

Renyi Neural Processes



Xuesong Wang

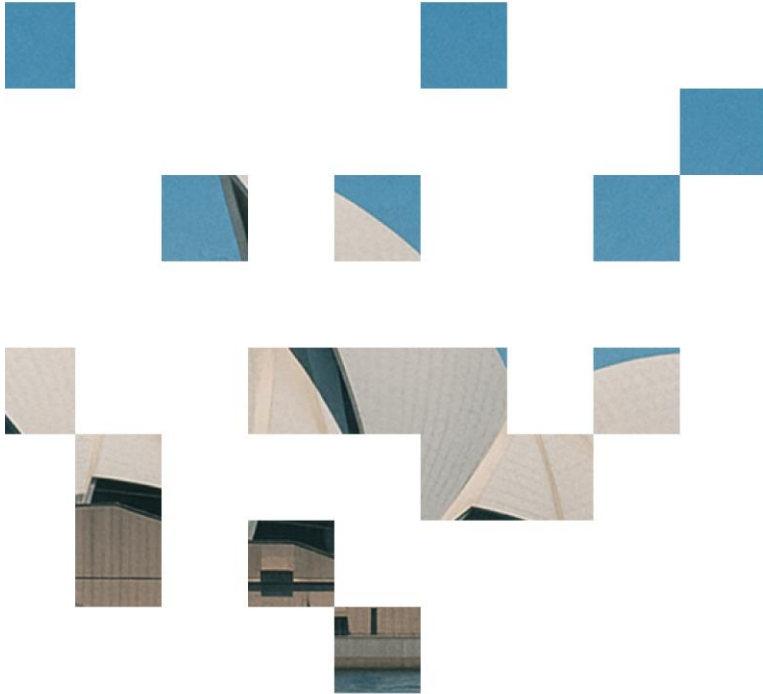


He Zhao

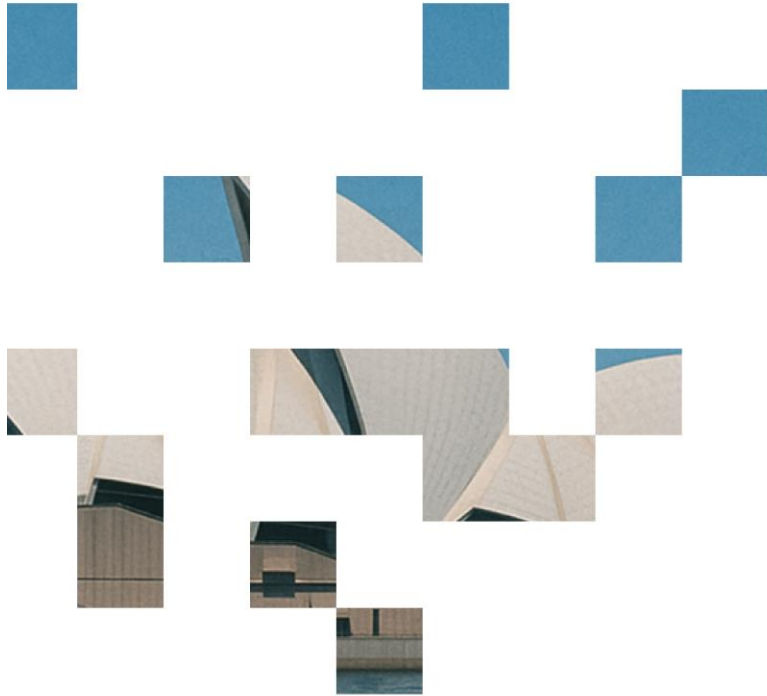


Edwin V. Bonilla

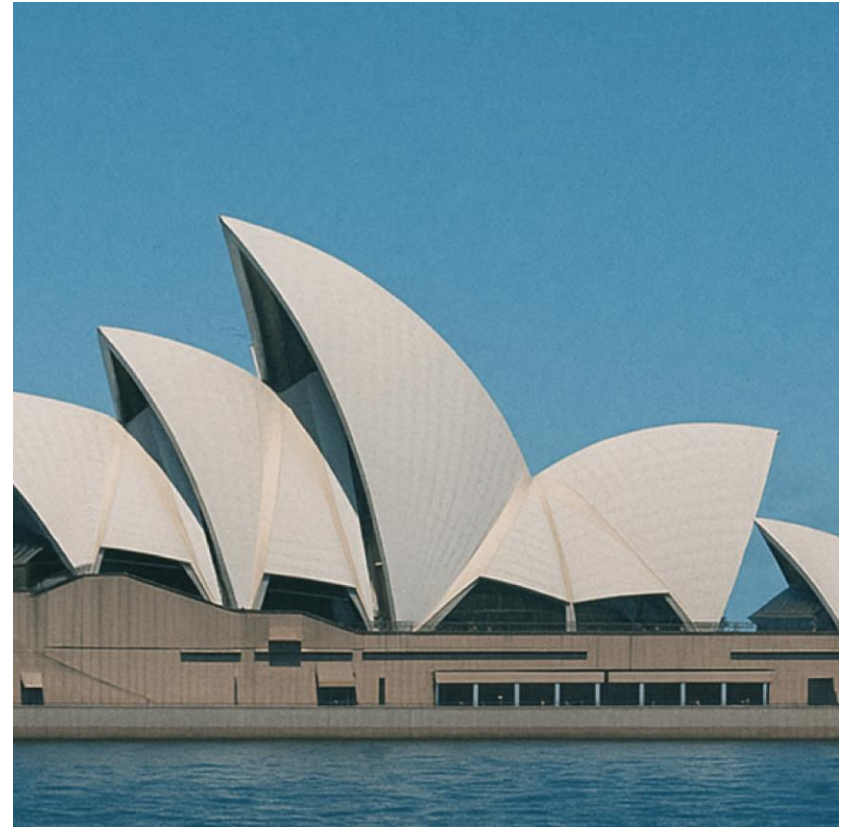
Where is this ?



