

Flow matching for atmospheric retrieval of exoplanets: Where reliability meets adaptive noise levels

Timothy D. Gebhard^{1,2,*}, Jonas Wildberger^{1,3}, Maximilian Dax¹, Annalena Kofler^{1,4},
Daniel Angerhausen², Sascha P. Quanz^{2,5}, and Bernhard Schölkopf^{1,6}

¹ Max Planck Institute for Intelligent Systems, Max-Planck-Ring 4, 72076 Tübingen, Germany

² ETH Zurich, Institute for Particle Physics & Astrophysics, Wolfgang-Pauli-Strasse 27, 8093 Zurich, Switzerland

³ ELLIS Institute Tübingen, Maria-von-Linden-Straße 2, 72076 Tübingen, Germany

⁴ Max Planck Institute for Gravitational Physics (Albert Einstein Institute), Am Mühlenberg 1, 14476 Potsdam, Germany

⁵ ETH Zurich, Department of Earth and Planetary Sciences, Sonneggstrasse 5, 8092 Zurich, Switzerland

⁶ ETH Zurich, Department of Computer Science, Universitätsstrasse 6, 8092 Zurich, Switzerland

Received 11 August 2024 / Accepted 16 October 2024

ABSTRACT

Context. Inferring atmospheric properties of exoplanets from observed spectra is key to understanding their formation, evolution, and habitability. Since traditional Bayesian approaches to atmospheric retrieval (e.g., nested sampling) are computationally expensive, a growing number of machine learning (ML) methods such as neural posterior estimation (NPE) have been proposed.

Aims. We seek to make ML-based atmospheric retrieval (1) more reliable and accurate with verified results, and (2) more flexible with respect to the underlying neural networks and the choice of the assumed noise models.

Methods. First, we adopted flow matching posterior estimation (FMPE) as a new ML approach to atmospheric retrieval. FMPE maintains many advantages of NPE, but provides greater architectural flexibility and scalability. Second, we used importance sampling (IS) to verify and correct ML results, and to compute an estimate of the Bayesian evidence. Third, we conditioned our ML models on the assumed noise level of a spectrum (i.e., error bars), and thus made them adaptable to different noise models.

Results. Both our noise-level-conditional FMPE and NPE models perform on a par with nested sampling across a range of noise levels when tested on simulated data. FMPE trains about three times faster than NPE and yields higher IS efficiencies. IS successfully corrects inaccurate ML results, identifies model failures via low efficiencies, and provides accurate estimates of the Bayesian evidence.

Conclusions. FMPE is a powerful alternative to NPE for fast, amortized, and parallelizable atmospheric retrieval. IS can verify results, helping to build confidence in ML-based approaches, while also facilitating model comparison via the evidence ratio. Noise level conditioning allows design studies for future instruments to be scaled up; for example, in terms of the range of signal-to-noise ratios.

Key words. methods: data analysis – methods: statistical – planets and satellites: atmospheres

1. Introduction

In exoplanet science, the term “atmospheric retrieval” refers to the inference of atmospheric parameters (e.g., abundances of chemical species, presence of clouds, etc.) from an observed planet spectrum (Madhusudhan 2018). These parameters offer essential insights into the formation and evolution of planets and are central to the study of planetary habitability and the search for life beyond the Solar System. Consequently, the atmospheric characterization of exoplanets is also a key driver of the development of future instruments and missions, such as the Habitable Worlds Observatory (HWO; NASA 2023) or the Large Interferometer for Exoplanets (LIFE; Quanz et al. 2022). From a data analysis perspective, the standard approach to atmospheric retrieval is to formulate the task as a Bayesian inference problem. In this framework, one starts with an observed (emission or transmission) planet spectrum $x \in \mathbb{R}^b$ in the form of a vector of flux values in b discrete wavelength bins. The atmospheric parameters of interest are denoted by $\theta \in \Theta$, where for simplicity we can assume $\Theta = \mathbb{R}^d$. Any a priori knowledge or belief about the value of θ , including known constraints, is encoded in the

prior $p(\theta)$. Moreover, we assume a likelihood function $p(x|\theta)$, which describes the probability of observing the spectrum x given some parameter values θ . In practice, the likelihood consists of two parts: (1) a forward simulator to transform a given θ into the corresponding spectrum, $\hat{x} = \text{simulator}(\theta)$, and (2) a way to quantify the probability that the true spectrum is $\hat{x}$ when we observe x . A common approach (see, e.g., Nasedkin et al. 2024a) is to assume that $x - \hat{x}$ follows a b -dimensional Gaussian distribution with mean zero and a covariance matrix $\Sigma \in \mathbb{R}^{b \times b}$. This assumption, which yields a simple expression for $p(x|\theta)$, is equivalent to assuming that the noise in each spectral bin comes from a Gaussian distribution. The likelihood and prior together define the Bayesian posterior distribution $p(\theta|x) = p(x|\theta) p(\theta) / p(x)$, which quantifies the belief about the value of θ after observing x . The normalizing constant $p(x)$ is called the evidence. The primary goal of Bayesian inference is now to summarize the posterior in terms of samples $\theta \sim p(\theta|x)$.

Traditionally, the posterior $p(\theta|x)$ is sampled with Markov chain Monte Carlo (MCMC; see, e.g., Hogg & Foreman-Mackey 2018, for a practical introduction) or nested sampling (Skilling 2006). These methods essentially run the forward simulator in a loop to explore the parameter space. The hardness of this problem scales exponentially with the number of atmospheric

* Corresponding author; tgebhard@tue.mpg.de

parameters, and in practice, millions of likelihood evaluations (i.e., simulator calls) may be required to estimate the posterior of a single exoplanet spectrum. Parallelizing this process is not trivial (see, e.g., the discussion in [Ashton et al. 2022](#)). As a result, atmospheric retrievals often imply a considerable computational cost even for the simplest forward models, with overall runtimes on the order of days to weeks. In addition, the simulations from an MCMC or nested sampling run typically cannot be reused to analyze another spectrum – or the same spectrum with different error bars – even when assuming the same prior distribution and simulator settings. Finally, the traditional setup is limited to tractable likelihoods where we can evaluate $p(x|\theta)$ explicitly. This can be restrictive if we want to go beyond the Gaussianity assumption and use realistic instrumental noise for which no analytical expression may be available.

Related work. Due to the computational limitations of traditional methods, alternative approaches have recently received increasing attention – especially those using machine learning (ML). The primary advantage of ML lies in its potential for amortization: while data generation and training may be expensive, the resulting ML model can be used for cheap inference for any number of retrievals without retraining. This amortizes the computational cost of training over inference runs. Approaches that fit into this category include deep belief networks ([Waldmann 2016](#)), generative adversarial networks ([Zingales & Waldmann 2018](#)), random forests ([Márquez-Neila et al. 2018](#); [Fisher et al. 2020](#); [Nixon & Madhusudhan 2020](#)), Monte Carlo dropout ([Soboczenski et al. 2018](#)), Bayesian neural networks ([Cobb et al. 2019](#)), and different deep learning architectures ([Yip et al. 2021](#); [Ardévol Martínez et al. 2022](#); [Giobergia et al. 2023](#); [Unlu et al. 2023](#); [Sweet 2024](#)). A number of these methods do not produce proper Bayesian posteriors, but rather samples from a “predictive distribution” ([Soboczenski et al. 2018](#)) whose exact statistical interpretation is not yet fully understood. This applies in particular to methods that are trained on noise-free data. Approaches that do yield a tractable posterior density commonly rely on simplifying assumptions regarding its form; for example, [Cobb et al. \(2019\)](#) and [Ardévol Martínez et al. \(2022\)](#) both assume the posterior to be a multidimensional Gaussian. A major advance in this regard is the work of [Vasist et al. \(2023\)](#), who bring discrete normalizing flows (DNFs) trained with neural posterior estimation (NPE) to the field of exoplanetary atmospheric retrieval. This approach is a simulation-based inference (SBI)¹ technique that foregoes explicit assumptions about the functional form of $p(\theta|x)$ and allows one to evaluate the posterior density at arbitrary values of θ instead of only providing sample access. DNFs trained with NPE were subsequently used by the winning entry to the 2023 edition of the ARIEL data challenge ([Aubin et al. 2023](#); see also [Changeat & Yip 2023](#)). More recently, the potential of ML has been recognized in non-amortized settings as well: [Yip et al. \(2024\)](#) combine DNFs with variational inference using a differentiable forward simulator, and [Ardévol Martínez et al. \(2024\)](#) train DNFs using sequential NPE (in particular, automatic posterior transformation as introduced by [Greenberg et al. 2019](#)). Neither approach produces a single model that can be reused for different input spectra; instead, they both speed up individual retrievals by reducing the number of likelihood evaluations required compared to traditional methods. Another class of non-amortized

retrieval methods are the approaches of [Himes et al. \(2022\)](#); [Hendrix et al. \(2023\)](#); [Tahseen et al. \(2024\)](#), and [Dahlbudding et al. \(2024\)](#), who do not predict a posterior directly, but instead aim to accelerate retrievals by replacing the computationally expensive simulator with a learned emulator, which can then be used with MCMC or nested sampling. In a similar vein, [Gebhard et al. \(2023a\)](#) propose to learn efficient parameterizations of pressure-temperature profiles to speed up retrievals by reducing the number of required parameters. Despite all these advances, significant limitations remain. Although faster than traditional methods, existing ML approaches lack the theoretical guarantees of stochastic samplers, which affects their reliability. In addition, existing ML methods are often trained on noise-free data, or assume a fixed noise level, which limits their practical applicability. These challenges motivate the improvements proposed in this work.

Contributions. We introduce three main ideas to the field of ML-based atmospheric retrieval of exoplanets²: (1) We adopt flow matching posterior estimation (FMPE; [Dax et al. 2023b](#)) with continuous normalizing flows (CNFs) as an alternative to DNFs trained with NPE. (2) We achieve several advantages compared to existing ML methods for atmospheric retrieval by combining FMPE and NPE with importance sampling (IS). In particular, IS (a) allows samples from an ML model to be reweighted toward the true posterior distribution, (b) provides the sampling efficiency as a quality metric for the ML-based proposal distribution, and (c) offers a way to estimate the Bayesian evidence. (3) We show that we can learn noise-level-conditional models that can handle spectra with different assumptions about the size of the error bars and return accurate posteriors without the need to retrain.

2. Method

Before presenting our methodical contributions, we briefly summarize NPE using DNFs, which we use as a baseline for CNFs trained with FMPE. We note that in Sections 2.1 and 2.2, we explicitly distinguish between the type of normalizing flow (DNF or CNF) and the training method (NPE or FMPE). In the rest of this paper, we simply write “NPE” to refer to “DNFs trained with NPE” and “FMPE” to refer to “CNFs trained with FMPE.” A schematic comparison of the two methods is found in Fig. 1.

2.1. Neural posterior estimation

NPE is an SBI technique introduced by [Papamakarios & Murray \(2016\)](#). The idea is to train a density estimator $q(\theta|x)$ to approximate the true posterior $p(\theta|x)$ by minimizing the loss function

$$\mathcal{L}_{\text{NPE}} = - \underbrace{\mathbb{E}_{\theta \sim p(\theta)}}_{\text{sample } \theta \text{ from prior}} \underbrace{\mathbb{E}_{x \sim p(x|\theta)}}_{\text{run the simulator}} \log q(\theta|x), \quad (1)$$

where $\mathbb{E}$ denotes the expectation value. Unlike in classic supervised learning, this loss function does not require access to the value of the true posterior $p(\theta|x)$. To approximate $p(\theta|x)$ well, we must choose a sufficiently flexible class of density estimators for $q(\theta|x)$. In practice, q is commonly implemented as a (conditional) DNF ([Tabak & Vanden-Eijnden 2010](#);

¹ Simulation-based inference (SBI) refers to Bayesian inference based on likelihood simulations, $x \sim p(x|\theta)$. This is in contrast to likelihood-based inference methods (e.g., MCMC) that rely on explicit evaluation of $p(x|\theta)$. For a review of SBI, see, e.g., [Cranmer et al. \(2020\)](#).

² Earlier versions of this work were presented at a AAAI 2024 workshop (see [Gebhard et al. 2023b](#)), AbSciCon 2024, and Exoplanets 5.

