

Machine learning based (likelihood-free) statistical inference

Tom Charnock
Guilhem Lavaux & Ben Wandelt

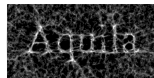
Alessandro Manzotti
Justin Alsing

Institut d'Astrophysique de Paris

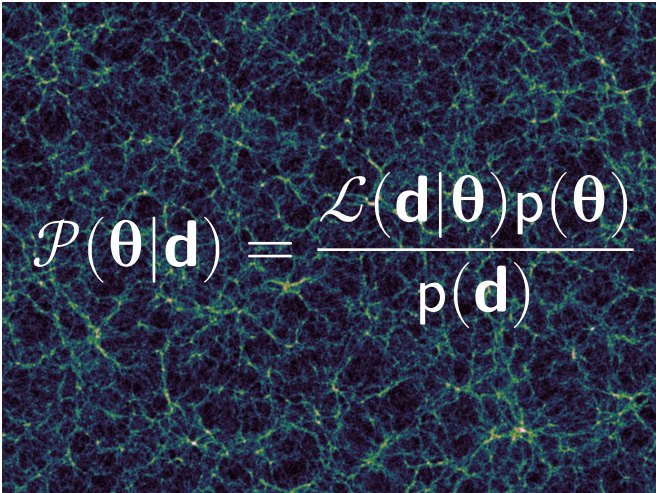
Physical Review D **97**, 083004

DOI [10.5281/zenodo.1175196](https://doi.org/10.5281/zenodo.1175196)

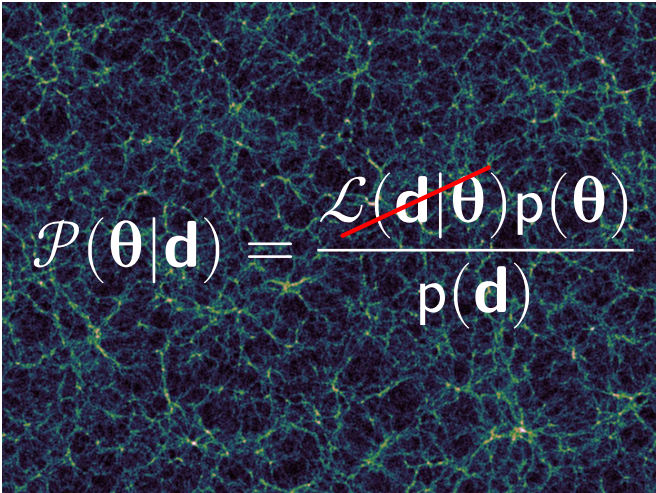
[github:information_maximiser](#)



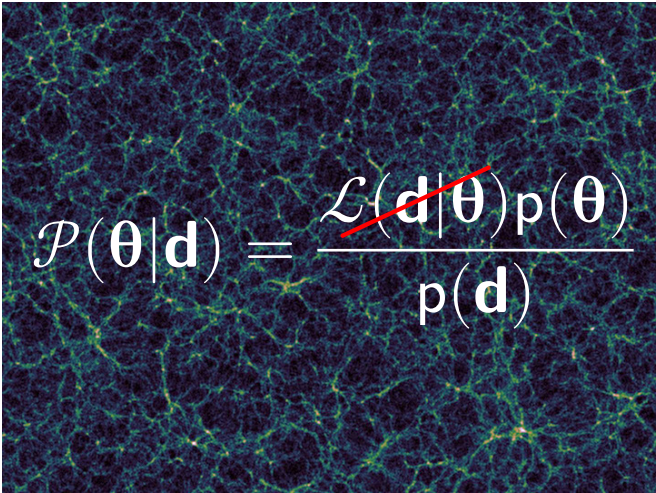
Statistical inference

A visualization of the cosmic web, showing a complex network of dark matter filaments and galaxy clusters. The filaments are represented by thin, interconnected lines of blue and green, while the clusters are shown as bright yellow and orange points. The background is a deep black, suggesting the vastness of space.
$$\mathcal{P}(\boldsymbol{\theta}|\mathbf{d}) = \frac{\mathcal{L}(\mathbf{d}|\boldsymbol{\theta})p(\boldsymbol{\theta})}{p(\mathbf{d})}$$

Approximate Bayesian computation

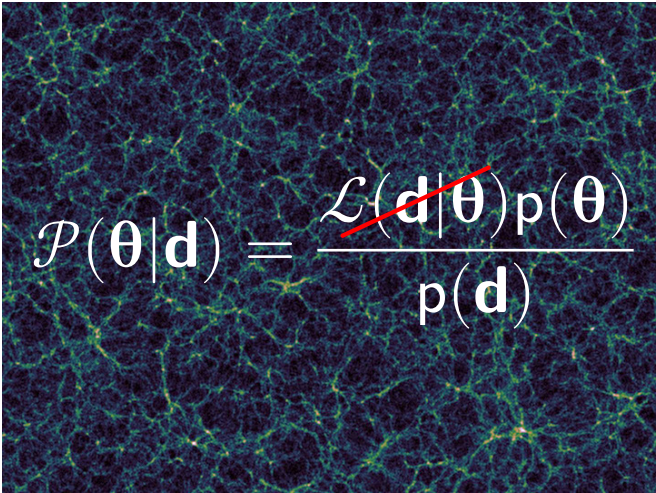
A background image of a cosmic web, showing a complex network of dark blue and green filaments with bright yellow and orange nodes, representing galaxy clusters and filaments in the universe.
$$\mathcal{P}(\boldsymbol{\theta}|\mathbf{d}) = \frac{\cancel{\mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}p(\boldsymbol{\theta})}{p(\mathbf{d})}$$

Approximate Bayesian computation

A visualization of the cosmic web, showing a complex network of dark matter filaments and galaxy clusters. The filaments are represented by thin, interconnected lines of blue and green, while the clusters are shown as denser, yellowish-orange regions. The background is a deep black, suggesting the vastness of space.
$$\mathcal{P}(\boldsymbol{\theta}|\mathbf{d}) = \frac{\cancel{\mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}p(\boldsymbol{\theta})}{p(\mathbf{d})}$$

Simulate data at different $\boldsymbol{\theta}$ drawn from $p(\boldsymbol{\theta})$

