
Integration Matters for Learning PDEs with Backward SDEs

Sungje Park

University of Southern California
sungjepa@usc.edu

Stephen Tu

University of Southern California
stephen.tu@usc.edu

Abstract

Backward stochastic differential equation (BSDE)-based deep learning methods provide an alternative to Physics-Informed Neural Networks (PINNs) for solving high-dimensional partial differential equations (PDEs), offering potential algorithmic advantages in settings such as stochastic optimal control, where the PDEs of interest are tied to an underlying dynamical system. However, standard BSDE-based solvers have empirically been shown to underperform relative to PINNs in the literature. In this paper, we identify the root cause of this performance gap as a discretization bias introduced by the standard Euler-Maruyama (EM) integration scheme applied to one-step self-consistency BSDE losses, which shifts the optimization landscape off target. We find that this bias cannot be satisfactorily addressed through finer step-sizes or multi-step self-consistency losses. To properly handle this issue, we propose a Stratonovich-based BSDE formulation, which we implement with stochastic Heun integration. We show that our proposed approach completely eliminates the bias issues faced by EM integration. Furthermore, our empirical results show that our Heun-based BSDE method consistently outperforms EM-based variants and achieves competitive results with PINNs across multiple high-dimensional benchmarks. Our findings highlight the critical role of integration schemes in BSDE-based PDE solvers, an algorithmic detail that has received little attention thus far in the literature.

1 Introduction

Numerical solutions to partial differential equations (PDEs) are foundational to modeling problems across a diverse set of fields in science and engineering. However, due to the curse of dimensionality of traditional numerical methods, application of classic solvers to high dimensional PDEs is computationally intractable. In recent years, motivated by the success of deep learning methods, both Physics-Informed Neural Networks (PINNs) [1, 2] and backward stochastic differential equation (BSDE) methods [3–5] have emerged as promising alternatives to classic techniques.

Despite the widespread popularity of PINNs methods, in this paper we focus on the use of BSDE-based methods for solving high-dimensional PDEs. The key difference between PINNs and BSDE methods is that while PINNs minimize the PDE residual directly on randomly sampled collocation points, BSDE methods reformulate PDEs as forward-backward SDEs (FBSDEs) and simulate the resulting stochastic processes to minimize the discrepancy between predicted and terminal conditions at the end of the forward SDE trajectory [3], or across an intermediate time-horizon via *self-consistency* losses [4, 5]. BSDE methods are especially well-suited for high-dimensional problems where there is underlying dynamics—such as in stochastic optimal control or quantitative finance—as the crux of these methods involving sampling over stochastic trajectories rather than over bounded spatial domains. Furthermore, BSDE methods offers a significant advantage in problems where the governing equations of the PDE are unknown and can only be accessed through simulation [6], as in model-free optimal control (cf. Appendix C for more details). In contrast, PINNs

methods require explicit knowledge of the PDE equations, which may be either impractical to obtain for various tasks or require a separate model learning step within the training pipeline.

Surprisingly, despite the aforementioned benefits of BSDE methods compared with PINNs, a thorough comparison between the two techniques remains largely absent from the literature. One notable exception is recent work by [5] which finds that on several benchmark problems, PDE solutions found by BSDE-based approaches significantly underperform the corresponding PINNs solutions. To address this gap, they propose a hybrid *interpolating* loss between the PINNs and BSDE losses. While promising, their result has two key disadvantages. First, the underlying cause of the performance gap between BSDE and PINNs methods is not elucidated. Second, their method introduces a new hyper-parameter (the horizon-length controlling the level of interpolation) which must be tuned for optimal performance, adding complexity to the already delicate training process [7].

In this work, we identify the key source of the performance gap between BSDE and PINNs methods as the standard Euler-Maruyama (EM) scheme used in BSDE methods for stochastic integration. Although simple to implement, we show that the EM scheme introduces a significant discretization bias in *one-step* BSDE losses, resulting in a discrepancy between the optimization objective and the true solution. We furthermore show that the EM discretization bias can only be made arbitrarily small by using *multi-step* BSDE losses, which we show both theoretically and empirically comes at a significant cost in performance. Our analysis reveals the interpolating loss of [5] as a method to find (via hyper-parameter tuning) the best suitable horizon length (i.e., number of self-consistency steps).

As an alternative, we propose interpreting both the forward and backward SDEs as Stratonovich SDEs—as opposed to Itô SDEs—and utilizing the stochastic Heun integration scheme for numerical integration. We prove that the use of the stochastic Heun method completely eliminates the non-vanishing bias issues which occur in the EM formulation for one-step BSDE losses. This removes all performance tradeoffs in the horizon-length, allowing us to utilize single-step self-consistency losses. The result is a practical BSDE-based algorithm that is competitive with PINNs methods without the need for interpolating losses. Surprisingly, prior to our work the role of stochastic integration has received little attention in the BSDE literature; we hope that our results inspire further algorithmic and implementation level improvements for BSDE solvers.

2 Related Work

In recent years, PINNs [1, 2, 8–10] has emerged as a popular method for solving high dimensional PDEs. PINNs methods parameterize the PDE solution as a neural network and directly minimize the PDE residual as a loss function, provides a mesh-free method that can easily incorporate complex boundary conditions and empirical data. However, the PINNs approach remains an incomplete solution and still suffers from notable issues including optimization challenges [11–15], despite a concerted effort to remedy these difficulties [7, 11, 12, 16–20]. Hence, the application of PINNs as a general purpose solver for complex high-dimensional PDEs remains an active area of research.

On the other hand, a complementary line of work proposes methods based on BSDEs to solve high-dimensional PDEs [3–5, 21–25]. These approaches reformulate PDEs as forward-backward SDEs to derive a trajectory-based loss. While the original deep BSDE methods [3, 21, 22] learn separate neural networks to predict both the value and gradient at each discrete time-step, follow up work [4, 5] uses *self-consistency*, i.e., the residual of stochastic integration along BSDE trajectories, to form a loss. In this work, we exclusively focus on self-consistency BSDE losses, as they generalize the original method while allowing for a single network to parameterize the PDE solution for all space and time, similarly to PINNs. Similar self-consistency losses have also been recently utilized to learn solutions to Fokker-Planck PDEs [26–28]. Further discussions on related works and the paper’s relationship to recent BSDE-based methods can be found in Appendix B.

The main purpose of our work is to understand the performance differences between PINNs and BSDE methods on high-dimensional PDEs. The most relevant work to ours is [5], which to the best of our knowledge is the only work in the literature that directly compares PINNs and BSDE methods in a head-to-head evaluation. As discussed in Section 1, one of our main contributions shows that the gap in performance between PINNs and BSDE methods observed in [5] is due to the choice of stochastic integration. Previous work [29] studying stochastic Runge-Kutta discretizations for BSDEs methods considers only the original BSDE losses instead of self-consistency methods, and hence does not uncover the issues identified in our work.

3 Background and Problem Setup

We consider learning approximate solutions to the following non-linear boundary value PDEs

$$R[u](x, t) := \partial_t u(x, t) + \frac{1}{2} \text{tr}(H(x, t) \cdot \nabla^2 u(x, t)) + \langle f(x, t), \nabla u(x, t) \rangle - h[u](x, t) = 0, \quad (3.1)$$

over domain $x \in \Omega \subseteq \mathbb{R}^d$, $t \in \mathcal{I} := [0, T]$ with boundary conditions (a) $u(x, T) = \phi(x) \ \forall x \in \Omega$ and (b) $u(x, t) = \phi_b(x, t) \ \forall x \in \partial\Omega, t \in \mathcal{I}$. Here, $u : \Omega \times \mathcal{I} \mapsto \mathbb{R}$ is a candidate PDE solution, $f : \Omega \times \mathcal{I} \mapsto \mathbb{R}^d$ is a vector-field, $h[u] := h_0(x, t, u(x, t), \nabla u(x, t))$ for some $h_0 : \Omega \times \mathcal{I} \times \mathbb{R} \times \mathbb{R}^d \mapsto \mathbb{R}$ captures the non-linear terms, $H(x, t) = g(x, t)g(x, t)^\top \in \mathbb{R}^{d \times d}$ for $g(x, t) \in \mathbb{R}^{d \times d}$ is a positive definite matrix-valued function, $\phi : \Omega \times \mathcal{I}$ and $\phi_b : \partial\Omega \times \mathcal{I}$ are boundary conditions, and both ∇ and ∇^2 denote spatial gradients and Hessians, respectively. For expositional clarity, we will assume that $\Omega = \mathbb{R}^d$ and drop the second boundary condition (b), noting that all subsequent arguments can be extended in a straightforward manner for bounded domains Ω .

Physics-Informed Neural Networks (PINNs). Under the assumption of knowledge of the operator $R[u]$ and the boundary condition ϕ , the standard PINNs methodology [1, 2, 8–10] for solving (3.1) works by parameterizing the solution $u(x, t)$ in a function class $\mathcal{U} := \{u_\theta(x, t) \mid \theta \in \Theta\}$ (e.g., θ represents the weights of a neural network), and minimizing the PINNs loss over $\mathcal{U}$:¹

$$L_{\text{PINNs}}(\theta; \lambda) := \mathbb{E}_{(x, t) \sim \mu}[(R[u_\theta](x, t))^2] + \lambda \cdot \mathbb{E}_{x \sim \mu'}[(u_\theta(x, T) - \phi(x))^2], \quad (3.2)$$

where μ is a measure over $\Omega \times \mathcal{I}$ and μ' is a measure over Ω . The choice of measures μ, μ' , in addition to the relative weight λ are hyper-parameters which must be carefully selected by the user. To simplify exposition further, we will assume that each $u_\theta(x, t) \in \mathcal{U}$ satisfies $u(\cdot, T) = \phi$ (e.g., as in [30, 31]), and hence the PINNs loss simplifies further to $L_{\text{PINNs}}(\theta) = \mathbb{E}_{(x, t) \sim \mu}[(R[u_\theta](x, t))^2]$.

Backward SDEs and self-consistency losses. While the PINNs loss has received much attention in the literature, a separate line of work has advocated for an alternative approach to solving PDEs based on backward SDEs [3, 4, 21–23]. The key idea is that given the *forward* (Itô) SDE:

$$dX_t = f(X_t, t)dt + g(X_t, t)dB_t, \quad X_0 = x_0, \quad (3.3)$$

where $(B_t)_{t \geq 0}$ is standard Brownian motion in $\mathbb{R}^d$, the corresponding *backward* (Itô) SDE:

$$dY_t = h(X_t, t, Y_t, Z_t)dt + Z_t^\top g(X_t, t)dB_t, \quad Y_T = \phi(X_T), \quad (3.4)$$

is solved by setting $Y_t = u(X_t, t)$ and $Z_t = \nabla u(X_t, t)$, where u is a solution to the PDE (3.1); this equivalence is readily shown with Itô's lemma. The relationship between the forward and backward SDE has motivated several different types of BSDE loss functions for solving (3.1). In this work, we focus on BSDE losses based on *self-consistency* [4, 5], which uses the residual of stochastic integration along the BSDE trajectories as supervision. Self-consistency losses are more practical than other BSDE variants as only one network is required and the weights can be shared across time (unlike e.g., the original BSDE losses [3, 21] which learn N models to predict both Y_t and Z_t at N discretization points, and require retraining for every new initial condition x_0). Specifically, we consider the following H -horizon (for $N = T/H \in \mathbb{N}_+$ and $t_n = nH$) self-consistency BSDE loss:²

$$L_{\text{BSDE}, H}(\theta) := \mathbb{E}_{x_0, B_t} \frac{1}{NH^2} \sum_{n=0}^{N-1} (u_\theta(X_{t_{n+1}}, t_{n+1}) - u_\theta(X_{t_n}, t_n) - S_\theta(t_n, t_{n+1}))^2, \quad (3.5)$$

where $S_\theta(t_0, t_1) := \int_{t_0}^{t_1} h_\theta(X_t, t)dt - \int_{t_0}^{t_1} \nabla u_\theta(X_t, t)^\top g(X_t, t)dB_t$ with $h_\theta(x, t) := h[u_\theta](x, t)$, and $x_0 \sim \mu_0$ is drawn from a distribution μ_0 over initial conditions for the forward SDE (3.3).

Euler-Maruyama integration. Unlike the PINNs loss (3.2), the BSDE loss (3.5) must be discretized with an appropriate stochastic integrator. The standard choice is to use the Euler-Maruyama (EM) method, selecting $N \in \mathbb{N}_+$ to define a step-size $\tau := T/N$ and time-points $t_n := n\tau$, and integrating the forward and backward SDEs as follows:

$$\hat{X}_{n+1} = \hat{X}_n + \tau f(\hat{X}_n, t_n) + \sqrt{\tau} g(\hat{X}_n, t_n) w_n, \quad w_n \sim \mathcal{N}(0, I_d), \quad \hat{X}_0 = x_0,$$

¹We leave consideration of the PINNs loss with non-square losses (e.g., [18]) to future work.

²We note that the most general form of self-consistency losses are due to [5] and take on the form $L_{\text{BSDE}}(\theta; \rho) = \mathbb{E}_{\substack{x_0 \sim \mu_0, \\ (t_s, t_f) \sim \rho, B_t}} \frac{1}{\Delta_t^2} (u_\theta(X_{t_f}, t_f) - u_\theta(X_{t_s}, t_s) - S_\theta(t_s, t_f))^2$ involving a time-pair distribution ρ over $\mathcal{I}^2$, where $\Delta_t := t_f - t_s$. We choose to present (3.5) as it more closely aligns with the discrete losses.

$$\hat{Y}_{n+1}^\theta = \hat{Y}_n^\theta + \tau h_\theta(\hat{X}_n, t_n) + \sqrt{\tau} \nabla u_\theta(\hat{X}_n, t_n)^\top g(\hat{X}_n, t_n) w_n, \quad \hat{Y}_0^\theta = u_\theta(x_0, 0). \quad (3.6)$$

With this discretization, the k -step EM-BSDE loss (for $N/k \in \mathbb{N}_+$) for $L_{\text{BSDE}}(\theta)$ is:

$$L_{\text{EM},k,\tau}(\theta) := \mathbb{E}_{x_0, w_n} \frac{k}{N\tau^2} \sum_{n=0}^{\frac{N}{k}-1} \left(u_\theta(\hat{X}_{(n+1)k}, t_{(n+1)k}) - u_\theta(\hat{X}_{nk}, t_{nk}) - (\hat{Y}_{(n+1)k}^\theta - \hat{Y}_{nk}^\theta) \right)^2. \quad (3.7)$$

In the *one-step* (or *single-step*) setting $k = 1$, (3.7) reduces to the self-consistent BSDE loss from [4], also discussed in [32, 33]; we use the shorthand $L_{\text{EM},\tau}(\theta)$ to denote this setting. We refer to the $k > 1$ case generally as *multi-step*, which is a form of interpolating loss from [5]. Another notable case is when $k = N$, which recovers the *full-horizon* losses used in the original BSDE works [3, 21].

4 Analysis of One-Step Self-Consistency Losses

In this section we conduct an analysis of both EM and Heun stochastic integration applied to the one-step self-consistency BSDE loss.

The Hölder space $C^{k,1}$. Let $f : M \mapsto M'$, where both M, M' are subsets of Euclidean space (with possibly different dimension). We say that f is $C^{k,1}(M, M')$ ($C^{k,1}$ when M, M' are clear) if f is both bounded and k -times continuously differentiable on M , and all j -th derivatives of f for $j \in \{0, \dots, k\}$ are Lipschitz continuous. The Hölder norm $\|f\|_{C^{k,1}(M, M')}$ is the smallest bound possible on $\|f\|$ and all the Lipschitz constants for $D^j f$, $j \in \{0, \dots, k\}$.

4.1 Analysis of Euler-Maruyama for BSDE

We first illustrate the bias when using EM to integrate the single-step consistency loss. To do this, we define the point-wise EM loss at resolution τ for a fixed $(x, t) \in \mathbb{R}^d \times \mathcal{I}$ as:

$$\ell_{\text{EM},\tau}(\theta, x, t) := \mathbb{E}_w \left(u_\theta(\hat{x}_{t+\tau}, t + \tau) - u_\theta(x, t) - \tau h_\theta(x, t) - \sqrt{\tau} \langle \nabla u_\theta(x, t), g(x, t) w \rangle \right)^2, \quad (4.1)$$

$$\hat{x}_{t+\tau} := x + \tau f(x, t) + \sqrt{\tau} g(x, t) w, \quad w \sim \mathcal{N}(0, I_d).$$

The point-wise EM loss (4.1) is related to the one-step EM-BSDE loss via $L_{\text{EM},\tau}(\theta) = \frac{1}{N\tau^2} \sum_{n=0}^{N-1} \mathbb{E}_{\hat{X}_n} [\ell_{\text{EM},\tau}(\theta, \hat{X}_n, t_n)]$. Our first result shows that the dominant error term (in τ) of the loss (4.1) suffers from an additive bias term that is introduced as a result of the EM integration.

Lemma 4.1. Suppose that f, g are bounded and u_θ is $C^{2,1}$. We have that

$$\tau^{-2} \cdot \ell_{\text{EM},\tau}(\theta, x, t) = (R[u_\theta](x, t))^2 + \frac{1}{2} \text{tr} \left[(H(x, t) \cdot \nabla^2 u_\theta(x, t))^2 \right] + O(\tau^{1/2}), \quad (4.2)$$

where the $O(\cdot)$ hides factors depending on d , the bounds on f, g , and $\|u_\theta\|_{C^{2,1}}$.

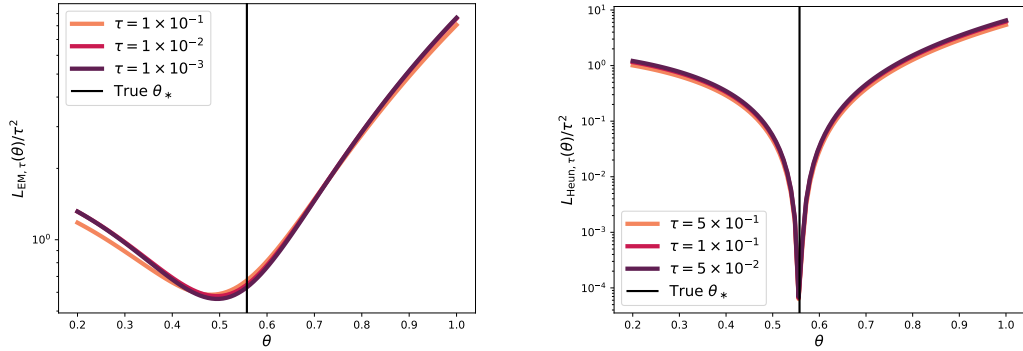
Lemma 4.1 can further be used, in conjunction with standard results on the order 1/2 strong convergence of EM integration [34], to show the following statement regarding the full loss $L_{\text{EM},\tau}(\theta)$.

Theorem 4.2. Suppose that $f, g, h_\theta \in C^{0,1}$, $u_\theta \in C^{2,1}$, and $\tau \leq 1$. We have that:

$$L_{\text{EM},\tau}(\theta) = \frac{1}{T} \int_0^T \mathbb{E} \left[(R[u_\theta](X_t, t))^2 + \frac{1}{2} \text{tr} \left[(H(X_t, t) \cdot \nabla^2 u_\theta(X_t, t))^2 \right] \right] dt + O(\tau^{1/2}), \quad (4.3)$$

where the $O(\cdot)$ hides constants that depend on d, T , and the Hölder norms of f, g, h_θ , and u_θ .

Theorem 4.2 implies that even if the function class $\mathcal{U}$ is expressive enough to contain a PDE solution u_{θ_*} to (3.1) satisfying $R[u_{\theta_*}] = 0$, in general $L_{\text{EM},\tau}(\theta_*) > \inf_{u_\theta \in \mathcal{U}} L_{\text{EM},\tau}(\theta)$, and hence optimizing $L_{\text{EM},\tau}(\theta)$ can lead to sub-optimal solutions *even in the limit of infinite simulated data*. Furthermore, this bias *cannot* be resolved by simply reducing the step-size τ , since the PDE residual term and the bias term are both the same order (cf. (4.3)). We illustrate this with a one-dimensional example in Figure 1a. Proofs for both Lemma 4.1 and Theorem 4.2 are given in Appendix D.2.



(a) Plot of $L_{EM,\tau}(\theta)$ at $\tau \in \{10^{-1}, 10^{-2}, 10^{-3}\}$ levels of discretization. (b) Plot of $L_{Heun,\tau}(\theta)$ at $\tau \in \{5 \times 10^{-1}, 10^{-1}, 5 \times 10^{-2}\}$ levels of discretization.

Figure 1: A plot of both $L_{EM,\tau}(\theta)$ and $L_{Heun,\tau}(\theta)$ at various levels of discretization. The PDE is a one dimensional Linear Quadratic Regulator HJB equation, where θ parameterizes a quadratic value function.

4.2 Stratonovich BSDEs and Stochastic Heun Integration

Our next step is to derive a new BSDE loss based on Heun integration. The starting point is to interpret the forward SDE as a Stratonovich SDE (in contrast to (3.3) which is an Itô SDE):

$$dX_t^\bullet = f(X_t^\bullet, t)dt + g(X_t^\bullet, t) \circ dB_t, \quad X_0^\bullet = x_0. \quad (4.4)$$

For any u that satisfies (3.1), we have $du(X_t^\bullet, t) = h^\bullet[u](X_t^\bullet, t)dt + \nabla u(X_t^\bullet, t)^\top g(X_t^\bullet, t) \circ dB_t$ with $h^\bullet[u](x, t) := h[u](x, t) - \frac{1}{2} \text{tr}(H(x, t) \nabla^2 u(x, t))$ by the Stratonovich chain rule, which motivates the following H -horizon self-consistency Stratonovich BSDE loss:

$$L_{S\text{-BSDE}, H}(\theta) := \mathbb{E}_{x_0, B_t} \frac{1}{NH^2} \sum_{n=0}^{N-1} \left(u_\theta(X_{t_{n+1}}^\bullet, t_{n+1}) - u_\theta(X_{t_n}^\bullet, t_n) - S_\theta^\bullet(t_n, t_{n+1}) \right)^2, \quad (4.5)$$

with $S_\theta^\bullet(t_0, t_1) := \int_{t_0}^{t_1} h_\theta^\bullet(X_t^\bullet, t)dt + \int_{t_0}^{t_1} \nabla u_\theta(X_t^\bullet, t)^\top g(X_t^\bullet, t) \circ dB_t$ where $h_\theta^\bullet(x, t) := h^\bullet[u_\theta](x, t)$. As (4.5) utilizes Stratonovich integration, the Euler-Maruyama scheme cannot be used for integration, as it converges to the Itô solution. Hence, we will consider the stochastic Heun integrator [34, 35] which has the favorable property of converging to the Stratonovich solution. We proceed first by defining the augmented forward and backward SDE process $Z_t^{\bullet, \theta} := (X_t^\bullet, Y_t^{\bullet, \theta})$:

$$\begin{aligned} d \begin{bmatrix} X_t^\bullet \\ Y_t^{\bullet, \theta} \end{bmatrix} &= \begin{bmatrix} f(X_t^\bullet, t) \\ h_\theta^\bullet(X_t^\bullet, t) \end{bmatrix} dt + \begin{bmatrix} g(X_t^\bullet, t) \\ \nabla u_\theta(X_t^\bullet, t)^\top g(X_t^\bullet, t) \end{bmatrix} \circ dB_t, \quad \begin{bmatrix} X_0 \\ Y_0^{\bullet, \theta} \end{bmatrix} = \begin{bmatrix} x_0 \\ u_\theta(x_0, 0) \end{bmatrix}, \\ &=: F_\theta(Z_t^{\bullet, \theta}, t)dt + G_\theta(Z_t^{\bullet, \theta}, t) \circ dB_t. \end{aligned} \quad (4.6)$$

The augmented SDE is discretized as follows using the stochastic Heun scheme:

$$\begin{aligned} \bar{Z}_{n+1}^{\bullet, \theta} &= \hat{Z}_n^{\bullet, \theta} + \tau F_\theta(\hat{Z}_n^{\bullet, \theta}, t_n) + \sqrt{\tau} G_\theta(\hat{Z}_n^{\bullet, \theta}, t_n) w_n, \quad w_n \sim \mathcal{N}(0, I_d), \\ \hat{Z}_{n+1}^{\bullet, \theta} &= \hat{Z}_n^{\bullet, \theta} + \frac{\tau}{2} \left(F_\theta(\hat{Z}_n^{\bullet, \theta}, t_n) + F_\theta(\bar{Z}_{n+1}^{\bullet, \theta}, t_{n+1}) \right) + \frac{\sqrt{\tau}}{2} \left(G_\theta(\hat{Z}_n^{\bullet, \theta}, t_n) + G_\theta(\bar{Z}_{n+1}^{\bullet, \theta}, t_{n+1}) \right) w_n, \end{aligned} \quad (4.7)$$

with $\hat{Z}_0^{\bullet, \theta} = (x_0, u_\theta(x_0, 0))$. This gives rise to the k -step Heun-BSDE loss defined as:

$$L_{\text{Heun}, k, \tau}(\theta) := \beta_k \mathbb{E}_{x_0, w_n} \sum_{n=0}^{\frac{N}{k}-1} \left(u_\theta(\hat{X}_{(n+1)k}^\bullet, t_{(n+1)k}) - u_\theta(\hat{X}_{nk}^\bullet, t_{nk}) - (\hat{Y}_{(n+1)k}^{\bullet, \theta} - \hat{Y}_{nk}^{\bullet, \theta}) \right)^2, \quad (4.8)$$

with $\beta_k := \frac{k}{N\tau^2}$. For the one-step $k = 1$ case, we use the shorthand $L_{\text{Heun}, \tau}(\theta)$. We now show that the one-step Heun-BSDE loss $L_{\text{Heun}, \tau}(\theta)$ avoids the undesirable bias term which appears for the one-step EM-BSDE loss $L_{EM, \tau}(\theta)$ (cf. Section 4.1). To do this, analogous to (4.1), we define the point-wise Heun loss at resolution τ for a fixed (x, t) as:

$$\ell_{\text{Heun}, \tau}(\theta, x, t) := \mathbb{E}_w (u_\theta(\hat{x}_{t+\tau}, t + \tau) - \hat{y}_{t+\tau}^\theta)^2,$$

$$\begin{aligned}\bar{z}_{t+\tau}^\theta &= z_t^\theta + \tau F_\theta(z_t^\theta, t) + \sqrt{\tau} G_\theta(z_t^\theta, t) w, \quad z_t^\theta = (x, u_\theta(x, t)), \\ \bar{z}_{t+\tau}^\theta &= z_t^\theta + \frac{\tau}{2} (F_\theta(z_t^\theta, t) + F_\theta(\bar{z}_{t+\tau}^\theta, t + \tau)) + \frac{\sqrt{\tau}}{2} (G_\theta(z_t^\theta, t) + G_\theta(\bar{z}_{t+\tau}^\theta, t + \tau)) w,\end{aligned}$$

noting that $\bar{z}_{t+\tau}^\theta = (\hat{x}_{t+\tau}, \hat{y}_{t+\tau}^\theta)$. Similar to $L_{\text{EM},\tau}(\theta)$, we have the following identity $L_{\text{Heun},\tau}(\theta) = \frac{1}{N\tau^2} \sum_{n=0}^{N-1} \mathbb{E}_{\hat{X}_n^\bullet} [\ell_{\text{Heun},\tau}(\theta, \hat{X}_n^\bullet, t_n)]$. Our next result illustrates that the point-wise Heun loss avoids the issues identified with the point-wise EM loss in Lemma 4.1.

