

From the theory of (stochastic) control to deep learning and back - Part 3

Lukasz Szpruch
University of Edinburgh, The Alan Turing Institute, London

Gradient flow perspective on control

We work with

- ▶ One hidden layer networks via gradient flow on $(\mathcal{P}_2(\mathbb{R}^d), \mathcal{W}_2)$

Gradient flow perspective on control

We work with

- ▶ One hidden layer networks via gradient flow on $(\mathcal{P}_2(\mathbb{R}^d), \mathcal{W}_2)$
- ▶ Neural ODEs via gradient flow on $(\mathcal{V}_2, \mathcal{W}_2^T(\mu, \nu))$, where

$$\mathcal{W}_q^T(\mu, \nu) := \left(\int_0^T \mathcal{W}_q(\mu_t, \nu_t)^q dt \right)^{1/q}.$$

$$\mathcal{V}_2 := \left\{ \nu : \nu(dt, da) = \nu_t(da)dt, \nu_t \in \mathcal{P}_2(\mathbb{R}^p), \int_0^T \int |a|^2 \nu_t(da) dt < \infty \right\},$$

Gradient flow perspective on control

We work with

- ▶ One hidden layer networks via gradient flow on $(\mathcal{P}_2(\mathbb{R}^d), \mathcal{W}_2)$
- ▶ Neural ODEs via gradient flow on $(\mathcal{V}_2, \mathcal{W}_2^T(\mu, \nu))$, where

$$\mathcal{W}_q^T(\mu, \nu) := \left(\int_0^T \mathcal{W}_q(\mu_t, \nu_t)^q dt \right)^{1/q}.$$

$$\mathcal{V}_2 := \left\{ \nu : \nu(dt, da) = \nu_t(da)dt, \nu_t \in \mathcal{P}_2(\mathbb{R}^p), \int_0^T \int |a|^2 \nu_t(da) dt < \infty \right\},$$

- ▶ Neural SDEs via gradient flow on $(\mathcal{V}_q^W, \rho_q)$, where

$$\rho_q(\mu, \mu') = \left(\mathbb{E}^W \left[|\mathcal{W}_q^T(\mu, \mu')|^q \right] \right)^{1/q}$$

$$\mathcal{V}_q^W := \left\{ \nu : \Omega^W \rightarrow \mathcal{M}_q : \mathbb{E}^W \int_0^T \int |a|^q \nu_t(da, dt) < \infty \text{ and } \nu_t \in \mathcal{F}_t^W, \forall t \in [0, T] \right\}$$

Gradient flow perspective on control

We work with

- ▶ One hidden layer networks via gradient flow on $(\mathcal{P}_2(\mathbb{R}^d), \mathcal{W}_2)$
- ▶ Neural ODEs via gradient flow on $(\mathcal{V}_2, \mathcal{W}_2^T(\mu, \nu))$, where

$$\mathcal{W}_q^T(\mu, \nu) := \left(\int_0^T \mathcal{W}_q(\mu_t, \nu_t)^q dt \right)^{1/q}.$$

$$\mathcal{V}_2 := \left\{ \nu : \nu(dt, da) = \nu_t(da)dt, \nu_t \in \mathcal{P}_2(\mathbb{R}^p), \int_0^T \int |a|^2 \nu_t(da)dt < \infty \right\},$$

- ▶ Neural SDEs via gradient flow on $(\mathcal{V}_q^W, \rho_q)$, where

$$\rho_q(\mu, \mu') = \left(\mathbb{E}^W \left[|\mathcal{W}_q^T(\mu, \mu')|^q \right] \right)^{1/q}$$

$$\mathcal{V}_q^W := \left\{ \nu : \Omega^W \rightarrow \mathcal{M}_q : \mathbb{E}^W \int_0^T \int |a|^q \nu_t(da, dt) < \infty \text{ and } \nu_t \in \mathcal{F}_t^W, \forall t \in [0, T] \right\}$$

- ▶ $(\mathcal{V}_q^W, \rho_q)$ is complete.

Gradient Flows for Regularized Stochastic Control Problems

Stochastic Control

For $\xi \in \mathbb{R}^d$ and $\mu \in \mathcal{V}_q^W$, consider the controlled process

$$X_t(\mu) = \xi + \int_0^t \Phi_r(X_r(\mu), \mu_r) dr + \int_0^t \Gamma_r(X_r(\mu), \mu_r) dW_r, \quad t \in [0, T].$$

Stochastic Control

For $\xi \in \mathbb{R}^d$ and $\mu \in \mathcal{V}_q^W$, consider the controlled process

$$X_t(\mu) = \xi + \int_0^t \Phi_r(X_r(\mu), \mu_r) dr + \int_0^t \Gamma_r(X_r(\mu), \mu_r) dW_r, \quad t \in [0, T].$$

Given F and g we define the objective functional

$$J^\sigma(\nu, \xi) := \mathbb{E}^W \left[\int_0^T \left[F_t(X_t(\nu), \nu_t) + \frac{\sigma^2}{2} \text{Ent}(\nu_t) \right] dt + g(X_T(\nu)) \middle| X_0(\nu) = \xi \right].$$

$$\text{Ent}(m) := \begin{cases} \int_{\mathbb{R}^d} m(x) \log \left(\frac{m(x)}{g(x)} \right) dx & \text{if } m \text{ is a.c. w.r.t. Lebesgue measure} \\ \infty & \text{otherwise} \end{cases}$$

and Gibbs measure g :

$$g(x) = e^{-U(x)} \quad \text{with } U \text{ s.t. } \int_{\mathbb{R}^d} e^{-U(x)} dx = 1.$$

Stochastic Control

For $\xi \in \mathbb{R}^d$ and $\mu \in \mathcal{V}_q^W$, consider the controlled process

$$X_t(\mu) = \xi + \int_0^t \Phi_r(X_r(\mu), \mu_r) dr + \int_0^t \Gamma_r(X_r(\mu), \mu_r) dW_r, \quad t \in [0, T].$$

Given F and g we define the objective functional

$$J^\sigma(\nu, \xi) := \mathbb{E}^W \left[\int_0^T \left[F_t(X_t(\nu), \nu_t) + \frac{\sigma^2}{2} \text{Ent}(\nu_t) \right] dt + g(X_T(\nu)) \mid X_0(\nu) = \xi \right].$$

$$\text{Ent}(m) := \begin{cases} \int_{\mathbb{R}^d} m(x) \log \left(\frac{m(x)}{g(x)} \right) dx & \text{if } m \text{ is a.c. w.r.t. Lebesgue measure} \\ \infty & \text{otherwise} \end{cases}$$

and Gibbs measure g :

$$g(x) = e^{-U(x)} \quad \text{with } U \text{ s.t. } \int_{\mathbb{R}^d} e^{-U(x)} dx = 1.$$

Example - Relaxed Control

$$\Phi_t(x, m) = \int \phi_t(x, a) m(da), \quad \text{and} \quad \Gamma_t(x, m)(\Gamma_t(x, m))^\top = \int \gamma_t(x, a) \gamma_t(x, a)^\top m(da)$$

Stochastic Control

- ▶ DPP and HJB equation
 - ▶ Solve nonlinear PDE or corresponding (2)BSDE
 - ▶ 'Linearise' with policy or value iteration and solve linear PDE [Kerimkulov et al., 2020, Gobet and Labart, 2010]
- ▶ Maximum principle
 - ▶ Solve the corresponding non-Markov FBSDE.

Stochastic Control

- ▶ DPP and HJB equation
 - ▶ Solve nonlinear PDE or corresponding (2)BSDE
 - ▶ 'Linearise' with policy or value iteration and solve linear PDE [Kerimkulov et al., 2020, Gobet and Labart, 2010]
- ▶ Maximum principle
 - ▶ Solve the corresponding non-Markov FBSDE.

Motivation to study gradient flow solution for stochastic control problem

- ▶ Aiming at solving high-dimensional control problems

Stochastic Control

- ▶ DPP and HJB equation
 - ▶ Solve nonlinear PDE or corresponding (2)BSDE
 - ▶ 'Linearise' with policy or value iteration and solve linear PDE [Kerimkulov et al., 2020, Gobet and Labart, 2010]
- ▶ Maximum principle
 - ▶ Solve the corresponding non-Markov FBSDE.

Motivation to study gradient flow solution for stochastic control problem

- ▶ Aiming at solving high-dimensional control problems
- ▶ Building algorithms for adaptive stochastic control (data-driven problems)

Stochastic Control

- ▶ DPP and HJB equation
 - ▶ Solve nonlinear PDE or corresponding (2)BSDE
 - ▶ 'Linearise' with policy or value iteration and solve linear PDE [Kerimkulov et al., 2020, Gobet and Labart, 2010]
- ▶ Maximum principle
 - ▶ Solve the corresponding non-Markov FBSDE.

Motivation to study gradient flow solution for stochastic control problem

- ▶ Aiming at solving high-dimensional control problems
- ▶ Building algorithms for adaptive stochastic control (data-driven problems)
- ▶ Bridging the gap between stochastic control and Reinforcement learning.

Stochastic Control

Let us now introduce the Hamiltonian

$$H_t^\sigma(x, y, z, m) := \Phi_t(x, m)y + \text{tr}(\Gamma_t^\top(x, m)z) + F_t(x, m) + \frac{\sigma^2}{2}\text{Ent}(m) .$$

Stochastic Control

Let us now introduce the Hamiltonian

$$H_t^\sigma(x, y, z, m) := \Phi_t(x, m)y + \text{tr}(\Gamma_t^\top(x, m)z) + F_t(x, m) + \frac{\sigma^2}{2}\text{Ent}(m).$$

We will also use the adjoint process

$$\begin{aligned} dY_t(\mu) &= -(\nabla_x H_t^0)(X_t(\mu), Y_t(\mu), Z_t(\mu), \mu_t) dt + Z_t(\mu) dW_t, \quad t \in [0, T], \\ Y_T(\mu) &= (\nabla_x g)(X_T(\mu)) \end{aligned}$$

Theorem 1 (Necessary condition for optimality)

Fix $\sigma > 0$. Fix $q > 2$. If $\nu \in \mathcal{V}_q^W$ is (locally) optimal for $J^\sigma(\cdot, \xi)$, $X(\nu)$ and $Y(\nu)$, $Z(\nu)$ are the associated optimally controlled state and adjoint processes respectively, then for any other $\mu \in \mathcal{V}_q^W$ it holds that

i)

$$\int \frac{\delta H^\sigma}{\delta m}(X_t(\nu), Y_t(\nu), Z_t(\nu), \nu_t, a)(\mu_t - \nu_t)(da) \geq 0 \text{ a.a. } (\omega, t) \in \Omega^W \times (0, T)$$

ii) For a.a. $(\omega, t) \in \Omega^W \times (0, T)$ there exists $\varepsilon > 0$ (small and depending on μ_t) such that

$$H^\sigma(X_t(\nu), Y_t(\nu), Z_t(\nu), \nu_t + \varepsilon(\mu_t - \nu_t)) \geq H^\sigma(X_t(\nu), Y_t(\nu), Z_t(\nu), \nu_t).$$

In other words, the optimal relaxed control $\nu \in \mathcal{V}_q^W$ locally minimizes the Hamiltonian.

