

μGUIDE: a framework for quantitative imaging via generalized uncertainty-driven inference using deep learning

Maeliss Jallais , Marco Palombo

Cardiff University Brain Research Imaging Centre (CUBRIC), School of Psychology, Cardiff University, Cardiff, United Kingdom; • School of Computer Science and Informatics, Cardiff University, Cardiff, United Kingdom

 https://en.wikipedia.org/wiki/Open_access

 Copyright information

Reviewed Preprint

v2 • October 30, 2024

Revised by authors

Reviewed Preprint

v1 • August 22, 2024

eLife Assessment

The authors proposed an **important** novel deep-learning framework to estimate posterior distributions of tissue microstructure parameters. The method shows superior performance to conventional Bayesian approaches and there is **convincing** evidence for generalizing the method to use data from different protocol acquisitions and work with models of varying complexity.

<https://doi.org/10.7554/eLife.101069.2.sa3>

Abstract

This work proposes μGUIDE: a general Bayesian framework to estimate posterior distributions of tissue microstructure parameters from any given biophysical model or signal representation, with exemplar demonstration in diffusion-weighted MRI. Harnessing a new deep learning architecture for automatic signal feature selection combined with simulationbased inference and efficient sampling of the posterior distributions, μGUIDE bypasses the high computational and time cost of conventional Bayesian approaches and does not rely on acquisition constraints to define model-specific summary statistics. The obtained posterior distributions allow to highlight degeneracies present in the model definition and quantify the uncertainty and ambiguity of the estimated parameters.

1 Introduction

Diffusion-weighted Magnetic Resonance Imaging (dMRI) is a promising technique for characterizing brain microstructure in-vivo using a paradigm called microstructure imaging [Novikov et al., 2018a, Alexander et al., 2019 , Jelescu et al., 2020 ]. Traditionally, microstructure imaging quantifies histologically meaningful features of brain microstructure by fitting a forward (biophysical) model voxel-wise to the set of signals obtained from images acquired with different sensitivities, yielding model parameter maps [Alexander et al., 2019 ].

Most commonly used techniques rely on a non-linear curve fitting of the signal and return the optimal solution, i.e. the best parameters guess of the fitting procedure. However, this may hide model degeneracy, that is all the other possible estimates that could explain the observed signal equally well [Jelescu et al., 2016]. Another crucial consideration in model fitting is accounting for the uncertainty in parameter estimates. This uncertainty serves various purposes, including assessing result confidence [Jones, 2003], quantifying noise effects [Behrens et al., 2003] or assisting in experimental design [Alexander, 2008].

Instead of attempting to remove the degeneracies, which has been the focus of a large number of studies [Palombo et al., 2023, De Almeida Martins et al., 2021, Sator et al., 2021, Jelescu et al., 2022, Warner et al., 2023, Uhl et al., 2024, Mougél et al., 2024, Olesen et al., 2022, Palombo et al., 2020, Howard et al., 2022, Jones et al., 2018, Vincent et al., 2020, Henriques et al., 2021, Afzali et al., 2021, Lampinen et al., 2023, Zhang et al., 2012-07, Guerreri et al., 2023, Gyori et al., 2022, Novikov et al., 2018b], we propose to highlight them and present all the possible parameter values that could explain an observed signal, providing users with more information to make more confident and explainable use of the inference results.

Posterior distributions are powerful tools to characterize all the possible parameter estimations that could explain an observed measurement, their uncertainty, and existing model degeneracy [Box and Tiao, 2011]. Bayesian inference allows for the estimation of these posterior distributions, traditionally approximating them using numerical methods, such as Markov-Chain-Monte-Carlo (MCMC) [Metropolis et al., 1953]. In quantitative MRI, these methods have been used for example to estimate brain connectivity [Behrens et al., 2003], optimize imaging protocols [Alexander, 2008], or infer crossing fibers by combining multiple spatial resolutions [Sotiropoulos et al., 2013]. However, these classical Bayesian inference methods are computationally expensive and time consuming. They also often require adjustments and tuning specific to each biophysical model [Harms and Roebroeck, 2018].

Harnessing a new deep learning architecture for automatic signal feature selection and efficient sampling of the posterior distributions using simulation-based inference [Cranmer et al., 2020, Lueckmann et al., 2017, Papamakarios et al., 2019], here we propose μ GUIDE: a general Bayesian framework to estimate posterior distributions of tissue microstructure parameters from any given biophysical model/signal representation. μ GUIDE extends and generalises previous work [Jallais et al., 2022] to any forward model and without acquisition constraints, providing fast estimations of posterior distributions voxel-wise. We demonstrate μ GUIDE using numerical simulations on three biophysical models of increasing complexity and degeneracy and compare the obtained estimates with existing methods, including the classical MCMC approach. We then apply the proposed framework to dMRI data acquired from healthy human volunteers and participants with epilepsy. μ GUIDE framework is agnostic to the origin of the data and the details of the forward model, so we envision its usage and utility to perform Bayesian inference of model parameters also using data from other MRI modalities (e.g. relaxation MRI) and beyond.

2 Results

2.1 Framework overview

The full architecture of the proposed Bayesian framework, dubbed μ GUIDE, is presented in Fig. 1. μ GUIDE allows to efficiently estimate full posterior distributions of tissue parameters. It is comprised of two modules that are optimized together to minimize the Kullback-Leibler divergence between the true posterior distribution and the estimated one for every parameters of a given forward model. The ‘Neural Posterior Estimator’ (NPE) module [Papamakarios et al.,

2017 [\[1\]](#) uses normalizing flows [\[Papamakarios et al., 2021 \[\\[2\\]\]\(#\)\]](#) to approximate the posterior distribution, while the ‘Multi-Layer Perceptron’ (MLP) module is used to reduce the data dimensionality and ensure fast and robust convergence of the NPE module.

The full posterior distribution contains a lot of useful information. To summarize and easily visualize this information, we propose three measures that quantify the best estimates and the associated confidence levels, and a way to highlight degeneracy. The three measures are the Maximum A Posteriori (MAP), which corresponds to the most likely parameter estimate; an uncertainty measure, which quantifies the dispersion of the 50% most probable samples using the interquartile range, relative to the prior range; and an ambiguity measure, which measures the Full Width at Half Maximum (FWHM), in percentage with respect to the prior range. **Fig. 2** [\[3\]](#) presents those measures on exemplar posterior distributions.

We show exemplar applications of μ GUIDE to three biophysical models of increasing complexity and degeneracy from the dMRI literature: Ball&Stick [\[Behrens et al., 2003 \[\\[4\\]\]\(#\)\]](#) (Model 1); Standard Model (SM) [\[Novikov et al., 2018a\] \(Model 2\)](#); and extended-SANDI [\[Palombo et al., 2020\] \(Model 3\)](#).

2.2 Evaluation of μ GUIDE on simulations

2.2.1 Comparison with MCMC

We performed a comparison between the posterior distributions obtained using μ GUIDE and MCMC, a classical Bayesian method. **Fig. 3A** [\[5\]](#) shows posterior distributions on three exemplar simulations with SNR = 50 using the model 2 (SM), obtained with 15000 samples. Sharper and less biased posterior estimations are obtained using μ GUIDE. **Fig. 3B** [\[6\]](#) presents histograms for each model parameter of the bias between the ground truth value used to simulate a signal, and the MAP of the posterior distributions obtained with either μ GUIDE or MCMC, on 200 simulations. Results indicate that the bias is similar or smaller using μ GUIDE. Overall, μ GUIDE posterior distributions are more accurate than the ones obtained with MCMC.

Moreover, it took on average 29.3 s to obtain the posterior distribution using MCMC on a GPU (NVIDIA GeForce GT 710) for one voxel, while it only took 0.02 s for μ GUIDE. μ GUIDE is about 1500 times faster than MCMC, which makes it more suitable for applying it on large datasets.

2.2.2 The importance of feature selection

Fig. 4 [\[7\]](#) shows the MAP extracted from the posterior distributions versus the ground truth parameters used to generate the diffusion signal with μ GUIDE and manually-defined summary statistics for the three models. Less biased MAPs with lower ambiguities and uncertainties are obtained with μ GUIDE, indicating that the MLP allows for the extraction of additional information not contained in the summary statistics, helping to solve the inverse problem with higher accuracy and precision. μ GUIDE generalizes the method developed in [\[Jallais et al., 2022 \[\\[8\\]\]\(#\)\]](#) to make it applicable to any forward model and any acquisition protocol, while making the estimates more precise and accurate thanks to the automatic feature extraction.

2.2.3 μ GUIDE highlights degeneracies

Fig. 5 [\[9\]](#) presents the posterior distributions of microstructure parameters for the three models obtained with μ GUIDE on exemplar noise-free simulations. Blue curves correspond to non-degenerate posterior distributions, while the red ones present at least one degeneracy for one of the parameters. As the complexity of the model increases, degeneracy in the model definitions appear. This figure showcases μ GUIDE ability to highlight degeneracy in the model parameter estimation.

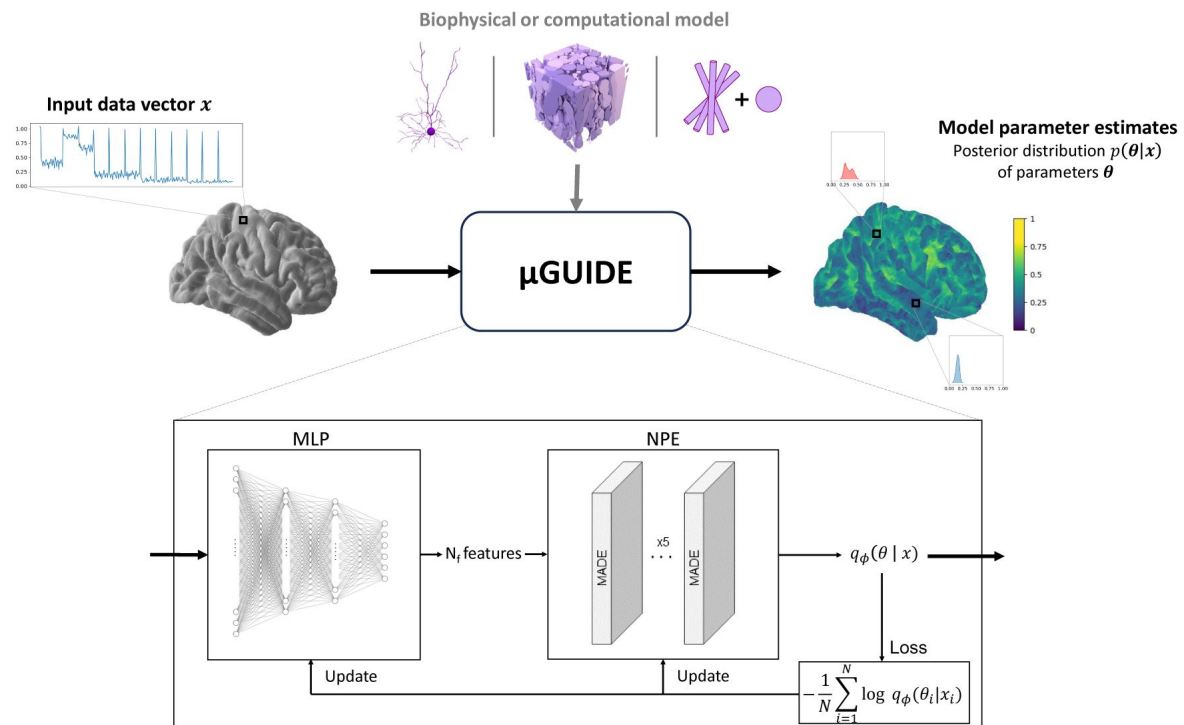


Figure 1

μGUIDE framework.

μGUIDE takes as input an observed data vector and relies on the definition of a biophysical or computational model [Ascoli et al., 2007, Callaghan et al., 2020, Jelescu et al., 2020]. It outputs a posterior distribution of the model parameters.

Based on a SBI framework, it combines a Multi-Layer Perceptron (MLP) with 3 layers and a Neural Posterior Estimator (NPE). The MLP learns a low-dimensional representation of $\mathbf{x}$, based on a small number of features (N_f), that can be either defined a priori or determined empirically during training. The MLP is trained simultaneously with the NPE, leading to the extraction of the optimal features that minimize the bias and uncertainty of $p(\boldsymbol{\theta}|\mathbf{x})$.

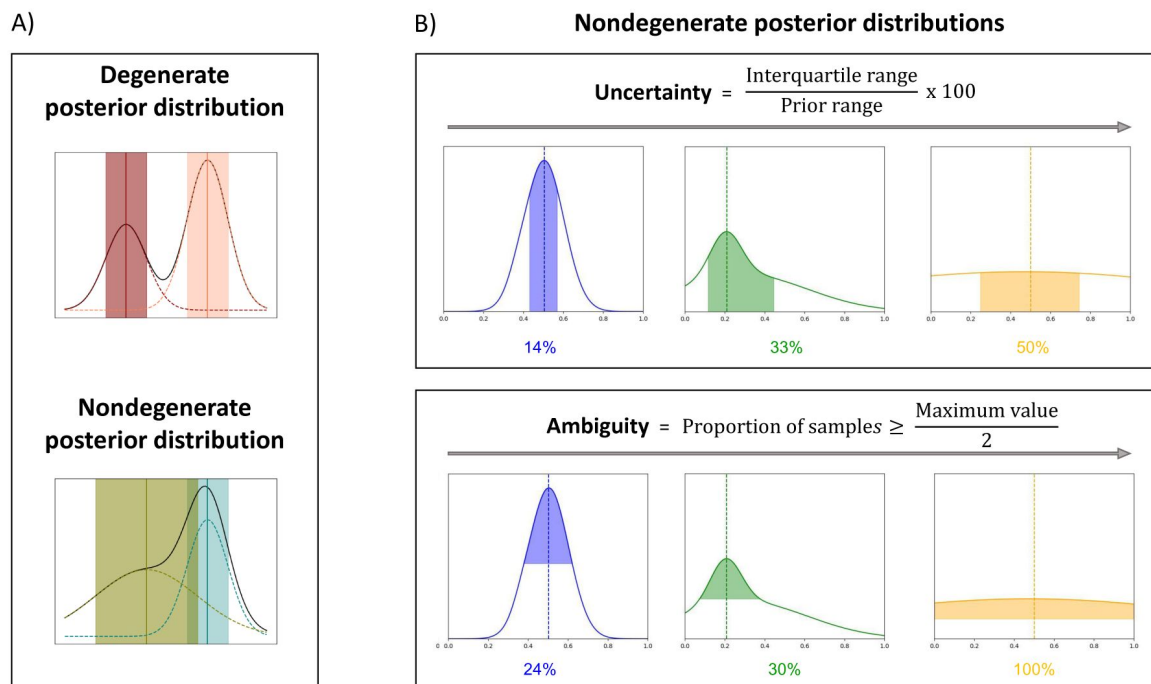


Figure 2

μGUIDE summarizes information contained in the estimated posterior distributions.