Where is this ?

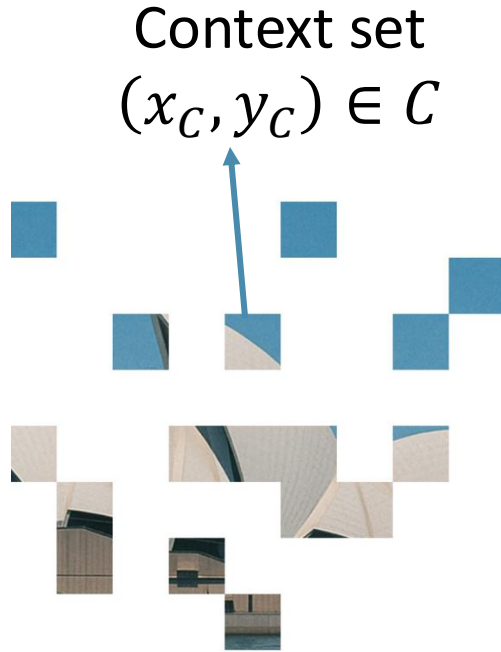


Sydney Opera House !



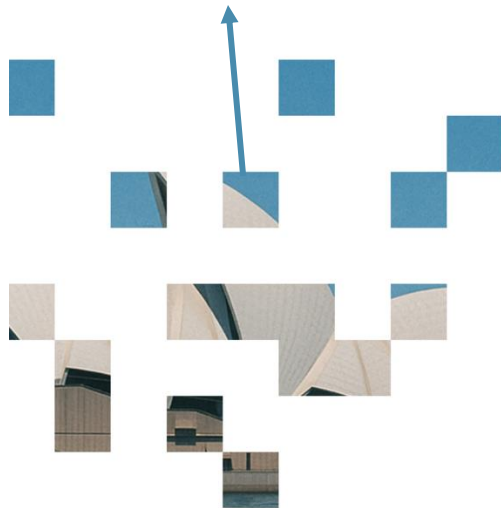


Neural Processes: a predictive model based on contexts

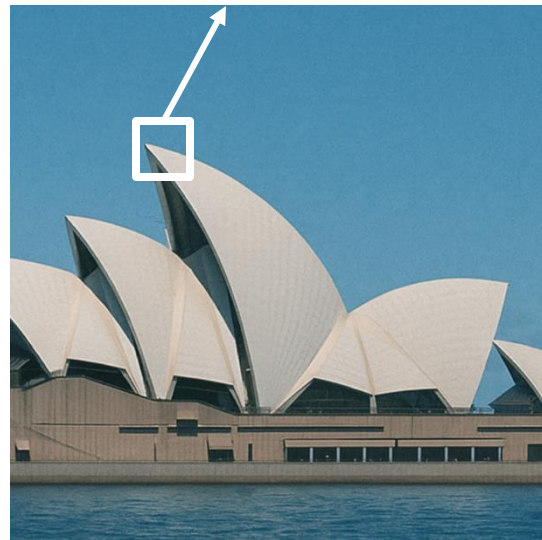


Neural Processes: a predictive model based on contexts

Context set
 $(x_C, y_C) \in \mathcal{C}$



Target set
 $(x_T, y_T) \in \mathcal{T}$



$$p(y_T | x_T, \mathcal{C}) = ?$$



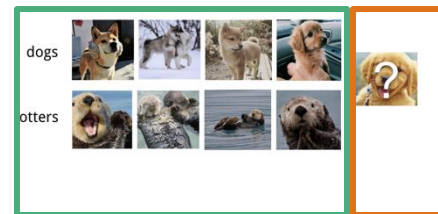
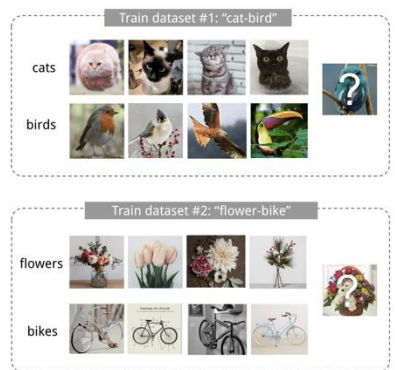
Why is this important ?

