

# Autoencoder based surrogate model for gravitational waves from BBH mergers

---

arXiv: 2602.00203

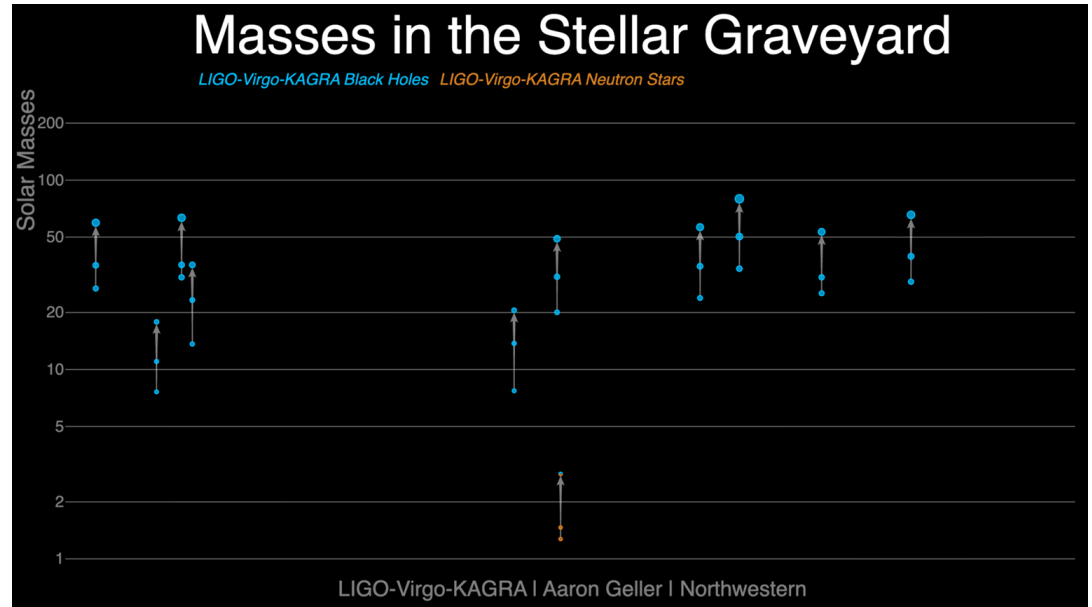
Anastasios Theodoropoulos  
CoCoNut Meeting - 08/09/2026

# Contents

- Introduction
  - Gravitational waves are heard not seen
  - Binary Black Hole mergers
  - Surrogate models
  - Autoencoders
- Data and model
  - Numerical Relativity data
  - General Idea
  - Synthetic dataset
  - Autoencoder Architecture
  - Final Architecture
- Results
  - Training
  - Accuracy
  - Generated waveforms
  - Parameter estimation
  - Speed-of-inference
  - Higher modes
- Conclusions