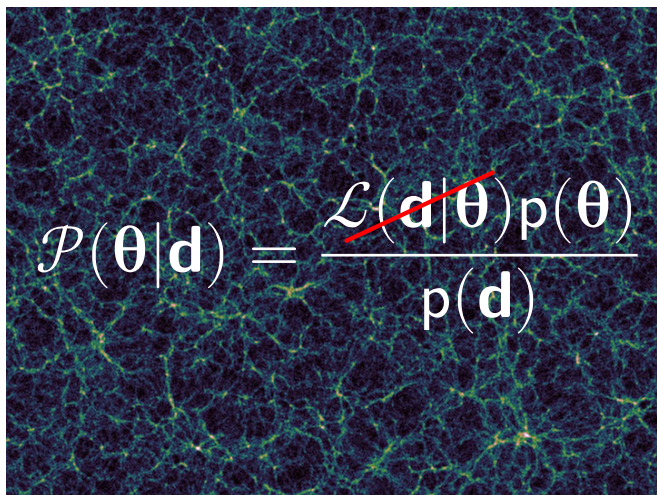
Approximate Bayesian computation

A background image showing a complex, interconnected network of green and yellow lines on a dark blue field, resembling a cosmic web or a neural network.
$$\mathcal{P}(\boldsymbol{\theta}|\mathbf{d}) = \frac{\cancel{\mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}p(\boldsymbol{\theta})}{p(\mathbf{d})}$$

Simulate data at different $\boldsymbol{\theta}$ drawn from $p(\boldsymbol{\theta})$

Accept or reject $\boldsymbol{\theta}$ using similarity of simulation and real data

Likelihood-free rejection sampling

A background image of a Cosmic Microwave Background (CMB) fluctuation map, showing a complex network of blue and yellow filaments against a dark background. Overlaid on this is the mathematical formula for likelihood-free rejection sampling.
$$\mathcal{P}(\boldsymbol{\theta}|\mathbf{d}) = \frac{\cancel{\mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}p(\boldsymbol{\theta})}{p(\mathbf{d})}$$

Simulate data at different $\boldsymbol{\theta}$ drawn from $p(\boldsymbol{\theta})$

Accept or reject $\boldsymbol{\theta}$ using similarity of simulation and real data

Why we need compressed summary statistics

Reduce dimensionality of data

Unlikely for simulation to hit real data in high dimensions

Use summary statistics such as power spectrum, C_ℓ , ect.

Why we need compressed summary statistics

Reduce dimensionality of data

Unlikely for simulation to hit real data in high dimensions

Use summary statistics such as power spectrum, C_ℓ , etc.

Inadequately compressed summary statistics

Upcoming surveys will observe huge amounts of raw data

Even the summary statistics will be $O(10^4)$ (LSST, EUCLID, etc.)

Huge amount of computing power needed

Still nearly impossible to do ABC

Compression for general likelihoods

Optimal compression for general likelihood functions

Fisher information

- Alsing & Wandelt, 2017

Amount of information data $\mathbf{d}$ contains about parameters $\boldsymbol{\theta}$

$$\begin{aligned}\mathbf{F}_{\alpha\beta} &= - \left\langle \frac{\partial^2 \ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}{\partial \theta_\alpha \partial \theta_\beta} \right\rangle \bigg|_{\boldsymbol{\theta} = \boldsymbol{\theta}^{\text{fid}}} \\ &= \left\langle \frac{\partial \ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}{\partial \theta_\alpha} \frac{\partial \ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}{\partial \theta_\beta} \right\rangle \bigg|_{\boldsymbol{\theta} = \boldsymbol{\theta}^{\text{fid}}}\end{aligned}$$

Optimal compression for general likelihood functions

Fisher information

- Alsing & Wandelt, 2017

Amount of information data $\mathbf{d}$ contains about parameters $\boldsymbol{\theta}$

$$\begin{aligned}\mathbf{F}_{\alpha\beta} &= - \left\langle \frac{\partial^2 \ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}{\partial \theta_\alpha \partial \theta_\beta} \right\rangle \bigg|_{\boldsymbol{\theta} = \boldsymbol{\theta}^{\text{fid}}} \\ &= \left\langle \frac{\partial \ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}{\partial \theta_\alpha} \frac{\partial \ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}{\partial \theta_\beta} \right\rangle \bigg|_{\boldsymbol{\theta} = \boldsymbol{\theta}^{\text{fid}}}\end{aligned}$$

Information inequality

Cramér-Rao bound for unbiased estimators $\langle \mathcal{T}_\alpha \rangle = \theta_\alpha$

$$\langle (\mathcal{T}_\alpha - \langle \mathcal{T}_\alpha \rangle) (\mathcal{T}_\beta - \langle \mathcal{T}_\beta \rangle) \rangle \geq \mathbf{F}_{\alpha\beta}^{-1}$$

Lower bound on the variances of sufficient statistics, $\mathcal{T}_\alpha$.

$$\langle (\mathcal{T}_\alpha - \langle \mathcal{T}_\alpha \rangle) (\mathcal{T}_\beta - \langle \mathcal{T}_\beta \rangle) \rangle \geq \left\langle \frac{\partial \mathcal{T}_\mu}{\partial \theta_\alpha} \right\rangle \mathbf{F}_{\mu\nu}^{-1} \left\langle \frac{\partial \mathcal{T}_\nu}{\partial \theta_\beta} \right\rangle$$

Optimal compression for geneal likelihood functions

Sufficient statistics

Expand the log-likelihood about some fiducial parameters, θ^{fid}

$$\ln \mathcal{L}(\mathbf{d}|\theta) = \ln \mathcal{L}(\mathbf{d}|\theta^{\text{fid}}) + \delta\theta_{\alpha}^{\text{T}} \frac{\partial \ln \mathcal{L}(\mathbf{d}|\theta^{\text{fid}})}{\partial \theta_{\alpha}} + \frac{1}{2} \delta\theta_{\alpha}^{\text{T}} \frac{\partial^2 \ln \mathcal{L}(\mathbf{d}|\theta^{\text{fid}})}{\partial \theta_{\alpha} \partial \theta_{\beta}} \delta\theta_{\beta} + \dots$$

To linear order the data only couples through the score function

$$\mathcal{S}_{\alpha} = \frac{\partial \ln \mathcal{L}(\mathbf{d}|\theta)}{\partial \theta_{\alpha}}$$