Let

$$\frac{\delta \mathbf{H}_t^\sigma}{\delta m}(a, \mu) = \frac{\delta \mathbf{H}_t^0}{\delta m}(a, \mu) + \frac{\sigma^2}{2} (U(a) + \log \mu_t(a) + 1) .$$

Let

$$\frac{\delta \mathbf{H}_t^\sigma}{\delta m}(a, \mu) = \frac{\delta \mathbf{H}_t^0}{\delta m}(a, \mu) + \frac{\sigma^2}{2} (U(a) + \log \mu_t(a) + 1) .$$

Consider

$$d\theta_{s,t} = - \left((\nabla_a \frac{\delta H_t^0}{\delta m})(X_{s,t}, Y_{s,t}, Z_{s,t}, \theta_{s,t}) + \frac{\sigma^2}{2} (\nabla_a U)(\theta_{s,t}) \right) ds + \sigma dB_s, \quad s \geq 0,$$

where

$$\begin{cases} \nu_{s,t} &= \mathcal{L}(\theta_{s,t} | \mathcal{F}_t^W), \\ X_{s,t} &= \xi + \int_0^t \Phi_r(X_{s,r}, \nu_{s,r}) dr + \int_0^t \Gamma_r(X_{s,r}, \nu_{s,r}(da)) dW_r, \quad t \in [0, T], \\ dY_{s,t} &= -(\nabla_x H_t^0)(X_{s,t}, Y_{s,t}, Z_{s,t}, \nu_{s,t}) dt + Z_{s,t} dW_t, \\ Y_{s,T} &= (\nabla_x g)(X_T). \end{cases}$$

Let

$$\frac{\delta \mathbf{H}_t^\sigma}{\delta m}(a, \mu) = \frac{\delta \mathbf{H}_t^0}{\delta m}(a, \mu) + \frac{\sigma^2}{2} (U(a) + \log \mu_t(a) + 1) .$$

Consider

$$d\theta_{s,t} = - \left((\nabla_a \frac{\delta H_t^0}{\delta m})(X_{s,t}, Y_{s,t}, Z_{s,t}, \theta_{s,t}) + \frac{\sigma^2}{2} (\nabla_a U)(\theta_{s,t}) \right) ds + \sigma dB_s, \quad s \geq 0,$$

where

$$\begin{cases} \nu_{s,t} &= \mathcal{L}(\theta_{s,t} | \mathcal{F}_t^W), \\ X_{s,t} &= \xi + \int_0^t \Phi_r(X_{s,r}, \nu_{s,r}) dr + \int_0^t \Gamma_r(X_{s,r}, \nu_{s,r}(da)) dW_r, \quad t \in [0, T], \\ dY_{s,t} &= -(\nabla_x H_t^0)(X_{s,t}, Y_{s,t}, Z_{s,t}, \nu_{s,t}) dt + Z_{s,t} dW_t, \\ Y_{s,T} &= (\nabla_x g)(X_T). \end{cases}$$

We have

$$\frac{d}{ds} J(\nu_{s,\cdot}) = -\mathbb{E}^W \int_0^T \left[\int \left| (\nabla_a \frac{\delta \mathbf{H}_t^\sigma}{\delta m})(\cdot, \nu_{s,t}) \right|^2 \nu_{s,t}(da) \right] dt \leq 0.$$

Assumption 1

Let $\nabla_a U$ be Lipschitz continuous in a and let there be $\kappa > 0$ such that:

$$(\nabla_a U(a') - \nabla_a U(a)) \cdot (a' - a) \geq \kappa |a' - a|^2, \quad a, a' \in \mathbb{R}^p.$$

Assumption 1

Let $\nabla_a U$ be Lipschitz continuous in a and let there be $\kappa > 0$ such that:

$$(\nabla_a U(a') - \nabla_a U(a)) \cdot (a' - a) \geq \kappa |a' - a|^2, \quad a, a' \in \mathbb{R}^p.$$

Assumption 2

Assume that there exists $\eta_1, \eta_2 \in \mathbb{R}$, $\bar{\eta} \in L^{q/2}(\Omega^W \times (0, T); \mathbb{R})$ and $\mathcal{E} : \mathcal{V}_q^W \times \mathcal{V}_q^W \rightarrow [0, \infty)$ s.t for any $a \in \mathbb{R}^p$, $\mu \in \mathcal{V}_2^W$, $t \in [0, T]$ we have

$$\left(\nabla_a \frac{\delta \mathbf{H}_t^0}{\delta m} \right) (a, \mu) a \geq \eta_1 |a|^2 - \eta_2 \mathcal{E}_t(\mu, \delta_0)^2 - \bar{\eta}_t$$

and for all $\mu, \mu' \in \mathcal{V}_q^W$ we have $\mathbb{E}^W \left[\int_0^T \mathcal{E}_t(\mu, \mu')^q dt \right] \leq \rho_q(\mu, \mu')^q$.

Assumption 1

Let $\nabla_a U$ be Lipschitz continuous in a and let there be $\kappa > 0$ such that:

$$(\nabla_a U(a') - \nabla_a U(a)) \cdot (a' - a) \geq \kappa |a' - a|^2, \quad a, a' \in \mathbb{R}^p.$$

Assumption 2

Assume that there exists $\eta_1, \eta_2 \in \mathbb{R}$, $\bar{\eta} \in L^{q/2}(\Omega^W \times (0, T); \mathbb{R})$ and $\mathcal{E} : \mathcal{V}_q^W \times \mathcal{V}_q^W \rightarrow [0, \infty)$ s.t for any $a \in \mathbb{R}^p$, $\mu \in \mathcal{V}_2^W$, $t \in [0, T]$ we have

$$\left(\nabla_a \frac{\delta \mathbf{H}_t^0}{\delta m} \right) (a, \mu) a \geq \eta_1 |a|^2 - \eta_2 \mathcal{E}_t(\mu, \delta_0)^2 - \bar{\eta}_t$$

and for all $\mu, \mu' \in \mathcal{V}_q^W$ we have $\mathbb{E}^W \left[\int_0^T \mathcal{E}_t(\mu, \mu')^q dt \right] \leq \rho_q(\mu, \mu')^q$.

Assumption 3

There exists $\eta_1, \eta_2 \in \mathbb{R}$ and $\mathcal{E} : \mathcal{V}_q^W \times \mathcal{V}_q^W \rightarrow [0, \infty)$ s.t all $t \in [0, T]$, for all a, a' and for all $\mu, \mu' \in \mathcal{V}_q^W$ we have $\mathbb{E}^W \left[\int_0^T \mathcal{E}_t(\mu, \mu')^q dt \right] \leq \rho_q(\mu, \mu')^q$ and

$$2 \left(\left(\nabla_a \frac{\delta \mathbf{H}_t^0}{\delta m} \right) (a', \mu') - \left(\nabla_a \frac{\delta \mathbf{H}_t^0}{\delta m} \right) (a, \mu) \right) (a' - a) \geq \eta_1 |a' - a|^2 - \eta_2 \mathcal{E}_t(\mu', \mu)^2.$$

Let $P_s \mu^0 := (P_{s,t} \mu^0)_{t \in [0, T]}$. Moreover note that due to uniqueness we have $P_{s+s'} \mu^0 = P_s (P_{s'} \mu^0)$.

Let $P_s \mu^0 := (P_{s,t} \mu^0)_{t \in [0, T]}$. Moreover note that due to uniqueness we have $P_{s+s'} \mu^0 = P_s (P_{s'} \mu^0)$.

Theorem 2

Let Assumptions 1 and 3 hold. Moreover, assume that $\lambda = \frac{q}{2} (\sigma^2 \kappa + \eta_1 - \eta_2) > 0$. Then there is $\mu^ \in \mathcal{V}_q^W$ such that for any $s \geq 0$ we have $P_s \mu^* = \mu^*$ and μ^* is unique. For any $\mu^0 \in \mathcal{V}_q^W$ we have that*

$$\rho_q(P_s \mu^0, \mu^*) \leq e^{-\frac{1}{q} \lambda s} \rho_q(\mu^0, \mu^*).$$

Theorem 3

Assume that for any $\mu^0 \in \mathcal{V}_q^W$ the MFLD has unique solution $P_s \mu^0$ and that it admits unique invariant measure $\mu^* \in \mathcal{V}_q^W$ such that for any $\mu^0 \in \mathcal{V}_q^W$, $\lim_{s \rightarrow \infty} \rho_q(P_s \mu^0, \mu^*) = 0$. Let

$$\mathcal{I}^\sigma := \left\{ \nu \in \mathcal{V}_q^W : \frac{\delta \mathbf{H}_t^\sigma}{\delta m}(a, \nu) \text{ is constant for a.a. } a, t, \omega^W \right\}. \quad (1)$$

Then

- i) We have $J^\sigma(\mu^*) < \infty$ and $\mathcal{I}^\sigma = \{\mu^*\}$. In other words, μ^* is the only control which satisfies the first order condition in (1).
- ii) The unique minimizer of J^σ is μ^* .