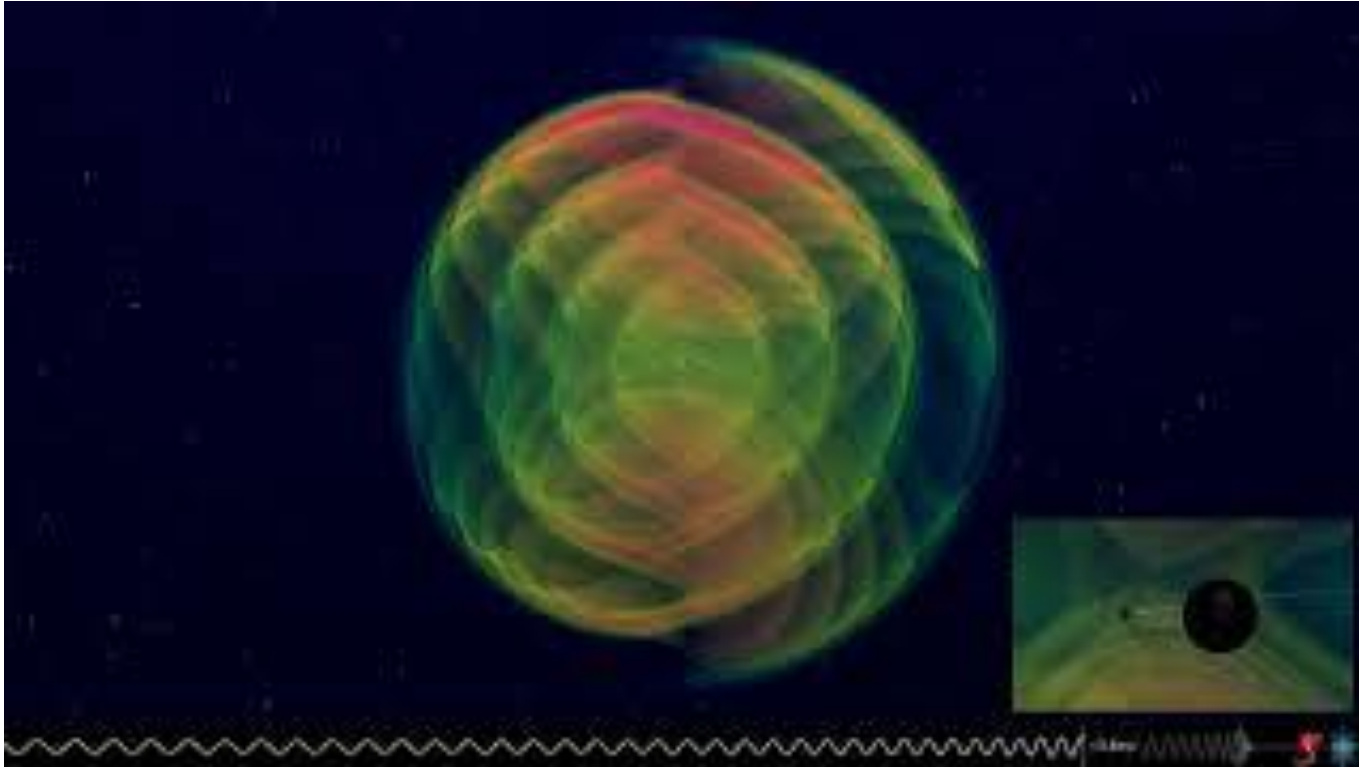
# Introduction

# Gravitational waves are heard not seen



Credit: Northwestern University

# Binary Black Hole mergers



# Surrogate models

## Post Newtonian:

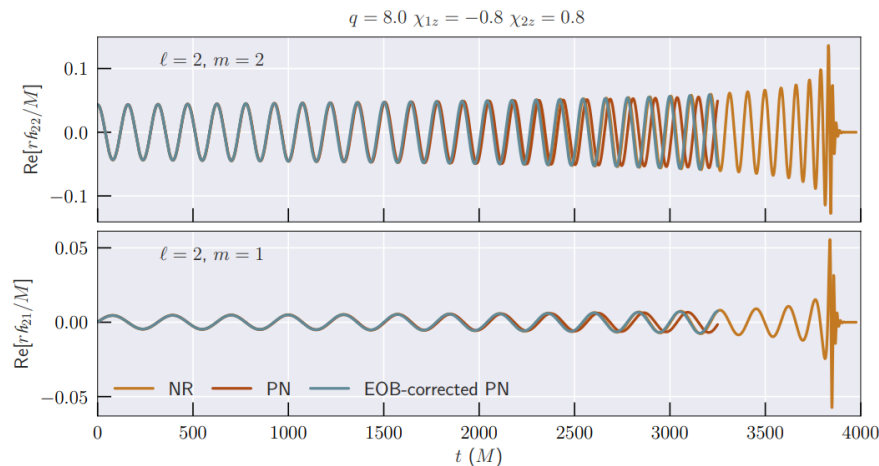
- Newtonian physics with added GR corrections
- Accurate during the inspiral

## Effective-one-body:

- Handles waveforms up to the merger

## Phenomenological:

- Reduced basis
- Interpolation

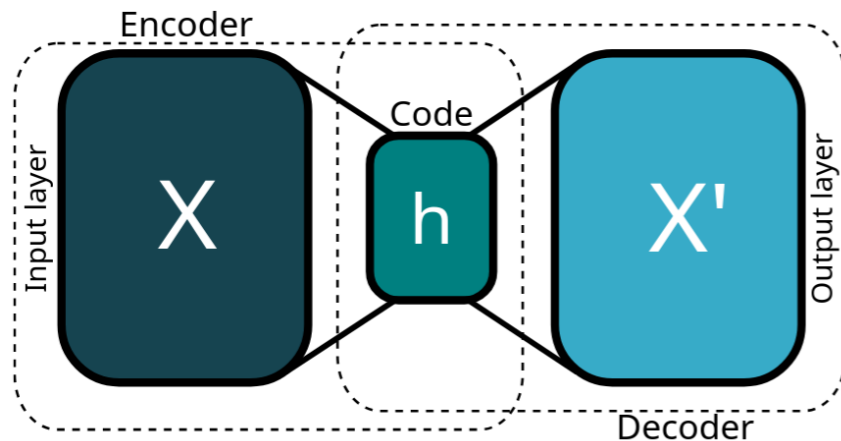
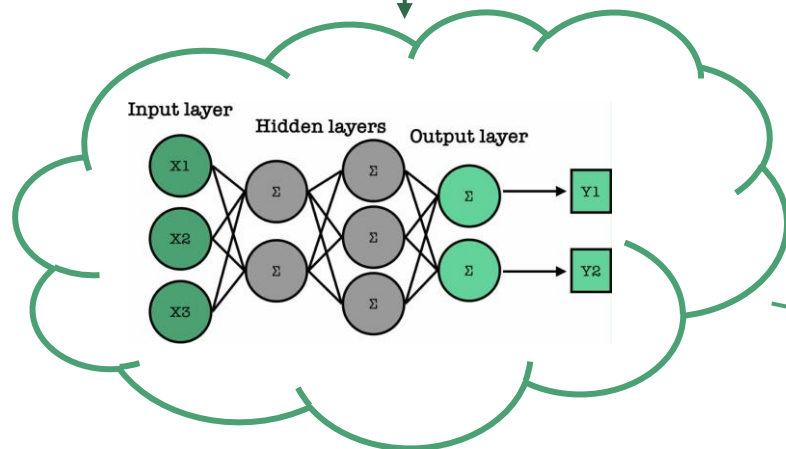