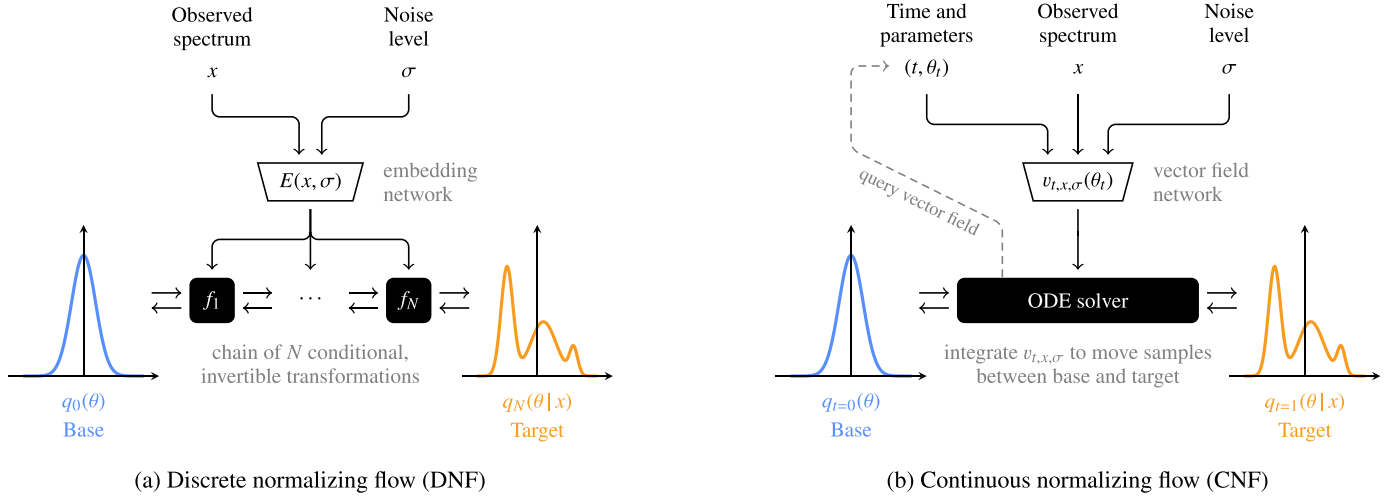


Fig. 1. Schematic comparison of DNFs and CNFs. DNFs (trained with NPE) were first applied to exoplanetary atmospheric retrieval by [Vasist et al. \(2023\)](#), while CNFs (trained with FMPE) and the idea of noise level conditioning are discussed in this work.

[Rezende & Mohamed 2015](#)). DNFs construct a bijection between a simple base distribution³ q_0 (e.g., a Gaussian) and a complex target distribution (e.g., a posterior) by applying a chain $\phi_x = f_{N,x} \circ \dots \circ f_{1,x}$ of invertible transformations $f_{i,x} : \mathbb{R}^b \times \Theta \rightarrow \Theta$. The density of this target distribution is

$$q(\theta|x) = q_0(\phi_x^{-1}(\theta)) \cdot \left| \det \frac{d\phi_x^{-1}(\theta)}{d\theta} \right|. \quad (2)$$

Each transformation $f_{i,x}$ in this chain is conditional on the observation (indicated by the subscript x) and has learnable parameters that are determined by minimizing Eq. (1) during training. The Jacobian determinant in Eq. (2) is introduced by the change of variables formula and accounts for the volume changes induced by the transformations $f_{i,x}$. It is required to ensure that the learned posterior distribution $q(\theta|x)$ integrates to 1. To keep the evaluation of $q(\theta|x)$ computationally feasible, we typically have to restrict our choice of $f_{i,x}$ to transformations that yield a triangular Jacobian matrix. This results in strong constraints on the architecture of the model, limiting its flexibility and expressiveness. For a more detailed introduction to DNFs, see, for example, the reviews by [Kobyzev et al. \(2021\)](#) or [Papamakarios et al. \(2021\)](#).

2.2. Flow matching posterior estimation

An alternative approach to density estimation is to parameterize the transformation from the base to the target distribution continuously using a time parameter $t \in [0, 1]$. This is the idea of a CNF ([Chen et al. 2018](#)). In the case in which the target distribution is a posterior, a CNF consists of a neural network that models a time-dependent vector field $v : \mathbb{R}^b \times [0, 1] \times \Theta \rightarrow \Theta$. This vector field describes the temporal derivative of the trajectory ϕ of a sample moving from the base to the target distribution,

$$\frac{d}{dt} \phi_{t,x}(\theta) = v_{t,x}(\phi_{t,x}(\theta)), \quad (3)$$

where the initial location of the sample under the base distribution is given by $\phi_{t=0,x}(\theta) = \theta_0$. Since x is fixed for a given retrieval, we also write $\theta_t \equiv \phi_{t,x}(\theta)$. Equation (3) is an

³ It is important to keep in mind that the base distribution q_0 of a normalizing flow is different from the prior distribution $p(\theta)$ of a retrieval.

ordinary differential equation (ODE) where the right-hand side is parameterized by a neural network. To obtain samples from the target distribution $q_{t=1,x}(\theta)$, we can draw samples from the base distribution $q_{t=0,x}(\theta) = q_0(\theta)$ and transform them by integrating the vector field over time t using a standard ODE solver. Intuitively, this process can be interpreted as moving the samples from the base distribution along the time-dependent vector field, which is illustrated in Fig. 2. The main advantage of this continuous approach to density estimation is that we can use an arbitrary neural network to parameterize the vector field v . Since we model only the derivative of v , there is no need to track a normalization constant. This results in greater architectural flexibility, improved scalability to more learnable parameters, and generally faster training ([Dax et al. 2023b](#)). On the downside, sampling $q(\theta|x)$ during inference requires many forward passes through the neural network that parameterizes the vector field $v_{x,t}$ to solve the ODE, instead of just a single forward pass in the case of DNFs.

To obtain an explicit expression for the density $q(\theta|x)$, we can solve the homogeneous continuity (or transport) equation,

$$\frac{\partial}{\partial t} q_{t,x}(\theta) + \text{div}(q_{t,x}(\theta) v_{t,x}(\theta)) = 0. \quad (4)$$

The continuity equation essentially states that probability is a conserved quantity: if $q_{t=0,x} = q_0$ integrates to 1, then so must any $q_{t,x}$, and in particular $q_{t=1,x} = q(\theta|x)$. Solving Eq. (4) yields

$$q(\theta|x) = q_1(\theta_1|x) = q_0(\theta_0) \cdot \exp \left\{ - \int_0^1 \text{div} v_{t,x}(\theta_t) dt \right\}, \quad (5)$$

which describes how the density of the target distribution can be evaluated in terms of the base distribution and the vector field. If we want to sample from the target distribution and evaluate the corresponding probability density, we can solve a joint ODE, which is more efficient than sampling and evaluating the density separately (see Appendix C of [Lipman et al. 2022](#) for details).

So far, we have discussed how to sample and evaluate $q(\theta|x)$, assuming that we already have a model for the vector field v , but not how we can actually learn one. In principle, Eq. (5) allows training via likelihood maximization (i.e., NPE), but evaluating the integral over $\text{div} v_{t,x}$ millions or billions of times during training is expensive. A much more efficient training scheme for CNFs has been proposed by [Lipman et al. \(2022\)](#), which

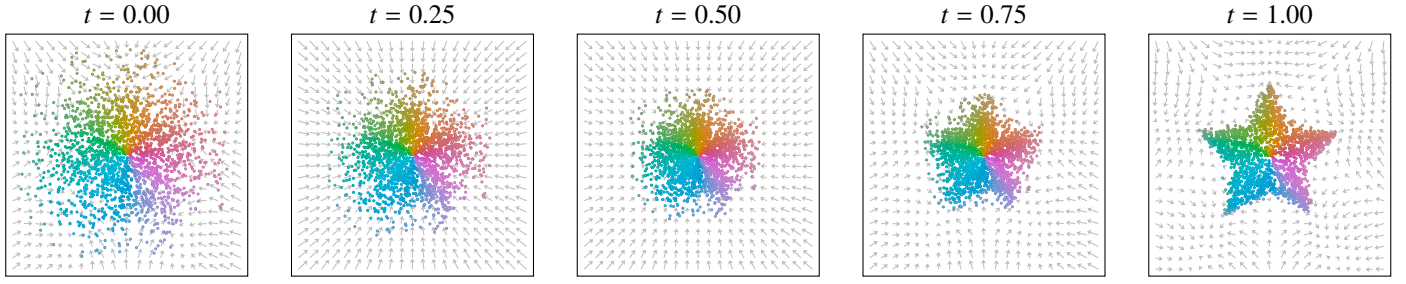


Fig. 2. Illustration of a time-dependent vector field (indicated by the gray arrows) continuously transforming samples from a standard 2D Gaussian at $t = 0$ into samples from a more complex, star-shaped distribution at $t = 1$. For simplicity, we show an unconditional example here; for an atmospheric retrieval, the vector field would depend not only on t but also on the observed spectrum x , and the assumed noise level σ . The size of the arrows has been rescaled for visual purposes.

was subsequently adapted for Bayesian inference by [Dax et al. \(2023b\)](#): flow matching. The idea is to train the model for v by regressing it directly onto a target vector field u . The trick is that the target vector field $u_t : [0, 1] \times \Theta \rightarrow \Theta$ can be chosen on a sample-conditional basis⁴. Sample-conditional means that the target u_t for a given training sample x depends on the corresponding true parameter value θ_1 . [Lipman et al. \(2022\)](#) show that in this case, there exist some simple choices for u that result in a training objective that is equivalent to likelihood maximization. For this work, we have used

$$u_t(\theta | \theta_1) = \frac{\theta_1 - (1 - \sigma_{\min})\theta}{1 - (1 - \sigma_{\min})t}, \quad (6)$$

which generates the Gaussian probability paths

$$p_t(\theta | \theta_1) = \mathcal{N}(\theta_1 t, (1 - (1 - \sigma_{\min})t)^2 \cdot I_d). \quad (7)$$

This probability path transforms a d -dimensional standard normal distribution at $t = 0$ into a Gaussian with a standard deviation $\sigma_{\min}$ centered on θ_1 at $t = 1$. When averaging over many $\theta_1 \sim p(\theta)$, regressing onto these target vector fields $u_t(\theta | \theta_1)$ results in a vector field $v_{t,x}(\theta)$ that transforms the base into the correct target distribution. Thus, the loss function we need to minimize is

$$\mathcal{L}_{\text{FMPE}} = \underbrace{\mathbb{E}_{t \sim p(t)}}_{\text{sample time } t} \underbrace{\mathbb{E}_{\theta_1 \sim p(\theta)}}_{\text{sample } \theta_1 \text{ from prior}} \underbrace{\mathbb{E}_{x \sim p(x | \theta_1)}}_{\text{run the simulator}} \underbrace{\mathbb{E}_{\theta_t \sim p_t(\theta_t | \theta_1)}}_{\text{sample } \theta_t \text{ from probability path}} \|r\|^2, \quad (8)$$

where r is shorthand for the residual vector field

$$r = r(t, \theta_1, x, \theta_t) \equiv v_{t,x}(\theta_t) - u_t(\theta_t | \theta_1).$$

[Dax et al. \(2023b\)](#) provide both theoretical and empirical arguments that this training objective results in posterior estimates that are mass-covering. This means that $q(\theta | x)$ always places some probability mass in regions where the true posterior $p(\theta | x)$ is nonzero, and thus yields conservative estimates, similar to NPE.

2.3. Extension to noise-level-conditional models

Previous ML approaches to atmospheric retrieval have typically trained their models using a fixed noise distribution. For example, [Vasist et al. \(2023\)](#) choose a spectral covariance matrix

⁴ We emphasize that “sample-conditional” (a term introduced in [Dax et al. 2023a](#)) is different from the usual conditioning on the observation x .

$\Sigma \in \mathbb{R}^{b \times b}$, with $\Sigma_{ij} = \delta_{ij}\sigma^2$ (i.e., the noise in different bins is uncorrelated), and a fixed value for σ . During training, they sample noise realizations, $n \sim \mathcal{N}(0, \Sigma)$, on the fly and add them to the simulated spectra in their training set. As a result, at inference time, all retrievals assume a likelihood that is determined by the Σ used during training. To perform an atmospheric retrieval with a different assumption for the noise distribution (e.g., during instrument design), the model must be retrained. We overcame this limitation by constructing models that are not only conditioned on the observed spectrum x , but that also receive as input a description of the assumed noise distribution that can be varied at inference time. One natural choice for such a description is the covariance matrix Σ . For our proof of concept, we limited ourselves to Gaussian noise with $\Sigma = \sigma^2 \cdot I_d$, with variable σ . In this case, the assumed noise distribution is fully described by a single number, the noise level σ , giving rise to the notion of noise-level-conditional models. However, extending our approach is straightforward. For example, future work may obtain an empirical estimate of the spectral covariance structure using a Gaussian process (see [Greco & Brandt 2016](#)) and condition on the kernel parameters. This would allow the assumed noise level to vary as a function of wavelength and include correlations between spectral bins.