Theorem 3

Assume that for any $\mu^0 \in \mathcal{V}_q^W$ the MFLD has unique solution $P_s \mu^0$ and that it admits unique invariant measure $\mu^* \in \mathcal{V}_q^W$ such that for any $\mu^0 \in \mathcal{V}_q^W$, $\lim_{s \rightarrow \infty} \rho_q(P_s \mu^0, \mu^*) = 0$. Let

$$\mathcal{I}^\sigma := \left\{ \nu \in \mathcal{V}_q^W : \frac{\delta \mathbf{H}_t^\sigma}{\delta m}(a, \nu) \text{ is constant for a.a. } a, t, \omega^W \right\}. \quad (1)$$

Then

- i) We have $J^\sigma(\mu^*) < \infty$ and $\mathcal{I}^\sigma = \{\mu^*\}$. In other words, μ^* is the only control which satisfies the first order condition in (1).
- ii) The unique minimizer of J^σ is μ^* .

Recall that if $\nu \in \mathcal{I}^\sigma$ then

$$\nu_t = Z^{-1} e^{\frac{-2}{\sigma^2} \frac{\delta H_t^0}{\delta m}(X_t, Y_t, Z_t, \nu_t, a)} \gamma(a), \quad Z = \int e^{\frac{-2}{\sigma^2} \frac{\delta H_t^0}{\delta m}(X_t, Y_t, Z_t, \nu_t, a)} \gamma(a) da,$$

Step 1: Show that $I^\sigma = \{\mu^*\}$

Step 1: Show that $I^\sigma = \{\mu^*\}$

Fix $t \in [0, T]$ and $\omega^W \in \Omega^W$.

Step 1: Show that $I^\sigma = \{\mu^*\}$

Fix $t \in [0, T]$ and $\omega^W \in \Omega^W$.

Let $b_{s,t}(a) := (\nabla_a \frac{\delta \mathbf{H}_t^0}{\delta m})(a, \mu_s) + \frac{\sigma^2}{2}(\nabla_a U)(a)$ and $\mu_{s,t} = \mathcal{L}(\theta_{s,t} | \mathcal{F}_t^W)$.

$\mu_{s,t}$ is a solution to

$$\partial_s \mu_{s,t} = \nabla_a \cdot \left(b_{s,t} \mu_{s,t} + \frac{\sigma^2}{2} \nabla_a \mu_{s,t} \right), \quad s \geq 0, \quad \mu_{s,0} = \mu_t^0 := \mathcal{L}(\theta_{0,t} | \mathcal{F}_t^W).$$

Step 1: Show that $I^\sigma = \{\mu^*\}$

Fix $t \in [0, T]$ and $\omega^W \in \Omega^W$.

Let $b_{s,t}(a) := (\nabla_a \frac{\delta \mathbf{H}_t^0}{\delta m})(a, \mu_s) + \frac{\sigma^2}{2}(\nabla_a U)(a)$ and $\mu_{s,t} = \mathcal{L}(\theta_{s,t} | \mathcal{F}_t^W)$.

$\mu_{s,t}$ is a solution to

$$\partial_s \mu_{s,t} = \nabla_a \cdot \left(b_{s,t} \mu_{s,t} + \frac{\sigma^2}{2} \nabla_a \mu_{s,t} \right), \quad s \geq 0, \quad \mu_{s,0} = \mu_t^0 := \mathcal{L}(\theta_{0,t} | \mathcal{F}_t^W).$$

The solution is unique and $P_s \mu^0 = \mu_{s,\cdot}$, so P_s is the solution operator.

Since μ^* is an invariant measure for almost all $t \in [0, T]$ and $\omega^W \in \Omega^W$ we have $(P_s \mu^*)_t = \mu_t^*$ and so $\partial_s \mu_{s,t}^* = 0$.

Since μ^* is an invariant measure for almost all $t \in [0, T]$ and $\omega^W \in \Omega^W$ we have $(P_s \mu^*)_t = \mu_t^*$ and so $\partial_s \mu_{s,t}^* = 0$.

Hence for almost all $t \in [0, T]$ and $\omega^W \in \Omega^W$ we have that μ_t^* is the solution to the stationary Kolmogorov–Fokker–Planck equation

$$0 = \nabla_a \cdot \left(\left((\nabla_a \frac{\delta \mathbf{H}_t^0}{\delta m})(\cdot, \mu^*) + \frac{\sigma^2}{2} (\nabla_a U) \right) \mu_t^* + \frac{\sigma^2}{2} \nabla_a \mu_t^* \right). \quad (2)$$

This implies that $\mu^* \in \mathcal{I}^\sigma$.

Since μ^* is an invariant measure for almost all $t \in [0, T]$ and $\omega^W \in \Omega^W$ we have $(P_s \mu^*)_t = \mu_t^*$ and so $\partial_s \mu_{s,t}^* = 0$.

Hence for almost all $t \in [0, T]$ and $\omega^W \in \Omega^W$ we have that μ_t^* is the solution to the stationary Kolmogorov–Fokker–Planck equation

$$0 = \nabla_a \cdot \left(\left((\nabla_a \frac{\delta \mathbf{H}_t^0}{\delta m})(\cdot, \mu^*) + \frac{\sigma^2}{2} (\nabla_a U) \right) \mu_t^* + \frac{\sigma^2}{2} \nabla_a \mu_t^* \right). \quad (2)$$

This implies that $\mu^* \in \mathcal{I}^\sigma$.

Consider now some $\nu \in \mathcal{I}^\sigma$. Then

$$\nu_t(a) = Z^{-1} e^{-\frac{2}{\sigma^2} \frac{\delta \mathbf{H}_t^0}{\delta m}(a, \nu)} g(a),$$

Since μ^* is an invariant measure for almost all $t \in [0, T]$ and $\omega^W \in \Omega^W$ we have $(P_s \mu^*)_t = \mu_t^*$ and so $\partial_s \mu_{s,t}^* = 0$.

Hence for almost all $t \in [0, T]$ and $\omega^W \in \Omega^W$ we have that μ_t^* is the solution to the stationary Kolmogorov–Fokker–Planck equation

$$0 = \nabla_a \cdot \left(\left((\nabla_a \frac{\delta \mathbf{H}_t^0}{\delta m})(\cdot, \mu_t^*) + \frac{\sigma^2}{2} (\nabla_a U) \right) \mu_t^* + \frac{\sigma^2}{2} \nabla_a \mu_t^* \right). \quad (2)$$

This implies that $\mu^* \in \mathcal{I}^\sigma$.

Consider now some $\nu \in \mathcal{I}^\sigma$. Then

$$\nu_t(a) = Z^{-1} e^{-\frac{2}{\sigma^2} \frac{\delta \mathbf{H}_t^0}{\delta m}(a, \nu)} g(a),$$

We see that almost all $t \in [0, T]$ and $\omega^W \in \Omega^W$ we have that ν_t solves (2). But the solution to (2) is unique and so $\nu = \mu^*$. This proves item i).

Step 2: Show that the unique minimizer of J^σ is μ^* .

Step 2: Show that the unique minimizer of J^σ is μ^* .

We will show by contradiction that μ^* is at least (locally) optimal.
Assume that μ^* is not the (locally) optimal control for J^σ .

Step 2: Show that the unique minimizer of J^σ is μ^* .

We will show by contradiction that μ^* is at least (locally) optimal.
Assume that μ^* is not the (locally) optimal control for J^σ .

Then for some $\mu^0 \in \mathcal{V}_2^W$ it holds that $J^\sigma(\mu^0) < J^\sigma(\mu^*)$.

Step 2: Show that the unique minimizer of J^σ is μ^* .

We will show by contradiction that μ^* is at least (locally) optimal. Assume that μ^* is not the (locally) optimal control for J^σ .

Then for some $\mu^0 \in \mathcal{V}_2^W$ it holds that $J^\sigma(\mu^0) < J^\sigma(\mu^*)$.

On the other we know that $\lim_{s \rightarrow \infty} P_s \mu^0 = \mu^*$. From the lower semi-continuity of J^σ we get

$$\begin{aligned} J^\sigma(\mu^*) - J^\sigma(\mu^0) &\leq \liminf_{s \rightarrow \infty} J^\sigma(P_s \mu^0) - J^\sigma(\mu^0) \\ &= - \liminf_{s \rightarrow \infty} \int_0^s \mathbb{E}^W \int_0^T \left[\int \left| \left(\nabla_a \frac{\delta \mathbf{H}^\sigma}{\delta m} \right) (a, (P_s \mu^0)_t) \right|^2 (P_s \mu^0)_t(da) \right] dt ds \\ &\leq 0. \end{aligned}$$

This is a contradiction and so μ^* must be (locally) optimal.

Can there be any other (locally) optimal control $\nu^* \in \mathcal{V}_2^W$ that is not $\nu^* \notin \mathcal{I}^\sigma$?

Can there be any other (locally) optimal control $\nu^* \in \mathcal{V}_2^W$ that is not $\nu^* \notin \mathcal{I}^\sigma$?

For any other (locally) optimal control $\nu^* \in \mathcal{V}_2^W$ we have for any $\nu \in \mathcal{V}_2^W$, due to necessary condition that

$$0 \leq \mathbb{E}^W \left[\int_0^T \int \frac{\delta \mathbf{H}_t^\sigma}{\delta m}(a, \nu^*)(\nu_t - \nu_t^*)(da) dt \right].$$

It is easy to show that this implies that $\nu^* \in \mathcal{I}^\sigma$

Can there be any other (locally) optimal control $\nu^* \in \mathcal{V}_2^W$ that is not $\nu^* \notin \mathcal{I}^\sigma$?

For any other (locally) optimal control $\nu^* \in \mathcal{V}_2^W$ we have for any $\nu \in \mathcal{V}_2^W$, due to necessary condition that

$$0 \leq \mathbb{E}^W \left[\int_0^T \int \frac{\delta \mathbf{H}_t^\sigma}{\delta m}(a, \nu^*)(\nu_t - \nu_t^*)(da) dt \right].$$

It is easy to show that this implies that $\nu^* \in \mathcal{I}^\sigma$

But we have already shown that $\mathcal{I}^\sigma = \{\mu^*\}$ and so the set of local minimizers is a singleton and thus μ^* is the global minimizer of J^σ and item ii) is proved.