fast adaptation by conditioning on a new context set

Why is this important ?

fast adaptation by conditioning on a new context set

- Meta learning & multitask learning



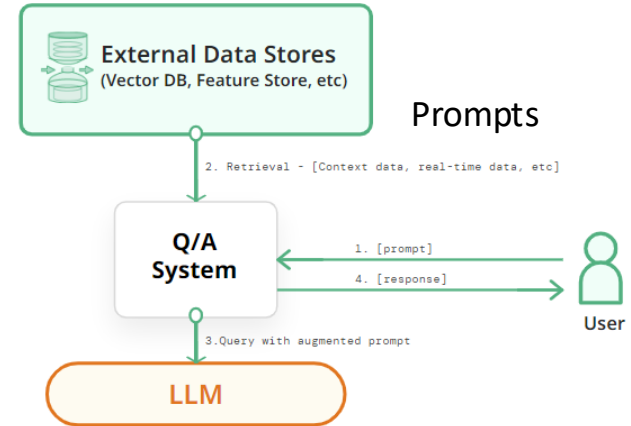
new task



Why is this important ?

fast adaptation by conditioning on a new context set

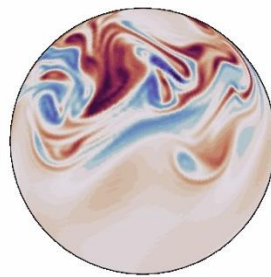
- Meta learning & multitask learning
- Retrieval-augmented generation (RAG) for LLMs



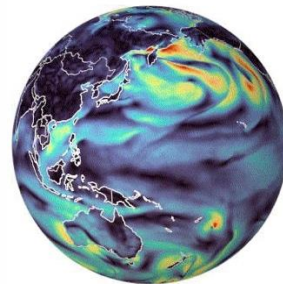
Why is this important ?

fast adaptation by conditioning on a new context set

- Meta learning & multitask learning
- Retrieval-augmented generation (RAG) for LLMs
- **Sim2real in scientific computing**
- ...



simulation



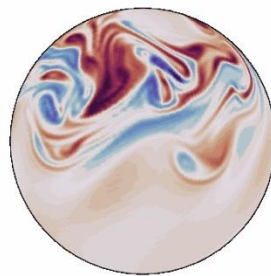
physical

Why is this important ?

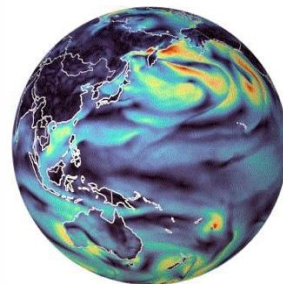
fast adaptation by conditioning on a new context set

- Meta learning & multitask learning
- Retrieval-augmented generation (RAG) for LLMs
- **Sim2real in scientific computing**
- ...

NPs empowers an interpretable and robust decision-making process.