Lemma 4.3. *Suppose that f , g , and h_θ are all in $C^{1,1}$, and u_θ is in $C^{3,1}$. We have that*

$$\tau^{-2} \cdot \ell_{\text{Heun},\tau}(\theta, x, t) = (R[u_\theta](x, t))^2 + O(\tau^{1/2}), \quad (4.9)$$

where the $O(\cdot)$ hides factors depending on d and the Hölder norms of f , g , h_θ , and u_θ .

Furthermore, analogously to Theorem 4.2, we can utilize Lemma 4.3 in conjunction with the order $1/2$ strong convergence of stochastic Heun (cf. Appendix F) to the Stratonovich solution to show the following relationship for the full loss $L_{\text{Heun},\tau}(\theta)$.

Theorem 4.4. *Suppose that f , g , and h_θ are all in $C^{1,1}$, $u_\theta \in C^{3,1}$, and $\tau \leq 1$. We have that*

$$L_{\text{Heun},\tau}(\theta) = \frac{1}{T} \int_0^T \mathbb{E}[(R[u_\theta](X_t^\bullet, t))^2] dt + O(\tau^{1/2}), \quad (4.10)$$

where the $O(\cdot)$ hides factors depending on d , T , and the Hölder norms of f , g , h_θ , and u_θ .

Therefore, unlike the situation with EM integration in (4.3), any additional bias terms only enter through a $O(\tau^{1/2})$ term which is of higher order than the leading PDE residue term (cf. (4.10)). In Figure 1b, we show the plot of $L_{\text{Heun},\tau}(\theta)$ on the same HJB PDE problem as in Figure 1a, and show that the bias issue in $L_{\text{EM},\tau}(\theta)$ is now resolved. The proofs of both Lemma 4.3 and Theorem 4.4 are given in Appendix D.3.

5 Trade-offs for Multi-Step BSDE Losses

In Section 4, we conducted a thorough analysis of the one-step self-consistency losses $L_{\text{EM},\tau}(\theta)$ and $L_{\text{Heun},\tau}(\theta)$. We now consider the other extreme: the full-horizon ($k = N$) losses $L_{\text{EM},N,\tau}(\theta)$ and $L_{\text{Heun},N,\tau}(\theta)$. The intermediate multi-step regime [5], where $1 < k < N$, serves as an extension to the cited method and is studied experimentally in Section 6. Due to space constraints, we defer the precise theorem statements arising from our analysis, in addition to the proofs, to Appendix E.

BSDE loss and Euler-Maruyama discretization. We start with the H -horizon BSDE loss (3.5). Using Jensen's inequality, we show (Proposition E.3) the relationship $L_{\text{BSDE},T}(\theta) \leq L_{\text{BSDE},\tau}(\theta) + O(\tau^{1/2})$. Thus, at the SDE level, the benefits of using the full-horizon loss $L_{\text{BSDE},T}(\theta)$ over the τ -horizon loss $L_{\text{BSDE},\tau}(\theta)$ are not clear, given that (a) the full-horizon loss is dominated by the latter τ -horizon loss (up to an order $\tau^{1/2}$ term), and (b) $L_{\text{BSDE},\tau}(\theta)$ does indeed vanish for an optimal $\theta_\star$.

The situation becomes more complex when factoring in EM discretization. Using the order $1/2$ strong convergence of EM, we show (Proposition E.4) that $L_{\text{EM},N,\tau}(\theta) = L_{\text{BSDE},T}(\theta) + O(\tau^{1/2})$. On the other hand, from Theorem 4.2 we also show (Proposition E.5) that $L_{\text{EM},\tau}(\theta) = L_{\text{BSDE},\tau}(\theta) + \text{Bias}(\theta) + O(\tau^{1/2})$ with the bias term $\text{Bias}(\theta) := \frac{1}{2T} \int_0^T \mathbb{E} \text{tr}((H(X_t, t) \cdot \nabla^2 u_\theta(X_t, t))^2) dt$ not vanishing as $\tau \rightarrow 0$. Hence, the loss $L_{\text{EM},N,\tau}(\theta)$ presents an advantage over $L_{\text{EM},\tau}(\theta)$ for sufficiently small discretization sizes in terms of bias. However, the inequality $L_{\text{BSDE},T}(\theta) \leq L_{\text{BSDE},\tau}(\theta) + O(\tau^{1/2})$ still holds, meaning that while $L_{\text{EM},N,\tau}(\theta)$ does not suffer from the bias issues identified in $L_{\text{EM},\tau}(\theta)$, the trade-off is that the loss $L_{\text{BSDE},T}(\theta)$ it approximates without bias is nearly dominated by another loss $L_{\text{BSDE},\tau}(\theta)$; this is precisely the loss that $L_{\text{EM},\tau}(\theta)$ attempts to approximate, but it does so in a way that introduces an irreducible bias term $\text{Bias}(\theta)$. Thus, neither of the EM-BSDE losses for $k = 1$ nor $k = N$ provides a completely satisfactory solution. In Section 6.2, we illustrate these issues empirically. Furthermore, in light of this analysis, we can interpret the interpolating loss of [5] as attempting to resolve this trade-off by finding the best intermediate multi-step $k \in \{1, \dots, N\}$.

Stratonovich BSDE and Heun discretization. In the setting of the Stratonovich BSDE and the Heun-BSDE loss, we first show (Proposition E.8) that $L_{\text{S-BSDE},T}(\theta) \leq L_{\text{S-BSDE},\tau}(\theta) + O(\tau^{1/2})$ holds at

the SDE level, analogous to the relationship between $L_{\text{BSDE},T}(\theta)$ and $L_{\text{BSDE},\tau}(\theta)$. Next, we use the order $1/2$ strong convergence of Heun to show (Proposition E.9) that $L_{\text{Heun},N,\tau}(\theta) = L_{\text{S-BSDE},T}(\theta) + O(\tau^{1/2})$; again analogous to the relationship between $L_{\text{EM},N,\tau}(\theta)$ and $L_{\text{BSDE},T}(\theta)$. However, unlike the one-step EM case, using Theorem 4.4 we show (Proposition E.10) that $L_{\text{Heun},\tau}(\theta) = L_{\text{S-BSDE},\tau}(\theta) + O(\tau^{1/2})$, from which we conclude $L_{\text{Heun},N,\tau}(\theta) \leq L_{\text{Heun},\tau}(\theta) + O(\tau^{1/2})$. Thus—unlike the EM setting—the relationship between the full-horizon and one-step case at both the SDE level and Heun-BSDE level is the same, suggesting questionable benefits of $L_{\text{Heun},N,\tau}(\theta)$ over $L_{\text{Heun},\tau}(\theta)$. In Section 6.2, we show that this conclusion is indeed reflected in practice.

6 Experiments

In this section, we compare the proposed Heun-based BSDE method against both standard PINNs, a variant of PINNs which uses the forward SDE to sample collocation points, and standard EM-based BSDE solvers on various high-dimensional PDE problems. Specifically, we compare the methods:

- (a) **PINNs:** The standard PINNs loss $L_{\text{PINNs}}(\theta)$ from (3.2) is minimized. Since we consider unbounded domains, the collocation measure μ over (x, t) is chosen by fitting a normal distribution over the spatial dimensions of the forward SDE trajectories prior to training.
- (b) **FS-PINNs:** The standard PINNs loss (3.2) is again minimized, where the measure μ over space-time is chosen by directly sampling trajectories from the forward SDE (3.3).
- (c) **EM-BSDE:** The self-consistency loss discretized with the standard Euler-Maruyama (EM) scheme, i.e., $L_{\text{EM},k,\tau}(\theta)$ from (3.7) as described in Section 3.
- (d) **EM-BSDE (NR):** A variant of EM-BSDE where we use the BSDE to propagate Y_t instead of setting it directly to $u_\theta(X_t, t)$ [cf. 4, 5]. We refer to this variant as the *no-reset* (NR) variant. Specifically, the backward SDEs in (3.6) is integrated as (starting from $\hat{Y}_0^\theta = u_\theta(x_0, 0)$):

$$\hat{Y}_{n+1}^\theta = \hat{Y}_n^\theta + \tau h_0(\hat{X}_n, t_n, \hat{Y}_n^\theta, \nabla u_\theta(\hat{X}_n, t_n)) + \sqrt{\tau} \nabla u_\theta(\hat{X}_n, t_n)^\top g(\hat{X}_n, t_n) w_n, \quad (6.1)$$

with $w_n \sim \mathcal{N}(0, I_d)$. The loss $L_{\text{EM},k,\tau}(\theta)$ remains the same except replacing (3.6) with (6.1).

- (e) **Heun-BSDE (Ours):** The self-consistency loss discretized with stochastic Heun integration, i.e., $L_{\text{Heun},k,\tau}(\theta)$ from (4.8) as described in Section 4.2.

We evaluate these methods on three PDE benchmark problems: (i) a Hamilton-Jacobi-Bellman (HJB) equation [4], (ii) a Black-Scholes-Barenblatt (BSB) equation [4], and (iii) a fully-coupled FBSDE from Bender & Zhang (BZ) [36]; the PDEs are detailed in Appendix G.1. In addition, we evaluate the methods on an optimal control pendulum swing-up problem to demonstrate application to a non-linear control problem (see Appendix G.6). To evaluate model performance, the analytical solution (available for all PDEs under consideration) is compared with the model output along 5 forward SDE trajectories, using the *relative L_2 error* (RL2) metric:

$$\text{RL2} := \sqrt{\frac{\sum_{i=0}^N (u_{\text{ref}}(X_{t_i}, t_i) - u_{\text{pred}}(X_{t_i}, t_i))^2}{\sum_{i=0}^N u_{\text{ref}}^2(X_{t_i}, t_i)}}. \quad (6.2)$$

Unless otherwise noted, we set $T = 1$ and $N = 50$ (i.e., $\tau = 0.02$). Model architectures and training details are described in Appendix G. Additionally, the code to reproduce our experiments is available at: <https://github.com/sungje-park/heunbsde>. For what follows, we report two main sets of results on (i) one-step self-consistency losses (Section 6.1) and (ii) multi-step self-consistency losses (Section 6.2).

Efficient sub-sampling BSDE implementation. In our experimental results, we consider both a full FSDE rollout algorithm, in addition to a batched, sub-sampled FSDE rollout variation of FS-PINNs, EM-BSDE, and Heun-BSDE, which we find performs similarly to the original algorithm while providing significant speed improvements. We sketch the details of the sub-sampled FSDE variation in Algorithm 1; further details on the two algorithms can be found in Appendix G.4.

6.1 One-Step Self-Consistency Losses

For our first set of results, we compare the PINNs baselines and the one-step ($k = 1$) EM-BSDE baseline with our one-step Heun-BSDE method. We solve each PDE instantiated with a state space

Algorithm 1 Batched, Sub-sampling BSDE Algorithm (Simplified)

Input: Neural network $\hat{u}_\theta(x, t)$, parameters θ , terminal function ϕ , time step Δt , trajectory length N , evaluation batch size B .

- 1: Sample initial state: $(x[0], t[0]) = (x_0, 0)$, with $x_0 \sim \mu$
- 2: Sample Brownian noise: $\xi[0 : N - 1] \sim \mathcal{N}(0, I_d)$
- 3: Evaluate network at initial state: $(u, u_x) = (\hat{u}_\theta(x[0], t[0]), \nabla_x \hat{u}_\theta(x[0], t[0]))$
- 4: /* Forward SDE rollout */
- 5: **for** $i = 0, \dots, N - 1$ **do**
- 6: /* Use either EM or Heun integration */
- 7: Propagate forward state (with (u, u_x) if coupled): $x[i + 1] = x[i] + \Delta x$
- 8: Propagate time: $t[i + 1] = t[i] + \Delta t$
- 9: Evaluate network at new state: $(u, u_x) = (\hat{u}_\theta(x[i + 1], t[i + 1]), \nabla_x \hat{u}_\theta(x[i + 1], t[i + 1]))$
- 10: **end for**
- 11: Stop gradient: $x[0 : N] = \text{SG}(x[0 : N])$
- 12: Random sub-sampling: $(x_i, x_{i+1}, t_i, t_{i+1}) = \text{perm}(x_i, x_{i+1}, t_i, t_{i+1})[0 : B - 1]$
- 13: /* Use either EM or Heun integration */
- 14: Propagate backward SDE at batched points: $y_{i+1} = u_i + \Delta y$
- 15: /* Use PINNs loss instead for FS-PINNs */
- 16: Compute self-consistency loss: $\mathcal{L}_{\text{sr}} = \frac{N}{B} \sum_{i=0}^{B-1} (\hat{u}_{i+1} - y_{i+1})^2$
- 17: Compute terminal loss: $\mathcal{L}_\phi = (u + \phi)^2 + \|u_x - \nabla_x \phi\|^2$

Cases	PINNs	FS-PINNs	EM-BSDE (NR)	EM-BSDE	Heun-BSDE
Full Algorithm (cf. Algorithm 2) Results					
100D HJB	.1260 $\pm$.0107	.0737 $\pm$.0110	.5214 $\pm$.0452	.3626 $\pm$.0113	.0493 $\pm$.0109
100D BSB	1.5066 $\pm$.2349	.0497 $\pm$.0031	.1855 $\pm$.0078	.3735 $\pm$.0470	.0535 $\pm$.0113
100D BZ	-	-	-	3.1259 $\pm$.1807	3.5619 $\pm$.2716
10D BZ	3.8566 $\pm$.0310	.0351 $\pm$.0041	.1309 $\pm$.0311	.1903 $\pm$.0066	.0228 $\pm$.0016
Batched Algorithm (cf. Algorithm 3) Results					
100D HJB	.1362 $\pm$.0276	.1828 $\pm$.0774	.5214 $\pm$.0452	.3831 $\pm$.0084	.0573 $\pm$.0106
100D BSB	3.0488 $\pm$ 1.5625	.0851 $\pm$.0027	.1855 $\pm$.0078	.3668 $\pm$.0244	.0472 $\pm$.0076
100D BZ	-	5.4502 $\pm$.1351	-	5.7330 $\pm$.2342	1.7973 $\pm$.1108
10D BZ	3.8495 $\pm$.1562	.0270 $\pm$.0017	.1309 $\pm$.0311	.1933 $\pm$.0022	.0236 $\pm$.0031

Table 1: Summary of RL2 errors averaged over three different initialization random seeds, $\pm$ one standard deviation. Settings that failed to converge to a satisfactory solution are denoted with -. The first set of results correspond to the full algorithm (see Appendix G.4 for a detailed description), whereas the second set of results correspond to the batched algorithm (cf. Algorithm 1).

of 100 dimensions using three different initialization seeds for training. The results are reported in

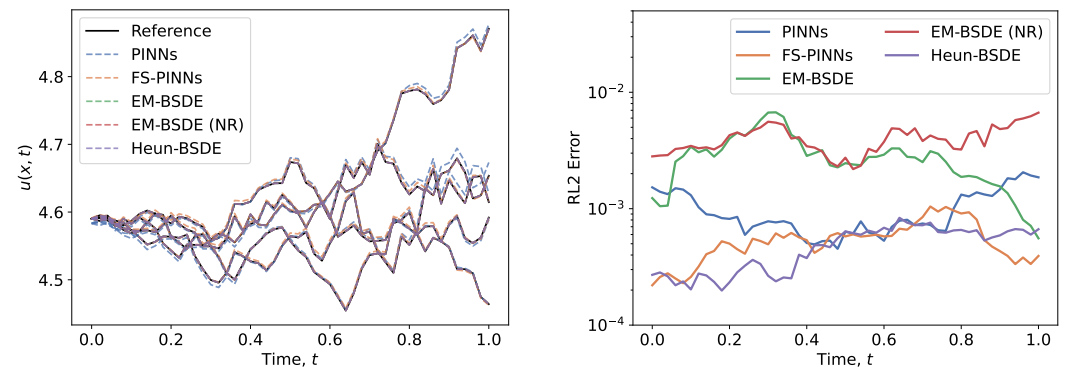


Figure 2: A plot of the 100D HJB reference and learned solutions for each model and the associated RL2 errors.

Table 1, which shows that for nearly all the cases, the Heun-BSDE method outperforms EM-BSDE methods (i.e., lower RL2 error) as predicted by our analysis in Section 4. Furthermore, Figure 2 illustrates the performance of all methods across time $t \in [0, 1]$ on the 100D HJB case, which also highlights the low RL2 error of the Heun-BSDE method. The one exception to the trend is the 100D BZ case, where all methods failed to converge to a high-quality solution. We hypothesize due to the high dimensionality of the problem involving fully-coupled SDEs, the optimization landscape for all methods is too complex to recover high-fidelity solutions. To evaluate this hypothesis, we further reduce the dimensionality of the BZ problem to 10D, which restores the relative performance of all methods (cf. Table 1, last row). Another key observation from Table 1 is that FS-PINNs and Heun-BSDE perform similarly across all cases, showing that parity between the BSDE and PINNs is restored through Heun integration. Finally, we note that the performance of PINNs is quite poor in comparison to FS-PINNs, which illustrates the relative impact of the sampling distribution μ for PINNs methods (cf. [37–39]).

To show that the gap between EM and Heun performance cannot be closed with finer discretization meshes, we re-run the 10D BSB example at varying discretization sizes. The results are reported in Figure 3, which show that the EM-BSDE methods only experience minimal improvement with smaller discretization size compared with the Heun-BSDE method. Figure 3 corroborates the findings in Section 4, which shows that the one-step EM-BSDE loss contains a bias term of the same order as the residual error which is not present in the one-step Heun-BSDE loss.

Computational considerations. Although Heun-BSDE outperforms EM-BSDE, it comes at a computational cost. In Table 2, we report the average training time for each method on a single NVIDIA A100 GPU; on average Heun-BSDE is approximately 6x slower than the EM-BSDE method for the batched algorithm. There are two major factors to this overhead. First, for the specific elliptic/parabolic PDEs we consider in (3.1), EM-BSDE does not require the computationally expensive Hessian computation $\nabla^2 u(x, t)$, which PINNs, FS-PINNs, and Heun-BSDE all do require. However, this does not necessarily hold true for all PDEs (e.g., Appendix G.6). Second, the Heun integration requires approximately double the compute of the EM integration—this is clearly reflected in the overhead between FS-PINNs and Heun-BSDE.

Method	Full	Batched
PINNs	1x	1x
FS-PINNs	2.64x	1.14x
EM-BSDE	2.83x	0.34x
Heun-BSDE	36.37x	2.03x

Table 2: A table of average training time overhead relative to PINNs for both the full and batched algorithm runs.

In Figure 4, we show the runtime-normalized RL2 performance demonstrating that while EM-BSDE shows strong convergence at first, its performance does not improve with more compute. Conversely, both FS-PINNs and Heun-BSDE achieve similar RL2 performance at equal runtimes.

6.2 Multi-Step Self-Consistency Losses

We next consider multi-step self-consistency BSDE losses [5] in order to evaluate the mathematical analysis conducted in Section 5. Specifically, we evaluate both the multi-step $L_{EM_k, \tau}(\theta)$ (cf. (3.7))

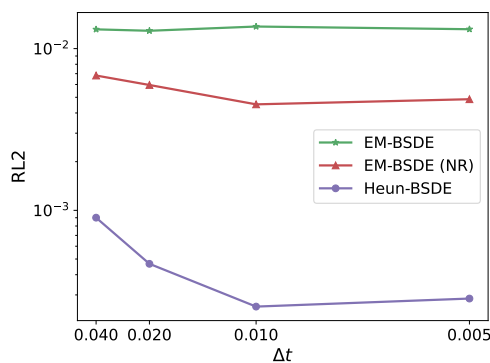


Figure 3: RL2 performance for 10D BSB at discretization step-sizes $\tau = N^{-1}$ for $N \in \{25, 50, 100, 200\}$.

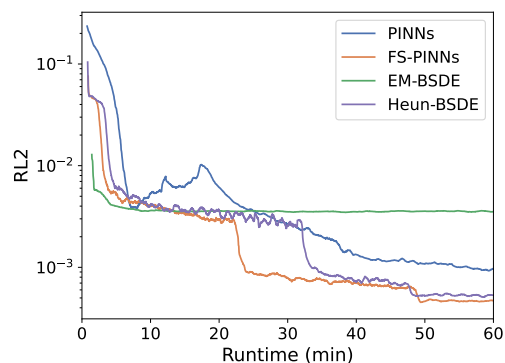
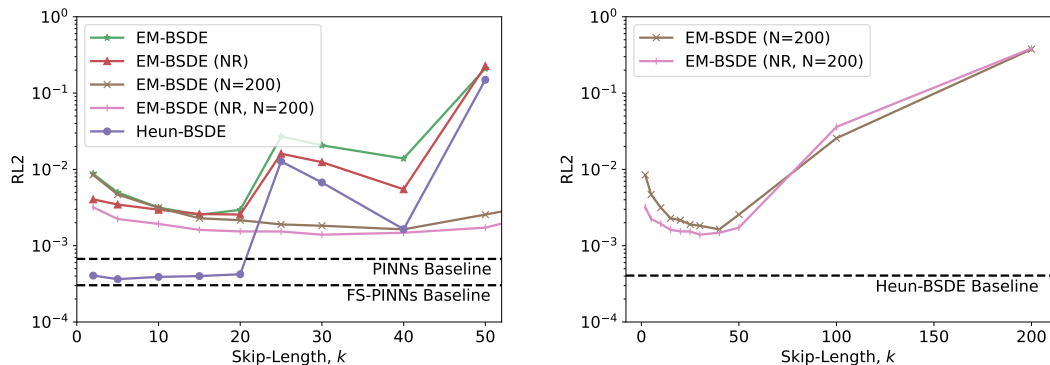


Figure 4: A plot of the RL2 performance versus runtime for the 100D HJB problem.

and $L_{\text{Heun},k,\tau}(\theta)$ (cf. (4.8)) for varying values of skip-length k . For multi-step losses, we also need to determine where both EM-BSDE and EM-BSDE (NR) will “reset” the value of Y_t to $u_\theta(X_t, t)$. Note that there are many degrees of freedom here in the multi-step formulation, so we simply pick one choice as a representative choice. For EM-BSDE, we set $\hat{Y}_{nk}^\theta = u_\theta(\hat{X}_{nk}, t_{nk})$, and use (6.1) to integrate between t_{nk} and $t_{(n+1)k}$. On the other hand, for EM-BSDE (NR), we directly use the value of $\hat{Y}_{nk}^\theta$ from (6.1). We also vary EM-BSDE and EM-BSDE (NR) with N , the number of discretization steps for the interval $[0, 1]$, varying between $N \in \{50, 200\}$ as well. We conduct this experiment on the 10D BSB setting, with the results reported in Figure 5.