A) Examples of degenerate and non-degenerate posterior distributions. Two Gaussian distributions are fitted to the obtained posterior distribution, where the means and standard deviations are represented by the vertical lines and shaded areas. A voxel is considered as degenerate if the derivative of the fitted Gaussian distributions changes signs more than once (i.e. multiple local maxima), and if the two Gaussian distributions are not overlapping (the distance between the two Gaussian means is inferior to the sum of their standard deviations). **B)** Presentation of the measures introduced to quantify a posterior distribution on exemplar non-degenerate posterior distributions. Maximum A Posteriori (MAP): is the most likely parameter estimate (dashed vertical lines). Uncertainty: measures the dispersion of the 50% most probable samples using the interquartile range, with respect to the prior range. Ambiguity: measures the Full Width at Half Maximum (FWHM), in percentage with respect to the prior range.

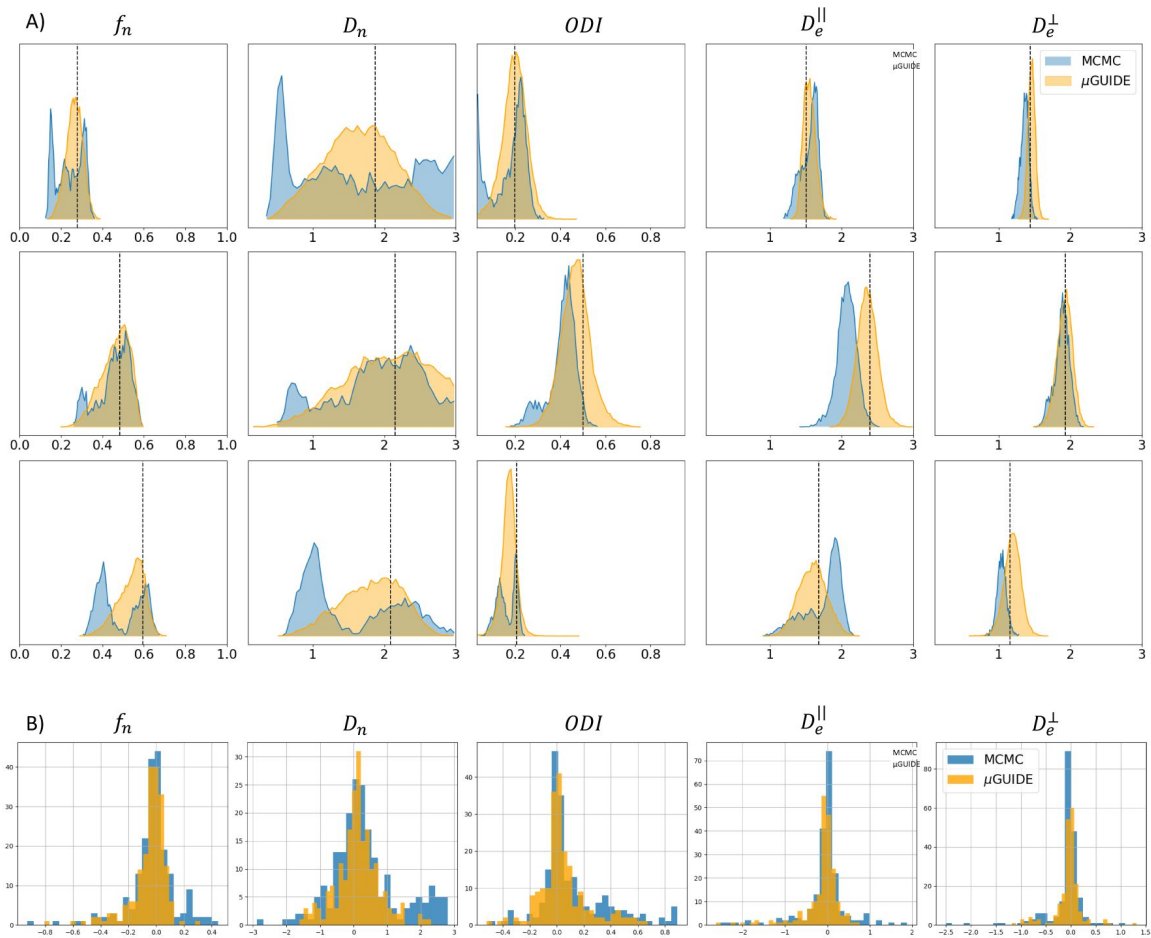


Figure 3

Comparison between μ GUIDE and MCMC.

A) Posterior distributions obtained using either μ GUIDE or MCMC on three exemplar simulations with model 2 (SM - SNR = 50). Names of the model parameters are indicated in the titles of the panels. **B)** Bias between the ground truth values used for simulating the diffusion signals, and the maximum-a-posteriori extracted from the posterior distributions using either μ GUIDE or MCMC. Sharper and less biased posterior distributions are obtained using μ GUIDE.

Tables 1 and **2** present the number of degenerate cases for each parameter in the three models, on 10000 simulations. **Table 1** considers noise-free simulations and the training and estimations were performed on CPU. **Table 2** reports results on noisy simulations (Rician noise with SNR = 50), with training and testing performed on a GPU (NVIDIA GeForce RTX 4090). The time needed for the inference and to estimate the posterior distributions on 10000 simulations, define if they are degenerate or not, and extract the MAP, uncertainty, ambiguity, are also reported. The more complex the model, the more degeneracies.

2.3 Application of μ GUIDE to real data

After demonstrating that the proposed framework provides good estimates in the controlled case of simulations, we applied μ GUIDE to both a healthy volunteer and a participant with epilepsy. The estimation of the posterior distributions is done independently for each voxel. To easily assess the values and the quality of the fitting, we are plotting the maximum-a-posterior, ambiguity and uncertainty maps, but the full posterior distributions are stored and available for all the voxels. Voxels presenting a degeneracy are highlighted with a red dot.

2.3.1 Healthy volunteer

We applied μ GUIDE to a healthy volunteer, using the Ball&Stick, SM and extended-SANDI models. **Fig. 6** presents the parametric maps of an exemplar set of model parameters for each model, alongside their degeneracy, uncertainty and ambiguity. The Ball&Stick model presents no degeneracy, the SM presents some degeneracy, mostly in voxels with high likelihood of partial voluming with cerebrospinal fluid (CSF) and at the white matter - gray matter boundaries. The extended-SANDI model is the model showing the highest number of degenerate cases, mostly localized within the white matter areas characterized by complex microstructure, e.g. crossing fibers. This result is expected, as the complexity of the models increases, leading to more combinations of tissue parameters that can explain an observed signal. Measures of ambiguity and uncertainty allow to quantify the confidence in the estimates and help interpreting the results.

2.3.2 Participant with epilepsy

Fig. 7 demonstrates μ GUIDE application to a participant with epilepsy, using the SM. Noteworthy, the axonal signal fraction estimates within the epileptic lesion show low uncertainty and ambiguity measures hence high confidence, while orientation dispersion index estimates show high uncertainty and ambiguity suggesting low confidence, cautioning the interpretation.

3 Discussion

3.1 Applicability of μ GUIDE to multiple models

The μ GUIDE framework offers the advantage of being easily applicable to various biophysical models and representations, thanks to its data-driven approach for data reduction. The need to manually define specific summary statistics that capture the relevant information for microstructure estimation from the multi-shell diffusion signal is removed. This also eliminates the acquisition constraints that were previously imposed by the summary statistics definition [Jallais et al., 2022]. The extracted features contain additional information compared to the summary statistics (see Appendix A), resulting in a notable reduction in bias (on average 5.2 fold lower), uncertainty (on average 2.6 fold lower) and ambiguity (on average 2.7 fold lower) in the estimated posterior distributions. Consequently, μ GUIDE improves parameters estimation over

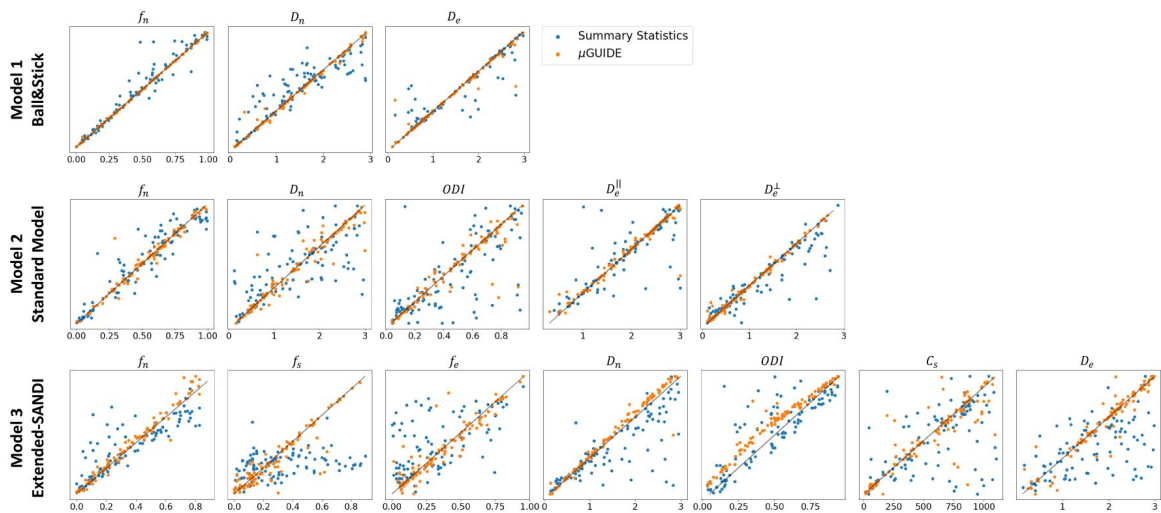


Figure 4

Fitting accuracy comparison between μ GUIDE's MLP-extracted features and manually-defined summary statistics.

Maximum-A-Posterioris extracted from the posterior distributions vs ground truth parameters used for generating the signal for the three models. Orange points correspond to the MAPs obtained using MLP-extracted features (μ GUIDE) and the blue ones to the MAPs with the manually-defined summary statistics. Only the non-degenerate posterior distributions were kept. The summary statistics used in those three models are the direction-averaged signal for the Ball&Stick model, the LEMONADE system of equations [Novikov et al., 2018b] for the SM, and the summary statistics defined in [Jallais et al., 2022] for the extended-SANDI model. Results are shown on 100 exemplar noise-free simulations with random parameter combinations. The optimal features extracted by the MLP allow to reduce the bias and variance of the obtained microstructure posterior distributions.

Figure 5

Exemplar posterior distributions of the microstructure parameters for the Ball&Stick, SM and extended SANDI models, obtained using μ GUIDE on exemplar noise-free simulations.

As the complexity of the model increases, degeneracies appear (red posterior distributions). μ GUIDE allows to highlight those degeneracies present in the model definition.

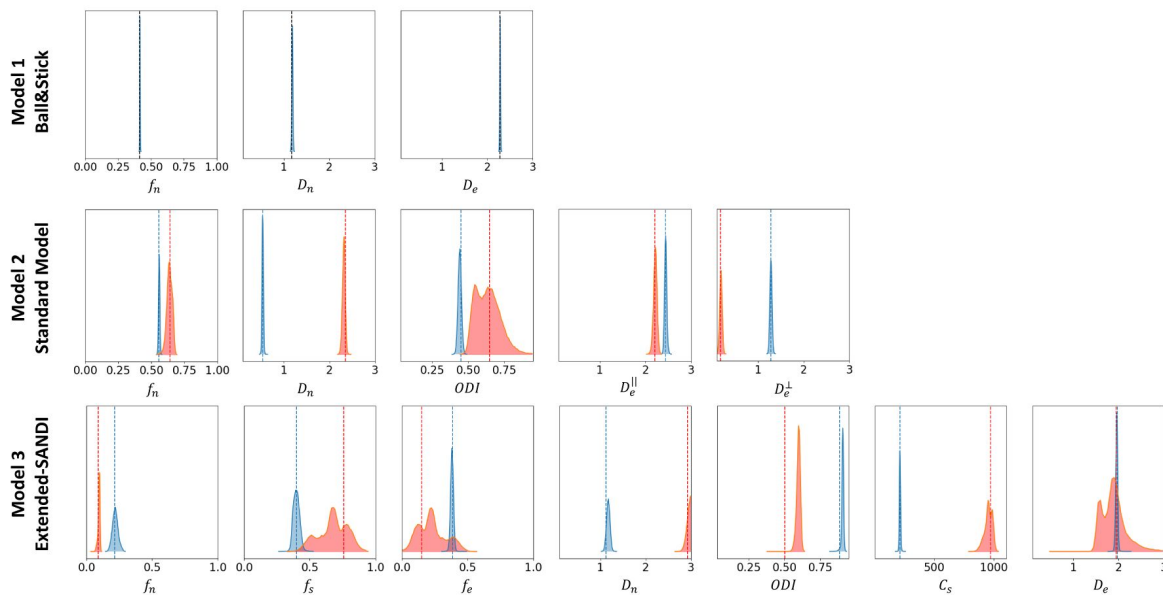


Table 1

Number of degenerate cases per parameter on 10000 noise-free simulations.

Training and estimations of the posterior distributions was performed on CPU. Time for training each model and time for estimating posterior distributions of 10000 noise-free simulations, define if they are degenerate or not, and extract the MAP, uncertainty, ambiguity are also reported.

Model (SNR = ∞)	Training Time (CPU)	Fitting Time (on 10000 simulations)	Number of Degeneracies (on 10000 simulations)							
			f_n	D_n	$D_e^{ }$	ODI	$D_e^{\perp}$	f_s	f_e	C_s
Model 1 Ball&Stick	11min	96s	0	0	0	-	-	-	-	-
Model 2 Standard Model	2h02	135s	4	34	23	33	8	-	-	-
Model 3 Extended-SANDI model	2h02	1412s	205	4	260	57	-	1395	2571	1011

current state-of-the-art methods (e.g. Jallais et al. [2022]), showing for example reduced bias (on average 5.2 fold lower) and dispersion (on average 6.4 fold lower) on the MAP estimates for each of the three example models investigated (see **Fig. 4** [↗](#)).

In this study, we presented applications of μ GUIDE to brain microstructure estimation using three well-established biophysical models, with increased complexity: the Ball&Stick model, the Standard Model, and an extended-SANDI model. However, our approach is not limited to brain tissue nor to diffusion-weighted MRI and can be extended to different organs by employing their respective acquisition encoding and forward models, such as NEXI for exchange estimates [Jallais et al., 2024 [↗](#)]; mcDESPOT for myelin water fraction mapping using quantitative MRI relaxation [Deoni et al., 2008 [↗](#)]; VERDICT in prostate imaging [Panagiotaki et al., 2014 [↗](#)]; or even adapted to different imaging modalities (e.g., EEG and MEG), where there is a way to link (via modelling or simulation) the observed signal to a set of parameters of interest. This versatility underscores the broad applicability of our proposed approach across various biological systems and imaging techniques.

It is important to note that μ GUIDE is still a model-dependent method, meaning that the training process is based on the specific model being used. Additionally, the number of features extracted by the MLP needs to be predetermined. One way to determine the number of features is by matching it with the number of parameters being estimated. Alternatively, a dimensionality-reduction study using techniques like t-distributed stochastic neighbour embedding (tSNE) [Van der Maaten and Hinton, 2008] can be conducted to determine the optimal number of features.