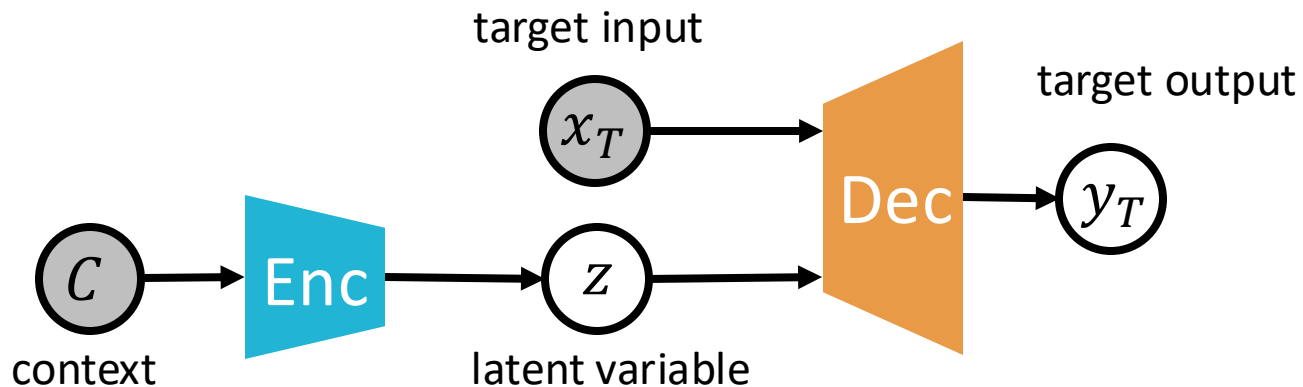


simulation

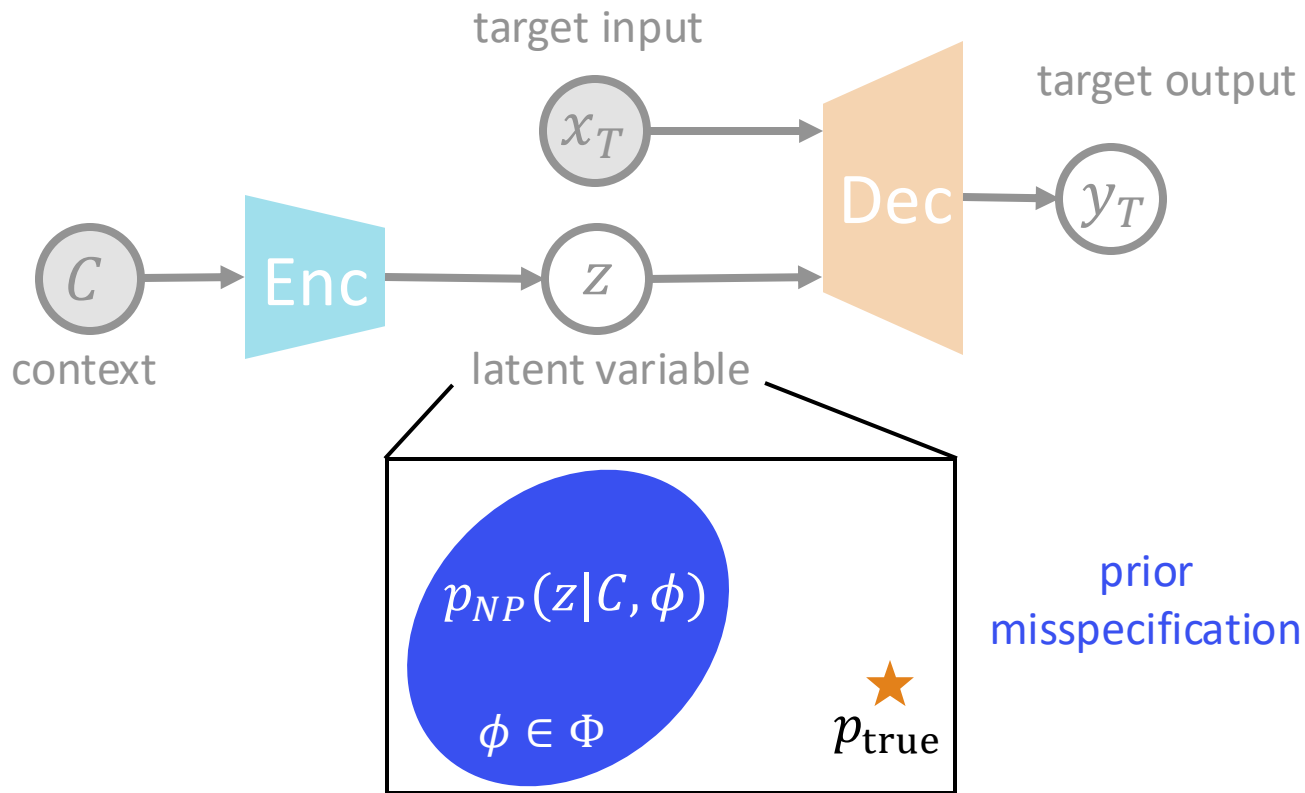


physical

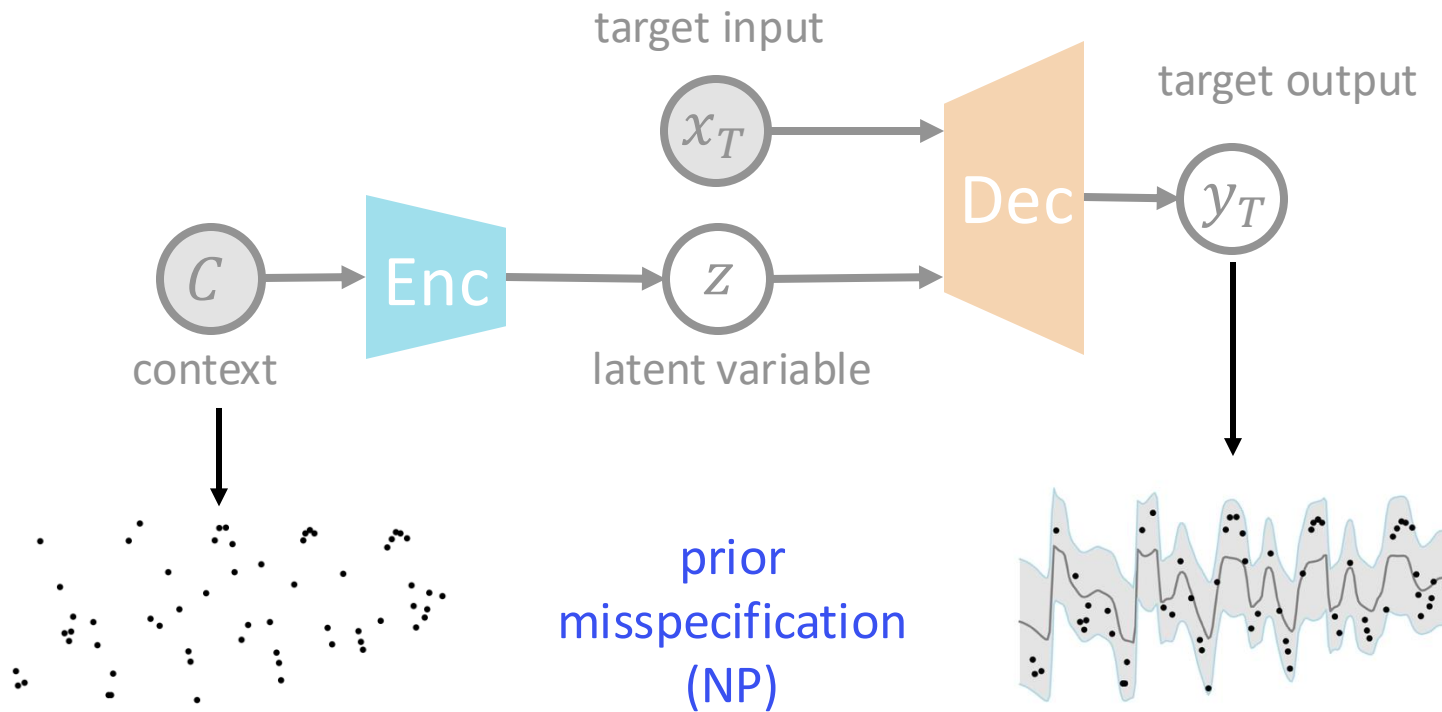
Challenge with encoding the context information & Prior misspecification



