

$$\min_u \int_0^T L(u(t), \dot{u}(t)) dt$$

$$L(u, v) = \frac{1}{2} |v - b(u)|^2$$

$$u(0) = u_0$$

$$u(T) = u_1$$

Standard approach : solve E.L eq.
 ↳ HAM, GHAM, etc.

Machine learning :

- represent solution by neural network
- view objective fun as loss for network parameters
- optimize parameters by stochastic gradient descent

Here : solution = $\underset{\substack{\text{funct of time}}}{u(t)}$ or $\underset{\substack{\text{funct of space}}}{u(t, x)}$

ex 1) $\int_0^T dt \int_{\Omega} | \partial_t u - \Delta u - \underbrace{d|\nabla u|^2}_{\text{non gradient}} - f(u) |^2 dx dt$

2) $\int_0^T dt \int_0^1 dx | \partial_t u - \partial_x^2 u - \underbrace{a \partial_x u}_{\text{shear non gradient}} - f(u) |^2 dx dt$

...

$u(t, x) \longrightarrow u^\theta(t, x)$ = neural network
 θ = parameters.

$$u^\theta(t, x) = \sigma_L \left(\dots \sigma_2 \left(W_2 \left(\sigma_1 \left(W_1 \begin{pmatrix} t \\ x \end{pmatrix} + b_1 \right) + b_2 \right) \right) \dots \right)$$

where: $\theta = (W_1, b_1, W_2, b_2, \dots)$

with: W_j = weight matrices eg $W_1 \in \mathbb{R}^{p \times (1+d)}$

b_j = bias vectors

σ_j = non linearities $\sigma_1: \mathbb{R}^p \rightarrow \mathbb{R}^q$ eg ReLU

$$\Rightarrow \int_0^T L(u^\theta(t), \dot{u}^\theta(t)) dt = L(\theta) \quad \text{Loss}$$

$$\approx \frac{1}{N} \sum_{j=1}^N L(u^\theta(t_j), \dot{u}^\theta(t_j)) = L^N(\theta) \quad \text{empirical loss}$$

$$t_j \sim \text{iid } U(0, T).$$

$$\text{also } x_j \sim \text{iid } U(\Omega) \quad \text{if space needed}$$



optimization

SGD.

PinN loss.

physics informed neural network.

$$\text{repeat: } \theta^{n+1} = \theta^n - \eta^n \nabla_{\theta} L(\theta^n) \quad n=0, 1, \dots$$

↪ learning rate.

until convergence.



$u^\theta(t, x)$ = analytical fn of (t, x, θ)

all derivatives computed exactly

approx? # of parameter θ is finite.

More ambitious calculations with ML?

SDE: $dX_t = b(X_t) dt + \sigma dw_t$

FPE: $\partial_t p_t = -\nabla \cdot (b p_t) + \frac{\sigma^2}{2} \Delta p_t$

let $p_{-t} = p_t^R$ time reversed

$$\partial_t p_t^R = + \nabla \cdot (b p_t^R) - \frac{\sigma^2}{2} \Delta p_t^R$$

$$= \nabla \cdot \left[\underbrace{\left(b - \sigma^2 \nabla \log p_t^R \right)}_{-b^R(x)} p_t^R \right] + \frac{\sigma^2}{2} \Delta p_t^R$$

~

$-b^R(x)$ time reversed drift

time reversed SDE:

$$dX_t^R = b^R(X_t^R) dt + \sigma dw_t$$

time reversal symmetry :

$$b^R(x) = b(x)$$

e.g. if $p_t = p_t^R$ (NESS)

$$\Rightarrow b^R(x) = -b(x) + \sigma^2 \nabla \log p(x)$$

$$= b(x) \quad ?$$

$$\Leftrightarrow b(x) = \frac{1}{2} \sigma^2 \nabla \log p(x)$$

ex:

$$b(x) = -\nabla U(x) \quad \frac{\sigma^2}{2} = k_B T = \beta^{-1}$$

$$\Rightarrow p(x) = Z^{-1} e^{-\beta U(x)}$$

$$\frac{\sigma^2}{2} \nabla \log p(x) = -\beta^{-1} \beta \nabla U(x) = -\nabla U(x)$$

In general

$$b(x) \neq b^R(x)$$

breakup of TRS.

Can we check TRS?

Can we estimate $\nabla \log p(x)$?

useful also in Stochastic Thermodynamics

on NESS.

$$\nabla \log p(x) = \arg \min_s \int_{\Sigma} |s(x) - \nabla \log p(x)|^2 p(x) dx$$

Fisher divergence

$$\begin{aligned} & \int_{\Sigma} |s(x) - \nabla \log p(x)|^2 p(x) dx \\ &= \int_{\Sigma} |s(x)|^2 p(x) dx - 2 \int_{\Sigma} s(x) \cdot \underbrace{\nabla \log p(x) p(x)}_{= \nabla p(x)} dx + \int_{\Sigma} |\nabla \log p(x)|^2 p(x) dx \\ & \quad + 2 \int_{\Sigma} \nabla \cdot s(x) p(x) dx \end{aligned}$$

" csh (no $s(x)$)

$\Rightarrow$

$$\nabla \log p(x) = \underset{s}{\operatorname{argmin}} \left(\int_{\Omega} (|s(x)|^2 + 2 \nabla \cdot s(x)) p(x) dx \right)$$

$$\approx \frac{1}{N} \sum_{j=1}^N [|s(x_j)|^2 + 2 \nabla \cdot s(x_j)]$$

$\Downarrow$ x_j drawn from $p(x)$

Empirical loss for $s(x)$

!

here

$s: \mathbb{R}^d \rightarrow \mathbb{R}^d$

high dim. fct.

$\hookrightarrow$

use power of approx of neural networks