

This is a repository copy of *A market resilient data-driven approach to option pricing*.

White Rose Research Online URL for this paper:

<https://eprints.whiterose.ac.uk/217234/>

Preprint:

Goswami, Anindya and Rana, Nimit (2024) A market resilient data-driven approach to option pricing. [Preprint]

<https://doi.org/10.48550/arXiv.2409.08205>

Reuse

Items deposited in White Rose Research Online are protected by copyright, with all rights reserved unless indicated otherwise. They may be downloaded and/or printed for private study, or other acts as permitted by national copyright laws. The publisher or other rights holders may allow further reproduction and re-use of the full text version. This is indicated by the licence information on the White Rose Research Online record for the item.

Takedown

If you consider content in White Rose Research Online to be in breach of UK law, please notify us by emailing eprints@whiterose.ac.uk including the URL of the record and the reason for the withdrawal request.

A MARKET RESILIENT DATA-DRIVEN APPROACH TO OPTION PRICING

ANINDYA GOSWAMI

Department of Mathematics, IISER Pune, India

NIMIT RANA

Department of Mathematics, University of York, UK

ABSTRACT. In this paper, we present a data-driven ensemble approach for option price prediction whose derivation is based on the no-arbitrage theory of option pricing. Using the theoretical treatment, we derive a common representation space for achieving domain adaptation. The success of an implementation of this idea is shown using some real data. Then we report several experimental results for critically examining the performance of the derived pricing models.

1. INTRODUCTION

Some products or services are often bought to protect the buyer from uncertainties. Warranty is one of them. Figuring out a fair price for a warranty is not an easy task. Since the financial markets consist of risky assets, traders purchase some warranty-type contracts, called options. Several different types of options are traded by numerous traders daily in every modern stock exchange. Determination of an option's fair price is one of the central questions in Mathematical Finance. Following the seminal work [2] by Black and Scholes the mathematical treatment for addressing the option pricing problem quickly flourished. In the last fifty years, thousands of peer-reviewed articles came into existence to address various mathematical, computational, or statistical aspects of fair pricing of a large spectrum of options under diverse scenarios. Consideration of increasingly realistic stochastic models of asset returns and then showing that the model could preclude arbitrage, before deriving a stochastic or PDE representation of an option's price is the state of the art for option pricing research. These works have tremendously enriched our understanding of fair pricing of options, beyond any doubt. Besides, option pricing research also played a vital role in inspiring deeper and more diverse investigation of stochastic calculus, differential equations, numerical simulations etc. Given this, it is natural to have a judgment that "pricing of option is not just about data science". Without disagreeing with that judgment, we propose a novel direction of research that translates some existing data science ideas into the language of stochastic modeling to derive some theoretical algorithms of data-driven option pricing that leverage from both worlds.

A solution to the option pricing problem is called data-driven if the approach assumes no theoretical model of asset dynamics and relies only on the observed data. However, it is possible to

E-mail addresses: anindya@iiserpune.ac.in , nimit.rana@york.ac.uk.

The authorship is distributed equally among all authors. The authors' names appear in the alphabetical order of their surnames.

utilize the general theory of fair pricing of options on an abstract underlying asset while deriving a data-driven option pricing model. In other words, despite the pricing solution being data-driven the derivation of the solution need not always be data-driven. For example, in a recent work [4], some theoretical constraints of fair price of options are utilized to come up with a data-driven model. The investigations on the prospect of various types of data-driven option pricing are not new in the literature. Nevertheless, the volume of work in this direction is much less than its counterpart in the theoretical option pricing. A comprehensive literature survey is beyond the scope of this paper. In the initial works of [13, 10, 15] the promise of data-driven option pricing is reported using the data of option contracts on the S&P 500. Further investigation on feature selection appeared in [3, 14]. Modularity is used in [9, 21] for improving the pricing quality. In the context of option pricing, modularity entails dividing the data set into disjoint modules, according to the moneyness and time-to-maturity parameters of the contracts. Some more recent works [8, 17, 18] build upon the above findings and a theoretical property of options contracts, called homogeneity hint [7].

In this paper, we first revisit the homogeneity hint property by producing theoretical proof in a wide general setting. Then we explain the shortcomings of the option price prediction models based only on the homogeneity hint. Such models can learn from one asset class data and be applied to another, provided their log returns have similar distributions under their risk-neutral measures. However, requiring this identity in two different assets is quite restrictive. For this reason, we formulate, in a parametric setting, a common representation space that bridges two different risk-neutral return distributions. The common representation space is used for connecting option prices on assets having no similarities. This is a fresh idea for the machine learning application of option pricing, where domain adaptation is implemented. We show theoretically, as well as empirically that a data-driven approach based on a common representation space can gain accuracy on test data having significant domain shifts. For the empirical study, we consider the option price data of two indices in the National Stock Exchange (NSE) of India. It is worth noting that options on these indices (NIFTY 50, and NIFTY Bank) have very high trading volume in respect of the global market.

We compare the performance of both approaches, the one that is based on homogeneity hint (HH) and the one that accounts for domain shift (DS) using the price data of options on NIFTY 50, and NIFTY Bank indices. We identify that option price data during the COVID lockdown period in India has a significant domain shift and is hence eligible to assess the performance of models for domain shift. The assessment shows a significant gain in accuracy for models that utilize a common representation space. We also propose an ensemble model using the models based on HH and the one that is to address DS. We put forward a notion of domain shift quotient, that measures the degree of domain shift and is utilized to set weights on constituent models in the ensemble model. We have performed extensive experiments to assess the performance of the ensemble model using empirical as well as synthetic data. In this paper, we also assess the performance of each of these three approaches for multi-source training models, where the model is trained on data coming from symbols NIFTY 50 and NIFTY Bank. We report empirical evidence of the reliability of the multi-source ensemble model. Option price prediction model with multi-source training is capable of giving reliable answers even for an underlying asset, which has no or little historical data. This overcomes the common hurdle faced by any data-driven approach.

It is worth mentioning that we include no macroeconomic or fundamental financial variables in the feature sets, to retain the interpretability of the model. Thus the proposed approaches do not lack the interpretability of the model which is a primary criticism that a data science model commonly receives. We also clarify that the goal of this paper is not merely to showcase the usefulness of an existing or newly developed machine-learning method. This rather lays the foundation of the theoretical solution to domain adaptation, a data science idea, in the context of option pricing. The full potential of a non-asset-specific model with the above-mentioned domain

adaptation may further be explored by extensive experimentation. In this connection, we recall that [19] reports a similar extensive experiment to study some other universal non-asset-specific relations captured by a deep learning model. However, keeping in mind that this is the first research on domain adaptation in the context of option pricing, we focus more on theory building than on large-scale empirical verification across all markets.

In this paper, we introduce a variable called *volatility scalar* that plays a key role in the formulation of common representation space. Theoretically, it represents the root mean square of the volatilities of the underlying during the remaining life of the option, provided the underlying follows diffusion dynamics. This variable is used for some type of scaling at the stage of defining features in the common representation space. We derive its expressions under the constant volatility as well as stochastic volatility scenarios. These expressions allow one to estimate the value of the volatility scalar. However, we retain simplicity in the implementation by assuming constant volatility while estimating the variable for investigating whether the notion of domain adaptation works even with such a crude assumption. It is also worth noting that since we do not calibrate a theoretical model for the option price prediction, the discussion of no-arbitrage calibration like [6] is irrelevant here. We also do not address the question of hedging options. A comprehensive literature survey on the statistical hedging problem can be found in [16].

The rest of the paper is organized as follows. In Section 2, we develop the theoretical framework of approaches based on HH and one for addressing DS. For achieving the domain adaptation, a common representation space is formulated using an approximate identity in Section 3. In this section, we develop an ensemble approach too. The empirical data is described in Section 4, using exploratory data analysis. Subsequently, we describe all proposed supervised learning approaches in Section 5. In Section 6, we report the performance of every model based on proposed approaches. Some concluding remarks are added to Section 7.

2. THEORETICAL FRAMEWORK

2.1. Non-parametric Settings. Let $(\Omega, \mathcal{F}, \mathbb{P}, \{\mathcal{F}_u\}_{u \geq 0})$ be a filtered probability space. In this section we assume that $S^{(t,s)} := \{S^{(t,s)}(u), u \geq t\}$ is an adapted right continuous with left limit (rcll) process that models the price of a risky asset for all future time $u \geq t$ and at present time $t(\geq 0)$ it is equal to a positive value s . Evidently, $S = \{S^{(t,s)} \mid t \geq 0, s > 0\}$ forms a family of adapted processes. With a minor but usual abuse of terminology, we often call S itself as a process. Let $B := \{B(u)\}_{u \geq 0}$ denote the price process of a risk-free asset. We further assume the following.

Model Assumptions:-

- M1. The market model, consisting of risky and risk-free assets with price processes $S^{(0,s)}$ and B respectively is free of arbitrage under admissible strategies for any $s > 0$.
- M2. The family of processes S satisfies

$$S^{(t,s)}(u) = sS^{(t,1)}(u), \quad (2.1)$$

for all $u \geq t \geq 0$ and $s > 0$ almost surely.

- M3. Either $S^{(0,s)}$ is Markov or there is an additional auxiliary process $V := \{V(u)\}_{u \geq 0}$ such that $\{(S^{(0,s)}(u), V(u))\}_{u \geq 0}$ is Markov w.r.t. the filtration $\{\mathcal{F}_u\}_{u \geq 0}$ for all $s > 0$.

It is worth noting that M1 is a standard technical assumption. On the other hand, M2 asserts affine structure, i.e. S is expressed as exponential. This is valid, for example, with Geometric Brownian motion (GBM), regime-switching GBM, Merton's Jump Diffusion model, etc. Additional V in M3 is not needed for GBM, but that is essential for other multi-factor models. For example, V is the Markov regime in a regime-switching model, instantaneous volatility in the Heston model, etc.

Definition 2.1 (Moneyness). *Let K be the strike price of a call option. If the present asset price is s , $p = K/s$ is called the present moneyness of the said option. Note that $p = 1$ represents at-the-money, $p > 1$ represents out-of-the-money and $p < 1$ represents in-the-money.*

Theorem 2.2. *Assume that the market contains two risky assets other than a risk-free asset, whose price processes are modeled by S_1 and S_2 , satisfying (M1)-(M3) with a common auxiliary process V . Let φ_i denote the European call option price function on the i th risky asset. That is, $\varphi_i(t, s_i, v; K_i, T)$ gives the present (time t) price of the call option on S_i having exercise price and maturity K_i and T respectively with $V(t) = v$ and $S_i(t) = s_i$. Also assume that the conditional laws of $S_1^{(t,1)}(T)$ and $S_2^{(t,1)}(T)$ given $V(t) = v$ are identical under their corresponding risk-neutral measures $\mathbb{P}_1$ and $\mathbb{P}_2$, respectively. If the contracts are chosen with an identical time to maturity and equal moneyness at present time t i.e., $\frac{K_1}{s_1} = \frac{K_2}{s_2} = p$ (say), we have*

$$s_1^{-1} \varphi_1(t, s_1, v; p s_1, T) = s_2^{-1} \varphi_2(t, s_2, v; p s_2, T). \quad (2.2)$$

The above theorem is applicable for a wide range of models. For example, when two geometric Brownian motions (GBM) with different drifts but the same diffusion coefficients are taken, their dynamics under respective risk neutral measures are identical. The requirement of common V is also not a real restriction. For example, consider two regime switching GBMs S_1 and S_2 , with Markov regime processes X_1 and X_2 respectively. In other words, for each i , X_i is a finite state continuous-time Markov chain that influences otherwise constant drift and volatility coefficients of a GBM S_i . Then the process $V := (X_1, X_2)$, is such that both (S_1, V) and (S_2, V) are Markov. It is also worth noting that due to the normalization condition in the equality of conditional law, the theorem remains applicable even if one asset's price is more than several orders of magnitude of another. The theorem is proved below.