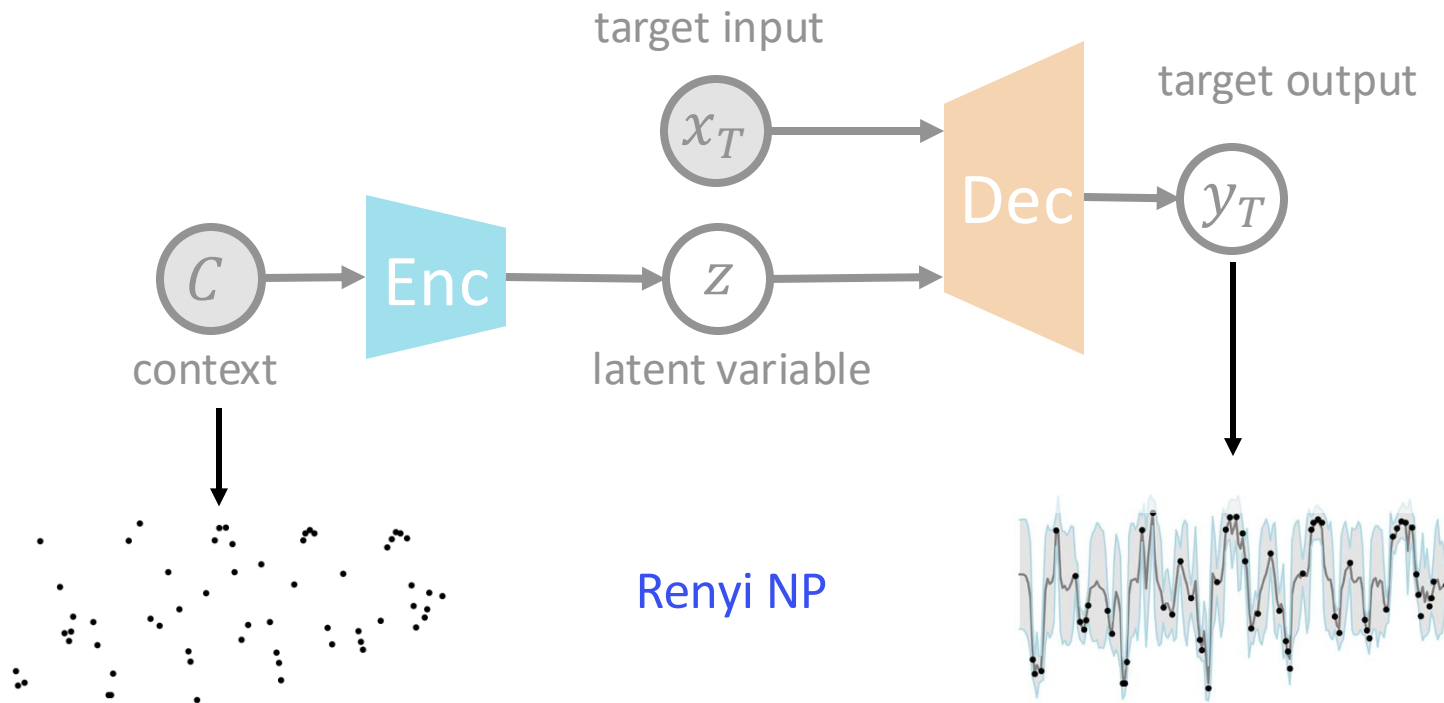
Challenge with encoding the context information & Prior misspecification



