

Modern ML for particle physics

0. Plan of attack

1. Bayesian neural networks - (errors and regularization)
2. Generative Models I: GANs + applications + VAE
3. Generative Models II: Normalizing flow + applications + (diffusion model)
4. Anomaly detection

I. Bayesian Neural Networks (nothing really bayesian about them)

So far: NNs are deterministic $\rightarrow$ one input $\rightarrow$ one output

However: in physics we are not only interested in results, but also in errors

Main idea:

deterministic weights $\theta \rightarrow$ distribution $\theta \sim p(\theta|x)$

1 input $\rightarrow$ distribution of output $\rightarrow \langle y \rangle + \sigma_y$
 $p(y|x)$

I. 1. Basics

Now, we are interested to evaluate $p(y|x)$, $y \approx f(x) \langle y(x) \rangle$
which is given by:

$$p(y|x) = \int d\theta \, p(y|\theta, x) p(\theta|x)$$

$\uparrow$ require knowledge of!

$\Rightarrow$ Variational approximation:

$$p(\theta|x) \simeq q_{\alpha}(\theta|x) \stackrel{\text{simple}}{=} q_{\alpha}(\theta)$$

$\uparrow$ parameter to adjust to match $p(\theta|x)$

KL:

$$\arg \min_{\alpha} KL[q_{\alpha}(\theta), p(\theta|x)]$$

②

where KL is given by:

$$KL[q_\lambda(\theta), p(\theta|x)] = \int d\theta q_\lambda(\theta) \log \frac{q_\lambda(\theta)}{p(\theta|x)}$$

Baye's Theorem: $p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$

Annotations: $p(x|\theta)$ is likelihood, $p(\theta)$ is prior, $p(x)$ is normalization, $p(\theta|x)$ is posterior.

$$= \int d\theta q_\lambda(\theta) \log \frac{q_\lambda(\theta) p(x)}{p(\theta) p(x|\theta)}$$

$$= KL[q_\lambda(\theta), p(\theta)] - \int d\theta q_\lambda(\theta) \cdot \log p(x|\theta) + \underbrace{\log p(x)}_{\text{intractable!}} \underbrace{\int d\theta q_\lambda(\theta)}_{=1}$$

From this, we can define:

① + ② = -ELBO $\rightarrow$ see later!

$$\mathcal{L}_{BNN} = - \int d\theta q_\lambda(\theta) \log p(x|\theta) + KL[q_\lambda(\theta), p(\theta)]$$

① neg-log likelihood averaged over $q_\lambda(\theta)$ $\leftarrow$ in practice see later!

② $q_\lambda(\theta)$ should not deviate too much from prior!

In practice: p & q are gaussian $\alpha \in \{\mu, \sigma\}$

$$q_\lambda(\theta) = q_{\mu\sigma}(\theta) = \frac{1}{\sqrt{2\pi\sigma_q^2}} \cdot e^{-\frac{(\theta - \mu_q)^2}{2\sigma_q^2}}$$

θ : mean μ_q and width σ_q

$\rightarrow$ same for $p(\theta) \sim$ usually $\mu_p = 0$

$$\rightarrow KL[q_\lambda(\theta), p(\theta)] = \frac{\sigma_q^2 - \sigma_p^2}{2\sigma_p^2} + \frac{\mu_q^2}{2\sigma_p^2} + \log \frac{\sigma_p}{\sigma_q}$$

$\rightarrow$ ② is regularization!

Deterministic limit:

$$\left. \begin{array}{l} \sigma_q \rightarrow 0 \\ \mu_q \rightarrow \theta_0 \end{array} \right\} \rightarrow q(\theta) = \delta(\theta - \theta_0)$$

$$\mathcal{L}_{BNN} = -\log(p(x|\theta_0)) + \underbrace{\frac{\theta_0^2}{2\sigma_p^2}}_{\text{L2 regularization!}} + \text{const}$$

usually λ free parameter!

Comments:

for K parameters of NN use product:

$$q_k(\theta) = \prod_{i=1}^k N(\theta_i | \mu_i, \sigma_i)$$

$\Rightarrow$ BNN has $2K$ parameters (Mean-field approx. [1601.00670])

with correlations $\rightarrow$ BNN $\sim K^2$ params $\rightarrow$ bad scaling

$\hookrightarrow$ deepness of NN compensates for approx?

$\rightarrow$ open question!