Proof of Theorem 2.2. It is given that the market contains two risky assets, whose price processes are modeled by S_1 and S_2 , satisfying (M1)-(M3) with an auxiliary process V . For each $i = 1, 2$, let $\mathbb{P}_i$ be an equivalent martingale measure (EMM) corresponding to S_i , i.e., S_i/B is martingale under $\mathbb{P}_i$, a probability measure equivalent to $\mathbb{P}$. Using $\mathbb{P}_i$, we express a present (time t) fair price of a European-style call option on the i th risky asset having the exercise price and maturity time as K_i and T respectively, as

$$\mathbb{E}_i \left(\frac{B(t)}{B(T)} (S_i(T) - K_i)^+ \mid \mathcal{F}_t \right)$$

where $\mathbb{E}_i$ is the expectation w.r.t. $\mathbb{P}_i$. Using (M3), the above is

$$\mathbb{E}_i \left(\frac{B(t)}{B(T)} (S_i(T) - K_i)^+ \mid S_i(t), V(t) \right) = \varphi_i(t, S_i(t), V(t); K_i, T),$$

where φ_i denotes the call option price function on the i th asset with maturity T and exercise price K_i . For each $i = 1, 2$, using (2.1), and the above equation, and by denoting $V(t) = v$, $S_i(t) = s_i$, we get

$$\begin{aligned} \varphi_i(t, s_i, v; K_i, T) &= \mathbb{E}_i \left(\frac{B(t)}{B(T)} (S_i^{(t, s_i)}(T) - K_i)^+ \mid S_i^{(t, s_i)}(t) = s_i, V(t) = v \right) \\ &= \mathbb{E}_i \left(\frac{B(t)}{B(T)} s_i (S_i^{(t, 1)}(T) - \frac{K_i}{s_i})^+ \mid s_i S_i^{(t, 1)}(t) = s_i, V(t) = v \right) \\ &= s_i \mathbb{E}_i \left(\frac{B(t)}{B(T)} (S_i^{(t, 1)}(T) - \frac{K_i}{s_i})^+ \mid S_i^{(t, 1)}(t) = 1, V(t) = v \right) \\ &= s_i \varphi_i \left(t, 1, v; \frac{K_i}{s_i}, T \right). \end{aligned} \quad (2.3)$$

Furthermore, it is given that the conditional laws of $S_1^{(t,1)}(T)$ and $S_2^{(t,1)}(T)$ given identical V values at time t under their corresponding risk-neutral measures $\mathbb{P}_1$, and $\mathbb{P}_2$ respectively are identical. Hence for the contracts with identical time to maturity and equal moneyness at time t , i.e., $\frac{K_1}{s_1} = p = \frac{K_2}{s_2}$, we have

$$\begin{aligned} & \mathbb{E}_1 \left(\frac{B(t)}{B(T)} (S_1^{(t,1)}(T) - p)^+ \mid S_1^{(t,1)}(t) = 1, V(t) = v \right) \\ &= \mathbb{E}_2 \left(\frac{B(t)}{B(T)} (S_2^{(t,1)}(T) - p)^+ \mid S_2^{(t,1)}(t) = 1, V(t) = v \right) \\ & \text{or, } \varphi_1(t, 1, v; p, T) = \varphi_2(t, 1, v; p, T). \end{aligned}$$

Using (2.3), the above equality is rewritten as

$$s_1^{-1} \varphi_1(t, s_1, v; p s_1, T) = s_2^{-1} \varphi_2(t, s_2, v; p s_2, T).$$

Hence, (2.2) holds. $\square$

Remark 2.3 (Homogeneity Hint Approach or $\mathcal{A}_{HH}$). *The above theorem is alluded to as homogeneity hint (see [7] and references therein) as the scale-free terms on both sides of (2.2) are equal. This also reflects the homogeneity of order one of the option prices in stock and strike prices. For predicting option price, the scale-free ratio of option and stock price can be considered as a response/target variable in a learning approach whereas another observed scale-free quantity, the array of order statistics of mean-adjusted historical log returns, can be taken as a predictor/feature variable. By doing so, from the input asset price data, the magnitude, temporal, and drift information are removed and the spatial distribution of return is retained. On the other hand from the predicted target value, which stands for the ratio of option and spot prices, the predicted option price is obtained by multiplying the target with the observed stock price. Such an approach has great flexibility. It can learn from one asset class data and be applied to another, provided their log returns have similar distributions. See [8] for more details. We call such approaches as the Homogeneity Hint Approach or $\mathcal{A}_{HH}$.*

However, requiring identical distributions of log returns in two different assets is quite restrictive, even under their risk-neutral measures. Unless the features and targets are further scaled, to factor out the large disparity of their distributions, the trained model fails to perform across assets. In other words, $\mathcal{A}_{HH}$ adapts poorly to domain shifts. Indeed such an approach may find feature values of the test data from another asset unfamiliar with reference to the training data and hence fail to perform for test dataset. In view of this, we wish to formulate a common representation space, i.e., a bridge between two different risk-neutral return distributions and utilize that for connecting their corresponding option prices. To this end, we adopt parametric models of asset price dynamics.

2.2. Parametric Settings. In this section and onwards we only consider the processes having continuous paths, unlike rcll process in the preceding section.

The simplest parametric model for the continuous-time dynamics of asset prices is geometric Brownian motion (GBM). The GBM which is also a Markov process satisfies (M1)-(M3), with no need of auxiliary process V and hence in particular (2.1) holds. Therefore, under GBM the theoretical call option price function $C(t, s; K, T; r, \sigma)$ satisfies

$$C(t, s; K, T; r, \sigma) = s C \left(t, 1; \frac{K}{s}, T; r, \sigma \right)$$

as a consequence of (2.3). The function C is also known as the Black-Scholes-Merton (BSM) formula for a European call option. Here we replace $(S(t), V(t))$ by only $S(t)$, since S itself is Markov. As a result, the additional variable V is dropped from the option price function. The

rationale behind such discrepancies in notation in other places is also evident from the context, so the discrepancies cause no ambiguity. Next, we consider an extension of BSM model for further discussion. Let $(\Omega, \mathcal{F}, \mathbb{F}, \mathbb{P})$ be a filtered probability space, with a Brownian motion $W = W(t)_{t \geq 0}$ adapted to $\mathbb{F}$. We consider $S = \{S(t)\}_{t \geq 0}$ an Itô Process that models a stock, i.e., S satisfies the following SDE

$$dS(t) = \mu(t)S(t)dt + \sigma(t)S(t)dW(t), \quad (2.4)$$

where $\{\mu(t)\}_{t \geq 0}$ and $\{\sigma(t)\}_{t \geq 0}$ are adapted and bounded processes.

Definition 2.4 (ρ -scaling of a process). *Given a positive valued process $S = \{S(t)\}_{t \geq 0}$, define $A := \{A(t)\}_{t \geq 0}$ such that $A(t)^\rho = S(t)$ for all $t \geq 0$ and for some $\rho > 0$. We call the process A as ρ -scaling of process S .*

The ρ -scaling is an easy recipe for obtaining a parametric family of processes, all having different risk-neutral return distributions. A more precise version of this statement is given in the following lemma.

Lemma 2.5. *Assume that S satisfies (2.4) and $\rho > 0$. (i) Both S and its ρ -scaling satisfy (M1)-(M3). (ii) The instantaneous volatility of the ρ -scaling of S is identical to the instantaneous volatility of S divided by ρ .*

Proof. (i) The strong solution of SDE (2.4) is

$$S(t) = S(0)e^{\int_0^t \left(\mu(u) - \frac{\sigma(u)^2}{2} \right) du + \int_0^t \sigma(u) dW(u)},$$

for an arbitrary initial value $S(0) > 0$. Thus S satisfies (2.1) and (M3) holds with no additional auxiliary process V , as S itself is Markov. Consequently, ρ -scaling of S satisfies

$$A(t) = S(t)^{\frac{1}{\rho}} = A(0)e^{\left[\rho^{-1} \int_0^t \left(\mu(u) - \frac{\sigma(u)^2}{2} \right) du + \int_0^t \frac{\sigma(u)}{\rho} dW(u) \right]} \quad (2.5)$$

as $S(0)^{\frac{1}{\rho}} = A(0)$. Therefore, A also satisfies (2.1) and (M3). Similarly, it is not hard to see that both S and A fulfill (M1), as an application of Girsanov's transformation and the fundamental theorem of fair pricing (See [11, Example 2 in Subsection 7.4]).

(ii) Using Itô's lemma, from (2.5), we get

$$\begin{aligned} dA(t) &= \left(\frac{\mu(t)}{\rho} - \frac{\rho-1}{2\rho^2} \sigma(t)^2 \right) A(t)dt + \frac{\sigma(t)}{\rho} A(t)dW(t) \\ \text{i.e., } \frac{1}{A(t)} dA(t) &= \left(\frac{\mu(t)}{\rho} - \frac{\rho-1}{2\rho^2} \sigma(t)^2 \right) dt + \frac{\sigma(t)}{\rho} dW(t). \end{aligned} \quad (2.6)$$

The left side is the instantaneous simple return of A . Thus the instantaneous volatility of A , i.e., the coefficient in $dW(t)$ term, is ρ^{-1} times the volatility of S . $\square$

Lemma 2.5 implies that although S and A satisfy (M1)-(M3), Theorem 2.2 is not applicable between them. In other words, an identity like (2.2) between the option prices on S and A does not hold, as their returns' distributions under the respective risk-neutral measures are not identical. However, given a pair of assets whose prices are modeled by (2.4) with different parameter values, one can choose a pair of ρ values so that the assets' ρ -scalings satisfy an identity like (2.2). This fact is stated in the following theorem.

Theorem 2.6. Assume that for each $i = 1, 2$, S_i solves (2.4) with $\mu = \mu_i$, $\sigma = \sigma_i$ and Brownian motion $W = W_i$. For a fixed $t \in [0, T)$, denote

$$\rho_i := \left(\frac{1}{T-t} \mathbb{E} \int_t^T \sigma_i^2(u) du \right)^{\frac{1}{2}}, \quad \forall i = 1, 2, \quad (2.7)$$

where $\mathbb{E}$ is the expectation with respect to $\mathbb{P}$, and denote the price function of a European call option on the ρ_i -scaling of i th asset, by ψ_i . Then

- (i) the risk-neutral distribution of log return of the ρ_i -scaling of i th assets are identical,
- (ii) for all $a_1 > 0$, $a_2 > 0$, $p > 0$

$$\frac{1}{a_1} \psi_1(t, a_1, pa_1, T) = \frac{1}{a_2} \psi_2(t, a_2, pa_2, T). \quad (2.8)$$

Proof of Theorem 2.6. It is given that S_i solves for all $t > 0$

$$dS_i(t) = \mu_i(t)S_i(t)dt + \sigma_i(t)S_i(t)dW_i(t),$$

where the Brownian motions W_1 and W_2 may or may not be independent. Now we fix a $t > 0$ and for each $i = 1, 2$, set ρ_i arbitrarily and $A_i(u) = S_i(u)^{\frac{1}{\rho_i}}$ for all $u \geq t$ and also denote the European call option price function on the asset having price process A_i as ψ_i . In view of (2.6), A_i solves the following SDE

$$dA_i(u) = A_i(u) \left(\frac{dB(u)}{B(u)} + \frac{\sigma_i(u)}{\rho_i} dW'_i(u) \right)$$

where, $W'_i(\cdot) := W_i(\cdot) + \int_0^\cdot \left(\frac{\mu_i(u)}{\sigma_i(u)} - \frac{\rho_i - 1}{2\rho_i} \sigma_i(u) \right) du - \int_0^\cdot \frac{\rho_i}{\sigma_i(u)} \frac{dB(u)}{B(u)}$. Hence

$$\frac{A_i}{B}(\cdot) = \frac{A_i}{B}(t) \mathcal{E} \left(\int_t^\cdot \frac{\sigma_i(u)}{\rho_i} dW'_i(u) \right) \quad (2.9)$$

where $\mathcal{E}(X)$ denotes the Doléans-Dade exponential of the semi-martingale X . For each $i = 1, 2$, let $\mathbb{P}'_i$ denote a probability measure equivalent to $\mathbb{P}$ such that W'_i is a Brownian motion under $\mathbb{P}'_i$. This exists as a consequence of Girsanov's Theorem. Then A_i/B is a martingale under $\mathbb{P}'_i$, and hence $\mathbb{P}'_i$ is an equivalent martingale measure (EMM) corresponding to A_i .