(a) Plot of RL2 performance as a function of the skip-length k , with the number of discretization steps also varying in $N \in \{50, 200\}$.

(b) Plot of RL2 performance as a function of the skip-length k for EM-BSDE and EM-BSDE (NR), with number of discretization steps set to $N = 200$.

Figure 5: A plot of RL2 performance of each model on the 10D BSB case at various skip lengths.

Figure 5 shows that while the Heun-BSDE performance decreases as the skip-length k increases, the performance of both EM-BSDE methods initially *improves* with skip-length k before then degrading, demonstrating the trade-off between minimizing the bias term and decreasing quality of the self-consistency loss identified in Section 5. Furthermore, this trade-off is present for EM-BSDE at both $N = 50$ and $N = 200$, illustrating again that these issues for EM are not mitigated with finer discretization step-sizes. Finally, although the EM-BSDE method overall improves its performance by tuning the skip-length k (consistent with the findings from [5]), the Heun-BSDE model at $k = 1$ still outperforms the best multi-step EM-BSDE model by a significant margin.

7 Conclusion and Future Work

We conducted a systematic study of discretization strategies for BSDE-based loss formulations used to solve high-dimensional PDEs. By comparing the commonly used Euler-Maruyama scheme with stochastic Heun integration, we demonstrated that the choice of discretization can significantly impact the accuracy of BSDE-based methods. Our theoretical analysis showed that EM discretization introduces a non-trivial bias to the single-step self-consistency BSDE loss which does not vanish as the step-size decreases. On the other hand, we show that this bias issue is not present when utilizing Heun discretization. Finally, the empirical results confirmed that the Heun scheme consistently outperforms EM in solution accuracy and performs competitively with PINNs.

Our work underscores the importance of stochastic integrator choice in BSDE solvers and suggest that higher-order schemes—though more computationally intensive—can offer substantial gains in performance. In future works, we aim to reduce Heun-BSDE’s computational costs through methods such as Hutchinson trace estimation [40], reversible Heun [41], and adaptive time-stepping. Furthermore, while this current work focuses on understanding and restoring performance parity between BSDE and PINNs methods, future work will utilize the advantages of Heun-BSDE to solve problems such as high-dimensional stochastic control problems in model-free settings (cf. [6]).

References

- [1] Maziar Raissi, Paris Perdikaris, and George Em Karniadakis. Physics informed deep learning (part i): Data-driven solutions of nonlinear partial differential equations. *arXiv preprint arXiv:1711.10561*, 2017.
- [2] Maziar Raissi, Paris Perdikaris, and George Em Karniadakis. Physics informed deep learning (part ii): Data-driven discovery of nonlinear partial differential equations. *arXiv preprint arXiv:1711.10566*, 2017.
- [3] Weinan E, Jiequn Han, and Arnulf Jentzen. Deep learning-based numerical methods for high-dimensional parabolic partial differential equations and backward stochastic differential equations. *Communications in Mathematics and Statistics*, 5(4):349–380, 2017.
- [4] Maziar Raissi. Forward–backward stochastic neural networks: deep learning of high-dimensional partial differential equations. In *Peter Carr Gedenkschrift: Research Advances in Mathematical Finance*, pages 637–655. World Scientific, 2024.
- [5] Nikolas Nüsken and Lorenz Richter. Interpolating between bsdes and pinns: Deep learning for elliptic and parabolic boundary value problems. *Journal of Machine Learning*, 2(1):31–64, 2023.
- [6] Yutian Wang and Yuan-Hua Ni. Deep bsde-ml learning and its application to model-free optimal control, 2022.
- [7] Sifan Wang, Shyam Sankaran, Hanwen Wang, and Paris Perdikaris. An expert’s guide to training physics-informed neural networks. *arXiv preprint arXiv:2308.08468*, 2023.
- [8] Maziar Raissi, Paris Perdikaris, and George Em Karniadakis. Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations. *Journal of Computational Physics*, 378:686–707, 2019.
- [9] Weinan E and Bing Yu. The deep ritz method: a deep learning-based numerical algorithm for solving variational problems. *Communications in Mathematics and Statistics*, 6(1):1–12, 2018.
- [10] Justin Sirignano and Konstantinos Spiliopoulos. Dgm: A deep learning algorithm for solving partial differential equations. *Journal of Computational Physics*, 375:1339–1364, 2018.
- [11] Aditi Krishnapriyan, Amir Gholami, Shandian Zhe, Robert Kirby, and Michael W Mahoney. Characterizing possible failure modes in physics-informed neural networks. In *Advances in Neural Information Processing Systems*, volume 34, pages 26548–26560. Curran Associates, Inc., 2021.
- [12] Pratik Rathore, Weimu Lei, Zachary Frangella, Lu Lu, and Madeleine Udell. Challenges in training pinns: a loss landscape perspective. In *Proceedings of the 41st International Conference on Machine Learning, ICML’24*. JMLR.org, 2024.
- [13] Sifan Wang, Xinling Yu, and Paris Perdikaris. When and why pinns fail to train: A neural tangent kernel perspective. *Journal of Computational Physics*, 449:110768, 2022.
- [14] Tao Wang, Bo Zhao, Sicun Gao, and Rose Yu. Understanding the difficulty of solving Cauchy problems with PINNs. In *Proceedings of the 6th Annual Learning for Dynamics & Control Conference*, volume 242 of *Proceedings of Machine Learning Research*, pages 453–465. PMLR, 15–17 Jul 2024.
- [15] Pi-Yueh Chuang and Lorena A Barba. Predictive limitations of physics-informed neural networks in vortex shedding. *arXiv preprint arXiv:2306.00230*, 2023.
- [16] Chenxi Wu, Min Zhu, Qinyang Tan, Yadhu Kartha, and Lu Lu. A comprehensive study of non-adaptive and residual-based adaptive sampling for physics-informed neural networks. *Computer Methods in Applied Mechanics and Engineering*, 403:115671, 2023.
- [17] Haixu Wu, Huakun Luo, Yuezhou Ma, Jianmin Wang, and Mingsheng Long. Ropinn: Region optimized physics-informed neural networks. *arXiv preprint arXiv:2405.14369*, 2024.

- [18] Chuwei Wang, Shanda Li, Di He, and Liwei Wang. Is L^2 physics informed loss always suitable for training physics informed neural network? In *Advances in Neural Information Processing Systems*, volume 35, pages 8278–8290. Curran Associates, Inc., 2022.
- [19] Zhiwei Gao, Liang Yan, and Tao Zhou. Failure-informed adaptive sampling for pinns. *SIAM Journal on Scientific Computing*, 45(4):A1971–A1994, 2023.
- [20] Rambod Mojtani, Maciej Balajewicz, and Pedram Hassanzadeh. Lagrangian pinns: A causality-conforming solution to failure modes of physics-informed neural networks. *arXiv preprint arXiv:2205.02902*, 2022.
- [21] Jiequn Han, Arnulf Jentzen, and Weinan E. Solving high-dimensional partial differential equations using deep learning. *Proceedings of the National Academy of Sciences*, 115(34):8505–8510, 2018.
- [22] Côme Huré, Huyên Pham, and Xavier Warin. Deep backward schemes for high-dimensional nonlinear pdes. *Mathematics of Computation*, 89(324):1547–1579, 2020.
- [23] Christian Beck, Weinan E, and Arnulf Jentzen. Machine learning approximation algorithms for high-dimensional fully nonlinear partial differential equations and second-order backward stochastic differential equations. *Journal of Nonlinear Science*, 29:1563–1619, 2019.
- [24] Akihiko Takahashi, Yoshifumi Tsuchida, and Toshihiro Yamada. A new efficient approximation scheme for solving high-dimensional semilinear pdes: Control variate method for deep bsde solver. *Journal of Computational Physics*, 454:110956, 2022.
- [25] Kristoffer Andersson, Adam Andersson, and C. W. Oosterlee. Convergence of a robust deep fbsde method for stochastic control. *SIAM Journal on Scientific Computing*, 45(1):A226–A255, 2023.
- [26] Zebang Shen, Zhenfu Wang, Satyen Kale, Alejandro Ribeiro, Amin Karbasi, and Hamed Hassani. Self-consistency of the fokker planck equation. In *Proceedings of Thirty Fifth Conference on Learning Theory*, volume 178 of *Proceedings of Machine Learning Research*, pages 817–841. PMLR, 02–05 Jul 2022.
- [27] Nicholas M Boffi and Eric Vanden-Eijnden. Probability flow solution of the fokker–planck equation. *Machine Learning: Science and Technology*, 4(3):035012, 2023.
- [28] Lingxiao Li, Samuel Hurault, and Justin M Solomon. Self-consistent velocity matching of probability flows. In *Advances in Neural Information Processing Systems*, volume 36, pages 57038–57057. Curran Associates, Inc., 2023.
- [29] Jean-François Chassagneux, Junchao Chen, and Noufel Frikha. Deep runge-kutta schemes for bsdes. *arXiv preprint arXiv:2212.14372*, 2022.
- [30] I.E. Lagaris, A. Likas, and D.I. Fotiadis. Artificial neural networks for solving ordinary and partial differential equations. *IEEE Transactions on Neural Networks*, 9(5):987–1000, 1998.
- [31] Aditya Singh, Zeyuan Feng, and Somil Bansal. Exact imposition of safety boundary conditions in neural reachable tubes. *arXiv preprint arXiv:2404.00814*, 2024.
- [32] Quentin Chan-Wai-Nam, Joseph Mikael, and Xavier Warin. Machine learning for semi linear pdes. *Journal of Scientific Computing*, 79(3):1667–1712, 2019.
- [33] Lorenc Kapllani and Long Teng. Deep learning algorithms for solving high-dimensional nonlinear backward stochastic differential equations. *Discrete and Continuous Dynamical Systems - Series B (DCDS-B)*, 29(4):1695–1729, 2024.
- [34] Peter E. Kloeden and Eckhard Platen. *Numerical Solution of Stochastic Differential Equations*. Springer, 1992.
- [35] W. Rümelin. Numerical treatment of stochastic differential equations. *SIAM Journal on Numerical Analysis*, 19(3):604–613, 1982.

- [36] Christian Bender and Jianfeng Zhang. Time discretization and markovian iteration for coupled fbsdes. *The Annals of Applied Probability*, 18(1):143–177, 2008.
- [37] Mohammad Amin Nabian, Rini Jasmine Gladstone, and Hadi Meidani. Efficient training of physics-informed neural networks via importance sampling. *Computer-Aided Civil and Infrastructure Engineering*, 36(8):962–977, 2021.
- [38] Arka Daw, Jie Bu, Sifan Wang, Paris Perdikaris, and Anuj Karpatne. Mitigating propagation failures in physics-informed neural networks using retain-resample-release (R3) sampling. In *Proceedings of the 40th International Conference on Machine Learning*, volume 202 of *Proceedings of Machine Learning Research*, pages 7264–7302. PMLR, 23–29 Jul 2023.
- [39] Zhengqi Zhang, Jing Li, and Bin Liu. Annealed adaptive importance sampling method in pinns for solving high dimensional partial differential equations. *Journal of Computational Physics*, 521:113561, 2025.
- [40] Zheyuan Hu, Zekun Shi, George Em Karniadakis, and Kenji Kawaguchi. Hutchinson trace estimation for high-dimensional and high-order physics-informed neural networks. *Computer Methods in Applied Mechanics and Engineering*, 424:116883, 2024.
- [41] Patrick Kidger, James Foster, Xuechen (Chen) Li, and Terry Lyons. Efficient and accurate gradients for neural sdes. In M. Ranzato, A. Beygelzimer, Y. Dauphin, P.S. Liang, and J. Wortman Vaughan, editors, *Advances in Neural Information Processing Systems*, volume 34, pages 18747–18761. Curran Associates, Inc., 2021.
- [42] Riu Naito and Toshihiro Yamada. An acceleration scheme for deep learning-based bsde solver using weak expansions. *International Journal of Financial Engineering*, 07(02):2050012, 2020.
- [43] Huyen Pham, Xavier Warin, and Maximilien Germain. Neural networks-based backward scheme for fully nonlinear pdes, 2021.
- [44] Christian Beck, Sebastian Becker, Patrick Cheridito, Arnulf Jentzen, and Ariel Neufeld. Deep splitting method for parabolic pdes. *SIAM Journal on Scientific Computing*, 43(5):A3135–A3154, 2021.
- [45] Nikola Kovachki, Zongyi Li, Burigede Liu, Kamyar Azizzadenesheli, Kaushik Bhattacharya, Andrew Stuart, and Anima Anandkumar. Neural operator: Learning maps between function spaces with applications to pdes. *Journal of Machine Learning Research*, 24(89):1–97, 2023.
- [46] Lorenz Richter, Leon Sallandt, and Nikolas Nüsken. From continuous-time formulations to discretization schemes: tensor trains and robust regression for bsdes and parabolic pdes. *arXiv preprint arXiv:2307.15496*, 2023.
- [47] Jiequn Han and Jihao Long. Convergence of the deep bsde method for coupled fbsdes. *Probability, Uncertainty and Quantitative Risk*, 5(1), 2020.
- [48] Terry J. Lyons. Differential equations driven by rough signals. *Revista Matemática Iberoamericana*, 14(2):215–310, 1998.
- [49] Jan R. Magnus. The moments of products of quadratic forms in normal variables. *Statistica Neerlandica*, 32(4):201–210, 1978.
- [50] Tuomas Hytönen, Jan Van Neerven, Mark Veraar, and Lutz Weis. *Analysis in Banach Spaces: Volume I: Martingales and Littlewood-Paley Theory*. Springer, 2016.
- [51] Dean S. Clark. Short proof of a discrete gronwall inequality. *Discrete Applied Mathematics*, 16(3):279–281, 1987.
- [52] Prajit Ramachandran, Barret Zoph, and Quoc V Le. Searching for activation functions. *arXiv preprint arXiv:1710.05941*, 2017.
- [53] Matthew Tancik, Pratul P. Srinivasan, Ben Mildenhall, Sara Fridovich-Keil, Nithin Raghavan, Utkarsh Singhal, Ravi Ramamoorthi, Jonathan T. Barron, and Ren Ng. Fourier features let networks learn high frequency functions in low dimensional domains. In *Proceedings of the 34th International Conference on Neural Information Processing Systems*, NIPS ’20, Red Hook, NY, USA, 2020. Curran Associates Inc.

- [54] James Bradbury, Roy Frostig, Peter Hawkins, Matthew James Johnson, Chris Leary, Dougal Maclaurin, George Necula, Adam Paszke, Jake VanderPlas, Skye Wanderman-Milne, and Qiao Zhang. JAX: composable transformations of Python+NumPy programs, 2018.
- [55] Haoyu Han and Heng Yang. On the nonsmooth geometry and neural approximation of the optimal value function of infinite-horizon pendulum swing-up. *arXiv preprint arXiv:2312.17467*, 2024.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: We make two main claims in the abstract and introduction: (i) the performance gap between existing BSDE and PINNs methods is explained by an irreducible bias introduced through the use of EM integration for BSDE, and (ii) reformulating a Stratonovich BSDE loss and integrating it via Heun removes the irreducible bias and restores parity between the two methods. Both claims (i) and (ii) are well supported theoretically (Section 4 and Section 5) and experimentally (Section 6).

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: Please see Appendix A for a discussion of the limitations of our work.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [Yes]

Justification: Both Lemma 4.1 and Lemma 4.3 have stated their full set of assumptions, and have accompanying proofs in Appendix D.2 and Appendix D.3. Furthermore, for the theoretical claims made in Section 5, the corresponding derivatives are given in Appendix E.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: Both Section 6 and Appendix G contain the necessary details to reproduce the experiments conducted in this paper, including architectural choices, learning rate schedules, batch sizes, benchmark PDE formulas and analytical solutions, and compute platforms used. We have also included our code in a github repository with a README.md file.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).

- (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: For the submission, we have included an anonymized version of the code in the supplementary material, including a README.md file. For the camera ready version of the paper, we have included a link to our github repository (<https://github.com/sungje-park/heunbsde>). Our code on github is released under the Apache 2.0 License.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: At the beginning of Section 6, we include a detailed description of the all methods being evaluated, including the loss functions and how the collocation/trajectory data is generated. The benchmark PDEs are described in detail in Appendix G.1, which includes the analytical solutions we use to evaluate model quality. We define within the main text the error metric RL2 (6.2) which we use to assess model fidelity, and describe the number of test trajectories we collect for computing RL2. Finally, as noted previous, details about the training pipeline are described in Appendix G.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: All our models are trained over 3 different training initialization seeds. In Table 1, Table 3, and Table 2, we report the average RL2/runtime $\pm$ one standard deviation. For the figures, we only report the average in order to avoid cluttering the diagrams, since the standard deviation is modest relative to the average (cf. Table 1).

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: In Appendix G, we describe our compute environment (single node NVIDIA A100 GPU in our internal cluster).

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: We confirm that our work conforms with the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.

- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. **Broader impacts**

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: Our work falls into the category of foundational research; we do not believe it poses any specific and/or unique potential harms.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. **Safeguards**

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: The assets resulting from our work do not pose a high risk for misuse.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. **Licenses for existing assets**

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [NA]

Justification: Our code does not use any existing assets beyond standard open-source machine learning frameworks, which we give proper attribute to.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [Yes]

Justification: We provide documentation via a README.md file for the code we release.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: Our paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: Our paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigor, or originality of the research, declaration is not required.

Answer: [NA]

Justification: LLMs do not constitute an important, original, or non-standard component of our work.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

A Limitations

We provide a concise, bulleted list of the limitations present in our work.

- (a) *Computational overhead of Heun-BSDE*: As discussed in Section 6.1 (cf. Table 2), the computational overhead of the Heun-BSDE method compared with the existing EM-BSDE methods is non-trivial. Furthermore, we also found the Heun integrator to be more susceptible to floating point imprecision (cf. Table 3), and hence we run our main experiments in float64, which further adds to the computation time. We mitigate these issues using batched computation and random sub-sampling (see Appendix G.4) which helped significantly reduce computation time, but there still remains a computational penalty. While we believe more sophisticated techniques (e.g., randomized trace estimation and more numerically stable Heun integrators discussed in Section 6.1) can help to reduce the computational overhead, we have not yet verified this hypothesis experimentally in our current work.
- (b) *Performance relative to PINNs*: Our results in Section 6.1 show that the proposed Heun-BSDE method restores parity with the FS-PINNs method in terms of the RL2 error. While this is a significant improvement over EM-BSDE, Heun-BSDE does not yet provide significant (e.g., orders of magnitude) performance improvement over the best PINNs method (cf. Table 1); further work is needed to determine whether or not such an improvement is possible. Hence, the current advantage of Heun-BSDE over PINNs is with its model-free capabilities (cf. Appendix C).
- (c) *Limitations of theoretical analysis*: While our theoretical analysis in Section 4 and Section 5 is fairly predictive of practice (cf. Section 6), our analysis is not without its own set of limitations. One limitation is that we only analyze the two extremes of horizon length: the one-step case (Section 4) and the full-horizon N -step case (Section 5); the intermediate regimes (i.e., skip-lengths k satisfying $1 < k < N$) are studied empirically (cf. Section 6.2). Another limitation is that we do not consider *fully-coupled* FS-BSDEs in our analysis (e.g., the Bender & Zhang (BZ) PDE, cf. Appendix G.1), where the forwards SDE (3.3) is allowed to depend on Y_t .

B Further Discussion on Related Works

While our work focuses on the impact that the widely used EM integration scheme has on the performance of BSDE-based solvers, several complementary enhancements to the BSDE loss have been proposed in literature [6, 24, 42–44]. Many of these improvements are naturally compatible with our proposed Heun-BSDE framework. For example, the Heun-BSDE method could be adapted to utilize a control variate [24], applied in operator splitting settings [44], and extended to fully non-linear PDEs [43], enabling direct comparison against their EM-BSDE baseline.

Furthermore, in addition to PINNs and EM-based BSDE methods discussed in the paper, there are various other deep-learning methods for solving PDEs such as Deep Ritz [9], Neural Operators [45], and tensor trains [46], in addition to various theoretical analyses developed for Itô-based BSDE formulations [47]. We leave extending these approaches and analyses to Stratonovich-based formulations as future research directions.

C Details Regarding Model-Free BSDE Formulation

Suppose that the drift term $f(x, t)$ from the forwards SDE (3.3) is *unknown*. Instead, suppose that our computational model takes as input a realization of $(B_t)_{t=0}^T$ and returns FSDE *trajectories* $(X_t)_{t=0}^T$, or more practically a sub-sampled trajectory $\{X_k\}_{k=0}^N$. Under this computational model, PINNs methods cannot be used to solve the PDE (3.1)—even if the other terms g , h , and ϕ are all known—since the residual term $R[u]$ cannot be computed without knowledge of the drift term f . However, in this settings, BSDE methods including the proposed Heun-BSDE method can still be used. This is because BSDE losses only require access to the FSDE trajectory (X_t) and the Brownian motion (B_t) used to generate it (cf. Equation (3.5)).

A specific example where the proposed computational model is realistic comes from model-free optimal control, and was first described in [6]. Consider the following deterministic continuous-time control-affine *fully-actuated* system:

$$\dot{X}_t = f(X_t, t) + \Phi(X_t, t)U_t, \quad X_t, U_t \in \mathbb{R}^d, \quad \text{rank}(\Phi(x, t)) = d \quad \forall (x, t). \quad (\text{C.1})$$

We assume that we do not know the drift term $f(x, t)$, but we are able to select control inputs U_t and obtain the resulting trajectories (X_t) ; this setting is often called the *model-free setting* in optimal control and reinforcement learning. In this framework, by setting the input U_t to be a nominal input $\bar{U}_t$ injected with excitation noise injected, i.e., $U_t = \bar{U}_t + \text{“Noise}_t\text{”}$, we can select the realization (B_t) of Brownian noise and observe the trajectories of forward SDE of the forms:

$$dX_t = [f(X_t, t) + \Phi(X_t, t)\bar{U}_t]dt + g(X_t, t) \diamond dB_t, \quad (\text{C.2})$$

where the $\diamond$ indicates the SDE is to be interpreted in terms of either Itô or Stratonovich, depending on the context. To rigorously define “Noise_t” to establish a connection between (C.1) and (C.2) is technical, requiring the use of e.g., rough path theory [48]. We will take a more practical approach inspired from [6] and observe how injecting Gaussian noise into discretizations of (C.1) induces stochastic discretizations of (C.2).

Concretely, we proceed as follows. We work with constant step-size integrators, and define integration times $t_k = k\tau$, $k \in \mathbb{N}$, for a step-size $\tau > 0$. Given a nominal control input $\bar{U}_t$ which is an open-loop (i.e., only time-dependent) signal,³ we form our control input U_t by injecting Gaussian noise as follows:

$$U_t = \bar{U}_t + w_{\lfloor t/\tau \rfloor} / \sqrt{\tau}, \quad w_k \sim \mathcal{N}(0, I_d). \quad (\text{C.3})$$

For what follows we define $\hat{U}_k := U_{t_k}$ and $\hat{\bar{U}}_k := \bar{U}_{t_k}$ for $k \in \mathbb{N}$. We will assume that our dynamics f, Φ are continuous in both (x, t) , in addition to the nominal signal $\bar{U}_t$ being continuous in t . However, since our signal U_t is discontinuous in t due to the addition of the Gaussian noise, the resulting vector field $F(x, t) := f(x, t) + \Phi(x, t)U_t$ is discontinuous on t . Hence, we will assume that the integration strategies will, in order to generate the $(k+1)$ -th iterate given the k -th iterate, only evaluate the vector field $F(x, t)$ on the half-open interval $[t_k, t_{k+1})$. In particular, we will interpret $F(x, t_k)$ using the right limit $F(x, t_k^+) = \lim_{t \rightarrow t_k^+} F(x, t)$ and $F(x, t_{k+1})$ using the left limit $F(x, t_{k+1}^-) = \lim_{t \rightarrow t_{k+1}^-} F(x, t)$.

