

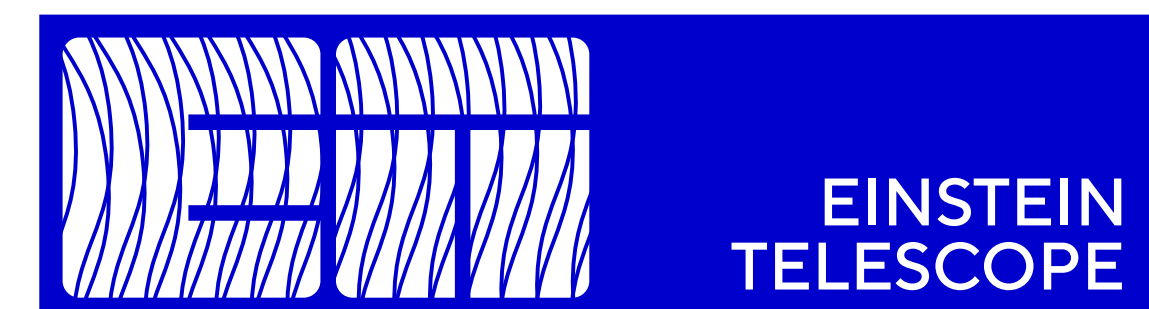
Fast and accurate parameter estimation of high-redshift sources with the Einstein Telescope

Filippo Santoliquido

Jacopo Tissino, Ulyana Dupletsa, Jan Harms, Marica Branchesi, Manuel Arca Sedda, Maximilian Dax, Annalena Kofler,
Stephen R. Green, Nihar Gupte, Isobel M. Romero-Shaw, Emanuele Berti

17 June 2025

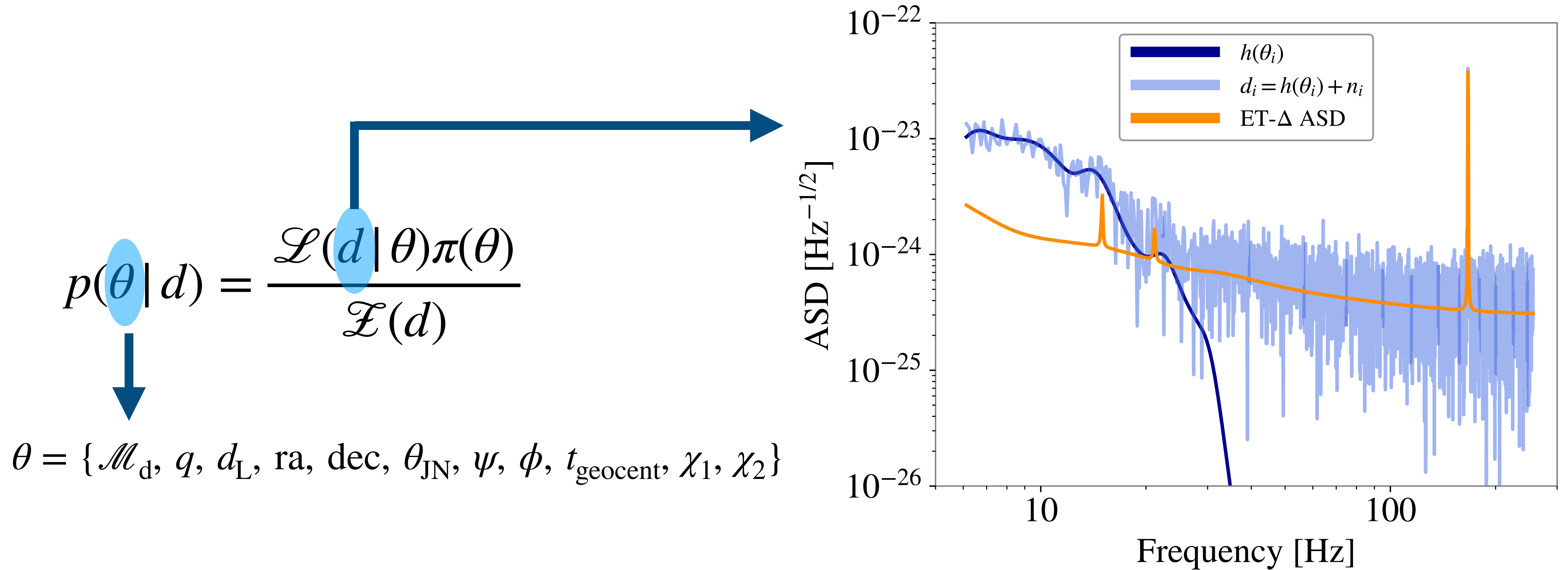
arXiv: [2504.21087](https://arxiv.org/abs/2504.21087)



Fast and accurate parameter estimation
of high-redshift sources with the Einstein Telescope

3 2
4 1

Parameter estimation



Parameter estimation

θ through **stochastic sampling** (e.g. nested sampling).
This requires tens of millions of likelihood evaluations ...

$$p(\theta | d) = \frac{\mathcal{L}(d | \theta) \pi(\theta)}{\mathcal{Z}(d)}$$

- Next generations detectors will observe $\sim 10^5$ sources per year
- Need of new methods for fast parameter estimation

... and evaluating likelihood takes time

Likelihood-free inference

Sampling parameters from the prior $\theta \sim \pi(\theta)$
and data from the likelihood $d \sim \mathcal{L}(d | \theta)$ is fast

$$p(\theta | d) = \frac{\mathcal{L}(d | \theta) \pi(\theta)}{p(d)}$$

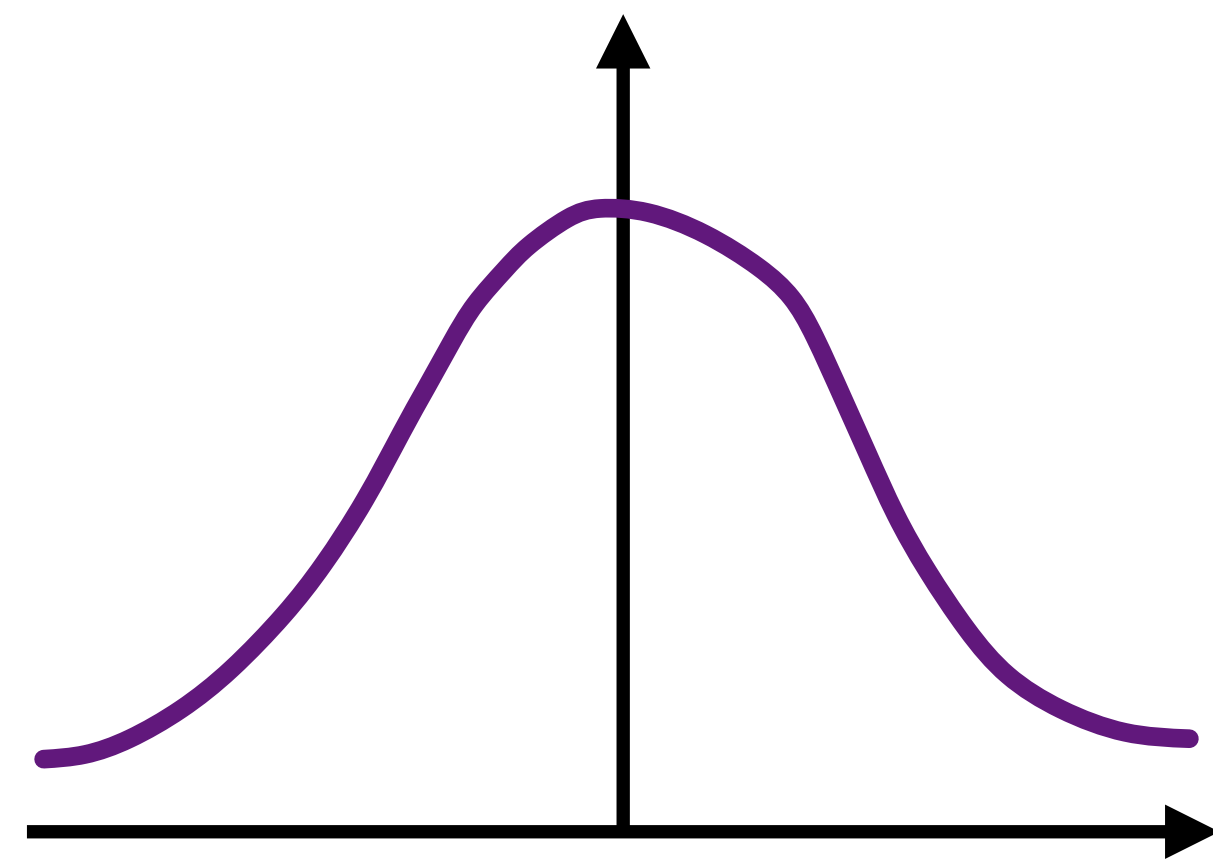
Using (θ, d) to construct with **deep learning** an estimator $q(\theta | d)$ of $p(\theta | d)$

Ref. [Lueckmann et al. 2017](#), [Greenberg et al. 2019](#), [Cranmer et al. 2020](#), [Chua et al. 2020](#), [Dax et al. 2023](#)

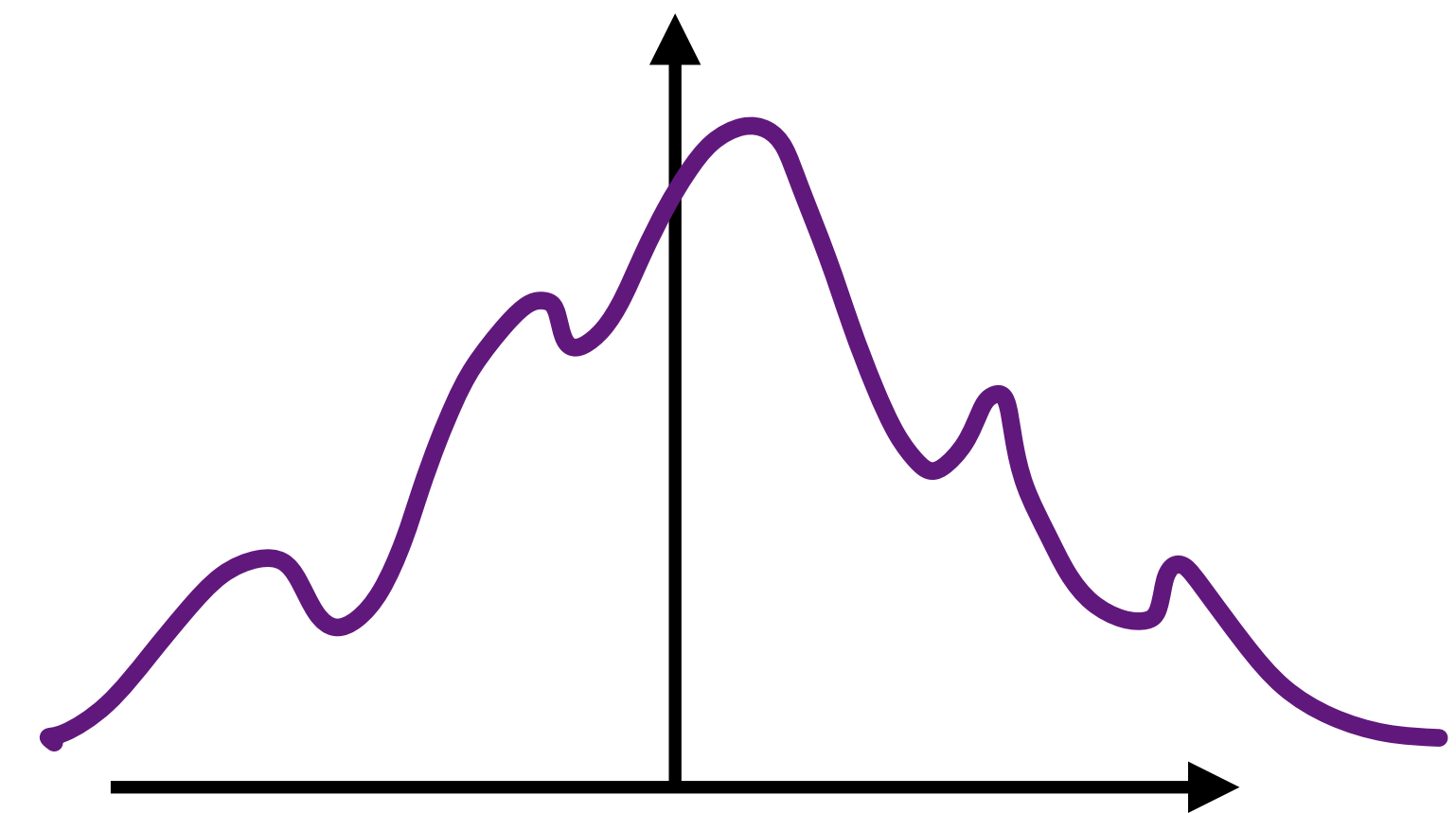


Normalising flows

Many advantages: they represent a complicated distribution q using a series of change of variables $f_d : u \rightarrow \theta$