To learn a noise-level-conditional model in practice, we modified the respective neural networks to take not only a spectrum x as input, but also an assumed noise level σ (cf. [Fig. 1](#)). The updated training procedure for each sample looks as follows:

1. Sample θ from the prior $p(\theta)$ and simulate the corresponding spectrum x , or draw (θ, x) from a pre-generated training set.
2. New: sample Σ from a hyper-prior, $\Sigma \sim p(\Sigma)$.
3. Sample $n \sim \mathcal{N}(0, \Sigma)$ and add it to the spectrum: $x' = x + n$.
4. Use (θ, x', Σ) to train the model.

This procedure yields a model that can perform multiple retrievals for different assumed error bars on the spectrum without retraining. The feasibility of such a conditional approach has previously been demonstrated, for example, by [Dax et al. \(2021\)](#), with further refinements by [Wildberger et al. \(2023\)](#). In their work, they apply an NPE-based approach to gravitational wave parameter inference and learn a model that can be conditioned on an estimate of the power spectral density, which describes the detector noise properties at the time of the event being analyzed.

2.4. Importance sampling

In practice, both NPE and FMPE can produce results that deviate from the true posterior. Using too small a network, too little training data, suboptimal hyperparameter choices, or not training to convergence are just some of the reasons why the predicted

posterior $q(\theta|x)$ may differ from the true posterior $p(\theta|x)$. One way to improve the quality of the predicted posterior post hoc is IS (see, e.g., Tokdar & Kass 2009, for a review). This is an established statistical procedure that is also used by nested sampling implementations such as MultiNest (Feroz et al. 2019) or nautilus (Lange 2023). To combine it with ML, the model must allow the density $q(\theta|x)$ to be evaluated for arbitrary values of θ . As has been discussed, this is the case for both FMPE and NPE. IS is not limited to amortized approaches, either; for example, Zhang et al. (2023) have recently described how the idea can be applied to sequential NPE (cf. Ardévol Martínez et al. 2024).

To improve FMPE or NPE results with IS, we proceeded as follows (see Dax et al. 2023a). Assume that we have drawn N samples $\theta_i \sim q(\theta|x)$ from the posterior estimate produced by an ML model. We can improve the quality of these samples (in the sense of moving their distribution closer to the true posterior) by reweighting them according to the ratio of the true posterior $p(\theta|x)$ and the “proposal” distribution $q(\theta|x)$. To this end, we assign each sample θ_i an importance weight w_i such that

$$w_i \propto p(\theta_i|x) / q(\theta_i|x).$$

Values θ_i that are more likely under the true posterior than under the proposal thus receive a high weight, while samples that are less likely under $p(\theta|x)$ than under $q(\theta|x)$ receive a low weight. Moreover, for a perfect proposal distribution, where $p(\theta|x) = q(\theta|x)$, all weights are 1. We cannot evaluate $p(\theta|x)$ directly, but by virtue of Bayes’ theorem, we know that $p(\theta|x) \propto p(x|\theta)p(\theta)$. Therefore, we can define the weights as

$$w_i = p(x|\theta_i)p(\theta_i) / q(\theta_i|x), \quad (9)$$

and normalize them such that $\sum_{i=1}^N w_i = N^5$. Equation (9) also reveals the downside of IS: computing the weights w_i requires additional likelihood evaluations. While these can be fully parallelized, $p(x|\theta)$ must still be tractable. At this point, we can also see why it is crucial that $q(\theta|x)$ must be mass-covering. Suppose the true posterior has a mode that is not covered by the proposal distribution; that is, the parameter space has a region where $q(\theta|x)$ is zero but $p(\theta|x)$ is not. In this case, reweighting will not result in samples that are closer to the true posterior distribution because values of θ that are never sampled cannot be up-weighted.

Using the weights, we define the effective sample size (ESS),

$$N_{\text{eff}} = \left[\sum_{i=1}^N w_i \right]^2 / \left[\sum_{i=1}^N w_i^2 \right] = \frac{N}{\text{var}(w_i) + 1}, \quad (10)$$

and from it the sampling efficiency ε ,

$$\varepsilon = N_{\text{eff}}/N \in [0, 1]. \quad (11)$$

The sampling efficiency ε is a useful diagnostic criterion to assess the quality of the proposal distribution. A high efficiency implies that $q(\theta|x)$ matches $p(\theta|x)$ well. The inverse is not necessarily true, however: ε is sensitive to even small deviations in just one dimension, meaning that a low efficiency does not automatically imply that $q(\theta|x)$ is a bad approximation of $p(\theta|x)$. In practice, $\varepsilon \gtrsim 1\%$ is often already considered a “good” value.

Finally, IS allows us to estimate the Bayesian evidence Z associated with a posterior. The evidence (also known as the

marginal likelihood) plays an important role when comparing models via the Bayes factor, and is defined as

$$Z = \int p(\theta) p(x|\theta) d\theta.$$

Using the (unnormalized) weights w_i , we can estimate Z as

$$Z = \int \frac{q(\theta|x)p(\theta)p(x|\theta)}{q(\theta|x)} d\theta \approx \frac{1}{N} \sum_{i=1}^N \underbrace{\frac{p(\theta_i)p(x|\theta_i)}{q(\theta_i|x)}}_{=w_i}. \quad (12)$$

The standard deviation of $\log Z$ is estimated as:

$$\sigma(\log Z) = \sqrt{(N - N_{\text{eff}}) / (N \times N_{\text{eff}})} = \sqrt{(1 - \varepsilon)/(N\varepsilon)}. \quad (13)$$

A derivation of Eq. (13) is found, for example, in the supplementary material of Dax et al. (2023a), who show the effectiveness of combining NPE with IS for another challenging astronomical parameter inference task; namely, gravitational wave analysis.

3. Retrieval setup and training dataset

We tested the applicability of our approach using simulated data. To maintain comparability with the existing literature, we used a setup very similar to that of Vasist et al. (2023), which in turn is based on a study of HR8799 e by Mollière et al. (2020). We refer to these papers for a more detailed discussion of the atmospheric model. For the simulator, we used version 2.6.7 of petitRADTRANS (Mollière et al. 2019) with the convenience wrapper from Vasist et al. (2023). The dimensionality of the atmospheric parameter space, over which we define a uniform prior, is $\dim(\Theta) = 16$ (see Appendix A for details). The simulator maps each $\theta \in \Theta$ to a simulated emission spectrum of a gas giant-type exoplanet. Each spectrum x has a spectral resolution of $R = \lambda/\Delta\lambda = 400$, covering a wavelength range of $0.95 \mu\text{m}$ to $2.45 \mu\text{m}$ in $\dim(x) = 379$ bins. We also experimented with a resolution of $R = 1000$, but found this to be computationally infeasible for the nested sampling baselines. A plot with an example spectrum is found in Fig. A.1.

We generated $2^{25} = 33\,554\,432$ spectra as a training dataset for our ML models (plus $2^{20} = 1\,048\,576$ spectra for validation) by randomly sampling θ from the prior and passing it to the simulator (see also Appendix E for an ablation study on the effect of the number of training samples). The total upfront cost for this is about 12 500 CPU hours, which can be fully parallelized. The noise was generated on the fly during training using the procedure from Section 2.3. In particular, we sampled noise realizations n as

$$n \sim \mathcal{N}(0, \sigma^2 I_d) \quad \text{with} \quad \sigma \sim \mathcal{U}(0.05, 0.50),$$

and passed both $x' = x + n$ (i.e., the noisy spectrum) and σ (i.e., the noise level) to the model. The unit of σ here and in the rest of the paper is $1 \times 10^{-16} \text{ W/m}^2/\mu\text{m}$. The range of the hyper-prior for σ was chosen somewhat arbitrarily to cover the (fixed) value of $\sigma = 0.125\,754 \times 10^{-16} \text{ W/m}^2 \mu\text{m}^{-1}$ used by Vasist et al. (2023)⁶.

⁵ For the practical implementation, it is numerically much more stable to perform these calculations in log-space; see our code for details.

⁶ We have confirmed with the authors of Vasist et al. (2023) that the value of $\sigma = 0.125\,754 \times 10^{-17} \text{ W/m}^2 \mu\text{m}^{-1}$ in the journal version of their paper was a typographical error.

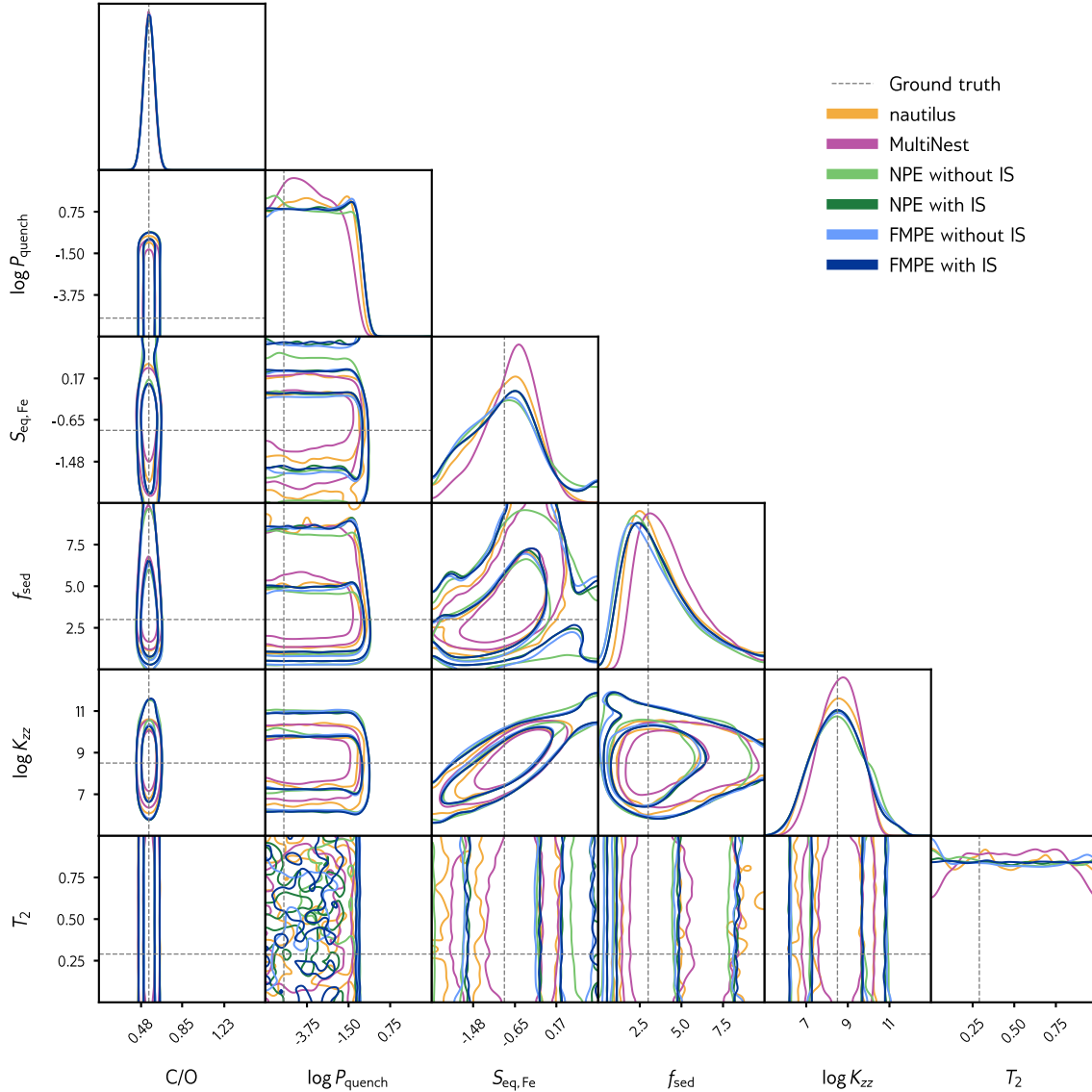


Fig. 3. Comparison of the 1D and 2D marginal posteriors for the noise-free benchmark spectrum from nested sampling (implemented by *nautilus* and *MultiNest*), FMPE, and NPE. For the latter two, we include the results with and without importance sampling. For visual purposes, we applied some light Gaussian smoothing to the histograms. Furthermore, we only show 6 selected parameters here; the full version featuring all 16 parameters is found in Fig. D.1.

4. Experiments and results

We trained an FMPE and an NPE model with approximately 310 million learnable parameters using stochastic gradient descent. Details about the training procedure and the network architectures are found in Appendix B. Training to convergence on an NVIDIA H100 GPU takes about 55 hours for the FMPE model, and about 163 hours for the NPE model. As a baseline, we also ran a number of traditional atmospheric retrievals using two different nested sampling implementations; namely, *nautilus* and (Py)MultiNest (see Appendix C for details).