Euler scheme. Consider the standard forward Euler scheme used to discretize (C.1) with constant step-size τ :

$$\begin{aligned} \hat{X}_{k+1} &= \hat{X}_k + \tau[f(\hat{X}_k, t_k) + \Phi(\hat{X}_k, t_k)\hat{U}_k] \\ &= \hat{X}_k + \tau[f(\hat{X}_k, t_k) + \Phi(\hat{X}_k, t_k)\hat{\bar{U}}_k] + \sqrt{\tau}\Phi(\hat{X}_k, t_k)w_k, \end{aligned}$$

which corresponds to the standard Euler-Maruyama discretization of the Itô variant of (C.2) with $g = \Phi$.

Heun scheme. Now consider the Heun scheme used to discretize (C.1), again with constant step-size τ :

$$\begin{aligned} \bar{X}_{k+1} &= \hat{X}_k + \tau[f(\hat{X}_k, t_k) + \Phi(\hat{X}_k, t_k)\hat{U}_k], \\ \hat{X}_{k+1} &= \hat{X}_k + \frac{\tau}{2} [f(\hat{X}_k, t_k) + \Phi(\hat{X}_k, t_k)\hat{U}_k + f(\bar{X}_{k+1}, t_{k+1}) + \Phi(\bar{X}_{k+1}, t_{k+1})\hat{U}_{k+1}^-], \end{aligned}$$

where $\hat{U}_{k+1}^- = \hat{\bar{U}}_{k+1} + w_k / \sqrt{\tau}$. Using the shorthand $\bar{F}(x, t) := f(x, t) + \Phi(x, t)\bar{U}_t$, we see that:

$$\begin{aligned} \bar{X}_{k+1} &= \hat{X}_k + \tau\bar{F}(\hat{X}_k, t_k) + \sqrt{\tau}\Phi(\hat{X}_k, t_k)w_k, \\ \hat{X}_{k+1} &= \hat{X}_k + \frac{\tau}{2} [\bar{F}(\hat{X}_k, t_k) + \bar{F}(\bar{X}_{k+1}, t_{k+1})] + \frac{\sqrt{\tau}}{2} [\Phi(\hat{X}_k, t_k) + \Phi(\bar{X}_{k+1}, t_{k+1})]w_k, \end{aligned}$$

which corresponds to the stochastic Heun discretization of the Stratonovich variant of (C.2), again with $g = \Phi$.

D Proofs of Main Results

D.1 Auxiliary results

We first recall a standard formula for the variance of Gaussian quadratic forms.

³A similar argument can also be made for state-dependent policies.

Proposition D.1. Let Q be a $d \times d$ symmetric matrix and $w \sim \mathcal{N}(0, I_d)$. Then,

$$\mathbb{E}_w(\text{tr}(Q) - w^\top Q w)^2 = 2\|Q\|_F^2.$$

Proof. We have that $\mathbb{E}_w(\text{tr}(Q) - w^\top Q w)^2 = \mathbb{E}_w(w^\top Q w)^2 - \text{tr}(Q)^2$. From [cf. 49, Lemma 2.2] we obtain the identity $\mathbb{E}_w(w^\top Q w)^2 = 2\text{tr}(Q^2) + \text{tr}(Q)^2$, which concludes the proof. $\square$

Proposition D.2. Let $(X_t)_{t=a}^b$ for $a \leq b$ denote an $\mathbb{R}^d$ -valued stochastic process, and let $r : \mathbb{R}^d \times [a, b] \mapsto \mathbb{R}$ be an L -Lipschitz function on its domain. Suppose that for all $t_1, t_2 \in [a, b]$ we have $\mathbb{E}\|X_{t_1} - X_{t_2}\|^2 \leq M|t_1 - t_2|$. Then,

$$\int_a^b \mathbb{E}[|r(X_t, t) - r(X_a, a)|]dt \leq L[M^{1/2}(b-a)^{3/2} + (b-a)^2]. \quad (\text{D.1})$$

Furthermore, suppose that for some $0 < \tau \leq 1$, we have $N := (b-a)/\tau \in \mathbb{N}_+$. Define $t_n := a + \tau n$ for $n \in \{0, \dots, N\}$. Then we have that the left-endpoint Riemann sum satisfies:

$$\left| \frac{1}{b-a} \int_a^b \mathbb{E}[r(X_t, t)]dt - \frac{1}{N} \sum_{n=0}^{N-1} \mathbb{E}[r(X_{t_n}, t_n)] \right| \leq L(1 + M^{1/2})\tau^{1/2}. \quad (\text{D.2})$$

Proof. First, we have:

$$\begin{aligned} \int_a^b \mathbb{E}[|r(X_t, t) - r(X_a, a)|]dt &\leq L \int_a^b \mathbb{E}\left\| \begin{bmatrix} X_t - X_a \\ t - a \end{bmatrix} \right\| dt \\ &\leq L \int_a^b \mathbb{E}\|X_t - X_a\|dt + L \int_a^b (t-a)dt \\ &\stackrel{(a)}{=} L \int_a^b \mathbb{E}\|X_t - X_a\|dt + \frac{L(b-a)^2}{2} \\ &\leq \frac{2LM^{1/2}}{3}(b-a)^{3/2} + \frac{L(b-a)^2}{2}, \end{aligned}$$

where (a) follows by Jensen's inequality and the second moment assumption on $(X_t)_t$:

$$L \int_a^b \mathbb{E}\|X_t - X_a\|dt \leq L \int_a^b \sqrt{\mathbb{E}\|X_t - X_a\|^2}dt \leq LM^{1/2} \int_a^b (t-a)^{1/2}dt = \frac{2LM^{1/2}}{3}(b-a)^{3/2}.$$

This establishes (D.1). We now turn to (D.2). We write:

$$\begin{aligned} \left| \int_a^b \mathbb{E}[r(X_t, t)]dt - \tau \sum_{n=0}^{N-1} \mathbb{E}[r(X_{t_n}, t_n)] \right| &= \left| \sum_{n=0}^{N-1} \int_{t_n}^{t_{n+1}} \mathbb{E}[r(X_t, t)]dt - \tau \sum_{n=0}^{N-1} \mathbb{E}[r(X_{t_n}, t_n)] \right| \\ &= \left| \sum_{n=0}^{N-1} \left(\int_{t_n}^{t_{n+1}} \mathbb{E}[r(X_t, t) - r(X_{t_n}, t_n)]dt \right) \right| \\ &\leq \sum_{n=0}^{N-1} \int_{t_n}^{t_{n+1}} \mathbb{E}[|r(X_t, t) - r(X_{t_n}, t_n)|]dt \\ &\stackrel{(a)}{\leq} NL[M^{1/2}\tau^{3/2} + \tau^2] \\ &\stackrel{(b)}{\leq} NL(M^{1/2} + 1)\tau^{3/2} = (b-a)L(M^{1/2} + 1)\tau^{1/2}, \end{aligned}$$

where (a) is from (D.1) and (b) is from the assumption that $\tau \leq 1$. This establishes (D.2). $\square$

Proposition D.3. Let $(X_t)_{t=a}^b$ for $a \leq b$ denote an $\mathbb{R}^d$ -valued stochastic process, and let $r : \mathbb{R}^d \times [a, b] \mapsto \mathbb{R}$ be a B -bounded and L -Lipschitz function on its domain. Suppose that for all $t_1, t_2 \in [a, b]$ we have $\mathbb{E}\|X_{t_1} - X_{t_2}\|^2 \leq M|t_1 - t_2|$, with $M \geq 1$. Then, if $b-a \leq 1$,

$$\left| \mathbb{E} \left(\int_a^b r(X_t, t)dt \right)^2 - (b-a)^2 \mathbb{E}[r^2(X_a, a)] \right| \leq [L^2M + 4BLM^{1/2}](b-a)^{5/2}. \quad (\text{D.3})$$

Proof. We first decompose:

$$\begin{aligned}\mathbb{E}\left(\int_a^b r(X_t, t) dt\right)^2 &= \mathbb{E}\left((b-a)r(X_a, a) + \int_a^b [r(X_t, t) - r(X_a, a)] dt\right)^2 \\ &= (b-a)^2 \mathbb{E}[r^2(X_a, a)] + \mathbb{E}\left(\int_a^b [r(X_t, t) - r(X_a, a)] dt\right)^2 \\ &\quad + 2(b-a) \mathbb{E}\left[r(X_a, a) \int_a^b [r(X_t, t) - r(X_a, a)] dt\right].\end{aligned}$$

Next, by Jensen's inequality,

$$\begin{aligned}\mathbb{E}\left(\int_a^b [r(X_t, t) - r(X_a, a)] dt\right)^2 &\leq (b-a) \int_a^b \mathbb{E}[(r(X_t, t) - r(X_a, a))^2] dt \\ &\leq (b-a)L^2 \int_a^b (\mathbb{E}[\|X_t - X_a\|^2] + (t-a)^2) dt \\ &\leq (b-a)L^2 \int_a^b (M(t-a) + (t-a)^2) dt \stackrel{(a)}{\leq} L^2 M(b-a)^3,\end{aligned}$$

where in (a) we use the assumptions that $M \geq 1$ and $b-a \leq 1$. On the other hand, by another application of Jensen's inequality,

$$\begin{aligned}2(b-a) \left| \mathbb{E}\left[r(X_a, a) \int_a^b [r(X_t, t) - r(X_a, a)] dt\right] \right| &\leq 2(b-a)B \int_a^b \mathbb{E}[|r(X_t, t) - r(X_a, a)|] dt \\ &\stackrel{(a)}{\leq} 2BL [M^{1/2}(b-a)^{5/2} + (b-a)^3] \\ &\stackrel{(b)}{\leq} 4BLM^{1/2}(b-a)^{5/2},\end{aligned}$$

where (a) uses Proposition D.2, specifically (D.1), and (b) uses the assumptions that $M \geq 1$ and $b-a \leq 1$. The claim now follows. $\square$

Proposition D.4. Consider the Itô SDE $(X_t)_{t=a}^b$ defined by $dX_t = f(X_t, t)dt + g(X_t, t)dB_t$, with

$$\sup_{(x,t) \in \mathbb{R}^d \times [a,b]} \max\{\|f(x, t)\|, \|g(x, t)\|_F\} \leq B.$$

For any $t_0, t_1 \in [a, b]$,

$$\mathbb{E}\|X_{t_1} - X_{t_0}\|^2 \leq 2B^2[(t_1 - t_0)^2 + |t_1 - t_0|].$$

Proof. Assume wlog that $t_1 \geq t_0$. We first decompose:

$$\begin{aligned}\mathbb{E}\|X_{t_1} - X_{t_0}\|^2 &= \mathbb{E}\left\|\int_{t_0}^{t_1} f(X_t, t) dt + \int_{t_0}^{t_1} g(X_t, t) dB_t\right\|^2 \\ &\leq 2\mathbb{E}\left\|\int_{t_0}^{t_1} f(X_t, t) dt\right\|^2 + 2\mathbb{E}\left\|\int_{t_0}^{t_1} g(X_t, t) dB_t\right\|^2.\end{aligned}$$

Next, we have:

$$\left\|\int_{t_0}^{t_1} f(X_t, t) dt\right\| \leq \int_{t_0}^{t_1} \|f(X_t, t)\| dt \leq (t_1 - t_0)B \implies \mathbb{E}\left\|\int_{t_0}^{t_1} f(X_t, t) dt\right\|^2 \leq B^2(t_1 - t_0)^2.$$

On the other hand, by Itô isometry,

$$\mathbb{E}\left\|\int_{t_0}^{t_1} g(X_t, t) dB_t\right\|^2 = \int_{t_0}^{t_1} \mathbb{E}\|g(X_t, t)\|_F^2 dt \leq B^2(t_1 - t_0).$$

The result now follows $\square$

Proposition D.5. Consider the Stratonovich SDE $(X_t^\bullet)_{t=a}^b$ defined by $dX_t^\bullet = f(X_t^\bullet, t) + g(X_t^\bullet, t) \circ dB_t$. Suppose that for $t \in [a, b]$, the map $x \mapsto g(x, t)$ is C^1 , and that

$$\sup_{(x,t) \in \mathbb{R}^d \times [a,b]} \max \left\{ \left\| f(x, t) + \frac{1}{2} \sum_{k=1}^m \partial_x g^k(x, t) g(x, t) \right\|, \|g(x, t)\|_F \right\} \leq B.$$

Then for any $t_0, t_1 \in [a, b]$,

$$\mathbb{E} \|X_{t_1}^\bullet - X_{t_0}^\bullet\|^2 \leq 2B^2 [(t_1 - t_0)^2 + |t_1 - t_0|].$$

Proof. Consider $\bar{f}(x, t) := f(x, t) + \frac{1}{2} \sum_{k=1}^m \partial_x g^k(x, t) g(x, t)$. The Itô SDE $dX_t = \bar{f}(X_t, t)dt + g(X_t, t)dB_t$ is pathwise equivalent to $(X_t^\bullet)_t$, i.e., $(X_t^\bullet(\omega))_t = (X_t(\omega))_t$ for a.e. ω , and hence the result follows from Proposition D.4. $\square$

D.2 Proofs of Lemma 4.1 and Theorem 4.2

We first restate and prove Lemma 4.1.

Lemma 4.1. Suppose that f, g are bounded and u_θ is $C^{2,1}$. We have that

$$\tau^{-2} \cdot \ell_{\text{EM}, \tau}(\theta, x, t) = (R[u_\theta](x, t))^2 + \frac{1}{2} \text{tr}[(H(x, t) \cdot \nabla^2 u_\theta(x, t))^2] + O(\tau^{1/2}), \quad (4.2)$$

where the $O(\cdot)$ hides factors depending on d , the bounds on f, g , and $\|u_\theta\|_{C^{2,1}}$.

Proof. We first introduce two pieces of notation: $O(\cdot)$ and $O^*(\cdot)$. The former $O(\cdot)$ hides constants that depend arbitrarily on the dimension d , the bounds on f and g , and $\|u_\theta\|_{C^{2,1}}$, whereas the latter $O^*(\cdot)$ in addition also hides constants that depend *polynomially* on $\|w\|$. The latter polynomial dependence is important when we take expectations of powers of $O^*(\cdot)$ terms, since $\mathbb{E}\|w\|^p$ is finite for any finite $p \in \mathbb{N}$.

Setting $\bar{\Delta} := (\hat{x}_{t+\tau} - x, \tau) \in \mathbb{R}^d \times \mathcal{I}$ and writing $u = u_\theta$, a second-order Taylor expansion yields:

$$\begin{aligned} u(\hat{x}_{t+\tau}, t + \tau) - u(x, t) &= Du(x, t)\bar{\Delta} + \frac{1}{2} \bar{\Delta}^\top D^2 u(x, t) \bar{\Delta} + O(\|\bar{\Delta}\|^3), \\ Du(x, t)\bar{\Delta} &= \langle \nabla u(x, t), \hat{x}_{t+\tau} - x \rangle + \partial_t u(x, t)\tau, \\ \frac{1}{2} \bar{\Delta}^\top D^2 u(x, t) \bar{\Delta} &= \frac{1}{2} ((\hat{x}_{t+\tau} - x)^\top \nabla^2 u(x, t) (\hat{x}_{t+\tau} - x) \\ &\quad + 2\tau \langle \partial_t \nabla u(x, t), \hat{x}_{t+\tau} - x \rangle + \tau^2 \partial_t^2 u(x, t)). \end{aligned}$$

Plugging in $\hat{x}_{t+\tau} - x = f(x, t)\tau + \sqrt{\tau}g(x, t)w$, we obtain:

$$\begin{aligned} &(u(\hat{x}_{t+\tau}, t + \tau) - u(x, t)) - (h(x, t)\tau - \sqrt{\tau} \langle \nabla u(x, t), g(x, t)w \rangle) \\ &= \tau \left[\langle \nabla u(x, t), f(x, t) \rangle + \partial_t u(x, t) - h(x, t) + \frac{1}{2} w^\top g(x, t)^\top \nabla^2 u(x, t) g(x, t)w \right] + O^*(\tau^{3/2}) \\ &= \tau \left[R[u](x, t) - \frac{1}{2} \text{tr}(H(x, t) \cdot \nabla^2 u(x, t)) + \frac{1}{2} w^\top g(x, t)^\top \nabla^2 u(x, t) g(x, t)w \right] + O^*(\tau^{3/2}), \end{aligned}$$

where in the last equality we used the definition of the PDE residual from (3.1). Hence,

$$\begin{aligned} \ell_{\text{EM}, \tau}(x, t) &= \mathbb{E}_w (u(\hat{x}_{t+\tau}, t + \tau) - u(x, t)) - (h(x, t)\tau - \sqrt{\tau} \langle \nabla u(x, t), g(x, t)w \rangle)^2 \\ &= \tau^2 \cdot \mathbb{E}_w \left(R[u](x, t) - \frac{1}{2} \text{tr}(H(x, t) \cdot \nabla^2 u(x, t)) + \frac{1}{2} w^\top g(x, t)^\top \nabla^2 u(x, t) g(x, t)w \right)^2 + O(\tau^{5/2}) \\ &= \tau^2 \left((R[u](x, t))^2 + \frac{1}{4} \mathbb{E}_w (\text{tr}(H(x, t) \cdot \nabla^2 u(x, t)) - w^\top g(x, t)^\top \nabla^2 u(x, t) g(x, t)w)^2 \right) + O(\tau^{5/2}) \\ &= \tau^2 \left((R[u](x, t))^2 + \frac{1}{2} \text{tr}((H(x, t) \cdot \nabla^2 u(x, t))^2) \right) + O(\tau^{5/2}), \end{aligned}$$

where the final equality follows from Proposition D.1. The claim now follows. $\square$

Next, we use Lemma 4.1, along with order 1/2 strong convergence for EM integration (Appendix F) to show the following result.

Theorem 4.2. Suppose that $f, g, h_\theta \in C^{0,1}$, $u_\theta \in C^{2,1}$, and $\tau \leq 1$. We have that:

$$L_{\text{EM},\tau}(\theta) = \frac{1}{T} \int_0^T \mathbb{E} \left[(R[u_\theta](X_t, t))^2 + \frac{1}{2} \text{tr} \left[(H(X_t, t) \cdot \nabla^2 u_\theta(X_t, t))^2 \right] \right] dt + O(\tau^{1/2}), \quad (4.3)$$

where the $O(\cdot)$ hides constants that depend on d, T , and the Hölder norms of f, g, h_θ , and u_θ .

Proof. To start, we have that:

$$\begin{aligned} L_{\text{EM},\tau}(\theta) &= \frac{1}{N} \sum_{n=0}^{N-1} \mathbb{E}_{\hat{X}_n} [\tau^{-2} \cdot \ell_{\text{EM},\tau}(\theta, \hat{X}_n, t_n)] \\ &\stackrel{(a)}{=} \frac{1}{N} \sum_{n=0}^{N-1} \mathbb{E}_{\hat{X}_n} \left[\underbrace{(R[u_\theta](\hat{X}_n, t_n))^2 + \frac{1}{2} \text{tr} \left[(H(\hat{X}_n, t_n) \cdot \nabla^2 u_\theta(\hat{X}_n, t_n))^2 \right]}_{=: \bar{R}_\theta(\hat{X}_n, t_n)} \right] + O(\tau^{1/2}), \end{aligned}$$

where (a) comes from Lemma 4.1 which holds since (i) $f, g \in C^{0,1}$ implies f, g are bounded, and (ii) $u_\theta \in C^{2,1}$. Our next observation is that the map $\bar{R}_\theta$ is also Lipschitz continuous over the domain $\mathbb{R}^d \times \mathcal{I}$ by our assumptions $f, g, h[u_\theta] \in C^{0,1}$ and $u_\theta \in C^{2,1}$. Let us call this Lipschitz constant $L_{\bar{R}_\theta}$, which depends only on the norms $\|f\|_{C^{0,1}}, \|g\|_{C^{0,1}}, \|h[u_\theta]\|_{C^{0,1}}, \|u_\theta\|_{C^{2,1}}$. Continuing from above,

$$\begin{aligned} \mathbb{E}_{\hat{X}_n} [\bar{R}_\theta(\hat{X}_n, t_n)] &\stackrel{(a)}{=} \mathbb{E}_{(B_t)_t} [\bar{R}_\theta(\hat{X}_n, t_n)] \\ &\stackrel{(b)}{=} \mathbb{E}_{(B_t)_t} [\bar{R}_\theta(X_{t_n}, t_n)] + \mathbb{E}_{(B_t)_t} [\bar{R}_\theta(\hat{X}_n, t_n) - \bar{R}_\theta(X_{t_n}, t_n)], \end{aligned}$$

where in (a) we consider the process $\{\hat{X}_n\}_n$ as being defined over $\{\Delta W_n\}_n := \{B_{t_{n+1}} - B_{t_n}\}_n$ in place of the process $\{\sqrt{\tau} w_n\}_n$ in (3.6), which is distributionally equivalent, in (b) we consider the forward SDE $(X_t)_t$ from (3.3) as being coupled with the process $\{\hat{X}_n\}_n$ via the same realization of both Brownian motion $(B_t)_t$ and $X_0 = \hat{X}_0 = x_0$. Hence we have:

$$\begin{aligned} |\mathbb{E}_{\hat{X}_n} [\bar{R}_\theta(\hat{X}_n, t_n)] - \mathbb{E}_{(B_t)_t} [\bar{R}_\theta(\hat{X}_n, t_n)]| &\leq \mathbb{E}_{(B_t)_t} [|\bar{R}_\theta(\hat{X}_n, t_n) - \bar{R}_\theta(X_{t_n}, t_n)|] \\ &\leq L_{\bar{R}_\theta} \mathbb{E}_{(B_t)_t} [\|\hat{X}_n - X_{t_n}\|] \\ &\leq L_{\bar{R}_\theta} \sqrt{\mathbb{E}_{(B_t)_t} [\|\hat{X}_n - X_{t_n}\|^2]}. \end{aligned}$$

Now, by definition, since the functions $f, g \in C^{0,1}$, then the pair (f, g) is EM-regular (Definition F.1). From Theorem F.2, we have that $\{\hat{X}_n\}_n$ is strong order 1/2 convergent towards $(X_t)_t$, and hence $\mathbb{E}_{(B_t)_t} [\|\hat{X}_n - X_{t_n}\|^2] \leq \mathbb{E}_{(B_t)_t} [\max_{n \in \{0, \dots, N\}} \|\hat{X}_n - X_{t_n}\|^2] \leq C^2 \tau$, where the constant C does not depend on τ , but can depend on d, T , and the norms $\|f\|_{C^{0,1}}$ and $\|g\|_{C^{0,1}}$. Consequently,

$$L_{\text{EM},\tau}(\theta) = \frac{1}{N} \sum_{n=0}^{N-1} \mathbb{E}_{(B_t)_t} [\bar{R}_\theta(X_{t_n}, t_n)] + O(\tau^{1/2}), \quad (\text{D.4})$$

Our last step is to approximate the sum $\frac{1}{N} \sum_{n=0}^{N-1}$ in (D.4) with the integral. To do this, we will use Proposition D.2. We already have $\bar{R}_\theta$ is Lipschitz over $\mathbb{R}^d \times \mathcal{I}$. Furthermore, since $f, g \in C^{0,1}$, they are both bounded over the domain, and hence Proposition D.4 shows that $\mathbb{E} \|X_{t_1} - X_{t_0}\|^2 \leq O(1) |t_1 - t_0|$ for any $t_0, t_1 \in \mathcal{I}$. Therefore by Proposition D.2 and the assumption $\tau \leq 1$,

$$\frac{1}{T} \int_0^T \mathbb{E} [\bar{R}_\theta(X_t, t)] dt = \frac{1}{N} \sum_{n=0}^{N-1} \mathbb{E} [\bar{R}_\theta(X_{t_n}, t_n)] + O(\tau^{1/2}) \stackrel{(a)}{=} L_{\text{EM},\tau}(\theta) + O(\tau^{1/2}),$$

where (a) is from (D.4). The result now follows. $\square$

D.3 Proofs of Lemma 4.3 and Theorem 4.4

We now restate and prove Lemma 4.3.