arXiv:1812.07865

# Autoencoders

## Autoencoder:

- Neural Network Architecture
- **Unsupervised** Learning
- Non-linear



$$L_1(c, \hat{c}) = \frac{1}{N} \sum_{i=1}^N |c_i - \hat{c}_i|$$

# Autoencoders

## Encoder:

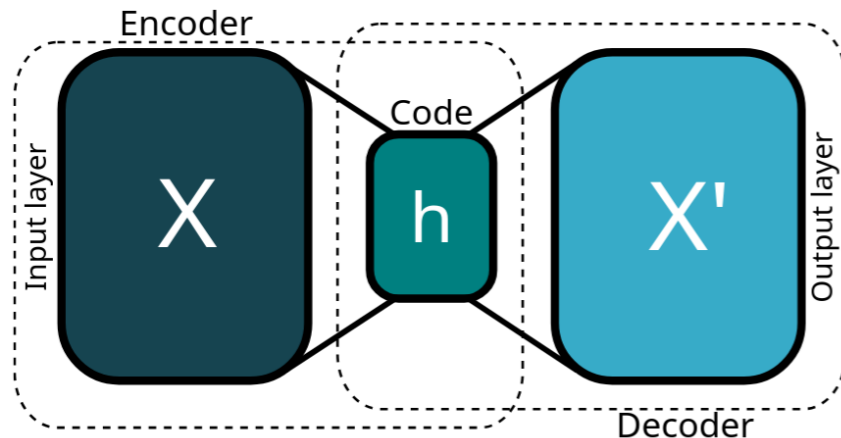
- Efficiently compresses **(encodes)** input data down to its essential features.

## Latent space:

- Features / Representation / **Code**

## Decoder:

- It reconstructs **(decodes)** the original input from this compressed representation **(code)**



# Data and model

---

# Numerical Relativity data

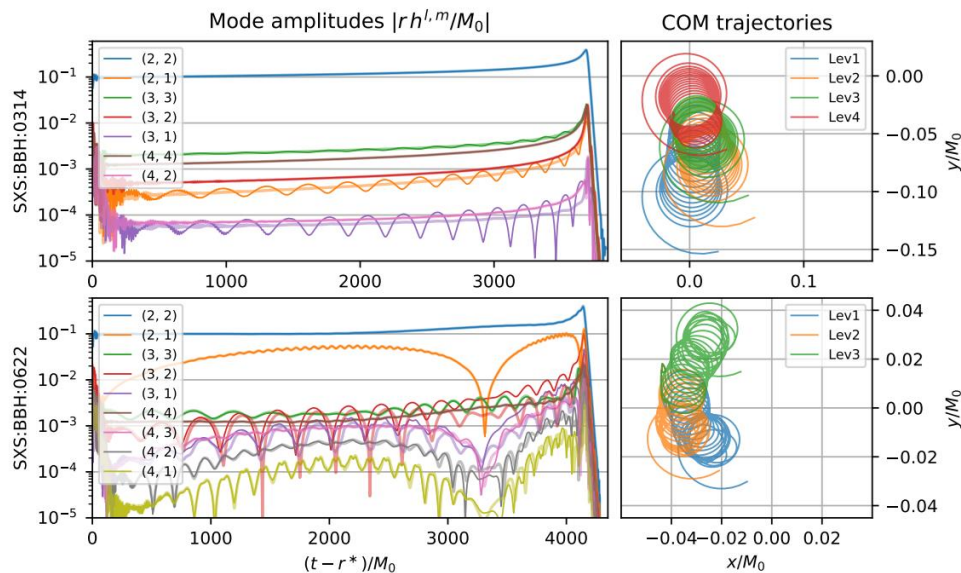
## Our samples:

- (2,2) mode
- SXS dataset

## Filtering:

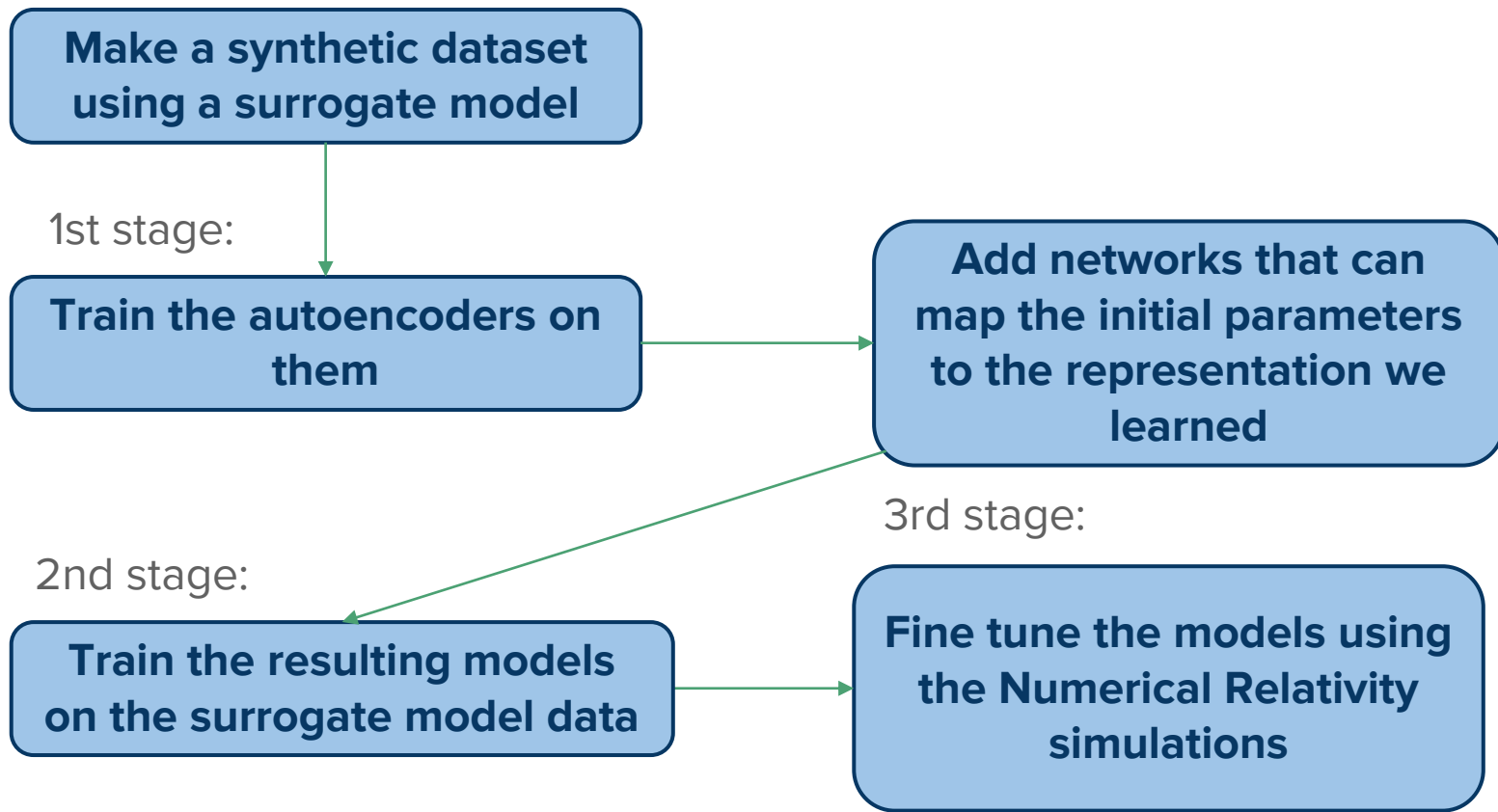
- Duration:  $d \geq 8192 M_{\text{sun}}$
- Minimal precession
- Mass ratio  $\leq 8$