Challenge with encoding the context information & Prior misspecification



Challenge with encoding the context information & Prior misspecification





Cause for prior misspecification

$$-\mathcal{L}_{VI}(\theta, \phi) = \mathbb{E}_{q_{\phi}(\mathbf{z}|C, T)} \log p_{\theta}(Y_T|X_T, \mathbf{z}) - D_{KL}(q_{\phi}(\mathbf{z}|C, T) \| p(\mathbf{z}|C))$$



Cause for prior misspecification

$$-\mathcal{L}_{VI}(\theta, \phi) = \overset{\text{Likelihood}}{\mathbb{E}_{q_{\phi}(\mathbf{z}|C,T)} \log p_{\theta}(Y_T|X_T, \mathbf{z})} - \overset{\text{Posterior}}{D_{KL}(\underbrace{q_{\phi}(\mathbf{z}|C,T)}_{\text{Posterior}} || \underbrace{p(\mathbf{z}|C)}_{\text{Prior}})}$$



Cause for prior misspecification

$$-\mathcal{L}_{VI}(\theta, \phi) = \mathbb{E}_{q_{\phi}(\mathbf{z}|C,T)} \log p_{\theta}(Y_T|X_T, \mathbf{z}) - D_{KL}(\underbrace{q_{\phi}(\mathbf{z}|C, T)}_{\text{Posterior}} \parallel \underbrace{p(\mathbf{z}|C)}_{\text{Prior}})$$

$q_{\phi}(\mathbf{z}|C)$
approximated prior
(parameter coupling)



Prior misspecification: $q_\phi(\mathbf{z}|C) \neq p(\mathbf{z}|C), \forall \phi \in \Phi$

What if we don't have to trust this approximated prior?



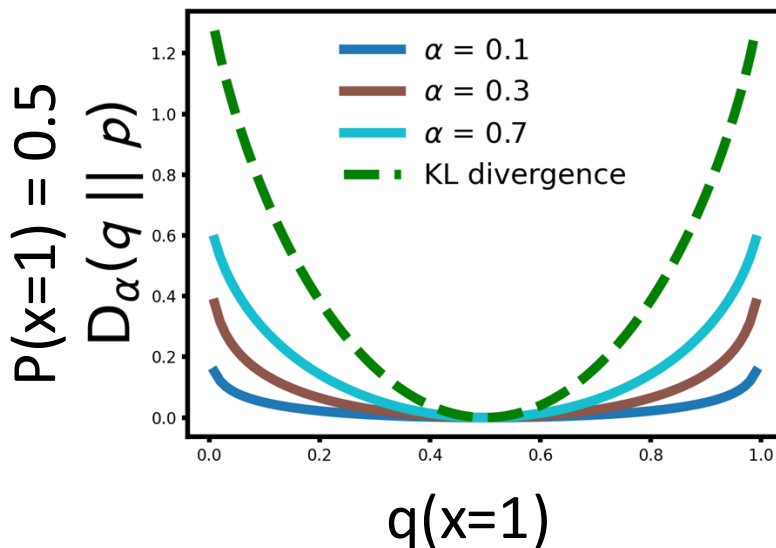
Renyi divergence (1961):

$$D_{\alpha}(q(\mathbf{z})||p(\mathbf{z})) = \frac{1}{\alpha - 1} \log \int q(\mathbf{z})^{\alpha} p(\mathbf{z})^{1-\alpha} d\mathbf{z}$$

- $D_{\alpha} \rightarrow D_{KL}$ as $\alpha \rightarrow 1$

Renyi divergence (1961):

$$D_{\alpha}(q(\mathbf{z})||p(\mathbf{z})) = \frac{1}{\alpha - 1} \log \int q(\mathbf{z})^{\alpha} p(\mathbf{z})^{1-\alpha} d\mathbf{z}$$



- Smaller α , less trust on prior
- Larger α , more trust on prior



Renyi neural processes (RNP)

$$-\mathcal{L}_{RNP}(\theta, \phi) = \frac{1}{1-\alpha} \log \mathbb{E}_{q_\phi(\mathbf{z}|C, T)} \left[\frac{p_\theta(Y_T|X_T, \mathbf{z}) q_\phi(\mathbf{z}|C)}{q_\phi(\mathbf{z}|C, T)} \right]^{1-\alpha}$$

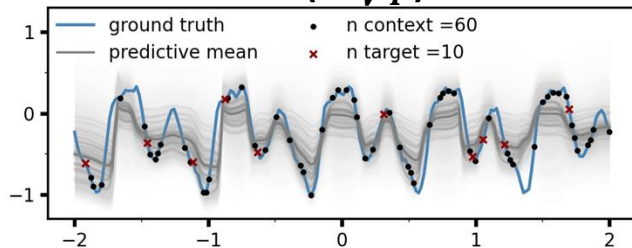
Renyi neural processes (RNP)

$$-\mathcal{L}_{RNP}(\theta, \phi) = \frac{1}{1-\alpha} \log \mathbb{E}_{q_\phi(\mathbf{z}|C, T)} \left[\frac{p_\theta(Y_T|X_T, \mathbf{z}) q_\phi(\mathbf{z}|C)}{q_\phi(\mathbf{z}|C, T)} \right]^{1-\alpha}$$