This is a set of sufficient statistics for the linearised log-likelihood

$$\mathcal{T}_{\alpha} = \mathcal{S}_{\alpha}$$

Optimal compression for general likelihood functions

Fisher information

$$\begin{aligned}\langle (\mathcal{T}_\alpha - \langle \mathcal{T}_\alpha \rangle) (\mathcal{T}_\beta - \langle \mathcal{T}_\beta \rangle) \rangle &= \left\langle \frac{\partial \ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}{\partial \theta_\alpha} \frac{\partial \ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}{\partial \theta_\beta} \right\rangle \bigg|_{\boldsymbol{\theta} = \boldsymbol{\theta}^{\text{fid}}} \\ &= \mathbf{F}_{\alpha\beta}\end{aligned}$$

$$\begin{aligned}\left\langle \frac{\partial \mathcal{T}_\alpha}{\partial \theta_\beta} \right\rangle &= \left\langle \frac{\partial^2 \ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}{\partial \theta_\alpha \partial \theta_\beta} \right\rangle \bigg|_{\boldsymbol{\theta} = \boldsymbol{\theta}^{\text{fid}}} \\ &= -\mathbf{F}_{\alpha\beta}\end{aligned}$$

Optimal compression for general likelihood functions

Fisher information

$$\begin{aligned}\langle (\mathcal{T}_\alpha - \langle \mathcal{T}_\alpha \rangle) (\mathcal{T}_\beta - \langle \mathcal{T}_\beta \rangle) \rangle &= \left\langle \frac{\partial \ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}{\partial \theta_\alpha} \frac{\partial \ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}{\partial \theta_\beta} \right\rangle \bigg|_{\boldsymbol{\theta} = \boldsymbol{\theta}^{\text{fid}}} \\ &= \mathbf{F}_{\alpha\beta}\end{aligned}$$

$$\begin{aligned}\left\langle \frac{\partial \mathcal{T}_\alpha}{\partial \theta_\beta} \right\rangle &= \left\langle \frac{\partial^2 \ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})}{\partial \theta_\alpha \partial \theta_\beta} \right\rangle \bigg|_{\boldsymbol{\theta} = \boldsymbol{\theta}^{\text{fid}}} \\ &= -\mathbf{F}_{\alpha\beta}\end{aligned}$$

Information inequality saturated

$$\langle (\mathcal{T}_\alpha - \langle \mathcal{T}_\alpha \rangle) (\mathcal{T}_\beta - \langle \mathcal{T}_\beta \rangle) \rangle \geq \left\langle \frac{\partial \mathcal{T}_\mu}{\partial \theta_\alpha} \right\rangle \mathbf{F}_{\mu\nu}^{-1} \left\langle \frac{\partial \mathcal{T}_\nu}{\partial \theta_\beta} \right\rangle$$

Massively optimised parameter estimation and data (MOPED) compression

- Heavens et. al, 2000

Assume a Gaussian likelihood ($\mu(\boldsymbol{\theta}) \equiv \boldsymbol{\mu}$ and $\partial \mathbf{C} / \partial \theta_{\alpha} = 0$)

$$-2 \ln \mathcal{L}(\mathbf{d} | \boldsymbol{\theta}) = (\mathbf{d} - \boldsymbol{\mu})^T \mathbf{C}^{-1} (\mathbf{d} - \boldsymbol{\mu}) + \ln |2\pi \mathbf{C}|$$

Massively optimised parameter estimation and data (MOPED) compression

- Heavens et. al, 2000

Assume a Gaussian likelihood ($\mu(\boldsymbol{\theta}) \equiv \boldsymbol{\mu}$ and $\partial \mathbf{C} / \partial \theta_{\alpha} = 0$)

$$-2 \ln \mathcal{L}(\mathbf{d} | \boldsymbol{\theta}) = (\mathbf{d} - \boldsymbol{\mu})^T \mathbf{C}^{-1} (\mathbf{d} - \boldsymbol{\mu}) + \ln |2\pi \mathbf{C}|$$

Compression function

$$\begin{aligned} \mathcal{T}_{\alpha} &= \left. \frac{\partial \ln \mathcal{L}(\mathbf{d} | \boldsymbol{\theta})}{\partial \theta_{\alpha}} \right|_{\boldsymbol{\theta} = \boldsymbol{\theta}^{\text{fid}}} \\ &= \boldsymbol{\mu}_{,\alpha}^T \mathbf{C}^{-1} \mathbf{d} \Big|_{\boldsymbol{\theta} = \boldsymbol{\theta}^{\text{fid}}} \quad \text{using} \quad \frac{\partial \boldsymbol{\mu}}{\partial \theta_{\alpha}} \equiv \boldsymbol{\mu}_{,\alpha} \end{aligned}$$

Massively optimised parameter estimation and data (MOPED) compression

- Heavens et. al, 2000

Assume a Gaussian likelihood ($\mu(\boldsymbol{\theta}) \equiv \boldsymbol{\mu}$ and $\partial \mathbf{C} / \partial \theta_{\alpha} = 0$)

$$-2 \ln \mathcal{L}(\mathbf{d} | \boldsymbol{\theta}) = (\mathbf{d} - \boldsymbol{\mu})^T \mathbf{C}^{-1} (\mathbf{d} - \boldsymbol{\mu}) + \ln |2\pi \mathbf{C}|$$

Compression function

$$\begin{aligned} \mathcal{T}_{\alpha} &= \left. \frac{\partial \ln \mathcal{L}(\mathbf{d} | \boldsymbol{\theta})}{\partial \theta_{\alpha}} \right|_{\boldsymbol{\theta} = \boldsymbol{\theta}^{\text{fid}}} \\ &= \boldsymbol{\mu}_{,\alpha}^T \mathbf{C}^{-1} \mathbf{d} \Big|_{\boldsymbol{\theta} = \boldsymbol{\theta}^{\text{fid}}} \quad \text{using} \quad \frac{\partial \boldsymbol{\mu}}{\partial \theta_{\alpha}} \equiv \boldsymbol{\mu}_{,\alpha} \end{aligned}$$

Fisher information

$$\begin{aligned} \mathbf{F}_{\alpha\beta} &= \left\langle \frac{\partial \ln \mathcal{L}(\mathbf{d} | \boldsymbol{\theta})}{\partial \theta_{\alpha}} \frac{\partial \ln \mathcal{L}(\mathbf{d} | \boldsymbol{\theta})}{\partial \theta_{\beta}} \right\rangle \Big|_{\boldsymbol{\theta} = \boldsymbol{\theta}^{\text{fid}}} \\ &= \text{Tr} \left[\boldsymbol{\mu}_{,\alpha}^T \mathbf{C}^{-1} \boldsymbol{\mu}_{,\beta} \right] \end{aligned}$$

How can we use this compression?