**Only 376 samples!**



DOI 10.1088/1361-6382/ab34e2

## General Idea - 3 stage training



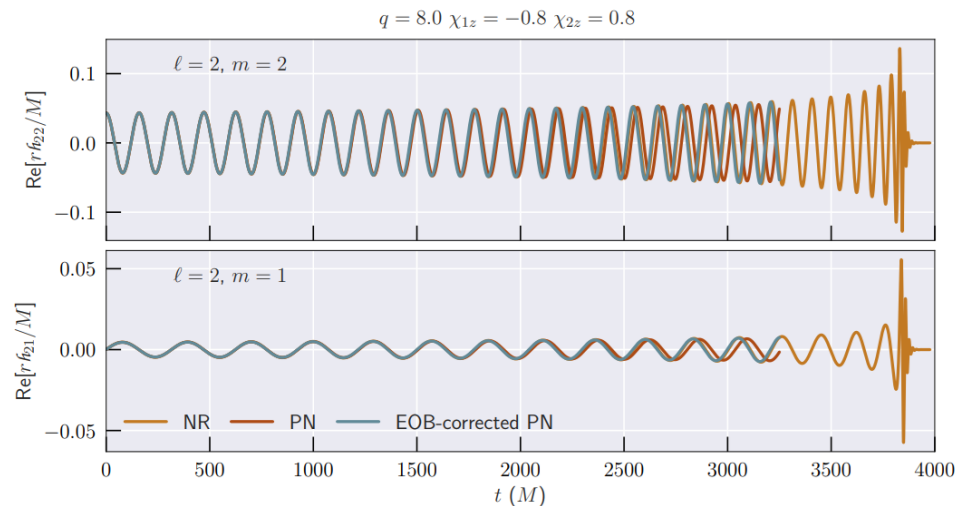
# Synthetic dataset

## Our samples:

- (2,2) mode
- NRHybSur3dq8

## Sampling:

- Uniformly
- Mass ratio  $\leq 8$
- $-0.8 < \text{spins} < 0.8$
- sample rate =  $2M^{-1}_{\odot}$



Credit: arXiv:1812.07865

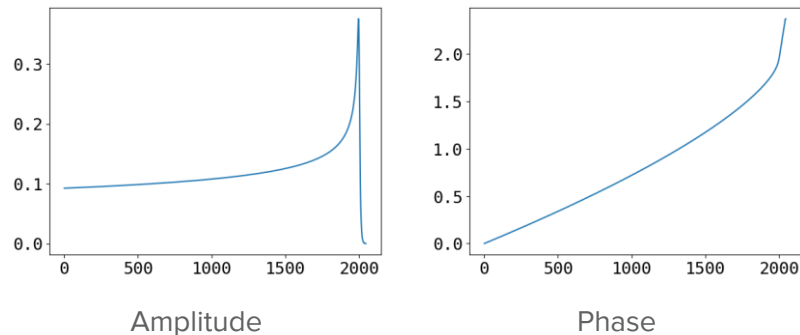
# Autoencoder Architecture

## Split into components:

Waveforms are split into phase and amplitude

## 1st step :

We construct an autoencoder for each component and train them separately



| Component       | Layer              | Units | Activation |
|-----------------|--------------------|-------|------------|
| Encoder         | Input              | 8192  | —          |
|                 | Dense 1            | 1024  | Leaky ReLU |
|                 | Dense 2            | 512   | Leaky ReLU |
|                 | Dense 3            | 256   | Leaky ReLU |
|                 | Dense 4            | 128   | Leaky ReLU |
|                 | Output (Latent)    | 64    | Linear     |
| Decoder         | Input (Latent)     | 64    | —          |
|                 | Dense 1            | 128   | Leaky ReLU |
|                 | Dense 2            | 256   | Leaky ReLU |
|                 | Dense 3            | 512   | Leaky ReLU |
|                 | Dense 4            | 1024  | Leaky ReLU |
|                 | Output             | 8192  | Linear     |
| Mapping Network | Input (Parameters) | 3     | —          |
|                 | Dense 1            | 128   | Leaky ReLU |
|                 | Dense 2            | 512   | Leaky ReLU |
|                 | Dense 3            | 512   | Leaky ReLU |
|                 | Dense 4            | 128   | Leaky ReLU |
|                 | Output (Latent)    | 64    | Linear     |