3.2 μ GUIDE: an efficient framework for Bayesian inference

One notable advantage of μ GUIDE is its amortized nature. With this approach, the training process is performed only once, and thereafter, the posterior estimations can be independently obtained for all voxels. This amortization enables efficient estimations of the posterior distributions. μ GUIDE outperforms in terms of speed conventional Bayesian inference methods such as MCMC, showing a ~ 1500 fold acceleration. The time savings achieved with μ GUIDE make it a highly efficient and practical tool for estimating posterior distributions in a timely manner.

This unlocks the possibility to process with Bayesian inference very large datasets in manageable time (e.g. approximately 6 months to process 10k dMRI datasets) and to include Bayesian inference in iterative processes that require the repeated computation of the posterior distributions (e.g., dMRI acquisition optimization Alexander [2008]).

In the dMRI community, the use of SBI methods to characterize full posterior distributions as well as quantify the uncertainty in parameter estimations was first introduced in Jallais et al. [2022] for a grey matter model. An application to crossing fibers has recently been proposed by Karimi et al. [2024]. Those approaches use different density estimators. This work and Jallais et al. [2022] rely on Masked Autoregressive FLOws (MAF) Papamakarios et al. [2017], while the work by Karimi et al. [2024] is based on Mixture Density Networks (MDN) [Bishop, 1994 [↗](#)]. MAFs have been found to show superior performance compared to MDNs [Goncalves et al., 2020 [↗](#), Patron et al., 2022 [↗](#)].

3.3 μ GUIDE quantifies confidence to guide interpretation

Quantifying confidence in an estimate is of crucial importance. As demonstrated by our pathological example, changes in the tissue microstructure parameters can help clinicians decide which parameters are the most reliable and better interpret microstructure changes within diseased tissue. On large population studies, the quantified uncertainty can be taken into account when performing group statistics and to detect outliers.

Model (SNR=50)	Training Time (GPU)	Fitting Time (on 10000 simulations)	Number of Degeneracies (on 10000 simulations)							
			f_n	D_n	$D_e^{ }$	ODI	$D_e^\perp$	f_s	f_e	C_s
Model 1 Ball&Stick	26min	79s	0	0	0	-	-	-	-	-
Model 2 Standard Model	42min	82s	75	71	117	109	29	-	-	-
Model 3 Extended-SANDI model	50min	238s	47	24	784	6	-	828	1047	56

Table 2

Number of degenerate cases per parameter on 10000 noisy simulations (Rician noise with SNR = 50).

Training and estimations of the posterior distributions was performed using a GPU. Time for training each model and time for estimating posterior distributions of 10000 noisy simulations, define if they are degenerate or not, and extract the MAP, uncertainty, ambiguity are also reported.

Multiple approaches have been used to try and quantify this uncertainty. Gradient descent often provides a measure of confidence for each parameter estimate. Alternative approaches use the shape of the fitted tensor itself as a measure of uncertainty for the fiber direction [Koch et al., 2002 [↗](#), Parker and Alexander, 2003 [↗](#)]. Other methods also rely on bootstrapping techniques to estimate uncertainty. Repetition bootstrapping for example depends on repeated measurements of signal for each gradient direction, but imply a long acquisition time and cost, and are prone to motion artifacts [Lazar and Alexander, 2005 [↗](#), Jones, 2003 [↗](#)]. In contrast, residual bootstrapping methods resample the residuals of a regression model. Yet, this approach is heavily dependent on the model and can lead to overfitting [Whitcher et al., 2008 [↗](#), Chung et al., 2006 [↗](#)]. In general, resampling methods can be problematic for sparse samples, as the bootstrapped samples tend to underestimate the true randomness of the distribution [Kauermann et al., 2009 [↗](#)]. We propose to quantify the confidence by estimating full posterior distributions, which also has the benefit of highlighting degeneracy. Model-fitting methods with different initializations, as done in e.g. Jelescu et al. [2016], also allow to highlight degeneracies. However, they only provide a partial description of the solution landscape, which can be interpreted as a partial posterior distribution. In contrast, Bayesian methods estimate the full posterior distributions, offering a more accurate and precise characterization of degeneracies and uncertainties. Hence, in this work we decided to use MCMC, a traditional Bayesian method, as benchmark.

Variance observed in the posterior distributions can be attributed to several factors. The presence of noise in the signal contributes to irreducible variance, decreasing the confidence in the estimates as the noise level increases (see Appendix B). Another source of variance can arise from the choice of acquisition parameters. Different acquisitions may provide varying levels of confidence in the parameter estimates. Under-sampled acquisitions or inadequate b-shells may fail to capture essential information about a tissue microstructure, such as soma or neurite radii, resulting in inaccurate estimates.

μ GUIDE can guide users in determining whether an acquisition is suitable for estimating parameters of a given model and *viceversa*, the variance and bias of the posterior distributions estimated with μ GUIDE can be used to guide the optimization of the data acquisition to maximize accuracy and precision of the model parameters estimates.

The presence of degeneracy in the solution of the inverse problem is influenced by the complexity of the model being used and the lack of sufficient information in the data. In recent years, researchers have introduced increasingly sophisticated models to better represent the brain tissue, such as SANDI [Palombo et al., 2020], NEXI [Jelescu et al., 2022 [↗](#)] and eSANDIX [Olesen et al., 2022 [↗](#)], that take into account an increasing number of tissue features. By applying μ GUIDE, it becomes possible to gain insights into the degree of degeneracy within a model and to assess the balance between model realism and the ability to accurately invert the problem. We have recently provided an example of such application for NEXI and SANDIX [Jallais et al., 2024 [↗](#)].

3.4 Summary

We propose a general Bayesian framework, dubbed μ GUIDE, to efficiently estimate posterior distributions of tissue microstructure parameters. For any given acquisition and signal model/representation, μ GUIDE improves parameters estimation and computational time over existing state-of-the-art methods. It allows to highlight degeneracy, and quantify confidence in the estimates, guiding results interpretation towards more confident and explainable diagnosis using modern deep learning. μ GUIDE is not inherently limited to dMRI and microstructure imaging. We envision its usage and utility to perform efficient Bayesian inference also using data from any modality where there is a way to link (via modelling or simulation) the observed measurements to a set of parameters of interest.

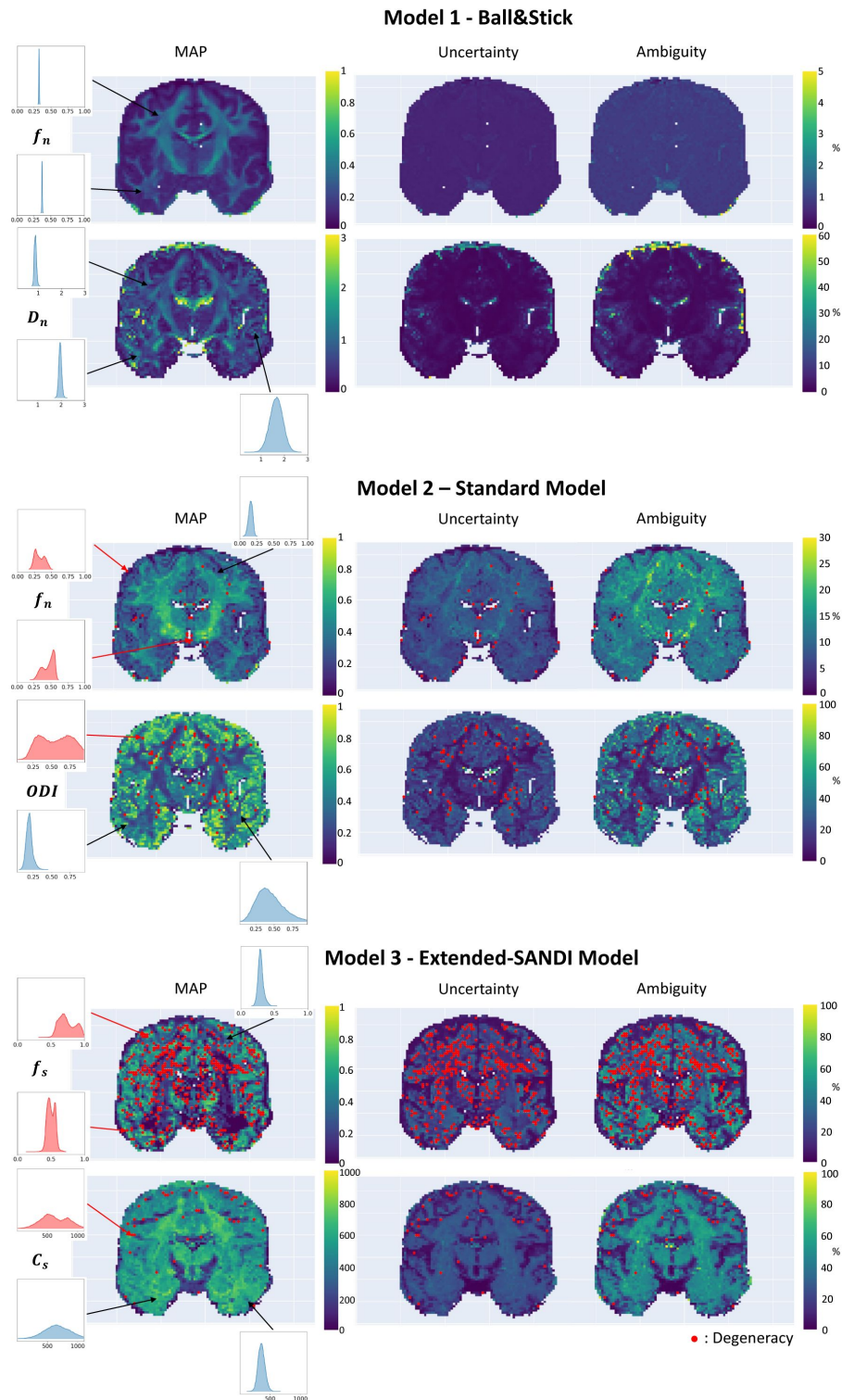


Figure 6

Parametric maps of the Ball&Stick (top), SM (middle) and extended-SANDI model (bottom), obtained using μ GUIDE.

Maximum-a-posteriori estimation, uncertainty and ambiguity measure maps are reported, overlaid with voxels considered degenerate (red dots).

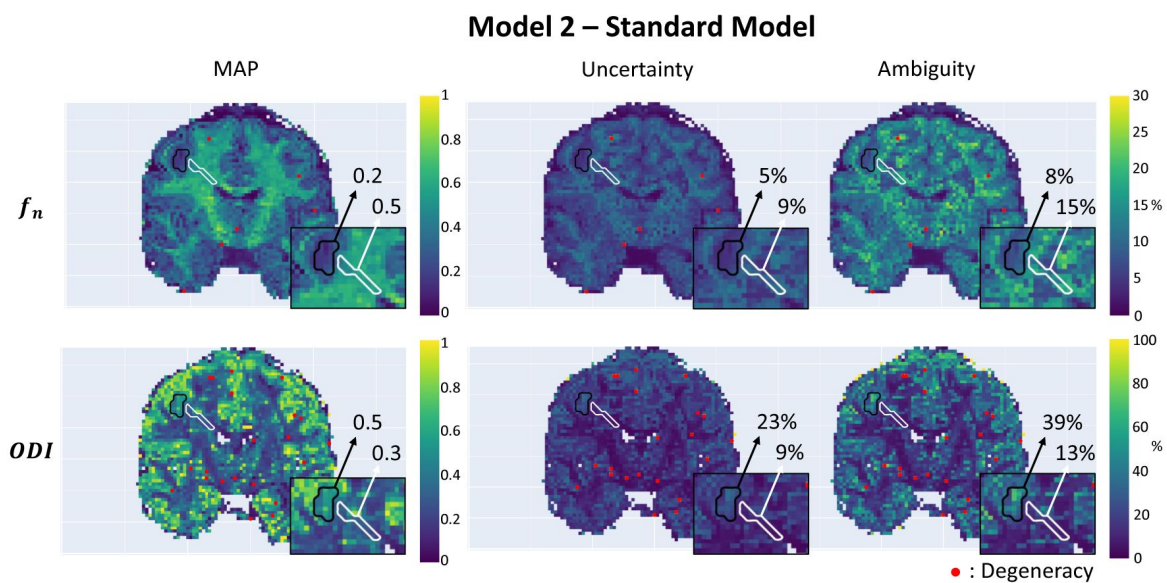


Figure 7

Parametric maps of a participant with epilepsy obtained using μ GUIDE with the SM, superimposed with the grey matter (black) and white matter (white) lesions segmentation.

Mean values of the maximum-a-posterior, uncertainty and ambiguity measures are reported in the two regions of interest. Lower MAP values are obtained in the lesions for the axonal signal fraction and the orientation dispersion index compared to healthy tissue. Higher uncertainty and ambiguity ODI values are reported, suggesting less stable estimations.

4 Methods

4.1 Solving the inverse problem using Bayesian inference

4.1.1 The inference problem

We make the hypothesis that an observed dMRI signal $\mathbf{x}_0$ can be explained (and generated) using a handful of relevant tissue microstructure parameters $\boldsymbol{\theta}_0$, following the definition of a forward model:

$$\mathbf{x}_0 = \mathcal{M}(\boldsymbol{\theta}_0)$$

The objective is, given this observation $\mathbf{x}_0$, to estimate the parameters $\boldsymbol{\theta}_0$ that generated it.

Forward models are designed to mimic at best a given biophysical phenomenon, for some given time and scale [Alexander, 2009 [\[1\]](#), Yablonskiy and Sukstanskii, 2010 [\[2\]](#), Jelescu and Budde, 2017 [\[3\]](#), Novikov et al., 2018a, Alexander et al., 2019 [\[4\]](#), Jelescu et al., 2020 [\[5\]](#)]. As a consequence, forward models are injection functions (every biologically plausible $\boldsymbol{\theta}_i$ generates exactly one signal $\mathbf{x}_i$), but do not always happen to be bijections, meaning that multiple $\boldsymbol{\theta}_i$ can generate the same signal $\mathbf{x}_i$. It can be impossible, based on biological considerations, to infer which solution $\boldsymbol{\theta}_i$ best reflects the probed structure. We refer to these models as ‘degenerate models’.

Point estimates algorithms, such as minimum least square or maximum likelihood estimation algorithms, allow to estimate one set of microstructure parameters that could explain an observed signal. In the case of degenerate models, the solution space can be multi-modal and those algorithms will hide possible solutions. When considering real-life acquisitions, i.e. noisy and/or under-sampled acquisitions, one also needs to consider the bias introduced with respect to the forward model, and the resulting variance in the estimates [Jones, 2003 [\[6\]](#), Behrens et al., 2003 [\[7\]](#)].

We propose a new framework that allows for the estimation of full posterior distributions $p(\boldsymbol{\theta} | \mathbf{x}_0)$, that is all the probable parameters that could represent the underlying tissue, along with an uncertainty measure and the interdependency of parameters. These posteriors can help interpreting the obtained results and make more informed decisions.