Recovering Markovian Control

Lemma 4

Assume that for any $m, m' \in \mathcal{V}_q^W$ it holds that

$$F(x, (1 - \alpha)m + \alpha m') \leq (1 - \alpha)F(x, m) + \alpha F(x, m') \text{ for all } \alpha \in [0, 1] \text{ and } x \in \mathbb{R}^d.$$

Further assume that there exists ϕ and γ , are such that

$$\Phi_t(x, m) = \int \phi_t(x, a) m(da) \quad \Gamma_t(x, m)(\Gamma_t(x, m))^\top = \int \gamma_t(x, a) \gamma_t(x, a)^\top m(da).$$

Define Markov control $\hat{\nu}_t(a, x) := \mathbb{E}^W[\nu_t(a) | X_t(\nu) = x]$. Then

$$J^\sigma(\nu, \xi) = J^\sigma(\hat{\nu}, \xi).$$

Define

$$\hat{\Phi}_t(x, \hat{\nu}_t) := \int \phi_t(x, a) \hat{\nu}_t(da, x),$$

$$\hat{\Gamma}_t(x, \hat{\nu}_t) \hat{\Gamma}_t(x, \hat{\nu}_t)^\top := \int \gamma_t(x, a) \gamma_t(x, a)^\top \hat{\nu}_t(da, x).$$

The mimicking theorem states that there exists a weak solution to

$$\hat{X}_t(\hat{\nu}) = \xi + \int_0^t \hat{\Phi}_r(\hat{X}_r(\hat{\nu}), \hat{\nu}_r) dr + \int_0^t \hat{\Gamma}_r(\hat{X}_r(\hat{\nu}), \hat{\nu}_r) d\hat{W}_r, \quad t \in [0, T],$$

and $\mathcal{L}(\hat{X}_t) = \mathcal{L}(X_t(\mu))$ for all $t \in [0, T]$. This is a controlled Markov process.

Define

$$\hat{\Phi}_t(x, \hat{\nu}_t) := \int \phi_t(x, a) \hat{\nu}_t(da, x),$$

$$\hat{\Gamma}_t(x, \hat{\nu}_t) \hat{\Gamma}_t(x, \hat{\nu}_t)^\top := \int \gamma_t(x, a) \gamma_t(x, a)^\top \hat{\nu}_t(da, x).$$

The mimicking theorem states that there exists a weak solution to

$$\hat{X}_t(\hat{\nu}) = \xi + \int_0^t \hat{\Phi}_r(\hat{X}_r(\hat{\nu}), \hat{\nu}_r) dr + \int_0^t \hat{\Gamma}_r(\hat{X}_r(\hat{\nu}), \hat{\nu}_r) d\hat{W}_r, \quad t \in [0, T],$$

and $\mathcal{L}(\hat{X}_t) = \mathcal{L}(X_t(\mu))$ for all $t \in [0, T]$. This is a controlled Markov process.

First note that, by convexity of entropy and due to Jensen's inequality

$$\int \log(\hat{\nu}_t(a, x)) \hat{\nu}_t(a, x) da \leq \int \mathbb{E}^W [\log(\nu_t(a)) \nu_t(a) | X_t = x] da.$$

Define

$$\begin{aligned}\hat{\Phi}_t(x, \hat{\nu}_t) &:= \int \phi_t(x, a) \hat{\nu}_t(da, x), \\ \hat{\Gamma}_t(x, \hat{\nu}_t) \hat{\Gamma}_t(x, \hat{\nu}_t)^\top &:= \int \gamma_t(x, a) \gamma_t(x, a)^\top \hat{\nu}_t(da, x).\end{aligned}$$

The mimicking theorem states that there exists a weak solution to

$$\hat{X}_t(\hat{\nu}) = \xi + \int_0^t \hat{\Phi}_r(\hat{X}_r(\hat{\nu}), \hat{\nu}_r) dr + \int_0^t \hat{\Gamma}_r(\hat{X}_r(\hat{\nu}), \hat{\nu}_r) d\hat{W}_r, \quad t \in [0, T],$$

and $\mathcal{L}(\hat{X}_t) = \mathcal{L}(X_t(\mu))$ for all $t \in [0, T]$. This is a controlled Markov process.

First note that, by convexity of entropy and due to Jensen's inequality

$$\int \log(\hat{\nu}_t(a, x)) \hat{\nu}_t(a, x) da \leq \int \mathbb{E}^W [\log(\nu_t(a)) \nu_t(a) | X_t = x] da.$$

Since $\mathcal{L}(\hat{X}_t) = \mathcal{L}(X_t(\nu))$ and F is convex, we have

$$\begin{aligned}J^\sigma(\hat{\nu}, \xi) &:= \mathbb{E}^{\hat{W}} \left[\int_0^T \left[F_t(\hat{X}_t, \hat{\nu}_t(\cdot, \hat{X}_t)) + \frac{\sigma^2}{2} \text{Ent}(\nu_t(\cdot, \hat{X}_t)) \right] dt + g(\hat{X}_T) \middle| \hat{X}_0 = \xi \right] \\ &= \mathbb{E}^W \left[\int_0^T \left[F_t(X_t, \hat{\nu}_t(\cdot, X_t)) + \frac{\sigma^2}{2} \text{Ent}(\hat{\nu}_t(\cdot, X_t)) \right] dt + g(X_T) \middle| X_0(\nu) = \xi \right] \\ &\leq \mathbb{E}^W \left[\int_0^T \int \mathbb{E} \left[F_t(X_t, \nu_t) + \frac{\sigma^2}{2} \text{Ent}(\nu_t) | X_t \right] dt + g(X_T) \middle| X_0 = \xi \right] = J^\sigma(\nu, \xi).\end{aligned}$$

Since we always have $J^\sigma(\nu, \xi) \leq J^\sigma(\hat{\nu}, \xi)$ the conclusion follows.

Regularity of Control

$$\begin{cases} \nu_t^* &= \operatorname{argmin}_{\nu \in \mathcal{V}_q} H_t^\sigma(X_t, Y_t, Z_t, \nu), \\ X_t &= \xi + \int_0^t \int \Phi_r(X_r, \nu_r^*) dr + \int_0^t \int \Gamma_r(X_r, \nu_r^*(da)) dW_r, \quad t \in [0, T], \\ Y_t &= (\nabla_x g)(X_T) + \int_t^T (\nabla_x H_r)(X_r, Y_r, Z_r, \nu_r^*) dr + \int_t^T Z_r dW_r. \end{cases}$$

If (Y, Z) were Markov, then we would have that $\nu^* \in \mathcal{V}_q^M$ (Markov).

Regularity of Control

$$\begin{cases} \nu_t^* &= \operatorname{argmin}_{\nu \in \mathcal{V}_q} H_t^\sigma(X_t, Y_t, Z_t, \nu), \\ X_t &= \xi + \int_0^t \int \Phi_r(X_r, \nu_r^*) dr + \int_0^t \int \Gamma_r(X_r, \nu_r^*(da)) dW_r, \quad t \in [0, T], \\ Y_t &= (\nabla_x g)(X_T) + \int_t^T (\nabla_x H_r)(X_r, Y_r, Z_r, \nu_r^*) dr + \int_t^T Z_r dW_r. \end{cases}$$

If (Y, Z) where Markov, then we would have that $\nu^* \in \mathcal{V}_q^M$ (Markov).

We proceed by iteration. Let $\nu^0 \in \mathcal{V}_q^{X, W^2}$. For $n \geq 0$ define

$$\begin{cases} \nu_t^{n+1} &= Z^{-1} e^{\frac{-2}{\sigma^2} \frac{\delta H_t^0}{\delta m}(X_t^n, Y_t^n, Z_t^n, \nu_t^n, a)} \gamma(a), \quad Z = \int e^{\frac{-2}{\sigma^2} \frac{\delta H_t^0}{\delta m}(X_t^n, Y_t^n, Z_t^n, \nu_t^n, a)} \gamma(a) da, \\ X_t^n &= \xi + \int_0^t \int \Phi_r(X_r^n, \nu_r^n) dr + \int_0^t \int \Gamma_r(X_r^n, \nu_r^n) dW_r, \quad t \in [0, T], \\ Y_t^n &= (\nabla_x g)(X_T^n) + \int_t^T (\nabla_x H_r^0)(X_r^n, Y_r^n, Z_r^n, \nu_r^n) dr + \int_t^T Z_r^n dW_r. \end{cases}$$

► See [Reisinger and Zhang, 2020] for related work.

We have

$$u^n(t, x) := Y_t^{n, (t, x)} \quad \text{and} \quad Z_t^{n, (t, X_t)} = \nabla_x u^n(t, X_t) \Gamma_t(X_t^n, \nu^n(X_t^n))$$

We have

$$u^n(t, x) := Y_t^{n,(t,x)} \quad \text{and} \quad Z_t^{n,(t,X_t)} = \nabla_x u^n(t, X_t) \Gamma_t(X_t^n, \nu^n(X_t^n))$$

For suitable test function one can show that

$$\left| \int \nabla_x \nu_t^{n+1}(x, a) da \right| \leq \frac{4}{\sigma^2} \sup_a \left\| \left(\nabla_x \frac{\mathbf{H}_t^0}{\delta m} \right) (x, \nu^n(x), a) \right\|,$$

$$\left| \int f(a) \nabla_x \nu_t^{n+1}(x, a) da \right| \leq \frac{4}{\sigma^2} \sup_a \left\| \left(\nabla_x \frac{\mathbf{H}_t^0}{\delta m} \right) (x, \nu^n(\cdot, x), a) \right\| \int f(a) \nu^{n+1}(a, x) da,$$

We have

$$u^n(t, x) := Y_t^{n,(t,x)} \quad \text{and} \quad Z_t^{n,(t,X_t)} = \nabla_x u^n(t, X_t) \Gamma_t(X_t^n, \nu^n(X_t^n))$$

For suitable test function one can show that

$$\left| \int \nabla_x \nu_t^{n+1}(x, a) da \right| \leq \frac{4}{\sigma^2} \sup_a \left\| \left(\nabla_x \frac{\mathbf{H}_t^0}{\delta m} \right) (x, \nu^n(x), a) \right\|,$$

$$\left| \int f(a) \nabla_x \nu_t^{n+1}(x, a) da \right| \leq \frac{4}{\sigma^2} \sup_a \left\| \left(\nabla_x \frac{\mathbf{H}_t^0}{\delta m} \right) (x, \nu^n(\cdot, x), a) \right\| \int f(a) \nu^{n+1}(a, x) da,$$

For large σ conclude by Arzelá-Ascoli theorem.