Lemma 4.3. Suppose that f , g , and h_θ are all in $C^{1,1}$, and u_θ is in $C^{3,1}$. We have that

$$\tau^{-2} \cdot \ell_{\text{Heun},\tau}(\theta, x, t) = (R[u_\theta](x, t))^2 + O(\tau^{1/2}), \quad (4.9)$$

where the $O(\cdot)$ hides factors depending on d and the Hölder norms of f , g , h_θ , and u_θ .

Proof. Similar to the proof of Lemma 4.1, we let $O(\cdot)$ hide constants that depend on d and the Hölder norms $\|f\|_{C^{1,1}}$, $\|g\|_{C^{1,1}}$, $\|h_\theta\|_{C^{1,1}}$, and $\|u_\theta\|_{C^{3,1}}$, and $O^*(\cdot)$ additionally hides constants that depend polynomially on $\|w\|$.

Our first step is to check that $h_\theta^* \in C^{1,1}$ under our assumptions. Recalling that $h_\theta^*(x, t) = h_\theta(x, t) - \frac{1}{2} \text{tr}(H(x, t) \nabla^2 u_\theta(x, t))$, this is ensured if $h_\theta, H \in C^{1,1}$ and $u_\theta \in C^{3,1}$, which holds since $g \in C^{1,1}$ implies $H \in C^{1,1}$. Furthermore, $\|h_\theta^*\|_{C^{1,1}} = O(1)$. For what follows, we drop the dependency in the notation on θ .

We next Taylor expand $\hat{z}_{t+\tau} - z_t$ up to order τ terms. To do this, we observe that by our assumptions on f, g, h^*, u , the functions F, G are both in $C^{1,1}$. Hence,

$$\begin{aligned} F(\bar{z}_{t+\tau}, t + \tau) &= F(z_t, t) + D_Z F(z_t, t)[\bar{z}_{t+\tau} - z_t] + \partial_t F(z_t, t)\tau + O(\|\bar{z}_{t+\tau} - z_t\|^2) + O(\tau^2) \\ &= F(z_t, t) + D_Z F(z_t, t)[G(z_t, t)w]\sqrt{\tau} + O^*(\tau). \end{aligned}$$

By a similar argument,

$$G(\bar{z}_{t+\tau}, t + \tau) = G(z_t, t) + D_Z G(z_t, t)[G(z_t, t)w]\sqrt{\tau} + O^*(\tau).$$

Hence,

$$\hat{z}_{t+\tau} - z_t = \left[F(z_t, t) + \frac{1}{2} D_Z G(z_t, t)[G(z_t, t)w]w \right] \tau + G(z_t, t)w\sqrt{\tau} + O^*(\tau^{3/2}). \quad (\text{D.5})$$

A straightforward computation yields:

$$D_Z G((x, y), t)[(\Delta_x, \Delta_y)] = \left[\nabla u(x, t)^\top Dg(x, t)[\Delta_x] + \Delta_x^\top \nabla^2 u(x, t)g(x, t) \right],$$

and therefore:

$$\frac{1}{2} D_Z G(z_t, t)[G(z_t, t)w]w = \frac{1}{2} \left[\nabla u(x, t)^\top Dg(x, t)[g(x, t)w]w + w^\top g(x, t)^\top \nabla^2 u(x, t)g(x, t)w \right].$$

Substituting the above into expression (D.5) for $\hat{z}_{t+\tau} - z_t$,

$$\begin{aligned} \hat{z}_{t+\tau} - z_t &= \left[h(x, t) - \frac{1}{2} \text{tr}(H(x, t) \nabla^2 u(x, t)) \right] \tau \\ &\quad + \frac{1}{2} \left[\nabla u(x, t)^\top Dg(x, t)[g(x, t)w]w + w^\top g(x, t)^\top \nabla^2 u(x, t)g(x, t)w \right] \tau \\ &\quad + \sqrt{\tau} \left[\frac{g(x, t)w}{\nabla u(x, t)^\top g(x, t)w} \right] + O^*(\tau^{3/2}). \end{aligned} \quad (\text{D.6})$$

Setting $\bar{\Delta} := (\hat{x}_{t+\tau} - x, \tau) \in \mathbb{R}^{d+1}$ we have:

$$\begin{aligned} u(\hat{x}_{t+\tau}, t + \tau) - u(x, t) &= Du(x, t)\bar{\Delta} + \frac{1}{2} \bar{\Delta}^\top D^2 u(x, t)\bar{\Delta} + O(\|\bar{\Delta}\|^3), \\ Du(x, t)\bar{\Delta} &= \langle \nabla u(x, t), \hat{x}_{t+\tau} - x \rangle + \partial_t u(x, t)\tau, \\ \frac{1}{2} \bar{\Delta}^\top D^2 u(x, t)\bar{\Delta} &= \frac{1}{2} ((\hat{x}_{t+\tau} - x)^\top \nabla^2 u(x, t)(\hat{x}_{t+\tau} - x) \\ &\quad + 2\tau \langle \partial_t \nabla u(x, t), \hat{x}_{t+\tau} - x \rangle + \tau^2 \partial_t^2 u(x, t)), \end{aligned}$$

from which we conclude,

$$\begin{aligned} u(\hat{x}_{t+\tau}, t + \tau) - u(x, t) &= \langle \nabla u(x, t), \hat{x}_{t+\tau} - x \rangle + \partial_t u(x, t)\tau + \frac{1}{2} (\hat{x}_{t+\tau} - x)^\top \nabla^2 u(x, t)(\hat{x}_{t+\tau} - x) + O(\tau^{3/2}) \end{aligned}$$

$$= \left[\langle \nabla u(x, t), f(x, t) + \frac{1}{2} Dg(x, t)[g(x, t)w]w \rangle + \partial_t u(x, t) + \frac{1}{2} w^\top g(x, t)^\top \nabla^2 u(x, t) g(x, t)w \right] \tau \\ + \sqrt{\tau} \langle \nabla u(x, t), g(x, t)w \rangle + O^*(\tau^{3/2}).$$

On the other hand, from (D.6),

$$\hat{y}_{t+\tau} - y_t = \left[h(x, t) - \frac{1}{2} \text{tr}(H(x, t) \nabla^2 u(x, t)) \right] \tau \\ + \frac{1}{2} \left[\nabla u(x, t)^\top Dg(x, t)[g(x, t)w]w + w^\top g(x, t)^\top \nabla^2 u(x, t) g(x, t)w \right] \tau \\ + \sqrt{\tau} \nabla u(x, t)^\top g(x, t)w + O^*(\tau^{3/2}).$$

Hence,

$$u(\hat{x}_{t+\tau}, t + \tau) - \hat{y}_{t+\tau} \\ = (u(\hat{x}_{t+\tau}, t + \tau) - u(x, t)) - (\hat{y}_{t+\tau} - y_t) \\ = \left[\langle \nabla u(x, t), f(x, t) \rangle + \partial_t u(x, t) + \frac{1}{2} \text{tr}(H(x, t) \nabla^2 u(x, t)) - h(x, t) \right] \tau + O^*(\tau^{3/2}) \\ = R[u](x, t)\tau + O^*(\tau^{3/2}).$$

To conclude,

$$\mathbb{E}_w(u(\hat{x}_{t+\tau}, t + \tau) - \hat{y}_{t+\tau})^2 = \mathbb{E}_w(R[u](x, t)\tau + O^*(\tau^{3/2}))^2 = (R[u](x, t))^2 \tau^2 + O(\tau^{5/2}).$$

□

Theorem 4.4. Suppose that f , g , and h_θ are all in $C^{1,1}$, $u_\theta \in C^{3,1}$, and $\tau \leq 1$. We have that

$$L_{\text{Heun}, \tau}(\theta) = \frac{1}{T} \int_0^T \mathbb{E}[(R[u_\theta](X_t^\bullet, t))^2] dt + O(\tau^{1/2}), \quad (4.10)$$

where the $O(\cdot)$ hides factors depending on d , T , and the Hölder norms of f , g , h_θ , and u_θ .

Proof. The proof follows the structure of Theorem 4.2 closely. We start with:

$$L_{\text{Heun}, \tau}(\theta) = \frac{1}{N} \sum_{n=0}^{N-1} \mathbb{E}_{\hat{X}_n^\bullet}[\tau^{-2} \cdot \ell_{\text{Heun}, \tau}(\theta, \hat{X}_n^\bullet, t_n)] \stackrel{(a)}{=} \frac{1}{N} \sum_{n=0}^{N-1} \mathbb{E}_{\hat{X}_n^\bullet}[(R[u_\theta](\hat{X}_n^\bullet, t_n))^2] + O(\tau^{1/2})$$

where (a) follows from Lemma 4.3. Next, we define $\bar{R}_\theta(x, t) := (R[u_\theta](x, t))^2$, and observe that $\bar{R}_\theta$ is Lipschitz over $\mathbb{R}^d \times \mathcal{I}$ due to our assumptions on f, g, h_θ, u_θ . Hence, we have the decomposition:

$$\mathbb{E}_{\hat{X}_n^\bullet}[\bar{R}_\theta(\hat{X}_n^\bullet, t_n)] \stackrel{(a)}{=} \mathbb{E}_{(B_t)_t}[\bar{R}_\theta(\hat{X}_n^\bullet, t_n)] \\ \stackrel{(b)}{=} \mathbb{E}_{(B_t)_t}[\bar{R}_\theta(X_{t_n}^\bullet, t_n)] + \mathbb{E}_{(B_t)_t}[\bar{R}_\theta(\hat{X}_{t_n}^\bullet, t_n) - \bar{R}_\theta(X_{t_n}^\bullet, t_n)],$$

where in (a) and (b) we take same steps as Theorem 4.2: (a) considers the process $\{\hat{X}_n^\bullet\}_n$ as being defined over $\{\Delta W_n\}_n$ (Brownian increments) in place of $\{\sqrt{\tau}w_n\}_n$ in (4.7), and (b) couples the SDE $(X_t^\bullet)_t$ together with $\{\hat{X}_n^\bullet\}_n$ via the same Brownian motion $(B_t)_t$ and initial condition $X_0^\bullet = \hat{X}_0^\bullet = x_0$. Hence, we have:

$$|\mathbb{E}_{\hat{X}_n^\bullet}[\bar{R}_\theta(\hat{X}_n^\bullet, t_n)] - \mathbb{E}_{(B_t)_t}[\bar{R}_\theta(X_{t_n}^\bullet, t_n)]| \leq O(1) \sqrt{\mathbb{E}_{(B_t)_t}[\|X_{t_m}^\bullet - \hat{X}_n^\bullet\|^2]}.$$

Since $f, g \in C^{1,1}$, then they are by definition Heun-regular (Definition F.3), and hence by Theorem F.4, $\mathbb{E}_{(B_t)_t}[\|X_{t_n}^\bullet - \hat{X}_n^\bullet\|^2] \leq \mathbb{E}_{(B_t)_t}[\max_{n \in \{0, \dots, N\}} \|X_{t_n}^\bullet - \hat{X}_n^\bullet\|^2] \leq C^2 \tau$, where C does not depend on τ but depends on d, T , and the Hölder norms on f, g . Therefore we have:

$$L_{\text{Heun}, \tau}(\theta) = \frac{1}{N} \sum_{n=0}^{N-1} \mathbb{E}_{(B_t)_t}[(R[u_\theta](X_{t_n}^\bullet, t_n))^2] + O(\tau^{1/2}). \quad (D.7)$$

Now we can finish up using the same ending as Theorem 4.2. The only thing we need to do differently is to control the second moment $\mathbb{E}\|X_{t_1}^\bullet - X_{t_0}^\bullet\|^2$. Since $f, g \in C^{1,1}$, Proposition D.5 yields that

$\mathbb{E}\|X_{t_1}^\bullet - X_{t_0}^\bullet\|^2 \leq O(1)|t_1 - t_0|$. From this inequality and the Lipschitz continuity of $\bar{R}_\theta$ over the domain $\mathbb{R}^d \times \mathcal{I}$, by Proposition D.2 and the assumption $\tau \leq 1$,

$$\begin{aligned} \frac{1}{T} \int_0^T \mathbb{E}_{(B_t)_t} [(R[u_\theta](X_t^\bullet, t))^2] dt &= \frac{1}{N} \sum_{n=0}^{N-1} \mathbb{E}_{(B_t)_t} [(R[u_\theta](X_{t_n}^\bullet, t_n))^2] + O(\tau^{1/2}) \\ &\stackrel{(a)}{=} L_{\text{Heun}, \tau}(\theta) + O(\tau^{1/2}), \end{aligned}$$

where (a) is from (D.7). The result now follows. $\square$

E Analysis of Multi-Step BSDE Losses

We now present the derivations supporting the analysis in Section 5. For what follows, we define the follow FS-PINNs loss:

$$L_{\text{FS-PINNs}}(\theta) := \mathbb{E}_{x_0 \sim \mu_0, B_t} \frac{1}{T} \int_0^T (R[u_\theta](X_t, t))^2 dt. \quad (\text{E.1})$$

E.1 BSDE loss and Euler-Maruyama discretization

Proposition E.1. *Suppose that $u_\theta \in C^2$. We have that:*

$$L_{\text{BSDE}, T}(\theta) \leq L_{\text{FS-PINNs}}(\theta). \quad (\text{E.2})$$

Proof. To start, we abbreviate $R_\theta(x, t) = R[u_\theta](x, t)$. By application of Itô's Lemma, we have:

$$u_\theta(X_T, T) - u_\theta(X_0, 0) - \int_0^T h_\theta(X_t, t) dt - \int_0^T \nabla u_\theta(X_t, t)^\top g(X_t, t) dB_t = \int_0^T R_\theta(X_t, t) dt,$$

which immediately yields the following identity:

$$L_{\text{BSDE}, T}(\theta) = \mathbb{E}_{x_0 \sim \mu_0, B_t} \left(\frac{1}{T} \int_0^T R_\theta(X_t, t) dt \right)^2. \quad (\text{E.3})$$

Thus, the BSDE loss is equal to averaging the square of the accumulation of the residual error R_θ along the forward SDE trajectories. Hence by Jensen's inequality, the BSDE loss is dominated by:

$$L_{\text{BSDE}, T}(\theta) \leq \mathbb{E}_{x_0 \sim \mu_0, B_t} \frac{1}{T} \int_0^T (R_\theta(X_t, t))^2 dt = L_{\text{FS-PINNs}}(\theta).$$

$\square$

Note that while the relationship in (E.2) is pointed out in [5, Section 5.2.3], the implications of this inequality are not discussed further in their work.

Proposition E.2. *Suppose that $f, g, h_\theta \in C^{0,1}$, $u_\theta \in C^{2,1}$, and $\tau \leq 1$. Then,*

$$L_{\text{BSDE}, \tau}(\theta) = L_{\text{FS-PINNs}}(\theta) + O(\tau^{1/2}), \quad (\text{E.4})$$

where the $O(\cdot)$ hides constants depending on the Hölder norms of f , g , h_θ , and u_θ .

Proof. Again, we abbreviate $R_\theta(x, t) = R[u_\theta](x, t)$. By our assumptions on $f, g, h_\theta \in C^{0,1}$ and $u_\theta \in C^{2,1}$, we have that $R_\theta \in C^{0,1}$, with $\|R_\theta\|_{C^{0,1}} = O(1)$. Also, since $f, g \in C^{0,1}$, by Proposition D.4 we also have that $\mathbb{E}\|X_{t_0} - X_{t_1}\|^2 \leq O(1)|t_0 - t_1|$ for $t_0, t_1 \in \mathcal{I}$. Hence by Proposition D.3,

$$\mathbb{E} \left(\int_{t_n}^{t_{n+1}} R_\theta(X_t, t) dt \right)^2 = \tau^2 \mathbb{E}[R_\theta^2(X_{t_n}, t_n)] + O(\tau^{5/2}). \quad (\text{E.5})$$

Furthermore, since $R_\theta \in C^{0,1}$, we can readily check that R_θ^2 is Lipschitz on its domain as well, and hence Proposition D.2 yield:

$$\frac{1}{T} \int_0^T \mathbb{E}[(R_\theta(X_t, t))^2] dt = \frac{1}{N} \sum_{n=0}^{N-1} \mathbb{E}[(R_\theta(X_{t_n}, t_n))^2] + O(\tau^{1/2}). \quad (\text{E.6})$$

Therefore,

$$\begin{aligned}
L_{\text{BSDE},\tau}(\theta) &= \mathbb{E}_{x_0 \sim \mu_0, B_t} \frac{1}{N\tau^2} \sum_{n=0}^{N-1} \left(\int_{t_n}^{t_{n+1}} R_\theta(X_t, t) dt \right)^2 \\
&= \mathbb{E}_{x_0 \sim \mu_0, B_t} \frac{1}{N\tau^2} \sum_{n=0}^{N-1} (\tau^2 R_\theta^2(X_{t_n}, t_n) + O(\tau^{5/2})) \quad [\text{using (E.5)}] \\
&= \mathbb{E}_{x_0 \sim \mu_0, B_t} \frac{1}{N} \sum_{n=0}^{N-1} R_\theta^2(X_{t_n}, t_n) + O(\tau^{1/2}) \\
&= \mathbb{E}_{x_0 \sim \mu_0, B_t} \frac{1}{T} \int_0^T (R_\theta(X_t, t))^2 dt + O(\tau^{1/2}) \quad [\text{using (E.6)}] \\
&= L_{\text{FS-PINNs}}(\theta) + O(\tau^{1/2}).
\end{aligned}$$

□

Proposition E.3. Suppose the assumptions of Proposition E.2 hold. Then,

$$L_{\text{BSDE},T}(\theta) \leq L_{\text{BSDE},\tau}(\theta) + O(\tau^{1/2}), \quad (\text{E.7})$$

where the $O(\cdot)$ hides constants depending on the Hölder norms of f , g , h_θ , and u_θ .

Proof. Follows immediately from Proposition E.1 and Proposition E.2. □

Proposition E.4. Suppose that $f, g, h_\theta \in C^{0,1}$, $u_\theta \in C^{1,1}$, and $\tau \leq 1$. We have that:

$$L_{\text{EM}_N,\tau}(\theta) = L_{\text{BSDE},T}(\theta) + O(\tau^{1/2}), \quad (\text{E.8})$$

where the $O(\cdot)$ hides constants depending on the d , T , and the Hölder norms of f , g , h_θ , and u_θ .

Proof. We consider the joint forward and backward SDE (cf. (3.3) and (3.4)) with $Z_t^\theta := (X_t, Y_t^\theta)$:

$$\begin{aligned}
d \begin{bmatrix} X_t \\ Y_t^\theta \end{bmatrix} &= \underbrace{\begin{bmatrix} f(X_t, t) \\ h_\theta(X_t, t) \end{bmatrix}}_{=: F_\theta(Z_t^\theta, t)} dt + \underbrace{\begin{bmatrix} g(X_t, t) \\ \nabla u_\theta(X_t, t)^\top g(X_t, t) \end{bmatrix}}_{=: G_\theta(Z_t^\theta, t)} dB_t, \quad \begin{bmatrix} X_0 \\ Y_0^\theta \end{bmatrix} = \begin{bmatrix} x_0 \\ u_\theta(x_0, 0) \end{bmatrix}.
\end{aligned}$$

Given our assumptions on f, g, h_θ, u_θ , we have that both $F_\theta, G_\theta \in C^{0,1}$. Hence, the pair (F_θ, G_θ) is EM-regular (Definition F.1). Therefore, by Theorem F.2, the EM-discretization $\{(\hat{X}_n, \hat{Y}_n^\theta)\}_n$ (cf. (3.6)), coupled with $(Z_t^\theta)_t$ through Brownian increments $\{\Delta W_n\}$, satisfies order 1/2 strong convergence to $(Z_t^\theta)_t$:

$$\left(\mathbb{E} \left[\max_{n \in \{0, \dots, N\}} \max\{\|\hat{X}_n - X_{t_n}\|^2, |\hat{Y}_n^\theta - Y_{t_n}^\theta|^2\} \right] \right)^{1/2} \leq C\tau^{1/2}.$$

Now, define $\hat{\Psi}_N := u_\theta(\hat{X}_N, T) - \hat{Y}_N^\theta$ and $\Psi_T := u_\theta(X_T, T) - Y_T^\theta$

$$\mathbb{E}[\hat{\Psi}_N^2] = \mathbb{E}[(\Psi_T + (\hat{\Psi}_N - \Psi_T))^2] = \mathbb{E}[\Psi_T^2] + \mathbb{E}[(\hat{\Psi}_N - \Psi_T)^2] + 2\mathbb{E}[\Psi_T(\hat{\Psi}_N - \Psi_T)].$$

Hence by Cauchy-Schwarz,

$$|\mathbb{E}[\hat{\Psi}_N^2] - \mathbb{E}[\Psi_T^2]| \leq \mathbb{E}[(\hat{\Psi}_N - \Psi_T)^2] + 2\sqrt{\mathbb{E}[\Psi_T^2]}\sqrt{\mathbb{E}[(\hat{\Psi}_N - \Psi_T)^2]}.$$

Since $u_\theta \in C^{1,1}$ the function is $\|u_\theta\|_{C^{1,1}}$ -Lipschitz and therefore:

$$\begin{aligned}
\mathbb{E}[(\hat{\Psi}_N - \Psi_T)^2] &\leq 2\mathbb{E}[(u_\theta(\hat{X}_N, T) - u_\theta(X_T, T))^2] + 2\mathbb{E}[(\hat{Y}_N^\theta - Y_T^\theta)^2] \\
&\leq 2\|u_\theta\|_{C^{1,1}}^2 \mathbb{E}[\|\hat{X}_N - X_T\|^2] + 2\mathbb{E}[(\hat{Y}_N^\theta - Y_T^\theta)^2] \\
&\leq 2C^2(1 + \|u_\theta\|_{C^{1,1}}^2)\tau.
\end{aligned}$$

Furthermore u_θ is also $\|u_\theta\|_{C^{1,1}}$ -bounded and hence:

$$\mathbb{E}[\Psi_T^2] \leq 2\|u_\theta\|_{C^{1,1}}^2 + 2\mathbb{E}[|Y_T^\theta|^2].$$

Next, since F_θ, G_θ are both bounded, then by Proposition D.4, we have that $\mathbb{E}[|Y_T^\theta|^2] = O(1)$. Putting these bounds together yields:

$$|\mathbb{E}[\hat{\Psi}_N^2] - \mathbb{E}[\Psi_T^2]| \leq O(\tau + \sqrt{\tau}) = O(\sqrt{\tau}), \quad (\text{E.9})$$

since $\tau \leq 1$. To finish the proof, we observe

$$\begin{aligned} L_{\text{EM},\tau}(\theta) &= \mathbb{E}_{x_0 \sim \mu_0, w_n} \frac{1}{T^2} \hat{\Psi}_N^2 \\ &= \mathbb{E}_{x_0 \sim \mu_0, B_t} \frac{1}{T^2} \hat{\Psi}_N^2 && \text{[coupling with Brownian increments]} \\ &= \mathbb{E}_{x_0 \sim \mu_0, B_t} \frac{1}{T^2} \Psi_T^2 + O(\tau^{1/2}) && \text{[using (E.9)]} \\ &= L_{\text{BSDE},T}(\theta) + O(\tau^{1/2}). \end{aligned}$$

□

Proposition E.5. Suppose that $f, g, h_\theta \in C^{0,1}$, $u_\theta \in C^{2,1}$, and $\tau \leq 1$. We have that:

$$\begin{aligned} L_{\text{EM},\tau}(\theta) &= L_{\text{BSDE},\tau}(\theta) + \text{Bias}(\theta) + O(\tau^{1/2}), \\ \text{Bias}(\theta) &:= \frac{1}{2T} \int_0^T \mathbb{E}[\text{tr}((H(X_t, t) \cdot \nabla^2 u_\theta(X_t, t))^2)] dt. \end{aligned} \quad (\text{E.10})$$

Here, $O(\cdot)$ hides factors depending on d, T , and the Hölder norms of f, g, h_θ , and u_θ .