Next, we choose the constants ρ_1 , and ρ_2 so that under the respective risk-neutral measures $\mathbb{P}'_1$, and $\mathbb{P}'_2$, the conditional laws of $A_1^{(t,1)}(T)$ and $A_2^{(t,1)}(T)$ are identical. That is, using (2.9)

$$\int_t^T \frac{\sigma_1(u)}{\rho_1} dW'_1(u) \stackrel{d}{=} \int_t^T \frac{\sigma_2(u)}{\rho_2} dW'_2(u) \quad (2.10)$$

where the equality " $\stackrel{d}{=}$ " is in the sense of distribution. Since the conditional distributions of $A_i(T)/A_i(t)$ under the corresponding risk neutral measures are identical, from Lemma 2.5 and Theorem 2.2 we get (2.8). Moreover, as σ_i is bounded, $\int_t^T \frac{\sigma_i(u)}{\rho_i} dW'_i(u)$ becomes Gaussian random variable with mean zero and variance $\rho_i^{-2} \mathbb{E} \int_t^T \sigma_i^2(u) du$ for each i . Hence the values of the parameters ρ_1 and ρ_2 specified by (2.7) are sufficient to ensure (2.10). Hence the proofs of (i) and (ii) are complete. $\square$

Although ρ in (2.7) depends on the law of σ , and values of t , and T , we don't write it here to avoid clumsy notation. However, we refer to the value in (2.7) as *volatility scalar* to avoid ambiguity. For the rest of the paper, ρ and ρ_i denote the volatility scalars of S and S_i respectively, unless otherwise mentioned.

Remark 2.7 (Estimation of the volatility scalar ρ). From (2.7), it is also evident that, if σ_1 and σ_2 are stationary, ρ_1, ρ_2 can be estimated from the past data. More precisely, ρ , the average of forward variance, see (2.7), is typically estimated by taking the past 20 days' data and we assume that the past 20 days' historical volatility is a good estimator of future volatility, at least in average. It is worth mentioning that ρ must be re-estimated as t and T change which is a sensible thing to do as in practice one needs to calibrate the model at every epoch (i.e., as t changes and then time to maturity $T - t$ also changes).

The selection of ρ , so that (2.7) holds becomes particularly simple if the volatility factor σ is constant in time. Indeed, the volatility scalar ρ becomes identical to σ . Even if the volatility process is neither constant nor stationary but follows a parametric diffusion equation, the volatility scalar may be estimated by estimating the parameters. We illustrate this by considering the Heston model. In that, the square of volatility follows a Cox–Ingersoll–Ross (CIR) process, i.e., for all $t \geq 0$

$$d\sigma^2(t) = \kappa(\theta - \sigma^2(t))dt + \xi\sigma(t)d\tilde{W}(t),$$

with $\sigma^2(0) > 0$ where $\kappa, \theta, \xi, \sigma(0)$ are positive parameters and $\tilde{W}$ is a standard Brownian motion. A direct calculation shows that

$$\frac{1}{T-t}\mathbb{E} \int_t^T \sigma^2(u)du = \theta - (\theta - \sigma^2(t)) \frac{1 - e^{-\kappa(T-t)}}{\kappa(T-t)}.$$

Therefore, $\rho = \sqrt{\left(\theta - (\theta - \sigma^2(t)) \frac{1 - e^{-\kappa(T-t)}}{\kappa(T-t)}\right)}$, a number between the long-run mean of volatility $\sqrt{\theta}$ and the current volatility $\sigma(t)$, can be estimated using the estimator of $\theta, \sigma(t)$ and the speed parameter κ of mean reversion.

On the other hand, for a special case of non-parametric settings, consider two stocks whose prices satisfy (2.4) with identical volatility processes, i.e., $\sigma_1 \stackrel{d}{=} \sigma_2$. Then (2.10) holds with any equal values of ρ_1 and ρ_2 .

In the aim of building a market-resilient learning approach for option price prediction, we follow the approach of common representation. First of all, we are required to select features that can help to extrapolate for successfully predicting option prices corresponding to atypical out-sample data. In this connection, Theorem 2.6 (i) is useful. In particular, it implies that the risk-neutral distribution of log return of the ρ -scaling of the asset is a market-resilient feature influencing the option pricing mechanism. Needless to say, the learning approach should also have features coming from contract parameters and risk-free assets. This feature representation from source and target domains is indistinguishable. Next, we are required to decide on a suitable target that depends solely on features and not on factors, not included in the feature for formulating the common representation space. There is no unique or obvious way to do that. For example, a relation between φ_i and ψ_i , the observable and unobserved option prices of the assets with prices S_i and A_i for each i , may be derived. Then a relation between φ_1 and φ_2 , the observable option prices of the assets with prices S_1 and S_2 can be derived, as Theorem 2.6 relates ψ_1 and ψ_2 . It seems that this is hard in general. If for each $i = 1, 2$,

$$\psi_i(t, a_i, pa_i, T) = F(\varphi_i(t, a_i^{\rho_i}, pa_i^{\rho_i}, T), a_i, \rho_i, p, T) \quad (2.11)$$

for some known F , then from Theorem 2.6 we have

$$\frac{1}{s_1^{1/\rho_1}} F(\varphi_1(t, s_1, ps_1, T), s_1^{1/\rho_1}, \rho_1, p, T) = \frac{1}{s_2^{1/\rho_2}} F(\varphi_2(t, s_2, ps_2, T), s_2^{1/\rho_2}, \rho_2, p, T) = \mathcal{U} \text{ (say)} \quad (2.12)$$

provided ρ_1 and ρ_2 are chosen as in (2.7). The target $\mathcal{U}$ is invariant for two different contracts with identical moneyness and time to maturity on two different assets, possibly with domain shifts. Also, $\mathcal{U} = \psi_i(t, a_i, pa_i, T)/a_i$, which depends solely on the A_i features and the features related to the

contract and risk-free asset. One may perhaps obtain some other targets based on a transformation involving implied volatility etc so that similar common representation is created. Assume that a learning approach $\mathcal{A}$ is trained to take centered return distribution of S_i^{1/ρ_i} (i.e., the returns of A_i after being subtracted by its mean) and variables and parameters p, T, r as inputs and to give $\mathcal{U}$ as target. Then $\mathcal{A}$ learns a fundamental property of option pricing from the common representation space, that holds across assets having a wide range of return distributions. Further, if the map $\varphi \mapsto \mathcal{U}$ is invertible, then for test data, the predicted target may be used for computing the option price. We explain this idea by making use of an approximation formula [1] in the next section.

3. APPLICATION OF AN IMPLIED VOLATILITY ESTIMATE IN DOMAIN ADAPTATION

3.1. A Common Representation Space. In this section, instead of an equality (2.11), which is hard to find, we use an approximate equality relation to derive an approximation of the equality (2.12). To this end, we refer to [1], which gives a simple approximation of implied volatility, the inversion of the BSM formula for near ATM options. The literature on the inversion of the BSM formula is reasonably rich. We cite [12] for various other improvements over the simple formula in [1] which we use in this section. For each asset $i = 1, 2$, the approximation formula for IV^{S_i, p_i, T_i} , the implied volatility of asset S_i corresponding to a call contract with moneyness and time to maturity p_i , and T_i respectively, is given by

$$IV^{S_i, p_i, T_i} \approx \sqrt{\frac{2\pi}{T_i}} \left(\frac{\varphi_i(0, s_i, p_i s_i, T_i)}{s_i(1 + p_i^*)/2} - \frac{(1 - p_i^*)}{(1 + p_i^*)} \right),$$

where $p_i \approx 1$, $p_i^* = p_i e^{-rT_i}$, s_i is the present price of i th stock. We recall that φ_1 and φ_2 are price functions of European call contracts on assets having prices S_1 and S_2 respectively which follow diffusion equations as in (2.4). Since, ρ_i is defined (see Definition 2.7) using the instantaneous volatility values on the time interval between present and maturity time, the implied volatilities should nearly be proportional to ρ_i , provided all other parameters are kept constants. That is if $p_1 \approx p_2$, and $T_1 \approx T_2$, we have

$$\frac{IV^{S_1, p_1, T_1}}{\rho_1} \approx \frac{IV^{S_2, p_2, T_2}}{\rho_2}.$$

However, the stylized facts and other theoretical validations of volatility smile indicate that the said approximation should worsen as T_1 , and p_1 depart from T_2 , and p_2 . Finally, by plugging in the approximate expressions in the above approximate equality we get

$$\sqrt{\frac{2\pi}{T_1}} \left(\frac{\varphi_1(0, s_1, p_1 s_1, T_1)}{s_1(1 + p_1^*)/2} - \frac{(1 - p_1^*)}{(1 + p_1^*)} \right) / \rho_1 \approx \sqrt{\frac{2\pi}{T_2}} \left(\frac{\varphi_2(0, s_2, p_2 s_2, T_2)}{s_2(1 + p_2^*)/2} - \frac{(1 - p_2^*)}{(1 + p_2^*)} \right) / \rho_2. \quad (3.1)$$

We denote the left and right sides by $\mathcal{U}(s_i, r, \sigma_i, p_i, T_i, \rho_i)$ with $i = 1, 2$ respectively. Also, we can approximate $\mathcal{U}(s_i, r, \sigma_i, p_i, T_i, \rho_i) = \frac{IV^{S_i, p_i, T_i}}{\rho_i} \approx IV^{A_i, p_i, T_i}$, the implied volatility corresponding to the asset A_i . The last approximate equality is due to the identical values of moneyness and time to maturity and the fact that the volatility scalar value corresponding to A is one. Hence, $\mathcal{U}(s_i, r, \sigma_i, p_i, T_i, \rho_i)$ may be regarded as primarily dependent on asset A_i features and the features related to contract and risk-free asset.