Generative modelling

Generative modelling

- ▶ Generative models such as GANs or VAEs demonstrated a great success in seemingly high dimensional setups.

Generative modelling

- ▶ Generative models such as GANs or VAEs demonstrated a great success in seemingly high dimensional setups.
- ▶ **Input:** Source distribution μ and target distribution ν i.e input-output data

Generative modelling

- ▶ Generative models such as GANs or VAEs demonstrated a great success in seemingly high dimensional setups.
- ▶ **Input:** Source distribution μ and target distribution ν i.e input-output data
- ▶ A generative model is a **transport map** T from μ to ν i.e T is a map that “pushes μ onto ν ”. We write $T_{\#}\mu = \nu$.

Generative modelling

- ▶ Generative models such as GANs or VAEs demonstrated a great success in seemingly high dimensional setups.
- ▶ **Input:** Source distribution μ and target distribution ν i.e input-output data
- ▶ A generative model is a **transport map** T from μ to ν i.e T is a map that “pushes μ onto ν ”. We write $T_{\#}\mu = \nu$.
- ▶ Parametrise transport map $T(\theta)$, $\theta \in \mathbb{R}^p$, e.g some network architecture or Heston model

Generative modelling

- ▶ Generative models such as GANs or VAEs demonstrated a great success in seemingly high dimensional setups.
- ▶ **Input:** Source distribution μ and target distribution ν i.e input-output data
- ▶ A generative model is a **transport map** T from μ to ν i.e T is a map that “pushes μ onto ν ”. We write $T_{\#}\mu = \nu$.
- ▶ Parametrise transport map $T(\theta)$, $\theta \in \mathbb{R}^p$, e.g some network architecture or Heston model
- ▶ Seek θ^* s.t $T(\theta^*)_{\#}\mu \approx \nu$.

Generative modelling

- ▶ Generative models such as GANs or VAEs demonstrated a great success in seemingly high dimensional setups.
- ▶ **Input:** Source distribution μ and target distribution ν i.e input-output data
- ▶ A generative model is a **transport map** T from μ to ν i.e T is a map that “pushes μ onto ν ”. We write $T_{\#}\mu = \nu$.
- ▶ Parametrise transport map $T(\theta)$, $\theta \in \mathbb{R}^p$, e.g some network architecture or Heston model
- ▶ Seek θ^* s.t $T(\theta^*)_{\#}\mu \approx \nu$.
- ▶ Need to make the choice of the metric

$$D(T(\theta)_{\#}\mu, \nu) := \sup_{f \in \mathcal{K}} \left\| \int f(x)(T(\theta)_{\#}\mu)(dx) - \int f(x)\nu(dx) \right\|$$

- ▶ $\mathcal{K}$ could be set of options we want to calibrate to, could be neural network

Generative modelling

- ▶ Generative models such as GANs or VAEs demonstrated a great success in seemingly high dimensional setups.
- ▶ **Input:** Source distribution μ and target distribution ν i.e input-output data
- ▶ A generative model is a **transport map** T from μ to ν i.e T is a map that “pushes μ onto ν ”. We write $T_{\#}\mu = \nu$.
- ▶ Parametrise transport map $T(\theta)$, $\theta \in \mathbb{R}^p$, e.g some network architecture or Heston model
- ▶ Seek θ^* s.t $T(\theta^*)_{\#}\mu \approx \nu$.
- ▶ Need to make the choice of the metric

$$D(T(\theta)_{\#}\mu, \nu) := \sup_{f \in \mathcal{K}} \left\| \int f(x)(T(\theta)_{\#}\mu)(dx) - \int f(x)\nu(dx) \right\|$$

- ▶ $\mathcal{K}$ could be set of options we want to calibrate to, could be neural network
- ▶ **The modelling choices are**
 1. **metric** D
 2. **parametrisation of** T
 3. **algorithm used for training!!!**

Generative modelling with causal transport

Fix $m^\star \in \mathcal{P}([0, T] \times \mathbb{R}^d)$ to be a target distribution.

Generative modelling with causal transport

Fix $m^\star \in \mathcal{P}([0, T] \times \mathbb{R}^d)$ to be a target distribution.

The aim of the generative model is to map some basic distribution, in our case $m^0 := \mathcal{L}(\xi) \otimes \mathcal{L}(W)$, into $m^\star$.

Generative modelling with causal transport

Fix $m^\star \in \mathcal{P}([0, T] \times \mathbb{R}^d)$ to be a target distribution.

The aim of the generative model is to map some basic distribution, in our case $m^0 := \mathcal{L}(\xi) \otimes \mathcal{L}(W)$, into $m^\star$.

There is a measurable map $G^\mu : \mathbb{R}^d \times C[0, T]^d \rightarrow C[0, T]^d$ such that $X(\mu) = G^\mu(\xi, (W_s)_{s \in [0, T]})$, and $X_t(\mu) := G_t^\mu(\xi, (W_{s \wedge t})_{s \in [0, T]})$

Generative modelling with causal transport

Fix $m^\star \in \mathcal{P}([0, T] \times \mathbb{R}^d)$ to be a target distribution.

The aim of the generative model is to map some basic distribution, in our case $m^0 := \mathcal{L}(\xi) \otimes \mathcal{L}(W)$, into $m^\star$.

There is a measurable map $G^\mu : \mathbb{R}^d \times C[0, T]^d \rightarrow C[0, T]^d$ such that $X(\mu) = G^\mu(\xi, (W_s)_{s \in [0, T]})$, and $X_t(\mu) := G_t^\mu(\xi, (W_{s \wedge t})_{s \in [0, T]})$

one can view solution map as a generative model that maps $\mathcal{L}(\xi) \otimes \mathcal{L}(W)$ into $(G_t^\mu)_\# m^0$. Note that by construction G^μ is a causal transport map

Generative modelling with causal transport

Fix $m^* \in \mathcal{P}([0, T] \times \mathbb{R}^d)$ to be a target distribution.

The aim of the generative model is to map some basic distribution, in our case $m^0 := \mathcal{L}(\xi) \otimes \mathcal{L}(W)$, into m^* .

There is a measurable map $G^\mu : \mathbb{R}^d \times C[0, T]^d \rightarrow C[0, T]^d$ such that $X(\mu) = G^\mu(\xi, (W_s)_{s \in [0, T]})$, and $X_t(\mu) := G_t^\mu(\xi, (W_{s \wedge t})_{s \in [0, T]})$

one can view solution map as a generative model that maps $\mathcal{L}(\xi) \otimes \mathcal{L}(W)$ into $(G_t^\mu)_\# m^0$. Note that by construction G^μ is a causal transport map

One then seeks μ^* such that $G_\#^{\mu^*} m^0$ is a good approximation of m^* optimisation problem e.g

$$J^\sigma(\nu, \xi) := \mathbb{E}^W \left[\int_0^T \left[\log \left(\frac{m_t(X_t(\nu))}{m_t^*(X_t(\nu))} \right) + \frac{\sigma^2}{2} \text{Ent}(\nu_t) \right] dt \middle| X_0(\nu) = \xi \right].$$

Generative modelling with causal transport

Fix $m^* \in \mathcal{P}([0, T] \times \mathbb{R}^d)$ to be a target distribution.

The aim of the generative model is to map some basic distribution, in our case $m^0 := \mathcal{L}(\xi) \otimes \mathcal{L}(W)$, into m^* .

There is a measurable map $G^\mu : \mathbb{R}^d \times C[0, T]^d \rightarrow C[0, T]^d$ such that $X(\mu) = G^\mu(\xi, (W_s)_{s \in [0, T]})$, and $X_t(\mu) := G_t^\mu(\xi, (W_{s \wedge t})_{s \in [0, T]})$

one can view solution map as a generative model that maps $\mathcal{L}(\xi) \otimes \mathcal{L}(W)$ into $(G_t^\mu)_\# m^0$. Note that by construction G^μ is a causal transport map

One then seeks μ^* such that $G_\#^{\mu^*} m^0$ is a good approximation of m^* optimisation problem e.g

$$J^\sigma(\nu, \xi) := \mathbb{E}^W \left[\int_0^T \left[\log \left(\frac{m_t(X_t(\nu))}{m_t^*(X_t(\nu))} \right) + \frac{\sigma^2}{2} \text{Ent}(\nu_t) \right] dt \middle| X_0(\nu) = \xi \right].$$

The above optimization problem does not fit within the framework of what I presented about neural SDE.

► See [Acciaio et al., 2019] for related work

Robust pricing and hedging via neural SDEs

Model Selection

Until recently, models in finance and economics were mostly conceived in a three step fashion:

- ▶ gathering statistical properties of the underlying time-series or the so called stylized facts

Model Selection

Until recently, models in finance and economics were mostly conceived in a three step fashion:

- ▶ gathering statistical properties of the underlying time-series or the so called stylized facts
- ▶ handcrafting a parsimonious model, which would best capture the desired market characteristics without adding any needless complexity and

Model Selection

Until recently, models in finance and economics were mostly conceived in a three step fashion:

- ▶ gathering statistical properties of the underlying time-series or the so called stylized facts
- ▶ handcrafting a parsimonious model, which would best capture the desired market characteristics without adding any needless complexity and
- ▶ calibration and validation of the handcrafted model.

Model Selection

Until recently, models in finance and economics were mostly conceived in a three step fashion:

- ▶ gathering statistical properties of the underlying time-series or the so called stylized facts
- ▶ handcrafting a parsimonious model, which would best capture the desired market characteristics without adding any needless complexity and
- ▶ calibration and validation of the handcrafted model.

