

IMPERIAL

IMPERIAL COLLEGE LONDON

DEPARTMENT OF MATHEMATICS

Effective Approximation Methods for Utility Maximization Problems with Partially Observable Drifts

Author: Yushan Zhang (CID: 06010127)

A thesis submitted for the degree of
MSc in Mathematics and Finance, 2024-2025

Declaration

The work contained in this thesis is my own work unless otherwise stated.

Acknowledgements

First and foremost, I would like to express my deepest gratitude to my supervisor, Professor Harry Zheng. His teaching in stochastic control has laid a solid academic foundation for me. His continuous encouragement and insightful guidance throughout my research and writing of this thesis have been a source of inspiration and motivation. Without his invaluable advice, my dissertation project would not have remained on the right track and progressed so smoothly.

I would also like to extend my sincere gratitude to all the academic staff of the MSc Mathematics and Finance program at Imperial College London for their dedication and generous support over the past year. The resources and support they have provided have not only enabled us to achieve our academic goals but also allowed us to grow in ways that go beyond academics.

Finally, my heartfelt thanks goes to my family and friends. My parents' unconditional support has given me the opportunity to study at such a world-class university, and every call from my family was filled with care. In particular, my father's presence during the last two months in the UK helped me to overcome challenges and psychological pressure. I am grateful as well to my friend, Charilaos Managias, whose many conversations with me made the research process more enjoyable.

Abstract

This paper concentrates on solving partially observable utility maximization problems in the context of trading risk-free and risky assets in financial markets. A class of stochastic control problems are studied, where the drift vector of stock prices process cannot be fully determined by observable information. The filtering theory is applied to estimate the drift vector and transform original problems into fully observable problems. Kalman-Bucy filters and Bayesian filters are both analyzed. The corresponding HJB equations, dual problems and forward-backward SDE systems are derived.

Numerical experiments are conducted to solve this class of problems. Two types of algorithms are tested for 1D and 2D problems, including the bound estimate approach and the deep learning approach. The former estimates lower bounds and upper bounds for solutions by the dual control theory and the stochastic maximum principle, and thus numerical solutions are midpoints of these two bounds. In contrast, the latter trains a neural network to directly solve HJB equations, so that numerical solutions are obtained.

Results show that both algorithms can accurately solve these problems with small errors. Moreover, by comparing value functions of partially observable problems and fully observable problems, it is empirically proved that unobservable information has positive values, and a way to measure the information value is also obtained.

Contents

1	Introduction	6
2	1D Partially Observable Problems under Kalman-Bucy Filters	8
2.1	Problem Formulation	8
2.2	Problem Transformation by 1D Kalman-Bucy Filters	9
2.3	The Bound Estimate Approach	11
2.3.1	Dual Control Problems	12
2.3.2	Stochastic Maximum Principle	13
2.3.3	Numerical Experiments	15
2.4	The Deep Learning Approach	19
3	2D Partially Observable Problems under Kalman-Bucy Filters	26
3.1	Problem Formulation	26
3.2	Problem Transformation by 2D Kalman-Bucy Filters	27
3.3	The Bound Estimate Approach	28
3.3.1	Dual Control Problems	29
3.3.2	Stochastic Maximum Principle	29
3.3.3	Numerical Experiments	30
4	1D Partially Observable Problems under Bayesian Filters	36
4.1	1D Bayesian Filters	36
4.2	The Bound Estimate Approach	40
4.2.1	Dual Control Problems and Stochastic Maximum Principle	41
4.2.2	Numerical Experiments	41
4.3	The Deep Learning Approach	43
5	2D Partially Observable Problems under Bayesian Filters	49
5.1	2D Bayesian Filters	49
5.2	The Bound Estimate Approach	53
5.2.1	Dual Control Problems and Stochastic Maximum Principle	53
5.2.2	Numerical Experiments	54
6	Information Value	57
6.1	Fully Observable Problems	57
6.2	Numerical Experiments	58
7	Conclusion	61
	Bibliography	63

List of Figures

2.1	The training loss curve of solving the HJB equation with $h_0 = 0.1, \Sigma = 0.2$ under the Kalman-Bucy filter for unconstrained problems	22
2.2	The training loss curve of solving the HJB equation with $h_0 = 0.1, \Sigma = 0.2$ under the Kalman-Bucy filter for unconstrained problems	22
2.3	$J(0, x_0, h_0)$ curves by varying Σ_0 under the Kalman-Bucy filter for unconstrained problems	23
2.4	$J(0, x_0, h_0)$ curves by varying Σ_0 under the Kalman-Bucy filter for constrained problems	23
2.5	The value function surface at $t = 0$ with $\Sigma_0 = 0.2$ under the Kalman-Bucy filter (left panel: unconstrained problems, right panel: constrained problems)	25
3.1	An example of the solution curve of $\Sigma(t), t \in [0, T]$	32
4.1	The training loss curve of solving the HJB equation with $h_0 = 0.1, \Sigma = 0.2$ under the Bayesian filter for unconstrained problems	45
4.2	The training loss curve of solving the HJB equation with $h_0 = 0.1, \Sigma = 0.2$ under the Bayesian filter for constrained problems	45
4.3	$J(0, x_0, h_0)$ curves by varying Σ_0 under the Bayesian filter for unconstrained problems	46
4.4	$J(0, x_0, h_0)$ curves by varying Σ_0 under the Bayesian filter for constrained problems	46
4.5	The value function surface at $t = 0$ with $\Sigma_0 = 0.2$ under the Bayesian filter (left panel: unconstrained problems, right panel: constrained problems) . .	48
6.1	Full information value and partial information value at $t = 0$ (Kalman-Bucy filter for unconstrained problems)	59
6.2	Full information value and partial information value at $t = 0$ (Kalman-Bucy filter for constrained problems)	59
6.3	Full information value and partial information value at $t = 0$ (Bayesian filter for unconstrained problems)	60
6.4	Full information value and partial information value at $t = 0$ (Bayesian filter for constrained problems)	60

List of Tables

2.1	Parameters to determine for 1D unconstrained problems under the Kalman-Bucy filter	16
2.2	Numerical results for 1D unconstrained problems by varying h_0 under the Kalman-Bucy filter	17
2.3	Numerical results for 1D unconstrained problems by varying Σ_0 under the Kalman-Bucy filter	17
2.4	Parameters to determine for 1D constrained problems under the Kalman-Bucy filter	18
2.5	Numerical results for 1D constrained problems by varying h_0 under the Kalman-Bucy filter	18
2.6	Numerical results for 1D constrained problems by varying Σ_0 under the Kalman-Bucy filter	18
2.7	Numerical results for 1D constrained problems by varying w under the Kalman-Bucy filter	18
2.8	Layers of the neural network $\hat{J}_\theta$	20
2.9	Numerical results for 1D unconstrained problems by varying h_0 under the Kalman-Bucy filter with PINN results included	24
2.10	Numerical results for 1D unconstrained problems by varying Σ_0 under the Kalman-Bucy filter with PINN results included	24
2.11	Numerical results for 1D constrained problems by varying h_0 under the Kalman-Bucy filter with PINN results included	24
2.12	Numerical results for 1D constrained problems by varying Σ_0 under the Kalman-Bucy filter with PINN results included	24
3.1	Parameters to determine for 2D unconstrained problems under the Kalman-Bucy filter	33
3.2	Numerical results for 2D unconstrained problems by varying $\mathbf{h}_0$ under the Kalman-Bucy filter	33
3.3	Numerical results for 2D unconstrained problems by varying Σ_0 under the Kalman-Bucy filter	34
3.4	Parameters to determine for 2D constrained problems under the Kalman-Bucy filter	35
3.5	Numerical results for 2D constrained problems by varying $\mathbf{h}_0$ under the Kalman-Bucy filter	35
3.6	Numerical results for 2D constrained problems by varying Σ_0 under the Kalman-Bucy filter	35
3.7	Numerical results for 2D constrained problems by varying w under the Kalman-Bucy filter	35
4.1	Numerical results for 1D unconstrained problems by varying h_0 under the Bayesian filter	42

4.2	Numerical results for 1D unconstrained problems by varying Σ_0 under the Bayesian filter	42
4.3	Numerical results for 1D constrained problems by varying h_0 under the Bayesian filter	43
4.4	Numerical results for 1D constrained problems by varying Σ_0 under the Bayesian filter	43
4.5	Numerical results for 1D constrained problems by varying w under the Bayesian filter	43
4.6	Numerical results for 1D unconstrained problems by varying h_0 under the Bayesian filter with PINN results included	47
4.7	Numerical results for 1D unconstrained problems by varying σ_0 under the Bayesian filter with PINN results included	47
4.8	Numerical results for 1D constrained problems by varying h_0 under the Bayesian filter with PINN results included	47
4.9	Numerical results for 1D constrained problems by varying Σ_0 under the Bayesian filter with PINN results included	47
5.1	Numerical results for 2D unconstrained problems by varying $\mathbf{h}_0$ under the Bayesian filter	55
5.2	Numerical results for 2D unconstrained problems by varying Σ_0 under the Bayesian filter	55
5.3	Numerical results for 2D constrained problems by varying $\mathbf{h}_0$ under the Bayesian filter	55
5.4	Numerical results for 2D constrained problems by varying Σ_0 under the Bayesian filter	56
5.5	Numerical results for 2D constrained problems by varying w under the Bayesian filter	56

Chapter 1

Introduction

Continuous-time utility maximization problems have been widely studied, especially in the field of quantitative financial economics. In these problems, asset prices are usually driven by SDEs and investors must determine the proportion of wealth allocated to each asset to achieve the best investment results.

A finite investment horizon $[0, T]$ (with $T < \infty$) is usually specified so that the optimal portfolio is constructed by maximizing the utility which is some function of the terminal wealth (the wealth at T), i.e. utility functions. Introductions to these problems can be seen in [1], [7], [13], and [17]. The investment horizon can also be infinite, where the terminal time point is $T = \infty$. This means that the trading never stops and the total discounted utility over the long period is treated as the objective function. For approaches that solve different kinds of problems with infinite investment horizons, we refer to [2], [6], [10], [13] and [14]. We assume that $T < \infty$ all the time in this paper.

When all parameters in the model can be directly observed and thus are known, the problem is called a fully observable problem. In this case, the problem is a standard stochastic control problem. There are standard approaches to solve these problems such as HJB equations and the dual control theory (see [4], [15] and [18]). However, in practice, we cannot avoid situations where some parameters are unknown. A problems like this is called a partially observable problem.

[9] considers a class of partially observable problems where the drift vector of stock prices at any time t (denoted by $\boldsymbol{\mu}(t)$) is not available. Therefore, $\boldsymbol{\mu}(t)$ can only be estimated using the information that can be collected from 0 to t (denoted by $I(t)$):

$$\hat{\boldsymbol{\mu}}(t) = \mathbb{E}[\boldsymbol{\mu}(t)|I(t)], t \in [0, T],$$

where $\hat{\boldsymbol{\mu}}(t)$ is the estimate of $\boldsymbol{\mu}(t)$. This is the idea of filtering theory. Although in general, the filter process $\hat{\boldsymbol{\mu}}(t)$ has no closed-form SDEs, there still exist three types of filters that are driven by closed-form SDEs, including the Kalman-Bucy filter (which assumes that $\boldsymbol{\mu}(t)$ is driven by a linear diffusion), the Wonham filter (which assumes that $\boldsymbol{\mu}(t)$ is driven by a continuous-time Markov chain whose state space is finite) and the Bayesian filter (which assumes that $\boldsymbol{\mu}(t)$ is a fixed random vector). By the filtering theory, partially observable problems can be transformed into fully observable problems.

Concentrating on estimating the unknown drift vector by the Kalman-Bucy filter, [18] considers the same class of problems as [9]. But [18] includes some extensions such as constraints that short selling is rejected. Moreover, [18] gives a numerical approach (the bound estimate approach) that computes the lower bound $\underline{L}$ and the upper bound $\overline{L}$ of the value function. It shows by a series of numerical experiments that $\underline{L}$ and $\overline{L}$ are close to each other, meaning that the midpoint $\frac{\underline{L} + \overline{L}}{2}$ is a good estimate of the value function.

This paper follows [18] and we provide more extensions and numerical experiments. First, we reproduce all numerical results given by [18], where 1D partially observable

problems are solved by the bound estimate approach. Second, we verify these numerical results by solving HJB equations by deep learning models (the PINN technique). Third, [18] only tests 1D problems in the numerical experiment, while 2D problems are also tested in this paper. Finally, [18] derives all results under the Kalman-Bucy filter, while we also test the Bayesian filter.

The remaining contents of this paper are organized as follows. Chapter 2 solves 1D partially observable problems under the Kalman-Bucy filter, where both the bound estimate approach and the deep learning approach are tested. Chapter 3 solves 2D partially observable problems under the Kalman-Bucy filter, where the bound estimate approach is tested. Chapter 4 solves 1D partially observable problems under the Bayesian filter, where both the bound estimate approach and the deep learning approach are tested. Chapter 5 solves 2D partially observable problems under the Bayesian filter, where the bound estimate approach is tested. Chapter 6 discusses the information value which is the difference between the value function of fully observable problems and that of partially observable problems. Chapter 7 concludes.

Chapter 2

1D Partially Observable Problems under Kalman-Bucy Filters

In this chapter we have a recap on 1D partially observable utility maximization problems discussed by Zhu and Zheng [18]. We conduct all numerical experiments done by these authors to ensure that we are able to truly reproduce what they obtain. Also, we use the deep learning approach, namely physics-informed neural network (PINN), to solve the same problems, so that results in [18] can be further verified.

2.1 Problem Formulation

Consider the situation where the financial market only has two assets, one is the risk-free bond with price process $\{S_0(t), t \in [0, T]\}$ and the other is a stock (risky) with price process $\{S(t), t \in [0, T]\}$, where $[0, T]$ is the investment horizon. The financial market is modelled by a filtered probability space $(\Omega, \mathbb{F}, \mathcal{F}, \mathbb{P})$, where $\mathcal{F} = \{\mathbb{F}(t), t \in [0, T]\}$ is the filtration that describes all information up to t (the full information). For any t , only part of the information in $\mathcal{F}(t)$ can be observed. Process S_0 is driven by (2.1.1):

$$dS_0(t) = r(t)S_0(t)dt, \quad (2.1.1)$$

where $\{r(t), t \in [0, T]\}$ is the risk-free rate process. In contrast, process S follows SDE (2.1.2)

$$dS(t) = \mu(t)S(t)dt + \sigma(t)S(t)dW(t), \quad (2.1.2)$$

where $\{\mu(t), t \in [0, T]\}$ and $\{\sigma(t), t \in [0, T]\}$ are the drift process and volatility process respectively.

To build a partially observable problem, we define $\mathcal{F}^S$ (the observable information) as the filtration generated by process S , and assume that both r and σ are $\mathcal{F}^S$ -measurable processes on $\Omega \times [0, T]$. Obviously, we have $\mathcal{F}^S \subset \mathcal{F}$, i.e. the information that we can observe is a subset of the full information (partially observable). In contrast, μ is assumed to be $\mathcal{F}$ -measurable, meaning that we cannot obtain the value of $\mu(t)$ only by the observable information $\mathcal{F}^S(t)$, i.e. μ is not observable. One way to solve this problem is that we can use $\hat{\mu}(t)$, the conditional expectation of $\mu(t)$ given $\mathcal{F}^S(t)$, as the estimate of $\mu(t)$:

$$\hat{\mu}(t) = \mathbb{E} [\mu(t) | \mathcal{F}^S(t)]. \quad (2.1.3)$$

This idea is based on the filtering theory.

Suppose that the investor trades these two assets and only follows self-financing trading strategies denoted by process $\{\pi(t), t \in [0, T]\}$, where $\pi(t)$ is the fraction of total wealth invested in the stock at t and must be $\mathcal{F}^S(t)$ -measurable, because the investment

decision must be made merely through the observable information. Accordingly, the fraction of wealth allocated to the bond at t is $1 - \pi(t)$. The set of all admissible trading strategies is given by

$$\mathcal{A} = \left\{ \pi : \pi \in \mathcal{H}^2(0, T; \mathbb{R}) \text{ and } \pi(t) \in K \text{ almost everywhere on } t \in [0, T] \right\}. \quad (2.1.4)$$

In (2.1.4), set $K \subset \mathbb{R}$ stands for constraints on π . For example, when short selling is not allowed, the investor has to require that $\pi(t) \geq 0$ almost everywhere on $t \in [0, T]$. In this case, the constraint set $K = \mathbb{R}^+$. For unconstrained cases, we have $K = \mathbb{R}$. Space $\mathcal{H}^2(0, T; \mathbb{R})$ in (2.1.4) is defined as

$$\mathcal{H}^2(0, T; \mathbb{R}) = \left\{ \zeta : \Omega \times [0, T] \rightarrow \mathbb{R} : \zeta \text{ is } \mathcal{F}^S\text{-measurable, and } \mathbb{E} \left[\int_0^T |\zeta(s)|^2 ds \right] < \infty \right\},$$

which contains all real-valued square-integrable processes that are $\mathcal{F}^S$ -measurable.

For any trading strategy π , since π is self-financing, the wealth process $\{X^\pi(t), t \in [0, T]\}$ is driven by the following SDE:

$$\begin{aligned} dX^\pi(t) &= \underbrace{\pi(t)X^\pi(t)}_{\text{wealth in the stock}} \underbrace{(\mu(t)dt + \sigma(t)dW(t))}_{\text{growth rate of the stock}} + \underbrace{(1 - \pi(t))X^\pi(t)}_{\text{wealth in the bond}} \underbrace{r(t)dt}_{\text{growth rate of the bond}} \\ &= X^\pi(t) \{[r(t) + \pi(t)(\mu(t) - r(t))]dt + \pi(t)\sigma(t)dW(t)\}, t \in [0, T], \end{aligned}$$

with initial wealth $X^\pi(0) = x > 0$. The utility is connected with wealth through the utility function $U : \mathbb{R}^+ \rightarrow \mathbb{R}$, and U must be continuous, increasing, concave with $U(0) = 0$, i.e. zero wealth leads to zero utility. It is true that other types of utility functions may also be applied in real-world problems such as the U-shaped utility functions discusses in [4]. However, these utility function are not discussed in this paper.

The goal is maximizing the expected terminal utility, namely $\mathbb{E}[U(X^\pi(T))]$, and the corresponding value function is

$$V(x) = \sup_{\pi \in \mathcal{A}} \mathbb{E}[U(X^\pi(T))].$$

A trading strategy π^* is optimal if it achieves the supremum:

$$V(x) = \mathbb{E}[U(X^{\pi^*}(T))],$$

with the optimal state process $\{X^{\pi^*}(t), t \in [0, T]\}$. To simplify notations, the optimal state process is also written as $\{X^*(t), t \in [0, T]\}$.

2.2 Problem Transformation by 1D Kalman-Bucy Filters

Directly solving the partially observable problem is difficult. However, it is possible to transform this problem into a fully observable one. Consider process $\{\hat{V}(t), t \in [0, T]\}$ defined by

$$\hat{V}(t) = W(t) + \int_0^t \frac{\mu(s) - \hat{\mu}(s)}{\sigma(s)} ds, t \in [0, T], \quad (2.2.1)$$

where $\hat{\mu}$ is the filter of μ (see (2.1.3)). Following [8], we have theorem 2.2.1.

Theorem 2.2.1. *$\hat{V}$ is an $\mathcal{F}^S$ -measurable Brownian motion under probability measure $\mathbb{P}$ provided that process $\{\sigma^{-1}(t), t \in [0, T]\}$ is uniformly bounded and*

$$\mathbb{E} \left[\int_0^T \left| \frac{\mu(s)}{\sigma(s)} \right|^2 ds \right] < \infty.$$

Proof: see [8] for the proof.

$\hat{V}$ can be used to transform the wealth process so that X^π is driven by a new SDE whose randomness is $\mathcal{F}^S$ -measurable, meaning that the transformed problem is fully observable. To be specific, by the fact that

$$d\hat{V}(t) = dW(t) + \frac{\mu(t) - \hat{\mu}(t)}{\sigma(t)} dt,$$

we obtain the new SDE:

$$\begin{aligned} dX^\pi(t) &= X^\pi(t) \{ [r(t) + \pi(t)(\mu(t) - r(t))] dt + \pi(t)\sigma(t)dW(t) \} \\ &= X^\pi(t) \left\{ [r(t) + \pi(t)(\mu(t) - \hat{\mu}(t) + \hat{\mu}(t) - r(t))] dt + \pi(t)\sigma(t) \left(d\hat{V}(t) - \frac{\mu(t) - \hat{\mu}(t)}{\sigma(t)} dt \right) \right\} \\ &= X^\pi(t) \left\{ [r(t) + \pi(t)(\hat{\mu}(t) - r(t))] dt + \pi(t)\sigma(t)d\hat{V}(t) \right\}. \end{aligned}$$

The form of $\hat{\mu}(t)$ depends on the filter that we choose. In this section, the Kalman-Bucy filter is applied. Let $H(t) = \mu(t)$ and suppose that H follows SDE

$$dH(t) = -\lambda (H(t) - \bar{H}) dt + \sigma_H \left(\rho W(t) + \sqrt{1 - \rho^2} dW_H(t) \right),$$

where $\lambda, \bar{H}, \sigma_H \in \mathbb{R}, \rho \in [-1, 1]$ and W_H is another Brownian motion that is independent of W . Moreover, the initial value $H(0)$ is independent of W and W_H . The Kalman-Bucy filter is the estimate

$$\hat{\mu}(t) = \hat{H}(t) = \mathbb{E} [H(t) | \mathcal{F}^S(t)]. \quad (2.2.2)$$

At $t = 0$, no information is available, so the filter is $\hat{H}(0) = \mathbb{E} [H(0)]$ initially. When the conditional distribution of $H(t)$ given $\mathcal{F}^S(t)$ is Gaussian, the Kalman-Bucy filter must follow the SDE below (see [12]):

$$d\hat{H}(t) = -\lambda (\hat{H}(t) - \bar{H}) dt + \left(\frac{\Sigma(t)}{\sigma(t)} + \rho\sigma_H \right) d\hat{V}(t), \quad (2.2.3)$$

where $\Sigma(t)$ is the conditional variance of the filter:

$$\Sigma(t) = \text{Var} [H(t) | \mathcal{F}^S(t)].$$

Since we know that the conditional mean is given by (2.2.2), the conditional variance can also be written as

$$\Sigma(t) = \mathbb{E} \left[\left(H(t) - \hat{H}(t) \right)^2 \middle| \mathcal{F}^S(t) \right].$$

Additionally, [12] shows that function $\Sigma(t)$ is determined by ODE

$$\frac{d\Sigma(t)}{dt} = \sigma_H^2 - 2\lambda\Sigma(t) - \left(\frac{\Sigma(t)}{\sigma(t)} + \rho\sigma_H \right)^2, \quad (2.2.4)$$

with some prespecified prior value $\Sigma(0) = \text{Var} [H(0)]$. Typically, for the filter, its initial mean $h_0 = \hat{H}(0)$ and initial variance $\Sigma_0 = \Sigma(0)$ must be determined by the prior knowledge (such as the practical experience) of users. When $\sigma(t)$ is a constant (i.e. $\sigma(t) = \sigma$), equation (2.2.4) has a closed-form solution:

$$\Sigma(t) = \sqrt{G}\sigma \frac{G_1 \exp\left(\frac{2\sqrt{G}}{\sigma}t\right) + G_2}{G_1 \exp\left(\frac{2\sqrt{G}}{\sigma}t\right) - G_2} - (\lambda\sigma + \rho\sigma_H)\sigma,$$

where

$$\begin{aligned} G &= (\lambda\sigma + \rho\sigma_H)^2 + \sigma_H^2 (1 - \rho^2), \\ G_1 &= \sqrt{G}\sigma + (\lambda\sigma + \rho\sigma_H)\sigma + \Sigma_0, \\ G_2 &= G_1 - 2\sqrt{G}\sigma. \end{aligned}$$

This solution is used in numerical experiments later. When $\lambda \in \mathbb{R}^+$, H is mean-reverting and becomes an Ornstein-Uhlenbeck process. There exists other filters that are special and useful, like Wonham filter and Bayesian filter, which we will discuss later. We provide discussions on problems under the Bayesian filter in chapter 4 and 5.

The value function of the transformed problem evaluated at t depends not only on the current state $X^\pi(t)$ but also on the estimate of $H(t)$ (i.e. $\hat{H}(t)$) based on all available information up to t . Consequently, the value function is represented by

$$J(t, x, h) = \sup_{\pi \in \mathcal{A}} \mathbb{E}_{t,x,h} [U(X^\pi(T))],$$

where the expectation operator $\mathbb{E}_{t,x,h}[\cdot] = \mathbb{E}[\cdot | X^\pi(t) = x, \hat{H}(t) = h]$ is a conditional expectation. Note that at each time t , we can only evaluate $\hat{H}(t)$ instead of $H(t)$ which is only measurable when the full information is given. To summarize, the newly constructed 1D fully observable problem is formulated as

$$\begin{cases} dX^\pi(t) = X^\pi(t) \left\{ \left[r(t) + \pi(t) (\hat{H}(t) - r(t)) \right] dt + \pi(t)\sigma(t)d\hat{V}(t) \right\}, X^\pi(0) = x, \\ d\hat{H}(t) = -\lambda (\hat{H}(t) - \bar{H}) dt + \left(\frac{\Sigma(t)}{\sigma(t)} + \rho\sigma_H \right) d\hat{V}(t), \hat{H}(0) = h_0, \\ J(t, x, h) = \sup_{\pi \in \mathcal{A}} \mathbb{E}_{t,x,h} [U(X^\pi(T))]. \end{cases} \quad (2.2.5)$$

We solve this problem by two different approaches. The first approach is numerically solving the HJB equation of problem (2.2.5). The value function $J(t, x, h)$ has three arguments which makes the finite difference method not the best choice for this application. In contrast, deep learning models can effectively handle solution functions whose dimension is 3 or higher in complex PDEs. This approach is typically called PINN. The second approach cannot solve problem (2.2.5) exactly but simply estimate a lower bound and an upper bound for the value function by duality control theory and stochastic maximum principle (SMP). When these two bounds are close enough to each other, we can use their average as a solution whose error is within our tolerance. This is the approach proposed by [18]. We discuss these two approaches in the following two sections respectively.

2.3 The Bound Estimate Approach

The HJB equation of value function $J(t, x, h)$ in this problem is

$$\begin{aligned} \frac{\partial J}{\partial t} + \sup_{\pi(t) \in K} \left\{ x [r(t) + \pi(t) (h - r(t))] \frac{\partial J}{\partial x} + \frac{1}{2} x^2 \pi^2(t) \sigma^2(t) \frac{\partial^2 J}{\partial x^2} - \lambda (h - \bar{H}) \frac{\partial J}{\partial h} + \right. \\ \left. x \pi(t) \sigma(t) \left(\frac{\Sigma(t)}{\sigma(t)} + \rho\sigma_H \right) \frac{\partial^2 J}{\partial x \partial h} + \frac{1}{2} \left(\frac{\Sigma(t)}{\sigma(t)} + \rho\sigma_H \right)^2 \frac{\partial^2 J}{\partial h^2} \right\} = 0. \end{aligned} \quad (2.3.1)$$

In general, directly solving (2.3.1) is difficult. It is impossible to find a suitable ansatz for (2.3.1). Also, martingale representation theorem cannot be applied. These troubles are caused by the complex form of the HJB equation and the existence of constraints. Consequently, we use the dual control theory and SMP to derive the bound estimate

approach that generates a lower bound $\underline{L}$ and an upper bound $\overline{L}$ of the value function. If $\underline{L}$ and $\overline{L}$ are close enough, then they are helpful in estimating the value function. In section 2.3.3, we show empirically that the interval $[\underline{L}, \overline{L}]$ is narrow enough. Therefore, the midpoint $\frac{\underline{L} + \overline{L}}{2}$ is a good estimate of the value function.

2.3.1 Dual Control Problems

The first step of deriving bound estimate is formally giving the definition of dual functions.