3.2. Discussion on the approximation error. Equation (3.1) expresses an approximate relation between these two option prices and conveys that $\mathcal{U}$ is nearly constant wrt s, σ , and the volatility scalar ρ . In particular, if prices of a pair of stocks satisfy (2.4) with $\sigma_1 \stackrel{d}{=} \sigma_2$, then (2.10) holds with $\rho_1 = \rho_2 = \rho$, for any arbitrary value. With $\rho_1 = \rho_2 = 1$, clearly $S_i = A_i$ and $\psi_i = \varphi_i$ for each

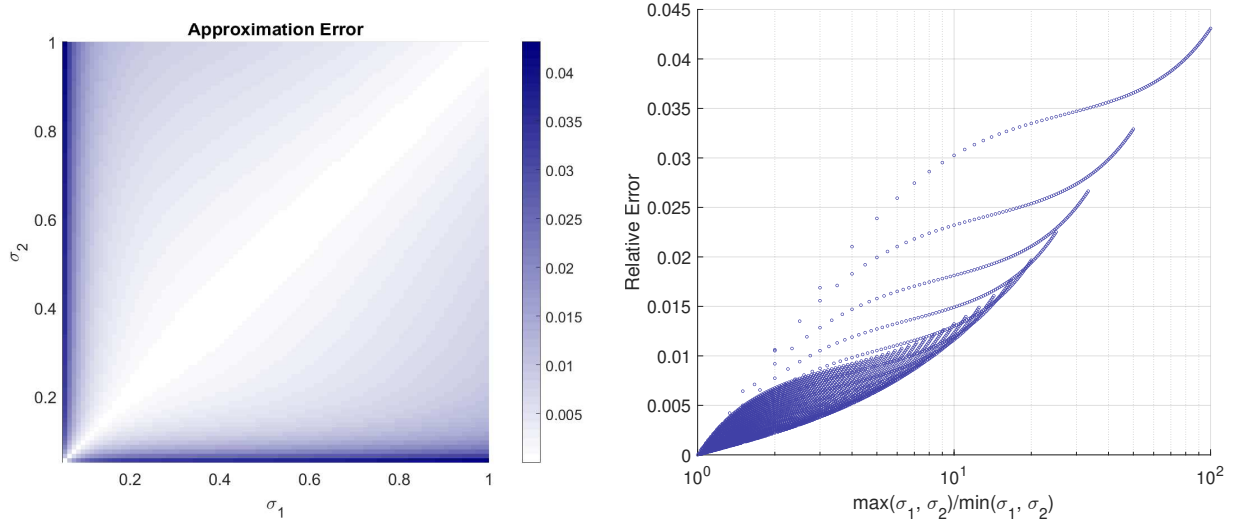
$i = 1, 2$. Consequently (2.8) gives $\frac{\varphi_1(0, s_1, p_1 s_1, T)}{s_1} = \frac{\varphi_2(0, s_2, p_2 s_2, T)}{s_2}$. On the other hand if $T_1 = T_2 = T$, $p_1 = p_2 = p$, Equation (3.1) gives

$$\left(\frac{\varphi_1(0, s_1, p s_1, T)}{s_1} - \frac{(1 - p^*)}{2} \right) / \rho \approx \left(\frac{\varphi_2(0, s_2, p s_2, T)}{s_2} - \frac{(1 - p^*)}{2} \right) / \rho$$

$$\text{i.e., } \frac{\varphi_1(0, s_1, p s_1, T)}{s_1} \approx \frac{\varphi_2(0, s_2, p s_2, T)}{s_2}.$$

Thus the relation in (3.1) is exact, i.e., ' $\approx$ ' can be replaced by ' $=$ ', when the volatility scalars are identical, i.e., $\rho_1 = \rho_2$ and the contract moneyness and time to maturities match.

The approximation error in (3.1) associated to a pair similar contracts (i.e., $T_1 = T_2 = T$, $p_1 = p_2 = p$) on assets having volatility processes σ_1 and σ_2 is measured by the relative error $\sup\{\frac{|\mathcal{U}_1 - \mathcal{U}_2|}{\max(\mathcal{U}_1, \mathcal{U}_2)} \mid s_1 > 0, s_2 > 0\}$ where, $\mathcal{U}_i = \mathcal{U}(s_i, r, \sigma_i, p, T, \rho_i)$. We compute this error for the ATM options, i.e., moneyness parameter $p = 1$ under the special case of BSM models with constant volatility coefficients. The time to maturity T is taken as 0.2 year. In Figure 1a, the horizontal and vertical axes represent the volatility coefficients values of two different assets. Each volatility value ranges on $[5\%, 100\%]$. The darkness of shades show the relative error values for each ordered pair (σ_1, σ_2) . The maximum error is less than 4.5%. If none of σ_1 and σ_2 is smaller than 9%, the error is less than 2%. This numerically validates that (3.1) is a reasonable approximation for practical purposes where the volatilities are not too small.



(A) Grid heatmap of the approximation error for a pair of volatilities (B) The approximation error vs disparity(max-min ratio) of volatilities

FIGURE 1. Visualization of relative error in the approximation relation (3.1)

It is evident from Figure 1a that the error is minimal when σ_1 and σ_2 are identical. This can be explained by appealing to the homogeneity hint, stated under the non-parametric setting in Theorem 2.2. A separate visualization of error's dependence on the degree of mismatch of the volatilities is given in the scatter plot Figure 1b under the same setting.

In this $\frac{\max(\sigma_1, \sigma_2)}{\min(\sigma_1, \sigma_2)}$, a measure of mismatch, is plotted on the horizontal axis and the corresponding relative error of the approximation (3.1) is plotted on the vertical axis. Every blue circle corresponds to a pair of volatilities (σ_1, σ_2) for which the relative error is computed. The horizontal and the vertical coordinates of each circle are $\frac{\max(\sigma_1, \sigma_2)}{\min(\sigma_1, \sigma_2)}$ and $\sup\{\frac{|\mathcal{U}_1 - \mathcal{U}_2|}{\max(\mathcal{U}_1, \mathcal{U}_2)} \mid s_1 > 0, s_2 > 0\}$ respectively,

where σ_i is the constant volatility of i th asset. It shows that if the largest volatility is not larger than double the smallest, the relative error is less than 1.5%. On the other hand, the error is less than 3% if the largest is not larger than 9 times the smallest.

3.3. Approach $\mathcal{A}_{DS}$ for Domain Shift. Assume that a learning approach takes risk-neutral return distribution of ρ_1 -scaling A_1 of asset prices S_1 as a feature variable. To be more precise, the order statistics of mean subtracted log returns of A_1 for the past few time periods are included in the feature set. The approach should also have features from contract parameters p and T , and risk-free assets' interest rate r . Also, set $\mathcal{U}(s_1, r, \sigma_1, p_1, T_1, \rho_1)$ as the target, which can be obtained from the observed option value φ_1 in the training data set. Consequently, for a test data set on another asset S_2 , after feeding in the feature values to the trained model, a value of the target $\mathcal{U}(s_2, r, \sigma_2, p_2, T_2, \rho_2)$ is predicted. Then the option price can be revived using the following formula

$$\varphi_2(t, s_2, p_2, T_2) = s_2 \left(\rho_2 \frac{1 + p_2^*}{2} \sqrt{\frac{T_2}{2\pi}} \mathcal{U}(s_2, r, \sigma_2, p_2, T_2, \rho_2) + \frac{1 - p_2^*}{2} \right), \quad (3.2)$$

irrespective of its volatility as long as ρ_2 is the volatility scalar for S_2 . Therefore, the above-mentioned learning approach is better adapted for the domain shift. We denote this as $\mathcal{A}_{DS}$. Of course, it is possible to develop a different learning model following this idea by choosing an approximation relation more accurate than (3.1) for possibly better performance. So, as far $\mathcal{A}_{DS}$ is concerned, we can imagine that although this simple approximation facilitates predicting option prices reasonably well for test data sets with domain shifts, would fail to outperform $\mathcal{A}_{HH}$ when domain shift is absent. This raises the need of an ensemble model.

3.4. Ensemble Model. The difference between the present and the past mean volatilities, normalized by the second one, may be considered as a crude measure of the domain shift. We call this measure the domain shift quotient or simply DSQ. To be more precise, if σ_0 and σ_i are the average historical volatility of the full training data set and the current historical volatility associated with i th test data respectively, the formula for the DSQ for the specific test data is $\frac{|\sigma_i - \sigma_0|}{\sigma_0}$. In this subsection, we propose an ensemble model where the ratio of weights given to the predictions from $\mathcal{A}_{DS}$ and $\mathcal{A}_{HH}$ models is $\lambda_1 \left(\frac{|\sigma_i - \sigma_0|}{\sigma_0} \right)^{\lambda_2}$, a power function of DSQ. We denote the resulting approach as $\mathcal{A}_{E(\lambda)}$ where $\lambda = (\lambda_1, \lambda_2)$. If $P_{DS}(i)$ and $P_{HH}(i)$ are the predicted prices of i -th option from the models with approaches $\mathcal{A}_{DS}$ and $\mathcal{A}_{HH}$ respectively trained with same data, then $P_E(i)$, the predicted price from an ensemble approach $\mathcal{A}_{E(\lambda)}$, is defined as

$$P_{E(\lambda)}(i) = \frac{1}{1 + \lambda_1 \left(\frac{|\sigma_i - \sigma_0|}{\sigma_0} \right)^{\lambda_2}} P_{HH}(i) + \frac{\lambda_1 \left(\frac{|\sigma_i - \sigma_0|}{\sigma_0} \right)^{\lambda_2}}{1 + \lambda_1 \left(\frac{|\sigma_i - \sigma_0|}{\sigma_0} \right)^{\lambda_2}} P_{DS}(i). \quad (3.3)$$

As a result, for a fixed λ , the weight to P_{HH} diminishes as the atypicality rises in the test data. The more atypical the return data, i.e., the larger the relative difference between the test and the training volatilities, the higher weight is assigned to the prediction P_{DS} in the ensemble approach $\mathcal{A}_{E(\lambda)}$. The parameter λ is estimated using the least square method from a sample of test data. The ensemble approach using the estimated λ is simply denoted by $\mathcal{A}_E$. The performance of the ensemble model may be evaluated for in-sample and out-of-sample data.

4. DATA

According to the Futures Industry Association (FIA), 84% of global equity options are now traded on Indian exchanges, up from just 15% a decade ago.¹ In the first quarter of 2024 alone, over 34.8 million equity index options were traded on India’s two main financial derivatives exchanges—the National Stock Exchange of India (NSE) and the Bombay Stock Exchange (BSE) India. This was double the volume from Q1 2023 and accounted for 73% of all futures and options traded globally. Despite this high trading volume, the market’s value remains relatively small compared to North American and European markets due to lower notional values and premiums per contract. For example, in March 2024, NSE’s total options premiums were \$150 billion, about a quarter of the \$598 billion seen across all US options exchanges. Nonetheless, the Indian market’s rapid growth and substantial trading volume makes it a crucial area for developing and applying data-driven option trading strategies.

For our study, we analyze daily contract price data for European call options on NIFTY50² and BANKNIFTY³. The abundance of data allows a learning model to better capture and depict market behavior and trading patterns. On the other hand, due to the presence of many traders’ participation, the traded prices are free from market inefficiency. The dataset, spanning January 1, 2015, to April 30, 2020, is obtained from the NSE’s historical archives⁴. The data-driven supervised learning approaches for pricing European-style call options, outlined in [8] reportedly performed worse in abrupt highly volatile market conditions. Our objective is to introduce a novel market-resilient approach for improving the prediction for out-sample data having atypical asset returns. To this end, we incorporate testing of the trained model with data from the COVID-19 lockdown period, when the financial market crashed, i.e., from January 2020 to April 2020.

Each data point includes the following information, (i) the daily closing prices of the underlying asset and a particular European call option, (ii) moneyness and the time to maturity of that call contract, (iii) daily 19 log returns of the underlying asset from the 20 most recent daily data, (iv) daily close price of the same call contract on the current and earlier day and v) prevailing risk-free interest rate. Moreover, we filter out all data points for which the same option contract did not exist the previous day. Our choice of a 20-day window aligns approximately with a calendar month, accounting for holidays. Since we compute the daily 19 log returns of the underlying asset, to have these log returns from January 1, 2015, to April 30, 2020, we consider the daily data of the underlying assets (specifically, NIFTY50 and BANKNIFTY indices) from October 1, 2014, to April 30, 2020.

We retain datapoints only for option contracts that are near at-the-money (ATM), specifically those where the ratio of the strike price to the spot price is between 96% and 104%. We refer to these as near-ATM option contracts. It is observed that numerous near-ATM contracts are traded daily, with varying times to maturity. Since, contracts with large or minimal times to maturity have significantly low trading volumes, we focus on studying contracts with time-to-maturity values between 3 and 45 days.

Train/Test Split: To develop a predictive model using XGBoost supervised learning, we partitioned the dataset into two segments, allocating roughly 80% for training and 20% for evaluating

¹Article by Bennett Voyles <https://www.fia.org/marketvoice/articles/premium-turnover-indian-options-hits-150-billion>

²NIFTY 50, abbreviated as NIFTY50, serves as a benchmark index for the Indian stock market. It reflects the weighted average performance of 50 of the largest companies listed on the National Stock Exchange of India.