... model complexity was undesirable

Classical Risk Models:

Pros:

- ▶ Interpretable parameters
- ▶ Relatively easy to calibrate with relatively small amount of data
- ▶ Several decades of underpinning research

Classical Risk Models:

Pros:

- ▶ Interpretable parameters
- ▶ Relatively easy to calibrate with relatively small amount of data
- ▶ Several decades of underpinning research

Cons:

- ▶ Lack of systematic framework for model selection
- ▶ Knighting uncertainty (Unknown unknowns)
- ▶ limited expressivity

Model Calibration

Classical Calibration:

- ▶ Pick a parametric model $(S_t(\theta))_{t \in [0, T]}$ (e.g an Itô process) with parameters $\theta \in \mathbb{R}^p$

Model Calibration

Classical Calibration:

- ▶ Pick a parametric model $(S_t(\theta))_{t \in [0, T]}$ (e.g an Itô process) with parameters $\theta \in \mathbb{R}^p$
- ▶ Parametric model induces martingale measure $\mathbb{Q}(\theta)$

Model Calibration

Classical Calibration:

- ▶ Pick a parametric model $(S_t(\theta))_{t \in [0, T]}$ (e.g an Itô process) with parameters $\theta \in \mathbb{R}^p$
- ▶ Parametric model induces martingale measure $\mathbb{Q}(\theta)$
- ▶ **Input Data:** prices of traded derivatives $p(\Phi_i)_{i=0}^M$ with corresponding payoffs $(\Phi_i)_{i=0}^M$

Model Calibration

Classical Calibration:

- ▶ Pick a parametric model $(S_t(\theta))_{t \in [0, T]}$ (e.g an Itô process) with parameters $\theta \in \mathbb{R}^p$
- ▶ Parametric model induces martingale measure $\mathbb{Q}(\theta)$
- ▶ **Input Data:** prices of traded derivatives $p(\Phi_i)_{i=0}^M$ with corresponding payoffs $(\Phi_i)_{i=0}^M$
- ▶ **Output :** θ^* such that $p(\Phi_i) \approx \mathbb{E}_{\mathbb{Q}(\theta^*)}[\Phi_i]$

Robust Price bounds

- ▶ There are infinitely many models that are consistent with the market

Robust Price bounds

- ▶ There are infinitely many models that are consistent with the market
- ▶ $\mathcal{M}$ - set of all martingale measures that are calibrated to data

Robust Price bounds

- ▶ There are infinitely many models that are consistent with the market
- ▶ $\mathcal{M}$ - set of all martingale measures that are calibrated to data
- ▶ Compute conservative bounds for the price

$$\sup_{Q \in \mathcal{M}} E[\Psi] \quad \text{and} \quad \inf_{Q \in \mathcal{M}} E[\Psi]$$

- ▶ Use duality theory to deduce (semi-static) hedging strategy

Robust Price bounds

- ▶ There are infinitely many models that are consistent with the market
- ▶ $\mathcal{M}$ - set of all martingale measures that are calibrated to data
- ▶ Compute conservative bounds for the price

$$\sup_{\mathbb{Q} \in \mathcal{M}} E[\Psi] \quad \text{and} \quad \inf_{\mathbb{Q} \in \mathcal{M}} E[\Psi]$$

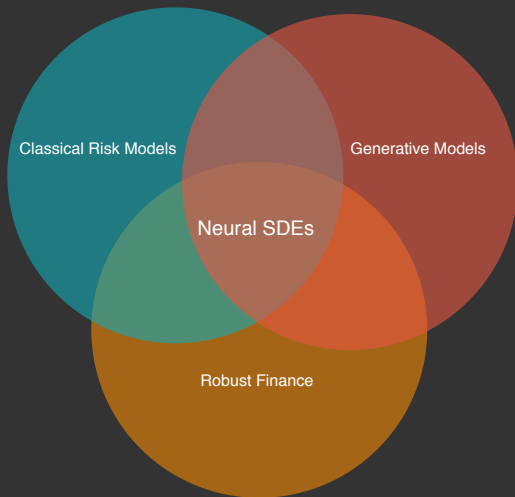
- ▶ Use duality theory to deduce (semi-static) hedging strategy
- ▶ The obtain bounds typically to wide to be of practical value

Robust Price bounds

- ▶ There are infinitely many models that are consistent with the market
- ▶ $\mathcal{M}$ - set of all martingale measures that are calibrated to data
- ▶ Compute conservative bounds for the price

$$\sup_{Q \in \mathcal{M}} E[\Psi] \quad \text{and} \quad \inf_{Q \in \mathcal{M}} E[\Psi]$$

- ▶ Use duality theory to deduce (semi-static) hedging strategy
- ▶ The obtain bounds typically to wide to be of practical value
- ▶ Challenges:
 - a) Incorporate prior information to restrict a search space $\mathcal{M}$
 - b) Design efficient algorithms for computing price bounds and corresponding hedges



Neural SDEs

- We build an Itô process $(X_t^\theta)_{t \in [0, T]}$, with parameters $\theta \in \mathbb{R}^p$

$$\begin{aligned}dS_t^\theta &= rS_t^\theta dt + \sigma^S(t, X_t^\theta, \theta) dW_t, \\dV_t^\theta &= b^V(t, X_t^\theta, \theta) dt + \sigma^V(t, X_t^\theta, \theta) dW_t, \\X_t^\theta &= (S_t^\theta, V_t^\theta),\end{aligned}$$

where σ^S, b^V, σ^V are given by neural networks (can be path-depedend)

Neural SDEs

- We build an Itô process $(X_t^\theta)_{t \in [0, T]}$, with parameters $\theta \in \mathbb{R}^p$

$$\begin{aligned}dS_t^\theta &= rS_t^\theta dt + \sigma^S(t, X_t^\theta, \theta) dW_t, \\dV_t^\theta &= b^V(t, X_t^\theta, \theta) dt + \sigma^V(t, X_t^\theta, \theta) dW_t, \\X_t^\theta &= (S_t^\theta, V_t^\theta),\end{aligned}$$

where σ^S, b^V, σ^V are given by neural networks (can be path-depedend)

- The model induces a martingale probability measure $\mathbb{Q}(\theta)$

Neural SDEs

- ▶ We build an Itô process $(X_t^\theta)_{t \in [0, T]}$, with parameters $\theta \in \mathbb{R}^p$

$$\begin{aligned}dS_t^\theta &= rS_t^\theta dt + \sigma^S(t, X_t^\theta, \theta) dW_t, \\dV_t^\theta &= b^V(t, X_t^\theta, \theta) dt + \sigma^V(t, X_t^\theta, \theta) dW_t, \\X_t^\theta &= (S_t^\theta, V_t^\theta),\end{aligned}$$

where σ^S, b^V, σ^V are given by neural networks (can be path-depedend)

- ▶ The model induces a martingale probability measure $\mathbb{Q}(\theta)$
- ▶ Solution map is an instance of casual transport
- ▶ See [Cuchiero et al., 2020] for neural SDEs with a prior on vol process.
- ▶ See [Arribas et al., 2020] for Sig-SDEs (neural SDE in a signature feature space)
- ▶ Neural SDEs are easy to work with e.g consistent change from $\mathbb{Q}$ to $\mathbb{P}$.

Neural SDEs

- i) **Calibration to market prices** Find model parameters θ^* such that model prices match market prices:

$$\theta^* \in \arg \min_{\theta \in \Theta} \sum_{i=1}^M \ell(\mathbb{E}^{\mathbb{Q}(\theta)}[\Phi_i], p(\Phi_i)).$$

Neural SDEs

- i) **Calibration to market prices** Find model parameters θ^* such that model prices match market prices:

$$\theta^* \in \arg \min_{\theta \in \Theta} \sum_{i=1}^M \ell(\mathbb{E}^{\mathbb{Q}(\theta)}[\Phi_i], p(\Phi_i)) .$$

- ii) **Robust pricing** Find model parameters $\theta^{l,*}$ and $\theta^{u,*}$ which provide robust arbitrage-free price bounds for an illiquid derivative, subject to available market data:

$$\begin{aligned} \theta^{l,*} &\in \arg \min_{\theta \in \Theta} \mathbb{E}^{\mathbb{Q}(\theta)}[\Psi] , & \text{subject to } & \sum_{i=1}^M \ell(\mathbb{E}^{\mathbb{Q}(\theta)}[\Phi_i], p(\Phi_i)) = 0 , \\ \theta^{u,*} &\in \arg \max_{\theta \in \Theta} \mathbb{E}^{\mathbb{Q}(\theta)}[\Psi] , & \text{subject to } & \sum_{i=1}^M \ell(\mathbb{E}^{\mathbb{Q}(\theta)}[\Phi_i], p(\Phi_i)) = 0 . \end{aligned}$$

where $\ell : \mathbb{R} \times \mathbb{R} \rightarrow [0, \infty)$ is a convex loss function such that $\min_{x \in \mathbb{R}, y \in \mathbb{R}} \ell(x, y) = 0$.

Neural SDEs

Let $\mathbb{Q}^N(\theta) := \frac{1}{N} \sum_{i=1}^N \delta_{X^{i,\theta}}$ be empirical approximation of $\mathbb{Q}(\theta)$.

From CLT,

$$\mathbb{P} \left(\mathbb{E}^{\mathbb{Q}(\theta)}[\Phi] \in \left[\mathbb{E}^{\mathbb{Q}^N(\theta)}[\Phi] - z_{\alpha/2} \frac{\sigma}{\sqrt{N}}, \mathbb{E}^{\mathbb{Q}^N(\theta)}[\Phi] + z_{\alpha/2} \frac{\sigma}{\sqrt{N}} \right] \right) \rightarrow 1$$

where $\sigma = \sqrt{\text{Var}[\Phi]}$.

Neural SDEs

Let $\mathbb{Q}^N(\theta) := \frac{1}{N} \sum_{i=1}^N \delta_{X^i, \theta}$ be empirical approximation of $\mathbb{Q}(\theta)$.

