax Temperature	0.05
<i>Target Parameters for Inference</i>	
Infection Probability (p_{infect})	0.6
Recovery Probability (p_{recover})	0.3
Waning Immunity Probability (p_{wane})	0.1
<i>Fixed Initial Conditions</i>	
Total Population	1280
Number of Initial Infected	160
Number of Initial Recovered	80

B Neural Network Architectures, Training and Evaluation

All the hyperparameter sweeps were performed on a computer with 15 GB of RAM and a NVIDIA T4 GPU. The AD GVI models train in seconds on a laptop GPU and `sbi` [15] notes that GPU acceleration is not necessary for most of its models.

AD GVI. The variational distribution is parametrised by a masked autoregressive flow [16, 17] implemented in `FlowJAX` [18]. The discrepancy loss $d(x_{\text{obs}}, \theta)$ is the MSE between x_{obs} and a batch of realisations $\{x_i\}_{i=1}^B$ sampled from the ABM with parameter θ . Throughout our experiments we sample 10 ABM realisations per optimisation step; two ELBO samples and 5 ABM samples per discrepancy loss evaluation (i.e. $B = 5$). Training is performed using AdamW [19] and we scale the number of ABM samples by increasing the number of optimisation steps. However, as we noted in Section 5, we could have instead fixed the number of optimisation steps and scaled the number of ABM samples per step.

SNVLI. We use the `sbi` package’s [15] implementation of SNVLI with default settings apart from using the objective based on the α -divergence instead of the default forward KL divergence.

Following the suggestions of Glöckler et al. [8] in the SNVI paper, we set $\alpha = 0.1$ and use the ‘Sticking the Landing’ gradient estimator [20]. Throughout our experiments we use two rounds of inference and we scale the number of ABM samples by increasing the number of ABM simulations sampled within each round. We could have alternatively scaled the number of rounds; however, this is less desirable since each round requires training two normalising flow models to convergence.

Evaluation. We generated 100 groundtruth realisations from the ABM for the NLPD calculation. For each trained model, we sampled 1000 parameters, used realisations generated by those parameters to fit a Kernel Density Estimate of the posterior predictive distribution and calculated the NLPD of the groundtruth realisations using that estimate. In order to reduce the dimensionality of the posterior predictive distribution, we only calculated the NLPD on the Infected and Recovered portions of the timeseries. This is justified since $N_{\text{Susceptible}}^t = N_{\text{Total}} - N_{\text{Infected}}^t - N_{\text{Recovered}}^t$. For completeness we also include all the posteriors and posterior predictive distributions in the next section of the Appendices.