4.1. Experiments on the benchmark spectrum

Comparison with nested sampling. As a first test case, we used all methods to perform a retrieval of the benchmark spectrum from Vasist et al. (2023) (see Appendix A for details). For the noise level, we assumed $\sigma = 0.125754$. We ran the retrieval twice with each method, once without adding any noise to the

simulated spectrum, and once with noise randomly drawn from $\mathcal{N}(0, \sigma^2)$ ⁷. For FMPE and NPE, we used IS with $2^{20} = 1\,048\,576$ proposal samples each. The main results of these retrievals are shown as a corner plot in Fig. 3 (subset of 6 parameters) and Fig. D.1 (all 16 parameters). We only include the plots for the noise-free retrievals, but we report that the results with noise look very similar and share all the relevant characteristics. The corresponding evidence, ESS, and sampling efficiencies are summarized in Table 1.

First, we observe that for well-constrained parameters (such as the C/O ratio or the metallicity), all methods are in excellent agreement. Second, without IS, the posteriors from FMPE and NPE show slight deviations for some parameters, while with IS, the results are virtually indistinguishable. Given the high sampling efficiencies and large ESSs (see Table 1), we are confident

⁷ We remind the reader that there is a difference between the actual noise in the observation x and the noise assumption in the likelihood $p(x|\theta)$.

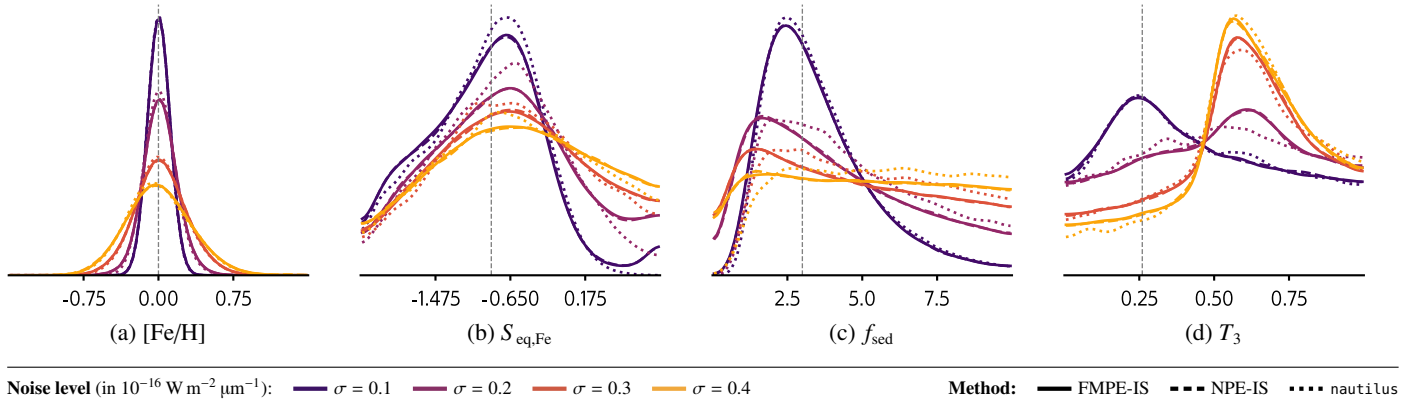


Fig. 4. Marginal posterior distributions for the noise-free benchmark spectrum for four atmospheric parameters at different assumed noise levels, σ (encoded by color), and for three different inference methods (encoded by line style). In each plot, the x axis spans the prior range, the dashed gray line ---- marks the respective ground truth value, and the y axis shows the density at an arbitrary scale.

Table 1. Estimate of the Bayesian log-evidence ($\log Z$), the ESS, and the sampling efficiency (ϵ) on the benchmark spectrum (with and without noise) for the different methods.

(a) Without noise			
Method	$\log Z$	ESS	ϵ
FMPE-IS	401.843 ± 0.002	199 761.8	19.1%
NPE-IS	401.832 ± 0.003	121 856.5	11.6%
MultiNest	407.182 ± 0.080	30 540	–
nautilus	401.844 ± 0.010	10 004.4	3.1%
(b) With noise			
Method	$\log Z$	ESS	ϵ
FMPE-IS	198.798 ± 0.003	111 010.8	10.6%
NPE-IS	198.792 ± 0.004	52 724.3	5.0%
MultiNest	203.235 ± 0.083	32 311	–
nautilus	198.737 ± 0.010	10 002.9	3.2%

Notes. MultiNest does not return a value for ϵ . Therefore, we simply report the default number of equal-weighted posterior samples as its ESS. The evidence here is the one without IS; with IS, the $\log Z$ estimates are slightly closer to the other methods.

that the IS-based results are a very good approximation of the true posterior. This is supported by the fact that the posteriors for FMPE and NPE also agree well with the nested sampling results from nautilus. For MultiNest, we find clear deviations from all other methods. These differences are particularly prominent for the quench pressure and the cloud scaling parameters. If we visually compare these results with Figure 2 from Vasist et al. (2023), we find that the deviations between MultiNest and the ML-based methods follow the same patterns. This is in line with previous findings that MultiNest seems to have a tendency to produce overconfident posteriors (cf., e.g., Ardévol Martínez et al. 2022; Dittmann 2024). Fourth, we observe that there are a few instances in which nautilus and MultiNest place no probability mass on a region of the parameter space where both ML methods agree that the posterior should be nonzero. One example of this effect are values of $\log K_{zz} > 11$. If the true posterior were indeed zero, as is suggested by the nested sampling baselines, we would expect that all ML-based proposal samples in this region would be strongly down-weighted

when applying IS. The fact that this is not the case leads us to believe that nautilus and MultiNest are slightly overconfident. This also explains why the peak around the true value is higher for these methods, since the probability mass that should have been placed at $\log K_{zz} > 11$ must be placed elsewhere. Finally, if we go beyond the corner plot and compare the log-evidence (see Table 1), we find good agreement between FMPE-IS, NPE-IS, and nautilus, while MultiNest again stands out and deviates more significantly.

Noise-conditional results. To demonstrate the effect of the assumed error bars on the posterior, we repeated the experiment from the last section using four different noise level assumptions: $\sigma = \{0.1, 0.2, 0.3, 0.4\}$. For simplicity, we limited ourselves to noise-free retrievals. Additionally, we used nautilus to create a nested sampling baseline for each case. We also ran MultiNest, which confirmed all the trends reported below, but as before deviated clearly from all the other methods. Therefore, we do not include it in the selection of the 1D marginal posteriors shown in Fig. 4: first, the results from FMPE and NPE with IS are again virtually indistinguishable, for all parameters and all assumed noise levels. Second, the agreement with the results from nautilus is generally good, with the same caveats as before: nautilus seems slightly overconfident in a few cases. Further analysis reveals evidence that this might be due to nautilus missing one or multiple posterior modes. Finally, we can see how the assumed noise level σ affects the resulting posterior: For some parameters, such as the metallicity, the 1D marginal posterior simply becomes broader as the size of the error bars is increased. This makes intuitive sense, as wider error bars mean that we should accept a broader range of potential explanations (i.e., parameter value combinations) for our observation. However, some parameters also show nontrivial changes in the posterior. For example, in Fig. 4d, the posterior does not become broader as we increase σ , but rather the mode of the distribution shifts away from the ground truth value. This effect is observed for all three methods, which suggests that this is the true behavior of the posterior and not a figment of our noise-conditional models. Our experiment shows how sensitive the result of an atmospheric retrieval can be to the noise level that one assumes in the likelihood function. This illustrates the value of learning an ML model that amortizes over different noise levels, as it allows retrievals to be repeated quickly and cheaply for different assumptions about the error bars to test the robustness of the values (or bounds) of the inferred atmospheric parameters.

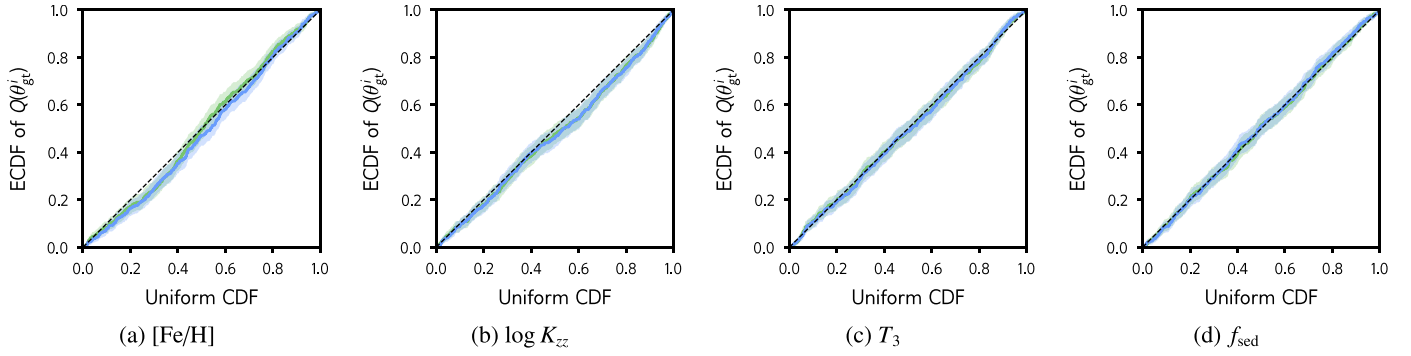


Fig. 5. P–P plot showing the ECDFs of $Q(\theta_i^{\text{gt}})$ against a uniform CDF for different atmospheric parameters and for both FMPE and NPE. The shaded regions indicate the respective 95% confidence interval computed by `scipy.stats.ecdf()`. Panels a and b show the results for the atmospheric parameters with the largest visual deviation from the diagonal (“worst case”), while panels c and d show the “best case.” All results are from the default test set.

Sampling efficiency as a diagnostic criterion. In Section 2.4, we argued that low values of the sampling efficiency ε can be used as a diagnostic criterion to identify cases in which we might not want to trust the results of our ML model without additional analysis. To illustrate this, we ran the following simple experiment: First, we performed an atmospheric retrieval on a version of the benchmark spectrum to which we added random Gaussian noise with $\sigma = 0.1$. However, to the ML model, we passed a noise level of $\sigma = 0.5$. This corresponds to a conservative retrieval in which the error bars are overestimated. Then, we repeated the experiment, but swapped the values of σ : we now added noise with $\sigma = 0.5$ to the spectrum but ran the retrieval assuming $\sigma = 0.1$. This is a misspecification that could yield unreliable results. Comparing the results from the two retrievals using our FMPE model with IS (100 000 proposal samples), we find that for the conservative retrieval, the sampling efficiency is $\varepsilon = 22.58\%$, while for the underestimated noise level, we get $\varepsilon = 1.53 \times 10^{-5}$. Thus, in this case, the efficiency ε has indeed flagged the problematic retrieval.

4.2. Large-scale analysis

To confirm that the promising results from Section 4.1 do not only hold for the benchmark spectrum, we performed a large number of retrievals to study the properties of our method for 500 different values of θ . To this end, we created two test sets: a “default” one, where we sampled θ directly from the Bayesian prior, and a “Gaussian” one, consisting of spectra from a neighborhood around the parameters θ_0 of the benchmark spectrum. The exact details and motivation for using two test sets are described in Appendix B.5.

Faithfulness of posteriors. It is important for the interpretability of an ML-based posterior that the estimated distribution is well calibrated; that is, neither broader (i.e., underconfident) nor tighter (i.e., overconfident) than the true posterior. To probe this for our models, we performed a retrieval for each spectrum in the default test set, drawing 100 000 posterior samples each. In the following, we considered every atmospheric parameter θ_i separately (where θ_i is the i -th entry of the parameter vector θ). For each retrieval, we determined the quantile $Q(\theta_i^{\text{gt}})$ of the ground-truth value θ_i^{gt} . In other words, for a given retrieval, $Q(\theta_i^{\text{gt}})$ is the fraction of posterior samples where the value of θ_i is less than the ground truth. If the posterior estimates are well

calibrated, we expect that these quantiles $Q(\theta_i^{\text{gt}})$ should follow a uniform distribution. For example, the ground-truth value should fall into the lower half of the posterior exactly 50% of the time. To test this, we created a P–P plot, where we plot the empirical cumulative density function (ECDF) of the $Q(\theta_i^{\text{gt}})$ values against the CDF of a uniform distribution. This should yield a diagonal line (up to statistical deviations). If the posteriors are underconfident, values of $Q(\theta_i^{\text{gt}})$ close to 0.5 are overrepresented compared to values near 0 or 1, and the resulting shape of the P–P plot looks like a sigmoid function: $\smile$. Conversely, for overconfident posteriors, we expect a shape that resembles a sinh-function: $\swarrow$. A representative selection of our results is shown in Fig. 5. The P–P plots look well calibrated for all atmospheric parameters, with no clear signs of under- or overconfidence. The differences between FMPE and NPE are negligible. We note that a similar experiment with nested sampling is computationally infeasible. From the runtimes listed in Table C.1, we can estimate that running 500 traditional retrievals would easily require a million CPU hours or more.