From CLT,

$$\mathbb{P} \left(\mathbb{E}^{\mathbb{Q}(\theta)}[\Phi] \in \left[\mathbb{E}^{\mathbb{Q}^N(\theta)}[\Phi] - z_{\alpha/2} \frac{\sigma}{\sqrt{N}}, \mathbb{E}^{\mathbb{Q}^N(\theta)}[\Phi] + z_{\alpha/2} \frac{\sigma}{\sqrt{N}} \right] \right) \rightarrow 1$$

where $\sigma = \sqrt{\mathbb{V}\text{ar}[\Phi]}$. We seek a random variable Φ^{cv} such that:

$$\mathbb{E}^{\mathbb{Q}^N(\theta)}[\Phi^{cv}] = \mathbb{E}[\Phi] \quad \text{and} \quad \mathbb{V}\text{ar}[\Phi^{cv}] < \mathbb{V}\text{ar}[\Phi] .$$

Stochastic Optimisation

Let $M = 1$ and the loss function

$$h(\theta) = \ell\left(\mathbb{E}^{\mathbb{Q}(\theta)}[\Phi^{\text{cv}}], \mathfrak{p}(\Phi)\right).$$

Then in the gradient step update we have

$$\partial_{\theta} h(\theta) = \partial_x \ell\left(\mathbb{E}^{\mathbb{Q}}[\Phi^{\text{cv}}(X^{\theta})], \mathfrak{p}(\Phi)\right) \mathbb{E}^{\mathbb{Q}}[\partial_{\theta} \Phi(X^{\theta})],$$

Since ℓ is typically not an identity function, a mini-batch estimator of $\partial_{\theta} h(\theta)$, obtained by replacing $\mathbb{Q}$ with $\mathbb{Q}^N$ given by

$$\partial_{\theta} h^N(\theta) := \partial_x \ell\left(\mathbb{E}^{\mathbb{Q}^N}[\Phi^{\text{cv}}(X^{\theta})], \mathfrak{p}(\Phi)\right) \mathbb{E}^{\mathbb{Q}^N}[\partial_{\theta} \Phi(X^{\theta})],$$

is a biased estimator of $\partial_{\theta} h$.

Stochastic Optimisation

Let $M = 1$ and the loss function

$$h(\theta) = \ell\left(\mathbb{E}^{\mathbb{Q}(\theta)}[\Phi^{\text{cv}}], \mathfrak{p}(\Phi)\right).$$

Then in the gradient step update we have

$$\partial_{\theta} h(\theta) = \partial_x \ell\left(\mathbb{E}^{\mathbb{Q}}[\Phi^{\text{cv}}(X^{\theta})], \mathfrak{p}(\Phi)\right) \mathbb{E}^{\mathbb{Q}}[\partial_{\theta} \Phi(X^{\theta})],$$

Since ℓ is typically not an identity function, a mini-batch estimator of $\partial_{\theta} h(\theta)$, obtained by replacing $\mathbb{Q}$ with $\mathbb{Q}^N$ given by

$$\partial_{\theta} h^N(\theta) := \partial_x \ell\left(\mathbb{E}^{\mathbb{Q}^N}[\Phi^{\text{cv}}(X^{\theta})], \mathfrak{p}(\Phi)\right) \mathbb{E}^{\mathbb{Q}^N}[\partial_{\theta} \Phi(X^{\theta})],$$

is a biased estimator of $\partial_{\theta} h$.

Lemma 5

For $\ell(x, y) = |x - y|^2$, we have

$$|\mathbb{E}^{\mathbb{Q}}[\partial_{\theta} h^N(\theta)] - \partial_{\theta} h(\theta)| \leq \frac{2}{N} (\text{Var}^{\mathbb{Q}}[\Phi^{\text{cv}}(X^{\theta})])^{1/2} (\text{Var}^{\mathbb{Q}}[\partial_{\theta} \Phi(X^{\theta})])^{1/2}.$$

Learning PDEs

Let

$$X_t = \sigma(t, (X_{s \wedge t})_{s \in [0, T]}) dW_t,$$

Learning PDEs

Let

$$dX_t = \sigma(t, (X_{s \wedge t})_{s \in [0, T]}) dW_t,$$

$$F_t := F(t, (X_{s \wedge t})_{s \in [0, T]}) = \mathbb{E} [\psi((X_s)_{s \in [0, T]}) | (X_{s \wedge t})_{s \in [0, T]} = (X_{s \wedge t})_{s \in [0, T]}]$$

Learning PDEs

Let

$$X_t = \sigma(t, (X_{s \wedge t})_{s \in [0, T]}) dW_t,$$

$$F_t := F(t, (X_{s \wedge t})_{s \in [0, T]}) = \mathbb{E} [\psi((X_s)_{s \in [0, T]}) | (X_{s \wedge t})_{s \in [0, T]} = (X_{s \wedge t})_{s \in [0, T]}]$$

Martingale representation theorem

$$F_t = \Psi - \int_t^T Z_s dW_s.$$

With functional Itô calculus

$$F_t = \psi((X_s)_{s \in [0, T]}) - \int_t^T \nabla_\omega \psi((X_{r \wedge s})_{r \in [0, T]}) dX_s.$$

Learning PDEs

Let

$$X_t = \sigma(t, (X_{s \wedge t})_{s \in [0, T]}) dW_t,$$

$$F_t := F(t, (X_{s \wedge t})_{s \in [0, T]}) = \mathbb{E} [\psi((X_s)_{s \in [0, T]}) | (X_{s \wedge t})_{s \in [0, T]} = (X_{s \wedge t})_{s \in [0, T]}]$$

Martingale representation theorem

$$F_t = \Psi - \int_t^T Z_s dW_s.$$

With functional Itô calculus

$$F_t = \psi((X_s)_{s \in [0, T]}) - \int_t^T \nabla_\omega \psi((X_{r \wedge s})_{r \in [0, T]}) dX_s.$$

- ▶ Can learn (parametric) path dependent PDEs
- ▶ We have unbiased approximation to the PDE by hybrid Monte Carlo/deep learning, see [Vidales et al., 2018]

Neural SDEs - Algorithm

Input: $\pi = \{t_0, t_1, \dots, t_{N_{\text{steps}}}\}$ time grid for numerical scheme.

Input: $(\Phi_i)_{i=1}^{N_{\text{prices}}}$ option payoffs.

Input: Market option prices $p(\Phi_j)$, $j = 1, \dots, N_{\text{prices}}$.

for epoch : 1 : N_{epochs} **do**

Generate N_{trn} paths $(x_{t_n}^{\pi, \theta, i})_{n=0}^{N_{\text{steps}}} := (s_{t_n}^{\pi, \theta, i}, v_{t_n}^{\pi, \theta, i})_{n=0}^{N_{\text{steps}}}$, $i = 1, \dots, N_{\text{trn}}$ using Euler scheme.

During one epoch: Freeze ξ , use Adam to update θ , where

$$\theta = \widehat{\argmin_{\theta}} \sum_{j=1}^{N_{\text{prices}}} \left(\mathbb{E}^{N_{\text{trn}}} \left[\Phi_j \left(X^{\pi, \theta} \right) - \sum_{k=0}^{N_{\text{steps}}-1} \bar{h}(t_k, \tilde{X}_{t_k}^{\pi, \theta}, \xi_j) \Delta \tilde{S}_{t_k}^{\pi, \theta} \right] - p(\Phi_j) \right)^2$$

During one epoch: Freeze θ , use Adam to update ξ , by optimising the sample variance

$$\xi = \widehat{\argmin_{\xi}} \sum_{j=1}^{N_{\text{prices}}} \mathbb{V}ar^{N_{\text{trn}}} \left[\Phi_j \left(X^{\pi, \theta} \right) - \sum_{k=0}^{N_{\text{steps}}-1} \bar{h}(t_k, X_{t_k}^{\pi, \theta}, \xi_j) \Delta \tilde{S}_{t_k}^{\pi, \theta} \right]$$

end for

return θ, ξ_j for all prices $(\Phi_i)_{i=1}^{N_{\text{prices}}}$.

Results

We calibrate (local) Stochastic Volatility model

$$\begin{aligned}dS_t &= rS_t dt + \sigma^S(t, S_t, V_t, \nu) S_t dB_t^S, & S_0 &= 1, \\dV_t &= b^V(V_t, \phi) dt + \sigma^V(V_t, \varphi) dB_t^V, & V_0 &= v_0, \\d\langle B^S, B^V \rangle_t &= \rho dt\end{aligned}$$

to European option prices

$$p(\Phi) := \mathbb{E}^{\mathbb{Q}(\theta)}[\Phi] = e^{-rT} \mathbb{E}^{\mathbb{Q}(\theta)} [(S_T - K)_+ | S_0 = 1]$$

for maturities of 2, 4, ..., 12 months and typically 21 uniformly spaced strikes between in $[0.8, 1.2]$.

Results

We calibrate (local) Stochastic Volatility model

$$\begin{aligned}dS_t &= rS_t dt + \sigma^S(t, S_t, V_t, \nu) S_t dB_t^S, & S_0 &= 1, \\dV_t &= b^V(V_t, \phi) dt + \sigma^V(V_t, \varphi) dB_t^V, & V_0 &= v_0, \\d\langle B^S, B^V \rangle_t &= \rho dt\end{aligned}$$

to European option prices

$$p(\Phi) := \mathbb{E}^{\mathbb{Q}(\theta)}[\Phi] = e^{-rT} \mathbb{E}^{\mathbb{Q}(\theta)} [(S_T - K)_+ | S_0 = 1]$$

for maturities of 2, 4, ..., 12 months and typically 21 uniformly spaced strikes between in $[0.8, 1.2]$.

As an example of an illiquid derivative for which we wish to find robust bounds we take the lookback option

$$p(\Psi) := \mathbb{E}^{\mathbb{Q}(\theta)}[\Psi] = e^{-rT} \mathbb{E}^{\mathbb{Q}(\theta)} \left[\max_{t \in [0, T]} S_t - S_T | X_0 = 1 \right].$$

Results

We calibrate (local) Stochastic Volatility model

$$\begin{aligned}dS_t &= rS_t dt + \sigma^S(t, S_t, V_t, \nu) S_t dB_t^S, & S_0 &= 1, \\dV_t &= b^V(V_t, \phi) dt + \sigma^V(V_t, \varphi) dB_t^V, & V_0 &= v_0, \\d\langle B^S, B^V \rangle_t &= \rho dt\end{aligned}$$

to European option prices

$$p(\Phi) := \mathbb{E}^{\mathbb{Q}(\theta)}[\Phi] = e^{-rT} \mathbb{E}^{\mathbb{Q}(\theta)} [(S_T - K)_+ | S_0 = 1]$$

for maturities of 2, 4, ..., 12 months and typically 21 uniformly spaced strikes between in $[0.8, 1.2]$.