4.1.2 The Bayesian formalism

The posterior distribution can be defined using Bayes’ theorem as follows:

$$p(\boldsymbol{\theta} | \mathbf{x}_0) = \frac{p(\mathbf{x}_0 | \boldsymbol{\theta}) p(\boldsymbol{\theta})}{p(\mathbf{x}_0)} , \quad (1)$$

where $p(\mathbf{x}_0 | \boldsymbol{\theta})$ is the likelihood of the observed data point, $p(\boldsymbol{\theta})$ is the prior distribution defining our initial knowledge of the parameter values, and $p(\mathbf{x}_0)$ is a normalizing constant, commonly referred to as the evidence of the data.

The evidence term is usually very hard to estimate, as it corresponds to all the possible realisations of $\mathbf{x}_0$, i.e. $p(\mathbf{x}_0) = \int_{\text{all } \mathbf{x}_0} p(\mathbf{x}_0 | \boldsymbol{\theta}) p(\boldsymbol{\theta}) d\mathbf{x}_0$. For simplification, methods usually estimate an unnormalized probability density function, i.e.

$$p(\boldsymbol{\theta} | \mathbf{x}_0) \propto p(\mathbf{x}_0 | \boldsymbol{\theta}) p(\boldsymbol{\theta}) . \quad (2)$$

To approximate these posterior distributions, traditional methods rely on the estimation of the likelihood $p(\mathbf{x}_0|\boldsymbol{\theta})$ of the observed data point $\mathbf{x}_0$ via an analytic expression. This likelihood function corresponds to an integral over all possible trajectories through the latent space, that is $p(\mathbf{x}_0|\boldsymbol{\theta}) = \int p(\mathbf{x}_0, \mathbf{z}|\boldsymbol{\theta})d\mathbf{z}$, where $p(\mathbf{x}_0, \mathbf{z}|\boldsymbol{\theta})$ is the joint probability density of observed data $\mathbf{x}_0$ and latent variables $\mathbf{z}$. For forward models with large latent spaces, computing this integral explicitly becomes impractical. The likelihood function is then intractable, rendering these methods unusable [Cranmer et al., 2020]. Models that do not admit a tractable likelihood are called *implicit models* [Diggle and Gratton, 1984].

To circumvent this issue, some techniques have been proposed to sample numerically from the likelihood function, such as Markov-Chain-Monte-Carlo (MCMC) [Metropolis et al., 1953]. Another set of approaches proposes to train a conditional density estimator to learn a surrogate of the likelihood distribution [Papamakarios et al., 2019, Lueckmann et al., 2019], the likelihoodratio [Cranmer et al., 2016, Gutmann et al., 2018] or the posterior distribution [Papamakarios and Murray, 2016, Lueckmann et al., 2017, Papamakarios et al., 2019], allowing to greatly reduce computation times. These methods are dubbed Likelihood-Free Inference (LFI), or SimulationBased Inference (SBI) methods [Cranmer et al., 2020, Tejero-Cantero et al., 2020]. In particular, there has been a growing interest towards deep generative modeling approaches in the machine learning community [Lueckmann et al., 2021]. They rely on specially tailored neural network architectures to approximate probability density functions from a set of examples. Normalizing flows [Papamakarios et al., 2021] are a particular class of such neural networks that have demonstrated promising results for SBI in different research fields [Goncalves et al., 2020, Greenberg et al., 2019].

While this work focuses on the estimate of the posterior distribution using a conditional density estimator, we show a comparison with MCMC, which are commonly used methods in the community. We will therefore introduce this method in the following paragraph.

4.1.3 Estimating the likelihood function

Well-established approaches for estimating the likelihood function are MCMC methods. These methods rely on a noise model to define the likelihood distribution, such as the Rician [Panagiotaki et al., 2012] or Offset Gaussian models [Alexander, 2009]. In this work, we will be using the Microstructure Diffusion Toolbox to perform the MCMC computations [Harms and Roebroek, 2018], which relies on the Offset Gaussian model. The log-likelihood function is then the following:

$$\log(p(\mathbf{x}|\boldsymbol{\theta})) = -\frac{\sum_{i=1}^m \left(\mathbf{x}_i - \sqrt{\mathcal{M}(\boldsymbol{\theta})_i^2 + \sigma^2} \right)^2}{2\sigma^2} - m \cdot \log(\sigma\sqrt{2\pi}) , \quad (3)$$

where $\mathcal{M}(\boldsymbol{\theta})$ is the signal obtained using the biophysical model, $\mathcal{M}(\boldsymbol{\theta})_i$ is the i th measurement of the signal, σ is the standard deviation of the Gaussian distributed noise, estimated from the reconstructed magnitude images [Dietrich et al., 2007], and m is the number of observations in the dataset.

MCMC methods allow to obtain posterior distributions using Bayes' formula (Eq. (2)) with the previously defined likelihood function (Eq. (3)) and some prior distributions, which are usually uniform distributions defined on biologically plausible ranges. They generate a multi-dimensional chain of samples which is guaranteed to converge towards a stationary distribution, which approximates the posterior distribution [Metropolis et al., 1953].

The need to compute the signal following the forward model at each iteration makes these sampling methods computationally expensive and time consuming. In addition, they require some adjustments specific to each model, such as the choice of burn-in length, thinning and the number of samples to store. Harms and Roebroek [2018] recommend to use the Adaptive Metropolis-Within-Gibbs (AMWG) algorithm for sampling dMRI models, initialized with a maximum likelihood estimator obtained from non-linear optimization, with 100 to 200 samples for burn-in and no thinning. Authors notably investigated the use of starting from the MLE and thinning. They concluded that starting from the MLE allows to start in the stationary distribution of the Markov chain, and has the advantage of removing salt- and pepper-like noise from the resulting mean and standard deviation maps. Their findings also indicate that thinning is unnecessary and inefficient, and they recommend using more samples instead. The recommended number of samples is modeldependent. Authors recommendations can be found in their paper.

4.1.4 Bypassing the likelihood function

An alternative method was proposed to overcome the challenges associated with approximating the likelihood function and the limitations of MCMC sampling algorithms. This approach involves directly approximating the posterior distribution by using a conditional density estimator, i.e. a family of conditional probability density function approximators denoted as $q_{\phi}(\boldsymbol{\theta}|\mathbf{x})$. These approximators are parameterized by ϕ and accept both the parameters $\boldsymbol{\theta}$ and the observation $\mathbf{x}$ as input arguments. Our posterior approximation is then obtained by minimizing its average Kullback-Leibler divergence with respect to the conditional density estimator for different choices of $\mathbf{x}$, as per [Papamakarios and Murray, 2016]:

$$\min_{\phi} \mathcal{L}(\phi) \quad \text{with} \quad \mathcal{L}(\phi) = \mathbb{E}_{\mathbf{x} \sim p(\mathbf{x})} [D_{\text{KL}}(p(\boldsymbol{\theta}|\mathbf{x}) \| q_{\phi}(\boldsymbol{\theta}|\mathbf{x}))], \quad (4)$$

which can be rewritten as

$$\begin{aligned} \mathcal{L}(\phi) &= \int D_{\text{KL}}(p(\boldsymbol{\theta}|\mathbf{x}) \| q_{\phi}(\boldsymbol{\theta}|\mathbf{x})) p(\mathbf{x}) d\mathbf{x}, \\ &= - \iint \log(q_{\phi}(\boldsymbol{\theta}|\mathbf{x})) p(\boldsymbol{\theta}|\mathbf{x}) p(\mathbf{x}) d\boldsymbol{\theta} d\mathbf{x} + C, \\ &= - \iint \log(q_{\phi}(\boldsymbol{\theta}|\mathbf{x})) p(\mathbf{x}, \boldsymbol{\theta}) d\boldsymbol{\theta} d\mathbf{x} + C, \\ &= - \mathbb{E}_{(\mathbf{x}, \boldsymbol{\theta}) \sim p(\mathbf{x}, \boldsymbol{\theta})} [\log(q_{\phi}(\boldsymbol{\theta}|\mathbf{x}))] + C, \end{aligned} \quad (5)$$

where C is a constant that does not depend on ϕ . Note that in practice we consider a N -sample Monte-Carlo approximation of the loss function:

$$\mathcal{L}(\phi) \approx \mathcal{L}^N(\phi) = -\frac{1}{N} \sum_{i=1}^N \log(q_{\phi}(\boldsymbol{\theta}_i|\mathbf{x}_i)), \quad (6)$$

where the N data points $(\boldsymbol{\theta}_i, \mathbf{x}_i)$ are sampled from the joint distribution with $\boldsymbol{\theta}_i \sim p(\boldsymbol{\theta})$ and $\mathbf{x}_i \sim p(\mathbf{x}|\boldsymbol{\theta}_i)$. We can then use stochastic gradient descent to obtain a set of parameters ϕ which minimizes $\mathcal{L}^N$.

If the class of conditional density estimators is sufficiently expressive, it can be demonstrated that the minimizer of Eq. (6) converges to $p(\boldsymbol{\theta}|\mathbf{y})$ when $N \rightarrow \infty$ [Greenberg et al., 2019]. It is worth noting that the parametrization ϕ , obtained at the end of the optimization procedure, serves as an amortized posterior for various choices of $\mathbf{x}$. Hence, for a particular observation $\mathbf{x}_0$, we can simply use $q_{\phi}(\boldsymbol{\theta}|\mathbf{x}_0)$ as an approximation of $p(\boldsymbol{\theta}|\mathbf{x}_0)$.

4.2 μ GUIDE framework

The full architecture of the proposed Bayesian framework, dubbed μ GUIDE, is presented in **Fig. 1**. The analysis codes underpinning the results presented here can be found upon publication in the Cardiff University data catalogue and on Github: <https://github.com/mjallais/uGUIDE> (both CPU and GPU are supported).

μ GUIDE is comprised of two modules that are optimized together to minimize the Kull-back-Leibler divergence between the true posterior distribution and the estimated one for every parameters of a given forward model. The ‘Neural Posterior Estimator’ (NPE) module uses normalizing flows to approximate the posterior distribution, while the ‘Multi-Layer Perceptron’ (MLP) module is used to reduce the data dimensionality and ensure fast and robust convergence of the NPE module. The following Sections provide more details about our implementation of each module.

4.2.1 Neural Posterior Estimator

In this study, the Sequential Neural Posterior Estimation (SNPE-C) algorithm [Papamakarios and Murray, 2016, Greenberg et al., 2019] with a single round is employed to train a neural network that directly approximates the posterior distribution. Thus, sampling from the posterior can be done by sampling from the trained neural network. Neural density estimators have the advantage of providing exact density evaluations, in contrast to Variational Autoencoders (VAE) Kingma and Welling [2019] or generative adversarial networks (GAN) Goodfellow et al. [2014], which are better suited for generating synthetic data.

The conditional probability density function approximators used in this project belong to a class of neural networks called normalizing flows [Papamakarios et al., 2021]. These flows are invertible functions capable of transforming vectors generated from a simple base distribution (e.g. the standard multivariate Gaussian distribution) into an approximation of the true posterior distribution. An autoregressive architecture for normalizing flows is employed, implemented via the Masked Autoregressive Flow (MAF) Papamakarios et al. [2017], which is constructed by stacking five Masked Autoencoder for Distribution Estimation models (MADE) Germain et al. [2015]. An explanation of how MAF and MADE work is provided in Appendix C.

To test that the predicted posteriors for a given model are not incorrect we use posterior predictive checks, which is described in more details in Appendix D.

4.2.2 Handling the large dimensionality of the data with Multi-Layer Perceptron

As the dimensionality of the input data $\mathbf{x}$ grows, the complexity of the corresponding inverse problem also increases. Accurately characterizing the posterior distributions or estimating the tissue microstructure parameters becomes more challenging. As a consequence, it is often necessary to rely on a set of low-dimensional features (or summary statistics) instead of the raw data for the inference task process [Blum et al., 2013, Fearnhead and Prangle, 2012, Papamakarios et al., 2019]. These summary statistics are features that capture the essential information within the raw data, allowing to reduce the size of the input vector. Learning a set of sufficient statistics before estimating the posterior distribution makes the inference easier and offers many benefits (see e.g. the Rao-Blackwell theorem).

A follow-up challenge lies in the choice of suitable summary statistics. For well-understood problems and data, it is possible to manually design these features using deterministic functions that condense the information contained in the raw signal into a set of handful summary statistics. Previous works, such as Novikov et al. [2018b] and Jallais et al. [2022], have proposed specific

summary statistics for two different biophysical models. However, defining these summary statistics is difficult and often requires prior knowledge of the problem at hand. In the context of dMRI, they also rely on acquisition constraints and are model-specific.

In this work, the proposed framework aims to be applicable to any forward model and be as general as possible. We therefore propose to learn the summary statistics from the high-dimensional input signals $\mathbf{x}$ using a neural network. This neural network is referred to as an embedding neural network. The observed signals are fed into the embedding neural network, whose outputs are then passed to the neural density estimator. The parameters of the embedding network are learned together with the parameters of the neural density estimator, leading to the extraction of optimal features that minimize the uncertainty of $p(\boldsymbol{\theta}|\mathbf{x})$. Here, we propose to use a MLP with three layers as a summary statistics extractor. The number of features N_f extracted by the MLP can be either defined a priori or determined empirically during training.

4.2.3 Training μ GUIDE

To train μ GUIDE we need couples of input vectors $\mathbf{x}$ and corresponding ground-truth values for the model parameters that we want to estimate, $\boldsymbol{\theta}$. The input $\mathbf{x}$ can be real or simulated data (e.g. dMRI signals); or a mixture of these two. We train μ GUIDE by stochastically minimizing the loss function defined in Eq. (6) [\[Kingma and Ba, 2015\]](#) using the Adam optimizer [\[Kingma and Ba, 2015\]](#) with a learning rate of 10^{-3} and a minibatch size of 128. We use 1 million simulations for each model, 5% of which are randomly selected to be used as a validation set. Training is stopped when the validation loss does not decrease for 30 consecutive epochs.

4.2.4 Quantifying the confidence in the estimates

The full posterior distribution contains a lot of useful information about a given model parameter best estimates; uncertainty; ambiguity and degeneracy. To summarize and easily visualize this information, we propose three measures that quantify the best estimates and the associated confidence levels, and a way to highlight degeneracy.

We start by checking whether a posterior distribution is degenerate, that is if the distribution presents multiple distinct parameter solutions, appearing as multiple local maxima (**Fig. 2** [\[Kingma and Ba, 2015\]](#)). To that aim, we fit two Gaussian distributions to the obtained posterior distributions. A voxel is considered as degenerate if the derivative of the fitted Gaussian distributions changes signs more than once (i.e. multiple local maxima), and if the two Gaussian distributions are not overlapping (the distance between the two Gaussian means is inferior to the sum of their standard deviations).

For non-degenerate posterior distributions, we extract three quantities:

1. The Maximum A Posteriori (MAP), which corresponds to the most likely parameter estimate.
2. An uncertainty measure, which quantifies the dispersion of the 50% most probable samples using the interquartile range, relative to the prior range.
3. An ambiguity measure, which measures the Full Width at Half Maximum (FWHM), in percentage with respect to the prior range.

Fig. 2 [\[Kingma and Ba, 2015\]](#) presents those measures on exemplar posterior distributions.

4.3 Application of μ GUIDE to biophysical modelling of dMRI data

We show exemplar applications of μ GUIDE to three biophysical models of increasing complexity and degeneracy from the dMRI literature. For each model, we compare the fitting quality of the posterior distributions obtained using the MLP and manually defined summary statistics.