Definition 2.3.1. The dual function of a continuous, increasing and concave utility function $U : \mathbb{R}^+ \rightarrow \mathbb{R}$ with $U(0) = 0$ is $\tilde{U} : \mathbb{R}^+ \rightarrow \mathbb{R}$ given by

$$\tilde{U}(y) = \sup_{x \in \mathbb{R}^+} [U(x) - xy], y \in \mathbb{R}^+.$$

Obviously, the dual function is continuous, increasing but convex. Below is an example that gives the dual function of the power utility function.

Corollary 2.3.2. *The dual function of the power utility function*

$$U(x) = \frac{1}{\beta} x^\beta, 0 < \beta < 1, \quad (2.3.2)$$

is given by

$$\tilde{U}(y) = -\frac{\beta - 1}{\beta} y^{\frac{\beta}{\beta - 1}}.$$

Proof: Let $\gamma_y : \mathbb{R}^+ \rightarrow \mathbb{R}$ be a function such that $\gamma_y(x) = U(x) - xy, x, y \in \mathbb{R}^+$. The first order derivative of γ is

$$\gamma'_y(x) = U'(x) - y = x^{\beta - 1} - y$$

whose value is 0 at $x^* = y^{\frac{1}{\beta - 1}}$. As a result, the supremum of γ_y is

$$\tilde{U}(y) = \sup_{x \in \mathbb{R}^+} \gamma_y(x) = \gamma_y(x^*) = U(x^*) - x^*y = -\frac{\beta - 1}{\beta} y^{\frac{\beta}{\beta - 1}}.$$

According to the dual control theory, the dual process of X^π (denoted by $Y^{(y,v)}$) is driven by SDE

$$\begin{cases} dY^{(y,v)}(t) = -Y^{(y,v)}(t) \left\{ [r(t) + \delta_K(v(t))] dt + \frac{1}{\sigma(t)} \left(\hat{H}(t) - r(t) + v(t) \right) d\hat{V}(t) \right\}, \\ Y^{(y,v)}(0) = y, \end{cases}$$

where $\hat{H}$ is the 1D Kalman-Bucy filter driven by (2.2.3), δ_K is the support function of set $-K$:

$$\delta_K(v) = \sup_{\pi \in K} (-\pi v), v \in \mathbb{R},$$

and v is the dual control process whose admissible values all stay in set

$$\mathcal{B} = \left\{ \zeta : \Omega \times [0, T] \rightarrow \mathbb{R} : \zeta \text{ is } \mathcal{F}^S - \text{measurable, and } \int_0^T [\delta_K(\zeta(s)) + |\zeta(s)|^2] ds < \infty \text{ almost surely} \right\}.$$

The corresponding dual value function is

$$\tilde{V}(t, x, h) = \inf_{y \in \mathbb{R}^+, v \in \mathcal{B}} \left\{ xy + \mathbb{E}_{t,y,h} \left[\tilde{U} \left(Y^{(y,v)}(T) \right) \right] \right\}.$$

Theorem 2.3.3. *Let $\tilde{K} = \{v : \pi v \geq 0, \forall \pi \in K\}$ be the positive polar cone of K . Then for both the unconstrained case ($K = \mathbb{R}$) and the constrained case ($K = \mathbb{R}^+$), we have*

$$\delta_K(v) = \begin{cases} 0, & \text{if } v \in \tilde{K} \\ \infty, & \text{if } v \in K/\tilde{K} \end{cases}. \quad (2.3.3)$$

Proof: When $K = \mathbb{R}$, we have $\tilde{K} = \{0\}$. Then for any $\pi \in K$, we have

$$\begin{cases} v \in \tilde{K} \implies -\pi v = 0 \implies \delta_K(v) = 0, \\ v \in K/\tilde{K} \implies \text{for any } M > 0 \text{ there exists } \pi \in K \text{ such that } -\pi v > M \implies \delta_K(v) = \infty, \end{cases}$$

meaning that (2.3.3) holds for the unconstrained case. Alternatively, when $K = \mathbb{R}^+$, we have $\tilde{K} = \mathbb{R}^+$. Then for any $\pi \in K$, we have

$$\begin{cases} v \in \tilde{K} \implies -\pi v \leq 0 \implies \delta_K(v) = 0, \\ v \in K/\tilde{K} \implies \text{for any } M > 0 \text{ there exists } \pi \in K \text{ such that } -\pi v > M \implies \delta_K(v) = \infty, \end{cases}$$

meaning that (2.3.3) also holds for the constrained case.

A dual control v that takes ∞ does not make sense, so we always have $v(t) \in \tilde{K}$ for any $t \in [0, T]$. This fact makes some seemingly complex formulae simple, as is shown by numerical experiments in section 2.3.3.

2.3.2 Stochastic Maximum Principle

SMP gives the necessary and sufficient optimality conditions for both the primal and dual problems by the solution to a system of SDEs. Although in general, these SDEs cannot be solved analytically, SMP makes it possible to solve the problem by Monte Carlo and optimization algorithms. Details of SMP can be found in many sources such as [11].

Results given by SMP are summarized in theorem 2.3.4). We use the bound estimate approach designed in [18] to conduct numerical experiments, so that we can reproduce what [18] obtains.

Theorem 2.3.4. *Let $\tilde{y} \in \mathbb{R}^+$ and $\tilde{v} \in \mathcal{B}$. Consider the system of SDE*

$$\begin{cases} dY^{(\tilde{y}, \tilde{v})}(t) = -Y^{(\tilde{y}, \tilde{v})}(t) \left\{ [r(t) + \delta_K(\tilde{v}(t))] dt + \left[\frac{1}{\sigma(t)} (\hat{H}(t) - r(t) + \tilde{v}(t)) \right] d\hat{V}(t) \right\}, \\ d\hat{P}(t) = \left[r(t)\hat{P}(t) + \frac{1}{\sigma(t)} \hat{Q}(t) (\hat{H}(t) - r(t)) \right] dt + \hat{Q}(t) d\hat{V}(t), \\ d\hat{H}(t) = -\lambda (\hat{H}(t) - \bar{H}) dt + \left(\frac{\Sigma(t)}{\sigma(t)} + \rho_{\sigma H} \right) d\hat{V}(t), \end{cases} \quad (2.3.4)$$

with initial or terminal conditions $Y^{(\tilde{y}, \tilde{v})}(0) = \tilde{y}$, $\hat{P}(T) = -\tilde{U}'(Y^{(\tilde{y}, \tilde{v})}(T))$, $\hat{H}(0) = h_0$. Then $(\tilde{y}, \tilde{v})$ is the optimal control for the dual control problem if and only if the solution to (2.3.4) (denoted by $(Y^{(\tilde{y}, \tilde{v})}, \hat{P}, \hat{Q})$) satisfies

- *condition 1: the initial value of $\hat{P}$ must match with X :*

$$\hat{P}(0) = X(0) = x_0, \quad (2.3.5)$$

- *condition 2: the optimal trading strategy (i.e. π^*) of the primal problem can be shown to be given by (2.3.6) and must obey constraints K :*

$$\pi^*(t) = \frac{\hat{Q}(t)}{\hat{P}(t)\sigma(t)} \in K, \quad (2.3.6)$$

- condition 3: for any $t \in [0, T]$, (2.3.7) holds almost surely:

$$\hat{P}(t)\delta_K(\tilde{v}(t)) + \frac{\hat{Q}(t)\tilde{v}(t)}{\sigma(t)} = 0, \quad \mathbb{P} - a.s. \quad (2.3.7)$$

Proof: see theorem 3.9 and 3.10 in [11].

Condition 1 (i.e. (2.3.5)) says that $\hat{P}$ and X match at $t = 0$. And by condition 2 (i.e. (2.3.6)), the SDE of $\hat{P}$ in (2.3.4) can be written as

$$\begin{aligned} d\hat{P}(t) &= \left[r(t)\hat{P}(t) + \frac{1}{\sigma(t)}\hat{Q}(t) \left(\hat{H}(t) - r(t) \right) \right] dt + \hat{Q}(t)d\hat{V}(t) \\ &= \left[r(t)\hat{P}(t) + \pi^*(t)\hat{P}(t) \left(\hat{H}(t) - r(t) \right) \right] dt + \pi^*(t)\sigma(t)\hat{P}(t)d\hat{V}(t) \\ &= \hat{P}(t) \left\{ \left[r(t) + \pi^*(t) \left(\hat{H}(t) - r(t) \right) \right] dt + \pi^*(t)\sigma(t)d\hat{V}(t) \right\}. \end{aligned} \quad (2.3.8)$$

which matches the SDE of X^π in (2.2.5) (with $\pi = \pi^*$). So process $\hat{P}$ matches the optimal state process X^* . We omit the ‘ $\sim$ ’ notation over y and v next to simplify notations. The weak duality relation implies that

$$\sup_{\pi \in \mathcal{A}} \mathbb{E}_{t,x,h} [U(X^\pi(T))] \leq \inf_{y \in \mathbb{R}^+, v \in \mathcal{B}} \left\{ x_0 y + \mathbb{E}_{t,y,h} \left[\tilde{U} \left(Y^{(y,v)}(T) \right) \right] \right\}$$

By the definition of supremum and infimum, we have

$$\sup_{\pi \in \mathcal{A}} \mathbb{E} [U(X^\pi(T))] \geq \mathbb{E} [U(X^\pi(T))],$$

and

$$\inf_{y \in \mathbb{R}^+, v \in \mathcal{B}} \left\{ x_0 y + \mathbb{E} \left[\tilde{U} \left(Y^{(y,v)}(T) \right) \right] \right\} \leq x_0 y + \mathbb{E} \left[\tilde{U} \left(Y^{(y,v)}(T) \right) \right].$$

So the value function has a lower bound $\underline{L}$ and an upper bound $\bar{L}$, where

$$\underline{L} = \mathbb{E} [U(X^\pi(T))] \quad \text{and} \quad \bar{L} = x_0 y + \mathbb{E} \left[\tilde{U} \left(Y^{(y,v)}(T) \right) \right]. \quad (2.3.9)$$

Note that two requirements must be satisfied (given by theorem 2.3.4):

$$\begin{cases} \hat{P}(T) = -\tilde{U}'(Y^{(y,v)}(T)) \\ \hat{P}(t)\delta_K(v(t)) + \frac{\hat{Q}(t)v(t)}{\sigma(t)} = 0, \implies \delta_K(v(t)) + \pi^*(t)v(t) = 0, \quad \mathbb{P} - a.s. \end{cases}$$

To avoid handling the ‘almost sure’ claim, we take integral for the second requirement and the integral must be 0. In this way, we obtain the combined requirement:

$$\mathbb{E} \left[w \left| \hat{P}(T) + \tilde{U}'(Y^{(y,v)}(T)) \right|^2 + (1-w) \int_0^T [\delta_K(v(t)) + \pi^*(t)v(t)] dt \right] = 0, \quad (2.3.10)$$

where $w \in (0, 1)$ is a given constant that measures the relative importance of these two requirements. Therefore, the goal is to find values of $y, \pi(t), v(t), t \in [0, T]$ such that terminal values $\hat{P}(T)$ and $Y(T)$ generated by system (2.3.4) satisfy (2.3.10). We have shown that $\delta_K(v(t)) \in \{0, \infty\}$ and $v(t) \in \tilde{K} \implies \pi(t)v(t) \geq 0$ for any $\pi \in K$, so the integral in (2.3.10) is nonnegative. The problem is then solved by minimizing the objective function

$$f(y, \pi, v) = \mathbb{E} \left[w \left| \hat{P}(T) + \tilde{U}'(Y^{(y,v)}(T)) \right|^2 + (1-w) \int_0^T [\delta_K(v(t)) + \pi(t)v(t)] dt \right]. \quad (2.3.11)$$

After obtaining y, π and v , bounds $\underline{L}$ and $\bar{L}$ in (2.3.9) can be computed.

2.3.3 Numerical Experiments

We introduce the numerical experiment in this section. We first assign values to all variables involved in this problem. These numbers are the same as those used in [18]: $r = 0.05, \sigma = 0.8, x_0 = 10.0, h_0 = 0.1, \bar{H} = 0.1, \sigma_H = 0.5, \lambda = 1.0, \rho = 0.0, \Sigma_0 = 0.2, T = 1.0$, where $r(t)$ and $\sigma(t)$ are assumed to be constants $r \in \mathbb{R}$ and $\sigma \in \mathbb{R}$ for all $t \in [0, T]$. The utility function U is set to be the power utility function (2.3.2) with $\beta = 0.5$.

SDEs in system (2.3.4) are approximated by discretization. The investment horizon $[0, T]$ is divided into N pieces with equal length:

$$[t_i, t_{i+1}] = [ih, (i+1)h], i = 0, 1, \dots, N-1,$$

where the step size is $h = \frac{T}{N}$. Then the approximation is given by

$$\begin{cases} Y_{i+1} = Y_i - rY_i h - Y_i \left[\frac{1}{\sigma} (\hat{H}_i - r + v_i) \right] (\hat{V}_{i+1} - \hat{V}_i), & Y_0 = y, \\ \hat{P}_{i+1} = \hat{P}_i + r\hat{P}_i h + \hat{P}_i \pi_i (\hat{H}_i - r) h + \hat{P}_i \pi_i \sigma (\hat{V}_{i+1} - \hat{V}_i), & \hat{P}_0 = x_0, \\ \hat{H}_{i+1} = \hat{H}_i + \lambda (\bar{H} - \hat{H}_i) h + \left[\frac{1}{\sigma} (\Sigma_i + \rho \sigma \sigma_H) \right] (\hat{V}_{i+1} - \hat{V}_i), & \hat{H}_0 = h_0, \end{cases}$$

for $i = 0, 1, \dots, N-1$, where $\{Y_i, \hat{P}_i, \hat{H}_i\}_{i=0}^N$ are approximations of $\{Y(ih), \hat{P}(ih), \hat{H}(ih)\}_{i=0}^N$, and the innovation term is 1D Gaussian:

$$\hat{V}_{i+1} - \hat{V}_i = \sqrt{h} \varepsilon.$$

We choose $N = 10$ as [18] suggests. The objective function is thus approximated by

$$\begin{aligned} \hat{f}(y, \{\pi\}_{i=0}^{N-1}, \{v\}_{i=0}^{N-1}) &= \mathbb{E} \left[w \left| \hat{P}_N + \tilde{U}'(Y_N) \right|^2 + (1-w)h \sum_{i=0}^{N-1} [\delta_K(v_i) + \pi_i v_i] \right] \\ &= \mathbb{E} \left[w \left| \hat{P}_N + \tilde{U}'(Y_N) \right|^2 + (1-w)h \sum_{i=0}^{N-1} \pi_i v_i \right], \end{aligned}$$

where we use the fact that $\delta_K(v_i) = 0$ for $v_i \in \tilde{K}$. Indeed, $v_i \in \tilde{K}$ must hold because otherwise $\delta_K(v_i) = \infty$ and then the minimization of $\hat{f}$ is impossible. The expectation above, and bounds $\underline{L}, \bar{L}$ are all estimated by Monte Carlo methods. When $\hat{f}$ is evaluated, $M_1 = 100$ sample paths for (2.3.4) are simulated. This number cannot be too large, although a large number of sample paths make the estimate of $\hat{f}$ accuracy. This is because running the optimization algorithm for $\hat{f}$ triggers the simulation process many times, which is extremely time-consuming when M_1 is too large. In contrast, we set the number of sample paths $M_2 = 100000$ when evaluating $\underline{L}$ and $\bar{L}$ through (2.3.9). Using a large M_2 causes no trouble since the simulation process is only triggered once.

Unconstrained Problems

For unconstrained problems we have $K = \mathbb{R}$ and $\tilde{K} = \{0\}$. Since $v_i \in \tilde{K}$ holds all the time, we must have $v_i = 0$ and $\delta_K(v_i) = 0$ for any index i . So the second part of $\hat{f}$ can be removed and thus $\hat{f}$ reduces to

$$\hat{f}(y, \{\pi_i\}_{i=0}^{N-1}) = \mathbb{E} \left[\left| \hat{P}_N + \tilde{U}'(Y_N) \right|^2 \right].$$

Two forms of parameterization are tested for π . Form I uses $\pi_i = a^\pi + b^\pi \hat{H}_i$, where $a^\pi, b^\pi \in \mathbb{R}$ are unknown constants that we should determine. Form II allows a^π to vary

	Number of Parameters	Details of Parameters
Form I	3	y, a^π, b^π
Form II	$N + 2$	$y, \{a_i^\pi\}_{i=0}^{N-1}, b^\pi$

Table 2.1: Parameters to determine for 1D unconstrained problems under the Kalman-Bucy filter

over time and thus the trading strategy is $\pi_i = a_i^\pi + b^\pi \hat{H}_i$. Therefore, we have 3 and $N + 2$ parameters to determine for form I and form II parameterization respectively (see Table 2.1).

As have been discussed before, we prefer cases where $\underline{L}$ and $\bar{L}$ are close to each other. So we compute their relative error (called Error-1):

$$\text{Error} - 1 = \left| \frac{\bar{L} - \underline{L}}{\underline{L}} \right|.$$

Since we know that the optimal value lies between $\underline{L}$ and $\bar{L}$, it is reasonable to estimate the optimal value by their average (called MV):

$$\text{MV} = \frac{\underline{L} + \bar{L}}{2}.$$

According to [18], the unconstrained problems have close-form solution given by

$$J(0, x_0, h_0) = U(x_0) \exp(A(0)h_0^2 + B(0)h_0 + C(0)),$$

where the ODEs for functions A, B, C and their numerical solutions are given in appendix C of [18]. The closed-form solutions serve as benchmarks (called BC) to measure if MV is accurate enough, and thus the relative error between MV and BC (called Error-2) is computed

$$\text{Error} - 2 = \left| \frac{\text{MV} - \text{BC}}{\text{BC}} \right|.$$

We also record the optimal value of the objective function $\hat{f}$ (called OB).

To see how results change as the prior distribution changes, we repeat the work for six different combinations. To be specific, let $h_0^- = 0.05, h_0^\circ = 0.1, h_0^+ = 0.2$ and $\Sigma_0^- = 0.1, \Sigma_0^\circ = 0.2, \Sigma_0^+ = 0.3$. We first fix $\Sigma_0 = \Sigma_0^\circ$ and run the algorithm for $h_0 = h_0^-, h_0^\circ, h_0^+$ one by one (see Table 2.2). Then we fix $h_0 = h_0^\circ$ and run the algorithm for $\Sigma_0 = \Sigma_0^-, \Sigma_0^\circ, \Sigma_0^+$ one by one (see Table 2.3).

Several key points can be observed from these results. First, BC always lies between $\underline{L}$ and $\bar{L}$, indicating that the bound estimate is effective. Second, almost all Error-1 values are less than 3%, so the interval $[\underline{L}, \bar{L}]$ is narrow. Additionally, Error-2 values are all less than 1%, which indicates that MV estimates BC accurately. Finally, MV increases with h_0 and Σ_0 .

Constrained Problems

For constrained problems we have $K = \mathbb{R}^+$ and $\tilde{K} = \mathbb{R}^+$, so $\hat{f}$ cannot reduce anymore, so both π and v need to be parameterized. We have to ensure that both $\{\pi_i\}_{i=0}^{N-1}$ and $\{v_i\}_{i=0}^{N-1}$ lie in $\mathbb{R}^+$. So form I parameterization is given by

$$\begin{cases} \pi_i = \left(a^\pi + b^\pi \hat{H}_i \right)^+, \\ v_i = \left(a^v + b^v \hat{H}_i \right)^+, \end{cases}$$

h_0	Form	$\underline{L}$	$\overline{L}$	Error-1	BC	MV	Error-2	OB
h_0^-	I	6.4756	6.5599	1.3012%	6.5521	6.5178	0.5242%	0.0010
h_0^-	II	6.4662	6.5637	1.5092%	6.5521	6.5150	0.5670%	0.0008
h_0°	I	6.4770	6.5736	1.4909%	6.5645	6.5253	0.5969%	0.0008
h_0°	II	6.4760	6.5841	1.6693%	6.5645	6.5300	0.5248%	0.0008
h_0^+	I	6.4736	6.6698	3.0300%	6.6300	6.5717	0.8791%	0.0015
h_0^+	II	6.4749	6.6698	3.0109%	6.6300	6.5724	0.8695%	0.0015

Table 2.2: Numerical results for 1D unconstrained problems by varying h_0 under the Kalman-Bucy filter

Σ_0	Form	$\underline{L}$	$\overline{L}$	Error-1	BC	MV	Error-2	OB
Σ_0^-	I	6.4822	6.5359	0.8284%	6.5353	6.5075	0.4261%	0.0003
Σ_0^-	II	6.4746	6.5382	0.9823%	6.5353	6.5031	0.4922%	0.0005
Σ_0°	I	6.4772	6.6054	1.9784%	6.5645	6.5413	0.3531%	0.0011
Σ_0°	II	6.4752	6.5961	1.8659%	6.5645	6.5356	0.4395%	0.0010
Σ_0^+	I	6.4756	6.6598	2.8448%	6.6050	6.5677	0.5647%	0.0014
Σ_0^+	II	6.4801	6.6711	2.9474%	6.6050	6.5756	0.4458%	0.0011

Table 2.3: Numerical results for 1D unconstrained problems by varying Σ_0 under the Kalman-Bucy filter

where $a^\pi, a^v, b^\pi, b^v \in \mathbb{R}$ are unknown constants that we should determine. Form II parameterization allows both a^π and a^v to vary over time and thus we have

$$\begin{cases} \pi_i = \left(a_i^\pi + b^\pi \hat{H}_i \right)^+, \\ v_i = \left(a_i^v + b^v \hat{H}_i \right)^+. \end{cases}$$

Therefore, we have 5 and $2N+3$ parameters to determine for form I and form II respectively (see Table 2.4).

The metrics that we compute for constrained problems are almost the same as those computed for unconstrained problems, except that BC and Error-2 are removed because no closed-form solutions are available for constrained problems. Since we only have Error-1 for this case, it is called Error now. In this part, parameter w in (2.3.11) is set to be $w = w^\circ = 0.5$. Results are listed in Tables 2.5 and 2.6. Key points observed for unconstrained problems can also be observed for constrained problems.

The effect of varying w is tested next. Let $w^- = 0.1, w^\circ = 0.5, w^+ = 0.9$. We fix $h_0 = h_0^\circ, \Sigma_0 = \Sigma_0^\circ$, and run the bound estimate approach for $w = w^-, w^\circ, w^+$ one by one. Results are listed in Table 2.7. These results show that both the optimal value (MV) and the relative error (Error) are not sensitive to w and no optimal w can be determined.

	Number of Parameters	Details of Parameters
Form I	5	$y, a^\pi, a^v, b^\pi, b^v$
Form II	$2N + 3$	$y, \{a_i^\pi\}_{i=0}^{N-1}, \{a_i^v\}_{i=0}^{N-1}, b^\pi, b^v$

Table 2.4: Parameters to determine for 1D constrained problems under the Kalman-Bucy filter

h_0	Form	$\underline{L}$	$\overline{L}$	Error	OB	MV
h_0^-	I	6.4843	6.5629	1.1983%	0.0003	6.5236
h_0^-	II	6.4824	6.5631	1.2301%	0.0003	6.5228
h_0°	I	6.4843	6.5837	1.5096%	0.0004	6.5340
h_0°	II	6.4867	6.5875	1.5296%	0.0008	6.5371
h_0^+	I	6.4985	6.6505	2.2866%	0.0007	6.5745
h_0^+	II	6.4920	6.6620	2.5517%	0.0006	6.5770

Table 2.5: Numerical results for 1D constrained problems by varying h_0 under the Kalman-Bucy filter

Σ_0	Form	$\underline{L}$	$\overline{L}$	Error	OB	MV
Σ_0^-	I	6.4843	6.5286	0.6782%	0.0002	6.5064
Σ_0^-	II	6.4846	6.5304	0.7019%	0.0004	6.5075
Σ_0°	I	6.4843	6.5863	1.5484%	0.0004	6.5353
Σ_0°	II	6.4854	6.5852	1.5160%	0.0007	6.5353
Σ_0^+	I	6.4813	6.6532	2.5838%	0.0008	6.5672
Σ_0^+	II	6.4893	6.6544	2.4816%	0.0008	6.5718

Table 2.6: Numerical results for 1D constrained problems by varying Σ_0 under the Kalman-Bucy filter

w	Form	$\underline{L}$	$\overline{L}$	Error	OB	MV
w^-	I	6.4843	6.5680	1.2738%	0.0001	6.5261
w^-	II	6.4873	6.6514	2.4671%	0.0002	6.5694
w°	I	6.4843	6.5837	1.5096%	0.0004	6.5340
w°	II	6.4867	6.5875	1.5296%	0.0008	6.5371
w^+	I	6.4852	6.5731	1.3370%	0.0012	6.5292
w^+	II	6.4860	6.6047	1.7976%	0.0012	6.5453

Table 2.7: Numerical results for 1D constrained problems by varying w under the Kalman-Bucy filter

2.4 The Deep Learning Approach

Differential equations like ODEs and PDEs are equations for functions and their derivatives. Due to the fact that neural networks can approximate functions and their derivatives accurately if trained in a correct way, we can solve many widely used differential equations by deep learning techniques.

We show how to solve problem (2.2.5) by PINN in this section. The main task is using a neural network to approximate the solution to HJB equation (2.3.1). Before the PINN algorithm can be implemented, we must handle the supremum in it. Taking out the terms related to $\pi(t)$ in the supremum, we obtain a quadratic function of $\pi(t)$:

$$S(\pi(t)) = \left(\frac{1}{2} x^2 \sigma^2(t) \frac{\partial^2 J}{\partial x^2} \right) \pi^2(t) + \left[x(h - r(t)) \frac{\partial J}{\partial x} + x(\Sigma(t) + \rho \sigma_H \sigma(t)) \frac{\partial^2 J}{\partial x \partial h} \right] \pi(t),$$

where the constraint $\pi(t) \in K$ must hold. When $K = \mathbb{R}$ (unconstrained problems), S achieves its supremum at

$$\pi^*(t) = - \frac{(h - r(t)) \frac{\partial J}{\partial x} + (\Sigma(t) + \rho \sigma_H \sigma(t)) \frac{\partial^2 J}{\partial x \partial h}}{x \sigma^2(t) \frac{\partial^2 J}{\partial x^2}},$$

which is the optimal trading strategy. However, when there exist short selling constraints, the optimal trading strategy is given by

$$\pi^*(t) = \left(- \frac{(h - r(t)) \frac{\partial J}{\partial x} + (\Sigma(t) + \rho \sigma_H \sigma(t)) \frac{\partial^2 J}{\partial x \partial h}}{x \sigma^2(t) \frac{\partial^2 J}{\partial x^2}}, 0 \right)^+$$

instead. In this way, the HJB equation becomes a standard PDE that is ready to be solved by PINN:

$$\begin{aligned} \frac{\partial J}{\partial t} + x[r(t) + \pi^*(t)(h - r(t))] \frac{\partial J}{\partial x} + \frac{1}{2} x^2 \pi^{*2}(t) \sigma^2(t) \frac{\partial^2 J}{\partial x^2} - \lambda(h - \bar{H}) \frac{\partial J}{\partial h} + \\ x \pi^*(t) \sigma(t) \left(\frac{\Sigma(t)}{\sigma(t)} + \rho \sigma_H \right) \frac{\partial^2 J}{\partial x \partial h} + \frac{1}{2} \left(\frac{\Sigma(t)}{\sigma(t)} + \rho \sigma_H \right)^2 \frac{\partial^2 J}{\partial h^2} = 0, \end{aligned} \quad (2.4.1)$$

with terminal condition

$$J(T, x, h) = U(x).$$

Variables here are all the same as settings in numerical experiments for the bound estimate approach. The range of t is $t \in [0, T]$. To efficiently solve (2.4.1) by PINN, we need to specify the ranges of x and h . These ranges are denoted by $x \in [X_{\min}, X_{\max}]$ and $h \in [0, H_{\max}]$, so that the solution region is $\mathcal{R} = [0, T] \times [X_{\min}, X_{\max}] \times [0, H_{\max}]$. Values of $X_{\min}$, $X_{\max}$ and $H_{\max}$ must be properly selected. If they lead to a region $\mathcal{R}$ that is too large, then running the algorithm can be extremely time-consuming, because achieving high accuracy requires large sample size and many iterations. But $\mathcal{R}$ cannot be too small either. For example, suppose that we have an investor whose initial wealth $X^*(0) = x_0$, and the initial value of the filter is $\hat{H}(0) = h_0$. At $t = 0$, we care about the value function $J(0, x_0, h_0)$, meaning that $(0, x_0, h_0)$ must stay in $\mathcal{R}$. With all these issues considered, we choose $X_{\min} = 0.7x_0$, $X_{\max} = 1.3x_0$ and $H_{\max} = 0.25$.