C Posteriors and Predictive Distributions

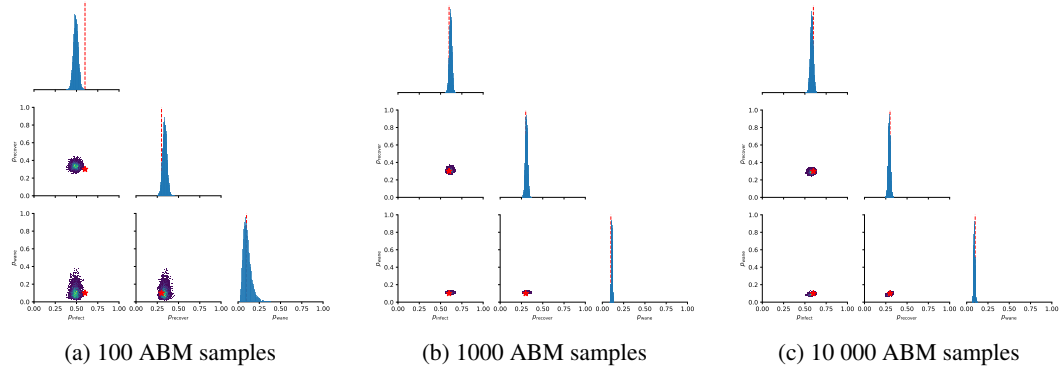


Figure 3: Posteriors for SNVLI

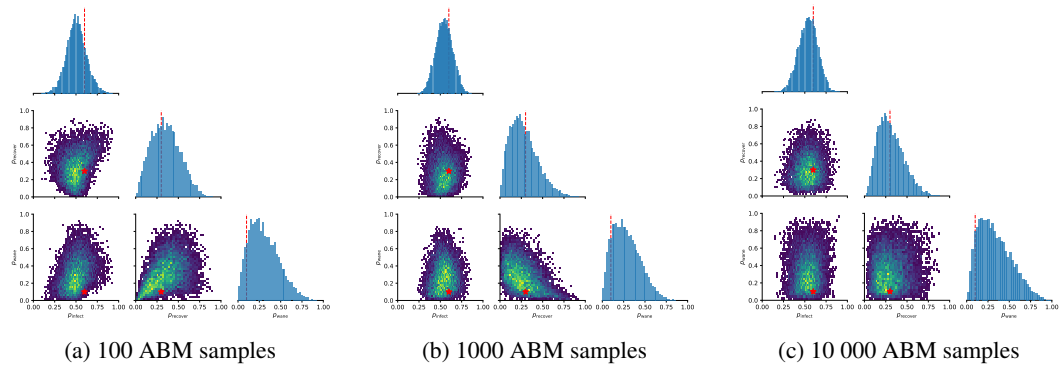


Figure 4: Posteriors for AD GVI ($w = 1e3$)

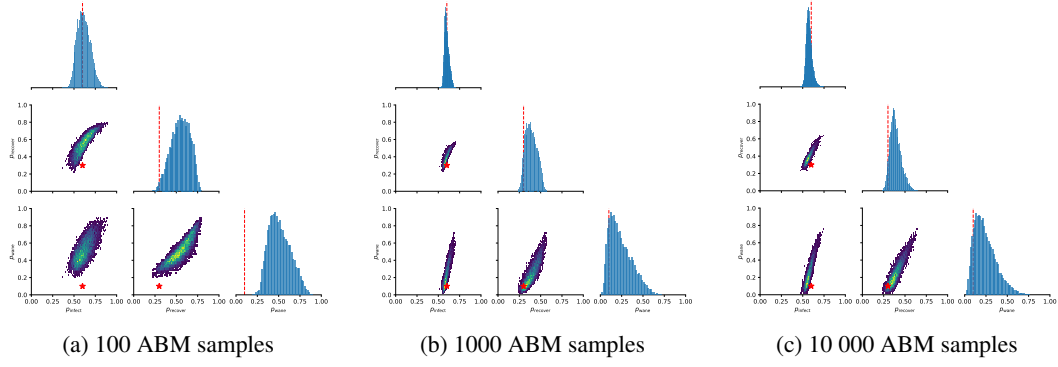


Figure 5: Posteriors for AD GVI ($w = 1e5$)

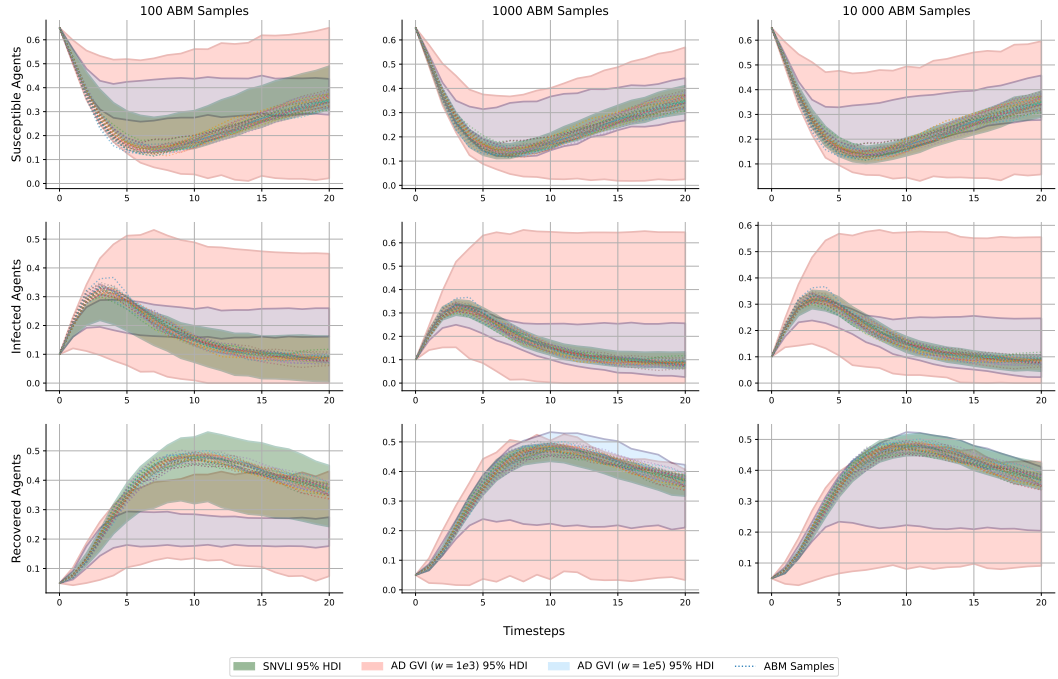


Figure 6: 95% Highest Density Intervals (HDIs) of the posterior predictive distributions with groundtruth ABM realisations overlaid.