$$\mathcal{N}^D(0,1)$$



$$q(\theta | d) = \mathcal{N}^D(0,1)(f_d^{-1}(\theta)) | \det J_{f_d}^{-1} |$$

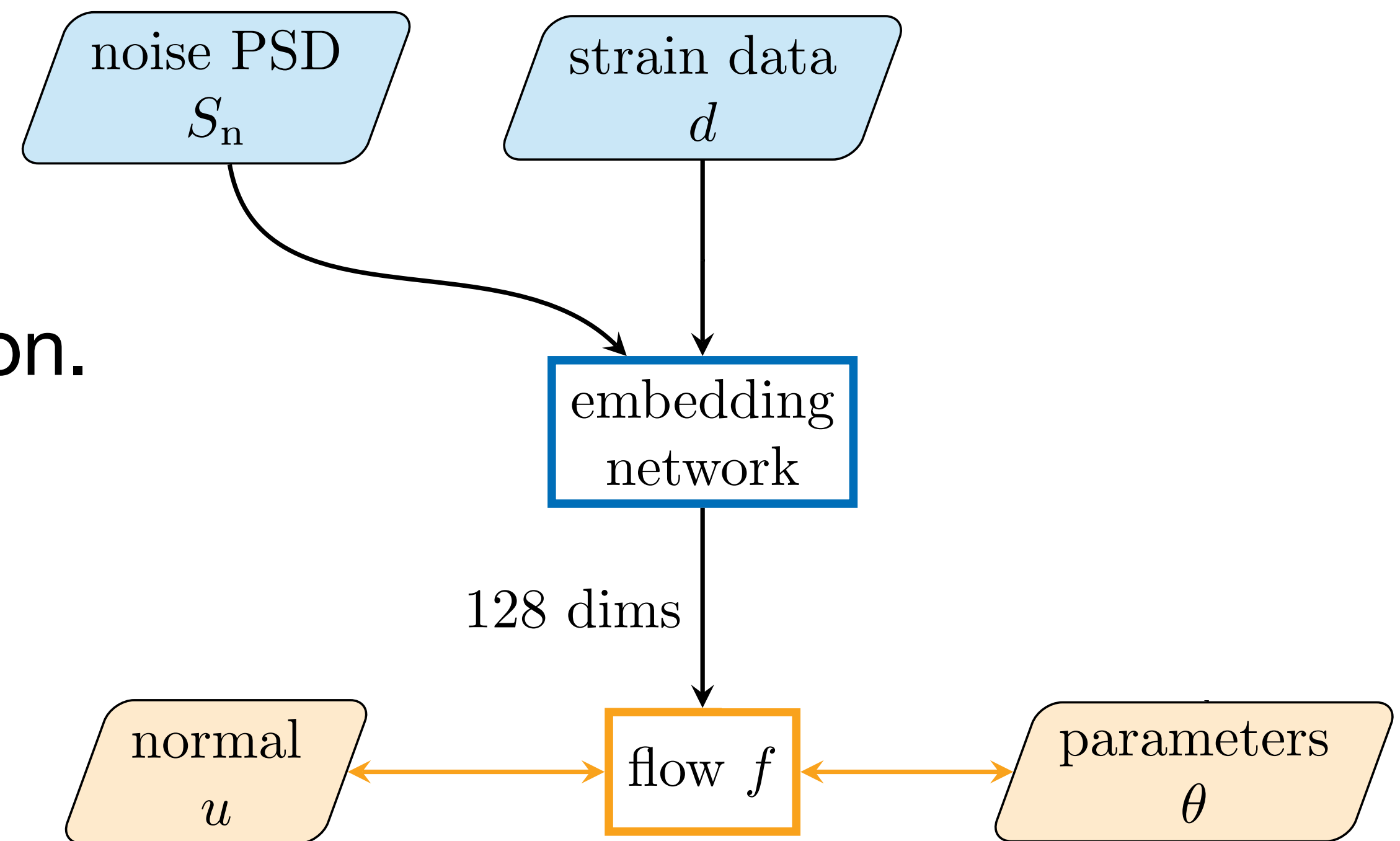
Rapidly evaluated and sampled from

Dingo



Dingo implements neural posterior estimation.

Training in **days**, inference in **minutes**

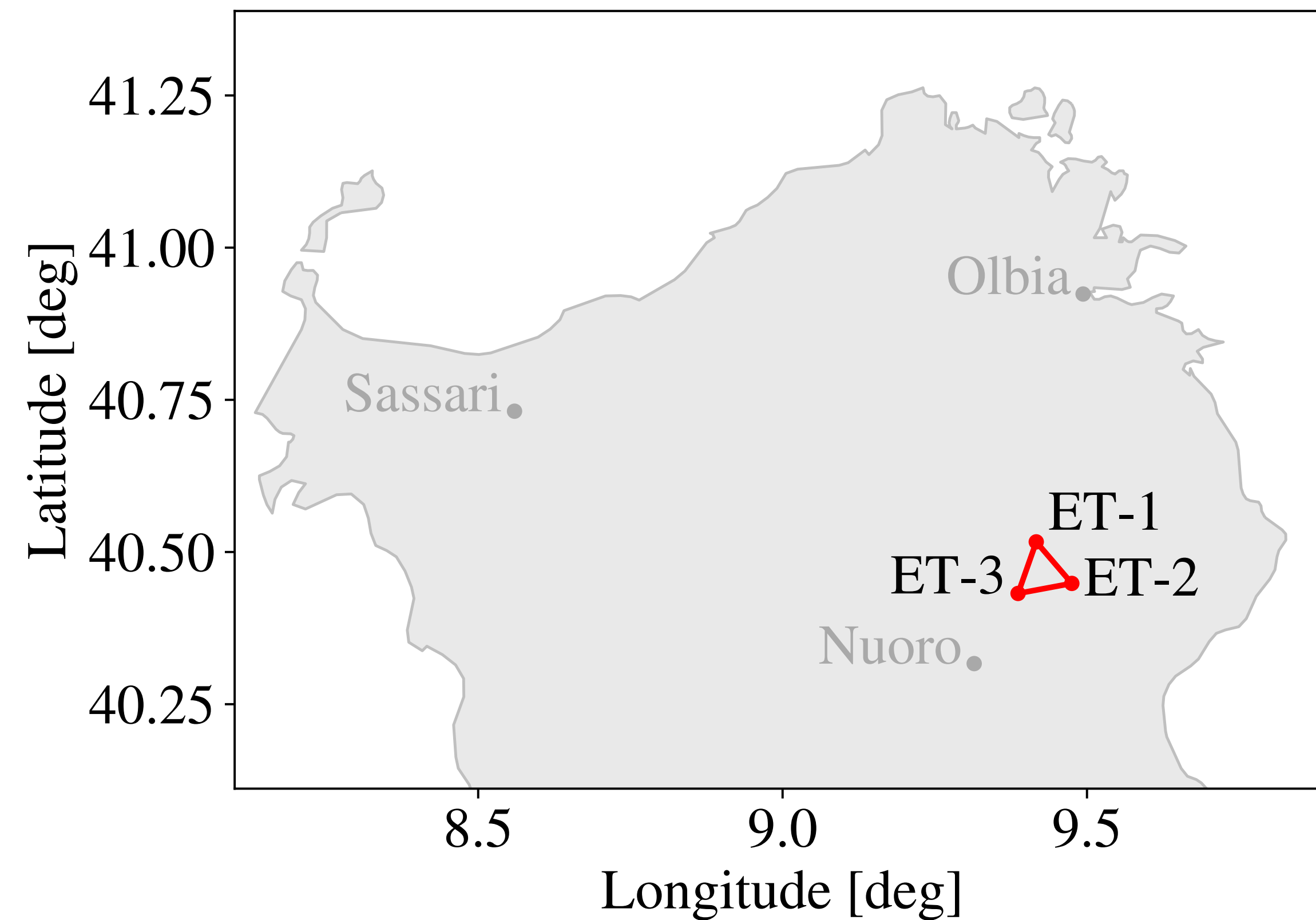


Adapted from [Dax et al. 2023](#)

Ref. [Green et al. 2020](#), [Green et al. 2021](#), [Dax et al. 2021](#), [Dax et al. 2022](#), [Wildberger et al. 2022](#), [Dax et al. 2024](#)

Einstein Telescope

I trained Dingo using the **HFLF-cryo ASD** with **ET- Δ** configuration placed in Sardinia