Calculated summary statistics from raw data

Calculate summary statistics (such as power spectrum) from raw data

Use MOPED (or similar) on summary statistics (assuming Gaussianity)

How can we use this compression?

Calculated summary statistics from raw data

Calculate summary statistics (such as power spectrum) from raw data

Use MOPED (or similar) on summary statistics (assuming Gaussianity)

Likelihood-free compression

Forward simulate to get realisations of summary statistics

Use MOPED (or similar) on summary statistics (assuming Gaussianity)

How can we use this compression?

Calculated summary statistics from raw data

Calculate summary statistics (such as power spectrum) from raw data

Use MOPED (or similar) on summary statistics (assuming Gaussianity)

Likelihood-free compression

Forward simulate to get realisations of summary statistics

Use MOPED (or similar) on summary statistics (assuming Gaussianity)

Don't have to approximate your likelihoods!

Non-linear likelihood-free compression

Non-linear likelihood-free compression

Start with unknown likelihood $\ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})$

Data $\mathbf{d}$ can be generated using a model $\mathcal{M}(\boldsymbol{\theta})$ with parameters $\boldsymbol{\theta}$

Non-linear likelihood-free compression

Start with unknown likelihood $\ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})$

Data $\mathbf{d}$ can be generated using a model $\mathcal{M}(\boldsymbol{\theta})$ with parameters $\boldsymbol{\theta}$

Transform data

Find a function, $\ell : \mathbf{d} \rightarrow \mathbf{x}$ which Gaussianizes the summaries, i.e.

$$-2\ln \mathcal{L}(\mathbf{x}|\boldsymbol{\theta}) = (\mathbf{x} - \boldsymbol{\mu}_\ell)^T \mathbf{C}_\ell^{-1} (\mathbf{x} - \boldsymbol{\mu}_\ell)$$

Non-linear likelihood-free compression

Start with unknown likelihood $\ln \mathcal{L}(\mathbf{d}|\boldsymbol{\theta})$

Data $\mathbf{d}$ can be generated using a model $\mathcal{M}(\boldsymbol{\theta})$ with parameters $\boldsymbol{\theta}$

Transform data

Find a function, $\ell : \mathbf{d} \rightarrow \mathbf{x}$ which Gaussianizes the summaries, i.e.

$$-2\ln \mathcal{L}(\mathbf{x}|\boldsymbol{\theta}) = (\mathbf{x} - \boldsymbol{\mu}_{\ell})^T \mathbf{C}_{\ell}^{-1} (\mathbf{x} - \boldsymbol{\mu}_{\ell})$$

What is ℓ ?

Non-linear likelihood-free compression

Start with unknown likelihood $\mathcal{L}(\mathbf{d}|\boldsymbol{\theta})$

Data $\mathbf{d}$ can be generated using a model $\mathcal{M}(\boldsymbol{\theta})$ with parameters $\boldsymbol{\theta}$

Transform data

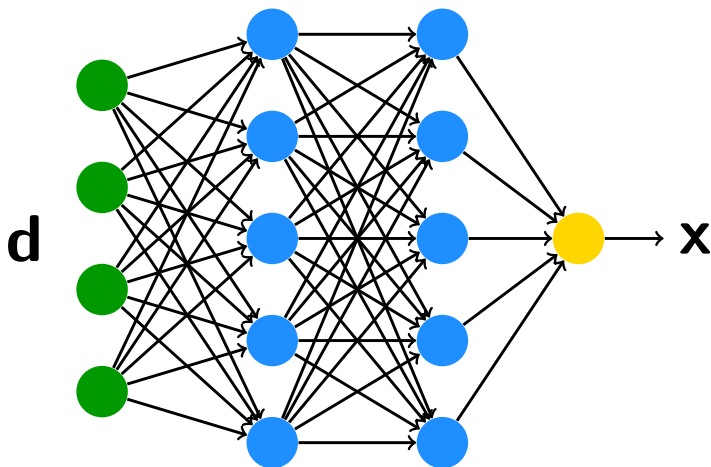
Find a function, $\ell : \mathbf{d} \rightarrow \mathbf{x}$ which Gaussianizes the summaries, i.e.

$$-2\ln\mathcal{L}(\mathbf{x}|\boldsymbol{\theta}) = (\mathbf{x} - \boldsymbol{\mu}_{\ell}(\boldsymbol{\theta}))^{\top} \mathbf{C}_{\ell}^{-1} (\mathbf{x} - \boldsymbol{\mu}_{\ell}(\boldsymbol{\theta}))$$

What is ℓ ?

A neural network!

Information maximising neural network



Information maximising neural network

Compression network

Fewer outputs than inputs (outputs can be as few as the number of model parameters)

Should be built for your data structure for best performance

Information maximising neural network

Compression network

Fewer outputs than inputs (outputs can be as few as the number of model parameters)

Should be built for your data structure for best performance

Loss function

$$\begin{aligned}\text{Loss} &= -\frac{1}{2}|\mathbf{F}|^2 \\ &= -\frac{1}{2}|\boldsymbol{\mu}_{\ell,\alpha}^{\top} \mathbf{C}^{-1} \boldsymbol{\mu}_{\ell,\beta}|^2\end{aligned}$$

Information maximising neural network

Compression network

Fewer outputs than inputs (outputs can be as few as the number of model parameters)

Should be built for your data structure for best performance

Loss function

$$\begin{aligned}\text{Loss} &= -\frac{1}{2}|\mathbf{F}|^2 \\ &= -\frac{1}{2}|\boldsymbol{\mu}_{\ell,\alpha}^{\top} \mathbf{C}^{-1} \boldsymbol{\mu}_{\ell,\beta}|^2\end{aligned}$$

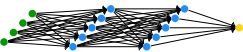
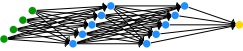
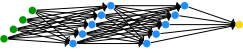
Simulations

Need high quality simulations of the real data

Only needed at the fiducial parameter values for $\mathbf{C}^{-1}$

Need the derivative of the simulation (analytic or numerical (surpressed sample variance)) for $\boldsymbol{\mu}_{\ell,\alpha}$