Proof. We have the following:

$$\begin{aligned} L_{\text{EM},\tau}(\theta) &\stackrel{(a)}{=} \frac{1}{T} \int_0^T \mathbb{E} \left((R[u_\theta](X_t, t))^2 + \frac{1}{2} \text{tr}((H(X_t, t) \cdot \nabla^2 u_\theta(X_t, t))^2) \right) dt + O(\tau^{1/2}) \\ &= L_{\text{FS-PINNs}}(\theta) + \text{Bias}(\theta) + O(\tau^{1/2}) \\ &\stackrel{(b)}{=} L_{\text{BSDE},\tau}(\theta) + \text{Bias}(\theta) + O(\tau^{1/2}), \end{aligned}$$

where (a) holds from Theorem 4.2, and (b) holds from Proposition E.2. □

E.2 Stratonovich BSDE and Heun discretization

We first define the Stratonovich variant of the FS-PINNs loss (E.1):

$$L_{\text{S-FS-PINNs}}(\theta) := \mathbb{E}_{x_0 \sim \mu_0, B_t} \frac{1}{T} \int_0^T (R[u_\theta](X_t^\bullet, t))^2 dt.$$

Proposition E.6. Suppose that $u_\theta \in C^2$. We have that:

$$L_{\text{S-BSDE},T}(\theta) \leq L_{\text{S-FS-PINNs}}(\theta).$$

Proof. We mimic the arguments in Proposition E.1. Abbreviating $R_\theta(x, t) = R[u_\theta](x, t)$ and using the Stratonovich chain rule, we have:

$$\begin{aligned} u_\theta(X_T^\bullet, T) - u_\theta(X_0^\bullet, 0) &= \int_0^T \left[R_\theta(X_t^\bullet, t) + h_\theta(X_t^\bullet, t) - \frac{1}{2} \text{tr}(H(X_t^\bullet, t) \nabla^2 u_\theta(X_t^\bullet, t)) \right] dt \\ &\quad + \int_0^T \nabla u_\theta(X_t^\bullet, t)^\top g(X_t^\bullet, t) \circ dB_t, \end{aligned}$$

and hence the following identity which parallels (E.3) holds:

$$L_{\text{S-BSDE},T}(\theta) = \mathbb{E}_{x_0 \sim \mu_0, B_t} \left(\frac{1}{T} \int_0^T R_\theta(X_t^\bullet, t) dt \right)^2.$$

Now we simply apply Jensen's inequality to conclude:

$$\begin{aligned} L_{\text{S-BSDE},T}(\theta) &= \mathbb{E}_{x_0 \sim \mu_0, B_t} \left(\frac{1}{T} \int_0^T R_\theta(X_t^\bullet, t) dt \right)^2 \\ &\leq \mathbb{E}_{x_0 \sim \mu_0, B_t} \frac{1}{T} \int_0^T (R_\theta(X_t^\bullet, t))^2 dt = L_{\text{S-FS-PINNs}}(\theta). \end{aligned}$$

□

Proposition E.7. Suppose that $f, h_\theta \in C^{0,1}$, $g \in C^{1,1}$, $u_\theta \in C^{2,1}$, and $\tau \leq 1$. Then,

$$L_{\text{S-BSDE},\tau}(\theta) = L_{\text{S-FS-PINNs}}(\theta) + O(\tau^{1/2}), \quad (\text{E.11})$$

where the $O(\cdot)$ hides constants depending on the Hölder norms of f , g , h_θ , and u_θ .

Proof. The proof is nearly identical to the proof of Proposition E.2, but with $L_{\text{S-BSDE},\tau}(\theta)$ taking the place of $L_{\text{BSDE},\tau}(\theta)$ and $L_{\text{S-FS-PINNs}}(\theta)$ taking the place of $L_{\text{FS-PINNs}}(\theta)$. The only notable difference is we need to establish the condition $\mathbb{E}\|X_{t_0}^\bullet - X_{t_1}^\bullet\|^2 \leq O(1)|t_0 - t_1|$ for $t_0, t_1 \in \mathcal{I}$. By our assumption that $f \in C^{0,1}$ and $g \in C^{1,1}$, we have that both $f(x, t) + \frac{1}{2} \sum_{k=1}^d \partial_x g^k(x, t)g(x, t)$ and $g(x, t)$ are bounded, and hence the condition $\mathbb{E}\|X_{t_0}^\bullet - X_{t_1}^\bullet\|^2 \leq O(1)|t_0 - t_1|$ holds by Proposition D.5. $\square$

Proposition E.8. Suppose the assumptions of Proposition E.7 hold. Then,

$$L_{\text{S-BSDE},T}(\theta) \leq L_{\text{S-BSDE},\tau}(\theta) + O(\tau^{1/2}),$$

where the $O(\cdot)$ hides constants depending on the Hölder norms of f , g , h_θ , and u_θ .

Proof. Follows immediately from Proposition E.6 and Proposition E.7. $\square$

Proposition E.9. Suppose that $f, h_\theta \in C^{0,1}$, $g \in C^{1,1}$, $u_\theta \in C^{2,1}$, and $\tau \leq 1$. We have that:

$$L_{\text{Heun}_N,\tau}(\theta) = L_{\text{S-BSDE},T}(\theta) + O(\tau^{1/2}), \quad (\text{E.12})$$

where the $O(\cdot)$ hides constants depending on the d , T , and the Hölder norms of f , g , h_θ , and u_θ .

Proof. We consider the joint forward/backward Stratonovich SDE $Z_t^{\bullet,\theta} = (X_t^\bullet, Y_t^{\bullet,\theta})$ from (4.6) of the form $dZ_t^{\bullet,\theta} = F_\theta(Z_t^{\bullet,\theta}, t)dt + G_\theta(Z_t^{\bullet,\theta}, t) \circ dB_t$. We first show that the pair (F_θ, G_θ) is Heun-regular (cf. Definition F.3). A sufficient condition is that (a) $F_\theta \in C^{0,1}$ and (b) $G_\theta \in C^{1,1}$. For condition (a), it is equivalent to both f and $h_\theta^\bullet(x, t) = h_\theta(x, t) - \frac{1}{2} \text{tr}(H(x, t) \nabla^2 u_\theta(x, t))$ are in $C^{0,1}$; the former is by assumption, and the latter holds since $h_\theta \in C^{0,1}$, $g \in C^{1,1}$, and $u_\theta \in C^{2,1}$ by assumption. Now for condition (b), it is equivalent to both g and $\nabla u_\theta(x, t)^\top g(x, t)$ are in $C^{1,1}$. The former is again by assumption, and the latter holds since $u_\theta \in C^{2,1}$ and $g \in C^{1,1}$. Hence by Theorem F.4, we have that Heun discretization $\{\hat{Z}_n^{\bullet,\theta}\}_n$ from (4.7), coupled with the SDE $(Z_t^{\bullet,\theta})_t$ through Brownian increments $\{\Delta W_n\}_n$, satisfies order 1/2 strong convergence to the SDE $(Z_t^{\bullet,\theta})_t$, i.e.,

$$\left(\mathbb{E} \left[\max_{n \in \{0, \dots, N\}} \max\{\|\hat{X}_n^\bullet - X_{t_n}^\bullet\|^2, |\hat{Y}_n^{\bullet,\theta} - Y_{t_n}^{\bullet,\theta}|^2\} \right] \right)^{1/2} \leq C\tau^{1/2},$$

where C does not depend on τ . The remainder of the proof proceeds nearly identically to Proposition E.4, with the only difference being that in order to argue $\mathbb{E}[|Y_T^{\bullet,\theta}|^2] = O(1)$, we utilize that $F_\theta \in C^{0,1}$ and $G_\theta \in C^{1,1}$ to invoke Proposition D.5. $\square$

Showing that $L_{\text{Heun},\tau}(\theta) = L_{\text{S-BSDE},\tau}(\theta) + O(\tau^{1/2})$. We start from (4.10) and following nearly identical arguments as in the derivation of (E.10).

Proposition E.10. Suppose that f , g , and h_θ are all in $C^{1,1}$, $u_\theta \in C^{3,1}$, and $\tau \leq 1$. We have:

$$L_{\text{Heun},\tau}(\theta) = L_{\text{S-BSDE},\tau}(\theta) + O(\tau^{1/2}),$$

where the $O(\cdot)$ hides factors depending on d , T , and the Hölder norms of f , g , h_θ , and u_θ .

Proof. We proceed similarly to Proposition E.5:

$$\begin{aligned} L_{\text{Heun},\tau}(\theta) &\stackrel{(a)}{=} \frac{1}{T} \int_0^T \mathbb{E}(R[u_\theta](X_t^\bullet, t))^2 dt + O(\tau^{1/2}) \\ &= L_{\text{S-FS-PINNs}}(\theta) + O(\tau^{1/2}) \\ &\stackrel{(b)}{=} L_{\text{S-BSDE},\tau}(\theta) + O(\tau^{1/2}), \end{aligned}$$

where (a) holds from Theorem 4.4, and (b) holds from Proposition E.7. $\square$

F Strong Convergence of Euler-Maruyama and Heun Integration

Let $\mathcal{I} := [0, T]$ denote a time interval, and consider functions $a : \mathbb{R}^d \times \mathcal{I} \mapsto \mathbb{R}^d$ and $b : \mathbb{R}^d \times \mathcal{I} \mapsto \mathbb{R}^{d \times m}$ which define the following SDE:

$$dX_t = a(X_t, t)dt + b(X_t, t) \diamond dB_t, \quad X_0 \sim \mathcal{D}_0. \quad (\text{F.1})$$

where $(B_t)_{t \geq 0}$ is m -dimensional Brownian motion and $\mathcal{D}_0$ is an arbitrary distribution over $\mathbb{R}^d$ with bounded second moments, i.e., $\mathbb{E}\|X_0\|^2 < \infty$. Here, the pair (a, b) is used instead of (f, g) to avoid confusion with the forward SDE (3.3), and the $\diamond$ notation denotes that the SDE (F.1) is either to be interpreted as an Itô or Stratonovich SDE. We write $b^k : \mathbb{R}^d \times \mathcal{I} \mapsto \mathbb{R}^d$ for $k \in [m]$ so that $b = (b^1, \dots, b^m)$, i.e., $b^k(t, x)$ is the k -th column of the matrix $b(t, x)$. We consider a discretization time $\tau \in (0, T]$ such that $N := \lfloor T/\tau \rfloor \in \mathbb{N}_+$. We denote the timesteps $\{t_n\}_{n=0}^N$ and Brownian increments $\{\Delta W_n\}_{n=0}^{N-1}$ as $t_n := n\tau$ and $\Delta W_n := B_{t_{n+1}} - B_{t_n}$. We will also often utilize the shorthand notation $a_n(x) := a(x, t_n)$, $b_n(x) := b(x, t_n)$, and $b_n^k(x) := b^k(x, t_n)$ for $n \in \{0, \dots, N\}$.

In this section, we review standard results regarding convergence of basic stochastic integration schemes (Euler-Maruyama for Itô, stochastic Heun for Stratonovich) for the SDE (F.1).

F.1 Euler-Maruyama Convergence

The Euler-Maruyama scheme for integrating the SDE (F.1) interpreted as an Itô SDE is the following discrete-time process:

$$\hat{X}_{n+1} = a_n(\hat{X}_n)\tau + b_n(\hat{X}_n)\Delta W_n, \quad \hat{X}_0 = X_0. \quad (\text{F.2})$$

The order 1/2 strong convergence of the Euler-Maruyama process (F.2) to the Itô SDE (F.1) is thoroughly documented in the literature. Concretely, we will state a result from [34]. First, we define the necessary regularity condition on the drift and diffusion terms

Definition F.1 (EM-regularity). *The pair (a, b) with $a : \mathbb{R}^d \times \mathcal{I} \mapsto \mathbb{R}^d$ and $g : \mathbb{R}^d \times \mathcal{I} \mapsto \mathbb{R}^{d \times m}$ is EM-regular if there exists finite K_1, K_2, K_3 such that for all $x, y \in \mathbb{R}^d$ and $s, t \in \mathcal{I}$,*

$$\begin{aligned} \|a(x, t)\| + \|b(x, t)\|_{\text{op}} &\leq K_1(1 + \|x\|), \\ \|a(x, t) - a(y, t)\| + \|b(x, t) - b(y, t)\|_{\text{op}} &\leq K_2\|x - y\|, \\ \|a(x, s) - a(x, t)\| + \|b(x, s) - b(x, t)\|_{\text{op}} &\leq K_3(1 + \|x\|)|s - t|^{1/2}. \end{aligned}$$

By definition of the Hölder class $C^{0,1}$, we have that if the pair (a, b) satisfies $a, b \in C^{0,1}$, then (a, b) is EM-regular, although Definition F.1 is a weaker assumption. With this notation of regularity in place, we have the following order 1/2 strong convergence result.

Theorem F.2 ([34, Theorem 10.2.2]). *Suppose the pair (a, b) is EM-regular (cf. Definition F.1). Then the Itô SDE $(X_t)_t$ defined in (F.1) and the Euler-Maruyama discretization $\{\hat{X}_n\}_n$ defined in (F.2) satisfy the following bound:*

$$\left(\mathbb{E} \max_{n \in \{0, \dots, N\}} \|X_{t_n} - \hat{X}_n\|^2 \right)^{1/2} \leq C\sqrt{\tau},$$

where the constant C does not depend on τ .

F.2 Stochastic Heun Convergence

The stochastic Heun discretization of the Stratonovich SDE $(X_t)_t$ defined in (F.1) is the discrete-time process with $\hat{X}_0 = X_0$ and:

$$\hat{Y}_{n+1} = \hat{X}_n + a_n(\hat{X}_t)\tau + b_n(\hat{X}_t)\Delta W_n, \quad (\text{F.3a})$$

$$\hat{X}_{n+1} = \hat{X}_n + \frac{1}{2} [a_n(\hat{X}_t) + a_{n+1}(\hat{Y}_{n+1})]\tau + \frac{1}{2} [b_n(\hat{X}_t) + b_{n+1}(\hat{Y}_{n+1})]\Delta W_n. \quad (\text{F.3b})$$

Analogous to the order 1/2 strong convergence results in Appendix F.1 for the EM discretization of the Itô SDE, we also have a similar result that holds for the Heun discretization (F.3) of the

Stratonovich SDE (F.1). While we consider such a result to be a folklore result, we were unable to find a specific theorem statement in the literature listing out a precise set of sufficient conditions on (a, b) for strong convergence to hold.⁴ Thus, the rest of this sub-section provides a result and mostly self-contained proof that builds on top of EM results stated in Appendix F.1.

We first start with a sufficient regularity condition, which adds a few extra assumptions to the EM-regular definition (Definition F.1).

Definition F.3 (Heun-regularity). *The pair (a, b) with $a : \mathbb{R}^d \times \mathcal{I} \mapsto \mathbb{R}^d$ and $b : \mathbb{R}^d \times \mathcal{I} \mapsto \mathbb{R}^{d \times m}$ is Heun-regular if for every $t \in \mathcal{I}$ and $k \in [m]$, the map $x \mapsto b^k(t, x)$ is $C^1(\mathbb{R}^d)$, and there exists finite $K_i, i \in [5]$ such that for all $x, y \in \mathbb{R}^d$ and $s, t \in \mathcal{I}$:*

$$\begin{aligned} \|a(x, t)\| + \|b(x, t)\|_{\text{op}} &\leq K_1, \\ \|a(x, t) - a(y, t)\| + \|b(x, t) - b(y, t)\|_{\text{op}} &\leq K_2\|x - y\|, \\ \|a(x, s) - a(x, t)\| + \|b(x, s) - b(x, t)\|_{\text{op}} &\leq K_3(1 + \|x\|)|s - t|, \\ \|\partial_x b^k(x, t) - \partial_x b^k(y, t)\|_{\text{op}} &\leq K_4\|x - y\|, \\ \|\partial_x b^k(x, s) - \partial_x b^k(x, t)\|_{\text{op}} &\leq K_5(1 + \|x\|)|s - t|^{1/2}. \end{aligned}$$

We note that from the definition of the Hölder classes $C^{0,1}$ and $C^{1,1}$ that if $a \in C^{0,1}$ and $b \in C^{1,1}$, then the pair (a, b) is Heun-regular. The following result is the main convergence result for Heun.

Theorem F.4. *Suppose that the pair (a, b) is Heun-regular (cf. Definition F.3). Then the Stratonovich SDE $(X_t)_t$ defined in (F.1) and the stochastic Heun discretization $\{\hat{X}_n\}_n$ defined in (F.3) satisfy:*

$$\left(\mathbb{E} \max_{n \in \{0, \dots, N\}} \|X_{t_n} - \hat{X}_n\|^2 \right)^{1/2} \leq C\sqrt{\tau}, \quad (\text{F.4})$$

where the constant C does not depend on τ .

F.2.1 Proof of Theorem F.4

Our proof of Theorem F.4 is based on the following reduction. By defining:

$$\bar{a}(x, t) := a(x, t) + \frac{1}{2} \sum_{k=1}^m \partial_x b^k(x, t) b^k(x, t),$$

the Itô SDE:

$$d\bar{X}_t = \bar{a}(\bar{X}_t, t)dt + b(\bar{X}_t, t)dB_t, \quad \bar{X}_0 = X_0, \quad (\text{F.5})$$

defines an identical process as the Stratonovich SDE $(X_t)_t$ from (F.1), i.e., $(X_t(\omega))_t = (\bar{X}_t(\omega))_t$ for almost every t, ω . Furthermore, we can consider an Euler-Maruyama discretization of (F.5):

$$\hat{\bar{X}}_{n+1} = \hat{\bar{X}}_n + \bar{a}_n(\hat{\bar{X}}_n)\tau + b_n(\hat{\bar{X}}_n)\Delta W_n, \quad \hat{\bar{X}}_0 = X_0. \quad (\text{F.6})$$

As the Itô SDE (F.5) and Stratonovich SDE (F.1) are identical, then we also have that if the pair $(\bar{a}, b)$ is EM-regular (cf. Definition F.1), then the EM discretization (F.6) is order 1/2 strongly convergent to the Stratonovich SDE (F.1). Hence, this reduces the problem to comparing the two discrete processes $\{\hat{X}_n\}_n$ from (F.3) and $\{\hat{\bar{X}}_n\}_n$ from (F.6). In particular, if we can show that:

$$\left(\mathbb{E} \max_{n \in \{0, \dots, N\}} \|\hat{X}_n - \hat{\bar{X}}_n\|^2 \right)^{1/2} \leq C\sqrt{\tau},$$

where again the two processes are coupled under the same Brownian motion $(B_t)_t$ and initial condition X_0 , then by triangle inequality we have the desired result Theorem F.4. One advantage of this proof strategy is that we only need to study the evolution of two discrete-time processes, which we can do with purely elementary (discrete-time) martingale techniques, avoiding the need for any stochastic calculus. Indeed, the main tools we utilize are the following two results.

⁴The closest statement we were able to find in the literature is [41, Theorem D.12], which shows order 1/2 strong convergence of the *reversible* Heun method, which is a modified version of the stochastic Heun method that is algebraically reversible.

Proposition F.5 (Doob's maximal inequality (vector-valued), cf. [50, Theorem 3.2.2]). Let $(X_n)_{n \in \mathbb{N}_+}$ denote a martingale taking values in a normed vector space X with norm $\|\cdot\|_X$. We have that for any $p \in (0, \infty]$ and $n \in \mathbb{N}_+$:

$$\mathbb{E} \left(\max_{i \in [n]} \|X_i\|_X^p \right)^{1/p} \leq \frac{p}{p-1} (\mathbb{E}[\|X_n\|_X^p])^{1/p}.$$

Proposition F.6 (Discrete Gronwall inequality, cf. [51]). Let $\{x_n\}_{n \in \mathbb{N}}$ and $\{\beta_n\}_{n \in \mathbb{N}}$ be non-negative sequences satisfying for some $\alpha > 0$:

$$x_n \leq \alpha + \sum_{k=0}^{n-1} \beta_k x_n, \quad n \in \mathbb{N}.$$

Then we have:

$$x_n \leq \alpha \exp \left(\sum_{k=0}^{n-1} \beta_k \right), \quad n \in \mathbb{N}.$$

Here, we interpret $\sum_{k=0}^{-1}$ to indicate zero.

Our first step shows that Heun-regularity of the pair (a, b) implies EM-regularity of the pair $(\bar{a}, b)$.

Proposition F.7. If the pair (a, b) is Heun-regular (cf. Definition F.3), then the pair $(\bar{a}, b)$ is EM-regular (cf. Definition F.1).

Proof. We let $K := \max_{i \in [5]} K_i$. We first check the growth condition on $\|\bar{a}(x, t)\|$:

$$\begin{aligned} \|\bar{a}(x, t)\| &= \left\| a(x, t) + \frac{1}{2} \sum_{k=1}^m \partial_x b^k(x, t) b^k(x, t) \right\| \\ &\stackrel{(a)}{\leq} K + \frac{1}{2} \sum_{k=1}^m \|\partial_x b^k(x, t) b^k(x, t)\| \stackrel{(b)}{\leq} K + \frac{K}{2} \sum_{k=1}^m \|\partial_x b^k(x, t)\|_{\text{op}} \stackrel{(c)}{\leq} K + \frac{K^2 m}{2}, s \end{aligned}$$

where (a) holds since $\|a(x, t)\| \leq K$, (b) holds since $\|b^k(x, t)\| \leq \|b(x, t)\|_{\text{op}} \leq K$, and (c) holds since $\|g^k(x, t) - g^k(y, t)\| \leq K\|x - y\|$ implies that $\|\partial_x g^k(x, t)\|_{\text{op}} \leq K$.

Next, we check the Lipschitz condition over x on $\bar{a}(x, t)$:

$$\begin{aligned} \|\bar{a}(x, t) - \bar{a}(y, t)\| &\leq \|a(x, t) - a(y, t)\| + \frac{1}{2} \sum_{k=1}^m \|\partial_x b^k(x, t) b^k(x, t) - \partial_x b^k(y, t) b^k(y, t)\| \\ &\stackrel{(a)}{\leq} K\|x - y\| + \frac{1}{2} \sum_{k=1}^m \|\partial_x b^k(x, t)[b^k(x, t) - b^k(y, t)]\| \\ &\quad + \frac{1}{2} \sum_{k=1}^m \|[\partial_x b^k(x, t) - \partial_x b^k(y, t)]b^k(y, t)\| \\ &\stackrel{(b)}{\leq} K\|x - y\| + \frac{K^2 m}{2} \|x - y\| + \frac{K^2 m}{2} \|x - y\| = (K + K^2 m) \|x - y\|. \end{aligned}$$

where (a) uses the Lipschitz condition $\|a(x, t) - a(y, t)\| \leq K\|x - y\|$, and (b) uses $\|\partial_x b^k(x, t) - \partial_x b^k(y, t)\|_{\text{op}} \leq K\|x - y\|$, $\|b^k(x, t)\| \leq K$, $\|\partial_x b^k(x, t)\|_{\text{op}} \leq K$, and $\|b^k(x, t) - b^k(y, t)\| \leq K\|x - y\|$,

Finally, we check the Hölder 1/2 condition over t on $\bar{a}(x, t)$:

$$\begin{aligned} \|\bar{a}(x, s) - \bar{a}(x, t)\| &\leq \|a(x, s) - a(x, t)\| + \frac{1}{2} \sum_{k=1}^m \|\partial_x b^k(x, s) b^k(x, s) - \partial_x b^k(x, t) b^k(x, t)\| \\ &\stackrel{(a)}{\leq} K(1 + \|x\|)|s - t| + \frac{1}{2} \sum_{k=1}^m \|\partial_x b^k(x, s)[b^k(x, s) - b^k(x, t)]\| \\ &\quad + \frac{1}{2} \sum_{k=1}^m \|[\partial_x b^k(x, s) - \partial_x b^k(x, t)]b^k(x, t)\| \end{aligned}$$

$$\begin{aligned}
&\stackrel{(b)}{\leq} K(1 + \|x\|)|s - t| + \frac{K^2 m}{2}(1 + \|x\|)|s - t| + \frac{K^2 m}{2}(1 + \|x\|)|s - t|^{1/2} \\
&\stackrel{(c)}{\leq} (K + K^2 m)\sqrt{T}(1 + \|x\|)|s - t|^{1/2},
\end{aligned}$$

where in (a) we use $\|a(x, s) - a(x, t)\| \leq K(1 + \|x\|)|s - t|$, in (b) we use $\|b^k(x, s) - b^k(x, t)\| \leq K(1 + \|x\|)|s - t|$, $\|\partial_x b^k(x, t)\|_{\text{op}} \leq K$, $\|\partial_x b^k(x, s) - \partial_x b^k(x, t)\| \leq K(1 + \|x\|)|s - t|^{1/2}$, and $\|b^k(x, t)\| \leq K$, and in (c) we use $|s - t| = |s - t|^{1/2} \cdot |s - t|^{1/2} \leq \sqrt{T}|s - t|^{1/2}$.

Since the growth, Lipschitz, and Hölder 1/2 conditions on $b(x, t)$ are immediate from the Heun regularity assumptions, this concludes the claim. $\square$

The following result shows that the discrete-time processes (F.3) and (F.6) are strongly convergent.

Lemma F.8. *Suppose that the pair (a, b) is Heun-regular (cf. Definition F.3) and that $\tau \leq 1$. Then, we have that the updates (F.3) and (F.6) satisfy:*

$$\left(\mathbb{E} \max_{n \in \{0, \dots, N\}} \|\hat{X}_n - \hat{\bar{X}}_n\|^2 \right)^{1/2} \leq C\sqrt{\tau},$$

where the constant C does not depend on τ .