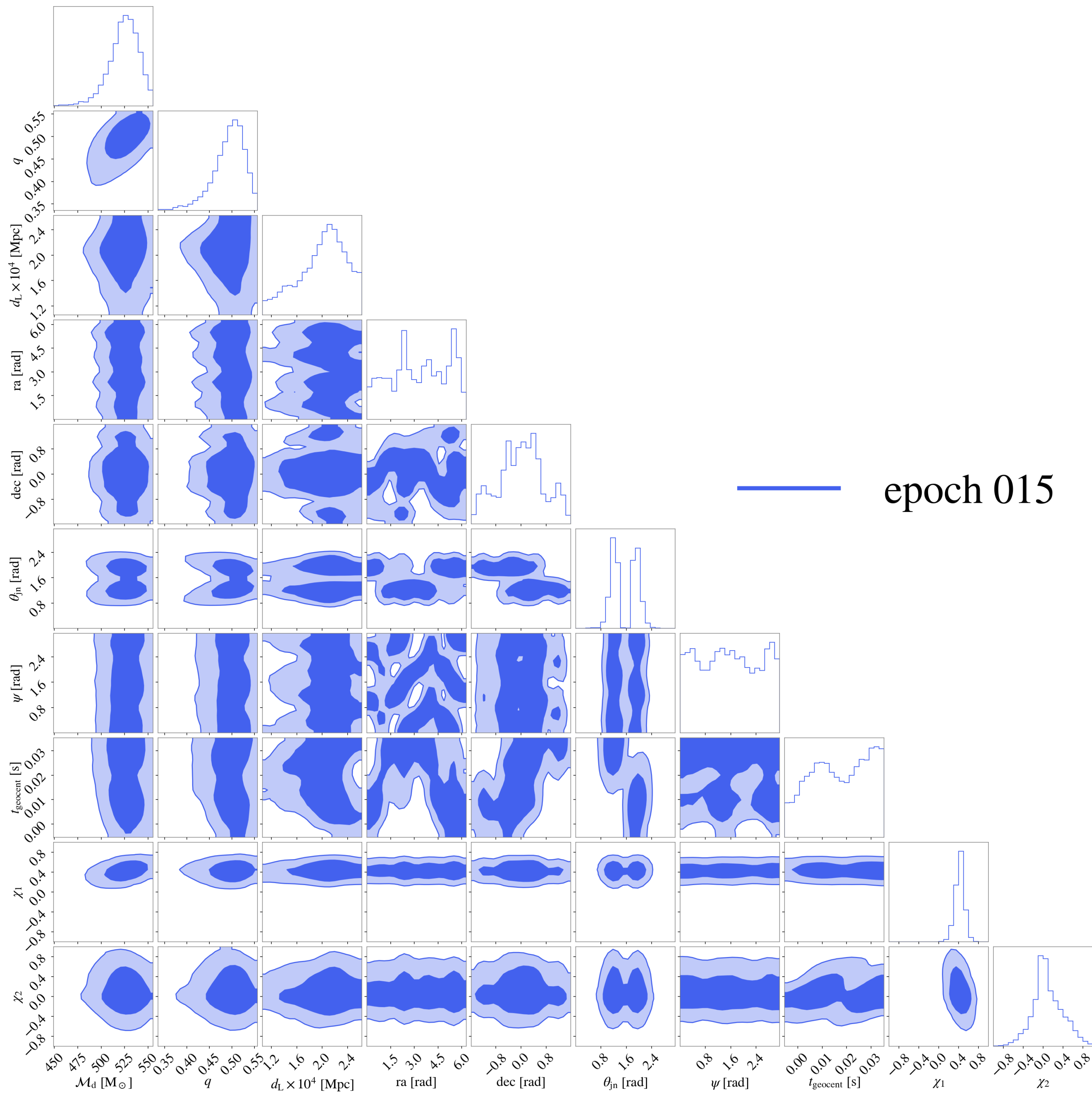


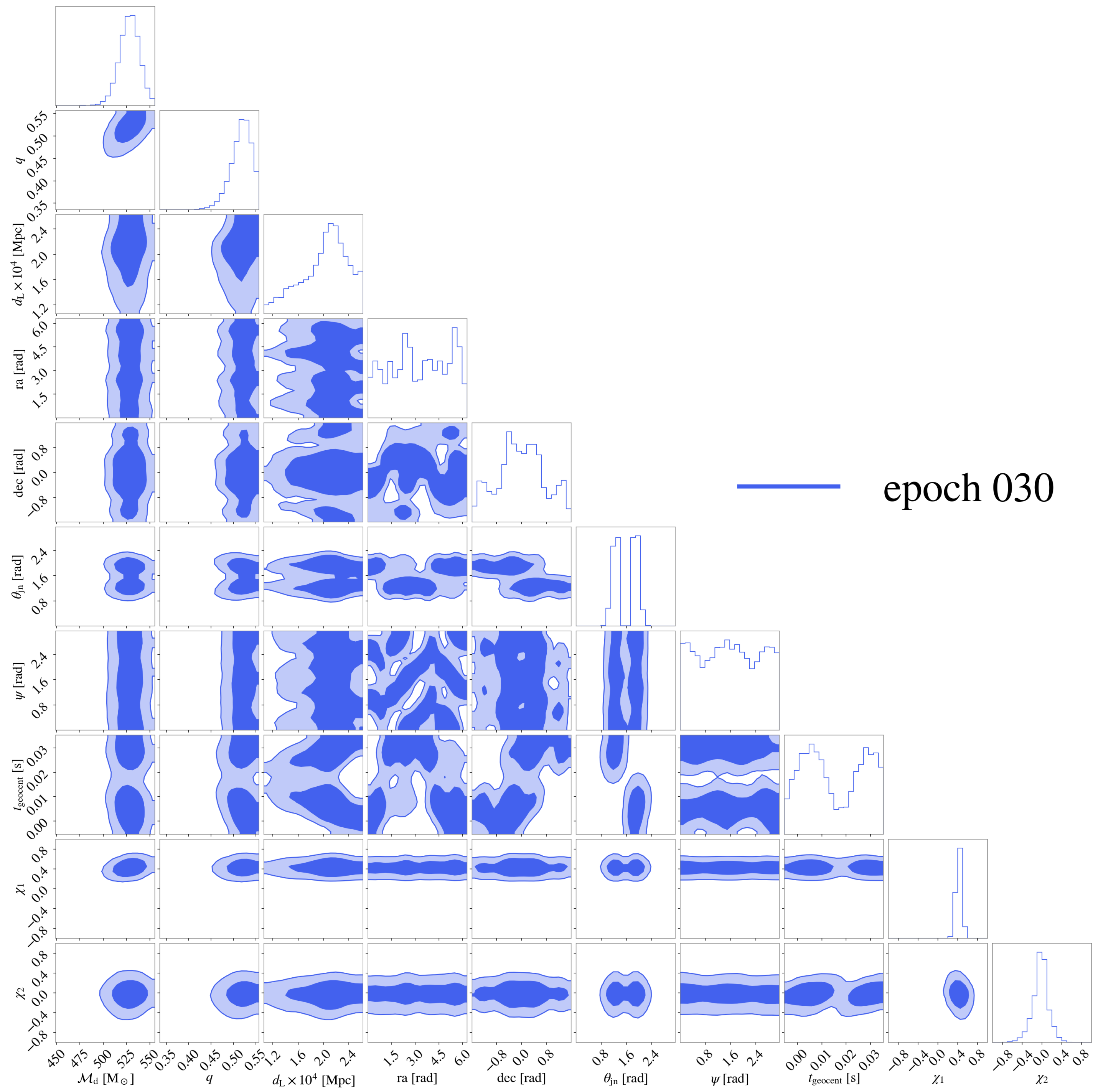
high-redshift sources

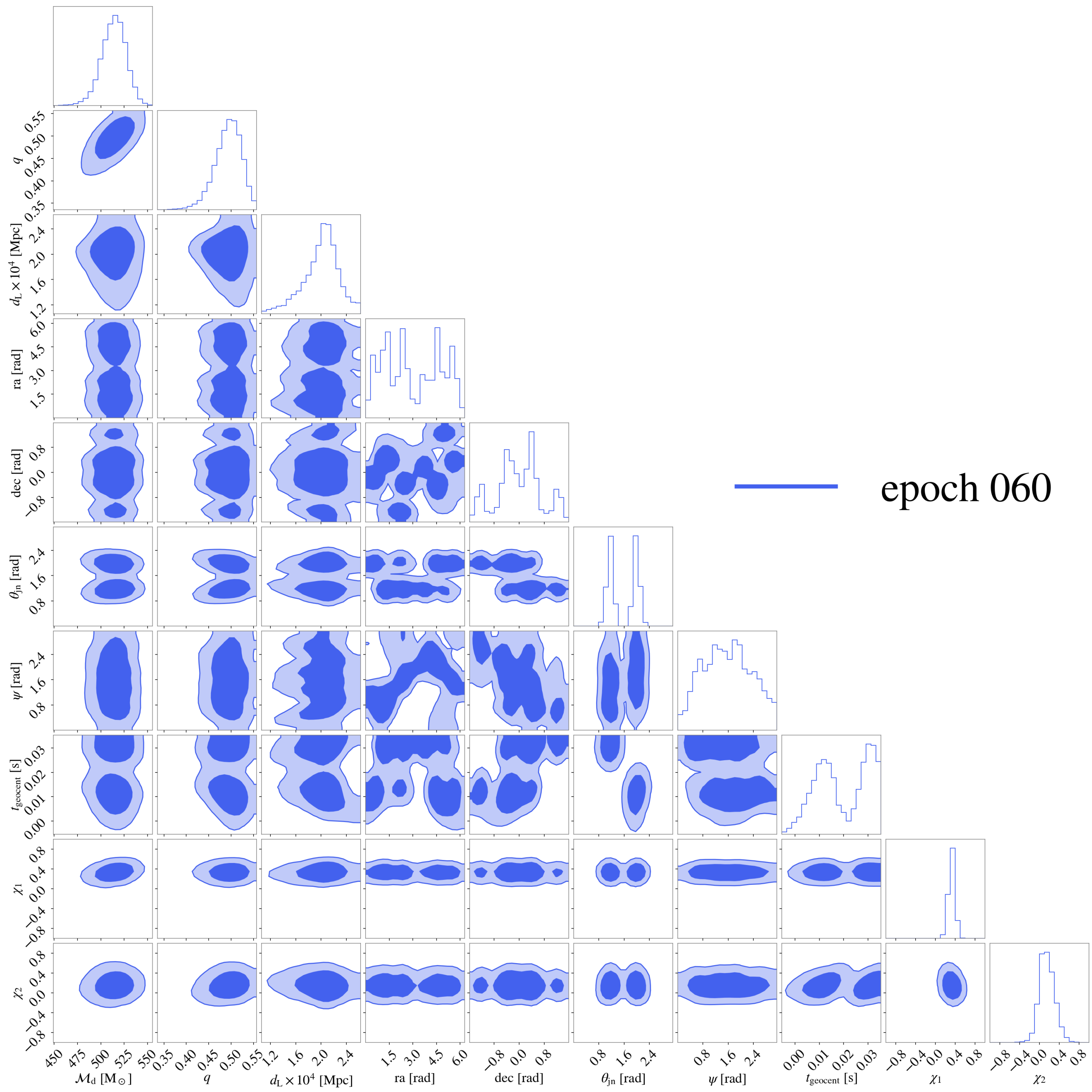
`'chirp_mass': UniformInComponentsChirpMass(minimum=40, maximum=1100)`

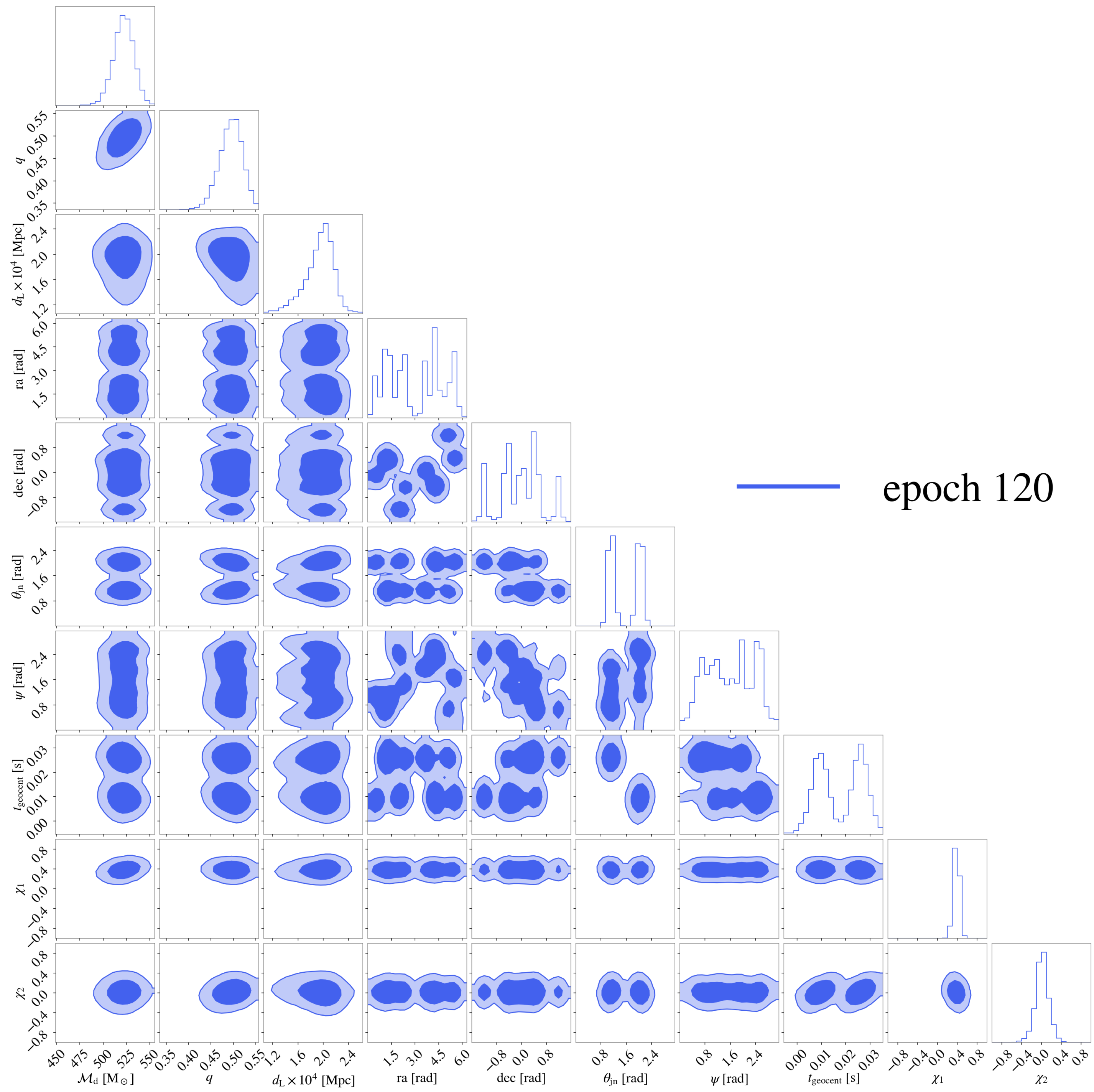
`'luminosity_distance': UniformSourceFrame(minimum=5_000.0, maximum=500_000.0)`

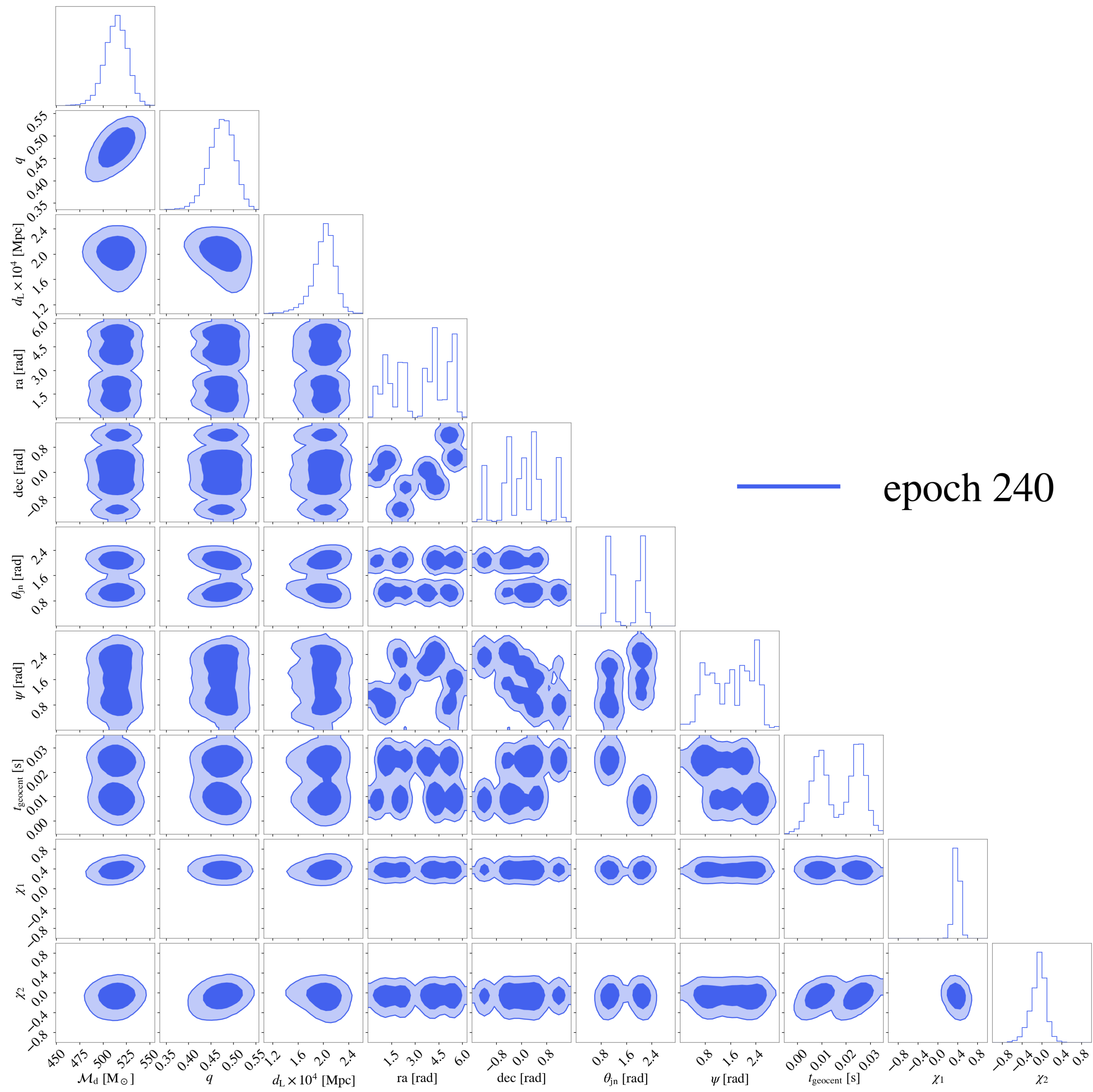
with $f \in [6, 256]$ Hz, $df = 1/8$ Hz, waveform approximant = IMRPhenomXPHM