Architecture Components Table

# Final Architecture - Adding the mapping network

## Mapping network:

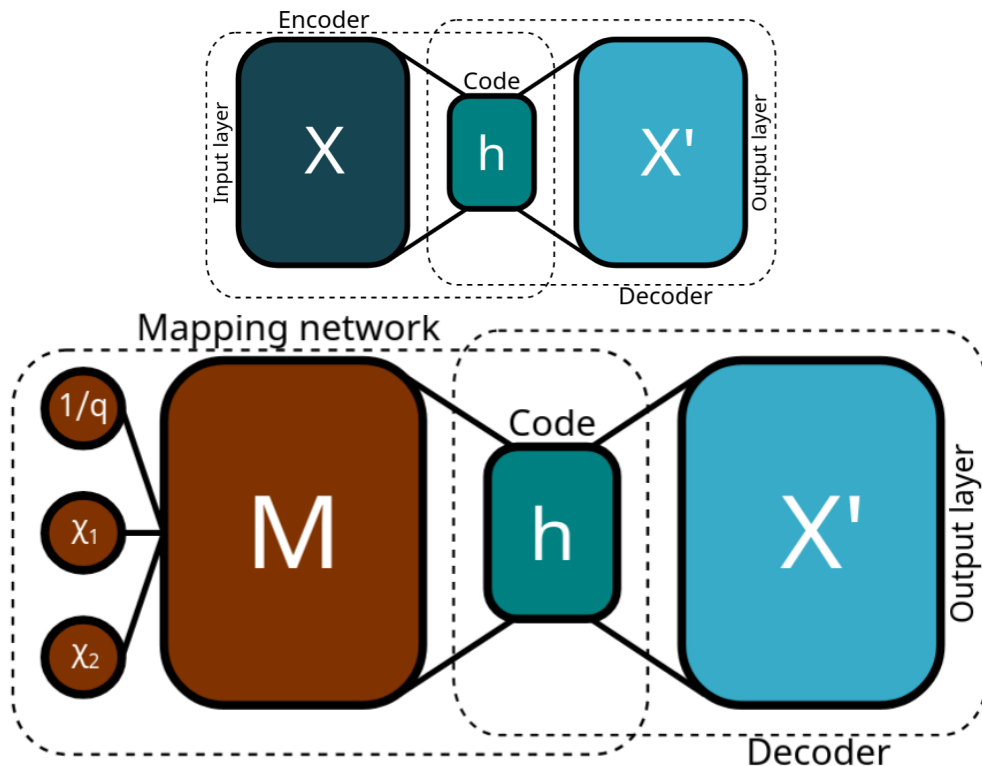
After the autoencoders are trained, we keep the decoder and replace the encoder with a mapping network

## Second step:

We retrain the full model, with the decoder weights free to change

## Fine tuning:

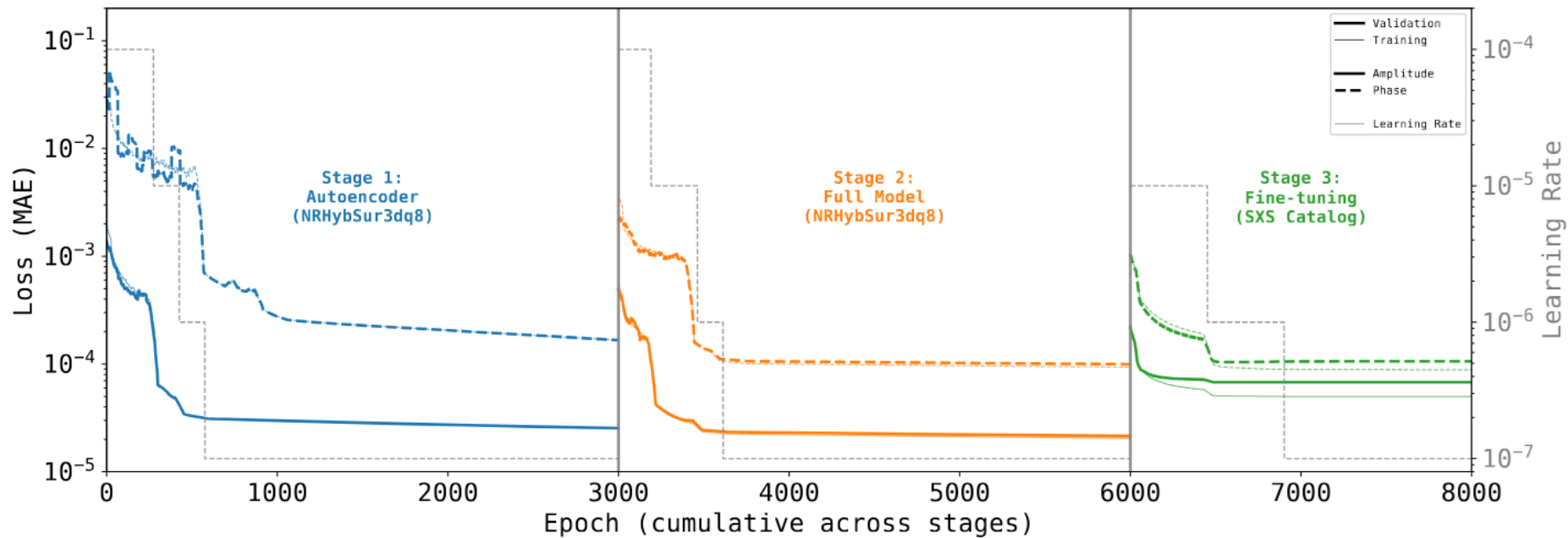
We retrain the full model on the NR waveforms, with a lower starting learning rate