Proof. To start the proof, we recall that $t_n = n\tau$, $N = \lfloor T/\tau \rfloor$ which is assumed to be a positive integer, and $\Delta W_n := B_{t_{n+1}} - B_{t_n}$. We will equivalently write $\Delta W_n = \sqrt{\tau}w_n$, where $\{w_n\}_{n=0}^{N-1}$ are i.i.d. $N(0, I_m)$ random vectors. In order to index the coordinates of both ΔW_n and w_n , we use the notation ΔW_n^i and w_n^i , for $i \in [m]$, to refer to the i -th coordinate of the vector. We let $K := 1 + \mathbb{E}\|X_0\|^2 + \max_{i \in [5]} K_i$ to denote a bound on all the parameters from both Heun regularity and the second moment of the initial condition X_0 . As with a_n, b_n , we also define $\bar{a}_n(x) := \bar{a}(x, t_n)$ for $n \in [N]$. Furthermore, we will drop the hat notation to reduce notational clutter and write $X_n, Y_n = \hat{X}_n, \hat{Y}_n$, and similarly $\bar{X}_n = \hat{\bar{X}}_n$; since we are not dealing with the SDEs (F.1) and (F.5), there is no risk of confusion with this notation.

To avoid keeping track of explicit dependence on constants besides τ , for a set of parameters Q we use C_Q to denote a finite constant that depends only on the parameters listed in Q , and $a \lesssim_Q b$ to denote $a \leq C_Q b$. For example, $C_{K,m}$ is a constant that depends only on (K, m) (its dependency may be arbitrary however). We also let $a \lesssim b$ denote $a \leq Cb$ where C is a universal positive constant.

With the aforementioned notation in place, we have the following discrete-time update rules for (F.6):

$$\bar{X}_{n+1} = \bar{X}_n + \bar{a}_n(\bar{X}_n)\tau + b_n(\bar{X}_n)\Delta W_n,$$

and for (F.3):

$$\begin{aligned}
Y_{n+1} &= X_n + a_n(X_n)\tau + b_n(X_n)\Delta W_n, \\
X_{n+1} &= X_n + \frac{1}{2}[a_n(X_n) + a_{n+1}(Y_{n+1})]\tau + \frac{1}{2}[b_n(X_n) + b_{n+1}(Y_{n+1})]\Delta W_n.
\end{aligned}$$

Let us define the filtration $\mathcal{F}_n := \sigma(w_0, \dots, w_{n-1})$ for $n \in \mathbb{N}_+$ with $\mathcal{F}_0$ the trivial σ -algebra. We observe a key property of both $\{X_n\}_n$ and $\{\bar{X}_n\}_n$ is that both X_n and $\bar{X}_n$ are $\mathcal{F}_n$ -measurable. Hence by the tower property we have for expressions A, B , $\mathbb{E}[A(X_n)B(w_n)] = \mathbb{E}[A(X_n)\mathbb{E}[B(w_n) | \mathcal{F}_n]]$ and $\mathbb{E}[B(w_n) | \mathcal{F}_n] = \mathbb{E}_{w_n \sim N(0, I_m)}[B(w_n)]$, a property we make heavy use of in our calculations. Another simple inequality we make use of is that for any $q \in \mathbb{N}_+$ and any set of vectors $v_1, \dots, v_q$,

$$\left\| \sum_{i=1}^q v_i \right\|^2 \leq q \sum_{i=1}^q \|v_i\|^2,$$

which follows by triangle inequality and Cauchy-Schwarz.

Heun update decomposition. To relate the Heun and EM updates, we write the Heun update as:

$$X_{n+1} = X_n + \bar{a}_n(X_n)\tau + b_n(X_n)\Delta W_n + E_n, \quad (\text{F.7})$$

where E_n contains the residual terms:

$$E_n := \underbrace{\frac{1}{2} [a_{n+1}(Y_{n+1}) - a_n(X_n)] \tau}_{=: E_n^a} + \underbrace{\frac{1}{2} [b_{n+1}(Y_{n+1}) - b_n(X_n)] \Delta W_n - \frac{\tau}{2} \sum_{i=1}^m \partial_x b_n^k(X_n) b_n^k(X_n)}_{=: E_n^b}.$$

We further decompose E_n^b as follows. We first write:

$$b_{n+1}(Y_{n+1}) - b_n(X_n) = (b_{n+1}(Y_{n+1}) - b_n(Y_{n+1})) + (b_n(Y_{n+1}) - b_n(X_n)).$$

Next, we use the Heun-regularity to expand the RHS above:

$$\begin{aligned} b_n^k(Y_{n+1}) - b_n^k(X_n) &= \partial_x b_n^k(X_n)(Y_{n+1} - X_n) + R_n^k \\ &= \partial_x b_n^k(X_n)(a_n(X_n)\tau + b_n(X_n)\Delta W_n) + R_n^k, \end{aligned}$$

where the remainder term R_n^k satisfies $\|R_n^k\| \leq K\|Y_{n+1} - X_n\|^2$. Hence,

$$\begin{aligned} &(b_n(Y_{n+1}) - b_n(X_n))\Delta W_n \\ &= \sum_{k=1}^m (b_n^k(Y_{n+1}) - b_n^k(X_n))\Delta W_n^k \\ &= \tau \sum_{k=1}^m \partial_x b_n^k(X_n) a_n(X_n) \Delta W_n^k + \sum_{k=1}^m \partial_x b_n^k(X_n) b_n(X_n) \Delta W_n \Delta W_n^k + \sum_{k=1}^m R_n^k \Delta W_n^k \end{aligned}$$

Now the middle term further decomposes as:

$$\begin{aligned} \sum_{k=1}^m \partial_x b_n^k(X_n) b_n(X_n) \Delta W_n \Delta W_n^k &= \sum_{k_1, k_2=1}^m \partial_x b_n^{k_1}(X_n) b_n^{k_2}(X_n) \Delta W_n^{k_1} \Delta W_n^{k_2} \\ &= \sum_{k_1, k_2=1}^m \partial_x b_n^{k_1}(X_n) b_n^{k_2}(X_n) (\Delta W_n^{k_1} \Delta W_n^{k_2} - \tau \mathbb{1}_{\{k_1=k_2\}}) \\ &\quad + \tau \sum_{k=1}^m \partial_x b_n^k(X_n) b_n^k(X_n). \end{aligned}$$

Combining these decompositions together,

$$\begin{aligned} E_n^b &= \frac{1}{2} [b_{n+1}(Y_{n+1}) - b_n(X_n)] \Delta W_n - \frac{\tau}{2} \sum_{i=1}^m \partial_x b_n^k(X_n) b_n^k(X_n) \\ &= \frac{1}{2} [b_{n+1}(Y_{n+1}) - b_n(Y_{n+1})] \Delta W_n + \frac{1}{2} [b_n(Y_{n+1}) - b_n(X_n)] \Delta W_n - \frac{\tau}{2} \sum_{i=1}^m \partial_x b_n^k(X_n) b_n^k(X_n) \\ &= \frac{1}{2} [b_{n+1}(Y_{n+1}) - b_n(Y_{n+1})] \Delta W_n + \frac{\tau}{2} \sum_{k=1}^m \partial_x b_n^k(X_n) a_n(X_n) \Delta W_n^k + \frac{1}{2} \sum_{k=1}^m R_n^k \Delta W_n^k \\ &\quad + \frac{1}{2} \sum_{k=1}^m \partial_x b_n^k(X_n) b_n(X_n) \Delta W_n \Delta W_n^k - \frac{\tau}{2} \sum_{i=1}^m \partial_x b_n^k(X_n) b_n^k(X_n) \\ &= \underbrace{\frac{1}{2} [b_{n+1}(Y_{n+1}) - b_n(Y_{n+1})] \Delta W_n}_{=: E_n^{b,1}} + \underbrace{\frac{\tau}{2} \sum_{k=1}^m \partial_x b_n^k(X_n) a_n(X_n) \Delta W_n^k}_{=: E_n^{b,2}} + \underbrace{\frac{1}{2} \sum_{k=1}^m R_n^k \Delta W_n^k}_{=: E_n^{b,3}} \\ &\quad + \underbrace{\frac{1}{2} \sum_{k_1, k_2=1}^m \partial_x b_n^{k_1}(X_n) b_n^{k_2}(X_n) (\Delta W_n^{k_1} \Delta W_n^{k_2} - \tau \mathbb{1}_{\{k_1=k_2\}})}_{=: E_n^{b,4}}. \end{aligned}$$

Thus, (F.7) becomes:

$$X_{n+1} = X_n + \bar{a}_n(X_n)\tau + b_n(X_n)\Delta W_n + E_n^a + \sum_{\ell=1}^4 E_n^{b,\ell}, \quad (\text{F.8})$$

which serves as the starting point for what follows.

Second moment bounds on error terms. Our next step is to bound the second moment of all the error terms separately. Here we make heavy use of the Heun-regularity conditions.

Bound on $\mathbb{E}\|E_n^a\|^2$: We write:

$$\begin{aligned}\|a_{n+1}(Y_{n+1}) - a_n(X_n)\| &= \|(a_{n+1}(Y_{n+1}) - a_{n+1}(X_n)) + (a_{n+1}(X_n) - a_n(X_n))\| \\ &\leq \|a_{n+1}(Y_{n+1}) - a_{n+1}(X_n)\| + \|a_{n+1}(X_n) - a_n(X_n)\| \\ &\leq K\|Y_{n+1} - X_n\| + K(1 + \|X_n\|)\sqrt{\tau} \\ &\leq K(\|a_n(X_n)\|\tau + \|b_n(X_n)\|\sqrt{\tau}\|w_n\|) + K(1 + \|X_n\|)\sqrt{\tau} \\ &\leq K^2(\tau + \sqrt{\tau}\|w_n\|) + K(1 + \|X_n\|)\sqrt{\tau}.\end{aligned}$$

Hence,

$$\mathbb{E}\|a_{n+1}(Y_{n+1}) - a_n(X_n)\|^2 \lesssim K^4(\tau^2 + \tau m) + K^2(1 + \mathbb{E}\|X_n\|^2)\tau \lesssim_{K,m} (1 + \mathbb{E}\|X_n\|^2)\tau.$$

which implies:

$$\mathbb{E}\|E_n^a\|^2 = \frac{\tau^2}{4}\mathbb{E}\|a_{n+1}(Y_{n+1}) - a_n(X_n)\|^2 \lesssim_{K,m} (1 + \mathbb{E}\|X_n\|^2)\tau^3.$$

Bound on $\mathbb{E}\|E_n^{b,1}\|^2$: We have:

$$\begin{aligned}\mathbb{E}\|E_n^{b,1}\|^2 &= \frac{1}{4}\mathbb{E}\|[b_{n+1}(Y_{n+1}) - b_n(Y_{n+1})]\Delta W_n\|^2 \\ &\lesssim K^2\tau^3\mathbb{E}[(1 + \|Y_{n+1}\|^2)\|w_n\|^2] \\ &\lesssim K^2\tau^3\mathbb{E}[(1 + \|X_n\|^2 + \|a_n(X_n)\|^2\tau^2 + \|b_n(X_n)\|_{\text{op}}^2\|w_n\|^2\tau)\|w_n\|^2] \\ &\lesssim K^4\tau^3\mathbb{E}[(1 + \|X_n\|^2 + \|w_n\|^2)\|w_n\|^2] \\ &\lesssim_{K,m} (1 + \mathbb{E}\|X_n\|^2)\tau^3.\end{aligned}$$

Bound on $\mathbb{E}\|E_n^{b,2}\|^2$: We have:

$$\begin{aligned}\mathbb{E}\|E_n^{b,2}\|^2 &= \frac{\tau^2}{4}\mathbb{E}\left\|\sum_{k=1}^m \partial_x b_n^k(X_n) a_n(X_n) \Delta W_n^k\right\|^2 = \frac{\tau^3}{4} \sum_{k=1}^m \mathbb{E}\|\partial_x b_n^k(X_n) a_n(X_n)\|^2 \\ &\leq \frac{\tau^3}{4} \sum_{k=1}^m \mathbb{E}\|\partial_x b_n^k(X_n)\|_{\text{op}}^2 \|a_n(X_n)\|^2 \lesssim_{K,m} \tau^3.\end{aligned}$$

Bound on $\mathbb{E}\|E_n^{b,3}\|^2$: Recall that the residual R_n^k satisfies $\|R_n^k\| \leq K\|Y_{n+1} - X_n\|^2$. We have:

$$\begin{aligned}\mathbb{E}\|E_n^{b,3}\|^2 &= \frac{1}{4}\mathbb{E}\left\|\sum_{k=1}^m R_n^k \Delta W_n^k\right\|^2 \\ &\lesssim \tau m \sum_{k=1}^m \mathbb{E}[\|R_n^k\|^2 |w_n^k|^2] \\ &\lesssim \tau K^2 m \sum_{k=1}^m \mathbb{E}[\|Y_{n+1} - X_n\|^4 |w_n^k|^2] \\ &\lesssim \tau K^2 m \sum_{k=1}^m \mathbb{E}[(\|a_n(X_n)\|\tau + \|b_n(X_n)\|_{\text{op}}\|w_n\|\sqrt{\tau})^4 |w_n^k|^2] \\ &\lesssim_{K,m} \tau^3.\end{aligned}$$

Bound on $\mathbb{E}\|E_n^{b,4}\|^2$: We have:

$$\mathbb{E}\|E_n^{b,4}\|^2 = \frac{\tau^2}{4}\mathbb{E}\left\|\sum_{k_1, k_2=1}^m \partial_x b_n^{k_1}(X_n) b_n^{k_2}(X_n) (w_n^{k_1} w_n^{k_2} - \mathbf{1}_{\{k_1=k_2\}})\right\|^2 \lesssim_{K,m} \tau^2.$$

Second moment bounds on the Heun process. We now use (F.8) to write for any $n \in [N]$:

$$\begin{aligned}\mathbb{E}\|X_n\|^2 &= \mathbb{E}\left\|X_0 + \sum_{i=0}^{n-1} \bar{a}_i(X_i)\tau + \sum_{i=0}^{n-1} b_i(X_i)\Delta W_i + \sum_{i=0}^{n-1} \left(E_i^a + \sum_{\ell=1}^4 E_i^{b,\ell}\right)\right\|^2 \\ &\lesssim \mathbb{E}\|X_0\|^2 + \tau^2 \mathbb{E}\left\|\sum_{i=0}^{n-1} \bar{a}_i(X_i)\right\|^2 + \mathbb{E}\left\|\sum_{i=0}^{n-1} b_i(X_i)\Delta W_i\right\|^2 + \mathbb{E}\left\|\sum_{i=0}^{n-1} E_i^a\right\|^2 + \sum_{\ell=1}^4 \mathbb{E}\left\|\sum_{i=0}^{n-1} E_i^{b,\ell}\right\|^2.\end{aligned}\tag{F.9}$$

We now focus on bounding these second moments, using the fact that for $n \in [N]$, we have $\tau n \leq \tau N = \tau \lceil T/\tau \rceil \leq T$. First, we have $\|a_i(X_i)\| \leq K + \frac{K^2 m}{2}$ and hence

$$\tau^2 \mathbb{E}\left\|\sum_{i=0}^{n-1} \bar{a}_i(X_i)\right\|^2 \lesssim \tau^2 n \sum_{i=0}^{n-1} \mathbb{E}\|\bar{a}_i(X_i)\|^2 \lesssim_{K,m} (n\tau)^2 \lesssim_{K,m,T} 1.$$

Next, since $\{b_i(X_i)\Delta W_i\}_i$ forms a martingale difference sequence (MDS),

$$\mathbb{E}\left\|\sum_{i=0}^{n-1} b_i(X_i)\Delta W_i\right\|^2 = \sum_{i=0}^{n-1} \mathbb{E}\|b_i(X_i)\Delta W_i\|^2 = \tau \sum_{i=0}^{n-1} \mathbb{E}\|b_i(X_i)\|_F^2 \lesssim_{K,m} n\tau \lesssim_{K,m,T} 1.$$

Next, we have the following bound using the second moment computations for the error terms:

$$\begin{aligned}\mathbb{E}\left\|\sum_{i=0}^{n-1} E_i^a\right\|^2 + \mathbb{E}\left\|\sum_{i=0}^{n-1} E_i^{b,1}\right\|^2 &\lesssim n \sum_{i=0}^{n-1} (\mathbb{E}\|E_i^a\|^2 + \mathbb{E}\|E_i^{b,1}\|^2) \lesssim_{K,m} n \sum_{i=0}^{n-1} (1 + \mathbb{E}\|X_i\|^2)\tau^3 \\ &\lesssim_{K,m} (n\tau)^2 \tau + n\tau^3 \sum_{i=0}^{n-1} \mathbb{E}\|X_i\|^2 \lesssim_{K,m,T} \tau + \tau^2 \sum_{i=0}^{n-1} \mathbb{E}\|X_i\|^2, \\ \mathbb{E}\left\|\sum_{i=0}^{n-1} E_i^{b,2}\right\|^2 + \mathbb{E}\left\|\sum_{i=0}^{n-1} E_i^{b,3}\right\|^2 &\lesssim n \sum_{i=0}^{n-1} (\mathbb{E}\|E_i^{b,2}\|^2 + \mathbb{E}\|E_i^{b,3}\|^2) \lesssim_{K,m} (n\tau)^2 \tau \lesssim_{K,m,T} \tau.\end{aligned}$$

Furthermore, since $\{E_i^{b,4}\}_i$ is an MDS,

$$\mathbb{E}\left\|\sum_{i=0}^{n-1} E_i^{b,4}\right\|^2 = \sum_{i=0}^{n-1} \mathbb{E}\|E_i^{b,4}\|^2 \lesssim_{K,m} (n\tau)\tau \lesssim_{K,m,T} \tau.$$

Combining these bounds together in (F.9), we obtain:

$$\mathbb{E}\|X_n\|^2 \lesssim_{K,m,T} 1 + \tau^2 \sum_{i=0}^{n-1} \mathbb{E}\|X_i\|^2.$$

By the discrete Gronwall lemma (Proposition F.6), we have:

$$\mathbb{E}\|X_n\|^2 \leq C_{K,m,T} \exp(n\tau^2 C'_{K,m,T}) \leq C_{K,m,T} \exp(\tau C''_{K,m,T}) \lesssim_{K,m,T} 1.$$

Hence, we have shown that:

$$\max_{n \in \{0, \dots, N\}} \mathbb{E}\|X_n\|^2 \lesssim_{K,m,T} 1.$$

Final result. Our goal now is to estimate $\mathbb{E}[\Delta_n^2]$ for $\Delta_n := \max_{i \in [n]} \|\delta_i\|$, where $\delta_i := \bar{X}_i - X_i$. We start with:

$$\begin{aligned}\delta_{n+1} &= \bar{X}_{n+1} - X_{n+1} \\ &= \delta_n + [\bar{a}_n(\bar{X}_n) - \bar{a}_n(X_n)]\tau + [b_n(\bar{X}_n) - b_n(X_n)]\Delta W_n - E_n.\end{aligned}$$

Hence for $n \in [N]$,

$$\|\delta_n\|^2 = \left\|\tau \sum_{i=0}^{n-1} [\bar{a}_i(\bar{X}_i) - \bar{a}_i(X_i)] + \sum_{i=0}^{n-1} [b_i(\bar{X}_i) - b_i(X_i)]\Delta W_i - \sum_{i=0}^{n-1} E_i\right\|^2$$

$$\begin{aligned} &\lesssim \tau^2 \left\| \sum_{i=0}^{n-1} [\bar{a}_i(\bar{X}_i) - \bar{a}_i(X_i)] \right\|^2 + \left\| \sum_{i=0}^{n-1} [b_i(\bar{X}_i) - b_i(X_i)] \Delta W_i \right\|^2 + \left\| \sum_{i=0}^{n-1} E_i \right\|^2 \\ &\lesssim_K \tau^2 n \sum_{i=0}^{n-1} \|\delta_i\|^2 + \left\| \sum_{i=0}^{n-1} [b_i(\bar{X}_i) - b_i(X_i)] \Delta W_i \right\|^2 + n \sum_{i=0}^{n-1} \left(\|E_i^a\|^2 + \sum_{\ell=1}^3 \|E_i^{b,\ell}\|^2 \right) + \left\| \sum_{i=0}^{n-1} E_i^{b,4} \right\|^2. \end{aligned}$$

Therefore,

$$\begin{aligned} \Delta_n^2 = \max_{k \in [n]} \|\delta_k\|^2 &\lesssim_K \tau^2 n \sum_{i=0}^{n-1} \|\delta_i\|^2 + n \sum_{i=0}^{n-1} \left(\|E_i^a\|^2 + \sum_{\ell=1}^3 \|E_i^{b,\ell}\|^2 \right) \\ &\quad + \max_{k \in [n]} \left\| \sum_{i=0}^{k-1} [b_i(\bar{X}_i) - b_i(X_i)] \Delta W_i \right\|^2 + \max_{k \in [n]} \left\| \sum_{i=0}^{k-1} E_i^{b,4} \right\|^2 \\ &\lesssim_K \tau^2 n \sum_{i=0}^{n-1} \Delta_i^2 + n \sum_{i=0}^{n-1} \left(\|E_i^a\|^2 + \sum_{\ell=1}^3 \|E_i^{b,\ell}\|^2 \right) \\ &\quad + \max_{k \in [n]} \left\| \sum_{i=0}^{k-1} [b_i(\bar{X}_i) - b_i(X_i)] \Delta W_i \right\|^2 + \max_{k \in [n]} \left\| \sum_{i=0}^{k-1} E_i^{b,4} \right\|^2. \quad (\text{F.10}) \end{aligned}$$

Now, using nearly identical arguments as in the second moment calculation of the error terms, in addition to the uniform bound on $\mathbb{E}\|X_n\|^2$ for $n \in \{0, \dots, N\}$, we have:

$$\mathbb{E} \left[n \sum_{i=0}^{n-1} \left(\|E_i^a\|^2 + \sum_{\ell=1}^3 \|E_i^{b,\ell}\|^2 \right) \right] \lesssim_{K,m,T} \tau.$$

On the other hand, since both $\{[b_i(\bar{X}_i) - b_i(X_i)] \Delta W_i\}_i$ and $\{E_i^{b,4}\}_i$ are both martingale difference sequences, using Doob's maximal inequality (Proposition F.5),

$$\begin{aligned} &\mathbb{E} \max_{k \in [n]} \left\| \sum_{i=0}^{k-1} [b_i(\bar{X}_i) - b_i(X_i)] \Delta W_i \right\|^2 + \mathbb{E} \max_{k \in [n]} \left\| \sum_{i=0}^{k-1} E_i^{b,4} \right\|^2 \\ &\lesssim \mathbb{E} \left\| \sum_{i=0}^{n-1} [b_i(\bar{X}_i) - b_i(X_i)] \Delta W_i \right\|^2 + \mathbb{E} \left\| \sum_{i=0}^{n-1} E_i^{b,4} \right\|^2 \\ &= \sum_{i=0}^{n-1} \mathbb{E} \left\| [b_i(\bar{X}_i) - b_i(X_i)] \Delta W_i \right\|^2 + \sum_{i=0}^{n-1} \mathbb{E} \|E_i^{b,4}\|^2 \\ &\lesssim_{K,m,T} \tau \sum_{i=0}^{n-1} \mathbb{E} \|\delta_i\|^2 + \tau \lesssim_{K,m,T} \tau \sum_{i=0}^{n-1} \mathbb{E} [\Delta_i^2] + \tau. \end{aligned}$$

Hence, taking expectation in (F.10) and combining the previous second moment estimates, we have:

$$\mathbb{E} [\Delta_n^2] \lesssim_{K,m,T} \tau \sum_{i=0}^{n-1} \mathbb{E} [\Delta_i^2] + \tau.$$

From the discrete Gronwall inequality (Proposition F.6),

$$\mathbb{E} [\Delta_n^2] \leq C_{K,m,T} \tau \exp(n\tau C'_{K,m,T}) \lesssim_{K,m,T} \tau,$$

which completes the proof. $\square$

We can now complete the proof of Theorem F.4.