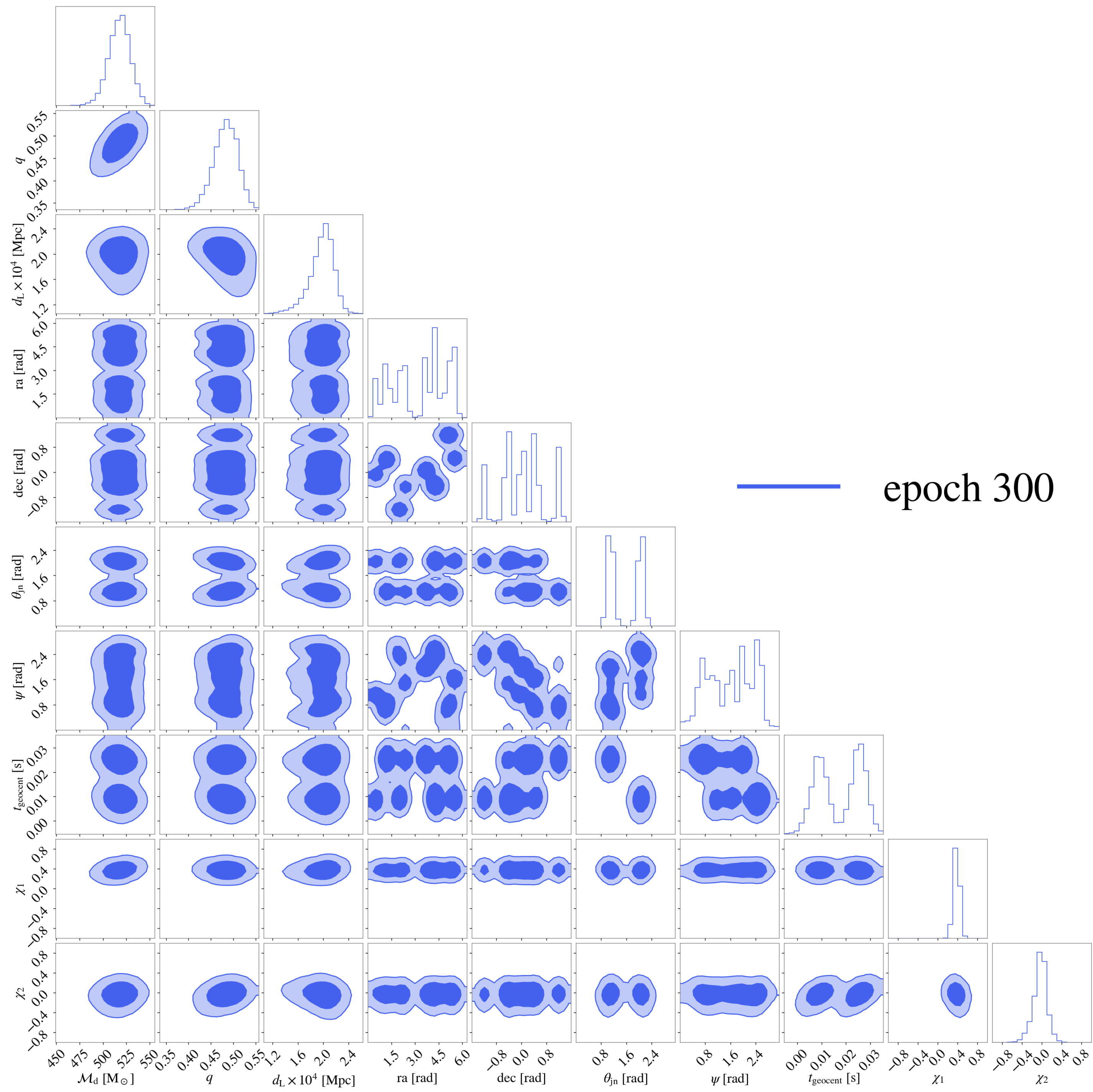










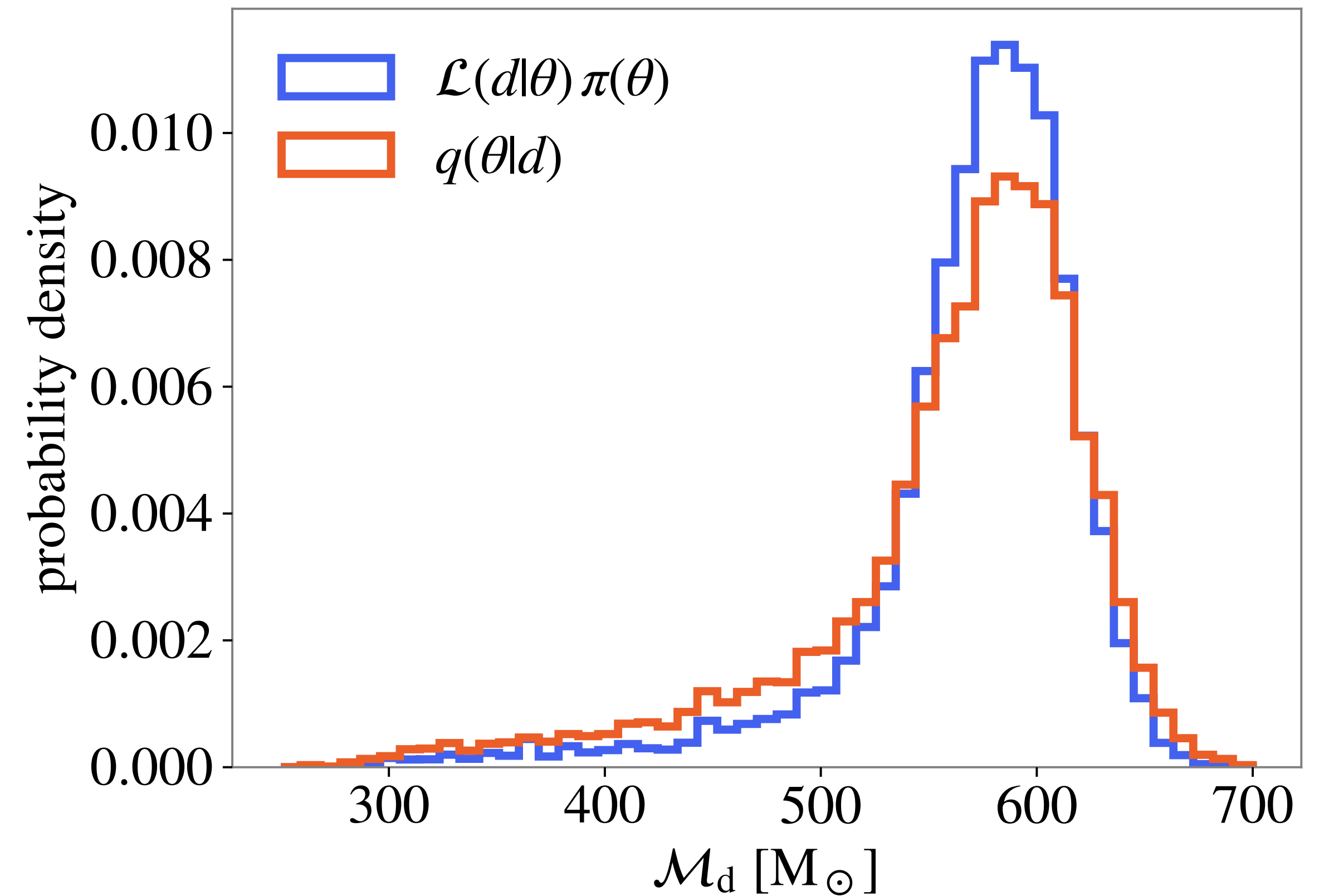


Importance sampling

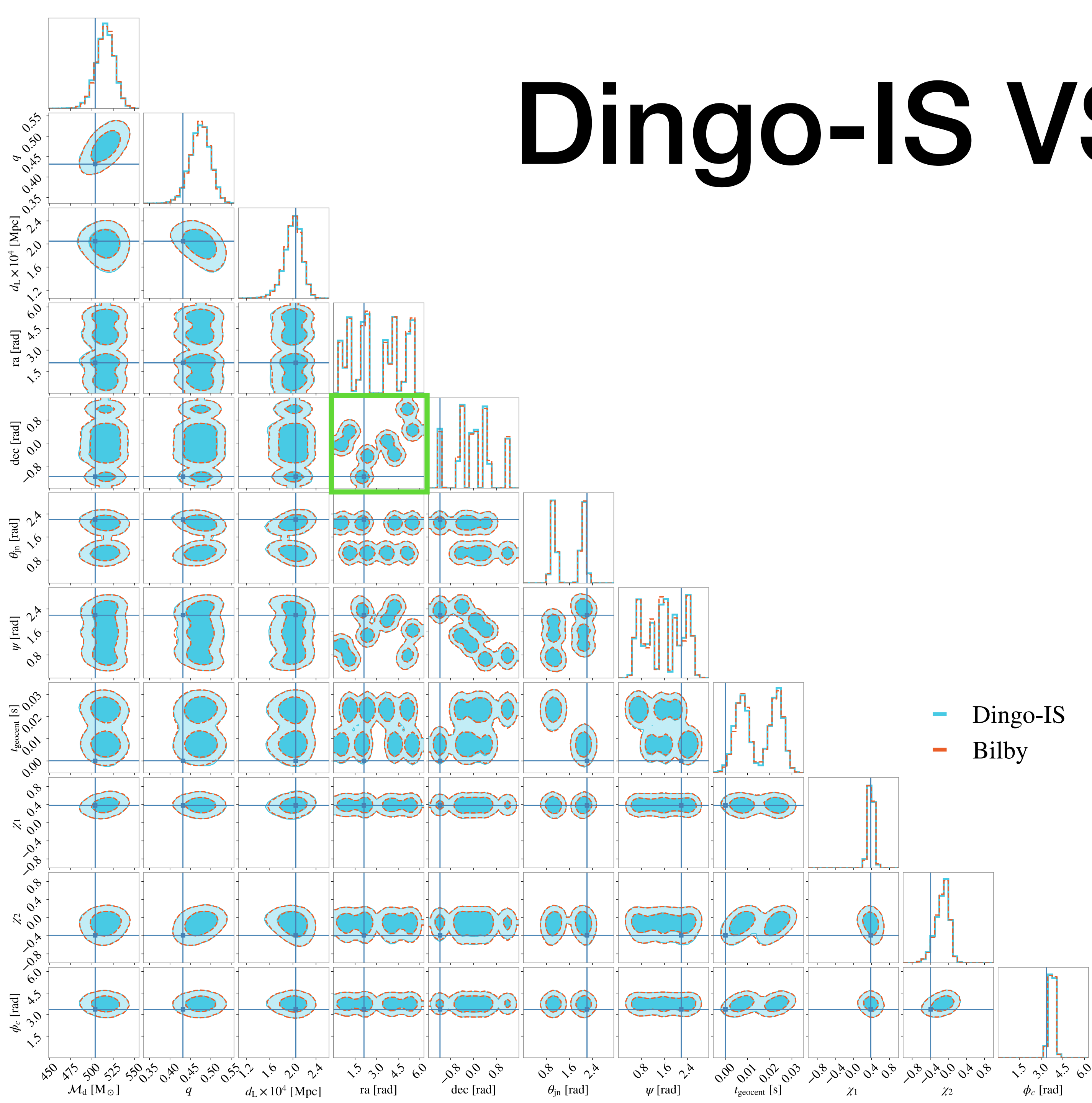
$$w_i = \frac{\mathcal{L}(d|\theta)\pi(\theta)}{q(\theta|d)}$$

Target

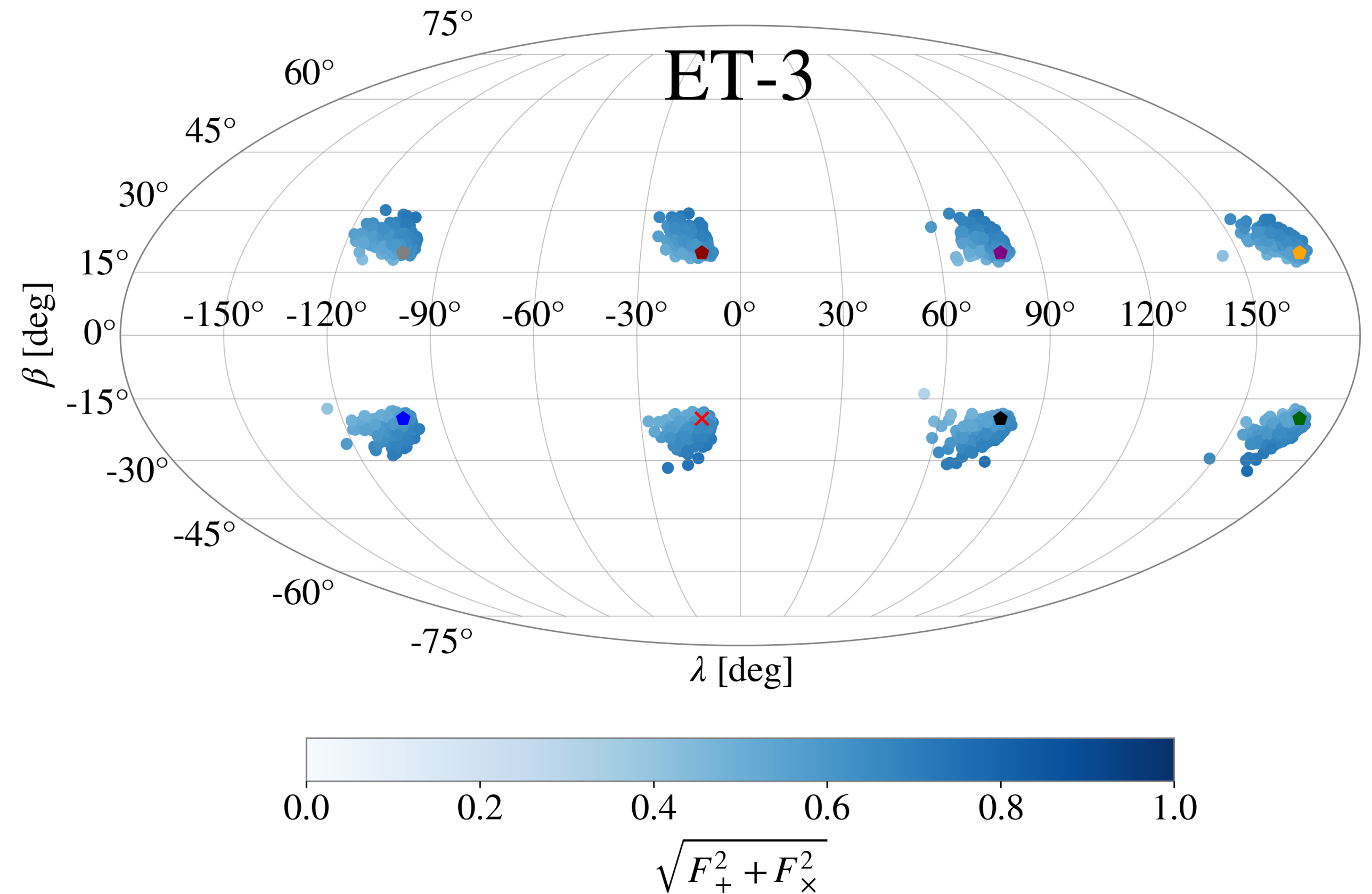
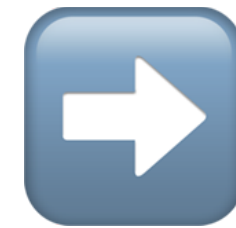
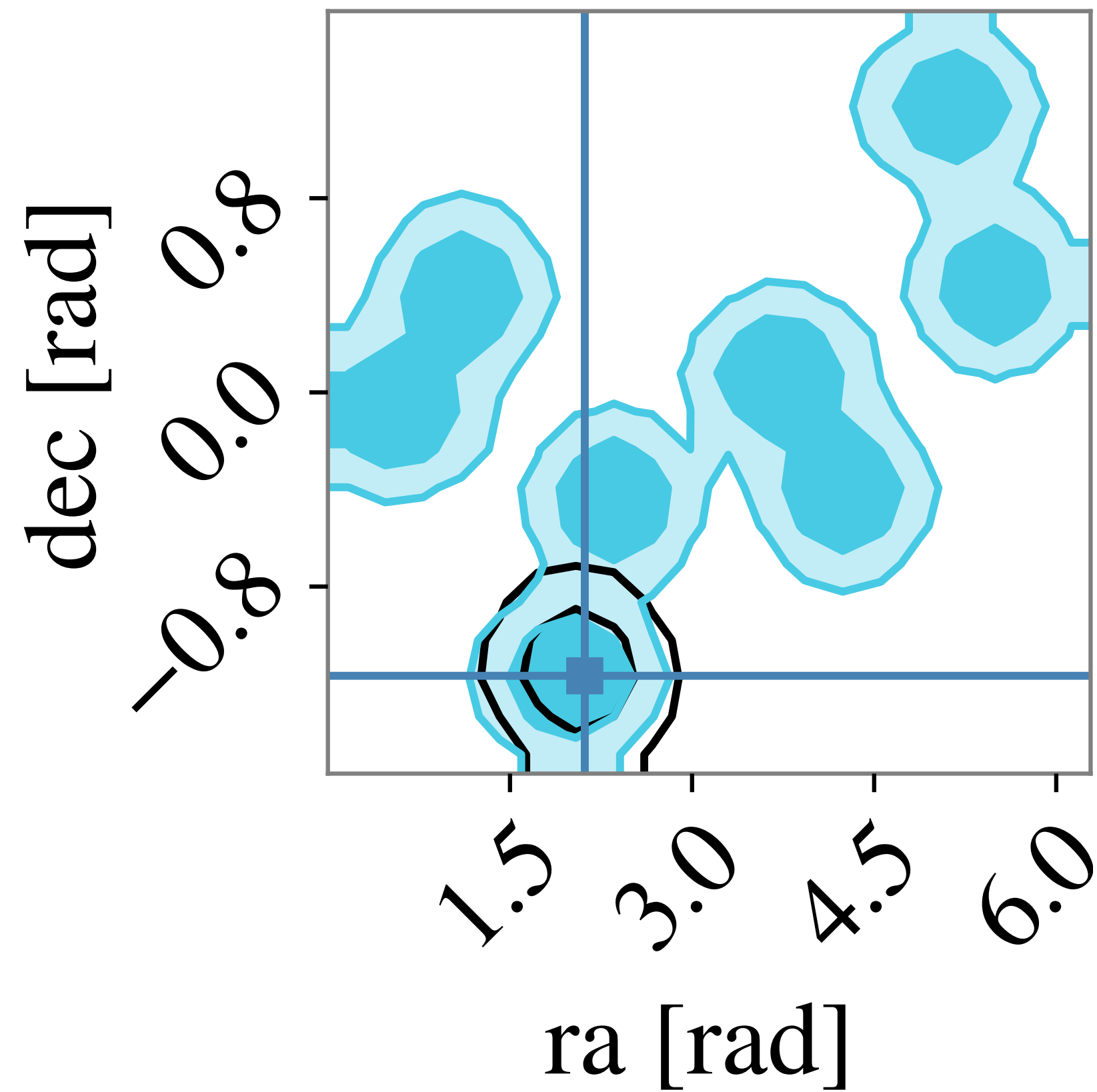
Dingo proposal