4.3.1 Biophysical models of dMRI signal.

Model 1: Ball&Stick [Behrens et al., 2003]. This is a two-compartment model (intra-neurite and extra-neurite space) where the dMRI signal from the brain tissue is modeled as a weighted sum, with weight f_{in} , of signals from water diffusing inside the neurites, approximated as sticks (i.e. cylinders of zero radius) with diffusivity D_{in} , and water diffusing within the extra-neurite space, approximated as Gaussian diffusion in an isotropic medium with diffusivity D_e . The direction of the stick is randomly sampled on a sphere. This model has the main advantage of being nondegenerate. We define the summary statistics as the direction-averaged signal (six b-shells, see Section 4.3.3).

Model 2: Standard Model (SM) [Novikov et al., 2018a]. Expanding on model 1, this model represents the dMRI signal from the brain tissue as a weighted sum of the signal from water diffusing within the neurite space, approximated as sticks with symmetric orientation dispersion following a Watson distribution and water diffusing within the extra-neurite space, modelled as anisotropic Gaussian diffusion. The microstructure parameters of this two-compartment model are the neurite signal fraction f , the intra-neurite diffusivity D_a , the orientation dispersion index ODI , and the parallel and perpendicular diffusivities within the extra-neurite space $D_e^{\parallel}$ and $D_e^{\perp}$. We use the LEMONADE [Novikov et al., 2018b] system of equations, which is based on a cumulant decomposition of the signal, to define six summary statistics.

Model 3: extended-SANDI [Palombo et al., 2020]. This is a three-compartment model (intra-neurite, intra-soma and extra-cellular space) where the dMRI signal from the brain tissue is modelled as a weighted sum of the signal from water diffusing within the neurite space, approximated as sticks with symmetric orientation dispersion following a Watson distribution; water diffusing within cell bodies (namely soma), modelled as restricted diffusion in spheres; and water diffusing within the extra-cellular space, modelled as isotropic Gaussian diffusion. The parameters of interest are the neurite signal fraction f_n , the intra-neurite diffusivity D_n , the orientation dispersion index ODI , the extra-cellular signal fraction f_e and isotropic diffusivity D_e , the soma signal fraction f_s , and a proxy of soma radius and diffusivity C_s , defined as [Jallais et al., 2022]:

$$C_s = \frac{2}{D_s \delta^2} \sum_{m=1}^{\infty} \frac{\alpha_m^{-4}}{\alpha_m^2 r_s^2 - 2} \cdot \left(2\delta - \frac{2 + e^{-\alpha_m^2 D_s (\Delta - \delta)} - e^{-\alpha_m^2 D_s \delta} - e^{-\alpha_m^2 D_s \Delta} + e^{-\alpha_m^2 D_s (\Delta + \delta)}}{\alpha_m^2 D_s} \right),$$

with r_s and D_s the soma radius and diffusivity respectively, and α_m the m th root of $(ar_s)^{-1} J_{\frac{3}{2}}(ar_s) = J_{\frac{3}{2}}(ar_s)$, with $J_n(x)$ the Bessel functions of the first kind. We use the six summary statistics defined in [Jallais et al., 2022], which are based on a high and low b-value signal expansion. Signal fractions follow the rule $f_n + f_s + f_e = 1$, leading to six parameters to estimate for this model.

Prior distributions $p(\theta)$ are defined as uniform distributions over biophysically plausible ranges. Signal fractions are defined within the interval $[0, 1]$, diffusivities between 0.1 and $3 \mu\text{m}^2 \text{ms}^{-1}$, ODI between 0.03 and 0.95 , and C_s between 0.15 and $1105 \mu\text{m}^2$ (which correspond to $r_s \in [1; 15] \mu\text{m}$ and fixed $D_s = 3 \mu\text{m}^2 \text{ms}^{-1}$).

The SM imposes the constraint $D_e^{\perp} < D_e^{\parallel}$. To generate samples uniformly distributed on the space defined by this condition, we are using two random variables u_0 and u_1 , both sampled uniformly between 0 and 1 , and then relate them to $D_e^{\parallel}$ and $D_e^{\perp}$ using the following equations:

$$\begin{cases} D_e^{\parallel} = \sqrt{(3.0 - 0.1)^2 \cdot u_0} + 0.1 \\ D_e^{\perp} = (D_e^{\parallel} - 0.1) \cdot u_1 + 0.1 \end{cases} \quad (7)$$

The extended-SANDI model requires for the signal fractions to sum to 1, that is $f_n + f_s + f_e = 1$. To uniformly cover the simplex $f_n + f_s + f_e = 1$, we define two new parameters k_1 and k_2 , uniformly sampled between 0 and 1, and use the following equations to get the corresponding signal fractions:

$$\begin{cases} f_n = k_2 \sqrt{k_1} \\ f_s = (1 - k_2) \sqrt{k_1} \\ f_e = 1 - \sqrt{k_1} \end{cases} \quad (8)$$

To ensure comparability of results, we extract the same number of features N_f using the MLP as the number of summary statistics for each model. We therefore use $N_f = 6$ for the Ball&Stick, the SM and the extended-SANDI models. Although the number of features predicted by the MLP is fixed to $N_f = 6$ for the three models, the characteristics of these six features can be very different, depending on the chosen forward model and the available data (see Appendix A). Training the MLP together with the NPE module allows to maximise inference performance in terms of accuracy and precision.

4.3.2 Validation in Numerical Simulations

We start by validating the proposed method using posterior predictive checks (PPC) and simulated signals from model 2 (see more details in Appendix D). Since PPC alone does not guarantee the correctness of the estimated posteriors, we further validated the obtained posterior distributions comparing them with the Adaptive Metropolis-Within-Gibbs MCMC [Roberts and Rosenthal, 2009]. We generated simulations following the same acquisition protocol as the real data (see Section 4.3.3), and added Gaussian noise to the real and imaginary parts of the simulated signal with a Signal-to-Noise-Ratio (SNR) of 50, and then used the magnitude of this noisy complex signal for our experiments. We then estimated the posterior distributions using both μ GUIDE and the MCMC method implemented in the MDT toolbox [Harms and Roebroek, 2018]. We initialized the sampling using a maximum likelihood estimator. We sampled 15200 samples from the distribution, the first 200 ones being used as burn-in, and no thinning. Similarly, we sampled 15000 samples from the estimated posterior distributions using μ GUIDE.

Then, we show that μ GUIDE can be applicable to any model. We use model 1 and 3 as examples of simpler (and non-degenerate) and more complex (and degenerate) models than model 2, respectively.

We compared the proposed framework to a state-of-the-art method for posterior estimation [Jallais et al., 2022]. This method relies on manually-defined summary statistics, while μ GUIDE automatically extracts some features using an embedded neural network. μ GUIDE was trained directly on noisy simulations. The manually-defined summary statistics were extracted from these simulated noisy signals and then used as training dataset for a MAF, similarly to [Jallais et al., 2022].

Finally, we used μ GUIDE to highlight degeneracy in all the models. While the complexity of the models increases, more degeneracy can be found. The degeneracy is inherent to the model definition, and is not induced by the noise. μ GUIDE allows to highlight those degeneracies and quantify the confidence in the obtained estimates.

The training was performed on $N = 10^6$ numerical simulations for each model, computed using the MISST package [lanus et al., 2017] and random combinations of the model parameters, each uniformly sampled from the previously defined ranges, with the addition of Rician distributed noise with SNR equivalent to the experimental data, i.e. 50.

4.3.3 dMRI Data Acquisition and Processing

We applied μ GUIDE to dMRI data collected from two participants: a healthy volunteer from the WAND dataset [McNabb et al., 2024] and an age-matched participant with epilepsy, acquired with the same protocol used for the MICRA dataset [Koller et al., 2021]. Data were acquired on a Connectome 3T scanner using a single-shot spin-echo, echo-planar imaging sequence with b-values = [200, 500, 1200, 2400, 4000, 6000] s mm⁻², [20, 20, 30, 61, 61, 61] uniformly distributed directions respectively and 13 non-diffusion-weighted images at 2 mm isotropic resolution. TR was set to 3000 ms, TE to 59 ms, and the diffusion gradient duration and separation to 7 ms and 24 ms respectively. Short diffusion times and TE were achieved thanks to the Connectom gradients, allowing to enhance the signal-to-noise ratio and sensitivity to small water displacements [Jones et al., 2018, Setsompop et al., 2013]. We considered the noise as Rician with an SNR of 50 for both subjects.

Data were preprocessed using a combination of in-house pipelines and tools from the FSL [Andersson et al., 2003, Andersson and Sotiropoulos, 2016, Smith, 2002, Smith et al., 2004] and MRTrx3 [Tournier et al., 2019] software packages. The preprocessing steps included brain extraction [Smith, 2002], denoising [Cordero-Grande et al., 2019, Veraart et al., 2016], drift correction [Vos et al., 2017, Sairanen et al., 2018], susceptibility-induced distortions [Andersson et al., 2003, Smith et al., 2004], motion and eddy current correction [Andersson and Sotiropoulos, 2016], correction for gradient non-linearity distortions [Glasser et al., 2013], and Gibbs ringing artefacts correction [Kellner et al., 2016].

4.3.4 dMRI Data Analysis

Diffusion signals were first normalized by the mean non-diffusion-weighted signals acquired for each voxel. Each voxel was then estimated in parallel using the μ GUIDE framework. For each observed signal $\mathbf{x}_v$ (i.e. for each voxel), we drew 50000 samples via rejection-sampling from $q_\phi(\theta_i | \mathbf{x}_v)$ for each model parameter θ_i , allowing to retrieve the full posterior distributions. If a posterior distribution was not deemed degenerate, the maximum-a-posteriori, uncertainty and ambiguity measures were extracted from the posterior distributions.

The manually-defined summary statistics of the SM are defined using a cumulant expansion, which is only valid for small b-values. We therefore only used the $b \leq 2500$ s mm⁻² data for this model. In order to obtain comparable results, we restricted the application of μ GUIDE to this range of b-values as well. An extra b-shell (b-value = 5000 s mm⁻²; 61 directions) was interpolated using *mapl* [Fick et al., 2016] for the extended-SANDI model when using the method developed by Jallais et al. [2022] based on summary statistics.

The training of μ GUIDE was performed as described in 4.3.2 and an example of training dataset and input signal vector is provided in Fig. 8.

All the computations were performed both on CPU and GPU (NVIDIA GeForce RTX 4090).

Acknowledgements

This work, Maëliss Jallais and Marco Palombo are supported by UKRI Future Leaders Fellowship (MR/T020296/2). We are thankful to Dr. Dmitri Sastin and Dr. Khalid Hamandi for sharing their dataset from a participant with epilepsy, and to Dr. Carolyn McNabb, Dr. Eirini Messar-itaki and Dr. Pedro Luque Laguna for preprocessing the data of the healthy participant from the WAND data. The WAND data were acquired at the UK National Facility for In Vivo MR Imaging of Human

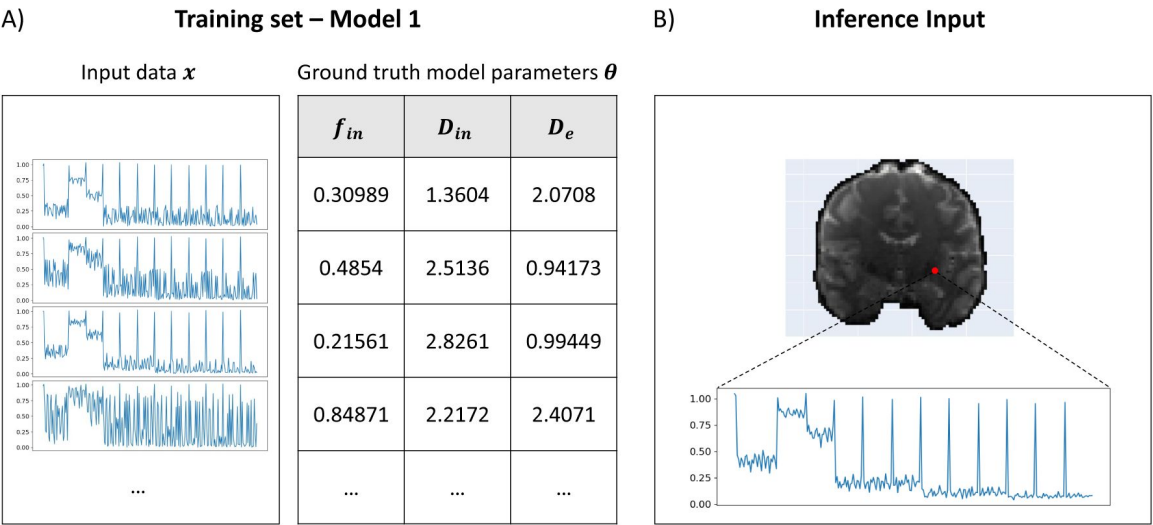


Figure 8

Example training set and input signals for μ GUIDE.

A) Examples of input synthetic data vectors and corresponding ground truth model parameters used in the training set of Model 1 (Ball&Stick). B) Example of input measured signals from a voxel in a healthy participant, used for inference.

Appendix

A. A Correlation analysis between features extracted by μ GUIDE and manually-defined summary statistics

Fig. 9 [↗](#) presents the correlation matrices obtained from the correlation between the MLP-extracted features from μ GUIDE and the manually-defined summary statistics defined in Section 4.3, considering noisy simulations (SNR=50). For each model, at least one feature extracted by μ GUIDE is not or weakly correlated with the summary statistics. Additional information, not contained in the summary statistics, is extracted by the MLP from the input signal, leading to reduced bias, uncertainty and ambiguity in the parameter estimates (see **Fig. 4** [↗](#)).

B. Impact of noise on the posterior distributions

Noise in the signal impacts the fitting quality of a biophysical model. **Fig. 10A** [↗](#) shows example posterior distributions for one combination of model 2 parameters, with varying noise levels (no noise, SNR = 50 and SNR = 25). **Fig. 10B** [↗](#) presents uncertainties values obtained on 1000 simulations with varying SNRs. We observe that, as the SNR reduces (i.e. as the noise increases), uncertainty increases. Noise in the signal contributes to irreducible variance. The confidence in the estimates therefore reduces as the noise level increases.

C. Masked Autoregressive Flows

The conditional probability density function approximators used in this project belong to a class of neural networks called Normalizing Flows (NF) Papamakarios et al. [2021]. NFs provide a general way of transforming complex probability distributions over continuous random variables into simple base distributions $p(z)$ (such as normal distributions) through a chain of invertible and differentiable transformations f_ϕ . By applying the change of variable formula, the target distribution $q_\phi(\theta | x)$ can be written as:

$$q_\phi(\theta | x) = p(f_\phi(\theta; x)) |\det J_{f_\phi}(\theta; x)|, \quad (9)$$

where $z = f_\phi(\theta; x)$ is invertible and differentiable (i.e. a diffeomorphism), and $J_{f_\phi}(\theta; x)$ is the Jacobian of $f_\phi(\theta; x)$. The forward direction allows for density evaluation, that is learning the mapping between the target and the base distributions, i.e. learning the parameters ϕ . The inverse direction allows to estimate a density estimator $q_\phi(\theta | x_0)$ by sampling points z from the base distribution and applying the inverse transform $f_\phi^{-1}(z; x_0)$.