Let $\hat{J}_\theta$ be the neural network that approximates J , where θ represents all parameters of the neural network, including weights and biases. The structure of $\hat{J}_\theta$ is shown in Table 2.8, where all hidden layers have 16 neurons and the activation function between adjacent layers is

$$\tanh(x) = \frac{e^x - e^{-x}}{e^x + e^{-x}}.$$

Layer Type	Input Size	Output Size
Linear	3	16
Tanh Activation	—	—
Linear	16	16
Tanh Activation	—	—
Linear	16	16
Tanh Activation	—	—
Linear	16	1

Table 2.8: Layers of the neural network $\hat{J}_\theta$

The input size of $\hat{J}_\theta$ is 3, (i.e. (t, x, h)), while the output size is 1 (the value function at (t, x, h)). Note that the activation function cannot be ReLU or leaky ReLU because they are piecewise linear and thus become 0 after taking the second order derivative. Due to advantages of automatic differentiation that has been implemented in software designed for deep learning models, derivatives of $\hat{J}_\theta$ can be immediately computed and be used to approximate derivatives of J .

The loss function is constructed by sampling data points at the boundary and interior of $\mathcal{R}$. We sample a dataset $\mathcal{D}_1 = \{T, x_i, h_i\}_{i=1}^B$ for the terminal condition and then this part of loss is given by

$$\mathcal{L}_1(\theta) = \frac{1}{B} \sum_{i=1}^B \left(\hat{J}_\theta(T, x_i, h_i) - U(x_i) \right)^2. \quad (2.4.2)$$

As for the interior of $\mathcal{R}$, we sample a dataset $\mathcal{D}_2 = \{\bar{t}_i, \bar{x}_i, \bar{h}_i\}_{i=1}^B$ and compute the corresponding loss $\mathcal{L}_2$:

$$\begin{aligned} \mathcal{L}_2(\theta) = \frac{1}{B} \sum_{i=1}^B & \left[\frac{\partial \hat{J}_i}{\partial t} + \bar{x}_i [r + \pi_i^* (h - r)] \frac{\partial \hat{J}_i}{\partial x} + \frac{1}{2} \bar{x}_i^2 \pi_i^{*2} \sigma^2 \frac{\partial^2 \hat{J}_i}{\partial x^2} - \lambda (h - \bar{H}) \frac{\partial \hat{J}_i}{\partial h} + \right. \\ & \left. \bar{x}_i \pi_i^* \sigma \left(\frac{\Sigma(\bar{t}_i)}{\sigma} + \rho \sigma_H \right) \frac{\partial^2 \hat{J}_i}{\partial x \partial h} + \frac{1}{2} \left(\frac{\Sigma(\bar{t}_i)}{\sigma} + \rho \sigma_H \right)^2 \frac{\partial^2 \hat{J}_i}{\partial h^2} \right]^2, \end{aligned} \quad (2.4.3)$$

where $\hat{J}_i = \hat{J}_\theta(\bar{t}_i, \bar{x}_i, \bar{h}_i)$, and $\frac{\partial \hat{J}_i}{\partial t}$ means the value of $\frac{\partial \hat{J}}{\partial t}$ evaluated at $(\bar{t}_i, \bar{x}_i, \bar{h}_i)$. Moreover, π_i^* in (2.4.3) means the value of π^* evaluated at $(\bar{t}_i, \bar{x}_i, \bar{h}_i)$. The total loss $\mathcal{L}$ is thus the sum of $\mathcal{L}_1$ and $\mathcal{L}_2$. In the sampling stage, all sampled data points follow uniform distribution in their corresponding ranges.

When training the neural network, we set $N = 100000$, $B = 2048$, $\eta = 5 \times 10^{-4}$. Both unconstrained problems and constrained problems are solved by running algo. 1. Training loss curves are shown in Figures 2.1 (unconstrained) and 2.2 (constrained) respectively. The loss becomes lower and lower during the training stage, which is what we expect.

Curves $J(0, x_0, h_0)$ (with $x_0 = 10.0$) for $\Sigma_0 = 0.1, 0.2, 0.3$ are shown in Figures 2.3 (unconstrained) and 2.4 (constrained) respectively. These curves show that the value function $J(0, x_0, h_0)$ increases with h_0 (because each single curve is an increasing function) and σ (because curves of $\sigma_0 = 0.1(k+1)$ are always above curves of $\sigma_0 = 0.1k$, $k = 1, 2$).

To compare results of the bound estimate approach and the deep learning approach, we add optimal values computed by PINN to Tables 2.2 ~ 2.6 and obtain Tables 2.9 ~ 2.12. For unconstrained problems, the error of PINN (called Error-PINN) is measured by the relative error between PINN and BC:

$$\text{Error - PINN} = \left| \frac{\text{PINN} - \text{BC}}{\text{BC}} \right|.$$

However, since BC is not available for constrained problems, we define Error–PINN as the relative error between PINN and MV for constrained problems:

$$\text{Error – PINN} = \left| \frac{\text{PINN} - \text{MV}}{\text{MV}} \right|.$$

It can be observed that PINN incurs small errors (for most cases Error–PINN is less than 1%) and thus can accurately fit the value function J .

We also vary x and h at the same time (with $\Sigma_0 = 0.2$, $x \in [0.7x_0, 1.3x_0]$ and $h \in [0, h_{\max}]$) and draw surfaces $J(0, x, h)$ (see Figure 2.5). These surfaces show that larger initial wealth x and larger initial guess h (of the drift) lead to greater values. This is intuitively correct, because more wealth and higher positive drift of stock returns have a higher probability to generate greater values in the stochastic environment of the financial market.

Algo. 1 PINN algorithm for HJB equation (2.4.1)

Input:

$\theta^{(0)}$: initial value of the neural network parameter

N : the number of iterations

B : the sample size parameter

η : the learning rate

Output:

$\theta^{(N)}$: the estimated optimal neural network parameter

for iteration $m = 1, 2, \dots, N$ **do**

 generate datasets $\mathcal{D}_1 = \{T, x_i, h_i\}_{i=1}^B$, $\mathcal{D}_2 = \{\bar{t}_i, \bar{x}_i, \bar{h}_i\}_{i=1}^B$

 compute the total loss function $\mathcal{L}(\theta^{(m-1)})$ through (2.4.2) and (2.4.3)

 update $\theta^{(m-1)}$ through gradient descent method and obtain $\theta^{(m)}$

end for

return $\theta^{(N)}$

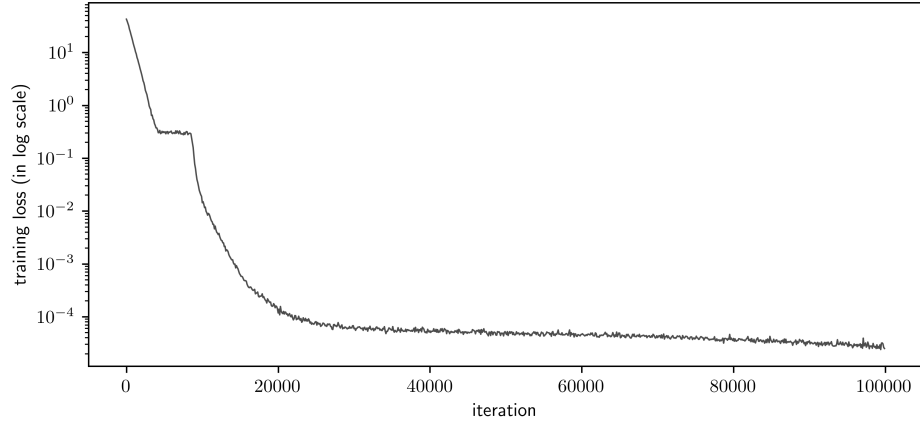


Figure 2.1: The training loss curve of solving the HJB equation with $h_0 = 0.1, \Sigma = 0.2$ under the Kalman-Bucy filter for unconstrained problems

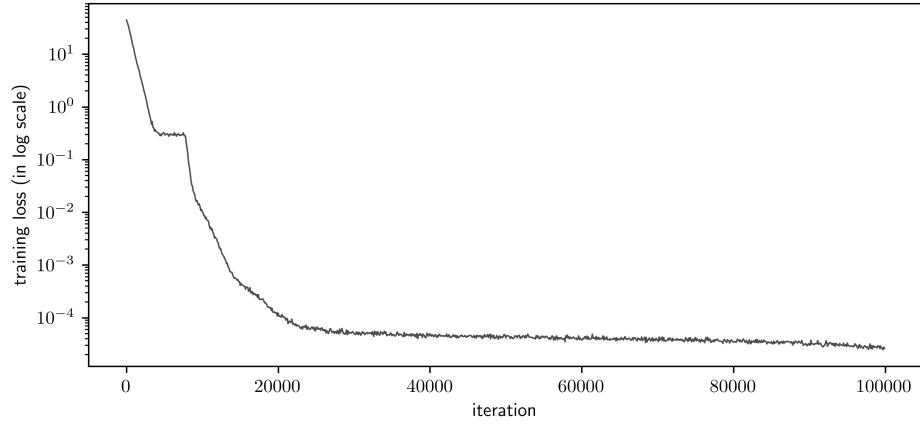


Figure 2.2: The training loss curve of solving the HJB equation with $h_0 = 0.1, \Sigma = 0.2$ under the Kalman-Bucy filter for unconstrained problems

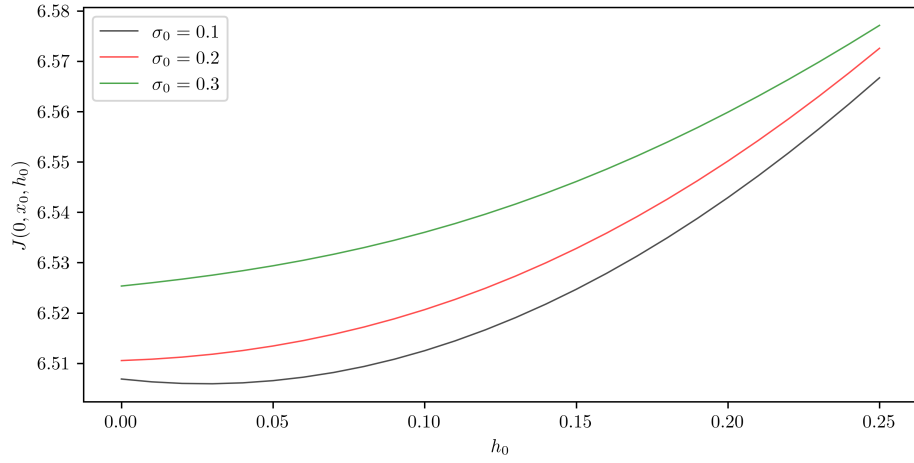


Figure 2.3: $J(0, x_0, h_0)$ curves by varying Σ_0 under the Kalman-Bucy filter for unconstrained problems

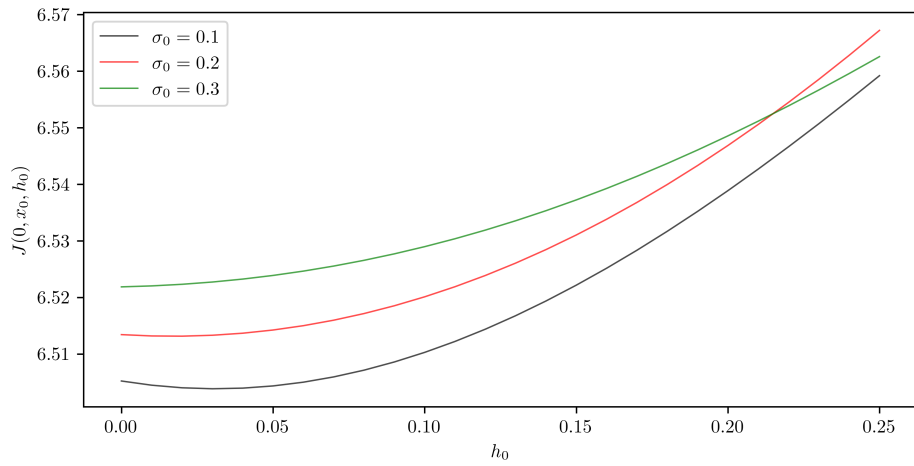


Figure 2.4: $J(0, x_0, h_0)$ curves by varying Σ_0 under the Kalman-Bucy filter for constrained problems

h_0	Form	LB	UB	Error-1	BC	MV	Error-2	OB	PINN	Error-PINN
h_0^-	I	6.4756	6.5599	1.3012%	6.5521	6.5178	0.5242%	0.0010	6.5187	0.5098%
h_0^-	II	6.4662	6.5637	1.5092%	6.5521	6.5150	0.5670%	0.0008	6.5187	0.5098%
h_0^0	I	6.4770	6.5736	1.4909%	6.5645	6.5253	0.5969%	0.0008	6.5232	0.6291%
h_0^0	II	6.4760	6.5841	1.6693%	6.5645	6.5300	0.5248%	0.0008	6.5232	0.6291%
h_0^+	I	6.4736	6.6698	3.0300%	6.6300	6.5717	0.8791%	0.0015	6.5444	1.2911%
h_0^+	II	6.4749	6.6698	3.0109%	6.6300	6.5724	0.8695%	0.0015	6.5444	1.2911%

Table 2.9: Numerical results for 1D unconstrained problems by varying h_0 under the Kalman-Bucy filter with PINN results included

σ_0	Form	LB	UB	Error-1	BC	MV	Error-2	OB	PINN	Error-PINN
Σ_0^-	I	6.4822	6.5327	0.7792%	6.5353	6.5075	0.4261%	0.0003	6.5155	0.3030%
Σ_0^-	II	6.4746	6.5317	0.8823%	6.5353	6.5031	0.4922%	0.0005	6.5155	0.3030%
Σ_0^0	I	6.4772	6.6054	1.9784%	6.5645	6.5413	0.3531%	0.0011	6.5232	0.6291%
Σ_0^0	II	6.4752	6.5961	1.8659%	6.5645	6.5356	0.4395%	0.0010	6.5232	0.6291%
Σ_0^+	I	6.4756	6.6598	2.8448%	6.6050	6.5677	0.5647%	0.0014	6.5489	0.8494%
Σ_0^+	II	6.4801	6.6711	2.9474%	6.6050	6.5756	0.4458%	0.0011	6.5489	0.8494%

Table 2.10: Numerical results for 1D unconstrained problems by varying Σ_0 under the Kalman-Bucy filter with PINN results included

h_0	Form	LB	UB	Error	OB	MV	PINN	Error-PINN
h_0^-	I	6.4843	6.5629	1.1983%	0.0003	6.5236	6.5111	0.1916%
h_0^-	II	6.4824	6.5631	1.2301%	0.0003	6.5228	6.5111	0.1794%
h_0^0	I	6.4843	6.5837	1.5096%	0.0004	6.5340	6.5191	0.2280%
h_0^0	II	6.4867	6.5875	1.5296%	0.0008	6.5371	6.5191	0.2754%
h_0^+	I	6.4985	6.6505	2.2866%	0.0007	6.5745	6.5451	0.4472%
h_0^+	II	6.4920	6.6620	2.5517%	0.0006	6.5770	6.5451	0.4850%

Table 2.11: Numerical results for 1D constrained problems by varying h_0 under the Kalman-Bucy filter with PINN results included

Σ_0	Form	LB	UB	Error	OB	MV	PINN	Error-PINN
Σ_0^-	I	6.4843	6.5286	0.6782%	0.0002	6.5064	6.5070	0.0092%
Σ_0^-	II	6.4846	6.5304	0.7019%	0.0004	6.5075	6.5070	0.0077%
Σ_0^0	I	6.4843	6.5863	1.5484%	0.0004	6.5353	6.5191	0.2479%
Σ_0^0	II	6.4854	6.5852	1.5160%	0.0007	6.5353	6.5191	0.2479%
Σ_0^+	I	6.4813	6.6532	2.5838%	0.0008	6.5672	6.5455	0.3304%
Σ_0^+	II	6.4893	6.6544	2.4816%	0.0008	6.5718	6.5455	0.4002%

Table 2.12: Numerical results for 1D constrained problems by varying Σ_0 under the Kalman-Bucy filter with PINN results included

optimal value function at $t = 0$, i.e. $J(0, x, h)$, given by PINN

optimal value function at $t = 0$, i.e. $J(0, x, h)$, given by PINN

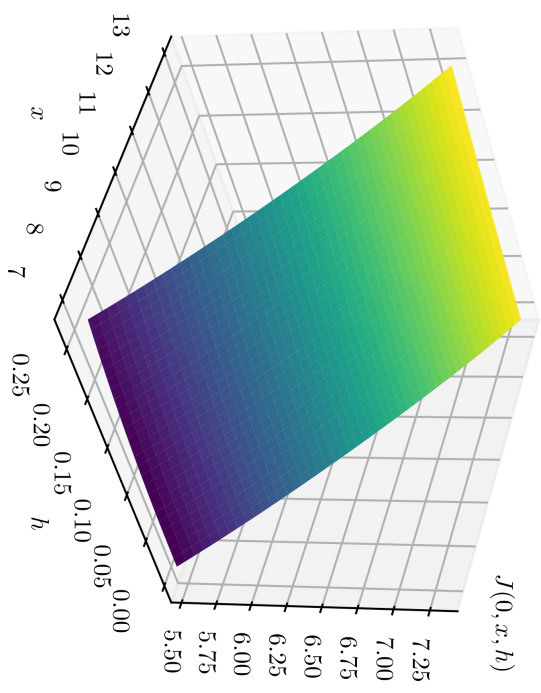
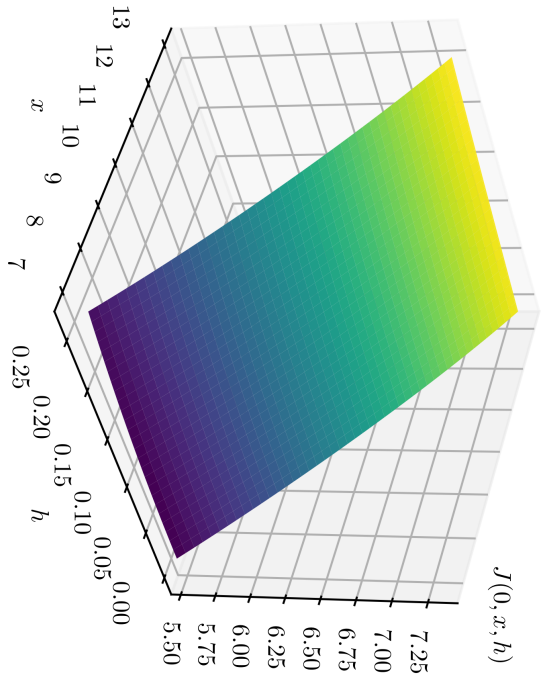


Figure 2.5: The value function surface at $t = 0$ with $\Sigma_0 = 0.2$ under the Kalman-Bucy filter
(left panel: unconstrained problems, right panel: constrained problems)

Chapter 3

2D Partially Observable Problems under Kalman-Bucy Filters

This chapter is an extension of chapter 2. Previously, only one stock exists in the financial market. However, in reality, many stocks can be traded. Therefore, we must consider high dimensional problems. Let D be the dimension, and thus D is also the total number of traded stocks. For our partially observable problems, there is no difference between the case where $D \geq 3$ and the case where $D = 2$. So it is enough to consider the 2D case, which is what this chapter mainly discusses. As what is done for 1D problems, theoretically, we give the HJB equation and the estimated bounds given by SMP. Numerical experiments follows these theoretical analysis to confirm all results. The bound estimate approach is tested.

3.1 Problem Formulation

When the investor trades two stocks, we obtain 2D partially observable utility maximization problems. Price of the risk-free bond still follows SDE (2.1.1). Stock prices follow SDE (3.1.1) which is in matrix-vector form:

$$d \begin{bmatrix} S^{(1)}(t) \\ S^{(2)}(t) \end{bmatrix} = \begin{bmatrix} S^{(1)}(t) & 0 \\ 0 & S^{(2)}(t) \end{bmatrix} \left(\begin{bmatrix} \mu^{(1)}(t) \\ \mu^{(2)}(t) \end{bmatrix} dt + \begin{bmatrix} \sigma^{(11)}(t) & \sigma^{(12)}(t) \\ \sigma^{(21)}(t) & \sigma^{(22)}(t) \end{bmatrix} d \begin{bmatrix} W^{(1)}(t) \\ W^{(2)}(t) \end{bmatrix} \right). \quad (3.1.1)$$

Or, in a compact form, (3.1.1) is written as

$$d\mathbf{S}(t) = \text{diag}(\mathbf{S}(t)) (\boldsymbol{\mu}(t)dt + \boldsymbol{\sigma}(t)d\mathbf{W}(t)), \quad (3.1.2)$$

where

$$\mathbf{S}(t) = \begin{bmatrix} S^{(1)}(t) \\ S^{(2)}(t) \end{bmatrix}, \boldsymbol{\mu}(t) = \begin{bmatrix} \mu^{(1)}(t) \\ \mu^{(2)}(t) \end{bmatrix}, \boldsymbol{\sigma}(t) = \begin{bmatrix} \sigma^{(11)}(t) & \sigma^{(12)}(t) \\ \sigma^{(21)}(t) & \sigma^{(22)}(t) \end{bmatrix}, \mathbf{W}(t) = \begin{bmatrix} W^{(1)}(t) \\ W^{(2)}(t) \end{bmatrix}.$$

Now we formulate the 2D partially observable problem. Let $\mathcal{F}^{\mathbf{S}}$ (the observable information) be the filtration generated by $\mathbf{S}$. The full information is denoted by filtration $\mathcal{F}$, and thus $\mathcal{F}^{\mathbf{S}} \subset \mathcal{F}$. Assume that the drift process $\boldsymbol{\mu}$ is $\mathcal{F}$ -measurable while r and $\boldsymbol{\sigma}$ are both $\mathcal{F}^{\mathbf{S}}$ -measurable. In this way, the problem becomes partially observable. 2D process $\{\boldsymbol{\pi}(t), t \in [0, T]\}$ stands for the self-financing trading strategy, where $\pi^{(d)}(t), d = 1, 2$ is the fraction of total wealth invested in stock d . There may be constraints on $\boldsymbol{\pi}$ which is denoted by set $K \subset \mathbb{R}^2$, i.e. we require $\boldsymbol{\pi}(t) \in K$ for any $t \in [0, T]$. When no constraint exists, K is the full 2D space $\mathbb{R}^2$. But when short selling is prohibited, $K = \mathbb{R}^{+2}$. The set of all admissible trading strategies is given by

$$\mathcal{A} = \{ \boldsymbol{\pi} : \boldsymbol{\pi} \in \mathcal{H}^2(0, T; \mathbb{R}^2) \text{ and } \boldsymbol{\pi}(t) \in K \text{ almost everywhere on } t \in [0, T] \},$$

where

$$\mathcal{H}^2(0, T; \mathbb{R}^2) = \left\{ \zeta : \Omega \times [0, T] \rightarrow \mathbb{R}^2 : \zeta \text{ is } \mathcal{F}^S \text{-measurable, and } \mathbb{E} \left[\int_0^T \|\zeta(s)\|^2 ds \right] < \infty \right\}.$$

The wealth process (when trading strategy π is executed) is denoted by $\{X^\pi(t), t \in [0, T]\}$, and is driven by SDE

$$\begin{aligned} dX^\pi(t) &= \pi^{(1)}(t) X^\pi(t) \frac{dS^{(1)}(t)}{S^{(1)}(t)} + \pi^{(2)}(t) X^\pi(t) \frac{dS^{(2)}(t)}{S^{(2)}(t)} + \left(1 - \pi^{(1)}(t) - \pi^{(2)}(t)\right) X^\pi(t) r(t) dt \\ &= X^\pi(t) \pi^\top(t) (\mu(t) dt + \sigma(t) d\mathbf{W}(t)) + X^\pi(t) \left(r(t) - \pi^\top(t) \begin{bmatrix} r(t) \\ r(t) \end{bmatrix} \right) dt \\ &= X^\pi(t) \left\{ \left[r(t) + \pi^\top(t) (\mu(t) - r(t) \mathbf{1}) \right] dt + \pi^\top(t) \sigma(t) d\mathbf{W}(t) \right\}. \end{aligned} \tag{3.1.3}$$

The goal of solving the problem is to find the optimal trading strategy π^* that achieves the supremum below:

$$V(x) = \sup_{\pi \in \mathcal{A}} \mathbb{E} [U(X^\pi(T))] = \mathbb{E} [U(X^*(T))].$$

The wealth process generated by executing π^* (namely $\{X^*(t), t \in [0, T]\}$) is called the optimal state process.

3.2 Problem Transformation by 2D Kalman-Bucy Filters

Since the drift process μ is not observable, we use $\hat{\mu}$, the conditional expectation of μ given the observable information $\mathcal{F}^S$, to estimate it:

$$\hat{\mu}(t) = \mathbb{E} [\mu(t) | \mathcal{F}^S(t)],$$

which is called a filter of μ . To convert the 2D partially problem into the fully observable problem, we must obtain a 2D Brownian motion that is $\mathcal{F}^S$ -measurable. This can be achieved by theorem 3.2.1 which is an extension of theorem 2.2.1.

Theorem 3.2.1. Define process $\{\hat{\mathbf{V}}(t), t \in [0, T]\}$ by

$$\hat{\mathbf{V}}(t) = \mathbf{W}(t) + \int_0^t \sigma^{-1}(s) (\mu(s) - \hat{\mu}(s)) ds. \tag{3.2.1}$$

Then $\hat{\mathbf{V}}$ is an $\mathcal{F}^S$ -measurable 2D Brownian motion under probability measure $\mathbb{P}$ provided that process $\{\sigma^{-1}(t), t \in [0, T]\}$ is uniformly bounded and

$$\mathbb{E} \left[\int_0^T \|\sigma^{-1}(s) \mu(s)\|^2 ds \right] < \infty.$$

Proof: see [8] for the proof.

By replacing $\mathbf{W}$ in (3.1.3) by $\hat{\mathbf{V}}$, we obtain the transformed SDE of the wealth process:

$$dX^\pi(t) = X^\pi(t) \left\{ \left[r(t) + \pi^\top(t) (\hat{\mu}(t) - r(t) \mathbf{1}) \right] dt + \pi^\top(t) \sigma(t) d\hat{\mathbf{V}}(t) \right\}. \tag{3.2.2}$$

Let $\mathbf{H}(t) = \mu(t)$. Then to construct the Kalman-Bucy filter process, we suppose that $\mathbf{H}$ is driven by SDE

$$d\mathbf{H}(t) = -\lambda (\mathbf{H}(t) - \bar{\mathbf{H}}) dt + \sigma_H \left(\rho d\mathbf{W}(t) + \sqrt{1 - \rho^2} d\mathbf{W}_H(t) \right),$$

where $\lambda, \sigma_H \in \mathbb{R}^{2 \times 2}$, $\bar{H} \in \mathbb{R}^2$, $\rho \in [-1, 1]$ and $\mathbf{W}_H$ is a 2D Brownian motion that is independent of $\mathbf{W}$. Moreover, the initial value $\mathbf{H}(0)$ is independent of $\mathbf{W}$ and $\mathbf{W}_H$. The Kalman-Bucy filter $\hat{\mathbf{H}}$ estimates $\mathbf{H}$ by conditional expectation

$$\hat{\mu}(t) = \hat{\mathbf{H}}(t) = \mathbb{E} [\mathbf{H}(t) | \mathcal{F}^S(t)].$$

It can be proved (see [12]) that $\hat{\mathbf{H}}$ is driven by SDE

$$d\hat{\mathbf{H}}(t) = -\lambda \left(\hat{\mathbf{H}}(t) - \bar{\mathbf{H}} \right) dt + \left(\sigma^{-1}(t) \Sigma(t) + \rho \sigma_H^\top \right)^\top d\hat{\mathbf{V}}(t), \quad (3.2.3)$$

provided that the conditional distribution of $\mathbf{H}(t)$ given $\mathcal{F}^S(t)$ is Gaussian. Process Σ follows ODE

$$\frac{d}{dt} \Sigma(t) = \sigma_H \sigma_H^\top - \lambda \Sigma(t) - \Sigma(t) \lambda^\top - \left(\sigma^{-1}(t) \Sigma(t) + \rho \sigma_H^\top \right)^\top \left(\sigma^{-1}(t) \Sigma(t) + \rho \sigma_H^\top \right).$$

The value of $\Sigma(t)$ is the conditional covariance matrix of $\mathbf{H}(t)$ given $\mathcal{F}^S(t)$:

$$\Sigma(t) = \text{Cov} [\mathbf{H}(t) | \mathcal{F}^S(t)].$$

Prior knowledge must be applied to specify initial values $\hat{\mathbf{H}}(0)$ and $\Sigma(0)$. To guarantee that H is an Ornstein-Uhlenbeck process with mean-reverting property, we can set λ to be a diagonal matrix with all diagonal entries being positive.