Sampling efficiency and proposal quality. We repeated the experiment from the last section using the Gaussian test set; that is, we performed 500 retrievals for both the FMPE and the NPE model, drawing 100 000 proposal samples each and applying IS⁸. In Fig. 6, we compare the corresponding sampling efficiencies for FMPE and NPE. First, there is a clear correlation ($\rho = 0.696$) between the two methods. This could be explained by the fact that a key challenge for both FMPE and NPE is the extraction of relevant information from the given spectra, which is independent of the actual density estimation task and equally difficult for both methods. Second, FMPE-IS yields on average a higher sampling efficiency than NPE-IS. This indicates that the posterior estimates from FMPE are closer to the true posterior than those from NPE. To confirm this, we looked at the Jensen–Shannon divergence (JSD) between the marginal posteriors for each atmospheric parameter with and without IS. The JSD is a measure of the similarity of distributions: the lower the JSD, the more similar the distributions are. Looking at the distribution of the JSD across the test set (Fig. 7), the average difference between the FMPE and FMPE-IS posteriors is smaller than the difference between the NPE and NPE-IS posteriors for all but one parameter (T_1), which supports the above hypothesis. The fact that T_1 yields the lowest JSD value is likely not a coincidence.

⁸ All results in this section hold qualitatively also for the default test set.

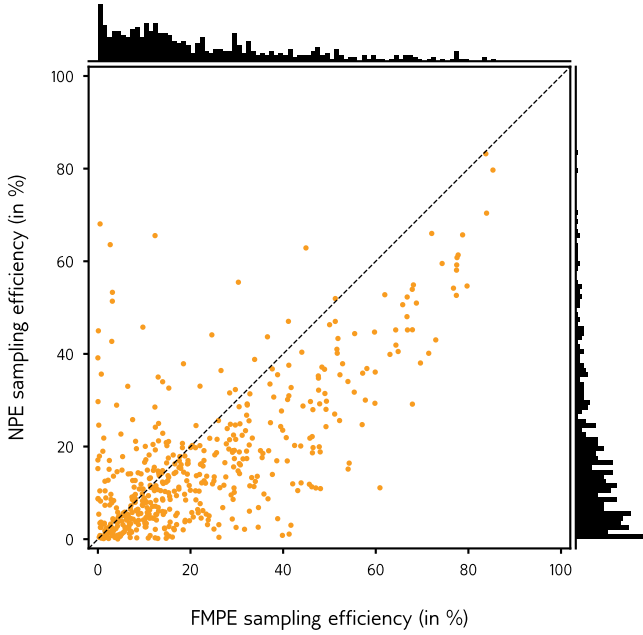


Fig. 6. Sampling efficiency of FMPE vs. NPE (Gaussian test set).

T_1 describes the temperature at the very upper end of the atmosphere. This property is notoriously difficult to constrain from emission spectra (cf., e.g., Sect. 4.4 of Gebhard et al. 2023a), and consequently the resulting posterior for T_1 is, in most cases, equal to its prior.

Bayesian evidence. We also looked at the estimates for the Bayesian (log)-evidence $\log Z$ produced by FMPE-IS and NPE-IS. Due to the prohibitive computational cost of running 500 retrievals with nested sampling, we cannot compare these values to those from traditional methods. However, we can report that in more than half of the cases, the $\log Z$ estimates from FMPE-IS and NPE-IS agree within one standard deviation, and in 98% of cases, the agreement is within three standard deviations.

4.3. Runtime comparison at inference time

Finally, we compared the runtimes of FMPE and NPE at inference time. For NPE, drawing $2^{16} = 65\,536$ posterior samples takes 1.161 ± 0.004 s, while sampling and evaluating log-probability densities takes 1.163 ± 0.002 s. In contrast, the FMPE model requires 16.262 ± 1.433 s for sampling and 258.08 ± 4.93 s for evaluating log-probabilities when using the settings from Appendix B.4. This significant difference is because NPE only needs one forward pass through the neural network, whereas FMPE requires multiple model evaluations by the ODE solver. All runtimes were measured on an NVIDIA H100 GPU and averaged over 10 runs.

5. Discussion

Amortization. Previous work has mainly considered the potential of ML-based methods for amortization in terms of the ability to efficiently process large numbers of different spectra that share similar instrument characteristics, as we can expect from future missions such as ARIEL (Tinetti et al. 2021). In this work, we have opened another perspective: the potential of using amortized models to perform atmospheric retrievals under different assumptions about the noise in the observed spectrum. This is useful in several ways. First, even for the same instrument, the

expected noise level in a spectrum is not always the same, but may depend, for example, on the host star of the target planet. Being able to handle these different scenarios without retraining makes a model significantly more useful in practice and increases the number of cases over which the training costs can be amortized. Second, being able to efficiently run retrievals for different noise levels can help us to understand how sensitive the inferred parameters and constraints are to these assumptions, and thus make retrievals more robust. Finally, amortization across different noise assumptions may become particularly valuable during the design phase of new instruments and missions for exoplanet science, such as HWO or LIFE. To understand the effect of specific design choices on the expected yield or the ability to constrain certain exoplanet properties, it is necessary to simulate surveys with different instrument configurations, which is expensive with traditional methods. For example, Konrad et al. (2022) explore the minimum requirements for LIFE and quantify how well various atmospheric parameters can be constrained as a function of the signal-to-noise ratio (S/N). Our method would allow this type of analysis to be performed at a much higher resolution (e.g., 100 S/N values instead of 4) with minimal additional computational cost.

Importance sampling. For a mass-covering posterior estimate $q(\theta|x)$, which can generally be assumed for FMPE and NPE (see Dax et al. 2023a,b), IS reweights samples from $q(\theta|x)$ to follow the true posterior $p(\theta|x)$ exactly in the asymptotic case. Of course, IS is not a magic bullet: If q is a poor approximation of p , the variance of the importance weights is high and the resulting ESS small. In this case, the IS posterior is not reliable. Conversely, a high IS efficiency ε is a strong indicator that $q(\theta|x)$ is close to the true posterior. Using ε as a quality criterion thus allows us to go beyond statistical arguments about the calibration of our posteriors in the average case and reason quantitatively about the reliability of our model outputs for individual retrievals. Furthermore, IS allows us to estimate the Bayesian evidence associated with a retrieval result, something that previous amortized ML approaches have not addressed. The evidence plays an important role in model comparison via the Bayes factor, which is given by the evidence ratio⁹. Access to the evidence therefore facilitates further studies; for example, to test if a planet spectrum is better explained by a model with or without clouds, or with or without a particular chemical species.

Runtime considerations. A key difference between FMPE and NPE is the trade-off between the time required to train a model and the time to do inference. Our FMPE model is substantially faster to train than the NPE model (by a factor of ~ 3), but also considerably slower at inference time; namely, by a factor of ~ 10 for sampling, and a factor of ~ 200 when evaluating log-probability densities. Of course, we must not forget that the training time is usually on the order of days, while the inference time is measured in seconds. Furthermore, evaluating log-probabilities is mainly needed for IS, in which case we have to consider two more things. First, FMPE yields on average higher sampling efficiencies than NPE, meaning that fewer proposal samples are required to achieve the same ESS. Second, when running IS, the computational cost of the inference procedure is dominated by the required additional likelihood evaluations, which usually significantly outweigh the cost of drawing proposal samples. It is also important to note that both sampling from the proposal and IS can be arbitrarily parallelized, which is a key advantage over nested sampling or MCMC.

⁹ For details and a note of caution, see, e.g., Jenkins & Peacock (2011).

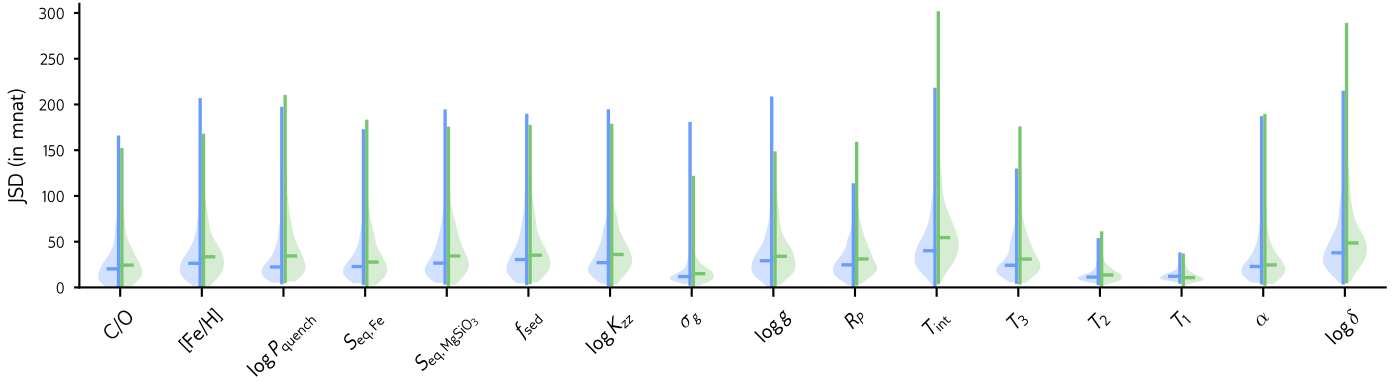


Fig. 7. Distribution of the JSD between the posterior marginals with and without IS on the Gaussian test set for FMPE and NPE. Lower is better. The vertical line indicates the support of the distribution, while the horizontal bar marks the median. We have only included retrievals with a sampling efficiency of $\varepsilon \geq 1\%$ to ensure the IS-based posterior is reliable.

While the true strength of amortized approaches is, of course, in repeated inference, the ability to easily run computations in parallel on hundreds or thousands of CPUs can lead to scenarios where generating training data, training a model, and running inference for only a single retrieval can still be preferable (in terms of wall time) to traditional methods.

Future directions. Going forward, we see multiple directions in which the ideas we have presented in this work can be built on and extended. First, while we have shown that models can be trained to amortize over different noise levels, we have not yet considered noise levels that vary as a function of wavelength or, more importantly, correlations of noise across different bins. The dangers of ignoring spectral covariance in atmospheric retrievals have been discussed, for example, in Greco & Brandt (2016) and Nasedkin et al. (2023). Extending our approach in this direction is straightforward, but will require close collaboration with the instrumentation community to ensure that the resulting noise models are rooted in reality. Second, our ideas can be transferred to similar problems in adjacent fields, such as inferring parameters of exoplanet interiors, or analyzing metal abundances in polluted white dwarfs, where ML has recently been applied (see, e.g., Haldemann et al. 2023; Baumeister & Tosi 2023; Badenas-Agusti et al. 2024). Third, our framework could be extended to adaptive priors: Some atmospheric parameters, such as the planet radius and surface gravity, can be constrained before a retrieval using coarse methods, or may be known from other sources (e.g., radial velocity measurements). In this case, we might want to use an “informed prior” (cf. Hayes et al. 2020). Prior conditioning (Dax et al. 2024) can integrate such information to simplify the inference task, while retaining the amortized properties of our framework. Finally, in this work, we have only considered setups in which the target spectrum of a retrieval was generated using the same atmospheric model that produced the training data. Future work should investigate in more detail how our approach performs on out-of-distribution data, or in the case of model mismatch, that is, when the assumed forward model is not a good fit for the actual atmosphere (e.g., the simulator assumes no clouds when the planet is cloudy).