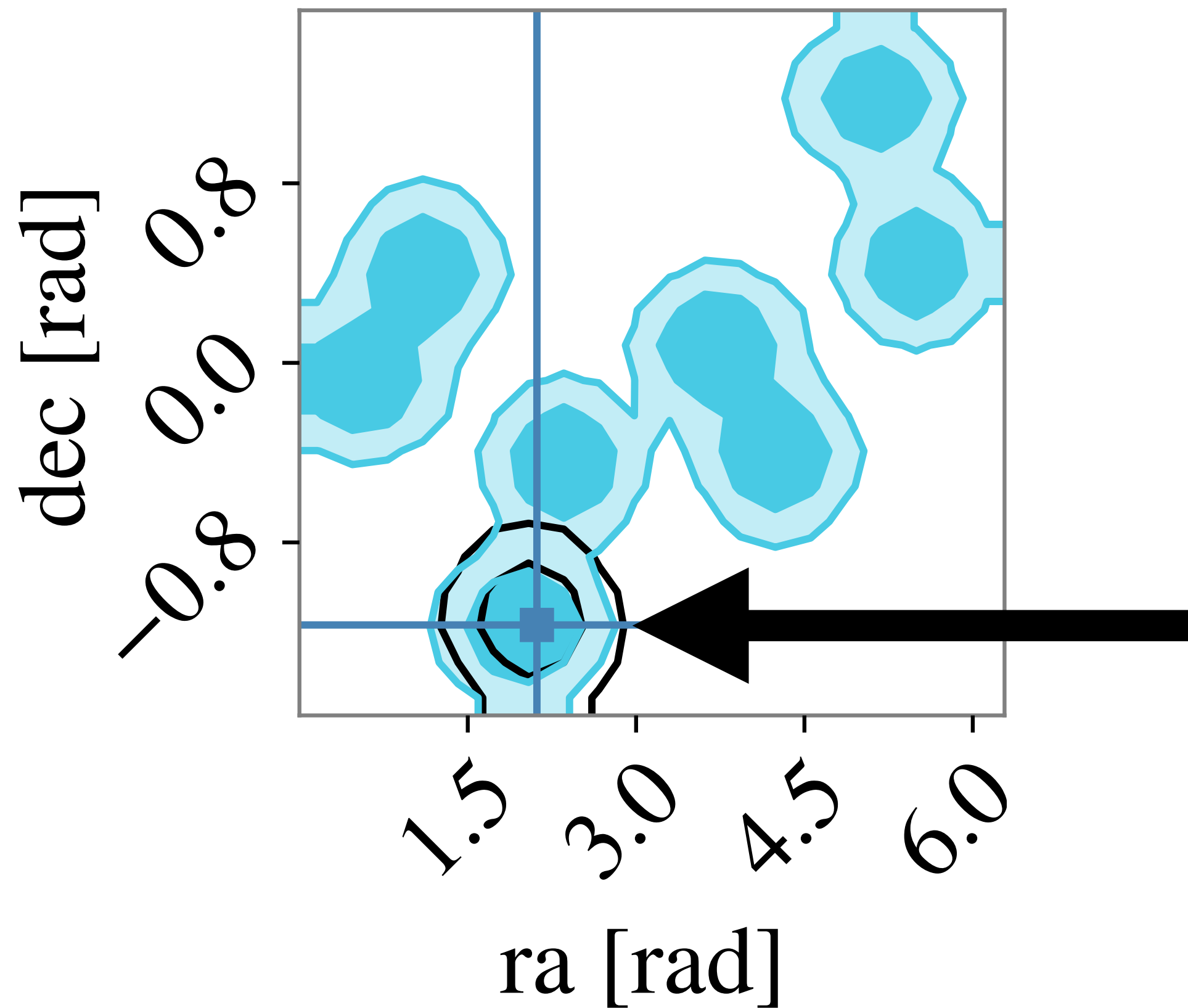
Dingo-IS VS Bilby



Sky modes

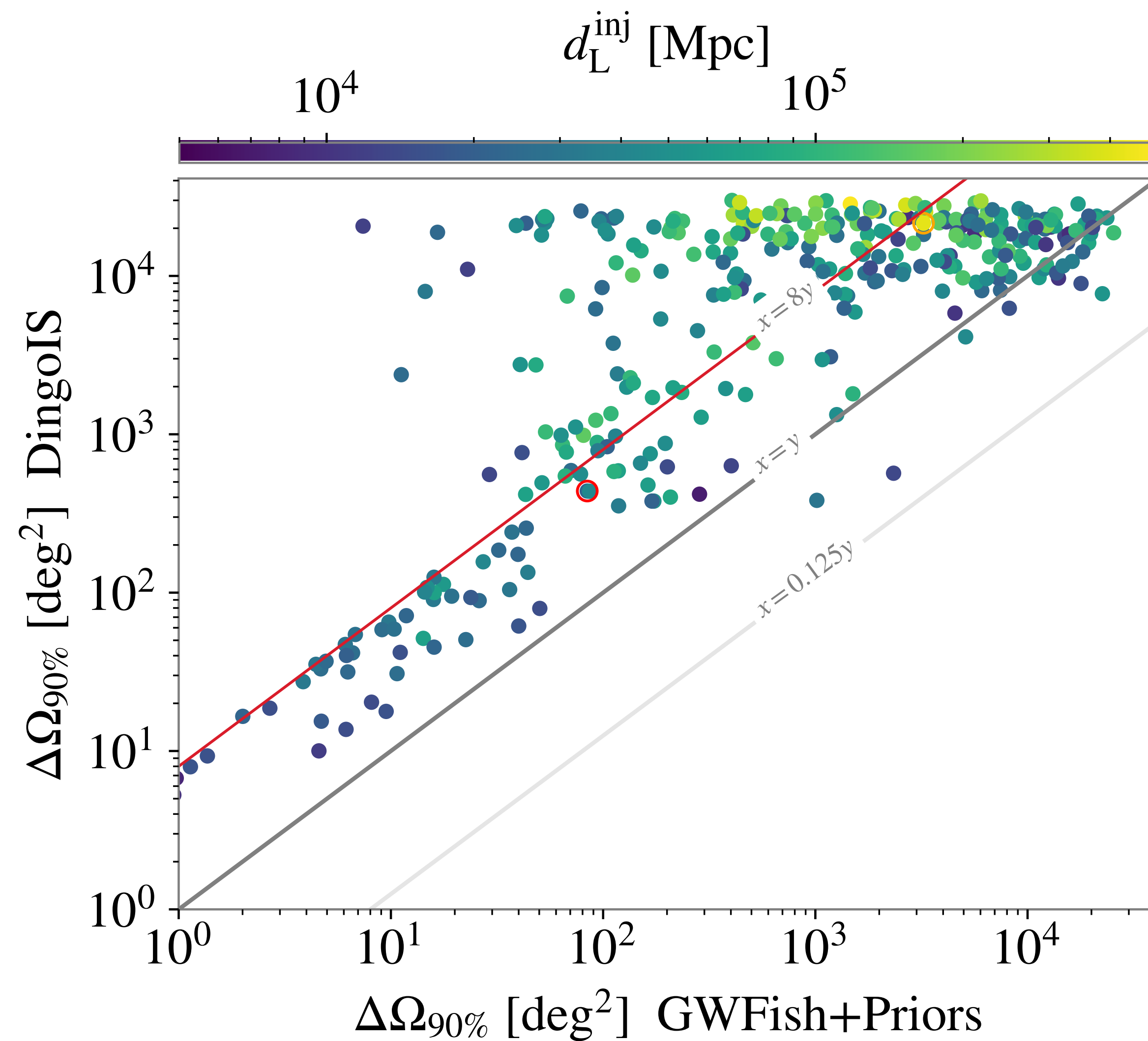


Sky modes



$\mathcal{L}(d | \theta)$ is approximated as a **multivariate Gaussian distribution** where the inverse of the covariance matrix is the **Fisher information matrix**

Sky localisation



Contributions

Fast and accurate parameter estimation for high-redshift sources

Look at multimodalities in sky localisation

This approach thrives where Fisher matrix approximation is less reliable

arXiv: [2504.21087](https://arxiv.org/abs/2504.21087)

Backup slides

Other priors

`'mass_ratio': UniformInComponentsMassRatio(minimum=0.125, maximum=1.0)`

`'dec': Cosine(minimum=- $\pi/2$ , maximum= $\pi/2$ )`

`'ra': Uniform(minimum=0, maximum= $2\pi$ )`

`'theta_jn': Sine(minimum=0.0, maximum= $\pi$ )`

`'psi': Uniform(minimum=0, maximum= $\pi$ )`

`'chi_1': AlignedSpin(a_prior=Uniform(minimum=0, maximum=0.9))`

`'chi_2': AlignedSpin(a_prior=Uniform(minimum=0, maximum=0.9))`

`'phase': Uniform(minimum=0, maximum= $2\pi$ )`

The loss function

Kullback–Leibler divergence

$$D_{\text{KL}}[p(\theta, s) \parallel q(\theta, s)] = \int ds \, p(s) \left[\int d\theta \, p(\theta | s) \log \frac{p(\theta | s)}{q(\theta | s)} \right]$$

Bayes' theorem

$$= \int ds \, p(s) \left[\int d\theta \, \frac{p(s | \theta)p(\theta)}{p(s)} \log \frac{p(\theta | s)}{q(\theta | s)} \right]$$