The transformation changes the source of randomness of both the wealth process (3.2.2) and the filter (3.2.3). To be specific, the randomness is changed from $\mathbf{W}$ ($\mathcal{F}$ -measurable) to $\hat{\mathbf{V}}$ ($\mathcal{F}^S$ -measurable). In this way, the transformed problem is fully observable. The new value function becomes

$$J(t, x, \mathbf{h}) = \sup_{\pi \in \mathcal{A}} \mathbb{E}_{t,x,\mathbf{h}} [U(X^\pi(T))].$$

To summarize, the newly constructed 2D fully observable problem is formulated as

$$\begin{cases} dX^\pi(t) = X^\pi(t) \left\{ \left[r(t) + \pi^\top(t) \left(\hat{\mathbf{H}}(t) - r(t) \mathbf{1} \right) \right] dt + \pi^\top(t) \sigma(t) d\hat{\mathbf{V}}(t) \right\}, X^\pi(0) = x, \\ d\hat{\mathbf{H}}(t) = -\lambda \left(\hat{\mathbf{H}}(t) - \bar{\mathbf{H}} \right) dt + \left(\sigma^{-1}(t) \Sigma(t) + \rho \sigma_H^\top \right)^\top d\hat{\mathbf{V}}(t), \hat{\mathbf{H}}(0) = \mathbf{h}_0, \\ J(t, x, \mathbf{h}) = \sup_{\pi \in \mathcal{A}} \mathbb{E}_{t,x,\mathbf{h}} [U(X^\pi(T))]. \end{cases} \quad (3.2.4)$$

3.3 The Bound Estimate Approach

The HJB equation of value function $J(t, x, \mathbf{h})$ in problem (3.2.4) is

$$\begin{aligned} \frac{\partial J}{\partial t} + \sup_{\pi(t) \in K} \left\{ x \left[r(t) + \pi^\top(t) (\mathbf{h} - r(t)) \right] \frac{\partial J}{\partial x} + \frac{1}{2} x^2 \pi^\top(t) \sigma(t) \sigma^\top(t) \pi(t) \frac{\partial^2 J}{\partial x^2} - (\mathbf{h} - \bar{\mathbf{H}})^\top \lambda^\top \frac{\partial J}{\partial \mathbf{h}} + \right. \\ \left. x \pi^\top(t) \sigma(t) \left(\sigma^{-1}(t) \Sigma(t) + \rho \sigma_H^\top \right)^\top \frac{\partial^2 J}{\partial x \partial \mathbf{h}} + \right. \\ \left. \frac{1}{2} \text{Tr} \left[\left(\sigma^{-1}(t) \Sigma(t) + \rho \sigma_H^\top \right) \left(\sigma^{-1}(t) \Sigma(t) + \rho \sigma_H^\top \right)^\top \frac{\partial^2 J}{\partial \mathbf{h}^2} \right] \right\} = 0, \end{aligned}$$

which has no closed-form solution. To numerically solve problem (3.2.4) by the bound estimate approach, the dual control theory and SMP must be applied.

3.3.1 Dual Control Problems

The dual process of X^π (denoted by $Y^{(y,v)}$) is driven by SDE

$$\begin{cases} dY^{(y,v)}(t) = -Y^{(y,v)}(t) \left\{ [r(t) + \delta_K(\mathbf{v}(t))] dt + \left[\boldsymbol{\sigma}^{-1}(t) \left(\hat{\mathbf{H}}(t) - r(t)\mathbf{1} + \mathbf{v}(t) \right) \right]^\top d\hat{\mathbf{V}}(t) \right\}, \\ Y^{(y,v)}(0) = y, \end{cases}$$

where $\hat{\mathbf{H}}$ is the 2D Kalman-Bucy filter driven by (3.2.3), δ_K is the support function of set $-K$:

$$\delta_K(\mathbf{v}) = \sup_{\boldsymbol{\pi} \in K} \left(-\boldsymbol{\pi}^\top \mathbf{v} \right), \mathbf{v} \in \mathbb{R}^2,$$

and $\mathbf{v}$ is the dual control process whose admissible values all stay in set

$$\mathcal{B} = \left\{ \boldsymbol{\zeta} : \Omega \times [0, T] \rightarrow \mathbb{R}^2 : \boldsymbol{\zeta} \text{ is } \mathcal{F}^S \text{-measurable, and } \int_0^T [\delta_K(\boldsymbol{\zeta}(s)) + \|\boldsymbol{\zeta}(s)\|^2] ds < \infty \text{ almost surely} \right\}.$$

The corresponding dual value function is

$$\tilde{V}(t, x, \mathbf{h}) = \inf_{y \in \mathbb{R}^+, \mathbf{v} \in \mathcal{B}} \left\{ xy + \mathbb{E}_{t,y,\mathbf{h}} \left[\tilde{U} \left(Y^{(y,\mathbf{v})}(T) \right) \right] \right\}.$$

Theorem 3.3.1. *Let $\tilde{K} = \{\mathbf{v} : \boldsymbol{\pi}^\top \mathbf{v} \geq 0, \forall \boldsymbol{\pi} \in K\}$ be the positive polar cone of K . Then for both the unconstrained case ($K = \mathbb{R}^2$) and the constrained case ($K = \mathbb{R}^{+2}$), we have*

$$\delta_K(\mathbf{v}) = \begin{cases} 0, & \text{if } \mathbf{v} \in \tilde{K} \\ \infty, & \text{if } \mathbf{v} \in K/\tilde{K} \end{cases}. \quad (3.3.1)$$

Proof: When $K = \mathbb{R}^2$, we have $\tilde{K} = \{\mathbf{0}\}$. Then for any $\boldsymbol{\pi} \in K$, we have

$$\begin{cases} \mathbf{v} \in \tilde{K} \implies -\boldsymbol{\pi}^\top \mathbf{v} = 0 \implies \delta_K(\mathbf{v}) = 0, \\ \mathbf{v} \in K/\tilde{K} \implies \text{for any } M > 0 \text{ there exists } \boldsymbol{\pi} \in K \text{ such that } -\boldsymbol{\pi}^\top \mathbf{v} > M \implies \delta_K(\mathbf{v}) = \infty, \end{cases}$$

meaning that (3.3.1) holds for the unconstrained case. Alternatively, when $K = \mathbb{R}^{+2}$, we have $\tilde{K} = \mathbb{R}^{+2}$. Then for any $\boldsymbol{\pi} \in K$, we have

$$\begin{cases} \mathbf{v} \in \tilde{K} \implies -\boldsymbol{\pi}^\top \mathbf{v} \leq 0 \implies \delta_K(\mathbf{v}) = 0, \\ \mathbf{v} \in K/\tilde{K} \implies \text{for any } M > 0 \text{ there exists } \boldsymbol{\pi} \in K \text{ such that } -\boldsymbol{\pi}^\top \mathbf{v} > M \implies \delta_K(\mathbf{v}) = \infty, \end{cases}$$

meaning that (3.3.1) also holds for the constrained case. The dual control process $\mathbf{v}$ is required to satisfy $\mathbf{v}(t) \in \tilde{K}$ for any $t \in [0, T]$ as in 1D problems.

3.3.2 Stochastic Maximum Principle

By SMP, we obtain the 2D analog of theorem 2.3.4:

Theorem 3.3.2. *Let $\tilde{y} \in \mathbb{R}^+$ and $\tilde{\mathbf{v}} \in \mathcal{B}$. Consider the system of SDE*

$$\begin{cases} dY^{(\tilde{y},\tilde{\mathbf{v}})}(t) = -Y^{(\tilde{y},\tilde{\mathbf{v}})}(t) \left\{ [r(t) + \delta_K(\tilde{\mathbf{v}}(t))] dt + \left[\boldsymbol{\sigma}^{-1}(t) \left(\hat{\mathbf{H}}(t) - r(t)\mathbf{1} + \tilde{\mathbf{v}}(t) \right) \right]^\top d\hat{\mathbf{V}}(t) \right\}, \\ d\hat{P}(t) = \left[r(t)\hat{P}(t) + \hat{\mathbf{Q}}^\top(t)\boldsymbol{\sigma}^{-1}(t) \left(\hat{\mathbf{H}}(t) - r(t)\mathbf{1} \right) \right] dt + \hat{\mathbf{Q}}^\top(t)d\hat{\mathbf{V}}(t), \\ d\hat{\mathbf{H}}(t) = -\lambda \left(\hat{\mathbf{H}}(t) - \bar{\mathbf{H}} \right) dt + \left(\boldsymbol{\sigma}^{-1}(t)\boldsymbol{\Sigma}(t) + \rho\boldsymbol{\sigma}_H^\top \right)^\top d\hat{\mathbf{V}}(t), \end{cases} \quad (3.3.2)$$

with initial or terminal conditions $Y^{(\tilde{y},\tilde{\mathbf{v}})}(0) = \tilde{y}, \hat{P}(T) = -\tilde{U}'(Y^{(\tilde{y},\tilde{\mathbf{v}})}(T)), \hat{\mathbf{H}}(0) = \mathbf{h}_0$. Then $(\tilde{y}, \tilde{\mathbf{v}})$ is the optimal control for the dual control problem if and only if the solution to (3.3.2) (denoted by $(Y^{(\tilde{y},\tilde{\mathbf{v}})}, \hat{P}, \hat{\mathbf{Q}})$) satisfies

- condition 1: the initial value of $\hat{P}$ must match with X :

$$\hat{P}(0) = X(0) = x_0,$$

- condition 2: the optimal trading strategy (i.e. π^*) of the primal problem can be shown to be given by (3.3.3) and must obey constraints K :

$$\pi^*(t) = \frac{\sigma^{-\top}(t)\hat{Q}(t)}{\hat{P}(t)} \in K, \quad (3.3.3)$$

- condition 3: for any $t \in [0, T]$, (3.3.4) holds almost surely:

$$\hat{P}(t)\delta_K(\tilde{\mathbf{v}}(t)) + \hat{Q}^\top(t)\sigma^{-1}(t)\tilde{\mathbf{v}}(t) = 0, \quad \mathbb{P} - a.s. \quad (3.3.4)$$

Proof: see theorem 3.9 and 3.10 in [11].

With 3.3.2 and the 2D analog of (2.3.8), i.e.

$$d\hat{P}(t) = \hat{P}(t) \left\{ \left[r(t) + \pi^{*\top}(t) \left(\hat{\mathbf{H}}(t) - r(t)\mathbf{1} \right) \right] dt + \pi^{*\top}(t)\sigma(t)d\hat{\mathbf{V}}(t) \right\},$$

we obtain that $\hat{P}$ matches the optimal state process X^* . A lower bound $\underline{L}$ and an upper bound $\bar{L}$ of the value function is given by the weak duality, i.e.

$$\underline{L} = \mathbb{E}[U(X^\pi(T))] \quad \text{and} \quad \bar{L} = x_0y + \mathbb{E}\left[\tilde{U}\left(Y^{(y,v)}(T)\right)\right]. \quad (3.3.5)$$

Processes π and $\mathbf{v}$ are required to satisfy (see theorem 3.3.2)

$$\mathbb{E}\left[w\left|\hat{P}(T) + \tilde{U}'(Y^{(\tilde{y},\tilde{\mathbf{v}})}(T))\right|^2 + (1-w)\int_0^T \left[\delta_K(\mathbf{v}(t)) + \pi^\top(t)\mathbf{v}(t)\right] dt\right] = 0, \quad (3.3.6)$$

with $0 < w < 1$. The goal is to find values of $y, \pi(t), \mathbf{v}(t), t \in [0, T]$ such that terminal values $\hat{P}(T)$ and $Y^{(\tilde{y},\tilde{\mathbf{v}})}(T)$ generated by system (3.3.2) satisfy (3.3.6). We have shown that $\delta_K(\mathbf{v}(t)) \in \{0, \infty\}$ and $\mathbf{v}(t) \in \tilde{K} \implies \pi^\top(t)\mathbf{v}(t) \geq 0$ for any $\pi \in K$, so the integral in (3.3.6) is nonnegative. The problem is then solved by minimizing the objective function

$$f(y, \pi, \mathbf{v}) = \mathbb{E}\left[w\left|\hat{P}(T) + \tilde{U}'(Y^{(\tilde{y},\tilde{\mathbf{v}})}(T))\right|^2 + (1-w)\int_0^T \left[\delta_K(\mathbf{v}(t)) + \pi^\top(t)\mathbf{v}(t)\right] dt\right]. \quad (3.3.7)$$

After obtaining y, π and $\mathbf{v}$, bounds $\underline{L}$ and $\bar{L}$ in (3.3.5) can be computed.

3.3.3 Numerical Experiments

We introduce the numerical experiment in this section. We first assign values to all variables involved in this problem. To be specific, we set $r = 0.05$, $x_0 = 10.0$, $\rho = 0.0$, $T = 1.0$, and

$$\mathbf{h}_0 = \begin{bmatrix} 0.1 \\ 0.1 \end{bmatrix}, \bar{\mathbf{H}} = \begin{bmatrix} 0.1 \\ 0.1 \end{bmatrix},$$

$$\boldsymbol{\lambda} = \begin{bmatrix} 1.0 & 0 \\ 0 & 1.0 \end{bmatrix}, \boldsymbol{\sigma} = \begin{bmatrix} 0.8 & 0 \\ 0 & 0.8 \end{bmatrix}, \boldsymbol{\sigma}_H = \begin{bmatrix} 0.5 & 0 \\ 0 & 0.5 \end{bmatrix}, \boldsymbol{\Sigma}_0 = \begin{bmatrix} 0.2 & 0.0 \\ 0.0 & 0.2 \end{bmatrix},$$

where $r(t)$ and $\sigma(t)$ are assumed to be constants $r \in \mathbb{R}$ and $\sigma \in \mathbb{R}^2$ for all $t \in [0, T]$. The utility function U is set to be the power utility function (2.3.2) with $\beta = 0.5$. These are called baseline settings.

SDEs in system (3.3.2) are approximated by discretization. The investment period $[0, T]$ is divided into N pieces with equal length:

$$[t_i, t_{i+1}] = [ih, (i+1)h], i = 0, 1, \dots, N-1,$$

where the step size is $h = \frac{T}{N}$. Then the approximation is given by

$$\begin{cases} Y_{i+1} = Y_i - rY_i h - Y_i \left[\sigma^{-1} \left(\hat{\mathbf{H}}_i - r\mathbf{1} + \mathbf{v}_i \right) \right]^\top \left(\hat{\mathbf{V}}_{i+1} - \hat{\mathbf{V}}_i \right), & Y_0 = y, \\ \hat{P}_{i+1} = \hat{P}_i + r\hat{P}_i h + \hat{P}_i \pi_i^\top \left(\hat{\mathbf{H}}_i - r\mathbf{1} \right) h + \hat{P}_i \pi_i^\top \sigma \left(\hat{\mathbf{V}}_{i+1} - \hat{\mathbf{V}}_i \right), & \hat{P}_0 = x_0, \\ \hat{\mathbf{H}}_{i+1} = \hat{\mathbf{H}}_i + \lambda \left(\bar{\mathbf{H}} - \hat{\mathbf{H}}_i \right) h + \left[\sigma^{-1} \left(\Sigma_i + \rho \sigma \sigma_H^\top \right) \right]^\top \left(\hat{\mathbf{V}}_{i+1} - \hat{\mathbf{V}}_i \right), & \hat{\mathbf{H}}_0 = \mathbf{h}_0, \end{cases} \quad (3.3.8)$$

for $i = 0, 1, \dots, N-1$, where $\{Y_i, \hat{P}_i, \hat{\mathbf{H}}_i\}_{i=0}^N$ are approximations of $\{Y(ih), \hat{P}(ih), \hat{\mathbf{H}}(ih)\}_{i=0}^N$, and the innovation term is 2D Gaussian:

$$\hat{\mathbf{V}}_{i+1} - \hat{\mathbf{V}}_i = \sqrt{h} \boldsymbol{\varepsilon} = \sqrt{h} \begin{bmatrix} \varepsilon^{(1)} \\ \varepsilon^{(2)} \end{bmatrix},$$

where $\varepsilon^{(1)}$ and $\varepsilon^{(2)}$ both follow 1D standard Gaussian distribution $\mathcal{N}(0, 1)$ and are independent. We choose $N = 10$ as [18] suggests. The objective function is thus approximated by

$$\begin{aligned} \hat{f}(y, \{\pi\}_{i=0}^{N-1}, \{\mathbf{v}\}_{i=0}^{N-1}) &= \mathbb{E} \left[w \left| \hat{P}_N + \tilde{U}'(Y_N) \right|^2 + (1-w)h \sum_{i=0}^{N-1} \left[\delta_K(\mathbf{v}_i) + \pi_i^\top \mathbf{v}_i \right] \right] \\ &= \mathbb{E} \left[w \left| \hat{P}_N + \tilde{U}'(Y_N) \right|^2 + (1-w)h \sum_{i=0}^{N-1} \pi_i^\top \mathbf{v}_i \right], \end{aligned} \quad (3.3.9)$$

where we use the fact that $\delta_K(\mathbf{v}_i) = 0$ for $\mathbf{v}_i \in \tilde{K}$. Indeed, $\mathbf{v}_i \in \tilde{K}$ must hold because otherwise $\delta_K(\mathbf{v}_i) = \infty$ and then the minimization of $\hat{f}$ is impossible. Expectation (3.3.9) and bounds $\underline{L}, \bar{L}$ are all estimated by Monte Carlo methods, and the settings for the Monte Carlo simulation are the same as those in 1D problems.

In the following analysis, we use notation ' $\odot$ ' as the operator for elementwise product, i.e. given $\mathbf{x}, \mathbf{y} \in \mathbb{R}^2$ we have

$$\mathbf{x} \odot \mathbf{y} = \begin{bmatrix} x^{(1)} \\ x^{(2)} \end{bmatrix} \odot \begin{bmatrix} y^{(1)} \\ y^{(2)} \end{bmatrix} = \begin{bmatrix} x^{(1)}y^{(1)} \\ x^{(2)}y^{(2)} \end{bmatrix}.$$

This operator makes some matrix-vector formulae simple to present.

Evaluation of $\Sigma(t)$

We need the value of Σ_i (the approximation of $\Sigma(t_i)$) when we do Monte Carlo simulation for (3.3.8). The matrix process Σ is driven by ODE

$$\frac{d}{dt} \Sigma(t) = \sigma_H \sigma_H^\top - \lambda \Sigma(t) - \Sigma(t) \lambda^\top - \hat{\Sigma}_R^\top(t) \hat{\Sigma}_R(t), \quad (3.3.10)$$

with $\Sigma(t), \lambda, \sigma_H \in \mathbb{R}^{2 \times 2}$ and

$$\hat{\Sigma}_R(t) = \sigma^{-1}(t) \left[\Sigma(t) + \rho \sigma(t) \sigma_H^\top \right],$$

Hence, we can compute Σ_i by numerically solving (3.3.10).

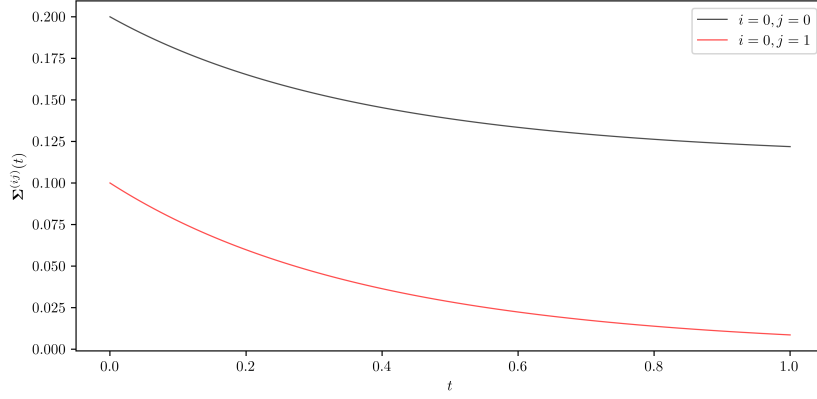


Figure 3.1: An example of the solution curve of $\Sigma(t), t \in [0, T]$

To accurately evaluate Σ , we divide the whole investment period $[0, T]$ into E pieces and calculate the value of Σ at endpoints of all pieces. In this paper we use $T = 1$ and set $E = 1000$ (the corresponding step size is $\delta t = \frac{T}{E} = 0.001$). In this way, all endpoints are $t_e = e\delta t, e = 0, 1, \dots, E$, and thus $\Sigma(t_e)$ is estimated by

$$\begin{cases} \Sigma(t_0) = \Sigma(0) \\ \Sigma(t_{e+1}) \approx \Sigma(t_e) + \left(\sigma_H \sigma_H^\top - \lambda \Sigma(t_e) - \Sigma(t_e) \lambda^\top - \hat{\Sigma}_R^\top(t_e) \hat{\Sigma}_R(t_e) \right) \delta t, \end{cases} \quad (3.3.11)$$

for $e = 0, 1, \dots, E - 1$, where

$$\hat{\Sigma}_R(t_e) = \sigma^{-1} \left[\Sigma(t_e) + \rho \sigma \sigma_H^\top \right].$$

When $\Sigma(t)$ needs to be evaluated but t is not an endpoint of any piece, we use linear interpolation to compute $\Sigma(t)$. For any $t \in [0, T]$, there exists an index e such that $t \in (t_e, t_{e+1})$. Then $\Sigma(t)$ is estimated by

$$\Sigma(t) \approx \Sigma(t_e) + \frac{t - t_e}{t_{e+1} - t_e} [\Sigma(t_{e+1}) - \Sigma(t_e)].$$

We provide a numerical example, where all variables follow baseline settings except that $\Sigma(0)$ is changed into

$$\Sigma(0) = \begin{bmatrix} 0.2 & 0.01 \\ 0.01 & 0.2 \end{bmatrix}.$$

Algorithm (3.3.11) generates curves of entries in $\Sigma(t), t \in [0, T]$. These curves are shown in Figure 3.1. Note that although $\Sigma(t) \in \mathbb{R}^{2 \times 2}$, Figure 3.1 only shows two curves $\Sigma^{(00)}(t)$ and $\Sigma^{(01)}(t)$ for $t \in [0, T]$. This is because $\Sigma^{(00)}(t) = \Sigma^{(11)}(t)$ (since $\Sigma^{(00)}(0) = \Sigma^{(11)}(0) = 0.2$) and $\Sigma^{(01)}(t) = \Sigma^{(10)}(t)$ (since $\Sigma(t)$ is a symmetric matrix).

Unconstrained Problems

For unconstrained problems we have $K = \mathbb{R}^2$ and $\tilde{K} = \{\mathbf{0}\}$. Since $\mathbf{v}_i \in \tilde{K}$ holds all the time, we must have $\mathbf{v}_i = \mathbf{0}$ for any index i . So the second part of $\hat{f}$ is removed and thus $\hat{f}$ reduces to

$$\hat{f}(y, \{\pi_i\}_{i=0}^{N-1}) = \mathbb{E} \left[\left| \hat{P}_N + \tilde{U}'(Y_N) \right|^2 \right].$$

	Number of Parameters	Details of Parameters
Form I	5	$y, \mathbf{a}^\pi, \mathbf{b}^\pi$
Form II	$2N + 3$	$y, \{\mathbf{a}_i^\pi\}_{i=0}^{N-1}, \mathbf{b}^\pi$

Table 3.1: Parameters to determine for 2D unconstrained problems under the Kalman-Bucy filter

Two forms of parameterization for the trading strategy π are tested. In general, the optimal trading strategy π is represented by

$$\pi_i = \begin{bmatrix} \pi_i^{(1)} \\ \pi_i^{(2)} \end{bmatrix} = \underbrace{\begin{bmatrix} a_i^{(\pi,1)} \\ a_i^{(\pi,2)} \end{bmatrix}}_{\mathbf{a}_i^\pi} + \underbrace{\begin{bmatrix} b_i^{(\pi,1)} \\ b_i^{(\pi,2)} \end{bmatrix}}_{\mathbf{b}_i^\pi} \odot \hat{H}_i = \mathbf{a}_i^\pi + \mathbf{b}_i^\pi \odot \hat{H}_i.$$

Form I requires

$$\mathbf{a}_0^{(\pi)} = \mathbf{a}_1^{(\pi)} = \dots = \mathbf{a}_{N-1}^{(\pi)} := \mathbf{a}^\pi \text{ and } \mathbf{b}_0^{(\pi)} = \mathbf{b}_1^{(\pi)} = \dots = \mathbf{b}_{N-1}^{(\pi)} := \mathbf{b}^\pi.$$

In contrast, form II only requires

$$\mathbf{b}_0^{(\pi)} = \mathbf{b}_1^{(\pi)} = \dots = \mathbf{b}_{N-1}^{(\pi)} := \mathbf{b}^\pi.$$

Therefore, we have 5 and $2N+3$ parameters to determine for form I and form II respectively (see Table 3.1).

To see how results change as the distribution of the initial guess changes, we repeat the test for six different combinations. To be specific, let

$$\mathbf{h}_0^- = \begin{bmatrix} 0.05 \\ 0.05 \end{bmatrix}, \mathbf{h}_0^\circ = \begin{bmatrix} 0.1 \\ 0.1 \end{bmatrix}, \mathbf{h}_0^+ = \begin{bmatrix} 0.2 \\ 0.2 \end{bmatrix}.$$

$$\Sigma_0^- = \begin{bmatrix} 0.2 & 0.01 \\ 0.01 & 0.2 \end{bmatrix}, \Sigma_0^\circ = \begin{bmatrix} 0.2 & 0 \\ 0 & 0.2 \end{bmatrix}, \Sigma_0^+ = \begin{bmatrix} 0.2 & -0.01 \\ -0.01 & 0.2 \end{bmatrix}.$$

We first fix $\Sigma_0 = \Sigma_0^\circ$ and run the algorithm for $\mathbf{h}_0 = \mathbf{h}_0^-, \mathbf{h}_0^\circ, \mathbf{h}_0^+$ one by one (see Table 3.2). Then we fix $\mathbf{h}_0 = \mathbf{h}_0^\circ$ and run the algorithm for $\Sigma_0 = \Sigma_0^-, \Sigma_0^\circ, \Sigma_0^+$ one by one (see Table 3.3). We observe similar results to those of 1D problems: $\underline{L}$ and $\bar{L}$ are close to each other with relative errors around 2%.