# Results

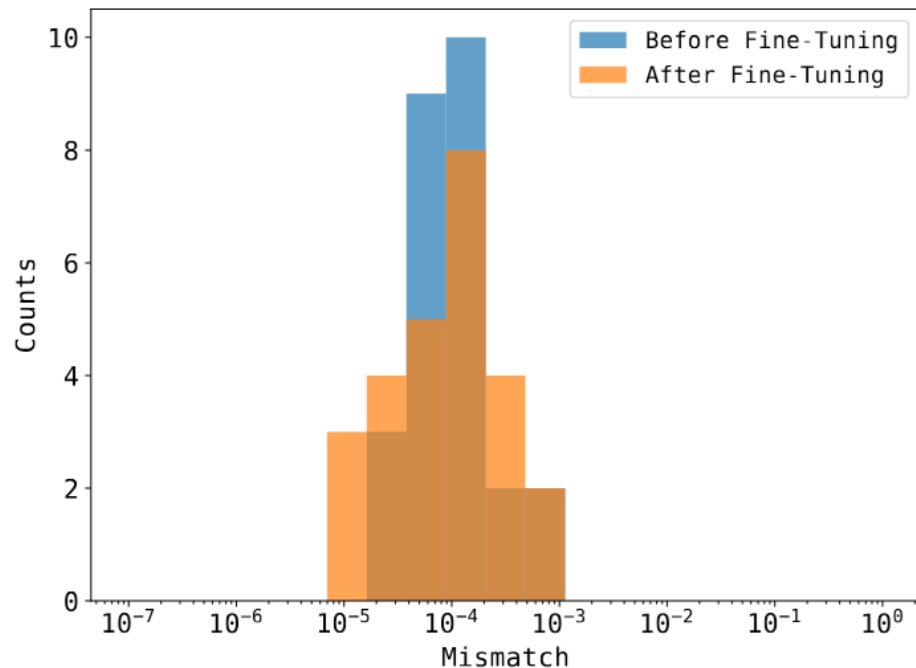
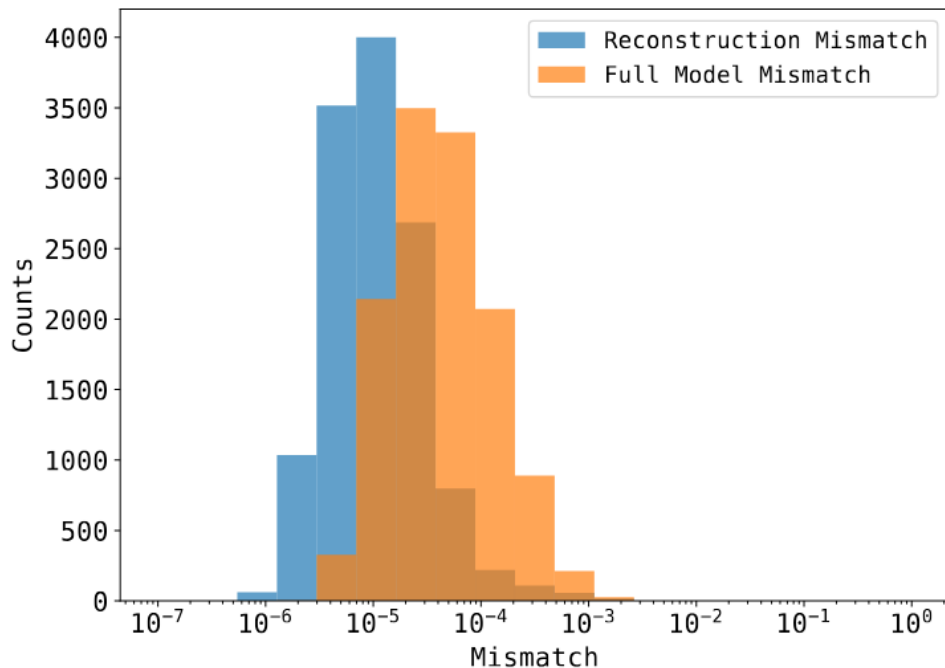
---