Proof of Theorem F.4. We have:

$$X_{t_n} - \hat{X}_n \stackrel{(a)}{=} \bar{X}_{t_n} - \hat{X}_n = (\bar{X}_{t_n} - \hat{\bar{X}}_n) + (\hat{\bar{X}}_n - \hat{X}_n),$$

where (a) holds since the Stratonovich SDE (X_t) (F.1) and the Itô SDE $(\bar{X}_t)$ (F.5) are identical for a.e. (ω, t) . Hence we have:

$$\mathbb{E} \max_{n \in \{0, \dots, N\}} \|X_{t_n} - \hat{X}_n\|^2 \lesssim \mathbb{E} \max_{n \in \{0, \dots, N\}} \|\bar{X}_{t_n} - \hat{\bar{X}}_n\|^2 + \mathbb{E} \max_{n \in \{0, \dots, N\}} \|\hat{\bar{X}}_n - \hat{X}_n\|^2.$$

Next, since the pair (a, b) is Heun-regular by assumption, then the pair $(\bar{a}, b)$ is EM-regular by Proposition F.7. Hence by Theorem F.2, we have the EM discretization $\{\hat{\bar{X}}_n\}_n$ is order 1/2 strongly convergent to the SDE $(\bar{X}_t)_t$, i.e., $\mathbb{E} \max_{n \in \{0, \dots, N\}} \|\bar{X}_{t_n} - \hat{\bar{X}}_n\|^2 \leq C\tau$. Furthermore, by Lemma F.8 we have $\mathbb{E} \max_{n \in \{0, \dots, N\}} \|\hat{\bar{X}}_n - \hat{X}_n\|^2 \leq C'\tau$ as well. Note in both cases, C, C' do not depend on τ , and hence the proof is complete. $\square$

G Experimental Details

We use each algorithm to train an 8-layer neural network with 64 neurons per layer and swish activation [52] to model the solution $u(x, t)$ of a PDE. The boundary condition is enforced by adding a boundary condition penalty $\mathbb{E}_{x \sim \mu'}[(u_\theta(x, T) - \phi(x))^2] + \mathbb{E}[\|\nabla u_\theta(x, T) - \nabla \phi(x)\|^2]$ involving both the zero-th and first-order values of ϕ [4], where the distribution μ' is taken over each method's approximation of the distribution of X_T . Additionally, following state-of-the-art PINN architectures practices [7], Fourier embeddings [53] with a 256 embedding dimension and skip connections on odd layers are used. We use a trajectory batch size of 64, translating to 64 realizations of the underlying Brownian motions. Additionally, we utilize a sub-sampling batch size of 1024 for the batched algorithm runs. We use the Adam optimizer with a multi-step learning rate schedule [4] of $10^{-3}, 10^{-4}$, and 10^{-5} at 50k, 75k, and 100k iterations, respectively. All models are trained on a single NVIDIA A100 GPU node in our internal cluster, using the jax library [54].

G.1 PDE Test Cases

Hamilton-Jacobi-Bellman (HJB) Equation. First, we consider the following Hamilton-Jacobi-Bellman (HJB) equation studied in [4]:

$$\partial_t u(x, t) = -\text{Tr}[\nabla^2 u(x, t)] + \|\nabla u(x, t)\|^2, \quad x \in \mathbb{R}^d, \quad t \in [0, T].$$

For the terminal condition $u(x, T) = \phi(x) = \ln(.5(1 + \|x\|^2))$, the analytical solution is given as

$$u(x, t) = -\ln\left(\mathbb{E}\left[\exp\left(-g(x + \sqrt{2}B_{T-t})\right)\right]\right). \quad (\text{G.1})$$

The HJB PDE is related to the forward-backward stochastic differential equation of the form:

$$\begin{aligned} dX_t &= \sigma dB_t, \quad t \in [0, T], \\ dY_t &= \|Z_t\|^2 dt + \sigma Z_t^\top dB_t, \quad t \in [0, T], \end{aligned}$$

where $T = 1$, $\sigma = \sqrt{2}$, $X_0 = 0$, and $Y_T = \phi(X_T)$. Additionally, the Stratonovich SDE is given as:

$$\begin{aligned} dX_t &= \sigma \circ dB_t, \quad t \in [0, T], \\ dY_t &= \left[\|Z_t\|^2 - \frac{1}{2} \text{Tr}[\sigma^2 \nabla^2 u(X_t, t)] \right] dt + \sigma Z_t^\top \circ dB_t, \end{aligned}$$

In our experiments, in order to compute the analytical solution (G.1), we approximate it using 10^5 Monte-Carlo samples.

Black-Scholes-Barenblatt (BSB) Equation. Next, we consider the 100D Black-Scholes-Barenblatt (BSB) equation from [4] of the form

$$\partial_t u(x, t) = -\frac{1}{2} \text{Tr}[\sigma^2 \text{diag}(x^2) \nabla^2 u(x, t)] + r(u(x, t) - \nabla u(x, t)^\top x), \quad x \in \mathbb{R}^d, \quad t \in [0, T],$$

where x^2 is understood to be coordinate-wise, and $\text{diag}(v)$ is a diagonal matrix with $\text{diag}(v)_i = v_i$. Given the terminal condition $u(x, T) = \phi(x) = \|x\|^2$, the explicit solution to this PDE is

$$u(x, t) = \exp((r + \sigma^2)(T - t)) \phi(x).$$

The BSB PDE is related to the following FBSDE

$$\begin{aligned} dX_t &= \sigma \text{diag}(X_t) dB_t, \quad t \in [0, T], \\ dY_t &= r(Y_t - Z_t^\top X_t) dt + \sigma Z_t^\top \text{diag}(X_t) dB_t, \quad t \in [0, T], \end{aligned}$$

where $T = 1$, $\sigma = .4$, $r = .05$, $X_0 = (1, .5, 1., 5, \dots, 1, .5)$, and $Y_T = \phi(X_T)$. The equivalent Stratonovich SDE is given as:

$$\begin{aligned} dX_t &= \frac{\sigma^2}{2} X_t dt + \sigma \operatorname{diag}(X_t) \circ dB_t, \\ dY_t &= \left[r(Y_t - Z_t^\top X_t) - \frac{\sigma^2}{2} (Z_t^\top X_t + \operatorname{Tr}[\operatorname{diag}(X_t^2) \nabla_x^2 u(X_t, t)]) \right] dt + \sigma Z_t^\top \operatorname{diag}(X_t) \circ dB_t, \end{aligned}$$

Fully-Coupled FBSDE. Finally, we consider a FBSDE with *coupled* forward and backward dynamics adapted from Bender & Zhang (BZ) [36]:

$$\begin{aligned} dX_t &= \sigma Y_t dB_t, \quad t \in [0, T], \\ dY_t &= \left[-rY_t + \frac{1}{2} e^{-3r(T-t)} \sigma^2 \left(D \sum_{j=0}^d \sin(X_{j,t}) \right)^3 \right] dt + Z_t^\top dB_t, \quad t \in [0, T), \end{aligned}$$

where $X_{j,t}$ denotes the j -th coordinate of $X_t \in \mathbb{R}^d$. Due to the dependence of the forward process (X_t) on (Y_t), this set of coupled FBSDE does not fit into the mathematical formulation set forth in (3.3) and (3.4). Nevertheless, we can still apply the BSDE methods described at the beginning of Section 6 by initializing $Y_0 = u_\theta(x, 0)$ and jointly integrating (X_t, Y_t) . We set $T = 1$, $r = .1$, $\sigma = .3$, $D = .1$, $X_0 = (\pi/2, \pi/2, \dots, \pi/2)$, and $Y_T(X_T) = \phi(X_T) = D \sum_{j=1}^d \sin(X_{j,T})$. The above FBSDE is induced from the following PDE

$$\partial_t u(x, t) = -\frac{1}{2} \sigma^2 u(x, t)^2 \nabla^2 u(x, t) + ru(x, t) - \frac{1}{2} e^{-3r(T-t)} \sigma^2 \left(D \sum_{j=0}^d \sin(x_j) \right)^2$$

with the analytical solution $u(x, t) = e^{-r(T-t)} D \sum_{j=0}^d \sin(x_j)$. Additionally, the equivalent Stratonovich SDE is given as:

$$\begin{aligned} dX_t &= \frac{\sigma^2}{2} Z_t Y_t + \sigma Y_t \circ dB_t, \\ dY_t &= \left[-rY_t + \frac{1}{2} e^{-3r(T-t)} \sigma^2 \left(D \sum_{j=0}^d \sin(X_{j,t}) \right)^3 - \frac{\sigma^2}{2} (Z_t^2 Y_t + \operatorname{Tr}[Y_t^2 \nabla_x^2 u(X_t, t)]) \right] dt + Z_t^\top \circ dB_t, \end{aligned}$$

Method	100D HJB		100D BSB	
	float32	float64	float32	float64
PINNs	0.1281 ± .0136	0.1281 ± .0171	2.9648 ± .8652	1.5066 ± .2349
FS-PINNs	0.0838 ± .0170	0.0702 ± .0074	0.0602 ± .0150	0.0497 ± .0031
EM-BSDE	0.3776 ± .0365	0.3820 ± .0219	0.2451 ± .0160	0.3735 ± .0470
EM-BSDE (NR)	0.4459 ± .0410	0.4676 ± .0153	0.1648 ± .0143	0.1855 ± .0078
Heun-BSDE	0.0675 ± .0053	0.0529 ± .0029	0.4587 ± .0261	0.0535 ± .0113

Table 3: float32 vs float64 Performance in 100D HJB/BSB

G.2 Sensitivity to Floating Point Precision

In BSDE-based losses, floating point errors can accumulate through out integration of the SDEs, leading to poor performance on the trained model. As seen in Table 3, the floating point error is especially apparent in the Heun loss on the 100D BSB case where the performance of the model is improved by a factor of 10 between float32 and float64. In addition, performance improvements were observed for the PINNs and FS-PINNs models as well. It is also noted that the EM-BSDE models performed slightly worse at a float64, which may be attributed to the bias term present in its loss. Overall, floating point sensitivity is more apparent in the BSB problem than the HJB problem. We attribute this to the non-trivial forward trajectory in the BSB problem. We leave more numerically stable implementations of the Heun solver in float32, such as PDE non-dimensionalization [7] and the use of the reversible Heun solver from [41], to future work.

G.3 Dimensionality Study

Additionally, we demonstrate scalability of the algorithms by re-running the HJB problem at various dimensions and plotting the RL2 error for each method. Figure 6 shows EM-BSDE underperforming both FS-PINNs and Heun-BSDE across all dimensions tested. Additionally, trajectory-based methods scale more effectively to high-dimensional problems with PINNs-based methods showing a slight advantage in lower dimensions.

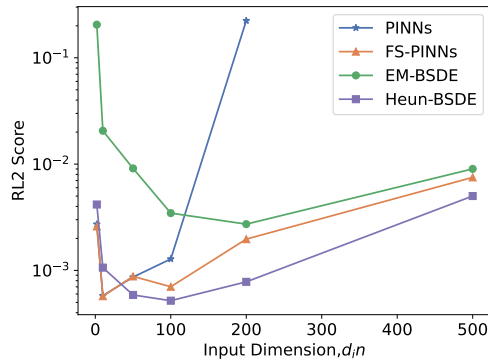


Figure 6: A plot of RL2 performance for the HJB problem at various dimensions $d_{in} = \{2, 10, 50, 100, 200, 500\}$

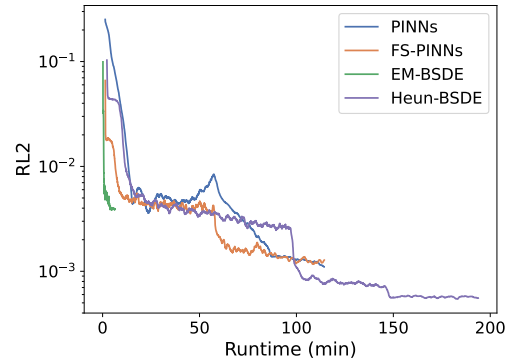


Figure 7: A plot of RL2 performance vs runtime at fixed iteration steps

G.4 Improved Algorithm Schemes

We describe the algorithmic differences between the full roll-out algorithm and the batched sub-sampling variation for the general BSDE loss; a similar algorithm is used FS-PINNs, replacing self-regularization loss with the PINNs loss.

Algorithm 2 Full BSDE Loss Algorithm

Input: Neural network $\hat{u}_\theta(x, t)$, parameters θ , terminal function ϕ , time step Δt , trajectory length N .
Output: Self-consistency loss $\mathcal{L}_{sr}$ and terminal loss $\mathcal{L}_\phi$

- 1: Sample initial state: $(x[0], t[0]) = (x_0, 0)$, with $x_0 \sim \mu$
- 2: Evaluate network at initial state: $(u, u_x) = (\hat{u}_\theta(x[0], t[0]), \nabla_x \hat{u}_\theta(x[0], t[0]))$
- 3: Initialize self-consistency losses: $\ell_{step}[0 : N-1] \leftarrow 0$
- 4: **for** $i = 0, \dots, N-1$ **do**
- 5: Sample Brownian noise: $\xi \sim \mathcal{N}(0, I_d)$
- 6: /* Forward SDE rollout */
- 7: Propagate forward state: $x[i+1] = x[i] + \Delta x$
- 8: **if** no resetting **then**
- 9: /* Use either EM or Heun integration (NR) */
- 10: Propagate backward state: $y[i+1] = y[i] + \Delta y$
- 11: **else**
- 12: /* Use either EM or Heun integration */
- 13: Propagate backward state: $y[i+1] = u + \Delta y$
- 14: **end if**
- 15: Propagate time: $t[i+1] = t[i] + \Delta t$
- 16: Evaluate network at new state: $(u, u_x) = (\hat{u}_\theta(x[i+1], t[i+1]), \nabla_x \hat{u}_\theta(x[i+1], t[i+1]))$
- 17: Record local residual loss: $\ell_{step}[i] = (u - y[i+1])^2$
- 18: **end for**
- 19: Compute self-consistency loss: $\mathcal{L}_{sr} = \sum_{i=0}^{N-1} \ell_{step}[i]$
- 20: Compute terminal loss: $\mathcal{L}_\phi = (u - \phi)^2 + \|u_x - \nabla_x \phi\|^2$
- 21: **return** $(\mathcal{L}_{sr}, \mathcal{L}_\phi)$

As shown in Algorithms 2 and 3, the new variation of the loss separates the forward and backward propagation which enables random sub-sampling in the loss evaluation. At full sampling ($B = N$), the batched algorithm recovers loss and gradient values consistent with the original algorithm with

Algorithm 3 Batched, Sub-sampling BSDE Loss Algorithm (Full description of Algorithm 1)

Input: Neural network $\hat{u}_\theta(x, t)$, parameters θ , terminal function ϕ , time step Δt , trajectory length N , evaluation batch B .
Output: Self-consistency loss $\mathcal{L}_{\text{sr}}$ and terminal loss $\mathcal{L}_\phi$

- 1: Sample initial state: $(x[0], t[0]) = (x_0, 0)$, with $x_0 \sim \mu$
- 2: Sample Brownian noise: $\xi[0 : N - 1] \sim \mathcal{N}(0, I_d)$
- 3: Evaluate network at initial state: $(u, u_x) = (\hat{u}_\theta(x[0], t[0]), \nabla_x \hat{u}_\theta(x[0], t[0]))$
- 4: /* Forward SDE rollout */
- 5: **for** $i = 0, \dots, N - 1$ **do**
- 6: /* Use either EM or Heun integration */
- 7: Propagate forward state: $x[i + 1] = x[i] + \Delta x$
- 8: Propagate time: $t[i + 1] = t[i] + \Delta t$
- 9: **if** coupled **then**
- 10: Evaluate network at new state: $(u, u_x) = (\hat{u}_\theta(x[i + 1], t[i + 1]), \nabla_x \hat{u}_\theta(x[i + 1], t[i + 1]))$
- 11: **end if**
- 12: **end for**
- 13: Stop gradient: $x[0 : N] = \text{SG}(x[0 : N])$
- 14: Separate states: $(x_i, x_{i+1}, t_i, t_{i+1}) = (x[0 : N - 1], x[1 : N], t[0 : N - 1], t[1 : N])$
- 15: /* Same permutations */
- 16: Random sub-sampling: $(x_i, x_{i+1}, t_i, t_{i+1}) = \text{perm}(x_i, x_{i+1}, t_i, t_{i+1})[0 : B - 1]$
- 17: Evaluate network at i points: $(u_i, u_{ix}) = (\hat{u}_\theta(x_i, t_i), \nabla_x \hat{u}_\theta(x_i, t_i))$
- 18: /* Use either EM or Heun Integration */
- 19: Compute backward state at batched point: $y_{i+1} = u_i + \Delta y$
- 20: Evaluate network at $i + 1$ points: $u_{i+1} = \hat{u}_\theta(x_{i+1}, t_{i+1})$
- 21: /* Use PINNs loss instead for FS-PINNs */
- 22: Compute self-consistency loss: $\mathcal{L}_{\text{sr}} = \frac{N}{B} \sum_{i=0}^{B-1} (u_{i+1} - y_{i+1})^2$
- 23: Evaluate network at T : $(u, u_x) = (\hat{u}_\theta(x[N], t[N]), \nabla_x \hat{u}_\theta(x[N], t[N]))$
- 24: Compute terminal loss: $\mathcal{L}_\phi = (u + \phi)^2 + \|u_x - \nabla_x \phi\|^2$
- 25: **return** $(\mathcal{L}_{\text{sr}}, \mathcal{L}_\phi)$

proper scaling. Additionally, the stop gradient on the forward SDE has negligible effects on model performance but can improve convergence as it fixes optimization to only the backward SDE or PDE.

Note that Algorithms 2 and 3 are simplified for readability and excludes some algorithmic details such as trajectory batching and loss weighting.

G.5 Behavior Policy Rollouts for HJB Optimal Control

Suppose our control system is a deterministic control-affine system: $\dot{x} = f(x) + g(x)u$. For a positive definite R , the HJB equation for stagewise cost $c(x) + \frac{1}{2}\|u\|_R^2$ and terminal cost c_T is:

$$\partial_t V + \langle \nabla_x V, f \rangle + c - \frac{1}{2} \|g^\top \nabla V\|_{R^{-1}}^2 = 0, \quad V(x, T) = c_T(x),$$

and the optimal control induced by V is $\pi_V(x, t) := -R^{-1}g^\top(x)\nabla V(x, t)$.

Now, suppose we have any rollout policy $\pi(x, t)$, and we consider Itô stochastic rollouts of the form:

$$dX_t^\pi = [f(X_t^\pi) + g(X_t^\pi)\pi(X_t^\pi, t)]dt + \sigma dB_t.$$

Now, the optimal value function $V^*(x, t)$ satisfies the following SDE:

$$\begin{aligned} dV^*(X_t^\pi, t) &= \left[\partial_t V^*(X_t^\pi, t) + \langle f(X_t^\pi) + g(X_t^\pi)\pi(X_t^\pi, t), \nabla V^*(X_t^\pi, t) \rangle + \frac{\sigma^2}{2} \text{tr}(\nabla^2 V^*(X_t^\pi, t)) \right] dt \\ &\quad + \sigma \langle \nabla V^*(X_t^\pi, t), dB_t \rangle \\ &= \left[\frac{1}{2} \|g^\top \nabla V^*\|_{R^{-1}}^2 - c + \langle g\pi, \nabla V^* \rangle + \frac{\sigma^2}{2} \text{tr}(\nabla^2 V^*) \right] dt + \sigma \langle \nabla V^*, dB_t \rangle, \end{aligned}$$

noting that the last expression is evaluated at (X_t^π, t) . Hence, the forward/backward Itô SDEs for a given value function V and behavior policy π are:

$$dX_t^\pi = [f(X_t^\pi) + g(X_t^\pi)\pi(X_t^\pi, t)]dt + \sigma dB_t, \tag{G.2a}$$

$$dY_t^V = \left[\frac{1}{2} \|g^\top \nabla V\|_{R^{-1}}^2 - c + \langle g\pi, \nabla V \rangle + \frac{\sigma^2}{2} \text{tr}(\nabla^2 V) \right] dt + \sigma \langle \nabla V, dB_t \rangle, \quad (\text{G.2b})$$

noting again that the expressions in (G.2b) are evaluated at (X_t^π, t) .

Similarly, we can define the forward/backward Stratonovich SDEs as:

$$dX_t^\pi = [f(X_t^\pi) + g(X_t^\pi)\pi(X_t^\pi, t)]dt + \sigma dB_t, \quad (\text{G.3a})$$

$$dY_t^V = \left[\frac{1}{2} \|g^\top \nabla V\|_{R^{-1}}^2 - c + \langle g\pi, \nabla V \rangle \right] dt + \sigma \nabla V^\top \circ dB_t. \quad (\text{G.3b})$$

Note that the Itô BSDE (G.2b) requires explicit Hessian computation $\nabla^2 V$ while the Stratonovich BSDE (G.3b) does not. This holds for all first-order PDEs, such as in deterministic HJB problems.

These forward/backward SDEs can be used in conjunction with the induced policy π_V from the current value function V . Some care must be taken though when setting up the BSDE losses. In particular, since both the forward SDE trajectory $(X_t^{\pi_V})_t$ and the $\pi_V(X_t^{\pi_V}, t)$ terms which appear the backward SDEs depend implicitly on V , stop-gradient operators should be placed so that gradients are not back-propagated through these values, which can destabilize training.

G.6 Pendulum Swing Up Experiment

In addition to the results above, we include a simple pendulum swing-up optimal control experiment inspired by [55]. Given the pendulum equations of motion,

$$x = \begin{bmatrix} \theta \\ \dot{\theta} \end{bmatrix}, \quad f(x, u) = \dot{x} = \begin{bmatrix} \dot{\theta} \\ -\frac{1}{ml^2} (b\dot{\theta} - mgl \sin \theta - u) \end{bmatrix},$$

we define a optimal control objective:

$$J^*(x_0) = \min_{u(t)} \int_0^T c(x(t), u(t)) dt + \Phi(x(T)),$$

where:

$$c(x, u) = \Phi(x) + ru^2, \quad \Phi(x) = q_1 \sin^2 \theta + q_1 (\cos \theta - 1)^2 + q_2 \dot{\theta}^2,$$

with $q_1, q_2, r > 0$. Observe that this problem setup exactly fits the setup in Appendix G.5, and hence both the forward/backward Itô and Stratonovich SDEs in (G.2) and (G.3) directly apply.

G.6.1 Pendulum Results

Metric	PINNs	FS-PINNs	EM-BSDE	Heun-BSDE
Cost	53.17	46.59	46.42	46.43
PDE Error	2.77	3.38	78.94	18.6

Table 4: Accumulated cost and average PDE error for the pendulum swing-up problem.

The results of the pendulum swing-up case are outlined in Table 4. We use the specific constants $m = 1, b = 0.1, l = 1, g = 9.8, q_1 = 10, q_2 = 1, r = 1$ in our experiment. It is observed that while the accumulated cost between the three trajectory-based methods remain similar, the lower PDE error on FS-PINNs and Heun-BSDE signify better learned solutions. Furthermore, in Figure 8, we observe that Heun-BSDE generally has the lowest PDE error with high errors only at the discontinuities.

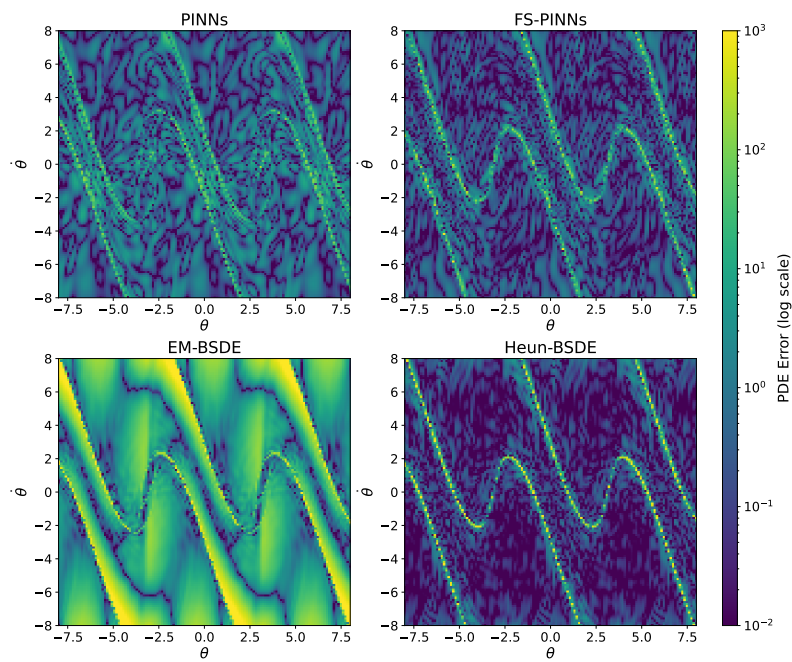


Figure 8: PDE error at $t = 0$ for the pendulum swing up case.