$\mathbf{h}_0$	Form	$\underline{L}$	$\bar{L}$	Error	OB	MV
$\mathbf{h}_0^-$	I	6.4742	6.6380	2.5299%	0.0023	6.5561
$\mathbf{h}_0^-$	II	6.4986	6.6611	2.5008%	0.0024	6.5798
$\mathbf{h}_0^\circ$	I	6.5209	6.6789	2.4235%	0.0016	6.5999
$\mathbf{h}_0^\circ$	II	6.5140	6.6840	2.6101%	0.0022	6.5990
$\mathbf{h}_0^+$	I	6.6175	6.8700	3.8161%	0.0018	6.7437
$\mathbf{h}_0^+$	II	6.6009	6.8386	3.6006%	0.0020	6.7197

Table 3.2: Numerical results for 2D unconstrained problems by varying $\mathbf{h}_0$ under the Kalman-Bucy filter

Σ_0	Form	$\underline{L}$	$\overline{L}$	Error	OB	MV
Σ_0^-	I	6.5276	6.6805	2.3423%	0.0017	6.6040
Σ_0^-	II	6.5261	6.6708	2.2176%	0.0016	6.5984
Σ_0°	I	6.5242	6.6620	2.1131%	0.0019	6.5931
Σ_0°	II	6.5221	6.6661	2.2087%	0.0019	6.5941
Σ_0^+	I	6.5205	6.6793	2.4359%	0.0019	6.5999
Σ_0^+	II	6.5349	6.6704	2.0731%	0.0015	6.6027

Table 3.3: Numerical results for 2D unconstrained problems by varying Σ_0 under the Kalman-Bucy filter

Constrained Problems

For constrained problems we have $K = \mathbb{R}^{+2}$ and $\tilde{K} = \mathbb{R}^{+2}$, so $\hat{f}$ cannot reduce anymore and the dual control $\mathbf{v}$ also needs to be parameterized. We have to ensure that both $\{\pi_i\}_{i=0}^{N-1}$ and $\{v_i\}_{i=0}^{N-1}$ lie in $\mathbb{R}^{+2}$. In general, the optimal trading strategy π and the dual control $\mathbf{v}$ are represented by

$$\pi_i = \begin{bmatrix} \pi_i^{(1)} \\ \pi_i^{(2)} \end{bmatrix} = \left(\underbrace{\begin{bmatrix} a_i^{(\pi,1)} \\ a_i^{(\pi,2)} \end{bmatrix}}_{\mathbf{a}_i^\pi} + \underbrace{\begin{bmatrix} b_i^{(\pi,1)} \\ b_i^{(\pi,2)} \end{bmatrix}}_{\mathbf{b}_i^\pi} \odot \hat{H}_i \right)^+ = \left(\mathbf{a}_i^\pi + \mathbf{b}_i^\pi \odot \hat{H}_i \right)^+$$

and

$$\mathbf{v}_i = \begin{bmatrix} v_i^{(1)} \\ v_i^{(2)} \end{bmatrix} = \left(\underbrace{\begin{bmatrix} a_i^{(\mathbf{v},1)} \\ a_i^{(\mathbf{v},2)} \end{bmatrix}}_{\mathbf{a}_i^{\mathbf{v}}} + \underbrace{\begin{bmatrix} b_i^{(\mathbf{v},1)} \\ b_i^{(\mathbf{v},2)} \end{bmatrix}}_{\mathbf{b}_i^{\mathbf{v}}} \odot \hat{H}_i \right)^+ = \left(\mathbf{a}_i^{\mathbf{v}} + \mathbf{b}_i^{\mathbf{v}} \odot \hat{H}_i \right)^+$$

respectively, where $(\mathbf{x})^+$ means applying $(\cdot)^+$ to vector $\mathbf{x} \in \mathbb{R}^2$ in an elementwise manner. Form I requires

$$\mathbf{a}_0^{(\pi)} = \mathbf{a}_1^{(\pi)} = \dots = \mathbf{a}_{N-1}^{(\pi)} := \mathbf{a}^\pi, \mathbf{b}_0^{(\pi)} = \mathbf{b}_1^{(\pi)} = \dots = \mathbf{b}_{N-1}^{(\pi)} := \mathbf{b}^\pi$$

and

$$\mathbf{a}_0^{(\mathbf{v})} = \mathbf{a}_1^{(\mathbf{v})} = \dots = \mathbf{a}_{N-1}^{(\mathbf{v})} := \mathbf{a}^{\mathbf{v}}, \mathbf{b}_0^{(\mathbf{v})} = \mathbf{b}_1^{(\mathbf{v})} = \dots = \mathbf{b}_{N-1}^{(\mathbf{v})} := \mathbf{b}^{\mathbf{v}}$$

In contrast, form II only requires

$$\mathbf{b}_0^{(\pi)} = \mathbf{b}_1^{(\pi)} = \dots = \mathbf{b}_{N-1}^{(\pi)} := \mathbf{b}^\pi \text{ and } \mathbf{b}_0^{(\mathbf{v})} = \mathbf{b}_1^{(\mathbf{v})} = \dots = \mathbf{b}_{N-1}^{(\mathbf{v})} := \mathbf{b}^{\mathbf{v}}.$$

Consequently, we have 9 and $4N + 5$ parameters to determine for form I and form II respectively (see Table 3.4).

All 6 combinations tested for unconstrained problems are also tested for constrained problems with $w = w^\circ = 0.5$. Results are listed in Tables 3.5 and 3.6. The effect of varying w is tested next. Let $w^- = 0.1, w^\circ = 0.5, w^+ = 0.9$. We fix $h_0 = h_0^\circ$ and $\Sigma_0 = \Sigma_0^\circ$ and run the bound estimate algorithm for $w = w^-, w^\circ, w^+$ one by one, where $w^- = 0.1, w^\circ = 0.5, w^+ = 0.9$. Results are listed in Table 3.7. We observe that $\underline{L}$ and $\overline{L}$ are close to each other with relative errors around 2% in almost all combinations. Moreover, relative errors for w^- and w° are greater than that for w^+ , so w^+ is a better choice than w^- or w° .

	Number of Parameters	Details of Parameters
Form I	9	$y, \mathbf{a}^\pi, \mathbf{b}^\pi, \mathbf{a}^v, \mathbf{b}^v$
Form II	$4N + 5$	$y, \{\mathbf{a}_i^\pi\}_{i=0}^{N-1}, \{\mathbf{a}_i^v\}_{i=0}^{N-1}, \mathbf{b}^\pi, \mathbf{b}^v$

Table 3.4: Parameters to determine for 2D constrained problems under the Kalman-Bucy filter

$\mathbf{h}_0$	Form	$\underline{L}$	$\overline{L}$	Error	OB	MV
$\mathbf{h}_0^-$	I	6.4860	6.6382	2.3466%	0.0011	6.5621
$\mathbf{h}_0^-$	II	6.4843	6.6517	2.5815%	0.0010	6.5680
$\mathbf{h}_0^\circ$	I	6.5413	6.6715	1.9905%	0.0009	6.6064
$\mathbf{h}_0^\circ$	II	6.5091	6.6767	2.5739%	0.0009	6.5929
$\mathbf{h}_0^+$	I	6.6036	6.8384	3.5559%	0.0016	6.7210
$\mathbf{h}_0^+$	II	6.5971	6.8297	3.5259%	0.0034	6.7134

Table 3.5: Numerical results for 2D constrained problems by varying $\mathbf{h}_0$ under the Kalman-Bucy filter

Σ_0	Form	$\underline{L}$	$\overline{L}$	Error	OB	MV
Σ_0^-	I	6.5099	6.6685	2.4362%	0.0011	6.5892
Σ_0^-	II	6.5068	6.6871	2.7717%	0.0009	6.5969
Σ_0°	I	6.4849	6.6631	2.7476%	0.0012	6.5740
Σ_0°	II	6.4845	6.6664	2.8058%	0.0011	6.5755
Σ_0^+	I	6.4846	6.6691	2.8452%	0.0012	6.5769
Σ_0^+	II	6.4861	6.6795	2.9814%	0.0013	6.5828

Table 3.6: Numerical results for 2D constrained problems by varying Σ_0 under the Kalman-Bucy filter

w	Form	$\underline{L}$	$\overline{L}$	Error	OB	MV
w^-	I	6.4847	6.7143	3.5406%	0.0006	6.5995
w^-	II	6.5223	6.6529	2.0036%	0.0002	6.5876
w°	I	6.4844	6.6750	2.9396%	0.0011	6.5797
w°	II	6.5483	6.6646	1.7760%	0.0006	6.6065
w^+	I	6.5483	6.6646	1.7760%	0.0006	6.6065
w^+	II	6.5267	6.6560	1.9821%	0.0091	6.5913

Table 3.7: Numerical results for 2D constrained problems by varying w under the Kalman-Bucy filter

Chapter 4

1D Partially Observable Problems under Bayesian Filters

Besides the Kalman-Bucy filter, other filters can also be applied to estimate the partially observable drift process $\{\mu(t), t \in [0, T]\}$. In this chapter, we give the solution to 1D partially observable problems under the Bayesian filter. Both the bound estimate approach and the deep learning approach are covered in numerical experiments.

4.1 1D Bayesian Filters

This section mostly follows [5] and is a brief recap on 1D Bayesian filter. In 1D problems, when Bayesian filters are applied, the drift process μ is treated as a random variable B that is independent of the Brownian motion W . The prior distribution of B is described by measure ν such that $\mathbb{E}[B^2] < \infty$. Initially, ν is determined by the prior belief, knowledge or experience on the problem. The stock volatility process is set to be a constant σ so that we can derive mathematically tractable results.

Define a 1D process U as

$$U(t) = \frac{1}{\sigma}Bt + W(t), t \in [0, T].$$

Since S and U share the same source of randomness, the filtration generated by process U (denoted by $\mathcal{F}^U$) is the same as $\mathcal{F}^S$.

Theorem 4.1.1. *Process S can be represented by U through (4.1.1):*

$$S(t) = S(0) \exp \left(\sigma U(t) - \frac{1}{2} \sigma^2 t \right). \quad (4.1.1)$$

Proof: the SDE of Y is

$$dU(t) = \frac{1}{\sigma}Bdt + dW(t).$$

It follows that

$$dS(t) = S(t) (Bdt + \sigma dW(t)) = \sigma S(t) dY(t).$$

By Itô's formula, we obtain

$$\begin{aligned} d \log S(t) &= \frac{1}{S(t)} dS(t) - \frac{1}{2S^2(t)} dS(t) dS(t) \\ &= \sigma dU(t) - \frac{1}{2} \sigma^2 dU(t) dU(t) \\ &= \sigma dU(t) - \frac{1}{2} \sigma^2 dt. \end{aligned} \quad (4.1.2)$$

Taking integral both sides of (4.1.2) on interval $[0, t]$, we obtain

$$\log S(t) = \log S(0) + \sigma \int_0^t dU(s) - \frac{1}{2}\sigma^2 \int_0^t ds = \log S(0) + \sigma U(t) - \frac{1}{2}\sigma^2 t,$$

so we have

$$S(t) = S(0) \exp \left(\sigma U(t) - \frac{1}{2}\sigma^2 t \right).$$

We next represent $\hat{H}$ by U so that Itô's formula and the SDE of U can be used to derive the SDE of $\hat{H}$.

Theorem 4.1.2. *Let $g : \mathbb{R} \rightarrow \mathbb{R}$ be a function with*

$$\int_{-\infty}^{\infty} |g(b)| \nu(db) < \infty.$$

Then the estimate of $g(B)$ is given by

$$\mathbb{E} [g(B)|\mathcal{F}^S(t)] = \mathbb{E} [g(B)|\mathcal{F}^U(t)] = \mathbb{E} [g(B)|U(t)] = \frac{\int_{-\infty}^{\infty} g(b) \exp \left(\frac{1}{\sigma} b U(t) - \frac{b^2}{2\sigma^2} t \right) \nu(db)}{\int_{-\infty}^{\infty} \exp \left(\frac{1}{\sigma} b U(t) - \frac{b^2}{2\sigma^2} t \right) \nu(db)}.$$

Proof: see [5] for the proof.

Let $g(x) = x$ for any $x \in \mathbb{R}$ in theorem 4.1.2, we obtain the filter of B that is used to solve our utility maximization problem:

$$\hat{H}(t) = \mathbb{E} [B|\mathcal{F}^S(t)] = f(t, U(t)),$$

where function $f : [0, T] \times \mathbb{R} \rightarrow \mathbb{R}$ is given by

$$f(t, u) = \frac{\int_{-\infty}^{\infty} b \exp \left(\frac{1}{\sigma} b u - \frac{b^2}{2\sigma^2} t \right) \nu(db)}{\int_{-\infty}^{\infty} \exp \left(\frac{1}{\sigma} b u - \frac{b^2}{2\sigma^2} t \right) \nu(db)}. \quad (4.1.3)$$

It can also be shown that $\hat{H}$ is a martingale with respect to measure $\mathbb{P}$ and filtration $\mathcal{F}^S$ (see [5]). In Bayesian filter, the $\mathcal{F}^S$ -measurable Brownian motion $\hat{V}$ defined by (2.2.1) satisfies

$$d\hat{V}(t) = dW(t) + \frac{B - \hat{H}(t)}{\sigma} dt$$

Then we can write the SDE of Y driven by $\hat{V}$:

$$dU(t) = \frac{1}{\sigma} B dt + dW(t) = \frac{1}{\sigma} B dt + \left(d\hat{V}(t) - \frac{B - \hat{H}(t)}{\sigma} dt \right) = d\hat{V}(t) + \frac{1}{\sigma} \hat{H}(t) dt. \quad (4.1.4)$$

By Itô's formula and the martingale property (see [5]) of $\hat{H}$, we obtain

$$d\hat{H}(t) = df(t, U(t)) = \frac{\partial f(t, U(t))}{\partial u} d\hat{V}(t), \quad (4.1.5)$$

which is the SDE of the 1D Bayesian filter.

It is necessary to derive the form of function f to calculate the partial derivative in (4.1.5), so we must specify a valid measure ν and then evaluate (4.1.3). We consider two cases. In the first case, the prior distribution of B is a discrete distribution with only two possible values $b_l \in \mathbb{R}$ (with probability $1 - p$) and $b_h \in \mathbb{R}$ (with probability p):

$$B \sim \begin{pmatrix} b_l & b_h \\ 1 - p & p \end{pmatrix}.$$

The corresponding measure ν is represented by Dirac measures at b_l and b_h :

$$\nu(db) = (1 - p)\delta_{b_l}(db) + p\delta_{b_h}(db),$$

where $r < b_l < b_h$ and $0 < p < 1$. The integral in the denominator of (4.1.3) is

$$\begin{aligned} & \int_{-\infty}^{\infty} \exp\left(\frac{1}{\sigma}bu - \frac{b^2}{2\sigma^2}t\right) \nu(db) \\ &= \int_{-\infty}^{\infty} \exp\left(\frac{1}{\sigma}bu - \frac{b^2}{2\sigma^2}t\right) [(1 - p)\delta_{b_l}(db) + p\delta_{b_h}(db)] \\ &= (1 - p) \exp\left(\frac{1}{\sigma}b_l u - \frac{b_l^2}{2\sigma^2}t\right) + p \exp\left(\frac{1}{\sigma}b_h u - \frac{b_h^2}{2\sigma^2}t\right). \end{aligned}$$

The integral in the numerator is evaluated in a similar way:

$$\begin{aligned} & \int_{-\infty}^{\infty} b \exp\left(\frac{1}{\sigma}bu - \frac{b^2}{2\sigma^2}t\right) \nu(db) \\ &= \int_{-\infty}^{\infty} b \exp\left(\frac{1}{\sigma}bu - \frac{b^2}{2\sigma^2}t\right) [(1 - p)\delta_{b_l}(db) + p\delta_{b_h}(db)] \\ &= (1 - p)b_l \exp\left(\frac{1}{\sigma}b_l u - \frac{b_l^2}{2\sigma^2}t\right) + pb_h \exp\left(\frac{1}{\sigma}b_h u - \frac{b_h^2}{2\sigma^2}t\right). \end{aligned}$$

So the analytic formula of $f(t, y)$ is

$$f(t, y) = \frac{(1 - p)b_l \exp\left(\frac{1}{\sigma}b_l u - \frac{b_l^2}{2\sigma^2}t\right) + pb_h \exp\left(\frac{1}{\sigma}b_h u - \frac{b_h^2}{2\sigma^2}t\right)}{(1 - p) \exp\left(\frac{1}{\sigma}b_l u - \frac{b_l^2}{2\sigma^2}t\right) + p \exp\left(\frac{1}{\sigma}b_h u - \frac{b_h^2}{2\sigma^2}t\right)}. \quad (4.1.6)$$

In the second case, the prior distribution of B is a 1D Gaussian distribution $\mathcal{N}(h_0, \Sigma_0)$. The corresponding measure ν is given by

$$\nu(db) = \frac{1}{\sqrt{2\pi\Sigma_0}} \exp\left[-\frac{(b - h_0)^2}{2\Sigma_0}\right] db.$$

Both integrals in (4.1.3) can be evaluated by completing the square. The integral in the denominator of (4.1.3) is

$$\begin{aligned} & \int_{-\infty}^{\infty} \exp\left(\frac{1}{\sigma}bu - \frac{b^2}{2\sigma^2}t\right) \nu(db) \\ &= \frac{1}{\sqrt{2\pi\Sigma_0}} \int_{-\infty}^{\infty} \exp\left(\frac{1}{\sigma}bu - \frac{b^2}{2\sigma^2}t\right) \exp\left[-\frac{(b - h_0)^2}{2\Sigma_0}\right] db \\ &= \frac{1}{\sqrt{2\pi\Sigma_0}} \int_{-\infty}^{\infty} \exp\left[-\left(\frac{t}{2\sigma^2} + \frac{1}{2\Sigma_0}\right)b^2 + \left(\frac{u}{\sigma} + \frac{h_0}{\Sigma_0}\right)b - \frac{h_0^2}{2\Sigma_0}\right] db \\ &= \frac{\iota}{\sqrt{\Sigma_0}} \exp\left(\frac{\kappa^2}{2\iota^2} - \frac{h_0^2}{2\Sigma_0}\right) \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi\iota}} \exp\left[-\frac{(b - \kappa)^2}{2\iota^2}\right] db, \end{aligned} \quad (4.1.7)$$

where we define

$$\frac{1}{2\iota^2} = \frac{t}{2\sigma^2} + \frac{1}{2\Sigma_0} \text{ and } 2\kappa = \frac{\frac{u}{\sigma} + \frac{h_0}{\Sigma_0}}{\frac{t}{2\sigma^2} + \frac{1}{2\Sigma_0}}.$$

Note that the integral in the last line of (4.1.7) is the integral of the probability density function of Gaussian distribution $\mathcal{N}(\kappa, \iota^2)$, and thus the integral is 1. Therefore, we have

$$\int_{-\infty}^{\infty} \exp\left(\frac{1}{\sigma}bu - \frac{b^2}{2\sigma^2}t\right) \nu(db) = \frac{\iota}{\sqrt{\Sigma_0}} \exp\left(\frac{\kappa^2}{2\iota^2} - \frac{h_0^2}{2\Sigma_0}\right).$$

The integral in the numerator is evaluated in the same way:

$$\begin{aligned} & \int_{-\infty}^{\infty} b \exp\left(\frac{1}{\sigma}bu - \frac{b^2}{2\sigma^2}t\right) \nu(db) \\ &= \frac{\iota}{\sqrt{\Sigma_0}} \exp\left(\frac{\kappa^2}{2\iota^2} - \frac{h_0^2}{2\Sigma_0}\right) \int_{-\infty}^{\infty} b \frac{1}{\sqrt{2\pi}\iota} \exp\left[-\frac{(b-\kappa)^2}{2\iota^2}\right] db. \end{aligned} \quad (4.1.8)$$

Note that the integral in the last line of (4.1.8) is the expectation of a random variable that follows Gaussian distribution $\mathcal{N}(\kappa, \iota^2)$, and thus the integral is κ . Therefore, we have

$$\int_{-\infty}^{\infty} b \exp\left(\frac{1}{\sigma}bu - \frac{b^2}{2\sigma^2}t\right) \nu(db) = \frac{\iota}{\sqrt{\Sigma_0}} \exp\left(\frac{\kappa^2}{2\iota^2} - \frac{h_0^2}{2\Sigma_0}\right) \kappa.$$

So the analytic formula of $f(t, y)$ is

$$f(t, u) = \frac{\frac{\iota}{\sqrt{\Sigma_0}} \exp\left(\frac{\kappa^2}{2\iota^2} - \frac{h_0^2}{2\Sigma_0}\right) \kappa}{\frac{\iota}{\sqrt{\Sigma_0}} \exp\left(\frac{\kappa^2}{2\iota^2} - \frac{h_0^2}{2\Sigma_0}\right)} = \kappa = \frac{\frac{u}{\sigma} + \frac{h_0}{\Sigma_0}}{\frac{t}{2\sigma^2} + \frac{1}{2\Sigma_0}}. \quad (4.1.9)$$

By (4.1.6) and (4.1.5) we derive the SDE of $\hat{H}$ when the prior distribution is a two-point discrete distribution:

$$d\hat{H}(t) = \frac{\partial}{\partial u} \left[\frac{(1-p)b_l \exp\left(\frac{1}{\sigma}b_l u - \frac{b_l^2}{2\sigma^2}t\right) + pb_h \exp\left(\frac{1}{\sigma}b_h u - \frac{b_h^2}{2\sigma^2}t\right)}{(1-p) \exp\left(\frac{1}{\sigma}b_l u - \frac{b_l^2}{2\sigma^2}t\right) + p \exp\left(\frac{1}{\sigma}b_h u - \frac{b_h^2}{2\sigma^2}t\right)} \right] \Big|_{u=U(t)} d\hat{V}(t).$$

To simplify notations, let

$$Z_1(y) = \exp\left(\frac{1}{\sigma}b_l y - \frac{b_l^2}{2\sigma^2}t\right), \quad Z_2(y) = \exp\left(\frac{1}{\sigma}b_h y - \frac{b_h^2}{2\sigma^2}t\right).$$

Then the partial derivative is represented by

$$\begin{aligned} & \frac{\partial}{\partial y} \left[\frac{(1-p)b_l \exp\left(\frac{1}{\sigma}b_l y - \frac{b_l^2}{2\sigma^2}t\right) + pb_h \exp\left(\frac{1}{\sigma}b_h y - \frac{b_h^2}{2\sigma^2}t\right)}{(1-p) \exp\left(\frac{1}{\sigma}b_l y - \frac{b_l^2}{2\sigma^2}t\right) + p \exp\left(\frac{1}{\sigma}b_h y - \frac{b_h^2}{2\sigma^2}t\right)} \right] \\ &= \frac{\partial}{\partial u} \left[\frac{(1-p)b_l Z_1(u) + pb_h Z_2(u)}{(1-p)Z_1(u) + pZ_2(u)} \right] \\ &= \frac{\left[\frac{1}{\sigma}(1-p)b_l^2 Z_1(u) + \frac{1}{\sigma}pb_h^2 Z_2(u)\right] [(1-p)Z_1(u) + pZ_2(u)]}{[(1-p)Z_1(u) + pZ_2(u)]^2} - \\ & \quad \frac{[(1-p)b_l Z_1(u) + pb_h Z_2(u)] \left[\frac{1}{\sigma}(1-p)b_l Z_1(u) + \frac{1}{\sigma}pb_h Z_2(u)\right]}{[(1-p)Z_1(u) + pZ_2(u)]^2} \\ &= \frac{1}{\sigma} \frac{(1-p)b_l^2 Z_1(u) + pb_h^2 Z_2(u)}{(1-p)Z_1(u) + pZ_2(u)} - \frac{1}{\sigma} \left[\frac{(1-p)b_l Z_1(u) + pb_h Z_2(u)}{(1-p)Z_1(u) + pZ_2(u)} \right]^2. \end{aligned}$$

Consequently, the SDE of $\hat{H}$ is

$$d\hat{H}(t) = \left\{ \frac{1}{\sigma} \frac{(1-p)b_\ell^2 Z_1(U(t)) + pb_h^2 Z_2(U(t))}{(1-p)Z_1(U(t)) + pZ_2(U(t))} - \frac{1}{\sigma} \left[\frac{(1-p)b_\ell Z_1(U(t)) + pb_h Z_2(U(t))}{(1-p)Z_1(U(t)) + pZ_2(U(t))} \right]^2 \right\} d\hat{V}(t). \quad (4.1.10)$$

By (4.1.9) and (4.1.5), we derive the SDE of $\hat{H}$ when the prior distribution is a 1D Gaussian distribution:

$$d\hat{H}(t) = \frac{\partial}{\partial u} \left(\frac{\frac{u}{\sigma} + \frac{h_0}{\Sigma_0}}{\frac{t}{\sigma^2} + \frac{1}{\Sigma_0}} \right) \Big|_{u=U(t)} d\hat{V}(t) = \frac{\frac{1}{\sigma}}{\frac{t}{\sigma^2} + \frac{1}{\Sigma_0}} d\hat{V}(t) = \frac{\frac{1}{\sigma}}{\frac{t}{\Sigma_S} + \frac{1}{\Sigma_0}} d\hat{V}(t), \quad (4.1.11)$$

where $\Sigma_S = \sigma^2$ is the variance rate of the stock.

In [18], the analysis requires that the SDE of $\hat{H}$ must have the following form:

$$d\hat{H}(t) = \hat{\mu}_{\hat{H}}(t, \hat{H}(t)) dt + \hat{\sigma}_{\hat{H}}(t, \hat{H}(t)) d\hat{V}(t). \quad (4.1.12)$$

Although (4.1.11) has this form, (4.1.10) does not follow this form. Indeed, the function $\hat{\sigma}_{\hat{H}}$ in (4.1.12) only depends on the value of $\hat{H}$ at t . But in (4.1.10), the coefficient of $d\hat{V}(t)$ depends on $U(t)$ which, by (4.1.4), is given by

$$U(t) = \hat{V}(t) + \frac{1}{\sigma} \int_0^t \hat{H}(s) ds.$$

Therefore, for any $t \in [0, T]$, this coefficient is determined by the whole path of $\hat{H}$ on interval $[0, t]$. For this reason, we avoid using the two-point discrete distribution as the prior distribution of B , and only the Gaussian prior distribution is applied. All results that are derived under the Kalman-Bucy filter can be converted into their Bayesian analogues through replacing (2.2.3) by (4.1.11).

4.2 The Bound Estimate Approach

Consider the partially observable utility maximization problem as before, and the corresponding fully observable problem is given by

$$\begin{cases} dX^\pi(t) = X^\pi(t) \left\{ \left[r(t) + \pi(t) (\hat{H}(t) - r(t)) \right] dt + \pi(t) \sigma(t) d\hat{V}(t) \right\}, X^\pi(0) = x, \\ d\hat{H}(t) = \left(\frac{\sigma^{-1}}{\Sigma_S^{-1}t + \Sigma_0^{-1}} \right) d\hat{V}(t), \\ J(t, x, h) = \sup_{\pi \in \mathcal{A}} \mathbb{E}_{t,x,h} [U(X^\pi(T))]. \end{cases} \quad (4.2.1)$$

The HJB equation of problem (4.2.1) is

$$\begin{aligned} \frac{\partial J}{\partial t} + \sup_{\pi(t) \in K} \left\{ x [r(t) + \pi(t) (h - r(t))] \frac{\partial J}{\partial x} + \frac{1}{2} x^2 \pi^2(t) \sigma^2 \frac{\partial^2 J}{\partial x^2} + \right. \\ \left. x \pi(t) \sigma \left(\frac{\sigma^{-1}}{\Sigma_S^{-1}t + \Sigma_0^{-1}} \right) \frac{\partial^2 J}{\partial x \partial h} + \frac{1}{2} \left(\frac{\sigma^{-1}}{\Sigma_S^{-1}t + \Sigma_0^{-1}} \right)^2 \frac{\partial^2 J}{\partial h^2} \right\} = 0. \end{aligned} \quad (4.2.2)$$

4.2.1 Dual Control Problems and Stochastic Maximum Principle

The dual control technique and SMP are used in the same way as in chapter 2, except that the drift of stock price must be changed into the 1D Bayesian filter. We only write what are affected by changing the filter, and all details that are not affected by this change are not presented, because they can be found in chapter 2. Theorem 4.2.1 is the main theoretical support for building the bound estimate approach.