Information maximising neural network

$\mathbf{d}_1^{\text{s fid}}$  $\mathbf{x}_1^{\text{s fid}}$ $\left(\frac{\partial \mathbf{d}_1^{\text{s fid}}}{\partial \theta_\alpha} \right)$ $\frac{\partial \mathbf{x}_1^{\text{s fid}}}{\partial \mathbf{d}_1^{\text{s fid}}}$			$\frac{\partial \mathbf{x}_i}{\partial \theta_\alpha} = \frac{\partial \mathbf{x}_i}{\partial \mathbf{d}_i} \frac{\partial \mathbf{d}_i}{\partial \theta_\alpha}$
$\mathbf{d}_2^{\text{s fid}}$  $\mathbf{x}_2^{\text{s fid}}$ $\left(\frac{\partial \mathbf{d}_2^{\text{s fid}}}{\partial \theta_\alpha} \right)$ $\frac{\partial \mathbf{x}_2^{\text{s fid}}}{\partial \mathbf{d}_2^{\text{s fid}}}$			
$\vdots$			
$\mathbf{d}_{n_s}^{\text{s fid}}$  $\mathbf{x}_{n_s}^{\text{s fid}}$ $\left(\frac{\partial \mathbf{d}_{n_s}^{\text{s fid}}}{\partial \theta_\alpha} \right)$ $\frac{\partial \mathbf{x}_{n_s}^{\text{s fid}}}{\partial \mathbf{d}_{n_s}^{\text{s fid}}}$			

$\mathbf{C}_f^{-1}$	$\longrightarrow \mathbf{F}_{\alpha\beta}$
$\mu_{f,\alpha}$	

Automatic physical inference

Automatic physical inference

Make simulations of data, $\left\{ \mathbf{d}_i^{\text{sfid}}, \frac{\partial \mathbf{d}_i^{\text{sfid}}}{\partial \theta_\alpha} \right\}$ at $\boldsymbol{\theta}^{\text{fid}}$

Automatic physical inference

Make simulations of data, $\left\{ \mathbf{d}_i^{\text{sfid}}, \frac{\partial \mathbf{d}_i^{\text{sfid}}}{\partial \theta_\alpha} \right\}$ at $\boldsymbol{\theta}^{\text{fid}}$

Train IMNN to maximise Fisher information

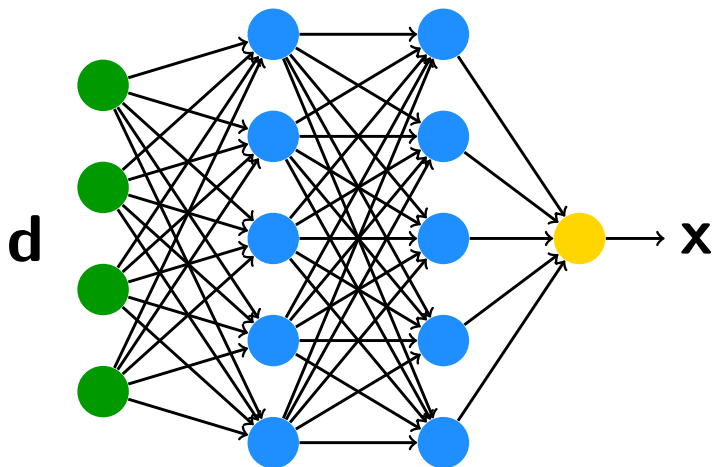
Automatic physical inference

Make simulations of data, $\left\{ \mathbf{d}_i^{\text{sfid}}, \frac{\partial \mathbf{d}_i^{\text{sfid}}}{\partial \theta_\alpha} \right\}$ at $\boldsymbol{\theta}^{\text{fid}}$

Train IMNN to maximise Fisher information

Compress real data

Automatic physical inference



Automatic physical inference

Make simulations of data, $\left\{ \mathbf{d}_i^{\text{sfid}}, \frac{\partial \mathbf{d}_i^{\text{sfid}}}{\partial \theta_\alpha} \right\}$ at $\boldsymbol{\theta}^{\text{fid}}$

Train IMNN to maximise Fisher information

Compress real data

Draw $\boldsymbol{\theta}$ from $p(\boldsymbol{\theta})$ and create simulation

Automatic physical inference

Make simulations of data, $\left\{ \mathbf{d}_i^{\text{sfid}}, \frac{\partial \mathbf{d}_i^{\text{sfid}}}{\partial \theta_\alpha} \right\}$ at $\boldsymbol{\theta}^{\text{fid}}$

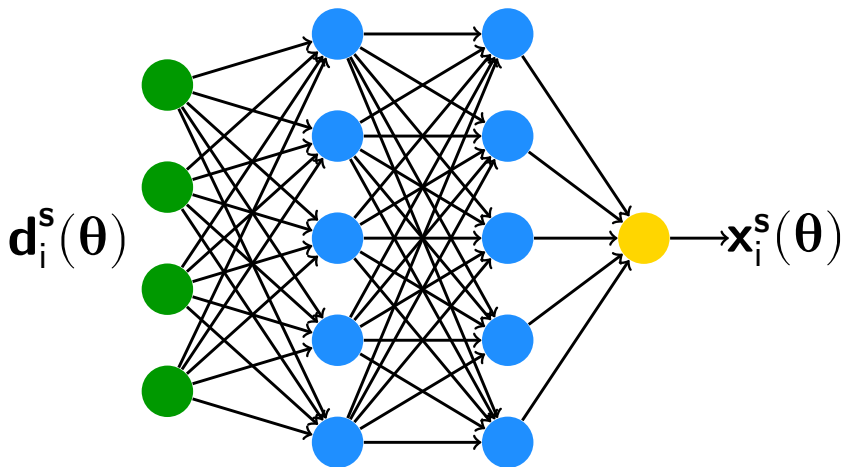
Train IMNN to maximise Fisher information

Compress real data

Draw $\boldsymbol{\theta}$ from $p(\boldsymbol{\theta})$ and create simulation

Compress simulation

Automatic physical inference



Automatic physical inference

Make simulations of data, $\left\{ \mathbf{d}_i^{\text{sfid}}, \frac{\partial \mathbf{d}_i^{\text{sfid}}}{\partial \theta_\alpha} \right\}$ at $\boldsymbol{\theta}^{\text{fid}}$