³The NIFTY Bank index, referred to as BANKNIFTY, consists of the most liquid and highly capitalized Indian banking stocks.

⁴Link to access the data -https://www1.nseindia.com/products/content/derivatives/equities/historical_fo.htm

the trained models. The split is such that the training set includes data dated before the oldest data of the test set. This is for avoiding information leakage as explained in [20]. Specifically, the training dataset spans 56 months, covering the period from January 01, 2015, to August 31, 2019.

Subsequently, the test dataset is subdivided into two parts to facilitate comprehensive evaluation:

- (1) **Typical:** The “typical” test dataset encompasses data from September 01, 2019, to December 31, 2019. This period is characterized by typical market dynamics, enabling model evaluation under market behaviour similar to that observed during the training phase.
- (2) **Atypical:** The “atypical” test dataset consists of data from January 01, 2020, to April 30, 2020. This period includes the COVID-19 Indian financial market crash, presenting a unique testing scenario to assess model robustness under an extremely volatile market. The dynamics during this period are significantly different from those observed during the training phase, thereby offering valuable insights into the model’s performance in unprecedented market scenarios.

Table 1 provides a breakdown of the number of data points at each stage of the model building and evaluation process, offering insight into the dataset’s composition and distribution across training and test phases.

TABLE 1. Train/Test split of European style call option contract dataset sizes

Dataset	Time Period	N50	BNF
<i>Raw</i>	01/01/2015 - 30/04/2020	1611549	576302
<i>Filtered</i>	01/01/2015 - 30/04/2020	27372	37284
<i>Train</i>	01/01/2015 - 31/08/2019	19182	26348
<i>Typical Test</i>	01/09/2019 - 31/12/2019	3692	3709
<i>Atypical Test</i>	01/01/2020 - 30/04/2020	3047	2774

It is important to highlight that the sum of data points in the training, typical test, and atypical test datasets do not match the number of data points in the filtered dataset. This discrepancy occurs because, during the feature generation step in our approach, we remove the rows of data associated with option contracts that do not have a previous day option closing price, which is a feature in our learning approach.

A brief exploratory data analysis is presented below. The Q-Q plot in Figure 2 shows that the log returns of daily closing values of NIFTY50 and BANKNIFTY indices do not possess identical distribution. Nevertheless they are not significantly different during the period that constitutes the training dataset. The Q-Q plot is linear and has a slightly different slope than the identity line. Next, we compare the statistical behavior of log-returns of daily closing values of NIFTY50 index during different time intervals. In Figure 3a the said comparison between the data from training window and the typical test window is presented. On the other hand, Figure 3b presents a comparison between the data from training and the atypical test windows. In 3a the Q-Q plot is aligned with the identity line but in 3b the slope of the Q-Q plot is visibly far from the identity line. Similar comparisons are reported in Figure 4 where we compare the distributions of log-returns of daily closing values of BANKNIFTY index during different time intervals. In 4a, except few outliers, the return distribution of the typical test dataset matches well with that of the training dataset. The points in the Q-Q plot approximately lie on the identity line. However, when the atypical test data is compared with the training data in Figure 4b, the distributions of log-returns

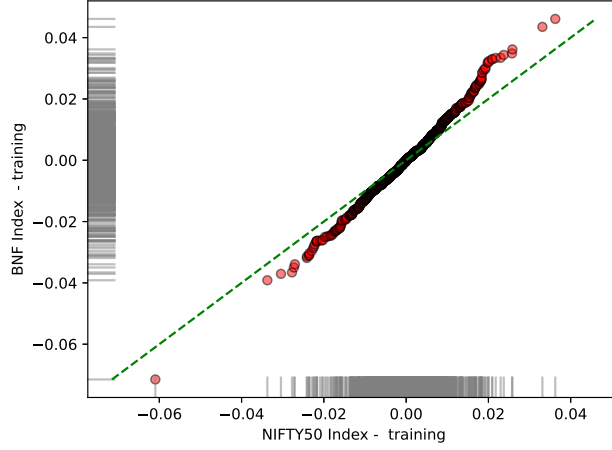


FIGURE 2. Q-Q plot for the log return distributions of the “Close” prices of the BANKNIFTY and NIFTY50 indices for training dataset from January 01, 2015 to August 31, 2019

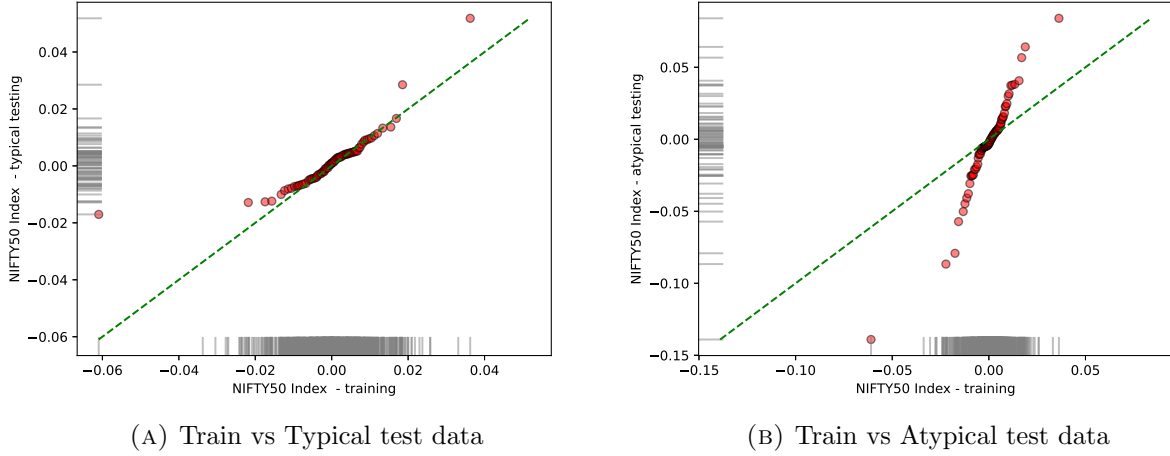


FIGURE 3. Q-Q plot for the log return distributions of the “Close” prices of NIFTY50 in training and testing data

of the closing values appear to differ drastically. Although the Q-Q plot is linear in this case, the slope vastly mismatches with that of the identity line. These observations justify the use of terminology “typical” and “atypical” for two non-overlapping test datasets.

5. DESCRIPTION OF SUPERVISED LEARNING APPROACH

In this section, we outline the steps involved in the supervised learning approach utilized in this study.

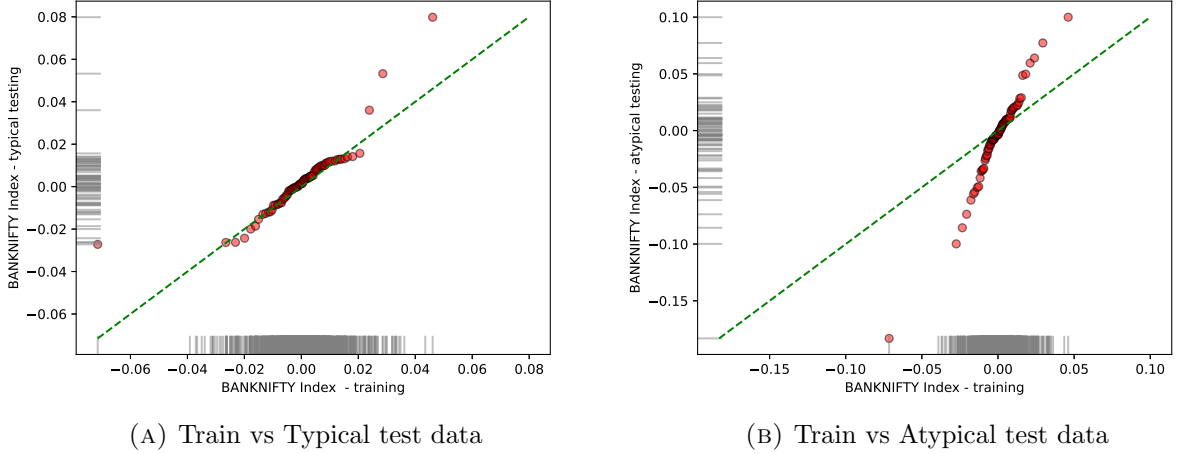


FIGURE 4. Q-Q plot for the log return distributions of the “Close” prices of BANKNIFTY in training and testing data

Step 1: Data Cleaning and Filtering In this essential step, we prepare the gathered raw data for European call option contracts to apply the ML algorithms. Initially, we eliminate rows containing blank ‘-’ entries in the ‘Underlying Value’ or ‘Strike Price’ columns. Additionally, any rows containing values of ‘0’ in the ‘Open’ or ‘Close’ columns are removed.

To further refine the selection, we follow the modularity approach of [8] and focus on option contracts that closely resemble at-the-money (ATM) contracts. This involves filtering contracts based on a criterion where the absolute quantity $|1 - \frac{S}{K}|$ is less than or equal to the pre-decided value of 0.04 as considered in [8], where S and K respectively represent the spot price of underlying and strike price of the option contract. However, it is important to note that contracts with significantly large or small time to maturities exhibit notably diminished trading volumes. Therefore, to conduct a more focused analysis, we opt to exclusively investigate contracts with time-to-maturity (‘TTM’) values falling within the range of not more than 45 days and not less than three days.

Step 2: Feature Generation Let Δ denote the time granularity in year unit and $\bar{S}(t) := S(\Delta * t)$ the price of the risky asset at present time $\Delta * t$, where t is an integer. The discrete time series $\{\bar{S}(t) \mid t = 0, 1, \dots\}$ gives daily data if $\Delta = 1/255$, by imagining 255 number of trading days in a year. Let n be a fixed natural number. For each $i \leq n$, the feature variable $F_i := R_{(i)}$, where, $\{R_{(i)} : i = 1, \dots, n\}$ is order statistics of $R_i := \ln(\bar{S}(t - n + i)/\bar{S}(t - n - 1 + i)) - \frac{1}{n} \ln(\bar{S}(t)/\bar{S}(t - n))$, the centered daily log return. The order statistics $R_{(i)}$ are determined by arranging R_i in ascending order for each sample. To clarify, the i^{th} order statistic, denoted by $x_{(i)}$, of a sample comprising various real values $(x_1, x_2, x_3, \dots, x_n)$ represents the i^{th} -smallest value, such that $x_{(i)} = x_{\beta(i)}$ where β is a permutation on $\{1, \dots, n\}$, and $x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(n)}$ is satisfied. We set $n = 19$ for sampling past daily data of nearly one month. At this stage, we compile a set of 23 features to be utilized in supervised learning models:

- The set comprises 19 order statistics of centered log return, denoted by F_1 through F_{19} , derived from index (or underlying asset) data.
- The time-to-maturity (‘TTM’) of the option contract under study, measured in days.
- Moneyness of the option contracts is calculated as the ratio $\frac{K}{S}$, representing the ratio between the strike price K of the option contract and the underlying value or Spot Price S being examined. Its reciprocal is used as a feature.

- Normalized previous option price, computed as $100 \times \frac{C_{t-1}}{S_{t-1}}$, where C_{t-1} denotes the previously reported close price of the option contract under consideration and S_{t-1} denotes the spot price at time $t - 1$. If the option contract was not traded on the previous day (i.e., C_{t-1} data is unavailable), we exclude this option from consideration.
- The ideal bank's interest rate is approximated by the three-month sovereign bond yield rates in the literature and added as a feature. This serves as an estimation for risk-free interest rates.