A main requirement is that the flow needs to be expressive enough to approximate any arbitrarily highly complex distribution. An interesting property of diffeomorphisms is that they are closed under composition, which means that a composition of K diffeomorphisms $f = f_1 \circ f_2 \circ \dots \circ f_K$ is also a diffeomorphism, and the Jacobian determinant is the product of the determinant of each component. Combining multiple transformations allows to increase the expressivity of the general flow. We obtain:

$$\begin{aligned} q_\phi(\theta | x) &= p(f_\phi(\theta; x)) \log \left| \det \prod_{k=1}^K J_{f_k}(\theta_i; x_i) \right| \\ &= p(f_\phi(\theta; x)) \sum_{k=1}^K \log |\det J_{f_k}(\theta_i; x_i)| \end{aligned} \quad (10)$$

Figure 9

Correlation matrices between features extracted by the MLP in μ GUIDE and manually-defined summary features for the three models.

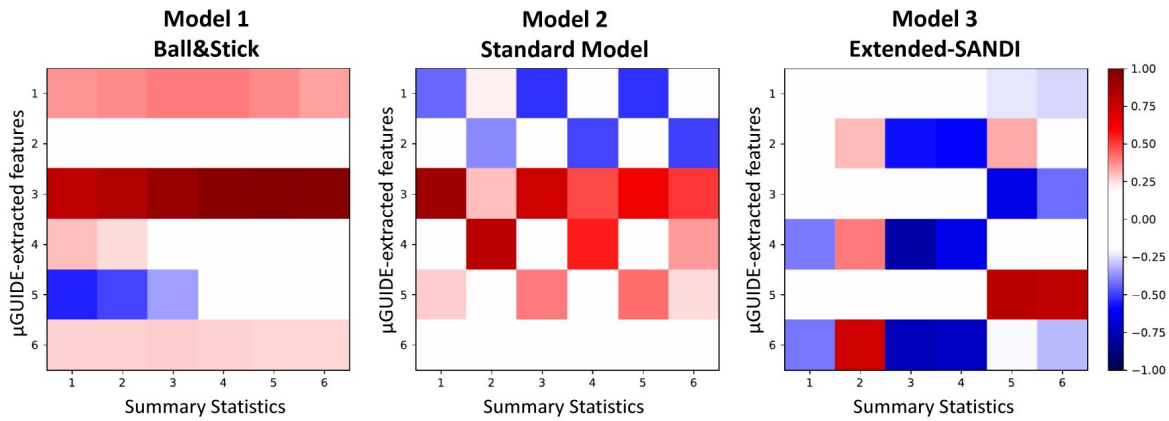
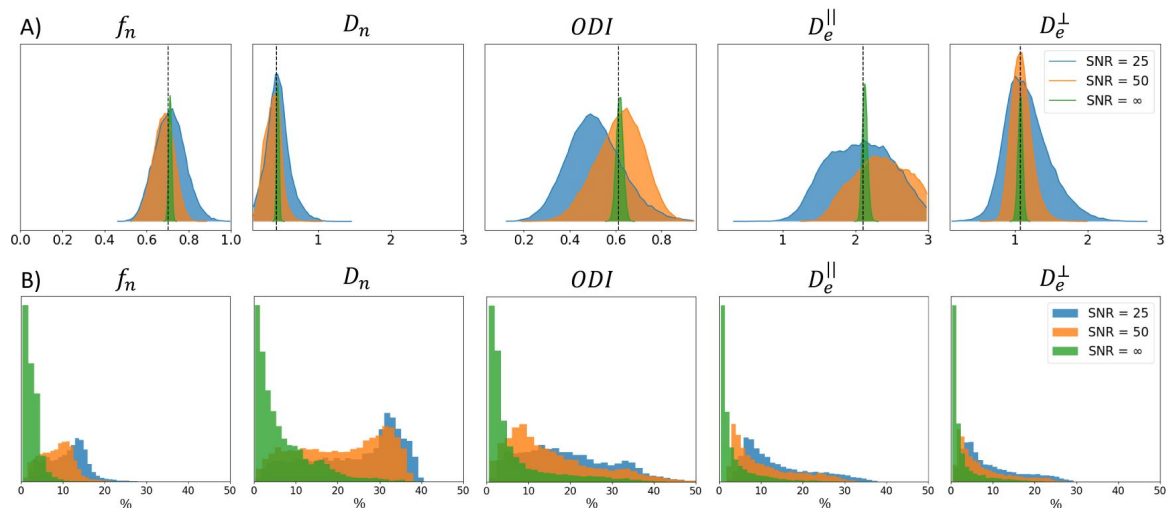


Figure 10

SNR Uncertainty comparison between signals with different noise levels: no noise, $SNR = 50$ and $SNR = 25$, using model 2. A) Posterior distributions obtained on one example parameter combination (vertical black dashed line) with the three noise levels. B) Histogram of the uncertainty obtained for 1000 signals with different noise levels (in %). Similar ground truths are used for each noise level.



Flows need to be flexible and expressive enough to model any desired distribution but also need to be computationally efficient, that is, computing the associated Jacobian determinants need to be tractable and efficient. Among a number of proposed architectures such as mixture density networks [Bishop, 1994 [↗](#)] or neural spline flows [Durkan et al., 2019 [↗](#)], we focused on Masked Autoregressive Flows (MAF) [Papamakarios et al., 2017 [↗](#)], which has shown state-of-the-art performance as well as the ability to estimate multi-modal posterior distributions [Goncalves et al., 2020 [↗](#), Patron et al., 2022 [↗](#), Papamakarios et al., 2021 [↗](#)].

Autoregressive flows are universal approximators and have the form $z'_i = \tau(z_{<i}; \mathbf{h}_i)$ where $\mathbf{h}_i = c_i(\mathbf{z}_{<i})$ [Papamakarios et al., 2021 [↗](#)]. τ is termed the *transformer* and is a strictly monotonic function parametrized by $\mathbf{h}_i$ and c_i the *i*-th *conditioner*. Each $\mathbf{h}_i$ and therefore each z'_i can be computed independently in parallel, helping to keep a low computation time. The conditioner constraints each output to depend only on variables with dimension indices less than *i*, which makes the Jacobian of the flow lower diagonal. Its determinant can then be obtained easily as the product of its diagonal elements.

To efficiently implement the conditioner, this method relies on the Masked Autoencoder for Distribution Estimation (MADE) [Germain et al., 2015 [↗](#)] architecture. To create a neural network that obey the autoregressive structure of the conditioner, a fully connected feedforward neural network is multiplied to binary masks, which removes some connections by assigning them a weigh of 0. The binary masks can easily be obtained by following a few simple steps (see **Fig. 11** [↗](#) for an illustration):

1. Label the input and output nodes between 1 and *D*, *D* being the dimension of the input vector *z*.
2. Randomly assign each hidden unit a number between 1 and *D*-1, which indicates the number of inputs it will be connected to.
3. For each hidden layer, connect the hidden units to units with inferior or equal labels.
4. Connect output units to units with strictly inferior labels.

For μ GUIDE's implementation, we use a combination of 5 MADEs.

As a result, the MAF architecture only needs a single forward pass through the flow and, combined with the low-cost computation of the determinant, allows for fast training and evaluation of the posterior distributions.

D Posterior Predictive Checks

Posterior Predictive Checks (PPC) are a common safety check to verify inference is not wrong. The idea is to compare input signals with generated signals from samples drawn from the posterior distributions. If the inference is correct, the generated signals should look similar to the input signal.

We performed the following steps:

- Sample *N* θ_i from the prior distribution: $\theta_i \sim p(\theta)$
- Generate the corresponding signals using the forward model: $\mathbf{x}_i = \mathbf{M}(\theta_i)$
- Perform the inference and estimate the posterior distributions $p(\theta_i | \mathbf{x}_i)$
- Sample *N_{pp}* samples $\theta_{i,s}$ from $p(\theta_i | \mathbf{x}_i)$
- Reconstruct the signals from the sampled $\theta_{i,s}$ using the forward model: $\mathbf{x}_{i,s} = \mathbf{M}(\theta_{i,s})$
- Compare the obtained $\mathbf{x}_{i,s}$ with $\mathbf{x}_i$.

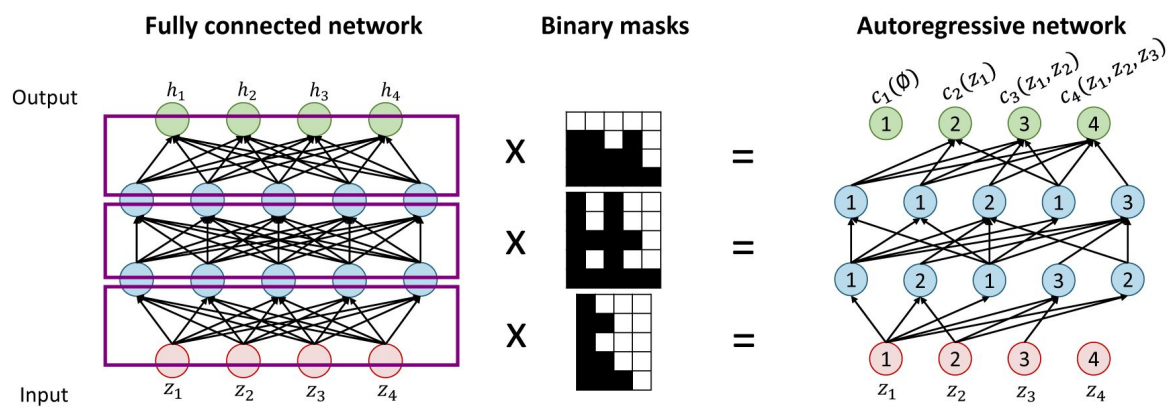


Figure 11

Schematic of MADE autoregressive network construction.

Fig. 12 [↗](#) presents results on model 2 (SM), on both noise-free and noisy signals (Rician noise with SNR=50) for $N=10$ random combinations of model parameters, and $N_{pp} = 100$. As dMRI data have a high dimensionality, we report the direction-average signal. Plain lines show the signals $\mathbf{x}_i$, and the shaded areas correspond to the area in which the corresponding $\mathbf{x}_{i,s}$ fall. $\mathbf{x}_i$ lie within the support of $\mathbf{x}_{i,s}$, indicating the inference is not wrong. Note that the support of $\mathbf{x}_{i,s}$ is bigger for noisy simulations, reflecting the wider posterior distributions obtained from the inference.

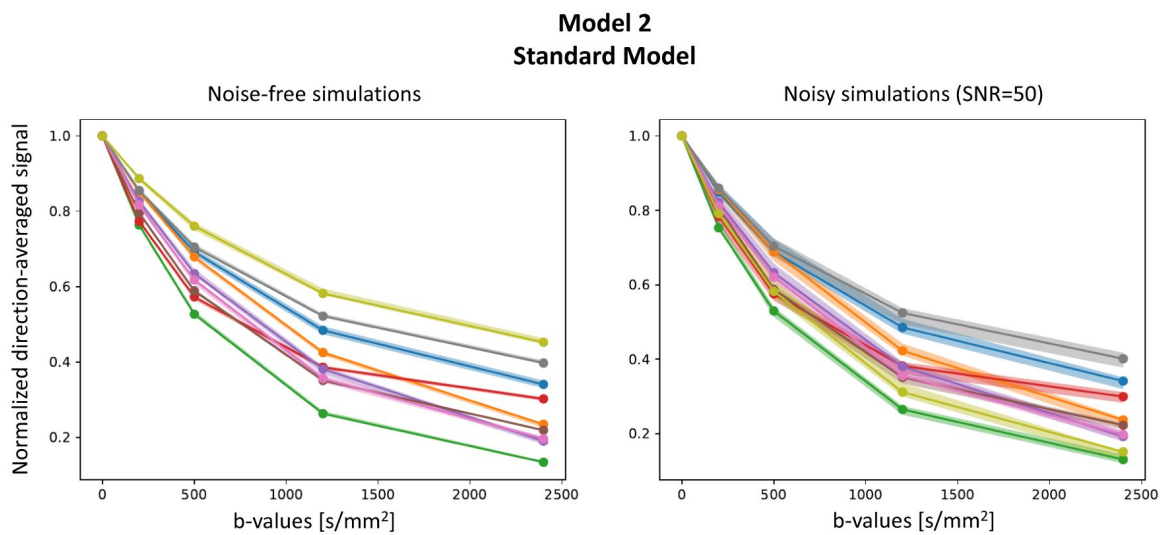


Figure 12

Posterior Predictive Checks.

Comparison between signals $\mathbf{x}_i$ generated using random parameter combinations and their reconstructions using samples from $p(\boldsymbol{\theta}_i | \mathbf{x}_i)$.

References

- Afzali Maryam, Nilsson Markus, Palombo Marco, Jones Derek K (2021) **SPHERIOUSLY? the challenges of estimating sphere radius non-invasively in the human brain from diffusion MRI** *NeuroImage* **237** <https://doi.org/10.1016/j.neuroimage.2021.118183>
- Alexander Daniel C. (2008) **A general framework for experiment design in diffusion MRI and its application in measuring direct tissue-microstructure features** *Magnetic Resonance in Medicine* **60**:439–448 <https://doi.org/10.1002/mrm.21646>
- Alexander Daniel C., Laidlaw David, Weickert Joachim (2009) **Modelling, fitting and sampling in diffusion MRI** *Visualization and Processing of Tensor Fields* Springer Berlin Heidelberg :3–20 <https://doi.org/10.1007/978-3-540-88378-4-1>
- Alexander Daniel C., Dyrby Tim B., Nilsson Markus, Zhang Hui (2019) **Imaging brain microstructure with diffusion MRI: practicality and applications** *NMR in Biomedicine* **32** <https://doi.org/10.1002/nbm.3841>
- Andersson Jesper L.R., Sotiropoulos Stamatios N. (2016) **An integrated approach to correction for off-resonance effects and subject movement in diffusion MR imaging** *NeuroImage* **125**:1063–1078 <https://doi.org/10.1016/j.neuroimage.2015.10.019>
- Andersson Jesper L.R., Skare Stefan, Ashburner John (2003) **How to correct susceptibility distortions in spin-echo echo-planar images: application to diffusion tensor imaging** *NeuroImage* **20**:870–888 [https://doi.org/10.1016/S1053-8119\(03\)00336-7](https://doi.org/10.1016/S1053-8119(03)00336-7)
- Ascoli Giorgio A, Donohue Duncan E, Halavi Maryam (2007) **Neuromorpho. org: a central resource for neuronal morphologies** *Journal of Neuroscience* **27**:9247–9251
- Behrens T.E.J., Woolrich M.W., Jenkinson M., Johansen-Berg H., Nunes R.G., Clare S., Matthews P.M., Brady J.M., Smith S.M. (2003) **Characterization and propagation of uncertainty in diffusion-weighted MR imaging** *Magnetic Resonance in Medicine* **50**:1077–1088 <https://doi.org/10.1002/mrm.10609>
- Bishop Christopher M. (1994) **Mixture density networks**
- Blum M. G. B., Nunes M. A., Prangle D., Sisson S. A. (2013) **A comparative review of dimension reduction methods in approximate bayesian computation** *Statistical Science* **28** <https://doi.org/10.1214/12-STS406>
- Box George EP, Tiao George C (2011) **Bayesian inference in statistical analysis** John Wiley & Sons
- Callaghan Ross, Alexander Daniel C., Palombo Marco, Zhang Hui (2020) **Config: Contextual fibre growth to generate realistic axonal packing for diffusion mri simulation** *NeuroImage* **220** <https://doi.org/10.1016/j.neuroimage.2020.117107>
- Chung SungWon, Lu Ying, Henry Roland G. (2006) **Comparison of bootstrap approaches for estimation of uncertainties of DTI parameters** *NeuroImage* **33**:531–541 <https://doi.org/10.1016/j.neuroimage.2006.07.001>