Theorem 4.2.1. *Let $\tilde{y} \in \mathbb{R}^+$ and $v \in \mathcal{B}$. Consider the system of SDE*

$$\begin{cases} dY^{(\tilde{y}, \tilde{v})}(t) = -Y^{(\tilde{y}, \tilde{v})}(t) \left\{ [r(t) + \delta_K(\tilde{v}(t))] dt + \left[\frac{1}{\sigma(t)} (\hat{H}(t) - r(t) + \tilde{v}(t)) \right] d\hat{V}(t) \right\}, \\ d\hat{P}(t) = \hat{P}(t) \left\{ [r(t) + \pi^*(t)(\hat{H}(t) - r(t))] dt + \pi^*(t)\sigma(t)d\hat{V}(t) \right\}, \\ d\hat{H}(t) = \left(\frac{\sigma^{-1}}{\Sigma_S^{-1}t + \Sigma_0^{-1}} \right) d\hat{V}(t), \end{cases} \quad (4.2.3)$$

with initial and terminal conditions $Y^{(\tilde{y}, \tilde{v})}(0) = \tilde{y}$, $\hat{P}(T) = -\tilde{U}'(Y^{(\tilde{y}, \tilde{v})}(T))$, $\hat{H}(0) = h_0$. Then $(\tilde{y}, \tilde{v})$ is the optimal control for the dual control problem if and only if the solution to (2.3.4) (denoted by $(Y^{(\tilde{y}, \tilde{v})}, \hat{P}, \hat{Q})$) satisfies

- condition 1: the initial value of $\hat{P}$ must match with X :

$$\hat{P}(0) = X(0) = x_0,$$

- condition 2: the optimal trading strategy (i.e. π^*) of the primal problem can be shown to be given by (4.2.4) and must obey the constraints:

$$\pi^*(t) = \frac{\hat{Q}(t)}{\hat{P}(t)\sigma(t)} \in K, \quad (4.2.4)$$

- condition 3: for any $t \in [0, T]$, (4.2.5) holds almost surely:

$$\hat{P}(t)\delta_K(\tilde{v}(t)) + \frac{\hat{Q}(t)\tilde{v}(t)}{\sigma(t)} = 0, \quad \mathbb{P} - a.s. \quad (4.2.5)$$

Proof: see theorem 3.9 and 3.10 in [11].

The objective function is still (2.3.11) without any change. The goal is to find optimal values of y, π and v to minimize f , i.e. make $f(y, \pi, v)$ as close to zero as possible.

4.2.2 Numerical Experiments

Assume that $r(t)$ is a constant $r \in \mathbb{R}$ for all $t \in [0, T]$. Variables related to the test are set to be $r = 0.05, \sigma = 0.8, x_0 = 10.0, h_0 = \hat{H}(0) = 0.1, \Sigma_0 = 0.2, T = 1.0$, and the utility function U is still the power utility function (2.3.2) with $\beta = 0.5$. The discrete approximation of (4.2.3) is given by

$$\begin{cases} Y_{i+1} = Y_i - rY_i h - Y_i \left[\frac{1}{\sigma} (\hat{H}_i - r + v_i) \right] (\hat{V}_{i+1} - \hat{V}_i), & Y_0 = y, \\ \hat{P}_{i+1} = \hat{P}_i + r\hat{P}_i h + \hat{P}_i \pi_i (\hat{H}_i - r) h + \hat{P}_i \pi_i \sigma (\hat{V}_{i+1} - \hat{V}_i), & \hat{P}_0 = x_0, \\ \hat{H}_{i+1} = \hat{H}_i + \left(\frac{\sigma^{-1}}{\Sigma_S^{-1}t_i + \Sigma_0^{-1}} \right) (\hat{V}_{i+1} - \hat{V}_i), & \hat{H}_0 = h_0, \end{cases}$$

for $i = 0, 1, \dots, N - 1$. Since we observe large relative errors between $\underline{L}$ and $\bar{L}$ when $M_1 = 100$ (which is the value used in problems under the Kalman-Bucy filter), we set $M_1 = 500$ to reduce errors. The value of M_2 is still 100000 as before.

Unconstrained Problems

Two forms of parameterization for trading strategy π and the dual control variable v are tested. For problems with no constraints, i.e. $K = \mathbb{R}$, we have shown that $v_i = 0$ for all index i and thus $\{v_i\}_{i=0}^{N-1}$ does not need to be parameterized. Lower bounds ($\underline{L}$) and upper bounds ($\bar{L}$) are computed by varying h_0 and Σ_0 . Results are listed in Tables 4.1 $\sim$ 4.2. Since no closed-form solution is available, we have no column called BC and thus only the relative error between $\underline{L}$ and $\bar{L}$ can be evaluated (called Error).

Several key points can be observed from these results. First, most values in column Error are smaller than 2%, so the interval $[\underline{L}, \bar{L}]$ is narrow. Second, the width of interval $[\underline{L}, \bar{L}]$ is not sensitive to h_0 or Σ_0 , which is different from what we observe in problems under the Kalman-Bucy filter. Finally, MV increases with h_0 and Σ_0 .

h_0	Form	$\underline{L}$	$\bar{L}$	Error	OB	MV
h_0^-	I	6.4881	6.5235	0.5456%	0.0006	6.5058
h_0^-	II	6.4743	6.5222	0.7400%	0.0009	6.4982
h_0°	I	6.4833	6.5231	0.6140%	0.0008	6.5032
h_0°	II	6.4844	6.5244	0.6166%	0.0006	6.5044
h_0^+	I	6.4955	6.5671	1.1015%	0.0010	6.5313
h_0^+	II	6.5060	6.5650	0.9066%	0.0008	6.5355

Table 4.1: Numerical results for 1D unconstrained problems by varying h_0 under the Bayesian filter

Σ_0	Form	$\underline{L}$	$\bar{L}$	Error	OB	MV
Σ_0^-	I	6.4877	6.5157	0.4313%	0.0005	6.5017
Σ_0^-	II	6.4936	6.5152	0.3335%	0.0004	6.5044
Σ_0°	I	6.4941	6.5695	1.1613%	0.0012	6.5318
Σ_0°	II	6.4951	6.5698	1.1503%	0.0010	6.5324
Σ_0^+	I	6.5003	6.6721	2.6424%	0.0023	6.5862
Σ_0^+	II	6.4924	6.6793	2.8796%	0.0022	6.5859

Table 4.2: Numerical results for 1D unconstrained problems by varying Σ_0 under the Bayesian filter

Constrained Problems

For constrained problems we need to parameterize both π and v in two different forms. Also, lower bounds ($\underline{L}$) and upper bounds ($\bar{L}$) are computed by varying h_0 , Σ_0 and w . Results are listed in Tables 4.3 $\sim$ 4.5. Key points that are observed for unconstrained problems can also be observed for constrained problems. But for w we observe different effects: w^- leads to lower values of Error and OB than w° and w^+ . So w^- is a better choice than w° and w^+ for 1D constrained problems under the Bayesian filter.

h_0	Form	$\underline{L}$	$\bar{L}$	Error	OB	MV
h_0^-	I	6.4843	6.5462	0.9462%	0.0006	6.5153
h_0^-	II	6.4787	6.6229	2.1778%	0.0053	6.5508
h_0°	I	6.5065	6.5427	0.5564%	0.0003	6.5246
h_0°	II	6.4847	6.6531	2.5310%	0.0058	6.5689
h_0^+	I	6.4843	6.5710	1.3190%	0.0007	6.5276
h_0^+	II	6.5068	6.6624	2.3913%	0.0072	6.5846

Table 4.3: Numerical results for 1D constrained problems by varying h_0 under the Bayesian filter

Σ_0	Form	$\underline{L}$	$\bar{L}$	Error	OB	MV
Σ_0^-	I	6.4843	6.5555	1.0980%	0.0006	6.5199
Σ_0^-	II	6.4943	6.6879	2.9811%	0.0062	6.5911
Σ_0°	I	6.4846	6.6587	2.6848%	0.0016	6.5717
Σ_0°	II	6.5009	6.6675	2.5627%	0.0065	6.5842
Σ_0^+	I	6.5021	6.7081	3.1682%	0.0011	6.6051
Σ_0^+	II	6.5042	6.7092	3.1518%	0.0018	6.6067

Table 4.4: Numerical results for 1D constrained problems by varying Σ_0 under the Bayesian filter

w	Form	$\underline{L}$	$\bar{L}$	Error	OB	MV
w^-	I	6.5001	6.5363	0.5580%	5.9554×10^{-5}	6.5181
w^-	II	6.4950	6.5472	0.8025%	0.0003	6.5211
w°	I	6.5155	6.5786	0.9679%	0.0004	6.5470
w°	II	6.5153	6.5820	1.0241%	0.0063	6.5487
w^+	I	6.5240	6.6544	1.9987%	0.0032	6.5892
w^+	II	6.5414	6.6789	2.1031%	0.0031	6.6102

Table 4.5: Numerical results for 1D constrained problems by varying w under the Bayesian filter

4.3 The Deep Learning Approach

This section solves HJB equation (4.2.2) by PINN techniques. Ideas and settings all follow those in section 2.4, except that the HJB equation is changed into (4.2.2).

Let $\hat{J}_\theta$ be the neural network whose structure is the same as the one in Table 2.8. We use $\hat{J}_\theta$ to approximate J . The solution region is still $\mathcal{R} = [0, T] \times [X_{\min}, X_{\max}] \times [0, H_{\max}]$. Two dataset $\mathcal{D}_1 = \{T, x_i, h_i\}_{i=1}^B$ and $\mathcal{D}_2 = \{\bar{t}_i, \bar{x}_i, \bar{h}_i\}_{i=1}^B$ are generated, where $\mathcal{D}_1$ tunes $\hat{J}_\theta$ so that $\hat{J}_\theta$ fits the terminal condition well, while $\mathcal{D}_2$ tunes $\hat{J}_\theta$ so that $\hat{J}_\theta$ makes the HJB equation hold in the interior of $\mathcal{R}$. Losses with respect to $\mathcal{D}_1$ and $\mathcal{D}_2$ are denoted by $\mathcal{L}_1$ and $\mathcal{L}_2$, given by

$$\mathcal{L}_1(\theta) = \frac{1}{B} \sum_{i=1}^B \left(\hat{J}_\theta(T, x_i, h_i) - U(x_i) \right)^2 \quad (4.3.1)$$

and

$$\mathcal{L}_2(\theta) = \frac{1}{B} \sum_{i=1}^B \left[\frac{\partial \hat{J}_i}{\partial t} + \bar{x}_i [r + \pi_i^* (h - r)] \frac{\partial \hat{J}_i}{\partial x} + \frac{1}{2} \bar{x}_i^2 \pi_i^{*2} \sigma^2 \frac{\partial^2 \hat{J}_i}{\partial x^2} + \right.$$

$$x_i \pi_i^* \sigma \left(\frac{\sigma^{-1}}{\Sigma_S^{-1} \bar{t}_i + \Sigma_0^{-1}} \right) \frac{\partial^2 \hat{J}_i}{\partial x \partial h} + \frac{1}{2} \left(\frac{\sigma^{-1}}{\Sigma_S^{-1} \bar{t}_i + \Sigma_0^{-1}} \right)^2 \frac{\partial^2 \hat{J}_i}{\partial h^2} \Bigg]^2, \quad (4.3.2)$$

respectively. The total loss $\mathcal{L}$ to be minimized is $\mathcal{L} = \mathcal{L}_1 + \mathcal{L}_2$.

When training the neural network, we set $N = 100000$, $B = 2048$, $\eta = 5 \times 10^{-4}$. Both unconstrained problems and constrained problems are solved by running algo. 2. Training loss curves are shown in Figures 4.1 (unconstrained) and 4.2 (constrained) respectively. The loss becomes lower and lower during the training stage, which is what we expect.

Curves $J(0, x_0, h_0)$ (with $x_0 = 10.0$) for $\Sigma_0 = 0.1, 0.2, 0.3$ are shown in Figures 4.3 (unconstrained) and 4.4 (constrained) respectively. These curves show that the value function $J(0, x_0, h_0)$ increases with h_0 (because each single curve is an increasing function) and σ (because curves of $\sigma_0 = 0.1(k+1)$ are always above curves of $\sigma_0 = 0.1k$, $k = 1, 2$).

To compare results of the bound estimate approach and the deep learning approach, we add optimal values computed by PINN to Tables 4.1 ~ 4.4 and obtain Tables 4.6 ~ 4.9. Since BC is not available for all problems under the Bayesian filter, we define Error–PINN as the relative error between PINN and MV for all problems:

$$\text{Error – PINN} = \left| \frac{\text{PINN} - \text{MV}}{\text{MV}} \right|.$$

It can be observed that PINN incurs small errors (for most cases Error–PINN is less than 1%) and thus can accurately fit the value function J .

We also vary x and h at the same time (with $\Sigma_0 = 0.2$, $x \in [0.7x_0, 1.3x_0]$ and $h \in [0, h_{\max}]$) and draw surfaces of $J(0, x, h)$ (see Figure 4.5). These surfaces show that larger initial wealth x and larger initial guess h (of the drift) lead to greater values. This is intuitively correct, because more wealth and higher positive drift of stock returns have a higher probability to generate greater values in the stochastic environment of the financial market.

Algo. 2 PINN algorithm for HJB equation (4.2.2)

Input:

$\theta^{(0)}$: initial value of the neural network parameter

N : the number of iterations

B : the sample size parameter

η : the learning rate

Output:

$\theta^{(N)}$: the estimated optimal neural network parameter

for iteration $m = 1, 2, \dots, N$ **do**

 generate datasets $\mathcal{D}_1 = \{T, x_i, h_i\}_{i=1}^B$, $\mathcal{D}_2 = \{\bar{t}_i, \bar{x}_i, \bar{h}_i\}_{i=1}^B$

 compute the total loss function $\mathcal{L}(\theta^{(m-1)})$ through (4.3.1) and (4.3.2)

 update $\theta^{(m-1)}$ through gradient descent method and obtain $\theta^{(m)}$

end for

return $\theta^{(N)}$

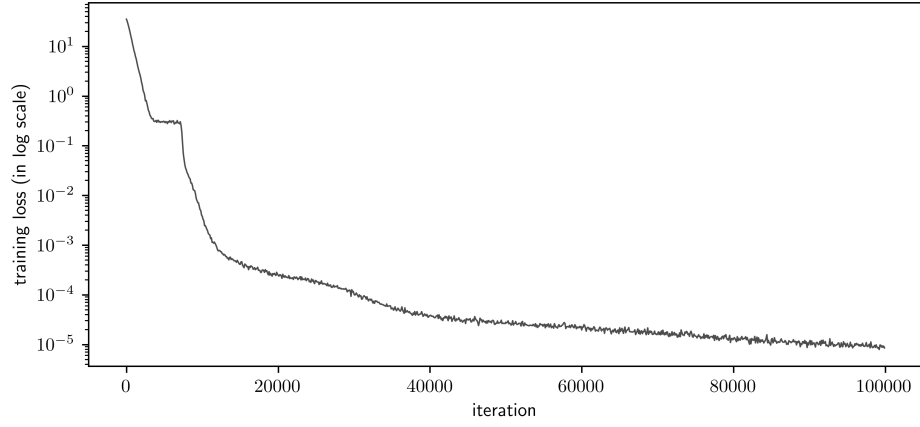


Figure 4.1: The training loss curve of solving the HJB equation with $h_0 = 0.1, \Sigma = 0.2$ under the Bayesian filter for unconstrained problems

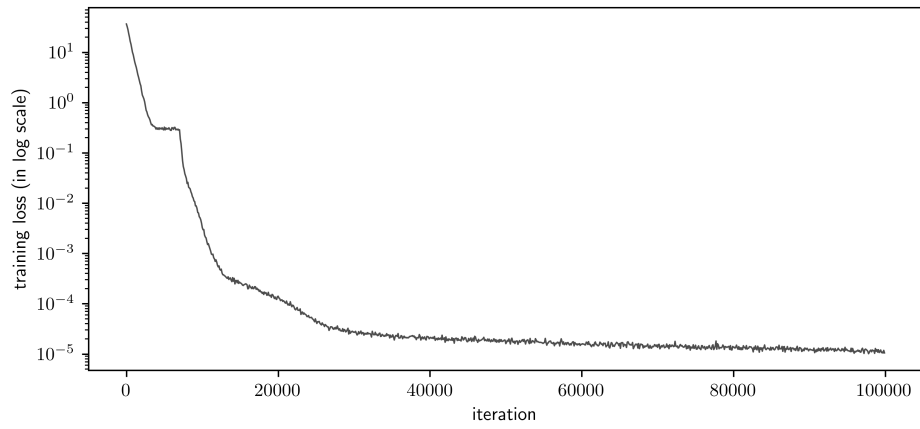


Figure 4.2: The training loss curve of solving the HJB equation with $h_0 = 0.1, \Sigma = 0.2$ under the Bayesian filter for constrained problems

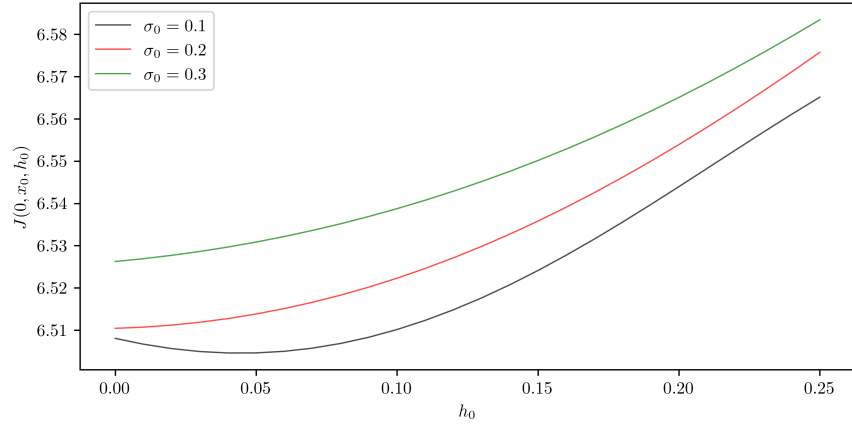


Figure 4.3: $J(0, x_0, h_0)$ curves by varying Σ_0 under the Bayesian filter for unconstrained problems

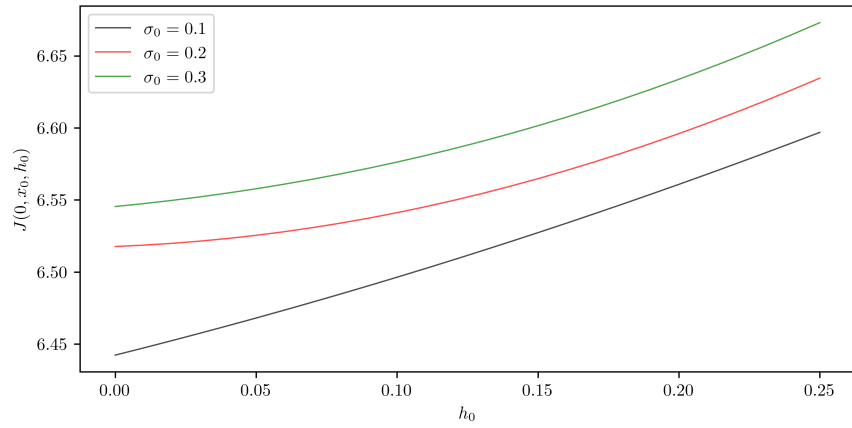


Figure 4.4: $J(0, x_0, h_0)$ curves by varying Σ_0 under the Bayesian filter for constrained problems

h_0	Form	$\underline{L}$	$\overline{L}$	Error	OB	MV	PINN	Error–PINN
h_0^-	I	6.4881	6.5235	0.5456%	0.0006	6.5058	6.5138	0.1230%
h_0^-	II	6.4743	6.5222	0.7400%	0.0009	6.4982	6.5138	0.2401%
h_0^0	I	6.4833	6.5231	0.6140%	0.0008	6.5032	6.5223	0.2937%
h_0^0	II	6.4844	6.5244	0.6166%	0.0006	6.5044	6.5223	0.2752%
h_0^+	I	6.4955	6.5671	1.1015%	0.0010	6.5313	6.5539	0.3460%
h_0^+	II	6.5060	6.5650	0.9066%	0.0008	6.5355	6.5539	0.2815%

Table 4.6: Numerical results for 1D unconstrained problems by varying h_0 under the Bayesian filter with PINN results included

Σ_0	Form	$\underline{L}$	$\overline{L}$	Error	OB	MV	PINN	Error–PINN
Σ_0^-	I	6.4877	6.5157	0.4313%	0.0005	6.5017	6.5102	0.1307%
Σ_0^-	II	6.4936	6.5152	0.3335%	0.0004	6.5044	6.5102	0.0892%
Σ_0^0	I	6.4941	6.5695	1.1613%	0.0012	6.5318	6.5223	0.1454%
Σ_0^0	II	6.4951	6.5698	1.1503%	0.0010	6.5324	6.5223	0.1546%
Σ_0^+	I	6.5003	6.6721	2.6424%	0.0023	6.5862	6.5388	0.7197%
Σ_0^+	II	6.4924	6.6793	2.8796%	0.0022	6.5859	6.5388	0.7152%

Table 4.7: Numerical results for 1D unconstrained problems by varying σ_0 under the Bayesian filter with PINN results included

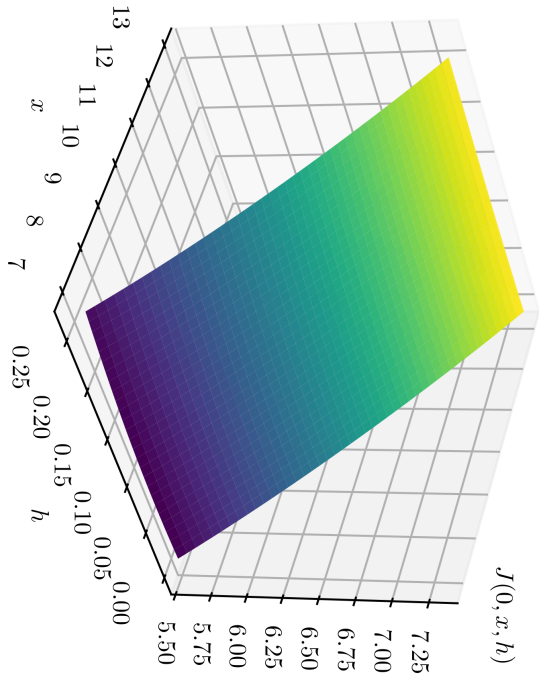
h_0	Form	$\underline{L}$	$\overline{L}$	Error	OB	MV	PINN	Error–PINN
h_0^-	I	6.4843	6.5462	0.9462%	0.0006	6.5153	6.5255	0.1566%
h_0^-	II	6.4787	6.6229	2.1778%	0.0053	6.5508	6.5255	0.3862%
h_0^0	I	6.5065	6.5427	0.5564%	0.0003	6.5246	6.5412	0.2544%
h_0^0	II	6.4847	6.6531	2.5310%	0.0058	6.5689	6.5412	0.4217%
h_0^+	I	6.4843	6.5710	1.3190%	0.0007	6.5276	6.5960	1.0479%
h_0^+	II	6.5068	6.6624	2.3913%	0.0072	6.5846	6.5960	0.1731%

Table 4.8: Numerical results for 1D constrained problems by varying h_0 under the Bayesian filter with PINN results included

Σ_0	Form	$\underline{L}$	$\overline{L}$	Error	OB	MV	PINN	Error–PINN
Σ_0^-	I	6.4843	6.5555	1.0980%	0.0006	6.5199	6.4964	0.3604%
Σ_0^-	II	6.4943	6.6879	2.9811%	0.0062	6.5911	6.4964	1.4368%
Σ_0^0	I	6.4846	6.6587	2.6848%	0.0016	6.5717	6.5412	0.4641%
Σ_0^0	II	6.5009	6.6675	2.5627%	0.0065	6.5842	6.5412	0.6531%
Σ_0^+	I	6.5021	6.7081	3.1682%	0.0011	6.6051	6.5763	0.4360%
Σ_0^+	II	6.5042	6.7092	3.1518%	0.0018	6.6067	6.5763	0.4601%

Table 4.9: Numerical results for 1D constrained problems by varying Σ_0 under the Bayesian filter with PINN results included

optimal value function at $t = 0$, i.e. $J(0, x, h)$, given by PINN



optimal value function at $t = 0$, i.e. $J(0, x, h)$, given by PINN

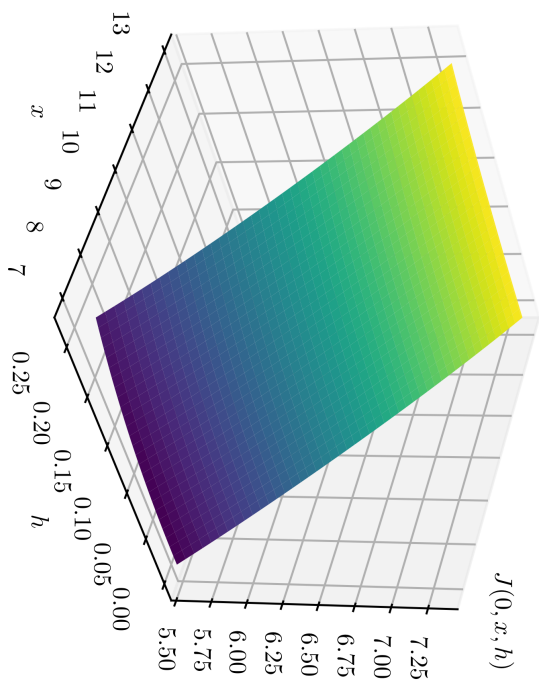


Figure 4.5: The value function surface at $t = 0$ with $\Sigma_0 = 0.2$ under the Bayesian filter
(left panel: unconstrained problems, right panel: constrained problems)

Chapter 5

2D Partially Observable Problems under Bayesian Filters

In chapter 3, we move from 1D partially observable problems to 2D partially observable problems under the Kalman-Bucy filter. Similarly, as an extension of chapter 4, this chapter mainly discusses 2D partially observable problems under the Bayesian filter. Problems whose dimension $D \geq 3$ can be solved in the same way, so we concentrate on $D = 2$. The form of 2D Bayesian filter is derived, followed by transforming partially observable problems into fully observable problems. Then the HJB equation, the dual problem and the theorem obtained by SMP are all discussed, so that the bound estimate approach can be introduced. This chapter ends with numerical experiments that test the bound estimate approach.

5.1 2D Bayesian Filters

The 2D Bayesian filter assumes that the drift process $\boldsymbol{\mu}$ of stock prices (3.1.1) (or (3.1.2)) is a 2D random vector $\mathbf{B}$, where $\mathbf{B}$ is independent of the 2D Brownian motion $\mathbf{W}$ and follows prior distribution ν such that for $d = 1, 2$, $\mathbb{E} \left[[B^{(d)}]^2 \right] < \infty$. The volatility matrix $\boldsymbol{\sigma}$ is set to be a constant as in the 1D case to ensure the analysis is mathematically tractable.

Define a 2D process $\mathbf{U}$ as

$$\mathbf{U}(t) = \boldsymbol{\sigma}^{-1} \mathbf{B}t + \mathbf{W}(t), t \in [0, T].$$

Since $\mathbf{S}$ and $\mathbf{U}$ share the same source of randomness, the filtration generated by process $\mathbf{U}$ (denoted by $\mathcal{F}^{\mathbf{U}}$) is the same as $\mathcal{F}^{\mathbf{S}}$.

Theorem 5.1.1. *Process $\mathbf{S}$ can be represented by $\mathbf{U}$ through (5.1.1):*

$$\begin{cases} S^{(1)}(t) = S^{(1)}(0) \exp \left[\left\langle \boldsymbol{\sigma}^{(1\cdot)}, \mathbf{U}(t) \right\rangle - \frac{1}{2} \left\| \boldsymbol{\sigma}^{(1\cdot)} \right\|^2 t \right] \\ S^{(2)}(t) = S^{(2)}(0) \exp \left[\left\langle \boldsymbol{\sigma}^{(2\cdot)}, \mathbf{U}(t) \right\rangle - \frac{1}{2} \left\| \boldsymbol{\sigma}^{(2\cdot)} \right\|^2 t \right] \end{cases}, \quad (5.1.1)$$

where for $d = 1, 2$, $\boldsymbol{\sigma}^{(d\cdot)}$ stands for the d^{th} row of matrix $\boldsymbol{\sigma}$, and $\langle \cdot, \cdot \rangle$ is the inner product of two vectors, i.e. $\langle \mathbf{x}, \mathbf{y} \rangle = \mathbf{x}^\top \mathbf{y} = \mathbf{y}^\top \mathbf{x}$.