Step 3: Target Variable We use regression to implement the Homogeneity Hint Approach $\mathcal{A}_{HH}$, as described in Remark 2.3. According to this approach (also see [8]) ‘ $\mathcal{A}_{HH}$ target’ is the ‘normalized price’ and is given by

$$\mathcal{A}_{HH} \text{ target} = 100 \times \frac{C}{S},$$

a scale-free target variable. For the regression implementation of the $\mathcal{A}_{DS}$ approach, as presented in Subsection 3.3, the scale-free target variable referred to as the ‘ $\mathcal{A}_{DS}$ target’ is defined by (3.1), i.e.,

$$\mathcal{A}_{DS} \text{ target} = \frac{1}{\rho} \sqrt{\frac{2\pi}{T}} \left(\frac{C}{S(1+p^*)/2} - \frac{(1-p^*)}{(1+p^*)} \right).$$

where C, S , and T denote the call option price, underlying value, and time-to-maturity in years, respectively. Here, $p = \frac{K}{S}$ is moneyness, $p^* = \frac{Ke^{-rT}}{S} = pe^{-rT}$ is discounted moneyness, where K is the strike price of the call option contract and r is the interest rate of an ideal bank estimated to ‘Yield03’ value divided by 100. The volatility scalar ρ is estimated using historical volatility, the annualized (assuming 255 trading days in a year) standard deviation computed from the set of feature variables F_1 to F_{19} , derived earlier.

Step 4: Ensemble Model At this step we implement the ensemble approach (3.3) formulated in Subsection 3.4. The parameters λ_1 and λ_2 are estimated from a sample of data using the least square method. Given a sample, we first have a train-test split. Following the training of both $\mathcal{A}_{HH}$ and $\mathcal{A}_{DS}$ on the training data, the performance of the ensemble model (3.3) on the test data is measured using RMSE for various values of λ in a square. The value of parameter λ that minimizes the RMSE, gives an estimator of ensembling parameter and hence yields an ensemble model for the given sample. Subsequently the performance of the same is reported for both in-sample and out-of-sample data.

To achieve this, we select one of the datasets: NIFTY50 (N50), BANKNIFTY (BNF), or the combined dataset (referred to as N50+BNF) as the sample. For the numerical least square estimation of λ_1 and λ_2 , we vary them within the range $[0, 5]$ with increment of 0.1. The square root of the sum of squared errors, normalized by the number of data points in the sample, is then computed for both typical and atypical test scenarios for the selected symbol. The values of λ_1 and λ_2 that minimize this quantity are identified as the optimal parameters for the chosen dataset.

Step 5: Testing and Error Metric After the training phase, it is imperative to evaluate our model's performance using unseen test data, ensuring a comprehensive understanding of prediction quality. This study employs Root Mean Square Error (RMSE) to assess regression model predictions. Assume that the total number of data points in the test set is N and e_i is the difference between the target value and the predicted target value of a model $\mathcal{M}$ at the i th entry. Then e_i denotes the prediction error for the i th test data for each $i \leq N$. Finally, the RMSE of the prediction of model $\mathcal{M}$ for the test data set is defined as

$$\text{RMSE} = \left(\frac{1}{N} \sum_{i=1}^N (e_i)^2 \right)^{\frac{1}{2}}. \quad (5.1)$$

It is crucial to emphasize that the evaluation specifically centers on gauging the regression model’s performance with respect to the scale-free target variable denoted as ‘normalized price’. In the context of ‘ $\mathcal{A}_{HH}$ target’ scenarios, RMSE formula are applied directly to the test variable ‘normalized price’. However, for scenarios involving ‘ $\mathcal{A}_{DS}$ target’, we utilize (3.2) to compute the predicted ‘normalized price’ and subsequently evaluate the regression model’s performance.

Step 6: Setting Numerical Values of Hyperparameters We ran the codes in Python version 3.10.4. XGBoost, a well-studied supervised ML algorithm developed by Chen and Guestrin [5], was utilized in this research. Table 2 presents the hyperparameter values used to train the XGBoost (Extreme Gradient Boosting) linear regression model described in this manuscript. While XGBoost offers numerous hyperparameters, we focused on the following key ones. For detailed information on each hyperparameter used, please refer to the XGBoost Documentation⁵.

TABLE 2. Hyperparameter values for XGBoost

Name	Value Set
nthread	4
objective	reg:linear
n_estimators	750
max_depth	7
learning_rate	0.03
min_child_weight	4
colsample_bytree	0.7
subsample	0.7

We conducted a cross-validated grid search using the scikit-learn GridSearchCV routine to determine the optimal values for these hyperparameters. However, we observed that the model performance did not vary significantly with changes in the hyperparameters.

6. MODEL PERFORMANCE

Our study involves analyzing and comparing the performance of three approaches, namely, $\mathcal{A}_{HH}$, $\mathcal{A}_{DS}$, and $\mathcal{A}_E$ for option price predictions. While the first two are supervised learning models, the last one is an ensemble model, whose ensembling parameter is subject to calibration using the sample. The first approach $\mathcal{A}_{HH}$ appeared in [8] whereas the other two are proposed in this article. We have collected historical option price data on two different indices N50 and BNF. Based on the data from each symbol we train three models using above mentioned three approaches. Moreover, we perform multi-source training using historical data from both symbols to get an additional combined model using each of the three approaches. Hence we obtain a total of 9 models. For each symbol of the training set, the ensembling parameter of the ensembling model is estimated from the training and test data of that symbol. We test each of these 9 models on all four disjoint test data samples.

⁵There are many resources available to read about XGBoost parameters; for example, one can refer to <https://xgboost.readthedocs.io/en/stable/parameter.html>

6.1. Performance comparison using RMSE. The model performances are presented in Table 3. In that, the models are classified into 3 classes based on the training data sets. The symbol/symbols of each training data are listed in the first column. The approaches used in the model are mentioned in the second column. The estimated values of ensembling parameters are mentioned next to each $\mathcal{A}_E$ approach in the second column. The option price prediction using historical volatility and the Black-Scholes-Merton formula is set as the benchmark for comparison. We compare the performance of all these models and the benchmark using two types of test data coming from each of the two symbols N50 and BNF. The performance, i.e., the RMSE of predictions, of each model is reported for all four disjoint test data samples. Given a training symbol and test data, the best-performed approach is indicated by highlighting the least RMSE value in orange. Given a test data, the least RMSE across all nine trained models is indicated with a star (*) sign.

TABLE 3. Performance of XGBoost regression for $\mathcal{A}_{HH}$, $\mathcal{A}_{DS}$, and Ensemble Approaches are assessed. RMSE values for models with training (01/01/2015 - 31/08/2019), typical testing (01/09/2019 - 31/12/2019), and atypical testing (01/01/2020 - 30/04/2020) periods are reported. RMSE values for the Ensemble approach are reported for optimal $\lambda_1, \lambda_2 \in [0, 5]$, adjusted in increments of 0.1, across varying training error sets.

RMSEs of option price prediction models using XGBoost regression					
Symbol(s) of Training data	Approach	Test - N50		Test - BNF	
		Typical	Atypical	Typical	Atypical
-	Benchmark - BSM	0.358	1.559	0.527	1.416
N50	$\mathcal{A}_{HH}$	0.176*	1.299	0.348	1.278
	$\mathcal{A}_{DS}$	0.235	0.749	0.302	0.697
	$\mathcal{A}_E(\lambda_1 = 2.6, \lambda_2 = 0.0)$	0.199	0.617*	0.214	0.569*
BNF	$\mathcal{A}_{HH}$	0.193	1.225	0.245	1.181
	$\mathcal{A}_{DS}$	0.269	0.759	0.272	0.741
	$\mathcal{A}_E(\lambda_1 = 0.8, \lambda_2 = 0.5)$	0.173	0.648	0.189*	0.598
N50+BNF	$\mathcal{A}_{HH}$	0.184	1.199	0.255	1.182
	$\mathcal{A}_{DS}$	0.251	0.736	0.288	0.715
	$\mathcal{A}_E(\lambda_1 = 1.3, \lambda_2 = 0.4)$	0.179	0.628	0.194	0.576

At the outset, we notice that for each test data, the column-wise reported RMSE values for all nine training models are smaller than the RMSE of the benchmark prediction. In fact, the lowest RMSE turns out to be smaller than 50% of the benchmark RMSE for each of the four test data sets. It is observed that on every atypical data, the $\mathcal{A}_{DS}$ approach outperforms $\mathcal{A}_{HH}$ approach, no matter what the training data set is. This validates the usefulness of the domain shift approach presented in this paper. On the other hand, we also observe that $\mathcal{A}_{HH}$ outperforms $\mathcal{A}_{DS}$ on the typical test data of each symbol when the training data is from the same symbol. This empirically validates the theoretical conclusion of Theorem 2.2. We explain this below. We recall that $\mathcal{A}_{HH}$ is founded on perfect equalities whereas $\mathcal{A}_{DS}$ is based on some approximations. So, in the special circumstances where Theorem 2.2 is applicable, $\mathcal{A}_{HH}$ gives a better result than $\mathcal{A}_{DS}$. On the other

hand for typical test data of the same symbol of training data, as anticipated, the assumption on similarity between return distributions of training and test is valid. Again, for the same reason, in the multi-source training scenario (see the row with heading N50+BNF), $\mathcal{A}_{HH}$ outperforms $\mathcal{A}_{DS}$ on both of the typical data. This justifies the benefit of combined training. In this scenario, however, the ensemble approach outperforms both $\mathcal{A}_{HH}$ and $\mathcal{A}_{DS}$ on each of the 4 test data. It is interesting to note that each of the 9 models performs better on typical test data of N50 than those of BNF. On the contrary, each of the 9 models performs better on atypical test data of BNF than those of N50. Such consistent performance of approaches on typical or atypical test data irrespective of the source of training data indicates the importance of consideration of domain shifts in the model building. It also indicates that while $\mathcal{A}_{HH}$ is insensitive to the domain shift, $\mathcal{A}_{DS}$ incorporates the domain adaptability reasonably well. So far prediction on BNF atypical test data is concerned, RMSE value of the ensemble model trained with N50 data is lower than that of the ensemble model trained with BNF data. This is another instance of the success of common representation space and domain adaptation.

6.2. Performance of multi-source training model. Now we contrast the performance of multi-source or combined training with single-source training models. We start with the following simple observation from Table 3. We consider the performance of models trained on data of one symbol and tested on typical/atypical data of another symbol with the approach $\mathcal{A}_{HH}$ or $\mathcal{A}_{DS}$. RMSEs of eight such cross-symbol experiments are reported in Table 3. We compare those with eight other corresponding RMSEs for the multi-source models. In seven out of eight cases, the RMSE of multi-source training is lower than those of cross-symbol experiments. Only for BNF atypical test data, the RMSE of multi-source training $\mathcal{A}_{DS}$ model is larger by 2.6% than that by N50 trained $\mathcal{A}_{DS}$ model. We also observe the following. Although $\mathcal{A}_{HH}$ approach performs poorer than $\mathcal{A}_{DS}$ on atypical test data, the performance of $\mathcal{A}_{HH}$ approach on atypical test data gets better in combined training model, compared to single-source training model. Finally, for every test data, the performance of the multi-source ensemble model is close to the best performance. To be more precise, the relative difference of its RMSE value is less than 3% from the RMSE of the best model on the same test data. This is empirical evidence of the reliability of the multi-source ensemble model.