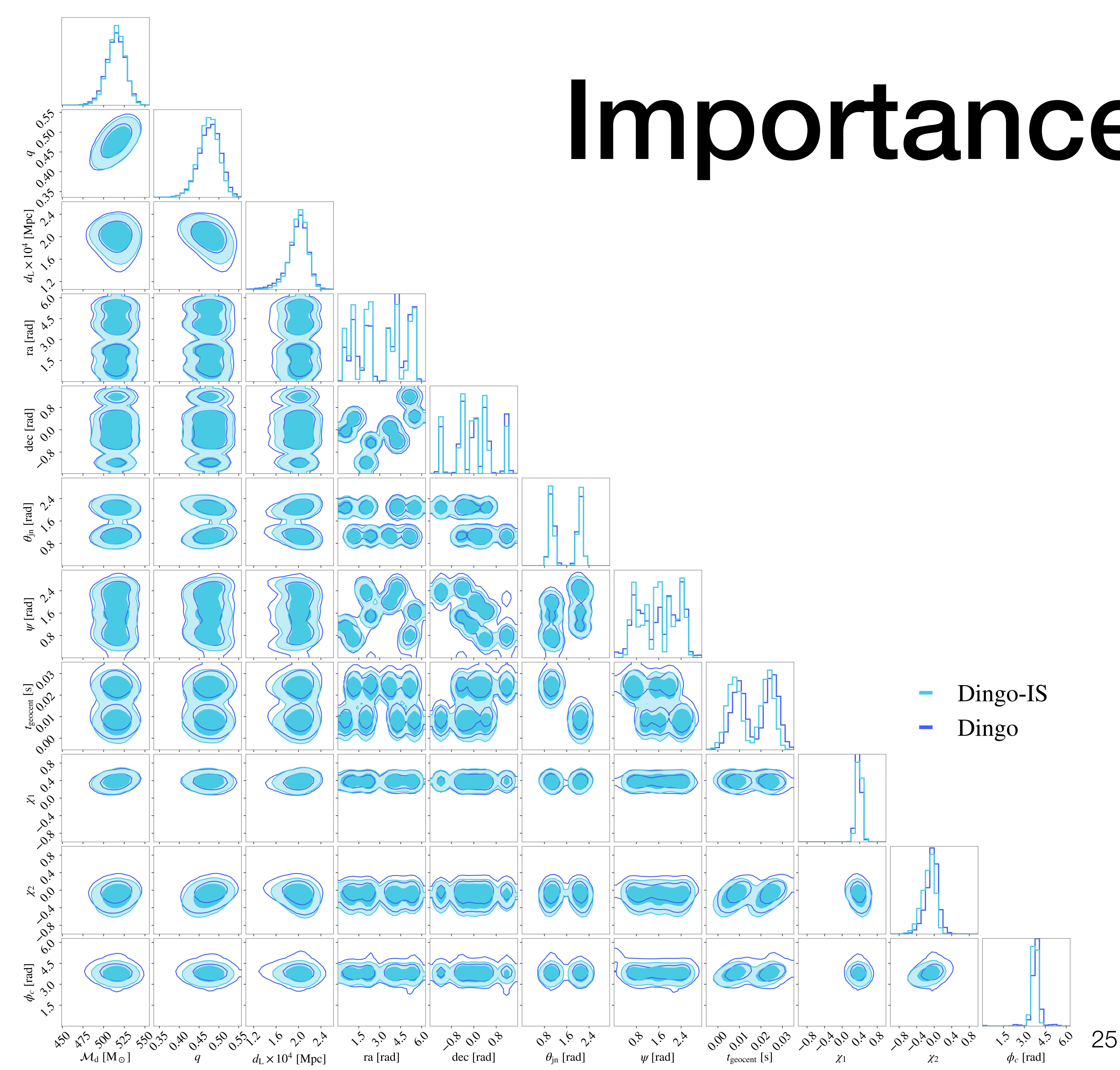
This term is not affected by the neural network

$$= \int ds \, p(s) \left[\int d\theta \, \frac{1}{p(s)} \left[p(\theta)p(s | \theta)\log p(\theta | s) - p(\theta)p(s | \theta)\log q(\theta | s) \right] \right]$$

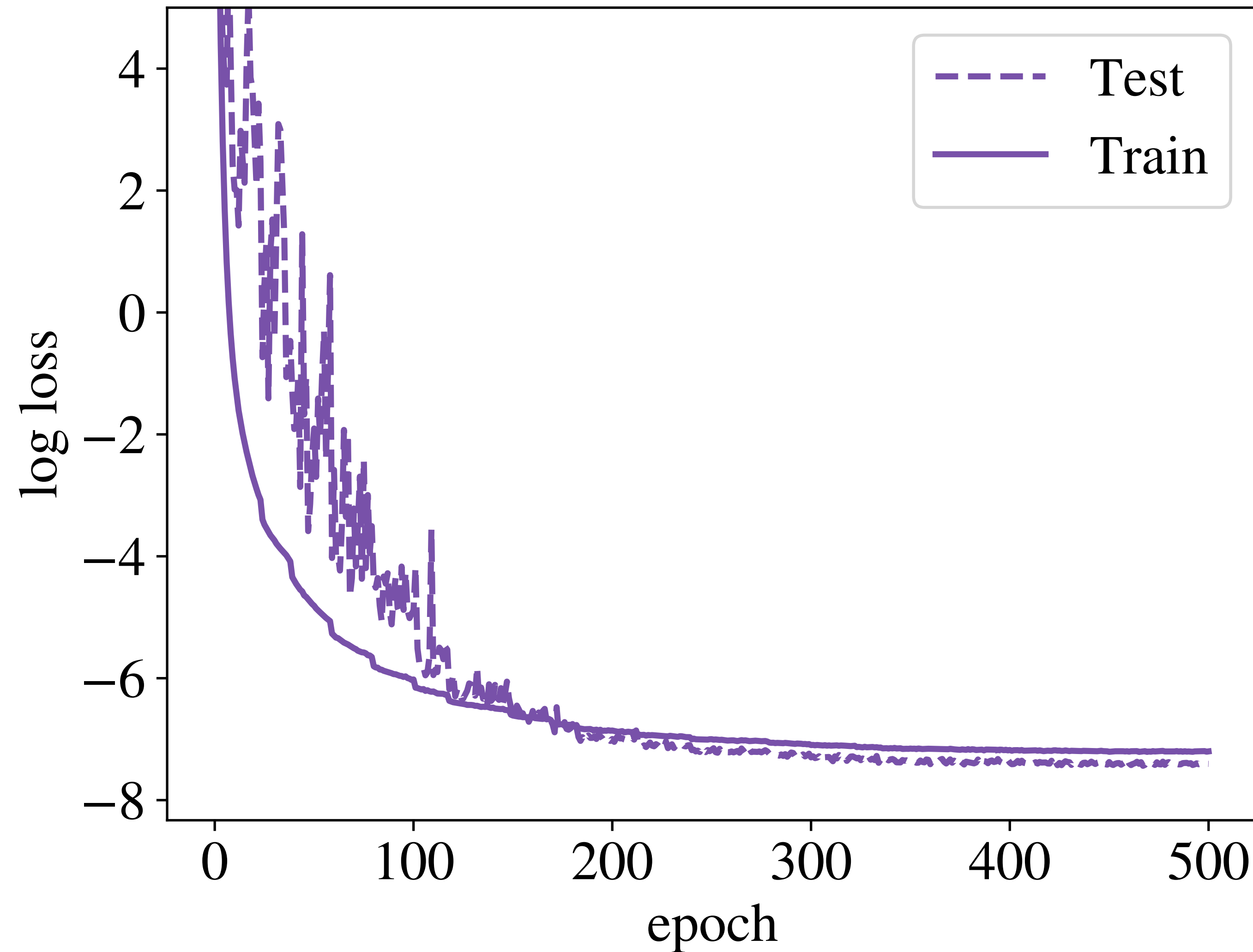
Neglecting constant values

$$\propto - \int ds \, d\theta \, p(\theta)p(s | \theta)\log q(\theta | s) = - \frac{1}{N_S} \sum_{i=1}^{N_S} \log q(\theta_i, s_i) = \mathbb{E}_{p(\theta)} \mathbb{E}_{p(s|\theta)} [-\log q(\theta | d)]$$

Importance sampling



Training



Antenna amplitude pattern functions

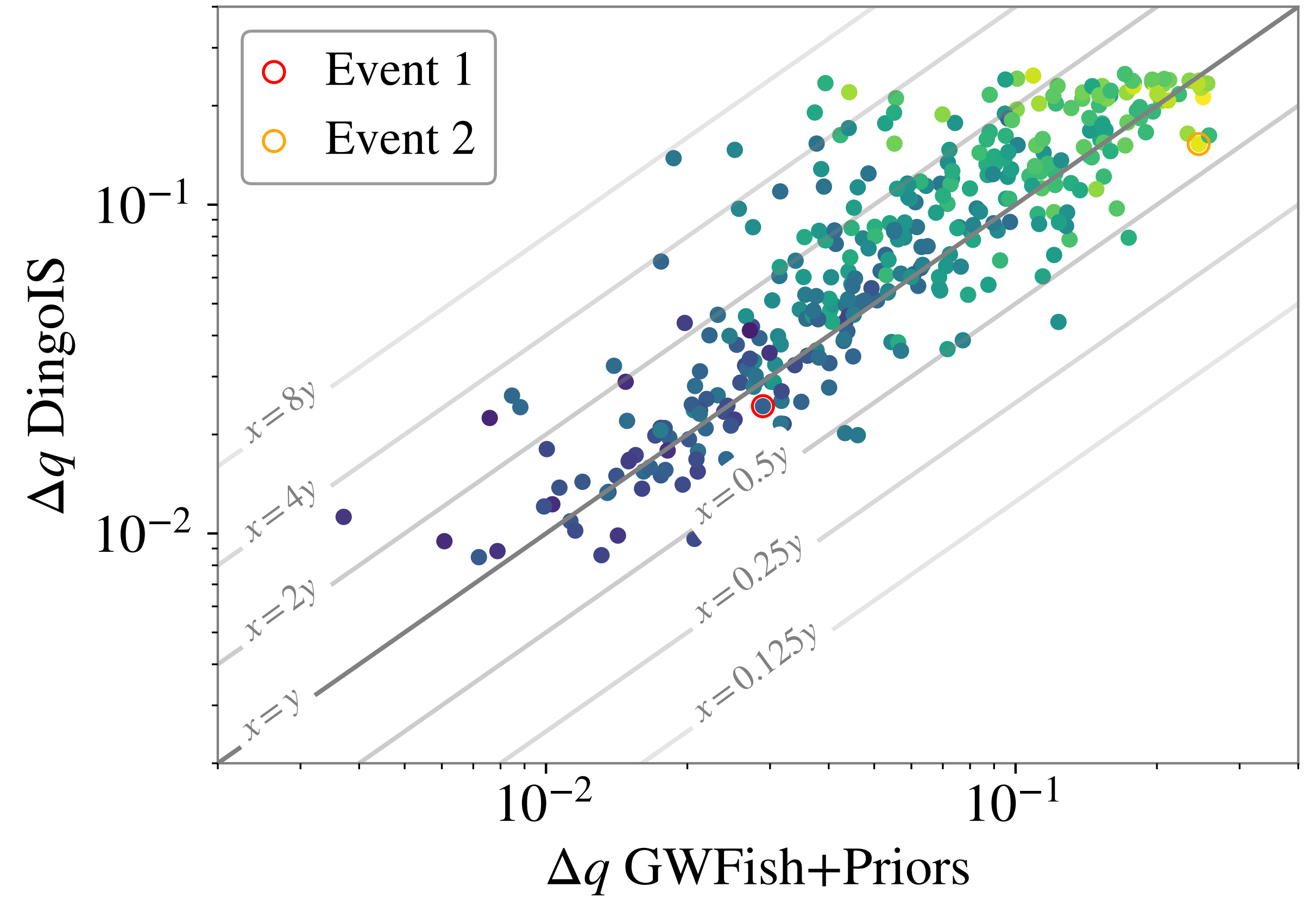
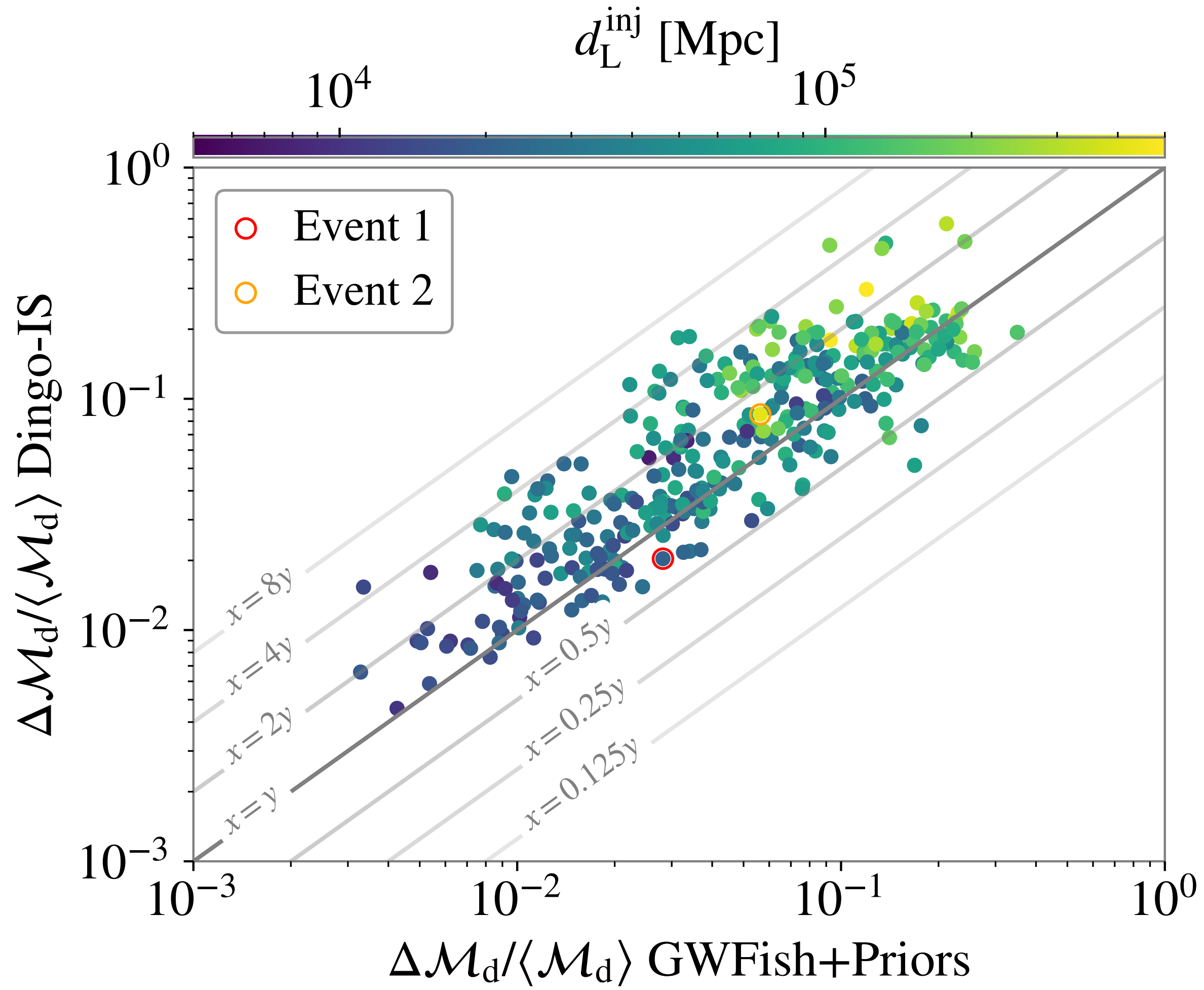
$$F_{+}^i(\beta, \lambda, \psi) = -\frac{\sqrt{3}}{4}[(1 + \cos^2 \beta)\sin 2\lambda \cos 2\psi + 2 \cos \beta \cos 2\lambda \sin 2\psi],$$

$$F_{\times}^i(\beta, \lambda, \psi) = +\frac{\sqrt{3}}{4}[(1 + \cos^2 \beta)\sin 2\lambda \sin 2\psi - 2 \cos \beta \cos 2\lambda \cos 2\psi],$$