Proof: The SDE of $\mathbf{U}$ is

$$d\mathbf{U}(t) = \boldsymbol{\sigma}^{-1} \mathbf{B}dt + d\mathbf{W}(t). \quad (5.1.2)$$

Left multiplying both sides of (5.1.2) by $\boldsymbol{\sigma}$ yields

$$\boldsymbol{\sigma} d\mathbf{U}(t) = \mathbf{B}dt + \boldsymbol{\sigma} d\mathbf{W}(t).$$

It follows that

$$d\mathbf{S}(t) = \text{diag}(\mathbf{S}(t)) (\mathbf{B}dt + \boldsymbol{\sigma}d\mathbf{W}(t)) = \text{diag}(\mathbf{S}(t))\boldsymbol{\sigma}d\mathbf{U}(t). \quad (5.1.3)$$

Consider the first entry in (5.1.3):

$$dS^{(1)}(t) = S^{(1)}(t) \left[\sigma^{(11)}dU^{(1)}(t) + \sigma^{(12)}dU^{(2)}(t) \right].$$

By Itô's formula, we obtain

$$\begin{aligned} d \log S^{(1)}(t) &= \sigma^{(11)}dU^{(1)}(t) + \sigma^{(12)}dU^{(2)}(t) - \frac{1}{2} \left\{ \left[\sigma^{(11)} \right]^2 + \left[\sigma^{(12)} \right]^2 \right\} dt \\ &= \sigma^{(11)}dU^{(1)}(t) + \sigma^{(12)}dU^{(2)}(t) - \frac{1}{2} \left\| \boldsymbol{\sigma}^{(1\cdot)} \right\|^2 dt. \end{aligned} \quad (5.1.4)$$

Taking integral on both sides of (5.1.4) on interval $[0, t]$, we obtain

$$\begin{aligned} \log S^{(1)}(t) &= \log S^{(1)}(0) + \sigma^{(11)}U^{(1)}(t) + \sigma^{(12)}U^{(2)}(t) - \frac{1}{2} \left\| \boldsymbol{\sigma}^{(1\cdot)} \right\|^2 t \\ &= \log S^{(1)}(0) + \left\langle \boldsymbol{\sigma}^{(1\cdot)}, \mathbf{U}(t) \right\rangle - \frac{1}{2} \left\| \boldsymbol{\sigma}^{(1\cdot)} \right\|^2 t, \end{aligned}$$

so we have

$$S^{(1)}(t) = S^{(1)}(0) \exp \left[\left\langle \boldsymbol{\sigma}^{(1\cdot)}, \mathbf{U}(t) \right\rangle - \frac{1}{2} \left\| \boldsymbol{\sigma}^{(1\cdot)} \right\|^2 t \right].$$

The second entry $S^{(2)}(t)$ can be handled in the same way. We next represent $\hat{\mathbf{H}}$ by $\mathbf{U}$ so that Itô's formula and the SDE of $\mathbf{Y}$ can be used to derive the SDE of $\hat{\mathbf{H}}$.

Theorem 5.1.2. *Let $\mathbf{g} : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ be a function with*

$$\int_{\mathbb{R}^2} |g(\mathbf{b})| \nu(d\mathbf{b}) < \infty.$$

Then the estimate of $g(\mathbf{B})$ is given by

$$\begin{aligned} \mathbb{E} [g(\mathbf{B}) | \mathcal{F}^{\mathbf{S}}(t)] &= \mathbb{E} [g(\mathbf{B}) | \mathcal{F}^{\mathbf{U}}(t)] = \mathbb{E} [g(\mathbf{B}) | \mathbf{U}(t)] \\ &= \frac{\int_{\mathbb{R}^2} g(\mathbf{b}) \exp \left(\left\langle \boldsymbol{\sigma}^{-1} \mathbf{b}, \mathbf{U}(t) \right\rangle - \frac{1}{2} \left\| \boldsymbol{\sigma}^{-1} \mathbf{b} \right\|^2 t \right) \nu(d\mathbf{b})}{\int_{\mathbb{R}^2} \exp \left(\left\langle \boldsymbol{\sigma}^{-1} \mathbf{b}, \mathbf{U}(t) \right\rangle - \frac{1}{2} \left\| \boldsymbol{\sigma}^{-1} \mathbf{b} \right\|^2 t \right) \nu(d\mathbf{b})}. \end{aligned}$$

Proof: see [5] for the proof.

Let $\mathbf{g}(\mathbf{x}) = \mathbf{x}$ for any $\mathbf{x} \in \mathbb{R}^2$ in theorem 4.1.2, we obtain the filter of $\mathbf{B}$ that is used to solve our utility maximization problem:

$$\hat{\mathbf{H}}(t) = \mathbb{E} [\mathbf{B} | \mathcal{F}^{\mathbf{S}}(t)] = \mathbf{f}(t, \mathbf{U}(t)),$$

where function $\mathbf{f} : [0, T] \times \mathbb{R}^2 \rightarrow \mathbb{R}^2$ is given by

$$\mathbf{f}(t, \mathbf{u}) = \frac{\int_{\mathbb{R}^2} \mathbf{b} \exp \left(\left\langle \boldsymbol{\sigma}^{-1} \mathbf{b}, \mathbf{u} \right\rangle - \frac{1}{2} \left\| \boldsymbol{\sigma}^{-1} \mathbf{b} \right\|^2 t \right) \nu(d\mathbf{b})}{\int_{\mathbb{R}^2} \exp \left(\left\langle \boldsymbol{\sigma}^{-1} \mathbf{b}, \mathbf{u} \right\rangle - \frac{1}{2} \left\| \boldsymbol{\sigma}^{-1} \mathbf{b} \right\|^2 t \right) \nu(d\mathbf{b})}. \quad (5.1.5)$$

It can also be shown that $\hat{\mathbf{H}}$ is a martingale with respect to measure $\mathbb{P}$ and filtration $\mathcal{F}^S$ (see [5]). In Bayesian filter, the $\mathcal{F}^S$ -measurable Brownian motion $\hat{\mathbf{V}}$ defined by (3.2.1) satisfies

$$d\hat{\mathbf{V}}(t) = d\mathbf{W}(t) + \boldsymbol{\sigma}^{-1} \left(\mathbf{B} - \hat{\mathbf{H}}(t) \right) dt.$$

Then we can write the SDE of U driven by $\hat{\mathbf{V}}$:

$$d\mathbf{U}(t) = \boldsymbol{\sigma}^{-1} \mathbf{B} dt + d\mathbf{W}(t) = \boldsymbol{\sigma}^{-1} \mathbf{B} dt + d\hat{\mathbf{V}}(t) - \boldsymbol{\sigma}^{-1} \left(\mathbf{B} - \hat{\mathbf{H}}(t) \right) dt = d\hat{\mathbf{V}}(t) + \boldsymbol{\sigma}^{-1} \hat{\mathbf{H}}(t) dt.$$

By Itô's formula and the martingale property (see [5]) of $\hat{\mathbf{H}}$, we obtain

$$\begin{aligned} d\hat{\mathbf{H}}(t) &= d\mathbf{f}(t, \mathbf{U}(t)) = \left(\left[\begin{array}{cc} \frac{\partial f^{(1)}}{\partial u^{(1)}} & \frac{\partial f^{(1)}}{\partial u^{(2)}} \\ \frac{\partial f^{(2)}}{\partial u^{(1)}} & \frac{\partial f^{(2)}}{\partial u^{(2)}} \end{array} \right] \bigg|_{u^{(1)}=U^{(1)}(t), u^{(2)}=U^{(2)}(t)} \right) d\hat{\mathbf{V}}(t) \\ &= (\mathbf{J}(\mathbf{f}, \mathbf{u})|_{\mathbf{u}=\mathbf{U}(t)}) d\hat{\mathbf{V}}(t), \end{aligned} \quad (5.1.6)$$

which is the SDE of the 2D Bayesian filter, where $\mathbf{J} \in \mathbb{R}^{2 \times 2}$ is the Jacobian matrix of $\mathbf{f}$ with respect to $\mathbf{u}$ evaluated at $\mathbf{u} = \mathbf{U}(t)$.

$$d\hat{\mathbf{V}}(t) = d\mathbf{W}(t) + \boldsymbol{\sigma}^{-1} \left(\mathbf{B} - \hat{\mathbf{H}}(t) \right) dt$$

$$d\mathbf{U}(t) = \boldsymbol{\sigma}^{-1} \mathbf{B} dt + d\mathbf{W}(t) = d\hat{\mathbf{V}}(t) + \boldsymbol{\sigma}^{-1} \hat{\mathbf{H}}(t) dt.$$

We now specify measure ν so that function $\mathbf{f}$ can be derived by evaluating (5.1.5), and then the SDE of $\hat{\mathbf{H}}$ (i.e. (5.1.6)) can be obtained. Two cases are considered. In the first case, the prior distribution of $\mathbf{B}$ is a discrete distribution with only two possible values $\mathbf{b}_\ell \in \mathbb{R}^2$ (with probability $1 - p$) and $\mathbf{b}_h \in \mathbb{R}^2$ (with probability p):

$$\mathbf{B} \sim \begin{pmatrix} \mathbf{b}_\ell & \mathbf{b}_h \\ 1 - p & p \end{pmatrix}.$$

The corresponding measure ν is represented by Dirac measures at $\mathbf{b}_\ell$ and $\mathbf{b}_h$:

$$\nu(d\mathbf{b}) = (1 - p)\delta_{\mathbf{b}_\ell}(d\mathbf{b}) + p\delta_{\mathbf{b}_h}(d\mathbf{b}),$$

where $r < b_\ell^{(1)} < b_h^{(1)}$, $r < b_\ell^{(2)} < b_h^{(2)}$ and $0 < p < 1$. The integral in the denominator of (5.1.5) is

$$\begin{aligned} & \int_{\mathbb{R}^2} \exp \left(\langle \boldsymbol{\sigma}^{-1} \mathbf{b}, \mathbf{u} \rangle - \frac{1}{2} \|\boldsymbol{\sigma}^{-1} \mathbf{b}\|^2 t \right) \nu(d\mathbf{b}) \\ &= \int_{\mathbb{R}^2} \exp \left(\langle \boldsymbol{\sigma}^{-1} \mathbf{b}, \mathbf{u} \rangle - \frac{1}{2} \|\boldsymbol{\sigma}^{-1} \mathbf{b}\|^2 t \right) [(1 - p)\delta_{\mathbf{b}_\ell}(d\mathbf{b}) + p\delta_{\mathbf{b}_h}(d\mathbf{b})] \\ &= (1 - p) \exp \left(\langle \boldsymbol{\sigma}^{-1} \mathbf{b}_\ell, \mathbf{u} \rangle - \frac{1}{2} \|\boldsymbol{\sigma}^{-1} \mathbf{b}_\ell\|^2 t \right) + p \exp \left(\langle \boldsymbol{\sigma}^{-1} \mathbf{b}_h, \mathbf{u} \rangle - \frac{1}{2} \|\boldsymbol{\sigma}^{-1} \mathbf{b}_h\|^2 t \right). \end{aligned}$$

The integral in the numerator is evaluated in a similar way:

$$\begin{aligned} & \int_{\mathbb{R}^2} \mathbf{b} \exp \left(\langle \boldsymbol{\sigma}^{-1} \mathbf{b}, \mathbf{u} \rangle - \frac{1}{2} \|\boldsymbol{\sigma}^{-1} \mathbf{b}\|^2 t \right) \nu(d\mathbf{b}) \\ &= \mathbf{b} \exp \left(\langle \boldsymbol{\sigma}^{-1} \mathbf{b}, \mathbf{u} \rangle - \frac{1}{2} \|\boldsymbol{\sigma}^{-1} \mathbf{b}\|^2 t \right) [(1 - p)\delta_{\mathbf{b}_\ell}(d\mathbf{b}) + p\delta_{\mathbf{b}_h}(d\mathbf{b})] \\ &= (1 - p) \mathbf{b}_\ell \exp \left(\langle \boldsymbol{\sigma}^{-1} \mathbf{b}_\ell, \mathbf{u} \rangle - \frac{1}{2} \|\boldsymbol{\sigma}^{-1} \mathbf{b}_\ell\|^2 t \right) + p \mathbf{b}_h \exp \left(\langle \boldsymbol{\sigma}^{-1} \mathbf{b}_h, \mathbf{u} \rangle - \frac{1}{2} \|\boldsymbol{\sigma}^{-1} \mathbf{b}_h\|^2 t \right). \end{aligned}$$

So the analytic formula of $\mathbf{f}(t, \mathbf{u})$ is

$$\mathbf{f}(t, \mathbf{u}) = \frac{(1-p)\mathbf{b}_\ell \exp(\langle \boldsymbol{\sigma}^{-1}\mathbf{b}_\ell, \mathbf{u} \rangle - \frac{1}{2}\|\boldsymbol{\sigma}^{-1}\mathbf{b}_\ell\|^2 t) + p\mathbf{b}_h \exp(\langle \boldsymbol{\sigma}^{-1}\mathbf{b}_h, \mathbf{u} \rangle - \frac{1}{2}\|\boldsymbol{\sigma}^{-1}\mathbf{b}_h\|^2 t)}{(1-p) \exp(\langle \boldsymbol{\sigma}^{-1}\mathbf{b}_\ell, \mathbf{u} \rangle - \frac{1}{2}\|\boldsymbol{\sigma}^{-1}\mathbf{b}_\ell\|^2 t) + p \exp(\langle \boldsymbol{\sigma}^{-1}\mathbf{b}_h, \mathbf{u} \rangle - \frac{1}{2}\|\boldsymbol{\sigma}^{-1}\mathbf{b}_h\|^2 t)}.$$

In the second case, the prior distribution of $\mathbf{B}$ is a 2D Gaussian distribution $\mathcal{N}(\mathbf{h}_0, \boldsymbol{\Sigma}_0)$. The corresponding measure ν is given by

$$\nu(d\mathbf{b}) = \frac{1}{2\pi\sqrt{\det(\boldsymbol{\Sigma}_0)}} \exp\left[-\frac{1}{2}(\mathbf{b} - \mathbf{h}_0)^\top \boldsymbol{\Sigma}_0^{-1}(\mathbf{b} - \mathbf{h}_0)\right] d\mathbf{b}.$$

Both integrals in (5.1.5) can be evaluated by completing the square. The integral in the denominator of (5.1.5) is

$$\begin{aligned} & \int_{\mathbb{R}^2} \exp\left(\langle \boldsymbol{\sigma}^{-1}\mathbf{b}, \mathbf{u} \rangle - \frac{1}{2}\|\boldsymbol{\sigma}^{-1}\mathbf{b}\|^2 t\right) \nu(d\mathbf{b}) \\ &= \frac{1}{2\pi\sqrt{\det(\boldsymbol{\Sigma}_0)}} \int_{\mathbb{R}^2} \exp\left(\langle \boldsymbol{\sigma}^{-1}\mathbf{b}, \mathbf{u} \rangle - \frac{1}{2}\|\boldsymbol{\sigma}^{-1}\mathbf{b}\|^2 t\right) \exp\left[-\frac{1}{2}(\mathbf{b} - \mathbf{h}_0)^\top \boldsymbol{\Sigma}_0^{-1}(\mathbf{b} - \mathbf{h}_0)\right] d\mathbf{b} \\ &= \frac{1}{2\pi\sqrt{\det(\boldsymbol{\Sigma}_0)}} \int_{\mathbb{R}^2} \exp\left[-\frac{1}{2}\mathbf{b}^\top (\boldsymbol{\Sigma}_0^{-1} + \boldsymbol{\sigma}^{-\top}\boldsymbol{\sigma}^{-1}t)\mathbf{b} + (\mathbf{u}^\top \boldsymbol{\sigma}^{-1} + \mathbf{h}_0^\top \boldsymbol{\Sigma}_0^{-1})\mathbf{b} - \frac{1}{2}\mathbf{h}_0^\top \boldsymbol{\Sigma}_0^{-1}\mathbf{h}_0\right] d\mathbf{b} \\ &= \frac{1}{2\pi\sqrt{\det(\boldsymbol{\Sigma}_0)}} \int_{\mathbb{R}^2} \exp\left(-\frac{1}{2}\mathbf{b}^\top \tilde{\boldsymbol{\Sigma}}^{-1}\mathbf{b} + \tilde{\mathbf{h}}^\top \tilde{\boldsymbol{\Sigma}}^{-1}\mathbf{b} - \frac{1}{2}\mathbf{h}_0^\top \boldsymbol{\Sigma}_0^{-1}\mathbf{h}_0\right) d\mathbf{b} \\ &= \sqrt{\frac{\det(\tilde{\boldsymbol{\Sigma}})}{\det(\boldsymbol{\Sigma}_0)}} \exp\left[\frac{1}{2}(\tilde{\mathbf{h}}^\top \tilde{\boldsymbol{\Sigma}}^{-1}\tilde{\mathbf{h}} - \mathbf{h}_0^\top \boldsymbol{\Sigma}_0^{-1}\mathbf{h}_0)\right] \int_{\mathbb{R}^2} \frac{1}{2\pi\sqrt{\det(\tilde{\boldsymbol{\Sigma}})}} \exp\left[-\frac{1}{2}(\mathbf{b} - \tilde{\mathbf{h}})^\top \tilde{\boldsymbol{\Sigma}}^{-1}(\mathbf{b} - \tilde{\mathbf{h}})\right] d\mathbf{b} \end{aligned} \quad (5.1.7)$$

where we define

$$\begin{cases} \tilde{\boldsymbol{\Sigma}}^{-1} = \boldsymbol{\Sigma}_0^{-1} + \boldsymbol{\sigma}^{-\top}\boldsymbol{\sigma}^{-1}t \\ \tilde{\mathbf{h}}^\top \tilde{\boldsymbol{\Sigma}}^{-1} = \mathbf{u}^\top \boldsymbol{\sigma}^{-1} + \mathbf{h}_0^\top \boldsymbol{\Sigma}_0^{-1} \end{cases} \implies \begin{cases} \tilde{\boldsymbol{\Sigma}} = (\boldsymbol{\Sigma}_0^{-1} + \boldsymbol{\sigma}^{-\top}\boldsymbol{\sigma}^{-1}t)^{-1} \\ \tilde{\mathbf{h}} = \tilde{\boldsymbol{\Sigma}}(\mathbf{u}^\top \boldsymbol{\sigma}^{-1} + \mathbf{h}_0^\top \boldsymbol{\Sigma}_0^{-1})^\top. \end{cases}$$

Note that the integral in the last line of (5.1.7) is the integral of the probability density function of 2D Gaussian distribution $\mathcal{N}(\tilde{\mathbf{h}}, \tilde{\boldsymbol{\Sigma}})$, and thus the integral is 1. Therefore, we have

$$\int_{\mathbb{R}^2} \exp\left(\langle \boldsymbol{\sigma}^{-1}\mathbf{b}, \mathbf{u} \rangle - \frac{1}{2}\|\boldsymbol{\sigma}^{-1}\mathbf{b}\|^2 t\right) \nu(d\mathbf{b}) = \sqrt{\frac{\det(\tilde{\boldsymbol{\Sigma}})}{\det(\boldsymbol{\Sigma}_0)}} \exp\left[\frac{1}{2}(\tilde{\mathbf{h}}^\top \tilde{\boldsymbol{\Sigma}}^{-1}\tilde{\mathbf{h}} - \mathbf{h}_0^\top \boldsymbol{\Sigma}_0^{-1}\mathbf{h}_0)\right].$$

The integral in the numerator is evaluated in the same way:

$$\begin{aligned} & \int_{\mathbb{R}^2} \mathbf{b} \exp\left(\langle \boldsymbol{\sigma}^{-1}\mathbf{b}, \mathbf{u} \rangle - \frac{1}{2}\|\boldsymbol{\sigma}^{-1}\mathbf{b}\|^2 t\right) \nu(d\mathbf{b}) \\ &= \sqrt{\frac{\det(\tilde{\boldsymbol{\Sigma}})}{\det(\boldsymbol{\Sigma}_0)}} \exp\left[\frac{1}{2}(\tilde{\mathbf{h}}^\top \tilde{\boldsymbol{\Sigma}}^{-1}\tilde{\mathbf{h}} - \mathbf{h}_0^\top \boldsymbol{\Sigma}_0^{-1}\mathbf{h}_0)\right] \int_{\mathbb{R}^2} \mathbf{b} \frac{1}{2\pi\sqrt{\det(\tilde{\boldsymbol{\Sigma}})}} \exp\left[-\frac{1}{2}(\mathbf{b} - \tilde{\mathbf{h}})^\top \tilde{\boldsymbol{\Sigma}}^{-1}(\mathbf{b} - \tilde{\mathbf{h}})\right] d\mathbf{b}. \end{aligned} \quad (5.1.8)$$

Note that the integral in the last line of (5.1.8) is the expectation of a random vector that follows 2D Gaussian distribution $\mathcal{N}(\tilde{\mathbf{h}}, \tilde{\boldsymbol{\Sigma}})$, and thus the integral is $\tilde{\mathbf{h}}$. Therefore, we have

$$\int_{\mathbb{R}^2} \mathbf{b} \exp\left(\langle \boldsymbol{\sigma}^{-1}\mathbf{b}, \mathbf{u} \rangle - \frac{1}{2}\|\boldsymbol{\sigma}^{-1}\mathbf{b}\|^2 t\right) \nu(d\mathbf{b}) = \tilde{\mathbf{h}} \sqrt{\frac{\det(\tilde{\boldsymbol{\Sigma}})}{\det(\boldsymbol{\Sigma}_0)}} \exp\left[\frac{1}{2}(\tilde{\mathbf{h}}^\top \tilde{\boldsymbol{\Sigma}}^{-1}\tilde{\mathbf{h}} - \mathbf{h}_0^\top \boldsymbol{\Sigma}_0^{-1}\mathbf{h}_0)\right].$$

So the analytic formula of $\mathbf{f}(t, \mathbf{u})$ is

$$\begin{aligned} \mathbf{f}(t, \mathbf{u}) &= \frac{\tilde{\mathbf{h}} \sqrt{\frac{\det(\tilde{\Sigma})}{\det(\Sigma_0)}} \exp \left[\frac{1}{2} \left(\tilde{\mathbf{h}}^\top \tilde{\Sigma}^{-1} \tilde{\mathbf{h}} - \mathbf{h}_0^\top \Sigma_0^{-1} \mathbf{h}_0 \right) \right]}{\sqrt{\frac{\det(\tilde{\Sigma})}{\det(\Sigma_0)}} \exp \left[\frac{1}{2} \left(\tilde{\mathbf{h}}^\top \tilde{\Sigma}^{-1} \tilde{\mathbf{h}} - \mathbf{h}_0^\top \Sigma_0^{-1} \mathbf{h}_0 \right) \right]} = \tilde{\mathbf{h}} \\ &= \tilde{\Sigma} \left(\mathbf{u}^\top \boldsymbol{\sigma}^{-1} + \mathbf{h}_0^\top \Sigma_0^{-1} \right)^\top = \tilde{\Sigma} \left(\boldsymbol{\sigma}^{-\top} \mathbf{u} + \Sigma_0^{-\top} \mathbf{h}_0 \right). \end{aligned} \quad (5.1.9)$$

For the same reason as that given in chapter 4, we only apply the Gaussian prior distribution. By (5.1.9) and (5.1.6), we derive the SDE of $\hat{\mathbf{H}}$:

$$d\hat{\mathbf{H}}(t) = (\mathbf{J}(\mathbf{f}, \mathbf{u})|_{\mathbf{y}=\mathbf{U}(t)}) d\hat{\mathbf{V}}(t) = \tilde{\Sigma} \boldsymbol{\sigma}^{-\top} d\hat{\mathbf{V}}(t) = \left(\Sigma_0^{-1} + \boldsymbol{\sigma}^{-\top} \boldsymbol{\sigma}^{-1} t \right)^{-1} \boldsymbol{\sigma}^{-\top} d\hat{\mathbf{V}}(t). \quad (5.1.10)$$

All results that are derived under the Kalman-Bucy filter can be converted into their Bayesian analogues through replacing (3.2.3) by (5.1.10).

5.2 The Bound Estimate Approach

Consider the partially observable utility maximization problem as before, and the corresponding fully observable problem is given by

$$\begin{cases} dX^\pi(t) = X^\pi(t) \left\{ \left[r(t) + \boldsymbol{\pi}^\top(t) \left(\hat{\mathbf{H}}(t) - r(t)\mathbf{1} \right) \right] dt + \boldsymbol{\pi}^\top(t) \boldsymbol{\sigma}(t) d\hat{\mathbf{V}}(t) \right\}, X^\pi(0) = x, \\ d\hat{\mathbf{H}}(t) = \left(\Sigma_0^{-1} + \boldsymbol{\sigma}^{-\top} \boldsymbol{\sigma}^{-1} t \right)^{-1} \boldsymbol{\sigma}^{-\top} d\hat{\mathbf{V}}(t), \hat{\mathbf{H}}(0) = \mathbf{h}_0, \\ J(t, x, \mathbf{h}) = \sup_{\boldsymbol{\pi} \in \mathcal{A}} \mathbb{E}_{t,x,\mathbf{h}} [U(X^\pi(T))]. \end{cases} \quad (5.2.1)$$

The HJB equation of problem (5.2.1) is

$$\begin{aligned} \frac{\partial J}{\partial t} + \sup_{\boldsymbol{\pi}(t) \in K} \left\{ x \left[r(t) + \boldsymbol{\pi}^\top(t) (\mathbf{h} - r(t)\mathbf{1}) \right] \frac{\partial J}{\partial x} + \frac{1}{2} x^2 \boldsymbol{\pi}^\top(t) \boldsymbol{\sigma} \boldsymbol{\sigma}^\top \boldsymbol{\pi}(t) \frac{\partial^2 J}{\partial x^2} + \right. \\ \left. x \boldsymbol{\pi}^\top(t) \left(\Sigma_0^{-1} + \boldsymbol{\sigma}^{-\top} \boldsymbol{\sigma}^{-1} t \right)^{-\top} \frac{\partial^2 J}{\partial x \partial \mathbf{h}} + \right. \\ \left. \frac{1}{2} \text{Tr} \left[\left(\Sigma_0^{-1} + \boldsymbol{\sigma}^{-\top} \boldsymbol{\sigma}^{-1} t \right)^{-1} \boldsymbol{\sigma}^{-\top} \boldsymbol{\sigma}^{-1} \left(\Sigma_0^{-1} + \boldsymbol{\sigma}^{-\top} \boldsymbol{\sigma}^{-1} t \right)^{-\top} \frac{\partial^2 J}{\partial \mathbf{h}^2} \right] \right\} = 0. \end{aligned}$$

5.2.1 Dual Control Problems and Stochastic Maximum Principle

The dual control technique and SMP are used in the same way as in chapter 3, except that the drift of stock prices must be changed into the 2D Bayesian filter. We only write what are affected by changing the filter, and all details that are not affected by this change are not presented, because they can be found in chapter 3. Theorem 5.2.1 is the main theoretical support for building the bound estimate approach.

Theorem 5.2.1. *Let $\tilde{\mathbf{y}} \in \mathbb{R}^+$ and $\tilde{\mathbf{v}} \in \mathcal{B}$. Consider the system of SDE*

$$\begin{cases} dY^{(\tilde{\mathbf{y}}, \tilde{\mathbf{v}})}(t) = -Y^{(\tilde{\mathbf{y}}, \tilde{\mathbf{v}})}(t) \left\{ \left[r(t) + \delta_K(\tilde{\mathbf{v}}(t)) \right] dt + \left[\boldsymbol{\sigma}^{-1}(t) (\hat{\mathbf{H}}(t) - r(t)\mathbf{1} + \tilde{\mathbf{v}}(t)) \right]^\top d\hat{\mathbf{V}}(t) \right\}, \\ d\hat{P}(t) = \hat{P}(t) \left\{ \left[r(t) + \boldsymbol{\pi}^{*\top}(t) \left(\hat{\mathbf{H}}(t) - r(t)\mathbf{1} \right) \right] dt + \boldsymbol{\pi}^{*\top}(t) \boldsymbol{\sigma}(t) d\hat{\mathbf{V}}(t) \right\}, \\ d\hat{\mathbf{H}}(t) = \left(\Sigma_0^{-1} + \boldsymbol{\sigma}^{-\top} \boldsymbol{\sigma}^{-1} t \right)^{-1} \boldsymbol{\sigma}^{-\top} d\hat{\mathbf{V}}(t), \end{cases} \quad (5.2.2)$$

with initial or terminal conditions $Y^{(\tilde{y}, \tilde{v})}(0) = \tilde{y}, \hat{P}(T) = -\tilde{U}'(Y^{(\tilde{y}, \tilde{v})}(T)), \hat{H}(0) = \mathbf{h}_0$. Then $(\tilde{y}, \tilde{v})$ is the optimal control for the dual control problem if and only if the solution to (5.2.2) (denoted by $(Y^{(\tilde{y}, \tilde{v})}, \hat{P}, \hat{Q})$) satisfies

- condition 1: the initial value of $\hat{P}$ must match with X :

$$\hat{P}(0) = X(0) = x_0,$$

- condition 2: the optimal trading strategy (i.e. π^*) of the primal problem can be shown to be given by (5.2.3) and must obey constraints K :

$$\pi^*(t) = \frac{\sigma^{-\top}(t)\hat{Q}(t)}{\hat{P}(t)} \in K, \quad (5.2.3)$$

- condition 3: for any $t \in [0, T]$, (5.2.4) holds almost surely:

$$\hat{P}(t)\delta_K(\tilde{v}(t)) + \hat{Q}^\top(t)\sigma^{-1}(t)\tilde{v}(t) = 0, \quad \mathbb{P} - a.s. \quad (5.2.4)$$

Proof: see theorem 3.9 and 3.10 in [11].