The option price prediction performance of multi-source models with approaches $\mathcal{A}_{HH}$ and $\mathcal{A}_{DS}$, summarized in Table 3, is illustrated more effectively through a histogram of residuals, defined as the differences between predicted and actual normalized price of options, as shown in Figure 5. The horizontal axis represents the range of residuals, which is divided into 100 intervals, while the vertical axis shows the relative frequency. The histograms in Figure 5 is constructed using predictions from the XGBoost model, trained on combined datasets, and tested on NIFTY50 (plot 5a) and BANKNIFTY (plot 5b) datasets respectively. These predictions are made using both the $\mathcal{A}_{DS}$ and $\mathcal{A}_{HH}$ approaches, under typical and atypical testing scenarios. Hence each plot has four histograms. The performance reported in Table 3 indicates that under typical testing scenarios, the $\mathcal{A}_{HH}$ approach yields better accuracy than $\mathcal{A}_{DS}$ approach, as corroborated by the histogram. Specifically in plot 5a, for N50 test set, the residuals from $\mathcal{A}_{HH}$ approach (represented by the ‘blue’ bars) are more tightly clustered around zero compared to the $\mathcal{A}_{DS}$ approach (‘orange’ bars). In contrast, for the BANKNIFTY test data, the difference in frequency levels between the two approaches is less pronounced, indicating a similar error margin. On the N50 typical test data residual distributions for both the approaches exhibit (Figure 5a) symmetric and low dispersion around zero. This confirms the absence of bias in the prediction. A similar low dispersion of the prediction residual is also visible for BANKNIFTY typical test data in Figure 5b.

The dispersion of prediction residual for atypical test data appears to be significantly larger than that of the typical data scenario for every approach in both Figures 5a and 5b. Additionally,

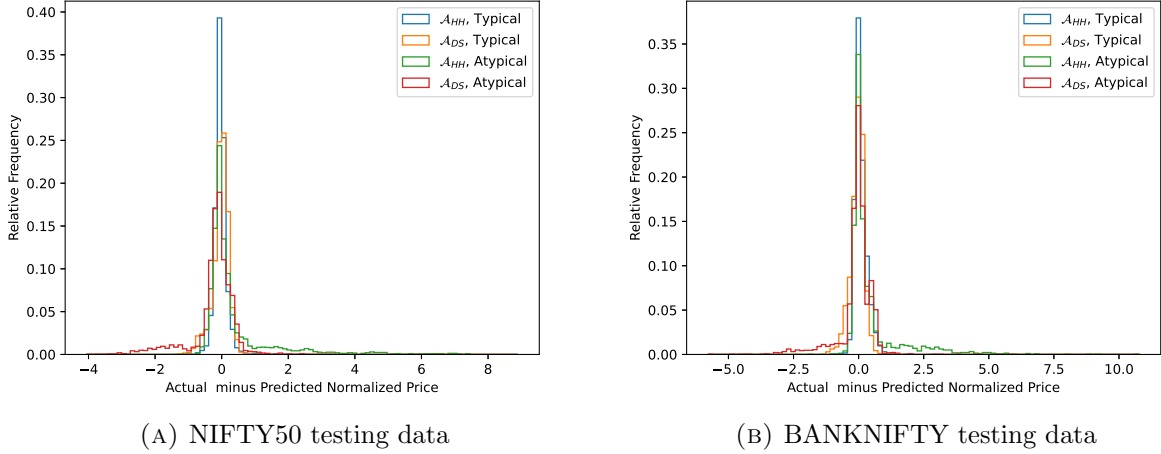


FIGURE 5. Histograms of subtraction of predicted from actual ‘normalized price’ are plotted for models trained on combined N50 and BNF training data. Both models with $\mathcal{A}_{HH}$, and $\mathcal{A}_{DS}$ approaches are tested on different indices with both typical and atypical scenarios.

an intriguing observation in the atypical testing scenario is that the $\mathcal{A}_{DS}$ approach leads to some outliers of overestimation of option prices relative to actual values, as evidenced by the leftward skew of the ‘red bars’, while the $\mathcal{A}_{HH}$ approach leads to outliers of underestimation of prices, as indicated by the rightward skew of the ‘green’ bars across both atypical testing indices. Hence the ensemble model performs significantly better on atypical data by neutralizing such polar outliers in $\mathcal{A}_{HH}$ and $\mathcal{A}_{DS}$ predictions. In summary, the histograms provide a clearer and deeper insight into the models’ performance across different testing scenarios than the RMSE errors offer in Table 3.

6.3. Performance of ensemble models. Next, we wish to gain better insight into the parametric ensemble models and the least square estimation of the ensemble parameters. To this end, we produce Figure 6, which displays contour plots illustrating the RMSE error as a function of the independent variables λ_1 and λ_2 . This helps in providing a clear rationale for why the optimal (λ_1, λ_2) values for the NIFTY50, BANKNIFTY, and combined datasets are those reported in Table 3. These plots offer a two-dimensional representation where points with the same RMSE values are connected by contour lines, effectively visualizing the RMSE landscape. The RMSE values are represented through a color gradient, with ‘blue’ indicating lower errors and ‘red’ signifying higher errors. The plots show the presence of unique minima in each case. For the ensemble approach using NIFTY50 dataset, the lowest RMSE values occur near $\lambda_1 = 2.6$ and $\lambda_2 = 0$, as shown by the darkest ‘blue’ region. This suggests that the domain shift measure (DSQ), introduced in Section 3, has a small impact on determining the weights between the $\mathcal{A}_{HH}$ and $\mathcal{A}_{DS}$ approaches. In contrast, the BANKNIFTY dataset shows that DSQ plays a more significant role, with the lowest RMSE observed around $(\lambda_1, \lambda_2) = (0.8, 0.5)$. The importance of DSQ in weight allocation is also pronounced in the combined dataset, where the minimum RMSE is observed near the optimal point $(\lambda_1, \lambda_2) = (1.3, 0.4)$.

6.4. Model performance on synthetic test data. Now we present the results of two experiments conducted to assess the performance of trained models on synthetic data. The use of synthetic data overcomes real-world data limitations for assessing domain adaptation in rare but large domain shifts. The synthetic data is simulated using Black–Scholes–Merton model of asset

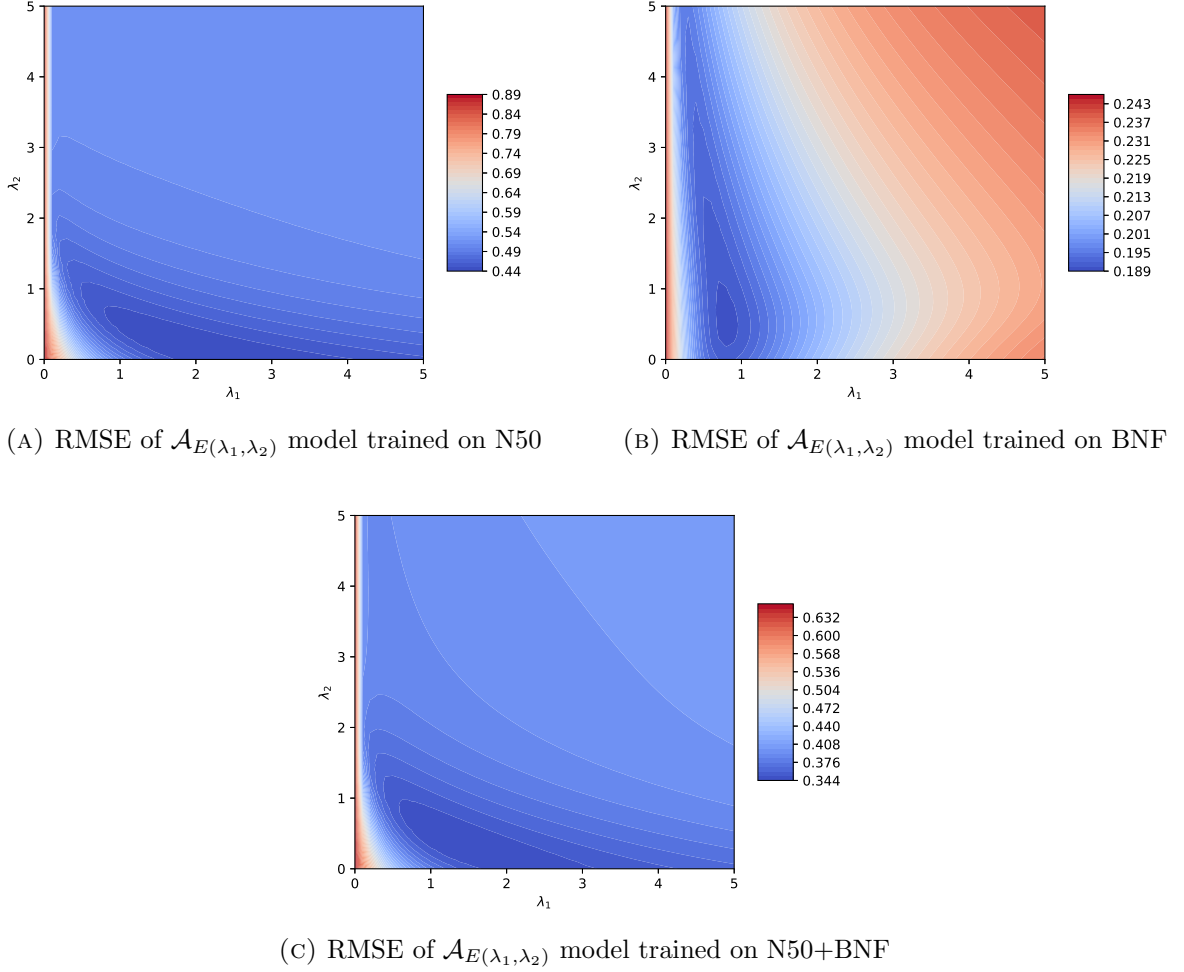


FIGURE 6. Contours for RMSE of prediction by models with $\mathcal{A}_{E(\lambda_1, \lambda_2)}$ approach are plotted against different values of $\lambda_1, \lambda_2 \in [0, 5]$ with increment of 0.1 for each training dataset.

prices and the formula for theoretical option prices. The primary goal of these experiments is to assess how the prediction quality of the proposed approaches varies as the test data shifts from the training dataset. It is important to clarify that these experiments are not intended to evaluate the overall performance of the models but to gain a comparative insight into model performances in extreme cases. These do not give a reliable performance measure in a typical scenario as both underlying asset returns and the option prices of test data are derived under specific mathematical assumptions that may not accurately reflect real-world prices.

We simulate Geometric Brownian motion (GBM) with a drift parameter $\mu = 0.1$ and vary the volatility parameter from 1% to 30% in increments of 1%. For each volatility level, we generate daily data of length 520 days. This test set is supplemented with prices of several near-ATM option contracts, with time to maturity in $\{10, 25, 40\}$ and risk-free interest rate of 5%, calculated using the Black-Scholes-Merton formula. We then test each of the trained models for each variant of the test data and plot the predictions' RMSE against the volatility parameter. The analysis is centered on comparing the performance of the models with $\mathcal{A}_{DS}$, $\mathcal{A}_{HH}$, and ensemble model $\mathcal{A}_E$ approaches, and trained on various datasets. Given the nature of simulated test data, although the distinction

between typical and atypical test scenarios is not direct, still all the synthetic data generated with $\sigma > 20\%$ may be viewed as atypical.