6. Summary and conclusions

In this work, we have established FMPE (Dax et al. 2023b) as a new amortized ML approach to atmospheric retrieval. FMPE retains the desirable properties of NPE, such as distributional

expressiveness (i.e., no simplifying assumptions about the form of the posterior), tractable posterior density, and simulation-based training, while offering a different trade-off between training and inference time. Due to the lack of architectural constraints on the neural networks, FMPE is also more scalable than NPE. In the amortized scenario (i.e., ignoring training costs), both FMPE and NPE are orders of magnitude faster than traditional samplers; however, their high degree of parallelizability can make them a powerful alternative even for a single retrieval. In terms of posterior accuracy, FMPE and NPE are both on a par with traditional nested sampling, especially when combining them with IS. We demonstrated the value of IS for atmospheric retrieval in several ways. IS not only reweights the proposal samples from an ML model to more closely match the true posterior; it also provides a metric (the sampling efficiency) for when the model outputs can be trusted. Moreover, IS produces an accurate estimate of the Bayesian evidence, which is important for model comparison, and which previous amortized ML approaches did not provide. Although this comes at an increased computational cost (for additional likelihood evaluations), we believe that these clear benefits justify IS becoming a standard tool for ML-based atmospheric retrieval. Finally, we have shown how to train ML models that are noise-level-conditional and enable amortized inference without retraining not only for different spectra, but also for the same spectrum with different error bars. This allows one to test the robustness of inferred parameter values with respect to the assumed noise level, and is also highly relevant for design studies of future instruments (e.g., HWO or LIFE). Such studies need to analyze the impact of different configurations on the ability to constrain atmospheric parameters, and will clearly benefit from being able to explore, for example, a wide range of S/N values at almost no additional computational cost.

Data availability

Our entire code is available online at github.com/timothygebhard/fm4ar, and our data and results can be downloaded from <https://doi.org/10.17617/3.LYSSVN>.

Acknowledgements. The authors thank Kai Hou Yip for the fast and constructive review. We also thank Malavika Vasist and Evert Nasedkin for their help with the data generation code, and Jean Hayoz and Emily Garvin for valuable comments on the manuscript. Parts of this work have been carried out within the framework of the National Centre of Competence in Research PlanetS supported by the Swiss National Science Foundation under grant 51NF40_20560.

The computational work for this manuscript was performed on the Atlas cluster at the Max Planck Institute for Intelligent Systems. TDG acknowledges funding through the Max Planck ETH Center for Learning Systems. MD thanks the Hector Fellow Academy for support. The authors thank the International Max Planck Research School for Intelligent Systems (IMPRS-IS) for supporting AK. SPQ acknowledges the financial support of the SNSF. This work has made use of many open-source Python packages, including corner (Foreman-Mackey 2016), dynesty (Speagle 2020), glasflow (Williams et al. 2024), matplotlib (Hunter 2007), multiprocessing (McKerns et al. 2012), nautilus-sampler (Lange 2023), normflows (Stimper et al. 2023), numpy (Harris et al. 2020), scipy (Virtanen et al. 2020), scikit-learn (Pedregosa et al. 2011), pydantic (Colvin et al. 2024), PyMultiNest (Buchner et al. 2014), pandas (McKinney 2010), petitRADTRANS (Mollière et al. 2019), torch (Paszke et al. 2019), and ultranest (Buchner 2021). The colorblind-friendly color schemes in our figures are based on Petroff (2021).

References

- Alei, E., Quanz, S. P., Konrad, B. S., et al. 2024, *A&A*, **689**, A245
- Ardévol Martínez, F., Min, M., Kamp, I., & Palmer, P. I. 2022, *A&A*, **662**, A108
- Ardévol Martínez, F., Min, M., Huppenkothen, D., Kamp, I., & Palmer, P. I. 2024, *A&A*, **681**, L14
- Ashton, G., Bernstein, N., Buchner, J., et al. 2022, *Nat. Rev. Methods Primers*, **2**, 44
- Aubin, M., Cuesta-Lazaro, C., Tregidga, E., et al. 2023, arXiv e-prints [arXiv:2309.09337]
- Ba, J. L., Kiros, J. R., & Hinton, G. E. 2016, arXiv e-prints [arXiv:1607.06450]
- Badenas-Agusti, M., Viaña, J., Vanderburg, A., et al. 2024, *MNRAS*, **529**, 1688
- Baumeister, P., & Tosi, N. 2023, *A&A*, **676**, A106
- Buchner, J. 2021, *JOSS*, **6**, 3001
- Buchner, J. 2022, *Phys. Sci. Forum*, **5**, 41
- Buchner, J., Georgakakis, A., Nandra, K., et al. 2014, *A&A*, **564**, A125
- Changeat, Q., & Yip, K. H. 2023, *RAS Tech. Instrum.*, **2**, 45
- Chen, R. T. Q. 2018, <https://github.com/rtqichen/torchdiffeq>
- Chen, R. T. Q., Rubanova, Y., Bettencourt, J., et al. 2018, NeurIPS 2018, [arXiv:1806.07366]
- Cobb, A. D., Himes, M. D., Soboczenski, F., et al. 2019, *AJ*, **158**, 33
- Colvin, S., Jolibois, E., Ramezani, H., et al. 2024, <https://github.com/pydantic/pydantic>
- Cranmer, K., Brehmer, J., & Louppe, G. 2020, *PNAS*, **117**, 30055
- Dahlbidding, D., Molaverdikhani, K., Ercolano, B., & Grassi, T. 2024, *MNRAS*, **533**, 3475
- Dauphin, Y. N., Fan, A., Auli, M., & Grangier, D. 2016, ICML 2017, [arXiv:1612.07803]
- Dax, M., Green, S. R., Gair, J., et al. 2021, *Phys. Rev. Lett.*, **127**, 241103
- Dax, M., Green, S. R., Gair, J., et al. 2023a, *Phys. Rev. Lett.*, **130**, 171403
- Dax, M., Wildberger, J., Buchholz, S., et al. 2023b, NeurIPS 2023, [arXiv:2305.17161]
- Dax, M., Green, S. R., Gair, J., et al. 2024, arXiv e-prints [arXiv:2407.09602]
- Dittmann, A. 2024, *Open J. Astrophys.*, **7**, 123872
- Durkan, C., Bekasov, A., Murray, I., & Papamakarios, G. 2019, NeurIPS 2019, [arXiv:1906.04032]
- Durkan, C., Bekasov, A., Murray, I., & Papamakarios, G. 2020, arXiv e-prints [arXiv:1906.04032]
- Feroz, F., Hobson, M. P., & Bridges, M. 2009, *MNRAS*, **398**, 1601
- Feroz, F., Hobson, M. P., Cameron, E., & Pettitt, A. N. 2019, *Open J. Astrophys.*, **2**, 10
- Fisher, C., Hoeijmakers, H. J., Kitzmann, D., et al. 2020, *AJ*, **159**, 192
- Foreman-Mackey, D. 2016, *JOSS*, **1**, 24
- Gebhard, T. D., Angerhausen, D., Konrad, B. S., et al. 2023a, *A&A*, **681**, A3
- Gebhard, T. D., Wildberger, J., Dax, M., et al. 2023b, arXiv e-prints [arXiv:2312.08295]
- Giobergia, F., Koudounas, A., & Baralis, E. 2023, arXiv e-prints [arXiv:2310.01227]
- Greco, J. P., & Brandt, T. D. 2016, *ApJ*, **833**, 134
- Greenberg, D. S., Nonnenmacher, M., & Macke, J. H. 2019, arXiv e-prints [arXiv:1905.07488], ICML 2019
- Haldemann, J., Ksoll, V., Walter, D., et al. 2023, *A&A*, **672**, A180
- Harris, C. R., Millman, K. J., van der Walt, S. J., et al. 2020, *Nature*, **585**, 357
- Hayes, J. J. C., Kerins, E., Awiphan, S., et al. 2020, *MNRAS*, **494**, 4492
- Hendrix, J. L. A. M., Louca, A. J., & Miguel, Y. 2023, *MNRAS*, **524**, 643
- Hendrycks, D., & Gimpel, K. 2016, arXiv e-prints [arXiv:1606.08415]
- Higson, E., Handley, W., Hobson, M., & Lasenby, A. 2018, *Stat. Comput.*, **29**, 891
- Himes, M. D., Harrington, J., Cobb, A. D., et al. 2022, *PSJ*, **3**, 91
- Hinton, G. E., Srivastava, N., Krizhevsky, A., Sutskever, I., & Salakhutdinov, R. R. 2012, arXiv e-prints [arXiv:1207.0580]
- Hogg, D. W., & Foreman-Mackey, D. 2018, *ApJS*, **236**, 11
- Hunter, J. D. 2007, *Comput. Sci. Eng.*, **9**, 90
- Jenkins, C. R., & Peacock, J. A. 2011, *MNRAS*, **413**, 2895
- Kobyzev, I., Prince, S. J., & Brubaker, M. A. 2021, *IEEE Trans. Patt. Anal. Mach. Intell.*, **43**, 3964
- Konrad, B. S., Alei, E., Quanz, S. P., et al. 2022, *A&A*, **664**, A23
- Lange, J. U. 2023, *MNRAS*, **525**, 3181
- Lipman, Y., Chen, R. T. Q., Ben-Hamu, H., et al. 2022, arXiv e-prints [arXiv:2210.02747], ICLR 2023
- Loshchilov, I., & Hutter, F. 2017, ICLR 2019, [arXiv:1711.05101]
- Madhusudhan, N. 2018, in *Handbook of Exoplanets*, eds. H. Deeg, & J. Belmonte (Cham: Springer International Publishing), 1
- McKerns, M. M., Strand, L., Sullivan, T., Fang, A., & Aivazis, M. A. G. 2012, arXiv e-prints [arXiv:1202.1056]
- McKinney, W. 2010, Proc. of the 9th Python in Science Conf., 56
- Mollière, P., Wardener, J. P., van Boekel, R., et al. 2019, *A&A*, **627**, A67
- Mollière, P., Stolker, T., Lacour, S., et al. 2020, *A&A*, **640**, A131
- Márquez-Neila, P., Fisher, C., Sznitman, R., & Heng, K. 2018, *Nat. Astron.*, **2**, 719
- NASA 2023, Habitable Worlds Observatory, <https://science.nasa.gov/astrophysics/programs/habitable-worlds-observatory>
- Nasedkin, E., Mollière, P., Wang, J., et al. 2023, *A&A*, **678**, A41
- Nasedkin, E., Mollière, P., & Blain, D. 2024a, *JOSS*, **9**, 5875
- Nasedkin, E., Mollière, P., Lacour, S., et al. 2024b, *A&A*, **687**, A298
- Nixon, M. C., & Madhusudhan, N. 2020, *MNRAS*, **496**, 269
- Papamakarios, G., & Murray, I. 2016, NeurIPS 2016, [arXiv:1912.01703]
- Papamakarios, G., Nalisnick, E., Rezende, D. J., Mohamed, S., & Lakshminarayanan, B. 2021, *J. Mach. Learn. Res.*, **22**, 1
- Paszke, A., Gross, S., Massa, F., et al. 2019, NeurIPS 2019, [arXiv:1912.01703]
- Pedregosa, F., Varoquaux, G., Gramfort, A., et al. 2011, *J. Mach. Learn. Res.*, **12**, 2825
- Petroff, M. A. 2021, arXiv e-prints [arXiv:2107.02270]
- Quanz, S. P., Ottiger, M., Fontanet, E., et al. 2022, *A&A*, **664**, A21
- Rezende, D. J., & Mohamed, S. 2015, ICML 2015, [arXiv:1505.05770]
- Skilling, J. 2006, *Bayesian Anal.*, **1**, 833
- Soboczenski, F., Himes, M. D., O’Beirne, M. D., et al. 2018, arXiv e-prints [arXiv:1811.03390]
- Speagle, J. S. 2020, *MNRAS*, **493**, 3132
- Stimper, V., Liu, D., Campbell, A., et al. 2023, *JOSS*, **8**, 5361
- Sweet, A. 2024, arXiv e-prints [arXiv:2406.10771]
- Tabak, E. G., & Vanden-Eijnden, E. 2010, *Commun. Math. Sci.*, **8**, 217
- Tahseen, T. P. A., Mendonça, J. M., Yip, K. H., & Waldmann, I. P. 2024, *MNRAS*, submitted
- Tinetti, G., Eccleston, P., Haswell, C., et al. 2021, arXiv e-prints [arXiv:2104.04824]
- Tokdar, S. T., & Kass, R. E. 2009, *WIREs Comput. Stat.*, **2**, 54
- Unlu, E. B., Forestano, R. T., Matchev, K. T., & Matcheva, K. 2023, arXiv e-prints [arXiv:2310.10521]
- Vasist, M., Rozet, F., Absil, O., et al. 2023, *A&A*, **672**, A147
- Virtanen, P., Gommers, R., Oliphant, T. E., et al. 2020, *Nat. Methods*, **17**, 261
- Waldmann, I. P. 2016, *ApJ*, **820**, 107
- Wildberger, J., Dax, M., Green, S. R., et al. 2023, *Phys. Rev. D*, **107**, 023021
- Williams, M. J., McGinn, J., Stak, F., & Veitch, J. 2024, <https://zenodo.org/doi/10.5281/zenodo.7108558>
- Yip, K. H., Changeat, Q., Nikolaou, N., et al. 2021, *AJ*, **162**, 195
- Yip, K. H., Changeat, Q., Al-Refaie, A., & Waldmann, I. P. 2024, *ApJ*, **961**, 30
- Zhang, K., Bloom, J. S., van der Walt, S., & Hernitschek, N. 2023, arXiv e-prints [arXiv:2312.03824]
- Zingales, T., & Waldmann, I. P. 2018, *AJ*, **156**, 268

Appendix A: Atmospheric parameters and priors

In table A.1, we provide an overview of the 16 atmospheric parameters we consider (see section 3), as well as their priors. The table also shows θ_0 , that is, the parameter values for the benchmark spectrum that we use in section 4.1 and that is shown in fig. A.1. We adopted our choice of parameters and priors from Vasist et al. (2023) (cf. their Table 1 and Table 2), who in turn based it on Mollière et al. (2020).