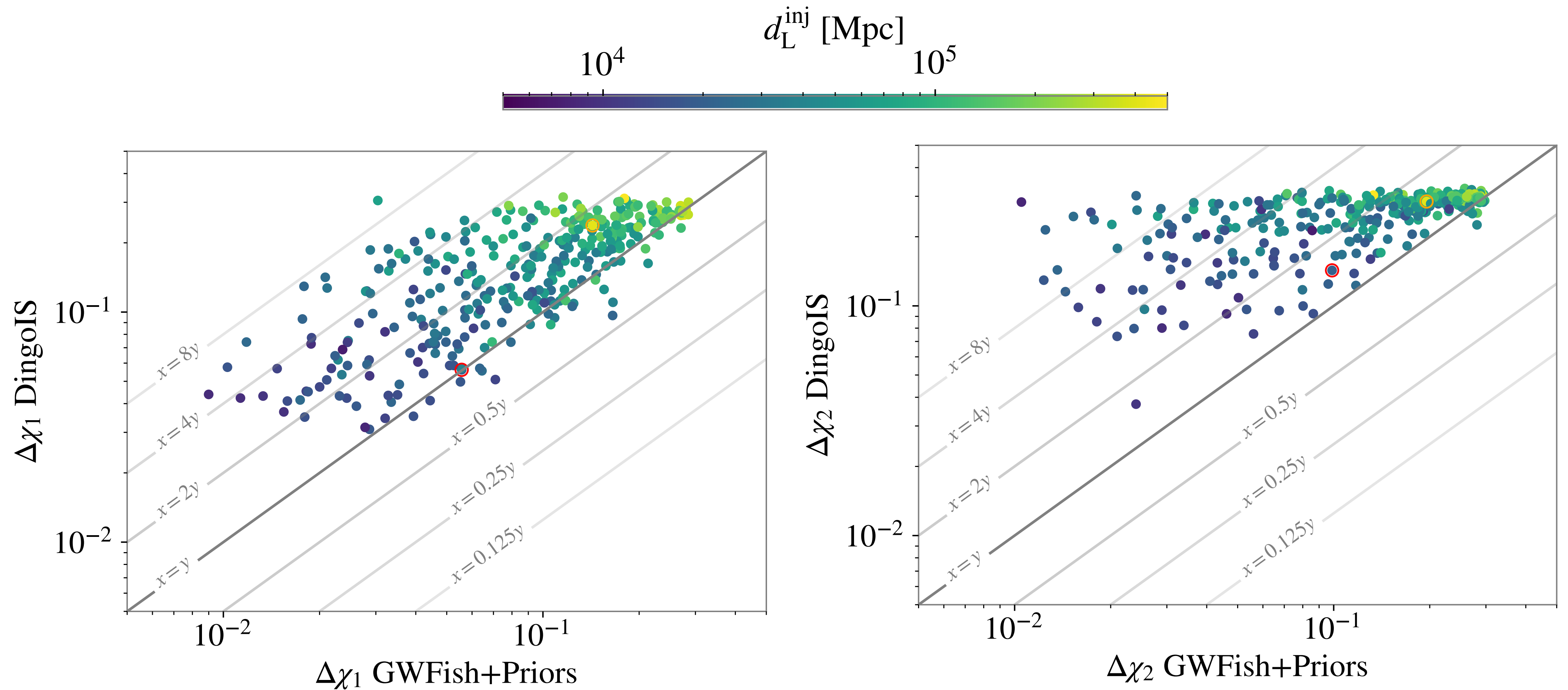
Antenna power pattern function

$$\sqrt{(F_+^i)^2 + (F_\times^i)^2} = \sqrt{\frac{3}{12} [(1 + \cos^2 \beta)^2 \sin^2 2\lambda + \cos^2 \beta \cos^2 2\lambda]}$$

Masses



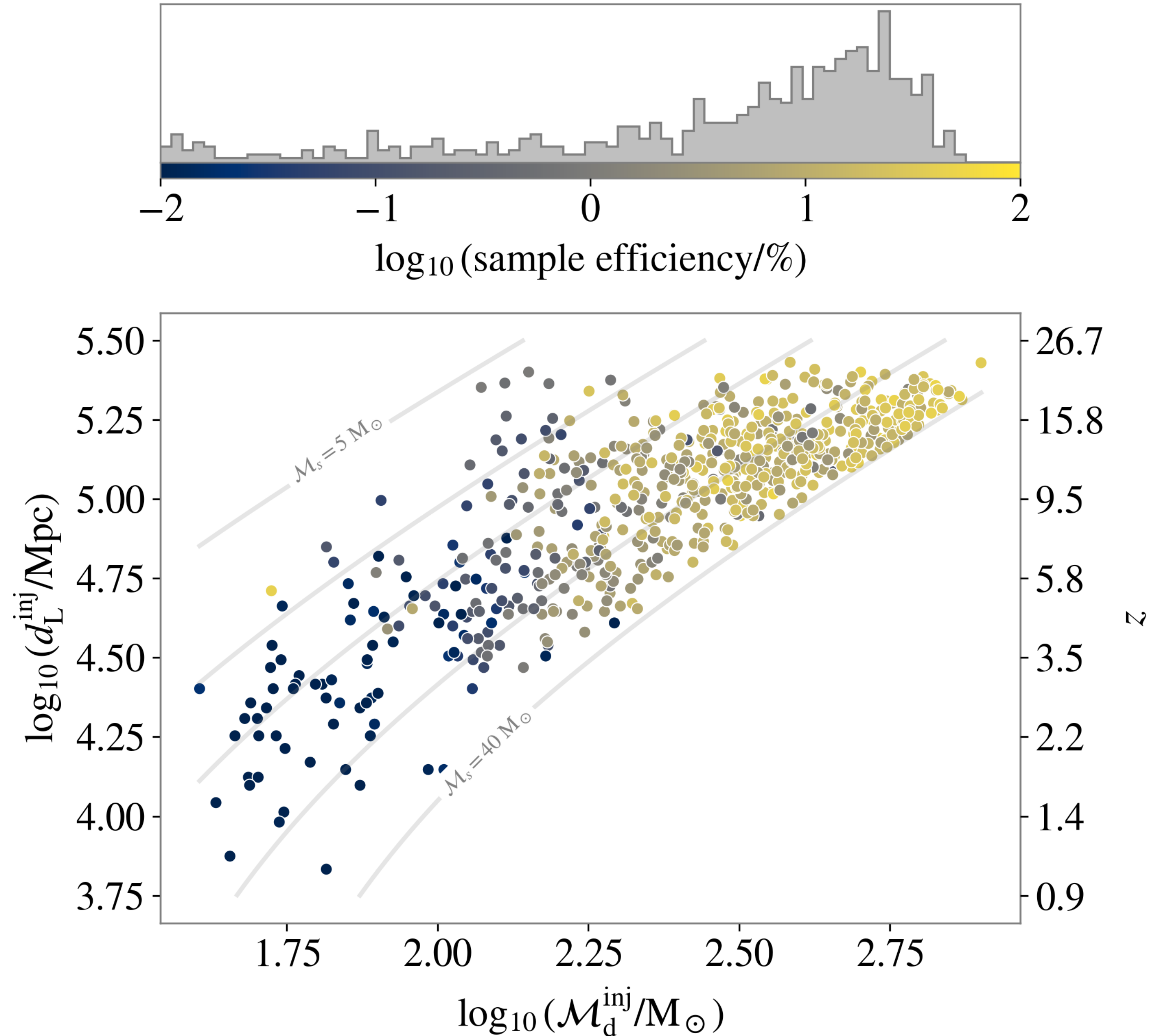
Aligned spins



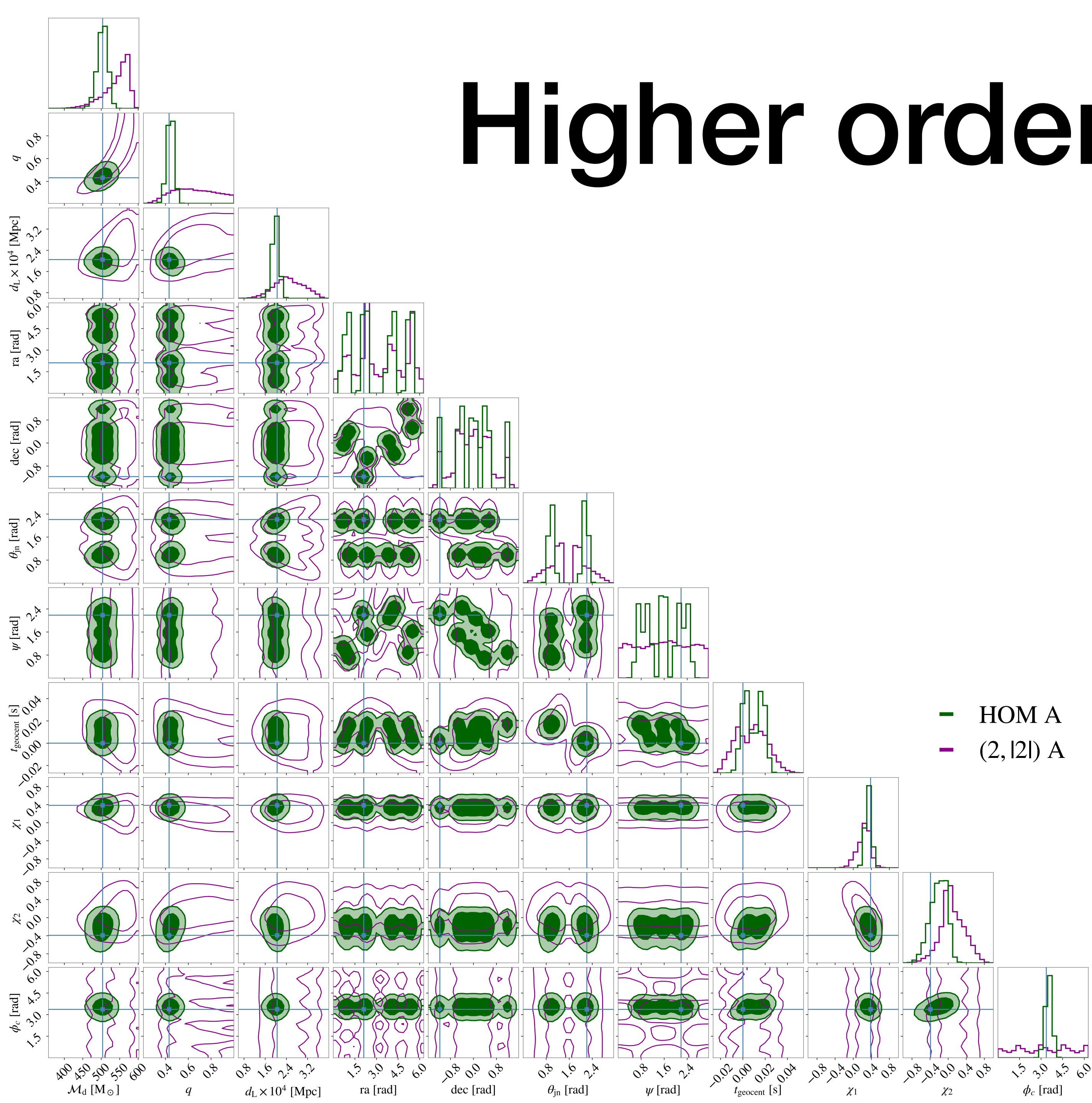
Astrophysical population

Binary black holes formed from Population III stars

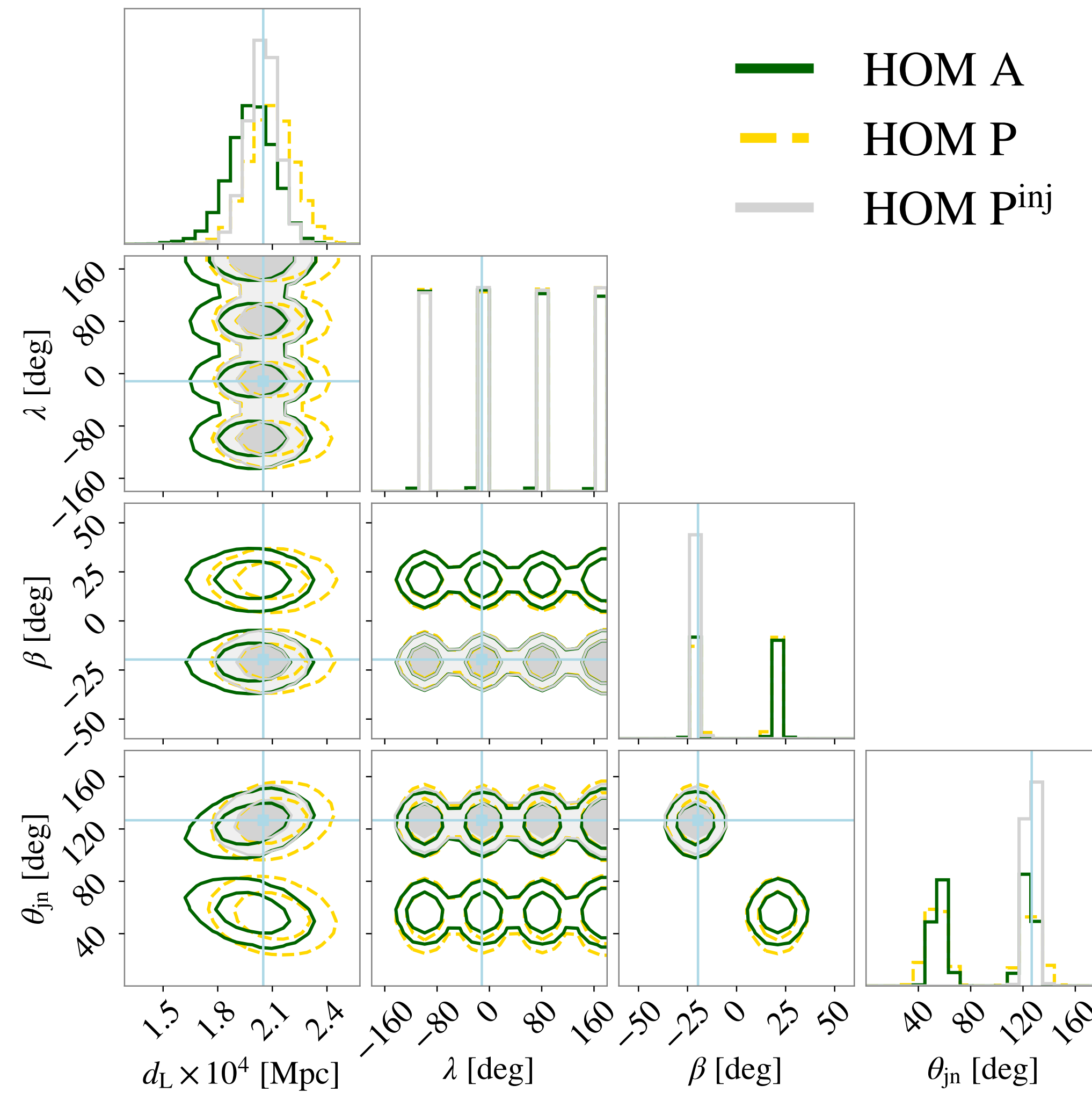
Ref. [Costa et al. 2023](#), [Santoliquido et al. 2023](#),
[Santoliquido et al. 2024](#), [Santoliquido et al. 2025](#)