The objective function is still (3.3.7) without any change. The goal is to find optimal values of y, π and v to minimize f , i.e. make $f(y, \pi, v)$ as close to zero as possible.

5.2.2 Numerical Experiments

Assume that $r(t)$ is a constant $r \in \mathbb{R}$ for all $t \in [0, T]$. Variables related to the test are set to be $r = 0.05$, $x_0 = 10.0$, $\rho = 0.0$, $T = 1.0$, and

$$\mathbf{h}_0 = \begin{bmatrix} 0.1 \\ 0.1 \end{bmatrix}, \bar{\mathbf{H}} = \begin{bmatrix} 0.1 \\ 0.1 \end{bmatrix},$$

$$\boldsymbol{\lambda} = \begin{bmatrix} 1.0 & 0 \\ 0 & 1.0 \end{bmatrix}, \boldsymbol{\sigma} = \begin{bmatrix} 0.8 & 0 \\ 0 & 0.8 \end{bmatrix}, \boldsymbol{\sigma}_H = \begin{bmatrix} 0.5 & 0 \\ 0 & 0.5 \end{bmatrix}, \boldsymbol{\Sigma}_0 = \begin{bmatrix} 0.2 & 0.0 \\ 0.0 & 0.2 \end{bmatrix}.$$

The utility function U is set to be the power utility function (2.3.2) with $\beta = 0.5$. The discrete approximation of (5.2.2) is given by

$$\begin{cases} Y_{i+1} = Y_i - rY_i h - Y_i \left[\sigma^{-1} \left(\hat{\mathbf{H}}_i - r\mathbf{1} + \mathbf{v}_i \right) \right]^\top \left(\hat{\mathbf{V}}_{i+1} - \hat{\mathbf{V}}_i \right), & Y_0 = y, \\ \hat{P}_{i+1} = \hat{P}_i + r\hat{P}_i h + \hat{P}_i \pi_i^\top \left(\hat{\mathbf{H}}_i - r\mathbf{1} \right) h + \hat{P}_i \pi_i^\top \sigma \left(\hat{\mathbf{V}}_{i+1} - \hat{\mathbf{V}}_i \right), & \hat{P}_0 = x_0, \\ \hat{\mathbf{H}}_{i+1} = \hat{\mathbf{H}}_i + \left(\boldsymbol{\Sigma}_0^{-1} + \sigma^{-\top} \sigma^{-1} t_i \right)^{-1} \sigma^{-\top} \left(\hat{\mathbf{V}}_{i+1} - \hat{\mathbf{V}}_i \right), & \hat{\mathbf{H}}_0 = \mathbf{h}_0, \end{cases}$$

for $i = 0, 1, \dots, N-1$. Since we observe large relative errors between $\underline{L}$ and $\bar{L}$ when $M_1 = 100$ (which is the value used in problems under the Kalman-Bucy filter), we set $M_1 = 500$ to reduce errors. The value of M_2 is still 100000 as before.

Unconstrained Problems

Two forms of parameterization for trading strategy π and the dual control variable v are tested. For problems with no constraints, i.e. $K = \mathbb{R}^2$, we have shown that $v_i = 0$ for all index i and thus $\{v_i\}_{i=0}^{N-1}$ does not need to be parameterized. Lower bounds ($\underline{L}$) and upper bounds ($\bar{L}$) are computed by varying $\mathbf{h}_0$ and $\boldsymbol{\Sigma}_0$. Results are listed in Tables 5.1

~ 5.2 . Since no closed-form solution is available, we have no column called BC and thus only the relative error between $\underline{L}$ and $\bar{L}$ can be evaluated (called Error).

First, most values in column Error are smaller than 2%, so the interval $[\underline{L}, \bar{L}]$ is narrow. Second, the width of interval $[\underline{L}, \bar{L}]$ is not sensitive to Σh_0 or Σ_0 , which is different from what we observe in problems under the Kalman-Bucy filter. Finally, MV increases with h_0 and Σ_0 .

h_0	Form	$\underline{L}$	$\bar{L}$	Error	OB	MV
h_0^-	I	6.4005	6.5516	2.3605%	0.0022	6.4760
h_0^-	II	6.4063	6.5863	2.8091%	0.0027	6.4963
h_0°	I	6.4993	6.5804	1.2498%	0.0010	6.5399
h_0°	II	6.5005	6.6037	1.5867%	0.0011	6.5521
h_0^+	I	6.7313	6.8581	1.8846%	0.0005	6.7947
h_0^+	II	6.7316	6.8658	1.9931%	0.0006	6.7987

Table 5.1: Numerical results for 2D unconstrained problems by varying h_0 under the Bayesian filter

Σ_0	Form	$\underline{L}$	$\bar{L}$	Error	OB	MV
Σ_0^-	I	6.4972	6.5877	1.3929%	0.0010	6.5425
Σ_0^-	II	6.4953	6.5874	1.4177%	0.0011	6.5414
Σ_0°	I	6.5107	6.5860	1.1559%	0.0009	6.5483
Σ_0°	II	6.5132	6.5898	1.1755%	0.0009	6.5515
Σ_0^+	I	6.4977	6.5923	1.4573%	0.0010	6.5450
Σ_0^+	II	6.4990	6.5820	1.2771%	0.0010	6.5405

Table 5.2: Numerical results for 2D unconstrained problems by varying Σ_0 under the Bayesian filter

Constrained Problems

For constrained problems we need to parameterize both π and v in two different forms. Also, lower bounds ($\underline{L}$) and upper bounds ($\bar{L}$) are computed by varying h_0 , Σ_0 and w . Results are listed in Tables 5.3 $\sim$ 5.5. Key points that are observed for unconstrained problems can also be observed for constrained problems. For w we observe that larger values of w cause larger relative errors between $\underline{L}$ and $\bar{L}$, so for 2D problems under the Bayesian filter, choosing $w = w^-$ is better than w° and w^+ .

h_0	Form	$\underline{L}$	$\bar{L}$	Error	OB	MV
h_0^-	I	6.3977	6.5511	2.3967%	0.0002	6.4744
h_0^-	II	6.4033	6.5591	2.4323%	0.0003	6.4812
h_0°	I	6.5043	6.5998	1.4674%	0.0002	6.5520
h_0°	II	6.4988	6.5950	1.4810%	0.0002	6.5469
h_0^+	I	6.7297	6.8603	1.9406%	0.0003	6.7950
h_0^+	II	6.7326	6.8536	1.7966%	0.0002	6.7931

Table 5.3: Numerical results for 2D constrained problems by varying h_0 under the Bayesian filter

Σ_0	Form	$\underline{L}$	$\overline{L}$	Error	OB	MV
Σ_0^-	I	6.5029	6.5842	1.2523%	0.0001	6.5436
Σ_0^-	II	6.5073	6.6040	1.4865%	0.0003	6.5557
Σ_0°	I	6.4928	6.5837	1.3989%	0.0002	6.5382
Σ_0°	II	6.4986	6.5881	1.3779%	0.0005	6.5433
Σ_0^+	I	6.4984	6.5829	1.2999%	0.0013	6.5407
Σ_0^+	II	6.4959	6.5990	1.5883%	0.0008	6.5475

Table 5.4: Numerical results for 2D constrained problems by varying Σ_0 under the Bayesian filter

w	Form	$\underline{L}$	$\overline{L}$	Error	OB	MV
w^-	I	6.5288	6.5791	0.7703%	6.0387×10^{-5}	6.5539
w^-	II	6.5039	6.5749	1.0913%	0.0001	6.5394
w°	I	6.5003	6.5758	1.1625%	0.0003	6.5380
w°	II	6.5077	6.5826	1.1509%	0.0003	6.5452
w^+	I	6.4982	6.5893	1.4018%	0.0006	6.5438
w^+	II	6.4987	6.5971	1.5149%	0.0004	6.5479

Table 5.5: Numerical results for 2D constrained problems by varying w under the Bayesian filter

Chapter 6

Information Value

Intuitively, people who can obtain more information tend to make better decisions than others. This is also true in financial market. Informed traders typically have private or superior information that give them extra advantages over other traders (see [3], for example). This chapter numerically shows that such an intuition is correct for our utility maximization problem by comparing value functions for both 1D partially observable problems and 1D fully observable problems. If information really has positive value, then the full information value must be greater than the partial information value. All numerical experiments in this chapter are done by PINN techniques.

6.1 Fully Observable Problems

The first step of measuring information value is solving fully observable problems. If the drift $H = \mu$ is fully observable, then the corresponding filter must have zero variance, i.e. $\Sigma(t) = 0$ for any $t \in (0, T]$. In fact, the filter is no longer needed, and it is enough to consider H only. The initial value $H(0)$ follows Gaussian distribution $\mathcal{N}(h_0, \Sigma_0)$.

The HJB equation of fully observable problems can be obtained by setting $\Sigma(t) = 0$ in (2.3.1), which leads to

$$\begin{aligned} \frac{\partial J^f}{\partial t} + \sup_{\pi(t) \in K} \left\{ x[r(t) + \pi(t)(h - r(t))] \frac{\partial J^f}{\partial x} + \frac{1}{2} x^2 \pi^2(t) \sigma^2(t) \frac{\partial^2 J^f}{\partial x^2} - \lambda(h - \bar{H}) \frac{\partial J^f}{\partial h} + \right. \\ \left. \rho \sigma_H \sigma(t) x \pi(t) \frac{\partial^2 J^f}{\partial x \partial h} + \frac{1}{2} \rho^2 \sigma_H^2 \frac{\partial^2 J^f}{\partial h^2} \right\} = 0, \end{aligned} \quad (6.1.1)$$

where J^f stands for the value function of fully observable problems. The optimal trading strategy $\pi^*(t)$ maximizes

$$S(\pi(t)) = \left(\frac{1}{2} x^2 \sigma^2(t) \frac{\partial^2 J^f}{\partial x^2} \right) \pi^2(t) + \left[x(h - r(t)) \frac{\partial J^f}{\partial x} + \rho \sigma_H \sigma(t) x \frac{\partial^2 J^f}{\partial x \partial h} \right] \pi(t),$$

so $\pi^*(t)$ is

$$\pi^*(t) = \begin{cases} - \frac{(h - r(t)) \frac{\partial J^f}{\partial x} + \rho \sigma_H \sigma(t) \frac{\partial^2 J^f}{\partial x \partial h}}{x \sigma^2(t) \frac{\partial^2 J^f}{\partial x^2}}, & K = \mathbb{R}, \\ \left(- \frac{(h - r(t)) \frac{\partial J^f}{\partial x} + \rho \sigma_H \sigma(t) \frac{\partial^2 J^f}{\partial x \partial h}}{x \sigma^2(t) \frac{\partial^2 J^f}{\partial x^2}}, 0 \right)^+, & K = \mathbb{R}^+. \end{cases}$$

Then (6.1.1) becomes a standard PDE

$$\frac{\partial J^f}{\partial t} + x[r(t) + \pi^*(t)(h - r(t))] \frac{\partial J^f}{\partial x} + \frac{1}{2} x^2 \pi^{*2}(t) \sigma^2(t) \frac{\partial^2 J^f}{\partial x^2} - \lambda(h - \bar{H}) \frac{\partial J^f}{\partial h} +$$

$$\rho\sigma_H\sigma(t)x\pi^*(t)\frac{\partial^2 J^f}{\partial x\partial h} + \frac{1}{2}\rho^2\sigma_H^2\frac{\partial^2 J^f}{\partial h^2} = 0. \quad (6.1.2)$$

The value function $J^f(t, x, h)$ can be determined by solving (6.1.2).

One problem is that $H(0) \sim \mathcal{N}(h_0, \Sigma_0)$ is a random variable, meaning that the value at 0 (i.e. $J^f(0, x_0, H(0))$) is also random. However, the partial information value $J(0, x_0, h_0)$ is a deterministic quantity. For this reason, $J^f(0, x_0, H(0))$ cannot be directly compared with $J(0, x_0, h_0)$. Therefore, we take expectation over the distribution of $H(0)$ and thus the full information value is given by $\mathbb{E}[J^f(0, x_0, H(0))]$. Then the value of information that are not observable is the following difference:

$$\mathbb{E}_{H(0) \sim \mathcal{N}(h_0, \Sigma_0)} [J^f(0, x_0, H(0))] - J(0, x_0, h_0). \quad (6.1.3)$$

The full information value can be estimated by Monte Carlo methods.

Similar analysis can be done for 1D problems under the Bayesian filter. In this case, the HJB equation of partially observable problems is (4.2.2). By setting $\Sigma(t) = 0, t \in (0, T]$, we obtain the HJB equation of fully observable problems. But (4.2.2) does not contain $\Sigma(t) = 0, t \in (0, T]$, so the HJB equation of fully observable problems is still (4.2.2). However, this does not mean that the full information value is equal to the partial information value. This is because in (6.1.3), even when $J^f = J$, the difference is still not zero due to the expectation operator.

6.2 Numerical Experiments

We test the Kalman-Bucy filter and the Bayesian filter for unconstrained problems and constrained problems. The initial mean h_0 takes its value from $\{0.1, 0.2, \dots, 0.9\}$ and the initial variance is fixed at $\Sigma_0 = 0.2$. Other variables are all the same as those used in previous chapters. For each value of h_0 , all HJB equations involved are solved by PINN techniques.

When neural networks are trained, we set $N = 200000, B = 1024, \eta = 1 \times 10^{-4}$. The variation of $\mathbb{E}_{H(0) \sim \mathcal{N}(h_0, \Sigma_0)} [J^f(0, x_0, H(0))]$ and $J(0, x_0, h_0)$ with respect to h_0 are drawn together (see Figures 6.1 ~ 6.4). To gain enough accuracy, when we use Monte Carlo methods to estimate the full information value, the sample size is set to be 100000.

Results show that full information always have greater value than partial information (the black curve is above the red curve in all cases tested), which is consistent with the intuition. The difference between these two curves measures how valuable the unobservable information is. It should be noted that, in fact, inequality

$$\mathbb{E}_{H(0) \sim \mathcal{N}(h_0, \Sigma_0)} [J^f(0, x_0, H(0))] \geq J(0, x_0, h_0)$$

can be proved theoretically (see [16] for the proof). But we find no published theoretical results on how to evaluate the difference $\mathbb{E}_{H(0) \sim \mathcal{N}(h_0, \Sigma_0)} [J^f(0, x_0, H(0))] - J(0, x_0, h_0)$.

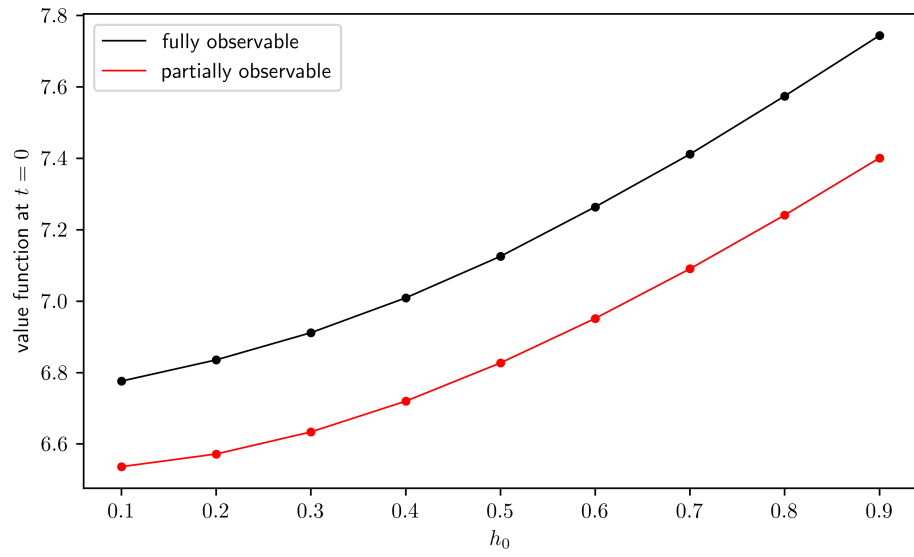


Figure 6.1: Full information value and partial information value at $t = 0$
(Kalman-Bucy filter for unconstrained problems)

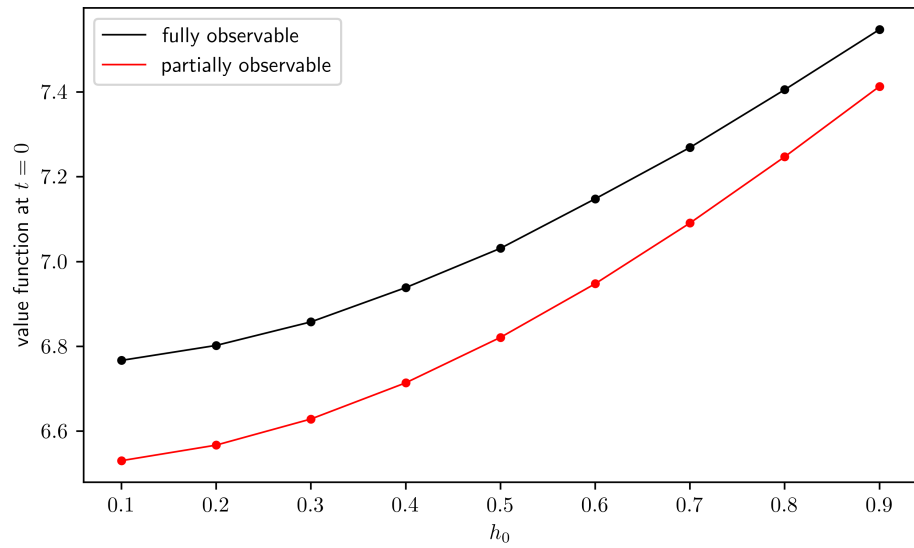


Figure 6.2: Full information value and partial information value at $t = 0$
(Kalman-Bucy filter for constrained problems)

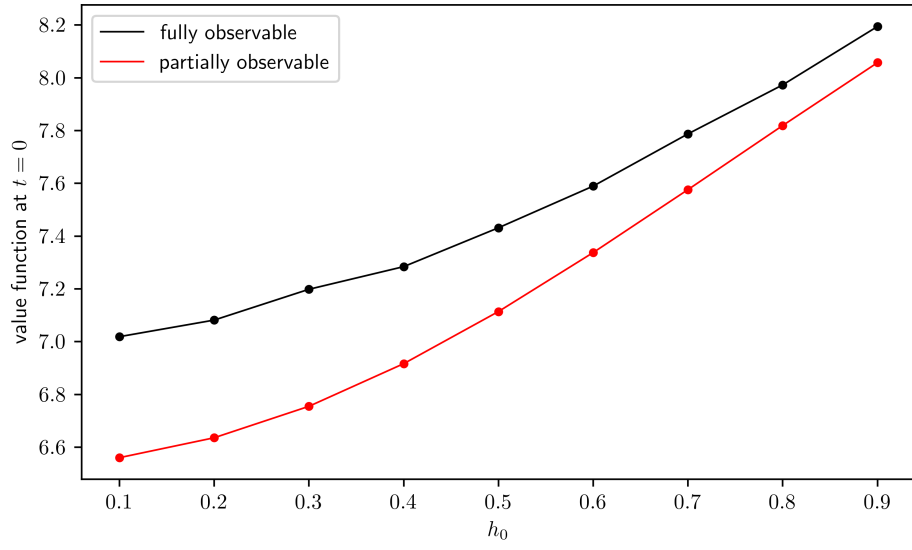


Figure 6.3: Full information value and partial information value at $t = 0$
(Bayesian filter for unconstrained problems)

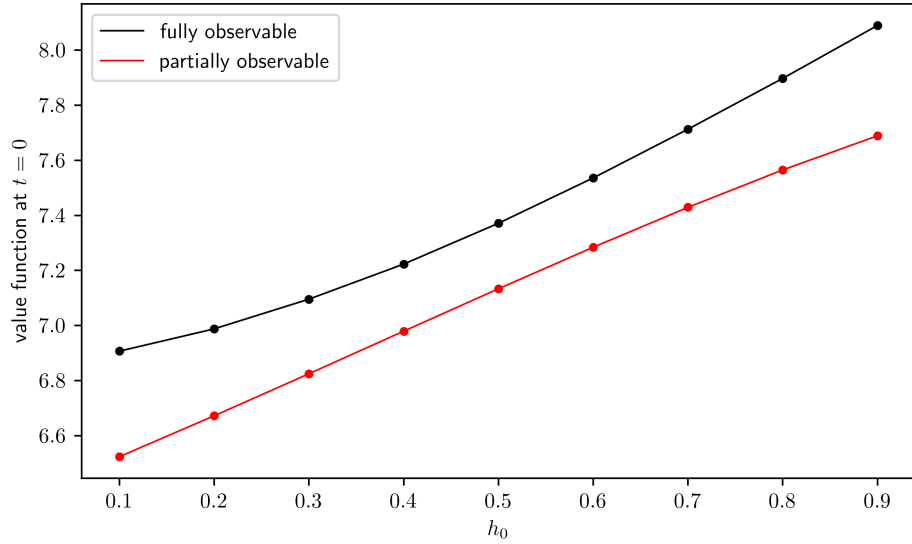


Figure 6.4: Full information value and partial information value at $t = 0$
(Bayesian filter for constrained problems)

Chapter 7

Conclusion

In this paper we derive theoretical results on transforming partially observable utility maximization problems into their fully observable analogues by the Kalman-Bucy filter and the Bayesian filter respectively. The drift is not observable, but when we estimate it by filters and then use the estimated drift to replace the original drift, everything becomes observable.

The fully observable problem is a standard stochastic control problem and can be solved in two ways. The first one is deriving HJB equations and find some approach to solve them. Alternatively, we can derive dual HJB equations and obtain a system of FBSDEs by SMP. Solving this system gives the solution to the original problem.

Due to the complexity embedded in these problems, numerical methods are always necessary. The deep learning approach (i.e. the PINN technique) is applied to solve HJB equations. In contrast, the system of FBSDEs is solved by the bound estimate approach, where we estimate lower bounds and upper bounds for value functions.

In numerical experiments, we consider 1D problems and 2D problems. The deep learning approach and the bound estimate approach are both tested for 1D problems, while for 2D problems we only test the bound estimate approach. Several key points are observed from results.

For 1D problems, the first point is that the bound estimate approach can effectively estimate value functions, because lower bounds and upper bounds are close to each other (with almost all relative errors less than 3%). Therefore, the average of any pair of bounds is a good estimate of the corresponding value function. Second, as the initial mean h_0 and Σ_0 become larger, the value function increases as well. Finally, results generated by the deep learning approach are close to those generated by the bound estimate approach (relative errors between value functions computed by these two approaches are less than 1% in almost all cases). This further verifies the correctness of the bound estimate approach.

We also obtain lower bounds and upper bounds of value functions for different kinds of 2D problems. It turns out that intervals formed by these bounds are narrow (with almost all relative errors less than 3%), meaning that the bound estimate approach works well for 2D problems.

Moreover, we show that information has positive value, because the value function generated by more information takes greater values consistently. When computing the value function of the fully observable problem, the Monte Carlo method is applied because the initial drift is random and follows the prior distribution.

Both approaches discussed in this paper have advantages and limitations. The deep learning approach can give value functions in a whole region. In other words, a carefully trained neural network $\hat{J}_\theta(t, x, h)$ can generate an output for each point in the solution region $\Omega = [0, T] \times [X_{\min}, X_{\max}] \times [0, H_{\max}]$. But when t, x or h is changed, the bound estimate approach must be executed again to compute new results. Also, A neural network

computes value functions with high speeds, while the bound estimate approach must run Monte Carlo simulations to compute bounds, which is time-consuming. Additionally, a neural network can handle high-dimensional problems efficiently, while the bound estimate approach suffers from the curse of dimensionality: when the investor trades D stocks and the total number of steps is N , the time complexity of the Monte Carlo simulation is approximately $\mathcal{O}(ND)$, meaning that running the bound estimate approach can be extremely slow when D is large. Although the deep learning approach has the advantages above, there are still some problems for it. The training process of neural networks is often computationally expensive and requires a significant amount of time. Also, the risk of overfitting cannot be fully avoided. Early stopping may work but it also brings the risk of underfitting in particular scenarios. An additional drawback of neural networks is that they are trained by data points sampled in the solution region Ω . Since Ω is fixed, when we need to compute value functions at points out of Ω , we have no option but to train the model again. For future research, finding more efficient approaches without disadvantages mentioned above is a possible direction.

Bibliography

- [1] N. M. APOLLINAIRE AND P. N. AMANDA, *Stochastic optimal control theory applied in finance*, Science, 7 (2022), pp. 59–67.
- [2] F. BARTALONI, *Intertemporal utility maximization problems with state constraints: existence theorems and dynamic programming*, (2014).
- [3] X. DAI, J. ZHANG, AND V. CHANG, *Noise traders in an agent-based artificial stock market*, Annals of Operations Research, (2023), pp. 1–30.
- [4] A. DAVEY AND H. ZHENG, *Deep learning methods for s shaped utility maximisation with a random reference point*, arXiv preprint arXiv:2410.05524, (2024).
- [5] C. DE FRANCO, J. NICOLLE, AND H. PHAM, *Bayesian learning for the markowitz portfolio selection problem*, International Journal of Theoretical and Applied Finance, 22 (2019), p. 1950037.
- [6] I. EKELAND AND T. A. PIRVU, *On a non-standard stochastic control problem*, arXiv preprint arXiv:0806.4026, (2008).
- [7] W. H. FLEMING AND T. PANG, *An application of stochastic control theory to financial economics*, SIAM journal on control and optimization, 43 (2004), pp. 502–531.
- [8] M. FUJISAKI, G. KALLIANPUR, AND H. KUNITA, *Stochastic differential equations for the non linear filtering problem*, (1972).
- [9] I. KARATZAS AND X.-X. XUE, *A note on utility maximization under partial observations 1*, Mathematical Finance, 1 (1991), pp. 57–70.
- [10] T. KNISPEL, *Asymptotics of robust utility maximization*, (2012).
- [11] Y. LI AND H. ZHENG, *Dynamic convex duality in constrained utility maximization*, Stochastics, 90 (2018), pp. 1145–1169.
- [12] R. S. LIPTSER AND A. N. SHIRYAEV, *Statistics of random processes II: Applications*, vol. 6, Springer Science & Business Media, 2013.
- [13] R. C. MERTON, *Lifetime portfolio selection under uncertainty: The continuous-time case*, The review of Economics and Statistics, (1969), pp. 247–257.
- [14] M. MONOYIOS, *Infinite horizon utility maximisation from inter-temporal wealth*, arXiv preprint arXiv:2009.00972, (2020).
- [15] M. NISIO, *Stochastic control theory: Dynamic programming principle*, vol. 72, Springer, 2014.
- [16] A. PAPANICOLAOU, *Backward sdes for control with partial information*, Mathematical Finance, 29 (2019), pp. 208–248.

- [17] H. PHAM, *Continuous-time stochastic control and optimization with financial applications*, vol. 61, Springer Science & Business Media, 2009.
- [18] D. ZHU AND H. ZHENG, *Effective approximation methods for constrained utility maximization with drift uncertainty*, *Journal of Optimization Theory and Applications*, 194 (2022), pp. 191–219.