Train IMNN to maximise Fisher information

Compress real data

Draw $\boldsymbol{\theta}$ from $p(\boldsymbol{\theta})$ and create simulation

Compress simulation

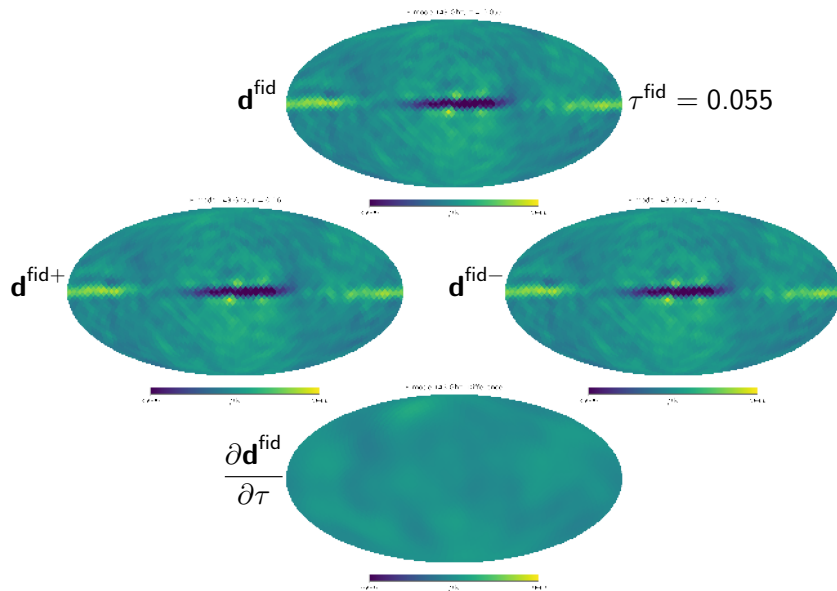
Accept or reject proposed sample to build approximate posterior distribution

Example time!

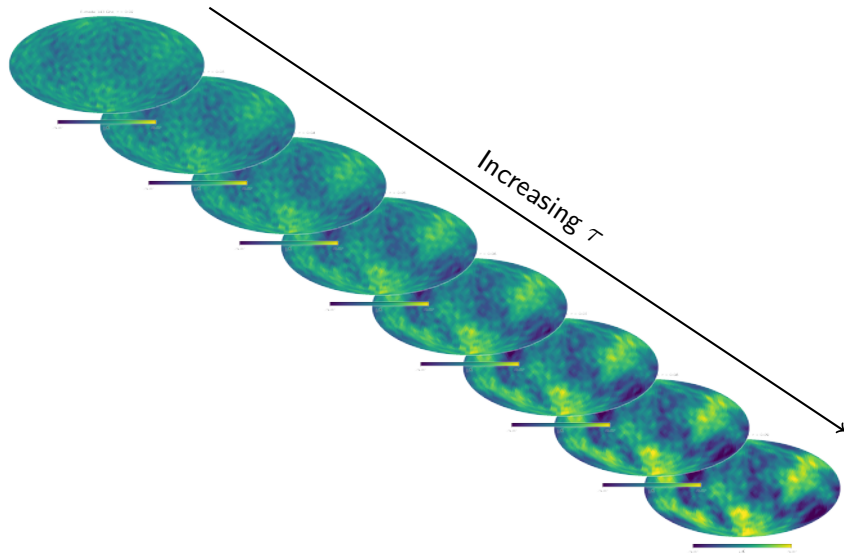
Can we infer τ from CMB polarisation maps?

Make simulations of data, $\{\mathbf{d}_i^{\text{sfid}}, \mathbf{d}_i^{\text{sfid-}}, \mathbf{d}_i^{\text{sfid+}}\}$ at θ^{fid}

CMB polarisation maps



Effect of τ on the maps

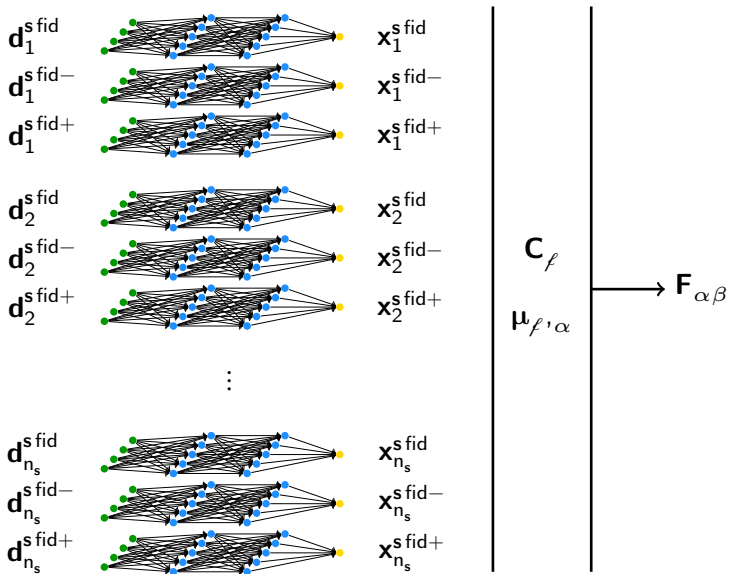


Can we infer τ from maps?

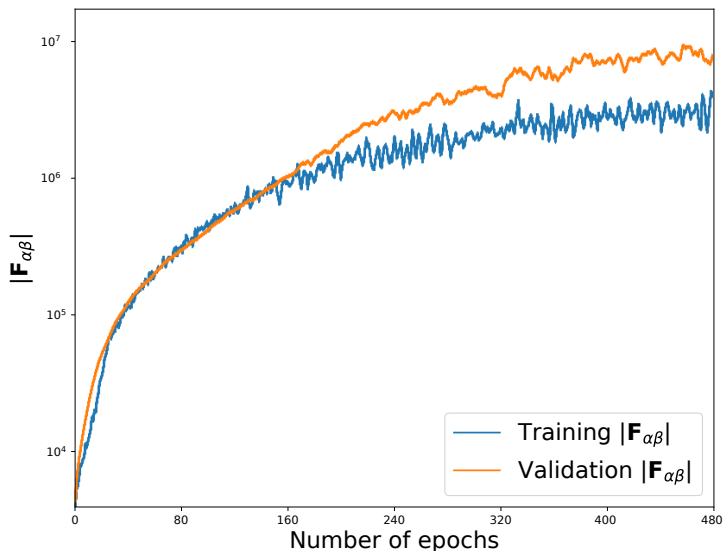
Make simulations of data, $\{\mathbf{d}_i^{\text{sfid}}, \mathbf{d}_i^{\text{sfid-}}, \mathbf{d}_i^{\text{sfid+}}\}$ at θ^{fid}

Train IMNN to maximise Fisher information

Information maximising neural network



How much does the Fisher grow?