We note two more things: The first comment concerns the name and interpretation of the parameters for the optical depth and the pressure-temperature (PT) profile. Tables 1 and 2 of Vasist et al. (2023) (and Tables 1 and 3 of Mollière et al. 2020) use a different convention than the actual Python implementation, which is the one that we adopt in table A.1. To clarify the differences and prevent any confusion, we explain the conventions and their connection in greater detail. The atmospheric model we use describes the PT profile using four temperature parameters in Kelvin, called T_0 through T_3 by Mollière et al. (2020) and Vasist et al. (2023), as well as the parameters α and δ , which relate the optical depth τ to the atmospheric pressure P via a power law,

$$\tau = \delta \cdot P^\alpha. \quad (\text{A.1})$$

The temperature T_0 , which is called T_{int} in the petitRADTRANS code, describes the interior temperature of the planet and appears as a free parameter in the Eddington approximation used for the photosphere (see eq. 2 in Mollière et al. 2020),

$$T(\tau)^4 = 3/4 \cdot T_0^4 \cdot (2/3 + \tau). \quad (\text{A.2})$$

For T_1 to T_3 , which are the support points of a cubic spline description of the PT profile at high altitudes, the priors defined by Mollière et al. (2020) and Vasist et al. (2023) are not independent; for example, the prior assumed for T_1 is $\mathcal{U}(0, T_2)$. From a physical perspective, this makes sense, as we want to enforce $T_{\text{connect}} > T_3 > T_2 > T_1$, where T_{connect} is the temperature of the uppermost photosphere layer at $\tau = 0.1$; see eq. (A.2). However, for sampling, it is more efficient to use a parameterization where we can draw all atmospheric parameters independently. Presumably for this reason, the petitRADTRANS method handling the simulation of the spectra (namely `emission_model_diseq()`) uses such an independent re-parameterization: The function expects four temperature arguments: T_{int} in Kelvin, and T_1, T_2, T_3 all dimensionless in $[0, 1]$. These inputs are then overwritten as

$$\begin{aligned} T_3 &\leftarrow \sqrt[4]{3/4 \cdot T_{\text{int}}^4 \cdot (0.1 + 2/3) \cdot (1 - T_3)}, \\ T_2 &\leftarrow T_3 \cdot (1 - T_2), \\ T_1 &\leftarrow T_2 \cdot (1 - T_1). \end{aligned}$$

The temperature values from Table 2 in Vasist et al. (2023) can be obtained by inserting the respective θ_0 values from table A.1 into the above conversion formulas. Finally, `emission_model_diseq()` expects an input argument `log_delta` in $[0, 1]$, which is, however, not simply the logarithm of the δ parameter in eq. (A.1). Instead, `log_delta` is used internally to compute δ via the transformation

$$\delta \leftarrow (10^6 \cdot 10^{-3+5 \cdot \log_delta})^{-\alpha}.$$

As before, plugging the θ_0 values from table A.1 into this transformation and taking the logarithm yields the $\log \delta$ value from Table 2 in Vasist et al. (2023). We have chosen to report the prior ranges and θ_0 values of the independent re-parameterization in table A.1 (i.e., the inputs to the `emission_model_diseq()`

function before internal conversions) to maintain consistency with the code and to use the same prior bounds for all retrievals.

The second comment concerns the parameters that Vasist et al. (2023) use to describe the cloud mass fraction abundance, namely $\log \tilde{X}_{\text{Fe}}$ and $\log \tilde{X}_{\text{MgSiO}_3}$. The meaning of the corresponding arguments to the `emission_model_diseq()` function, `log_X_cb_Fe(c)` and `X_cb_MgSiO3(c)`, changed in version 2.4.8 of petitRADTRANS. To get the same behavior as before, we use the arguments `eq_scaling_Fe(c)` and `eq_scaling_MgSiO3(c)` instead and label the parameters as $S_{\text{eq,Fe}}$ and $S_{\text{eq,MgSiO}_3}$, respectively. The prior ranges and θ_0 values remain unchanged.

Appendix B: Model and training details

In this appendix, we describe the model configurations and details of the training procedure we used for the experiments in section 4. For the full technical details, we recommend a look at our code.

B.1. Flux and parameter normalization

To rescale the flux values that are passed to the neural networks, we follow Vasist et al. (2023) and use a softclip function, where each value x_i is mapped to $x_i \mapsto x_i / (1 + |x_i/\beta|) \in [-\beta, \beta]$ with bound $\beta = 100$. The atmospheric parameters are standardized by subtracting the mean μ and dividing by the standard deviation σ . Both μ and σ are computed from the corresponding prior.

B.2. Model configurations

FMPE model. There are three parts to the FMPE model: a context encoder, a (t, θ_t) -encoder, and a network modeling the vector field. All three parts are implemented as deep residual neural networks (“resnet”) that use layer normalization (Ba et al. 2016), GELU activation functions (Hendrycks & Gimpel 2016), and dropout regularization (Hinton et al. 2012) with $p = 0.1$. The context encoder concatenates the normalized flux values and assumed error bars for each spectral bin and passes the result through a 3-layer resnet with 1024 units each to compute the context embedding. The (t, θ_t) -encoder first applies a sine/cosine positional encoding to t before stacking it with the original (t, θ_t) -tuple and passing it through another 3-layer resnet. Then, the embedded context and the embedded (t, θ_t) -values are combined using a gated linear unit (GLU; Dauphin et al. 2016) on the latter. The result is finally passed to the vector field network, which consists of 18 residual layers, with layer sizes decreasing from 8192 to 16 (i.e., the number of atmospheric parameters). The total number of learnable parameters for this configuration is 313 443 376. Our architecture and hyperparameters are based on a number of preliminary experiments to identify a combination that works well; however, a more systematic optimization is left for future work.

NPE model. The model consists of two parts: an encoder to compress the context to a lower-dimensional representation, and the actual DNF. For the context encoder, we normalize the flux values (see above) and concatenate them with the error bars before passing them to a 15-layer resnet with layer sizes decreasing from 4096 to 256. We use the same activation and regularization layers as for the FMPE model. The output of the context encoder is passed to each step of the DNF. We use a neural spline flow (NSF; Durkan et al. 2019) as implemented by the `glasflow` library (Williams et al. 2024), which is itself based on `nflows` (Durkan et al. 2020). The NSF consists of 16 steps, each consisting of 4 blocks with 1024 units. The number of bins for the splines is 16. This configuration results in a total

Table A.1: Overview of the 16 atmospheric parameters considered in this work, with the respective priors used for all retrievals, as well the parameter values θ_0 of the “benchmark spectrum” (cf. section 4.1).

Parameter	Prior	θ_0 value	Meaning
C/O	$\mathcal{U}(0.1, 1.6)$	0.55	Carbon-to-oxygen ratio
[Fe/H]	$\mathcal{U}(-1.5, 1.5)$	0.00	Metallicity
$\log P_{\text{quench}}$	$\mathcal{U}(-6.0, 3.0)$	-5.00	Pressure below which the CO, CH ₄ and H ₂ O abundances are taken as vertically constant
$S_{\text{eq,Fe}}$	$\mathcal{U}(-2.3, 1.0)$	-0.86	Scaling factor for equilibrium cloud abundances (Fe)
$S_{\text{eq,MgSiO}_3}$	$\mathcal{U}(-2.3, 1.0)$	-0.65	Scaling factor for equilibrium cloud abundances (MgSiO ₃)
f_{sed}	$\mathcal{U}(0.0, 10.0)$	3.00	Sedimentation parameter for both cloud types
$\log K_{zz}$	$\mathcal{U}(5.0, 13.0)$	8.50	Vertical eddy diffusion coefficient (in cm ² s ⁻¹)
σ_g	$\mathcal{U}(1.05, 3.0)$	2.00	Width of the (log-normal) cloud particle size distribution
$\log g$	$\mathcal{U}(2.0, 5.5)$	3.75	Surface gravity
R_p	$\mathcal{U}(0.9, 2.0)$	1.00	Planet radius (in Jupiter radii)
T_{int}	$\mathcal{U}(300, 2300)$	1063.60	Interior temperature of the planet (in Kelvin)
T_3	$\mathcal{U}(0.0, 1.0)$	0.26	Inner support point for pressure-temperature spline at high altitudes
T_2	$\mathcal{U}(0.0, 1.0)$	0.29	Middle support point for pressure-temperature spline at high altitudes
T_1	$\mathcal{U}(0.0, 1.0)$	0.32	Outer support point for pressure-temperature spline at high altitudes
α	$\mathcal{U}(1.0, 2.0)$	1.39	Exponent in $\tau = \delta P^\alpha$ (i.e., the assumed power law for the optical depth τ)
$\log \delta$	$\mathcal{U}(0.0, 1.0)$	0.48	Proportionality factor in $\tau = \delta P^\alpha$ (see above)

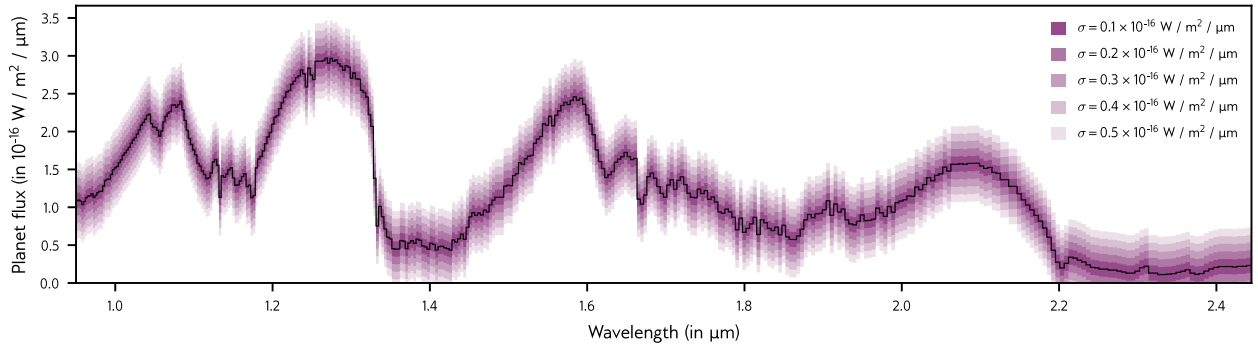


Fig. A.1: Noise-free benchmark spectrum (black line) obtained by running the simulator for θ_0 from table A.1. The colored bands illustrate different noise levels up to the largest one used for training our noise-level-conditional models (see section 3).

of 310 120 848 learnable parameters. As for FMPE, the NPE model may benefit from more hyperparameter optimization, which we leave for future work.

B.3. Training procedure

We train both our FMPE and our NPE model for 600 epochs using stochastic gradient descent in the form of AdamW (Loshchilov & Hutter 2017), using a batch size of 16 384, an initial learning rate of 5×10^{-5} , and a CosineAnnealingLR learning rate scheduler. Additionally, we apply gradient clipping on the ℓ_2 norm of the gradients. We only save the model that achieves the lowest validation loss. For FMPE, we train with automatic mixed precision (AMP) as it speeds up training significantly, while we found no such advantage for NPE. One training epoch takes 328 ± 35 s/epoch (total: 55 hours) for the FMPE model, while the average for the NPE model is 975 ± 313 s/epoch (total: 163 hours). All our models are trained on a single NVIDIA H100 GPU.

For the time prior $p(t)$ in the first expectation value of eq. (8), we found empirically that sampling $t^4 \sim \mathcal{U}(0, 1)$ works well in practice. Compared to sampling $t \sim \mathcal{U}(0, 1)$, this approach places more emphasis on learning the vector field at later times, where $v_{t,x}$ is typically more complicated than at early times: Intuitively, at $t \approx 0$, each sample only needs to be moved approximately in the right direction, whereas at $t \approx 1$, each sample

needs to be moved to its exact final location under the target distribution. Finally, we set $\sigma_{\min} = 1 \times 10^{-4}$ for the $\sigma_{\min}$ in eq. (6).