# Training

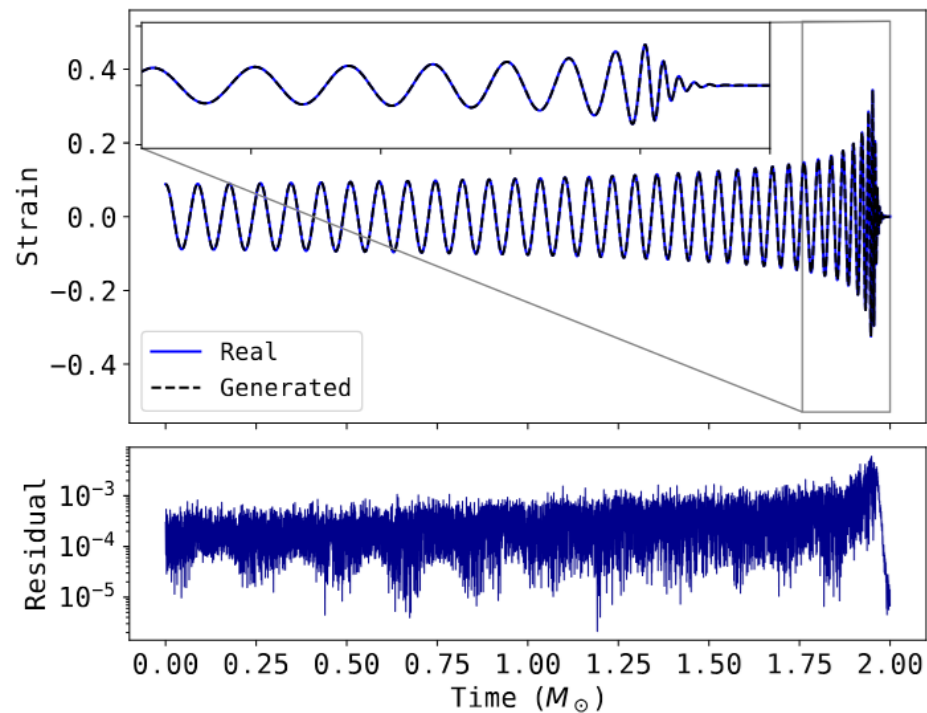
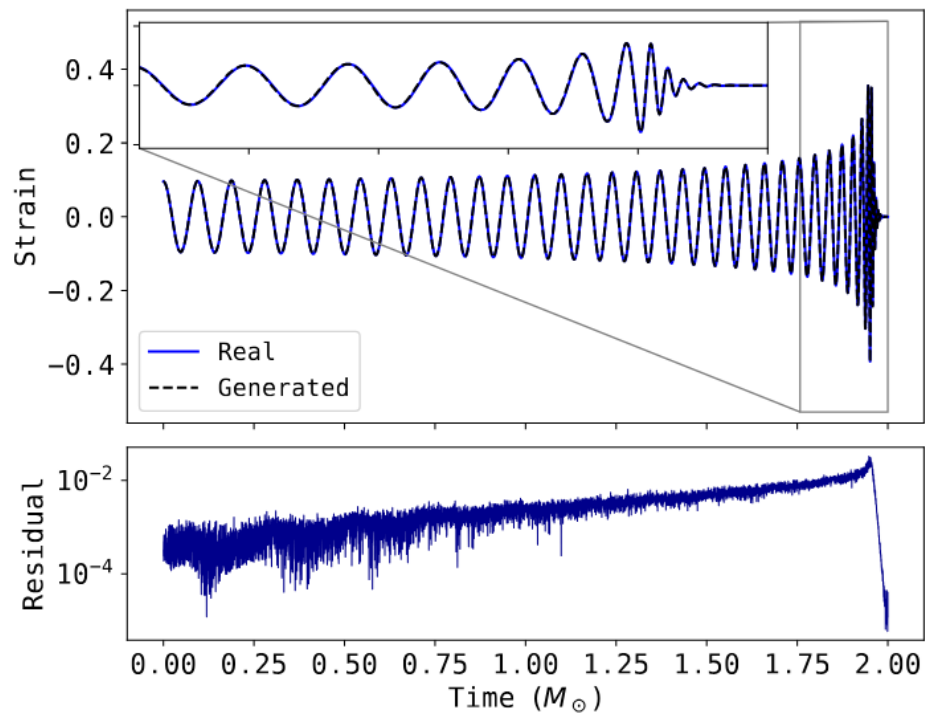


# Accuracy

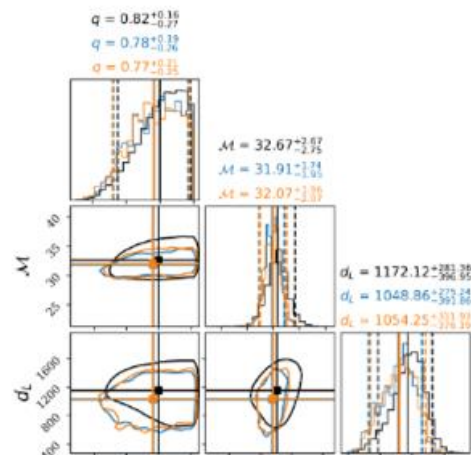
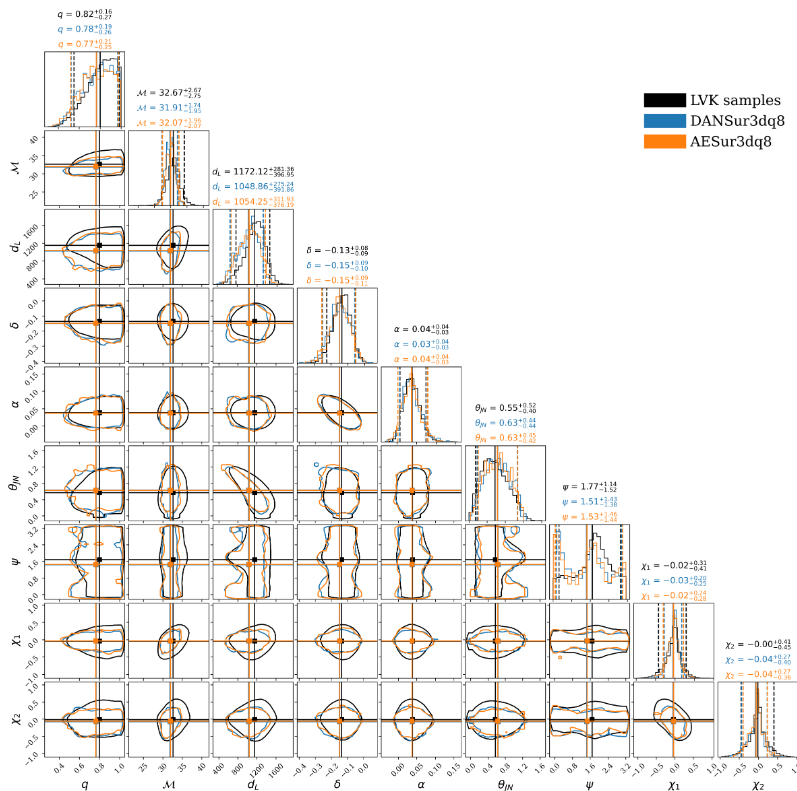
$$\mathcal{M}(h, \hat{h}) = 1 - \mathcal{O}(h, \hat{h}) = 1 - \frac{\langle h | \hat{h} \rangle}{\sqrt{\langle h | h \rangle \langle \hat{h} | \hat{h} \rangle}}$$