- Cordero-Grande Lucilio, Christiaens Daan, Hutter Jana, Price Anthony N., Hajnal Jo V. (2019) **Complex diffusion-weighted image estimation via matrix recovery under general noise models** *NeuroImage* **200**:391–404 <https://doi.org/10.1016/j.neuroimage.2019.06.039>
- Cranmer Kyle, Pavez Juan, Louppe Gilles (2016) **Approximating likelihood ratios with calibrated discriminative classifiers** *arXiv*
- Cranmer Kyle, Brehmer Johann, Louppe Gilles (2020) **The frontier of simulation-based inference** *Proceedings of the National Academy of Sciences* <https://doi.org/10.1073/pnas.1912789117>
- Martins João P. De Almeida, Nilsson Markus, Lampmen Bjorn, Palombo Marco, While Peter T., Westin Carl-Fredrik, Szczepankiewicz Filip (2021) **Neural networks for parameter estimation in microstructural MRI: Application to a diffusion-relaxation model of white matter** *NeuroImage* **244** <https://doi.org/10.1016/j.neuroimage.2021.118601>
- Deoni Sean CL, Rutt Brian K, Arun Tarunya, Pierpaoli Carlo, Jones Derek K (2008) **Gleaning multicomponent t1 and t2 information from steady-state imaging data** *Magnetic Resonance in Medicine: An Official Journal of the International Society for Magnetic Resonance in Medicine Wiley Online Library* **60**:1372–1387
- Dietrich Olaf, Raya Jose G., Reeder Scott B., Reiser Maximilian F., Schoenberg Stefan O. (2007) **Measurement of signal-to-noise ratios in MR images: Influence of multichannel coils, parallel imaging, and reconstruction filters** *Journal of Magnetic Resonance Imaging* **26**:375–385 <https://doi.org/10.1002/jmri.20969>
- Diggle Peter J., Gratton Richard J. (1984) **Monte carlo methods of inference for implicit statistical models** *Journal of the Royal Statistical Society: Series B (Methodological)* **46**:193–212 <https://doi.org/10.1111/j.2517-6161.1984.tb01290.x>
- Durkan Conor, Bekasov Artur, Murray Iain, Papamakarios George (2019) **Neural spline flows** *Advances in neural information processing systems* **32**
- Fearnhead Paul, Prangle Dennis (2012) **Constructing summary statistics for approximate bayesian computation: Semi-automatic approximate bayesian computation** *Journal of the Royal Statistical Society Series B: Statistical Methodology* **74**:419–474 <https://doi.org/10.1111/j.1467-9868.2011.01010.x>
- Fick Rutger H.J., Wassermann Demian, Caruyer Emmanuel, Deriche Rachid (2016) **MAPL: Tissue microstructure estimation using laplacian-regularized MAP-MRI and its application to HCP data** *NeuroImage* **134**:365–385 <https://doi.org/10.1016/j.neuroimage.2016.03.046>
- Germain Mathieu, Gregor Karol, Murray Iain, Larochelle Hugo (2015) **Made: Masked autoencoder for distribution estimation** *International conference on machine learning PMLR* :881–889
- Glasser Matthew F. *et al.* (2013) **The minimal preprocessing pipelines for the human connectome project** *NeuroImage* **80**:105–124 <https://doi.org/10.1016/j.neuroimage.2013.04.127>
- Goncalves Pedro J *et al.* (2020) **Training deep neural density estimators to identify mechanistic models of neural dynamics** *eLife* **9** <https://doi.org/10.7554/eLife.56261>

Goodfellow Ian J., Pouget-Abadie Jean, Mirza Mehdi, Xu Bing, Warde-Farley David, Ozair Sherjil, Courville Aaron, Bengio Yoshua (2014) **Generative adversarial networks** *arXiv*

Greenberg David S., Nonnenmacher Marcel, Macke Jakob H. (2019) **Automatic posterior transformation for likelihood-free inference** *arXiv*

Guerreri Michele, Epstein Sean, Azadbakht Hojjat, Zhang Hui (2023) **Resolving quantitative MRI model degeneracy with machine learning via training data distribution design** *arXiv*

Gutmann Michael U., Dutta Ritabrata, Kaski Samuel, Corander Jukka (2018) **Likelihood-free inference via classification** *Statistics and Computing* **28**:411–425 <https://doi.org/10.1007/s11222-017-9738-6>

Gyori Noemi G., Palombo Marco, Clark Christopher A., Zhang Hui, Alexander Daniel C. (2022) **Training data distribution significantly impacts the estimation of tissue microstructure with machine learning** *Magnetic Resonance in Medicine* **87**:932–947 <https://doi.org/10.1002/mrm.29014>

Harms Robbert L., Roebroek Alard (2018) **Robust and fast markov chain monte carlo sampling of diffusion MRI microstructure models** *Frontiers in Neuroinformatics* **12** <https://doi.org/10.3389/fninf.2018.00097>

Henriques Rafael N., Palombo Marco, Jespersen Sune N., Shemesh Noam, Lundell Henrik, Ianus Andrada (2021) **Double diffusion encoding and applications for biomedical imaging** *Journal of Neuroscience Methods* **348** <https://doi.org/10.1016/j.jneumeth.2020.108989>

Howard Amy Fd et al. (2022) **Estimating axial diffusivity in the NODDI model** *NeuroImage* **262** <https://doi.org/10.1016/j.neuroimage.2022.119535>

Ianus Andrada, Shemesh Noam, Alexander Daniel C., Drobnjak Ivana (2017) **Double oscillating diffusion encoding and sensitivity to microscopic anisotropy: Double oscillating diffusion encoding** *Magnetic Resonance in Medicine* **78**:550–564 <https://doi.org/10.1002/mrm.26393>

Jallais Maëliss, Rodrigues Pedro L. C., Gramfort Alexandre, Wassermann Demian (2022) **Inverting brain grey matter models with likelihood-free inference: a tool for trustable cytoarchitecture measurements** *Machine Learning for Biomedical Imaging* **1**:1–28 <https://doi.org/10.59275/j.melba.2022-a964>

Jallais Maëliss, Palombo Marco, Jelescu Ileana, Uhl Quentin (2024) **Shining light on degeneracies and uncertainties in the NEXI and SANDIX models with pGUIDE** *ISMRM 2024*

Jelescu Ileana O., Budde Matthew D. (2017) **Design and validation of diffusion MRI models of white matter** *Frontiers in Physics* **5** <https://doi.org/10.3389/fphy.2017.00061>

Jelescu Ileana O., Veraart Jelle, Fieremans Els, Novikov Dmitry S. (2016) **Degeneracy in model parameter estimation for multi-compartmental diffusion in neuronal tissue: Degeneracy in model parameter estimation of diffusion in neural tissue** *NMR in Biomedicine* **29**:33–47 <https://doi.org/10.1002/nbm.3450>

Jelescu Ileana O., Palombo Marco, Bagnato Francesca, Schilling Kurt G. (2020) **Challenges for biophysical modeling of microstructure** *Journal of Neuroscience Methods* **344** <https://doi.org/10.1016/j.jneumeth.2020.108861>

- Jelescu Ileana O., Skowronski Alexandre de, Geffroy Françoise, Palombo Marco, Novikov Dmitry S. (2022) **Neurite exchange imaging (NEXI): A minimal model of diffusion in gray matter with inter-compartment water exchange** *NeuroImage* **256** <https://doi.org/10.1016/j.neuroimage.2022.119277>
- Jones Derek K. (2003) **Determining and visualizing uncertainty in estimates of fiber orientation from diffusion tensor MRI** *Magnetic Resonance in Medicine* **49**:7–12 <https://doi.org/10.1002/mrm.10331>
- Jones D.K. *et al.* (2018) **Microstructural imaging of the human brain with a ‘super-scanner’: 10 key advantages of ultra-strong gradients for diffusion MRI** *NeuroImage* **182**:8–38 <https://doi.org/10.1016/j.neuroimage.2018.05.047>
- Karimi Hazhar Sufi, Pal Arghya, Ning Lipeng, Rathu Yogesh (2024) **Likelihood-free posterior estimation and uncertainty quantification for diffusion MRI models** *Imaging Neuroscience* **2**:1–22 https://doi.org/10.1162/imag_a_00088
- Kauermann Göran, Claeskens Gerda, Opsomer J. D. (2009) **Bootstrapping for penalized spline regression** *Journal of Computational and Graphical Statistics* **18**:126–146 <https://doi.org/10.1198/jcgs.2009.0008>
- Kellner Elias, Dhital Bibek, Kiselev Valerij G., Reiser Marco (2016) **Gibbs-ringing artifact removal based on local subvoxel-shifts** *Magnetic Resonance in Medicine* **76**:1574–1581 <https://doi.org/10.1002/mrm.26054>
- Kingma Diederik, Ba Jimmy (2015) **Adam: A method for stochastic optimization** *International Conference on Learning Representations (ICLR)*
- Kingma Diederik P., Welling Max (2019) **An introduction to variational autoencoders** *Foundations and Trends® in Machine Learning* **12**:307–392 <https://doi.org/10.1561/22000000056>
- Koch Martin A., Norris David G., Hund-Georgiadis Margret (2002) **An investigation of functional and anatomical connectivity using magnetic resonance imaging** *NeuroImage* **16**:241–250 <https://doi.org/10.1006/nimg.2001.1052>
- Koller Kristin *et al.* (2021) **MICRA: Microstructural image compilation with repeated acquisitions** *NeuroImage* **225** <https://doi.org/10.1016/j.neuroimage.2020.117406>
- Lampinen Björn, Szczepankiewicz Filip, Lätt Jimmy, Knutsson Linda, Mörtensson Johan, Björkman-Burtscher Isabella M., Westén Danielle Van, Sundgren Pia C., Ståhlberg Freddy, Nilsson Markus (2023) **Probing brain tissue microstructure with MRI: principles, challenges, and the role of multidimensional diffusion-relaxation encoding** *NeuroImage* **282** <https://doi.org/10.1016/j.neuroimage.2023.120338>
- Lazar Mariana, Alexander Andrew L. (2005) **Bootstrap white matter tractography (BOOT-TRAC)** *NeuroImage* **24**:524–532 <https://doi.org/10.1016/j.neuroimage.2004.08.050>
- Lueckmann Jan-Matthis *et al.* (2017) **Flexible statistical inference for mechanistic models of neural dynamics** *Advances in Neural Information Processing Systems* Curran Associates, Inc. **30**
- Lueckmann Jan-Matthis, Bassetto Giacomo, Karaletsos Theofanis, Macke Jakob H., Ruiz Francisco, Zhang Cheng, Liang Dawen, Bui Thang (2019) **Likelihood-free inference with emulator networks** *[i]Proceedings of The 1st Symposium on Advances in Approximate Bayesian Inference[i]*, volume 96 of *[i]Proceedings of Machine Learning Research[i]* PMLR :32–53

Lueckmann Jan-Matthis, Boelts Jan, Greenberg David, Goncalves Pedro, Macke Jakob, Banerjee Arindam, Fukumizu Kenji (2021) **Benchmarking simulation-based inference** [i] *Proceedings of The 24th International Conference on Artificial Intelligence and Statistics* [i], volume 130 of [i] *Proceedings of Machine Learning Research* [i] PMLR :343–351

McNabb Carolyn B *et al.* (2024) **The welsh advanced neuroimaging database: an open-source state-of-the-art resource for brain research** *ISMRM 2024*

Metropolis Nicholas, Rosenbluth Arianna W., Rosenbluth Marshall N., Teller Augusta H., Teller Edward (1953) **Equation of state calculations by fast computing machines** *The Journal of Chemical Physics* **21**:1087–1092 <https://doi.org/10.1063/1.1699114>

Mougel Eloïse, Valette Julien, Palombo Marco (2024) **Investigating exchange, structural disorder, and restriction in gray matter via water and metabolites diffusivity and kurtosis time-dependence** *Imaging Neuroscience* **2**:1–14

Novikov Dmitry S., Fieremans Els, Jespersen Sune N., Kiselev Valerij G. (2018) **Quantifying brain microstructure with diffusion MRI: Theory and parameter estimation: Brain microstructure with dMRI: Theory and parameter estimation** *NMR in Biomedicine* <https://doi.org/10.1002/nbm.3998>

Novikov Dmitry S., Veraart Jelle, Jelescu Ileana O., Fieremans Els (2018) **Rotationally-invariant mapping of scalar and orientational metrics of neuronal microstructure with diffusion MRI** *NeuroImage* **174**:518–538 <https://doi.org/10.1016/j.neuroimage.2018.03.006>

Olesen Jonas L., Østergaard Leif, Shemesh Noam, Jespersen Sune N. (2022) **Diffusion time dependence, power-law scaling, and exchange in gray matter** *NeuroImage* **251** <https://doi.org/10.1016/j.neuroimage.2022.118976>

Palombo Marco, Ianus Andrada, Guerreri Michele, Nunes Daniel, Alexander Daniel C., Shemesh Noam, Zhang Hui (2020) **SANDI: A compartment-based model for non-invasive apparent soma and neurite imaging by diffusion MRI** *NeuroImage* **215** <https://doi.org/10.1016/j.neuroimage.2020.116835>

Palombo Marco *et al.* (2023) **Joint estimation of relaxation and diffusion tissue parameters for prostate cancer with relaxation-VERDICT MRI** *Scientific Reports* **13** <https://doi.org/10.1038/s41598-023-30182-1>

Panagiotaki Eleftheria, Schneider Torben, Siow Bernard, Hall Matt G., Lythgoe Mark F., Alexander Daniel C. (2012) **Compartment models of the diffusion MR signal in brain white matter: A taxonomy and comparison** *NeuroImage* **59**:2241–2254 <https://doi.org/10.1016/j.neuroimage.2011.09.081>

Panagiotaki Eleftheria, Walker-Samuel Simon, Siow Bernard, Johnson S. Peter, Rajkumar Vineeth, Pedley R. Barbara, Lythgoe Mark F., Alexander Daniel C. (2014) **Noninvasive quantification of solid tumor microstructure using VERDICT MRI** *Cancer Research* **74**:1902–1912 <https://doi.org/10.1158/0008-5472.CAN-13-2511>

Papamakarios George, Murray Iain (2016) **Fast ϵ -free inference of simulation models with bayesian conditional density estimation** *Advances in neural information processing systems* **29**

Papamakarios George, Pavlakou Theo, Murray Iain (2017) **Masked autoregressive flow for density estimation** *Advances in neural information processing systems* **30**