As an example of an illiquid derivative for which we wish to find robust bounds we take the lookback option

$$p(\Psi) := \mathbb{E}^{\mathbb{Q}(\theta)}[\Psi] = e^{-rT} \mathbb{E}^{\mathbb{Q}(\theta)} \left[\max_{t \in [0, T]} S_t - S_T | X_0 = 1 \right].$$

We generate synthetic data using Heston model.

Calibration to market prices

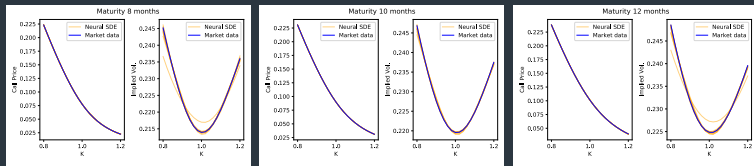


Figure: Vanilla option prices and implied volatility curves of the 10 calibrated Neural SDEs vs. the market data for different maturities.

Robust pricing

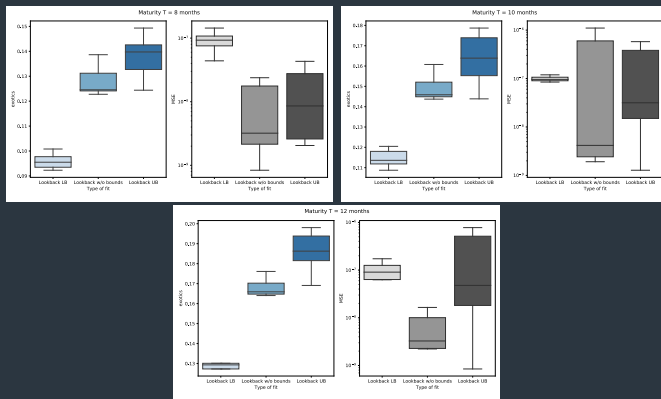


Figure: Exotic option price are in blue; Calibration error i in grey. The three box-plots in each group arise respectively from aiming for a lower bound, ad hoc and upper bound price of illiquid derivative. Each box plot comes from 10 different runs of Neural SDE calibration.

Control Variate effect on training

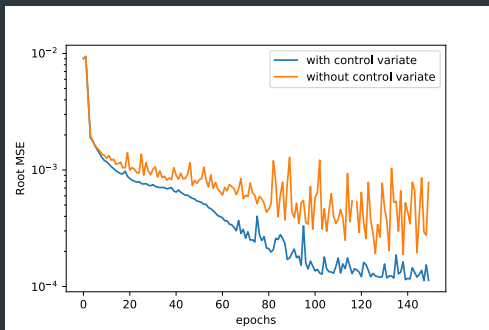


Figure: Root Mean Squared Error of calibration to Vanilla option prices with and without hedging strategy parametrisation

Joint SPX and VIX calibration with neural SDEs

Consider the Neural SDE

$$\begin{aligned}dS_t^\theta &= S_t^\theta \sigma(t, V_t^\theta; \theta) dW_t, \\dV_t^\theta &= a(t, V_t^\theta; \theta) dt + b(t, V_t^\theta; \theta) dB_t, \\ \rho &= \langle dW, dB \rangle_t.\end{aligned}$$

It can be shown that the VIX dynamics at time $t \in [0, T]$ can be expressed as

$$\text{VIX}_t^2 := \frac{1}{\Delta\tau} \mathbb{E} \left[\int_t^{t+\Delta\tau} \sigma_s^2 ds \middle| \mathcal{F}_t \right] = -\frac{2}{\Delta\tau} \mathbb{E} \left[\log \left(\frac{S_{t+\Delta\tau}}{S_t} \right) \middle| \mathcal{F}_t \right], \quad \Delta\tau = \frac{30}{365}$$

The VIX future with maturity maturity T is then given by

$$F_{t,T}^{\text{VIX}} := \mathbb{E} [\text{VIX}_T | \mathcal{F}_t]$$

VIX options are defined as

$$C_t^{\text{VIX}}(T, K) := \mathbb{E} \left[(\text{VIX}_T - K)^+ \middle| \mathcal{F}_t \right], \quad P_t^{\text{VIX}}(T, K) := \mathbb{E} \left[(K - \text{VIX}_T)^+ \middle| \mathcal{F}_t \right].$$

joint work with: Antoine Jacquier, Marc Sabate Vidales, David Siska, Zan Zuric.

Calibration to market data

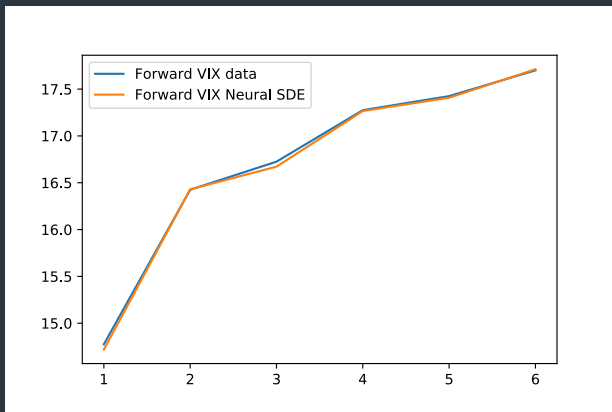


Figure: Calibration to market data (data source: OptionMetrics) containing SPX options, VIX options and VIX future for $T = 1, \dots, 6$ months

Calibration to market data

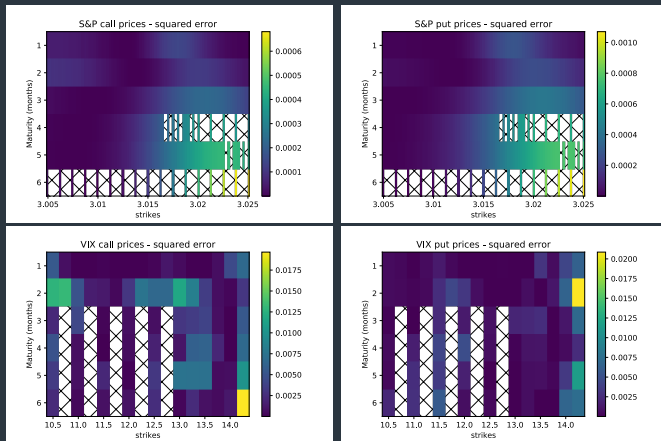


Figure: Calibrated neural SDE errors on SPX options and VIX options. Hatches correspond to combinations of Maturity/Strike for which there was not market data available

Key messages of this mini course

- ▶ Training neural nets is a sampling problem

Key messages of this mini course

- ▶ Training neural nets is a sampling problem
- ▶ Gradient flow view on training neural networks provides mathematical framework to study machine learning

Key messages of this mini course

- ▶ Training neural nets is a sampling problem
- ▶ Gradient flow view on training neural networks provides mathematical framework to study machine learning
- ▶ Probabilistic numerical analysis provides quantitative bounds that do not suffer from the curse of dimensionality

Key messages of this mini course

- ▶ Training neural nets is a sampling problem
- ▶ Gradient flow view on training neural networks provides mathematical framework to study machine learning
- ▶ Probabilistic numerical analysis provides quantitative bounds that do not suffer from the curse of dimensionality
- ▶ Machine learning perspective leads to new algorithms and mathematical tools for (stochastic) control problems/quantitative finance.

Some open questions

- ▶ Problem dependent choices of the metrics/Riemannian structures may lead to more efficient sampling methods

Some open questions

- ▶ Problem dependent choices of the metrics/Riemannian structures may lead to more efficient sampling methods
- ▶ Novel approaches to the 'regularisation by noise' for Transport PDEs

Some open questions

- ▶ Problem dependent choices of the metrics/Riemannian structures may lead to more efficient sampling methods
- ▶ Novel approaches to the 'regularisation by noise' for Transport PDEs
- ▶ Link between class of functions we aim to learn and the the distribution over parameters space of neural networks (need for non-asymptotic theory)

Some open questions

- ▶ Problem dependent choices of the metrics/Riemannian structures may lead to more efficient sampling methods
- ▶ Novel approaches to the 'regularisation by noise' for Transport PDEs
- ▶ Link between class of functions we aim to learn and the the distribution over parameters space of neural networks (need for non-asymptotic theory)
- ▶ Trained neural networks models are random functions hence we need estimates for the uncertainty

Some open questions

- ▶ Problem dependent choices of the metrics/Riemannian structures may lead to more efficient sampling methods
- ▶ Novel approaches to the 'regularisation by noise' for Transport PDEs
- ▶ Link between class of functions we aim to learn and the the distribution over parameters space of neural networks (need for non-asymptotic theory)
- ▶ Trained neural networks models are random functions hence we need estimates for the uncertainty
- ▶ ...there are many many more.

References I

- [Acciaio et al., 2019] Acciaio, B., Backhoff-Veraguas, J., and Carmona, R. (2019). Extended mean field control problems: stochastic maximum principle and transport perspective. *SIAM Journal on Control and Optimization*, 57(6):3666–3693.
- [Arribas et al., 2020] Arribas, I. P., Salvi, C., and Szpruch, L. (2020). Sig-sdes model for quantitative finance. *arXiv preprint arXiv:2006.00218*.
- [Cuchiero et al., 2020] Cuchiero, C., Khosrawi, W., and Teichmann, J. (2020). A generative adversarial network approach to calibration of local stochastic volatility models. *arXiv preprint arXiv:2005.02505*.
- [Gobet and Labart, 2010] Gobet, E. and Labart, C. (2010). Solving bsde with adaptive control variate. *SIAM Journal on Numerical Analysis*, 48(1):257–277.
- [Kerimkulov et al., 2020] Kerimkulov, B., Šiška, D., and Szpruch, L. (2020). Exponential convergence and stability of howard’s policy improvement algorithm for controlled diffusions. *SIAM Journal on Control and Optimization*, 58(3):1314–1340.
- [Reisinger and Zhang, 2020] Reisinger, C. and Zhang, Y. (2020). Regularity and stability of feedback relaxed controls. *arXiv preprint arXiv:2001.03148*.
- [Vidales et al., 2018] Vidales, M. S., Šiška, D., and Szpruch, L. (2018). Martingale functional control variates via deep learning. *arXiv:1810.05094*.