# Generated waveforms



# Parameter estimation



## Comparison:

When comparing with our previous work: **DAnsUr3dq8** Osvaldo Freitas et al. (2024) and LVK posterior samples. The results are **compatible**

# Speed-of-inference (in ms)

## Timing:

1. Model compilation
2. 3 warm-up computations
3. Average of 10 computations

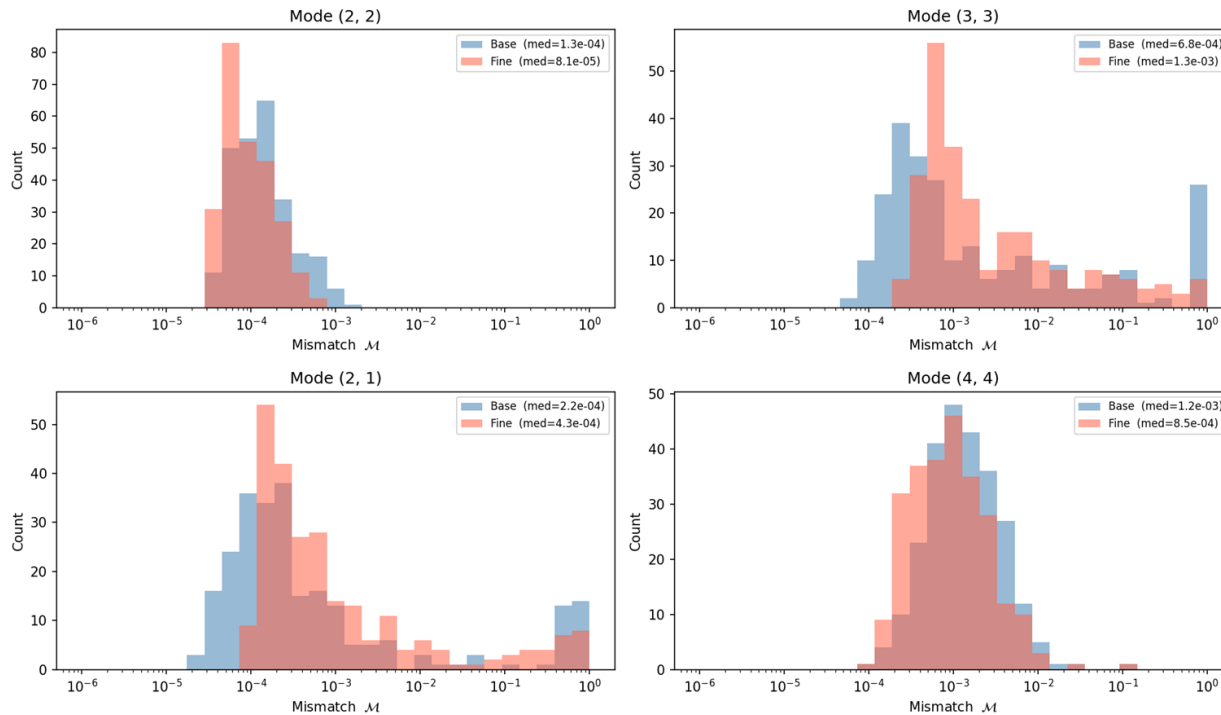
## Comparison:

1. Faster than the traditional method
2. Slower than our **PCA** based model

|         | AESur3dq8 |      | DANSur3dq8 |      | NRHybSur3dq8        |
|---------|-----------|------|------------|------|---------------------|
| Samples | CPU       | GPU  | CPU        | GPU  | CPU                 |
| $10^2$  | 6.38      | 0.46 | 1.0        | 0.18 | 881                 |
| $10^3$  | 38.0      | 1.01 | 3.0        | 0.21 | 8,830               |
| $10^4$  | 349       | 9.05 | 125        | 1.8  | 88,200              |
| $10^5$  | 3,682     | 86.6 | 1,222      | 17   | 882,000*            |
| $10^6$  | 37387     | 868  | 12,217     | 172  | $8.8 \times 10^6$ * |

# Higher modes - Ongoing work

Base vs Fine-tuned — mismatch distributions



# Conclusions

- We construct an autoencoder based surrogate model
- We compare the generated waveforms with Numerical Relativity waveforms from the **SXS** dataset
- We perform a **bilby** parameter estimation, showing the compatibility of our results with the current models
- Future tests/extensions include:
  - Multi-modal surrogates (Talk from Osvaldo Freitas)
  - Precessing spins
  - BNS mergers
  - Implementation in gw detection pipelines

# Thank you for your attention!

---

This work is supported by the Spanish Agencia Estatal de Investigación (grant PID2024-159689NB-C21) funded by MICIU/AEI/10.13039/501100011033 and by FEDER / EU and from the Universitat de València through the Atracció de Talent fellowship.