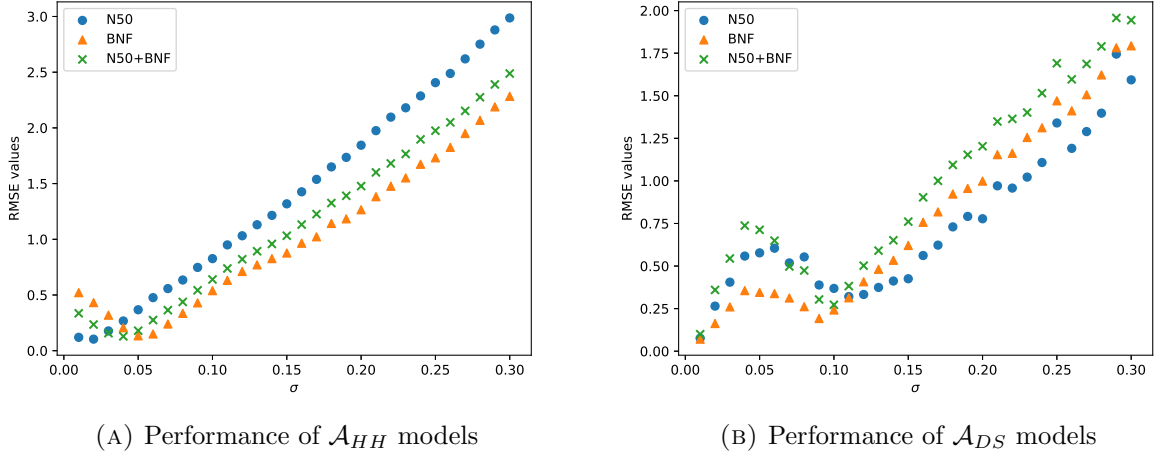
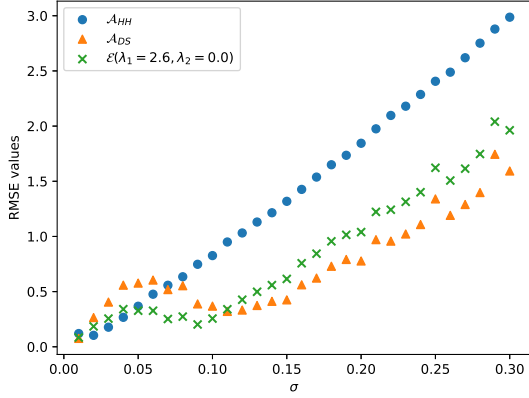


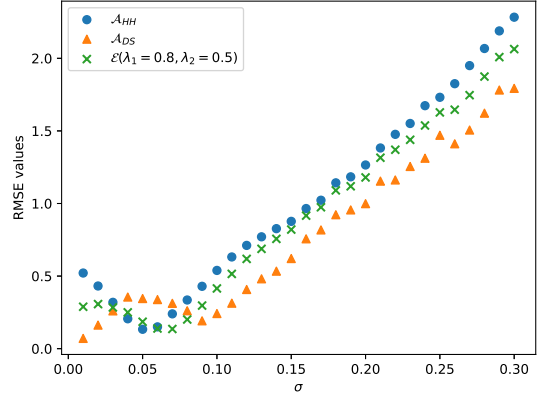
FIGURE 7. Model performance on a range of parametrized synthetic data

6.4.1. Experiment 1. In the first experiment, we evaluate and compare the models' prediction quality on synthetic test data for various training datasets. We perform this comparison for synthetic test data having varying volatility levels, and models with each approach, $\mathcal{A}_{HH}$ and $\mathcal{A}_{DS}$. Indeed, for every synthetic data with volatility parameter σ not less than 5%, among $\mathcal{A}_{HH}$ models, the BNF trained model has the least RMSE, whereas N50 trained model has the largest RMSE. This can be explained using the historical volatility values of respective training data. As the historical volatility of BNF (14.05%) is higher than that of N50 (11.38%), the test data with high σ is less unfamiliar to BNF-trained data than to N50-trained data. The same consistent comparison is not visible in the $\mathcal{A}_{DS}$ based models. In contrast, the $\mathcal{A}_{DS}$ approach aims to mitigate such errors. Figure 7b provides supporting evidence of the absence of performance dependence on the type or volume of training data. Moreover, for $\sigma < 15\%$ no monotonic growth of RMSE with increasing σ of synthetic test data is visible for $\mathcal{A}_{DS}$ models. On the other hand, as noted in [8], the effectiveness of approach $\mathcal{A}_{HH}$ diminishes with domain shift, i.e., the RMSE values of each $\mathcal{A}_{HH}$ model steadily increase for synthetic test data with larger volatility values. Although the RMSE of prediction by $\mathcal{A}_{DS}$ grows with larger σ beyond $\sigma = 15\%$, the growth is slower than that by $\mathcal{A}_{HH}$. This growth of RMSE for $\mathcal{A}_{DS}$ prediction may be attributed to the growth of relative error in approximation relation (3.1) with the increase of disparity(max-min ratio) of volatilities as illustrated in Figure 1b. These results align with the intent of making the $\mathcal{A}_{DS}$ approach less sensitive to discrepancies between the return distributions of the training and testing datasets.

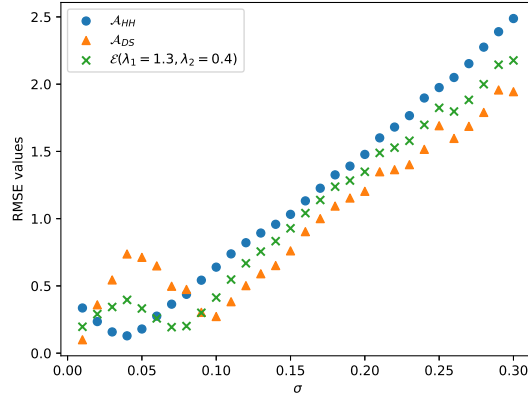
6.4.2. Experiment 2. In another experiment, we compare the RMSE values for all three approaches $\mathcal{A}_{HH}$, $\mathcal{A}_{DS}$ and the ensemble model on the afore mentioned collection of synthetic data. Separate comparison plots are presented for models trained on the NIFTY50, BANKNIFTY and combined datasets. The estimated values of λ_1 and λ_2 , as outlined in Step 4 in Section 5, are used for each selected sample dataset. As anticipated, above a volatility level (i.e., above 8%) of the test data, the approach $\mathcal{A}_{DS}$ outperforms $\mathcal{A}_{HH}$ for every training set (see each subplot of Fig 8). Additionally, performance graphs for the ensemble model are presented, utilizing the estimated parameters λ_1 and λ_2 . In every case (i.e., subplots 8a, 8b, 8c) the performance of the ensemble model is never worst and is between the performance of $\mathcal{A}_{HH}$ and $\mathcal{A}_{DS}$ models except a small range of the volatility values of the synthetic data, where it is the best.



(A) Performance of N50 trained model



(B) Performance of BNF trained model



(C) Performance of multi-source trained model

FIGURE 8. Comparison of approaches $\mathcal{A}_{HH}$, $\mathcal{A}_{DS}$ and $\mathcal{A}_E$ for each of three different trained set using a parametric family of synthetic data

7. CONCLUSION

This work presents several interesting elements and fresh ideas for option pricing applications of machine learning. This paper puts forward an interesting idea of utilizing a formula for implied volatility estimates to create a common representation space so that a better domain adaptation is achieved. This method also allows training models on several options markets (i.e., options with different underlying securities), since this could improve out-of-sample performance in each considered market. Finally, we formulate a novel ensemble model, depending on a domain shift quotient, that works for both typical and atypical test data. We believe, this is a promising avenue for further research and could potentially bring some new insights into the field of data-driven option pricing. Moreover, we also assess the models using a family of synthetic test data for gaining a better understanding of the models.

Acknowledgments: The first author acknowledges the financial support from the National Board for Higher Mathematics (NBHM), Department of Atomic Energy, Government of India, under grant number 02011-32-2023-R&D-II-13347 and the Department of Science and Technology under the Project-Related Personal Exchange Grant DST/INTDAAD/P-12/2020. The support and the

resources provided by ‘PARAM Brahma Facility’ under the National Supercomputing Mission, Government of India at the Indian Institute of Science Education and Research (IISER) Pune are gratefully acknowledged. The financial support by the London Mathematical Society through Scheme 5 ‘Collaborations with Developing Countries’, with Ref 52335, is greatly acknowledged by the last author.

Declarations:

- Conflict of interest: The authors declare no conflict of interest.
- Availability of data and materials: Data sharing is not applicable to this article as no new data were created or analyzed in this study. All the empirical studies reported in this article are based on openly available data.
- Authors’ contributions: All authors contributed equally towards the manuscript.

REFERENCES

- [1] M. BHARADIA, N. CHRISTOFIDES, AND G. SALKIN, *Computing the Black-Scholes implied volatility*, in Advances in Futures and Options Research, P. P. B. et al., ed., JAI Press, 1996, pp. 15–29.
- [2] F. BLACK AND M. SCHOLES, *The pricing of options and corporate liabilities*, Journal of Political Economy, 81 (1973), p. 637–654.
- [3] C. BOEK, P. LAJBCYGIER, M. PALANISWAMI, AND A. FLITMAN, *A hybrid neural network approach to the pricing of options*, in Proceedings of ICNN’95 - International Conference on Neural Networks, vol. 2, 1995, pp. 813–817 vol.2.
- [4] Y. CAO, X. LIU, AND J. ZHAI, *Option valuation under no-arbitrage constraints with neural networks*, European J. Oper. Res., 293 (2021), pp. 361–374.
- [5] T. CHEN AND C. GUESTRIN, *XGBoost: A scalable tree boosting system*, in Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, KDD ’16, New York, NY, USA, 2016, Association for Computing Machinery, p. 785–794.
- [6] M. R. FENGLER, *Arbitrage-free smoothing of the implied volatility surface*, Quant. Finance, 9 (2009), pp. 417–428.
- [7] R. GARCIA AND R. GENÇAY, *Pricing and hedging derivative securities with neural networks and a homogeneity hint*, Journal of Econometrics, 94 (2000), pp. 93–115.
- [8] A. GOSWAMI, S. RAJANI, AND A. TANKSALE, *Data-driven option pricing using single and multi-asset supervised learning*, International Journal of Financial Engineering, 08 (2021), p. 2141001.
- [9] N. GRADOJEVIC, R. GENÇAY, AND D. KUKOLJ, *Option pricing with modular neural networks*, IEEE Trans Neural Netw., 20 (2009), pp. 626–637.
- [10] J. M. HUTCHINSON, A. W. LO, AND T. POGGIO, *A nonparametric approach to pricing and hedging derivative securities via learning networks*, The Journal of Finance, 49 (1994), pp. 851–889.
- [11] G. KALLIANPUR AND R. L. KARANDIKAR, *Introduction to option pricing theory*, Birkhäuser Boston, Inc., Boston, MA, 2000.
- [12] S. LI, *A new formula for computing implied volatility*, Appl. Math. Comput., 170 (2005), pp. 611–625.
- [13] M. MALLIARIS AND L. SALCHENBERGER, *A neural network model for estimating option prices*, Applied Intelligence, 3 (1993), pp. 193–206.
- [14] M. MALLIARIS AND L. SALCHENBERGER, *Using neural networks to forecast the S&P 100 implied volatility*, Neurocomputing, 10 (1996), pp. 183–195. Financial Applications, Part I.
- [15] M. QI AND G. MADDALA, *Option pricing using artificial neural networks: The case of S&P 500 index call options*, in Proceedings of 3rd International Conference on Neural Networks in the Capital Markets, World Scientific, 1996, pp. 78–91.
- [16] J. RUF AND W. WANG, *Neural networks for option pricing and hedging: a literature review*, Journal of Computational Finance, 24 (2020), pp. 1–46.
- [17] A. SHARMA, C. VERMA, AND P. SINGH, *Enhancing option pricing accuracy in the Indian market: A CNN-BiLSTM approach*, Computational Economics, (2024).
- [18] K. SHUBHAM, V. TIWARI, AND K. S. PATEL, *Predictive learning methods to price european options using ensemble model and multi-asset data*, International Journal of Artificial Intelligence Tools, 32 (2023), p. 2350034.
- [19] J. SIRIGNANO AND R. CONT, *Universal features of price formation in financial markets: perspectives from deep learning*, Quant. Finance, 19 (2019), pp. 1449–1459.
- [20] W. WANG AND J. RUF, *A note on spurious model selection*, Quantitative Finance, 22 (2022), pp. 1797–1800.
- [21] J. YAO, Y. LI, AND C. L. TAN, *Option price forecasting using neural networks*, Omega, 28 (2000), pp. 455–466.