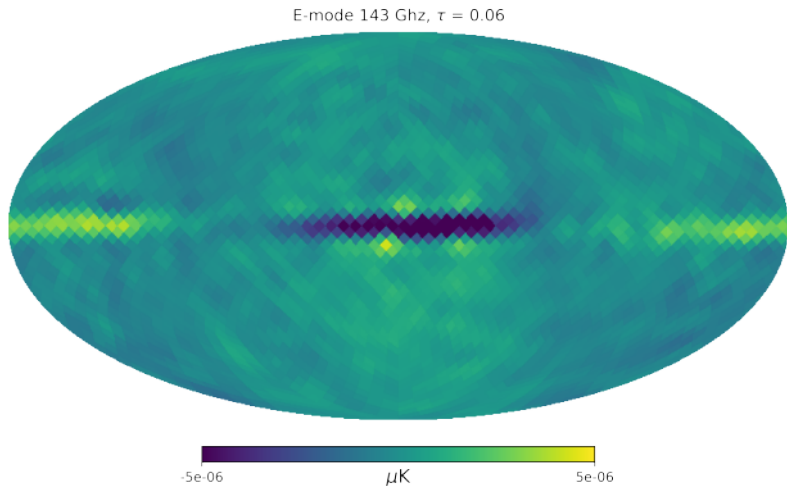
Can we infer τ from maps?

Make simulations of data, $\{\mathbf{d}_i^{\text{sfid}}, \mathbf{d}_i^{\text{sfid-}}, \mathbf{d}_i^{\text{sfid+}}\}$ at θ^{fid}

Train IMNN to maximise Fisher information

Compress real data

Our *real* data at $\tau = 0.06$



Can we infer τ from maps?

Make simulations of data, $\{\mathbf{d}_i^{\text{sfid}}, \mathbf{d}_i^{\text{sfid-}}, \mathbf{d}_i^{\text{sfid+}}\}$ at θ^{fid}

Train IMNN to maximise Fisher information

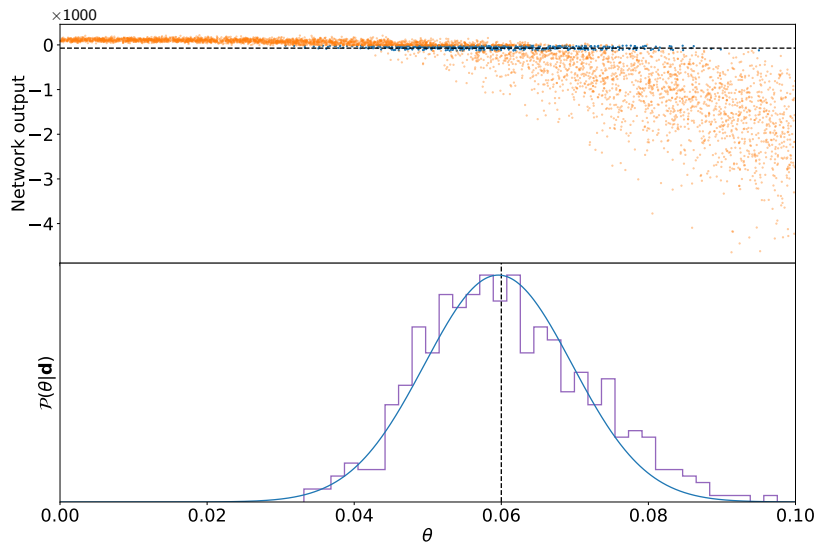
Compress real data

Draw θ from $p(\theta)$ and create simulation

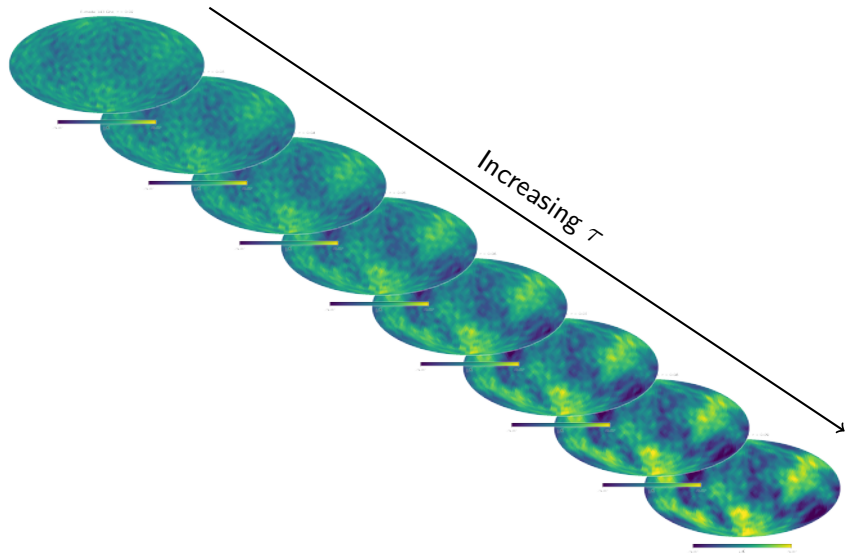
Compress simulation

Accept or reject proposed sample to build approximate posterior distribution

Approximate posterior distribution

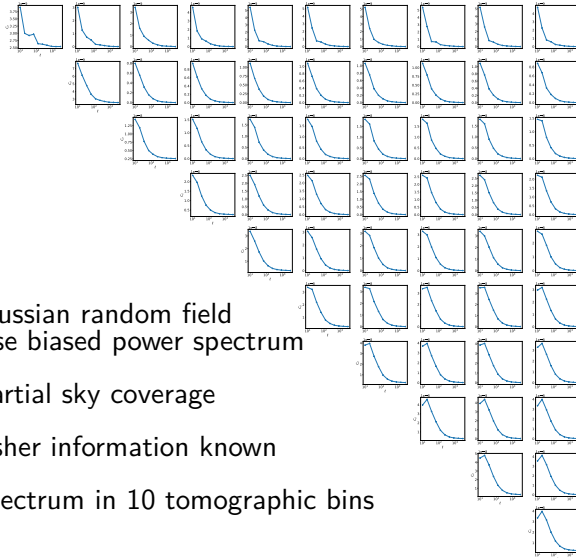


Effect of τ on the maps



Another quick example!