- Papamakarios George, Sterratt David, Murray Iain (2019) **Sequential neural likelihood: Fast likelihood-free inference with autoregressive flows** *The 22nd international conference on artificial intelligence and statistics* PMLR :837–848
- Papamakarios George, Nalisnick Eric, Rezende Danilo Jimenez, Mohamed Shakir, Lakshminarayanan Balaji (2021) **Normalizing flows for probabilistic modeling and inference** *The Journal of Machine Learning Research* **22**:2617–2680
- Parker Geoff J. M., Alexander Daniel C., Taylor Chris, Noble J. Alison (2003) **Probabilistic monte carlo based mapping of cerebral connections utilising whole-brain crossing fibre information** *Information Processing in Medical Imaging* Springer Berlin Heidelberg **2732**:684–695 <https://doi.org/10.1007/978-3-540-45087-057>
- Patron Jose Pedro Manzano, Kypraios Theodore, Sotiropoulos Stamatios N (2022) **Amortised inference in diffusion mri biophysical models using artificial neural networks and simulation-based frameworks** *ISMRM 2022*
- Roberts Gareth O., Rosenthal Jeffrey S. (2009) **Examples of adaptive MCMC** *Journal of Computational and Graphical Statistics* **18**:349–367 <https://doi.org/10.1198/jcgs.2009.06134>
- Sairanen Viljami, Leemans Alexander, Tax Chantal MW (2018) **Fast and accurate slice-wise outlier detection (solid) with informed model estimation for diffusion mri data** *NeuroImage* **181**
- Setsompop K. *et al.* (2013) **Pushing the limits of in vivo diffusion MRI for the human connectome project** *NeuroImage* **80**:220–233 <https://doi.org/10.1016/j.neuroimage.2013.05.078>
- Slator Paddy J. *et al.* (2021) **Combined diffusion-relaxometry microstructure imaging: Current status and future prospects** *Magnetic Resonance in Medicine* **86**:2987–3011 <https://doi.org/10.1002/mrm.28963>
- Smith Stephen M. (2002) **Fast robust automated brain extraction** *Human Brain Mapping* **17**:143–155 <https://doi.org/10.1002/hbm.10062>
- Smith Stephen M *et al.* (2004) **Advances in functional and structural mr image analysis and implementation as fsl** *NeuroImage* **23**:S208–S219
- Sotiropoulos S. N., Jbabdi S., Andersson J. L., Woolrich M. W., Ugurbil K., Behrens T. E. J. (2013) **RubiX: Combining spatial resolutions for bayesian inference of crossing fibers in diffusion MRI** *IEEE Transactions on Medical Imaging* **32**:969–982 <https://doi.org/10.1109/TMI.2012.2231873>
- Tejero-Cantero Alvaro, Boelts Jan, Deistler Michael, Lueckmann Jan-Matthis, Durkan Conor, Goncalves Pedro J., Greenberg David S., Macke Jakob H. (2020) **SBI - a toolkit for simulation-based inference** *arXiv*
- Tournier J-Donald, Smith Robert, Raffelt David, Tabbara Rami, Dhollander Thijs, Pietsch Maximilian, Christiaens Daan, Jeurissen Ben, Yeh Chun-Hung, Connelly Alan (2019) **Mrtrix3: A fast, flexible and open software framework for medical image processing and visualisation** *NeuroImage* **202**

Uhl Quentin, Pavan Tommaso, Molendowska Malwina, Jones Derek K, Palombo Marco, Jelescu Ileana (2024) **Quantifying human gray matter microstructure using neurite exchange imaging (nexi) and 300 mt/m gradients** *Imaging Neuroscience*

Maaten Laurens Van der, Hinton Geoffrey (2008) **Visualizing data using t-SNE** *Journal of machine learning research* **9**

Veraart Jelle, Novikov Dmitry S., Christiaens Daan, Ades-aron Benjamin, Sijbers Jan, Fieremans Els (2016) **Denoising of diffusion MRI using random matrix theory** *NeuroImage* **142**:394–406 <https://doi.org/10.1016/j.neuroimage.2016.08.016>

Vincent Mélissa, Palombo Marco, Valette Julien (2020) **Revisiting double diffusion encoding MRS in the mouse brain at 11.7t: Which microstructural features are we sensitive to?** *NeuroImage* **207** <https://doi.org/10.1016/j.neuroimage.2019.116399>

Vos Sjoerd B., Tax Chantal M. W., Luijten Peter R., Ourselin Sebastien, Leemans Alexander, Froeling Martijn (2017) **The importance of correcting for signal drift in diffusion MRI** *Magnetic Resonance in Medicine* **77**:285–299 <https://doi.org/10.1002/mrm.26124>

Warner William, Palombo Marco, Cruz Renata, Callaghan Ross, Shemesh Noam, Jones Derek K., Dell'Acqua Flavio, Ianus Andrada, Drobnjak Ivana (2023) **Temporal diffusion ratio (TDR) for imaging restricted diffusion: Optimisation and pre-clinical demonstration** *NeuroImage* **269** <https://doi.org/10.1016/j.neuroimage.2023.119930>

Whitcher Brandon, Tuch David S., Wisco Jonathan J., Sorensen A. Gregory, Wang Liqun (2008) **Using the wild bootstrap to quantify uncertainty in diffusion tensor imaging** *Human Brain Mapping* **29**:346–362 <https://doi.org/10.1002/hbm.20395>

Yablonskiy Dmitriy A., Sukstanskii Alexander L. (2010) **Theoretical models of the diffusion weighted MR signal** *NMR in Biomedicine* **23**:661–681 <https://doi.org/10.1002/nbm.1520>

Zhang Hui, Schneider Torben, Wheeler-Kingshott Claudia A., Alexander Daniel C. (2012) **NODDI: Practical in vivo neurite orientation dispersion and density imaging of the human brain** *NeuroImage* **61**:1000–1016 <https://doi.org/10.1016/j.neuroimage.2012.03.072>

Editors

Reviewing Editor

Jing Sui

Beijing Normal University, Beijing, China

Senior Editor

Aleksandra Walczak

École Normale Supérieure - PSL, Paris, France

Reviewer #1 (Public review):

The authors proposed a framework to estimate the posterior distribution of parameters in biophysical models. The framework has two modules: the first MLP module is used to reduce data dimensionality and the second NPE module is used to approximate the desired posterior distribution. The results show that the MLP module can capture additional information compared to manually defined summary statistics. By using the NPE module, the repetitive

evaluation of the forward model is avoided, thus making the framework computationally efficient. The results show the framework has promise in identifying degeneracy. This is an interesting work.

Comment on revised version:

The authors have addressed all the raised concerns and made appropriate modifications to the manuscript. The changes have improved the clarity, methodology, and overall quality of the paper. Given these improvements, I believe the paper now meets the standards for publication in this journal.

<https://doi.org/10.7554/eLife.101069.2.sa2>

Reviewer #2 (Public review):

Summary:

The authors improve the work of Jallais et al. (2022) by including a novel module capable of automatically learning feature selection from different acquisition protocols inside a supervised learning framework. Combining the module above with an estimation framework for estimating the posterior distribution of model parameters, they obtain rich probabilistic information (uncertainty and degeneracy) on the parameters in a reasonable computation time.

The main contributions of the work are:

- (1) The whole framework allows the user to avoid manually defining summary statistics, which may be slow and tedious and affect the quality of the results.
- (2) The authors tested the proposal by tackling three different biophysical models for brain tissue and using data with characteristics commonly used by the diffusion-MR-microstructure research community.
- (3) The authors validated their method well with the state-of-the-art.
- (4) The methodology allows the quantification of the inherent model's degeneration and how it increases with strong noise.

The authors showed the utility of their proposal by computing complex parameter descriptors automatically in an achievable time for three different and relevant biophysical models.

Importantly, this proposal promotes tackling, analyzing, and considering the degenerated nature of the most used models in brain microstructure estimation.

<https://doi.org/10.7554/eLife.101069.2.sa1>

Author response:

The following is the authors' response to the original reviews.

Reviewer #1 (Public Review):

The authors proposed a framework to estimate the posterior distribution of parameters in biophysical models. The framework has two modules: the first MLP module is used to reduce data dimensionality and the second NPE module is used to approximate the

desired posterior distribution. The results show that the MLP module can capture additional information compared to manually defined summary statistics. By using the NPE module, the repetitive evaluation of the forward model is avoided, thus making the framework computationally efficient. The results show the framework has promise in identifying degeneracy. This is an interesting work.

We thank the reviewer for the positive comments made on our manuscript.

Reviewer #1 (Recommendations For The Authors):

I have some minor comments.

(1) The uGUIDE framework has two modules, MLP and NPE. Why are the two modules trained jointly? The MLP module is used to reduce data dimensionality. Given that the number of features for different models is all fixed to 6, why does one need different MLPs? This module should, in principle, be general-purpose and independent of the model used.

The MLP must be trained together with the NPE module to maximise inference performance in terms of accuracy and precision. Although the number of features predicted by the MLP was fixed to six, the characteristics of these six features can be very different, depending on the chosen forward model and the available data, as we showed in Appendix 1 Figure 1. Training the MLP independently of the NPE would result in suboptimal performance of μ GUIDE, with potentially higher bias and variance of the predicted posterior distributions. We have now added these considerations in the Methods section.

(2) The authors mentioned at L463 that all the 3 models use 6 features. From L445 to L447, it seems model 3 has 7 unknown parameters. How can one use 6 features to estimate 7 unknowns?

Thank you for pointing out the lack of clarity regarding the parameters to estimate in this section. Model 3 is a three-compartment model, whose parameters of interest are the signal fraction and diffusivity from water diffusing in the neurite space (f_n and D_n), the neurites orientation dispersion index (ODI), the signal fraction in cell bodies (f_s), a proxy to soma radius and diffusivity (C_s), and the signal fraction and diffusivity in the extracellular space (f_e and D_e). The signal fractions are constrained by the relationship $f_n + f_s + f_e = 1$, hence f_e is calculated from the estimated f_n and f_s . This leaves us with 6 parameters to estimate: f_n , D_n , ODI, f_s , C_s , D_e . We clarified it in the revised version of the paper.

(3) L471, Rician noise is not a proper term. Rician distribution is the distribution of pixel intensities observed in the presence of noise. And Rician distribution is the result of magnitude reconstruction. See "Noise in magnitude magnetic resonance images" published in 2008. I assume that real-valued Gaussian noise is added to simulated data.

We apologize for the confusion. We added Gaussian noise to the real and imaginary parts of the simulated signals and then used the magnitude of this noisy complex signal for our experiments. We rephrased the sentence for more clarity.

(4) L475, why thinning is not used in MCMC? In figure 3, the MCMC results are more biased than uGUIDE, is it related to no thinning in MCMC?

We followed the recommendations by Harms et al. (2018) for the MCMC experiments. They analysed the impact of thinning (among other parameters) on the estimated posterior distributions. Their findings indicate that thinning is unnecessary and inefficient, and they recommend using more samples instead. For further details, we refer the reviewer to their

publication, along with the theoretical works they cite. We have now added this note in the Methods section.

(5) Did the authors try model-fitting methods with different initializations to get a distribution of the parameters? Like the paper "Degeneracy in model parameter estimation for multi-compartmental diffusion in neuronal tissue". For the in vivo data, it is informative to see the model-fitting results.

No, we did not try model-fitting methods with different initializations because such methods provide only a partial description of the solution landscape, which can be interpreted as a partial posterior distribution. Although this approach can help to highlight the problem of degeneracy, it does not provide a complete description of all potential solutions. In contrast, MCMC estimates the full posterior distribution, offering a more accurate and precise characterization of degeneracies and uncertainties compared to model-fitting methods with varying initializations. Hence, we decided to use MCMC as benchmark. We have now added these considerations to the Discussion section.

Reviewer #2 (Public Review):

Summary:

The authors improve the work of Jallais et al. (2022) by including a novel module capable of automatically learning feature selection from different acquisition protocols inside a supervised learning framework. Combining the module above with an estimation framework for estimating the posterior distribution of model parameters, they obtain rich probabilistic information (uncertainty and degeneracy) on the parameters in a reasonable computation time.

The main contributions of the work are:

(1) The whole framework allows the user to avoid manually defining summary statistics, which may be slow and tedious and affect the quality of the results.

(2) The authors tested the proposal by tackling three different biophysical models for brain tissue and using data with characteristics commonly used by the diffusion-MRmicrostructure research community.

(3) The authors validated their method well with the state-of-the-art.

The main weakness is:

(1) The methodology was tested only on scenarios with a signal-to-noise ratio (SNR) equal to 50. It is interesting to show results with lower SNR and without noise that the method can detect the model's inherent degenerations and how the degeneration increases when strong noise is present. I suggest expanding the Figure in Appendix 1 to include this information.

The authors showed the utility of their proposal by computing complex parameter descriptors automatically in an achievable time for three different and relevant biophysical models.

Importantly, this proposal promotes tackling, analysing, and considering the degenerated nature of the most used models in brain microstructure estimation.

We thank the reviewer for these positive remarks.

Concerning the main weakness highlighted by the reviewer: In our submitted work, we presented results both without noise and with a signal-to-noise ratio (SNR) equal to 50 (similar to the SNR in the experimental data analysed). Figure 5 shows exemplar posterior distributions obtained in a noise-free scenario, and Table 1 reports the number of degeneracies for each model on 10000 noise-free simulations. These results highlight that the presence of degeneracies is inherent to the model definition. Figures 3, 6 and 7 present results considering an SNR of 50. We acknowledge that results with lower SNR have not been included in the initial submission. To address this, we added a figure in the appendix illustrating the impact of noise on the posterior distributions. Specifically, Figure 1A of Appendix 2 shows posterior distributions estimated from signals generated using an exemplar set of model parameters with varying noise levels

(no noise, SNR=50 and SNR=25). Figure 1B presents uncertainties values obtained on 1000 simulations for each noise level. We observe that, as the SNR reduces, uncertainty increases. Noise in the signal contributes to irreducible variance. The confidence in the estimates therefore reduces as the noise level increases.

Reviewer #2 (Recommendations For The Authors):

Some suggestions:

Panel A of Figure 2 may deserve a better explanation in the Figure's caption.

We agree that the description of panel A of figure 2 was succinct and added more explanation in the figure's caption.

The caption of Figure 3 should mention that the panel's titles are the parameters of the used biophysical models.

We added in the caption of figure 3 that the names of the model parameters are indicated in the titles of the panels. We apologise for the confusion it may have created.

In equation (3), the authors should indicate the summation index.

We apologise for not putting the summation index in equation 3. We added it in the revised version.

In line 474, the authors should discuss if the systematic use of the maximum likelihood estimator as an initializer for the sampling does not bias the computed results.

Concerning the MCMC estimations, we followed the recommendations from Harms et al. (2018). They investigated the use of starting from the maximum likelihood estimator (MLE). They concluded that starting from the MLE allows to start in the stationary distribution of the Markov chain, removing the need for some burn-in. Additionally, they showed that initializing the sampling from the MLE has the advantage of removing salt- and pepper-like noise from the resulting mean and standard deviation maps. We have now added this note in the Methods section.

<https://doi.org/10.7554/eLife.101069.2.sa0>