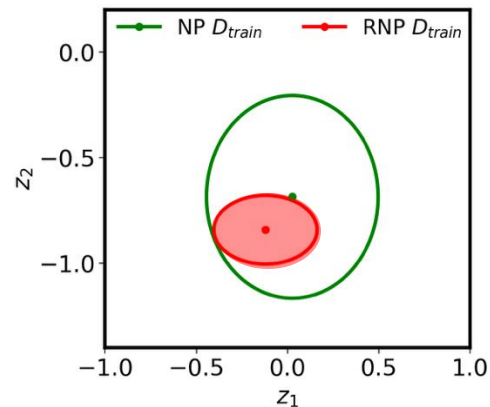
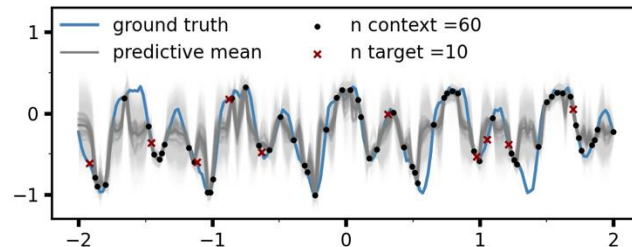
- $\alpha = 0$, $\mathcal{L}_{RNP} = \mathcal{L}_{ML}$ (maximizing likelihood estimation)
- $\alpha \rightarrow 1$, $\mathcal{L}_{RNP} = \mathcal{L}_{VI}$ (variational inference)
- $\alpha \in (0, 1)$, RNP mitigates the prior misspecification via α

Vanilla NP vs RNP on a GP-regression benchmark

NP ($\mathcal{L}_{VI}$)

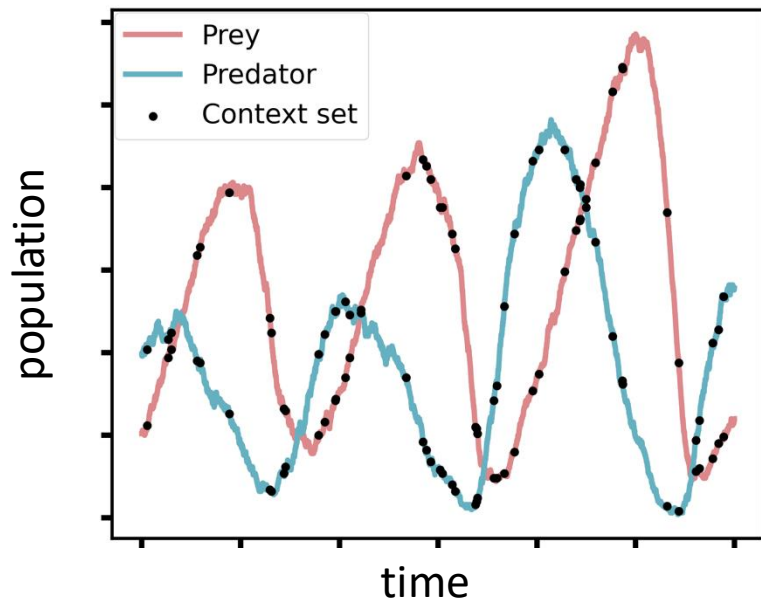


RNP ($\alpha = 0.7$)