Cosmic shear



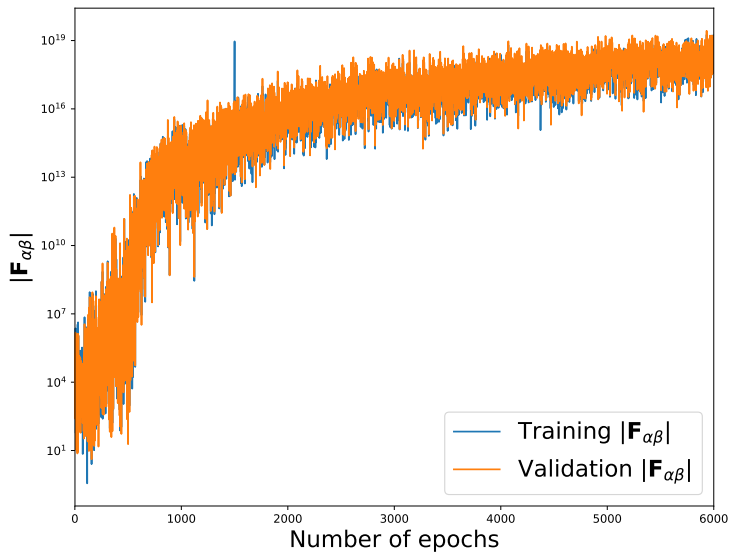
Draw Gaussian random field
from noise biased power spectrum

Mimic partial sky coverage

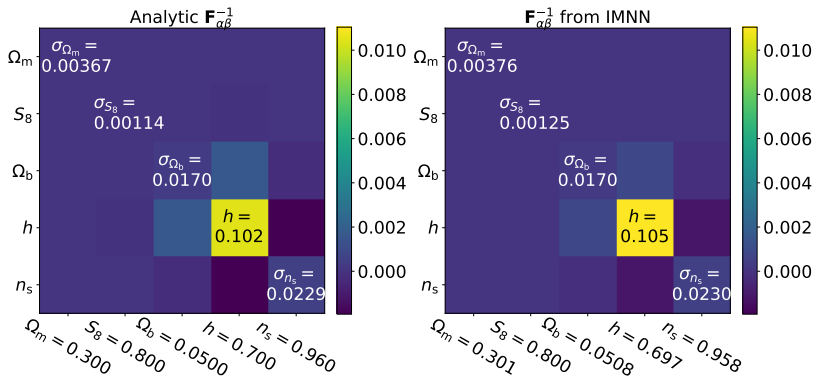
Exact Fisher information known

Power spectrum in 10 tomographic bins

How much does the Fisher grow?



Fisher comparison

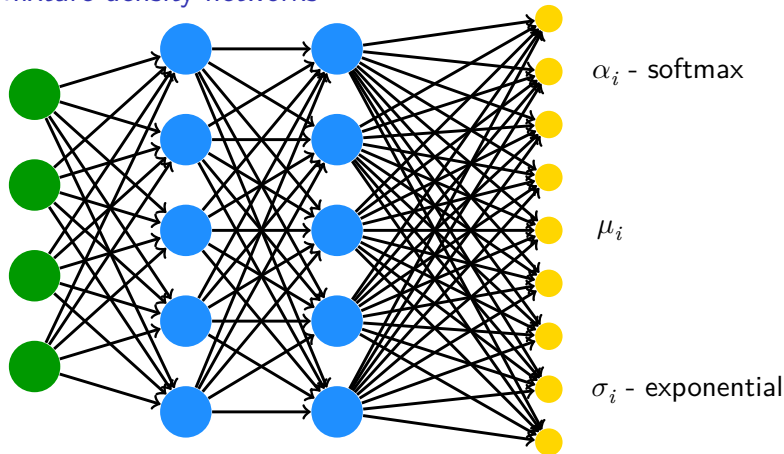


Density estimation likelihood-free estimation

Density estimation likelihood-free inference (DELFI)

- Alsing, Feeney & Wandelt, 2018

Mixture density networks



$$\mathcal{L}(\mathbf{x}|\boldsymbol{\theta}, \mathbf{w}) = \prod_{i=1}^{n_{\mathcal{N}}} \alpha_i(\mathbf{w}) \mathcal{N}(\mathbf{x}|\boldsymbol{\theta}, \mu_i(\mathbf{w}), \sigma_i(\mathbf{w}))$$

Super-efficient inference

Pretrained on the Fisher information

Really quick to produce samples to get MDN into right region

Taught actively

Geometric mean of each round of updates can be used as an updated proposal distribution

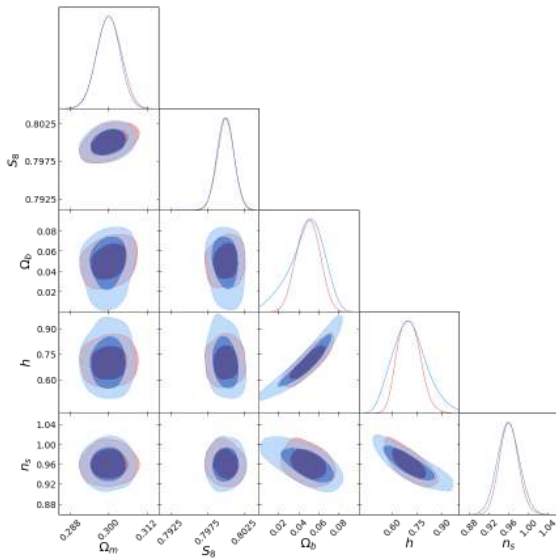
Reduce necessary simulations

Full MCMC $\sim 10^6$ simulations

DELFI $\sim \text{few} \times 10^3$ simulations

MCMC vs DELFI constraints

Note these results are MOPED compression (not IMNN)



Totally ignorant inference

What I've just presented

An optimal compression scheme for general likelihoods

A non-linear compression network trained with the Fisher information

A couple of amazing (albeit toy) examples of how well the IMNN works

A way to drastically reduce the number of simulations needed to do inference using mixture density networks

A framework with which to do totally ignorant statistical inference

Machine learning based (likelihood-free) statistical inference

Tom Charnock
Guilhem Lavaux & Ben Wandelt

Alessandro Manzotti
Justin Alsing

Institut d'Astrophysique de Paris

Physical Review D **97**, 083004

DOI [10.5281/zenodo.1175196](https://doi.org/10.5281/zenodo.1175196)

[github:information_maximiser](#)