B.4. Inference procedure (for FMPE)

Using the FMPE model at inference time requires solving a differential equation; see eqs. (3) and (5). For this, we use the dopri5 integrator as implemented by torchdiffeq (Chen 2018). One critical parameter here is the tolerance: We found that a (relative and absolute) tolerance of 2×10^{-4} is sufficient for sampling, while a value of 5×10^{-5} is required for evaluating log-probabilities. Lower tolerances increase the runtime of the ODE solver, but do not improve the results (e.g., in terms of IS efficiency).

B.5. Test set generation

To generate a test set for our large-scale analysis, we randomly sample 500 values of θ from the Bayesian prior $p(\theta)$ and simulate the corresponding spectra, to which we add a random noise realization like in section 3. We call this our “default” test set. When manually inspecting the spectra in the default test set, we noticed that a relatively large fraction (over 20%) have a mean noise-free flux value that is below the smallest assumed noise

level. This means that, after adding noise, many of these spectra are virtually indistinguishable from pure noise, thus yielding the same posteriors. Similarly, there are also a considerable number of spectra with an average planet flux so high that they are effectively noise-free given the considered noise levels. We concluded that this might not be ideal for evaluation purposes and therefore decided to create another test set, consisting of spectra from a “neighborhood” around the benchmark spectrum: We first map θ_0 into the unit hypercube $[0, 1]^{16}$, using the inverse cumulative density functions of the priors. Based on the result u_0 , we generate values of θ by sampling $z \sim \mathcal{N}(0, 0.1^2 \cdot I_{16})$, clipping $u = u_0 + z$ to the unit hypercube, and mapping u back to the original parameter space. After running the simulator, we add a random noise realization in the same way as before. This results in spectra that are similar to the benchmark spectrum in terms of the mean and the total variation of the flux. Given that the z -vectors are drawn from a Gaussian, we call the result the “Gaussian” test set.

Appendix C: Nested sampling baselines

We use two Python-based implementations of nested sampling to create the baselines for our experiments in section 4.1:

nautilus (Lange 2023) is a recent implementation of nested sampling that utilizes IS and makes use of neural networks to explore the target space more efficiently and improve the overall runtime. In our experiments, we reliably found *nautilus* to be significantly faster than its competitors. We set the number of live points to 4000 and otherwise use the default settings of version 1.0.3 of the package. In particular, we leave the target number of posterior samples at 10 000, and note that increasing this value often results in super-linear increases in runtime. This is because the ESS does not grow monotonically with the number of proposals: drawing a posterior sample that increases the variance of the importance weights can also reduce the ESS.

MultiNest (Feroz et al. 2009, 2019) is perhaps the most popular nested sampling implementation in the atmospheric retrieval community (see, e.g., Alei et al. 2024 or Nasedkin et al. 2024b for recent examples) and was also used as a baseline by Vasist et al. (2023). It is typically run using the PyMultiNest wrapper (Buchner et al. 2014). We set the number of live points to 4000 and otherwise use the default settings of version 2.12 of the package. PyMultiNest does not provide an option to control the target number of posterior samples.

All nested sampling experiments were parallelized over 96 CPU cores. We also experimented with two more nested sampling implementations, which we did not end up including as baselines:

dynesty (Speagle 2020) implements dynamic nested sampling to focus more on estimating the posterior instead of the Bayesian evidence (see Higson et al. 2018). After initial experiments using default settings resulted in poor estimates of the posterior (assessed by comparison with the results from other methods), we increased the number of live points to 5000 and set `sample="rwalk"` and `walks=100`. Despite

a runtime of over 9 weeks, the resulting posterior was again in strong disagreement with the posteriors obtained by all other methods, leading us not to include the *dynesty* results as a baseline.

UltraNest (Buchner 2021, 2022) implements a number of new techniques to improve the correctness and robustness of the results. However, after running *UltraNest* with default settings for more than four weeks on our benchmark retrieval, the logs indicated that the method was still far from convergence and thus not feasible as a baseline.

For easier comparison, the (partial) results of all four nested sampling implementations on the noise-free benchmark spectrum from section 4.1 are summarized in table C.1. We note further that *dynesty* and *UltraNest* are in good agreement with *nautilus* and *MultiNest* for test retrievals using only four atmospheric parameters. This leads us to believe that the issues we encountered in the full 16-dimensional case are likely due to a suboptimal choice of hyperparameters, and that a better configuration might yield better results. However, due to the high computational cost and long runtimes, optimizing the sampler settings is a substantial challenge in itself that is beyond the scope of this work.

Appendix D: Additional results

Figure D.1 contains the full 16-parameter version of fig. 3.

Appendix E: Effect of training set size

We obtained the main results of this paper using ML models that were trained on a dataset of $2^{25} = 33\,554\,432$ (“32M”) simulated spectra. To understand how the performance of FMPE and NPE depends on and scales with the amount of training data, we perform an ablation study where we retrain our models using subsets of the original training set consisting of $2^{24} = 16\,777\,216$ (“16M”), $2^{23} = 8\,388\,608$ (“8M”), $2^{22} = 4\,194\,304$ (“4M”), $2^{21} = 2\,097\,152$ (“2M”) and $2^{20} = 1\,048\,576$ (“1M”) spectra, respectively. We train each model using the same settings as in the main part (see appendix B for details). For each model, we then run 100 retrievals on our two test sets (default and Gaussian). In total, this corresponds to 2 model types $\times$ 2 test sets $\times$ 5 new training set sizes $\times$ 100 test spectra = 2000 additional retrievals. For each retrieval, we used 100 000 proposal samples for IS and computed the sampling efficiency ε as a metric for the quality of the proposal distribution. The results of this ablation study are shown in fig. E.1. As expected, we observe a clear trend for both methods where the sampling efficiency increases with the size of the training set, which implies that the accuracy of the proposal distribution also increases with the amount of training data.

Table C.1: Comparison of the four different nested sampling implementations on the noise-free benchmark retrieval from section 4.1.

Method	Live points	Likelihood calls	$\log Z$	N_{total}	N_{eff}	ε (in %)	Wall time (in h)
dynesty ⁽¹⁾	5000	18 730 100	401.497 ± 0.080	206 506	41 755.9	20.2	1548.1
MultiNest	4000	17 786 920	407.182 ± 0.080	30 540	—	—	326.2
nautilus	4000	1 042 944	401.844 ± 0.010	331 008	10 004.4	3.1	19.5
UltraNest ⁽²⁾	400	> 34 319 279	> 213.0	—	—	—	> 813.5

Notes. (1) While the $\log Z$ estimate from dynesty is similar to the other methods (including FMPE-IS / NPE-IS), the 1D and 2D marginal posteriors deviate strongly from all other methods. (2) We aborted the UltraNest run when the number of likelihood evaluations exceeded the size of our ML training set and the sampler was still far from convergence (based on its $\log Z$ estimate).

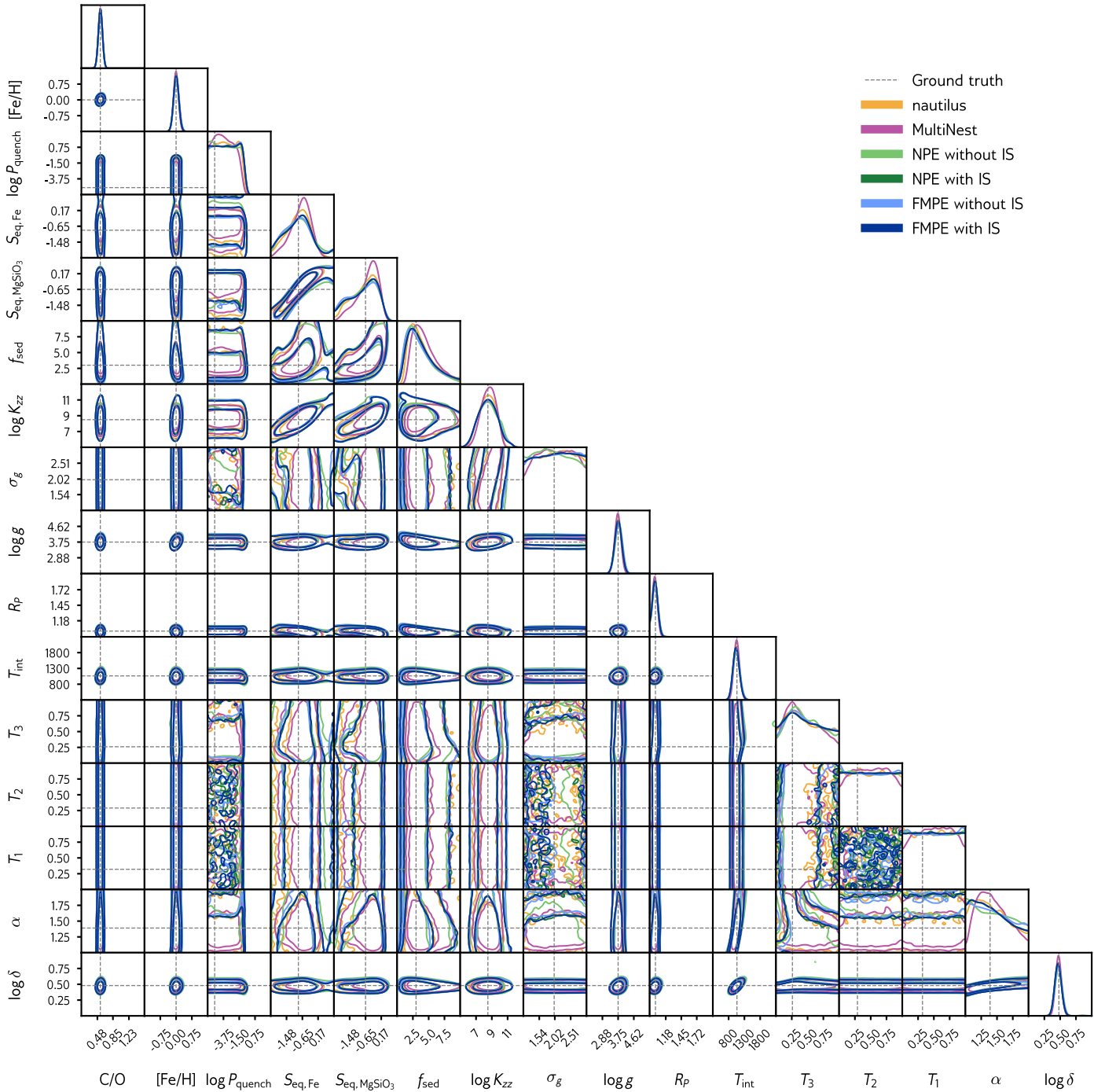
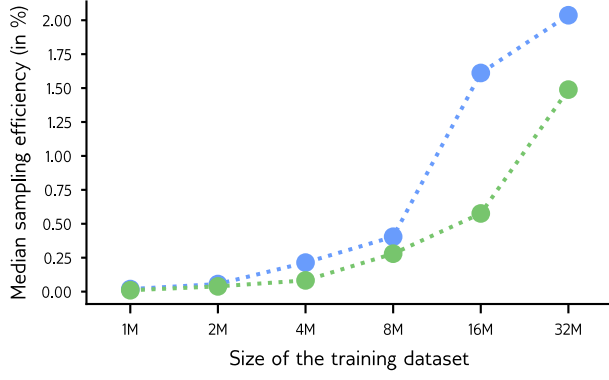
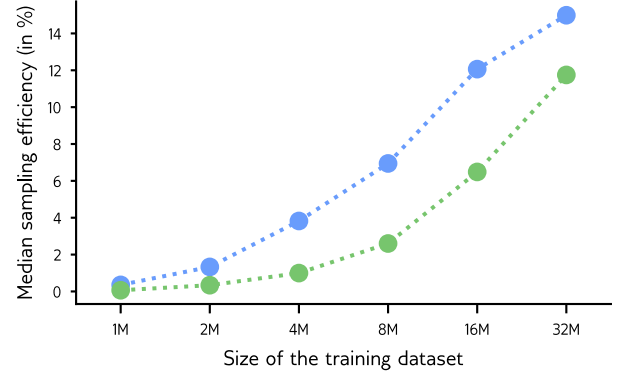


Fig. D.1: Cornerplot comparing the posteriors for the noise-free benchmark spectrum for all methods and all 16 parameters (cf. fig. 3).



(a) Results on default test set.



(b) Results on Gaussian test set.

Fig. E.1: Sampling efficiency as a function of the number of spectra used for training, for both FMPE and NPE.