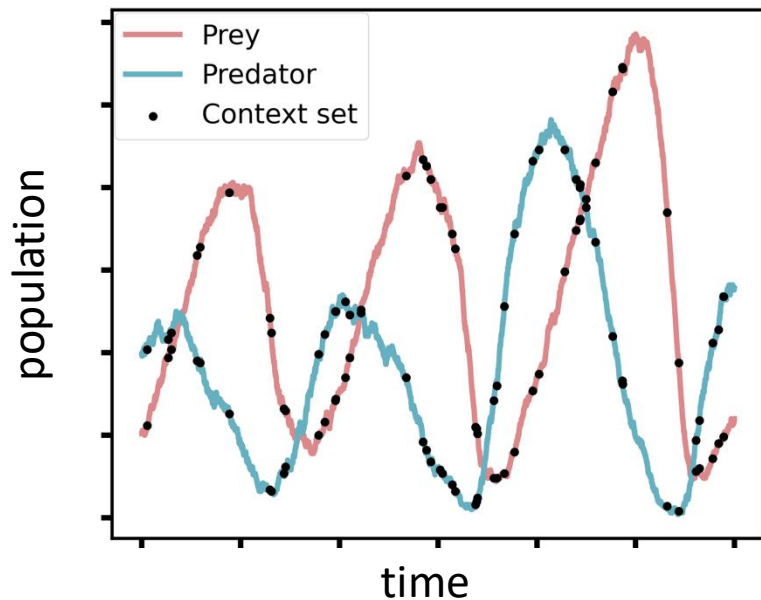
Better variance estimation

Simulation to real-world example

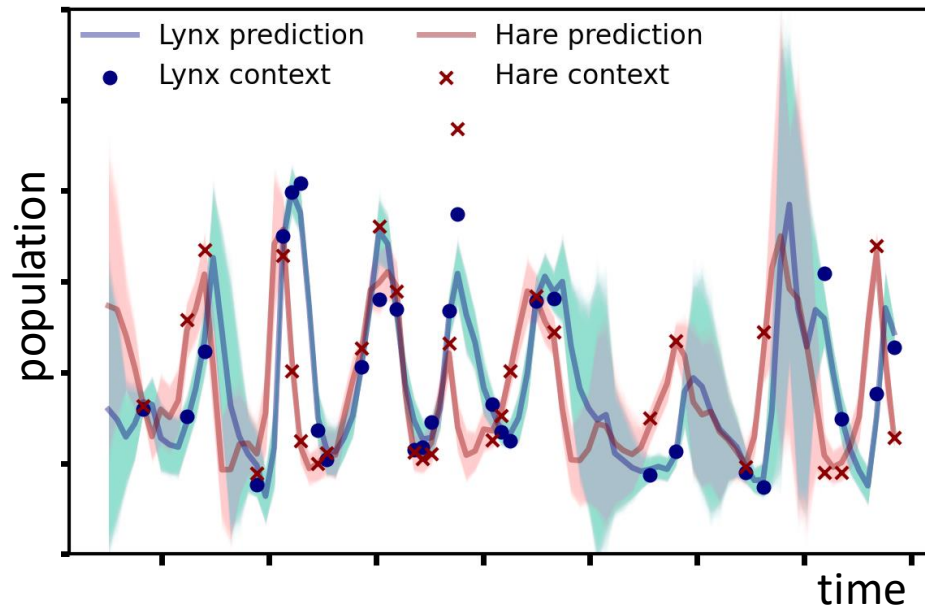


Sim: Prey-predator data D_{train}

Simulation to real-world example



Sim: **Prey**-**predator** data D_{train}

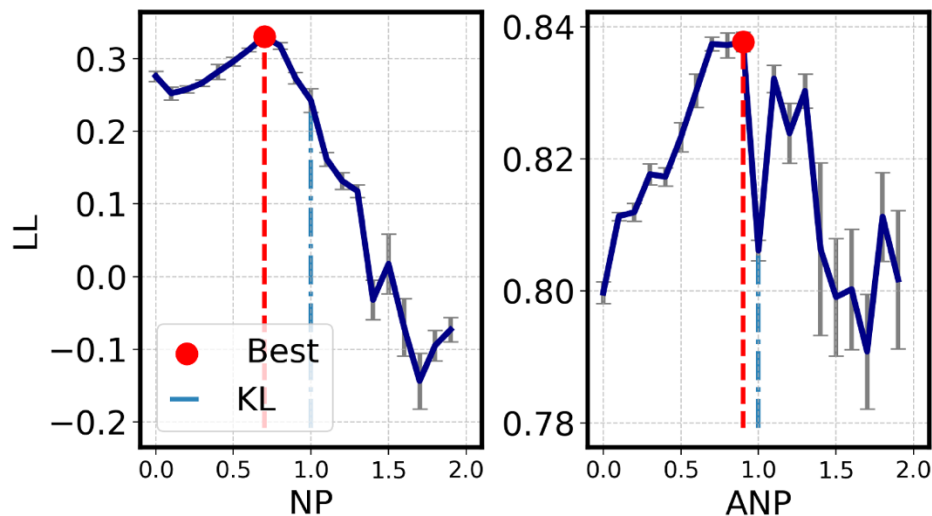


Real: **Hare**  - **lynx**  data D_{test}

$$LL_{RNP} = -3.63 \pm 0.09 \uparrow$$

$$LL_{NP} = -4.44 \pm 0.41$$

GP regression



- $\alpha = 0.7$ empirically generalizes well
- Optimal α can be found using cross-validation
- Heuristic: annealing α from 1 to 0



Renyi neural processes:

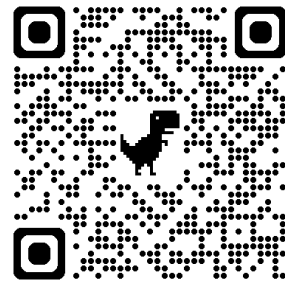
Our contributions:

- Identification of prior misspecification in neural processes
- Renyi divergence to mitigate prior misspecification
- Unification of two NP objectives with RNP

Poster session 5 (ID [43943](#)) (11 a.m. — 1:30 p.m. PDT)

Email: xuesong.wang@data61.csiro.au

Paper



Github