I.2. Example: Bayesian regression [2206.14831]

we are interested in fitting amplitudes A

training data: $\{x, A(x)\}$

We can define $p(A|x) \sim$ prob of A @ pos x ($\sim p(y|x)$)

then we can define:

$$\langle A \rangle = \int dA A p(A|x) \quad \text{with} \quad p(A|x) = \int d\theta p(A|\theta, x) p(\theta|x)$$

Now, using our approximation $p(\theta|x) \approx q_k(\theta)$

Now we want to describe this as mean and σ of prediction so we write:

$$\begin{aligned} (*) \quad \langle A \rangle &= \int dA A p(A|x) \quad \text{network output} \\ &= \int dA \int d\theta A p(A|\theta, x) q_k(\theta) \\ &= \int d\theta q_k(\theta) \underbrace{\int dA p(A|x, \theta) A}_{\equiv \bar{A}(x, \theta)} \\ &= \int d\theta q_k(\theta) \bar{A}(x, \theta) \quad \leftarrow \text{first network output!} \end{aligned}$$

(4)

Variance of A:

$$\begin{aligned}
\sigma_{\text{tot}}^2 &= \langle (A - \langle A \rangle)^2 \rangle \\
&= \int dA (A - \langle A \rangle)^2 p(A|x) \\
&= \int dA \int d\theta (A - \langle A \rangle)^2 q_\theta(\theta) \cdot p(A|x, \theta) \\
&= \int d\theta q_\theta(\theta) \left[\int dA A^2 p(A|x, \theta) - 2\langle A \rangle \int dA A p(A|x, \theta) + \langle A \rangle^2 \int dA p(A|x, \theta) \right] \\
&= \int d\theta q_\theta(\theta) \left[\underbrace{\bar{A}^2(x, \theta) - \bar{A}(x, \theta)^2}_{\sigma_{\text{model}}^2 \equiv \sigma_{\text{stoch}}^2} + \underbrace{(\bar{A}(x, \theta) - \langle A \rangle)^2}_{\sigma_{\text{pred}}^2} \right]
\end{aligned}$$

De Pires two kind of uncertainties:

1.

$$\sigma_{\text{pred}}^2 = \int d\theta q_\theta(\theta) [\bar{A}(x, \theta) - \langle A \rangle]^2$$

↳ following (*) this vanishes for $q(\theta) \rightarrow \delta(\theta - \theta_0)$ ↳ with precise training data $\rightarrow$ better training

↳ decreases with more data!

↳ represents statistical uncertainty (epistemic)

2.

$$\begin{aligned}
\sigma_{\text{stoch}}^2 &\equiv \langle \sigma_{\text{stoch}}^2(\theta) \rangle \quad (**) \\
&= \int d\theta q_\theta(\theta) \sigma_{\text{stoch}}^2(\theta) \\
&= \int d\theta q_\theta(\theta) (\bar{A}^2(x, \theta) - \bar{A}(x, \theta)^2) \equiv \langle \sigma_{\text{model}}^2(\theta) \rangle_\theta
\end{aligned}$$

↳ already occurs without sampling

↳ does not vanish for $q(\theta) \rightarrow \delta(\theta - \theta_0)$

↳ reaches constant

↳ represents systematic uncertainty (aleatoric)

(stoch training data, bad hyper, model expressivity)

(*) and (**) are network outputs:

$$\text{BNN: } x, \theta \rightarrow \begin{pmatrix} \bar{A}(\theta) \\ \sigma_{\text{model}}(\theta) \end{pmatrix} \quad \text{in practice: } \log \sigma^2$$

We obtain now the outputs as:

$$\langle A \rangle \approx \frac{1}{L} \sum_{i=1}^L \bar{A}(\theta_i) \quad \sigma_{\text{stoch}}^2 \approx \frac{1}{L} \sum_{i=1}^L \sigma_{\text{stoch}}^2(\theta_i)$$

$$\sigma_{\text{pred}}^2 \approx \frac{1}{L} \sum_{i=1}^L (\bar{A}(\theta_i) - \langle A \rangle)^2 \quad \text{with } \theta_i \sim q_{\lambda}(\theta)$$

usually $L \approx 10$ for evaluating and getting errorbars!

Likely we do not expect $\bar{A}(x, \theta)$ to be gaussian we can assume this for $\sigma_{\text{model}}(\theta)$. (Simplification) ($q_{\lambda}(\theta)$ is gaussian)!

$$\begin{aligned} \mathcal{L}_{\text{BNN}} = \int d\theta \, q_{\lambda}(\theta) \sum_j & \left[\frac{\bar{A}_j(\theta) - A_j^{\text{truth}}}{2\sigma_{\text{model}}^2(x_j, \theta)} + \log \sigma_{\text{model}} \right] \\ & + \underbrace{\frac{\sigma_q^2 - \sigma_p^2 + (\mu_q - \mu_p)^2}{2\sigma_p^2} + \log \frac{\sigma_q}{\sigma_p}}_{\text{KL}(q_{\lambda}(\theta), p(\theta))} \end{aligned}$$

→ minimize μ_q and σ_q

→ data has no uncertainty → NN constructs one!

w/o sampling θ we obtain:

$$\mathcal{L}_{\text{heteroskedastic}} = \sum_j \left[\frac{(\bar{A}_j(\theta_0) - A_j^{\text{truth}})^2}{2\sigma_{\text{model},j}^2(\theta_0)} + \log \sigma_{\text{model},j}(\theta_0) \right] + \frac{\sigma_0^2}{2\sigma_p^2}$$

→ show plots from paper → check-out for yourself!

↳ more details in thesis of Michel Luchmann

[Doi: [10.11588/hidok.00032414](https://doi.org/10.11588/hidok.00032414)]