Higher order modes

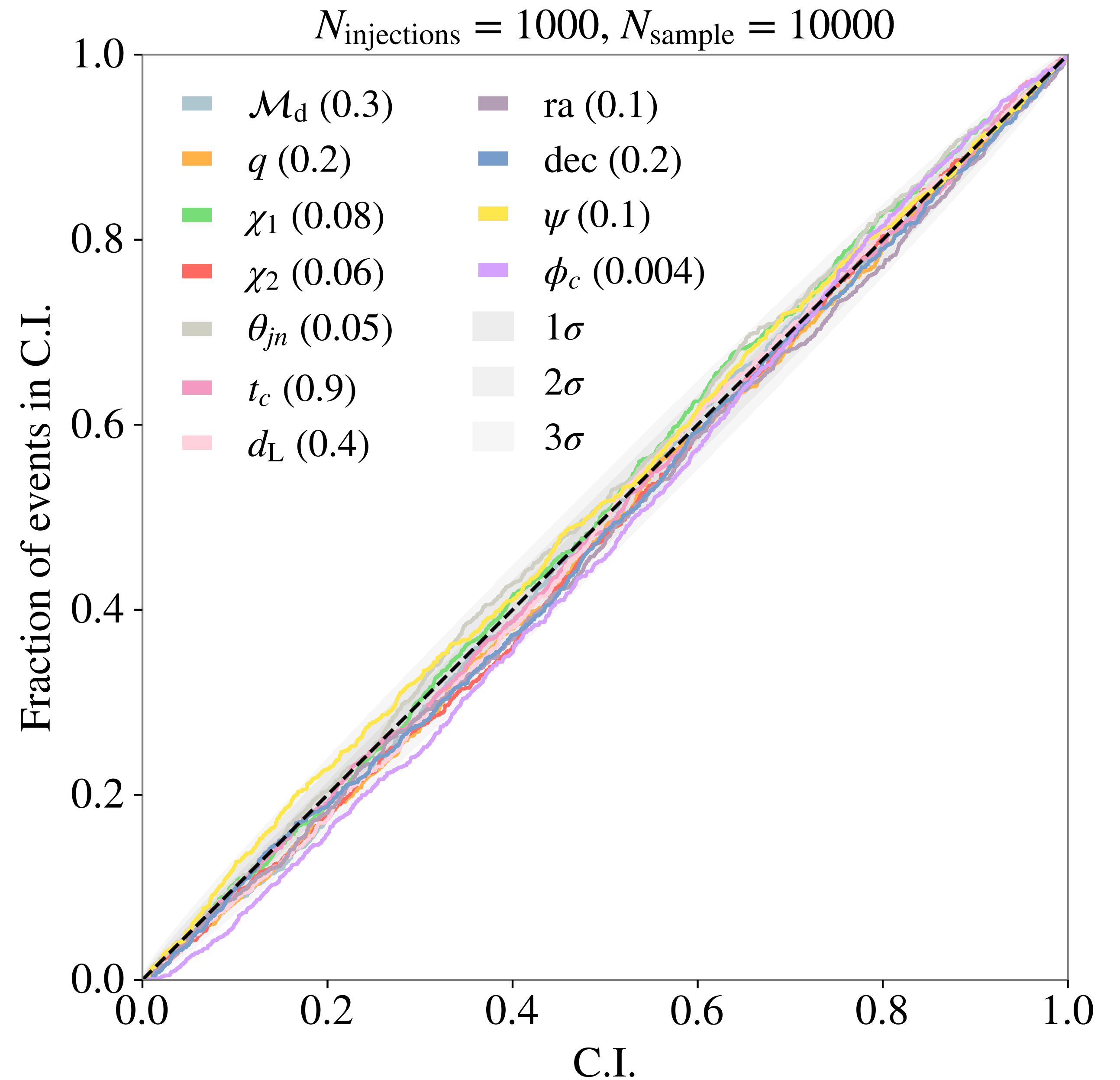
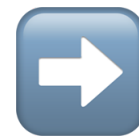


Precessing spins



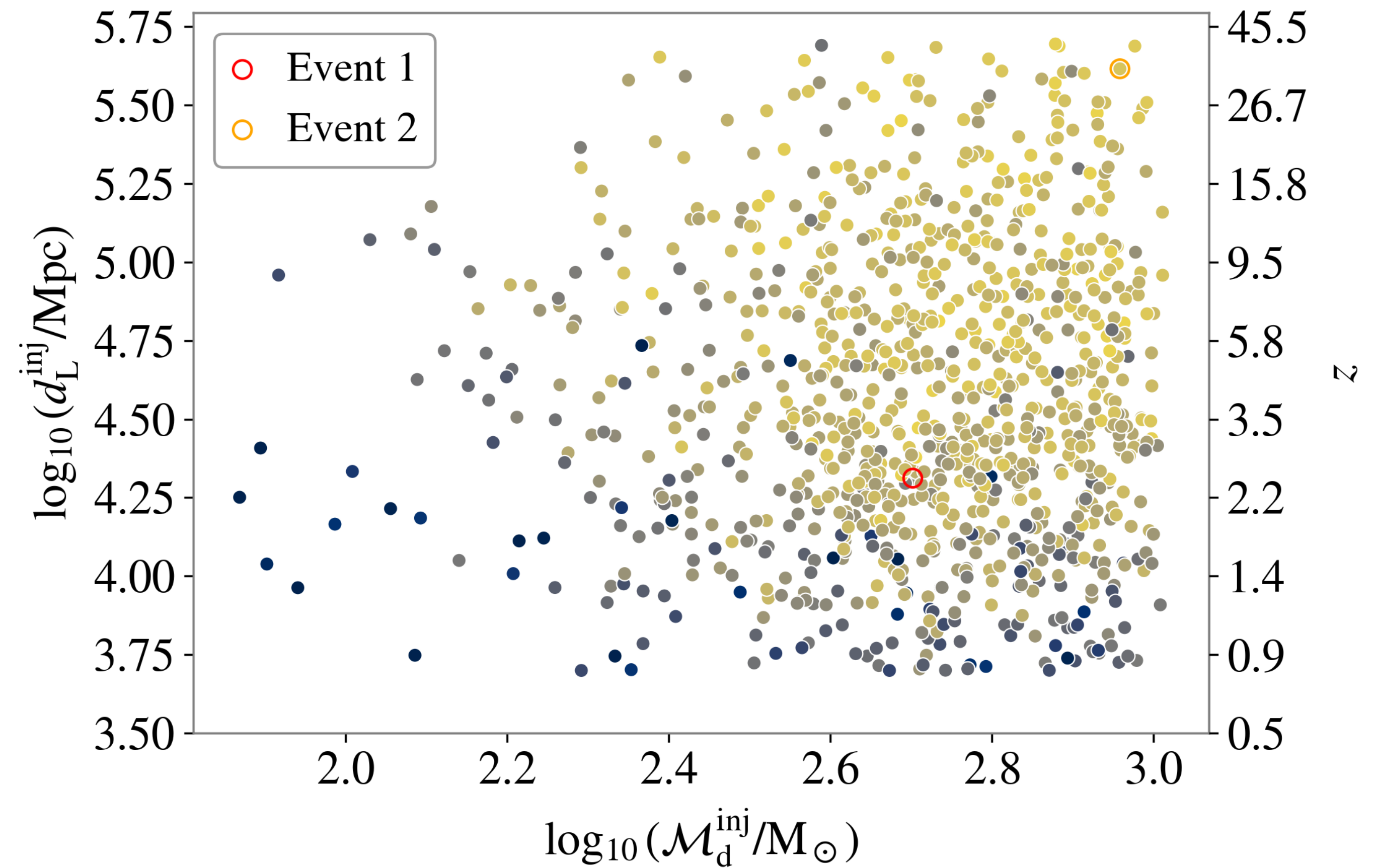
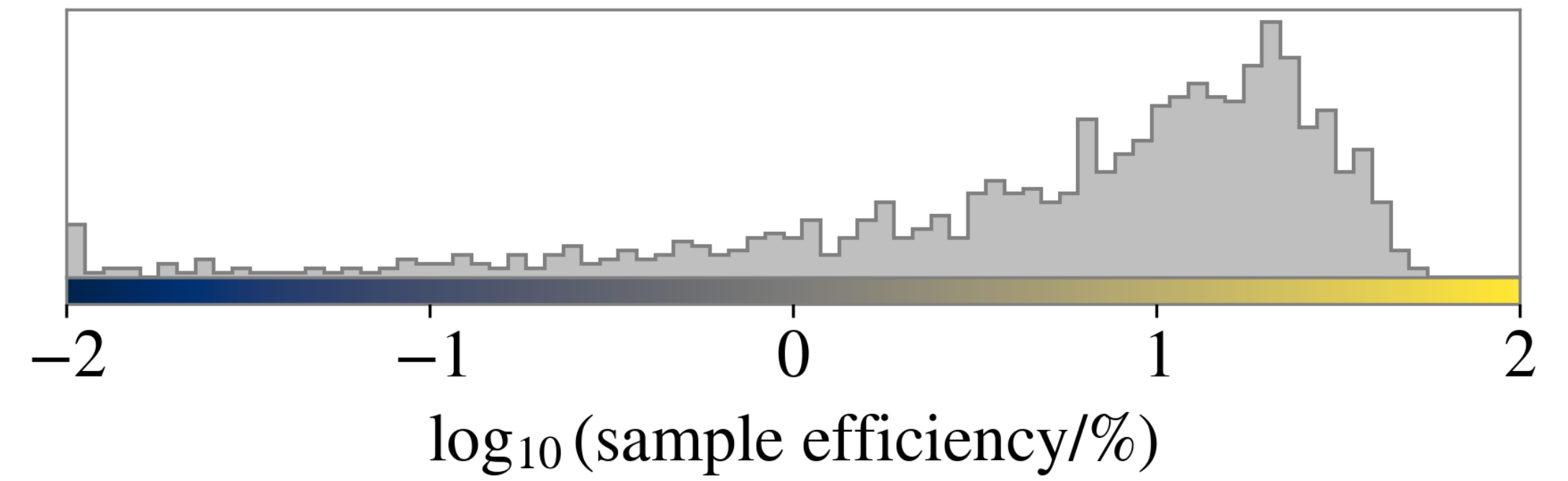
Validating results

1000 random injections
sampled from the priors



Wide redshift range accuracy ➡

85% of sources with sample efficiency > 1%



Sky modes in LISA

β_L, λ_L	injected sky location			
Sky mode	<i>Full</i>	<i>Frozen</i>	<i>Low-f</i>	<i>Frozen low-f</i>
<i>reflected:</i> $-\beta_L, \lambda_L$	<i>t</i> -dep.	degen.	<i>t</i> -dep.	degen.
<i>antipodal:</i> $-\beta_L, \lambda_L + \pi$	<i>f</i> -dep. + $\Delta\Phi_R$	<i>f</i> -dep.	$\Delta\Phi_R$	degen.
$\beta_L, \lambda_L + \pi/2$	<i>t-f</i> -dep.	<i>f</i> -dep.	<i>t</i> -dep.	degen.
$\beta_L, \lambda_L + \pi$	<i>t-f</i> -dep.	<i>f</i> -dep.	<i>t</i> -dep.	degen.
$\beta_L, \lambda_L - \pi/2$	<i>t-f</i> -dep.	<i>f</i> -dep.	<i>t</i> -dep.	degen.
$-\beta_L, \lambda_L + \pi/2$	<i>t-f</i> -dep.	<i>f</i> -dep.	<i>t</i> -dep.	degen.
$-\beta_L, \lambda_L - \pi/2$	<i>t-f</i> -dep.	<i>f</i> -dep.	<i>t</i> -dep.	degen.

Ref. [Vishal et al. 2020](#), [Sylvain et al. 2021](#), [Singh and Bulik 2021](#), [Singh and Bulik 2022](#)

Energy cost

- Hardware: **GPU** NVIDIA A100 80GB, **CPU** AMD EPYC 7513 32-core
- Dingo-IS:
 - Training: 970 kWh {1 GPU + 32 CPUs running for 6 days}
 - Inference: 45 kWh {8 CPUs, 1.7 minutes per source on average}
 - Total: **1015 kWh**
- Bilby:
 - Inference: **4000 kWh** {4 CPUs, ~5 hours per source on average}